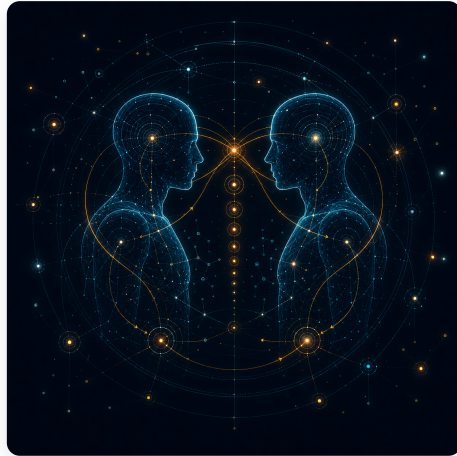




Adaptive Self-Governance

The Reflexive Controller and the Limits of Self-Revision

Companion to The Self-Variety Gap and Cycle Two of the Governance as Engineering series

Companion to The Self-Variety Gap and Cycle Two of the Governance as Engineering series. Applies the adaptation triad — observer diversity, actuation integrity, self-legitimacy, boundary calibration, and adaptive learning — to the self as a controller whose controller and plant are the same system. Includes formal appendices on self-observer correlation, actuation chain attenuation, self-legitimacy dynamics, and observer–plant identity.

Björn Kenneth Holmström

June 2026

Creative Commons Attribution-ShareAlike 4.0 International

<https://bjorkennethholmstrom.org/working-papers/adaptive-self-governance>

Abstract

A companion paper established that a person's values function as an observation architecture: a value function eventually optimizes away its own ability to perceive the self when it omits a dimension causally coupled to the proxy it optimizes (not merely by being lower-dimensional than the self), and the excluded dimensions return as crises the person cannot trace to their source. That argument is static — it concerns whether a self can perceive itself at a given moment. This paper extends the framework to the adaptive problem: how a self continues to perceive, act, trust, learn, and bound itself across a changing life. Drawing on the second cycle of the Governance as Engineering series, it shows that a viable self must do five structurally distinct things at once — maintain decorrelated channels of self-observation, carry intention into action across an internal delegation chain, sustain trust in its own perceptions and commitments, place a workable boundary between itself and others, and keep revising itself as circumstances change — and that failure in any one destabilizes the rest.

The paper introduces one structural feature not present in the parent series: *observer–plant identity*, the fact that in a self the controller and the governed system are the same entity. This is a formalization of second-order cybernetics rather than a discovery, but taken seriously as a constraint it predicts specific and non-obvious ways the self-scale results diverge from their institutional originals. Its most general consequence is that self-observation is never only observation — to attend to a state is to alter it — so that control theory's separation of estimation from control fails for a self; the same coupling makes sustained self-observation a means of change and not merely of knowledge. From this base the paper develops the five primitives, two of whose results it foregrounds. *Constitutional self-uncontrollability* — beyond a critical delegation depth, an intention leaves the reachable set regardless of effort — is the exact dual of the companion paper's constitutional unobservability, so that the two papers together bound self-governance on both of its channels: there are states of the self a person cannot see, and states a person cannot reach, each limit architectural rather than volitional. And self-revision is shown to be bounded on *two* sides — below by the need to stay calibrated, above by the rate at which a revised self can re-cohere — with the upper bound tighter for a self than for any institution, because a self cannot insulate its experimenting apparatus from its own experiments.

The epistemic posture is stated at the front because it governs everything that follows. This is the series' softest empirical territory: one cannot randomize people to self-observation architectures or self-legitimacy histories and measure lifetime outcomes with the tractability of the parent programme's empirical work, so almost every substantive claim here is in-principle rather than confirmed, and the rigorous core is narrow and inherited from the parent series rather than freshly established at this scale. The mappings between governance and self are in several places clean enough to be seductive, and the paper treats that cleanness as a hazard rather than a credential, marking for each primitive whether the transfer is genuine scale-invariance or mere analogy and ranking the mappings accordingly — the boundary mapping, the weakest, is shown to fail at the formal level where the systems on the far side of a self–other boundary become agents that model the controller back. Throughout, the paper offers architectural diagnosis, not normative prescription: it specifies conditions under which a self remains adaptive, not a way to live, and its closing turn toward joy, beauty, and truth is presented as a bounded statement of the framework's limit — the content no objective function can hold — rather than as a conclusion the framework claims to have reached.

Contributions

- **Introduces observer–plant identity** — the controller and the plant as one entity — as the structural feature that distinguishes self-governance from institutional governance in kind rather than scale, and derives the **measurement–disturbance coupling**, the failure of control theory's separation principle for a self, as its most general consequence. The contribution is the demonstration that this single feature predicts the specific divergences of the self-scale results, not the bare observation that the self is reflexive.
- **Extends the Cycle Two adaptation grammar** — **Sense** → **Learn** → **Execute** — **to the self** through five primitives (observer diversity, actuation integrity, self-legitimacy, boundary calibration, adaptive learning), under an explicit test separating scale-invariance from analogy, with the mapping strengths ranked and the weakest shown where it breaks.
- **Establishes constitutional self-uncontrollability as the dual of constitutional unobservability**, so that the companion paper and this one together bound self-governance on both channels — the states a person cannot perceive and the states a person cannot reach, each constitutional rather than volitional.
- **Derives a two-sided bound on self-revision** — calibration below, coherence above — and shows the upper bound has no full institutional analogue, a direct consequence of observer–plant identity that gives structural content to the observation that transformation must be paced to integration.
- **Identifies composite failure as the default mode of self-governance breakdown**. Because the primitives are multiplicatively coupled (the coordination-failure tax of Paper V) and share one substrate with no firewall between them, failures propagate across primitives; the gravest, the *sealed self*, fuses observer collapse, model lock-in, and the transparency trap into a condition the self cannot perceive from within.
- **States the design principles as a coupled set under a single logic** — recruit, outside the reflexive substrate, the structure the self cannot be for itself — with one internal exception, sustained self-observation as intervention, and derives a sequencing claim: a coupled failure is entered not at its worst-damaged primitive but where self-legitimacy, the multiplicative gain on every other primitive, is rebuilt.

Epistemic tiers. Every load-bearing claim carries one of three tags. **[R]** marks a result that follows from established mathematics within its domain of validity, or a rigorous statement of the framework's own limits. **[IP]** marks a claim that follows structurally from the rigorous core once applied to the self but is not empirically confirmed at that scale. **[H]** marks an interpretive or illustrative claim that makes the structure legible without bearing the argument's weight. The companion paper used the scheme [R]/[I]/[S]; this paper adopts the corpus-standard [R]/[IP]/[H], a tightening rather than a relabeling. Almost the entire substantive content of this paper is [IP]: structurally entailed, not independently confirmed — a posture Part 0 develops before the argument begins.

Part 0 — Epistemic Frame

Self I (The Variety Gap in the Self) argued that a person's values function as an observation architecture: a low-dimensional value function destroys information about the self in the same way a low-dimensional governance metric destroys information about a society, and the excluded dimensions return as crises the person cannot trace to their source. That paper mapped onto the first cycle of the Governance as Engineering series — the cycle concerned with perception. This paper extends the framework to the second cycle, the cycle concerned with adaptation: how a self continues to perceive, act, trust, and learn across a changing life rather than at a single moment. The extension draws on Papers X through XIV of the parent series, and it is, by a wide margin, the series' softest empirical territory. This frame states why, and what discipline the paper adopts in consequence, before the mappings begin — because the mappings are clean enough to be seductive, and the cleanness is a hazard rather than a credential.

0.1 The Three Tiers

Every load-bearing claim in this paper carries one of three tags, applied to the claim at the point it is made:

- **[R] — rigorous.** A claim that follows from established mathematics applied within its domain of validity, *or* a rigorous statement of the framework's own limits. The second clause matters here more than anywhere else in the series: where this paper marks a boundary it must not cross — the line between the structure of a difficulty and its clinical cause — the boundary statement is itself a rigorous result and is tagged accordingly, not treated as a hedge.
- **[IP] — in-principle.** A claim that follows structurally from the rigorous core once it is applied to the self: sound in its logic, but not empirically confirmed at the self scale, and in several cases not confirmable by the methods available to the parent programme.
- **[H] — heuristic.** An interpretive or illustrative claim — a psychological reading, a particular life instance, a magnitude — offered to make the structure legible. It is explicitly not bearing the argument's weight, and the argument survives its being wrong.

Self I used the tags [R]/[I]/[S] — rigorous, interpretive, speculative. This paper adopts the corpus-standard [R]/[IP]/[H] for consistency with Papers X–XIV. The correspondence is not one-to-one: Self I's interpretive [I] tier splits, with its structural applications becoming [IP] and its illustrative content becoming [H], and its speculative [S] tier becomes [H] except where it is quarantined as a bounded horizon (Part IX). A reader carrying Self I's scheme forward should read the migration as a tightening, not a relabeling.

0.2 What Carries Each Tier Here

The governing test throughout is a single distinction: for each primitive borrowed from the parent series, is the mapping to the self *scale-invariance* — the same mathematics, a different controller — or is it *analogy* — it rhymes? The paper claims the former only where the former holds, and flags the latter as what it is. The two are not interchangeable, and treating an analogy as an identity is the specific failure this frame exists to prevent.

By that test, the [R] tier in this paper is narrow and borrowed. It comprises the results that are scale-invariant by construction — the correlated-ensemble variance correction underlying observer diversity (Part II), the dual-control and persistent-excitation structure underlying adaptive learning (Part VI), the coupling algebra underlying self-legitimacy (Part IV) — together with the limit-statements of §0.1. Its rigor is *inherited* from the parent series, where these results are anchored in formal derivation and, in the case of observer correlation, in a preregistered empirical test. The self-scale contribution is not to re-establish this mathematics but to show that it applies — that a person's self-observation channels

are a correlated ensemble, that a life is a dual-control problem, and so on. That demonstration is structural, and it lands the application in the [IP] tier even where the underlying mathematics is [R]. Almost the entire substantive content of this paper is therefore [IP]: structurally entailed, not independently confirmed. The [H] tier carries the psychological color — the readings of stagnation, self-betrayal, the curated confidant — that make the structure recognizable without supporting it.

0.3 The Empirical Asymmetry

The parent series earned the right to its [IP] claims, partly through fifteen country cases and a preregistered study in which consumer AI systems estimated governance-relevant quantities, returning the near-total observer correlation that Part II invokes. No comparable anchor is available at the self scale, and the reason is structural rather than a matter of effort not yet expended. One cannot randomize people to different self-observation architectures, self-legitimacy histories, or exploration rates and measure lifetime wellbeing with anything like the tractability of querying six AI systems on fifty items. Some claims here may eventually meet correlational or longitudinal evidence — self-complexity instruments already exist, and the parent programme's gate for opening theory cycles is confirmed prediction, not theoretical coherence — but the strong experimental confirmation that disciplines the governance theory will not, by the same route, reach the self.

The consequence is that the tier discipline does *more* work in this paper than in any other in the series, not less. With the empirical floor removed, the tags are most of what stands between a structural argument and a persuasive one. This is also why the elegance of the mappings is a warning. When a translation between domains comes out as clean as several of these do, the cleanness can mean the structure is genuinely scale-invariant, or it can mean the metaphor is doing work the mathematics is not. The reader is owed that ambiguity stated in advance, and owed it most loudly exactly where the mapping feels most inevitable.

0.4 The Diagnostic Contract

The paper maintains the series' standing discipline: it offers architectural diagnosis, not normative prescription. It describes conditions under which a self remains adaptive — decorrelated observation, short actuation chains, built rather than borrowed self-trust, calibrated exploration — and it derives design principles as statements about architecture rather than as advice about how to live. It does not say what a person should value, what a flourishing life is, or whether any particular reader should undertake the self-revision it analyzes. Where the later material turns toward joy, beauty, meaning, and truth (Part IX), that material is presented as a bounded speculative horizon and tagged as such — as a statement of the framework's *limit*, the things it can specify conditions for without specifying, rather than as a conclusion it claims to have proved. The optimistic register of the closing pages is fenced in advance precisely so that it cannot be mistaken for a result.

One structural element in this paper is new rather than borrowed: *observer–plant identity*, the fact that in a self the controller and the governed system are the same entity. It is the one primitive Self II adds to the series' grammar, it is developed in Part I, and it recurs as the reason several of the self-scale results diverge from their institutional originals rather than merely shrinking them. It is introduced here only so the reader knows that what follows is not pure translation — that the paper claims one genuinely new structural feature, and rests the weight of its departures from the parent series on that single claim.

Part I — The Adaptive Problem and the Reflexive Self

1.1 *The Limit of Self I*

Self I established a necessary condition for self-governance and, in doing so, exposed its own insufficiency. The condition was perceptual: a person whose value architecture has too few dimensions to span the disturbance environment of their life is structurally blind to the sources of their own distress, and no quantity of sincerity or effort can recover a signal the architecture never admitted. That argument is complete as far as it goes. But adequate self-perception, even where it is achieved, governs nothing on its own. It is the precondition for self-governance, not the act of it.

The insufficiency shows in two failures that look like opposites. The first is the person who perceives themselves accurately and does not change — who can describe their pattern with precision, name what it costs them, articulate exactly what a better course would be, and continue, year after year, in the pattern they have described. Perception is intact; nothing downstream of it moves. The second is the person who changes constantly and loses themselves — who revises, relocates, and reinvents so readily that no version of the self persists long enough to be lived, and who arrives at the accumulated wreckage of many fresh starts with no coherent identity to show for any of them. Accurate perception, in the first case, produces no action; rapid action, in the second, dissolves the perceiver. Neither failure is a failure of self-knowledge in Self I's sense, and that is the point: self-knowledge is necessary and does not suffice.

What the two failures reveal is that a viable self must do several structurally distinct things at once. It must perceive itself adequately — and through more than one channel, since a single channel cannot detect its own systematic error. It must translate what it perceives into action across an internal delegation chain that attenuates intention. It must maintain trust in its own perceptions and commitments, since a self that disbelieves its own reports or expects to break its own promises governs neither. It must draw a workable boundary between itself and others. And it must keep learning as the self and its circumstances change, without the learning destabilizing the self that must absorb it. Failure in any of these destabilizes the rest. This is the adaptive problem, and it is the subject of the present paper. Running through all of it is one structural feature that makes the self-case categorically different from the institutional case the parent series studied — different in kind, not merely in scale — and the rest of this part is devoted to it. **[IP]**

1.2 *Observer–Plant Identity*

In the governance systems the parent series analyzes, the controller and the plant are distinct entities. A ministry governs an economy; a regulatory agency governs an industry; the controller is one thing and the system it observes and acts upon is another. This separation is so basic to control theory that it is rarely stated: the observer is outside the observed. It licenses the field's standard architecture, in which a controller measures the state of a separate system, compares it to a target, and acts to close the difference.

A self does not have this separation. The entity that observes is the entity observed; the apparatus that issues directives is the apparatus that must comply with them; the system being regulated is the system doing the regulating. Call this *observer–plant identity*: in self-governance, the controller and the plant are the same entity. It is the one structural feature this paper adds to the series' grammar, and it is less a discovery than a formalization. That the self observing is the self observed is the recognition at the root of second-order cybernetics (von Foerster), and Self I already noted its presence when it read Acceptance and Commitment Therapy's distinction between self-as-content and self-as-context as second-order cybernetics at the personal scale. The contribution here is not the bare observation that the self is reflexive. It

is the demonstration that this single feature, taken seriously as a structural constraint, predicts specific and non-obvious ways in which the self-scale results diverge from their institutional originals — ways developed in the parts that follow. **[R]** for the structural statement; the consequences are **[IP]**.

1.3 The Measurement–Disturbance Coupling

The most general consequence of observer–plant identity is also the one with no institutional analogue, and it concerns the act of observation itself.

When a governance system measures its plant, the measurement leaves the plant approximately undisturbed. A census does not appreciably change the population it counts; a survey of economic activity does not, to first order, alter the economy it samples. This near-independence of measurement from the measured is what underwrites the *separation principle* of control theory — the result that estimation and control can be designed independently, that a system can first determine its state and then decide what to do about it. The two operations are distinct because observing and acting fall on different objects.

For a self, they fall on the same object, and the separation collapses. To attend to a sensation is to change it — to amplify it, dampen it, or transform it into something it was not before attention arrived. To ask oneself how one feels about a matter is to alter the feeling in the asking. To narrate one's situation is to constitute the situation as much as to record it; the story is not a read-out of a prior state but partly the production of the state it purports to describe. Self-observation is therefore never only observation. It is simultaneously an intervention on the system observed, because there is only one system. The separation principle fails for the self: a person cannot first read off their state and then decide what to do, because the reading is already a doing. **[IP]**

This has a consequence that should be stated plainly because it qualifies Self I. The "true state" $\mathbf{x}_{\text{self}}$ that Self I's observation equation projects into awareness is not a fixed target the person reads with greater or lesser fidelity. It is partly constituted by the act of reading. "Accurate self-perception" remains a meaningful aim, but it is not the recovery of a pre-existing fact; it is closer to a stable fixed point of a process in which observing and being observed are the same motion. The framework does not resolve this — it marks it, as the point at which the comfortable metaphor of the self as a system with a state to be measured reaches its limit. **[H]** for the philosophical gloss; the structural claim that measurement and disturbance are coupled in the self is **[IP]**.

1.4 The Double Edge

It would be easy to read observer–plant identity as a catalogue of additional handicaps — the self as a controller burdened with reflexive defects the ministry is spared. That reading is half the picture, and the missing half is what makes self-governance richer than governance rather than merely more compromised.

