

Restorability Boundaries

Flyxion

Independent Researcher

Abstract

This volume develops a framework linking distinction, information, entropy, admissibility, repair, restorability, continuation, and reconstruction. Distinctions induce the partitions on which information and entropy are defined. Admissibility, given two independent characterizations that are shown to agree in type but not yet in content, governs which transitions between states are reachable at all. Repair, and the operators that persist as ecologies rather than isolated fixes, keep admissible structure viable under damage. Restorability boundaries mark where that viability fails, by four formally distinct mechanisms rather than one. Continuation resolves an apparent contradiction between admissibility as checked and admissibility as authored by showing each describes a different level of the same structure. Reconstruction closes the framework with a resource-independence result: some collapsed distinctions are not merely expensive to recover but structurally unrecoverable by any allocation of time, space, or energy. The volume's central open problem — whether its two admissibility constructions induce the same continuation family, a sharper question than whether they are equal as quantities — is stated precisely, given a formal argument for skepticism, and left open.

Preface

Persistence depends on the preservation and restoration of distinctions. That is the central claim of this volume, and the chapters that follow develop it through a sequence of increasingly structured constructions. Distinguishability induces information; information induces entropy; entropy constrains admissibility; admissibility supports repair; repair determines restorability; restorability governs continuation; and reconstruction studies what survives, and what does not, when distinctions are lost.

The mathematics is the argument. Distinguishability spaces, coding deficits, and the four operations of Chapter 1 are not preliminaries to a later, looser discussion — they are the objects the rest of the volume is built from, checked at each stage by substitution rather than assumed by analogy. When two constructions share a name or a formula, this volume treats that as a question to be settled, not a license to identify them: a type test first, then a test of whether one is literally a special case of the other. Most of the time, the answer turns out to be no. The volume’s single largest open problem — whether two independently-motivated notions of admissibility induce the same family of permitted continuations — is a direct consequence of taking that discipline seriously rather than resolving it by convenience.

Four appendices collect material that would otherwise interrupt the chapters that need it: a notation index, the shared machinery of quotients and fibers underlying Chapters 1, 2, and 6, the continuation and reachability structures used throughout the second half of the volume, and the persistence and reconstruction results that close it.

Contents

Preface	1
1 Distinction	3
1.1 Introduction	3
1.2 The Distinguishability Space	3
1.3 Two Deficits	4
1.4 The Four Operations	5
1.5 Toward Admissibility	6
2 Information	9
2.1 Introduction	9
2.2 Partitions and Surprisal	9
2.3 The Compression Theorem	10
2.4 Representational Entropy	10
2.5 The Representation Limit Theorem	11
2.6 The Fundamental Information Theorem	11
3 Admissibility	13
3.1 Two Definitions	13
3.2 Testing the Identification	13
3.2.1 A Structural Reason for Skepticism	15
3.3 The Architecture Gate as a Special Case	16
3.4 Continuation Needs No New Machinery	16
4 Constraint Populations	18
4.1 From a Single Region to a Population	18
4.2 Constraint Populations Are Sympoietic Systems	18
4.3 Two Independent Pathologies	19
4.4 Feasibility Constrains Measured Diversity — Sometimes Perversely	20
4.5 Marginal Diversity Contribution	22
4.6 Worked Examples	23
4.7 Connections to the Rest of This Volume	23

5	Entropy	26
5.1	Introduction: Five Candidates, Not One	26
5.2	Entropy A: Shannon/Distinction Entropy	27
5.3	Entropy B: Representational Entropy	27
5.4	The Distinction–Entropy Duality: A and B Are Not Independent .	27
5.5	Entropy C: Historical Entropy Reduces to B	27
5.6	Entropy Production, the Second Law, and Reachability	28
5.7	Entropy D: Repair Entropy	30
5.8	Entropy E: The RSVP Field, Left Open	32
5.9	Summary	33
6	Repair	34
6.1	Introduction	34
6.2	Three Levels of Modification	34
6.3	From Structures to Repairs to Ecologies	35
6.4	A Hierarchy of Equivalence	35
6.5	Operator Ecologies as Constraint Networks	36
6.6	Keystone Operators, Extinction, and Health	36
6.7	History, Compression, and Repairability	40
6.8	A Worked Example: Misalignment, Extinction, Regeneration . . .	42
6.9	Toward the Restorability Boundary	44
7	The Restorability Boundary	45
7.1	Introduction	45
7.2	The Signal–Diagnosis–Repair–Continuation Chain	46
7.3	Two Kinds of Restorability Boundary	47
7.3.1	The architecture gate	47
7.3.2	History indexing	48
7.4	Four Failure Modes	51
7.4.1	Expiatory Gap: throughput failure	51
7.4.2	Diagnostic Collapse: discriminability failure	52
7.4.3	Repair Deficit: operator failure	53
7.4.4	Continuation Failure: composition failure	56
7.5	Relating the Four Modes	57
7.6	Worked Examples Across the Four Modes	59
7.7	Open Problems for This Chapter	60
8	Continuation	61
8.1	Introduction: A Short Chapter, Predicted in Advance	61
8.2	Trajectories and Coordination Geometries	61
8.3	Checking Versus Authoring	62
8.4	Geometric Repair as Level-3 Modification	62

8.5	Geometric Repair and Operator-Ecology Regeneration	63
8.5.1	The Induced Ecology, Type-Corrected	63
8.5.2	The Regeneration Proposition	64
8.5.3	The Reachability Gain Vector	64
8.5.4	Signed Extension: Non-Monotone Repair	66
9	Reconstruction	68
9.1	Introduction: The Question This Chapter Actually Asks	68
9.2	Two Kinds of Materialization	68
9.3	The Resource Independence Theorem	69
9.4	The Distinction Budget	70
9.5	Historical Observability	71
9.6	Admissibility-Relevant Fibers and Selective Retention	72
9.7	The Allocation Problem and a Design Principle	72
	Conclusion	74
	Appendix A: Notation and Symbol Index	76
	Appendix B: Quotients, Fibers, and Projections	79
	Appendix C: Continuation and Reachability Structures	82
	Appendix D: Persistence and Reconstruction	85
	Appendix E: Worked Examples	90
	References	94

Chapter 1

Distinction

1.1 Introduction

Three independent lines of work, developed for unrelated purposes, converge on the same mathematical object once their concerns are stated precisely enough to compare. A line of work proceeding from process-primary computation establishes that an observable state is a projection of a richer history space, and that distinct histories routinely collapse to the same observable state under that projection. A line of work proceeding from kernel statistics demonstrates that distinct continuous distributions become mutually singular after an appropriate embedding, separated in the embedding space even where they overlapped substantially before it. A line of work proceeding from archive-based compression identifies a notion of reachable complexity: the shortest description an archive can currently produce, which may be strictly longer than the shortest description that exists in principle. The first of these is concerned with distinctions that are destroyed. The second is concerned with distinctions that are amplified. The third is concerned with distinctions that are, for now, simply inaccessible. None of the three names the object they share. This chapter makes it explicit.

1.2 The Distinguishability Space

Definition 1.2.1 (Distinguishability space). *A distinguishability space is a pair $(X, \sim)$ where X is a set and $\sim$ is an equivalence relation on X , interpreted as: $x \sim y$ exactly when x and y cannot be told apart within the current representational system. The equivalence classes $[x]_{\sim}$ are its fibers.*

A finer relation has smaller fibers and expresses more distinctions; a coarser relation has larger fibers and expresses fewer. The two extremes are the discrete relation, under which every element is distinguishable from every other, and

the indiscrete relation, under which nothing is distinguishable from anything. The canonical instance recurring throughout this book is a projection $\sigma : H \rightarrow S$ from a history space H to an observable state space S , where $h \sim h'$ exactly when $\sigma(h) = \sigma(h')$: the fiber above a state s is the set of every history that would have produced it.

Principle 1.2.2 (Distinguishability Principle). *A representational system is characterized not by what objects it contains but by what distinctions it can express.*

An immediate consequence of principle 1.2.2 is that no finitary representational system can express every distinction present in what it represents, so every such system carries some deficit relative to what a maximally expressive system could achieve. Making that deficit precise is this chapter's first formal task.

1.3 Two Deficits

Let T denote an ontology in the sense used throughout this book: a description language, a template hierarchy, or a representational system generating a distinguishability relation on some set of objects.

Definition 1.3.1 (Coding deficit). *Let $L_T(x)$ be the length of the shortest description of x reachable within T , and let $L^*(x)$ be the shortest description of x achievable by any ontology. The coding deficit is $\delta_T(x) = L_T(x) - L^*(x)$.*

Definition 1.3.2 (Distinguishability deficit). *Let $D_T(x)$ be the number of distinct objects separable from x by descriptions reachable within T , and let $D^*(x)$ be the maximum such capacity achievable by any ontology. The distinguishability deficit is $\Delta_T(x) = D^*(x) - D_T(x)$.*

These measure different things: $\delta_T(x)$ is a description-length excess, how many bits are wasted; $\Delta_T(x)$ is an inaccessible distinction capacity, how many distinctions involving x cannot be expressed in T at all. They are related rather than interchangeable.

Claim 1.3.3 (Deficit Correspondence). *Under the assumption that descriptions and distinctions are related by a faithful encoding, $\Delta_T(x) \leq C \cdot \delta_T(x)$ for a system-dependent constant $C > 0$, and $\delta_T(x) = 0$ if and only if $\Delta_T(x) = 0$: local optimality in coding corresponds exactly to exhausting every expressible distinction.*

Both deficits are non-negative for every ontology and every object, and both vanish exactly when the ontology is locally optimal at that object. Throughout the rest of this book, δ_T is used where the argument concerns description

length and Δ_T where the argument concerns fiber structure directly; Claim 1.3.3 licenses treating a claim proved for one as evidence, though not proof, for the corresponding claim about the other.

1.4 The Four Operations

Representational systems do not stay fixed. principle 1.2.2 raises an immediate question: what are the possible ways a distinguishability structure can change? Four operations answer it, organized by a two-level architecture.

Objects $\longleftrightarrow$ Distinctions $\longleftrightarrow$ Ontologies.

Projection and embedding operate at the level of *distinctions*: they modify $\sim$ within a single fixed space, moving between finer and coarser relations on the same underlying set. Revision and transport operate at the level of *ontologies*: revision changes the generating ontology of a space outright, and transport connects two distinct spaces to one another.

Definition 1.4.1 (Projection). *A projection on $(X, \sim)$ is a surjective map $\sigma : X \rightarrow S$ inducing a coarsening $\sim'$ of $\sim$: $\sigma(x) = \sigma(y)$ implies $x \sim' y$.*

Claim 1.4.2 (Projection Deficit Bound). *For a projection $\sigma : X \rightarrow S$ with fiber $F_s = \sigma^{-1}(s)$ above a state s with $|F_s| > 1$, the deficit increase satisfies $\Delta\delta(\sigma) \geq \log |F_s|$: projecting to S discards at least $\log |F_s|$ bits of distinguishing information, transferred from the realized history into the now-inaccessible remainder rather than destroyed outright.*

Definition 1.4.3 (Embedding). *An embedding of $(X, \sim)$ into $(Y, \sim')$ is an injective map $\phi : X \rightarrow Y$ such that $\phi(x) \sim' \phi(y)$ implies $x \sim y$. It is faithful when the converse holds as well.*

Claim 1.4.4 (Embedding Refinement). *A faithful embedding never increases the coding deficit: $\delta_{\phi(X)}(\phi(x)) \leq \delta_X(x)$. A non-faithful embedding may introduce spurious distinctions — elements separated in the target space that were equivalent in the source — and the refinement guarantee fails for it in general.*

Definition 1.4.5 (Revision). *An ontological revision is a transition $(X, \sim_T) \rightarrow (X, \sim_{T'})$ replacing an ontology T with a new ontology T' , changing which descriptions are reachable rather than merely reorganizing $\sim$ within a fixed T . A revision is productive when $\delta_{T'}(x) < \delta_T(x)$.*

Claim 1.4.6 (Revision Monotonicity). *For a sequence of ontologies $T_0 \subset T_1 \subset T_2 \subset \dots$ each extending the last, the coding deficit is non-increasing at every step, strictly*

decreasing wherever the revision is productive. The canonical failure mode is ontological lock-in: the deficit is positive and stable, no productive revision is available, and the system cannot escape its current template hierarchy because the path to a better ontology does not lie within the admissibility structure of the current one.

Definition 1.4.7 (Transport). A transport map is a function $\tau : (X, \sim_X) \rightarrow (Y, \sim_Y)$ with bounded distortion $\epsilon(\tau)$, preserving equivalence in the forward direction and controlling the measure of pairs that collapse together in Y without having been equivalent in X .

Claim 1.4.8 (Transport Stability). For a transport map τ with distortion at most ϵ and a fiber-constant invariant I , $|I(x) - I(\tau(x))| \leq C\epsilon$ for a constant C depending on the Lipschitz properties of I . An analogy τ is justified for inference exactly when its distortion is controlled and the relevant invariants have bounded Lipschitz constants — the precise condition under which reasoning by analogy is valid rather than merely suggestive.

Claim 1.4.9 (Representational Reduction). Every representational transformation induces a transformation on distinguishability structure, classified by its effect on fibers: a transformation that merges fibers is a projection, one that separates fibers is an embedding, one that changes the generating ontology is a revision, and one that connects distinct spaces is a transport. These four operations are not imposed as a taxonomy; they are the exhaustive qualitative possibilities for fiber behavior.

Open Problem 1.4.10 (Classification). It remains open whether every representational transformation factors into a finite composition of projection, embedding, revision, and transport, and, if so, whether a single expression bounds how the coding deficit accumulates across an arbitrary such composition.

1.5 Toward Admissibility

The deficit apparatus of section 1.3 connects directly to the admissibility geometry used throughout the rest of this book, though the connection needs to be stated carefully rather than assumed.

Claim 1.5.1 (Admissibility via deficit). An admissibility relation $\mathcal{A}$ specifies which transitions between states are reachable from a given state. States outside $\mathcal{A}$ from an ontology T are exactly those whose description would require $\delta_T(x) > 0$ from within T , and whose distinction structure is correspondingly inaccessible: $\Delta_T(x) > 0$.

This is a graded characterization: $\delta_T(x)$ and $\Delta_T(x)$ take real, non-negative values, and admissibility failure is a matter of degree rather than a single threshold crossed once and for all.

Remark 1.5.2 (An open connection, not an identity). Chapter 6 uses a different and cruder formalization: an architecture gate $\chi_\pi(x_{\text{obs}}) \in \{0, 1\}$, testing simple set membership in $\text{Im}(\pi)$ for a fixed prediction map $\pi : \mathcal{M} \rightarrow \Sigma$, with no graded notion of degree once outside the image. This chapter’s δ_T/Δ_T machinery is a natural candidate for refining that binary gate into something continuous, and a partial bridge — though not a full identification — can now be given explicitly.

Definition 1.5.3 (Architectural representation deficit). For a metric observation space (Σ, d_Σ) and prediction map $\pi : \mathcal{M} \rightarrow \Sigma$, define

$$\delta_\pi^{\text{arch}}(x) = \inf_{m \in \mathcal{M}} d_\Sigma(\pi(m), x).$$

Claim 1.5.4 (Gate as thresholded architectural deficit). If $\text{Im}(\pi)$ is closed, then

$$\chi_\pi(x) = \mathbf{1}_{\{\delta_\pi^{\text{arch}}(x) > 0\}}.$$

Proof. If $x \in \text{Im}(\pi)$, some m satisfies $\pi(m) = x$, so the infimum is zero. Conversely, if the infimum is zero and $\text{Im}(\pi)$ is closed, then x lies in $\text{Im}(\pi)$. The stated equality follows from the definition of χ_π . $\square$

This is a real result, but a narrower one than it may first appear: Claim 1.5.4 shows the binary gate is the zero/nonzero threshold of a graded quantity, δ_π^{arch} , built from the same kind of infimum-distance construction used throughout this chapter. It does *not* identify δ_π^{arch} with this chapter’s coding deficit δ_T . A full compatibility theorem would require a map $J : T \mapsto \pi_T$ from ontologies to prediction maps and constants $c_1, c_2 > 0$ such that

$$c_1 \delta_{\pi_T}^{\text{arch}}(x) \leq \delta_T(x) \leq c_2 \delta_{\pi_T}^{\text{arch}}(x),$$

or showing it fails, is left as an explicit open problem rather than assumed by continuing to use both vocabularies side by side in later chapters.

Remark 1.5.5 (Preserving distinctions is not preserving competence). *Nothing developed so far should suggest that a distinction-preserving transformation leaves a system’s operational capacity unchanged. A bijective relabeling of X — not a projection, not a quotient, not a fiber collapse, and therefore invisible to δ_T , Δ_T , and every entropy measure built from them in later chapters — can still destroy the alignment between a learned operator repertoire and the space that repertoire was trained to act on. Representational preservation and operational continuity are independent properties: a transformation can hold the first perfectly while destroying the second completely.¹ This distinction reappears, worked through in full, in Chapter 5.*

¹A mirror keyboard — a piano relabeled so pitch runs high to low rather than low to high — is the cleanest instance: every distinction in the note space survives exactly, yet a trained pianist’s fingering competence can be rendered nearly unusable.

Summary

This chapter develops distinctions and the two deficits they support, along with the four operations that act on distinguishability structure and a first, partial connection to admissibility. It does not develop the further step of showing how distinctions induce the partitions that information theory operates over, nor state an entropy functional on the resulting distinguishability space. The connection is real — a distinguishability relation $\sim$ induces a partition of X , and a partition supports an entropy measure in the ordinary way — and connecting the resulting entropy back to the coding deficit of definition 1.3.1, rather than treating them as two unrelated quantities that happen to share units, is developed in Chapter 2.

Chapter 2

Information

2.1 Introduction

Chapter 1 closed by naming two things it had not yet done. It had not shown how the distinguishability relation $\sim$ on a space $(X, \sim)$ induces the partition structure that information theory presupposes, and it had not connected the resulting information measures back to the coding deficit δ_T of definition 2.1.1 rather than treating them as two quantities that happen to be measured in the same units. This chapter does both, and is checked against Chapter 1's actual theorems rather than merely gestured at them.

Definition 2.1.1 (Coding deficit). *For an ontology T and object x , the coding deficit is $\delta_T(x) = L_T(x) - L^*(x)$, the excess of the shortest description of x reachable within T over the shortest description achievable by any ontology.*

2.2 Partitions and Surprisal

A distinguishability relation $\sim$ on X is, by definition 2.2.1, equivalently a partition of X into fibers.

Definition 2.2.1 (Distinguishability space). *A distinguishability space is a pair $(X, \sim)$ where $\sim$ is an equivalence relation on X , interpreted as observational indistinguishability; its equivalence classes are fibers.*

Suppose an observer knows only that a system lies somewhere within a partition $\mathcal{P} = \{P_1, \dots, P_n\}$ induced by $\sim$, and an observation identifies one cell P_i with probability p_i . The informational content of that event is the Shannon surprisal $I(P_i) = -\log p_i$. Read distinction-theoretically rather than purely statistically, this is not an imported foreign quantity: it measures how strongly the

observation refines the distinction structure already given by $\mathcal{P}$. An observation that identifies a rare cell refines more than an observation that identifies a common one, exactly as a finer $\sim$ expresses more distinctions than a coarser one. The Shannon entropy of the partition, $H(\mathcal{P}) = -\sum_i p_i \log p_i$, is the expected surprisal — the expected refinement an observation of this system delivers.

2.3 The Compression Theorem

Claim 2.3.1 (Compression Theorem). *For a surjective projection $\pi : X \rightarrow Y$, $H(Y) \leq H(X)$.*

The argument is immediate given Chapter 1’s own Fiber Monotonicity result: a surjective map merges fibers, the induced partition on Y is coarser than the one on X , and a coarser partition can express no more distinctions than a finer one. Entropy inherits the same monotonicity that fiber containment already established. Every representation compresses; every model compresses; every theory compresses. The question this chapter inherits from Chapter 1 is never whether compression occurs but which distinctions survive it.

2.4 Representational Entropy

Compression’s effect on a single point, rather than on a whole partition, is worth isolating as its own quantity.

Definition 2.4.1 (Representational entropy). *Let $\pi : X \rightarrow M$ be a representation. The representational entropy at $m \in M$ is $S_\pi(m) = \log |\pi^{-1}(m)|$, the log-size of the fiber collapsed beneath m .*

Large $S_\pi(m)$ indicates many underlying states rendered indistinguishable by the representation: high compression at m , with a correspondingly large blind spot.

Claim 2.4.2 (Representational entropy bounds the coding-deficit increase). *For a projection $\sigma : X \rightarrow S$ with fiber $F_s = \sigma^{-1}(s)$, Chapter 1’s Projection Deficit Bound states $\Delta\delta(\sigma) \geq \log |F_s|$. Since $\log |F_s| = S_\sigma(s)$ by definition 2.4.1 applied to the same fiber, this is exactly*

$$\Delta\delta(\sigma) \geq S_\sigma(s).$$

This is the bridge Chapter 1 was owed, and it is a checked identity of quantities rather than an assumed correspondence: both $\log |F_s|$ in Chapter 1’s bound and $S_\pi(m)$ in definition 2.4.1 are defined as the log-cardinality of the same fiber,

so Claim 2.4.2 follows by substitution, not by analogy. It licenses treating representational entropy as a lower bound on how much coding deficit a projection is forced to introduce, and treating a large deficit increase under a specific projection as evidence, though not proof, of large representational entropy at the corresponding point.

Remark 2.4.3 (What this bridge does not yet establish). *Claim 2.4.2 relates δ_T 's behavior under one specific operation, projection, to S_π at a single point. It does not relate the two deficits of Chapter 1 — δ_T and Δ_T — to $H(\mathcal{P})$ in general, and it says nothing yet about revision or transport, the two ontology-level operations of Chapter 1. Those connections, if they exist, are not established here.*

2.5 The Representation Limit Theorem

Claim 2.5.1 (Representation Limit Theorem). *No finite representation can preserve every distinction of an infinite domain: for $|X| = \infty$ and $|M| < \infty$, any $\pi : X \rightarrow M$ assigns infinitely many elements of X to at least one fiber, by the pigeonhole principle, and so collapses infinitely many distinctions.*

This is the first formal appearance of what recurs throughout this book under the name projection failure. Loss under a bounded representation of an unbounded domain is not an implementation defect to be engineered away; it is a mathematical necessity following directly from cardinality alone, independent of how cleverly the representation is designed.

2.6 The Fundamental Information Theorem

Claim 2.6.1 (Fundamental Information Theorem). *Every informational quantity can be expressed as a property of a distinction structure: information requires distinguishable alternatives, which define a partition; probability distributions are defined over partition cells; entropy is computed from those probabilities; mutual information measures shared refinement between partitions; compression measures distinction collapse; channel capacity measures transmissible distinctions. Each depends solely on partition structure.*

Remark 2.6.2 (The inverted hierarchy). *The traditional hierarchy places objects first and derives information from them. Claim 2.6.1 inverts this: distinctions precede information, and information precedes objects. An object, on this reading, is an informational artifact generated by a distinction-preserving compression — not a primitive that information is subsequently extracted from.*

Summary

This chapter develops information as a Shannon-style and representational quantity: surprisal, partition entropy, and representational entropy, all grounded directly in the distinguishability apparatus of Chapter 1 and checked against it rather than merely motivated by analogy to it. It does not touch the entropy used to define admissibility: the monotonic, thermodynamic-style field S satisfying $\dot{S} \geq 0$ along admissible trajectories in the RSVP formalization. Whether that S coincides with, is derived from, or is genuinely independent of the $H(\mathcal{P})$ and $S_\pi(m)$ developed here is addressed in Chapter 4, which also checks whether “entropy” is being used consistently across every context it appears in, rather than assuming any two uses coincide because they share a word.

This chapter has nothing to say about a chapter titled “Constraint”: no independent technical substrate for such a chapter, distinct from admissibility, has been found. Admissibility is the working technical term throughout the relevant material; “constraint” is never operationalized there as a separate object. That chapter is accordingly absent from this volume.

Chapter 3

Admissibility

3.1 Two Definitions

An admissibility relation specifies which transitions between states are reachable from a given state. Two independent constructions give this relation formal content.

Definition 3.1.1 (Coding deficit). *For an ontology T and object x , the coding deficit is $\delta_T(x) = L_T(x) - L^*(x)$, the excess of the shortest description of x reachable within T over the shortest description achievable by any ontology.*

Definition 3.1.2 (Deficit-based admissibility). *For an ontology T , a state x is inadmissible from T exactly when $\delta_T(x) > 0$.*

Definition 3.1.3 (Entropy-monotonicity admissibility). *For a field triple $(\Phi, \mathbf{v}, S)$ over configuration space $\mathcal{F}$, a configuration γ_2 is admissible from γ_1 when there is a continuous path in $\mathcal{F}$ connecting them along which the entropy density S is non-decreasing, $\dot{S} \geq 0$, and regularity holds. This is a directed reachability relation, a partial order compatible with entropy monotonicity.*

Both definitions describe the same underlying idea: which states can be reached from which. That shared ambition does not by itself make them the same relation.

3.2 Testing the Identification

Claim 3.2.1 (Type compatibility). *definitions 3.1.2 and 3.1.3 pass the type test: both are relations on pairs of states, or equivalently on paths between them, not objects of different kinds.*

Passing the type test is necessary but not sufficient. The stronger requirement is a definition-by-substitution test: can one definition be obtained from the other by specializing variables, the way historical entropy is representational entropy evaluated at a particular projection.

A schema general enough to compare the two precisely treats admissibility as membership in a continuation family.

