

SRIDHAR MAHADEVAN

DISCOVERING
COMPOSITIONAL WORLDS
FROM INTERACTION

VOLUME IV OF THE CATEGORICAL AI RESEARCH PROGRAM

Contents

<i>Preface</i>	11
<i>Trail Map</i>	19
<i>I Foundations of Categorical Identification</i>	23
<i>o The Infant's Problem</i>	25
<i>o.1 Two complementary hypotheses</i>	26
<i>o.2 Structured experience creates reusable structure</i>	27
<i>o.3 What is a compositional world?</i>	28
<i>o.4 Language as a discovered world</i>	30
<i>o.5 Embodied action as a discovered world</i>	31
<i>o.6 Social interaction as a discovered world</i>	32
<i>o.7 Three worlds, one learner</i>	33
<i>o.8 Particle physics as a discovered world</i>	33
<i>o.9 A developmental contract for machine learning</i>	35
<i>o.10 What the categorical theory must provide</i>	36
<i>Further Reading</i>	37
<i>1 A Categorical and Homotopical Toolkit</i>	39
<i>1.1 Categories and typed composition</i>	39
<i>1.2 Doctrines: the structure known in advance</i>	40

1.3	<i>Size, categories of categories, and indexed hypotheses</i>	43
1.4	<i>Universal properties, limits, and colimits</i>	43
1.5	<i>Algebraic theories, sketches, and generative worlds</i>	44
1.6	<i>Adjunctions and Kan extensions</i>	47
1.7	<i>Functor categories and online information</i>	49
1.8	<i>Simplicial sets and decision nerves</i>	49
1.9	<i>Horns, composition, and quasi-categories</i>	51
1.10	<i>Weak equivalence and localization</i>	52
1.11	<i>Homotopy-coherent diagrams</i>	53
1.12	<i>Homotopy limits, colimits, and derived Kan extensions</i>	53
1.13	<i>Stable infinity-categories</i>	54
1.14	<i>Repair spaces</i>	55
1.15	<i>Tangent categories and LINCS</i>	55
1.16	<i>Infinitesimal intrinsic models</i>	58
1.17	<i>Reading the derived skeletal-continuity theorem</i>	62
1.18	<i>What this chapter does not assume</i>	63
2	<i>The ORACLE Program</i>	65
2.1	<i>The categorical prior</i>	65
2.2	<i>Prediction before decision</i>	74
2.3	<i>From universal imitation to ORACLE</i>	74
2.4	<i>The ORACLE–UOCL enrichment square</i>	75
2.5	<i>Claim discipline and the theorem ladder</i>	75
3	<i>Categorical Identification Under a Prior</i>	77
3.1	<i>Presentation protocols</i>	77
3.2	<i>The hypothesis fibration</i>	78
3.3	<i>Universal query theories</i>	79
3.4	<i>The finite-transcript obstruction</i>	80
3.5	<i>A constrained positive regime</i>	80
3.6	<i>Coinductive realization</i>	81

3.7	<i>When prediction becomes decision</i>	81
3.8	<i>Identification under a categorical prior</i>	82
3.9	<i>The fragment cover of the core sketch</i>	82
3.10	<i>Local identification does not automatically glue</i>	83
3.11	<i>Observers must be jointly separating</i>	84
3.12	<i>Possible and impossible language worlds</i>	84
3.13	<i>Assimilation, ordinary accommodation, and doctrinal repair</i>	85
3.14	<i>The ORACLE readiness criterion</i>	86

II *Categorical Core Knowledge* 87

4	<i>Core Knowledge and Piagetian Construction</i>	89
4.1	<i>Core knowledge as a doctrine of possible worlds</i>	89
4.2	<i>Three meanings of categorical truth</i>	91
4.3	<i>The five-obligation test</i>	91
4.4	<i>Prior gain as a relative identification statement</i>	92
4.5	<i>Six systems, six categorical burdens</i>	93
4.6	<i>Cross-system composition is the decisive test</i>	94
4.7	<i>What would count against the hypothesis?</i>	94
4.8	<i>Does existing machinery already solve ORACLE?</i>	95
4.9	<i>Piagetian construction of a compositional world</i>	101
4.10	<i>Development is not merely clock time</i>	102
4.11	<i>Assimilation as conservative lifting</i>	102
4.12	<i>Two levels of accommodation</i>	103
4.13	<i>Equilibration as a repair process</i>	103
4.14	<i>Developmental composition of the core systems</i>	104
4.15	<i>A grounded research program</i>	104

III *Universal Online Categorical Learning* 107

5	<i>Learning Inside Fixed Doctrines</i>	109
5.1	<i>The doctrinal audit</i>	109
5.2	<i>System identification and predictive state representations</i>	111
5.3	<i>Topos World Models from documents</i>	113
5.4	<i>Automata and machine inference as coalgebra learning</i>	115
5.5	<i>Reinforcement learning inside the MDP–Bellman doctrine</i>	123
5.6	<i>Causal discovery as identification in a causal category</i>	126
5.7	<i>Language learning inside grammar and sequence doctrines</i>	127
5.8	<i>JEPA, world models, and learned quotients</i>	130
5.9	<i>What the fixed-doctrine analysis establishes</i>	135
6	<i>The UOCL Machine</i>	139
6.1	<i>From a semantic section to a learning machine</i>	139
6.2	<i>State, probes, and dependent update</i>	140
6.3	<i>Assimilation and accommodation as control flow</i>	141
6.4	<i>Active identification and separating probes</i>	142
6.5	<i>Learning generators, relations, and models</i>	143
6.6	<i>Executions realize ORACLE sections</i>	145
6.7	<i>Composing UOCL learners</i>	145
6.8	<i>Limits, colimits, and multimodal fusion</i>	147
6.9	<i>Failure modes</i>	149
7	<i>Persistent Categorical Identification</i>	151
7.1	<i>Non-anticipation and two forms of stabilization</i>	151
7.2	<i>Coherent persistence</i>	152
7.3	<i>Finite active identification</i>	152
7.4	<i>Presentation invariance</i>	152
7.5	<i>Tangent UOCL: learning how a world can vary</i>	153
7.6	<i>The cost of categorical learning</i>	156
7.7	<i>Two enrichments of categorical learning</i>	156

8	<i>Probably Approximately Categorically Correct Learning</i>	159
8.1	<i>The PAC template</i>	159
8.2	<i>The PACC declaration</i>	160
8.3	<i>Ordinary PAC learning is the discrete case</i>	161
8.4	<i>Finite PACC bounds</i>	162
8.5	<i>Coverage is part of the categorical prior</i>	163
8.6	<i>Three strengths of categorical approximation</i>	164
8.7	<i>Toward a categorical dimension theory</i>	164
8.8	<i>PACC evolvability</i>	165
8.9	<i>Persistent and online PACC</i>	166
8.10	<i>What PACC does and does not promise</i>	167
	 <i>IV Universal Online Decision Learning</i>	 169
9	<i>Online and Persistent Universal Decision Learning</i>	171
9.1	<i>Claim discipline and the theorem ladder</i>	173
9.2	<i>Scientific questions for universal online decision learning</i>	174
9.3	<i>Prefix semantics, nerves, homotopy coherence, and repair</i>	175
9.4	<i>Time, revelation, and the least available past</i>	175
9.5	<i>Prefix-indexed UDL and non-anticipation</i>	176
9.6	<i>Online Kan invariance</i>	177
9.7	<i>Growing diagrams and the Beck–Chevalley defect</i>	178
9.8	<i>Persistent online abstractions</i>	178
9.9	<i>Decision nerves and compositional depth</i>	179
9.10	<i>Simplicial online UDL</i>	180
9.11	<i>The online-to-offline comparison</i>	181
9.12	<i>Homotopy-coherent online UDL</i>	182
9.13	<i>Observers and the origin of regret</i>	186
9.14	<i>Information-structured regret beyond the classical chain</i>	188
9.15	<i>Running example: crossword solving as universal completion</i>	191
9.16	<i>Running example: Sudoku and flow-based recurrent repair</i>	193

<i>V</i>	<i>Global-Clock Online Decision Learning</i>	199
10	<i>Convex and Regularized Universal Action</i>	201
10.1	<i>Universal decision learning</i>	201
10.2	<i>LINCS discipline</i>	201
10.3	<i>The categorical OCO declaration</i>	202
10.4	<i>Regularized universal action</i>	203
10.5	<i>The FTRL action normal form</i>	203
10.6	<i>Finite enriched accumulation</i>	204
10.7	<i>Quadratic FTRL and typed tangents</i>	206
10.8	<i>Proximal FTRL and tangent repair</i>	207
11	<i>Universal Bandit Decision Learning</i>	211
11.1	<i>The bandit information channel</i>	211
11.2	<i>Pathwise loss and barycentric recovery</i>	212
11.3	<i>Exploration as infinitesimal admissibility</i>	213
11.4	<i>The expectation-level regret-transfer theorem</i>	214
11.5	<i>Entropic realization and the classical rate</i>	214
11.6	<i>Bandit convex optimization changes the repair target</i>	215
<i>VI</i>	<i>Information Beyond a Global Clock</i>	217
12	<i>Information Categories and Decentralized Decisions</i>	219
12.1	<i>Comparator sketches as descent data</i>	219
12.2	<i>The comparator spectrum on a finite horizon</i>	220
12.3	<i>Approximate descent and observer transfer</i>	221
12.4	<i>Running example: asynchronous event categories</i>	222
12.5	<i>Decentralized teams, games, and nonclassical information</i>	224
12.6	<i>Running example: asynchronous distributed minimization</i>	224

13	<i>Agentic Safety and Universal Online Decision Learning</i>	227
13.1	<i>From model safety to information-structure safety</i>	227
13.2	<i>Safety under a changing information shape</i>	228
13.3	<i>Running example: stale authority under asynchronous execution</i>	229
13.4	<i>Audit is an observer, not the execution</i>	230
13.5	<i>A contemporary incident as an information-structure failure</i>	230
13.6	<i>Temporal behavior types as trust contracts</i>	231
13.7	<i>Safety invariants beyond scalar regret</i>	233
	 VII Coherent Repair and Structural Amplification	 235
14	<i>Homotopy and Tangent Repair</i>	237
14.1	<i>Registered repair problems</i>	237
14.2	<i>Comparison-guided localization</i>	238
14.3	<i>Assimilation, accommodation, and doctrinal change</i>	239
14.4	<i>A repair ledger</i>	240
14.5	<i>Tangent repair, stability, and convergence</i>	241
14.6	<i>Tangent repair certificates</i>	243
14.7	<i>A three-axis stability doctrine</i>	244
15	<i>Invariance, Amplification, and Approachability</i>	245
15.1	<i>The decision quotient</i>	246
15.2	<i>When invariance fails</i>	246
15.3	<i>Boosting as free convex completion</i>	247
15.4	<i>The weak-learning interface</i>	248
15.5	<i>When amplification becomes persistent</i>	248
15.6	<i>Duality, games, and approachability</i>	249
15.7	<i>Universal equilibrium objects</i>	249
15.8	<i>Approachability as persistent consistency</i>	250
15.9	<i>Approachability under changing information</i>	250

VIII	<i>Persistent Structure Across a Lifetime</i>	253
16	<i>Persistent Structure Across a Lifetime</i>	255
16.1	<i>Environment morphisms and knowledge packages</i>	255
16.2	<i>A transport theorem</i>	256
16.3	<i>Transport defects and partial reuse</i>	257
16.4	<i>Persistent compositional knowledge</i>	258
16.5	<i>Usefulness by future factorization</i>	258
16.6	<i>The persistence core</i>	259
16.7	<i>The compositional storehouse</i>	260
16.8	<i>Progress without a single task score</i>	261
17	<i>Synthesis and the Next ORACLE Theorem Ladder</i>	263
17.1	<i>What the current theory establishes</i>	263
17.2	<i>The first next rung: finite doctrine stabilization</i>	265
17.3	<i>The open frontier</i>	266
17.4	<i>The ORACLE persistence conjecture</i>	269
17.5	<i>From the tetralogy to a research program</i>	270
17.6	<i>Coda: frontier models as collaborators</i>	270
A	<i>Claim ledger</i>	275
	<i>Bibliography</i>	279
	<i>Glossary of Notation</i>	289
	<i>Subject Index</i>	293

Preface

A CHILD CONFRONTS THE PROBLEM AT THE CENTER OF THIS BOOK. The world does not arrive already divided into objects, actions, causes, and rules. It arrives through many senses at once, changing as the child acts within it. William James memorably described this first encounter as “one great blooming, buzzing confusion” [James, 1890]. The thought is almost frightening: each of us once faced the task of making an intelligible world from that confusion, yet none of us can remember how our brain managed to do it.

Evolution had millions of years to shape a learner capable of meeting this challenge. Its achievement was not to provide the newborn with a finished theory of a particular world. It was, rather, to equip the child with abstractions through which a theory could be progressively constructed by observation and interaction. The learner begins neither with a blank slate nor with a complete model, but with forms of organization that make experience learnable.

Spelke’s core-knowledge hypothesis gives empirical content to these forms of organization. Evidence from infants and other learners supports six early emerging systems: representations of cohesive physical objects and their mechanical interactions; approximate number; the geometry of navigable space; object form; goal-directed agents; and social beings with attention, emotion, and mental states [Spelke, 2022, 2023]. These systems are not finished theories of the world. They are automatic, domain-specific interfaces between perception and explicit belief that guide learning and continue to function throughout life. They make the blooming, buzzing confusion initially probeable.

Piaget gave this construction a dynamic form. A child assimilates a new experience when it can be incorporated into an existing schema; the child accommodates when the schema itself must change to make sense of experience [Piaget, 1952]. Learning is therefore not simply the accumulation of observations. It is the continuing construction, application, and repair of structures that organize observations and actions into a coherent whole. Abstracted from its developmental setting, the Piagetian conjecture is that an intelligent learner progressively builds a

compositional theory of its world.

This book begins where *Infinitesimal Creativity* left us. That book studied how a learner can vary, compare, and repair compositional structure through infinitesimal probes and double involution. Here we move one level deeper. Before a learner can creatively transform a theory, how can it discover the theory's objects, transformations, compositions, and universal regularities from interaction? And before it can perform an infinitesimal repair, how can it discover which variations the world admits?

Discovery may also be more compact than a catalogue of the world. A child playing with construction pieces can learn a small algebraic theory of joining, juxtaposing, and separating whose composites generate indefinitely many assemblies. ORACLE must therefore distinguish learning a category of encountered structures from learning a sketch of generators and relations that presents those structures and transports across models.

Our categorical hypothesis is deliberately spare:

Design principle

The learner does not initially know the world, but it knows that the world is compositional. Formally, it knows what a category is while having to discover which category describes the world it inhabits.

The stronger refinement developed in this book replaces the unknown category by an unknown tangent category. The learner must then identify not only a base category $\mathcal{C}$, but tangent structure

$$(T, p, 0, +, \ell, c)$$

describing admissible infinitesimal variation and its coherence. This need not be an innate commitment imposed on every UOCL instance. It is a structured hypothesis class whose additional probes make infinitesimal learning possible.

This is not the claim that an infant consciously knows category theory. Rather, categorical structure plays the role of an innate form of organization. Objects can persist, arrows can be composed, identities can be recognized, and different paths can be compared. Interaction then supplies a presentation from which a particular category may be identified: its objects, morphisms, composites, equations, and universal properties, but only to the extent that the learner's admissible probes can distinguish them.

We call the resulting formal learner ORACLE: an *Online Rational Agent for Categorical Learning*. Here rationality means coherence with

the categorical structure supported by evidence, not prior possession of an optimal policy. Assimilation extends a currently viable categorical model with new observations and composites. Accommodation repairs the presentation—or changes the class of models—when the new evidence cannot be coherently absorbed. The familiar name is appropriate: the learner’s knowledge is tested by the structural questions it can answer.

Spelke’s results also sharpen what this learner may know before interaction. The categorical prior need not be an empty doctrine of objects and arrows. It may include a small, typed family of core observers—for objects, number, space, form, agency, and sociality—that determine which distinctions can be made first. The category of the world is still unknown; core knowledge provides some of the probes through which its presentation is revealed. A central question for ORACLE is then how initially distinct systems become linked by composition without losing the invariants that made early learning possible.

This formulation is a categorical generalization of Gold’s identification of languages from presentations and of predictive system identification. It also generalizes universal imitation games, where inductive inference by initial algebras was contrasted with coinductive inference by final coalgebras. ORACLE permits the unknown target itself to range over a world of categories. A universal coalgebra is one important case, not the boundary of the theory.

Universal Online Categorical Learning supplies the algorithmic bridge from this semantic program to decision. UOCL specifies how hypotheses are selected, probes are chosen, evidence is assimilated, and categorical declarations are repaired under progressive revelation. Its open learners form compositional systems: modality-specific learners can run in parallel, be wired in sequence, or be fused by limits enforcing a shared object, spatial, linguistic, or social doctrine. Tangent UOCL asks the same learner to identify how its categorical world can vary. In some domains it may seek the stronger explanation of a differential category whose coalgebras generate that tangent geometry; this is a separate identification claim, not an automatic consequence of possessing derivatives. Universal Online Decision Learning is a second enrichment, obtained when a learned category is equipped with information, action, consistency, consequence, and observation maps. Their intersection—tangent UODL—creates the typed infinitesimal conditions under which DIAL can diagnose and repair a decision mechanism rather than merely assume those conditions in advance.

Exact categorical identification is not always statistically attainable or task-relevant. PACC UOCL therefore adds a controlled relaxation: a learned world is probably approximately categorically correct rel-

ative to a declared distribution of equivalence-invariant categorical probes. Ordinary PAC learning is recovered when the world and probe language are discrete; structural and coherent PACC additionally test composition, universal properties, horns, and filler spaces.

This generality contains familiar learning paradigms as deliberately restricted cases. System identification searches a category of dynamical models; automata and machine inference learn transition coalgebras up to observable behavior; reinforcement learning fixes an MDP doctrine and adds consequential action; causal discovery searches a category of causal models; grammatical inference and Transformer language modeling restrict the unknown world to classes of languages or conditional sequence laws. Prometheus and Odyssey instead compile document presentations into local, action-oriented predictive-state models and glue them into provenance-bearing Topos World Models. Chapter 5 makes these reductions explicit while distinguishing genuine categorical structure from the vacuous act of treating any hypothesis set as a discrete category. Online convex optimization supplies a compact test suite rather than the architecture of the book. Hazan’s progression from regularized leader methods through bandit feedback and boosting occupies the global-clock case; later parts treat non-linear information structures, homotopy-coherent repair, safety, and lifelong transport.

THE FOUR BOOKS FORM A SYSTEMATIC ASCENT FROM STATIC MODELING TO DYNAMIC DISCOVERY. *Categories for Artificial General Intelligence* gave a functorial semantics of cognition: morphisms represented compositional steps of thought, and functors represented systematic translations between conceptual systems. Its categories, however, were available before the translation began. *Machine Learning by Enforcing Compositionality* made learning depend on sketches—finite declarations by generators, relations, cones, and cocones. Learning became the completion and comparison of models satisfying a declared sketch. The model was unknown, but the type of sketch was fixed for the learning problem; its constraints need not select a unique model unless additional hypotheses make them separating.

Infinitesimal Creativity then asked how an identified compositional structure could vary and be repaired. Tangent categories typed the admissible infinitesimal directions, while involution algebroids and double involution supplied an algebraic mechanism for creative variation and reconstruction. The base structure on which those variations acted was still assumed to have been identified. The present book makes the remaining assumption explicit: the doctrine—the structural grammar specifying what kinds of sketches, models, morphisms, and universal constructions are admissible—may itself be unknown.

Here *doctrine* is used in a deliberately broad Lawvere-style sense. A

doctrine is a structured theory of categories and their models, together with its admitted structure-preserving maps. Many doctrines can be presented by a 2-monad, an essentially algebraic theory, or a specified class of sketches, but the word is not restricted to any one of these realizations. Lawvere theories, monoidal categories, causal models, Markov decision processes, toposes, and tangent categories exemplify doctrines at different levels of structure.

The technical object that gathers the four books is a hypothesis fibration. Suppressing the separate transcript coordinate for the moment, write

$$p : \mathcal{E} \longrightarrow \mathcal{B}, \quad b \longmapsto \mathcal{E}_b, \quad (u : b \rightarrow b') \longmapsto (u^* : \mathcal{E}_{b'} \rightarrow \mathcal{E}_b).$$

The base $\mathcal{B}$ records admissible doctrines and their declared revisions; the fiber $\mathcal{E}_b$ contains the hypotheses meaningful under doctrine b ; and reindexing translates a hypothesis along a change of doctrine. In the full theory, presentation prefixes refine this base, as in the Grothendieck construction $\int H_{\text{doc}} \rightarrow \mathbf{Doc}$ and its pullback to the observed presentation path. An ORACLE learner is therefore a coherent section along that path. A global section $\sigma : \mathcal{B} \rightarrow \mathcal{E}$, when one exists, is a stronger semantic ideal rather than an assumption made by every learning problem.

Volume	Organizing structure	What was learned	ORACLE realization
Book 1	Functors $F : \mathcal{C} \rightarrow \mathcal{D}$	Compositional translations between categories taken as given	The projection p , reindexing u^* , and cartesian comparison maps are functorial; a coherent section selects related models in doctrine-indexed fibers
Book 2	Sketches S	Models of a declared finite presentation of objects, arrows, equations, cones, and cocones	Presentation prefixes accumulate generators and relations online; doctrine b specifies the admissible sketch type, while $\mathcal{E}_b$ contains its candidate realizations
Book 3	Tangent categories and involution algebroids	Admissible local variation, diagnosis, and creative repair of identified structure	Tangent UOCL equips hypotheses and comparison maps with compatible $(T, p, 0, +, \ell, c)$ -structure; when an involution-algebroid realization is available, tangent repair operationalizes double involution
Book 4	Doctrines and hypothesis fibrations	The structural grammar in which presentations, models, translations, and variations make sense	The base $\mathcal{B}$ becomes learnable; accommodation changes fibers, while the Grothendieck construction keeps the doctrine, presentation, hypothesis, and transport witness in one total category

Table 1: The categorical AI tetralogy. Each volume internalizes an assumption left fixed by the preceding one. ORACLE does not replace the earlier constructions; it organizes their selection, interaction, and revision.

FROM STATIC SKETCHES TO FIBERED PRESENTATIONS. In the second book, a learner worked relative to the selected sketch: it could complete or repair that declaration, but the ambient language of admissible declarations was held fixed. ORACLE distinguishes the growing presentation from its doctrine. A learner may begin under a coarse commitment such as “category” and move, when licensed by evidence, to a more structured doctrine such as “differential category” or “stable ∞ -category.” Reindexing and transport record which earlier conclusions survive the refinement.

FROM FIXED VARIATION TO LEARNED VARIATION. The third book’s

infinitesimal mechanisms presupposed a base on which variation was defined. Tangent UOCL learns the base category and tangent structure jointly. In a tangent enhancement of the displayed hypothesis fibration, the learner must identify not merely directions in a preselected model but the projection, zero, addition, lift, and canonical flip that make those directions coherent. This is what creates the conditions under which DIAL's infinitesimal repair has an identified rather than stipulated semantics.

BEYOND TRANSLATION. A functor in the first book mapped between already available worlds. A section of the ORACLE fibration is dependent: the category in which the selected hypothesis lives varies with the doctrine. The learning problem is consequently not only to translate between worlds, but to discover which world of models the interaction supports and which changes of world preserve established knowledge.

This also gives the Piagetian dynamic a precise address. *Assimilation* is update within a fixed fiber: completing a sketch, estimating a model, or moving within its learned tangent geometry. *Accommodation* changes the base object from b to b' , then transports the settled part of the old hypothesis into $\mathcal{E}_{b'}$. The two operations are not competing learning rules. They are vertical and base-changing components of one fibered learning process.

The tetralogy can therefore be summarized by four questions: Which compositional translation? Which presentation and model? Which admissible variation and repair? Which doctrine makes those questions meaningful? ORACLE provides the meta-theory that makes the earlier three programs operational: it asks how to discover the appropriate functors, sketches, and tangent structures, and how to revise all three when interaction exposes a doctrinal obstruction.

The present book assumes none of the earlier volumes' terminology. Chapter 1 develops the working language of categories, Kan extensions, simplicial nerves, homotopy-coherent diagrams, stable ∞ -categories, and tangent structure. Later chapters return to these constructions with sharper hypotheses and decision-theoretic interpretations.

The present draft preserves the exact finite UODL and FTRL results as controlled mathematical witnesses, but places them after a new UOCL theory of effective categorical and tangent identification. In this sense the fourth book supplies the acquisition theory that makes the DIAL process feasible. It also records theorem boundaries: prediction is not realization, categorical identification is not causal identification, and a tangent comparison does not establish convergence without stability assumptions.

The central organizing principle of the book is Piagetian: assimilate what fits the current compositional world; accommodate when the world itself must be revised.

Further reading

Piaget’s account of sensorimotor intelligence supplies the developmental model of assimilation, accommodation, and the progressive construction of schemata [Piaget, 1952]; James supplies the classic description of the infant’s initially undifferentiated experience [James, 1890]. Spelke’s *What Babies Know* and its subsequent précis synthesize the evidence for six systems of core knowledge and emphasize their role in guiding learning across the lifespan [Spelke, 2022, 2023]. *Infinitesimal Creativity* provides the immediate categorical point of departure [Mahadevan, 2026b]. Gold’s theory of identification in the limit and the universal-imitation formulation of inductive and coinductive inference provide the learning-theoretic lineage [Gold, 1967, Mahadevan, 2024]. For established online convex optimization, Hazan’s tutorial provides the principal map used by the UODL specialization [Hazan, 2023]. The broader categorical AI program is developed in *Categories for Artificial General Intelligence* and the LINC book [Mahadevan, 2026a,c]. These sources supply context rather than prerequisites: Chapter 1 reintroduces the categorical language needed below.

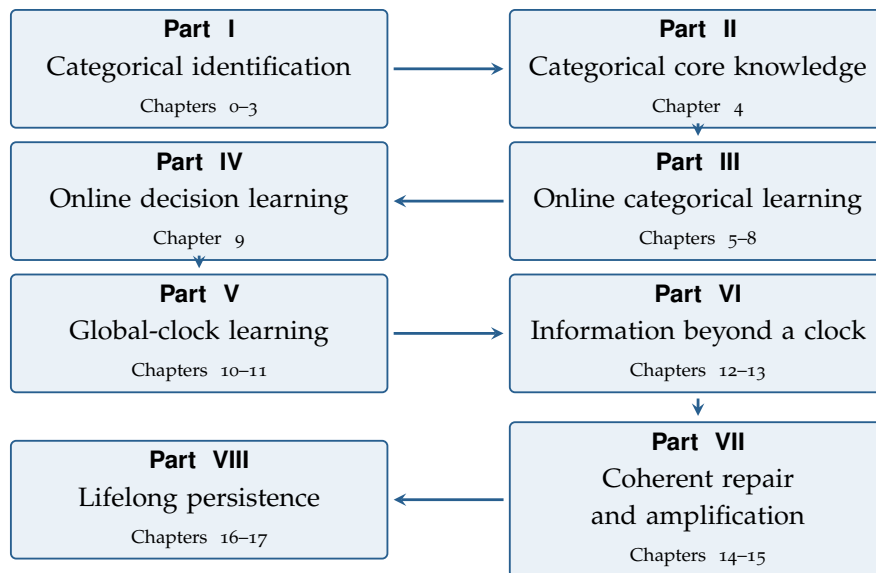
Trail Map

THE PROBLEM OF DISCOVERING COMPOSITIONAL WORLDS FROM INTERACTION IS FORMULATED THROUGH CATEGORICAL AND TANGENT IDENTIFICATION, ONLINE DECISION LEARNING, COHERENT REPAIR, STRUCTURAL AMPLIFICATION, SAFETY, AND LIFELONG TRANSPORT. Chapter 0 supplies the developmental and empirical motivation, and Chapter 1 supplies the language. Chapters 2–3 define ORACLE, presentation protocols, categorical identification in the limit, the Grothendieck hypothesis fibration, fragmentwise observability and gluing, and the boundary between assimilation and accommodation. Part II asks whether the Spelke–Piaget proposal has non-vacuous categorical content. Part III turns the ORACLE semantics into UOCL by first locating established learning paradigms inside the framework, then developing selection, probing, repair, learning of generators and relations, persistent identification, discovery of tangent structure, and the optional search for a differential presentation that realizes that geometry. It then relaxes exact identification to PACC learning through statistically covered categorical probes. Part IV introduces UODL as the independent specialization obtained by adding information, action, consistency, and observation maps.

Part V studies global-clock realizations, moving from online convex optimization to restricted bandit feedback. Part VI leaves the centralized clock behind for decentralized, asynchronous, and endogenous information and for the safety of agentic teams. Parts VII–VIII develop repair, structural amplification, and lifelong transport.

Particle physics provides a cross-part scientific thread. The same collider declaration appears first as a restricted Yoneda-style identification problem, then as active UOCL, PACC approximation, localized theory repair, and finally transport and persistence across experimental contexts. Its role is to keep presentation, response semantics, observational equivalence, and ontological claims visibly distinct.

From identifying a world to preserving useful structure across a lifetime.



The ORACLE architecture. A result is marked as a research target when its exact mathematical obligation is stated without a completed proof.

Chapter	Theme	Central question	Present status
0	Developmental grounding	How do core knowledge and Piagetian repair motivate ORACLE?	Empirical foundation
1	Categorical and homotopical toolkit	What language supports categorical identification and decision?	Tutorial foundation
2	The ORACLE program	What does a learner know before it discovers its categorical world?	Program and semantics
3	Identification under a categorical prior	When do Gold-style presentations converge, when do fragmentwise probes glue, and when must the doctrine change?	Definitions and local-to-global theorem
4	Categorical core knowledge	What would make the Spelke–Piaget proposal categorically non-vacuous?	Comparative criteria and developmental theory
5	Established paradigms as UOCL	Which categorical priors and query quotients are assumed by familiar learners?	Comparative specialization theory
6–8	UOCL algorithms and approximation	How are categorical and tangent hypotheses selected, probed, repaired, and stabilized; when can a differential realization be identified; and when is approximate identification statistically reliable?	Foundational theory
9	UODL specialization	When does a learned category support persistent online decision?	Foundational theory
10–11	Global-clock realizations	How do OCO, regularized action, and bandits instantiate UODL?	Exact finite theory
12–13	Information beyond a clock	What survives decentralized, asynchronous, or endogenous information?	Theorems and research targets
14–15	Coherent repair and amplification	When do defects admit repair, and when do weak structures compose into stronger ones?	Theorems and research targets
16	Lifelong persistence	Which learned constructions survive changes of task and environment?	Foundational research program
17	Synthesis	What has been established, and what is the next ORACLE theorem ladder?	Current conclusion

Guiding question

Under which categorical hypothesis classes, presentation protocols, probe languages, and convergence criteria can an online learner identify a compositional world and its admissible variations well enough to predict, decide, repair, and transport its knowledge?

Part I

Foundations of Categorical Identification

O

The Infant's Problem

PIAGET, CORE KNOWLEDGE, AND COMPOSITIONAL WORLDS

LONG BEFORE A CHILD CAN EXPLAIN THE WORLD, THE CHILD MUST MAKE THE WORLD EXPLAINABLE. Sensation does not announce its own ontology. A moving patch of color does not say that it is one object, that it persists when hidden, that it can be grasped, or that another person is reaching for it. A stream of speech does not arrive with reliable spaces between words, labels for parts of speech, or a grammar. A smile does not carry a typed declaration of attention, affiliation, or intent. Nevertheless, children rapidly begin to organize these changing fragments into objects, actions, places, quantities, utterances, agents, and relationships.

This is the motivating problem of the book. An intelligent learner does not merely fit a response to a fixed task. It discovers a world in which many tasks become expressible. The discovered world is *compositional* when local pieces can be connected into larger wholes, when the result of a connection can be reused, and when incompatible connections can be detected and repaired.

The purpose of this chapter is empirical and conceptual. Piaget supplies a dynamic account of construction through assimilation and accommodation. Spelke supplies evidence that construction begins with a small collection of domain-specific systems of core knowledge rather than an undifferentiated blank slate [Piaget, 1952, Spelke, 2022, 2023]. Modern machine learning supplies concrete systems against which these ideas can be tested. The categorical language begins in Chapter 1; here we first establish what that language must explain.

Guiding question

How can a learner equipped with a small repertoire of organizing capacities turn fragmentary interaction into persistent structures that support prediction, action, correction, and transfer across an open-ended lifetime?

0.1 Two complementary hypotheses

Piaget and Spelke answer different parts of the infant's problem. Their ideas are most useful here when treated as complementary rather than competing.

The Piagetian dynamics. Assimilation incorporates a new encounter into an existing schema. A child who already treats cups as graspable may assimilate a newly encountered mug into that organization. Accommodation changes the schema when the encounter cannot be absorbed without distortion: the mug may be too hot to grasp, may have an unfamiliar handle, or may turn out to be a picture rather than a physical container. Learning alternates between using a structure and revising the conditions under which it applies.

The core-knowledge prior. Spelke identifies six early-emerging systems concerning physical objects, approximate number, navigable space, object form, goal-directed agents, and social beings. These systems are automatic and domain-specific; they stand between raw perception and explicit concepts and beliefs, guide learning, and continue to operate in adults [Spelke, 2022, 2023]. Core knowledge does not tell the child which particular objects, people, languages, or rooms inhabit the world. It supplies initial distinctions through which evidence can become organized.

The two hypotheses separate *where organization begins* from *how organization changes*. Core systems provide typed entry points. Assimilation and accommodation provide a developmental update cycle. ORACLE will turn this separation into a learning architecture.

Drescher's *Made-Up Minds* provides an important computational bridge between Piaget and this program [Drescher, 1991]. His schema mechanism learns from experience by expressing discoveries in its current representational vocabulary and by extending that vocabulary with new concepts when necessary. It thereby turns constructivism into an explicit artificial cognitive architecture rather than leaving it solely as a developmental description. ORACLE is not a categorical restatement of that particular mechanism. It asks for foundations beneath the broader constructivist idea: what kinds of schemas compose, when two

presentations express the same learned world, how a failed assimilation is localized, and which universal property licenses an accommodation. In this sense, the book seeks categorical foundations for compositional constructivism.

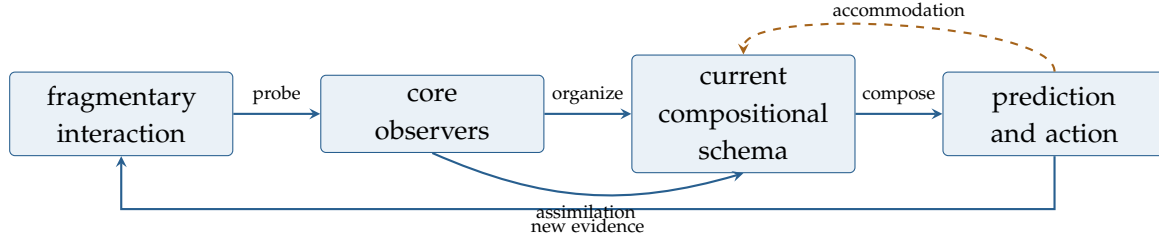


Figure 1: A developmental reading of ORACLE. Core observers turn interaction into typed evidence. Assimilation extends a viable schema; accommodation repairs the schema or its typing when prediction and action expose a defect.

0.2 Structured experience creates reusable structure

A recent study by Bowler and collaborators provides an unusually direct experimental bridge between this developmental picture, artificial networks, and animal learning [Bowler et al., 2026]. The study asks not merely whether prior training helps, but whether the *organization* of that training determines which later strategy becomes learnable. Recurrent neural networks and mice learned a temporal delayed non-match-to-sample task. Two odor cues could be short or long; the learner had to respond after non-matching pairs and withhold a response for a matching pair.

The critical manipulation was the curriculum. In the fully structured shaping condition, learners first experienced both orders of the non-matching primitive: short-long and long-short. An impoverished curriculum presented only short-long before the full task. Both curricula offered experience, reward, and a genuine subtask. They did not, however, make the same latent organization identifiable. Exposure to both orders forced a distinction between detecting cue structure and tracking elapsed time. Exposure to only one order admitted a simpler shortcut—roughly, respond after a long cue—that earned substantial reward but failed on the reversed composition.