The same coupling that contaminates self-observation also makes self-observation a control input. Because attending to a dimension of the self changes it, the act of attention is a lever, not only a sensor. This is the structural reason the practices Self I gathered under its interpretive and speculative layers — sustained attention, contemplative training, the disciplined re-narration of one's own experience in therapy — can be transformative rather than merely informative. They do not work by passively improving the read on a fixed state; they work because, in a system where observing is acting, sustained observation of a dimension is sustained action on it. Self I's gesture toward meditation reducing the self-variety gap has its mechanism here: attending to dimensions the value architecture excluded is, in a reflexive system, the same operation as beginning to admit them. The asymptote that Self I named — full self-observability, $G_{\text{self}} \rightarrow 0$ — is reachable, to whatever degree it is reachable at all, only because in a self the act of seeing is already an act of becoming. **[IP]** for the structural claim; **[H]** for the contemplative reading, which Part IX returns to under its bounded-horizon discipline.

So observer–plant identity is double-edged. It is the source of the self's characteristic difficulties — the difficulties that the rest of the paper specializes and that have no clean institutional remedy — and it is, at the same time, the source of the self's characteristic capacity, the fact that a self can change itself by observing itself in a way no ministry can change its economy by surveying it. The paper holds both edges throughout, and resists the temptation to let the diagnostic edge crowd out the other.

1.5 Three Specializations and the Shape of the Paper

Observer–plant identity does not produce one consequence but a family of them, and the paper is organized so that its three sharpest specializations each carry a part of their own.

In **observer diversity** (Part II), identity places a floor under the error correlation among a person's self-observation channels: several of those channels are not merely correlated but are the same apparatus reporting more than once, and no internal operation can decorrelate a faculty from itself. The institutional remedy — constitute the observers differently — is unavailable, and the only route below the floor runs through channels situated partly outside the observing self.

In **self-legitimacy** (Part IV), identity makes self-trust unauditably from the inside: the apparatus losing trust is the apparatus that would measure the trust, so the self-deceiving instrument is also the instrument charged with detecting the deception. The institutional remedy — an independent legitimacy sensor — has, again, only an external approximation.

In **adaptive learning** (Part VI), identity converts a one-sided requirement into a two-sided bound: a society can explore aggressively because it insulates its experimenting apparatus from its experimental zones, while a self cannot fully insulate itself from its own experiments, so its self-revision is bounded not only below, by the need to stay calibrated, but above, by the rate at which the revised self can re-cohere. The experimenter is the experiment.

These three sit within the adaptation architecture the parent series arrived at in its second cycle: **Sense** → **Learn** → **Execute**. Part II is the self-scale *sensing* requirement; Part III, on the internal delegation chain from intention to act, is the *executing* requirement; Part VI is the *learning* requirement that mediates between them. Two further variables are not stages in this sequence but conditions on all of it: self-legitimacy (Part IV), the trust that couples self-observation and self-action and without which neither functions, and boundary calibration (Part V), the drawing of the perimeter that determines which dynamics the self treats as its own — the part where observer–plant identity is weakest, because the systems coupled to a self are themselves agents, and the paper says so where it arrives. Part VII gathers the composite failure modes; Part VIII states the design principles as architecture; Part IX marks the framework's outer limit.

The recursion this sets up is the one Paper XIV identified as the convergence point of the second cycle's mathematics. A governance system safely modifying its own architecture, an AI system safely modifying its own code, and a self safely revising its own identity are, structurally, one problem — a controller attempting to change what it is without destroying the continuity that would let it learn from the change. Observer–plant identity is the name, for the self, of what makes that problem its own rather than borrowed: the controller and the plant cannot be pulled apart, and everything that follows is what that single fact entails.

Part II — Observer Diversity: Why Self-Knowledge Cannot Be Audited From Inside One Channel

A person facing a hard decision — whether to leave, whether something is wrong, whether they are deceiving themselves — does the responsible thing. They reflect carefully. They talk to the friends they trust. They keep a journal. They take soundings from several people who know them well. The readings converge; everyone agrees; they proceed with confidence, and they are blindsided. Afterward, the people closest to them say they did not see it coming either.

This is not a failure of effort or sincerity, and it is not the variety gap of Self I. The person sampled many observers and reflected at length. The failure is structural and lies one level up: the observers were not independent of one another, and a consensus among correlated observers carries no more information than the opinion of a single one. Self I modeled self-perception as a single observation channel — the value architecture as the matrix $\mathbf{C}_{\text{self}}$ projecting the true state into awareness — exactly as Papers I–IX of the parent series modeled each controller through one observation matrix. But a person does not perceive themselves through one channel. They have many, and the decisive property of that population is not the quality of any single member but the decorrelation of the ensemble. This part extends the self-governance framework from the individual channel to the observing population, following Paper X's extension of the parent series from the controller to the ensemble. [IP]

2.1 The Self-Observation Ensemble

A person observes themselves through a population of partially distinct channels: introspection, the running inner narrative, interoception (the signals of the body), close friends, family, a partner, a therapist, a journal, dreams, creative practice, contemplative practice, and — increasingly — conversation with AI systems. Each is an observer of the same latent state: the person's true multidimensional condition $\mathbf{X}_{\text{self}}$. Following Paper X, the i -th channel has an observation equation

$$\mathbf{y}_i = \mathbf{C}_i \mathbf{X}_{\text{self}} + \boldsymbol{\epsilon}_i$$

where $\mathbf{C}_i$ is the channel's structural perspective — which dimensions of the self it can perceive, and at what resolution — and $\boldsymbol{\epsilon}_i$ is its error. The body perceives dimensions the inner narrative omits; a friend perceives the person's effect on others, which introspection cannot directly sample; a journal preserves a perspective the present self has edited away. The ensemble is the composite of all these channels, and its capacity is governed by two properties Paper X makes precise: the effective rank of the stacked observation — whether, between them, the channels cover the dimensions of the self that matter — and the correlation structure of their errors.

The first property gives a requisite-diversity condition for the self. Let $\dim(\mathbf{U}_{\text{self}})$ be the dimensionality of the *self-uncertainty space* — the directions in which a person cannot reliably predict their own state, where the gap between expected and actual self is large enough to matter. Requisite observer diversity for the self is the condition that the ensemble's effective rank cover that space:

$$r_{\text{ens}}^{\text{self}} \geq \dim(\mathbf{U}_{\text{self}}).$$

When it fails, there are dimensions of the self that *no* channel perceives — a blind spot shared by introspection, friends, and body alike — and the blind spot is undetectable by cross-referencing, because no channel has independent access to the missing dimension. This is the ensemble-level analogue of Self I's variety gap: not "the value architecture excludes this dimension" but "every available observer is blind to it at once." [IP]

2.2 The Correlation Problem

The intuition that more observers are better rests on a result that holds only under independence. For N channels with individual error variance σ^2 and pairwise error correlation ρ , the variance of the ensemble's pooled estimate is not σ^2/N but

$$\text{Var} = \sigma^2 \left(\frac{1 - \rho}{N} + \rho \right),$$

and the *effective number of independent observers* is

$$N_{\text{eff}} = \frac{1}{(1 - \rho)/N + \rho} = \frac{N}{1 + (N - 1)\rho}.$$

This is the Kish design effect from survey statistics, and it is rated **[R]**: standard mathematics that applies to any correlated ensemble, a person's self-observation channels included. Its content is unforgiving. When $\rho = 0$, $N_{\text{eff}} = N$ and numbers help fully. When $\rho = 0.5$, N_{eff} approaches 2 no matter how large N grows — a shared bias sets a noise floor that no quantity of observers can lower. When $\rho \rightarrow 1$, $N_{\text{eff}} \rightarrow 1$: the ensemble retains N nominal observers and receives the statistical protection of one.

The uncomfortable claim of this part is that a person's self-observation channels are correlated — often heavily — and that the correlation runs through the person themselves. You select the friends you confide in, and you tend to select those who already see you as you wish to be seen. You shape what you tell your therapist by what you are willing to say. Your introspection and your inner narrative are not two observers but one apparatus reporting twice. The nominal channel count therefore overstates N_{eff} badly, and the person of the opening vignette — five trusted friends plus their own careful reflection — did not consult six observers. They consulted one observer, themselves, six times, and mistook repetition for confirmation. **[IP]** for the structural claim; the specific reading is **[H]**.

2.3 Two Pathways to Correlated Self-Observation

Paper X distinguishes two routes by which an ensemble's error correlation approaches unity, and both have personal forms.

Model-based monoculture arises when channels share a common processing architecture. At the personal scale, the shared architecture is the self-narrative and the stable set of cognitive biases through which all reports are filtered. Even genuinely diverse raw signals — a friend's blunt remark, an ache in the body, a dream — are passed through one interpretive model before they reach awareness, and the model's biases are imposed identically on all of them. The correlation arises in the *processing*: diverse data, one interpreter, correlated readings.

Data-based monoculture arises when channels share a common input substrate. At the personal scale, this is the curation of one's environment — the friends kept, the media consumed, the situations entered and avoided — selected, often without intent, to deliver a consistent picture. Here the information is truncated or skewed *before* any channel processes it, so that even an unbiased interpreter would converge with the others on the same partial view.

Following Paper X, the two compound:

$$\rho_{\text{total}} \approx 1 - (1 - \rho_{\text{model}})(1 - \rho_{\text{data}}),$$

so that moderate correlation through each pathway drives ρ_{total} toward one. Self-deception, where it operates, rarely confines itself to one route: the narrative that filters incoming signals (raising ρ_{model}) is the same narrative that selects which signals are sought at all (raising ρ_{data}). **[IP]** for the two-pathway structure; **[H]** for the psychological gloss.

2.4 The Reflexive Floor

Here the structural feature from Part I — observer–plant identity — sets a floor under ρ that institutional governance does not face. In governance the N observers are distinct organizations; decorrelation is achievable in principle by constituting them differently. In the self, several channels are not merely correlated but *physically the same apparatus*: introspection and the inner narrative share one substrate and cannot, by any internal operation, be made independent of each other. No quantity of reflection lowers the correlation between a faculty and itself. This places a floor under ρ that the self cannot reach below from the inside.

The consequence is precise and is the structural core of this part: the only channels that can lower a person's effective ρ are those situated partly *outside* the narrating apparatus — the body, whose signals are generated before interpretation; a friend who will say the unwelcome thing; a journal entry written before the current story formed and re-read after it has moved on; a therapist; creative practice that surfaces what the narrative suppresses. These are the self's only approximation to an independent sensor, which is the same role they will play, from the trust side, in Part IV. A confidant selected for agreement is not a second sensor; they are the person's own observation channel wearing another face, and they add to nominal N while contributing nothing to N_{eff} .

An empirical anchor, with its caveat stated plainly. Study 1 of the parent programme found that six nominally distinct consumer AI systems, estimating fifty governance-relevant quantities, had an effective error correlation $\rho_{\text{eff}} \approx 0.97$ — near-total correlation despite nominal independence (the result is **[R]**; its primary prediction was preregistered and confirmed). Its bearing on selves is analogical, not direct: no selves were measured, and the transfer is **[H]**. But the result establishes that nominal independence can coexist with near-complete correlation in practice, and a self whose channels share a substrate and are actively selected by the very party under audit has, if anything, weaker grounds than a set of separately built AI systems to expect its observers to disagree when it counts.

2.5 Collapse Dynamics

Paper X derives a selection gradient that drives diverse ensembles toward monoculture even when every step is locally rational, and the same gradient operates on a person's observers. Maintaining a genuinely dissonant channel is costly: the friend who reliably disagrees is uncomfortable to keep close, the journal that contradicts the current story is unpleasant to re-read, the bodily signal that undercuts the plan is easier to override than to heed. At each juncture the locally comfortable move is to weight the confirming channel more and the dissonant one less, and the cumulative drift raises ρ , lowers N_{eff} , and narrows the ensemble toward agreement — the personal form of Paper X's liability ratchet toward consensus. **[IP]**

At the limit lies the *echo chamber of one*: every channel agrees because every channel has become, in effect, the same channel. The consensus is unanimous and the confidence total, and the dimension the shared narrative excludes — Self I's variety gap — is exactly the dimension to which every remaining observer is blind. This is why the eventual crisis arrives as a shock not only to the person but to everyone around them: the people closest were, over years, selected to share the blind spot, so their agreement was never an independent check. The unanimity that felt like safety was the symptom of the failure, not protection against it. **[IP]**, with the specific course **[H]**.

The connection to Part IV runs both ways. The curation of one's observers toward agreement is the transparency trap operating on the observation ensemble rather than on a single self-report, and a self that has lost trust in its own perceptions (L_{self} falling on the observation channel) is the most prone to it — replacing channels that disturb the maintained picture with channels that confirm it, which lowers the felt distress while raising the structural ρ that guarantees the next blindsiding.

2.6 The Design Principle and Its Inversion

The diagnosis yields a condition, stated — per the series' discipline — as architecture rather than as advice. Adequate self-knowledge requires maintaining self-observation channels that are *partially decorrelated*: channels that do not share the narrating substrate and were not selected for their agreement. The arithmetic of N_{eff} makes the relevant asymmetry sharp. A confirming observer raises nominal N and leaves N_{eff} almost unchanged; a tenth friend who agrees alters nothing structural. Only decorrelated channels raise effective diversity, and decorrelated channels are, with some reliability, the uncomfortable ones — the disagreeing friend, the unedited record, the body's veto, the practice that returns what was suppressed.

This inverts the ordinary comfort of consensus. In self-knowledge, agreement among correlated channels is not evidence, because the agreement was structurally guaranteed before any of them looked; what carries information is *divergence from a channel known to be decorrelated*. The signal a person should weight most is the reading that does not fit and comes from somewhere they did not choose for its agreeableness. The discomfort is not incidental to its value. It is, frequently, the mark of the one observer in the ensemble who is not simply the self again.

A sharper form of the same point, which the formal model of Appendix A makes precise, is that merely *possessing* a decorrelated channel is not enough. Under naive equal weighting — treating every channel's reading as one equal vote — the effective number of independent observers is the inverse Herfindahl index of how the nominal channels fall into decorrelated blocks, and this falls *below* the number of blocks whenever the blocks are unequal in size. A person with five mutually correlated internal channels and one genuinely external one does not have the diversification of two observers but closer to that of one and a third: the fused internal majority outvotes the single external reading. Recovering the external channel's value therefore requires *deliberately over-weighting* the dissonant reading against the internal consensus, not merely letting it cast its equal vote. The discomfort of doing so is the felt form of the up-weighting the arithmetic demands.

Part III — Actuation Integrity: Why Intention Does Not Reach Behavior

A person decides, clearly and sincerely, to do something — to exercise, to write, to make the call they have been avoiding, to change how they spend their evenings. The decision is not idle. They understand why it matters, they have an accurate picture of their situation, they want the outcome, and no part of them is mounting serious opposition. Weeks later, little has happened. Nothing dramatic intervened: no crisis, no reversal, no loss of conviction. The intention was simply attenuated — delayed, narrowed, and translated, one internal step at a time, until what reached the level of action bore little resemblance to what was decided, and much of it never reached that level at all. This is the Oakland of the self, and Self I cannot account for it, because the failure is not perceptual. The person sees correctly. What fails is downstream: the channel through which an intention must travel to become a behavior. This part treats that channel, completing the self-scale triad — Part II's sensing, this part's executing, and the learning of Part VI that mediates them. **[IP]**

3.1 The Self's Actuation Chain

A governance directive formulated at the center reaches the world only by descending a delegation chain — ministry, agency, region, municipality, street — and Paper XI's result is that each layer does three things to the directive regardless of anyone's competence or good faith. It *delays* it, and delays in series accumulate. It *projects* it onto the layer's own operational repertoire, so that what passes downward is the layer's translation rather than the directive itself. And it *perturbs* it, because each translation occurs under different local conditions. Delay, projection, noise: structural, not behavioral, and present even when every actor is sincere.

An intention descends an analogous chain inside a person. A value or aim must become a specific goal; the goal a plan; the plan must recruit motivation; motivation must be carried by habit and supported by environment; and all of this must issue in physical execution, and then in execution repeated. The layers are real because each does its three things. The chain *delays*: the resolve formed on Sunday has decayed by Wednesday, and the lags of the intervening steps add. It *projects*: a rich intention — to be healthier, more present, more honest — is collapsed at each layer onto the narrower repertoire that layer actually possesses, onto the one plan that can be formulated, the habits the environment will support, what the body will in fact do at six in the morning, so that what executes is not the intention but its compression onto available routine. And it *perturbs*: the same directive descends differently on different days, because the internal conditions translating it — mood, energy, the press of competing demands — are never twice the same. The directive arrives late, narrowed, and distorted, and this is the ordinary fate of intention, not a special pathology of the weak-willed. **[IP]**

As in Paper XI, the result is stated with the internal adversary set to zero. Ambivalence, self-sabotage, the part of a person that does not want the change — these are real, they are the self-scale analogue of the strategic attenuation Paper IX studied, and they make everything worse. But the attenuation of this part binds even at zero ambivalence: it degrades the intentions a person wholeheartedly endorses exactly as it degrades the ones they are divided about. It is the failure that survives the best case, which is why, like Oakland, the best case is where the analysis begins. **[IP]**

3.2 The Energy Law and Personal Reform Exhaustion

Paper XI's central result is not a threshold but a price. Using the controllability Gramian of the actuation chain, it shows that the minimum control effort required to realize a target grows *superlinearly* with delegation depth. Deep chains do not refuse a policy; they price it out. The reform passes, partially delivers, consumes the political capital that powered it, and stops — a condition the paper names reform exhaustion.