Definition 3.2.2 (Continuation family and trajectory-relative admissibility). *A continuation family assigns to each state x a set $C(x)$ of trajectories beginning at x considered permitted. A state x_i occurring within a trajectory τ is admissible relative to τ and C when the suffix of τ beginning at x_i is a member of $C(x_i)$.*

Proposition 3.2.3 (Both definitions are continuation families). *$C_S(x) = \{\tau : \dot{S}(\tau) \geq 0\}$ is the continuation family induced by definition 3.1.3: checking membership requires nothing beyond evaluating a static condition at each instant. $C_\delta(x) = \{\tau : \delta_T(\tau) \leq \epsilon\}$, for a threshold ϵ , is the continuation family induced by definition 3.1.2, built from a distortion bound rather than an entropy condition. Recognizing both as instances of definition 3.2.2 does not identify them — different constructions of $C(x)$ remain different constructions — but replaces a comparison between two differently-typed-looking objects with a question about two set-valued constructions of the same type.*

Open Problem 3.2.4 (Central open problem). *The relevant question is not whether $\delta_T(x) = 0$ is equivalent to the existence of an entropy-nondecreasing path to x , stated bare. It is whether the induced continuation families coincide, nest, or merely intersect:*

$$C_\delta(x) \stackrel{?}{=} C_S(x), \quad C_\delta(x) \stackrel{?}{\subseteq} C_S(x), \quad C_\delta(x) \cap C_S(x) \stackrel{?}{\neq} \emptyset.$$

This is sharper and more tractable than a bare quantity comparison, since it asks about set relations between two well-typed objects rather than an equality between two differently-motivated scalars. Every later use of admissibility in this volume — the admissible future region $\Gamma_t(x)$, the architecture gate χ_π , and any treatment of continuation or coordination — inherits whichever answer this eventually receives. One structural correspondence is worth noting without being upgraded to a proof: the Second Law's $\dot{D} \leq 0$ and the entropy condition's $\dot{S} \geq 0$ are the same monotonic inequality under a sign flip, suggestive of S playing the role of consumed distinction capacity, which would suggest $C_\delta(x) \subseteq C_S(x)$ if it held. Suggestive is the correct word; nothing here upgrades it to a proof.

Remark 3.2.5 (On resolving this prematurely). *A framework this dependent on admissibility has an obvious incentive to declare open Problem 3.2.4 resolved as soon as a plausible bridge appears. That incentive is a reason for caution, not license. Resolving it requires an explicit map with checked functoriality between C_δ and C_S , not the convenience of having many prior results click into a single coherent structure once it is assumed.*

3.2.1 A Structural Reason for Skepticism

A general fact about conserved quantities gives this skepticism a formal basis, together with a caveat about exactly how far it reaches.

Theorem 3.2.6 (Constants of motion are generically non-injective on trajectory space). *Let F be a constant of motion along a trajectory ($\dot{F} = 0$ identically) and let $\mathcal{T}$ be the space of dynamically realizable trajectories. The level sets of F generically contain more than one distinct trajectory: except in fully integrable systems with as many independent constants of motion as degrees of freedom, knowing the value of F alone does not determine which trajectory, among those sharing that value, is realized.*

This is close to a restatement of ordinary partial integrability, but its consequence here is not incidental. A history channel earns its keep, in the sense δ_T and the distinguishability apparatus of Chapter 1 both depend on, by being *injective* over a fiber — separating predecessors that would otherwise be conflated. A constant of motion is, by construction, the opposite: it is valuable to physics precisely because it is shared by many different trajectories, which is what makes it a useful invariant rather than a complete description of the motion. Distinguishing and conserving are not two routes to the same destination; they are built to do opposite jobs.

Remark 3.2.7 (Scope of the theorem). *theorem 3.2.6 is stated for a strict constant of motion, $\dot{F} = 0$. The entropy condition entering definition 3.1.3 is monotonic, $\dot{S} \geq 0$, not conserved, and the theorem as proved does not directly cover it. The structural point plausibly extends — fixing “ S is non-decreasing” as a constraint on a trajectory is a coarse condition compatible with a large multiplicity of distinct trajectories, in the same spirit as a fixed level set — but that extension is not proved here. What can be said without overreach is that S belongs to the family of quantities theorem 3.2.6 concerns — built to be shared across trajectories, not to separate them — and that this family has been shown, at least in the conserved case, to be structurally unsuited to the job δ_T does.*

A related question in classical mechanics sharpens this further. A first pass might claim that no treatment of classical mechanics asks anything resembling a distinguishability question; constrained Hamiltonian systems are the counterexample, since a singular Legendre transform there produces genuinely non-injective structure, formalized in Dirac–Bergmann constraint theory. Physics’s own unresolved *problem of time* in general relativity is, in substance, the question of whether a distinction collapsed by a gauge constraint carries physical content, and the literature has not settled it after decades of direct engagement. The one identified formal difference is that the standard physics criterion (the Dirac-observable condition) is context-free and algebraic, applied once and for

all, while admissibility-relevance here is explicitly relative to a continuation family, so that the same fiber can matter under one continuation family and not another.

Remark 3.2.8 (A speculative reframing). *It is tempting, offered here only as speculation, to wonder whether open Problem 3.2.4 has resisted resolution for a structurally similar reason. If δ_T 's admissibility is inherently relative to an ontology and a continuation family, while the entropy condition is posed, at least on its surface, as a context-free field condition, then demanding an equivalence between them may be attempting the same kind of category error the problem of time keeps encountering: seeking a single, context-free answer to a question that may only admit context-relative ones. This reframes why open Problem 3.2.4 is hard without resolving it, and carries no more authority here than the underlying dispute in the foundations of general relativity carries elsewhere. A reframing of a difficulty is useful even when it settles nothing.*

3.3 The Architecture Gate as a Special Case

A partial, weaker bridge survives independently of open Problem 3.2.4's resolution.

Proposition 3.3.1 (Gate as thresholded architectural deficit). *For a prediction map $\pi : \mathcal{M} \rightarrow \Sigma$ with closed image, the binary architecture gate satisfies $\chi_\pi(x) = \mathbf{1}_{\{\delta_\pi^{\text{arch}}(x) > 0\}}$ for the architectural representation deficit $\delta_\pi^{\text{arch}}(x) = \inf_{m \in \mathcal{M}} d_\Sigma(\pi(m), x)$.*

This shows χ_π is a coarsening of a graded quantity built the same way δ_T is built — an infimum-distance construction — without showing δ_π^{arch} and δ_T are the same quantity. proposition 3.3.1 and definition 3.1.2 are siblings under the same construction pattern, not proven identical. Admissibility notions throughout this volume are repeatedly built as either a deficit (an infimum distance from a target set) or a monotonicity condition (a non-decreasing quantity along a path), and it remains an open question whether these are two families or one family with two descriptions.

3.4 Continuation Needs No New Machinery

The clearest positive result available is not an identification of definitions 3.1.2 and 3.1.3 with each other, but a reduction internal to definition 3.1.3 alone.

Theorem 3.4.1 (Trajectory-admissibility reduces to static admissibility, checked pointwise). *For a trajectory $x(t)$ and a coordination geometry $\mathcal{G}(t)$ with CLIO projection $\pi_{\mathcal{G}}$, “ $x(t)$ is continuation-admissible” is equivalent to $S_{\text{RSVP}}(\pi_{\mathcal{G}}(x(t))) \leq 0$ holding*

at every t : the same static condition of definition 3.1.3, evaluated at each moment along the path, introduces no further structure. When this condition is violated, the coordination geometry reorganizes — geometric repair — until it is restored.

Theorem 3.4.1 is proved by direct construction, not asserted by resemblance. Its consequence for continuation is specific. A theory of continuation does not need to introduce a second notion of admissibility, or a new axiom governing how admissibility changes over time; it needs only to define trajectories properly, use theorem 3.4.1, and examine what its violation condition — geometric repair — has in common with operator ecologies (Chapter 5), since both describe a system reorganizing itself to restore a condition that has been violated. That comparison, not new foundational machinery, is what a theory of continuation is actually for.

Remark 3.4.2 (This also explains the admissible future region). *The admissible future region $\Gamma_t(x) = \overline{\mathcal{R}}_t(x)$ of Chapter 5 is not a third notion of admissibility. It is definition 3.1.3’s reachability relation, restricted to the effective (quotiented) state space rather than the full state space, for the specific case where entropy production is realized as distinction coarsening. This confirms that the inventory of admissibility notions in this volume — deficit-based, entropy-based, and the architecture gate — is complete, with $\Gamma_t(x)$ falling under the second rather than constituting a fourth.*

Summary

Two independent constructions of admissibility — a deficit-based relation and an entropy-monotonicity condition — pass a type test but not a substitution test; whether they induce the same continuation family is the volume’s central open problem, sharpened here from a bare quantity comparison into a question about set relations between two well-typed objects, and left open. The architecture gate is shown to be a coarsening of a deficit-style construction without being identified with the specific deficit of Chapter 1. Continuation requires no independent foundational treatment beyond the material developed here, applied pointwise.

Chapter 4

Constraint Populations

4.1 From a Single Region to a Population

The admissibility relations developed earlier in this book fix a single region — a viability manifold $A \subseteq X$ — and ask whether a trajectory remains inside it. This chapter considers instead a *population* of admissibility constraints $\chi_1, \dots, \chi_n$, each χ_i specifying its own admissible region $A_i \subseteq X$, with the population's joint admissible region given by

$$A = \bigcap_{i=1}^n A_i.$$

Nothing about this construction is new on its own; intersecting constraints is the ordinary way multiple requirements combine, and each χ_i individually is admissibility-typed in exactly the sense already developed — it fails the type test to call it a distinct primitive. What is new is treating the χ_i themselves as objects with their own histories and dynamics rather than as fixed, given sets: a building code that is amended, a type system that is extended, a norm that shifts. Each χ_i is a compressed operator $\chi_i = C(H_i)$ arising from some history of prior cases, and each is subject to its own recursive modification, $\chi_i(t+1) = G_i(\chi_i(t), \dots)$. The question this chapter asks is what happens once that update rule is allowed to depend on the state of the *other* constraints in the population, not only on χ_i 's own trajectory.

4.2 Constraint Populations Are Sympoietic Systems

Definition 4.2.1 (Coupled recursive update). *A collection of objects $o_1, \dots, o_n$ exhibits coupled recursive update when each o_i 's own update rule at time $t+1$ depends not only on o_i 's own history but on the joint state of the whole collection: $o_i(t+1) =$*

$J_i(o_i(t), E_{t+1})$, where $E_{t+1} = H(E_t, o_1(t), \dots, o_n(t))$ aggregates the current state of every object in the collection. This is a Level-3 phenomenon in the sense that the rule governing o_i 's future behavior, not merely o_i 's state, is the object being modified — and it is a genuinely collective one: the population as a whole can exhibit coupled recursive update even where no individual o_i , considered in isolation from the others, does.

Proposition 4.2.2 (Constraint populations are sympoietic). *Let $\chi_1, \dots, \chi_n$ be a constraint population sharing a common state space X , with each χ_i 's update depending on the joint state of the population in the sense of definition 4.2.1: $\chi_i(t+1) = J_i(\chi_i(t), E_{t+1})$, $E_{t+1} = H(E_t, \chi_1(t), \dots, \chi_n(t))$. Then the population exhibits coupled recursive update at the collective level directly, by definition 4.2.1 applied with $o_i = \chi_i$.*

Proof. Immediate from definition 4.2.1: the population's update structure is stated in exactly that definition's coupled form, with χ_i substituted for the general object o_i . $\square$

Remark 4.2.3 (The same pattern recurs for operator ecologies). Chapter 6 develops this same coupling structure again, independently, for a different population: the operators of an operator ecology, each compressed from its own history and each capable of having its update depend on the state of the ecology around it. Nothing in that later development is assumed here, and nothing here is assumed there; the two are separate instances of definition 4.2.1; a reader who meets Chapter 6's version after this one will recognize the same shape appearing at a different scale, with operators playing the role constraints play here.

Remark 4.2.4 (Why this costs nothing new to state). A body of statutory law that shifts only through the accumulated, individually modest effect of many separate judicial rulings, no one of which represents or intends to modify the law's overall shape, is a direct instance: each ruling is a local, low-level update, and the shape of the law that emerges is a collective Level-3 phenomenon, with case law standing in for the aggregate of individually unremarkable local changes. What was informally described as one constraint's violation being "absorbed" by neighboring constraints tightening or loosening elsewhere is simply proposition 4.2.2 in operation: χ_i 's update rule J_i is permitted to respond to the joint state E_{t+1} , which encodes the current state of every other constraint in the population, not merely χ_i 's own history. No new mechanism is required beyond the coupling already built into definition 4.2.1.

4.3 Two Independent Pathologies

Definition 4.3.1 (Feasibility). *A constraint population is feasible if $A = \bigcap_i A_i \neq \emptyset$.*

Definition 4.3.2 (Constraint diversity). *Let $\bar{\rho} \in [0, 1]$ denote the average pairwise correlation among the χ_i 's failure events — the frequency with which they are jointly*

violated by the same states — and define

$$D_\chi = \frac{n}{1 + (n-1)\bar{\rho}}.$$

Corollary 4.3.3 (Constraint monoculture). *As $\bar{\rho} \rightarrow 1$, $D_\chi \rightarrow 1$ regardless of n : a single failure mode that defeats one χ_i defeats the population as a whole, and many nominally distinct constraints provide no more protection against inadmissible drift than a single constraint would.*

Remark 4.3.4 (Against a folk intuition). *Ordinary usage suggests that constraints cooperating — pulling in the same direction, reinforcing one another — is the healthy condition, and constraints competing is the pathological one. This is not quite right, and in one respect is backwards. Constraints always violated by the same states and always satisfied by the same states are exactly the $\bar{\rho} \rightarrow 1$ case: however many there are, they function as one constraint wearing many names, and the population inherits the fragility of a single point of failure. What ordinary usage calls competition — constraints binding on different states, ruling out different failure modes — is closer to the $\bar{\rho} \rightarrow 0$ case identified as healthy diversity. Genuine pathology is not disagreement between constraints; it is either infeasibility or redundancy dressed up as multiplicity.*

4.4 Feasibility Constrains Measured Diversity — Sometimes Perversely

Lemma 4.4.1 (Two-constraint infeasibility bound). *Let χ_1, χ_2 be a population of two constraints on a finite state space, with $\bar{\rho}$ their pairwise correlation computed as the standard ϕ coefficient. If the population is infeasible, $\bar{\rho} \leq 0$.*

Proof. Infeasibility means the contingency-table cell $n_{00} = 0$. The ϕ coefficient is $\phi = (n_{11}n_{00} - n_{10}n_{01})/\sqrt{n_{1\cdot}n_{0\cdot}n_{\cdot 1}n_{\cdot 0}}$; with $n_{00} = 0$ the numerator reduces to $-n_{10}n_{01} \leq 0$, and the denominator is non-negative. $\square$

Two genuinely incompatible constraints can never register as redundant under this measure: infeasibility, at $n = 2$, structurally pushes measured diversity up, not down. Whatever pathology infeasibility represents, it is not the pathology of redundancy, and the two cannot be conflated even in principle at this population size. The next result shows this protection does not survive a third constraint.

Proposition 4.4.2 (Diversity can mask redundancy in the presence of conflict). *There exists a population of four constraints, jointly infeasible, in which three of the four are pairwise identical, and yet the population's average pairwise correlation is exactly $\bar{\rho} = 0$, giving $D_\chi = n = 4$: the score a population of four genuinely independent constraints would report.*

Proof. Let $X = \{1, \dots, 8\}$. Let $\chi_1 = \chi_2 = \chi_3$ each admit $\{1, 2, 3, 4\}$ (violated exactly on $\{5, 6, 7, 8\}$), and let χ_4 admit $\{5, 6, 7, 8\}$ (violated exactly on $\{1, 2, 3, 4\}$). The joint admissible set is $\{1, 2, 3, 4\} \cap \{5, 6, 7, 8\} = \emptyset$: the population is infeasible. Among the $\binom{4}{2} = 6$ pairs, the three pairs within $\{\chi_1, \chi_2, \chi_3\}$ have violation indicators agreeing everywhere, giving $\phi = +1$ each. The three pairs crossing to χ_4 have violation indicators that are exact complements, giving $\phi = -1$ each. The average over all six pairs is $\bar{\rho} = (3 \cdot (+1) + 3 \cdot (-1))/6 = 0$ exactly, giving $D_\chi = 4/(1 + 3 \cdot 0) = 4$. $\square$

This is the chapter's central finding. The positive correlation contributed by real redundancy and the negative correlation contributed by genuine conflict cancel in a simple average, and the resulting scalar D_χ reports the population as maximally diverse despite three-quarters of it being a single constraint wearing three names. Population health cannot be read off D_χ alone whenever redundancy and conflict are simultaneously present: lemma 4.4.1 and proposition 4.4.2 together show the measure requires a feasibility check alongside it, not merely as an independent second axis, but because feasibility failures can actively corrupt what the diversity score appears to say once three or more constraints are in play.

Proposition 4.4.3 (Feasible populations span both high and low diversity). *Feasibility does not determine D_χ : there exist feasible populations with D_χ well above n and feasible populations with D_χ at its minimum of 1.*

Proof. On $X = \{1, 2, 3, 4\}$: let χ_1 admit $\{1, 2, 3\}$ and χ_2 admit $\{2, 3, 4\}$. Then $A = \{2, 3\} \neq \emptyset$, and the violation indicators give $\phi = -1/3$, so $D_\chi = 2/(1 - 1/3) = 3$, above $n = 2$. By contrast, let $\chi_1 = \chi_2$ both admit $\{1, 2, 3\}$: $A = \{1, 2, 3\} \neq \emptyset$, but the two constraints are violated by exactly the same state in every case, giving $\phi = 1$ and $D_\chi = 1$ despite $n = 2$. $\square$

Remark 4.4.4 (D_χ has no finite upper bound). *Unlike the impression the equal-and-opposite cancellation of proposition 4.4.2 might give, $D_\chi = n/(1 + (n - 1)\bar{\rho})$ is not bounded above by n , or by any finite value, once $\bar{\rho}$ is allowed to go negative: as $\bar{\rho} \rightarrow -1/(n - 1)$ from above, the denominator approaches 0^+ and $D_\chi \rightarrow \infty$. proposition 4.4.3's own example already exceeds n . The value $D_\chi = n$ obtained in proposition 4.4.2 is not a ceiling being reached; it is the specific value $\bar{\rho} = 0$ happens to produce, arising there from an exact cancellation between the redundant trio's positive correlation and the conflicting pair's negative correlation. Nothing about the measure forbids scores above n , and a reader computing examples independently should expect to encounter them.*

Feasibility must be checked directly, $A \neq \emptyset$, and diversity must be checked directly, via $\bar{\rho}$ computed with attention to whether conflicting and redundant

subgroups might be present simultaneously. Neither can be inferred from the other, and in populations of three or more constraints, a favorable diversity score cannot be trusted at face value without also checking whether it is the product of exactly this kind of cancellation.

4.5 Marginal Diversity Contribution

proposition 4.4.2 shows a population-level score can be misleading. It does not yet say which member of the population is responsible. This section supplies that.

Definition 4.5.1 (Marginal diversity contribution). *For a constraint population $C = \{\chi_1, \dots, \chi_n\}$ and $c \in C$, the marginal diversity contribution of c is*

$$\Delta_D(c) = D_\chi(C) - D_\chi(C \setminus \{c\}),$$

where $D_\chi(C \setminus \{c\})$ is computed over the remaining $n - 1$ constraints, with $\bar{\rho}(C \setminus \{c\})$ recomputed as the average pairwise correlation among only those $n - 1$ constraints — excluding every pair that involved c , not merely subtracting c 's own contribution from the original average.

Definition 4.5.2 (Classification). *c is diversity-positive if $\Delta_D(c) > 0$ (removing c would decrease measured diversity: c is contributing genuine coverage), diversity-neutral if $\Delta_D(c) = 0$, and diversity-negative, or invasive, if $\Delta_D(c) < 0$ (removing c would increase measured diversity: c 's presence is actively suppressing the population's diversity score, typically because it duplicates coverage another constraint already provides).*

Corollary 4.5.3 (The masking population, decomposed). *For the population of proposition 4.4.2, $\Delta_D(\chi_1) = -5$ and $\Delta_D(\chi_4) = +3$.*

Proof. $D_\chi(C) = 4$ by proposition 4.4.2. Removing χ_1 leaves $\{\chi_2, \chi_3, \chi_4\}$: the surviving pair (χ_2, χ_3) has $\phi = +1$, and the two pairs crossing to χ_4 each have $\phi = -1$, giving $\bar{\rho} = (1 - 1 - 1)/3 = -1/3$ and $D_\chi(C \setminus \{\chi_1\}) = 3/(1 - 2/3) = 9$, so $\Delta_D(\chi_1) = 4 - 9 = -5$. Removing χ_4 leaves $\{\chi_1, \chi_2, \chi_3\}$, all three pairs at $\phi = +1$, giving $\bar{\rho} = 1$ and $D_\chi(C \setminus \{\chi_4\}) = 3/(1 + 2) = 1$, so $\Delta_D(\chi_4) = 4 - 1 = 3$. $\square$

The population-level score $D_\chi(C) = 4$ gives no reason to suspect anything is wrong. The marginal decomposition immediately does: three of the four constraints are invasive by definition 4.5.2, each with $\Delta_D = -5$ by the symmetry of the construction, and the one constraint actually responsible for the population's apparent diversity, χ_4 , is identified directly, with a marginal contribution larger than its share of the headcount would suggest. This is the masking pathology,

resolved rather than merely diagnosed: reading $D_\chi(C)$ alone cannot distinguish this population from four genuinely independent constraints, but reading every $\Delta_D(c)$ can.

Remark 4.5.4 (The same construction predicts, rather than only diagnoses). *definition 4.5.1 is stated for c already a member of C , answering which existing members are worth keeping. The same computation, applied to a candidate c not yet in C by evaluating $\Delta_D(c)$ within $C' = C \cup \{c\}$, answers a different question with the same arithmetic: whether introducing c would help, do nothing, or hurt. $\Delta_D(c)$ computed in $C \cup \{c\}$ is exactly $D_\chi(C \cup \{c\}) - D_\chi(C)$, since removing c from $C \cup \{c\}$ recovers C — the pruning question and the introduction question are the same construction viewed from opposite directions, not two separate tools.*

4.6 Worked Examples

Constitutional and statutory law. A constitution changes slowly and sets the outer admissible region within which statutes, which change quickly, further narrow the admissible space. This is a two-timescale sympoietic pair in the sense of proposition 4.2.2: statutory revision is the fast, individually modest update; constitutional interpretation shifts more slowly and in response to the accumulated pattern of statutory and judicial activity.

Building codes and inspection regimes. A code is a constraint; an inspection regime is a verification channel. Multiple inspectors trained identically, using the same checklist and assumptions, add little effective diversity however many of them there are, while inspectors drawing on independent training or independent professional traditions raise D_χ meaningfully even at the same headcount.

Type systems and runtime checks. A static type system and a runtime assertion checking the same property are a low- D_χ pair: both fail together on the same category of error. A static type system paired with a runtime check for a genuinely different class of failure is a high- D_χ pair, catching different failure modes at different times.

Gene regulatory networks. Transcription factor networks routinely exhibit precisely this structure: feedback loops in which one regulatory constraint's activation state governs another's threshold, a pattern recurring across organisms because sympoietic constraint coupling of this kind is a repeatedly re-discovered solution rather than a special case.

4.7 Connections to the Rest of This Volume

This chapter sits at the intersection of three threads, two of which lie ahead rather than behind it. Admissibility, developed in the preceding chapter, supplies the single-constraint apparatus — viability manifolds and the boundary within which structure becomes possible — that this chapter extends from one constraint to a population of them. The coupled-recursive-update structure of definition 4.2.1 anticipates Chapter 6’s own treatment of operator ecologies, where the same pattern recurs with operators as the coupled objects rather than constraints, and where the compression apparatus underlying $\chi_i = C(H_i)$ here reappears as Chapter 6’s treatment of operators as compressed histories. Chapter 8’s coordination geometry, concerned with how a single geometry reorganizes to restore a violated admissibility condition, will turn out to be a special case of what this chapter develops: a coordination geometry is a constraint population of size one, and geometric repair is what coupled recursive update looks like when $n = 1$. None of these connections is claimed as a theorem belonging to the later chapters specifically; they are recorded here as the shape a fuller unification would need to take, to be confirmed or revised once those chapters exist.

Open Problems

Open Problem 4.7.1 (Ecological stability and extinction). *Under what conditions does a constraint in a population become effectively vacuous — never binding, because other constraints already exclude every state it would have ruled out — and should a vacuous constraint be understood as having gone extinct in the sense Chapter 6’s treatment of operator ecologies will later give that word?*

Remark 4.7.2 (Invasive constraints, resolved). *The criterion this population lacked when first raised is definition 4.5.1 applied via remark 4.5.4: whether a candidate χ_{n+1} will be absorbed, be diversity-neutral, or destabilize the population’s diversity score is $\Delta_D(\chi_{n+1})$ computed within $C \cup \{\chi_{n+1}\}$, positive, zero, or negative respectively — computable before introducing the constraint from its pairwise correlations with each existing member alone. What remains genuinely open is narrower: rejection (infeasibility) is a separate failure mode from destabilization (a negative but still feasible Δ_D), and nothing here characterizes which candidates produce which of the two, only that diversity impact and feasibility impact must both be checked, neither implying the other, in keeping with proposition 4.4.3.*

Open Problem 4.7.3 (Dynamics of J_i). *proposition 4.2.2 assumes each constraint’s update rule J_i exists without specifying its form. A general theory of how aggressively a constraint should tighten or loosen in response to violations elsewhere is left to future work.*

Summary

A population of admissibility regions is not adequately described by its size, nor by a single scalar diversity score taken at face value. Feasibility and diversity are independent axes that must each be checked directly; for populations of three or more, they can actively corrupt each other's apparent reading, with genuine redundancy and genuine conflict cancelling to report a maximally diverse population that is mostly one constraint under several names. This is not a restatement of single-region admissibility; it is what single-region admissibility looks like once more than one region is held simultaneously, coupled through exactly the sympoietic machinery Chapter 6 already built for a different purpose.

Chapter 5

Entropy

5.1 Introduction: Five Candidates, Not One

Chapter 2 developed two quantities under headings that both, in ordinary usage, fall under “entropy”: the Shannon entropy of a partition and the representational entropy of a projection fiber. Whether either coincides with the entropy used to define admissibility is not assumed here. Widening the search turns up at least three further candidates, for five in total.