The consequences persisted. Shaped RNN activity occupied fewer dimensions, followed less tangled trajectories, and exhibited dynamical motifs that separated cue detection from timing; they generalized trial time more effectively across contexts and were more robust to noise and temporal jitter. Networks without the appropriate shaping often converged to rigid error patterns. The animal comparison showed the same qualitative dependence: mice receiving the standard two-order shaping curriculum outperformed mice receiving short-long-only shaping after eight days of the full task, and medial entorhinal

cortex recordings showed stronger cross-context temporal structure. The principal behavioral comparison used 16 fully shaped and 4 single-order-shaped mice, so the study is evidence for the mechanism rather than a final measurement of its generality.

The ORACLE interpretation is sharper than the generic claim that curriculum matters. A curriculum is part of the *presentation* of the world. Different presentations can expose the same observations and rewards while inducing different observational quotients. The single-order presentation failed to separate a task-generating temporal construction from a reward-compatible shortcut. The two-order presentation supplied a counterposed probe: any hypothesis had to account for the same temporal relation under both orders. That pressure favored a reusable decomposition whose components could be recombined when the matching case was introduced.

Design principle

A useful curriculum is a sequence of separating presentations. It should not only make early behavior easy to reward; it should force distinctions that identify reusable generators and rule out shortcuts that cannot survive later composition.

This result also gives empirical content to persistence. Early structure did not merely accelerate optimization toward the same endpoint. It altered the dynamical organization and constrained which later strategy was reached. Conversely, poorly structured experience created a locally adequate strategy that both mice and RNNs found difficult to revise. In Piagetian terms, a curriculum can prepare schemas for conservative accommodation, or it can make later accommodation expensive by allowing brittle assimilation to harden. Nothing in the experiment proves that the brain represents a category. It does show that the ordering and compositional coverage of experience can determine whether transferable structure emerges at all—exactly the phenomenon a categorical theory of lifelong learning must explain.

0.3 *What is a compositional world?*

The word *world* does not mean the total physical universe. It means a stable domain of distinctions and transformations available to a learner. Natural language is one such world; embodied manipulation is another; social interaction is a third. They overlap, but none is reducible to a single input-output function.

For now, we use four informal questions to recognize a compositional world:

1. What are the recurring entities or states that the learner treats as stable enough to revisit?
2. Which transformations connect them, and when can two transformations be performed in sequence?
3. Which different histories count as equivalent for prediction or action?
4. Which missing constructions must be supplied to answer a new question, and which contradictions require the organization itself to change?

These questions foreshadow objects, morphisms, composition, equations, universal properties, and repair. They do not assume that there is one unique category hidden behind every domain. A categorical model is a declaration of which distinctions and composites matter to a specified family of probes. Two declarations may organize the same interaction at different scales.

Design principle

ORACLE does not seek an exhaustive inventory of appearances. It seeks the smallest persistent organization that supports the composites, comparisons, predictions, and repairs demanded by interaction.

0.3.1 Two small examples of compositional compression

Two elementary examples make this principle concrete before we turn to the larger developmental worlds.

LEGO: learn the construction theory, not every construction. Suppose a child encounters a small collection of LEGO pieces. The visible configurations proliferate rapidly: two pieces can be attached in several ways, assemblies can be joined to other assemblies, and a structure can be taken apart and rebuilt along a different sequence. Memorizing every finished configuration would miss the reusable knowledge. A more compact learner discovers pieces and interfaces as types, attachment and separation as generating operations, and relations saying when distinct building histories produce the same assembly.

This is already an ORACLE problem. The learner receives only a growing presentation of particular pieces, attempted joins, successful separations, and observed equivalences. Its hypothesis is a generative theory whose composites describe assemblies not yet encountered. Joining is generally typed and partial—not every pair of pieces connects at every interface—so the lesson is richer than an unstructured binary

operation. Assimilation expresses a new model using the existing operations. Accommodation adds an interface type, restricts a purported join, or revises a relation when a construction fails.

Diversity: learn tests, not hidden states. Rivest and Schapire’s diversity representation gives the same compression principle a predictive form [Rivest and Schapire, 1994]. For a finite-state system, an experiment consists of an action sequence followed by an observable test. Two experiments are identified when they return the same answer in every state. The learner can therefore represent the system by the resulting classes of tests and by the way actions transform those classes, rather than by enumerating or naming every hidden state.

Categorically, actions act on tests by precomposition, and observational equivalence forms a quotient of the generated test family. A compact, separating quotient answers every admitted predictive question even when the underlying state space is enormous. The qualification *admitted* is essential: tests sufficient for prediction may omit distinctions needed for control, causality, or safety. Chapter 5 develops the construction formally; here it supplies a simple moral. ORACLE should learn generators, transformation rules, and a query-relative quotient before it tries to reconstruct an exhaustive world state.

Design principle

Compositional compression has two dual faces. A generating theory builds many possible wholes from a few operations, while a separating test family distinguishes only the aspects of those wholes required by the declared questions. ORACLE must learn both without confusing compactness with completeness.

0.4 Language as a discovered world

Consider language before considering text. The infant does not receive a curated corpus of sentences. Speech is a temporally extended acoustic stream, mixed with gesture, gaze, repeated routines, and the physical consequences of other people’s actions. Even the boundaries of the candidate units are part of the learning problem. Eight-month-old infants can exploit transitional statistics to segment continuous speech, demonstrating that structure can be recovered before explicit instruction or a mature lexicon [Saffran et al., 1996].

Segmentation, however, is only the beginning. A learner must discover units at several scales and the lawful ways in which those units interact. Phonetic contrasts participate in syllables; syllables in words; words in phrases and utterances. Meaning introduces further structure.

The expressions “cup,” “blue cup,” and “bring me the blue cup” are not merely three strings of increasing length. They support reusable operations: identifying a kind, restricting a referent, specifying a goal, and recruiting another agent.

An ORACLE declaration of this world might therefore treat linguistic units, meanings, or information states as objects, and grammatical, semantic, or pragmatic transformations as morphisms. Composition records when one transformation can lawfully follow another. Equations record distinct expressions or derivations that are equivalent for the probes currently at issue.

The Piagetian distinction appears immediately. A child assimilates a novel utterance when its pieces and use fit an existing construction. The child accommodates when a familiar form behaves differently from the current schema—when a plural is irregular, a word changes meaning with context, or an apparently grammatical composition fails pragmatically. Accommodation is not simply a larger parameter update. It may revise the inventory of units, the typing of a transformation, or an equation previously taken to hold.

Modern sequence models make the contrast sharp. They can absorb enormous amounts of distributional structure, yet success on familiar mixtures does not by itself establish systematic reuse on novel compositions. The SCAN experiments, for example, were designed to separate interpolation from such systematic compositional generalization [Lake and Baroni, 2018]. For ORACLE, the empirical question is not merely whether the next fragment can be predicted, but whether the learner has identified persistent constructions that continue to work when familiar parts are recombined.

This is a modeling choice, not a claim that words must always be objects or that sentences must always be morphisms. A different probe language can induce a different, equally useful categorical declaration.

0.5 *Embodied action as a discovered world*

Learning to reach, grasp, stand, and walk presents a different kind of fragmentation. Here an action changes the evidence that will arrive next. A turn of the head changes the visual field; a reach changes both proprioception and the location of an object; a failed step changes balance and the set of available recoveries. The learner is simultaneously discovering a world and perturbing its presentation.

Spelke's systems for objects, form, and navigable space supply plausible initial observers. They distinguish a bounded persisting thing from a passing texture, a support relation from coincidence, and a stable layout from the learner's changing viewpoint. Those distinctions do not contain the skill of grasping a particular cup. They make the relevant regularities available for learning.

Embodied composition can be seen at three scales:

1. sensorimotor fragments, such as orienting the hand or shifting weight;
2. reusable skills, such as reaching, grasping, releasing, or recovering balance; and
3. activities assembled from skills, such as picking up a cup and carrying it without spilling.

The crucial word is *when*. Grasping composes with lifting only when the grasp is secure; stepping composes with transferring weight only while a support condition holds. A compositional action world must therefore encode domains of admissibility and not just concatenate motor commands.

Learned world models in machine learning separate, at least conceptually, a predictive model of dynamics from the controller that uses it [Ha and Schmidhuber, 2018]. ORACLE sharpens that separation. Learning that an action leads to an observation is not yet learning which actions compose, which histories are equivalent, or which failure calls for a local repair. Assimilation refines a skill inside a stable organization; accommodation may split one skill into context-dependent variants, introduce a new intermediate state, or withdraw an unsafe composite.

o.6 Social interaction as a discovered world

The social world is not a single dynamical system observed from outside. It is populated by other learners with private information, changing goals, and their own models of the interaction. Gaze direction, turn-taking, imitation, joint attention, refusal, and repair are fragments from which the infant must infer what kinds of agents and relationships are present.

Spelke’s distinction between goal-directed agents and social beings is important here. Efficient motion toward an object may support a prediction of goal-directed action; reciprocal attention or affect supports different expectations about social engagement. Neither system gives the infant a complete theory of a particular person. They provide typed probes through which repeated interaction can be organized.

Machine theory-of-mind systems illustrate one computational analogue: a model observes agents’ behavior and learns to predict their future actions from their inferred characteristics and mental states [Rabinowitz et al., 2018]. Prediction is valuable, but it is not yet causal, normative, or trustworthy interpretation. An ORACLE learner must also preserve the provenance of an inference: which observations support the attribution, which other agents supplied information, and which conclusions must be repaired when an expectation fails.

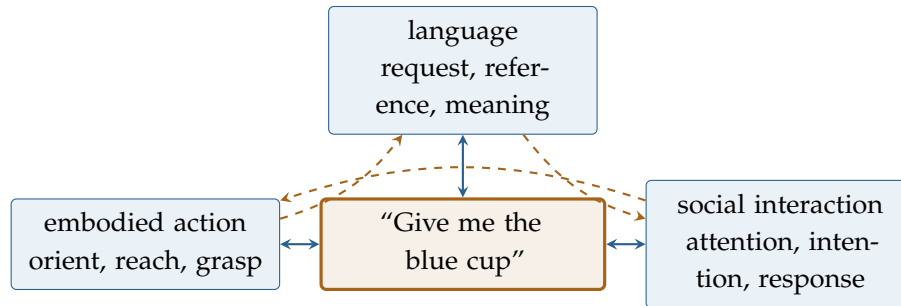


Figure 2: A single episode composes structures learned in several worlds. The central achievement is not three independent predictions but a coherent cross-domain composite whose failure can be localized and repaired.

Social accommodation is especially delicate. If another person acts unexpectedly, the learner might revise a belief about that person's goal, a general rule about social conduct, or its own interpretation of the episode. These repairs have different consequences. A viable theory must localize the defect rather than globally overwriting everything learned about the person or the practice.

0.7 *Three worlds, one learner*

The three examples are not isolated modules. Imagine that a caregiver says, "Give me the blue cup." The episode recruits language to identify an action and its argument, core object and form systems to locate the referent, spatial and motor organization to reach and grasp it, and social knowledge to interpret the utterance as a request. Success depends on composites that cross the boundaries between learned worlds.

This intersection exposes why an inventory of separate task solvers is not a theory of development. The child must retain enough typing to know which linguistic construction selects which object, which social relation licenses which response, and which motor transformation realizes it. Conversely, action supplies evidence that can repair language or social interpretation: the caregiver's correction may reveal that the wrong cup was selected, not that grasping has failed.

The example also suggests a criterion for lifelong utility. A construction earns persistence when it participates reliably in many composites across contexts and domains. Such a construction is valuable not because it solved one benchmark, but because it becomes part of the learner's reusable account of its world.

0.8 *Particle physics as a discovered world*

Developmental learning is not the only process that discovers compositional worlds from interaction. Science does so collectively and over much longer timescales. Particle physics supplies an unusually clear example because its objects are known through the transformations,

collisions, and detector responses in which they participate.

At the kinematic level, relativistic elementary particle types are classified by irreducible positive-energy unitary representations of the Poincaré group (more precisely, of the appropriate covering group), with mass and spin or helicity emerging as representation-theoretic invariants [Wigner, 1939]. The Standard Model adds much more than this classification. Its internal gauge structure, matter representations, symmetry breaking, and interaction terms form an algebraically presented compositional theory. The gauge-theoretic line begins with Yang–Mills theory and includes the electroweak synthesis exemplified by Weinberg’s lepton model [Yang and Mills, 1954, Weinberg, 1967]. Tensor products and their decompositions constrain composite systems and possible channels; interaction vertices generate processes; conservation and gauge laws impose relations among them.

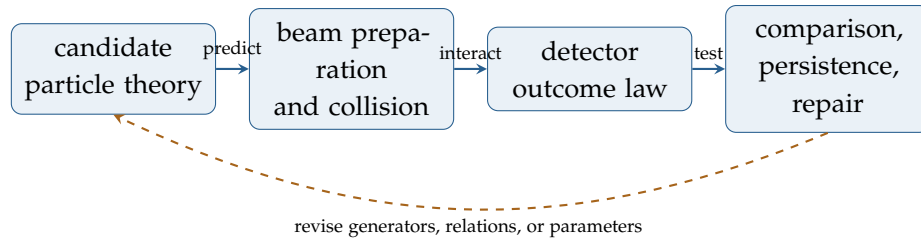
In this broad categorical sense, the Standard Model is not a catalogue of every possible physical state. It is a compact theory that generates allowed processes and predicts their observable consequences. This is the same pair of compression principles introduced in Section 0.3.1: LEGO illustrates generation from reusable operations, while diversity representations illustrate identification through a compact family of separating experiments. Particle physics combines both at the scale of a mature physical science.

An accelerator can accordingly be viewed as a *restricted physical Yoneda instrument*. Let $\mathcal{H}$ be a category of candidate physical theories and $\mathcal{P}$ a category of experimentally realizable beam preparations, energies, collision channels, and detector configurations, with its morphisms encoding refinement and restriction of probes. A candidate H induces a response functor

$$\text{Resp}_H : \mathcal{P}^{\text{op}} \longrightarrow \mathbf{Stoch}, \quad (1)$$

assigning to each probe a law over detector outcomes. By varying probes, physicists construct a response profile of H . The Yoneda analogy is that an object is characterized relationally by how every admissible probe maps into, acts on, or responds to it.

The qualification *restricted* is essential. The Yoneda lemma assumes all morphisms from all objects and exact access to their action. A collider offers only physically constructible probes, finite noisy samples, bounded energy, and detector-mediated observations. Moreover, Equation (1) is a statistical response functor, not automatically a representable hom-functor. Two theories with equivalent responses on $\mathcal{P}$ are experimentally indistinguishable under the current doctrine even if a richer probe category could separate them.



Collider science as online compositional identification. A compact theory generates predictions for controlled probes; detector statistics test those predictions; localized discrepancies update parameters or repair the theory's particles, interactions, and relations.

This example also clarifies scientific assimilation and accommodation. A new measurement can refine masses, couplings, or background models inside an accepted theory. A stable unexplained response may instead demand a new particle type, interaction generator, symmetry, or relation. Responsible repair is localized: detector calibration and nuisance models are tested before the entire particle ontology is replaced. Science advances by growing and revising a generative theory under increasingly separating interactions, not by storing the state of the universe.

Design principle

The categorical analogue of scientific understanding is not exhaustive state reconstruction. It is a compact generative theory together with a sufficiently separating, explicitly delimited family of experimental response functors.

0.9 A developmental contract for machine learning

The developmental sources do not prescribe a unique algorithm. They impose a contract on any proposed realization of ORACLE. The informal commitments and their observable consequences are summarized below.

Developmental idea	Categorical reading	Observable machine-learning test
Core knowledge	A small initial family of typed observers, invariants, and admissible distinctions	Learn new domains from sparse interaction without reconstructing objecthood, agency, space, or quantity from scratch
Assimilation	Extend a current presentation while preserving its validated composites and equations	Incorporate a compatible example without erasing unrelated capabilities
Accommodation	Repair objects, arrows, typings, or equations when evidence cannot be coherently incorporated	Localize contradiction and revise the responsible warrant rather than globally retraining
Composition	Reuse identified transformations inside larger constructions, including across learned worlds	Generalize to novel combinations and expose where an attempted composition is undefined
Development	Grow a sequence of observationally adequate presentations whose useful structure persists	Improve prediction and action while retaining and transporting earlier constructions

These obligations distinguish developmental learning from scale alone. Current systems offer important pieces—statistical discovery, learned representations, predictive world models, and increasingly general action policies—but the pieces are rarely coupled to an explicit account of structural identification and repair. The broader proposal that machines should learn and think through human-like capacities provides a useful point of comparison [Lake et al., 2017]; ORACLE focuses that proposal on the discovery and preservation of compositional organization.

Design principle

A learner has not discovered a compositional world merely because it predicts samples from that world. Discovery is evidenced by systematic reuse, well-typed composition, localized repair, and persistence under continued interaction.

0.10 What the categorical theory must provide

The examples leave five mathematical obligations for the rest of the book.

1. A *presentation language* must state candidate entities, transformations, composites, and equations without requiring the whole world in advance.
2. An *identification criterion* must say when successive presentations count as learning the same world, especially when only a restricted family of questions can be asked.
3. An *online semantics* must represent the fact that information is re-

vealed through a growing, action-dependent history.

4. A *theory of coherent repair* must distinguish a bad estimate inside a sound declaration from evidence that forces the declaration itself to change.
5. A *persistence principle* must explain why a learned construction should survive and transport to future problems.

Categories, functors, natural transformations, Kan extensions, nerves, horns, and tangent structure will be introduced because they answer these obligations, not as decoration imposed on the examples. Chapter 1 now develops that language. The chapters after it return to the infant's problem in an abstract form: what can be identified from interaction, what is visible to a probe, how incomplete compositions are filled, and when a defect requires accommodation rather than further assimilation.

Further Reading

Piaget's account of sensorimotor development and the construction of object permanence remains the classical source for assimilation and accommodation [Piaget, 1952]. Spelke's synthesis develops the empirical case for core systems of objects, number, space, form, agents, and social beings [Spelke, 2022, 2023]. Saffran, Aslin, and Newport provide a landmark demonstration of statistical learning in infant speech segmentation [Saffran et al., 1996]. Bowler and collaborators connect behavioral shaping to compositional abstraction in RNNs and mice, and relate the resulting strategies to dynamical organization in the medial entorhinal cortex [Bowler et al., 2026].

Drescher turns Piagetian constructivism into a working schema mechanism that can assimilate experience into an existing vocabulary and construct new concepts [Drescher, 1991]. It is the closest historical precursor to the book's developmental ambition. The categorical identification, equivalence, universal construction, and coherent-repair apparatus developed here should be read as a proposed foundation for that broader style of compositional discovery, not as a claim about the formal structure of Drescher's implementation.

For modern machine learning, Lake and collaborators articulate a research program centered on human-like learning capacities [Lake et al., 2017], while Lake and Baroni test systematic compositional generalization directly [Lake and Baroni, 2018]. Ha and Schmidhuber illustrate learned predictive world models for control [Ha and Schmidhuber, 2018]; Rabinowitz and collaborators study learned models of other agents [Rabinowitz et al., 2018]. These works do not implement ORACLE. They isolate capabilities and failure modes that a categorical theory of developmental learning must explain.

For the particle-physics example, Wigner gives the representation-theoretic classification of relativistic particle types [Wigner, 1939]. Yang and Mills introduce non-Abelian gauge invariance [Yang and Mills, 1954], while Weinberg’s lepton model is a foundational step in the electroweak theory [Weinberg, 1967]. These sources establish the physical structures; the restricted-Yoneda interpretation developed here is an ORACLE reading of experimental theory formation.

A Categorical and Homotopical Toolkit

THE LANGUAGE REQUIRED FOR CATEGORICAL IDENTIFICATION AND DECISION

ORACLE and Universal Online Decision Learning combine several mathematical languages. Ordinary category theory describes typed composition and universal extension. Algebraic theories and sketches separate compact generators and relations from their functorial models. Simplicial sets encode histories of every finite compositional depth. Categorical homotopy theory distinguishes presentation from semantic content and retains coherent spaces of possible repairs. Tangent categories describe infinitesimal variation. This chapter introduces exactly the portion of each language required by the book.

It is a tutorial, not a replacement for standard references. Mac Lane and Riehl provide systematic introductions to ordinary category theory [Mac Lane, 1971, Riehl, 2016]; Riehl develops the homotopical and enriched material used here [Riehl, 2014]. The purpose of the chapter is to make later declarations readable and to expose every assumption used by the first identification and derived UODL theorems.

Guiding question

What is the smallest mathematical language in which an online learner can construct a composite, preserve it across revelation, discover that its presentation was nonessential, and repair it coherently when new evidence arrives?

1.1 Categories and typed composition

Definition 1.1 (Category). A category $\mathcal{C}$ consists of objects $X, Y, \dots$, morphisms $f : X \rightarrow Y$, identity morphisms 1_X , and an associative composition

$$\mathcal{C}(Y, Z) \times \mathcal{C}(X, Y) \longrightarrow \mathcal{C}(X, Z), \quad (g, f) \longmapsto g \circ f,$$

with identities acting neutrally.

The definition says which processes compose. It does not say that every pair of processes does, nor that a legal composite is useful, causal, or optimal. In ORACLE, the category itself may be the unknown predictive model. In UODL, its objects may be information states and its morphisms may be admissible decisions, evidence transports, or repairs.

Example 1.2 (A decision history). A composable pair

$$x_0 \xrightarrow{a} x_1 \xrightarrow{b} x_2$$

has a composite $b \circ a : x_0 \rightarrow x_2$. Associativity makes the result of a longer finite history independent of parenthesization. It does not imply that the composite was known before the two steps were observed.

Definition 1.3 (Functor). A functor $F : \mathcal{C} \rightarrow \mathcal{D}$ maps objects and morphisms while preserving identities and composition:

$$F(1_X) = 1_{F(X)}, \quad F(g \circ f) = F(g) \circ F(f).$$

A functor is a compositional translation. A diagram of shape $\mathcal{J}$ in $\mathcal{D}$ is simply a functor $D : \mathcal{J} \rightarrow \mathcal{D}$. The indexing category records the formal wiring; the functor supplies a realization.

Definition 1.4 (Natural transformation). Given $F, G : \mathcal{C} \rightarrow \mathcal{D}$, a natural transformation $\eta : F \Rightarrow G$ assigns a component $\eta_X : F(X) \rightarrow G(X)$ to every X so that, for each $f : X \rightarrow Y$,

$$\begin{array}{ccc} F(X) & \xrightarrow{F(f)} & F(Y) \\ \eta_X \downarrow & & \downarrow \eta_Y \\ G(X) & \xrightarrow{G(f)} & G(Y) \end{array}$$

commutes.

Naturality is the first model of coherent change. Independently replacing the objects of a decision system does not define a repair unless the replacements respect the morphisms connecting them.

1.2 Doctrines: the structure known in advance

The word *doctrine* occurs throughout this book because saying that a learner expects “a category” is rarely precise enough. One learner may expect finite products, another a tensor product, another tangent structure, and another only a declared family of observable questions. A doctrine names this ambient package of admissible structure and the maps allowed to preserve it.

The first reading habit throughout the book is *type first*. Before interpreting a formula, identify the source and target of every map.

Definition 1.5 (Categorical doctrine). A *categorical doctrine* $\mathcal{D}$ is a specified 2-category $\mathbf{Doc}_{\mathcal{D}}$ together with a forgetful 2-functor

$$U_{\mathcal{D}} : \mathbf{Doc}_{\mathcal{D}} \longrightarrow \mathbf{Cat}.$$

Its objects are categories equipped with the structure declared by $\mathcal{D}$, its 1-cells are the admitted structure-preserving functors, and its 2-cells are the admitted transformations. The doctrine also fixes which of its 1-cells count as equivalences. When the structure is algebraic, $\mathbf{Doc}_{\mathcal{D}}$ is often presented as the 2-category of pseudoalgebras for a 2-monad on $\mathbf{Cat}$; a sketch gives another common presentation.

Thus a doctrine is not one particular world and is not normally one particular theory. It specifies the *kind of structure* a theory or world may carry, how that structure is transported, and the resolution at which two such structures are identified.

Definition 1.6 (Theory and model relative to a doctrine). A $\mathcal{D}$ -*theory* is a small object $\mathbb{T}$ of $\mathbf{Doc}_{\mathcal{D}}$. If $\mathcal{V}$ is a semantic object of the same doctrine, a *model* of $\mathbb{T}$ in $\mathcal{V}$ is an admitted 1-cell

$$M : \mathbb{T} \longrightarrow \mathcal{V}.$$

Consequently the category of models is the hom-category $\mathbf{Doc}_{\mathcal{D}}(\mathbb{T}, \mathcal{V})$.

The hierarchy is easiest to see by comparing mathematical, developmental, and scientific examples.

Finite-product doctrine. Its objects are categories with finite products and its maps preserve those products (strictly or coherently, according to the declared convention). A Lawvere theory is a small theory in this doctrine with a distinguished generator. The Lawvere theory of groups is one such theory; a product-preserving functor $\mathbb{T}_{\text{Grp}} \rightarrow \mathbf{Set}$ is a group. Hence groups do not constitute the doctrine: they are models of a theory living within it.

Symmetric-monoidal doctrine. Its objects are symmetric monoidal categories and its maps are admitted symmetric monoidal functors. The category $(\mathbf{Vect}_k, \otimes, k)$ is an instance of this doctrine, and its objects are vector spaces. A PROP is a particularly useful small theory in the doctrine; a symmetric monoidal functor from that PROP into $\mathbf{Vect}_k$ realizes its abstract operations as multilinear structure. Associators, unitors, and braidings record the coherence required when tensor products are not treated as literally strict.

Core-knowledge doctrine. Spelke's hypothesis admits a categorical reading as a claim about the kinds of world models for which a child is prepared. The proposed doctrine is heterogeneous: it combines

constraints on objects, approximate magnitude, space, form, goal-directed agents, and social-linguistic interaction, together with interfaces among them. A child does not begin knowing which particular objects or people exist, which language is spoken, or which regularities hold locally. Those are theories and models to be identified within the restricted category of possible worlds. Section 4.1 develops this proposal and its empirical burden.

Symmetry doctrine for matter. Particle physics treats admissible entities and processes through group actions, representations, equivariant maps, tensorial composition, and invariance laws. At the kinematic level, relativistic elementary particle types are classified by suitable irreducible representations of the Poincaré group; internal gauge symmetries add further representation data [Wigner, 1939]. Schematically, the doctrine says that physical theories and their maps must respect such symmetry and compositional structure. A particular choice of spacetime and gauge groups, matter representations, interaction generators, and relations belongs to a theory within that doctrine. Particle species and scattering processes then appear as particular objects and morphisms in its models. Section 0.8 explains how collider interaction probes this structured hypothesis space.

These examples separate four levels that later repair mechanisms must not confuse:

$$\boxed{\text{doctrine}} \supset \boxed{\text{theory or sketch}} \xrightarrow{\text{model}} \boxed{\text{semantic category}},$$

$$\boxed{\text{semantic category}} \ni \boxed{\text{particular objects}}.$$

Changing a group while retaining the group laws changes a model. Changing the equations changes the theory. Abandoning finite products as the governing notion of combination changes the doctrine. Likewise, learning a new word or object need not alter the child’s core doctrine, and discovering a new particle need not abandon symmetry. In either case, doctrinal repair is the stronger move: it changes the prior notion of which worlds and transports are admissible, rather than adding another inhabitant to the current world.

We use two qualified variants of the term. A *preservation doctrine* specifies which designated limits, colimits, tensors, or other constructions a realization must preserve. A *query doctrine* specifies an admissible family of probes, its closure rules, and the resulting observational equivalence of hypotheses. A *core doctrine* may bundle structural, preservation, and query commitments into the learner’s prior. These are deliberate extensions of the same organizing idea: each states what structure is typed, what maps respect it, and what comparisons count.

Boundary

“Doctrine” is not a decorative synonym for “assumption.” Every doctrine used below must expose at least its structured objects, admitted maps, and equivalences. If a theorem also depends on chosen probes or preservation requirements, those must be declared separately. This is what makes assimilation within a doctrine distinguishable from accommodation that repairs the doctrine itself.

1.3 *Size, categories of categories, and indexed hypotheses*

The phrase “the category of all categories” requires a size convention. Choose Grothendieck universes $\mathcal{U} \in \mathcal{V}$. The category $\mathbf{Cat}_{\mathcal{U}}$ of $\mathcal{U}$ -small categories is generally not $\mathcal{U}$ -small, but it is an object of the larger $\mathbf{Cat}_{\mathcal{V}}$. This prevents the ambient hypothesis space from being mistaken for one of its own small objects.

Categories, functors, and natural transformations form a 2-category. Hence a learner may compare candidate worlds by functors and compare those comparisons by natural transformations; equality of hypothesis codes is not the intended semantics.

Definition 1.7 (Grothendieck construction). For a functor $H : \mathcal{T}^{\text{op}} \rightarrow \mathbf{Cat}$, the Grothendieck construction $\int_{\mathcal{T}} H$ has objects (T, x) with $x \in H(T)$. A morphism $(T, x) \rightarrow (T', x')$ consists of a map $u : T \rightarrow T'$ and a morphism $x \rightarrow H(u)(x')$ in $H(T)$. The projection

$$\int_{\mathcal{T}} H \longrightarrow \mathcal{T}$$

is a fibration whose fiber over T is equivalent to $H(T)$.

For ORACLE, $\mathcal{T}$ is the category of interaction transcripts and $H(T)$ is the category of hypotheses consistent with transcript T . The construction turns a changing family of hypothesis categories into one typed total space.

1.4 *Universal properties, limits, and colimits*

A universal property characterizes an object through its maps to or from all competing objects. Such characterizations are invariant under unique isomorphism and therefore separate structural role from implementation.

Definition 1.8 (Limit and colimit). For a diagram $D : \mathcal{J} \rightarrow \mathcal{C}$, a limit is a terminal cone

$$\lim_{\mathcal{J}} D \longrightarrow D.$$

A colimit is an initial cocone

$$D \longrightarrow \operatorname{colim}_{\mathcal{J}} D.$$

Limits enforce simultaneous compatibility. Products, pullbacks, and equalizers are familiar finite limits. Colimits freely assemble compatible pieces. Coproducts, pushouts, and coequalizers are familiar finite colimits.

Construction	Universal reading	UODL reading
Product	maps into all components	satisfy several observations together
Pullback	pairs agreeing over a shared image	compatible evidence or repairs
Equalizer	inputs on which two maps agree	exact consistency locus
Coproduct	freely choose among components	retain alternative candidates
Pushout	glue along a shared part	extend a declaration by an overlap
Coequalizer	identify two declared routes	quotient presentation differences

Table 1.1: Finite universal constructions and their recurring decision roles.

Boundary

A universal property is not automatically a causal property. A pullback may encode compatibility and a Kan extension may encode optimal extension without identifying the effect of an intervention. Causal meaning requires an additional intervention semantics.

1.5 Algebraic theories, sketches, and generative worlds

A category records which objects and processes exist and how they compose. It need not reveal a small vocabulary from which those objects and processes are generated. Learning a world category and learning its *theory of construction* are therefore different problems. The second can be much more compact: a few operations and equations may account for indefinitely many composites.

A signature names sorts and primitive operations. An algebraic theory also specifies how those operations may be substituted and which composite expressions are equal. Lawvere expressed a one-sorted finitary algebraic theory as a category $\mathbb{L}$ with finite products, generated by a distinguished object X and its finite powers

$$1, X, X^2, X^3, \dots$$

A morphism $X^n \rightarrow X^m$ is an abstract m -tuple of n -ary operations, composition is substitution, and equality of morphisms records equational laws.

Definition 1.9 (Lawvere theory and model). A one-sorted *Lawvere theory* is a small finite-product category $\mathbb{L}$ generated by one object X . A model of $\mathbb{L}$ in a category $\mathcal{C}$ with finite products is a product-preserving functor

$$M : \mathbb{L} \longrightarrow \mathcal{C}.$$

Natural transformations between such functors are homomorphisms of models.

The distinction between theory and model is essential for learning. The theory of groups contains multiplication, unit, inverse, and their equations. A particular group is one model of that theory. Likewise, a learned parameter vector or one encountered physical assembly is a realization, not the generating theory shared across realizations.

A completed category of formal operations is often too large to learn or write down directly. A sketch presents it economically.

Definition 1.10 (Sketch and completion). A *sketch* S consists of a graph of generating objects and arrows, specified commutative diagrams, and designated cones or cocones. A model sends this data into a semantic category and preserves the declared equations and universal constructions. Relative to a preservation doctrine $\mathcal{D}$, write

$$\mathrm{Th}_{\mathcal{D}}(S)$$

for the category freely completed under the structure demanded by $\mathcal{D}$, modulo the declared relations.

Thus a sketch is neither an informal drawing nor a partial model. It is a finite presentation whose completion generates a theory and whose structure-preserving functors supply its semantics. Finite-product sketches present algebraic theories; finite-limit sketches can also specify typed domains of partially defined operations.

Definition 1.11 (PROP and algebra). A *PROP* is a strict symmetric monoidal category whose objects are the natural numbers and whose tensor on objects is addition. A morphism $m \rightarrow n$ represents an operation with m inputs and n outputs. An algebra of a PROP $\mathbb{P}$ in a symmetric monoidal category $\mathcal{V}$ is a symmetric monoidal functor

$$A : \mathbb{P} \longrightarrow \mathcal{V},$$

up to the chosen strictness or coherence convention.

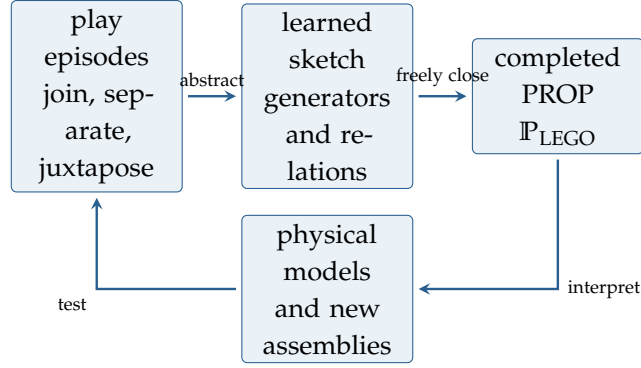


Figure 1.1: Theory learning in the LEGO world. The learner seeks not only a catalogue of assemblies but a compact presentation whose composites generate new construction programs and whose models can be tested in play.

PROPs are useful when juxtaposition is available but copying and deletion are not automatic. They are therefore natural theories of physical construction, circuits, resources, and multi-input/multi-output processes. Lawvere theories and PROPs are neighboring doctrines, not interchangeable names: Cartesian structure freely supplies copying and deletion, while a general symmetric monoidal theory does not.

1.5.1 The LEGO theory

Suppose a child encounters pieces that can be placed side by side, joined, and taken apart. A minimal one-sorted signature counts connected components. It contains parallel juxtaposition $\otimes$, an attachment generator

$$j : 2 \longrightarrow 1,$$

and a detachment generator

$$s : 1 \longrightarrow 2.$$

If brick shapes or interfaces must be typed separately, the corresponding structure is a colored PROP rather than a one-sorted PROP. Symmetry exchanges independent inputs. Interchange says that independent attachments can be performed in either order. Where joining preserves enough alignment and identity information to be exactly reversible, the child may discover the equations

$$s \circ j = 1_2, \quad j \circ s = 1_1.$$

Closing these generators under composition and tensor produces a presented PROP $\mathbb{P}_{\text{LEGO}}$. Its morphisms are formal construction and deconstruction programs; an algebra

$$A : \mathbb{P}_{\text{LEGO}} \longrightarrow \mathcal{V}_{\text{phys}}$$

interprets those programs as actual configurations and manipulations.

Proposition 1.12 (Presented LEGO PROP). *For any typed symmetric monoidal signature Σ of LEGO operations and set E of well-typed equations, there is a PROP $\mathbb{P}_{\Sigma,E} = \text{FPROP}(\Sigma)/E$ such that symmetric monoidal functors from it to $\mathcal{V}$ are precisely interpretations of the generators in $\mathcal{V}$ satisfying E .*

Proof. Form the free strict symmetric monoidal category on the typed generators, then quotient each hom-set by the smallest symmetric-monoidal congruence containing E . A symmetric monoidal interpretation of the generators extends uniquely through the free construction and factors through the quotient exactly when it satisfies the equations. $\square$

Boundary

Physical detachment is not automatically a two-sided inverse. If attachment forgets orientation, damages a piece, permits several decompositions, or is defined only for compatible interfaces, then the appropriate structure may be a partial inverse, a groupoid of reversible moves, a localization, or a finite-limit theory of admissible attachment domains. The theory is learned from the observed laws; invertibility must not be built in merely because an action is colloquially called *taking apart*.