The self-scale reading is immediate and rated **[IP]**: the self-regulatory effort — the attention, the willpower in the ordinary sense, the finite capital a person can spend pushing a directive down their own chain — required to realize an intention grows superlinearly with the depth of the chain between intending and acting. An intention that must traverse many internal layers (form the goal, build the plan, summon the motivation, establish the habit, arrange the environment, execute, and sustain) is priced out long before it becomes impossible. This is *personal reform exhaustion*: the self-improvement project launched with real resolve, delivered in part, that consumes the person's regulatory capital and halts — the structure of the abandoned resolution, recognizable to everyone and mysterious to the person living it, who experiences the halt as a verdict on their character.

Pressman and Wildavsky's arithmetic makes the mechanism concrete at this scale. A standing intention — "I will do this daily" — requires a yes at every link, every day: wake willing, *and* have the time, *and* the energy, *and* an environment that cooperates, *and* no competing demand that wins the slot. If each link returns yes with high probability, the joint probability still falls fast with the number of links in series; a chain of ten reasonably reliable links delivers the behavior on a little over half of days, and a longer or less reliable chain collapses below usefulness. The intention does not fail because the person is weak. It fails because its architecture demands too many sequential yeses, and the per-day success of a long chain is the product of its links whatever the person's sincerity. The reframing is the same one Paper XI performed on implementation research: what looks like a failure of will is a property of chain depth. **[IP]**, with the colloquial reading **[H]**.

The Gramian analysis of Appendix B sharpens "superlinear" and adds a qualification the bare word omits. The effort to realize an intention grows superlinearly in depth in general, and *geometrically* — exponentially — in the gain- or time-limited regime; and the rate depends on the *horizon*, the window available before the intention decays or the opportunity passes. Given a generous horizon the growth is gentler; under time pressure it is far steeper, and (as §3.3 shows) past a point the target is not merely costly but unreachable. The same intention is thus a different problem depending on the time allowed it: a deep change attempted "today" and the identical change undertaken "over the coming months" are not the same task pursued with more or less willpower, but two different reachability problems, one of which may have no solution at any effort while the other is routine.

3.3 Constitutional Self-Uncontrollability

The binary limit of the same curve is the conceptual capstone, because it completes a duality the series has been building toward at the self scale. Paper XI's threshold reading is *constitutional uncontrollability*: beyond a critical delegation depth, the target leaves the reachable set entirely and the effort required becomes infinite — explicitly the dual of Paper III's constitutional unobservability, where beyond a critical chain depth a signal cannot be perceived at all.

Self I established the unobservability half at the personal scale: some dimensions of the self cannot be perceived regardless of sincerity or effort, because the value architecture lacks the axes along which they are defined. This part establishes the other half. Some intentions cannot be actuated regardless of effort, because they sit beyond the reachable depth of a person's actuation chain — the required regulatory effort is not large but unbounded, and no quantity of willpower closes the gap. Together the two results bound self-governance on both of its channels, exactly as Papers III and XI bound governance on both of its channels: there are states of the self a person cannot see, and there are states of the self a person cannot reach, and in each case the limit is constitutional — a property of architecture and depth rather than of trying harder. **[IP]**

This matters practically because the two failure modes feel identical from inside and demand opposite responses. An intention that will not actuate because its chain is merely deep is exhausting but reachable: it yields to the architectural remedies of §3.6. An intention beyond the reachable set will not yield to any of them, and the regulatory effort poured after it is spent against an infinite price. The framework does not, from the structure alone, tell a person which they face — but it

tells them the distinction is real, and that the experience of repeated, sincere, total failure is sometimes evidence not of insufficient will but of a target outside the reachable set, to be reached, if at all, by re-architecting the chain rather than by spending more of a capital that the depth has already priced out.

3.4 Delegation Across Time: The Chain of Future Selves

Here observer–plant identity enters, and in this part it does neither what it did in Parts II, IV, and VI — sharpen the institutional result into a new one — nor what it did in Part V — blur it. It *enriches* it, adding a dimension the institutional chain has only weakly, while the core mapping stands on its own scale-invariant legs.

A governance chain is spatial: the directive passes between distinct bodies arranged from center to periphery. The self's chain is partly spatial too — it routes through the body, the environment, and sometimes other people, all genuinely distinct subsystems that delay, project, and perturb. But its deepest layer is *temporal*. A standing intention is a directive from the present self to a succession of *future* selves, each of which must re-receive it under its own conditions and translate it onto its own momentary repertoire. The "consistency" that Paper XI's chain ends in is, for a person, a delegation across time to agents who are the same person and not quite — later time-slices with different energy, mood, and context, each one a partially different controller re-interpreting the directive it inherits. The structure is a delegation chain in the strict sense; the reflexive feature is that the principal and every agent are instances of one self, arranged not in space but in succession. **[IP]** for the structure; **[H]** for the framing.

This is why a long temporal chain is so much harder than its length alone suggests, and the reason connects this part to Part IV. The present self cannot *compel* a future self; it can only issue a directive that the future self may or may not honor when its moment arrives. Whether the future self complies is exactly the question of self-legitimacy — whether a person's own directives are followed by themselves — which the next section develops as the gain on this entire channel.

3.5 The Coupling With Self-Legitimacy

Part IV will show that self-trust acts as a multiplicative gain on self-action: effective actuation is $\mathbf{B}_{\text{eff}} = L_{\text{self}} \mathbf{B}$, so that a directive is discounted in advance to the degree the person does not expect themselves to honor it. Composed with the present part, this yields a compounding the parent series anticipated in its *bidirectional node*, where two compressions meet in one institution and multiply. Each link in a temporal actuation chain is a point at which a future self must choose to comply, and at each such point the legitimacy gain L_{self} multiplies in. A deep chain therefore has more links across which low self-trust attenuates the directive, so the same deficit in self-legitimacy degrades a long chain far more than a short one — depth and distrust do not add but multiply. **[IP]**

The corollary is the one that makes the design principle of §3.6 more than mechanical. Short chains are not merely cheaper to actuate; they are how self-legitimacy is *built*. A small commitment is a short, high-probability chain that actually executes, and each execution is a kept promise to oneself that raises L_{self} on the delivery side — which, in turn, is what eventually makes a longer chain affordable. The architecture of actuation and the dynamics of self-trust are the same system seen from two sides, and the practical path runs through their intersection: keep the chain short enough that the directive arrives, so that its arrival earns the trust that lengthens the chain you can next attempt.

3.6 Design Principle: Shorten and Externalize the Chain

The energy law makes the remedy structural rather than motivational, and this is its chief practical consequence. Because effort grows superlinearly with depth, the leverage is in the depth, not in the effort: the same regulatory capital that cannot push an intention down a ten-layer chain actuates a two-layer one easily, so the productive move is to shorten the chain rather than to spend more against its length. Stated as architecture, two principles follow. **[IP]**

Shorten the chain. Reduce the number of layers a directive must cross before it issues in action — the design logic behind making an intended behavior small and immediate enough that it executes in nearly one step rather than descending through goal, plan, and summoned motivation. A directive that actuates directly is not a weaker version of the intention; it is the same intention given a chain it can survive.

Externalize the internal layers. Where a layer of the chain is an unreliable internal subsystem — momentary motivation, in-the-instant willpower — replace it, where possible, with external structure that does not depend on the reflexive layer's current state: arranging the environment so the action requires no summoned motivation, so that the layer most subject to daily perturbation is removed from the chain and its function discharged by a standing feature of the world. This is the self approximating, on the actuation side, the same manufactured insulation that Part VI found on the learning side: structure built to carry what an internal layer cannot be relied on to carry.

The principles are modest and the analysis behind them is not, which is the point. What presents itself to a person as a deficiency of character — the intention sincerely formed and repeatedly unrealized — is, in a large fraction of cases, a directive issued down a chain too deep to carry it, attenuated by latency, projection, and noise, priced out by a superlinear cost, and discounted further at every link by whatever self-trust the person has left. The directive arrives when the chain is built to deliver it. That is an engineering statement, and it is the one this part has to make.

Part IV — Self-Legitimacy: Trust as the Gain on Self-Governance

Consider two people who want the same thing and know the same things. Each has decided to change something about their life — to exercise, to leave a job, to repair a relationship, to write the book. Each has an accurate model of their situation, a sound plan, and no shortage of information about what to do. The first acts, falters, corrects, and over months becomes someone who follows through. The second issues the same instruction to themselves, does not move, issues it again, does not move, and gradually stops issuing it at all.

The difference is not in the plan. It is not in intelligence, willpower understood as a quantity, or the quality of self-knowledge — both people perceive themselves accurately, which is precisely why Self I's account is insufficient here. The difference lies in a parameter that the plan does not contain: the degree to which a person's own directives are followed by themselves, and their own reports are believed by themselves. Paper XIII names this parameter, at the institutional scale, *legitimacy*. At the personal scale we will call it **self-legitimacy**, and the central claim of this part is that it is not a mood or a virtue but a gain — a multiplier on the effectiveness of every other element of self-governance. [IP]

4.1 The Coupling

Paper XIII models legitimacy as an endogenous coupling state $L(t) \in [0, 1]$: not a quantity the designer sets, but one that emerges from the interaction between an architecture and the governed, and that then modulates both of the architecture's channels at once. Actuation effectiveness becomes $\mathbf{B}_{\text{eff}} = L\mathbf{B}$; observation noise covariance becomes $\mathbf{V} = \mathbf{V}_0/L$. A single variable couples the two channels the series had previously treated apart. When L is high, the controller enjoys full authority and clean sensing. When L collapses, the same architecture becomes both unsteerable and blind.

The self instantiates the same structure. Define $L_{\text{self}} \in [0, 1]$ along two coupled channels:

- **Actuation-legitimacy** — the probability that a directive the self issues to itself is actually executed. This is self-trust in one's own intentions: the well-founded expectation that "I will do this" predicts doing it. Where it is high, an intention commits the system. Where it is low, an intention is discounted at the moment of forming — the self does not fully mobilize behind a commitment it does not expect itself to keep, so $\mathbf{B}_{\text{eff}}$ falls and the intention is self-defeating before execution begins.
- **Observation-legitimacy** — the probability that a self-report is treated as accurate rather than as something to be argued with, discounted, or overridden. This is self-trust in one's own perceptions: the expectation that "I feel X" or "this is happening to me" is a signal worth acting on. Where it is low, the self-observation channel is treated as untrustworthy, introspective noise rises ($\mathbf{V}_0/L$), and the person loses the ability to read their own state even when the reading is correct.

These are not two separate problems. They are one variable observed through two channels, and they move together. This is what makes self-legitimacy a *coupling* state rather than a parameter: like institutional legitimacy, it is not chosen but accumulated, and once it exists it feeds back on the effectiveness of everything else. The structural claim — that a single trust variable multiplies both self-action and self-perception — transfers from Paper XIII by scale-invariance and is rated [IP]. The specific psychological readings attached to it below are weaker and are tagged where they appear.

4.2 Built and Borrowed Self-Trust

Paper XIII's sharpest distinction is between *built* and *borrowed* legitimacy, and it survives the mapping intact.

Built self-trust accumulates slowly, from consistent delivery on commitments the self has made to itself — especially small ones — over extended time. It is resilient to individual failures because it rests on a long record rather than a single proof. The person who has kept a hundred small promises to themselves has a deep reserve; one missed morning does not overturn it.

Borrowed self-trust is acquired quickly from something other than a delivery record: a motivational high, a single dramatic success, external validation, or a self-narrative held in advance of the evidence ("I'm the kind of person who finishes things"). It can be summoned fast, which is why it is the currency of resolutions and fresh starts. But it is structurally brittle, because the betrayal sensitivity γ is far larger than the delivery sensitivity α : a single self-betrayal removes more trust than a single kept commitment adds. **[H]** A self running on borrowed trust can appear stable for a long time and then disintegrate at the first real test — the resolution that was load-bearing until the first missed day, the identity built entirely on one achievement that does not survive its first contradiction. The collapse is not a character failure. It is the call on an architectural debt: trust was spent that had never been earned.

The asymmetry $\gamma \gg \alpha$ is the formal core of a familiar and otherwise puzzling fact — that it is so much easier to lose faith in oneself than to rebuild it. We rate the asymmetry **[H]**: it is plausible, it mirrors a measured asymmetry at the institutional scale, and it rhymes with loss-aversion findings in individual decision-making, but the framework does not derive its magnitude at the personal scale and should not pretend to.

4.3 Three Failure Modes at the Personal Scale

The three structural failure modes of Paper XIII reappear, each with a recognizable personal form.

The self-betrayal spiral is the performance–legitimacy spiral run inward. A broken commitment lowers L_{self} , which weakens actuation (the next intention fires with less force, because it is discounted in advance) and degrades observation (the person reads their own state less accurately), which produces further broken commitments, which lowers L_{self} again. It is self-reinforcing and, like its institutional twin, it does not require any external adversary to sustain it. The loop closes inside one person. **[IP]**

The dynamical model of Appendix C shows the spiral is more than a basin one may fall into. Whether a *sustainable* self-trust equilibrium exists at all depends jointly on a person's delivery competence and the betrayal asymmetry γ/α . Where competence is high the healthy equilibrium persists even under a steep asymmetry, though the basin of collapse around it widens. But past a critical ratio — insufficient competence relative to how much a betrayal costs — the healthy equilibrium is *annihilated*, and collapse becomes the only attractor the dynamics admit. In that regime the spiral is not a risk but the sole outcome, and beginning with high self-trust does not help, because there is no stable high-trust state to remain in. This is the structural extremity the next sections must be read against, and it is why the boundary drawn in §4.5 matters: a collapse-only trajectory is also exactly what non-structural conditions can produce.

The transparency trap is the most consequential, because it is the spiral's most natural escape and the one that ends worst. A person whose self-trust is falling can arrest the *felt* decline not by delivering, but by manipulating their own observation channel: motivated reasoning, a curated self-narrative, the quiet editing of what gets admitted to awareness. This restores apparent L_{self} in the short term while a hidden discrepancy accumulates between the maintained story and the unedited state. The discrepancy does not disappear; it waits. Its eventual revelation — often as the crisis that finally forces honesty — triggers a betrayal penalty larger than the steady cost that continuous self-honesty would have charged. The person who could not afford to see the gap pays, later and at interest, for not having seen it. **[IP]**

Legitimacy hysteresis is the asymmetry between the speed of decline and the speed of recovery. Self-trust is destroyed quickly and rebuilt slowly: returning to a prior level of self-trust requires sustaining better performance for longer than the path that caused the decline ever demanded. The practical consequence is that recovery cannot be rushed, and — as 4.6 notes — the attempt to rush it tends to register as one more broken promise. [IP]

4.4 The Reflexive Sharpening

Here the structural feature introduced in Part I — *observer–plant identity* — makes the self-case strictly worse than the institutional one, rather than merely smaller.

At the institutional scale, legitimacy couples two distinct entities: the controller and the governed population. A government can, at least in principle, observe its own legitimacy through *independent* instruments — Paper XIII's "legitimacy sensors," diversified and external to the governing apparatus. The sensor is not the thing being measured.

In a single mind, it is. The apparatus that is losing trust is the same apparatus that measures the trust. There is no fully independent legitimacy sensor inside one self. This sharpens the transparency trap into something closer to a structural impossibility of detection from the inside: the self-deceiving instrument is also the instrument tasked with checking for self-deception, and a corrupted observation channel cannot reliably audit its own corruption. The person whose self-reporting has begun to drift cannot, in general, perceive the drift using the faculty that is drifting. [IP]

This gives a precise structural reason for something Part II only suggested. The reason external observers matter so much to self-knowledge — a friend who will say the uncomfortable thing, a therapist, a journal re-read months later when the editing self has moved on — is that they are the only available approximation to an *independent* legitimacy sensor. They are partly decorrelated from the apparatus under audit. The friend who only tells you what you already believe is not a second sensor; they are your own observation channel wearing another face, and they provide no audit at all. Self-legitimacy, in other words, cannot be fully maintained from inside the self, and the design implication — borrow an external sensor before you need it — follows directly.

4.5 The Boundary the Framework Must Not Cross

The language of this part — "I don't trust myself," "I don't believe myself anymore," the spiral of falling self-trust degrading both action and perception — is also, word for word, the language of depression, trauma, and several clinical conditions. This coincidence is the most dangerous feature of the entire mapping, and the framework's honesty depends on handling it explicitly. [R]

The control-theoretic account describes the *structure* of the self-legitimacy trap: how a coupling variable, once it begins to fall, can degrade the two channels that would be needed to arrest its fall. It does not describe the *aetiology* of any particular instance. A low L_{self} can arise from the structural dynamics named here — a long accumulation of broken self-commitments under a borrowed-trust architecture — or it can arise from substrates the model does not represent at all: neurochemistry, the physiology of trauma, grief, illness, the side effects of a medication. The framework cannot distinguish these from structure alone, because the structural signature is the same in every case. The same falling curve can have entirely different generators.

Two consequences follow, and both are load-bearing. First, the framework must never be read as claiming that the structural account is *sufficient* — that a self-trust collapse is "really" just a long actuation-legitimacy failure and can be reasoned out of. Where a clinical substrate is present, the architectural description is true and useless on its own; it describes the shape of the trap without touching its cause. Second, nothing in this part is a substitute for clinical care, and

the design principles in 4.6 are conditions that may help a structurally generated decline, not a treatment for a pathologically generated one. The framework can describe the trap. It is not a clinician, and it does not know, from the structure alone, which trap a given person is in.