- A. **Shannon/distinction entropy**, $H(\mathcal{P}) = -\sum_i p_i \log p_i$, developed in section 5.2.
- B. **Representational entropy**, $S_\pi(m) = \log |\pi^{-1}(m)|$, developed in section 5.3.
- C. **Historical entropy**, the multiplicity of histories collapsing to a single state under a history-to-state projection, examined in section 5.5.
- D. **Repair entropy**, $S_R(F \mid T)$, the diagnostic uncertainty over fault locations given a retained trace, already used without full justification in Chapter 6, examined in section 5.7.
- E. **The RSVP entropy field**, $S(x, t)$, entering the admissibility relation of the RSVP formalization via the monotonicity constraint $\dot{S} \geq 0$, examined in section 5.8.

The goal of this chapter is to determine, for each pair, whether the relationship is identity, derivation, structural analogy, or genuine independence — and to say plainly which of these has actually been checked against a proof and which remains a resemblance.

5.2 Entropy A: Shannon/Distinction Entropy

For a partition $\mathcal{P} = \{P_1, \dots, P_n\}$ with cell probabilities p_i , $H(\mathcal{P}) = -\sum_i p_i \log p_i$ is the expected surprisal of an observation identifying a cell, interpreted distinction-theoretically as the expected refinement of the distinction structure that an observation delivers.

5.3 Entropy B: Representational Entropy

For a representation $\pi : X \rightarrow M$, the representational entropy at m is $S_\pi(m) = \log |\pi^{-1}(m)|$, the log-size of the fiber collapsed beneath m , already shown in Chapter 2 to bound Chapter 1's coding-deficit increase from below under projection.

5.4 The Distinction–Entropy Duality: A and B Are Not Independent

Claim 5.4.1 (Distinction–Entropy Duality). *Every distinction simultaneously produces information and entropy. Information measures separation between partition cells; entropy measures multiplicity within partition cells. For a partition $\mathcal{P}$ with $N_D = |\mathcal{P}|$ cells, $H = \log N_D$; for the multiplicity $\Omega(P_i) = |P_i|$ within a cell, $S(P_i) = k \log \Omega(P_i)$ for a constant k . Both arise from the same partition, viewed from opposite directions: refinement increases N_D while reducing average $\Omega(P_i)$, and coarsening does the reverse.*

Remark 5.4.2 (This is a reduction, not a new object). $S(P_i) = k \log \Omega(P_i)$ in Claim 5.4.1 has exactly the form of section 5.3's representational entropy: $\Omega(P_i) = |P_i|$ is the cardinality of a partition cell, and a partition cell is a fiber of the quotient map $q : X \rightarrow X/\sim$ in the sense of Chapter 1. Up to the constant k , $S(P_i)$ is $S_q(P_i)$. The Distinction–Entropy Duality is therefore not introducing a third independent quantity alongside A and B; it is the statement that A and B are dual descriptions of one partition structure, counting across cells in one case and within a cell in the other. This chapter treats A and B as related rather than as two of the five genuinely separate candidates from here on.

5.5 Entropy C: Historical Entropy Reduces to B

Claim 5.5.1 (Irreversibility Theorem). *Irreversibility emerges whenever a history projection $\pi_H : \mathcal{H} \rightarrow X$ is many-to-one: for $h_1 \neq h_2$ with $\pi_H(h_1) = \pi_H(h_2) = x$,*

knowledge of the present state x does not determine which history occurred, since π_H^{-1} is not unique.

Claim 5.5.2 (Historical entropy is representational entropy). *Historical entropy, informally the multiplicity of histories collapsing to a state, is $\log |\pi_H^{-1}(x)|$ for the history projection π_H of Claim 5.5.1. This is exactly section 5.3's representational entropy $S_{\pi_H}(x)$, applied to the specific case where the representation in question is a history-to-state projection. Historical entropy is not a fourth candidate; it is Entropy B evaluated at a particular kind of π .*

This reduction is a genuine derivation rather than an analogy: both quantities are defined as the log-cardinality of a fiber under a projection, and π_H is a projection in exactly Chapter 1's sense. Nothing about the reduction required inventing new machinery.

5.6 Entropy Production, the Second Law, and Reachability

Three further results build on Claim 5.4.1 without introducing new entropy candidates.

Claim 5.6.1 (Entropy Production Theorem). *Coarsening of distinction structure — merging cells P_i and P_j into $P_i \cup P_j$ — increases $\log \Omega$ in the merged cell relative to either constituent, since $|P_i| + |P_j| > \max(|P_i|, |P_j|)$. Entropy increases wherever distinction structure coarsens.*

Claim 5.6.2 (Second Law, reinterpreted). *In the absence of repair or distinction-generating dynamics, recoverable distinction capacity $D(t)$ tends to decrease: $\dot{D} \leq 0$. Entropy increase is the shadow cast by distinction loss.*

The relation between entropy and reachability deserves more care than stating the two are monotonically linked. Coarsening observational distinctions does not, by itself, shrink the set of states an underlying system can *physically* reach; it shrinks the set of reachable states an observer can *tell apart*. The quotient connecting the two must be built explicitly rather than asserted.

Definition 5.6.3 (Effective reachable set). *Let $\mathcal{R}(x)$ be the physically reachable set from x , and let $q_t : X \rightarrow X/\sim_t$ be the quotient map induced by the distinguishability relation at time t . The effective reachable set is $\overline{\mathcal{R}}_t(x) = q_t(\mathcal{R}(x))$, with effective reachability volume $V_{\text{eff}}(x, t) = \bar{\mu}_t(\overline{\mathcal{R}}_t(x))$ for a measure $\bar{\mu}_t$ on the quotient (cardinality, in the finite case).*

Claim 5.6.4 (Effective Reachability Contraction). *Suppose the distinction relation coarsens from t_1 to t_2 : $x \sim_{t_1} y \Rightarrow x \sim_{t_2} y$. Then there is a canonical surjection $c_{21} : X/\sim_{t_1} \rightarrow X/\sim_{t_2}$ with $q_{t_2} = c_{21} \circ q_{t_1}$, so $\overline{\mathcal{R}}_{t_2}(x) = c_{21}(\overline{\mathcal{R}}_{t_1}(x))$, and consequently $|\overline{\mathcal{R}}_{t_2}(x)| \leq |\overline{\mathcal{R}}_{t_1}(x)|$ in the finite case.*

Proof. Since $\sim_{t_2}$ is coarser, each $\sim_{t_1}$ -class lies within a unique $\sim_{t_2}$ -class; sending the former to the latter defines c_{21} . Then $\overline{\mathcal{R}}_{t_2}(x) = q_{t_2}(\mathcal{R}(x)) = c_{21}(q_{t_1}(\mathcal{R}(x))) = c_{21}(\overline{\mathcal{R}}_{t_1}(x))$, and a surjective image cannot contain more distinct elements than its domain. $\square$

Claim 5.6.4 proves distinction loss implies effective future contraction. It does not prove distinction loss implies *physical* future contraction — that stronger claim requires dynamics coupling representational loss back into executable action, and holds only conditionally.

Definition 5.6.5 (Distinction-dependent controllability). *A controlled system is distinction-dependent when the executable action set at x depends only on its currently distinguished class: $\mathcal{A}_t(x) = \mathcal{A}_t([x]_{\sim_t})$. Coarsening is action-monotone when $\sim_{t_1} \subseteq \sim_{t_2} \Rightarrow \mathcal{A}_{t_2}(x) \subseteq \mathcal{A}_{t_1}(x)$ for every x .*

Claim 5.6.6 (Physical Reachability Contraction under Action Dependence). *If distinction coarsening is action-monotone, the physically reachable set generated by executable actions satisfies $\mathcal{R}_{t_2}(x) \subseteq \mathcal{R}_{t_1}(x)$.*

Proof. Every trajectory available at t_2 is composed from actions in $\mathcal{A}_{t_2}$, each of which, by action-monotonicity, was already available at t_1 . Hence every t_2 -trajectory is also a t_1 -trajectory. $\square$

Claims 5.6.4 and 5.6.6 together give this section's actual result: entropy production provably contracts *effective* reachability unconditionally, and contracts *physical* reachability under the further assumption that action availability is itself distinction-dependent. The two should not be conflated; the second is a substantive cybernetic hypothesis about how representation couples to action, not a free consequence of the first.

Definition 5.6.7 (Reachability admissibility). *A state y is admissible from x at time t when its effective class is reachable: $x \rightsquigarrow_t y \iff q_t(y) \in \overline{\mathcal{R}}_t(x)$. The admissible future region is $\Gamma_t(x) = \overline{\mathcal{R}}_t(x)$.*

Claim 5.6.8 (Entropy–Admissibility Corollary). *If entropy production is realized as monotone coarsening of the distinction relation, the effective admissible future region is non-increasing: $\sim_{t_1} \subseteq \sim_{t_2} \Rightarrow \Gamma_{t_2}(x) = c_{21}(\Gamma_{t_1}(x))$, and in the finite case $|\Gamma_{t_2}(x)| \leq |\Gamma_{t_1}(x)|$.*

This is immediate from Claim 5.6.4 under definition 5.6.7’s identification of the admissible future region with the effective reachable set. It gives the chapter sequence this section was actually able to establish: Entropy $\rightarrow$ effective reachability $\rightarrow$ admissibility, deliberately avoiding an unsupported direct arrow from entropy to repair. What repair does with a contracted admissible region is Chapter 6’s business, not this chapter’s.

Remark 5.6.9 (Decorrelated dimensionality, not operator count). *A useful comparison, and a useful caution, is provided by Posani et al. (2026), who study the dimensionality and separability of neural population representations across the cortical hierarchy. Their dimensionality measure is the participation ratio*

$$\text{PR}(C) = \frac{(\sum_i \lambda_i)^2}{\sum_i \lambda_i^2},$$

where the λ_i are eigenvalues of a response-covariance matrix C . The participation ratio is not identical to any quantity defined in this book, and it is worth stating precisely why, rather than leaving the resemblance to do unearned work. It is not representational entropy: $\text{PR}(C) \neq \log |\pi^{-1}(m)|$. It is not reachability volume: $\text{PR}(C) \neq \mu(\mathcal{R}_{\mathcal{E}}(x))$. No definition-by-substitution identifies these constructions, and they take different inputs — a covariance operator over observed responses, against a set of attainable future states.

The comparison is nevertheless structurally informative, because of what the participation ratio discounts. If covariance is concentrated in a small number of eigen-directions, many recorded units still yield a low effective dimensionality; a high value requires variance distributed across several independently contributing directions. The lesson for an operator ecology is that repertoire cardinality alone need not measure generative diversity: $|\text{Reach}(\mathcal{E})|$ can be large while the operators it contains are nearly redundant, so that the induced future region $\mu(\mathcal{R}_{\mathcal{E}}(x))$ remains small. This is not a new observation so much as a sharpening of one already present elsewhere in this corpus: a resilience discount for correlated checks, on the principle that a hundred people relying on one hidden assumption do not constitute a hundred independent checks. The participation ratio and that resilience discount are not the same quantity, but both implement the same correction — correlated contributions count for less than independent ones — from two unrelated domains arriving at it separately. Chapter 6 develops a repertoire-level analogue of this correction directly.

5.7 Entropy D: Repair Entropy

Chapter 6 already used $S_R(F \mid T)$, the repair entropy of an operator F given a retained trace T , to state the Anti-Alignment Theorem, without independently justifying why this quantity deserves the name entropy rather than merely borrowing it by analogy. That justification can now be given in full, upgrading the

reduction from a structural resemblance to an ordinary conditional-entropy construction.

Definition 5.7.1 (Fault hypothesis space). *For an operator F and observed symptom s , let $\Omega_F(s)$ be the finite set of fault hypotheses capable of producing s . A retained trace T determines an evidence map $e_T : \Omega_F(s) \rightarrow E_T$, where $e_T(f)$ records the facts about hypothesis f distinguishable using T . This induces a partition $\mathcal{P}_T = \{e_T^{-1}(a) : a \in E_T\}$ of the fault hypothesis space, and the live hypotheses after observing the actual trace value a_T are $\mathcal{H}_{\text{fault}}(F, T) = e_T^{-1}(a_T)$.*

Definition 5.7.2 (Repair entropy). *For $p(f \mid s, T)$ a probability distribution over $\mathcal{H}_{\text{fault}}(F, T)$,*

$$S_R(F \mid T) = - \sum_{f \in \mathcal{H}_{\text{fault}}(F, T)} p(f \mid s, T) \log p(f \mid s, T),$$

which under a uniform prior becomes $S_R(F \mid T) = \log |\mathcal{H}_{\text{fault}}(F, T)|$.

Claim 5.7.3 (Trace Refinement Theorem). *Suppose T_2 retains every diagnostically relevant fact retained by T_1 , and possibly more. Then $\mathcal{P}_{T_2}$ refines $\mathcal{P}_{T_1}$ and $H(F_{\text{fault}} \mid s, T_2) \leq H(F_{\text{fault}} \mid s, T_1)$, hence under the uniform-prior convention $S_R(F \mid T_2) \leq S_R(F \mid T_1)$, strictly whenever the additional information separates two hypotheses that remained equivalent under T_1 and both had positive probability.*

Proof. Since T_2 contains all diagnostically relevant information in T_1 , there is a forgetful map $q : E_{T_2} \rightarrow E_{T_1}$ with $e_{T_1} = q \circ e_{T_2}$. If $e_{T_2}(f) = e_{T_2}(g)$ then $e_{T_1}(f) = e_{T_1}(g)$, so every cell of $\mathcal{P}_{T_2}$ lies inside a cell of $\mathcal{P}_{T_1}$: $\mathcal{P}_{T_2}$ refines $\mathcal{P}_{T_1}$. Equivalently, the variables form a Markov chain $F_{\text{fault}} \rightarrow T_2 \rightarrow T_1$, with T_1 obtained from T_2 by discarding information. Conditional entropy is monotone under additional conditioning, giving the stated inequality; under a uniform distribution, conditional entropy is the log-cardinality of the surviving hypotheses, giving the second inequality. Strictness follows whenever refinement splits a positive-probability cell. $\square$

Claim 5.7.4 (Anti-Alignment, as a corollary). *Let $T_{\min}$ be a minimal trace sufficient to enact an operator F , and $T \supseteq T_{\min}$ retain additional diagnostic information. Then $S_R(F \mid T) \leq S_R(F \mid T_{\min})$, strictly whenever the discarded information excludes at least one otherwise live hypothesis.*

Proof. Apply Claim 5.7.3 with $T_1 = T_{\min}$, $T_2 = T$. $\square$

Claim 5.7.4 makes Chapter 6's Anti-Alignment result a direct consequence of the partition structure developed in this chapter and Chapter 2, rather than an imported standalone claim resting on analogy. Repair entropy is accordingly

reclassified: it is not a fifth candidate alongside A through C, but Entropy A itself, applied to the specific partition a retained trace induces on a space of fault hypotheses.

Remark 5.7.5 (Anti-Alignment is a monotonicity lemma, not an optimization result). *Claim 5.7.4 shows $S_R(F \mid T)$ is weakly decreasing as T grows, and reads naturally as license to retain as much trace as possible. That reading silently assumes retention is free. Making this assumption explicit requires a retention cost functional $K(T)$, with the actual objective*

$$J(T) = S_R(F \mid T) + \lambda K(T)$$

for some trade-off parameter $\lambda > 0$. Under this objective the optimum is generally not $T = T_{\max}$ but a minimal admissibility-sufficient trace T^ : enough retained history to support every fault-diagnosis judgment that will actually be needed, and no more. Claim 5.7.4 remains correct as a monotonicity statement about S_R alone; it does not by itself determine T^* , since it says nothing about K . $K(T)$ and T^* are developed in Chapter 6, where retention cost is native; Anti-Alignment should be read as one term of a larger optimization, not as a standalone argument for maximal retention.*

5.8 Entropy E: The RSVP Field, Left Open

The RSVP formalization defines an admissibility relation directly in terms of a field triple $(\Phi, \mathbf{v}, S)$, where $S : M \rightarrow \mathbb{R}_{\geq 0}$ is required to satisfy $\dot{S} \geq 0$ along admissible trajectories — entropy monotonicity is part of what admissibility *means* in that formalization, not a consequence derived after the fact.

Open Problem 5.8.1. *No result here derives $S(x, t)$ from the partition- and fiber-based entropies A through D above. The RSVP fields are motivated by the distinction-repair-reachability chain rather than derived from it, and the specific functional form relating reachability volume to S is a heuristic argument, not a uniqueness theorem. There is a plausible structural correspondence worth flagging rather than asserting: the Second Law of Claim 5.6.2 gives $\dot{D} \leq 0$ for a distinction-capacity quantity D , and $\dot{S} \geq 0$ in the RSVP formalization is the same monotonic inequality under a sign flip, suggestive of S playing the role of consumed distinction capacity rather than distinction capacity itself. This is a resemblance in the direction and character of two monotonicity constraints, not a proof that S and D are the same quantity, functionally related, or defined over the same underlying space. Establishing or refuting this correspondence is left open.*

Remark 5.8.2 (Reconstructive and thermodynamic irreversibility are distinct). *Reconstructive irreversibility — the fiber non-injectivity of a transition function, a purely*

combinatorial property — is distinct from thermodynamic irreversibility: the former concerns neither entropy, phase-space volume, nor coarse-graining directly. This distinction, drawn from fiber and distinguishability apparatus essentially identical to entropies A through D, is independent evidence for the caution urged in open Problem 5.8.1. It does not resolve the open problem. It strengthens the case that leaving it open is the right call.

5.9 Summary

Candidate	Status	Relation
A: Shannon/distinction	Checked	Dual to B (Claim 5.4.1)
B: Representational	Checked	Dual to A; C is an instance of B
C: Historical	Checked	Reduces to B (Claim 5.5.2)
D: Repair	Checked	Is an instance of A (Claim 5.7.4)
E: RSVP field	Open	Motivated, not derived; see open Problem 5.8.1

Of five candidates, four are reduced above to a single underlying partition/fiber structure with genuine derivations, and the fifth — the one actually load-bearing for admissibility — remains explicitly open. That asymmetry matters: whatever is eventually said about admissibility inherits an unresolved dependency on Entropy E, however cleanly A through D have been tied together.

Summary

The distinction/information/entropy sequence developed across Chapters 1–2 and this chapter does not extend cleanly into a further stage called “History.” History (Entropy C) is not a separate node; it is Entropy B evaluated at a particular projection. What connects entropy to the second half of this volume is the Entropy–Admissibility Corollary of Claim 5.6.8, not any direct entropy-to-repair theorem. Reachability, not history or recovery as separate stages, is the actual bridge — but a bridge with two distinctly typed ends rather than one shared object. The executable repertoire $\text{Reach}(\mathcal{E})$ of Chapter 6 is a set of *operators*; the admissible future region $\Gamma_t(x)$ and reachability volume $V_{\text{eff}}(x, t)$ of this chapter are sets and measures of *states*. They fail a basic type test and cannot be identical. What connects them is a generator relation: an executable repertoire supplies the operators from which state-to-state trajectories are composed, and the reachable-state set is a derived closure of that repertoire, not the repertoire itself. That closure map, and its monotonicity, is constructed in Chapter 6.

Chapter 6

Repair

6.1 Introduction

A structure that persists over time is easy to mistake for a static fact, something that simply continues to be the case once it has been established. This chapter argues that persistence is never a static fact. A structure persists because something is continually acting to keep it in existence against the ordinary pressure of decay, damage, and drift, and the moment that action stops, the structure does not immediately vanish, but it has already stopped being maintained, and what follows is only a matter of time. This chapter develops the vocabulary needed to say precisely what that “something” is, what it means for it to persist in turn, and what quantities govern how much repair capacity a system actually has available to it.

6.2 Three Levels of Modification

Before asking what maintains a structure, it is worth being precise about what kind of change is even available to a system that modifies itself, since discussions of adaptation and repair routinely conflate phenomena that differ not in how much behavior changes but in what kind of object is doing the changing.

Definition 6.2.1 (Levels of adaptation). *Let X be a state space, Θ a parameter space, and $\mathcal{F}$ a space of maps $X \rightarrow X$. A system exhibits Level 1 (state update) behavior when $x_{t+1} = F(x_t)$ for fixed F : the system moves through X under an unchanging rule. It exhibits Level 2 (parameter update) behavior when $x_{t+1} = F_{\theta_t}(x_t)$ for fixed $F : \Theta \times X \rightarrow X$ with $\theta_t \in \Theta$ varying: behavior changes, but only within the range expressible by the fixed parametric family F . It exhibits Level 3 (rule update) behavior when $F_{t+1} = G(F_t, x_t, E_t)$ for some environment E_t : the mapping governing future*

behavior is itself the object being modified, not merely a parameter feeding into a fixed mapping.

Repair, in the sense this chapter develops, is a Level 3 phenomenon. A system that merely moves through its state space (Level 1) or adjusts parameters within a fixed family (Level 2) is not repairing anything; it is only running. Repair requires that the rule governing future behavior — the operator itself — be a target of modification, which is why the operator ecologies developed in this chapter are ecologies of Level 3 objects, not of states or parameters.

6.3 From Structures to Repairs to Ecologies

The argument of this section is developed at length, with a worked example of two initially identical cities that diverge over a decade purely in whether ordinary maintenance — road resurfacing, apprentice training, records-keeping, governance continuity — continues. Its conclusion is stated here rather than its argument reproduced: state equivalence between two systems is insufficient to predict their futures, because two systems can agree in every recorded particular while already diverging in what each is still capable of doing about that state. Structural equivalence fails for the same reason, one level down. The natural first correction — structures persist, and repair is something that occasionally happens to them — is shown there to have the explanatory priority backwards. The corrected view is that repair occurs, and structures persist because of it; and one level further, that repair occurs because an operator ecology persists, and structures are, in the resulting phrase, frozen operator ecologies: visible residues of a network of constraints that, for as long as it holds, keeps producing them.

6.4 A Hierarchy of Equivalence

Operator Ecology develops a hierarchy of increasingly demanding notions of sameness between two systems under damage, each an improvement on the last, each still shown insufficient in turn until the final one.

Definition 6.4.1 (Equivalence hierarchy). *Two systems are state equivalent at time t if they occupy the same point in state space at t . They are structurally equivalent if their underlying architectures agree, independent of current state. They are behaviorally equivalent if they are bisimilar in the coalgebraic sense: indistinguishable by any observation of their input-output behavior, including future behavior, from the current state onward. They are regeneratively equivalent if they possess sufficiently similar operator ecologies to recreate the same future capacities after damage, independent of whether their current state or current behavior agree at all.*

State and structural equivalence fail for the reason section 6.3 already gave: two systems, or two moments of one system, can agree on both while their capacity to persist has already diverged. Behavioral equivalence is a genuine improvement, since it tracks not a static fact but an ongoing capacity, and bisimulation gives that capacity a precise mathematical name. But behavioral equivalence is defined over a system’s *current* behavior, and has nothing left to say once that behavior has already been damaged or lost: two systems can be behaviorally equivalent right up until the moment of damage and then diverge sharply in what they can do afterward, a divergence behavioral equivalence alone cannot see coming because it was never asking about recovery. Regenerative equivalence is the notion built to see it coming: it is not a claim about what a system currently does, but about what it retains the capacity to do again.

6.5 Operator Ecologies as Constraint Networks

The natural first attempt to formalize an operator ecology is as an inventory: the set of operators a system currently has available. This undercounts what actually explains regenerative equivalence, since an inventory says nothing about why any particular operator remains executable, or what would happen to the rest of the inventory if one entry were lost.

Definition 6.5.1 (Operator ecology). *Following Juarrero’s account of causation as constraint rather than efficient push, an operator ecology $\mathcal{E}$ is not the set of operators it contains but the network of enabling constraints among them: the structure that determines which operators remain executable, and which become executable or cease to be executable as other parts of the network change.*

This definition is what allows the ecological vocabulary of the next section — keystone, extinction, health — to be more than metaphor. An inventory has no notion of one entry depending on another; a constraint network does, and the dependency structure is precisely what the following quantities measure.

6.6 Keystone Operators, Extinction, and Health

Definition 6.6.1 (Reachable repertoire and keystone impact). *For an operator ecology $\mathcal{E}$, let $\text{Reach}(\mathcal{E})$ denote the reachable executable repertoire: the operators actually executable given $\mathcal{E}$ ’s current constraint graph. The impact of an operator o is*

$$K(o) = |\text{Reach}(\mathcal{E})| - |\text{Reach}(\mathcal{E} \setminus \{o\})|,$$

the size of the repertoire lost by removing o and everything its removal disconnects. A keystone operator is one with unusually large $K(o)$ relative to how often it is itself invoked.

Apprenticeship training in the two-cities example of section 6.3 is a keystone operator in exactly this sense: it is invoked rarely, but its removal disconnects every operator that depended on the next generation of practitioners it would otherwise have supplied.

Definition 6.6.1's impact measure $K(o)$ counts lost operators. It does not, by itself, say anything about the states or futures those operators would have made reachable, which is a different and better question to ask of a keystone loss. Answering it requires building the map from an executable repertoire to a reachable-state set explicitly, since Chapter 5's admissible future region $\Gamma_t(x)$ is a set of states, not of operators, and the two constructions fail a basic type test: $\text{Reach}(\mathcal{E})$ is not the same kind of object as $\Gamma_t(x)$, and no theorem in this book identifies them.

Definition 6.6.2 (Operator-induced reachable set). *Let X be a state space and $\mathcal{O}_{\mathcal{E}} = \text{Reach}(\mathcal{E})$ the currently executable repertoire of an ecology $\mathcal{E}$. For $x \in X$ and a budget $n \in \mathbb{N}$, define*

$$\mathcal{R}_{\mathcal{E}}^{(n)}(x) = \{o_k \circ \dots \circ o_1(x) : 0 \leq k \leq n, o_i \in \mathcal{O}_{\mathcal{E}}, \text{ every intermediate composition executable}\},$$

with unrestricted reachable set $\mathcal{R}_{\mathcal{E}}(x) = \bigcup_{n \geq 0} \mathcal{R}_{\mathcal{E}}^{(n)}(x)$ and, where X carries a measure μ , induced reachability volume $V_{\mathcal{E}}(x) = \mu(\mathcal{R}_{\mathcal{E}}(x))$.