The LEGO example displays three distinct identification targets. The child may learn a category of configurations and transformations, a small sketch of generators and relations presenting that category, and a functorial model connecting the formal operations to perception and action. Two different sketches can present equivalent theories, so successful theory learning cannot mean recovering a privileged naming scheme or generating list. It must be stated up to a doctrine-relative equivalence of theories or their model semantics.

1.6 Adjunctions and Kan extensions

Definition 1.13 (Adjunction). Functors $L : \mathcal{C} \rightleftarrows \mathcal{D} : R$ form an adjunction $L \dashv R$ when there are natural equivalences

$$\mathcal{D}(LX, Y) \cong \mathcal{C}(X, RY).$$

The left adjoint L preserves colimits; the right adjoint R preserves limits.

Kan extensions generalize extension, aggregation, restriction, and quantification. Let $J : \mathcal{A} \rightarrow \mathcal{B}$ and $F : \mathcal{A} \rightarrow \mathcal{D}$.

Definition 1.14 (Kan extensions). A left Kan extension of F along J is a functor $\text{Lan}_J F : \mathcal{B} \rightarrow \mathcal{D}$ universal among maps from F to functors

restricted along J :

$$\text{Lan}_J \dashv J^*.$$

A right Kan extension is universal in the opposite direction:

$$J^* \dashv \text{Ran}_J.$$

The pointwise formulas for Kan extensions are indexed by comma categories.

Definition 1.15 (Comma category). Given functors

$$\mathcal{A} \xrightarrow{S} \mathcal{C} \xleftarrow{T} \mathcal{B},$$

the *comma category* $(S \downarrow T)$ has objects (a, b, α) , where $a \in \mathcal{A}$, $b \in \mathcal{B}$, and $\alpha : S(a) \rightarrow T(b)$. A morphism

$$(f, g) : (a, b, \alpha) \longrightarrow (a', b', \alpha')$$

consists of $f : a \rightarrow a'$ and $g : b \rightarrow b'$ satisfying

$$T(g) \circ \alpha = \alpha' \circ S(f).$$

Thus a morphism in a comma category is precisely a commuting square in $\mathcal{C}$ between the two displayed comparison arrows.

An object $b \in \mathcal{B}$ may be regarded as a functor $b : \mathbf{1} \rightarrow \mathcal{B}$. Consequently, $(J \downarrow b)$ has objects $(a, \alpha : J(a) \rightarrow b)$, while $(b \downarrow J)$ has objects $(a, \beta : b \rightarrow J(a))$. Each carries an evident projection π to $\mathcal{A}$, selecting the object a .

When the required (co)limits exist, Kan extensions are calculated pointwise:

$$(\text{Lan}_J F)(b) \cong \text{colim}_{(a, \alpha) \in (J \downarrow b)} F(a), \quad (\text{Ran}_J F)(b) \cong \lim_{(a, \beta) \in (b \downarrow J)} F(a).$$

Equivalently, these are the colimit and limit of $F \circ \pi$. The arrow orientation matters: maps $J(a) \rightarrow b$ contribute to left extension, whereas maps $b \rightarrow J(a)$ contribute to right extension.

Design principle

UODL reads the left Kan stage as universal candidate generation or evidence extension and the right Kan stage as universal consistency. This is a typed factorization, not a claim that every learning algorithm is secretly a Kan extension.

Definition 1.16 (Universal Decision Learning). Given

$$\mathcal{A} \xrightarrow{J} \mathcal{B} \xrightarrow{K} \mathcal{Q}, \quad F : \mathcal{A} \rightarrow \mathcal{D},$$

the UDL semantic object is

$$\text{U}(F) = \text{Ran}_K \text{Lan}_J F$$

whenever the displayed extensions exist.

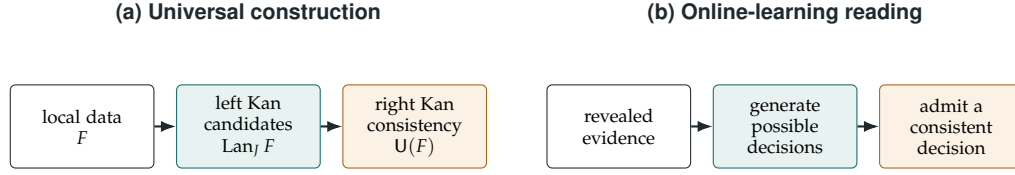


Figure 1.2: The recurring visual dictionary for UDL. Left Kan extension freely generates or aggregates candidates from available data; right Kan extension retains the candidates compatible with the declared consistency shape. The construction is universal until an observer supplies task-specific meaning.

1.7 Functor categories and online information

For categories $\mathbb{T}$ and $\mathcal{D}$, the functor category $[\mathbb{T}, \mathcal{D}]$ has time-indexed diagrams as objects and natural transformations as morphisms. This is the minimal setting for persistence.

Let $\mathbb{T}$ be a poset of times. The information available at t is its principal past

$$\downarrow t = \{s \in \mathbb{T} \mid s \leq t\}.$$

An action at t is online when it factors through restriction to $\downarrow t$. Equivalently, two complete histories with identical prefixes through t must induce the same action at t .

Definition 1.17. A decision family A_t is non-anticipating when

$$h|_{\downarrow t} = h'|_{\downarrow t} \implies A_t(h) = A_t(h').$$

PERSISTENCE IS STRONGER THAN REPEATED COMPUTATION. A collection of unrelated objects, one for each time, does not say what survived. A functor $\mathbb{T} \rightarrow \mathcal{D}$ registers the transport maps and their composition law. Later chapters allow those transports to be coherent only up to higher homotopy.

1.8 Simplicial sets and decision nerves

The simplex category Δ has finite nonempty ordinals $[n] = \{0 < \dots < n\}$ as objects and monotone maps as morphisms.

Definition 1.18 (Simplicial set). A simplicial set is a functor

$$X : \Delta^{\text{op}} \rightarrow \text{Set}.$$

Its elements in X_n are n -simplices. Face maps delete vertices or compose adjacent steps; degeneracy maps insert identities.

Definition 1.19 (Nerve). The nerve of a category $\mathcal{C}$ is the simplicial set

$$(N\mathcal{C})_n = \text{Fun}([n], \mathcal{C}).$$

Thus an n -simplex is a composable chain of n morphisms.

The first diagram below shows the first nontrivial case. Two successive decisions form a functor $[2] \rightarrow \mathcal{C}$. The corresponding 2-simplex remembers not only the two steps but also their composite and the witness that the triangular boundary coheres.

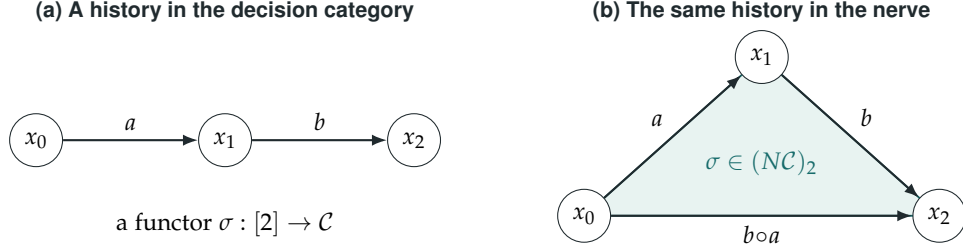


Figure 1.3: A composable two-step decision history becomes a 2-simplex in the nerve. Its faces retain the two local decisions and their composite. In an ordinary nerve the composite is unique; in a quasi-category the space of coherent fillers may contain more information.

The nerve makes compositional depth intrinsic. Its n -skeleton $\text{sk}_n X$ is generated by simplices of dimension at most n , and

$$X \cong \text{colim}_{n \geq 0} \text{sk}_n X.$$

External time and compositional depth are different indices. A learner may receive evidence without constructing a longer plan, or construct a longer internal plan before receiving new evidence.

In the simplest online protocol, one new decision morphism arrives per clock tick. A history observed from time 0 through time t is then a functor $[t] \rightarrow \mathcal{C}$, hence a t -simplex of the decision nerve. Its faces are the shorter histories obtained by omitting a time or composing adjacent steps, and its degeneracies insert identity decisions. The next diagram depicts this diagonal case $n = t$.

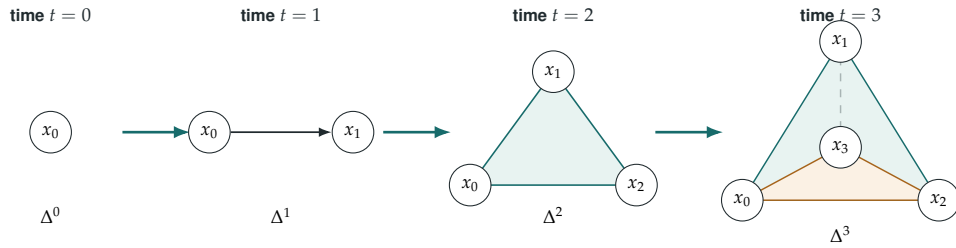


Figure 1.4: The growing simplex of a single online history. With one decision per round, each clock tick extends the observed chain and raises its compositional dimension. The simplicial closure retains every face, so later structure contains its shorter subhistories rather than overwriting them.

Design principle

If $\mathcal{C}$ is fixed, the clock does not enlarge the mathematical nerve $\mathcal{NC}$; it enlarges the learner's accessible simplicial subobject $X_t \subseteq \mathcal{NC}$, generated by histories revealed through time t . Multiple trajectories contribute multiple simplices and their shared faces. If the decision category itself is learned, both $\mathcal{C}$ and its accessible nerve may change, and that extra repair must be declared.

The category of simplices Δ/X has simplices $\Delta^n \rightarrow X$ as objects. It is the natural domain on which evidence, actions, losses, or warrants can decorate a decision nerve.

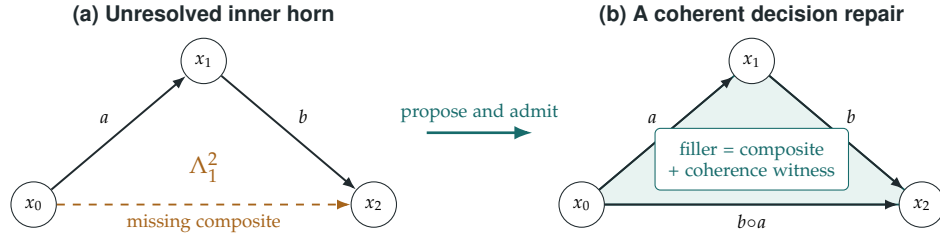
1.9 Horns, composition, and quasi-categories

The i -th horn $\Lambda_i^n \subset \Delta^n$ contains every face except the i -th. A filler extends a map $\Lambda_i^n \rightarrow X$ to $\Delta^n \rightarrow X$.

For $n = 2$, the inner horn Λ_1^2 contains

$$x_0 \xrightarrow{a} x_1 \xrightarrow{b} x_2$$

but omits the composite edge $x_0 \rightarrow x_2$. A filler supplies the missing composite together with its compositional witness.



Following Section 1.1 of Riehl and Verity [2022], we use quasi-categories as our concrete model of ∞ -categories.

Definition 1.20 (Quasi-category). A *quasi-category* is a simplicial set X in which every inner horn can be filled. Explicitly, for every $n \geq 2$, every $0 < i < n$, and every simplicial map $u : \Lambda_i^n \rightarrow X$, there is an extension $\bar{u} : \Delta^n \rightarrow X$ satisfying

$$\bar{u}|_{\Lambda_i^n} = u.$$

The filler is required to exist, but it need not be unique.

The nerve of an ordinary category has unique inner horn fillers. It fills all horns precisely when the category is a groupoid. A general quasi-category allows spaces of coherent compositions without requiring decisions to be reversible.

Example 1.21 (Rubik's Cube and Sokoban). Rubik's Cube gives a concrete world in which outer-horn filling is natural. Let X be the set of legal cube configurations and let G be the group generated by the legal face turns. The associated action groupoid

$$G \ltimes X$$

has configurations as objects and a morphism $(g, x) : x \rightarrow g \cdot x$ for each move sequence g . Its inverse is $(g^{-1}, g \cdot x)$. Consequently, the nerve $N(G \ltimes X)$ is a Kan complex: inner horns are filled by composing move sequences, while outer horns are filled by reversing moves to recover a missing side of a compositional history.

Figure 1.5: Horn filling as decision repair. The observed local steps specify a boundary but not a coherent composite. A filler completes the history. In a quasi-category a filler must exist but need not be unique, so the repair can remain a structured space of alternatives rather than a single forced action.

Sokoban has a sharply different geometry. Its objects may again be taken to be legal board configurations and its morphisms to be executable action sequences. Local walking moves are often reversible, but a push need not be. If a crate is pushed into a non-goal corner, no legal action sequence can “unpush” it. The resulting decision category still has composites, so its ordinary nerve has unique inner-horn fillers, but it is not a groupoid and some outer horns have no fillers.

The contrast is structural rather than merely a difference in difficulty. A Rubik’s Cube learner can seek a latent group action, generators, inverses, and relations. A Sokoban learner must also represent irreversible reachability and detect actions that destroy future possibilities. Thus horn filling turns reversibility into a property of the learned compositional world rather than an informal attribute attached to particular actions.

The classical algebraic counterpart of this contrast is the Krohn–Rhodes prime decomposition theorem [Krohn and Rhodes, 1965]. It represents a finite-state machine by a cascade of permutation components, whose dynamics are group-like and reversible, and reset components, whose dynamics discard information. Chapter 5 returns to the theorem as a possible learning principle: discover a compact hierarchy of reversible and irreversible generators instead of reconstructing a flat global state space.

CONVENTION. Unless another model is explicitly declared, the term ∞ -category in this book means a quasi-category. Thus its higher composition is encoded by compatible inner-horn fillers rather than by strictly unique composites.

Boundary

The term *Kan* appears in two related but distinct constructions. A Kan extension universally extends a diagram. A Kan horn-filling condition belongs to simplicial homotopy theory. UODL uses both, but one does not imply the other.

Inner horns ask whether observed local steps compose coherently. Outer horns ask whether a missing step can be recovered by solving backward. Rubik’s Cube generally permits both; Sokoban need not.

1.10 Weak equivalence and localization

Strict isomorphism is often too rigid for learned representations. Two models may encode the same observable decisions through different coordinates, factorizations, or resolutions.

Definition 1.22 (Relative category). A relative category is a pair $(\mathcal{D}, W)$, where $\mathcal{D}$ is a category and W is a declared class of weak equivalences containing all identities.

Localization formally makes every map in W invertible:

$$\mathcal{D} \longrightarrow \mathcal{D}[W^{-1}].$$

The localization is best regarded as an ∞ -category when the paths and higher witnesses relating weakly equivalent presentations matter.

Design principle

Weak equivalence is semantic data. UODL must state which observables a weak equivalence preserves. Choosing a familiar model structure does not by itself justify identifying two decision systems.

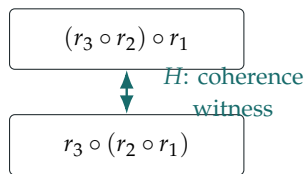
1.11 Homotopy-coherent diagrams

A strictly commuting diagram equates parallel composites. A diagram commuting up to homotopy instead supplies a path between them. A coherent diagram also supplies higher homotopies between the different ways of combining such paths. This compatibility continues through every dimension.

A category enriched in simplicial sets has a simplicial mapping object $\mathcal{C}(X, Y)$ rather than a bare set of arrows. When these mapping objects are Kan complexes, the homotopy-coherent nerve produces a quasi-category. Ordinary categories appear as the discrete special case.

WHY RETAIN THE HIGHER DATA? Suppose three consecutive repairs can be composed in several orders. Equality in the homotopy category says only that the resulting composites represent the same class. A coherent online learner must also retain witnesses comparing the orders and the higher compatibility among those witnesses.

(a) Coherence data



(b) Persistent online trace

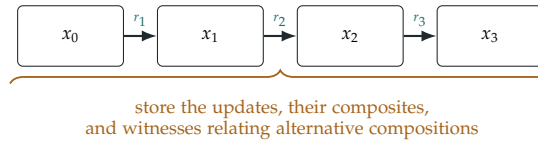


Figure 1.6: Homotopy coherence prevents persistence from collapsing to a list of states. The abstract witness H records how two composites agree; the online trace retains that witness so later revisions can reuse and audit the route by which the current decision was assembled.

1.12 Homotopy limits, colimits, and derived Kan extensions

Ordinary limits and colimits are not generally invariant under object-wise weak equivalence of diagrams. Homotopy limits and colimits correct them so that weakly equivalent input diagrams have equivalent

universal outputs. Riehl develops these constructions through enriched weighted (co)limits and derived functors [Riehl, 2014].

Definition 1.23 (Homotopy-coherent UDL). In a localized decision target $\mathcal{D}_\infty$, define

$$U^h(F) = \text{Ran}_K^h \text{Lan}_J^h F.$$

The superscript h records Kan extension in the homotopy-coherent semantics. In a compatible model presentation these are total derived Kan extensions.

Left Kan extension is a left adjoint and therefore preserves homotopy colimits. Right Kan extension is computed pointwise by homotopy limits. The asymmetry becomes decisive online: incremental candidate generation commutes with filtered revelation under broad conditions, whereas right-stage consistency requires an interchange between filtered colimits and the relevant limits.

1.13 Stable infinity-categories

Definition 1.24 (Pointed ∞ -category). An ∞ -category $\mathcal{C}$, in the quasi-categorical sense fixed in Definition 1.20, is *pointed* when it has a zero object: an object $0 \in \mathcal{C}$ that is both initial and terminal. Equivalently, for every object X , both mapping spaces

$$\text{Map}_{\mathcal{C}}(0, X) \quad \text{and} \quad \text{Map}_{\mathcal{C}}(X, 0)$$

are contractible.

Definition 1.25 (Stable ∞ -category). A pointed ∞ -category is *stable* when it has finite limits and finite colimits and a commutative square is a pullback exactly when it is a pushout.

Stability implies that finite limits and finite colimits agree. Consequently, in a presentable stable ∞ -category, filtered colimits commute with finite limits.

Example 1.26 (Why stability enters the first derived theorem). Suppose the pointwise right Kan extension uses only finite comma-category shapes. Its values are finite limits. If the online evidence object is the filtered homotopy colimit of its finite stages, stability permits that filtered colimit to pass through the right-Kan limits. This is the key interchange in the derived skeletal-continuity theorem.

Boundary

Stability is a sufficient hypothesis, not part of the definition of UODL. Decision spaces such as spaces, convex sets, or probability objects need not be stable. Later work should identify weaker target-specific exactness conditions rather than silently assuming stability everywhere.

1.14 Repair spaces

Let $p : E \rightarrow X$ be a simplicial family of decorated decisions. A partial decorated horn $e_\Lambda : \Lambda_i^n \rightarrow E$ over a proposed simplex $\sigma : \Delta^n \rightarrow X$ determines a derived lifting space

$$\mathrm{hofib}_{(e_\Lambda, \sigma)} \left[\mathrm{Map}(\Delta^n, E) \rightarrow \mathrm{Map}(\Lambda_i^n, E) \times_{\mathrm{Map}(\Lambda_i^n, X)} \mathrm{Map}(\Delta^n, X) \right].$$

This is the space of compatible fillers.

Repair-space shape	Structural reading	What a scalar can miss
$\emptyset$	no admissible completion	source of incompatibility
contractible	canonical up to coherent homotopy	internal witnesses
several components	qualitatively different repair classes	which alternatives remain possible
nontrivial higher homotopy	coherent families of transformations	higher relations among repairs

Table 1.2: A repair space carries more information than a success flag or penalty value.

1.15 Tangent categories and LINC

The algebraic probes used below come from synthetic differential geometry. Fix a ground field k , usually $\mathbb{R}$.

Definition 1.27 (Weil algebra and Weil probe). A *Weil algebra* over k is a finite-dimensional commutative local k -algebra W whose residue field is k . Equivalently, there is an augmentation $\epsilon_W : W \rightarrow k$ with nilpotent maximal ideal

$$\mathfrak{m}_W = \ker(\epsilon_W), \quad W/\mathfrak{m}_W \cong k, \quad \mathfrak{m}_W^r = 0$$

for some finite r . The infinitesimal object represented contravariantly by W is its *Weil probe* $\mathbb{D}_W$. For an object X , the corresponding *Weil prolongation* is denoted

$$T^W X := X^{\mathbb{D}_W}$$

whenever the indicated internal hom exists; more generally, T^W denotes the associated prolongation functor.

The basic example is the algebra of dual numbers

$$W_1 = k[\varepsilon]/(\varepsilon^2).$$

Its probe is written $\mathbb{D} = \mathbb{D}_{W_1}$. A map $\mathbb{D} \rightarrow X$ based at x represents a first-order tangent vector at x , and $T^{W_1}X$ is the ordinary tangent object TX . Larger Weil algebras encode several infinitesimal directions, higher jets, or prescribed relations among infinitesimals. The nilpotence is what truncates Taylor expansion to finite order.

The notation $\mathbb{D}_W$ is genuinely synthetic. For example, the only ordinary real number satisfying $d^2 = 0$ is zero, but the internal object $\mathbb{D} = \{d \mid d^2 = 0\}$ in a suitable smooth topos has enough generalized elements to detect derivatives [Kock, 2006]. Thus a Weil probe is not being treated as a small subset of classical points; it is a test object through which infinitesimal variation is observed.

A tangent category equips a category $\mathcal{C}$ with an endofunctor $T : \mathcal{C} \rightarrow \mathcal{C}$ and natural maps abstracting the zero section, projection, addition, lift, and canonical flip of ordinary tangent bundles [Cockett and Cruttwell, 2014]. These data satisfy coherence axioms ensuring that infinitesimal variation composes.

1.15.1 *Differential categories are not tangent categories with extra axioms*

Three related abstractions must be kept distinct. A *differential category* is an additive symmetric monoidal category carrying a coalgebra modality $!$ and a deriving transformation

$$d_A : !A \otimes A \longrightarrow !A$$

satisfying categorical analogues of the differential laws. Its native semantics is linear and resource-sensitive. A *Cartesian differential category* instead equips each map $f : A \rightarrow B$ with a derivative $D[f] : A \times A \rightarrow B$, coherently with products, composition, and linearity in the direction argument. A tangent category begins from the geometry of tangent bundles and their iteration; it need not supply a global operator $D[f]$ on every arrow.

These are connected constructions, not interchangeable definitions. The coKleisli category of a differential category is Cartesian differential. Under finite-biproduct and coreflexive-equalizer hypotheses, its coEilenberg–Moore category of $!$ -coalgebras is a tangent category [Cockett et al., 2020]. Conversely, suitable differential objects or bundles over a fixed base in a tangent category form a Cartesian differential

category under mild limit assumptions [Cockett and Cruttwell, 2014]. Thus differential programming and differential geometry can be related without being collapsed.

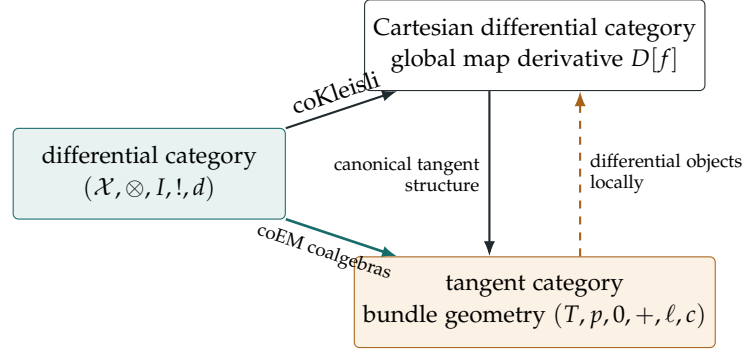


Figure 1.7: Differential and tangent structures answer different questions. The two routes from a differential category expose global differentiation of maps and geometric infinitesimal variation, respectively. Arrows summarize constructions and require the hypotheses stated in the text; they are not claims that the notions are equivalent.

For ORACLE this distinction creates two identification levels. A learner may identify the tangent geometry visible in the world without identifying a differential presentation that generates it. Alternatively, it may posit $(\mathcal{X}, !, d)$ as the latent theory and regard the observed tangent world as its category of coalgebras. The second target is stronger: it must explain not only which infinitesimal variations exist, but why they arise from a particular resource-sensitive differential mechanism.

UODL separates variation of evidence from variation of the decision mechanism. For an ordinary universal construction, the desired comparison is

$$T U(F) \longrightarrow U(TF).$$

In homotopical semantics the tangent functor must preserve weak equivalences or admit a derived functor T^h , giving

$$T^h U^h(F) \longrightarrow U^h(T^h F).$$

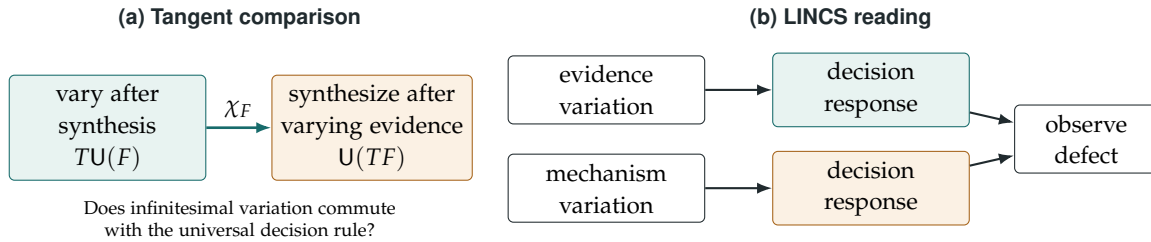


Figure 1.8: A tangent lift asks whether two routes to first-order decision change agree. LINC S keeps evidence variation distinct from variation of the action mechanism, constructs the comparison χ_F , and lets an observer measure its defect. The comparison is universal; causal meaning needs extra semantics.

Boundary

A discrete change of connected component in a repair space is not an infinitesimal motion. Tangent methods describe variation inside a local branch; component creation, deletion, or selection requires a registered repair.

LINCS adds a discipline around these comparisons: declare the typed structure, construct the relevant comparison map, observe its defect, propose a repair, and admit the repair through an independent criterion. A norm or loss may observe a defect, but does not define the structural defect by itself.

1.16 Infinitesimal intrinsic models

Witsenhausen’s intrinsic model isolates a decision problem by declaring its decision makers, their action spaces, and the information field available to each decision event [Witsenhausen, 1971, 1975]. A decision law must be measurable with respect to its declared information field. The causal question is consequently structural: can those information-local laws be assembled into a uniquely solvable closed loop? The Universal Decision Model recasts this architecture categorically [Mahadevan, 2021].

For tangent purposes, a bare sub- σ -algebra is the wrong object to differentiate. Present the information field of agent α by an observation map

$$y_\alpha : H \longrightarrow Y_\alpha, \quad \mathcal{I}_\alpha = \sigma(y_\alpha), \quad H = \Omega \times \prod_{\beta \in A} U_\beta,$$

and factor its decision law as

$$u_\alpha = \pi_\alpha(y_\alpha(\omega, u)).$$

This presentation exposes maps that can be tangent-lifted while retaining the measurability represented by $\mathcal{I}_\alpha$.

Definition 1.28 (Fixed-stratum infinitesimal intrinsic model). Let $\mathfrak{Int}_{\text{sm}}$ be an internal moduli object of smoothly presented intrinsic models in a setting admitting Weil prolongations. For a Weil algebra W , let $\mathbb{D}_W$ denote its infinitesimal probe. A *fixed-stratum infinitesimal intrinsic model* at $\mathcal{M}$ is a map

$$\widetilde{\mathcal{M}} : \mathbb{D}_W \longrightarrow \mathfrak{Int}_{\text{sm}}, \quad \widetilde{\mathcal{M}}(0) = \mathcal{M},$$

in which the agent set, action and observation types, information-dependency relation, and causal precedence relation remain fixed. The

observation maps, decision mechanisms, exogenous state, and objective may vary through the probe. Equivalently, it is a generalized element of the Weil prolongation $T^W \mathcal{I}nt_{\text{sm}}$ lying over $\mathcal{M}$.

1.16.1 Infinitesimal decision classes

Witsenhausen's classification separates two logically independent questions. The first asks whether the decision sites cooperate or compete; the second asks how information can flow among them. Tangent lifting should preserve this separation. A problem does not become a game merely because its information structure is nonclassical, and a team need not have classical information.

On the *objective axis*, an intrinsic model is a *team problem* when all decision sites share a single objective $J : \Omega \times \prod_{\alpha} U_{\alpha} \rightarrow \mathbb{R}$. An *infinitesimal team problem* carries one tangent objective TJ , so an admissible variation is evaluated by the same first variation $\dot{J}$ at every site. A *game problem* instead partitions the sites among players $p \in P$ and supplies player-indexed objectives J_p . Its infinitesimal variant carries the family $(TJ_p)_{p \in P}$. Thus the tangent of a team optimality condition is a common first-order optimality condition, whereas the tangent of a game solution concept is a coupled system of playerwise conditions. In particular, linearizing a Nash equilibrium is not the same operation as differentiating a common team optimum.

On the *information axis*, write $\beta \rightsquigarrow \alpha$ when the action at β is permitted to affect the observation at α . The following fixed-stratum variants mirror the usual intrinsic-model classes [Witsenhausen, 1971, 1975].

Infinitesimal static. Every y_{α} factors through the projection $H \rightarrow \Omega$.

No decision can change any decision site's information, and hence $D_u y_{\alpha} = 0$.

Infinitesimal sequential. There is an ordering $(\alpha_1, \dots, \alpha_N)$ such that y_{α_k} factors through $\Omega \times \prod_{j < k} U_{\alpha_j}$. Its tangent equations can therefore be solved in that order.

Infinitesimal classical. The model is sequential and its information fields are nested along a witnessing order: $\mathcal{I}_{\alpha_k} \subseteq \mathcal{I}_{\alpha_{k+1}}$. In a smooth presentation we record this nesting by comparison maps $r_k : Y_{\alpha_{k+1}} \rightarrow Y_{\alpha_k}$ satisfying $y_{\alpha_k} = r_k y_{\alpha_{k+1}}$. Tangent nesting is the identity $Ty_{\alpha_k} = Tr_k Ty_{\alpha_{k+1}}$.

Infinitesimal strictly classical. In addition, the next site observes the preceding action: the joint map $(y_{\alpha_k}, u_{\alpha_k})$ factors through $y_{\alpha_{k+1}}$. Its tangent lift carries both Ty_{α_k} and Tu_{α_k} forward.

Infinitesimal partially nested. This is Witsenhausen's quasiclassical case: the model is sequential and every structural influence $\beta \rightsquigarrow \alpha$ is

accompanied by information inclusion $\mathcal{I}_\beta \subseteq \mathcal{I}_\alpha$. Smooth comparison maps present these inclusions and their tangent lifts preserve them. In the strict version, α also observes the action u_β .

Infinitesimal nonclassical. The model is sequential but has a signaling edge $\beta \rightsquigarrow \alpha$ for which the required upstream information is not available at α . The tangent system retains this information defect: a perturbation can propagate through u_β , although the downstream decision mechanism cannot condition on all information that generated it.

Infinitesimal nonsequential. No fixed global order witnesses causal execution. A tangent response, when one exists, must be obtained from the coupled equation rather than from causal forward substitution.

These names describe structural strata, not properties of a single Jacobian. For example, an allowed influence edge remains part of the declaration even when its derivative happens to vanish at one base point. Conversely, a first-order perturbation may vary a permitted channel but may not silently add a new one.

Crossword solving will provide a running instance of these distinctions in section 9.15. Each clue is a decision site, crossing letters are information channels, and all sites share the objective of completing one grid. Perturbing the confidence assigned to candidate words stays within a fixed information stratum. Adding a candidate answer, fixing a new crossing, or changing the clue-incidence graph changes the declared combinatorial structure and therefore calls for repair rather than an ordinary tangent update.

Proposition 1.29 (Stratified invariance of intrinsic type). *Let $\widetilde{\mathcal{M}} : \mathbb{D}_W \rightarrow \text{Int}_{\text{sm}}$ be fixed-stratum. If $\mathcal{M}$ is a team, game, static, sequential, classical, strictly classical, quasiclassical, strictly quasiclassical, nonclassical, or nonsequential intrinsic model, then its tangent lift has the corresponding infinitesimal type. Changing membership in one of these structural classes requires a change of stratum rather than an ordinary tangent vector.*

Proof. The fixed-stratum declaration preserves the player partition, the number and indexing of objectives, the influence and precedence relations, and the factorization maps presenting information inclusion. Weil prolongation is functorial: it preserves identities and compositions. It therefore carries each recorded factorization $y_\beta = r_{\beta\alpha}y_\alpha$ to $Ty_\beta = Tr_{\beta\alpha}Ty_\alpha$, while retaining the same witnessing order and objective indexing. Each defining structural predicate is thus preserved. Altering one of those witnesses changes the declaration itself. $\square$

The taxonomy sharpens the solvability calculation below. In the static case $L_{\mathcal{M}} = 0$, so local variations do not circulate through other

decisions and $\dot{u} = b_{\text{ev}} + b_{\text{act}}$. In the sequential case, $L_{\mathcal{M}}$ is nilpotent after reordering. Classical and partially nested systems add information-preservation constraints to that triangular propagation. A nonclassical system may still be triangular and uniquely solvable, but its variations expose signaling without the corresponding inheritance of upstream information—the infinitesimal trace of the nonclassical information pattern.

For a first-order probe, assemble the intrinsic equations into

$$u = \Phi_{\mathcal{M}}(\omega, u), \quad \Phi_{\mathcal{M},\alpha}(\omega, u) = \pi_{\alpha}(y_{\alpha}(\omega, u)).$$

At a base solution, differentiation gives

$$(I - L_{\mathcal{M}})\dot{u} = b_{\text{ev}} + b_{\text{act}}, \quad L_{\mathcal{M}} = D_u \Phi_{\mathcal{M}}. \quad (1.1)$$

Here b_{ev} collects variation of exogenous evidence and the observation maps, whereas b_{act} collects variation of the decision mechanisms. Thus the intrinsic model gives a concrete realization of the LINC separation between varying what an agent knows and varying how the agent acts on what it knows.

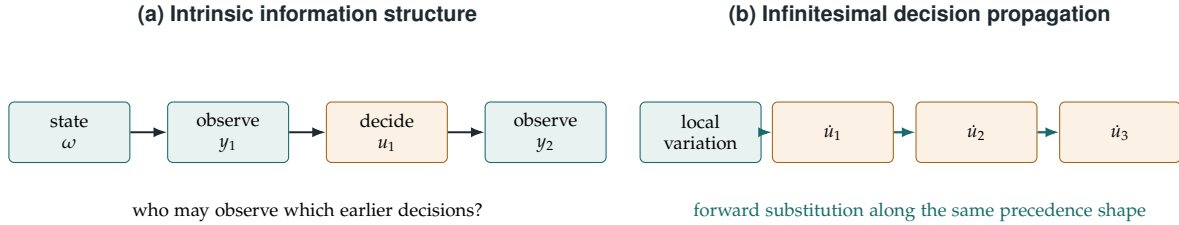


Figure 1.9: An infinitesimal intrinsic model preserves the information architecture while propagating first-order variation through it. Observation maps generate information fields; their tangent lifts transmit evidence variation, while tangent decision laws transmit mechanism variation.

Theorem 1.30 (Acyclic infinitesimal intrinsic solvability). *Let $A = \{\alpha_1, \dots, \alpha_N\}$ be finite. Suppose the exogenous-state, action, and observation spaces are finite-dimensional smooth spaces near a closed-loop solution, and every y_{α} and π_{α} is continuously differentiable there. Assume the fixed information-dependency relation admits the displayed ordering, with y_{α_k} depending only on ω and decisions u_{α_j} for $j < k$. Then $L_{\mathcal{M}}$ is strictly block lower triangular, so*

$$L_{\mathcal{M}}^N = 0, \quad (I - L_{\mathcal{M}})^{-1} = I + L_{\mathcal{M}} + \dots + L_{\mathcal{M}}^{N-1}.$$

Consequently, every admissible pair $(b_{\text{ev}}, b_{\text{act}})$ determines a unique infinitesimal closed-loop response $\dot{u}$, computable in causal order by forward substitution.

Proof. If $j \geq k$, the stated information restriction makes the block derivative $D_{u_{\alpha_j}} \Phi_{\mathcal{M},\alpha_k}$ zero. Hence $L_{\mathcal{M}}$ is strictly block lower triangular in the chosen order. Every such N -block matrix is nilpotent of degree at most N , and the finite geometric sum displayed above is a two-sided inverse of $I - L_{\mathcal{M}}$. Applying that inverse to equation (1.1) gives existence and

uniqueness. Its triangular form is exactly forward substitution through the decision events. $\square$

This theorem is the infinitesimal shadow of causal implementability. In a cyclic intrinsic model, the corresponding local criterion is invertibility of $I - L_{\mathcal{M}}$. Failure of invertibility records an infinitesimal solvability obstruction: a perturbation may be underdetermined, inconsistent, or poised at a bifurcation.