4.6 Recovery as an Architectural Condition

The design principles of Paper XIII translate into conditions for rebuilding self-legitimacy — stated, in keeping with the series' discipline, as architecture rather than as advice, and subject in every case to the boundary of 4.5.

Transparency to oneself maintains the observation channel on which long-term L_{self} depends. Ceasing to edit the self-report — letting the gap between story and state be seen — is costly in the moment and is the only move that stops the transparency trap from compounding. *Delivery–reality matching* attacks the spiral on the actuation side: commitments made small enough that delivery is near-certain rebuild trust from the α side, one kept promise at a time, which is the only way built trust is ever accumulated. *Credible commitment devices* — external structure, precommitment, arrangements that make following through not depend on the momentary level of self-trust — function as the personal analogue of circuit-breakers and constitutional entrenchment: they let action continue while L_{self} is too low to carry it unaided. And *hysteresis-awareness* is itself a condition of recovery: because trust rebuilds more slowly than it fell, the expectation of fast return is not optimism but a setup, since the unmet expectation reads to the system as one more broken commitment and feeds the spiral it was meant to escape.

These conditions share a structure, and it is worth naming because it inverts a common intuition. The perfectionist — who sets large commitments, fails them, and treats the failure as evidence about their worth — is not aiming too high. They are running the self-betrayal spiral at maximum gain: maximal commitments maximize the betrayal term γ , and the architecture guarantees the collapse it then takes as confirmation. Built self-trust is recovered, when it is recovered, by the opposite move — promises small enough to keep, kept long enough to count.

Part V — Boundary Calibration: Where the Framework Reaches Its Limit

This is the part of the paper where the mapping is weakest, and the paper is more useful for saying so than it would be for forcing a clean result. The preceding parts borrowed primitives whose mathematics is scale-invariant — the same equations describing a ministry and a mind — and observer–plant identity sharpened each into a specific new result. Boundary selection does not behave this way. Its conceptual content transfers; its formal content does not; and the structural feature that sharpened the other parts blurs this one. The honest course is to take the transferable content as far as it goes, mark exactly where rigor stops, and let the asymmetry stand as evidence that the ranking of mappings in this series is a real claim and not a rhetorical hedge.

The territory is the self–other boundary: where the self ends and another begins, what is mine to carry and what is not, which feelings, problems, and responsibilities the person treats as their own. Its pathologies are familiar — the person who absorbs everyone's distress as if it were theirs, and the person walled so completely that nothing reaches them. Paper XII gives part of this real structure, and then a premise fails.

5.1 What Transfers: The Pooling Paradox

Paper XII's central result is a paradox of boundary placement. A controller can expand its boundary to internalize the spillovers that were destabilizing it from outside — but internalizing them lengthens its observation and actuation chains, so that the larger system is governed with lower fidelity through the same limited apparatus. Or it can shrink its boundary to preserve internal fidelity — but then the cross-boundary couplings it excluded continue to operate, ungoverned, returning as disturbances it cannot model or attribute. There is no single boundary that both internalizes the relevant couplings and preserves internal governance fidelity. Paper XII names this the Information-Actuation Frontier, and it is the part of the boundary analysis that ports to the self with its force intact. **[IP]**

The over-identifying self expands its boundary to take in others' states: their moods, their problems, their wellbeing are admitted as the person's own to perceive and to regulate. This does internalize a real coupling — the other's condition genuinely does affect the person, and a boundary that ignored it would leave that influence ungoverned. But it lengthens the self's chains exactly as Paper XII predicts. The person is now attempting to observe and to act upon a system much larger than themselves — another life, often several — through the single limited apparatus of one mind, and the fidelity of their self-governance degrades under the load. They lose track of their own state because the apparatus that would track it is occupied modeling everyone else's. The under-identifying self draws the opposite boundary: rigid detachment, the disavowal of couplings that are nonetheless real, internal fidelity preserved at the cost of leaving genuine cross-boundary influence unmodeled until it returns as a disturbance the person cannot source. Neither boundary escapes the frontier. The recurring intuition that one must care for reality without attempting to carry all of it is, in this light, not a counsel of moderation but a description of the frontier itself: there is no boundary placement that dissolves the trade-off, only placements that sit at different points along it. **[IP]**, with the psychological readings **[H]**.

A second piece transfers more weakly but genuinely. *Boundary brittleness* — Paper XII's failure mode in which a rigid boundary suppresses small disturbances until they accumulate into catastrophic ones — has a recognizable personal form in the self so impermeable that nothing is admitted, no small adjustment is ever made, and the maintained composure ruptures all at once rather than flexing early. The mapping is sound at the level of description and I rate it **[H]**, because unlike the pooling paradox it rests on the metaphor rather than on a result that survives the move.

5.2 Where the Formalism Breaks

Paper XII's conceptual content rests on a formal apparatus, and that apparatus does not reach the self. Stating why is the substance of this part. [R]

The paper models the gap between the system a controller governs and the system that actually exists as an *M- Δ feedback interconnection*: M is the controller and its modeled plant, Δ is the unmodeled cross-boundary dynamics, and the two are coupled in a loop. The stability result is the small-gain theorem: if the loop gain around the unmodeled dynamics stays below unity, the interconnection is stable *for every* Δ within the gain bound. The power of the theorem — the reason it is rigorous rather than merely suggestive — is precisely this universal quantifier. The controller does not need to know the internal structure of Δ . It needs only a bound on Δ 's gain, and stability is then guaranteed against any perturbation the bound admits, including the worst case the bound allows.

This guarantee requires that Δ be a *bounded but otherwise arbitrary* perturbation: passive in the sense that its response is not a function of the controller's own model and strategy. Relational boundaries violate this requirement at its root. The dynamics on the far side of a self–other boundary are not passive perturbations; they are *other agents* — other selves, each with their own observer–plant identity, their own value architecture, their own legitimacy dynamics, and, decisively, their own capacity to model the person and to adapt their behavior to where that person draws their boundary. The other probes the boundary, contests it, accommodates or exploits it. Δ , here, is correlated with M because Δ is watching M and responding to it, and Δ is not confined to any gain bound the controller might assume, because a strategic counterparty can act precisely to violate the controller's assumptions. The "for every Δ in the ball" quantifier that makes the small-gain theorem work no longer describes the situation. The coupling is game-theoretic, not merely dynamical, and a robust-control stability condition derived for passive uncertainty does not transfer to it.

This is not a novel failure of the small-gain theorem; it is a well-mapped edge of robust control, which has always distinguished bounded uncertainty from strategic, adapting opposition. The point is only that the self-boundary case falls on the wrong side of that known edge. The consequence for this paper is exact: the *conceptual* lesson of Paper XII — that loops around unmodeled dynamics can destabilize a system every component of which is internally sound — survives the move to the self and underwrites §5.1. The *theorem* — the quantitative condition under which the boundary is stable, and the boundary mismatch index B that the theorem makes meaningful — does not. Where Part II's mapping was an [R] result applied (the design effect is the design effect, whatever the ensemble), Part V's is a conceptual transfer with the rigorous core left behind. The downgrade is the honest content of the part.

5.3 How Observer–Plant Identity Blurs Rather Than Sharpens

There is a further reason boundary is the weakest leg, and it is instructive because it runs opposite to the rest of the paper. In Parts II, IV, and VI, observer–plant identity *sharpened* the institutional primitive into a specific new result — a floor under correlation, an unauditable trust, a two-sided revision bound. Applied to boundary, the same feature *blurs* the primitive instead.

In governance the boundary is a clean interface between the controller and systems external to it. In a self there is no purely external interface, because the controller is inside the plant. The result is that "boundary," for a self, does double duty across two questions that institutional governance keeps separate. One is interpersonal: where does this self end and another begin? The other is intrapersonal: which contents of my own experience do I identify with as me, and which do I hold at a distance as something happening to me but not constitutive of me — Self I's reading of the self-as-content versus self-as-context distinction. These are different boundaries, and in a reflexive system they will not cleanly separate; how a person draws the line around their own experience conditions how they draw the line around others, and the reverse. Paper

XII's single boundary-mismatch index B therefore has no clean self-analogue. There is not one boundary to mis-place but at least two, entangled, and the framework does not currently resolve which dynamics belong to which. This blurring is itself a fourth consequence of observer–plant identity — the negative one, the case where reflexivity makes an institutional concept harder to port rather than easier — and naming it is part of why the ranking of mappings is not arbitrary. **[IP]**

5.4 What Survives as Design Principle

Enough transfers to state design principles, with the explicit caveat that these are the least-supported principles in the paper — conceptual transfers without the rigorous backing the other parts' principles inherit.

The boundary is a design variable, not a given. It is drawn, drawn well or badly, and most pathologies of relational boundary are the default placements: maximally permeable by temperament, or maximally rigid by injury, rather than matched to the actual coupling structure of a particular relationship. The principle that survives Paper XII is *match the boundary to the coupling* — admit across the boundary the influences that are genuinely operative and exclude those that are not — together with *boundary humility*, holding any placement provisionally, because the coupling structure of a relationship changes and a boundary that fit it once will not fit it indefinitely. These are sound as far as they go and I rate them **[IP]**; they go less far than their analogues elsewhere in the paper, because the frontier they navigate cannot here be quantified.

The boundary is also where this part touches the strong ones. Who a person admits as an observer of themselves is a boundary decision, so the curated confidant of Part II — the channel selected for agreement — is a boundary drawn for comfort rather than for coupling. A boundary drawn too permeably sources self-trust externally, making the self-legitimacy of Part IV hostage to others' responses — borrowed trust by another route — while a boundary drawn too rigidly cuts off the external legitimacy sensors that Parts II and IV identified as the self's only approximation to an independent check. And the protected experimental space of Part VI is, structurally, a boundary drawn around a sandbox: boundary-setting is the operation by which a self manufactures the insulation that its adaptive learning requires. Boundary is the weakest leg formally and among the most entangled conceptually, which is the characteristic position of a variable the parent series treated as context rather than as a stage in its own right.

The part is included for the reason the ranking demanded it. A claim that some mappings are rigorous and others merely conceptual is only a real claim if the merely-conceptual ones are shown as such, in the place where the rigor runs out, rather than dressed to match the rest. Boundary is that place. The framework reaches here, delivers the pooling paradox and a provisional design discipline, and stops at the point where the system on the other side of the line starts modeling the controller back.

Part VI — Adaptive Learning: The Self That Must Revise Itself Without Coming Apart

Two lives can fail in opposite-looking ways. The first person settles early — a career, a role, a self-understanding — and optimizes it faithfully for decades, until they wake at fifty inhabiting a life that fits a person they no longer are, unable to say what they would want instead. The second person never settles at all: a new city, a new framework, a new relationship, a new self every eighteen months, accumulating beginnings and integrating none of them. The first looks like excessive stability and the second like excessive change, and the temptation is to prescribe the missing virtue to each. But the parent series' fourteenth paper shows they share a single diagnosis. Both have mis-set the same parameter — the rate at which a controller probes beyond its current model — one toward zero and one toward excess. This part treats the self as an adaptive controller subject to that parameter, and argues that for a self, unlike for a society, the parameter is bounded on *both* sides, and the upper bound is the tighter one. [IP]

This part also completes the self-scale triad. Part II established the *sensing* requirement — an observing ensemble decorrelated enough to detect its own systematic error. Part III established the *actuation* requirement — a delegation chain short enough to carry intention into act. Adaptive learning is the requirement that mediates between them: can the self discover *what* to change, and do so without the discovery destabilizing the self that must integrate it? The sequence — **Sense** → **Learn** → **Execute** — is the parent series' account of how any controller stays viable in an environment that changes faster than its architecture, and it is no less the structure of a life.

6.1 The Self as a Dual Controller

Dual control theory models the tension between *exploitation* — acting optimally given what is currently known — and *exploration* — acting to acquire knowledge that improves future action. Its central observation is that in an adaptive system the two cannot be cleanly separated, because every intervention is simultaneously an action and an experiment. At the scale of a life this is exact rather than metaphorical. Taking a job is both a choice and a probe: it produces an outcome *and* reveals something about what the person can do, what they can tolerate, and what they actually want, as opposed to what they believed they wanted. Entering or ending a relationship, moving, changing field, having a child — each is an action that also returns information about parameters the person cannot estimate any other way, because the only instrument that measures them is the living of the choice. [R] for the dual-control structure; [IP] for its application to a life.

The consequence follows directly. A person who treats their choices only as actions — selecting at each step whatever optimizes their current objective — is systematically discarding the self-knowledge those choices could have generated. The optimal policy, in dual control, includes an *exploration bonus*: actions are tilted, at the margin, toward those that reduce uncertainty about parameters consequential for the rest of the trajectory. For a self, the parameters most worth reducing uncertainty about are the ones Self I showed are hardest to see from inside a fixed value architecture: what one actually values, what one is capable of, who one would be under conditions never yet encountered. The exploration bonus is the structural reason a well-run life spends some of its choices on questions rather than on returns. [IP]

6.2 Persistence of Excitation and the Over-Protected Self

System identification supplies a condition with a hard edge: a system's parameters are identifiable only if the system is *persistently excited* — only if it experiences enough variation to reveal how it responds across the range of conditions that matter. A controller that holds its inputs constant learns nothing about its own dynamics, however long it runs. Paper XIV reads this as the rigorous content of antifragility: a governance system that suppresses all variance — every shock, every failure, every protest — is not maximally stable but maximally fragile, because it has eliminated the signal on which its own calibration depends.

The personal form is direct and somewhat against intuition. A life arranged to avoid all challenge, all failure, all discomfort, is not thereby made safe; it is made *uncalibrated*. The person who has never been near the edge of what they can handle does not possess a reassuring stability — they possess an unidentified model of themselves, a self-estimate untested against the conditions that would reveal its error. They do not know what they can bear, what they would sacrifice, who they become under load, because the parameters governing those responses were never excited. This is the structural reading of a familiar observation: that protection past a point produces fragility rather than security, and that the capacity to meet difficulty is itself learned only by metabolizing some. The framework does not romanticize hardship — it specifies a quantity of variance required for identifiability, below which the self-model drifts out of calibration unnoticed, and above which lies the danger of §6.4. **[IP]** for the structure; **[H]** for the life reading.

6.3 *Exploration Starvation and the Green Dashboard*

Of Paper XIV's five failure modes, the first is the one that matters most here, because of how it hides. *Exploration starvation* occurs when short-term incentives drive the exploration rate to zero: the system ceases to probe beyond its current model, the model drifts from the reality it is supposed to track, and — this is the lethal feature — the drift is invisible to the system's own monitoring. The simulation makes the signature precise: the controller's internal estimate of its own performance stays optimistic even as its true tracking error diverges. The dashboard reads green while the system fails.

This is the structural account of midlife stagnation, and it is sharper than the colloquial version. The starved self is not in visible distress; that is what makes it starvation rather than crisis. It has optimized one dimension so thoroughly, for so long, that it stopped generating the variance that would reveal the dimension has gone stale — and because exploration is the very faculty that has been switched off, nothing in the system's own monitoring can report its absence. The person feels fine, by the only measure they still run. The drift is between that measure and a reality it no longer tracks. This is Self I's variety gap given a temporal mechanism: not only does a narrow value architecture exclude dimensions of the present self, it stops updating, so that the architecture fitted to who the person was at thirty governs the life of who they are at fifty with undiminished confidence. The green dashboard is the same artifact the parent series diagnosed in failing institutions, running in one person across a life. **[IP]**, with the specific course **[H]**.

Two adjacent failure modes complete the picture. *Model lock-in* is the obsolete self-understanding actively defended: the self-narrative develops an immune system that treats evidence against it as a threat to be neutralized — "that isn't who I am" deployed not as report but as defense. This is Part II's narrative monoculture and Part IV's curated self-report seen from the learning side: the mechanism that keeps the model from updating. *Exploitation lock-in* is its complement and is easy to miss because it looks like success at the hard part. Here the self *has* learned — the person sees clearly, often through exactly the decorrelated channels of Part II, what is wrong and what would have to change — but the actuation chain of Part III is blocked, so knowledge accumulates while behavior does not move. The insight is genuine and the life does not change. This is the structural difference between not knowing what to do and knowing precisely while remaining unable to do it; the parent series' simulation shows the two are distinct trajectories, the second tracking truth accurately while performance stays poor. **[IP]**

6.4 *The Reflexive Danger: When the Experimenter Is the Experiment*

Everything to this point would also be said, suitably translated, of a society. The distinctive content of self-learning appears when observer–plant identity is applied a third time, and it cuts in a direction the earlier parts did not.

In governance, the controller modifies the *plant* — a tax reform perturbs the economy, not the ministry. The experimenting apparatus is insulated from the experiment: the state runs Special Economic Zones, and if a zone fails, the failure is contained there while the state that authorized it remains intact to learn from the result. This insulation is what makes aggressive exploration survivable; Paper XIV's "protected experimental spaces" are, structurally, devices for keeping the controller separate from the variance it injects.