This is the closure map Chapter 5 asked for: it takes a repertoire of operators and returns the set of states composable from them, which is the right kind of object to compare against $\Gamma_t(x)$. A basic monotonicity result requires ruling out operators that disable one another's compositions.

Definition 6.6.3 (Extension-monotone ecology). *An operator ecology is extension-monotone when $\mathcal{O}_{\mathcal{E}_1} \subseteq \mathcal{O}_{\mathcal{E}_2}$ implies every composition executable in $\mathcal{E}_1$ remains executable in $\mathcal{E}_2$.*

Claim 6.6.4 (Repertoire–Reachability Monotonicity). *For extension-monotone ecologies, $\text{Reach}(\mathcal{E}_1) \subseteq \text{Reach}(\mathcal{E}_2) \Rightarrow \mathcal{R}_{\mathcal{E}_1}(x) \subseteq \mathcal{R}_{\mathcal{E}_2}(x)$ for every x , and consequently $V_{\mathcal{E}_1}(x) \leq V_{\mathcal{E}_2}(x)$ whenever μ is monotone under inclusion.*

Proof. Let $y \in \mathcal{R}_{\mathcal{E}_1}(x)$, reached by an executable composition $y = o_k \circ \dots \circ o_1(x)$ with every $o_i \in \text{Reach}(\mathcal{E}_1)$. Repertoire inclusion places every o_i in $\text{Reach}(\mathcal{E}_2)$, and extension-monotonicity ensures the same composition remains executable there, so $y \in \mathcal{R}_{\mathcal{E}_2}(x)$. The volume inequality follows from monotonicity of μ . $\square$

Without extension-monotonicity the inclusion can fail outright: an ecology's operators can interact so that adding one disables a composition that a smaller repertoire supported, and Claim 6.6.4 is not claimed to hold in that regime.

Chapter 5 flags, via a comparison to Posani et al. (2026)'s work on cortical population dimensionality, that $|\text{Reach}(\mathcal{E})|$ is a poor proxy for generative diversity whenever the repertoire's operators are redundant with one another. That comparison is illustrative there, not formal; a repertoire-native version of the same correction can be built directly from this section's own constructions.

Definition 6.6.5 (Effective repertoire dimension). *Suppose the local effect of each executable operator $o \in \text{Reach}(\mathcal{E})$ at x is represented by a vector $v_o(x) \in T_x X$ in a linearized tangent space. For non-negative invocation weights w_o , define*

$$C_{\mathcal{E}}(x) = \sum_{o \in \text{Reach}(\mathcal{E})} w_o v_o(x) v_o(x)^{\top}.$$

If $\lambda_1, \dots, \lambda_n$ are the eigenvalues of $C_{\mathcal{E}}(x)$, the effective repertoire dimension is

$$D_{\text{eff}}(\mathcal{E}, x) = \frac{(\sum_i \lambda_i)^2}{\sum_i \lambda_i^2}.$$

Claim 6.6.6 (Bounds on effective repertoire dimension). *For any nonzero positive semidefinite $C_{\mathcal{E}}(x)$,*

$$1 \leq D_{\text{eff}}(\mathcal{E}, x) \leq \text{rank } C_{\mathcal{E}}(x).$$

Equality on the left holds when all operator effects lie in one effective direction; equality on the right holds when the nonzero eigenvalues are equal.

Proof. Let $r = \text{rank } C_{\mathcal{E}}(x)$ with nonzero eigenvalues $\lambda_1, \dots, \lambda_r > 0$. Since $(\sum_i \lambda_i)^2 \geq \sum_i \lambda_i^2$, the lower bound follows. By Cauchy–Schwarz, $(\sum_{i=1}^r \lambda_i)^2 \leq r \sum_{i=1}^r \lambda_i^2$, giving the upper bound. The equality conditions are the standard ones for the two inequalities. $\square$

One hundred collinear operator effects give $D_{\text{eff}} = 1$ regardless of $|\text{Reach}(\mathcal{E})|$; k orthogonal, equally weighted effects give $D_{\text{eff}} = k$. This gives the correlated-operators concern a precise, checkable form rather than leaving it at the level of intuition.

Remark 6.6.7 (A local proxy, not a replacement). $D_{\text{eff}}(\mathcal{E}, x)$ measures local effective action dimension at a single state x , under a linearization. It is not $\mu(\mathcal{R}_{\mathcal{E}}(x))$, and Claim 6.6.4's monotonicity is not claimed to transfer to it automatically. Noncommutativity among operators, state-dependent executability, and the compounding effects of long compositions can all make ecologies with identical $D_{\text{eff}}(\mathcal{E}, x)$ generate different global reachable regions. This quantity is offered as a cheaper, locally computable predictor of generative diversity, not as a substitute for definition 6.6.2's actual reachable set when the two disagree.

Definition 6.6.2 also gives keystone impact a second, more consequential measure than $K(o)$.

Definition 6.6.8 (Generative keystone impact). *The generative impact of o at x is $K_V(o; x) = \mu(\mathcal{R}_\mathcal{E}(x)) - \mu(\mathcal{R}_{\mathcal{E} \setminus \{o\}}(x))$.*

Remark 6.6.9 (Inventory impact and generative impact can diverge). *There is no reason for $K(o)$ and $K_V(o; x)$ to rank operators identically. Repertoire impact counts lost executable operators; generative impact measures the future region their loss disconnects. A single operator may account for only one lost repertoire element under K while controlling access to a vast region of state space under K_V , or the reverse. The former is architectural; the latter is dynamical. This is why definition 6.6.11's ecological health is defined below in terms of recoverable futures rather than repertoire size: health was already, implicitly, tracking something closer to K_V than to K , without the intermediate construction that makes the relationship precise.*

Definition 6.6.10 (Extinction). *An operator is extinct within an ecology $\mathcal{E}$ when no executable realization of it remains, regardless of whether a description of it survives.*

Extinction in this sense is deliberately independent of documentation. City B's apprenticeship program, in the example of section 6.3, can be extinct while a charter document describing it still sits, accurately and uselessly, in an archive: the description survives, the executable operator does not. This is the same distinction Chapter 7 makes between representation and executability in its architecture gate, and the same distinction that recurs, independently, in the enactment/representation asymmetry of section 6.7 below.

Definition 6.6.11 (Ecological health). *Writing $\Gamma(\mathcal{E})$ for the set of futures recoverable from $\mathcal{E}$ under plausible further damage, the health of an operator ecology is $H(\mathcal{E}) = \mu(\Gamma(\mathcal{E}))$ for some measure μ . Extinction is what decreases H ; restoration is what increases it; keystone loss is what causes it to drop sharply rather than gradually.*

Health, so defined, does not measure how much of an ecology is currently running. It measures how much could still be regenerated if some further part of what is currently running were lost — a forward-looking quantity, in keeping with regenerative equivalence being a forward-looking notion. That phrasing has so far described what health tracks without saying precisely what an act of repair does to it; the reachable-set construction of definition 6.6.2 now makes that precise, and identifies $\Gamma(\mathcal{E})$ above with Chapter 5's admissible future region under the natural reading $\Gamma(\mathcal{E}) = \mathcal{R}_\mathcal{E}(x)$ for the ecology's current state x .

Definition 6.6.12 (Reachability-restoring repair). *Let $\Gamma^*(x)$ be a target admissible future region and $\Gamma_d(x)$ the region available after damage d . An operator r is a reachability-restoring repair when $\mu(\Gamma_{r(d)}(x)) > \mu(\Gamma_d(x))$, and complete relative to the target when $\Gamma_{r(d)}(x) = \Gamma^*(x)$, partial otherwise.*

Claim 6.6.13 (Repair Existence as Reachability). *A damaged system admits repair relative to target $\Gamma^*(x)$ if and only if there exists $r \in \mathcal{O}_\mathcal{E}$ such that $\mu(\Gamma_{r(d)}(x)) > \mu(\Gamma_d(x))$.*

Proof. Immediate from definition 6.6.12: the forward direction supplies such an operator by definition; the reverse direction identifies any executable operator satisfying the inequality as a repair. $\square$

Claim 6.6.13 is close to definitional, but it fixes the vocabulary the rest of this book needs: admissibility describes the future region; damage contracts it; repair restores some of it. A further distinction, between fixing a structure and restoring the ecology that fixes structures, is worth making formal rather than leaving implicit.

Definition 6.6.14 (Regenerative operator). *An operator g is regenerative for $\mathcal{E}$ when it increases not merely the current admissible future region but the ecology's future capacity to produce reachability-restoring repairs: $\text{Reach}(g(\mathcal{E})) \supsetneq \text{Reach}(\mathcal{E})$, or more generally $\mu(\mathcal{R}_{g(\mathcal{E})}(x)) > \mu(\mathcal{R}_\mathcal{E}(x))$ under plausible subsequent damage.*

An operator satisfying definition 6.6.12 fixes a structure. An operator satisfying definition 6.6.14 restores the ecology that fixes structures — retraining an apprentice rather than repaving a road. This is the formal difference behind section 6.3's claim that structures are frozen operator ecologies: a repair changes $\Gamma_d(x)$; a regeneration changes what $\text{Reach}(\mathcal{E})$ itself can produce, which is a strictly higher-leverage act whenever it succeeds, and exactly the kind of act a keystone-operator loss under remark 6.6.9 disables first.

6.7 History, Compression, and Repairability

The preceding sections describe an operator ecology as it stands at a moment: which operators are executable, how they depend on one another, how much future capacity they jointly support. This section adds a dimension the preceding ones leave implicit: an operator does not appear from nowhere. In every case, it is the compressed residue of a history of attempts, corrections, and successes.

Definition 6.7.1 (Compression). *Let H denote a history: an ordered trajectory of states, attempts, corrections, and outcomes. Let C be a compression map taking a history to an operator, $C(H) = F$, where F is characterized independently of C by its behavior — what it does when invoked — irrespective of how F came to exist.*

Once F exists, two different things can be done with it. It can be *enacted*: invoked and executed on some occasion. And it can, where defined, be *represented*: given a symbolic, inspectable description sufficient to reconstruct its behavior

without executing it. These are not the same capacity, and the domain of enactment is not in general contained in the domain of representation — an operator can be fully executable while admitting no complete symbolic account of what it does, in the sense Polanyi’s tacit knowledge already names. This is the same gap definition 6.6.10 exploited in the other direction: there, representation survived without enactment; here, enactment can exceed what representation is available to capture at all.

The central result of *History Before Function* bears directly on repair capacity, and it is worth stating the dependency honestly rather than re-deriving it twice: Chapter 5 develops the same result from this book’s own partition apparatus, rather than importing it by analogy, and this section defers to that proof instead of duplicating it.

Definition 6.7.2 (Repair entropy). *For an operator F and a retained trace T , $S_R(F \mid T)$ measures the number of distinct fault explanations consistent with a symptom s and the information retained in T . Chapter 5 gives this a full partition-based construction (Fault hypothesis space, repair entropy as conditional Shannon entropy) rather than the informal version stated here.*

Claim 6.7.3 (Anti-Alignment, proved in Chapter 5). *Let $T_{\min}$ be the minimal trace sufficient to enact F at all. Then $S_R(F \mid T) \leq S_R(F \mid T_{\min})$ for every retained trace T that contains $T_{\min}$, strictly whenever some fact in $T \setminus T_{\min}$ excludes an otherwise-live fault location. This is Chapter 5’s Anti-Alignment result, proved there as a corollary of its Trace Refinement Theorem; this chapter uses the result but does not reprove it.*

The consequence for operator ecologies is direct. An ecology under pressure to compress — to retain only what is needed to keep operators running, discarding whatever does not affect current executability — is, by Claim 6.7.3, systematically discarding exactly the information that would have been needed to repair those operators once they break. Efficient enactment and future repairability are not merely different goals; they sit, by construction, in tension with each other, and an ecology optimized purely for the former is an ecology that has made itself harder to fix.

Remark 6.7.4 (Why this matters for Chapter 7). *Chapter 7 introduces a history-indexed restorability boundary, $\partial_{R(X, \mathcal{E}, \mathcal{R}, H_t)}$, on the grounds that diagnosing recurrent, worsening, or structurally stubborn conditions requires a remembered repair history rather than a single snapshot. This chapter’s Anti-Alignment result is the reason that requirement is not merely convenient but load-bearing: an ecology that retains only $T_{\min}$ for each of its operators has, in effect, already discarded the H_t that Chapter 7 needs, not through neglect but as an unavoidable consequence of having compressed efficiently. A repair theory that only asks whether an operator exists, without asking how much of its generating history survives alongside it, will systematically overstate the repair capacity of a maximally compressed ecology.*

Remark 6.7.5 (The bridge from repertoire to admissible futures, in one place). *Three constructions have now been built across this chapter and the next two, and it is worth naming the chain that connects them once, rather than leaving a reader to reassemble it from separate remarks.*

$$\text{Reach}(\mathcal{E}) \longrightarrow \mathcal{R}_{\mathcal{E}}(x) \longrightarrow \Gamma_t(x).$$

$\text{Reach}(\mathcal{E})$ is this chapter’s executable repertoire: a set of operators, nothing more. $\mathcal{R}_{\mathcal{E}}(x)$, this chapter’s operator-induced reachable set, is the closure of that repertoire under composition: the states actually reachable from x by chaining executable operators together. $\Gamma_t(x)$, Chapter 5’s admissible future region, is a further restriction of a reachable set — there, $\overline{\mathcal{R}}_t(x)$, the effective reachable set under a distinguishability quotient — to the states an observer’s current admissibility structure actually licenses. Repertoire, then generated futures, then admissible futures: each stage is a genuine restriction of the one before it, not a synonym for it, and the type-test discipline used throughout this book applies to all three pairwise: none are shown identical, each is shown to generate or restrict its neighbor.

6.8 A Worked Example: Misalignment, Extinction, Regeneration

Chapter 1 closed with a remark that a distinction-preserving transformation need not preserve operational competence, and deferred the full illustration to this chapter. A single narrative, developed in three stages, separates three phenomena this chapter’s apparatus keeps distinct but that are easy to confuse in practice: an intact repertoire that no longer aligns with its target space, a repertoire with pieces genuinely and permanently removed, and the appearance of a new ecology that restores generative capacity rather than merely patching what was lost.

Stage one: misalignment

Consider a keyboard instrument relabeled so that pitch runs high to low across the keys rather than low to high — every note distinguishable from every other exactly as before, no fiber of the original distinguishability space collapsed, nothing lost in the sense Chapter 1 measures. A trained performer’s fingering repertoire, $\text{Reach}(\mathcal{E})$, is entirely intact: every motion the performer knows how to execute is still executable. What fails is the correspondence between that repertoire and the relabeled target space. A motion trained to produce a rising line now produces a falling one; a leap trained to reach a particular interval now

reaches a different one. No operator has been lost, and no operator has become inexecutable in the sense definition 6.6.10 requires. The repertoire and the space it acts on have simply come apart.

Remark 6.8.1 (Misalignment is neither extinction nor repair deficit). *Misalignment is not covered by any failure mode developed so far. It is not extinction, since every operator remains executable. It is not repair deficit in Chapter 7’s sense, since $R(d, x)$ for the relabeling “deviation” is not the right question — there is no damage to repair, only a representation to relearn. It is a third phenomenon: $\text{Reach}(\mathcal{E})$ intact, $\mathcal{R}_{\mathcal{E}}(x)$ well defined, but no longer usefully predictive of what a given operator, applied by a performer trained on the old correspondence, will actually produce.*

Stage two: extinction

The pianist Paul Wittgenstein lost his right arm in the First World War. This is extinction in the precise sense of definition 6.6.10: an entire sub-repertoire of his trained operators, everything requiring two hands, became permanently inexecutable. Nothing was relabeled; nothing needed relearning in the sense of stage one. A substantial fraction of $\text{Reach}(\mathcal{E})$ simply ceased to exist, $\text{Reach}(\mathcal{E}') \subsetneq \text{Reach}(\mathcal{E})$, and by definition 6.6.8 the loss carried generative impact far beyond the raw count of operators removed: an enormous fraction of the existing piano repertoire, and the compositional conventions behind it, assumed two independent hands, so the loss disconnected far more of $\mathcal{R}_{\mathcal{E}}(x)$ than the operator count alone would suggest.

Stage three: regeneration

Wittgenstein’s response was not to search for a repair operator that would restore two-handed capacity — none existed — nor to make do with the existing one-handed repertoire, which was thin and largely pedagogical. He commissioned an entirely new body of concert repertoire for the left hand alone, from composers including Maurice Ravel, Sergei Prokofiev, Benjamin Britten, and Paul Hindemith. The result was not a set of individual pieces compensating for specific missing capabilities. It was a new ecology: new techniques for implying two independent voices with one hand, new pedagogical traditions, a new compositional vocabulary specifically suited to the constraint, each piece extending what the next commission could attempt.

Claim 6.8.2 (The commissioned repertoire is a regenerative operator). *The commissions satisfy definition 6.6.14 rather than definition 6.6.12 alone: each addition to the repertoire increased not merely the immediately available performable material but the ecology’s own capacity to generate further left-hand-idiomatic material, in the sense*

that later commissions could draw on techniques the earlier ones had established. This is the distinction between fixing a structure and restoring the ecology that fixes structures, realized in a single, real, well-documented case rather than illustrated hypothetically.

Remark 6.8.3 (Why the progression matters). *The three stages are easy to conflate — all three involve “a pianist whose usual technique stops working” — and this chapter’s vocabulary is precisely what keeps them apart. Misalignment leaves the repertoire untouched and breaks only its correspondence to a relabeled space. Extinction removes part of the repertoire outright, with no relabeling involved. Regeneration is neither a fix nor a loss but the appearance of new generative capacity in response to one. A theory of repair that could not distinguish these three would not yet have earned the right to a general vocabulary of operators and ecologies; that it can distinguish them, using nothing beyond the definitions already given, is this example’s actual claim.*

6.9 Toward the Restorability Boundary

This chapter has argued that structures persist because operator ecologies persist, that the operative notion of sameness under damage is regenerative rather than behavioral equivalence, and that an ecology’s repair capacity depends not only on which operators remain executable but on how much of their generating history has been retained alongside them. None of this guarantees that repair capacity is unlimited. An ecology can lose a keystone operator outright; it can retain operators whose generating history has been compressed past the point of repairability; it can face a rate of damage that outpaces its capacity to respond regardless of how well-resourced any individual repair is. The next chapter is about what happens at that limit: not whether repair occurs, but where and why it stops.

Open Problem 6.9.1. *This chapter has not formally unified the Level 3 rule-update dynamics of definition 6.2.1, $F_{t+1} = G(F_t, x_t, E_t)$, with the executable reachability apparatus of definition 6.6.1 and Chapter 7’s $\text{Reach}(\mathcal{O}, x)$. Both describe an operator changing in response to conditions, but one is a dynamical update rule and the other is a set-valued reachability relation, and the relationship between G and the repair operators $o \in \mathcal{O}$ that realize it has not been made precise.*

Chapter 7

The Restorability Boundary

7.1 Introduction

A system under observation is rarely either fully known or fully unknown. More often it occupies an intermediate condition in which some deviations from its expected trajectory can be detected, diagnosed, and corrected, while others cannot. The boundary between these two conditions is the subject of this chapter. It is not a single mechanism but a coordinate at which several distinct mechanisms can independently place a system beyond the reach of its surrounding corrective ecology. Distinguishing those mechanisms from one another is the chapter's primary task, since each implies a different diagnosis of failure and a different remedy, and conflating them risks recommending a remedy that addresses the wrong mechanism entirely.

Restorability should not be understood as the return of a system to some prior, passively held state. In many of the systems this chapter draws on, the maintained state is itself an active process rather than a default condition. A neuron's resting potential is not simply what the membrane does when left alone; it is continuously re-established against a thermodynamic gradient by the sodium-potassium pump, which exchanges three sodium ions outward for two potassium ions inward at continuous metabolic cost. The resting potential is therefore not a passive baseline but a continuously repaired state, and if the pump's energy supply fails, the gradient does not remain stable while some other fault develops elsewhere: the distinction the gradient exists to maintain simply dissolves. The absence of repair, in such systems, does not merely block recovery from a prior failure. It can dissolve the very state that repair was maintaining. This is one reason the four failure modes developed below are given concrete non-computational instances wherever possible, rather than being treated only through the stellar-reconstruction case study: restorability is a property of maintained physical and biological systems as much as it is a property of inferential

ones.

This should not be overstated into a claim that all persistence is maintenance. A crystal lattice, a fossil, or a written inscription persists largely because it occupies a low-energy or kinetically trapped configuration, not because some operator continuously re-establishes it against an active gradient. It is more accurate to say that persistence generally decomposes into two components, $P = S \cup M$: persistence by stored energetic stability, S , and persistence by active maintenance, M . Systems can sit close to either pole — a crystal lives almost entirely in S , a membrane potential almost entirely in M — but many systems of later interest to this book, including institutional and civilizational ones, are mixtures rather than pure cases. A DNA sequence is stored information in the sense of S , yet its persistence across cell divisions depends on active proofreading and repair enzymes in the sense of M ; an archive is similarly a stored record that persists only under continuous physical and institutional maintenance. The interesting empirical claim, developed further in later chapters, is not that maintenance replaces storage as the general form of persistence, but that civilizational and biological structures cluster far more heavily toward M than intuition, which is calibrated on physically inert objects, tends to assume.

7.2 The Signal–Diagnosis–Repair–Continuation Chain

Every corrective process examined in this corpus can be decomposed into four stages through which a deviation must successfully pass before the system that produced it is considered restored. A deviation must first be registered as signal: some observable trace of the deviation must reach the surrounding ecology at all. That signal must then be diagnosed: the ecology must be able to distinguish the presence and character of the deviation from ordinary background variation. A diagnosed deviation must then be repaired: some operator capable of acting on the diagnosis must exist and be executable under current conditions. And a locally executed repair must finally be continued: the corrected local structure must compose with the surrounding structure in a way that does not reintroduce the deviation elsewhere or fail to persist once attention moves on.

Definition 7.2.1 (Restorability boundary). *Given a system X and a corrective ecology $\mathcal{E}$ acting on X through the chain of section 7.2, the restorability boundary $\partial_R(X, \mathcal{E})$ is the set of states of X beyond which $\mathcal{E}$ can no longer complete the chain: signal that is generated but not diagnosable, diagnosis that is accurate but not actionable, or repair that is executable locally but does not continue globally.*

The restorability boundary is deliberately defined without reference to which stage of the chain fails. This is the sense in which it is the general object: a system

found beyond ∂_R is uncorrectable, but the reason it is uncorrectable is left open until one of the four failure modes below is identified as the operative cause.

7.3 Two Kinds of Restorability Boundary

Definition 7.2.1 as stated is under-specified in two ways that matter before any of the four failure modes can be applied responsibly. Both corrections are adopted here from *Persistent Anomalies and the Geometry of Ontology Revision* (Flyxion, 2026), a standalone monograph developing a general theory of repair algebras and discrepancy hierarchies. This chapter cites that work rather than absorbing it: what follows imports exactly the two primitives needed to make definition 7.2.1 well posed, and leaves the monograph’s repair-algebra formalism, its discrepancy-type hierarchy, and its Persistence–Obstruction Schema as further reading rather than reproducing them here.

7.3.1 The architecture gate

The first under-specification is that definition 7.2.1 does not distinguish a target that remains representable in principle from a target that cannot be represented by the current architecture at all. Let $\pi : \mathcal{M} \rightarrow \Sigma$ be the prediction map of the representational architecture in use, so that $\text{Im } \pi \subseteq \Sigma$ is the set of observations that architecture can in principle produce.

Definition 7.3.1 (Architecture gate). *For an observation $x_{\text{obs}} \in \Sigma$, the architecture gate is*

$$\chi_\pi(x_{\text{obs}}) = \begin{cases} 0, & x_{\text{obs}} \in \text{Im } \pi, \\ 1, & x_{\text{obs}} \notin \text{Im } \pi. \end{cases}$$

A restorability boundary encountered at $\chi_\pi = 0$ is an inferential boundary: the target remains representable, but the available observer or repair ecology cannot presently recover it. A restorability boundary encountered at $\chi_\pi = 1$ is an ontological boundary: no admissible trajectory in the current representational architecture fits the observation, and no amount of additional repair within that architecture will resolve it.

This distinction matters because the four failure modes developed later in this chapter — more throughput, more diagnostic capacity, more operators, better global composition — are all remedies internal to a fixed architecture. None of them can close an ontological boundary. A system that persistently treats $\chi_\pi = 1$ conditions as if they were $\chi_\pi = 0$ conditions will apply an unbounded amount of the wrong kind of repair, which is precisely the trap the monograph’s central historical example is chosen to illustrate.

Example 7.3.2 (Mercury and Banik: an ontological/inferential contrast). *The nineteenth-century anomaly in Mercury’s perihelion precession survived every repair available within Newtonian celestial mechanics: hypothesized planets, atmospheric drag corrections, refined solar oblateness estimates. The residual was not merely undiagnosed; it was $x_{\text{obs}} \notin \text{Im } \pi_{\text{N}}$, where π_{N} is the Newtonian prediction map — no parameter choice within Newtonian mechanics could produce the observed precession. This is an ontological boundary, and it was resolved not by a stronger repair within Newtonian mechanics but by general relativity expanding the representational architecture itself. By contrast, the Banik reconstruction of example 7.4.5 is an inferential boundary only: the age-likelihood architecture is fixed throughout, and the crossover to the noise floor marks a point at which that specific reconstruction ceases to carry reliable information, not a demonstrated inadequacy of the underlying representational framework. It should be stated carefully that the noise floor is evidence for an inferential boundary; it is not evidence against an ontological one, since this chapter has not attempted to show that no enlarged inference within a similar general representation could recover more. The two cases are worth holding side by side precisely because they present a superficially similar symptom — a residual that will not go away — for two structurally unrelated reasons.*

7.3.2 History indexing

The second under-specification is that definition 7.2.1 treats unrestorability as an instantaneous condition. But whether a deviation is a first encounter, a repeated failure, a temporarily exhausted capacity, or a structurally stubborn residual cannot be determined from a single snapshot; it requires a remembered record of what has already been tried.