To connect the construction to UDL, let $F_{\mathcal{M}}$ encode the local laws compatible with the declared information fields, and let the right Kan stage encode their closed-loop compatibility. The canonical map

$$\chi_{\mathcal{M}} : TU(F_{\mathcal{M}}) \longrightarrow U(TF_{\mathcal{M}})$$

then compares solving before linearization with assembling the tangent-local laws. Under the fixed-shape preservation hypotheses for tangent Kan extension, the two routes agree. Witsenhausen's information structure supplies the causal interpretation that categorical universality alone does not provide.

The nerve supplies the temporal reading. A causal ordering of decision events determines a simplex in the decision nerve; its faces are shorter executable prefixes. The infinitesimal intrinsic model decorates that same simplex with tangent evidence and action maps. Coherent tangent transport therefore asks that these decorations restrict and compose along every face.

Boundary

An infinitesimal intrinsic model does not add or delete an information edge, refine or coarsen an information field, split a decision maker, or change a causal precedence relation. These operations move between structural strata and require a registered LINC repair. Likewise, singularity of $I - L_{\mathcal{M}}$ diagnoses local solvability failure; by itself it does not identify an intervention or a causal effect.

1.17 Reading the derived skeletal-continuity theorem

The later finite-consistency theorem compares incremental and all-at-once homotopy-coherent UDL. Its comparison has the form

$$\gamma^h : \operatorname{hocolim}_t \operatorname{Ran}_K^h \operatorname{Lan}_J^h \tilde{F}_t \longrightarrow \operatorname{Ran}_K^h \operatorname{Lan}_J^h \operatorname{hocolim}_t \tilde{F}_t.$$

Each symbol now has a declared role:

1. $\tilde{F}_t$ is the evidence available at finite stage t , extended to a common diagram domain.

Online construction versus hindsight construction

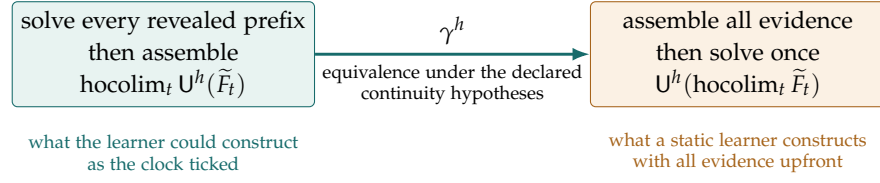


Figure 1.10: The online–offline comparison is a map between two orders of construction, not yet a numerical regret. When γ^h is an equivalence, incremental universal construction recovers the all-at-once semantic object. An observer is still required to turn any remaining defect into a scalar.

2. The homotopy colimit retains the progressively assembled evidence up to semantic weak equivalence.
3. Lan_J^h generates candidates and preserves the filtered homotopy colimit because it is a left adjoint.
4. Ran_K^h enforces consistency through pointwise homotopy limits.
5. Finite consistency shapes and a stable target allow the filtered colimit to commute with those finite limits.
6. Under these hypotheses, γ^h is an equivalence.

Admission contract

A claimed derived online-to-offline theorem must identify the weak equivalences, the localization, the common diagram domain, the existence of both derived Kan extensions, the pointwise right-Kan limit shapes, and the exact colimit–limit interchange being used. Without these declarations, the superscript h is only suggestive notation.

1.18 What this chapter does not assume

This tutorial intentionally leaves several choices open.

1. It does not require every decision target to be a model category.
2. It does not require every decision history to be reversible.
3. It does not assume that a learned theory has a unique presentation.
4. It does not identify homotopy equivalence with equal numerical loss.
5. It does not infer causal meaning from categorical universality.
6. It does not assume that tangent lifting preserves convergence.
7. It does not assume that every online update is an infinitesimal repair.

These boundaries are part of the framework. The chapters that follow state which special structures are used in each result and which obligations remain open.

Further reading

For an economical introduction to ordinary categories, functors, natural transformations, universal properties, and Kan extensions, see [Riehl \[2016\]](#); the classical reference is [Mac Lane \[1971\]](#). [Lawvere \[1963\]](#) introduced functorial semantics for algebraic theories. Ehresmann’s sketches and their model semantics are developed in [Ehresmann \[1968\]](#), [Barr and Wells \[1999\]](#); PROPs as one-sorted symmetric monoidal theories and their composition are treated by [Lack \[2004\]](#). [Richter \[2020\]](#) develops category theory and then crosses systematically to simplicial objects, nerves, classifying spaces, and homotopy-theoretic applications. This makes it a particularly direct bridge from the first half of this chapter to its simplicial and homotopical half.

For the basic theory of ∞ -categories, see [Riehl and Verity \[2022\]](#); its Section 1.1 supplies the inner-horn definition used here. Riehl develops the enriched and model-categorical machinery behind derived Kan extensions and homotopy-coherent UDL [[Riehl, 2014](#)]. Enriched universal constructions are developed by [Kelly \[1982\]](#). The tangent-categorical and synthetic differential background enters through [Cockett and Crutwell \[2014\]](#) and [Kock \[2006\]](#). The passage from differential categories through coKleisli and coEilenberg–Moore constructions to Cartesian differential and tangent categories is developed by [Cockett et al. \[2020\]](#). Witsenhausen’s information-field formulation and intrinsic model are developed in [Witsenhausen \[1971, 1975\]](#), with a categorical generalization in [Mahadevan \[2021\]](#). These references develop the mathematical machinery; the interpretation of that machinery as persistent online decision semantics is specific to UODL.

2

The ORACLE Program

LEARNING A COMPOSITIONAL WORLD FROM INTERACTION

ORACLE BEGINS ONE LEVEL BEFORE DECISION MAKING. An Online Rational Agent for Categorical Learning interacts with an unknown categorical environment and constructs a predictive model of its compositional structure. It knows the doctrine of categories, functors, natural transformations, and universal properties, but this is not its only prior. Chapter 0 suggested a stronger developmental hypothesis: the learner begins with a small categorical grammar of objects, magnitude, space, shape, agency, and social relation. It does not know which particular world realizes that grammar, or how its parts are composed in the world it inhabits.

This is a structural analogue of system identification. A predictive-state representation need not supply an optimal policy; it supplies a sufficient model for predicting the consequences of tests. Likewise, ORACLE first asks which compositions, equations, factorizations, extensions, and universal constructions can be predicted from an expanding interaction history. Action selection is a subsequent specialization.

Design principle

The primitive learning problem is not to optimize inside a fixed category. It is to identify, up to the distinctions exposed by admissible probes, the core-structured category in which later prediction and decision problems are posed.

2.1 *The categorical prior*

2.1.1 *From bare categories to structured worlds*

Choose Grothendieck universes $\mathcal{U} \in \mathcal{V}$. The hypothesis space cannot literally be the category of all categories: size must first be controlled. The 2-category $\mathbf{Cat}_{\mathcal{U}}$ of $\mathcal{U}$ -small categories is an object only at the larger

universe level $\mathcal{V}$.

Size discipline alone gives a prior that is far too weak. It says that the world has objects and composable arrows, but says nothing about which distinctions make fragmentary experience learnable. We therefore posit a small categorical sketch

$$\mathbb{T}_{\text{core}},$$

whose sorts, arrows, specified diagrams, and designated cones express the learner's initial organizing capacities. A realization in a candidate world $\mathcal{C}$ is a structure-preserving interpretation

$$M_{\mathcal{C}} : \mathbb{T}_{\text{core}} \longrightarrow \mathcal{C}.$$

Here “structure-preserving” is relative to the declarations in the sketch; it does not mean that every imaginable limit or colimit must be preserved. When a fragment includes an observer into an auxiliary category, this notation abbreviates the corresponding multi-sorted diagram of categories rather than forcing every observer target to lie inside $\mathcal{C}$.

Definition 2.1 (Core-structured world). A *core-structured world* is a pair $(\mathcal{C}, M_{\mathcal{C}})$, where $\mathcal{C}$ is a $\mathcal{U}$ -small category and $M_{\mathcal{C}} : \mathbb{T}_{\text{core}} \rightarrow \mathcal{C}$ interprets the declared core-knowledge sketch. A morphism of core-structured worlds is a functor equipped with coherent comparison data preserving the declared core structure. Structured natural transformations are its 2-morphisms.

Write $\mathbf{Cat}_{\mathcal{U}}^{\text{core}}$ for the resulting 2-category and

$$U : \mathbf{Cat}_{\mathcal{U}}^{\text{core}} \longrightarrow \mathbf{Cat}_{\mathcal{U}}$$

for the forgetful 2-functor. The ORACLE hypothesis space is now a sub-2-category

$$\mathfrak{H}_{\text{core}} \subseteq \mathbf{Cat}_{\mathcal{U}}^{\text{core}},$$

and the true environment is an unknown structured object

$$(\mathcal{C}_{\star}, M_{\star}) \in \mathfrak{H}_{\text{core}}.$$

The bare category $\mathcal{C}_{\star}$ is only its image under U . Thus the categorical prior is not a guess about the inventory of the world. It is a restriction on the kinds of organization through which an inventory can be discovered.

Guiding question

Which part of $\mathbb{T}_{\text{core}}$ must be fixed for identification to be possible, which part is merely observational, and which part may itself be revised through accommodation?

2.1.2 The generative-theory refinement

The fixed core sketch and the theory learned from experience must be distinguished. The former constrains what kinds of worlds are initially considered possible. The latter is a compact, revisable account of how the particular world is generated. A child playing with construction pieces may first identify configurations and transformations, then discover that joining, juxtaposition, and separation generate a much larger family of constructions subject to reusable equations.

Definition 2.2 (Generatively presented world). A *generatively presented world* consists of

$$(\mathcal{C}, \mathcal{S}, \mathcal{D}, \mathbb{A}, G), \quad \mathbb{A} = \text{Th}_{\mathcal{D}}(\mathcal{S}), \quad G : \mathbb{A} \longrightarrow \mathcal{C},$$

where $\mathcal{C}$ is a world category, $\mathcal{S}$ is a sketch of generators and relations, $\mathcal{D}$ is a preservation doctrine, $\mathbb{A}$ is the theory presented by that sketch, and G is a structure-preserving interpretation.

ORACLE may therefore seek either $\mathcal{C}_*$ alone or the stronger generative target

$$(\mathcal{C}_*, \mathcal{S}_*, \mathcal{D}_*, \mathbb{A}_*, G_*).$$

The stronger target explains indefinitely many composites using a finite presentation. It also creates new ambiguities: different signatures can present equivalent theories, and different theories can have equivalent model categories on the semantic worlds currently available to the learner. Literal recovery of a privileged generating list is therefore not the invariant objective. The learner must declare whether it seeks equivalence of completed theories, equivalence of their model semantics, or only factorization of a specified query doctrine.

This refinement sharpens Piagetian accommodation. Adding a newly encountered composite inside the old theory is assimilation. Adding a generator, imposing an equation, restricting an operation's domain, or changing from a Cartesian to a resource-sensitive monoidal doctrine changes the theory by which experience is organized. Chapter 6 makes this distinction algorithmic.

2.1.3 Running example: the collider declaration

The particle-physics example of Section 0.8 will run through the book in a fixed notation. Let $\mathbf{H}_{\text{phys}}$ be a declared category of admissible theory hypotheses, let $\mathcal{P}$ be a category of experimental probes, and let

$$\text{Resp}_H : \mathcal{P}^{\text{op}} \longrightarrow \mathbf{Stoch}$$

assign to each hypothesis H its predicted outcome law. A probe may encode beam species, energy, polarization, event selection, and detector context; its morphisms encode admitted changes of experimental

resolution or setup. We package the identification problem as

$$\mathfrak{C}_{\text{coll}} = (\mathbf{H}_{\text{phys}}, \mathbb{S}_{\text{phys}}, \mathfrak{D}_{\text{phys}}, \mathcal{P}, \text{Resp}, \mathcal{Q}),$$

where the sketch and doctrine delimit the permitted generators, relations, symmetries, and models, while $\mathcal{Q}$ declares which comparisons of responses count as questions.

For a probe subcategory $\mathcal{P}_0 \hookrightarrow \mathcal{P}$, define

$$H \simeq_{\mathcal{P}_0} H' \iff \text{Resp}_H|_{\mathcal{P}_0} \simeq \text{Resp}_{H'}|_{\mathcal{P}_0} \quad (2.1)$$

in the answer equivalence declared by $\mathcal{Q}$. Thus ORACLE does not initially recover “the true ontology.” It refines a quotient of hypotheses by the experiments performed and the distinctions the observer can resolve. The word *Yoneda* remains qualified here: collider responses form a restricted, noisy probe family, not necessarily the full representable embedding of a category.

Design principle

A scientific running example should expose all six parts of its declaration: hypothesis category, generating sketch, preservation doctrine, probe category, response semantics, and observational quotient. Without them, “learning a theory” conflates presentation, prediction, and ontology.

2.1.4 The tangent refinement

The unknown world may carry more than compositional structure. A tangent category equips $\mathcal{C}$ with an endofunctor T , projection and zero maps $p : T \Rightarrow 1$ and $0 : 1 \Rightarrow T$, fiberwise addition, a vertical lift $\ell : T \Rightarrow T^2$, and a canonical flip $c : T^2 \Rightarrow T^2$, subject to the tangent-category axioms reviewed in Chapter 1. We write such a world as

$$\mathcal{C}^{\text{tan}} = (\mathcal{C}, T, p, 0, +, \ell, c).$$

Its tangent objects do not merely append more observations. They declare which infinitesimal variations are meaningful, how they return to a base object, how they add, and how iterated variations cohere.

Let $\mathbf{TanCat}_{\mathcal{U}}$ denote the 2-category of $\mathcal{U}$ -small tangent categories, strong tangent functors, and tangent transformations. There is a forgetful 2-functor

$$U_{\text{tan}} : \mathbf{TanCat}_{\mathcal{U}} \longrightarrow \mathbf{Cat}_{\mathcal{U}}.$$

Tangent ORACLE restricts its hypothesis space to a suitable sub-2-category $\mathfrak{H}_{\text{tan}}$ and seeks an unknown object $\mathcal{C}_{\star}^{\text{tan}}$ up to the equivalence exposed jointly by base and tangent probes.

This is a genuine strengthening of the learning problem. Identifying $\mathcal{C}_\star$ does not in general identify a tangent structure over it: the fiber $U_{\text{tan}}^{-1}(\mathcal{C}_\star)$ may contain several inequivalent tangent doctrines. Conversely, a tangent prior can make learning easier by ruling out base-category hypotheses whose admissible variations do not fit the observed infinitesimal behavior.

Nor does identifying a tangent category necessarily identify a differential category behind it. One possible prior starts with a differential category $(\mathcal{X}, \otimes, I, !, d)$ and exposes a tangent world through the coEilenberg–Moore category of $!$ -coalgebras, provided the required limits exist [Cockett et al., 2020]. This is a *differential realization* of the tangent hypothesis, not extra structure that every tangent world must possess. Figure 1.15.1 therefore gives ORACLE two legitimate targets: the observable infinitesimal geometry, or a latent differential presentation that generates it.

Design principle

ORACLE discovers what composes. Tangent ORACLE also discovers how the compositional world can vary. Differential-realization ORACLE asks the stronger explanatory question of which resource-sensitive differential mechanism generates that geometry. None of these targets should be inferred from success on queries that forget the relevant structure.

2.1.5 The Spelke sketch

Spelke’s six systems do not all assert the same mathematical kind of structure [Spelke, 2022, 2023]. Treating each as one more object of $\mathcal{C}$ would erase their empirical content. The core sketch instead combines internal objects and relations with observers into auxiliary categories. The following declarations are a first abstraction, not a claim that cognitive development implements these particular mathematical models.

Physical objects: persistence, paths, and contact. The physical-object fragment declares a category $\mathcal{O}$ of object states, an abstract interval or path construction, endpoint maps

$$(d_0, d_1) : P(x) \longrightarrow x \times x,$$

and a class W of weak equivalences. An arrow in W says that two presentations preserve the same object for the probes at issue, despite a change of viewpoint, partial occlusion, or admissible deformation. Homotopies between paths express continuous alternatives that preserve this identity.

This relative category $(\mathcal{O}, W)$ is the minimal homotopical prior. A Quillen model structure may be imposed in a sufficiently complete realization to supply systematic factorizations and lifting arguments, but it should not be assumed without need [Riehl, 2014]. Weak equivalence by itself does not express solidity, boundedness, or contact causation. For these, the sketch also declares a region or boundary observer and a contact subobject

$$\text{Contact} \hookrightarrow \mathcal{O} \times \mathcal{O}.$$

Nonpenetration can be represented by forbidden fillers or exclusion conditions, while locality requires a physical-interaction arrow to factor through contact. Continuity, exclusion, and locality are therefore distinct parts of the prior.

Approximate number: magnitude before counting. A natural numbers object would supply exact recursion and counting, which is stronger than the approximate-number system. The core declaration should instead begin with an ordered commutative magnitude structure

$$(Q, \oplus, \leq)$$

regarded as a thin symmetric monoidal category $\mathbf{Q}$, together with a lax monoidal approximate-cardinality observer

$$\nu : \mathbf{FinSet} \longrightarrow \mathbf{Q}.$$

A family of resolution-dependent relations $\sim_\epsilon$ records that two magnitudes may be indistinguishable at one scale and distinguishable at another. Exact counting may arise later by constructing a natural numbers object $\mathbb{N}$ together with a comparison $\mathbb{N} \rightarrow Q$. In this reading, the passage from rough magnitude to counting is an accommodation of the numerical fragment, not something built silently into its starting semantics.

Spatial geometry: coherent changes of viewpoint. Let $\mathcal{V}$ be a groupoid of viewpoints and admissible movements, and let $\mathcal{G}$ be a category of geometric realizations. A spatial observer has the form

$$E : \mathcal{V} \longrightarrow \mathcal{G}.$$

Functoriality says that successive changes of viewpoint compose coherently. The surrounding room is not identified with one retinal image; it is the structure invariantly reconstructed from a diagram of such images. Local charts, embeddings, and gluing conditions may enrich this fragment when the presentation exposes metric or topological information.

Object form: shape separated from identity. The shape system is represented by a functor

$$S : \mathcal{O} \longrightarrow \mathbf{Shape}.$$

Its target might record geometric form, a homotopy type, a group-action orbit, or a persistent descriptor. The essential declaration is controlled invariance: viewpoint changes that preserve an object should induce specified equivalences of shape, while a genuine deformation need not. The functor separates the question “is it the same persisting object?” from “does it have the same form?” without making those questions independent.

Goal-directed agents: coalgebra plus intention. Autonomous evolution suggests an F -coalgebra

$$\gamma_a : X_a \longrightarrow FX_a$$

for each candidate agent. Coalgebra alone, however, describes behavior rather than intention. A goal-directed fragment must additionally declare goals and a comparison of paths toward them. Efficient behavior may be expressed by a cost-enriched hom-category, a preferred object in a goal-directed comma category, or another universal comparison appropriate to the realization. The prior says that some observed trajectories admit coherent goal-directed explanations; it does not identify prediction of motion with proof of intent.

Social beings: indexed relations, not merely labels. A poset can represent hierarchy, information refinement, trust, or affiliation, but cannot by itself express reciprocal communication and changing attention. A more general declaration is an indexed category

$$S : \mathcal{K}^{\text{op}} \longrightarrow \mathbf{Cat},$$

whose fiber over a context records identities and social relations. Its reindexing functors record contextual change. Groupoidal structure may track persistent identity, while poset enrichment records ordered evidence or trust. Temporal indexing can then distinguish a durable relationship from a momentary attentive or emotional state.

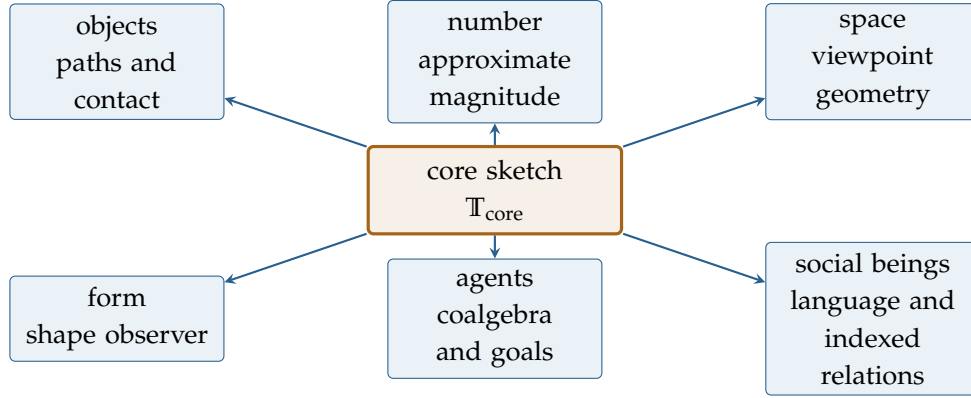


Figure 2.1: The Spelke sketch is a heterogeneous categorical grammar. Its six fragments constrain the possible organizations of evidence without naming the particular objects, spaces, agents, or social relations in the learned world. Language is acquired within the social fragment and makes the other systems communicable.

Language: a constrained coalgebra of communication. Language belongs inside the social fragment because it communicates reference, information, goals, and intentions between agents. A coalgebraic presentation writes a candidate language as

$$\gamma_L : L \longrightarrow F_{\text{lang}}L,$$

where a state records the residual linguistic and pragmatic possibilities after a fragment has been heard. When a final coalgebra exists, the unique behavior map $L \rightarrow \nu F_{\text{lang}}$ realizes those indefinitely extensible responses. Productive syntax retains a complementary algebraic aspect: expressions are assembled compositionally even as their observable behavior unfolds coinductively.

Moro’s study of *impossible languages* makes the restriction in this prior empirically substantive [Moro, 2023]. The class of formally specifiable grammars is larger than the class the human brain can acquire as language. ORACLE should therefore not search all F_{lang} -coalgebras, but a declared subcategory

$$\mathbf{Coalg}_{\text{human}}(F_{\text{lang}}) \hookrightarrow \mathbf{Coalg}(F_{\text{lang}}).$$

On this reading, universal grammar specifies the admissible signature, endofunctor, and invariants; interaction identifies the particular language coalgebra. A speech-sensitive observer first converts a structured acoustic stream into candidate linguistic evidence, while dialogue probes test whether the inferred behavior coherently changes shared social information. Language is thus not a seventh isolated core system. It is a learned coalgebraic interface through which the other core systems become communicable.

2.1.6 A layered prior and three levels of repair

Calling $\mathbb{T}_{\text{core}}$ a prior must not turn an empirical hypothesis into an infallible metaphysics. Magnets, remote communication, and virtual

objects illustrate why an initial expectation of contact or solidity may require reinterpretation. The declaration is therefore stratified:

1. *typing commitments* determine which evidence can initially count as an object, magnitude, place, form, agent, or social relation;
2. *observational invariants* specify distinctions that current probes are permitted to ignore;
3. *structural hypotheses* propose paths, contact conditions, equivariances, goals, and relations that further interaction may refute; and
4. *doctrinal commitments* determine which kinds of repair remain expressible inside the current core sketch.

This stratification yields three levels of Piagetian change. Assimilation changes the learned presentation while preserving its interpretation of the core sketch. Ordinary accommodation replaces $M_t : \mathbb{T}_{\text{core}} \rightarrow \mathcal{C}_t$ by a revised interpretation M_{t+1} inside the same doctrine. Foundational accommodation changes the doctrine itself:

$$(\mathbb{T}_{\text{core}}, M_t, \mathcal{C}_t) \longrightarrow (\mathbb{T}'_{\text{core}}, M_{t+1}, \mathcal{C}_{t+1}).$$

The last transition is the most consequential because it changes the hypothesis fibration in which subsequent learning occurs.

Design principle

Core knowledge restricts the space of discoverable worlds without specifying the discovered world. Its value must be measured by the identifications it makes possible and the coherent repairs it continues to permit.

2.1.7 Identification under a core doctrine

ORACLE is not asked to recover a literal code for $(\mathcal{C}_*, M_*)$. Its target is an equivalence class in $\mathfrak{H}_{\text{core}}$, because equivalences preserving the declared structure have the same core-categorical content. This is finer than bare categorical equivalence when the same category admits inequivalent interpretations of objecthood, agency, or social relation.

The restriction from $\mathbf{Cat}_{\mathcal{U}}$ to $\mathbf{Cat}_{\mathcal{U}}^{\text{core}}$ can make identification easier: extensions that violate the core doctrine are no longer admissible competing hypotheses. But the gain is not free. Every positive theorem must now state which fragment of the doctrine is known, how it is exposed by presentation, and whether convergence is required for the underlying category, the interpretation, or both. These are the categorical analogues of asking which universal grammar the learner possesses and which language remains to be identified.

2.2 Prediction before decision

Let $i : \mathcal{A} \rightarrow \mathcal{C}_*$ be a category of admissible probes. The interpretation M_* distinguishes a core family among them and types the answers as physical, numerical, spatial, shape, agentive, or social evidence. The response of an object x to those probes is the restricted nerve

$$N_{\mathcal{A}}(x) = \mathcal{C}_*(i(-), x) : \mathcal{A}^{\text{op}} \longrightarrow \mathbf{Set}.$$

Two objects are observationally equivalent when these presheaves are naturally isomorphic. For dense i , the restricted nerve is fully faithful. Otherwise it deliberately identifies distinctions unavailable to the learner.

The same idea operates one level higher. The Yoneda embedding of the ambient 2-category of structured hypotheses sends

$$(\mathcal{C}, M_{\mathcal{C}}) \longmapsto \mathbf{Cat}_{\mathcal{U}}^{\text{core}}(-, (\mathcal{C}, M_{\mathcal{C}})).$$

Thus structured category-shaped experiments probe a world through its core-preserving functors and transformations. Forgetting the core structure recovers ordinary category-shaped experiments $\text{Fun}(\mathcal{B}, \mathcal{C})$. ORACLE learns a structured category through its responses to such probes, just as ordinary Yoneda learns an object through all arrows into it.

Boundary 2.3. Yoneda supplies a complete extensional semantics, not an automatic decision procedure. Equality of presented morphisms, representability, or the existence of a universal object may remain undecidable. Effective ORACLE theorems must restrict both the hypothesis class and the query language.

2.3 From universal imitation to ORACLE

Universal imitation games distinguished static, dynamic, and evolutionary forms of imitation and contrasted inductive inference through initial algebras with coinductive inference through final coalgebras [Mahadevan, 2024]. Their central move was to replace the recovery of a finite description by agreement with an indefinitely extensible behavioral presentation.

ORACLE retains that coinductive insight but removes the requirement that the unknown world already be presented as one fixed universal coalgebra. A final coalgebra is now one especially important target object in $\mathfrak{H}_{\text{core}}$, equipped with the corresponding agent interpretation. More generally, the learner may have to discover the objects, arrows, equations, information shapes, and universal constructions that make the observed behavior compositional at all.

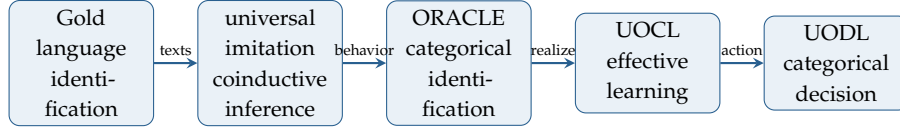


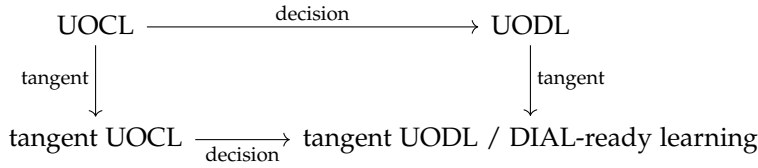
Figure 2.2: The conceptual lineage. ORACLE generalizes coinductive discovery of a universal coalgebra to identification of an unknown categorical world. UOCL makes the identification process algorithmic; UODL adds a decision declaration to the learned predictive structure.

2.4 The ORACLE–UOCL enrichment square

The distinction between modeling and acting can be expressed as

interaction $\longrightarrow$ categorical prediction $\longrightarrow$ categorical realization
 $\longrightarrow$ online categorical learning $\longrightarrow$ universal online decision.

ORACLE governs the semantic identification problem. UOCL supplies effective or explicitly realized hypothesis selection, active probing, and repair. From UOCL there are two independent enrichments. Tangent UOCL learns admissible infinitesimal variation. UODL adds information, action, consequence, and consistency maps to which left- and right-Kan constructions can be applied. Their intersection is tangent UODL:



UOCL realizes ORACLE; neither enrichment is implied by the other. DIAL-ready learning requires both, together with observers that can distinguish the infinitesimal defects on which repair depends. OCO, bandits, decentralized teams, safety, tangent repair, and lifelong transport are therefore target classes and query fragments on which ORACLE can be constrained sufficiently to support decision.

2.5 Claim discipline and the theorem ladder

The program separates six layers. A claim may be *presentational*, concerning what evidence is revealed; *predictive*, concerning answers to admissible probes; *categorical*, concerning typed composition; *developmental*, concerning preservation or revision of the core doctrine; *decision-theoretic*, concerning the selection of actions; or *causal*, when intervention semantics justify that stronger reading. No implication between adjacent layers is automatic.

The book develops seven theorem families: identification, realization, invariance, comparison, repair, stability, and transport. Identification is the ORACLE semantic problem; realization, effective probing, and online repair become UOCL questions. Each family also has a tangent

form once variation structure is part of the hypothesis, and a decision form once an action mechanism and observer have been declared.

Book part	Mathematical role	Governing question
I	Categorical identification	Which category is determined by an online presentation and a family of probes?
II	Categorical core knowledge	Which developmental restrictions have non-vacuous identification and composition consequences?
III	Online categorical learning	How are base and tangent hypotheses selected, probed, repaired, and stabilized?
IV	Decision specialization	When does a predictive category support UODL?
V	Global-clock test cases	How do OCO and bandits instantiate the specialization?
VI	Information beyond a clock	What survives decentralized, asynchronous, or endogenous information?
VII	Coherent repair	When must the categorical presentation itself change?
VIII	Lifelong structure	Which discoveries transport across worlds and tasks?

Further reading

Piaget’s theory of assimilation and accommodation motivates the layered repair semantics [Piaget, 1952]. Spelke’s account of core knowledge supplies the six empirical fragments abstracted by the Spelke sketch [Spelke, 2022, 2023]. Riehl develops the relative, model-categorical, and homotopical machinery used to distinguish weak equivalence from stronger physical structure [Riehl, 2014]. Moro uses experimentally motivated impossible languages to investigate the constraints separating possible human languages from arbitrary formal grammars [Moro, 2023]; this motivates the restricted category of humanly admissible language coalgebras. Gold introduced identification in the limit and the decisive dependence of learnability on the mode of presentation [Gold, 1967]. *Universal Imitation Games* develops the initial-algebra/final-coalgebra contrast that ORACLE generalizes [Mahadevan, 2024]. The categorical AI, LINCS, and DIAL books supply the decision, repair, and creative-learning infrastructure used after categorical identification [Mahadevan, 2026a,c,b].

3

Categorical Identification Under a Prior

PRESENTATIONS, CONVERGENCE, PROBES, GLUING, AND ACCOMMODATION

Hereafter, $\mathcal{C}_\star$ may abbreviate the structured target $(\mathcal{C}_\star, M_\star)$ when the interpretation is clear. A presentation may reveal facts about the underlying category, its core interpretation, or both.

3.1 Presentation protocols

Let $\mathcal{T}$ be a category of finite interaction transcripts. A morphism $T \rightarrow T'$ says that T' extends T without changing the evidence already registered. A presentation of a target $\mathcal{C}_\star$ is a directed path

$$T_0 \longrightarrow T_1 \longrightarrow T_2 \longrightarrow \cdots$$

that is sound for $\mathcal{C}_\star$ and fair for a declared collection of tests. Soundness says that no registered fact is false in the target; fairness says that every required test is eventually exposed.

Definition 3.1 (Categorical presentation modes). A *categorical text* presents positive witnesses: objects, arrows, composites, equations, cones, factorizations, or fillers that exist. A *categorical informant* may additionally certify negative facts, such as inequality of two paths or failure of a candidate universal property. An *active presentation* permits the learner to choose the next admissible probe.

Negative evidence is subtle. The absence of a morphism or universal object cannot usually be certified by waiting. It requires a finite refutation, a decision procedure for the presentation class, or a trusted informant.

The structured prior makes the transcript typed. Write

$$\text{Ev}(T) = \begin{pmatrix} \text{Ev}_{\text{obj}}(T), \text{Ev}_{\text{num}}(T), \text{Ev}_{\text{space}}(T), \\ \text{Ev}_{\text{form}}(T), \text{Ev}_{\text{agent}}(T), \text{Ev}_{\text{social}}(T) \end{pmatrix}$$

for the evidence sorted by the six fragments of $\mathbb{T}_{\text{core}}$. These components are not independent data streams. A single episode may expose an

object path, its shape under a change of viewpoint, an agent’s goal, and a communicated description of the same event. The transcript must therefore also register *cross-fragment warrants*: typed arrows and commuting diagrams saying which observations are claimed to concern the same object, place, agent, or interaction.

Definition 3.2 (Core-typed transcript). A *core-typed transcript* is a finite evidence object T equipped with a partial interpretation

$$e_T : \mathbb{T}_T \longrightarrow \mathcal{C}_T, \quad \mathbb{T}_T \hookrightarrow \mathbb{T}_{\text{core}},$$

where $\mathbb{T}_T$ is the fragment of the core sketch exposed so far. An extension $T \rightarrow T'$ preserves the registered interpretation but may enlarge $\mathbb{T}_T$, add witnesses, or record a declared repair.

The qualifier “partial” matters. An infant need not solve objecthood, number, space, agency, and language simultaneously before any of them can guide learning. The exposed subsketch grows with interaction, while its overlaps record the first compositional links among the fragments.

3.2 The hypothesis fibration

Assign to each transcript the category of models consistent with it:

$$H : \mathcal{T}^{\text{op}} \longrightarrow \mathbf{Cat}, \quad T \longmapsto H(T).$$

An evidence extension induces a forgetting functor $H(T') \rightarrow H(T)$. The Grothendieck construction produces

$$p : \int_{\mathcal{T}} H \longrightarrow \mathcal{T},$$

whose fiber over T is the current category of admissible hypotheses. Unless explicitly stated otherwise, these hypotheses lie in $\mathfrak{H}_{\text{core}}$, and restriction preserves the fragment of the core doctrine declared fixed by the presentation protocol.

Definition 3.3 (ORACLE learner). An ORACLE learner is a section, strict or pseudofunctorial as appropriate, of the hypothesis fibration along the observed transcript path. It chooses $(\hat{\mathcal{C}}_t, \hat{M}_t) \in H(T_t)$ together with comparison maps that make successive repairs compatible with evidence restriction.

This strengthens the usual index-valued learner. ORACLE does not merely emit unrelated conjectures; it records how one categorical hypothesis was revised into the next.

There are now two distinct sources of motion in the fibration. An evidence extension changes the base transcript while keeping the doctrine fixed. An accommodation may change the doctrine itself. If **Doc**

is a category of admissible core sketches and sketch morphisms, the fuller base is the category $\mathcal{T}_{\text{Doc}}$ of pairs $(T, \mathbb{T})$. Its arrows include both ordinary evidence extensions and registered doctrine revisions. The corresponding indexed hypothesis category

$$\mathbf{H}_{\text{doc}} : \mathcal{T}_{\text{Doc}}^{\text{op}} \longrightarrow \mathbf{Cat}$$

separates “the current model was wrong” from “the current language of models was too restrictive.” Chapter 3 develops this distinction.

3.3 Universal query theories

Let $\mathcal{Q}$ be an effective class of admissible questions. A query may ask for an equation between composites, a factorization, a limit or colimit, a lifting object, a Kan extension, or representability of a presheaf. Define

$$\text{Th}_{\mathcal{Q}}(\mathcal{C}) = \{(q, \text{Ans}_{\mathcal{C}}(q)) : q \in \mathcal{Q}\}.$$

Write $\mathcal{C} \simeq_{\mathcal{Q}} \mathcal{D}$ when their $\mathcal{Q}$ -theories agree. This is the operational target of learning.