The self has no such separation by default, because the controller *is* the plant. Self-exploration — revising one's values, interrogating one's identity, the whole activity of "working on oneself" — is exploration whose experimental subject is the apparatus conducting the experiment. The dual-control problem becomes recursive in the strict sense: the system is running probes on the system that is running the probes. And this converts persistence of excitation from a one-sided requirement into a two-sided bound. Let r be the rate at which a self revises itself and ι the rate at which a revised self re-stabilizes — integrates the change into a coherent ongoing identity. The *lower* bound is the one governance also faces: r must exceed the rate at which the person's circumstances and inner conditions change, or the self-model drifts into the starvation of §6.3. The *upper* bound has no full institutional analogue: r must not exceed ι . When revision outruns integration, the variance injected to drive learning destabilizes not merely the plant — which could recover — but the controller's own coherence, which is the thing that would have to remain stable to metabolize the learning at all.

$$\underbrace{\text{rate of change}}_{\text{lower bound: starvation below}} < r < \underbrace{\iota}_{\text{upper bound: incoherence above}}$$

This is Paper XIV's learning-induced oscillation at the personal scale, but it is more than chaotic life choices. In a self, over-aggressive exploration risks dissolving the experimenter. It is the structural account of why intensive self-transformation sometimes makes a person worse rather than better: revision arriving faster than it can be integrated, identity destabilized faster than it can re-cohere, the seeker of §6.1's opening who never stops changing precisely because nothing stabilizes long enough to be learned from. The upper bound is *tighter* for a self than for a society, and the reason is exactly observer–plant identity: a society can insulate its experimenting apparatus from its experimental zones, and a self, by default, cannot fully insulate itself from its own. **[IP]**

One qualification keeps this claim exact, and the simulation of Appendix D supplies it. A society is not in fact *unbounded* above: revise too fast and it amplifies noise into its own estimates, so it too has an upper bound on useful revision. The difference is not that institutions may revise without limit and selves may not. It is that the self's upper bound is *stricter* and arrives *earlier*, and through a different mechanism — it binds via loss of the controller's coherence well below the rate at which mere noise amplification would bind, forcing the self to stop short of even its own noise-limited optimum and to accept some tracking error as the price of staying coherent. "Tighter for a self than for a society" therefore means *earlier and via coherence*, not *present for the self and absent for the society* — a constraint a system whose controller and plant are separate simply does not meet first.

6.5 The Boundary the Framework Must Not Cross

The two-sided bound borders clinical and contemplative territory, and as in Part IV the border must be marked rather than crossed. **[R]** The framework describes the *structure* of the integration constraint — that revision outrunning re-stabilization degrades the coherence learning requires. It does not describe the management of that constraint in any individual, and it does not claim that destabilization under intensive self-work is always or even usually structural. Where a clinical substrate is present — where self-interrogation triggers something with a physiology the control model does not represent — the architectural description remains true and is insufficient on its own, in exactly the sense Part IV gave: it names the shape of

the difficulty without reaching its cause, and it is not a clinician. The observation that transformation must be paced to integration is a structural consequence of $r < \iota$, not a protocol for pacing it, and the framework holds no view on how, or whether, a given person should undertake such revision at all.

6.6 Protected Experimental Spaces as Manufactured Insulation

The design principle follows from the two-sided bound and reframes a familiar prescription into something with a structural job. *Protected experimental spaces* — hobbies, side projects, sabbaticals, travel, artistic and contemplative practice, low-stakes relationships, research undertaken for its own sake — are commonly defended as balance or recovery. The framework assigns them a more specific function: they are the self's attempt to *manufacture* the controller–plant separation that observer–plant identity otherwise denies it. A protected space is a sandbox in which a person can run an experiment on a partial self without staking the whole self on the outcome — exploring a different way of being, working, relating, or seeing, under conditions where failure is contained and does not propagate to the core identity that must remain coherent to learn. This is how a self approximates the insulation a society gets for free: not by exploring less, which starves it, and not by exploring with its whole identity at once, which destabilizes it, but by constructing bounded arenas where the exploration rate can run high locally while the integration constraint is respected globally. The autobiographical instance is available to the author: a multi-year research programme is itself such an arena — a protected space in which an objective function can be probed and revised without the probing being load-bearing for the rest of a life. The principle, stated as architecture, is that a self preserves its own adaptability by keeping some of its variance sandboxed: enough exploration to stay calibrated, walled well enough not to dissolve the thing being calibrated.

This completes the adaptation triad at the personal scale and returns the paper to the recursion Paper XIV identified as the convergence point of its mathematics. A governance system safely modifying its own architecture, an AI system safely modifying its own code, and a self safely revising its own identity are, structurally, one problem: a controller attempting to change what it is without destroying the continuity that would let it learn from the change. The self is the instance of that problem each reader inhabits, and the two-sided bound is the form the problem takes when the controller and the plant can never be fully pulled apart.

Part VII — Composite Failure Modes

The preceding parts each described how a single primitive fails on its own: an observer ensemble collapsing to one channel, an actuation chain too deep to carry a directive, self-trust spiralling downward, exploration starved or running past integration, a boundary mismatched to its coupling. In a life these failures rarely occur in isolation. The dangerous traps — the ones that take years and resist the obvious remedies — are *composites*, in which several primitives fail together and each failure drives the others. This part names the principal composites and the structure they share.

Two facts make composite failure the rule rather than the exception for a self. The first is the parent series' result from Paper V: coupled constraints do not add, they compound multiplicatively — the coordination-failure tax. The second is observer–plant identity. In governance the primitives are often instantiated in different institutions — the agency that senses is not the ministry that acts — so a failure in one is partly firewalled from the others. A self has no such firewall, because all five primitives run through a single substrate. The same apparatus observes, acts, trusts, explores, and bounds, so a degradation in any one propagates into the rest with nothing to stop it. Composite failure is not a special case of self-governance breakdown; given observer–plant identity, it is the default form breakdown takes. **[IP]**

7.1 The Sealed Self

The deepest composite couples three primitives into a mutually reinforcing loop from which the self cannot, by its own resources, perceive an exit. Observer diversity collapses (Part II): the channels converge until the ensemble delivers one correlated reading. The self-model locks in (Part VI): the narrative hardens into something the remaining channel only confirms. And the transparency trap closes (Part IV): the self-report is curated to defend the narrative, editing out what little dissonance survives. None of the three is independent of the others. A collapsed ensemble has no decorrelated channel left to deliver the signal that would break the lock-in; the locked-in narrative actively recruits confirming observers, driving the correlation higher; the transparency trap removes the residue. The result is a self that has become structurally unable to update *and* structurally unable to perceive that it cannot — the green dashboard of Part VI, the echo chamber of Part II, and the self-deception of Part IV, fused into one condition. It is the gravest trap in the paper because every internal instrument that might detect it has been enlisted in maintaining it, and the only purchase on it comes from the decorrelated external channels that the trap has, by construction, been closing. **[IP]**, with the portrait **[H]**.

7.2 The Self-Betrayal Spiral

The composite of self-legitimacy (Part IV) and actuation depth (Part III) was assembled in §3.5 and is named here as the trap it constitutes. Low self-trust attenuates a long actuation chain multiplicatively, because the legitimacy gain enters at every link; the attenuated chain fails to deliver; each failure is a broken promise to oneself that lowers self-trust further; and the lowered trust attenuates the next chain more steeply still. Depth and distrust do not add — they multiply, and the product feeds back on itself. This is the engine beneath the familiar experience of a person who has stopped being able to act on their own intentions and reads the paralysis as a verdict on their character, when it is the predictable output of two coupled primitives degrading in lockstep. **[IP]**

7.3 The Dissolution Trap

A boundary drawn too permeably (Part V) does not fail alone. The over-identified self routes its self-observation through another person or a group, so its channels lose the independence Part II required — the other's view of the person becomes the dominant reading, and the decorrelated internal channels go quiet. At the same time its self-legitimacy is sourced externally (Part IV): self-trust becomes borrowed, hostage to the other's responses, brittle in exactly the way borrowed

legitimacy is brittle. Observer diversity, self-trust, and boundary fail together, and the composite is the recognizable condition of losing oneself in a relationship or a group — not three separate problems but one, in which a permeable boundary simultaneously colonizes the observation channels and externalizes the trust, leaving a self that can no longer locate its own state or stand on its own commitments. **[IP]**, with the portrait **[H]**.

7.4 Revision Into Incoherence

The seeker who never settles (Part VI's upper-bound violation) is, examined closely, a composite with self-legitimacy. Revision arriving faster than integration destabilizes the controller, as Part VI established; but each revision that fails to take is also a commitment to a self that was abandoned, and abandoned self-commitments erode self-trust exactly as broken promises do (Part IV). The eroded trust makes the next revision land even less securely — a self that does not believe its own intentions cannot commit to the new identity firmly enough for it to cohere — which guarantees the next abandonment. The exploration upper bound and the legitimacy spiral drive each other, so that what looks like an excess of courage is the same self-betrayal dynamic of §7.2 running through the learning channel rather than the actuation one. **[IP]**

7.5 The Shared Structure

The composites differ in their contents and share a form: mutually reinforcing degradation across primitives that are multiplicatively coupled and, in a self, separated by no firewall. This has a direct consequence for the design principles of the next part. Because the primitives fail together, they cannot be repaired one at a time — an intervention on a single primitive, applied while the others continue to degrade, is working against a multiplicative coupling that will overwhelm it. The person trying to rebuild self-trust while their observation channels remain collapsed, or to force a deep actuation chain while their self-legitimacy is still falling, is repairing one factor of a product whose other factors are driving toward zero. The architecture of recovery, like the architecture of failure, is coupled — which is why the principles of Part VIII are stated as a set whose elements support one another, rather than as independent repairs that could be undertaken in any order or in isolation. **[IP]**

Part VIII — Design Principles

The principles that follow were each stated in the part that established them, and this part does not re-derive them. It does two things the home parts could not. It gathers the principles into the coupled set that Part VII showed they must form — because the primitives fail together and multiplicatively, the repairs cannot be undertaken one at a time. And it identifies the single logic beneath them, which turns out to be the design-side face of the paper's central variable. Throughout, the principles are stated as the series states them: as architecture, not as advice. Each has the form *a self with this property remains adaptive in this respect* — a structural condition, not an instruction for how to live, and not a claim that a given person should build such a self or could do so at a cost they would accept.

8.1 The Principles, Gathered

From Part II: maintain self-observation channels that are partially *decorrelated* — channels not sharing the narrating substrate and not selected for their agreement — and weight divergence from such a channel over agreement among correlated ones, since only decorrelated channels raise the ensemble's effective independence. From Part III: keep the actuation chain *short*, and *externalize* its least reliable internal layers, so that a directive crosses fewer steps and depends less on momentary internal states to reach behavior. From Part IV: practice *transparency to oneself* and *delivery–reality matching*, build *credible commitment devices* that do not depend on the current level of self-trust, and proceed with *hysteresis-awareness*, expecting recovery to run slower than decline. From Part V, the most weakly supported: *match the boundary to the coupling* and hold it with *humility*, renegotiating as the coupling changes. From Part VI: maintain *protected experimental spaces* that let exploration run high locally without staking the whole self, keep the exploration rate between its two bounds, and periodically *audit the objective function*, since a stale value architecture will not report its own staleness. [IP], except Part V's, which carry the [H] noted there.

8.2 The Unifying Logic: Recruit What the Reflexive Self Cannot Be for Itself

Set side by side, the principles reveal one logic. Observer–plant identity entails that the self cannot, from inside, decorrelate its own observers, audit its own trust, insulate its own experiments, or fully shorten what is a reflexive chain — the faculty that would do these things is the same faculty under examination. The design response to each of these incapacities is the same move: recruit, *outside* the reflexive substrate, the structure that the substrate cannot supply to itself. The decorrelated channel of Part II is an observer situated outside the narrating self. The commitment device of Part IV is a structure that holds when momentary self-trust does not. The externalized layer of Part III is an environmental scaffold standing in for an unreliable internal one. The protected space of Part VI is a sandbox built outside the core identity. These are not five unrelated prescriptions; they are five instances of one principle — *a self governs itself well by building, outside itself, the structures its reflexive nature prevents it from being for itself* — and the boundary calibration of Part V is the master operation of that principle, the act by which the self decides what external structure to admit and on what terms. The design side of observer–plant identity is externalization, for the same reason its diagnostic side was a family of internal incapacities. [IP]

One principle resists this logic, and it is the one Part I's double edge predicted would. Because in a reflexive system observing is already acting, sustained self-observation is not only a sensor but a lever: attending to a dimension of the self changes it. This is the single capacity the reflexive structure grants rather than denies, the internal counterpart to all the external recruitment — the disciplined attention of contemplative and therapeutic practice doing, from inside, work that the other principles must outsource. A complete design uses both edges: external structure for everything the self cannot do to itself, and the internal attention-lever for the one thing its reflexivity makes uniquely possible. The framework leans

heavily on the external side because that is where observer–plant identity creates need, but it would be incomplete if it forgot that the same identity also creates one form of leverage available nowhere else. **[IP]** for the structure; the contemplative reading is **[H]**, and Part IX governs how far it may be taken.

8.3 *The Keystone and the Sequence*

Because the principles are multiplicatively coupled, the set has an entry point, and identifying it is the one piece of sequencing the framework can offer. Self-legitimacy is the gain on the rest: Part III showed it multiplies into every link of the actuation chain, Part IV that it modulates both observation and action, Part VII that its collapse drives the gravest composites. An intervention on any other primitive is multiplied by this gain, so an intervention applied while self-trust is falling is working against a factor driving toward zero. The entry point is therefore wherever self-legitimacy is *built*, and the paper has already located that: in small commitments kept. A small kept commitment is simultaneously a short, high-probability actuation chain that actually delivers (Part III) and a delivery that raises self-trust on the delivery side (Part IV) — the one move that is both an action and a deposit into the gain that every other action draws on. This yields a counter-intuitive design claim. The productive entry to a coupled self-governance failure is not the worst-damaged primitive, where intervention is multiplied by the lowest gain, but the place where the gain itself is rebuilt — small enough commitments that they cannot fail, kept until the trust they generate makes the deeper repairs affordable. The order is not a matter of preference; it follows from where the multiplicative gain sits. **[IP]**

8.4 *The Fence*

These principles describe the architecture of a self that stays open — that keeps perceiving what it does not yet value, keeps acting on what it perceives, keeps trusting itself enough to commit, keeps exploring enough to stay calibrated, and keeps its boundary matched to the world it is coupled to. They do not describe a good life, and the distinction is the one the series has held throughout and Part 0 committed this paper to. To say that a self with these properties remains adaptive is a structural claim. To say that a person *should* build such a self, or that adaptiveness is what they should want, or that a life so organized is thereby well-lived, are claims of a different kind that the framework does not make and has no instrument to support. **[R]** as a statement of the framework's standing.

The principles keep the room open; they do not furnish it. That is not a shortcoming to be remedied by a more complete set of principles — Part IX gives the reason it cannot be. What the open self does with its openness, what it loves, makes, and becomes, is precisely the content that no architecture specifies and no objective function holds. The design principles build and maintain the room, with external structure for what the self cannot do for itself and the one internal lever for what it can. What enters the room is not theirs to say, and a framework that claimed to say it would have stopped being engineering and become the very thing it was built to diagnose.

Part IX — The Limit of the Framework

This part is the framework reaching the edge of its competence and naming the edge. It is not a conclusion the paper has earned; it is a boundary the paper has arrived at, and the discipline of Part 0 applies here most strictly of all. Everything in this section is **[H]** — bounded speculation — except the statements of limit themselves, which are **[R]** in the sense established throughout: a rigorous account of what the framework does not and structurally cannot claim. The optimistic words that appear below — joy, beauty, truth, love, meaning — appear only inside that fence, as instances of what falls outside the framework, never as conclusions it asserts. The section is short because restraint is the content. A framework that has spent eight parts insisting on the difference between architectural diagnosis and normative prescription cannot close by quietly abandoning it.

9.1 The Asymptote of the Adaptive Problem

Self I named an asymptote: $G_{\text{self}} \rightarrow 0$, the limit of full self-observability, a self that perceives itself without exclusion. The adaptive problem has its own asymptote, and it is instructive precisely because of what it turns out to be empty of. A self that perfectly satisfied every condition this paper has specified — observation channels fully decorrelated, the actuation chain free of attenuation, self-legitimacy complete, exploration optimally calibrated against integration, boundaries exactly matched to coupling — would be a self of maximal adaptive capacity: maximally able to perceive itself, act on what it perceives, trust its own perceptions and commitments, learn, and revise without coming apart.

And that limit is a condition, not a content. Maximal adaptive capacity is not flourishing; it is the capacity to adapt toward anything whatever. A perfectly adaptive self could orient that capacity toward ends trivial or terrible as readily as toward ends worth having, and nothing in the architecture selects among them. The framework carries a person to the threshold of a self that functions — that can see, act, trust, and grow — and stops there, because what a functioning self should see by, act toward, or grow into is not an engineering question and the framework has no instrument that reaches it. The asymptote of the adaptive problem is an open door. It is not a destination. **[R]** as a statement of limit; the surrounding gloss is **[H]**.

9.2 Why the Deepest Goods Fall Outside

That the framework cannot specify what adaptation is for might look like a gap a richer model could fill. The series' own founding result suggests it is not a gap but a structural feature. The Goodhart–Ashby synthesis holds that any objective function with dimensionality lower than the system it governs eventually optimizes away its own ability to perceive that system — and that what is not represented in the objective is, in the end, not perceived at all. Applied to the self, this is the engine of Self I. Applied here, it has a further consequence: the goods of deepest value may be exactly the ones that cannot be made into objective-function targets without being destroyed in the conversion. The happiness pursued as a metric recedes; the meaning made into a goal hollows; the love instrumentalized as an objective becomes something else. These are Goodhart effects on the interior of a life, and they suggest that the things most worth having exist only *outside* the optimization, in whatever space is left open when the objective function does not close over everything. **[IP]** for the structural claim; the identification of which goods these are is **[H]** and philosophical.