Definition 7.3.3 (History-indexed restorability boundary). *Let $H_t = (\varepsilon_0, R_1, \varepsilon_1, \dots, R_n, \varepsilon_n)$ denote the remembered sequence of residuals and attempted repairs up to time t , for a repair algebra $\mathcal{R}$ available to the ecology $\mathcal{E}$. The history-indexed restorability boundary is*

$$\partial_{R_t} = \partial_{\mathcal{R}}(X, \mathcal{E}, \mathcal{R}, H_t).$$

Without H_t , an ecology can detect that a deviation is currently unrestorable, but it cannot distinguish a first encounter from a repeated failure, a temporary capacity exhaustion from a persistent obstruction, or an untried repair from an invariant residual — distinctions the four failure modes below depend on whenever they are used to diagnose *recurrent*, *worsening*, or *structurally stubborn* conditions rather than momentary ones. The four modes are accordingly understood throughout the rest of this chapter as evaluated relative to H_t : the Expiatory Gap becomes $G_t = D_{\text{sys}}(H_t) - D_{\text{max}}(H_t)$; diagnostic collapse becomes the history-conditioned mutual information $I_t(F; E \mid H_t) \rightarrow 0$; repair deficit records which operators have already been attempted, found unavailable, or rendered

non-executable within H_t ; and continuation failure records whether repeated local sections fail to glue across changing covers as the history unfolds. This chapter does not revise the four definitions given below to make this history-dependence explicit in every equation, since doing so before the underlying repair-history formalism is developed further would overstate what has actually been established; it is stated here as the governing convention under which those definitions should be read.

A history-indexed boundary is not yet a diagnostic procedure. Nothing so far says how, from H_t alone, an ecology should actually conclude that a residual is persistent rather than merely unresolved so far. That inferential step needs its own construction.

Definition 7.3.4 (Residual orbit and persistent residual). *Let $\mathcal{R}$ be a repair algebra acting on a discrepancy space $\mathcal{D}$. For a repair word $w = R_n \circ \dots \circ R_1 \in \langle \mathcal{R} \rangle$, write $\varepsilon_w = w(\varepsilon_0)$. The repair orbit of ε_0 is $\text{Orb}_{\mathcal{R}}(\varepsilon_0) = \{w(\varepsilon_0) : w \in \langle \mathcal{R} \rangle\}$. For a discrepancy norm $\|\cdot\|$, define $P(\varepsilon_0, \mathcal{R}) = \inf_{w \in \langle \mathcal{R} \rangle} \|w(\varepsilon_0)\|$. The discrepancy is repairable to tolerance η when $P(\varepsilon_0, \mathcal{R}) \leq \eta$, and persistent above tolerance η when $P(\varepsilon_0, \mathcal{R}) > \eta$.*

H_t records only attempted words, not every word in $\langle \mathcal{R} \rangle$, so it can only approximate $P(\varepsilon_0, \mathcal{R})$ from above.

Definition 7.3.5 (Empirical persistence bound). *For a history H_t containing attempted repair words $w_1, \dots, w_n$, define $\hat{P}_t = \min_{1 \leq k \leq n} \|w_k(\varepsilon_0)\|$.*

Claim 7.3.6 (Monotonicity of the empirical persistence bound). *As repair history grows, $\hat{P}_{t+1} \leq \hat{P}_t$, and $P(\varepsilon_0, \mathcal{R}) \leq \hat{P}_t$ for every finite history.*

Proof. The attempted set at $t + 1$ contains the set attempted by t , so a minimum over the larger set cannot exceed the minimum over the smaller one. Since the full repair monoid contains every attempted word, the infimum over the full monoid cannot exceed the minimum over any finite attempted subset. $\square$

Claim 7.3.6 exposes a limitation worth stating plainly: *repeated failure does not by itself prove structural persistence unless the attempted repair family is sufficiently exhaustive.* “Every repair tried so far has failed” and “every repair available in the architecture must fail” are different claims, and only a completeness condition closes the gap between them.

Claim 7.3.7 (Persistence Identification under Exhaustive Search). *Suppose the attempted repair words become dense in $\langle \mathcal{R} \rangle$ under a topology for which $w \mapsto \|w(\varepsilon_0)\|$ is continuous. Then $\hat{P}_t \rightarrow P(\varepsilon_0, \mathcal{R})$.*

Remark 7.3.8 (Sharpening Mercury and Banik). *This distinction refines example 7.3.2’s contrast. The Newtonian astronomers who exhausted planetary hypotheses, drag corrections, and oblateness estimates against Mercury’s perihelion were not merely accumulating $\hat{P}_t$; the historical record suggests their attempted repair family came close to the density condition of Claim 7.3.7 within the Newtonian algebra $\mathcal{R}_N$, which is what licenses reading $P_{\mathcal{R}_N}(\varepsilon_{\text{Mercury}}) > 0$ as a genuine ontological boundary rather than an artifact of insufficient search. Nothing comparable has been shown for the Banik reconstruction: its noise floor reflects $\hat{P}_t$ for one specific reconstruction procedure, with no argument that the attempted procedures are dense in any relevant sense, which is precisely why example 7.3.2 was careful to call it an inferential boundary only.*

Remark 7.3.9 (Why H_t is load-bearing, not merely convenient). *Chapter 6 gives a reason H_t cannot be treated as an optional refinement that a well-run ecology could reasonably do without. Its Anti-Alignment result (Chapter 6, Claim 5.7, imported there from History Before Function) shows that the minimal trace $T_{\min}$ needed merely to enact an operator is generically the worst-positioned trace for repairing that operator later, strictly worse whenever the discarded surplus carried diagnostic information. An ecology under ordinary pressure to compress — to retain only what current executability requires — is therefore not neutral with respect to H_t : it is systematically discarding exactly the record this chapter’s history-indexed boundary depends on, as an unavoidable consequence of efficient operation rather than through neglect. This sharpens what “without H_t ” means in the paragraph above. It is not merely that an ecology which happens not to keep records cannot make the finer diagnostic distinctions listed there; it is that keeping the enactment of its operators efficient and keeping H_t available are, by Chapter 6’s result, in tension by construction. A history-indexed restorability boundary is therefore not a bookkeeping addition to an otherwise complete theory of repair. It names a resource — retained, undiscarded history — that ordinary operation continuously depletes.*

Remark 7.3.10 (What remains open). *This section does not claim that the four failure modes of section 7.4 reduce to, or are systematically organized by, the architecture gate or the discrepancy-type hierarchy of Persistent Anomalies. The relation between the two classifications remains open. An in-architecture numerical error can still produce an Expiatory Gap if observations simply arrive faster than they can be processed; an out-of-architecture categorical mismatch can remain highly discriminable and therefore need not involve diagnostic collapse at all; a transient discrepancy can encounter a temporary repair deficit because the relevant operator exists but happens not to be executable at the moment it is needed. None of these cases is naturally derived from a fixed cell of $\{\chi_\pi\} \times \{\text{discrepancy type}\} \times \{\text{persistence character}\}$, which suggests the four modes track operational coordinates —demand versus capacity, discriminability, operator availability, local-to-global composability— that may be orthogonal to, rather than subordinate to, the architecture-level classification. Establishing whether the two*

classifications are orthogonal, partially reducible, or connected only under additional assumptions is left as a separate comparison problem.

The four modes' mutual independence, established above at the level of definitions, does not mean they cannot cause one another dynamically. One such coupling admits a direct formal treatment.

Claim 7.3.11 (Throughput-Induced Information Bound). *Let signal arrive at rate D_{sys} while the diagnostic ecology processes at most D_{max} , and let $\rho = \min(1, D_{\text{max}}/D_{\text{sys}})$ be the retained fraction. Under independent random thinning at retention probability ρ — discarded items uncorrelated with which observations are diagnostically valuable — and conditional independence of feedback items, $I(\tilde{F}; E) \leq \rho I(F; E)$ for the thinned feedback $\tilde{F}$. Consequently, as $D_{\text{max}}/D_{\text{sys}} \rightarrow 0$, $I(\tilde{F}; E) \rightarrow 0$.*

Claim 7.3.11 gives a first formal coupling arrow, Expiatory Gap $\Rightarrow$ Diagnostic Collapse under lossy thinning, while preserving their analytical distinction: the reverse implication does not hold, since $I(F; E) \rightarrow 0$ does not imply $D_{\text{sys}} > D_{\text{max}}$ — a low-volume but systematically misleading channel can collapse diagnostically with no throughput problem at all. This resolves one direction of one coupling among the four modes; the remaining couplings, and the relation to the architecture gate raised above, are not addressed by this result and remain open.

7.4 Four Failure Modes

7.4.1 Expiatory Gap: throughput failure

The Expiatory Gap names the condition in which a system generates signal, complexity, or consequence faster than the surrounding ecology can absorb and process it, independent of whether that ecology could in principle diagnose the signal correctly given sufficient time.

Definition 7.4.1 (Expiatory Gap). *Let D_{sys} denote the rate at which a system generates diagnosable content and D_{max} denote the maximum rate at which the receiving ecology can process such content while retaining diagnostic reliability. The Expiatory Gap is*

$$G = D_{\text{sys}} - D_{\text{max}}.$$

Claim 7.4.2. *The Expiatory Gap is, in principle, closable by rate control alone: any combination of reduced emission rate, increased receiver bandwidth, staging, buffering, or temporal dilation that brings effective throughput within D_{max} restores the diagnostic reliability of the channel.*

Claim 7.4.2 is the load-bearing assumption behind the “natural damping” argument developed for this failure mode, and it is exactly the assumption that fails for the next failure mode.

7.4.2 Diagnostic Collapse: discriminability failure

Diagnostic collapse names the condition in which a feedback channel remains active, and may even remain high in volume, while ceasing to carry information about the error state it is meant to track.

Definition 7.4.3 (Diagnostic collapse). *Let F denote feedback supplied by a corrective ecology and E denote the true error state of the system under correction. The channel is diagnostically active when $I(F; E) > \varepsilon$ for some operationally meaningful threshold ε . Diagnostic collapse is the condition*

$$I(F; E) \longrightarrow 0$$

independent of the volume, frequency, or apparent sophistication of F .

Claim 7.4.4. *Diagnostic collapse is not closable by rate control. Increasing the volume of feedback, or the time allotted to produce it, leaves $I(F; E)$ unchanged when the feedback source lacks the capacity to discriminate correct from incorrect states. Closing the gap requires a qualitatively new diagnostic channel: additional expertise, a new class of observation, or a restructured overlap graph between the system’s active frontier and the population capable of checking it.*

Example 7.4.5 (The oldest star as a diagnostic-collapse boundary). *Banik et al. (2026) reconstruct the latent age distribution $P(A)$ of the oldest stars in a sample of 155,600 Milky Way subgiants, using the individual age likelihoods $\mathcal{L}_i(A)$ obtained from LAMOST DR7 spectroscopy and Gaia astrometry. The reconstructed $P(A)$ never reaches zero at any finite age, since positivity constraints and finite sampling guarantee a nonzero posterior at every grid point. The age of the oldest star, A_* , is instead located by examining the ratio $P(A)/\sigma(A)$ across Markov Chain Monte Carlo samples: below some age this ratio declines approximately linearly, and beyond it the ratio flattens into a noise floor that does not further decline no matter how the reconstruction is refined. The crossover between these two regimes, found at $A_* = 13.73^{+0.18}_{-0.15}$ Gyr, is a directly measured instance of definition 7.2.1: not a point at which signal ceases, but a point at which structured variation in the reconstruction gives way to variation indistinguishable from the reconstruction process’s own uncertainty.*

The Banik example is worth stating as a general principle in its own right, since it is the clearest available instance of the distinction this chapter depends on.

Claim 7.4.6 (The boundary is discriminability, not disappearance). *The restorability boundary is generally not the point at which a signal disappears, but the point at which the distinction between signal and reconstruction uncertainty ceases to be recoverable. A magnitude threshold asks whether a quantity has fallen to zero; a discriminability threshold asks whether variation in that quantity can still be distinguished from the noise generated by the process used to recover it. The two coincide only by coincidence.*

The four failure modes of this chapter are analytically distinct but not claimed to be dynamically independent, and a taxonomy that implied otherwise would misdescribe how these failures actually arise in practice. Sustained throughput overload can itself induce diagnostic collapse, since an ecology forced to process signal faster than $D_{\max}$ for long enough may lose the capacity to sort useful content from noise even after the overload subsides. Conversely, diagnostic collapse can present as apparent throughput overload, since an ecology that can no longer discriminate relevant from irrelevant signal will tend to treat a constant volume of input as increasingly unmanageable, not because more of it is arriving but because less of what arrives can be filtered. The four modes should therefore be read as four independent quantities that can each run out, any of which may trigger or mask the depletion of another, rather than as four mutually exclusive diagnoses of a single failure.

7.4.3 Repair Deficit: operator failure

Repair deficit names the condition in which diagnosis succeeds — the deviation is correctly identified and localized — but no operator capable of acting on that diagnosis exists or remains executable.

Definition 7.4.7 (Repair deficit). *A diagnosed deviation d at coordinate x suffers a repair deficit when the operator ecology $\mathcal{O}$ contains no operator $o \in \mathcal{O}$ such that o is executable at x and o corrects d .*

This definition is formally correct but comparatively thin next to the quantities available for the other three failure modes: a rate for the Expiatory Gap, a mutual information for diagnostic collapse, a gluing kernel for continuation failure. A more comparable formulation defines an executable reachability relation over the operator ecology directly.

Definition 7.4.8 (Executable reachability). *Let $\text{Reach}(\mathcal{O}, x)$ denote the set of deviations correctable from coordinate x by some composition of operators in $\mathcal{O}$ that remain executable at x under current conditions. Repair deficit for a diagnosed deviation d at x is then the condition*

$$d \notin \text{Reach}(\mathcal{O}, x).$$

Definition 7.4.8 makes repair deficit directly comparable to the other three failure modes and, more importantly, makes explicit that repair deficit is a property of the pair $(\mathcal{O}, x)$ rather than of d alone: the same deviation can be a repair deficit at one coordinate and fully correctable at another, depending on which operators remain executable there.

Repair deficit is the point at which this chapter’s chain interfaces directly with the operator-ecology framework of Chapter 6. A correctly diagnosed deviation that outpaces available operators is, in ecological terms, the loss of a keystone operator: the niche is identified with precision, but nothing remains to occupy it. This suggests that repair deficit should ultimately be treated jointly with that material rather than developed independently here.

Definition 7.4.8 treats “executable at x ” as a primitive, but repair deficit can arise for at least two structurally different reasons that this primitive conflates. Let $\text{Exec}(x) \subseteq \mathcal{O}$ denote the operators that can actually run at coordinate x under current conditions.

Definition 7.4.9 (Structural and energetic repair deficit). *A deviation d at x suffers a structural repair deficit when no operator capable of correcting d exists in the ecology at all, that is, there is no $o \in \mathcal{O}$ that corrects d regardless of executability. It suffers an energetic repair deficit when a corrective operator exists but cannot run at x : there is some $o \in \mathcal{O}$ that corrects d , yet $o \notin \text{Exec}(x)$.*

Example 7.4.10 (Energetic repair deficit: the sodium-potassium pump). *The neuronal membrane potential is not a passive baseline but a continuously repaired state, maintained by the sodium-potassium pump exchanging three sodium ions outward for two potassium ions inward at continuous metabolic cost. The corrective operator here is never structurally absent: the pump exists as long as the cell exists. What can fail is executability. If ATP production fails, the pump satisfies $o_{\text{pump}} \in \mathcal{O}$ but $o_{\text{pump}} \notin \text{Exec}(x)$, and the deviation — the gradual collapse of the ion gradient toward equilibrium — falls out of $\text{Reach}(\mathcal{O}, x)$ not because no operator was ever capable of correcting it, but because the operator’s precondition failed. This is repair deficit of the energetic kind, and it is worth keeping conceptually separate from the central-pattern-generator example of example 7.4.13, which illustrates the structural side of the same definition: redundancy among operators that remain executable, rather than the failure of a single operator’s executability.*

Open Problem 7.4.11. *Definition 7.4.9 may not be exhaustive. A third candidate, not yet formalized here, is kinetic repair deficit: an operator o with $o \in \text{Exec}(x)$ that nonetheless fails to correct d because its completion time exceeds the timescale over which d compounds or propagates. This would be a recurrence of the Expiatory Gap’s rate-mismatch logic ($D_{\text{sys}} > D_{\text{max}}$ in section 7.4.1) relocated from the diagnostic channel to the repair operator itself, and its existence raises a question this chapter does not yet*

resolve: whether throughput failure is properly a phenomenon unique to the Signal $\rightarrow$ Diagnosis boundary, or a failure pattern (rate exceeding capacity) that can recur at any stage of the chain, of which the Expiatory Gap as defined in section 7.4.1 is only the first instance.

Reachability as stated in definition 7.4.8 is binary: a deviation is either reachable or it is not. This conceals a distinction that matters in practice, since a deviation reachable by exactly one operator and a deviation reachable by ten independent operators are both formally free of repair deficit, yet occupy very different positions with respect to future risk. The repair reserve refines definition 7.4.8 accordingly.

Definition 7.4.12 (Repair reserve). *For a diagnosed deviation d at coordinate x , the repair reserve is*

$$R(d, x) = |\{ o \in \mathcal{O} : o \text{ executable at } x, o \text{ corrects } d \}|.$$

Repair deficit as defined in definition 7.4.8 is then the special case $R(d, x) = 0$, and $d \in \text{Reach}(\mathcal{O}, x)$ if and only if $R(d, x) > 0$. Values of $R(d, x)$ greater than one indicate overlapping, mutually substitutable corrective pathways rather than a single point of failure.

Example 7.4.13 (Central pattern generators and repair reserve). *Locomotor gait in most vertebrates is not produced by a single controller but by a nested hierarchy of corrective operators acting at different scales: cellular oscillators, local spinal circuits, limb-level coordination, gait-level coordination, and cortical modulation. Each layer corrects deviations at its own scale before they can propagate to the next. Consequently $R(d, x)$ for a typical gait deviation is well above one even when a high-level operator is lost: an animal that loses cortical control can often still walk, because lower-level operators remain executable and continue to correct the deviation. Repair deficit for gait — the condition $R(d, x) = 0$ — typically requires the loss of operators across several layers simultaneously, not the loss of any single controller. The example illustrates that healthy biological systems do not merely avoid repair deficit; they operate with $R(d, x) \gg 1$ as a matter of course, and it is the erosion of that reserve, rather than its first crossing of zero, that constitutes the early warning signal of an approaching repair deficit.*

Remark 7.4.14 (Successful repair can still collapse fault identity). *$R(d, x) > 0$ certifies that some executable operator corrects d ; it says nothing about whether the act of repairing preserves the identity of what was wrong. This concern is not idle: a repair operator is itself a transition in the sense of Chapter 9's framework, and repair operators built to accommodate anomalous evidence with minimal disruption generically collapse fibers — distinct prior faults $f_1 \neq f_2$ routinely satisfy $R(f_1) = R(f_2)$, precisely because a minimality criterion is more responsible for the repaired result than the specific defect*

that triggered it. Admissibility can therefore be restored while historical recoverability is lost in the same stroke: $\text{repair} \not\Rightarrow \text{historical recoverability}$. This is not a failure the four modes of this chapter currently register, since a repair satisfying $R(d, x) > 0$ counts as successful by definitions 7.4.8 and 7.4.12 regardless of whether it preserves fault identity. It bears directly on the residual-orbit machinery of section 7.3: a repair word w that reduces $\|\varepsilon\|$ to an acceptable tolerance has not thereby shown that the identity of the original discrepancy survives in H_t , and definition 7.3.4's apparatus tracks the residual's magnitude along the orbit, not whether successive repairs are individually distinguishable from one another after the fact. This is recorded here as an open phenomenon this chapter's taxonomy does not yet have a formal home for, rather than folded into one of the four existing modes by assumption.

Principle 7.4.15 (Repair–Recovery Separation).

$\text{Repair} \not\Rightarrow \text{Historical Recoverability}$.

Restoring admissibility and preserving the identity of what was wrong are independent achievements. An operator satisfying $R(d, x) > 0$ certifies the first; nothing in this chapter's taxonomy, on its own, certifies the second. This is named separately from remark 7.4.14 because it recurs beyond this chapter: Chapter 6's minimality-driven repair operators and Chapter 9's fiber-collapsing repair transitions are independent instances of the same separation, not three unrelated observations that happen to share a conclusion.

7.4.4 Continuation Failure: composition failure

Continuation failure names the condition in which local repair succeeds but the corrected structure does not compose with its surroundings into a globally coherent whole.

Let $\mathcal{F}$ be a sheaf of local structure over a covered space $(X, \{U_i\})$, and let each local repair yield a section $s_i \in \mathcal{F}(U_i)$. An earlier version of this section wrote continuation failure as the nonvanishing of a kernel from $\prod_i \mathcal{F}(U_i)$ to $\Gamma(X, \mathcal{F})$, but that arrow is not generally canonical: global sections restrict to local ones, not the reverse, so a map running the other way cannot be the right object to take a kernel of. The correct formulation distinguishes two failure modes the earlier notation conflated.

Definition 7.4.16 (Disagreement map). Define $\delta : \prod_i \mathcal{F}(U_i) \rightarrow \prod_{i,j} \mathcal{F}(U_i \cap U_j)$ by $\delta((s_i)_i) = (s_i|_{U_i \cap U_j} - s_j|_{U_i \cap U_j})_{i,j}$ for an abelian sheaf, or the corresponding pair of restrictions in the set-valued case.

Claim 7.4.17 (Local Compatibility Criterion). A family $(s_i)_i$ of local repairs is compatible on overlaps exactly when $\delta((s_i)_i) = 0$. If $\mathcal{F}$ is a sheaf, every compatible family glues uniquely to a global section $s \in \Gamma(X, \mathcal{F})$ with $s|_{U_i} = s_i$ for all i .

Proof. $\delta((s_i)_i) = 0$ states precisely that s_i and s_j agree on every overlap. Existence and uniqueness of the global section is the sheaf gluing axiom. $\square$

Definition 7.4.18 (Compatibility failure and descent failure). *A family of locally successful repairs exhibits compatibility failure when $\delta((s_i)_i) \neq 0$: the repairs disagree already on some overlap. For a presheaf not known to satisfy the sheaf axiom, a compatible family ($\delta((s_i)_i) = 0$) exhibits descent failure when no global section restricts to it. For a genuine sheaf, descent failure cannot occur — Claim 7.4.17 rules it out by definition, which is why obtaining a nontrivial global obstruction despite pairwise compatibility requires working with a presheaf, with higher cocycle data, or with nonlinear constraints not captured by pairwise equality alone.*

Claim 7.4.19 (Cohomological Continuation Obstruction). *Suppose local repairs determine a Čech 1-cocycle $g = (g_{ij}) \in \check{Z}^1(\mathcal{U}, \mathcal{G})$ of transition corrections on overlaps. A globally coherent repair exists whenever the class $[g] \in \check{H}^1(\mathcal{U}, \mathcal{G})$ vanishes. A nonzero class is an obstruction to continuation.*

Proof. A global repair corresponds to local reparameterizations h_i satisfying $g_{ij} = h_j h_i^{-1}$ on each overlap, i.e., to g being a coboundary. Such h_i exist exactly when the cohomology class of g vanishes. $\square$

Continuation failure in the sense this chapter needs is therefore best stated as compatibility failure ($\delta((s_i)_i) \neq 0$) for the basic case, with the cohomological obstruction of Claim 7.4.19 available for the genuinely global case where pairwise-compatible local repairs still fail to assemble.

Example 7.4.20 (Cardiac synchronization as continuation failure). *The heartbeat is not produced by a single clock winding down but by a population of cellular oscillators in the sinoatrial node continually re-synchronizing through gap-junction coupling. Each oscillator, taken locally, is a valid section: it continues to fire rhythmically on its own terms. A coherent heartbeat requires these local sections to glue into a single global rhythm through sufficient electrical coupling between cells. When coupling weakens, arrhythmia can result even though every individual oscillator remains functional and every local repair, in isolation, continues to succeed. The failure is not that signal generation stops or that diagnosis degrades; it is compatibility failure in the sense of definition 7.4.18: adjacent oscillators' phases disagree across the coupling that should reconcile them, $\delta((s_i)_i) \neq 0$, and local success no longer implies global coherence.*

Open Problem 7.4.21. *The treatment of the “existential expiatory gap” in sheaf-theoretic terms develops precisely this gluing obstruction, but files it under the throughput-failure vocabulary of section 7.4.1. Determine whether that material is more accurately relocated to this section, and whether its interpretation changes once continuation failure is treated as a distinct failure mode rather than a symptom of an unmanaged Expiatory Gap.*

7.5 Relating the Four Modes

Table 7.1 places the four failure modes side by side. The central methodological claim of this chapter is that no single remedy addresses all four: a system exhibiting diagnostic collapse will not be repaired by measures designed for throughput failure, and a system exhibiting continuation failure may show perfectly healthy local diagnosis and repair at every individual site while still failing as a whole.

Failure mode	Lost quantity	Formal signature	Remedy
Expiatory Gap	Throughput	$G = D_{\text{sys}} - D_{\text{max}}$	Rate control
Diagnostic collapse	Information	$I(F; E) \rightarrow 0$	New diagnostic capacity
Repair deficit	Operators	$d \notin \text{Reach}(\mathcal{O}, x)$	New operators
Continuation failure	Composability	Compatibility failure, $\delta \neq 0$	Global re-coordination

Table 7.1: The four failure modes of the restorability boundary, organized by which quantity is exhausted rather than by which stage of the chain is affected. This framing generalizes beyond this chapter: for any observed instance of a system failing to recover, the diagnostic question is which of throughput, information, operators, or composability is actually scarce, since the answer determines which of the four remedy classes applies and rules out the other three.