Definition 3.4 (Categorical identification in the limit). An ORACLE learner identifies $\mathcal{C}_{\star}$ from a presentation class Π relative to $\mathcal{Q}$ if, for every presentation $(T_t) \in \Pi(\mathcal{C}_{\star})$, there is an N such that

$$t \geq N \implies \hat{\mathcal{C}}_t \simeq_{\mathcal{Q}} \mathcal{C}_{\star}.$$

If the hypotheses may continue to change while this condition holds, the identification is *behaviorally correct*. If the hypotheses stabilize coherently to one equivalence class, it is *explanatory*.

A weaker pointwise criterion permits the stabilization time to depend on the query. The distinction between one global N and the family (N_q) is the categorical analogue of uniform versus pointwise predictive convergence.

The query language itself is typed by the prior. Let $\mathcal{Q}_i$ denote questions internal to one core fragment and let $\mathcal{Q}_{ij}$ contain questions about its overlap with another. For example, an object query may ask whether two appearances represent the same persisting object, while an object–agent query asks whether an observed path is compatible with a declared goal. A language query may ask whether a continuation belongs to the behavior of an admissible grammar coalgebra; a language–social query asks whether an utterance coherently updates shared social information.

This gives a hierarchy

$$\mathcal{Q}_{\text{frag}} = \bigcup_i \mathcal{Q}_i \subseteq \mathcal{Q}_{\text{pair}} = \mathcal{Q}_{\text{frag}} \cup \bigcup_{i,j} \mathcal{Q}_{ij} \subseteq \mathcal{Q}_{\text{core}}.$$

Agreement on each isolated fragment need not imply agreement on $\mathcal{Q}_{\text{core}}$: two worlds can contain equivalent object, space, and language components yet attach words to objects or agents to goals differently. Categorical identification is therefore not just simultaneous identification of six marginal theories. It must identify their gluing.

3.4 *The finite-transcript obstruction*

Proposition 3.5 (Observational extension obstruction). *Suppose a finite transcript T admits models $\mathcal{C}, \mathcal{D} \in \mathbf{H}(T)$ with $\mathcal{C} \not\equiv_{\mathcal{Q}} \mathcal{D}$. No learner using only T can be guaranteed to answer every query in $\mathcal{Q}$ correctly.*

Proof. Choose a query on which the two models disagree. The learner receives the same transcript whether the target is $\mathcal{C}$ or $\mathcal{D}$, and therefore returns the same answer in both cases. That answer is wrong for at least one admissible target. $\square$

The proposition is elementary but fixes the burden of every positive ORACLE theorem. A presentation must eventually separate competing models. Positive texts often fail because one can add a hidden object, arrow, equation, or universal witness without contradicting any finite positive evidence.

The prior changes the obstruction but does not remove it. It can exclude an extension that violates solidity, approximate magnitude, viewpoint coherence, goal-directedness, social indexing, or linguistic admissibility. But whenever two *core-preserving* extensions survive the same transcript and disagree on an admissible query, the argument applies unchanged.

Corollary 3.6 (Prior-relative extension obstruction). *Fix a core doctrine $\mathbb{T}_{\text{core}}$. If a finite core-typed transcript has two models in $\mathbf{Cat}_{\mathcal{U}}^{\text{core}}$ that agree on all exposed fragment and overlap data but differ on $\mathcal{Q}_{\text{core}}$, then the doctrine and transcript do not yet identify the target.*

Thus core knowledge earns its keep by shrinking the fiber $\mathbf{H}(T)$, not by turning finite observation into omniscience. Its empirical and mathematical value is measured by which otherwise indistinguishable extensions it rules out.

3.5 *A constrained positive regime*

Exact finite identification becomes possible when the target category is finite, the learner has a known finite inventory of objects and arrows, and it may query domain, codomain, identity, composition, and equality.

Proposition 3.7 (Finite categorical identification). *Under the preceding assumptions, a fair active categorical informant identifies the target category after finitely many queries, up to an isomorphism preserving the declared inventory.*

Proof. There are finitely many structure-table entries: identity assignments, composable pairs, their composites, and equalities among the named arrows. Fairness reveals every entry after finite time. The completed tables determine a category on the inventory, and soundness makes it isomorphic to the target. $\square$

This proposition is intentionally modest. Removing the known inventory, allowing finite presentations with undecidable word problems, or restricting the learner to positive text reopens the identification problem.

The same proof admits a structured refinement. Suppose the finite inventory also names the realization of every sort, arrow, diagram, and designated cone in a finite core sketch. Querying the structure tables and the interpretation tables identifies not only $\mathcal{C}_\star$ but the pair $(\mathcal{C}_\star, M_\star)$, up to a core-preserving isomorphism. The extra conclusion requires the extra tables: identifying the ambient category alone does not identify which objects realize agents, magnitudes, or social contexts.

3.6 Coinductive realization

A predictive theory need not terminate in a finite syntax. In the coalgebraic case, finite observations are unfoldings of an indefinitely continuing behavior, and behavioral equivalence is the natural success criterion. Universal imitation games use precisely this final-coalgebra perspective [Mahadevan, 2024].

ORACLE generalizes the move: the realized target may be a universal coalgebra, a category presented by generators and relations, a homotopy-coherent category, or an indexed family of local categories. Coinduction is therefore not replaced. It becomes one realization doctrine inside categorical identification.

3.7 When prediction becomes decision

Before prediction becomes decision, the semantic ORACLE section must be realized by an online learner. Part III supplies UOCL state, probe, selection, and repair mechanisms and studies their persistent stabilization. To pass from UOCL to UODL, enrich a learned hypothesis with a decision declaration: an information category, an action mechanism, consistency maps, and an observer. Part IV studies which constructions remain valid as information arrives online.

Further reading

Gold’s language-learning paradigm supplies the presentation-relative notion of identification used here [Gold, 1967]. Universal imitation games motivate the coinductive behavioral criterion [Mahadevan, 2024]. The use of the Grothendieck construction separates the intrinsic family of hypotheses from any one external enumeration of them; Chapter 1 reviews the relevant categorical machinery.

3.8 Identification under a categorical prior

Chapter 2 replaced the bare prior “the world is a category” by a heterogeneous core sketch. Chapter 3 then made presentations and convergence precise. We now combine those two ideas. The central question is not merely whether evidence identifies a category, but whether fragmentary observers identify a *realization of the prior and the way its fragments compose*.

3.9 The fragment cover of the core sketch

Let I index the six core fragments and choose subsketches

$$\{\mathbb{T}_i \hookrightarrow \mathbb{T}_{\text{core}}\}_{i \in I}$$

whose union generates $\mathbb{T}_{\text{core}}$. Pairwise and higher intersections $\mathbb{T}_{i_0 \dots i_k}$ encode the warrants joining the fragments. For a core-structured world $(\mathcal{C}, M)$, restriction gives local realizations

$$M_i = M|_{\mathbb{T}_i} \quad \text{and} \quad M_{i_0 \dots i_k} = M|_{\mathbb{T}_{i_0 \dots i_k}}.$$

The resulting Čech-shaped diagram is the learner’s decomposition of the world into initially probeable systems.

Definition 3.8 (Fragmentwise observational equivalence). Two core-structured worlds are *fragmentwise equivalent* for a family of observers $(P_i)_{i \in I}$ when every $P_i M_i$ is equivalent and the equivalences agree on all exposed overlaps. They are *core-equivalent* when the induced structured models of $\mathbb{T}_{\text{core}}$ are equivalent.

The distinction isolates a fundamental ORACLE risk. Local success can coexist with global error. A learner might track objects, infer goals, and segment utterances correctly while incorrectly gluing a name to an object or an utterance to an intention.

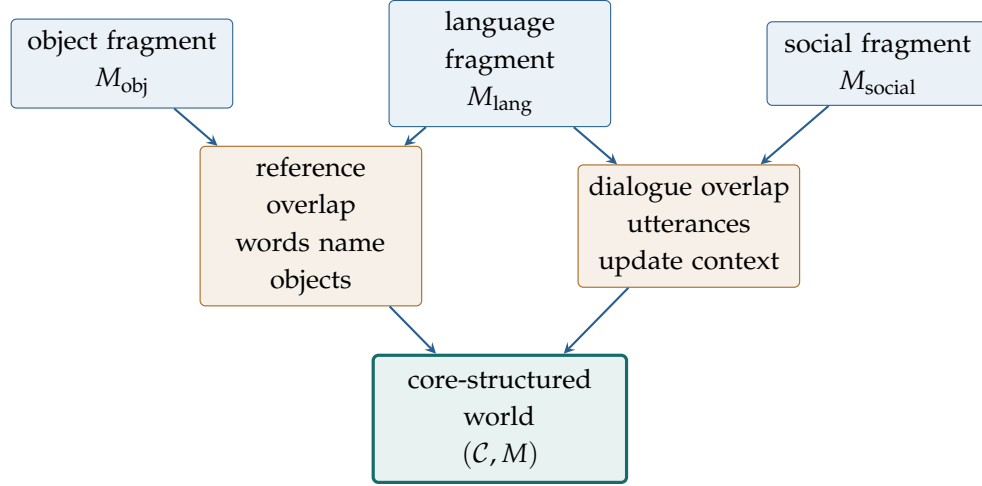


Figure 3.1: Fragmentwise identification is only the first layer. Overlap warrants determine how locally learned systems are glued into one core-structured world. Effective descent makes compatible local data both globally realizable and unique up to coherent equivalence.

3.10 Local identification does not automatically glue

Let $\text{Mod}_{\mathcal{C}}(\mathbb{T})$ be the category of realizations of a sketch $\mathbb{T}$ in $\mathcal{C}$. Restriction along the fragment cover defines a comparison functor

$$\Delta : \text{Mod}_{\mathcal{C}}(\mathbb{T}_{\text{core}}) \longrightarrow \lim_{[k] \in \Delta} \prod_{i_0, \dots, i_k} \text{Mod}_{\mathcal{C}}(\mathbb{T}_{i_0 \dots i_k}). \quad (3.1)$$

The right side is the category of compatible local realizations and coherent overlap data.

Definition 3.9 (Effective core descent). The fragment cover has *effective core descent* in $\mathcal{C}$ when Δ is an equivalence. It has *separated core descent* when Δ is fully faithful.

Effective descent says that compatible fragments can be glued and that the gluing is unique up to coherent equivalence. Separated descent gives uniqueness when a gluing exists, but not existence. Neither condition follows merely from calling the six systems a cover; it is a substantive property of the chosen sketch semantics.

Theorem 3.10 (Fragment-to-core identification). *Fix a target $(\mathcal{C}_*, M_*)$, a fragment cover with effective core descent, and fair presentations of every fragment and every finite overlap. Suppose the learner identifies each local realization and its overlap comparisons coherently in the limit. Then it identifies M_* in the limit up to core-equivalence.*

Proof. After local and overlap stabilization, the learner determines an object of the limit category on the right of equation (3.1). Effective core descent supplies a global realization, and full faithfulness makes its equivalence class unique. Hence every query invariant under core-equivalence eventually receives the target answer. $\square$

The theorem is deliberately conditional. It exposes three independent obligations: local learnability, learnability of the overlaps, and a descent theorem for the declared core semantics. The second is easily overlooked and is often the scientifically interesting part.

3.11 *Observers must be jointly separating*

Core systems become evidence through observers. Combine the fragment observers into

$$P = \prod_{i \in I} P_i : \text{Mod}_{\mathcal{C}}(\mathbb{T}_{\text{core}}) \longrightarrow \prod_{i \in I} \mathcal{E}_i.$$

Even effective descent cannot recover distinctions erased by P .

Definition 3.11 (Jointly separating core probes). The family (P_i) is *jointly separating* on a hypothesis class $\mathfrak{H}$ when $P(M) \simeq P(N)$, including all registered overlap comparisons, implies $M \simeq N$ for $M, N \in \mathfrak{H}$.

This is the structured analogue of observability in system identification. It is relative to a hypothesis class: acoustic evidence may separate the candidate languages in one declared family while failing completely on a larger family containing behaviorally indistinguishable grammars.

Proposition 3.12 (Unobservable core distinction). *If the core probes are not jointly separating on $\mathfrak{H}$, no ORACLE learner receiving only their outputs can identify every target in $\mathfrak{H}$ up to core-equivalence.*

Proof. Choose inequivalent $M, N \in \mathfrak{H}$ with equivalent observed outputs. They generate the same presentation to the learner, so any eventual answer distinguishing them fails on at least one target. $\square$

The role of active interaction is now precise: it enlarges the observer family by letting the learner choose probes intended to separate the remaining fiber of hypotheses.

3.12 *Possible and impossible language worlds*

Moro’s impossible languages provide a particularly sharp instance of a prior-relative hypothesis class [Moro, 2023]. A purely formal learner may range over all coalgebras of a language endofunctor, $\mathbf{Coalg}(F_{\text{lang}})$. A humanly structured ORACLE instead ranges over the declared subcategory

$$\mathbf{Coalg}_{\text{human}}(F_{\text{lang}}) \hookrightarrow \mathbf{Coalg}(F_{\text{lang}}).$$

The prior does not tell the learner which human language is present. It says which formally describable continuations are admissible candidates for human language at all.

This restriction has two effects. First, it can turn an otherwise nonidentifiable presentation into an identifiable one by deleting formal competitors that violate the universal-grammar doctrine. Second, it creates a diagnostic obligation: if sustained evidence can be modeled only outside the admissible subcategory, ORACLE must decide whether the evidence was mistyped, the current language model was wrong, or the doctrine itself needs revision. The slogan “the brain is a sieve” becomes categorical: the sieve is an admissibility subcategory together with the observers through which its objects are presented.

Language also demonstrates why overlaps matter. Acquisition is not only identification internal to the language coalgebra. Reference links language to objects; spatial expressions link it to geometry; numerals link it to magnitude; imperatives and explanations link it to agency; conversation links it to indexed social state. An impossible *world model* can arise from individually admissible fragments glued by inadmissible cross-fragment warrants.

3.13 *Assimilation, ordinary accommodation, and doctrinal repair*

For a transcript T and doctrine $\mathbb{T}$, write $H_{\mathbb{T}}(T)$ for the fiber of consistent hypotheses. New evidence $T \rightarrow T'$ induces

$$r_{T,T'} : H_{\mathbb{T}}(T') \longrightarrow H_{\mathbb{T}}(T).$$

This supports three different responses.

1. *Assimilation* chooses an object of $H_{\mathbb{T}}(T')$ whose restriction is coherently related to the current hypothesis.
2. *Ordinary accommodation* changes the chosen realization inside the same nonempty fiber, possibly revising cross-fragment warrants.
3. *Doctrinal accommodation* moves along an admissible sketch map $a : \mathbb{T} \rightarrow \mathbb{T}'$ because the old doctrine no longer supports an adequate realization of the evidence.

An empty fiber is decisive but not the only possible trigger for doctrinal repair. Under approximate or noisy semantics, a fiber may be formally nonempty while every object carries a persistent, observer-visible defect. Conversely, failure of the current learner to find a model does not prove that the fiber is empty.

Design principle

Revise the learned world before revising the language of possible worlds. Revise the doctrine only when a registered obstruction rules out adequate assimilation within the current hypothesis fiber.

This principle is conservative, not dogmatic. It prevents every surprising observation from rewriting the core prior, while retaining Piagetian accommodation when the prior itself blocks coherent learning.

3.14 *The ORACLE readiness criterion*

The passage to decision learning should occur only after the predictive structure required by a decision declaration has been identified to the resolution of its admissible queries.

Definition 3.13 (Decision-ready ORACLE model). An ORACLE hypothesis is *decision-ready* for a declaration $\mathbb{D}$ when:

1. the required core fragments and their overlaps have been identified up to the equivalence respected by $\mathbb{D}$;
2. the information, action, and consistency objects of $\mathbb{D}$ are realized in the learned world;
3. the observer family is separating for every distinction on which the action readout depends; and
4. unresolved presentation defects are either irrelevant to $\mathbb{D}$ or explicitly represented as repair obligations.

Decision readiness is task-relative and weaker than complete categorical identification. It marks the exact boundary between UOCL and UODL: a learner may act once the relevant predictive quotient is stable, even while unrelated parts of its compositional world remain open.

Further reading

Gold's distinction between positive texts and richer informants remains the prototype for presentation-relative learnability [Gold, 1967]. Spelke motivates the heterogeneous fragment cover and the problem of composition across core systems [Spelke, 2022, 2023]. Moro's investigation of impossible languages supplies the motivating case in which a biological prior restricts a much larger formal hypothesis class [Moro, 2023]. The language of descent, indexed categories, and homotopy-coherent gluing used here is reviewed in Chapter 1; later chapters turn failures of strict gluing into explicit repair spaces.

Part II

Categorical Core Knowledge

4

Core Knowledge and Piagetian Construction

CORE KNOWLEDGE, CATEGORICAL TRUTH, ASSIMILATION, AND ACCOMMODATION

THERE IS A DANGER IN TRANSLATING A PSYCHOLOGICAL THEORY INTO CATEGORY THEORY. Every collection of states and transformations can be placed in some category after the fact. That observation alone says nothing about infant cognition and supplies no evidence for Spelke's hypothesis. A categorical reconstruction becomes scientifically meaningful only when its structure excludes alternatives, predicts invariances, and explains how separately available capacities compose.

This chapter therefore asks a deliberately demanding question: in what sense could there be a *categorical truth* to the core-knowledge hypothesis? The answer cannot be that a category can be fitted to the data. It must be that a comparatively small categorical doctrine captures stable organization in the data and earns its restrictions through successful identification, composition, and repair.

4.1 Core knowledge as a doctrine of possible worlds

The categorical interpretation can now be stated sharply. In the terminology of Section 1.2, Spelke's hypothesis is not primarily a claim that an infant begins with a particular world model. It is a claim about the *doctrine that defines the category of world models the child is prepared to learn*. Core knowledge constrains the admissible objects, transformations, observers, invariances, and cross-system comparisons before the contents of one particular world have been identified.

On this reading, evolution supplies neither an empty hypothesis space nor a complete theory of the environment. It supplies structural commitments about what kinds of worlds are possible for the learner: objects persist and move continuously; magnitudes combine approximately; viewpoints transform spatially; shapes can be compared across appearances; some entities behave as goal-directed agents; and social

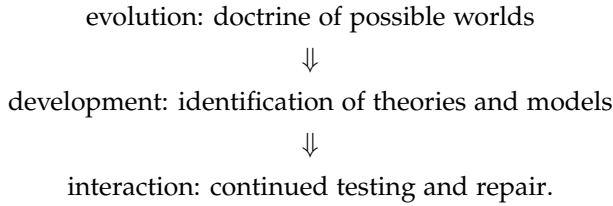
interaction includes communication through an admissible class of languages. Development must still discover which objects, places, agents, words, rules, and causal regularities populate the encountered world.

The six systems are heterogeneous, so “the core doctrine” should not be mistaken for one undifferentiated algebraic theory. Let $\mathfrak{D}_i$ denote the doctrine associated with the i th core system and let $\mathfrak{D}_{ij}$ specify an admitted interface between two systems. The structured hypothesis space has the schematic form

$$\mathbf{World}_{\text{core}} = \{(\mathcal{C}_i, \phi_{ij}) \mid \mathcal{C}_i \in \mathbf{Doc}_{\mathfrak{D}_i}, \phi_{ij} \text{ satisfies the declared compatibility laws}\}.$$

When the required 2-limits exist, this may be expressed as a pseudo-limit of the fragment doctrines over their interfaces. More generally it is an indexed family of doctrines equipped with gluing data. A child must therefore learn both modality-specific models and the comparison maps that make a seen object graspable, nameable, countable, located in space, and available to joint attention.

This yields a clean division of labor:



Assimilation updates a realization while preserving settled structure. Ordinary accommodation revises a theory, sketch, or realization within the inherited doctrine. Foundational accommodation becomes necessary only when a registered obstruction gives evidence that the doctrine of possible worlds itself is inadequate. Chapter 4 formalizes these three cases.

Design principle

Core knowledge is a prior over *kinds of compositional worlds*, not a catalogue of facts about one world. ORACLE begins with a doctrine, discovers theories and models within it, and revises that doctrine only when failures cannot be repaired inside its category of admissible worlds.

This interpretation is an empirical proposal rather than a consequence of category theory. It earns explanatory force only if the core doctrine rules out learnable alternatives, predicts the invariances infants actually preserve, and explains why the available presentation identifies structures that a materially weaker doctrine would leave ambiguous.

4.2 Three meanings of categorical truth

The phrase “categorical truth” has three progressively stronger readings.

1. *Representational adequacy*: the objects, transformations, and invariances posited by a core system admit a faithful categorical model.
2. *Explanatory adequacy*: universal properties and functorial constraints compress regularities that would otherwise require unrelated rules for different tasks or modalities.
3. *Developmental adequacy*: using the doctrine changes what can be identified from the presentations available to a learner, and its permitted repairs match the distinction between assimilation and accommodation.

The first reading is necessary but weak. The second says why category theory is useful. The third is the ORACLE claim: the prior contributes to learnability rather than merely redescribing a learned result.

Definition 4.1 (Categorical core-knowledge claim). A *categorical core-knowledge claim* is a tuple

$$(\mathbb{T}_i, \mathfrak{H}_i, P_i, Q_i, \mathcal{R}_i)$$

consisting of a fragment sketch, a hypothesis class, an evidence observer, a query language, and an admissible repair class. It is *empirically contentful* only if replacing $\mathbb{T}_i$ by a weaker doctrine changes at least one observable identification, invariance, composition, or repair prediction.

This formulation separates mathematical validity from empirical truth. A theorem may correctly describe models of $\mathbb{T}_i$ even if evidence later shows that infants do not employ the predicted distinction. Conversely, an observed competence does not select a unique categorical explanation without comparison to weaker or competing doctrines.

4.3 The five-obligation test

A candidate categorification of a core system should satisfy five obligations.

Restriction.

The doctrine defines a proper, non-vacuous subcategory of possible models.

Invariance.

Presentation changes that preserve the hypothesized competence induce declared equivalences rather than arbitrary retraining.

Composition.

Local capacities have typed interfaces and coherent overlaps with other core systems.

Identification.

The restricted hypothesis class removes ambiguities that cannot be removed from the same evidence under a weaker prior.

Repair.

Violations distinguish a mistaken realization from evidence that calls the doctrine itself into question.

These are joint obligations. A tiny hypothesis class may make identification easy by stipulation while misrepresenting the observed invariances. A very expressive category may fit every transcript while excluding nothing. A collection of accurate fragment models may still fail to explain how number, space, objects, agents, and language become one world.

Design principle

A categorical prior earns empirical meaning by the alternatives it rules out, the equivalences it preserves, and the cross-system predictions it forces. Mere encodability is not evidence for categorical core knowledge.

4.4 Prior gain as a relative identification statement

Let $j : \mathfrak{H}_{\text{core}} \hookrightarrow \mathfrak{H}_{\text{amb}}$ include the models satisfying a proposed core fragment into a weaker ambient hypothesis class. For an evidence transcript T , write $H_{\text{amb}}(T)$ and $H_{\text{core}}(T)$ for the corresponding consistent fibers.

Definition 4.2 (Query-effective prior). The core restriction is *effective* for a query q at T when all objects of $H_{\text{core}}(T)$ give the same answer to q , while two objects of $H_{\text{amb}}(T)$ give different answers. It is *presentation-stable* when this property is invariant under the declared equivalences of transcripts.

Proposition 4.3 (Identification gain from a core restriction). *If a core restriction is query-effective at T , then q is identified from T relative to $\mathfrak{H}_{\text{core}}$ but not relative to $\mathfrak{H}_{\text{amb}}$.*

Proof. Agreement throughout the core-consistent fiber determines the answer to q . In the ambient fiber, the two disagreeing hypotheses produce the same available transcript, so the observational extension obstruction of Chapter 3 applies. $\square$

The proposition is elementary, but it turns a broad developmental intuition into a comparative burden. For each proposed core fragment, one should exhibit presentations and queries for which the restriction matters. If no such contrast exists, the categorical prior is idle for the learning problem at hand.

4.5 Six systems, six categorical burdens

The declarations of Chapter 2 can now be read as research hypotheses rather than finished models. Table 4.5 records the categorical content and the corresponding burden for each system.

Core system	Candidate categorical content	What would make it non-vacuous?
Objects	Relative category of persisting states; paths, boundaries, contact, and exclusion	Identity judgments should be invariant under declared occlusions and viewpoint changes, but not under transformations that violate cohesion, continuity, or contact constraints.
Number	Ordered commutative magnitude object with resolution-dependent indistinguishability	Approximate comparisons should respect order and combination coherently across presentations while remaining distinct from exact natural-number recursion.
Space	Groupoid of viewpoints acting through a geometric observer	Route composition and reorientation should be predicted from transformations among views rather than memorized separately for each view.
Form	Shape functor separated from object identity	Shape equivalence should transport across identity-preserving viewpoint changes and fail in controlled ways under genuine deformation.
Agents	Behavioral coalgebras together with goal and efficiency comparisons	Self-propelled trajectories should support counterfactual predictions that motion-only coalgebras or object-path models do not determine.
Social beings and language	Indexed social categories linked to an admissible subcategory of language coalgebras	Communication should update shared context coherently, while formally possible but humanly impossible grammars remain outside the learnable class.

The table also exposes where the current abstractions are weakest. Objects and space already have well-developed notions of invariance, paths, and gluing. Approximate number requires a principled account of psychophysical resolution rather than a convenient enriched order. Agency requires more than behavioral coalgebra: intention and efficiency must yield predictions not already present in trajectory fitting. Social cognition cannot be reduced to a poset if reciprocal communication, attention, and belief change are essential.

Table 4.1: The six Spelke systems become categorical hypotheses only when their formal declarations impose discriminating observational and developmental burdens.

4.6 *Cross-system composition is the decisive test*

Spelke's systems are often described as distinct, but human knowledge is not a permanent disjoint union of six modules. Infants and children learn that a bounded object occupies a place, that several objects have a magnitude, that an agent moves an object, and that another person can use a word to direct attention toward it. The categorical proposal is strongest precisely here: typed composition and descent should explain how partial systems form a coherent world without erasing their local invariants.

Let $\mathbb{T}_i$ be the fragment sketches and $\mathbb{T}_{ij}$ their interfaces. The comparison

$$\text{Mod}(\mathbb{T}_{\text{core}}) \longrightarrow \lim \left(\prod_i \text{Mod}(\mathbb{T}_i) \rightrightarrows \prod_{i,j} \text{Mod}(\mathbb{T}_{ij}) \rightrightarrows \cdots \right)$$

is therefore not just technical packaging. Its failure modes distinguish three possibilities: a fragment was learned incorrectly; the overlap warrant was learned incorrectly; or the proposed core sketch does not admit the observed composition. Chapter 3 proved that effective descent turns local and overlap identification into global identification. The empirical question is whether the developmental systems actually exhibit the compatibility and uniqueness that such descent predicts.

4.7 *What would count against the hypothesis?*

A useful categorical theory must state its negative evidence. The proposed core doctrine would be weakened if:

1. the declared invariances fail to match stable infant discriminations;
2. a substantially weaker, untyped hypothesis class identifies the same queries from the same presentations;
3. cross-system judgments vary freely rather than satisfying coherent overlap conditions;
4. observed developmental repairs routinely cross boundaries that the doctrine declares fixed; or
5. multiple inequivalent core realizations remain observationally indistinguishable even under the interactions available to the learner.

No one failure refutes all categorical modeling. It identifies which claimed fragment, observer, or universal property lacks empirical support. This is why the layered prior of Chapter 2 is essential: typing commitments, observational invariants, structural hypotheses, and doctrinal commitments should not all be protected equally from revision.

4.8 Does existing machinery already solve ORACLE?

The strongest contemporary objection to this program is not that categorical structure is too abstract. It is that large Transformer systems already appear to have solved the problem. They converse fluently, construct programs, use tools, repair errors, and produce novel compositions over language, code, and images. Measurements of frontier agents show rapidly increasing competence on multi-step software tasks, even though success remains dependent on task duration, domain, scaffold, and evaluation protocol [METR, 2026, OpenAI, 2026b]. Any theory of compositional learning that begins by discounting this evidence has protected itself from its most important empirical fact.

The appropriate response is to make ORACLE *architecture-neutral*. A Transformer, a recurrent network, a symbolic engine, or a hybrid agent may implement an ORACLE learner. The issue is not whether its internal tensors literally contain objects and arrows. The issue is whether the learning system satisfies the identification, invariance, persistence, and repair obligations that the categorical theory declares.

There is direct evidence that this distinction is substantive rather than philosophical. Liu et al. [2023] use Krohn–Rhodes theory to show that shallow Transformers can exploit the algebraic structure of finite automata, compiling recurrent transitions into parallel shortcut computations. This is a genuine compositional achievement. It also sharpens the open question: efficiently executing a machine whose structure can be decomposed is not yet the same as discovering that decomposition from an open-ended presentation or preserving it through later revision. We return to this representation–identification gap in Section 5.7.1.

Definition 4.4 (Compositional performance). Let P be a finite presentation and $\mathcal{Q}_0$ a tested query family. A learner has *compositional performance on $(P, \mathcal{Q}_0)$* when it returns acceptable answers to the queries in $\mathcal{Q}_0$, including queries whose answers require composing fragments present in P .

This definition deliberately includes the achievements of frontier models. It does not require the learner to expose a symbolic derivation, nor does it infer from success what representation the learner uses.

Definition 4.5 (Compositional identification). A learner *compositionally identifies* a presented world, relative to a doctrine $\mathbb{T}$, observer family $\mathcal{O}$, equivalence class $\mathcal{E}$, and query language $\mathcal{Q}$, when its limiting hypothesis determines the $\mathcal{Q}$ -theory of the world up to $\mathcal{E}$, remains invariant under admissible re-presentations, and transports settled answers through compatible continuations and declared repairs.

Compositional performance is therefore local to a presentation and a test. Compositional identification is a claim about the world class

selected by an open-ended evidence process. Neither definition entails that the learner must recover a unique hidden implementation. Identification is only up to the equivalences and observers declared in advance.

Proposition 4.6 (Finite performance does not imply identification). *Suppose two admissible worlds $M, N \in H_T$ induce the same observer-visible answers on a finite presentation P and tested query family $\mathcal{Q}_0$, but disagree on some admissible continuation or query $q \in \mathcal{Q}$. Then success on $(P, \mathcal{Q}_0)$ cannot establish that a learner has identified M rather than N .*

Proof. The learner receives identical observer-visible evidence and is judged by identical answers in the two cases. Consequently every output distribution and every score determined only by $(P, \mathcal{Q}_0)$ is the same whether the underlying world is M or N . The distinguishing query q , or evidence from a continuation on which M and N separate, is absent. Finite test success therefore witnesses performance on the common observable fragment, not identification of either world. $\square$

The proposition is not a special criticism of Transformers. It is an information-theoretic warning about every learner, including a categorical one. It prevents the book from treating fluent output as proof of an identified categorical model, while equally preventing it from treating the lack of an explicit symbolic trace as proof that no such model has been acquired.

4.8.1 What would a Transformer have to demonstrate?

The comparison must hold training data and interaction fixed and vary the obligation being tested. A serious Transformer challenge should include at least the following contrasts.

1. *Unseen composition:* answer universal-property queries whose components appeared during learning but whose composite did not.
2. *Re-presentation:* preserve the answer after renaming, refactoring, changing viewpoint, or replacing a presentation by an equivalent one.
3. *Compatible continuation:* retain a settled query theory when new objects, arrows, or observations extend the world.
4. *Localized contradiction:* distinguish a faulty realization from a faulty overlap warrant and from evidence against the categorical doctrine.
5. *Underdetermination:* abstain, or return the residual hypothesis fiber, when the observations do not identify an answer.

6. *Transport*: reuse the discovered structure under a new task, observer, or modality without relearning the relevant compositions from scratch.

These tests differ from asking whether a model can solve a familiar crossword, write a function, or describe an object. They ask whether success persists under controlled changes to the presentation and evidence process. Ordinary benchmark scores may be poor witnesses for this purpose: recent audits have shown that apparently precise coding evaluations can contain underspecified prompts, overly restrictive tests, and incomplete behavioral coverage [OpenAI, 2026b].

The challenge should not be framed as a return to a rigid symbolic doctrine. Neural networks trained with an appropriate compositional curriculum can exhibit human-like systematic generalization [Lake and Baroni, 2023]; conversely, strong interpolation can coexist with failure on controlled systematic splits [Lake and Baroni, 2018]. This is evidence that architecture alone does not decide the issue. The presentation, training objective, observer, and admissible query transformations all matter.

4.8.2 *Beyond tokens: JEPA and predictive grounding*

A second objection strengthens the first. Perhaps token-predictive Transformers are only an early implementation, while joint-embedding predictive architectures already address the deeper complaint that learned composition is not grounded in a physical world. A JEPA predicts a target representation from a context representation rather than reconstructing every target pixel or token. I-JEPA applies this principle to spatial blocks of an image [Assran et al., 2023]; V-JEPA 2 combines large-scale video prediction with an action-conditioned latent predictor that can be used for model-predictive robot control [Assran et al., 2025]. This program is explicitly motivated as a route from self-supervised perception to world models, planning, and autonomous intelligence [LeCun, 2022].

The basic architecture has a direct categorical reading. Let $\mathcal{W}$ be a candidate world category, let $O : \mathcal{W} \rightarrow \mathcal{X}$ be a sensory observation functor, and let $E : \mathcal{X} \rightarrow \mathcal{Z}$ be a learned encoder. The predictor operates in the latent category $\mathcal{Z}$, comparing a predicted target $\hat{z}_T$ with the target-encoder value z_T .

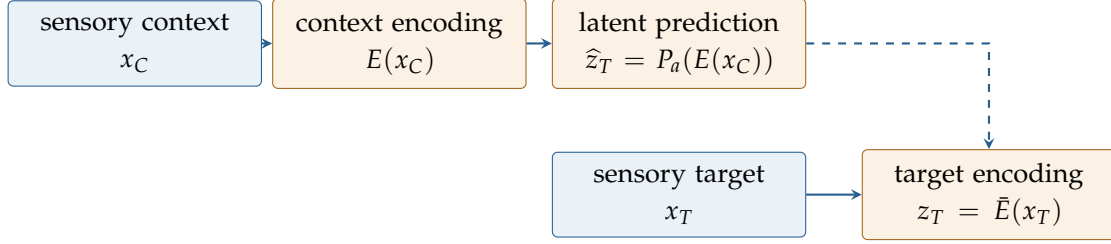


Figure 4.1: A JEPA as a latent predictive diagram. The learned comparison can discard sensory detail, but the training loss alone does not determine whether the retained quotient supports the queries, interventions, and invariances of the discovered world.

Prediction in $\mathcal{Z}$ is a powerful form of grounding only relative to what the encoder preserves. Suppressing unpredictable texture may expose object motion and physical regularity; suppressing a distinction needed for a later counterfactual or social query destroys relevant structure. The choice of latent representation is therefore a choice of observational quotient, not a neutral window onto the world.

Definition 4.7 (Query-sufficient grounded representation). For a declared query family $\mathcal{Q}$, the latent representation E is $\mathcal{Q}$ -sufficient when every query $q : \mathcal{W} \rightarrow \mathcal{Y}$ in $\mathcal{Q}$ factors, up to a declared equivalence, as

$$q \simeq \bar{q} \circ E \circ O$$

for some latent query $\bar{q} : \mathcal{Z} \rightarrow \mathcal{Y}$. It is *interventionally coherent* when each admissible world action $a : w \rightarrow w'$ has a latent predictor P_a for which the observation, encoding, and prediction square commutes up to the declared comparison.

This definition gives the JEPA proposal a demanding but fair ORACLE test. A low representation-prediction loss is evidence that some target information is predictable from context. It is not by itself evidence that all $\mathcal{Q}$ -relevant distinctions factor through the representation, that different viewpoints induce coherent maps, or that action-conditioned rollouts identify the correct intervention. Conversely, if these factorization and coherence conditions do hold, JEPA may be an especially natural computational substrate for an ORACLE learner.

The six Spelke systems sharpen the comparison. Object continuity and contact, viewpoint geometry, and persistent shape should become tests of latent invariance and action coherence. Approximate number requires preservation of magnitude structure rather than merely decodability by a probe. Goal-directed agency requires counterfactual sensitivity to actions. Social and linguistic structure require communicative compositions that video prediction alone need not identify. Thus “grounded” is not a single property: it is indexed by the core fragment, observer, action class, and query family.

JEPA also reveals why the Transformer-versus-world-model opposition is too coarse. Contemporary JEPAs use Transformer encoders and predictors, while multimodal systems can align latent video models

with language models. The scientific comparison concerns objectives, presentations, and preserved structure—not brand names for network components.