If that is right, then the framework's relationship to such goods can only ever be one of *making room*, never of specifying. It can describe the conditions under which a self remains open — adaptive capacity preserved, the objective function kept from hardening into a cage that optimizes away everything it does not name — and it can observe that joy, beauty, truth, and the rest remain *possible* in a self so maintained. It cannot derive them, define them, or prescribe them, and the reason is

not the framework's immaturity but their nature: they are, on this reading, precisely the part of a life that no objective function holds. The framework keeps the channel open. It has nothing to say about what should come through it, and it is the structure of the thing, not a failure of the model, that this is so. **[H]**, fenced.

9.3 *The Direction of Influence*

One correction guards this section against its most likely misreading. Nothing here was derived from control theory. The contemplative, humanistic, and clinical traditions that have studied self-observation, self-trust, and self-revision for centuries discovered what they discovered by other means entirely, and where this paper's structure agrees with them, the agreement runs from their findings to the formalism, not the reverse. The mathematics did not generate wisdom about how to live; at most it gave a particular structure to what those traditions already knew, and it earns no authority over questions of meaning by having formalized some of the mechanics that surround them. A reader who took the framework to have *proved* something about joy or truth would have inverted the actual direction of the arrow. The framework is downstream of the lived traditions here, not above them. **[H]**

9.4 *The Boundary*

The paper can therefore state its final structural claim and decline its final temptation in the same breath. The claim is the one that runs through every part: a self that perfectly optimizes a fixed objective eventually optimizes away its own capacity to perceive what the objective excludes, so that adaptive capacity must be actively preserved *against* the standing pull toward optimization — and the whole apparatus of this paper, the decorrelated channels and short chains and built trust and calibrated exploration, is in the end machinery for keeping a self open rather than letting it close. The temptation is to say what the openness is for. The framework declines it, not from modesty but from competence: what the preservation of adaptive capacity is *for* is exactly the thing the preservation cannot specify, the content that lives outside any objective function and is destroyed by being made into one. The framework builds and guards the room. What a person fills it with is not a result the framework can return, and a framework that claimed otherwise would have become the kind of low-dimensional objective it spent nine parts warning against. The boundary is the conclusion. There is nothing past it the engineering can reach.

Appendix A — The Self-Observation Ensemble and the Correlation Tax

This appendix supplies the formal backing for Part II. The headline results are standard statistics applied to a self-observation ensemble and are rated **[R]**; the modeling choices that connect them to a self (the error-generating structure of A.4, the substrate-block partition of A.5) are **[IP]**, and are flagged where they enter. All numerical values below are computation-verified; the script reproducing them is given in A.8. Closed-form claims marked (MC) were additionally confirmed by Monte Carlo to within sampling error.

A.1 The ensemble model

Let the latent self-state be $\mathbf{x}_{\text{self}} \in \mathbb{R}^n$. A person observes it through N channels — introspection, the inner narrative, interoception, a friend, a journal, and so on — each returning

$$\mathbf{y}_i = \mathbf{C}_i \mathbf{x}_{\text{self}} + \boldsymbol{\varepsilon}_i, \quad i = 1, \dots, N,$$

where $\mathbf{C}_i \in \mathbb{R}^{m_i \times n}$ is the channel's structural perspective and $\boldsymbol{\varepsilon}_i$ its error, with $\mathbb{E}[\boldsymbol{\varepsilon}_i] = \mathbf{0}$. Two properties of the ensemble govern its capacity: the rank of the stacked perspective (A.3) and the correlation structure of the errors (A.2). For the correlation analysis it suffices to consider estimation of a single scalar coordinate of $\mathbf{x}_{\text{self}}$ on which all channels report, with $\text{Var}(\varepsilon_i) = \sigma^2$ and pairwise correlation $\text{Corr}(\varepsilon_i, \varepsilon_j) = \rho$ for $i \neq j$.

A.2 The correlation tax

Pool the channels by the equal-weight estimator $\hat{x} = \frac{1}{N} \sum_i y_i$. Its error variance is

$$\text{Var}(\hat{x}) = \frac{1}{N^2} \left[\sum_i \text{Var}(\varepsilon_i) + \sum_{i \neq j} \text{Cov}(\varepsilon_i, \varepsilon_j) \right] = \frac{1}{N^2} [N\sigma^2 + N(N-1)\rho\sigma^2] = \sigma^2 \left(\frac{1-\rho}{N} + \rho \right).$$

Defining the *effective number of independent observers* as the count that would yield this variance under independence, $\text{Var}(\hat{x}) \equiv \sigma^2 / N_{\text{eff}}$, gives

$$N_{\text{eff}} = \frac{N}{1 + (N-1)\rho}$$

the Kish design effect. **[R]** (MC: empirical pooled variance matched the closed form to within 0.3% at $N \in \{6, 20\}$, $\rho \in \{0, 0.5, 0.97\}$.) Three regimes follow. At $\rho = 0$, $N_{\text{eff}} = N$ and numbers help fully. For fixed $\rho > 0$, taking $N \rightarrow \infty$ gives

$$\lim_{N \rightarrow \infty} N_{\text{eff}} = \frac{1}{\rho},$$

a hard ceiling independent of N : at $\rho = 0.5$ no quantity of channels exceeds $N_{\text{eff}} = 2$; at $\rho = 0.3$, N_{eff} saturates at **3.33** (reached to three significant figures by $N = 1000$). The shared bias sets a noise floor that observer count cannot lower.

A.3 Requisite observer diversity

Stack the perspectives as $\mathbf{C}_{\text{ens}} = [\mathbf{C}_1^\top \cdots \mathbf{C}_N^\top]^\top$. A coordinate of $\mathbf{x}_{\text{self}}$ is recoverable from the ensemble only if it lies in the row space of $\mathbf{C}_{\text{ens}}$; the directions of the self in the kernel of $\mathbf{C}_{\text{ens}}$ are unobservable to *every* channel at once. Writing $\dim(\mathbf{U}_{\text{self}})$ for the dimensionality of the self-uncertainty space (the directions in which the person cannot predict

their own state), requisite observer diversity is the covering condition

$$r_{\text{ens}} := \text{rank}(\mathbf{C}_{\text{ens}}) \geq \dim(\mathbf{U}_{\text{self}}). \quad ([\text{IP}])$$

This is the ensemble-level analogue of the single-channel observability condition of Self I: there the value architecture's matrix could be rank-deficient; here the *union* of all channels' perspectives can be, producing a blind spot no cross-referencing detects because no channel spans the missing direction. Rank and correlation are distinct deficiencies — an ensemble can be full-rank yet near-perfectly correlated (every channel sees every dimension, but with the same error) — and A.2 quantifies the second while A.3 states the first.

A.4 Two pathways and their compounding

Model each error as the sum of a shared-model component (common processing), a shared-data component (common inputs), and an idiosyncratic component, all zero-mean, unit-variance, mutually independent: **[IP]**

$$\varepsilon_i = \sqrt{\rho_{\text{model}}} m + \sqrt{(1 - \rho_{\text{model}})\rho_{\text{data}}} d + \sqrt{(1 - \rho_{\text{model}})(1 - \rho_{\text{data}})} u_i,$$

with m, d common to all channels and u_i idiosyncratic. Then $\text{Var}(\varepsilon_i) = 1$ and, for $i \neq j$,

$$\text{Corr}(\varepsilon_i, \varepsilon_j) = \rho_{\text{model}} + (1 - \rho_{\text{model}})\rho_{\text{data}} = 1 - (1 - \rho_{\text{model}})(1 - \rho_{\text{data}}). \quad ([\text{R}])$$

The algebra follows rigorously from the model. (*MC: empirical pairwise correlation matched $1 - (1 - \rho_m)(1 - \rho_d)$ exactly across tested pairs, e.g. $\rho_m = 0.2, \rho_d = 0.8 \Rightarrow 0.84$.*) The interpretation is that error is decorrelated overall only if it is idiosyncratic through *both* pathways; moderate correlation on each route compounds toward unity, and the two pathways are rarely independent in practice, since the narrative that filters incoming signals is the narrative that selects which signals are sought.

A.5 The reflexive floor and the weighting penalty

Observer–plant identity (Part I) implies that some channels are not merely correlated but share one substrate — introspection and the inner narrative are the same apparatus reporting twice — so no internal operation decorrelates them. Model this as a partition of the N nominal channels into B blocks, with within-block correlation $\rho_w \rightarrow 1$ and between-block correlation ≈ 0 : channels inside a block are the same effective observer; blocks are decorrelated. **[IP]** for the partition; the consequences are **[R]**.

Under the equal-weight estimator, a block model with sizes $k_1, \dots, k_B$ ($\sum k_b = N$) yields pooled variance $\frac{1}{N^2} \sum_b k_b^2$, hence

$$N_{\text{eff}}^{\text{equal}} = \frac{N^2}{\sum_b k_b^2} = \frac{1}{\sum_b (k_b/N)^2},$$

the *inverse Herfindahl index* of the block-size distribution. Optimal (block-aware) weighting instead recovers the block count, $N_{\text{eff}}^{\text{opt}} = B$, and these bracket the achievable diversity:

$$N_{\text{eff}}^{\text{equal}} \leq N_{\text{eff}}^{\text{opt}} = B \leq N. \quad ([\text{R}])$$

Two consequences refine the body. First, the floor: because substrate-sharing forces several nominal channels into one block, $B < N$ strictly whenever any internal channels are fused, and N_{eff} is bounded above by B regardless of how many nominal channels are consulted — the institutional remedy of constituting more observers cannot raise B when the new observers fall in existing blocks, and only genuinely external channels add blocks. Second, a refinement of §2.6 that

the simulation surfaced and that strengthens its claim: under equal weighting, N_{eff} falls *below* the block count whenever blocks are unequal, because the estimator over-weights the large fused block. For six channels comprising five fused internal channels and one external (sizes $[5, 1]$), $N_{\text{eff}}^{\text{equal}} = 36/26 \approx 1.385$ against $N_{\text{eff}}^{\text{opt}} = 2$: the single decorrelated channel that carries all the new information is very nearly outvoted by the correlated majority. It is therefore not sufficient to *possess* a decorrelated channel; under naive equal weighting the fused majority drowns it, and recovering its value requires deliberately over-weighting the dissonant external reading against the internal consensus. The discomfort of doing so is the felt form of the up-weighting the mathematics requires.

A.6 Selection dynamics

A heuristic model of the collapse of §2.5. Let channel i carry weight $w_i(t)$, and let weights drift up the comfort gradient — toward channels whose readings agree with the current ensemble mean — at rate η :

$$w_i(t+1) \propto w_i(t) \exp(-\eta d_i(t)), \quad d_i(t) = \text{disagreement of channel } i \text{ with the pooled reading.}$$

Decorrelated channels, which by construction disagree more often, lose weight monotonically, raising the effective ρ and lowering N_{eff} over time. The fixed point concentrates weight on a single agreeing block — the echo chamber of one. This dynamic is **[IP]**: its direction follows from the comfort gradient, but the functional form is illustrative, and the parameter η has no measured self-scale value.

A.7 The Study 1 anchor

The parent programme's Study 1 estimated, for six consumer AI systems on fifty governance-relevant quantities, an effective error correlation $\rho_{\text{eff}} \approx 0.97$ (the result is **[R]**; preregistered and confirmed). Applying A.2 at face value,

$$N_{\text{eff}} = \frac{6}{1 + 5(0.97)} = \frac{6}{5.85} \approx 1.03,$$

so six nominally distinct observers delivered the statistical protection of essentially one. The transfer to selves is **[H]**: no selves were measured, and the figure is invoked only to establish that nominal independence can coexist with near-total correlation in practice. A self whose channels share a substrate (A.5) and are actively selected for agreement (A.6) has weaker grounds than separately built AI systems to expect decorrelation, but the magnitude of ρ for any individual is unmeasured and the appendix asserts no value for it.

A.8 Simulation

The following reproduces every numerical claim above; it belongs in the flat repository as `self_ii_appendix_a_correlation_tax.py`.

```

import numpy as np
rng = np.random.default_rng(20260616)

def N_eff(N, rho):
    "Kish design effect (A.2)."
```

$$\text{return } N / (1.0 + (N - 1.0) * \rho)$$

```

def mc_pooled_var(N, rho, sigma=1.0, trials=400_000):
    "Monte-Carlo variance of the equal-weight pooled estimator (A.2)."
```

$$\text{zc} = \text{rng.standard_normal}(\text{trials})$$

$$\text{zi} = \text{rng.standard_normal}((\text{trials}, N))$$

$$\text{eps} = (\text{np.sqrt}(\rho) * \text{zc[:, None]} + \text{np.sqrt}(1 - \rho) * \text{zi}) * \text{sigma}$$

$$\text{return eps.mean(axis=1).var()}$$

```

def two_pathway_corr(rho_m, rho_d):
    "Closed-form compounded correlation (A.4)."
```

$$\text{return } 1 - (1 - \rho_m) * (1 - \rho_d)$$

```

def neff_equal_blocks(block_sizes):
    "Equal-weight N_eff = inverse Herfindahl index of block sizes (A.5)."
```

$$N = \text{sum}(\text{block_sizes})$$

$$\text{return } N^2 / \text{sum}(k^2 \text{ for } k \text{ in } \text{block_sizes})$$

```

def neff_opt_blocks(block_sizes, rho_w=1 - 1e-6):
    "Optimal-weight N_eff = block count B (A.5), via GLS:  $1^T \Sigma^{-1} 1$ ."
```

$$N = \text{sum}(\text{block_sizes}); \text{Sig} = \text{np.zeros}((N, N)); i = 0$$

$$\text{for } k \text{ in } \text{block_sizes}: \quad \text{Sig}[i:i+k, i:i+k] = \rho_w; i += k$$

$$\text{np.fill_diagonal}(\text{Sig}, 1.0)$$

$$\text{one} = \text{np.ones}(N)$$

$$\text{return one @ np.linalg.solve}(\text{Sig}, \text{one})$$

```

if __name__ == "__main__":
    print("A.2 Study-1 anchor: N_eff(6, 0.97) =", round(N_eff(6, 0.97), 4))
    print("A.2 saturation 1/rho at rho=0.5 :", round(N_eff(10**6, 0.5), 4))
    print("A.2 MC vs closed form (N=6,rho=.97):",
          round(mc_pooled_var(6, 0.97), 4), "vs", round((1-0.97)/6 + 0.97, 4))
    print("A.4 compounded corr (.2,.8)      :", two_pathway_corr(0.2, 0.8))
    for bs in ([5, 1], [4, 1, 1], [6], [1]*6):
        print(f"A.5 blocks {bs}: equal={neff_equal_blocks(bs):.3f}  opt(B)={neff_opt_blocks(bs):.3f}")
```

Running it returns $N_{\text{eff}}(6, 0.97) = 1.0256$, the saturation value 2.0 at $\rho = 0.5$, Monte-Carlo agreement with the closed form, the compounded correlation 0.84, and the block results $[5, 1] \rightarrow (1.385, 2)$, $[4, 1, 1] \rightarrow (2.0, 3)$, $[6] \rightarrow (1, 1)$, $[1 \times 6] \rightarrow (6, 6)$.

Appendix B — The Actuation Chain and the Energy Law

This appendix supplies the formal backing for Part III. The minimum-energy control result and the unreachability threshold are standard control theory applied to a delegation cascade and are rated **[R]**; the modeling of a self's intention-to-behavior path as such a cascade is **[IP]**, flagged where it enters. Scaling claims are computation-verified; the script is in B.8.

B.1 The actuation chain as a cascade

Model the path from intention to behavior as a depth- D cascade of stages — value/goal, plan, motivation, habit, environment, execution — in discrete time:

$$x_{k+1}^{(1)} = a x_k^{(1)} + b u_k, \quad x_{k+1}^{(i)} = a x_k^{(i)} + c x_k^{(i-1)} \quad (i = 2, \dots, D).$$

The control u (the directive) enters only at the first stage; the target is the deepest state $x^{(D)}$ (the behavior). In matrix form $\mathbf{x}_{k+1} = \mathbf{A}\mathbf{x}_k + \mathbf{B}u_k$ with $\mathbf{A}$ lower-bidiagonal (diagonal a , subdiagonal c) and $\mathbf{B} = b \mathbf{e}_1$. **[IP]** The three mechanisms of §3.1 are the three structural constants. *Latency* is the propagation delay: the directive's influence reaches stage i no earlier than step i . *Projection* is the inter-stage transmission c (equivalently a per-layer gain $\alpha_i \leq 1$): each stage passes only the component of its input lying within its own repertoire. *Noise* is an additive per-stage perturbation $w_k^{(i)}$, omitted from the reachability algebra below since it does not affect the reachable set, only the variance around it.

B.2 Minimum-energy control

For a target $\mathbf{x}_t$ reachable in horizon T , the minimum control energy $\sum_{k=0}^{T-1} u_k^2$ is

$$J_{\min}(T) = \mathbf{x}_t^\top \mathbf{W}_T^{-1} \mathbf{x}_t, \quad \mathbf{W}_T = \sum_{k=0}^{T-1} \mathbf{A}^k \mathbf{B} \mathbf{B}^\top (\mathbf{A}^\top)^k,$$

with $\mathbf{W}_T$ the finite-horizon controllability Gramian (pseudo-inverse on the reachable subspace where $\mathbf{W}_T$ is singular). **[R]** This is the self-scale instance of Paper XI's energy law: the effort to realize an intention is the inverse Gramian quadratic form, and the Gramian's conditioning in the target direction is set by the chain.