Example 7.5.1 (Scale-dependent neural organization). *Posani et al. (2026) provide an instructive case in which the apparent organization of a system changes with descriptive scale. Analyzing neural responses across the cortical hierarchy, they find that individual neurons are rarely cleanly categorical: local response profiles are heterogeneous and commonly mix several task-related variables. At the population level, however, those distributed responses can remain highly separable, so that relevant variables are recoverable from the collective geometry even though they are not assigned to sharply specialized local units.*

This is not itself a measurement of restorability, and Chapter 5’s comparison already establishes that the paper’s participation ratio should not be identified with this book’s reachability quantities. What the example exposes instead is a scale dependence a theory of restorability must permit for. At the single-unit scale, the absence of categorical specialization may look like weak localization: there is no unique component whose inspection or replacement corresponds to one represented variable. At the population scale, the same system may possess substantial diagnostic separability, because information is distributed across decorrelated response directions rather than concentrated in any one of them. A failure can therefore appear differently under different covers or levels of description. Damage to one locally mixed component need not destroy the population-level distinction, since other components participate in representing it; conversely, a pertur-

bation aligned with a low-dimensional population direction can erase a globally recoverable distinction while leaving many individual units active and, by any local measure, undamaged.

Remark 7.5.2 (The slogan this example licenses). *example 7.5.1 supports a compact statement worth having on hand: restorability depends on separability, not categoricity. A system does not need to sort its states into clean, labeled categories to support restoration — Posani et al.’s own finding is that cortex mostly does not. It needs enough geometric structure, at whatever scale diagnosis is actually performed, that the distinctions repair depends on remain recoverable. Diagnostic collapse (section 7.4.2) is the failure of exactly this condition, stated earlier in this chapter in information-theoretic terms; this example gives it a second, geometric reading.*

Remark 7.5.3 (A graded refinement of the architecture gate). *example 7.5.1 also sharpens the architecture gate of definition 7.3.1. That gate is binary: $\chi_\pi(x) \in \{0, 1\}$, according to whether $x \in \text{Im}(\pi)$. Posani’s result shows the interesting structure can live entirely on the $\chi_\pi = 0$ side of that boundary and still admit further distinctions: a represented state can be highly separable, weakly separable, or effectively inseparable from its neighbors, depending on the local geometry of the representation at the scale it is queried. An architecture can contain a phenomenon — $x \in \text{Im}(\pi)$ — while still lacking the geometry to isolate it from nearby alternatives. This does not revise definition 7.3.1; it observes that the gate answers a coarser question than the one diagnosis usually needs answered, and that a graded notion of diagnostic accessibility inside $\text{Im}(\pi)$ is a real further distinction this chapter has not formalized.*

Claim 7.5.4 (Scale relativity of restorability). *A system may lie inside the restorability boundary relative to one descriptive scale and outside it relative to another. For fine and coarse observation maps $q_f : X \rightarrow X_f$ and $q_c : X \rightarrow X_c$ with a scale map $p : X_f \rightarrow X_c$ relating them, the corresponding effective corrective ecologies $\mathcal{E}_f$ and $\mathcal{E}_c$ need not satisfy $p(\partial_{R_f}) = \partial_{R_c}$, and restorability at one scale need not imply restorability at the other.*

Coarse-graining can move a system relative to the restorability boundary in either direction, and for different reasons in each. It can hide a local failure by merging it with states that remain functionally equivalent at the coarser scale — the population-level separability of example 7.5.1, which survives damage no single fine-scale unit could survive by itself. Or it can destroy the very distinction a diagnosis or repair needs, by collapsing exactly the fine-scale difference the coarse scale was never built to track — the reverse failure, in which population-level structure conceals rather than compensates for a fine-scale problem. Neither direction is the default; which one occurs depends on which distinctions a given coarsening removes, which is why Claim 7.5.4 states an inequality rather

than a direction. The operative restorability boundary for a system is accordingly not a property of the system alone. It depends on which scale diagnosis and repair are actually attempted at, a dependency none of the four failure modes developed earlier in this chapter make explicit on their own.

7.6 Worked Examples Across the Four Modes

The four failure modes are illustrated across this chapter with examples chosen for concreteness rather than uniformity of domain: the stellar reconstruction of example 7.4.5 demonstrates diagnostic collapse in a statistical inference setting, the sodium-potassium pump discussed in the introduction demonstrates that continuous repair can constitute the maintained state itself rather than a response to its failure, the central-pattern-generator hierarchy of example 7.4.13 demonstrates that repair deficit is generally a property of a coordinate within a distributed operator ecology rather than of a single missing controller, and the cardiac synchronization example of example 7.4.20 demonstrates that continuation failure can occur with every local component fully functional. The scale-dependent example of example 7.5.1 adds a fifth lesson orthogonal to the other four: restorability boundaries are not scale-invariant, and a system's apparent position relative to one can reverse depending on the descriptive resolution at which repair is attempted. Taken together, these examples are intended to make the taxonomy feel like an empirical regularity across domains rather than a formal device native only to reconstruction theory.

7.7 Open Problems for This Chapter

Several matters remain unresolved in this skeleton and are flagged here rather than resolved prematurely. The relationship between repair deficit and the operator-ecology material has been sharpened through the reachability formulation of definition 7.4.8, but still needs full integration with that essay's keystone-operator vocabulary rather than a side reference. The sheaf-theoretic material currently housed in the companion Expiatory Gap essay's §20 needs a decision about relocation to section 7.4.4 now that continuation failure has its own worked example and no longer needs to borrow that essay's treatment by default. And while each failure mode now has at least one concrete non-computational instance, a fully worked civilizational or institutional example — as opposed to a biological one — for repair deficit and continuation failure would strengthen the chapter's claim to generality beyond the inferential and physiological domains it currently draws on.

Chapter 8

Continuation

8.1 Introduction: A Short Chapter, Predicted in Advance

The Admissibility chapter closed by stating that a chapter on continuation would not need new foundational machinery, only trajectories formalized properly, its pointwise-reduction theorem restated, and a comparison between geometric repair and operator ecologies. This chapter does those three things and one more: it resolves an apparent conflict between two sources this corpus draws on, which the Admissibility chapter’s brevity did not have room to address.

8.2 Trajectories and Coordination Geometries

Definition 8.2.1 (Continuation trajectory). *A continuation trajectory is a pair sequence $(x(t), \mathcal{G}(t))$, where $x(t)$ is a state in the sense of earlier chapters and $\mathcal{G}(t)$ is a coordination geometry supplying a CLIO projection $\pi_{\mathcal{G}(t)}$ at each t , following the construction in Coordination Geometry.*

Claim 8.2.2 (Pointwise reduction). *For fixed $\mathcal{G}(t)$ at each instant, “ $x(t)$ is continuation-admissible” is equivalent to $S_{\text{RSVP}}(\pi_{\mathcal{G}(t)}(x(t))) \leq 0$, the entropy-monotonicity admissibility condition of the Admissibility chapter, evaluated at that instant. This introduces no structure beyond that chapter’s static relation.*

Claim 8.2.2 is stated here exactly as it was in the Admissibility chapter, for a reason that becomes important in section 8.3: it is a claim about checking admissibility at a *fixed* $\mathcal{G}(t)$, and says nothing yet about how $\mathcal{G}(t)$ itself changes.

8.3 Checking Versus Authoring

Histories Become Operators makes a claim that, read against Claim 8.2.2 without care, looks like a direct contradiction: admissibility is not a passive condition fixed in advance by a trace, but something “the continuer partly authors in the act of continuing.” Reading is *admissibility-constituting*, on this account, not *admissibility-revealing*. If Claim 8.2.2 already reduces continuation-admissibility to checking a static inequality, what is left for a continuer to author?

Claim 8.3.1 (Resolution: checking and authoring operate at different levels). *Claim 8.2.2 is a claim about the relation between $x(t)$ and $\mathcal{G}(t)$ at a single, frozen instant: given $\mathcal{G}(t)$, whether $x(t)$ is admissible is a fact to be checked, not authored. It is silent about how $\mathcal{G}(t+1)$ is obtained from $\mathcal{G}(t)$, $x(t)$, and the accumulated history H_t . That transition — geometric repair, in Coordination Geometry’s own terms, triggered whenever $S_{\text{RSVP}}(\pi_{\mathcal{G}(t)}(x(t))) > 0$ — is exactly where *Histories Become Operators*’ authorship claim applies. Checking is pointwise and static; authoring is the reorganization of the geometry doing the checking, and is neither pointwise nor static.*

The apparent contradiction dissolves once “continuation-admissibility” is recognized as ambiguous between two questions: whether the present state satisfies the present geometry’s condition (checking, settled by Claim 8.2.2), and how the geometry itself came to be the one now doing the checking (authoring, not addressed by that theorem at all). Both sources are right about the question each was actually answering.

8.4 Geometric Repair as Level-3 Modification

Chapter 6 already gave a name to exactly this kind of transition.

Claim 8.4.1 (Geometric repair is a Level-3 modification). *Recall Chapter 6’s three-level adaptation hierarchy: Level 1 state update $x_{t+1} = F(x_t)$, Level 2 parameter update $x_{t+1} = F_{\theta_t}(x_t)$, and Level 3 rule update $F_{t+1} = G(F_t, x_t, E_t)$, in which the mapping governing future behavior is itself the object modified. Geometric repair, $\mathcal{G}(t) \rightarrow \mathcal{G}(t+1)$ triggered by an admissibility violation, is a Level 3 modification: $\mathcal{G}(t)$ plays the role of F_t , and the trigger condition plays the role of E_t .*

This is not a new observation about geometric repair so much as a recognition that Chapter 6’s hierarchy already had room for it. It also sharpens what “authoring” means in Claim 8.3.1: authorship is Level 3 activity, in a sense already made precise, not a vague appeal to agency or interpretation.

8.5 Geometric Repair and Operator-Ecology Regeneration

Chapter 6 defined a regenerative operator as one that increases an ecology's future capacity to produce repairs, rather than merely repairing a single structure. Geometric repair looks, on its face, like exactly this kind of operation applied to coordination geometries instead of operator ecologies. Whether it actually is requires the same discipline this book has applied to every other resemblance of this shape.

Remark 8.5.1 (Family resemblance, not identity). *Chapter 6's regenerative operator g acts on an operator ecology $\mathcal{E}$, with the payoff measured in $\text{Reach}(g(\mathcal{E}))$ or induced reachability volume. Geometric repair acts on a coordination geometry $\mathcal{G}$, with the payoff measured by whether $S_{\text{RSVP}}(\pi_{\mathcal{G}'}(x)) \leq 0$ is restored. These fail the definition-by-substitution test this book applies before treating any two same-shaped constructions as identical: passing a type test — both are operators that act on the acting-structure rather than on the state directly, which is a genuine structural parallel — is not the same as passing a substitution test, and no map has been constructed showing a coordination geometry is a special case of an operator ecology, or the reverse. They are cousins under Chapter 6's Level 3 category, not shown to be the same construction wearing different notation.*

8.5.1 The Induced Ecology, Type-Corrected

An admissible reprojection $\pi_{\mathcal{G}} : X \rightarrow X'$ is a map between spaces, in the sense of Chapter 1's projection and embedding operations. It is not, on its own, an element of an operator repertoire in Chapter 6's sense: the operators composed in $\mathcal{R}_{\mathcal{E}}(x) = \bigcup_n \mathcal{R}_{\mathcal{E}}^{(n)}(x)$ must be self-maps on one fixed space, since each stage of a composition $o_k \circ \dots \circ o_1(x)$ has to land back inside the domain of the next operator. Taking $\{\pi_{\mathcal{G}}\}$ itself as the repertoire of an induced ecology therefore fails a type test before any proposition about it can be stated.

What is genuinely Level 3 in definition 6.2.1's sense is not $\pi_{\mathcal{G}}$ but $\mathcal{G}$: $\pi_{\mathcal{G}}$ is the fixed rule checked against at an instant, and the transition $\mathcal{G}(t) \rightarrow \mathcal{G}(t+1)$, triggered by an admissibility violation, is the object being modified. The induced ecology is accordingly built one level up, on the augmented space $Y = X \times \mathfrak{G}$ of state-geometry pairs, with the geometry-repair transitions themselves as its operators.

Definition 8.5.2 (Induced coordination ecology). *For a coordination geometry $\mathcal{G}$ with candidate repair transitions $\rho : \mathfrak{G} \rightarrow \mathfrak{G}$ available to it, the induced ecology is*

$$\mathcal{E}_{\mathcal{G}} = \{ \rho : \mathfrak{G} \rightarrow \mathfrak{G} \mid \rho \text{ a candidate geometry-repair transition available given } \mathcal{G}'\text{'s current constraint network} \}$$

acting on $Y = X \times \mathfrak{G}$ by $\rho(x, \mathcal{G}) = (x, \rho(\mathcal{G}))$: a repair operator reconfigures which $\pi_{\mathcal{G}}$ is doing the checking and leaves x untouched.

Since $\mathcal{R}_{\mathcal{E}_{\mathcal{G}}}^{(0)}(x, \mathcal{G}) = \{(x, \mathcal{G})\}$ by definition, the identity (no-repair) case is already present in $\mathcal{E}_{\mathcal{G}}$ without needing to be posited as a separate hypothesis.

8.5.2 The Regeneration Proposition

Proposition 8.5.3 (Geometric repair induces an ecology). *Let $\mathcal{G}$ be a coordination geometry whose candidate repair transitions are closed under composition. Then $\mathcal{E}_{\mathcal{G}}$ of definition 8.5.2 is an operator ecology in the sense of Chapter 6. A transition ρ satisfying theorem 3.4.1 — restoring $S_{\text{RSVP}}(\pi_{\mathcal{G}(t)}(x(t))) \leq 0$ at the current instant — is reachability-restoring in Chapter 6’s sense. It is regenerative in $\mathcal{E}_{\mathcal{G}}$ if and only if it additionally enlarges the future repertoire itself, $\text{Reach}(\rho(\mathcal{E}_{\mathcal{G}})) \supsetneq \text{Reach}(\mathcal{E}_{\mathcal{G}})$, rather than merely restoring the present-instant admissibility verdict.*

The “if and only if” in proposition 8.5.3 is the open content, not a settled equivalence: nothing yet characterizes, independently of the definition itself, which reachability-restoring transitions satisfy it. A transition can restore $S_{\text{RSVP}} \leq 0$ at the present instant while leaving $\text{Reach}(\mathcal{E}_{\mathcal{G}})$ unchanged — a repair without regeneration, exactly the gap this section develops an invariant for rather than assumes away.

Remark 8.5.4 (The converse is not automatic). *Given an operator ecology $\mathcal{E}$, constructing a coordination geometry realizing it would require every operator to be expressible as an admissible reprojction of a common space, with ecological composition corresponding to geometric composition. Chapter 6’s own Wittgenstein example is a candidate obstruction: the commissioned left-hand repertoire is discrete, symbolic, and built by commission rather than by continuous reprojction of a shared state space, with no evident $\pi_{\mathcal{G}}$ family generating it by composition. The expected result is a correspondence between a restricted class of coordination geometries and a geometrically realizable class of operator ecologies, not a full equivalence.*

8.5.3 The Reachability Gain Vector

Chapter 6’s effective repertoire dimension $D_{\text{eff}}(\mathcal{E}, x)$ is built from local effect vectors $v_o(x) \in T_x X$, one per operator. A geometry-repair transition ρ does not act on x at all, so it has no natural home in $T_x X$; the fix is to stop representing an operator’s effect as a tangent vector and represent the *reachable set it changes* as a vector instead, which is well-typed for every operator, state-level or geometry-level alike.

Definition 8.5.5 (Reachability gain vector). Fix a measure space (X, μ) . For an operator o and state x , let $\chi_S \in L^2(X, \mu)$ denote the indicator function of a measurable set $S \subseteq X$. The reachability gain vector of o at x is

$$v_o(x) = \frac{\chi_{\mathcal{R}_{\mathcal{E} \cup \{o\}}(x)} - \chi_{\mathcal{R}_{\mathcal{E} \setminus \{o\}}(x)}}{\|\chi_{\mathcal{R}_{\mathcal{E} \cup \{o\}}(x)} - \chi_{\mathcal{R}_{\mathcal{E} \setminus \{o\}}(x)}\|}.$$

For a geometry-repair transition ρ , $\mathcal{R}$ is replaced by Γ : $v_\rho(x)$ is the normalized indicator of $\Gamma_{\rho(d)}(x) \setminus \Gamma_d(x)$.

Proposition 8.5.6 (Consistency with reachability-restoration). $v_\rho(x)$ is well defined and nonzero exactly when ρ is reachability-restoring in Chapter 6’s sense.

Proof. $v_\rho(x)$ is defined by normalizing $\chi_{\Gamma_{\rho(d)}(x) \setminus \Gamma_d(x)}$, which is nonzero exactly when $\Gamma_{\rho(d)}(x) \setminus \Gamma_d(x)$ has positive measure, i.e. when $\mu(\Gamma_{\rho(d)}(x)) > \mu(\Gamma_d(x))$ under the monotonicity assumption $\Gamma_d(x) \subseteq \Gamma_{\rho(d)}(x)$ adopted for this construction. This is exactly Chapter 6’s reachability-restoring condition. $\square$

Proposition 8.5.6 depends on treating repair as monotone: $\Gamma_d(x) \subseteq \Gamma_{\rho(d)}(x)$, so the signed difference $\chi_{\Gamma_{\rho(d)}(x)} - \chi_{\Gamma_d(x)}$ is nonnegative and the gain region is unambiguous. Repairs that both gain and lose reachable states are not covered by definition 8.5.5 as stated; they are addressed in section 8.5.4.

Remark 8.5.7 (Recovering Chapter 6’s tangent-vector construction). In the smooth setting Chapter 6’s original $v_o(x) \in T_x X$ was built for, applying o for small time ε moves x along $v_o(x)$ to first order, so the newly reachable region under o is a shrinking sliver in direction $v_o(x)$; rescaled, its indicator converges weakly to a point mass along $v_o(x)$ as $\varepsilon \rightarrow 0$. definition 8.5.5 is accordingly a generalization of Chapter 6’s construction, not a replacement for it — the original recovers as a limiting case. Rigorously establishing that limit is a short lemma not carried out here.

Claim 8.5.8 (Bounds theorem transfers). Define $C_\mathcal{E}(x) = \sum_o w_o v_o(x) v_o(x)^\top$ as an operator on $L^2(X, \mu)$, with $D_{\text{eff}}(\mathcal{E}, x) = (\sum_i \lambda_i)^2 / \sum_i \lambda_i^2$ over its eigenvalues. Then $1 \leq D_{\text{eff}}(\mathcal{E}, x) \leq \text{rank } C_\mathcal{E}(x)$, exactly as in the finite-dimensional case.

Proof. The proof of Chapter 6’s bounds theorem uses only Cauchy–Schwarz applied to the finitely many nonzero eigenvalues of $C_\mathcal{E}(x)$; it requires finite rank, not finite ambient dimension. Since $\text{rank } C_\mathcal{E}(x)$ is bounded by the (finite) number of contributing operators, the argument goes through verbatim on $L^2(X, \mu)$. $\square$

8.5.4 Signed Extension: Non-Monotone Repair

Dropping the monotonicity assumption of section 8.5.3 altogether, a repair ρ acting on damage d at x generally both gains and loses reachable states:

$$\text{gain}_\rho = \Gamma_{\rho(d)}(x) \setminus \Gamma_d(x), \quad \text{loss}_\rho = \Gamma_d(x) \setminus \Gamma_{\rho(d)}(x),$$

the positive and negative parts of the signed difference $\delta\chi_\rho = \chi_{\Gamma_{\rho(d)}(x)} - \chi_{\Gamma_d(x)}$, an ordinary Jordan decomposition and nothing further.

Definition 8.5.9 (Unnormalized gain and loss vectors). $g_\rho(x) = \chi_{\text{gain}_\rho}$, $l_\rho(x) = \chi_{\text{loss}_\rho}$, left unnormalized so that $\|g_\rho\|^2 = \mu(\text{gain}_\rho)$ and $\|l_\rho\|^2 = \mu(\text{loss}_\rho)$ carry magnitude rather than being discarded by normalization.

Definition 8.5.10 (Signed reachability covariance).

$$C(x) = \sum_{\rho} w_{\rho} (g_{\rho}(x)g_{\rho}(x)^{\top} - l_{\rho}(x)l_{\rho}(x)^{\top}).$$

$C(x)$ is Hermitian but not in general positive semidefinite. Its trace,

$$\text{tr } C(x) = \sum_{\rho} w_{\rho} (\mu(\text{gain}_\rho) - \mu(\text{loss}_\rho)),$$

is the weighted sum of every candidate repair's net reachability change, aggregating what Chapter 6's reachability-restoring condition checks one repair at a time.

Definition 8.5.11 (Net effective dimension). $D_{\text{eff}}^{\text{net}}(\mathcal{E}, x) = (\text{tr } C(x))^2 / \text{tr}(C(x)^2)$.

Claim 8.5.12 (Extended bounds). $0 \leq D_{\text{eff}}^{\text{net}}(\mathcal{E}, x) \leq \text{rank } C(x)$. The upper bound holds unconditionally; the lower bound of 1 from Claim 8.5.8 need not hold once $C(x)$ is indefinite, and $D_{\text{eff}}^{\text{net}} = 0$ exactly when $\text{tr } C(x) = 0$.

Proof. The Cauchy–Schwarz step $(\sum_i \lambda_i)^2 \leq r \sum_i \lambda_i^2$, for $r = \text{rank } C(x)$, holds for any real λ_i regardless of sign, giving the upper bound unchanged. The lower bound $(\sum_i \lambda_i)^2 \geq \sum_i \lambda_i^2$ used in the positive-semidefinite case relied on every cross term $2\lambda_i\lambda_j$ being nonnegative; with mixed signs this can fail, and the inequality can be violated down to $\text{tr } C(x) = 0$, at which point $D_{\text{eff}}^{\text{net}}$ is 0 by definition of the ratio at a zero numerator. $\square$

Pure-gain repair ($l_\rho \equiv 0$ for every ρ) recovers Claim 8.5.8 exactly as the special case $D_{\text{eff}}^{\text{net}} = D_{\text{eff}}$; nothing in the pure-gain theorem required revision, only extension around it.

Definition 8.5.13 (Strongly regenerative). *A geometry-repair transition ρ is strongly regenerative in $\mathcal{E}_G$ if adding it strictly increases $D_{\text{eff}}^{\text{net}}(\mathcal{E}_G, x)$, as distinct from weakly regenerative (proposition 8.5.3’s bare condition $\text{Reach}(\rho(\mathcal{E}_G)) \supsetneq \text{Reach}(\mathcal{E}_G)$).*

A transition that trades one reachable region for a differently shaped but equal-measure region is weakly regenerative — the repertoire count increases — while contributing zero net gain, and is not strongly regenerative. This is the same divergence Chapter 6 already names between $K(o)$ and $K_V(o; x)$: inventory impact counts lost executable operators, while generative impact measures the future region their loss disconnects, recurring here inside the regeneration criterion itself rather than only in keystone-impact accounting.

Open Problem 8.5.14. *The weights w_ρ in definition 8.5.10 are estimates of the same kind as Appendix D’s retention-cost weights and Chapter 9’s allocation-problem weights, not a new estimation gap specific to this construction. No general estimation method is supplied here, consistent with the corresponding open problems in those chapters.*

Open Problem 8.5.15. *$\text{tr}(C(x)^2)$ mixes gain and loss vectors across distinct operators ρ, ρ' through cross terms $g_\rho \cdot g_{\rho'}$, $l_\rho \cdot l_{\rho'}$, and $g_\rho \cdot l_{\rho'}$. By the shape of the construction this should discount correlated repairs — two candidate reprojections gaining overlapping regions — the same way Chapter 6’s original D_{eff} discounts collinear operator effects, but this has not been separately verified here.*

Open Problem 8.5.16 (Marginal criterion). *definition 8.5.13 is stated as a whole-ecology diagnostic. A clean per-operator criterion would instead compare $D_{\text{eff}}^{\text{net}}(\mathcal{E}_G \cup \{\rho\}, x)$ against $D_{\text{eff}}^{\text{net}}(\mathcal{E}_G, x)$ directly as a marginal quantity attributable to ρ alone, in the manner of Chapter 6’s keystone impact $K(o)$. This marginal form is not developed here.*

What This Chapter Adds

Continuation, as this chapter treats it, is neither a new foundational layer beneath admissibility nor a mere restatement of it. It is the setting in which two true but previously separate claims about admissibility — that it is checked, and that it is authored — turn out to describe different levels of the same structure rather than competing theories. That distinction, and the identification of authorship with Level 3 modification already defined in Chapter 6, is this chapter’s actual content. Whether geometric repair and regenerative operators are the same construction remains open, and is flagged as such rather than resolved by the convenience of the resemblance.

Chapter 9

Reconstruction

9.1 Introduction: The Question This Chapter Actually Asks

It is tempting to treat reconstruction as this book's concluding technical convenience: distinctions were made in Chapter 1, information and entropy quantified their consequences, repair and continuation kept structures viable, and reconstruction is simply the operation that recovers a prior state from what remains. This chapter argues that framing understates the actual difficulty by exactly the amount that matters. Recovery is not always a matter of degree, cheap at one end and expensive at the other. In the generic case, a state's successor does not encode which of several possible predecessors actually produced it, and no quantity of computational resource closes that gap, because the gap is not a resource-bounded search problem in the first place.

9.2 Two Kinds of Materialization

Definition 9.2.1 (Generating input). *Given a structure X produced by a process P , a generating input for X is a state s , independent of the system's own execution history, from which P can produce X without reference to any intermediate state of the system that produced the original instance of X .*

Definition 9.2.2 (Fiber). *For a transition function $F : X \times U \rightarrow X$ and a state $y \in X$, the fiber of y , restricted to admissible predecessor pairs $P \subseteq X \times U$, is $\text{Fib}_P(y) = \{(x, u) \in P : F(x, u) = y\}$. F is non-injective over P at y when $|\text{Fib}_P(y)| > 1$.*

Definition 9.2.3 (Reconstructive irreversibility). *A transition into y is reconstructively irreversible over $\text{Fib}_P(y)$ when F is non-injective there and no generating input exists from which the predecessor's identity could otherwise be recovered.*

The distinction this chapter needs is not a spectrum running from cheap reconstruction to expensive reconstruction to impossible reconstruction. It is categorical.