4.8.3 *Why a Kan extension is not yet an ORACLE learner*

There is an equally serious objection from within category theory. Mac Lane’s famous maxim that every concept is a Kan extension suggests that the universal construction might already contain the whole ORACLE program [Mac Lane, 1971]. Given an indexing functor $i : \mathcal{A} \rightarrow \mathcal{B}$, boundary data $F : \mathcal{A} \rightarrow \mathcal{C}$, and an ambient category $\mathcal{C}$, the left and right Kan extensions construct universal ways of extending F along i . What more could a universal categorical learner require?

The answer is contained in the word *given*. A Kan extension solves a conditional completion problem after the indexing shape, boundary data, codomain, admissible morphisms, and direction of universality have been declared. ORACLE begins without knowing which such declaration is warranted. It must discover or repair:

1. which observations denote the same object and which arrows compose;
2. which presentation shape records the relevant evidence;
3. which observer and query language determine empirical agreement;
4. which equivalences identify two realizations as the same world; and
5. whether new evidence extends the current model or invalidates part of the declaration itself.

Thus universal completion and categorical identification answer different questions. A Kan extension asks for the best extension *under this declaration*. ORACLE asks which declaration and which equivalence class of worlds are supported by interaction.

Proposition 4.8 (Kan-insufficiency principle). *Let M and N be inequivalent admissible worlds whose observer-visible restrictions induce equivalent boundary diagrams $F_M \simeq F_N : \mathcal{A} \rightarrow \mathcal{C}$. Suppose their current left and right Kan extensions along $i : \mathcal{A} \rightarrow \mathcal{B}$ agree on every registered query, but an admissible continuation or query separates M from N . Then no learner whose inference is restricted to those current Kan extensions can identify which of M or N generated the presentation.*

Proof. Kan extension is determined, up to its universal equivalence, by the declared indexing functor and boundary diagram. Equivalent visible boundary data therefore yield equivalent current universal completions and identical answers to every query that factors through

them. By assumption, the evidence or query that separates M from N lies outside this common observable completion. Consequently an extension-only learner receives no distinguishing witness. $\square$

This is another indistinguishability theorem, not a defect peculiar to Kan extensions. The construction cannot create evidence absent from its boundary diagram. Several further gaps follow.

<i>Relativity.</i>	Universality is relative to the chosen categories, morphisms, and 2-cells. A perfectly computed extension can be the correct answer to a wrongly declared learning problem.
<i>Probes.</i>	The extension need not select the next probe that separates the remaining hypothesis fiber.
<i>Persistence.</i>	Independently universal extensions at times t and $t + 1$ need not preserve settled knowledge. That requires coherent comparison maps, base-change conditions, and a repair policy.
<i>Effectivity.</i>	Abstract existence does not imply that an extension can be computed, estimated from finite noisy data, or represented within the learner's resources.
<i>Causal meaning.</i>	Left and right universality provide extremal completions, but an observer and intervention semantics are still needed to interpret their agreement or discrepancy causally.

Kan extensions nevertheless remain central rather than dispensable. Their place in ORACLE is one layer of a larger architecture:

declaration discovery and repair $\longrightarrow$ universal extension $\longrightarrow$ observer-relative comparison.

The first layer determines what is being extended. The second constructs the universal candidates. The third determines what their comparison warrants for the learner. Online persistence adds coherent transport among successive instances of all three layers.

One may reply that declaration discovery is itself expressible as a Kan extension in a higher category of presentations and doctrines. This may be formally correct; it does not remove the epistemic content. It relocates that content into the choice of higher category, observer, admissible repairs, and equivalence relation. Mac Lane's maxim establishes expressive universality, not the empirical sufficiency of an undeclared extension problem.

Design principle

Every concept may be expressible as a Kan extension, but a Kan extension does not by itself identify which concept is warranted by the evidence.

Design principle

A categorical prior has substantive developmental content only if it predicts an identification, invariance, transport, or repair distinction not already implied by unrestricted sequence or latent-state prediction on the same evidence.

There are then three legitimate outcomes. A Transformer system may satisfy the ORACLE obligations, and a JEPA may identify a query-sufficient, action-coherent world model; in either case the theory supplies semantics for what has been learned. A system may satisfy local prediction while failing persistence, re-presentation, underdetermination, or repair, in which case ORACLE localizes the missing structure. Or the proposed obligations may predict nothing beyond ordinary sequence or latent prediction, in which case the categorical account is a redescription and must be weakened, revised, or rejected.

The elephant in the room thus becomes a theorem and evaluation program rather than a contest between neural and categorical slogans. The central question is whether frontier systems discover persistent compositional worlds, learn query-sufficient predictive quotients of them, or reconstruct locally adequate fragments of worlds on demand.

Guiding question

Can two learners with matched compositional performance be separated by a finite family of re-presentation, continuation, repair, abstention, and transport tests derived from the declared categorical prior?

4.9 *Piagetian construction of a compositional world*

CORE KNOWLEDGE CONSTRAINS THE FIRST HYPOTHESES; PIAGETIAN LEARNING DESCRIBES HOW THOSE HYPOTHESES CHANGE. ORACLE needs both. A fixed Spelke sketch without developmental dynamics becomes an inflexible ontology. Accommodation without a structured prior becomes unrestricted model replacement. The combined theory asks how a learner extends, repairs, and occasionally revises a core-structured world while retaining what earlier interaction has genuinely

established.

4.10 *Development is not merely clock time*

Let $\mathbb{D}$ be a category of developmental contexts. Its objects may record time, available actions, perceptual resolution, social setting, and the fragment of the core sketch currently exposed. An arrow $d \rightarrow d'$ means that the learner can transport established structure from context d into the richer context d' . This is not required to be a total order: motor, linguistic, spatial, and social capacities can develop along partially independent directions.

A developmental presentation is a functor

$$E : \mathbb{D} \longrightarrow \mathcal{T}_{\text{Doc}}, \quad d \longmapsto (T_d, \mathbb{T}_d),$$

into the transcript-and-doctrine base introduced in Chapter 3. The learner selects a compatible section of the pulled-back hypothesis fibration. Thus development records both an expanding presentation and the genealogy of the hypotheses used to interpret it.

4.11 *Assimilation as conservative lifting*

For an evidence extension $u : T \rightarrow T'$ under a fixed doctrine $\mathbb{T}$, restriction gives

$$u^* : H_{\mathbb{T}}(T') \longrightarrow H_{\mathbb{T}}(T).$$

Suppose the learner currently holds $M \in H_{\mathbb{T}}(T)$.

Definition 4.9 (Categorical assimilation). *An assimilation of T' into M is a choice $M' \in H_{\mathbb{T}}(T')$ together with a registered comparison $u^*M' \rightarrow M$ that preserves a declared subtheory of settled queries. It is *conservative* when that comparison is an equivalence on the settled subtheory.*

Assimilation need not leave every belief unchanged. It enriches a model with new objects, arrows, composites, or overlap warrants while requiring the validated part of the old interpretation to survive. The declaration of the settled subtheory prevents “preservation” from becoming an all-or-nothing condition.

Proposition 4.10 (Persistence under conservative assimilation). *Let $\mathcal{Q}_0$ be a settled query subtheory. If M' conservatively assimilates T' into M , then every answer in $\text{Th}_{\mathcal{Q}_0}(M)$ is preserved by M' .*

Proof. By definition the comparison $u^*M' \rightarrow M$ is an equivalence for the semantics of every query in $\mathcal{Q}_0$. Query answers invariant under that equivalence therefore agree. $\square$

The content lies in choosing $\mathcal{Q}_0$ before seeing the desired outcome. If the learner is free to declare every overwritten conclusion unsettled after the fact, the persistence claim is vacuous.

4.12 Two levels of accommodation

Let

$$\text{Lift}_u(M) = \{M' \in H_{\mathbb{T}}(T') : u^*M' \simeq M \text{ on } \mathcal{Q}_0\}$$

be the compatible-assimilation fiber.

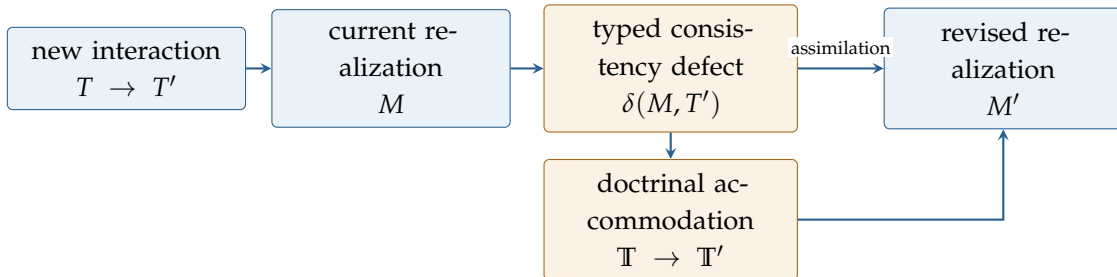
Proposition 4.11 (Assimilation obstruction). *If $\text{Lift}_u(M)$ is empty, the new evidence cannot be conservatively assimilated into M relative to $\mathcal{Q}_0$.*

This obstruction does not yet say which accommodation is warranted. *Ordinary accommodation* chooses another object of $H_{\mathbb{T}}(T')$: the realization changes but the doctrine does not. *Doctrinal accommodation* changes $\mathbb{T}$ along an admissible sketch morphism and transports the surviving interpretation into a different hypothesis fibration. The first response says “my theory of this world was wrong.” The second says “my language of possible worlds was too restrictive or wrongly typed.”

An empty compatible-assimilation fiber is evidence against the current realization, not automatically against the core doctrine. Doctrinal repair requires the stronger finding that adequate models remain absent, or retain a persistent observable defect, throughout the relevant old-doctrine fiber.

4.13 Equilibration as a repair process

Piagetian equilibration is not a return to a static endpoint. It is the process of restoring coherence after interaction exposes a mismatch among a realization, its observers, and its cross-fragment warrants. Let $\delta(M, T')$ be a typed defect object, not necessarily a scalar. Its components may include a failed lift, an unfillable horn, an incompatible overlap, or an observer-visible discrepancy.



Repair is successful when it maps the defect to a declared zero or admissible residual object while preserving the settled query subtheory.

Figure 4.2: A categorical Piagetian cycle. New evidence is first tested against the current realization. A typed defect guides assimilation or ordinary accommodation; only a persistent obstruction at the level of the hypothesis fiber warrants doctrinal accommodation.

This formulation connects Piagetian development to the later homotopy and tangent repair chapters without pretending that all cognitive disequilibrium is numerical optimization.

4.14 *Developmental composition of the core systems*

The six core systems need not become globally integrated at once. Let $I_d \subseteq I$ denote the fragments and overlaps available at developmental context d . An arrow $d \rightarrow d'$ may enlarge this cover, increase observer resolution, or introduce a new cross-fragment warrant. The partial core world at d is a compatible family

$$(M_i, M_{ij}, M_{ijk}, \dots)_{i,j,k \in I_d}.$$

This gives a concrete interpretation of developmental construction. Learning a word for a persisting object adds a language–object overlap. Learning that another agent can use the word to redirect attention adds a language–object–social coherence. Learning a numeral adds an object–number–language interface. The growth is simplicial: new higher compatibilities witness that previously separate capacities now participate in one compositional world.

Theorem 4.12 (Developmental gluing under conservative growth). *Assume each developmental cover has effective core descent. If every transition $d \rightarrow d'$ conservatively transports the identified fragments and overlap comparisons, then the corresponding global realizations form a functor $M : \mathbb{D} \rightarrow \text{Mod}(\mathbb{T}_{\text{core}})$ up to the declared coherence.*

Proof. At each d , effective descent supplies a global realization unique up to equivalence. Conservative transport defines a morphism between the local descent data. Functoriality and uniqueness of descent induce the global comparison, and the same uniqueness supplies identity and composition coherence. $\square$

The theorem states the clean case. Development becomes scientifically interesting when the transport fails: the defect then localizes whether an old fragment, a newly introduced overlap, or the doctrine of gluing requires repair.

4.15 *A grounded research program*

The categorical Spelke–Piaget proposal can now be evaluated through a common protocol for each core system.

1. Specify the weakest fragment sketch that expresses the hypothesized competence.

2. Define the presentations realistically available to the learner and the observers converting interaction into typed evidence.
3. State the equivalences under which competence should remain invariant.
4. Exhibit queries for which the fragment provides a presentation-stable identification gain over a weaker ambient class.
5. Specify overlaps with other systems and test whether their local models admit effective or approximate descent.
6. Predict which new evidence should be assimilated, which should revise a realization, and which would count against the doctrine itself.

This is primarily a theoretical book, so these items are theorem obligations rather than an experimental protocol implemented here. They nevertheless anchor the formalism in developmental claims. A categorical model that cannot state these contrasts has not yet captured the truth of the core-knowledge hypothesis; it has only supplied notation for it.

Guiding question

Is there a smallest core sketch whose realizations explain the observed early invariances, whose overlaps reconstruct later compositional competence, and whose weakening provably destroys identifiable developmental queries?

Further reading

Piaget's account of sensorimotor intelligence motivates the distinction among assimilation, accommodation, and equilibration [Piaget, 1952]. Spelke's core-knowledge program supplies the fragment-specific developmental constraints and the challenge of later composition [Spelke, 2022, 2023]. The hypothesis-fibration and 3. Later parts develop homotopy, tangent, safety, and transport doctrines for the repair obligations exposed by this developmental model.

Moro's impossible languages illustrate a domain in which biological admissibility is proposed to define a proper subclass of formal structures [Moro, 2023]. The Transformer architecture was introduced by Vaswani et al. [2017]. The systematic-generalization results of Lake and Baroni [2018, 2023] caution against both identifying benchmark composition with world identification and declaring neural architectures incapable of systematicity. Current agent time-horizon measurements and coding-evaluation audits illustrate both the extraordinary rate

of capability improvement and the fragility of the tests used to describe it [METR, 2026, OpenAI, 2026b]. The JEPA program provides a complementary account of abstraction through prediction in learned representation spaces [LeCun, 2022, Assran et al., 2023, 2025]. Mac Lane’s treatment of Kan extensions supplies the classical universality claim; the distinction developed here is between that expressive universality and identification from a presentation [Mac Lane, 1971].

Part III

Universal Online Categorical Learning

5

Learning Inside Fixed Doctrines

A DOCTRINAL AUDIT OF ESTABLISHED LEARNING PARADIGMS

No learning algorithm begins without a doctrine. Before observing data, it already presupposes which objects and maps may constitute a world, how observations compose, which solution constructions exist, and which equivalences make two candidate worlds indistinguishable. These commitments are often distributed across a model class, a loss, a feedback protocol, and the hypotheses of a convergence theorem. Chapter 1 collects them under one name: the learner’s doctrine.

Established methods become tractable by fixing a sharply restricted doctrine before learning begins. Reinforcement learning ordinarily fixes a Markovian, reward-enriched world in which Bellman operators can be formed; causal discovery fixes a world of graphs or structural mechanisms with intervention semantics; automata learning fixes a coalgebraic machine type; and Transformer language modeling fixes a sequential presentation and continuation-prediction interface. Their achievements show how powerful the right doctrine can be. They do not show that the doctrine itself was discovered from interaction.

The central claim of this chapter is therefore not that every learning algorithm is secretly doing category theory. It is that a learning paradigm becomes a *query-relative UOCL specialization* only after its structural doctrine, presentation, solution construction, and success quotient are made explicit. Most existing methods learn inside that declaration. ORACLE asks the additional question of how the declaration can be identified, compared, composed, and repaired.

5.1 *The doctrinal audit*

Recall that an ORACLE problem identifies a compositional world only up to the equivalence detected by a declared query doctrine. A doctrinal audit asks what had to be fixed before the familiar learner could even state its objective.

Definition 5.1 (Doctrinal UOCL specialization). A learning paradigm is a *query-relative UOCL specialization* when it specifies:

1. a structural doctrine $\mathcal{D}$ and a category $\mathbf{H}_{\mathcal{D}}$ of admissible theories, models, and structure-preserving comparisons;
2. a presentation protocol $\mathbf{Pres}$ revealing finite evidence;
3. a query doctrine $\mathcal{Q}$ of questions posed to hypotheses;
4. an internal solution construction, such as conditioning, minimization, a fixed point, final semantics, or predictive completion;
5. observational equivalence $\simeq_{\mathcal{Q}}$ induced by those questions;
6. a non-anticipatory update, selection, or posterior mechanism; and
7. a finite, limiting, probabilistic, or regret-based success criterion, together with a boundary separating repairs internal to $\mathcal{D}$ from evidence against the doctrine itself.

It is *structurally faithful* when the morphisms and compositions of $\mathbf{H}_{\mathcal{D}}$ affect at least one admitted query, solution, update, repair, or transport claim.

The last condition prevents a vacuous recoding. Any set can be regarded as a discrete category, but doing so explains nothing about learning. A useful specialization identifies compositions that constrain predictions, relate models, transport solutions, or localize revisions.

Definition 5.2 (Learning within a fixed doctrine). A specialization is *fixed-doctrine* when every learner state and update remains in $\mathbf{H}_{\mathcal{D}}$. It becomes *doctrine-learning* only when the learner state also records $\mathcal{D}_t$ and updates may follow declared doctrine maps $\mathcal{D}_t \rightarrow \mathcal{D}_{t+1}$, while transporting the structure that remains warranted.

Fixed-doctrine learning is not a defect. It is the source of the strong existence, convergence, and sample-complexity results available in mature fields. The distinction concerns the scope of what has been learned. A proof that an algorithm converges inside $\mathcal{D}$ cannot by itself establish that interaction selected $\mathcal{D}$ from competing doctrines.

Proposition 5.3 (Specialization criterion). *Suppose a learning method supplies the items of Definition 5.1, and its update maps respect reindexing of presentations and the declared observational equivalence. Then it determines a UOCL instance in the sense of Definition 6.1. If the update also conservatively transports settled query answers, it determines a persistent UOCL instance on that settled doctrine.*

Proof. Take $\mathbf{H}_{\mathcal{D}}$ as the hypothesis fiber, pull it back along the presentation protocol, and use the method's updates as lift selectors. Reindexing supplies non-anticipation, while compatibility with $\simeq_{\mathcal{Q}}$ makes

the query quotient well defined. The stated conservativity condition is precisely persistence of settled answers. $\square$

The proposition is an interface theorem, not a claim of algorithmic equivalence. It says exactly what must be exhibited before calling an existing method a UOCL specialization.

Proposition 5.4 (Internal success does not identify the doctrine). *Let $\mathfrak{D}$ and $\mathfrak{D}'$ admit hypotheses M and M' that induce the same presentation law and agree on every query in $\mathcal{Q}$. No learner whose input and success criterion factor only through that presentation and query doctrine can determine whether the underlying world was $(\mathfrak{D}, M)$ or $(\mathfrak{D}', M')$.*

Proof. The learner receives identically distributed evidence and every scored answer is equal in the two cases. Its observable execution therefore has the same law under both hypotheses. Any claimed doctrinal distinction would require an additional probe, intervention, or prior not present in the declaration. $\square$

This elementary obstruction is the recurring lesson of the chapter. Bellman convergence can certify an optimum in an MDP without certifying Markovity; likelihood can select a sequence law without identifying a grammar; and observational fit can select a causal graph only up to the equivalence exposed by the admitted data and assumptions.

5.2 System identification and predictive state representations

Fixed doctrine. System identification supplies the cleanest precursor to UOCL. A learner observes inputs and outputs of an unknown dynamical system and constructs a model sufficient to predict its future behavior. The hypothesis category may contain state-space models, stochastic automata, or controlled processes; morphisms express coordinate changes, simulations, or observational quotients. The presentation is a growing trajectory.

The declaration already assumes that inputs and outputs are typed, temporal composition is meaningful, the dynamics belong to the admitted model class, and the chosen observations can expose the distinctions demanded by the queries. Linear identification may additionally assume finite dimension, controllability, observability, or a noise family. These are not conclusions of the estimator. They define the doctrine in which estimation takes place.

Predictive state representations sharpen the analogy. Rather than postulate a hidden state with privileged ontology, a PSR represents state by predictions of future, action-conditional observation tests [Littman et al., 2001]. Let a history be $h = a_1 o_1 \cdots a_t o_t$, and let a test $\tau = a'_1 o'_1 \cdots a'_k o'_k$ ask for the probability of receiving the specified observations after

Paradigm	Fixed structural doctrine	Internal solution construction	Query quotient and repair boundary
Automata	Typed transition coalgebras with fixed alphabet, outputs, and often finite carriers	Final behavior, minimization, and counterexample refinement	Behavioral quotient; revise the endofunctor, types, or finiteness assumption when it fails
System identification and PSRs	Controlled stochastic processes with composable kernels and admitted tests	Conditioning and finite predictive state update	Predictive equivalence; revise stationarity, observability, or test sufficiency
Topos World Models	Context-indexed local models with restriction, provenance, and descent	Typed extraction followed by local update and sheaf gluing	Contextual predictive quotient; revise the cover, gluing law, or causal warrant
Reinforcement learning	Markov stochastic dynamics enriched by reward, policies, order, and Bellman structure	Planning, Bellman fixed points, or stochastic approximation	Bisimulation or policy-sufficient quotient; revise Markovity, observability, stationarity, or value structure
Causal discovery	DAGs or structural mechanisms with declared intervention semantics and independence assumptions	Constraint, score, or functional-model identification	Markov or interventional equivalence; revise acyclicity, sufficiency, modularity, or variable ontology
Sequence modeling	Token and prefix categories with conditional continuation laws and positional structure	Likelihood fitting, attention, and autoregressive generation	Predictive sequence equivalence; revise tokenization, grammar, grounding, or memory assumptions
JEPA and learned world models	Latent quotients with predictive targets, readouts, and sometimes action dynamics	Representation prediction and model-based planning	Representation- or control-sufficient quotient; revise the quotient when later probes need discarded distinctions

Table 5.1: A doctrinal audit of established learners. Their strongest guarantees hold inside the structural commitments in the second column. The final column marks the point at which improving a model within the doctrine may no longer be an adequate repair.

issuing the specified actions. For a family of core tests $\mathcal{T} = \{\tau_i\}_{i \in I}$, the predictive state is

$$s_{\mathcal{T}}(h) = (\Pr(\tau_i \mid h))_{i \in I}. \quad (5.1)$$

An action–observation pair extends the presentation from h to hao , and conditioning updates $s_{\mathcal{T}}(h)$ to $s_{\mathcal{T}}(hao)$. Thus a PSR has almost exactly the interface of Definition 5.1: controlled stochastic systems form the hypothesis category, histories form the presentation, future tests form the query doctrine, and conditioning supplies the non-anticipatory update.

The induced observational quotient is especially important. Two histories are predictively equivalent when

$$h \sim_{\mathcal{T}} h' \iff \Pr(\tau \mid h) = \Pr(\tau \mid h') \text{ for every } \tau \in \mathcal{T}. \quad (5.2)$$

The learner identifies the quotient visible through $\mathcal{T}$, not a unique hidden-state ontology. Enlarging the test family refines the quotient; finding a finite sufficient family amounts to finding a compact presentation of the relevant predictive world.

This test-based view has a deterministic precursor in the diversity representation of Rivest and Schapire. Their learner represents a fi-

nite automaton by equivalence classes of experiments rather than by enumerating its global states [Rivest and Schapire, 1994]. A PSR may therefore be read as a stochastic and, in common finite-dimensional realizations, linear analogue: Boolean outcomes of deterministic tests are replaced by conditional probabilities of action–observation tests. The analogy concerns the predictive interface; it does not identify the two learning algorithms or their assumptions.

Proposition 5.5 (PSRs are query-relative UOCL). *Fix a class of controlled stochastic processes and a test family $\mathcal{T}$ closed under the one-step extensions required by the update. If the PSR update is well defined on equivalence classes $[h]_{\mathcal{T}}$, then PSR identification is a UOCL specialization whose identified object is the action of one-step extensions on the predictive quotient of histories.*

Proof. The hypothesis class, presentation, and query doctrine have just been specified. Equation (5.2) supplies observational equivalence. Closure under the required extensions makes conditioning a well-defined update on equivalence classes, while predictive accuracy or limiting identification supplies the success criterion. The specialization criterion, Proposition 5.3, then applies. $\square$

This also marks the boundary between UOCL and UODL. The actions occurring in a PSR test are controlled input symbols used to interrogate dynamics. They do not by themselves say which action should be chosen. Preferences, consequences, feasibility, and a comparison observer must still be added to obtain a decision declaration. Conversely, a query family sufficient for one-step prediction may fail to preserve what is needed for long-horizon control.

Doctrinal boundary. A failed parameter fit is ordinarily repaired inside the system doctrine. A persistent failure of every admitted realization may instead challenge stationarity, the state dimension, the observation interface, the noise model, or the existence of a finite sufficient test family. PSR learning becomes doctrine-learning only if such alternatives are themselves typed and available to the update mechanism.

5.3 *Topos World Models from documents*

Prometheus and Odyssey provide a less conventional UOCL specialization. The presentation is not initially an action–observation trajectory. It is a growing corpus of papers, reports, manuals, reviews, data summaries, and other research artifacts. A compiler extracts typed fragments—entities, contexts, claims, candidate actions or interventions, reported outcomes, and provenance—and organizes them into local action-oriented predictive-state models. Prometheus introduced this

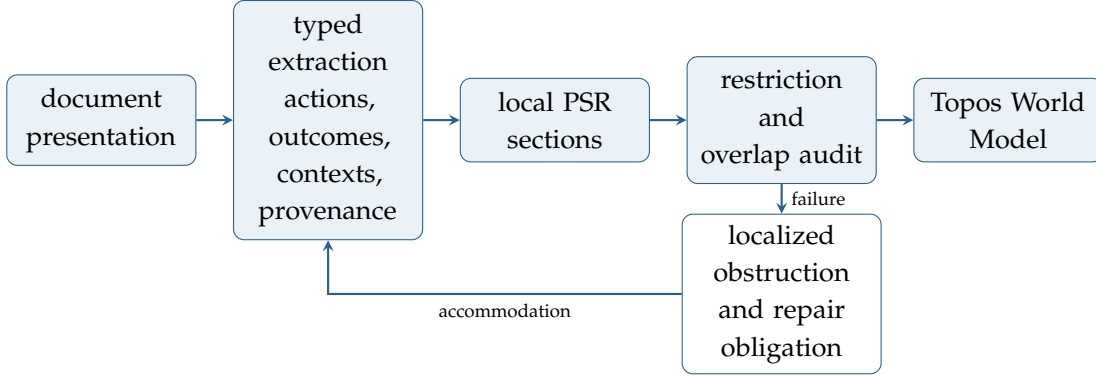


Figure 5.1: Prometheus/Odyssey as a UOCL pipeline. Documents provide an online presentation; typed extraction compiles them into action-oriented local PSR sections. Compatibility yields a glued Topos World Model, while failed overlaps return localized repair obligations rather than forcing a false global model.

construction as a Topos World Model; Odyssey places it inside a larger foundry whose intermediate artifacts, admission decisions, and repair obligations remain inspectable [Mahadevan, 2026e,d].

Fixed doctrine. The pipeline presupposes a site or context category, local predictive-model fibers, restriction maps, a provenance discipline, and a gluing semantics. Its internal constructions are typed extraction, local PSR update, and descent. The resulting guarantees are therefore conditional on documents being compilable into those types and on compatibility over the chosen cover being the right notion of global coherence.

The word *action-oriented* is deliberate. A document fragment such as “under condition U , action a was followed by outcome o ” can constrain a test $\tau = (a, o)$, even when the learner did not execute a . Text can therefore present predictive and causal claims about actions. It cannot, by extraction alone, establish that those claims identify causal effects. Interventional semantics still require provenance, design or identification assumptions, and consistency across sources. UOCL cleanly records this boundary by typing textual support, predictive adequacy, and causal warrant as different doctrines.

Let $\mathcal{C}$ be a category of overlapping contexts and let

$$\mathcal{S} : \mathcal{C}^{\text{op}} \longrightarrow \mathbf{PSR} \quad (5.3)$$

assign to each context U a local predictive model $\mathcal{S}(U)$. Restriction maps express what one local model says on a shared subcontext. When compatible local sections glue, the result is not one unqualified global state vector but a coherent atlas of context-indexed predictive states. This is the topos-theoretic refinement of the PSR idea: local predictive quotients become objects in a presheaf or sheaf world rather than coordinates in a single globally privileged model.

Proposition 5.6 (Persistent gluing of predictive state). *Suppose $\{U_i \rightarrow U\}$ is a cover, the local predictive sections $s_i \in \mathcal{S}(U_i)$ agree on every overlap,*

and restriction commutes with the admitted PSR updates. If $\mathcal{S}$ is a sheaf, the local states determine a unique glued section $s \in \mathcal{S}(U)$. Any later compatible update that preserves a settled local query preserves its glued answer.

Proof. Existence and uniqueness of s are the sheaf gluing axiom. Computation of restriction with update makes the updated family compatible. Its unique gluing restricts to the updated local sections; hence any query unchanged on the cover is unchanged in the glued section. $\square$

The UOCL reading is now precise. Documents are successive presentation fragments; the local PSRs are hypotheses in context-indexed fibers; new documents induce updates and possible refinements of the cover; restriction maps express comparison; gluing is coherent assimilation; and an incompatible overlap registers an accommodation problem. Prometheus/Odyssey therefore extends classical PSR identification in two directions at once: from executed trajectories to heterogeneous textual presentations, and from one predictive state space to a sheaf of locally valid, provenance-bearing predictive worlds. It is a concrete instance of UOCL, while remaining only one possible algorithmic realization of the broader ORACLE program.

Doctrinal boundary. Changing an extracted claim or refining a local PSR is internal repair. Changing the context cover may still be a controlled structural refinement. Rejecting the assumed restriction maps, sheaf condition, provenance types, or separation between textual support and causal warrant changes the doctrine under which the Topos World Model was assembled.

5.4 Automata and machine inference as coalgebra learning

The longstanding program of inferring finite-state machines, sequential machines, and computable devices from examples or experiments is an especially literal instance of UOCL. Fix an input alphabet A and an output object O . A deterministic Moore machine is a coalgebra

$$c : X \longrightarrow O \times X^A \quad (5.4)$$

for the endofunctor $F(X) = O \times X^A$. When $O = 2$, this is a deterministic language acceptor: the first component records acceptance and the second gives the successor reached by each input symbol. A finite-state learning problem therefore asks for a finite carrier X , an F -coalgebra structure on it, and a designated initial state that reproduce the observed behavior.

Fixed doctrine. The alphabet, output type, endofunctor, deterministic transition law, and usually finiteness of the carrier are supplied

before inference. Final semantics determines behavioral equivalence; minimization or an observation table supplies the internal solution construction. Angluin-style membership and equivalence queries are not generic interaction but a particularly strong presentation protocol. Together these commitments explain both the precision of automata-learning theorems and the narrowness of their target.

This formulation exposes the universal object that the learner can actually identify. When a final coalgebra $(\nu F, \zeta)$ exists, every candidate has a unique behavior map

$$\text{beh}_c : X \longrightarrow \nu F, \quad \zeta \circ \text{beh}_c = F(\text{beh}_c) \circ c. \quad (5.5)$$

For deterministic acceptors, $\text{beh}_c(x)$ is the language accepted from state x ; for Moore machines, it is the complete input-indexed output behavior. Classical automata inference consequently learns the image of reachable states in final semantics, usually represented by the minimal automaton, rather than the accidental names or redundant realization of hidden states. Angluin's L^* algorithm makes the presentation protocol unusually explicit: membership queries reveal local behavior, while equivalence queries return counterexamples that separate an incorrect conjecture [Angluin, 1987].

5.4.1 Diversity representations: learning tests instead of states

Rivest and Schapire give a particularly revealing alternative to literal state reconstruction [Rivest and Schapire, 1994]. Let B^* be the free monoid of action strings and let P be a family of Boolean predicates on the state set X . A test is a pair $ap \in B^*P$, with semantics

$$\text{val}(ap)(x) = p(\delta(x, a)).$$

Two tests are equivalent when they have the same value at every state. The *diversity*

$$D = |B^*P / \equiv|, \quad ap \equiv a'p' \iff \text{val}(ap) = \text{val}(a'p'), \quad (5.6)$$

measures the number of distinct observable tests, not the number of global states. Prefixing a test by an action descends to this quotient and produces an update graph. The values of the quotient tests at the current state, together with that graph, suffice to simulate every admitted experiment.

The Rubik's Cube example makes the point vivid. The modeled environment had about 10^{19} global states but diversity 54; their implementation learned the test representation while visiting only a minute fraction of the state space. What was compact was not a list of cube configurations. It was a generating family of observations together with the action-induced rules for transforming them.

Proposition 5.7 (Categorical test-quotient construction). *Let $\mathcal{C}$ be a category, $X \in \mathcal{C}$, Ω an observation object, $G \subseteq \mathcal{C}(X, X)$ a family of action generators, and $P \subseteq \mathcal{C}(X, \Omega)$ a family of observations. Let $\mathbb{A}$ be the submonoid of endomorphisms generated by G , and let*

$$\mathcal{T} = \{p \circ a \mid a \in \mathbb{A}, p \in P\} \subseteq \mathcal{C}(X, \Omega).$$

For any congruence $\sim$ on $\mathcal{T}$ preserved by precomposition with the generators, each $g \in G$ induces a well-defined update $U_g([t]) = [t \circ g]$ on $\mathcal{T}/\sim$. Hence a finite separating quotient $\mathcal{T}/\sim$, its generator updates, and its current evaluation form a compact predictive presentation of the admitted experiments.

Proof. If $t \sim t'$, preservation by precomposition gives $t \circ g \sim t' \circ g$, so U_g is independent of representatives. Iterating the U_g 's evaluates every word in the generated action monoid. Separation says that no distinction visible to the admitted tests is lost. $\square$

This proposition suggests the categorical generalization that the original construction leaves implicit. Replace the one-object action monoid by a typed action category or algebraic theory. Its actions act contravariantly on observable predicates by precomposition, producing a test presheaf; quotient that presheaf by the observational congruence detected in interaction. The learning target is then a compact presentation of the resulting action on tests. In theory-bearing UOCL, one may go further and learn generators and relations for both the action theory and its test module. This avoids constructing a global state model unless some later query actually requires one.

Boundary

Low diversity is relative to the chosen actions and observations. A compact test quotient can be sufficient for prediction while omitting distinctions needed for control, causality, safety, or transport. UOCL must therefore record the query doctrine under which the quotient is separating.

5.4.2 Krohn–Rhodes decomposition: learning cascades instead of states

The diversity construction compresses what the learner must observe. Krohn–Rhodes theory asks a complementary question: once a finite transformation system has been identified, from which elementary machines can its dynamics be composed? For a deterministic automaton with state set X and input alphabet A , let

$$M_A = \langle \delta_a : X \rightarrow X \mid a \in A \rangle \subseteq X^X$$

be its transformation monoid. A monoid M *divides* a monoid N , written $M \prec N$, when M is a homomorphic image of a submonoid of

N . Division is the appropriate notion of simulation here: the larger machine may contain additional states and distinctions, provided a stable subsystem maps onto the behavior to be represented.

Theorem 5.8 (Krohn–Rhodes prime decomposition). *For every finite transformation monoid M there are prime components $P_1, \dots, P_k$ such that*

$$M \prec P_1 \wr P_2 \wr \dots \wr P_k,$$

where each P_i is either a finite simple group or a flip-flop component. Equivalently, every finite automaton admits a homomorphic representation by a feedback-free cascade of permutation and reset machines.

The formulation suppresses a useful technical distinction. The three-element flip-flop monoid consists of an identity together with the two constant transformations of a two-state set. It is therefore a one-bit set-reset device, not a counter in the usual numerical sense. More general statements first decompose a finite monoid into group and aperiodic factors; the aperiodic factors can in turn be divided by iterated wreath products of flip-flops. The theorem is an existence and simulation result, not a claim that the decomposition is unique or easy to infer [Krohn and Rhodes, 1965].

Example 1.21 gives this algebraic dichotomy a simplicial reading. A group factor, regarded as a one-object category, has a Kan nerve: its transitions are reversible and its outer horns can be filled. A reset factor has noninvertible arrows; its nerve retains inner composition but generally fails the outer-horn condition. A wreath product does more than place these factors side by side. It arranges them as a directed cascade in which the input received by a later component may depend on the states and outputs of earlier components. Krohn–Rhodes theory therefore separates a finite compositional world into reversible memory, irreversible memory update, and hierarchical dependency.