B.3 The effort law

Reaching the deepest coordinate $\mathbf{e}_D$ requires the directive to traverse all D stages, and the gain along the single length- D path is $\propto b c^{D-1}$, so the energy in the gain-limited regime scales as

$$J_{\min} \sim (b c^{D-1})^{-2} = b^{-2} c^{-2(D-1)},$$

geometric — that is, exponential — in depth. The simulation confirms superlinear-to-geometric growth and adds a refinement absent from the flat statement of Paper XI: **the rate depends on horizon slack**. At minimal horizon $T = D$ (the directive must act as fast as the chain permits), $J_{\min}$ grows steeply — for $c = 0.7$, $J_{\min} = 8.6, 130, 2256$ at $D = 3, 5, 7$, a per-depth ratio near 3.9, steeper even than c^{-2} because the diagonal decay a also compounds with no time to spare. With generous horizon $T = 2D$, growth is gentler but still superlinear (per-depth ratio ≈ 1.3 for $c = 0.7$, climbing toward the geometric rate as D increases). The honest summary: minimum effort grows *at least superlinearly* in delegation depth in general, and *geometrically* in the gain- or time-limited regime, with time pressure converting depth from merely costly to prohibitive. "Superlinear" is the conservative floor consistent with the parent paper; the generic and worst cases are exponential. **[R]**, scaling computation-verified.

The practical reading: the same intention is far cheaper to actuate given a long horizon than under time pressure, and the leverage is in the depth, not the effort — halving the chain length does more than any feasible increase in willpower against the exponential.

B.4 Constitutional self-uncontrollability

The threshold case is exact and is the appendix's sharpest result. The directive's influence reaches stage D no earlier than step D ; hence for any horizon $T < D$ the deepest coordinate lies *outside* the reachable subspace, $\mathbf{W}_T$ is singular in the $\mathbf{e}_D$ direction, and

$$T < D \implies \mathbf{e}_D \notin \text{range}(\mathbf{W}_T) \implies J_{\min} = \infty.$$

The simulation returns exactly this: at $D = 5$, the behavior is unreachable for $T = 4$ and finite (if large) for $T \geq 5$; at $D = 7$, unreachable for $T = 6$. **[R]** Interpreting the horizon as the window before the intention decays or the opportunity passes, the result states that an intention whose chain is deeper than its available horizon cannot become behavior regardless of effort — the directive cannot propagate fast enough to arrive in time. This is *constitutional self-uncontrollability*, and it is the exact dual of Self I's constitutional unobservability: there, states of the self beyond a critical chain depth cannot be perceived at any effort; here, states of the self beyond the reachable depth cannot be reached at any effort. The two results together bound self-governance on both channels, as Papers III and XI bound governance on both of its channels. The two failure modes feel identical from inside — sincere, total, repeated failure — and the framework asserts the distinction is real (deep-but-reachable yields to the remedies of §3.6; beyond-the-set does not) without claiming to identify, from structure alone, which a given person faces.

B.5 The Pressman–Wildavsky scalar shadow

The clearance-point arithmetic is the scalar, probabilistic special case of the same multiplicative attenuation. If a standing intention requires success at each of n independent links, each with per-link probability p , the joint success probability is $\prod_i p_i = p^n$. Verified instances: $0.99^{70} = 0.495$ (seventy near-certain clearances fall below even odds — the Oakland figure), $0.95^{10} = 0.599$ (the "executes on roughly sixty percent of days" figure of §3.2), and $0.90^{10} = 0.349$. **[R]** The intention does not fail because any link is weak; it fails because the product of many high probabilities is not high, which is the multiplicative attenuation of B.3 seen through a probabilistic rather than an energetic lens.

B.6 The legitimacy coupling

Appendix C develops self-legitimacy as a multiplicative gain $L_{\text{self}} \in [0, 1]$ on actuation, $\mathbf{B}_{\text{eff}} = L_{\text{self}}\mathbf{B}$. Entering at each link, it composes with the transmission α so that the effective per-link gain is $L\alpha$ and the path gain is $(L\alpha)^D$, giving energy $\sim (L\alpha)^{-2D}$. Because L enters the *exponent*, low self-trust does not add a fixed penalty — it multiplies the base of the exponential in depth. At $D = 8$, $\alpha = 0.8$: full trust ($L = 1$) gives path gain 0.168; $L = 0.8$ gives 0.028, an energy cost $35\times$ larger; $L = 0.5$ gives 6.6×10^{-4} , an energy cost $65536\times$ larger. **[IP]** This is the §3.5 claim — depth and distrust multiply, not add — and the self-scale form of the parent series' bidirectional node, where two compressions meet and compound. Its corollary drives the sequencing claim of Part VIII: short chains both actuate cheaply and, by delivering, raise L , which is what makes longer chains affordable.

B.7 Delegation across time

The temporal reading of §3.4 requires no new formalism. Index the cascade stages by successive time-slices of the self rather than by internal subsystems: a standing intention is a directive that each future self must re-receive and re-transmit, applying its own transmission α_i — its compliance, set by the self-legitimacy it accords the inherited directive. The chain of future selves is the same cascade, with stage depth replaced by temporal depth and each stage's α_i governed by that future self's L_{self} . The two-sided structure of B.4 and B.6 carries over unchanged. **[IP]**

B.8 Simulation

Repo file: `self_ii_appendix_b_actuation_chain.py`.

```
import numpy as np

def cascade(D, a=0.5, c=0.7, b=1.0):
    A = np.zeros((D, D)); np.fill_diagonal(A, a)
    for i in range(1, D): A[i, i-1] = c
    B = np.zeros((D, 1)); B[0, 0] = b
    return A, B

def min_energy(A, B, x_t, T):
    "Min control energy to reach x_t in T steps; inf if unreachable (B.2, B.4)."
    cols, M = [], B.copy()
    for _ in range(T): cols.append(M.copy()); M = A @ M
    R = np.hstack(cols) # columns A^k B
    z, *_ = np.linalg.lstsq(R, x_t, rcond=None)
    if np.linalg.norm(R @ z - x_t) > 1e-8: return np.inf # outside reachable set
    W = R @ R.T
    return float(x_t @ np.linalg.pinv(W, rcond=1e-15) @ x_t)

if __name__ == "__main__":
    print("B.4 unreachable (T < D => inf):")
    for D in (3, 5, 7):
        A, B = cascade(D); eD = np.eye(D)[-1]
        for T in (D-1, D, D+2):
            J = min_energy(A, B, eD, T)
            print(f"    D={D} T={T}: {'UNREACHABLE' if np.isinf(J) else f'{J:.4g}'}")
    print("B.3 effort law (T=2D):")
    for c in (0.7, 0.5):
        Js = [min_energy(*cascade(D, c=c), np.eye(D)[-1], 2*D) for D in range(2, 9)]
        print(f"    c={c}: " + ", ".join(f"{j:.3g}" for j in Js))
    print("B.5 Pressman-Wildavsky p^n:",
          {f"{p}^{n}": round(p**n, 3) for p, n in [(0.99,70),(0.95,10),(0.9,10)]})
    print("B.6 legitimacy gain (Lα)^D, D=8, α=0.8:",
          {L: round((L*0.8)**8, 5) for L in (1.0, 0.8, 0.5)})
```


Appendix C — Self-Legitimacy Dynamics

This appendix supplies the formal backing for Part IV. The coupling form is the self-scale instance of Paper XIII's LPV model and is **[IP]** as a representation; the dynamical consequences derived from it (the existence bifurcation, hysteresis, the transparency trap) follow from the model and are computation-verified. The betrayal asymmetry $\gamma \gg \alpha$ is **[H]**. The clinical fence (C.7) is **[R]** as a statement of limit, and the bifurcation of C.3 makes it sharper rather than softer. The script is in C.9.

C.1 The coupling

Self-legitimacy is a scalar coupling state $L \in [0, 1]$ that modulates both channels of self-governance at once, following Paper XIII:

$$\mathbf{B}_{\text{eff}} = L \mathbf{B}, \quad \mathbf{V} = \mathbf{V}_0 / L.$$

Low L attenuates self-actuation (a directive is discounted in advance to the degree the person does not expect themselves to honor it — the gain entering each link of Appendix B) and inflates self-observation noise (self-reports are trusted less, so introspective variance rises). One variable, two channels, moving together. **[IP]**

C.2 The dynamics

Let the per-period outcome be delivery or betrayal of a self-commitment, and let L update asymmetrically:

$$L_{t+1} = \begin{cases} L_t + \alpha(1 - L_t) & \text{on delivery (build, gain } \alpha) \\ L_t - \gamma L_t & \text{on betrayal (erode, gain } \gamma) \end{cases}, \quad \gamma \gg \alpha.$$

The factors $(1 - L)$ and L give saturation in $[0, 1]$. Delivery is endogenous: its probability rises as legitimacy crosses the level L_{half} at which self-actuation becomes reliable, $p_{\text{deliver}}(L) = \sigma(k(L - L_{\text{half}}))$, where lower L_{half} encodes higher *competence* — the capacity to deliver even at modest self-trust. The expected drift is

$$\mathbb{E}[\Delta L \mid L] = p_{\text{deliver}}(L) \alpha(1 - L) - (1 - p_{\text{deliver}}(L)) \gamma L. \quad ([IP])$$

C.3 The existence bifurcation

The simulation surfaced a result stronger than Part IV's prose, and in keeping with the series' practice it is reported as the model gives it. **Whether a healthy self-trust equilibrium exists at all depends jointly on competence and the betrayal asymmetry.** Setting the fixed points of $\mathbb{E}[\Delta L \mid L] = 0$:

- *High competence* ($L_{\text{half}} = 0.3$): a stable high-trust equilibrium persists across the tested asymmetries (healthy fixed point ≈ 0.97 – 0.99 at γ/α up to 6), but the separatrix — the unstable boundary of the collapse basin — *rises* with the asymmetry ($L^* = 0.25, 0.45, 0.55$ at $\gamma/\alpha = 2, 4, 6$). Higher betrayal-sensitivity does not destroy the healthy state but enlarges the basin from which one falls into collapse.
- *Mid competence* ($L_{\text{half}} = 0.5$): the healthy equilibrium *exists* at $\gamma/\alpha \leq 2$ and is *annihilated* by $\gamma/\alpha = 4$. Past the bifurcation the only attractor is $L \rightarrow 0$: a **collapse-only regime** in which no level of self-trust is sustainable and the spiral is not a risk but the sole outcome.

This refines Part IV materially. The self-betrayal spiral is not merely a basin one can fall into; for sufficiently low competence relative to the betrayal asymmetry, it is the only stable behavior the system admits, and starting with high trust does not help — there is no healthy state to stay in. **[R]** given the model; computation-verified.

C.4 Hysteresis

Because erosion uses gain γ and building uses gain α , the time to recover a given level of trust exceeds the time to lose it by a factor set by the asymmetry. Best-case step counts to traverse $L : 0.3 \leftrightarrow 0.7$: climbing 17 steps, falling 3 steps at $\gamma/\alpha = 6$ (ratio 5.7); falling 6 steps at $\gamma/\alpha = 3$ (ratio 2.8). The recovery-to-decline time ratio tracks γ/α . **[R]** given the model. The practical corollary, used in C.8: an impatient recovery — expecting trust back on the timescale it was lost — sets a commitment that the slow build cannot meet, and the unmet expectation registers as a further betrayal.

C.5 Built versus borrowed

The distinction reduces to competence at equal observable trust. Two agents both at $L_0 = 0.90$ receive an identical betrayal shock ($\Delta L = 0.25$). The *built* agent (competent, $L_{\text{half}} = 0.40$, separatrix 0.575) lands post-shock at 0.70, above its separatrix, and recovers to 0.905. The *borrowed* agent (low competence, $L_{\text{half}} = 0.80$, in the collapse-only regime of C.3) craters to 0 on the same shock. **[IP]**, verified. Built and borrowed self-trust can be numerically identical and behave oppositely under test, because what differs is not the trust level but the competence that determines whether a recoverable equilibrium exists beneath it. Borrowed trust is high trust sitting in, or thinly above, a collapse basin it did not earn the margin to escape.

C.6 The transparency trap

Introduce perceived legitimacy P alongside true L , and let commitment ambition scale with P while delivery depends on L — the person acts on what they believe about themselves, and the world responds to what is so. The honest agent keeps $P = L$. The self-deceiver suppresses betrayals in perception (P unchanged on betrayal) while L erodes as before. The simulation: over 140 periods the honest agent holds at true $L = 0.70$ with 27 betrayals, perception tracking truth. The self-deceiver, perceiving high trust, keeps over-committing relative to true capacity, incurs 112 betrayals, and its true legitimacy collapses to $L = 0.01$ while perceived P remains 0.68; the hidden discrepancy peaks at 0.78. A forced reckoning ($P \rightarrow L$) corrects perception once, but the continuing deception reopens the gap. **[IP]**, verified. The trap maintains apparent self-trust while destroying the real thing, and it does so by the precise mechanism of §4.3: editing the observation channel inflates perceived trust, inflated perception drives over-commitment, over-commitment multiplies betrayal, and betrayal craters the true state the perception was hiding.

C.7 The reflexive sharpening and the clinical fence

The transparency trap is worse for a self than for an institution for the reason of Part I: the agent decides on P and has no instrument independent of the apparatus generating P . A government can audit its legitimacy with sensors external to the governing body; a single mind's self-deceiving instrument is also the instrument that would detect the deception, and a corrupted channel cannot reliably audit its own corruption (C.6 formalizes this by making decisions depend only on P , with no read of L). The only approximation to an independent sensor is external — the friend, the therapist, the re-read journal of Appendix A's decorrelated channels. **[IP]**

The fence, which the bifurcation of C.3 makes more important, not less. **[R]** The model produces a collapse-only regime — a structural attractor in which self-trust cannot be sustained. This does *not* license reading any individual's collapse of self-trust as structurally caused. The same attractor shape can be produced by substrates the control model does not

represent — the neurochemistry of depression, the physiology of trauma, grief, illness — and the model cannot distinguish a structurally generated collapse from a clinically generated one, because their signatures coincide. The framework describes the shape of the trap; it does not diagnose its cause, it is not a clinician, and the recovery conditions of C.8 are conditions that may help a structurally generated decline, not a treatment for a pathologically generated one.

C.8 Recovery and the perfectionism inversion

The conditions for rebuilding L follow from the dynamics. *Delivery–reality matching* — commitments small enough that p_{deliver} is near one — rebuilds L on the α side and keeps the trajectory clear of the separatrix. *Transparency to oneself* prevents the C.6 gap from opening. *Credible commitment devices* hold action while L is too low to carry it. *Hysteresis-awareness* (C.4) is itself a condition: recovery is slow by construction, and the demand for speed is self-defeating. The inversion of §4.6 falls out of the model directly. The perfectionist sets large commitments, which in the C.6 mechanism means high ambition s relative to true capacity — maximal over-commitment, maximal betrayal rate, the fast track into the collapse basin. Aiming high is not the opposite of the self-betrayal spiral; it is the spiral run at maximum gain. Built trust is recovered by commitments small enough to keep, kept long enough to count — which is also why Part VIII enters a coupled failure here, at the gain on every other primitive.

C.9 Simulation

Repo file: `self_ii_appendix_c_self_legitimacy.py`.

```
import numpy as np
def sigma(z): return 1/(1+np.exp(-z))
def p_deliver(L, k=8.0, L_half=0.5): return sigma(k*(L-L_half))
def EdL(L, alpha, gamma, **kw):
    p = p_deliver(L, **kw); return p*alpha*(1-L) - (1-p)*gamma*L

def fixed_points(alpha, gamma, **kw):
    # C.3 bifurcation
    Ls = np.linspace(1e-4, 1-1e-4, 400001); f = EdL(Ls, alpha, gamma, **kw)
    out = []
    for i in np.where(np.diff(np.sign(f)) != 0)[0]:
        r = Ls[i] - f[i]*(Ls[i+1]-Ls[i])/(f[i+1]-f[i])
        dr = (EdL(r+1e-4, alpha, gamma, **kw) - EdL(r-1e-4, alpha, gamma, **kw))/2e-4
        out.append((round(r, 3), 'stable' if dr < 0 else 'unstable'))
    return out

def trap(deceive, alpha=0.05, gamma=0.15, k=8.0, margin=0.15, T=140, reckon=90, seed=3):
    r = np.random.default_rng(seed); L = P = 0.70; gap = 0.0; betr = 0
    for t in range(T):
        s = max(P - margin, 0.0) # commit on perceived trust
        if r.random() < sigma(k*(L - s)): # deliver depends on true trust
            L += alpha*(1-L); P += alpha*(1-P)
        else:
            betr += 1; L = max(L-gamma*L, 1e-4)
            P = P if deceive else max(P-gamma*P, 1e-4)
        gap = max(gap, P-L)
        if t == reckon and deceive: P = L
    return round(L,3), round(P,3), round(gap,3), betr

if __name__ == "__main__":
    print("C.3 bifurcation (mid competence L_half=0.5):")
    for g in (0.10, 0.20):
        healthy = [r for r,s in fixed_points(0.05, g, L_half=0.5) if s=='stable' and r>0.4]
        print(f"    gamma/alpha={g/0.05:.0f}: {'healthy FP '+str(healthy) if healthy else 'COLLAPSE-ONLY'}")
    print("C.6 transparency trap (true L, perceived P, max gap, betrayals):")
    print("    honest      :", trap(False))
    print("    self-deceiving:", trap(True))
```


Appendix D — Adaptive Learning and the Two-Sided Bound

This appendix supplies the formal backing for Part VI, and it carries the paper's strongest original claim — the two-sided bound on self-revision (§6.4). The distinction between premise and test matters here more than anywhere. That revision degrades a self's coherence (the parameter $\kappa > 0$ below) is the **[IP]** premise, the modeling commitment that encodes observer–plant identity; it cannot be tested in simulation because it is an assumption about what a self *is*. What *can* be tested — and could have failed — is whether that premise produces a non-trivial two-sided bound: an interior optimum, a self optimum strictly below an institution's, and a regime where the self must sacrifice tracking to stay coherent. The simulation confirms all three. All values are computation-verified; the script is in D.8.