Claim 9.2.4 (Non-reducibility). *No monotonic transformation of cost, however extreme, converts a reconstruction-cost materialization into a reconstructive-irreversibility materialization or the reverse. The two are distinguished by the existence of a generating input and by whether a retained history channel is injective on the relevant fiber, neither of which is a cost.*

A chaotic Hamiltonian system illustrates why this distinction is not merely definitional bookkeeping: such a system can require reconstruction precision growing exponentially in t to recover an earlier state to any fixed tolerance, making expensive reconstruction entirely rational, while never once crossing into the regime definition 9.2.3 describes — because a Hamiltonian flow’s fibers are always singletons (Claim 9.2.5 below), so a generating input always exists in principle, however hard it is to use. Expense and the existence of a generating input are independent properties.

Claim 9.2.5 (Hamiltonian dynamics never collapse fibers). *For a Hamiltonian system with a flow map Φ_t smooth enough to guarantee existence and uniqueness of solutions, Φ_t is a diffeomorphism for each t : bijective, with smooth inverse. Every fiber under Φ_t is therefore a singleton, and $\Phi_{-t}(y)$ is always a generating input for the predecessor state.*

This is the sharpest available contrast case: a system in which reconstruction is, by construction, never reconstructively irreversible, however expensive it may be to actually carry out. Most systems this book has been concerned with — repair algebras, operator ecologies, diagnostic ecologies under throughput pressure — are not Hamiltonian in this sense. Their transitions are generically non-injective, and it is worth being precise about what that costs them.

9.3 The Resource Independence Theorem

Claim 9.3.1 (Resource independence). *Let $\text{Fib}_P(y)$ be subject to reconstructive irreversibility. Then for any candidate reconstruction procedure and any allocation of time, space, or energy to it, no such allocation enables recovery of the predecessor’s identity.*

Proof. By definition 9.2.3, reconstructive irreversibility holds exactly when no generating input exists for the predecessor’s identity. A reconstruction procedure, however much time, space, or energy it is allocated, is still a function from some input to an output; enlarging its resource bound changes which functions

are computable within that bound, but does not change whether an input exists to compute from. Since no state or artifact independent of the system's own execution history determines the predecessor's identity, there is no input, of any kind, from which any procedure — regardless of allocated resources — could recover it. The claim does not depend on complexity-theoretic bounds; it depends only on the absence of a domain element to compute from, which no resource allocation supplies. $\square$

Time, space, and energy are, within limits, fungible with one another: time-space tradeoffs and energy-time tradeoffs are ordinary subjects of complexity theory and physical computation respectively. Continuation is not a fourth item on that list. Claim 9.3.1 is the precise sense in which it fails to be fungible with any of the three: more time does not help, more space does not help, more energy does not help, not because the problem is merely very hard, but because there is no input for additional resources to be spent on.

Claim 9.3.2 (Continuation is not a quantity of memory). *Possessing unbounded memory does not by itself guarantee that a system's retained structure is injective over any given admissibility-relevant fiber. Injectivity is a property of what is stored — of the retention map applied at the moment of collapse — not of how much storage is available to store it in.*

A system with arbitrarily large memory can still apply a retention policy that discards exactly the information distinguishing two predecessors — recording that an edit occurred without recording what the edit was — regardless of unused capacity. This is why more storage is not the fix for a design that retains the wrong thing: the fix is retaining a different thing, which is a design decision, not a resource allocation.

9.4 The Distinction Budget

Definition 9.4.1 (Distinguishability preservation). *For a history channel update $G(h_t, x, u)$ and a fiber $\text{Fib}_P(y)$, the augmented dynamics preserve distinguishability over the fiber at h_t when $(x, u) \mapsto G(h_t, x, u)$ is injective on $\text{Fib}_P(y)$.*

Definition 9.4.2 (Distinction budget). *Given a system's trajectory up to time t , its distinction budget B_t is the set of collapsed fibers over which predecessor identity remains recoverable — either because a generating input still exists, or because the retained history channel is injective over that fiber.*

Claim 9.4.3 (Monotonic shrinkage). *Absent a positive retention action at the moment a fiber collapses, B_t can only lose members as t increases, never regain them.*

Proof. A fiber leaves B_t exactly when neither a generating input nor an injective retained channel is available for it. Both conditions concern facts fixed at or before the moment of collapse, and neither can be altered afterward: a generating input that has ceased to exist cannot be un-destroyed, and a channel that failed to be injective at the moment of retention cannot become injective retroactively, since its value was already fixed at that moment. $\square$

Remark 9.4.4 (A structural resemblance, not a special case). *Claim 9.4.3 has the flavor of the data processing inequality from information theory — no processing of a signal can increase the mutual information it carries about its source. The resemblance is structural, not formal: this chapter’s setting concerns discrete distinguishability over admissible predecessors, not mutual information over a channel, and no claim is made that one result is a special case of the other. What the two share is the one-directional shape: a process can only destroy distinguishing information already present, never manufacture it after the fact.*

This gives reconstruction, as this book treats it, a genuinely different character from repair. Chapter 6 and Chapter 7 both concern operators that restore admissibility to a damaged state, and Chapter 6’s Anti-Alignment result showed retained history can be more or less adequate to that task. Nothing in either chapter’s machinery is monotonic in the strong sense B_t is: a damaged operator ecology can regain reachable futures through regeneration, but a fiber that has left B_t is gone in a stronger sense — no operator, however constructed, restores it, because Claim 9.3.1 rules out recovery by resource expenditure and Claim 9.4.3 rules out recovery by waiting.

9.5 Historical Observability

Definition 9.5.1 (Historical observability). *Given a class Q of possible future admissibility judgments a system might need to make, its historical observability with respect to Q at time t is the proportion of Q whose judgments remain decidable given the current distinction budget B_t .*

Historical observability is not a fixed property of hardware or raw storage capacity. By Claim 9.3.1, it cannot be increased after the fact by any allocation of time, space, or energy once a relevant fiber has left B_t ; it can only be preserved, at the moment of collapse, by retaining the right thing there. This is the sense in which continuation is the resource that determines which of a system’s future admissibility judgments remain answerable at all, independent of how much of the other three resources the system is given. A system can have unlimited time, space, and energy and still be structurally unable to answer a question whose answer depended on a distinction its own dynamics discarded before anyone thought to ask it.

9.6 Admissibility-Relevant Fibers and Selective Retention

Not every collapsed fiber is worth defending against Claim 9.4.3’s shrinkage. A system that tried to keep every non-injective transition it ever executed inside B_t would exhaust its retention budget long before the retention bought it anything.

Definition 9.6.1 (Admissibility-relevant fiber). *A fiber $\text{Fib}_P(y)$ is admissibility-relevant if there exists some later reachable state x_j and continuation family C such that the admissibility of x_j relative to C differs depending on which member of $\text{Fib}_P(y)$ actually occurred.*

Claim 9.6.2 (Selective Retention). *A system that wishes to preserve exactly the admissibility judgments its continuation families support, at minimum retention cost, must allocate injective history channels precisely over admissibility-relevant fibers, and need not allocate them elsewhere.*

Claim 9.6.2 has two failure modes, not one. Under-retention leaves an admissibility-relevant fiber undefended, so some future judgment becomes undecidable — typically not visibly, since systems under-provisioned this way tend to resolve the ambiguity silently by defaulting to whichever fiber member is most recent or most common, rather than by raising a visible error. Over-retention spends retention budget on fibers no downstream judgment ever consults. Chapter 5’s reinterpretation of Anti-Alignment as a monotonicity lemma inside a retention-cost optimization, rather than a standalone argument for maximal retention, is this claim’s consequence stated in that chapter’s own vocabulary: the optimum is a minimal admissibility-sufficient trace, not the fullest trace affordable.

9.7 The Allocation Problem and a Design Principle

Definition 9.7.1 (Retention budget and candidate fibers). *Given candidate fibers $F_1, \dots, F_n$ with relevance weights $w_i > 0$ and retention costs $c_i > 0$, and a total retention budget β , the allocation problem is to choose $S \subseteq \{1, \dots, n\}$ maximizing $\sum_{i \in S} w_i$ subject to $\sum_{i \in S} c_i \leq \beta$.*

This is an instance of the weighted knapsack problem, known to be NP-hard in general. That is stated plainly rather than smoothed over: definition 9.7.1 gives the allocation decision a precise combinatorial shape it did not have before, but a precise shape is not the same as a tractable one.

Principle 9.7.2 (Admissibility-Aware Externalization). *Given finite retention capacity, a system should allocate history channels to admissibility-relevant fibers in order*

of relevance per unit retention cost, w_i/c_i , rather than uniformly, rather than in order of raw reconstruction difficulty, and rather than by implementation convenience.

This is stated as a design principle, not a theorem: its force depends on the accuracy of weights and costs a designer supplies, which can be estimated but not guaranteed correct in advance of deployment. A designer following it imperfectly will produce exactly Claim 9.6.2's two failure modes — under-retention where relevance was underestimated, over-retention where it was overestimated — both detectable after the fact by audit even when neither was foreseeable before it.

Open Problem 9.7.3. *Four gaps remain open. The greedy heuristic's approximation gap for retention-allocation problems specifically, as opposed to knapsack problems in general, has not been established. The weights and costs definition 9.7.1 requires are estimates in any real system, and no general estimation method is supplied beyond the audit apparatus gestured at above. Whether a system's space of possible fibers is fixed in advance or expands as evidence arrives bears directly on whether candidate fibers can even be enumerated at design time. And a broader conjecture — that any persistent system under finite retention capacity, shaped by selective pressure toward correct future functioning, will exhibit retention correlating with admissibility-relevance rather than raw reconstruction cost, with departures observable as the two failure signatures above — has been illustrated with motivating candidates (legal and financial audit trails, version control's near-unconditional retention explained by near-zero marginal cost rather than uniform relevance, and biological memory consolidation's bias toward salient events) but not tested by any systematic survey. It is stated in a form intended to be testable by existing methods, not one requiring new theoretical machinery, but the survey itself remains undone.*

What This Chapter Adds to the Book's Spine

Reconstruction, in the sense this chapter has developed it, is not the operation that undoes damage; that is repair's job, treated in Chapter 6 and Chapter 7. Reconstruction is the prior question of what remains available to be acted on at all — and this chapter's central result is that the answer is not always *everything, at a price*. Some distinctions are gone in a stronger sense than expensive: no allocation of time, space, or energy, however large, brings them back, because the resource that would need to be spent was never the kind of thing money buys. What a system retains at the moment a fiber collapses is the only chance it gets.

Conclusion

This volume set out to build a proof spine, not a topic list. It is worth being precise about which of those two things the preceding eight chapters actually delivered. A topic list would have eight independent essays, each competent on its own terms, connected by nothing stronger than shared vocabulary. What follows instead is a chain of results that genuinely depend on one another: Chapter 2's information measures are checked against Chapter 1's deficit apparatus by substitution, not analogy; Chapter 5's entropy inventory reduces three of five candidate quantities to one underlying partition structure and leaves the other two — repair entropy and the RSVP field — with different, explicitly stated epistemic status; Chapter 6's repair theory supplies the generator-closure map Chapter 5's admissible future region needed and never had; Chapter 7's four failure modes are shown independent of the architecture gate rather than subordinate to it, on the strength of explicit counterexamples rather than assumption; Chapter 7 resolves an apparent contradiction between two sources this corpus draws on by showing they describe different levels of the same structure; and Chapter 9 closes the spine with a result — resource independence — that could not have been stated without Chapters 1 through 6 already having built the distinction/fiber vocabulary it depends on.

Three findings recur across chapters rather than staying local to the one that introduced them, and are worth naming together here rather than only where each first appeared. The **Repair-Recovery Separation Principle** (Chapter 7) — that restoring admissibility and preserving the historical identity of what was wrong are independent achievements — reappears in Chapter 6's minimality-driven operators and Chapter 10's fiber-collapsing repair transitions as three instances of one phenomenon, not three unrelated observations that happen to share a conclusion. The distinction between **checking and authoring** admissibility (Chapter 8) turned out to dissolve an apparent contradiction between two independently-developed accounts by showing each was answering a different question at a different level, a pattern the methodology used throughout — type test, then substitution test, before any two same-named constructions are treated as identical — was built specifically to catch. And the book's **single central open problem**, whether the two working notions of admissibility devel-

oped in this corpus induce the same continuation family, was sharpened rather than closed: from a badly-typed comparison of two unrelated-looking quantities into a well-typed question about two set-valued constructions, strengthened by a formal argument for skepticism borrowed from an unrelated result about constants of motion, and left open throughout on the explicit grounds that a book already this dependent on the answer has every incentive to resolve it prematurely, which is exactly why it should not.

What is absent is also worth stating plainly. A chapter originally planned under the title “Constraint” was found, on direct inspection of its presumed source material, to have no technical content distinct from Admissibility, and the gap was left rather than filled with something that did not belong there. Whether Coordination deserves dedicated treatment beyond what Chapters 5 and 7 already give it remains undecided. And the relationship among the three reachability-flavored constructions in this volume — the executable repertoire, the operator-induced reachable set, and the admissible future region — is stated as a generator chain in Chapter 6, but the pairwise relationships among effective reachability, physical reachability, and admissibility proper were each established under different and non-interchangeable assumptions, a distinction easily lost on a careless reading.

The discipline that produced all of this was not, in the end, a technique for proving things. It was a standing refusal to let two constructions share a name, a formula, or a chapter without first asking whether they were actually the same thing. Most of the time in this volume, they were not.

Appendix A: Notation and Symbol Index

This index collects every symbol used recurrently in this volume, in order of first appearance, each with a one-line description and a pointer to its formal definition. It is meant to be consulted rather than read straight through.

Distinction and Information

Symbol	Meaning
$(X, \sim)$	A distinguishability space: a set with an equivalence relation of observational indistinguishability.
$\delta_T(x)$	Coding deficit of x under ontology T : excess description length over the optimum, $L_T(x) - L^*(x)$.
$\Delta_T(x)$	Distinguishability deficit of x under T : inaccessible distinction capacity, $D^*(x) - D_T(x)$.
$L_T(x), L^*(x)$	Shortest description of x reachable within T , and the shortest achievable by any ontology.
$D_T(x), D^*(x)$	Number of objects separable from x within T , and the maximum achievable by any ontology.
σ, ϕ, ρ, τ	The four operations on distinguishability structure: projection, embedding, revision, transport.
$H(\mathcal{P})$	Shannon entropy of a partition, $-\sum_i p_i \log p_i$.
$S_\pi(m)$	Representational entropy at m under a representation π : $\log \pi^{-1}(m) $, the log-size of the collapsed fiber.
$N_D, \Omega(P_i)$	Number of partition cells (distinction count), and multiplicity within a cell — dual readings of one partition structure.

Entropy and Reachability

Symbol	Meaning
$S(x, t)$	The RSVP entropy field: a monotonic quantity, $\dot{S} \geq 0$, entering the entropy-based admissibility relation.
$D(t)$	Recoverable distinction capacity; the Second Law of this volume states $\dot{D} \leq 0$.
$\overline{\mathcal{R}}_t(x)$	Effective reachable set: the physically reachable set from x , quotiented by the distinguishability relation at time t .
$V_{\text{eff}}(x, t)$	Effective reachability volume, $\bar{\mu}_t(\overline{\mathcal{R}}_t(x))$.
$\Gamma_t(x)$	Admissible future region: the effective reachable set, read as the set of states an observer's current admissibility structure licenses.
$S_R(F \mid T)$	Repair entropy: uncertainty over fault locations consistent with a symptom, given retained trace T .
$\mathcal{H}_{\text{fault}}(F, T), \Omega_F(s)$	The live fault hypotheses given T , and the full fault hypothesis space for symptom s .
$\mathcal{P}_T$	The partition of a hypothesis space induced by a retained trace T .

Admissibility

Symbol	Meaning
$\pi : \mathcal{M} \rightarrow \Sigma$	A model's prediction map, with $\text{Im}(\pi)$ the observations it can in principle produce.
$\chi_\pi(x)$	The architecture gate: 0 if $x \in \text{Im}(\pi)$ (inferential boundary), 1 otherwise (ontological boundary).
$\delta_\pi^{\text{arch}}(x)$	Architectural representation deficit: $\inf_m d_\Sigma(\pi(m), x)$, a graded quantity whose zero/nonzero threshold reproduces χ_π .
$C(x)$	Continuation family: the set of trajectories from x considered permitted.
$C_S(x), C_\delta(x)$	The continuation families induced by the entropy-monotonicity and deficit-bounded admissibility constructions, respectively — the two objects at the center of this volume's central open problem.
$\pi_{\mathcal{G}}$	The CLIO projection supplied by a coordination geometry $\mathcal{G}$.

Repair and Operator Ecologies

Symbol	Meaning
$\mathcal{E}$	An operator ecology: a constraint network determining which operators remain executable.
$\text{Reach}(\mathcal{E})$	The reachable executable repertoire of $\mathcal{E}$: a set of operators.
$\mathcal{R}_{\mathcal{E}}(x)$	Operator-induced reachable set: the states reachable from x by composing operators in $\text{Reach}(\mathcal{E})$. A different, generated object from $\text{Reach}(\mathcal{E})$ itself.
$K(o)$	Keystone impact (inventory): $ \text{Reach}(\mathcal{E}) - \text{Reach}(\mathcal{E} \setminus \{o\}) $.
$K_V(o; x)$	Generative keystone impact: the corresponding loss in $\mu(\mathcal{R}_{\mathcal{E}}(x))$, which need not track $K(o)$.
$D_{\text{eff}}(\mathcal{E}, x)$	Effective repertoire dimension: a participation-ratio-style measure of decorrelated operator diversity at x , bounded between 1 and $\text{rank}(C_{\mathcal{E}}(x))$.

Restorability and History

Symbol	Meaning
$\partial_{\mathcal{R}}, \partial_{\mathcal{R}_t}$	The restorability boundary, and its history-indexed form $\partial_{\mathcal{R}}(X, \mathcal{E}, \mathcal{R}, H_t)$.
H_t	Repair history: the remembered sequence $(\varepsilon_0, R_1, \varepsilon_1, \dots, R_n, \varepsilon_n)$ of residuals and attempted repairs up to time t .
$\text{Orb}_{\mathcal{R}}(\varepsilon_0)$	Repair orbit: the set of residuals reachable from ε_0 by any word in the repair algebra $\mathcal{R}$.
$P(\varepsilon_0, \mathcal{R})$	Persistent residual: $\inf_w \ w(\varepsilon_0)\ $ over the full repair monoid.
$\hat{P}_t$	Empirical persistence bound: the same infimum, taken only over words actually attempted by time t ; an upper bound on P .

Reconstruction

Symbol	Meaning
$F : X \times U \rightarrow X$	A transition function; $\text{Fib}_P(y)$ its fiber over y , restricted to admissible predecessor pairs P .
B_t	Distinction budget: the set of collapsed fibers over which predecessor identity remains recoverable at time t . Monotonically non-increasing absent a retention decision at the moment of collapse.

Appendix B: Quotients, Fibers, and Projections

Every chapter of this volume depends, at some point, on the same basic construction: a map that collapses some states together, and the question of what survives that collapse. This appendix develops that construction once, with enough generality to support its uses in Chapters 1, 2, 6, and 10 without repetition.

Equivalence Relations and Quotients

Definition 9.7.4 (Quotient by an equivalence relation). *For a set X and an equivalence relation $\sim$ on X , the quotient $X/\sim$ is the set of equivalence classes $[x]_\sim = \{y \in X : y \sim x\}$, with the quotient map $q : X \rightarrow X/\sim$ sending $x \mapsto [x]_\sim$.*

Definition 9.7.5 (Coarsening order). *For equivalence relations $\sim_1, \sim_2$ on X , $\sim_1$ is finer than $\sim_2$ (equivalently, $\sim_2$ is coarser) when $x \sim_1 y \implies x \sim_2 y$ for all x, y . This defines a partial order on the set of equivalence relations on X , with the discrete relation (every element its own class) as the finest element and the indiscrete relation (one class) as the coarsest.*

A coarsening $\sim_1 \subseteq \sim_2$ induces a canonical surjection $c_{21} : X/\sim_1 \rightarrow X/\sim_2$ sending each $\sim_1$ -class to the $\sim_2$ -class containing it, well defined because $\sim_1 \subseteq \sim_2$ guarantees every $\sim_1$ -class lies inside a single $\sim_2$ -class.

Fibers and Projections

Definition 9.7.6 (Projection and fiber). *A projection is a surjective map $\pi : X \rightarrow Y$. The fiber over $y \in Y$ is $\pi^{-1}(y) = \{x \in X : \pi(x) = y\}$. A projection induces an equivalence relation on X by $x \sim_\pi x' \iff \pi(x) = \pi(x')$, whose classes are exactly the fibers, so every projection is a quotient map for the equivalence relation it induces, and conversely every quotient map $X \rightarrow X/\sim$ is a projection with fibers equal to the $\sim$ -classes.*

Claim 9.7.7 (Fiber monotonicity). *If $\sigma : X \rightarrow Y$ factors as $X \xrightarrow{q} X/\sim_1 \xrightarrow{c} Y$ for a coarsening $\sim_1 \subseteq \sim_\sigma$, then each fiber of σ is a union of $\sim_1$ -classes, and $|\sigma^{-1}(y)| \geq |q^{-1}([x]_{\sim_1})|$ for any x with $\sigma(x) = y$. Coarser factoring maps have fibers containing at least as much as finer ones.*

Definition 9.7.8 (Restricted fiber). *For a transition function $F : X \times U \rightarrow X$ and a subset $P \subseteq X \times U$ of admissible predecessor pairs, the restricted fiber of y is $\text{Fib}_P(y) = \{(x, u) \in P : F(x, u) = y\}$. F is injective over P at y when $|\text{Fib}_P(y)| \leq 1$.*

This is the fiber notion used in Chapter 9: a restriction of the ordinary fiber of F to the pairs an admissibility structure has already licensed, which is what allows a system to be non-injective in general while remaining injective on the specific fibers that matter to it.

Pullbacks

Definition 9.7.9 (Pullback). *For maps $f : X \rightarrow Z$ and $g : Y \rightarrow Z$, the pullback $X \times_Z Y = \{(x, y) \in X \times Y : f(x) = g(y)\}$ is the set of pairs agreeing on their image in Z , with projections to X and Y that commute with f and g .*

Pullbacks are the construction underlying the disagreement map of Chapter 7: given local sections $s_i \in \mathcal{F}(U_i)$ over a cover $\{U_i\}$, the overlap $U_i \cap U_j$ is where the pullback of the two restriction maps is tested for agreement, and the sheaf gluing condition is exactly the statement that a family agreeing on every pairwise pullback determines a unique global section.

A Disciplined Notion of Projection

The projections of this appendix are required only to be constant on the equivalence classes they collapse. A stronger and useful notion requires, in addition, that nothing *beyond* those classes gets collapsed, and that singular cases be given structured treatment rather than excluded.

Definition 9.7.10 (Constraint-leveraged inference operator). *Let X be a state space, $\mathcal{A}$ a system of local admissibility constraints on X inducing an equivalence relation $\sim_{\mathcal{A}}$, and M a compressed invariant space. A map $\pi_{\mathcal{A}} : X \rightarrow M$ is a constraint-leveraged inference operator when it satisfies three conditions: it is constant on $\sim_{\mathcal{A}}$ -classes (an ordinary quotient map); its fibers are exactly the $\sim_{\mathcal{A}}$ -classes, $\pi_{\mathcal{A}}(x) = \pi_{\mathcal{A}}(x') \iff x \sim_{\mathcal{A}} x'$ (obstruction faithfulness — no accidental collapsing beyond what the equivalence relation itself demands); and it extends to a class of controlled singular inputs $\hat{X} \supset X$, with singularities treated as structured morphism types carrying definite invariant data rather than as excluded configurations (singular extension).*

Remark 9.7.11 (This sharpens, rather than replaces, the ordinary quotient map). *Every quotient map $q : X \rightarrow X/\sim$ of this appendix already satisfies obstruction faithfulness by construction — $q(x) = q(x')$ iff $x \sim x'$ is definitional. What definition 9.7.10 adds is the third condition, singular extension, which the ordinary quotient map construction is silent about: it says nothing about how to treat inputs where the equivalence relation itself becomes singular or ill-defined. This is exactly the situation the Admissibility chapter’s discussion of constrained Hamiltonian systems encounters — a singular Legendre transform producing genuinely non-injective structure — treated there only by analogy to physics. definition 9.7.10 gives that situation a formal name and a required treatment (structured morphism types, not exclusion) rather than leaving it as an unformalized precedent.*

Induced Partitions

Claim 9.7.12 (Composability of quotients). *For a chain of coarsenings $\sim_1 \subseteq \sim_2 \subseteq \dots \subseteq \sim_n$, the induced surjections compose: $c_{n1} = c_{n,n-1} \circ \dots \circ c_{21}$, and the composite factors $X/\sim_1 \rightarrow X/\sim_n$ through every intermediate quotient.*

This is the fact underlying Chapter 7’s effective reachability contraction: distinguishability coarsening over time, $\sim_{t_1} \subseteq \sim_{t_2} \subseteq \dots$, produces a chain of quotients through which any reachable set’s image can only lose elements as it passes through successive coarsenings, since a surjective image can never contain more distinct elements than its domain.

Summary Table

Construction	Where it is used
Quotient $X/\sim$	The distinguishability space of Chapter 1; partitions underlying Chapter 2’s entropy measures.
Fiber $\pi^{-1}(y)$	Representational entropy $S_\pi(m)$ (Chapter 2); the architecture gate χ_π (Chapter 7).
Restricted fiber $\text{Fib}_P(y)$	Reconstructive irreversibility and the distinction budget B_t (Chapter 9).
Coarsening order	The Second Law and effective reachability contraction (Chapter 7).
Pullback	The disagreement map and sheaf gluing condition for continuation failure (Chapter 7).