This suggests an explicit UOCL objective. First infer a finite observational quotient $\hat{M}_Q$ from a separating query doctrine Q , as in the diversity construction. Then search for a cascade presentation

$$\hat{M}_Q \prec P_1 \wr \dots \wr P_k \tag{5.7}$$

that preserves the admitted tests. Instead of penalizing the number of global states, the learner can penalize the number and complexity of the prime factors, their dependency pattern, and especially the number of nontrivial group factors. The classical minimum number of group layers is the Krohn–Rhodes group complexity. In ORACLE it becomes one candidate measure of the structural depth required to explain an observed finite world.

Ronca et al. [2023] provide a modern learning-theoretic step in this direction. They treat automata cascades as modular hypothesis classes

and derive sample-complexity bounds that, up to logarithmic factors, scale with the number of components and the maximum complexity of a component rather than with the potentially exponential global state space. Their result does not yet solve the ORACLE problem of discovering the right prime cascade from an unrestricted interaction presentation. It does demonstrate that a correct compositional prior can change the statistical scale of automata learning.

There is also a genuine categorical lineage rather than a merely suggestive analogy. Wells [1980] established a Krohn–Rhodes theorem for finite categories and set-valued functors. Wells later decomposed category-valued functors by categorical wreath products, explicitly motivating state-transition systems with structured, typed states [Wells, 1988]. Thérien [1991] defined a two-sided wreath construction, or block product, directly on categories. These results suggest replacing a one-object transformation monoid by a typed action category and replacing a flat state set by a functor of structured state fibers.

5.4.3 Toward stochastic and coalgebraic prime decomposition

The preceding account raises a question with direct consequences for compositional learning: can prime decomposition survive when transitions are stochastic? There is already an important partial answer. Carlsson and Yu [2015] define a probabilistic automaton by extending the action of a finite deterministic transition semigroup from states and letters to probability distributions. Convolution then gives a compact semigroup of distributions on the finite transition semigroup, and the classical Krohn–Rhodes theorem supplies its underlying prime decomposition. Their result shows that probability does not erase the group–reset architecture when randomness is distributed over a finite deterministic action semigroup.

That qualification matters. A general finite-state controlled Markov system is more naturally a coalgebra

$$c : X \longrightarrow (\mathcal{D}X)^A,$$

where $\mathcal{D}$ is the finite-distribution monad and each action determines a stochastic kernel $K_a : X \rightarrow \mathcal{D}X$. Its transitions compose in the Kleisli category $\mathbf{Kl}(\mathcal{D})$, not in the finite transformation monoid X^X . Even when X and A are finite, the stochastic matrices generated by composition can form an infinite, indeed continuously parameterized, semigroup. The finiteness hypothesis that drives the classical induction has disappeared.

There is nevertheless a clean bridge for rational systems.

Proposition 5.9 (Finite stochastic dilation). *Let X and A be finite, and suppose every entry of every controlled kernel K_a is rational. Then there*

are a finite noise set E , the uniform distribution u_E , and deterministic maps $f_{a,e} : X \rightarrow X$ such that

$$K_a(x, -) = \sum_{e \in E} u_E(e) \delta_{f_{a,e}(x)}.$$

Consequently the transition dynamics admit an exact randomized realization by a finite cascade of group and flip-flop components.

Proof. Choose a common denominator N for the finitely many entries of the kernels and let $E = \{1, \dots, N\}$. For each pair (x, a) , partition E into subsets of cardinality $NK_a(x, y)$, one subset for each $y \in X$, and set $f_{a,e}(x) = y$ on the corresponding subset. Uniformly sampling e gives exactly $K_a(x, -)$. Apply Theorem 5.8 to the finite transformation monoid generated by the maps $f_{a,e}$, then sample the noise inputs and forget them. $\square$

Thus a finite rational MDP has a Krohn–Rhodes realization at the level of its transition dynamics. Rewards may be carried as output maps or included in an enlarged state–output system. The construction is useful but not yet intrinsic: different deterministic dilations of the same kernels may have very different prime cascades. A genuine stochastic theorem should therefore define a wreath or cascade product inside $\mathbf{KI}(\mathcal{D})$, a stochastic notion of division or simulation, and a complexity invariant independent of the chosen noise realization.

Once a deterministic dilation has been fixed, the Transformer construction of Liu et al. [2023] supplies shallow parallel simulators for its finite semiautomaton, including constant-depth shortcuts when its group factors are solvable. The resulting pipeline

$$\begin{aligned} \text{stochastic kernel} &\longrightarrow \text{deterministic dilation} \\ &\longrightarrow \text{prime cascade} \longrightarrow \text{Transformer shortcut} \end{aligned}$$

is an exact realization scheme for rational kernels. It also compounds the nonuniqueness: behavioral success of the final Transformer need not identify either an intrinsic stochastic factorization or the dilation through which it was compiled.

Real-valued kernels suggest an approximate version. On a finite state and action space, approximate each kernel in total variation by a rational kernel. If the one-step error is at most η , a standard coupling argument bounds the discrepancy of length- T trajectory distributions by at most $T\eta$. Proposition 5.9 then gives a finite prime cascade for the rational approximation. This is a natural meeting point with PACC UOCL: the learner need not identify a unique stochastic mechanism, but should find, with high probability, a cascade whose admitted finite-horizon queries are within the declared tolerance.

PSRs point toward a second generalization. Their unnormalized predictive updates are linear operators on a space of tests, followed by normalization; the relevant factors need not themselves be stochastic state transitions. Plotkin and Plotkin [2015] develop decompositions of linear automata using triangular products together with wreath products, and characterize these products as terminal objects among appropriate cascade connections. This does not by itself yield a decomposition theorem for PSRs. It suggests that an ORACLE learner should decompose the observable operator system modulo predictive equivalence, while preserving positivity and normalization, rather than first reconstructing a latent MDP.

For an arbitrary endofunctor F , an unrestricted Krohn–Rhodes theorem for F -coalgebras is too much to expect. The functor selects the admissible branching, observation, and transition structure, and hence also determines what a cascade and a prime could mean. A workable theorem must be relative to such data as a finitary or accessible F , finite behavioral quotients, a simulation factorization system, and a distributive law supporting cascade composition. Final semantics supplies behavioral equivalence, but not by itself a prime factorization of realizations.

Design principle

Stochastic compositional learning should factor observable dynamics before it enumerates latent states. Deterministic dilation gives an exact starting point for finite rational kernels; an intrinsic theory must make the factors invariant under changes of noise realization and relative to the learner’s query doctrine.

This produces a concrete ladder of research problems for ORACLE:

1. exact randomized cascade realizations for finite rational kernels;
2. PACC cascade identification for real-valued kernels, evaluated by finite-horizon queries;
3. an intrinsic Kleisli–Krohn–Rhodes theorem for stochastic automata and MDPs;
4. positivity-preserving linear or operator decompositions for PSRs; and
5. functor-relative prime decompositions for finite behavioral quotients of F -coalgebras.

The conceptual payoff is larger than an automata generalization. A learned world model would come with a compositional anatomy:

reversible symmetries, irreversible resets, stochastic mixing, and predictive linear structure. Those factors are reusable hypotheses for transfer, accommodation, and causal intervention, rather than merely a smaller encoding of one task.

Design principle

A finite UOCL learner should seek not only the smallest observational realization, but the smallest compositional explanation of that realization. Diversity controls the number of separating tests; cascade complexity controls how the induced transformations are assembled from reversible and irreversible components.

Guiding question

Under which presentation and query conditions can an online learner identify a categorical wreath-product decomposition up to division or behavioral equivalence? Can the decomposition be chosen persistently, so that new evidence refines local factors without reconstructing the entire cascade?

Proposition 5.10 (Final-semantics obstruction). *Let the admitted query doctrine factor through the behavior maps into a final F -coalgebra. If two pointed coalgebras have the same final behavior, no learner receiving only answers to those queries can distinguish them. Consequently, the strongest identifiable target is their behavioral quotient; identification of a particular internal machine requires intensional probes or an additional realization prior.*

Proof. Every admitted query is, by assumption, a function of final behavior. Equal final behaviors therefore induce identical answers under every possible interaction transcript. A learner driven by such transcripts must evolve identically on the two targets, so it cannot be guaranteed to select distinct internal realizations. Quotienting by equality of final behavior removes exactly this observational ambiguity. $\square$

Turing-machine inference fits the same pattern, but with an important qualification. A Turing machine has a finite description, yet its operational coalgebra of configurations is generally infinite: a configuration records the control state, tape contents, and head position, and one transition exposes its next observable event and configuration. Learning the computed language, partial function, or stream behavior is therefore coinductive identification of this configuration dynamics. It is not automatically identification of a unique program. Distinct machines can compute the same behavior, and passive finite data cannot in general settle all future computation. Computability, finite description,

halting conventions, and the permitted presentation must be declared as categorical priors and success conditions. Gold’s limiting recursion and language identification mark the classical boundary between what stabilizes from a presentation and what remains unidentifiable [Gold, 1965, 1967].

This is also the closest established ancestor of universal imitation games. There, the target was a universal coalgebra recovered coinductively from its observable behavior [Mahadevan, 2024]. ORACLE retains that branch but allows the learner to revise the endofunctor, type system, and ambient category when the observed interaction cannot be represented by the initial machine doctrine. Automata learning is thus not merely an analogy: it is a sharply bounded coalgebra-learning laboratory in which presentations, separating queries, behavioral equivalence, minimization, and impossibility are all visible.

Doctrinal boundary. Adding states or correcting transitions remains internal to the fixed coalgebraic doctrine. Evidence for stochastic evolution, unbounded memory, new input types, partial transitions, or a different behavioral endofunctor calls for doctrinal accommodation. General ORACLE permits those changes; classical machine inference normally does not.

5.5 Reinforcement learning inside the MDP–Bellman doctrine

Fixed doctrine. Reinforcement learning does not begin merely with rewards and experience. It normally begins inside a doctrine declaring that the environment can be represented by a class $\mathcal{M}$ of Markov decision processes

$$M = (S, A, P, R, \gamma)$$

with measurable or finite state and action objects, stochastic transition kernels, an additive reward interpretation, an objective such as discounted return, and an admissible policy class. Choosing morphisms such as MDP homomorphisms, simulations, bisimulations, or policy-preserving abstractions then forms a hypothesis category $\mathbf{MDP}_{\mathcal{M}}$ [Puterman, 1994, Sutton and Barto, 2018].

The decisive additional commitment is Bellman structure. For a discounted control problem, the doctrine must support an operator of the form

$$(\mathcal{B}Q)(s, a) = R(s, a) + \gamma \int_S \sup_{a' \in A(s')} Q(s', a') P(ds' \mid s, a). \quad (5.8)$$

The reward must be integrable, the supremum must be meaningful in the chosen value object, and the relevant function space must support the required fixed point. In a finite discounted MDP with bounded

rewards, $\mathcal{B}$ is a contraction on bounded Q -functions in the supremum norm, so it has a unique fixed point. Other objectives require different existence and stability assumptions.

Definition 5.11 (MDP–Bellman doctrine). *An MDP–Bellman doctrine specifies the stochastic category of state transitions, reward and policy objects, the return aggregation rule, an ordered or normed value object, an admitted Bellman operator, and the equivalence under which its solutions represent the same decision behavior.*

These are prerequisites of ordinary RL, not discoveries made by the learning algorithm. Q-learning estimates the fixed point only after the Bellman equation has been declared to characterize optimal behavior. Its classical convergence guarantees additionally require stationary transition and reward laws, sufficient visitation of state–action pairs, and step-size conditions. Those assumptions belong to the presentation and update portions of the doctrinal audit.

An interacting learner receives a prefix of state, action, reward, and successor observations. Transition and reward queries ask which MDPs remain consistent with that prefix; value and policy queries add decision semantics. Model-based reinforcement learning identifies P and R , or a decision-sufficient quotient of them, and then invokes the internal planning construction supplied by the doctrine.

Model-based reinforcement learning is consequently a direct specialization: its UOCL component identifies an MDP model or an observational quotient, and its UODL component uses that model to select consequential actions. The environment’s response makes the presentation endogenous because an action changes which evidence arrives next.

Model-free reinforcement learning requires a more careful statement. A Q-learning agent need not identify P or R separately. It seeks a control-sufficient quotient represented by values or a policy. Calling this full identification of the MDP would be false. It is UOCL only relative to the smaller query doctrine whose answers are the value or action distinctions the algorithm preserves; as a decision process it lies naturally in UODL.

Doctrinal audit. The structural doctrine is the MDP–Bellman package; trajectories provide the endogenous presentation; Bellman optimization or stochastic approximation is the internal solution construction; value, return, and policy form the privileged query doctrine; and bisimulation or policy equivalence determines the visible quotient. The doctrine has been learned only if Markovity, reward aggregation, Bellman adequacy, and the policy/value interfaces are themselves among the hypotheses that interaction can compare.

5.5.1 Options as persistent compositional objects

The options framework gives classical reinforcement learning a revealing fragment of compositional persistence. An ordinary task policy $\pi : S \rightarrow \mathcal{DA}$ is usually valued only through the return it obtains on its parent task. Even when part of that policy performs a recognizable subtask, the learned parameters do not by themselves declare where the behavior is applicable, when it has completed, or how another decision maker may invoke it. The useful composite remains implicit in the solution.

An option promotes such a composite into a named, action-like object

$$o = (I_o, \pi_o, \beta_o),$$

where $I_o \subseteq S$ is its initiation set, π_o is its closed-loop policy, and $\beta_o : S \rightarrow [0, 1]$ is its termination rule [Sutton et al., 1999]. Executing o generates a distribution over termination states, accumulated rewards, and durations. This induced semi-Markov model is the option’s external interface. A planner or learner can therefore treat the entire temporally extended behavior as one action while retaining primitive actions at the lower level.

Categorically, an option is not merely a long path. Its policy selects a controlled family of paths from each admissible starting state, while its termination rule quotients their internal histories into an observable exit distribution. The option model is consequently a macro-morphism in a stochastic action category. Two options compose when the exit support of the first lies within the initiation region of the second; Kleisli composition integrates over the intermediate termination state. The triple (I_o, π_o, β_o) is thus an interface declaration that makes a learned behavior available for later composition.

This explains why an option can outlive the task on which it was discovered. A “navigate to doorway” option learned while delivering one object may be reused while searching, cleaning, or escaping, even when none of those tasks contains the original delivery objective. What persists is not the parent policy but a promoted substructure with a name and invocation contract. In ORACLE language, the learner has moved a useful composite from an ephemeral realization into its storehouse of compositional knowledge.

The promotion is not automatically valid across tasks. A policy optimized under the parent reward may exploit incidental dynamics; its initiation set may omit states encountered elsewhere; and its termination predicate may not match the boundary required by a new composition. Reusing a label despite these changes is syntactic persistence, not semantic transport. Genuine persistence requires an observer-relative equivalence under which the option model, safety

conditions, and interface obligations survive the change of task or are repaired locally.

Design principle

Persistent skill learning should promote a discovered behavior into a named compositional object only together with its domain of applicability, termination interface, behavioral model, and transport conditions. The label makes reuse addressable; the declared interface makes reuse meaningful.

Options therefore sit between conventional RL and ORACLE. They show that temporally extended knowledge can be learned once and invoked in unrelated future objectives. ORACLE generalizes the idea beyond a fixed MDP doctrine: it asks how the option itself is identified, when two realizations express the same reusable skill, how option interfaces compose across modalities and world models, and when failed reuse demands accommodation of the skill’s declaration rather than another value update.

Boundary

The category of MDPs is not determined until its morphisms and observational quotient are declared. State relabelings, homomorphisms, bisimulations, and policy-preserving abstractions answer different identification questions. UOCL makes that choice part of the learning problem rather than hiding it in the word “model.”

This perspective also locates the difference from classical RL. RL normally fixes the MDP doctrine and learns within it. General UOCL may accommodate by revising that doctrine: partial observability, non-stationarity, asynchronous information, non-Markovian composition, a failure of scalar additive reward, or the absence of the assumed Bellman fixed point may force the learner to leave $\mathbf{MDP}_{\mathcal{M}}$ itself. Re-running Q-learning with more samples cannot repair a world that lies outside the doctrine that makes Q-learning meaningful.

5.6 Causal discovery as identification in a causal category

Fixed doctrine. A causal discovery system begins with a category of causal hypotheses: for example directed acyclic graphs, structural causal models, or mechanism families, together with maps preserving the variables and interventions of interest. Its presentation may contain observational samples, interventions, or both. Its query doctrine may ask for conditional independences, interventional distributions, coun-

terfactuals, or transport across environments [Pearl, 2009, Peters et al., 2017].

That category is already highly structured. Depending on the method, its doctrine may include acyclicity, a causal Markov factorization, faithfulness, independent noise, causal sufficiency, modular mechanisms, a fixed variable ontology, or a restricted functional family. Constraint-based search, score-based search, and functional causal modeling use different internal solution constructions because they inherit different subsets of these commitments. None infers causal direction from an unconstrained category of all joint distributions.

The quotient is decisive. Observational data alone often identify only a Markov-equivalence class. Interventional evidence can refine that quotient, but only relative to the available intervention targets and causal assumptions. UOCL describes this as a change in the separating power of the presentation and query doctrine. It does not turn categorical universality into causal identification: intervention semantics and assumptions such as causal sufficiency remain additional declarations.

Causal discovery also illustrates doctrinal accommodation. A contradiction may require changing an edge, a mechanism, or a latent-variable realization within the current causal category. More severe evidence may reject acyclicity, causal sufficiency, modularity, or the chosen variable ontology. The latter is a repair of the hypothesis doctrine, not a better estimate inside the original model class.

Doctrinal audit. The structural doctrine supplies mechanisms and interventions; observational and experimental samples form the presentation; conditional-independence, interventional, or counterfactual inference supplies the solver; and Markov or interventional equivalence supplies the quotient. ORACLE begins where a causal learner can compare those background commitments rather than merely choose another graph satisfying them.

5.7 *Language learning inside grammar and sequence doctrines*

Grammar doctrine. Gold’s identification-in-the-limit paradigm is an especially crisp UOCL specialization. The learner begins with a class of possible languages or grammars, receives a text or informant, and must eventually stabilize on the target language according to a declared equivalence [Gold, 1967]. Composition enters through concatenation, parsing, syntactic substitution, translation, or coalgebraic generation, depending on the grammar doctrine. The alphabet, grammar formalism, hypothesis enumeration, presentation type, and criterion of limiting stabilization are all fixed before the learner begins. Gold’s impossibility

results show precisely how identification changes when that doctrinal declaration changes.

Sequence-model doctrine. A Transformer supplies a powerful realizer for a different, usually weaker, query doctrine. Given token prefixes, it estimates conditional continuation laws and generates completions [Vaswani et al., 2017]. We may therefore place its hypotheses in a category $\mathbf{SeqMod}_\Sigma$ of parameterized sequence models and compare them by maps that preserve selected continuation behavior. Pretraining becomes online or batch assimilation of a large language presentation.

This doctrine assumes a tokenization, a category of finite prefixes and their extensions, an autoregressive factorization of sequence probability, a positional or equivariant treatment of order, and a shared parameterized conditional law. Likelihood optimization and attention provide the internal solution construction. Next-token and completion probabilities provide the dominant probes. These assumptions are extraordinarily broad compared with a finite grammar class, but they are still a doctrine rather than an absence of one.

This makes the user’s phrase “a category of possible languages” precise, but with an important caveat. The Transformer architecture does not by itself declare grammar morphisms, identify a unique language-generating coalgebra, or guarantee conservative persistence after parameter updates. It realizes a query-relative UOCL learner when those structures and its predictive quotient are supplied. Fluent generation demonstrates extraordinary competence on that quotient without settling whether the underlying compositional world has been explanatorily identified.

Doctrinal boundary. Changing weights while preserving the token, prefix, and continuation interfaces is internal learning. Revising the tokenizer, memory structure, grammar ontology, multimodal grounding maps, or the claim that continuation prediction is the adequate query family changes the doctrine. Scaling a realizer within the old declaration does not by itself perform that repair.

5.7.1 Krohn–Rhodes shortcuts and the identification gap

Liu et al. [2023] make the connection between Transformers and Krohn–Rhodes decomposition precise. Let $\mathcal{A} = (Q, \Sigma, \delta)$ be a finite semiautomaton. Reading a length T word normally appears to require the recurrent computation

$$q_t = \delta(q_{t-1}, \sigma_t), \quad 1 \leq t \leq T.$$

A shallow Transformer can instead represent each symbol by the induced transformation $\delta(-, \sigma_t) : Q \rightarrow Q$ and compose these transfor-

mations hierarchically. This replaces temporal iteration by a parallel computation over the transition semigroup.

Their first result gives every finite semiautomaton an exact Transformer simulator of depth $O(\log T)$ with resources polynomial in the sequence length and automaton size. Krohn–Rhodes theory identifies a broader constant-depth class. When every group appearing in the transition semigroup is solvable, the semiautomaton decomposes into reset memories and solvable group factors that admit constant-depth Transformer implementations in T , although the dependence on $|Q|$ may be large. Conversely, a constant-depth simulator for semiautomata containing non-solvable group structure would imply the unexpected circuit-complexity collapse $TC^0 = NC^1$. Thus prime algebraic structure predicts which recurrent computations possess shallow parallel shortcuts.

This result supplies a concrete mechanism for compositional performance. Attention can implement prefix aggregation and memory lookup; feedforward blocks implement transformations of the prime factors; residual connections carry information through their cascade. A Transformer need not store and update the visible automaton state in the manner of an RNN. It can reparameterize the global dynamics into a different but behaviorally equivalent computation.

The same result exposes the distinction central to ORACLE. The theorem is a simulation theorem: for a specified automaton and horizon, suitable Transformer parameters exist. It does not say that interaction identifies the transition semigroup, recovers a minimal Krohn–Rhodes cascade, or chooses the same factors under an equivalent presentation. The paper’s experiments show that gradient training can find shortcut solutions, but also that these solutions can be brittle under incomplete supervision, distribution shift, and extrapolation to unseen sequence lengths. Behavioral agreement on a training horizon therefore need not reveal a persistent compositional model.

Design principle

A Transformer shortcut is evidence that algebraic structure can compress execution. ORACLE additionally asks whether the learner discovered a query-sufficient factorization, whether that factorization is invariant under re-presentation, and whether its prime components persist and transport when the horizon, task, or evidence stream changes.

This suggests a diagnostic program rather than an architectural verdict. Given a trained sequence model, probe for the reversible group factors, irreversible reset factors, and cascade dependencies predicted

by the smallest observational automaton. Then change the presentation or extend the horizon and ask whether the same factors continue to explain behavior. A stable match would be evidence of categorical identification; a newly learned shortcut with the same input–output accuracy would witness performance without persistence.

5.8 JEPA, world models, and learned quotients

Fixed doctrine. World-model and joint-embedding predictive architectures replace direct reconstruction with prediction in a learned representation space [Ha and Schmidhuber, 2018, LeCun, 2022, Assran et al., 2023]. Their hypothesis world contains an encoder, a latent transition or predictor, and a family of readouts. Masked sensory fragments or action-conditioned trajectories provide the presentation. The query doctrine asks whether a latent future, relation, or planning-relevant feature can be predicted.

The doctrine presupposes that a useful quotient of observations exists, that the selected context determines the chosen target at the intended resolution, and that prediction in the latent category preserves the distinctions needed by downstream probes. Representation learning selects a realization of this quotient; it does not by itself prove that the quotient is sufficient for unanticipated decisions, interventions, or modalities.

Categorically, the encoder proposes a quotient: distinctions in observation space that no admitted latent query detects are intentionally identified. This is not a defect if the quotient is sufficient for the declared tasks. It becomes a UOCL failure when a later query, intervention, or modality needs a distinction that the representation destroyed and no conservative refinement can recover it. Accommodation then means changing the representation doctrine, not merely fitting the predictor more accurately.

Doctrinal audit. The structural doctrine contains encoders, latent predictors, readouts, and their admitted compositions; masking or action-conditioned sensory experience supplies the presentation; representation prediction supplies the solver; and latent or control-sufficient equivalence supplies the quotient. A newly important distinction erased by the encoder is evidence against the current doctrine of representation, not simply another prediction residual.

5.8.1 RoboDreamer and compositional robot imagination

RoboDreamer makes the compositional-generalization problem unusually visible in an embodied setting [Zhou et al., 2024]. A language

instruction is parsed into lower-level components such as an action phrase and a spatial relation phrase. Rather than condition one monolithic video model on the whole instruction, the method composes component-conditioned diffusion models. In the paper’s idealized density notation, if $L = \{l_1, \dots, l_N\}$ is the parsed instruction and τ a video trajectory, then

$$p_\theta(\tau \mid L) \propto \prod_{i=1}^N p_\theta(\tau \mid l_i)^{1/N}. \quad (5.9)$$

The learned score fields implement this product composition during diffusion sampling. New commands can therefore recombine previously learned components even when the complete instruction was absent from training.

The project demonstrations make the intended generalization concrete. They include tabletop commands such as “move green can near water bottle,” separate galleries for seen and unseen task combinations, multimodal goals specified by images or sketches, and RLBench comparisons between synthesized video plans and executions. The reported experiments use RT-1 demonstrations for video generation and evaluate whether generated plans align with both familiar and held-out instructions; the paper separately reports successful closed-loop robot execution in simulation. These distinctions matter: a plausible generated video, an executable plan, and successful physical-world control are three different query doctrines.

RoboDreamer is a revealing fixed-doctrine UOCL instance. Its hypothesis world is a family of factorized conditional video generators. Its presentation consists of language-labeled manipulation videos together with optional goal images or sketches. Its decisive probes are novel combinations of action, object, and spatial-relation components. Its observational quotient identifies models that produce equally task-aligned video futures on those probes. The approach thus instantiates a central ORACLE wager: an explicit compositional prior can turn sparse coverage of a combinatorial task space into useful predictions on unseen composites.

Its doctrine declares in advance that instructions admit the selected factor types and that the corresponding conditional score fields compose by the product rule. Training identifies the component realizations inside that declaration. Evidence that parsing is not stable, that component scores do not compose semantically, or that video equivalence fails to preserve executable action would challenge the doctrine rather than merely its fitted parameters.

The boundary is equally instructive. Equation (5.9) composes probability factors; it does not by itself identify a category of objects and actions, prove that parsing respects semantic composition, or guaran-

tee persistence when new objects, relations, or dynamics are learned. An ORACLE enrichment would register types for the parsed components, compare language composition with physical trajectory composition, test naturality across viewpoints and modalities, and localize failures among parsing, generation, inverse dynamics, and execution. RoboDreamer therefore provides strong evidence for structured compositional bias while leaving open the categorical identification and coherent-repair questions of this book.

5.8.2 *PoE-World and the categorical frontier*

PoE-World supplies a second, complementary sign that compositional world modeling has become an active empirical program [Piriyakulkij et al., 2025]. Instead of composing diffusion score fields for language-conditioned robot videos, it represents stochastic dynamics as an exponentially weighted product of small programmatic experts synthesized by a large language model. Each expert captures a restricted regularity—for example, how a particular object attribute changes under an action or contact relation—and weighted combination yields the transition model. The learned model is then embedded in a model-based planner. Experiments on Pong and Montezuma’s Revenge test whether a model induced from few observations can predict complex dynamics and transfer to levels containing novel arrangements of familiar entities and relations.

Doctrinal reading. PoE-World fixes a category of small programmatic experts, an exponentially weighted product operation, and a planner that consumes the resulting stochastic dynamics. The learner searches for experts and weights inside that architecture. ORACLE would additionally allow evidence to challenge the expert vocabulary, the product law, or the assumption that the induced transition model preserves the planning queries.

RoboDreamer and PoE-World implement different notions of composition, yet their common wager is important. A world model should not be a single opaque predictor when reusable parts of its dynamics can be discovered and recombined. Related work in robotics, computer vision, generative modeling, object-centric learning, neuro-symbolic modeling, and program synthesis explores many versions of that wager. This literature gives ORACLE both empirical motivation and a demanding test bed: compositional bias can improve data efficiency, interpretability, planning, and generalization beyond combinations seen in training.

The distinction is that composition in these systems is ordinarily an architectural or representational operation: multiply experts, add score fields, concatenate modules, bind object slots, or compose generated

trajectories. In the representative literature surveyed here, the learner is not asked to identify a category or sketch together with its objects, morphisms, equations, universal constructions, and observational quotient. Consequently, the architecture alone does not say when two learned factorizations express the same world; whether a recombination respects typed composition and coherence; which universal property licenses transport to a new context; or how a failed component can be repaired while preserving unrelated knowledge. These are precisely the additional obligations of UOCL.

This is a difference of foundation rather than a dismissal of empirical progress. ORACLE can treat learned experts, object slots, language factors, program fragments, and video generators as candidate presentations or realizers of categorical structure. It then asks whether their compositions support invariant queries and conservative revision. We therefore make the bounded claim that a categorical semantics is absent from the representative compositional world-model methods analyzed in this chapter, not the global historical claim that category theory has never appeared anywhere in the broader literature.

Design principle

Architectural compositionality specifies *how components are combined*. Categorical compositionality additionally specifies their types, the laws and coherences that combination must satisfy, the equivalences under which two world models count as the same, and the repairs that preserve already verified structure. ORACLE is aimed at this second problem while remaining able to use the first kind of model as an implementation.

5.8.3 COMBO: composition under decentralized observation

COMBO moves the same empirical program into embodied multi-agent planning [Zhang et al., 2025]. Its agents receive partial egocentric RGBD observations, yet must reason about the consequences of joint action in a shared world. A generative reconstruction module first estimates an overall world state. The framework then assigns three roles to vision–language models: an Action Proposer generates candidate actions, an Intent Tracker predicts what the other agents may do, and an Outcome Evaluator scores imagined states. A tree search combines these modules with a video world model to plan over longer cooperative horizons.

The compositional world model exploits the factorization of a joint action $a = (a_1, \dots, a_n)$. If X is a future video and x_0 the estimated

current state, its idealized product takes the form

$$P_{\theta}(X \mid x_0, a_1, \dots, a_n) \propto P_{\theta}(X) \prod_{i=1}^n \frac{P_{\theta}(X \mid x_0, a_i)}{P_{\theta}(X)}. \quad (5.10)$$

The corresponding diffusion scores are composed at inference time. This lets one learned dynamics model simulate joint actions for varying team sizes. The reported evaluation uses three embodied cooperation benchmarks with two to four agents, including cooking and visually specified puzzle tasks in ThreeDWorld.

COMBO is especially revealing through the Witsenhausen lens developed later in this book. Each agent has a local information field, the available actions depend on that information, and another agent's intention is latent rather than directly observed. The estimated global image is therefore not simply "the world." It is a comparison object assembled from partial views, temporal memory, and generative completion. Likewise, the Intent Tracker's output is a defeasible belief about another decision mechanism, not an observation of that mechanism.

An ORACLE formulation would keep these warrants typed. Egocentric observation categories would map into a candidate shared-world category; overlap maps would test whether local reconstructions agree; agent-indexed action objects would combine into a joint-action object; and rollout would ask for a coherent extension of those local data through time. When a predicted joint outcome fails, repair could then be localized to perception, cross-view gluing, intent inference, action composition, dynamics, or evaluation. COMBO currently implements these functions as a powerful planning architecture, but it does not require the comparison maps to satisfy a categorical universal property or certify which previously established conclusions survive a repair.

Doctrinal reading. COMBO fixes a decentralized observation doctrine, a factorization of joint actions, a generative comparison object for the shared state, and role-specific interfaces for proposing, tracking, and evaluating. Tree search solves the planning problem inside that package. A persistent inconsistency among local views may instead require revising the shared-state construction, the intent ontology, or the joint-action factorization itself.

Thus COMBO adds an important dimension to the distinction above. Compositional world modeling is not only about recombining objects or skills; in a team it must also reconcile multiple information histories and models of other agents. That is exactly where UOCL, Witsenhausen information structures, and the book's later theory of coherent repair can provide a foundation beyond architectural modularity.

The same analysis applies to multimodal foundation models. Cross-modal alignment supplies comparison maps between local presentation

categories; grounding requires those maps to commute with action and observation in the world, not merely to place paired samples near one another in a latent metric.

5.9 *What the fixed-doctrine analysis establishes*

Table 5.1 reveals a common pattern, but UOCL is not valuable merely because it can rename familiar components. Its contribution begins by turning the older paradigms' implicit doctrines into explicit, comparable declarations:

1. It types the hypothesis category and its morphisms rather than naming only a parameter space.
2. It separates the world from the presentation through which the world is revealed.
3. It makes the query-relative quotient explicit, preventing prediction, control, causal, and explanatory identification from being conflated.
4. It distinguishes assimilation inside a doctrine from accommodation of the doctrine itself.
5. It asks whether learned answers persist under revision and transport, not merely whether instantaneous error becomes small.
6. It separates an internal solver—Bellman iteration, conditioning, likelihood optimization, final semantics, descent, or planning—from the doctrine that guarantees the solver is meaningful.

Design principle

An established learner is a special case of UOCL only after its hypothesis doctrine, presentation, queries, solution construction, equivalence, update, and convergence notion have been declared. Its theorem proves success *relative to that declaration*. It does not prove that the declaration is the uniquely correct account of the world.

This chapter also reverses the comparison. Automata, MDPs, causal models, grammars, PSRs, and latent world models are not competitors to ORACLE. They are candidate doctrines, often extraordinarily successful ones. Their theorems demonstrate how much becomes learnable once the right world class and solver are fixed. Their systematic failures reveal possible doctrinal boundaries: the world may not be finite-state, Markovian, causally sufficient, autoregressive at the chosen representation, globally gluable, or adequately captured by the current latent quotient.

ORACLE therefore adds a new outer loop. The inner learner estimates a model, predictive state, value function, causal graph, grammar, or representation inside $\mathfrak{D}_t$. The outer learner compares the residual obstructions against alternative doctrines and, when warranted, transports settled structure along

$$\mathfrak{D}_t \longrightarrow \mathfrak{D}_{t+1}.$$

The difficult problem is not merely choosing a larger hypothesis class. It is declaring which evidence licenses that move and which earlier conclusions survive it.

Guiding question

For each established paradigm, what is the weakest structural doctrine, smallest separating query doctrine, and minimal internal solver that recover its successful predictions? Which observable obstruction would justify leaving that doctrine while preserving the distinctions needed for later decision, causal repair, and cross-task transport?

Further reading

Gold supplies the classical language-identification template [Gold, 1965, 1967]. Angluin gives the canonical active-learning result for regular languages from membership and equivalence queries [Angluin, 1987]; Rivest and Schapire replace global-state enumeration by a diversity representation built from equivalence classes of tests [Rivest and Schapire, 1994]. Krohn and Rhodes give the prime decomposition of finite machines [Krohn and Rhodes, 1965]; Ronca et al. [2023] develop its modern sample-complexity consequences for learnable cascades; and categorical extensions are given by Wells [1980, 1988], Thérien [1991]. The stochastic direction begins with the support-semigroup decomposition of Carlsson and Yu [2015]; Plotkin and Plotkin [2015] provide a related decomposition theory for linear automata using triangular and wreath products. Liu et al. [2023] connect Krohn–Rhodes factors to shallow Transformer simulation and document both learnable shortcuts and their generalization failures. Rutten develops the coalgebraic account of automata, behavior, bisimulation, and coinduction [Rutten, 2000]. Predictive state representations isolate prediction from hidden-state ontology and control [Littman et al., 2001]. Puterman and Sutton–Barto develop the MDP and reinforcement-learning frameworks [Puterman, 1994, Sutton and Barto, 2018]; the options framework promotes closed-loop behaviors into temporally extended actions with explicit initiation and termination interfaces [Sutton et al., 1999]. Pearl and Peters, Janzing, and Schölkopf develop causal models and causal discovery [Pearl,

2009, Peters et al., 2017]. Transformer sequence modeling is introduced by Vaswani et al. [2017]; world models and joint-embedding predictive approaches are represented by Ha and Schmidhuber [2018], LeCun [2022], Assran et al. [2023]. RoboDreamer provides an embodied example in which language components and multimodal goals compose conditional video models for unseen robot tasks [Zhou et al., 2024]; PoE-World composes LLM-synthesized programmatic experts into a stochastic world model and tests few-observation planning and generalization in Atari environments [Piriyakulkij et al., 2025]; and COMBO combines partial-view state estimation, agent-wise diffusion factors, VLM-based intent and action modules, and tree search for embodied cooperation [Zhang et al., 2025]. The next two chapters turn this comparison into the explicit UOCL machine and its persistence theory.

6

The UOCL Machine

EFFECTIVE IDENTIFICATION BY ASSIMILATION, PROBING, AND REPAIR

ORACLE states a semantic problem: identify an unknown compositional world from a presentation, up to the equivalence respected by an admissible query language. It does not yet say how a learner computes its next hypothesis, chooses an informative interaction, or repairs a declaration that no longer fits the evidence. Universal Online Categorical Learning (UOCL) supplies this missing algorithmic layer.