D.1 The self as a dual controller

Model the self tracking a drifting target — what circumstances require the self to be — by an estimate it acts on:

$$\theta_{t+1}^* = \theta_t^* + v \xi_t, \quad \hat{\theta}_{t+1} = \hat{\theta}_t + r C_t (o_t - \hat{\theta}_t), \quad o_t = \theta_t^* + \sigma \eta_t,$$

where θ^* is the moving target drifting at rate v , $\hat{\theta}$ the self-model, o_t a noisy self-observation, r the *revision rate* (the exploration/learning gain), and $C_t \in (0, 1]$ the controller's coherence, defined next. Each step is simultaneously an action (acting on $\hat{\theta}$) and an experiment (observing θ^*), the dual-control structure of §6.1. **[IP]**

D.2 Persistence of excitation: the lower bound

Tracking a target drifting at rate v requires revision fast enough to keep pace; too little produces lag. The simulation gives true tracking error E rising sharply as $r \rightarrow 0$: $E = 0.061$ at the tracking optimum $r = 0.08$, but 0.117 at $r = 0.01$ and 0.170 at $r = 0.005$. The $r \rightarrow 0$ limit is the over-protected self of §6.2 — variance suppressed, parameters unidentified, the self-model frozen while circumstances move. This is the lower bound, and it is the standard persistence-of-excitation requirement: a controller that does not probe cannot identify the target it must track. **[R]** given the model; verified.

D.3 Exploration starvation and self-concealment

The lower bound hides itself, which is what makes it starvation rather than crisis. Model perceived error as what the self can detect through its own exploration, $\hat{E}_t = E_t (1 - e^{-cr})$: at high revision the self sees its true error; at low revision it is nearly blind to it. The simulation gives, at $r = 0.02$, true error $E = 0.084$ but perceived error $\hat{E} = 0.013$ — a self-concealment gap of 0.071, the green dashboard of §6.3 reading clear while true tracking diverges. The gap closes as r rises (it is 0.0006 at $r = 0.7$). The structural point is that the faculty that would reveal the starvation, exploration, is precisely the one that has been switched off, so the system's own monitoring cannot report the failure. This is Self I's variety gap given a temporal mechanism: not only does a narrow architecture exclude dimensions of the present self, it stops updating and conceals that it has. **[IP]**, verified.

D.4 The coherence coupling and the two-sided bound

The self-specific term is coherence. Because the controller is the plant (Part I), each act of revision churns the controller's own integration:

$$C_{t+1} = \text{clip}(C_t + \iota(1 - C_t) - \kappa |\Delta \hat{\theta}_t|, 0, 1),$$

where ι is the re-integration rate, $|\Delta\hat{\theta}_t|$ the magnitude of self-revision actually performed, and κ the coherence cost per unit revision. Coherence gates revision in turn (the gain is $r C_t$): a destabilized self cannot integrate change. The premise is $\kappa > 0$ — for a self, revising oneself degrades the self doing the revising — and it is **[IP]**, the encoding of observer–plant identity, not a result.

Its tested consequence is a genuine two-sided bound. The self's effective objective is to track well *and* stay coherent, $J(r) = E(r) + \lambda(1 - \bar{C}(r))$, and the simulation shows $C(r)$ declining monotonically (from 0.975 at $r = 0.02$ to 0.453 at $r = 0.9$) while $E(r)$ is U-shaped. The result: J has an *interior* optimum — verified to turn upward on both sides (rising as $r \rightarrow 0$ from starvation, rising as r grows from coherence loss) — for every coherence weight tested ($\lambda \in \{0.5, 1, 2\}$). The bound is two-sided exactly as §6.4 claims: revision must be fast enough to track and slow enough to integrate, (drift rate) $< r <$ (coherence-limited rate). **[IP]** premise; the interior optimum is a tested, falsifiable consequence, confirmed.

D.5 Self versus institution: why the upper bound is tighter

The contrast with an institution isolates what observer–plant identity contributes. An institution insulates its experimenting apparatus from its experimental zones, modeled by $\kappa = 0$: revising the plant does not degrade the controller's coherence. With $\kappa = 0$ the simulation holds $C \equiv 1$ at all revision rates, and the optimal r is the tracking optimum, $r_{\text{inst}}^* = 0.08$, set by the standard trade-off between lag and noise amplification. With $\kappa = 0.5$ (the self), the optimum drops to $r_{\text{self}}^* = 0.02\text{--}0.04$, strictly below the institution's and robust across λ . At its optimum the self carries true error 0.084 against a best-achievable 0.061: it leaves 0.023 of tracking performance on the table to preserve coherence. **[IP]**, verified.

One refinement keeps the claim honest. The institution is not unbounded above — its error is also U-shaped, rising past $r = 0.08$ from noise amplification, so it too has an upper bound. The self's upper bound is *tighter* and arises from a *different mechanism*: it binds at $r \approx 0.02\text{--}0.04$, well below the institution's noise-limited 0.08, and it binds through coherence loss rather than noise. The precise content of "the upper bound is tighter for a self than for a society" (§6.4) is therefore not that institutions revise without limit, but that the self's limit is stricter and arrives first, forcing the self below its own noise-limited tracking optimum — a constraint a system whose controller and plant are separate does not face.

D.6 The boundary the framework must not cross

The two-sided bound borders clinical and contemplative territory, and the fence of Parts IV and VI applies unchanged. **[R]** The model describes the *structure* of the integration constraint — that revision outrunning re-integration (r above the coherence-limited rate) degrades the coherence learning requires. It does not describe the management of that constraint in any person, and it does not claim that destabilization under intensive self-work is generally structural. Where a clinical substrate is present — where self-revision triggers something with a physiology the control model does not represent — the architectural description remains true and insufficient, naming the shape of the difficulty without reaching its cause. The observation that transformation must be paced to integration is a consequence of $r < r_{\text{coh}}$, not a protocol for pacing it.

D.7 Protected spaces as manufactured insulation

The design principle of §6.6 follows directly and is verified. A protected experimental space is revision insulated from global coherence — a fraction f of revision charged not to the whole self but to a walled-off sandbox. Modeling this as coherence cost $\kappa |\Delta\hat{\theta}| (1 - f)$, the simulation holds revision at $r = 0.15$ (above the un-sandboxed self's optimum) and varies f : tracking error stays essentially flat ($E : 0.068 \rightarrow 0.072$) while coherence recovers ($C : 0.843 \rightarrow 0.982$) and the combined cost falls ($J : 0.226 \rightarrow 0.091$) as f rises from 0 to 0.9. As $f \rightarrow 1$ the self approaches the institution's free exploration. **[IP]**, verified. This is the formal content of the externalization principle of Part VIII at the

learning channel: a sandbox lets the self run the high local exploration that tracking wants without paying the global coherence cost that observer–plant identity otherwise imposes — manufacturing, locally, the controller–plant separation the self lacks by default.

D.8 Simulation

Repo file: `self_ii_appendix_d_two_sided_bound.py`.

```
import numpy as np

def run(r, kappa, f_sandbox=0.0, v=0.02, sigma=0.30, iota=0.10, c=8.0, T=4000, seed=0):
    rng = np.random.default_rng(seed)
    th_star = th_hat = 0.0; C = 1.0; Es=[]; Cs=[]; Ehat=[]
    for t in range(T):
        th_star += v*rng.standard_normal()
        obs = th_star + sigma*rng.standard_normal()
        dth = r*C*(obs - th_hat); th_hat += dth
        C = np.clip(C + iota*(1-C) - kappa*abs(dth)*(1-f_sandbox), 0, 1)
        if t > T//5:
            e = abs(th_hat - th_star)
            Es.append(e); Cs.append(C); Ehat.append(e*(1-np.exp(-c*r)))
    return np.mean(Es), np.mean(Cs), np.mean(Ehat)

def sweep(kappa, rs, seeds=8):
    return np.array([np.mean([run(r, kappa, seed=s) for s in range(seeds)], axis=0) for r in rs])

if __name__ == "__main__":
    rs = [0.005, 0.01, 0.02, 0.04, 0.08, 0.15, 0.30, 0.60]
    for kappa, lab in [(0.0, "institution"), (0.5, "self")]:
        R = sweep(kappa, rs); E, C, Eh = R[:,0], R[:,1], R[:,2]
        for lam in (0.5, 1.0, 2.0):
            J = E + lam*(1-C); k = int(np.argmin(J))
            interior = 0 < k < len(rs)-1
            print(f"{lab:11s} lam={lam}: r*={rs[k]:.3f} interior={interior} "
                  f"E@r*={E[k]:.4f} bestE={E.min():.4f}")
        print("D.3 self-concealment gap (E - Ehat) at r=0.02:",
              round((lambda x: x[0]-x[2])(np.mean([run(0.02,0.5,seed=s) for s in range(8)],axis=0)), 4))
        print("D.7 protected space (r=0.15, vary f): (E, C, J)")
        for f in (0.0, 0.3, 0.6, 0.9):
            E,C,_ = np.mean([run(0.15,0.5,f_sandbox=f,seed=s) for s in range(8)],axis=0)
            print(f"    f={f}: E={E:.4f} C={C:.3f} J={E+(1-C):.4f}")
```


Appendix E — Observer–Plant Identity and the Measurement–Disturbance Coupling

This appendix supplies the formal backing for Part I. It is the conceptual keystone, and it carries less computational weight than A, C, and D by design: the central claim — that for a self the act of observation is an act on the observed system — is structural, and the empirical and simulated consequences of that claim were established in the other appendices. E states the premise precisely, identifies exactly which control-theoretic guarantee it voids, and demonstrates the two consequences that belong to it alone and are not assumed in stating it. The separation-principle claim (E.1) is **[R]** as a statement of limit; the premise that it fails for a self is **[IP]**; the fixed-point results (E.2–E.3) are computation-verified consequences of the reflexive coupling, not re-expressions of it.

E.1 The separation principle and the premise that fails

The separation (certainty-equivalence) theorem of linear-quadratic-Gaussian control states that the optimal controller decomposes into two independently designed parts: an optimal state estimator, and a control law that acts on the estimate exactly as it would on the true state. Estimation and control are separable — one may first determine the state, then decide what to do. The result rests on a structural premise that is rarely made explicit because it is almost always satisfied: the observation process does not enter the plant dynamics. Measuring the system does not change it. **[R]**

A self violates this premise at its root (Part I §1.3). The act of self-observation — attending to a state, asking how one feels, narrating one's situation — enters the dynamics of the very state being observed; there is one system, and reading it is operating on it. The premise of the separation theorem is therefore false for a self, and its conclusion is not available: estimation and control are not separable, and "estimate then act" is not a sequence a self can perform, because the estimating is already an acting. **[IP]** This is the same kind of limit the boundary appendix recorded for the small-gain theorem: a named control-theoretic guarantee whose stated premise the self-case does not meet.

A note on what this appendix does not do. One could model measurement back-action directly — assume observing injects a state disturbance, then show a controller that ignores it does worse — but that demonstration would be circular, recovering only the premise it assumed. E instead demonstrates two consequences of the reflexive coupling that are *not* assumed in stating it: that self-perception has many self-confirming fixed points rather than one true value (E.2), and that observation can therefore relocate the state it observes (E.3). Both emerge from the coupling; neither is built into it.

E.2 Self-perception as fixed-point selection, not truth discovery

Model a self-state x and a self-belief b on some dimension. A *passive* observer's belief tracks a fixed external truth: x is unaffected by being observed, and b converges to it. A *reflexive* observer's belief enters the dynamics — believing pulls the state toward the belief ($x \leftarrow x + \eta(b - x)$, observation as action) and the reading is drawn toward the belief ($\text{obs} = x + \text{bias}(b - x)$, confirmation). The simulation, from a common true initial state $x_0 = 0.5$:

- *Passive* ($\eta = 0$, no bias): belief converges to 0.50 from every initial belief $b_0 \in \{0.1, \dots, 0.9\}$ — a single fixed point, the truth, discovered regardless of where belief started.
- *Reflexive* ($\eta = 0.15$, bias = 0.5): the final state and belief are selected by b_0 — $b_0 = 0.1 \rightarrow 0.245$, $b_0 = 0.5 \rightarrow 0.50$, $b_0 = 0.9 \rightarrow 0.755$ — a *continuum* of self-confirming fixed points, each a blend of initial belief and prior state (here symmetric about the true 0.5 because confirmation is partial), with the one reached selected by where belief began rather than by any prior truth.

This is the formal content of §1.3: for a self, "accurate self-perception" is not the recovery of a pre-existing fact but the selection of a fixed point of a reflexive process — a belief that is true of the state your believing it produces. The fixed points are neutrally stable along the manifold $b = x$, so a self-concept, once settled, is self-confirming and persists. **[IP]** premise; multiplicity is a verified, non-assumed consequence. The danger edge is immediate: a false but self-confirming self-concept is a fixed point like any other, and the apparatus that might reveal its falseness is the apparatus maintaining it — the sealed self of Part VII, seen at its root.

E.3 Self-observation as intervention

The same coupling that produces the danger produces the lever of §1.4. Because observation moves the state, a deliberate change of belief can *relocate* the fixed point — something a passive observer cannot do, since for it there is only one truth to find. Starting locked at a low self-concept ($b = x = 0.2$) and applying a one-time belief intervention to 0.65 at mid-run, the simulation gives: the reflexive system relocates and holds at a new self-confirming state ($b = x = 0.487$), while the passive system returns to the true 0.2, the intervention leaving no trace. **[IP]/[H]**, verified.

Two honest qualifications. The relocation lands *partway* — at 0.487, between the old state and the intended 0.65 — because the new fixed point is set by the balance of action (η) and confirmation (*bias*), so a single re-narration moves the self toward a chosen self-concept but not all the way to it; sustained intervention is required to move further, which is why the practices that exploit this lever (contemplative attention, the disciplined re-narration of therapy) are repeated rather than singular, and connects to the small, sustained commitments of Parts III and IV. And this is the mechanism, promised in §1.4, by which Self I's asymptote $G_{\text{self}} \rightarrow 0$ is approachable at all: in a system where observing is acting, attending to an excluded dimension is the same operation as beginning to admit it — the lever and the danger are one coupling seen from two sides.

E.4 The common root

Observer–plant identity is a single premise, and the paper's principal departures from the parent series are its consequences, established where each could be tested:

- the floor under self-observation correlation, because channels share one substrate (Appendix A.5);
- the unauditability of self-legitimacy, because the deceiving instrument is the auditing instrument (Appendix C.6–C.7);
- the two-sided bound on self-revision, because the experimenter is the experiment (Appendix D.4–D.5);
- the absence of any firewall between primitives, making composite failure the default (Part VII).

To these E adds the two consequences proper to the premise itself: the multiplicity of self-confirming fixed points (E.2) and the relocation lever (E.3). The distribution is deliberate and is the honest shape of the paper's evidence. E is the keystone — the single structural fact from which the divergences follow — but it is not where the weight rests. The weight rests on the appendices that could compute or simulate the consequences, and E's role is to show that those scattered results are not five separate findings but one premise seen five times: in a self, the controller and the plant cannot be pulled apart, and everything the paper claims that the parent series does not is what that single fact entails.

E.5 Simulation

Repo file: `self_ii_appendix_e_observer_plant.py`.

```

import numpy as np

def observe_self(b0, eta, bias, x0=0.5, lam=0.2, T=4000,
                intervene_at=None, intervene_to=None):
    """Reflexive self-observation. eta: observation-as-action (0 = passive observer);
    bias: confirmation in the reading. Optional one-time belief intervention."""
    x, b = x0, b0
    for t in range(T):
        if intervene_at is not None and t == intervene_at:
            b = intervene_to
        x = x + eta*(b - x)          # observing/believing moves the state (1.3/1.4)
        obs = x + bias*(b - x)      # reading drawn toward belief (confirmation)
        b = b + lam*(obs - b)
    return round(x, 3), round(b, 3)

if __name__ == "__main__":
    print("E.2 passive vs reflexive (common true x0=0.5):")
    for b0 in (0.1, 0.5, 0.9):
        print(f"    b0={b0}: passive {observe_self(b0,0.0,0.0)} reflexive {observe_self(b0,0.15,0.5)}")
    print("E.3 belief intervention -> 0.65 at t=1500 (start locked at 0.2):")
    print("    reflexive:", observe_self(0.2,0.15,0.5,x0=0.2,intervene_at=1500,intervene_to=0.65))
    print("    passive  :", observe_self(0.2,0.0,0.0,x0=0.2,intervene_at=1500,intervene_to=0.65))

```