Appendix C: Continuation and Reachability Structures

Three closely related but formally distinct constructions recur through the second half of this volume: an executable repertoire, a reachable set generated from it, and an admissible future region built from that. A separate object, the continuation family, organizes how admissibility itself is stated. This appendix develops both in one place, together with the relationships among them, so that later chapters can use all four without re-deriving their connections each time.

Trajectory Space and Continuation Families

Definition 9.7.13 (Trajectory space). *For a state space X , let $\mathcal{T}$ denote the space of trajectories: sequences (or, in continuous time, paths) $\tau : [0, \infty) \rightarrow X$ or $\tau : \mathbb{N} \rightarrow X$ with $\tau(0) = x$ for some $x \in X$, the trajectory's basepoint. Write $\tau|_{\geq t}$ for the suffix of τ from time t onward, itself a trajectory with basepoint $\tau(t)$.*

Definition 9.7.14 (Continuation family). *A continuation family is a map $C : X \rightarrow \mathcal{P}(\mathcal{T})$ assigning to each $x \in X$ a set $C(x)$ of trajectories with basepoint x , considered permitted. A state x_i occurring within a trajectory τ is admissible relative to τ and C when the suffix of τ beginning at x_i is a member of $C(x_i)$.*

Two continuation families organize the admissibility discussion of Chapter 3: $C_S(x) = \{\tau : \dot{S}(\tau) \geq 0\}$, built from entropy monotonicity, and $C_\delta(x) = \{\tau : \delta_T(\tau) \leq \epsilon\}$, built from a distortion bound. Neither closure property below is automatic for either.

Definition 9.7.15 (Closure under restriction). *C is closed under restriction if $\tau \in C(x)$ implies $\tau|_{\geq t} \in C(\tau(t))$ for every t in τ 's domain.*

Definition 9.7.16 (Closure under concatenation). *C is closed under concatenation if, whenever $\tau_1 \in C(x)$ has endpoint y and $\tau_2 \in C(y)$, the concatenation $\tau_1 \cdot \tau_2 \in C(x)$.*

Remark 9.7.17 (Neither closure property is free). C_S is closed under restriction — a suffix of a non-decreasing path is itself non-decreasing — but its closure under concatenation depends on continuity of S at the join point, a substantive assumption rather than a consequence of $\dot{S} \geq 0$ on each piece. C_δ need not be closed under either property in general: a suffix can have higher deficit relative to its own basepoint if T is state-dependent, and concatenation can exceed a threshold ϵ if deficits are sub-additive rather than bounded by the maximum of the pieces.

Open Problem 9.7.18. Characterizing the conditions under which C_S and C_δ are closed under concatenation, and whether either can be modified to guarantee it without changing its extension on single trajectories, remains open.

Morphisms and Equivalence

Definition 9.7.19 (Continuation morphism). For continuation families C_1 on X_1 and C_2 on X_2 , a continuation morphism is a pair (f, ϕ) where $f : X_1 \rightarrow X_2$ and ϕ sends each $\tau \in C_1(x)$ to a trajectory $\phi(\tau) \in C_2(f(x))$, compatibly with restriction: $\phi(\tau|_{\geq t}) = \phi(\tau)|_{\geq t}$.

Definition 9.7.20 (Continuation equivalence). States $x, y \in X$ are continuation-equivalent, $x \sim_C y$, when there exist mutually inverse continuation morphisms between $C(x)$ and $C(y)$: that is, when $C(x) \cong C(y)$.

Remark 9.7.21 (Distinct from behavioral equivalence). Chapter 6 develops a hierarchy of equivalence notions — state, structural, behavioral, regenerative — comparing systems' capacity to recreate future behavior after damage. $\sim_C$ compares permitted continuation structure directly, and the two need not coincide: two states can be continuation-equivalent while behaviorally distinct, if their permitted trajectory sets are isomorphic without their actual transition dynamics agreeing, or the reverse.

Open Problem 9.7.22 (The central open problem, in morphism form). Does there exist a continuation morphism $(f, \phi) : C_\delta \rightarrow C_S$ with $f = \text{id}_X$? If such a morphism exists and is invertible, $C_\delta \cong C_S$ and the two admissibility notions coincide. If a morphism exists in one direction only, one family refines the other. If none exists in either direction, the two are formally incomparable — stronger evidence for genuine separation than an unproved equality alone.

Reachability Structures

A second cluster of constructions — an operator repertoire, a generated reachable set, and an admissible future region — appear across Chapters 3, 5, and 6 and are related by generation and restriction rather than by identity.

Definition 9.7.23 (Executable repertoire). *For an operator ecology $\mathcal{E}$, $\text{Reach}(\mathcal{E})$ is the set of operators currently executable given $\mathcal{E}$'s constraint network.*

Definition 9.7.24 (Operator-induced reachable set). *For $x \in X$, $\mathcal{R}_{\mathcal{E}}(x) = \bigcup_{n \geq 0} \mathcal{R}_{\mathcal{E}}^{(n)}(x)$, where $\mathcal{R}_{\mathcal{E}}^{(n)}(x)$ is the set of states reachable from x by composing at most n operators from $\text{Reach}(\mathcal{E})$, every intermediate composition executable.*

Definition 9.7.25 (Admissible future region). *For a distinguishability quotient $q_t : X \rightarrow X/\sim_t$ and physical reachable set $\mathcal{R}(x)$, the effective reachable set is $\overline{\mathcal{R}}_t(x) = q_t(\mathcal{R}(x))$, and the admissible future region is $\Gamma_t(x) = \overline{\mathcal{R}}_t(x)$.*

Claim 9.7.26 (The three constructions form a generator chain).

$$\text{Reach}(\mathcal{E}) \longrightarrow \mathcal{R}_{\mathcal{E}}(x) \longrightarrow \Gamma_t(x).$$

$\text{Reach}(\mathcal{E})$ is a set of operators. $\mathcal{R}_{\mathcal{E}}(x)$ is the closure of that repertoire under composition: the states reachable by chaining executable operators together. $\Gamma_t(x)$ is a further restriction of a reachable set — here, $\overline{\mathcal{R}}_t(x)$ — to the states an observer's current admissibility structure licenses. Each stage is a genuine restriction or generation of the one before it, not a synonym for it: $\text{Reach}(\mathcal{E})$ takes operators as elements, $\mathcal{R}_{\mathcal{E}}(x)$ and $\Gamma_t(x)$ take states, so the first arrow fails a type test outright and cannot be an identity; the second arrow is a quotient, not an inclusion.

Claim 9.7.27 (Repertoire–reachability monotonicity). *For extension-monotone ecologies — $\mathcal{O}_{\mathcal{E}_1} \subseteq \mathcal{O}_{\mathcal{E}_2}$ implies every composition executable in $\mathcal{E}_1$ remains executable in $\mathcal{E}_2$ — $\text{Reach}(\mathcal{E}_1) \subseteq \text{Reach}(\mathcal{E}_2)$ implies $\mathcal{R}_{\mathcal{E}_1}(x) \subseteq \mathcal{R}_{\mathcal{E}_2}(x)$ for every x , and consequently $V_{\mathcal{E}_1}(x) \leq V_{\mathcal{E}_2}(x)$ whenever μ is monotone under inclusion.*

Claim 9.7.28 (Effective reachability contraction). *For distinguishability coarsening $\sim_{t_1} \subseteq \sim_{t_2}$, $\overline{\mathcal{R}}_{t_2}(x) = c_{21}(\overline{\mathcal{R}}_{t_1}(x))$ for the canonical surjection c_{21} induced by the coarsening, and $|\overline{\mathcal{R}}_{t_2}(x)| \leq |\overline{\mathcal{R}}_{t_1}(x)|$ in the finite case.*

Remark 9.7.29 (Reading the diagram). *Growth in a repertoire propagates forward through both arrows: more operators generate at least as large a reachable set, which (holding the distinguishability quotient fixed) generates at least as large an admissible future region. Entropy production acts on the second arrow only, coarsening the quotient that turns a reachable set into an admissible one, and is silent about the first: an ecology can gain operators while an observer's distinguishability structure simultaneously coarsens, with the two effects pulling the composite chain in opposite directions.*

Appendix D: Persistence and Reconstruction

Two bodies of machinery, developed separately in Chapters 6 and 10, concern the same underlying question from different sides: given a retained history, what can still be known, and what determines whether retaining more of it would help. This appendix collects both, and supplies one piece neither chapter states formally: a criterion for exactly which fibers a retention policy needs to protect.

Residual Orbits and Persistence

Definition 9.7.30 (Residual orbit and persistent residual). *Let $\mathcal{R}$ be a repair algebra acting on a discrepancy space $\mathcal{D}$. For a repair word $w = R_n \circ \dots \circ R_1 \in \langle \mathcal{R} \rangle$, write $\varepsilon_w = w(\varepsilon_0)$. The repair orbit of ε_0 is $\text{Orb}_{\mathcal{R}}(\varepsilon_0) = \{w(\varepsilon_0) : w \in \langle \mathcal{R} \rangle\}$. For a discrepancy norm $\|\cdot\|$, the persistent residual is $P(\varepsilon_0, \mathcal{R}) = \inf_{w \in \langle \mathcal{R} \rangle} \|w(\varepsilon_0)\|$.*

Definition 9.7.31 (Empirical persistence bound). *For a history H_t containing attempted repair words $w_1, \dots, w_n$, $\hat{P}_t = \min_{1 \leq k \leq n} \|w_k(\varepsilon_0)\|$.*

Claim 9.7.32 (Monotonicity of the empirical persistence bound). *As repair history grows, $\hat{P}_{t+1} \leq \hat{P}_t$, and $P(\varepsilon_0, \mathcal{R}) \leq \hat{P}_t$ for every finite history.*

H_t records only attempted words, so it can only approximate $P(\varepsilon_0, \mathcal{R})$ from above; repeated failure does not by itself prove structural persistence unless the attempted repair family is sufficiently exhaustive.

Claim 9.7.33 (Persistence identification under exhaustive search). *If the attempted repair words become dense in $\langle \mathcal{R} \rangle$ under a topology for which $w \mapsto \|w(\varepsilon_0)\|$ is continuous, then $\hat{P}_t \rightarrow P(\varepsilon_0, \mathcal{R})$.*

This distinguishes “every repair tried so far has failed” from “every repair available in the architecture must fail” — the second is much stronger, and is what separates Mercury’s perihelion (where the Newtonian repair algebra was plausibly searched near-exhaustively) from a stellar-age reconstruction’s noise floor (where no comparable density argument has been made).

Generating Inputs, Fibers, and Reconstructive Irreversibility

Definition 9.7.34 (Generating input and fiber). *Given a structure X produced by a process P , a generating input for X is a state s , independent of the system's own execution history, from which P can produce X without reference to any intermediate state of the original instance. For a transition function $F : X \times U \rightarrow X$ and admissible predecessor pairs $P \subseteq X \times U$, the restricted fiber of y is $\text{Fib}_P(y) = \{(x, u) \in P : F(x, u) = y\}$; F is injective over P at y when $|\text{Fib}_P(y)| \leq 1$.*

Definition 9.7.35 (Reconstructive irreversibility). *A transition into y is reconstructively irreversible over $\text{Fib}_P(y)$ when F is non-injective there and no generating input exists from which the predecessor's identity could otherwise be recovered.*

Claim 9.7.36 (Non-reducibility). *No monotonic transformation of cost converts a reconstruction-cost materialization (a generating input exists; cost is the only issue) into a reconstructive-irreversibility materialization (no generating input exists at all), or the reverse. The two are distinguished by the existence of a generating input, which is not a cost.*

Claim 9.7.37 (Hamiltonian dynamics never collapse fibers). *For a Hamiltonian system with a flow map Φ_t smooth enough to guarantee existence and uniqueness of solutions, Φ_t is a diffeomorphism for each t , so every fiber under Φ_t is a singleton and $\Phi_{-t}(y)$ is always a generating input.*

This is the sharpest available contrast case: reconstruction can be arbitrarily expensive in a chaotic Hamiltonian system while never once becoming reconstructively irreversible, since expense and the existence of a generating input are independent properties.

Resource Independence

Theorem 9.7.38 (Resource independence). *Let $\text{Fib}_P(y)$ be subject to reconstructive irreversibility. Then for any candidate reconstruction procedure and any allocation of time, space, or energy to it, no such allocation enables recovery of the predecessor's identity.*

The proof turns on the absence of a domain element to compute from, not on any complexity bound: enlarging a procedure's resource allocation changes which functions are computable within that bound, but does not supply an input where none exists. Time, space, and energy are, within limits, fungible with one another; resource independence is the precise sense in which continuation fails to be fungible with any of the three.

Claim 9.7.39 (Continuation is not a quantity of memory). *Injectivity of a retained history channel is a property of what is stored, not of how much storage is available. Unbounded memory does not by itself guarantee that a retained structure is injective over any given fiber.*

The Distinction Budget

Definition 9.7.40 (Distinguishability preservation and distinction budget). *For a history channel update $G(h_t, x, u)$ and a fiber $\text{Fib}_P(y)$, the augmented dynamics preserve distinguishability over the fiber at h_t when $(x, u) \mapsto G(h_t, x, u)$ is injective on $\text{Fib}_P(y)$. Given a trajectory up to time t , the distinction budget B_t is the set of collapsed fibers over which predecessor identity remains recoverable — either a generating input still exists, or the retained history channel is injective over that fiber.*

Claim 9.7.41 (Monotonic shrinkage). *Absent a positive retention action at the moment a fiber collapses, B_t can only lose members as t increases, never regain them.*

A fiber leaves B_t exactly when neither a generating input nor an injective retained channel is available for it, and both conditions are fixed at or before the moment of collapse: nothing afterward can restore either.

Definition 9.7.42 (Historical observability). *Given a class Q of possible future admissibility judgments, historical observability with respect to Q at time t is the proportion of Q whose judgments remain decidable given B_t .*

By resource independence, historical observability cannot be increased after the fact by any allocation of time, space, or energy once a relevant fiber has left B_t ; it can only be preserved, at the moment of collapse, by retaining the right thing there.

Selective Retention

The preceding sections establish that some fibers are worth defending against monotonic shrinkage and others are not, without yet saying which. This section supplies the criterion.

Definition 9.7.43 (Admissibility-relevant fiber). *A fiber $\text{Fib}_P(y)$ is admissibility-relevant if there exists some later state x_j , reachable from y , and some continuation family C , such that the admissibility of x_j relative to C differs depending on which member of $\text{Fib}_P(y)$ actually occurred.*

Example 9.7.44. *If two distinct reasoning traces — one that correctly ruled out an unsafe action, and one that merely ran out of budget before considering it — are compressed into the same summarized state, and a later step re-proposes the same action, whether that proposal is admissible depends on which predecessor occurred: the fiber is admissibility-relevant. If two different low-level operations happen to produce identical log entries, and nothing downstream ever depends on which specific operation produced the entry, the corresponding fiber is collapsed but not admissibility-relevant.*

Theorem 9.7.45 (Selective Retention). *A system that wishes to preserve exactly the admissibility judgments its continuation families support, at minimum retention cost, must allocate injective history channels precisely over admissibility-relevant fibers, and need not allocate them elsewhere.*

Sketch. Necessity: if a fiber is admissibility-relevant, some future state’s admissibility verdict depends on which predecessor occurred. A non-injective history channel over that fiber leaves the predecessor’s identity unrecoverable, so the verdict cannot be determined. Sufficiency and minimality: by definition 9.7.43, no future admissibility judgment depends on distinguishing predecessors within a non-admissibility-relevant fiber, so injective retention there preserves a distinction with no consumer, at nonzero cost, and can be dropped without any judgment becoming unsupportable. $\square$

Corollary 9.7.46 (Two failure modes). *A retention architecture fails theorem 9.7.45 in exactly two ways: under-retention, leaving an admissibility-relevant fiber undefended so that some future judgment becomes undecidable; or over-retention, maintaining injective channels over fibers that are not admissibility-relevant, paying a cost no future judgment consumes.*

Under-retention is invisible until the specific future state that needed the distinction actually arrives and cannot be evaluated. Over-retention is visible directly, as wasted storage or compute, which is why naive economic accounts of retention notice and penalize it while missing under-retention entirely.

Remark 9.7.47 (Consequence for Chapter 5’s retention-cost objective). *Chapter 5 poses the objective $J(T) = S_R(F \mid T) + \lambda K(T)$ without specifying what the retention cost functional $K(T)$ should actually penalize. theorem 9.7.45 answers this: K should be understood as counting injective retention maintained over fibers that are not admissibility-relevant, and the minimal admissibility-sufficient trace T^* that objective seeks is exactly the trace achieving injectivity over the admissibility-relevant fibers of definition 9.7.43 and nowhere else.*

Summary

Reconstruction, across both chapters this appendix draws together, is not the operation that undoes damage — that is repair’s province. It is the prior question of what remains available to be acted on at all. The persistent residual $P(\varepsilon_0, \mathcal{R})$ asks whether a discrepancy would survive unlimited repair attempts; the distinction budget B_t asks whether a predecessor’s identity would survive unlimited resource expenditure. Both answers can be no, for reasons that have nothing to do with how much repair or resource is available. Selective Retention closes the gap between the two: it says which fibers a system must protect from the second failure to guarantee it never faces a version of the first that could have been prevented.

Appendix E: Worked Examples

The preceding chapters illustrate their formal claims with examples drawn from astrophysics, neurophysiology, and cortical population recordings. This appendix develops one extended example instead, returned to across three escalating tiers, because a single running example makes the relationships between chapters — which nothing else in this book states in one place — visible in a way scattered one-chapter illustrations cannot.

Setup: A Lane of Vehicles

Consider a lane of vehicles indexed $i = 1, 2, \dots$, each with state $x_i \in X$ (position, velocity) evolving under local sensing rather than centralized control. No vehicle has access to a global state; each has access only to its own sensors and, where extended below, a small number of communication channels to neighbors and infrastructure.

Tier 1: Stigmergic Local Control

Each vehicle does not track its absolute position or speed as the quantity governing safe following. It tracks the gap $g_i(t)$ to the vehicle ahead and the closing rate $\dot{g}_i(t)$: whether that gap is shrinking or growing. This is stigmergy in the classical sense — coordination through local, indirect signals in the shared environment, the mechanism behind ant pheromone trails and Reynolds flocking rules — rather than through any shared global representation.

Definition 9.7.48 (Pairwise admissibility). *For adjacent vehicles $i, i + 1$ with joint state $(x_i, x_{i+1}) \in X \times X$, a continuation family on the product space, $C : X \times X \rightarrow \mathcal{P}(\mathcal{T})$, determines whether a joint trajectory maintains $g_i(t) \geq g_{\min}$ for a safety threshold $g_{\min}$.*

This is an ordinary application of the continuation family of Appendix C, not a generalization of it: admissibility here is a property of the pair, but the pair

itself is simply a state in the product space $X \times X$, and definition 9.7.14 already covers states in any space, product or otherwise. No new formal machinery is needed for a fixed pair.

Remark 9.7.49 (Where this framing breaks). *The product-space framing depends on each vehicle having exactly one relevant neighbor. Merging lanes, variable following distance, or a vehicle with two adjacent neighbors on either side all break the fixed pair into a variable-size neighborhood, at which point $X \times X$ is no longer the right domain and something closer to a local system — an assignment of continuation families varying coherently across overlapping neighborhoods, in the spirit of the sheaf-theoretic machinery Chapter 7 uses for continuation failure — would be needed instead. This is not developed here.*

The safety condition $g_i(t) \geq g_{\min}$ does not require any vehicle to categorize its situation into discrete named states — congested, free-flowing, merging. It requires only that the local geometry of gap and closing rate remain separable enough to support one admissibility judgment. This is remark 7.5.2’s claim in a different domain: restorability depends on separability, not categoricity. A vehicle does not need a clean category for its situation any more than a cortical neuron needs one for the variable it encodes.

Correction under this scheme is also distributed in exactly the sense example 7.4.13 develops: no single controller adjusts the whole lane, and a local perturbation — one vehicle braking — is absorbed by neighboring vehicles’ local responses without any global recomputation, the repair reserve $R(d, x)$ of definition 7.4.12 realized as neighboring vehicles’ independent capacity to respond, not as a single central controller’s bandwidth.

Tier 2: Diagnostic Escalation and Geometric Repair

Vehicles in this system have no steering wheel. Passengers hold tablets; drones are stationed along the road; a central console displays a ranked list of pending questions; towers with human pilots can assume direct control during accidents or extreme ambiguity. Four distinct mechanisms operate across this hierarchy, each realizing a different piece of earlier chapters.

Ambiguous objects and diagnostic collapse

When an onboard camera cannot classify an object with confidence, its diagnostic channel is approaching $I(F; E) \rightarrow 0$ locally (section 7.4.2). The system’s response is not to apply more processing to the same camera feed; it routes the question to a passenger’s tablet, a structurally different diagnostic channel.

Chapter 7 states the general remedy for diagnostic collapse as requiring a qualitatively new channel, not more of the same one; this is that remedy with a concrete instantiation rather than an abstract one.

Drone control as geometric repair

A passenger who takes control of a local drone ahead does not correct their own vehicle’s trajectory. They modify the shared coordination geometry other vehicles are responding to — exactly geometric repair in the sense of theorem 3.4.1, and a Level 3 modification in definition 6.2.1’s hierarchy: the object being changed is the rule governing future behavior, not a state or a parameter within a fixed rule.

Remark 9.7.50 (Checking versus authoring, instantiated). *A camera’s classification check against the current coordination geometry is checking, in Claim 8.3.1’s sense: evaluated against a geometry held fixed at that instant. A passenger’s decision to take control of a drone, reorganizing what the geometry permits going forward, is authoring. The two happen at different levels of the same system, exactly as Chapter 8 argues they must.*

Tier 3: The Salience Screen

The central console does not surface every ambiguity uniformly. It ranks pending questions by how much they matter, which is precisely Admissibility-Aware Externalization (principle 9.7.2): allocate attention to admissibility-relevant fibers (definition 9.7.43) in order of relevance per unit cost of driver attention, not by raw ambiguity or uniform coverage. A question about an object that no downstream decision depends on — the non-admissibility-relevant case of Appendix D’s own example — does not belong on the screen even if the classifier’s confidence is low; a question that determines whether a merge is safe belongs there even if the classifier is otherwise fairly confident.

Tier 4: Tower Pilots and the Architecture Gate

Tower pilots are invoked far less often than passenger tablet queries, and their intervention is qualitatively different from every tier below. The condition that triggers it — “accidents or extremely ambiguous situations” — is doing real work: it distinguishes an inferential boundary, where more asking within the existing sensor-and-passenger architecture would eventually resolve the ambiguity, from an ontological boundary, where the architecture itself cannot rep-

resent what has happened and a different control modality is required regardless of how much more is asked of it. This is the architecture gate χ_π of definition 7.3.1, with the same two-sided character as Mercury’s perihelion versus Banik’s stellar-age reconstruction (example 7.3.2): most ambiguous situations are inferential, resolved by tiers 1 through 3; a genuine accident can be ontological, outside what the vehicle’s own sensors and passengers’ judgment can represent at all.

Remark 9.7.51 (A keystone at the top of the hierarchy). *Tower-pilot intervention is rare and high-leverage in exactly remark 6.6.9’s sense: its low invocation frequency does not indicate low importance. Its unavailability would disconnect a disproportionate share of the system’s correctable state space — exactly the situations tiers 1 through 3 cannot resolve — which is a generative keystone impact K_V far larger than its usage frequency alone would suggest, the same asymmetry apprenticeship training carries in the two-cities argument of Chapter 6.*

Summary Table

Tier	Formal object instantiated
Gap and closing rate	Pairwise admissibility on a product-space continuation family; separability over categoricity.
Passenger tablet query	A new diagnostic channel resolving diagnostic collapse, not more capacity on the same channel.
Drone control	Geometric repair; a Level-3 modification; the authoring side of the checking/authoring distinction.
Salience screen	Admissibility-Aware Externalization: ranking by relevance per unit cost, not raw ambiguity.
Tower pilot	The architecture gate: inferential versus ontological boundary, with keystone-scale generative impact despite rare invocation.

Each tier of this single system realizes a construction introduced independently, in a different chapter, for a different reason. That convergence is not itself a proof of anything; it is offered as evidence that the formal vocabulary developed across this volume describes a recurring structural pattern rather than a set of unrelated devices each suited only to the domain that motivated it.

References

Bibliography

- [1] Banik, I., Kudakolawa Kaluarachchige, T., Cookson, S., & Desmond, H. (2026). The age of the Universe from a large sample of the oldest Galactic stars. *Monthly Notices of the Royal Astronomical Society*, in press. arXiv:2607.00764.
- [2] Posani, L., Wang, S., Muscinelli, S. P., Paninski, L., & Fusi, S. (2026). Rarely categorical, highly separable representations along the cortical hierarchy. *Nature*. <https://doi.org/10.1038/s41586-026-10668-4>
- [3] Flyxion. *Distinguishability Geometry*. Companion essay.
- [4] Flyxion. *The Admissibility Field*. Companion essay.
- [5] Flyxion. *Constraint Before Content*. Companion monograph.
- [6] Flyxion. *Coordination Geometry: From Synchrony to Structure*. Companion essay.
- [7] Flyxion. *Operator Ecology*. Companion essay.
- [8] Flyxion. *History Before Function: Compression as the Origin of Operators*. Companion essay.
- [9] Flyxion. *Recursive Continuation: Autopoiesis, Sympoiesis, and the Compression of Self-Modification*. Companion essay.
- [10] Flyxion. *Histories Become Operators: Biological Notes on Retained Admissibility*. Companion essay.
- [11] Flyxion. *The Ecology of Distinctions*. Companion monograph.
- [12] Flyxion. *Persistent Anomalies and the Geometry of Ontology Revision: Repair, Admissibility, and Non-Abelian Tears*. Companion monograph.
- [13] Flyxion. *What Survives Reconstruction*. Companion monograph.

- [14] Flyxion. *The Expiatory Gap: Civilizational Incompressibility and the Limits of Superintelligent Optimization*. Companion essay.
- [15] Flyxion. *Generative Continuation Geometry*. Companion essay.