The semantic-to-algorithmic inclusion is strict in intent:

ORACLE program $\supset$ UOCL algorithms.

An ORACLE section can exist without an effective procedure for constructing it. A UOCL learner realizes such a section by computable or otherwise specified update, probe, and repair rules. Tangent UOCL and UODL then form two independent enrichments: one learns infinitesimal variation, while the other adds consequential action. Their intersection supports DIAL.

6.1 From a semantic section to a learning machine

Let **Doc** be a category of admissible categorical doctrines and let

$$p : \int H_{\text{doc}} \longrightarrow \mathbf{Doc}$$

be the hypothesis fibration of Chapter 3. For a doctrine $\mathbb{T}$, the fiber $H_{\text{doc}}(\mathbb{T})$ contains its admissible realized worlds. A presentation prefix determines an extension $u_t : \mathbb{T}_t \rightarrow \mathbb{T}_{t+1}$, and reindexing gives $u_t^* : H_{\text{doc}}(\mathbb{T}_{t+1}) \rightarrow H_{\text{doc}}(\mathbb{T}_t)$.

The semantic ORACLE formulation asks for a coherent sequence of objects in these fibers. An algorithm must also resolve three choices hidden by that statement: which lift to select when several fit, which probe to issue when the presentation is active, and what to change when no acceptable lift exists.

Definition 6.1 (UOCL instance). A *UOCL instance* consists of:

1. a presentation category **Pres** and a doctrine map $d : \mathbf{Pres} \rightarrow \mathbf{Doc}$;
2. the pulled-back hypothesis fibration $p_d : \int H_d \rightarrow \mathbf{Pres}$;
3. a query indexed category $\mathcal{Q}$ and a declared observational equivalence $\simeq_{\mathcal{Q}}$ in each hypothesis fiber;
4. a category **Probe**(T) of admissible probes at each presentation prefix T , together with response extensions; and
5. admissibility predicates for conservative lifts and for changes of doctrine.

This data says what counts as evidence, a hypothesis, a successful answer, an interaction, and a legitimate repair. It still does not choose among them.

Definition 6.2 (Lift fiber). For $u : T \rightarrow T'$ and $M \in H_d(T)$, define

$$\text{Lift}_u(M) = \{(M', \alpha) \mid M' \in H_d(T'), \alpha : u^* M' \rightarrow M \text{ is admissible}\}.$$

An element is *conservative* on a settled query family $\mathcal{Q}^{\text{set}}$ when every component of α observed by that family is an equivalence.

6.2 State, probes, and dependent update

A useful machine state must retain more than a current model. We write

$$S_t = (T_t, \mathbb{T}_t, \widehat{M}_t, \mathcal{Q}_t^{\text{set}}, \mathcal{F}_t, \Delta_t),$$

where T_t is the transcript, $\mathbb{T}_t = d(T_t)$ is the current doctrine, $\widehat{M}_t$ is the selected realization, $\mathcal{Q}_t^{\text{set}}$ records answers the learner has warranted, $\mathcal{F}_t$ is the residual version fiber, and Δ_t is the registered defect or repair context. This is a typed state: changing T_t changes the fiber in which the next state lives.

Definition 6.3 (UOCL algorithm). A *UOCL algorithm* for an instance consists of three partial realizers:

$$\begin{aligned} \Pi_t &: S_t \longrightarrow \mathbf{Probe}(T_t), \\ \text{Sel}_u &: (S_t, u) \dashrightarrow \text{Lift}_u(\widehat{M}_t), \\ \text{Rep} &: (S_t, u, \Delta) \dashrightarrow (a : \mathbb{T}_{t+1} \rightarrow \mathbb{T}'_{t+1}, S'_{t+1}). \end{aligned}$$

Here Π_t is a probe policy, Sel chooses an admissible assimilation lift, and Rep performs ordinary or doctrinal accommodation when assimilation is unavailable or inadequate. The realizers may be deterministic, randomized, or set-valued, but their type and admissibility obligations must be explicit.

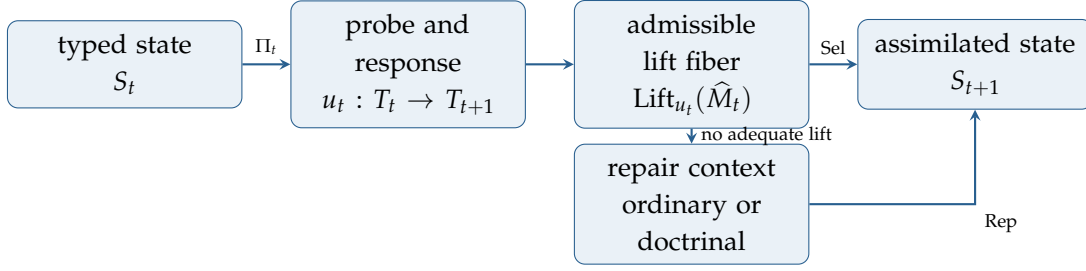


Figure 6.1: One UOCL step. Assimilation selects a conservative lift in the current doctrine. A registered obstruction redirects the update through accommodation; abstention remains available when neither route is warranted.

The partiality is essential. A learner may abstain because the evidence is insufficient, because its selection problem is not effectively solvable, or because a registered obstruction requires accommodation. Treating all three cases as an arbitrary guessed model would erase the distinction between ignorance and inconsistency.

Proposition 6.4 (One-step query soundness). *Suppose $(M_{t+1}, \alpha_t) \in \text{Lift}_{u_t}(M_t)$ is selected and α_t is conservative on $\mathcal{Q}_t^{\text{set}}$. If M_{t+1} satisfies the extended transcript, then the update preserves every settled answer and is sound for the new evidence.*

Proof. Soundness for the new evidence is the defining membership condition for $H_d(T_{t+1})$. Reindexing compares the new realization with the old one, and conservativity makes that comparison an equivalence under every settled query. Hence no previously warranted answer changes. $\square$

6.3 Assimilation and accommodation as control flow

The Piagetian distinction becomes an operational branch rather than a metaphor. Given a presentation arrow u_t , the learner first forms the lift fiber. If it contains an adequate conservative lift, selection is assimilation. If the fiber is nonempty but the chosen realization must change in an unsettled region, the step is ordinary accommodation. If a certified defect shows that the present doctrine admits no adequate lift, repair moves along a doctrine map and is doctrinal accommodation.

One typed UOCL update

1. Choose an admissible probe $p_t \leftarrow \Pi_t(S_t)$, or receive the next datum, and extend the transcript by its response $u_t : T_t \rightarrow T_{t+1}$.
2. Construct or approximate $\text{Lift}_{u_t}(\hat{M}_t)$.
3. If an adequate conservative lift can be selected, return the assimilated state and its comparison map.
4. Otherwise, if a registered obstruction licenses repair, apply Rep

and return an accommodated state.

5. If neither transition is warranted, abstain and retain the residual fiber and defect certificate.

The algorithm schema does not prescribe enumeration, gradient descent, proof search, or Bayesian updating. Those are possible realizers of its three categorical interfaces. UOCL is universal at the level of the typed learning problem, not a claim that one numerical routine solves every instance.

6.4 Active identification and separating probes

Let $\pi_0 \mathcal{F}_t$ denote the current observational-equivalence classes. A probe p is *separating* for $M, N \in \mathcal{F}_t$ if some possible response gives different query answers for the two classes. A probe policy is *fairly separating* when it eventually issues such a probe for every pair that remains observationally distinct and viable.

This makes active learning categorical: the learner does not merely seek a large scalar information gain. It chooses a morphism that refines a hypothesis fiber relative to the declared queries. Numerical utilities may rank such probes, but they are observers of the refinement, not its definition.

6.4.1 Running example: enlarging the collider probe family

At time t , a collider learner has interrogated only a finite subcategory $i_t : \mathcal{P}_t \hookrightarrow \mathcal{P}$ of the declaration $\mathfrak{C}_{\text{coll}}$ from Section 2.1.3. Its current hypothesis object is therefore the observational quotient

$$\mathbf{H}_{\text{phys}} / \simeq_{\mathcal{P}_t}.$$

Choosing the next beam energy or decay channel is an active UOCL step when it separates surviving classes rather than merely adding more events from an already redundant probe.

Proposition 6.5 (Probe enlargement refines collider identification). *If $\mathcal{P}_t \hookrightarrow \mathcal{P}_{t+1}$ is an inclusion of exact probe doctrines, then*

$$H \simeq_{\mathcal{P}_{t+1}} H' \implies H \simeq_{\mathcal{P}_t} H'.$$

Consequently the later observational quotient refines the earlier one: exact probe enlargement can split an equivalence class but cannot merge two that were already distinguished.

Proof. Restrict the equivalence of response functors along $\mathcal{P}_t \hookrightarrow \mathcal{P}_{t+1}$. Functorial restriction preserves the declared equivalence. $\square$

This monotonicity is intentionally idealized. With finite samples, changing calibrations, or approximate tests, an empirical quotient can fluctuate. Chapter 8 replaces exact response equivalence by a risk bound, while Chapter 14 records which layer should be repaired when a new channel conflicts with the current theory.

6.5 Learning generators, relations, and models

The basic UOCL target can be enriched once more. A learner may identify not only a category of encountered objects and arrows, but a compact algebraic theory explaining how that category is generated. This is stronger than memorizing a growing composition table and different from learning tangent structure. It asks for reusable syntax, equations, and functorial semantics.

Definition 6.6 (Theory-bearing UOCL hypothesis). A *theory-bearing UOCL hypothesis* is a tuple

$$\mathfrak{H} = (\mathcal{S}, \mathcal{D}, \mathbb{A}, M, \mathcal{C}), \quad \mathbb{A} = \text{Th}_{\mathcal{D}}(\mathcal{S}), \quad M : \mathbb{A} \longrightarrow \mathcal{C},$$

where $\mathcal{S}$ is a sketch, $\mathcal{D}$ is its preservation doctrine, $\mathbb{A}$ is the presented algebraic theory, $\mathcal{C}$ is a candidate world category, and M is a structure-preserving interpretation. Its queries may address the world, the presentation, the completed theory, the interpretation, or comparison maps among them.

The added levels separate four questions:

1. Which configurations and transformations have been encountered?
2. Which generators suffice to express them?
3. Which relations identify apparently different construction paths?
4. In which semantic worlds do those formal operations have valid models?

For the LEGO learner of Section 1.5, a newly observed tower can be assimilated by factoring its construction through known joining and tensor operations. Discovering that a separation reverses an attachment adds an equation, or perhaps a partial-inverse domain, to the sketch. That is theory-level accommodation.

Definition 6.7 (Generative adequacy). Let $\mathcal{O}_t \subseteq \mathcal{C}$ be the observation subcategory generated by the transcript at time t . A theory-bearing hypothesis is *generatively adequate at t* when every declared object and arrow of $\mathcal{O}_t$ is in the essential image of the interpretation of a well-typed expression in $\mathbb{A}$, and every equality registered by the transcript

is respected by M . It is *factorization-complete* for a query doctrine when every queried future composite in the doctrine admits such an expression.

Generative adequacy is presentation-relative and deliberately one-sided. It says the learned operations can express what has been observed; it does not say the inferred relations are complete, the generators are minimal, or the physical interpretation is faithful. Those are separate probes. A compact but false theory may generate every training example by equating too many paths.

Rivest–Schapire diversity representations provide a finite-state instance of this principle (Section 5.4.1). Their compact object is generated by action-transformed tests modulo observational equivalence, even when the associated global state space is enormous. This suggests a useful UOCL design bias: learn a presentation of the probes and their compositional update laws before attempting to enumerate the objects they distinguish. Algebraic theories become operationally valuable here because they generate future experiments, not merely because they compress past observations.

Particle physics supplies a richer instance (Section 0.8). There the learned presentation is not an enumeration of physical states: particle types, interaction vertices, symmetries, and relations generate response predictions for entire families of collider probes. The experimental response functor in the declaration $\mathcal{C}_{\text{coll}}$ is the semantic model against which the theory is assimilated or repaired. Proposition 6.5 then says how exact evidence narrows the observational quotient without pretending that a finite probe family determines a unique presentation.

Proposition 6.8 (Semantic presentation invariance). *Let $F : \mathcal{S} \rightarrow \mathcal{S}'$ be a map of sketches that induces, naturally over every registered semantic category $\mathcal{C}$, an equivalence*

$$F^* : \text{Mod}_{\mathfrak{D}}(\mathcal{S}', \mathcal{C}) \simeq \text{Mod}_{\mathfrak{D}}(\mathcal{S}, \mathcal{C}).$$

Then no UOCL query that factors through these model categories can distinguish $\mathcal{S}$ from $\mathcal{S}'$.

Proof. Each admitted query sends equivalent model objects and their morphisms to equivalent answers. Naturality makes the comparison stable under change of registered semantic category. Hence both presentations lie in the same observational class for every query factoring through model semantics. $\square$

This proposition supplies the right success criterion: learn a theory up to the equivalence visible in its model semantics, not a privileged string of generators. It also exposes a new nonidentifiability result. If two nonequivalent presentations have indistinguishable models under

all available semantic probes, interaction cannot choose between them without enlarging the doctrine.

Design principle

Prefer the smallest warranted generating theory to an inventory of observed composites, but treat minimality, completeness of relations, and semantic faithfulness as separate claims. A compact presentation is an explanation only relative to its model and query doctrine.

6.6 Executions realize ORACLE sections

Proposition 6.9 (UOCL realization). *Every well-typed execution of a UOCL algorithm determines a section of the hypothesis fibration along its realized presentation path, together with comparison maps satisfying the algorithm's declared coherence laws. Conversely, an abstract ORACLE section is realized by a UOCL algorithm only after effective or explicitly specified selection, probing, and repair data have been supplied.*

Proof. At time t , the execution supplies an object $M_t \in H_d(T_t)$. Each successful assimilation supplies a comparison $u_t^* M_{t+1} \rightarrow M_t$; a repair step supplies the corresponding comparison after change of base. The typing and coherence conditions therefore give a section along the path. The converse fails without additional data because existence of a lift does not provide a method for finding one, and existence of a separating probe or repair does not choose it uniformly. $\square$

This proposition locates a basic computability boundary. Universal properties can characterize a correct answer while equality, equivalence, extension, or selection in the relevant category remains undecidable. UOCL must therefore declare its representation of categories and its effective fragment; it cannot obtain an algorithm merely by appealing to Yoneda or Kan extension.

6.7 Composing UOCL learners

An infant does not identify one compositional world through one homogeneous stream. Visual tracking, speech, manipulation, navigation, and social interaction each induce a partial learner with its own presentation and query doctrine. Development depends on their combination: a heard noun is aligned with a seen object; grasping tests visual solidity; another person's gaze selects both a spatial target and a linguistic referent. UOCL should therefore compose learners, not merely compose the hypotheses produced by one learner.

Fong, Spivak, and Tuyéras provide the useful prototype. Their supervised learner from A to B is a tuple (P, I, U, r) : a parameter space, implementation, update, and request map. Equivalence classes of such learners form a symmetric monoidal category, so systems can be connected sequentially or placed in parallel [Fong et al., 2019]. The request map is essential: it translates a desired downstream correction into information that an upstream learner can use.

For UOCL the interfaces and internal state must be categorified. Write a *categorical learning interface* as

$$\mathbb{A} = (\mathbf{Pres}_A, \mathcal{Q}_A, \simeq_A),$$

consisting of a typed presentation protocol, a query doctrine, and its observational equivalence. The following definition is deliberately strict; the coherent weakening follows immediately afterward.

Definition 6.10 (Open UOCL learner). An *open UOCL learner* $L : \mathbb{A} \rightsquigarrow \mathbb{B}$ is a tuple

$$L = (\mathcal{H}_L, I_L, U_L, R_L),$$

where:

1. $\mathcal{H}_L$ is an indexed category of internal hypotheses and residual version fibers over $\mathbf{Pres}_A$;
2. I_L implements a current hypothesis on an A -presentation as a partial B -declaration with warranted query answers;
3. U_L is a typed UOCL update, returning an admissible comparison map after new B -evidence or a registered defect; and
4. R_L translates a downstream B -probe, contradiction, or unfilled obligation into an admissible upstream A -probe or repair obligation.

Two open learners are equivalent when an indexed equivalence of their hypothesis categories intertwines I, U, R and preserves their declared query equivalences.

Thus $\mathcal{H}_L$, rather than a set P , is the learner's internal state space. The implementation is a categorical declaration rather than only a function $A \rightarrow B$. Update may assimilate or accommodate. Most importantly, R_L generalizes the request function: it carries a failure of fit backward through a composite without reducing that failure to a numerical gradient.

Proposition 6.11 (Strict compositionality of open UOCL). *Suppose the admitted indexed hypothesis categories have chosen products, reindexing preserves them, and the class of realizers is closed under typed wiring. Then categorical learning interfaces and equivalence classes of open UOCL learners form a category $\mathbf{UOCL}_{\text{str}}$. Pointwise parallel composition equips it with a symmetric monoidal product.*

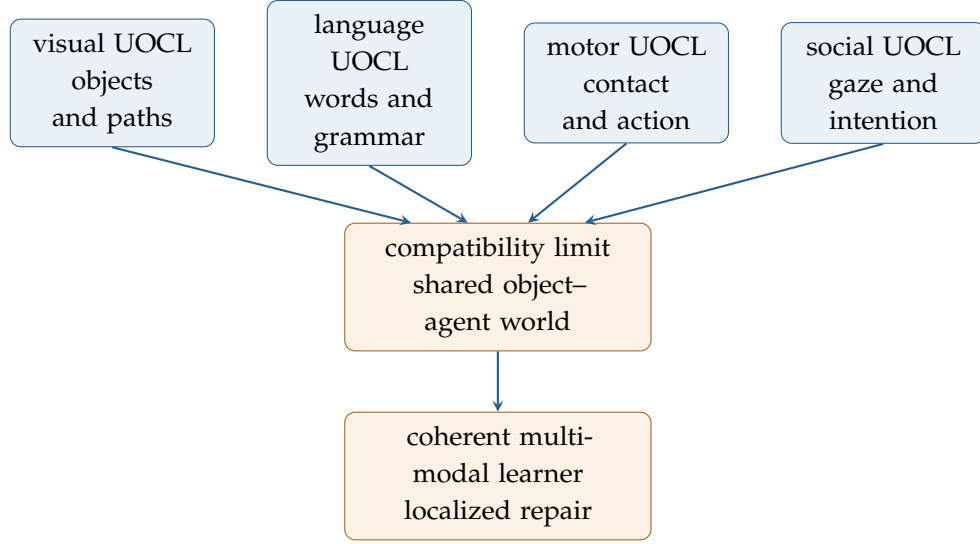


Figure 6.2: Multimodal UOCL is more than parallel execution. Independent learners supply a monoidal product; observers into a shared object-agent doctrine impose a compatibility limit. A failed comparison identifies which modality or cross-modal alignment requires accommodation.

Proof. For $L : \mathbb{A} \rightsquigarrow \mathbb{B}$ and $K : \mathbb{B} \rightsquigarrow \mathbb{C}$, take the composite hypothesis category to be the indexed product $\mathcal{H}_L \times \mathcal{H}_K$. Compose the implementations forward. A C -obligation is first translated by R_K to a B -obligation and then by R_L to an A -obligation; the two updates use these translated obligations and retain their comparison maps. The identity interface has the terminal hypothesis category and simply forwards declarations and obligations. Associativity and unit laws follow from those of wiring and the chosen products, modulo the stated learner equivalence. Taking products of interfaces, hypothesis categories, and realizers gives parallel composition; the product associator, unitors, and swap supply the symmetric monoidal laws. $\square$

Without chosen products the same construction is generally associative only up to coherent equivalence. Doctrinal accommodation may also change the interface itself. The natural home is then a bicategory, double category, or equipment of learners, interface changes, and comparison 2-cells. The strict category is the smallest testable fragment, not the final higher-categorical claim.

6.8 Limits, colimits, and multimodal fusion

The symmetric monoidal product places learners side by side but imposes no agreement. Genuine multimodal fusion begins when local hypotheses are observed in a shared doctrine. For a visual learner and a linguistic learner, suppose there are comparison functors

$$O_V : \mathcal{H}_V \longrightarrow \mathcal{K}, \quad O_L : \mathcal{H}_L \longrightarrow \mathcal{K},$$

where $\mathcal{K}$ expresses shared claims such as object identity, persistence, agency, or reference. Their compatible joint state is the indexed pullback

$$\mathcal{H}_{V \bowtie L} = \mathcal{H}_V \times_{\mathcal{K}} \mathcal{H}_L. \quad (6.1)$$

Proposition 6.12 (Universal multimodal fusion). *If the pullback in Equation 6.1 exists in every presentation fiber, is preserved by reindexing, and the two updates preserve the compatibility comparison, then it defines a fused UOCL learner. Every joint learner equipped with compatible maps to the visual and linguistic learners factors uniquely through it, up to the declared equivalence.*

Proof. Fiberwise pullbacks provide compatible joint hypotheses and their projection maps. Preservation by reindexing makes these fibers an indexed category; update preservation closes it under online learning steps. The factorization and its uniqueness are exactly the pullback universal property, interpreted fiberwise and then assembled by reindexing coherence. $\square$

This universal property distinguishes fusion from concatenating embeddings. It gives a minimal joint learner containing both local hypotheses subject to the declared cross-modal agreements. A contradiction can then be localized: the visual hypothesis, the linguistic hypothesis, or their comparison map may need repair. An undifferentiated joint parameter vector provides no analogous structural diagnosis.

Other universal constructions have different meanings. Products support parallel observation and joint queries. Equalizers enforce agreement between two proposed translations. Coproducts retain tagged alternative learners. Pushouts glue learned worlds along an already identified common subworld, and more general colimits assemble a world from overlapping fragments. None is an automatic operation: existence, effective construction, preservation of warranted answers, and avoidance of spurious identifications must be proved for the selected UOCL doctrine.

The tangent refinement adds one decisive compatibility test. If the tangent functor preserves the pullback used for fusion, then

$$T(\mathcal{H}_V \times_{\mathcal{K}} \mathcal{H}_L) \simeq T\mathcal{H}_V \times_{TK} T\mathcal{H}_L.$$

In that case, fusing first and learning infinitesimal variation commute. If the comparison fails, cross-modal perturbations contain information absent from the separate tangent learners; the fusion mechanism itself must be learned or accommodated. This supplies a concrete way in which the category of UOCL learners creates the conditions required by DIAL.

6.9 Failure modes

Four failures should remain distinct:

1. *underdetermination*: multiple query-inequivalent hypotheses fit the available transcript;
2. *ineffectivity*: an admissible lift exists but the learner lacks a realizer that can find it;
3. *realization defect*: the doctrine is adequate, but the current model or gluing data must change; and
4. *doctrinal obstruction*: no adequate realization exists in the current categorical prior.

Only the fourth licenses revision of the language of possible worlds. This separation will later prevent a failed decision optimizer from being misdiagnosed as evidence against the learned world itself.

Further reading

Gold’s identification-in-the-limit paradigm supplies the presentation and convergence discipline underlying UOCL [Gold, 1967]. The Grothendieck construction and indexed-category viewpoint are reviewed in Chapter 1. Lawvere’s functorial semantics, Ehresmann sketches, and PROPs supply the theory-bearing refinement [Lawvere, 1963, Ehresmann, 1968, Lack, 2004]. Piaget motivates the operational separation of assimilation and accommodation [Piaget, 1952], while universal imitation games supply the coinductive special case in which behavioral prediction, rather than categorical reconstruction, is the target [Mahadevan, 2024]. The category Learn , the (P, I, U, r) interface, and the symmetric monoidal composition of supervised learners are developed by Fong et al. [2019]; their construction is the direct prototype for the open UOCL interface introduced here.

7

Persistent Categorical Identification

COHERENT STABILIZATION UNDER CHANGING PRESENTATIONS

An online categorical learner should not merely fit every finite prefix. It should retain warranted structure as the presentation grows, revise locally when possible, and give equivalent accounts of equivalent evidence streams. Persistence is therefore a property of an entire execution and its comparison maps, not of the final hypothesis alone.

7.1 Non-anticipation and two forms of stabilization

Definition 7.1 (Non-anticipatory UOCL). A UOCL algorithm is *non-anticipatory* when its state, probe, selection, and repair at time t factor through the presentation prefix T_t . Two streams with isomorphic prefixes induce equivalent states before their first differing response.

The definition is the categorical analogue of adaptedness. It does not require a single global clock: a prefix may be a finite down-set in a partially ordered event category.

Definition 7.2 (Behavioral and explanatory stabilization). Let q be an admissible query. An execution *behaviorally stabilizes* on q when its answer is eventually constant up to the declared answer equivalence. It *explanatorily stabilizes* when, in addition, the hypothesis objects and comparison maps supporting that answer eventually lie in one coherent equivalence class.

Behavioral stabilization is analogous to predictive-state identification: the learner can answer the relevant tests without recovering a unique internal presentation. Explanatory stabilization is stronger and is needed when later transport or repair depends on the discovered construction rather than only its observed output.

7.2 Coherent persistence

For a chain $T_0 \rightarrow T_1 \rightarrow \dots$, write $\alpha_t : u_t^* M_{t+1} \rightarrow M_t$. Persistence requires the comparison for a composite presentation extension to agree with the composite of the local comparisons, strictly or up to a declared coherent equivalence.

Proposition 7.3 (Iterated conservative persistence). *If every α_t is conservative on a query family $\mathcal{Q}_0$, and conservativity is closed under reindexing and composition, then every answer in $\mathcal{Q}_0$ is preserved along every finite continuation of the execution.*

Proof. The comparison from M_n to M_0 is the reindexed composite of the α_t . Closure under reindexing and composition makes this composite an equivalence under each observer in $\mathcal{Q}_0$. $\square$

The proposition is modest but foundational. It states the categorical invariant that an implementation must maintain. Approximate versions need an observer-valued defect calculus and a composition law controlling how errors accumulate.

7.3 Finite active identification

Theorem 7.4 (Finite UOCL identification bound). *Suppose the initial residual fiber has $N < \infty$ observational-equivalence classes, the true class is never eliminated, probe responses are sound, and the policy selects a probe that eliminates at least one false class whenever more than one remains. Then UOCL behaviorally identifies the target after at most $N - 1$ separating probe responses.*

Proof. Each separating response strictly decreases the finite number of viable false equivalence classes and never deletes the true one. At most $N - 1$ such decreases leave a single class, which answers every admissible query as the target does. $\square$

The theorem does not promise explanatory identification. Nor does it cover infinite fibers without a topology, rank, measure, or compactness condition. It shows exactly where the work moves in harder settings: one must prove separation, sound elimination, and a well-founded decrease.

7.4 Presentation invariance

The same world may be revealed in different orders, by different generators, or at different granularities. A persistent learner should not confuse these encodings with different worlds.

Definition 7.5 (UOCL presentation invariance). Let $F : \mathbf{Pres} \simeq \mathbf{Pres}'$ be an equivalence of presentation protocols and let the induced hypothesis fibrations, query families, and probe categories be pseudonaturally equivalent. A UOCL algorithm is *presentation invariant* along F when its probe, selection, and repair rules are equivariant under these equivalences and the resulting executions correspond up to coherent observational equivalence.

Proposition 7.6 (Transport of UOCL executions). *Under the hypotheses of the definition, every execution over $\mathbf{Pres}$ transports to an observationally equivalent execution over $\mathbf{Pres}'$. Behavioral stabilization and conservative persistence are preserved.*

Proof. Transport each state and comparison through the pseudonatural equivalence. Equivariance transports the selected probes, lifts, and repairs, while coherence identifies transported composites. The query equivalence preserves answers and therefore preserves both stabilization and conservativity. $\square$

This is the online form of Kan invariance at the algorithmic level. Kan invariance of isolated constructions is necessary but not sufficient: the selection and repair mechanisms must also respect the change of presentation.

7.5 Tangent UOCL: learning how a world can vary

Ordinary UOCL identifies a category relative to base queries. Tangent UOCL (TUOCL) takes its hypotheses in **TanCat**, or in a declared category of models carrying tangent structure. Its transcript may therefore contain both base evidence and *tangent evidence*: observations of infinitesimal probes, their projections, sums, lifts, and iterated interchange.

Definition 7.7 (Tangent UOCL instance). A *tangent UOCL instance* is a UOCL instance whose hypothesis fibration

$$p_{\text{tan}} : \int H_{\text{tan}} \longrightarrow \mathbf{Pres}$$

has tangent-category hypotheses, whose comparison maps are strong or explicitly lax tangent morphisms, and whose query doctrine decomposes as $\mathcal{Q}^{\text{base}} \cup \mathcal{Q}^{\text{tan}}$. The tangent queries test the structure maps $(T, p, 0, +, \ell, c)$ and the coherence laws on which later infinitesimal calculations depend.

For a comparison $F : \mathcal{C} \rightarrow \mathcal{D}$, tangent compatibility includes a coherent comparison

$$\tau_F : FT_{\mathcal{C}} \Longrightarrow T_{\mathcal{D}}F,$$

invertible in the strong case, satisfying the projection, zero, addition, lift, and flip axioms. Thus tangent assimilation preserves not only a learned composite but also the admissible variations of that composite. Tangent accommodation repairs T or one of its coherence maps when new infinitesimal evidence cannot be absorbed by the current tangent realization.

Proposition 7.8 (Base evidence cannot identify tangent structure). *Suppose two tangent hypotheses $\mathbb{C}_1^{\text{tan}}$ and $\mathbb{C}_2^{\text{tan}}$ have equivalent underlying categories but are not equivalent in **TanCat**. Any presentation and query doctrine that factors through U_{tan} leaves them observationally indistinguishable. Hence categorical identification of the base does not imply tangent identification.*

Proof. Every admitted observation depends only on the common image under U_{tan} , so corresponding base queries have equivalent answers. Their disagreement lies entirely in structure forgotten by that functor and cannot be separated without a tangent-sensitive probe. $\square$

The proposition is the tangent analogue of the finite-transcript obstruction. It explains why differentiating a learned model after the fact is not the same as learning the world’s infinitesimal geometry: automatic differentiation supplies a tangent mechanism for a presentation, but does not establish that the mechanism is identified or invariant under an equivalent presentation.

7.5.1 Differential presentations of tangent worlds

TUOCL need not assume that its tangent hypotheses are primitive. Let $\mathbf{DiffCat}_{\text{adm}}$ be a declared class of differential categories satisfying the hypotheses needed for the coalgebra construction. Write

$$\text{Coalg}_! : \mathbf{DiffCat}_{\text{adm}} \longrightarrow \mathbf{TanCat}$$

schematically for the assignment sending a differential category to its coEilenberg–Moore category of coalgebras with the induced tangent structure [Cockett et al., 2020]. A differential-realized TUOCL instance factors its tangent hypotheses through this assignment. Its tangent queries observe $\text{Coalg}_!(\mathbb{X})$; stronger differential queries may also test the modality $!$, the deriving transformation d , and their coherence.

Proposition 7.9 (Tangent evidence need not identify a differential realization). *Suppose admissible differential hypotheses $\mathbb{X}_1$ and $\mathbb{X}_2$ are inequivalent as differential categories, while $\text{Coalg}_!(\mathbb{X}_1)$ and $\text{Coalg}_!(\mathbb{X}_2)$ are equivalent for every admitted tangent query. Any transcript and query doctrine factoring through $\text{Coalg}_!$ leaves the two differential realizations observationally indistinguishable.*

Proof. Every admitted answer is computed after applying $\text{Coalg}_\dagger$. By hypothesis the resulting tangent models agree on the entire tangent query doctrine, so no such answer can separate their preimages. A separating observation must inspect differential-presentation structure not retained by the realized tangent world. $\square$

The result is deliberately conditional: it does not assert that such pairs exist for every chosen doctrine. It identifies the exact obstruction whenever the coalgebra realization fails to be query-faithful. Accordingly, learning admits a three-rung ladder:

base category $\rightsquigarrow$ tangent geometry $\rightsquigarrow$ differential realization.

Each step requires new probes and a new persistence claim. DIAL needs the middle rung to type infinitesimal defects. A learned differentiation engine or resource semantics may require the third; ordinary predictive success guarantees neither.

Proposition 7.10 (Tangent conservative persistence). *Let a TUOCL execution have strong tangent comparison maps whose underlying base comparisons are conservative on $\mathcal{Q}_0^{\text{base}}$, and whose tangent comparison cells are conservative on $\mathcal{Q}_0^{\text{tan}}$. If both classes are closed under reindexing and composition, then all settled base and tangent answers persist along every finite continuation.*

Proof. Apply Proposition 7.3 to the base comparisons. Strong tangent coherence identifies the tangent comparison for a composite with the composite of the reindexed tangent cells. Closure then makes every settled tangent observer invert that composite as well. $\square$

This yields a sharper convergence ladder. A learner may stabilize behaviorally on base queries while its inferred tangent structure continues to change. *Tangent stabilization* requires eventual constancy of both the base hypothesis and its tangent structure up to tangent equivalence.

Definition 7.11 (DIAL-ready UOCL model). A UOCL hypothesis is *DIAL-ready* for a decision declaration when it is decision-ready, its relevant world and mechanism objects carry identified tangent structure, evidence and action maps admit coherent tangent lifts, tangent observers separate the repair-relevant defects, and unresolved tangent comparisons are represented as explicit obligations rather than silently treated as derivatives.

DIAL readiness therefore does not require a differential-category realization. When such a realization is claimed, however, its modality and deriving transformation become additional learned warrants rather than automatic consequences of the tangent axioms.

This criterion gives the fourth book a precise relation to *Infinitesimal Creativity*. DIAL specifies how infinitesimal defects can guide localized

creative repair. TUOCL explains how an agent might acquire, test, preserve, and accommodate the tangent structure that makes those defects meaningful.

7.6 *The cost of categorical learning*

Classical sample complexity is only one coordinate of UOCL complexity. A useful accounting vector is

$$\mathbf{C}_{\text{UOCL}} = (N_{\text{base}}, N_{\text{tan}}, N_{\text{revision}}, N_{\text{doctrine}}, C_{\text{update}}, C_{\text{memory}}).$$

It records base probes, tangent-sensitive probes, ordinary accommodation, doctrinal accommodation, computational update cost, and the size of persistent warrants. Two learners with identical predictive error can differ radically: one may reconstruct a stable reusable theory, while the other repeatedly relearns answers after destructive revision.

This vector also clarifies the role of a categorical prior. A stronger prior may reduce probe and revision complexity while increasing the risk of doctrinal obstruction. The gain must therefore be evaluated together with the cost and correctness of accommodation.

7.7 *Two enrichments of categorical learning*

Let $\mathbb{D}$ be a decision declaration specifying information objects, actions, consistency maps, consequences, and observers. There is a forgetful passage

$$U : \mathbf{UODL}_{\mathbb{D}} \longrightarrow \mathbf{UOCL}$$

that discards the action readout and its consequence semantics. Its essential image consists of UOCL learners whose hypotheses are decision-ready for $\mathbb{D}$. The passage is not generally reversible: a predictive category may omit an action object, fail to identify distinctions on which action depends, or support no safe intervention semantics.

Active UOCL probes and UODL actions must also be distinguished. A probe is licensed for its ability to separate hypotheses. A decision is evaluated by its consequences and may endogenously change the later presentation. In special cases one event serves both roles, but then the identification value, decision value, and safety obligation remain separately typed.

Independently, there is a forgetful passage

$$U_{\text{tan}} : \mathbf{TUOCL} \longrightarrow \mathbf{UOCL}$$

that discards tangent structure. Neither this functor nor the decision forgetful functor factors through the other in general. Their pullback

over **UOCL** is the class of tangent decision learners: the natural home for LINCOS comparisons and DIAL repair.

Design principle

First identify the categorical state update. Then identify the infinitesimal geometry and add a decision declaration as separately justified enrichments. Do not infer tangent or action semantics from predictive success alone.

Part IV now specializes UOCL to decision-bearing worlds. Its later tangent sections operate on the intersection just defined. It introduces the principal-past semantics, nerves, homotopy-coherent persistence, and observer comparisons needed before the later OCO, bandit, decentralized, safety, and lifelong instances can be studied.

Further reading

Gold provides the classical distinction between finite-prefix consistency and identification in the limit [Gold, 1967]. Universal imitation games motivate behavioral stabilization under coinductive presentations [Mahadevan, 2024]. The categorical notions of equivalence, indexed structure, and coherent transport used here are developed systematically by Mac Lane and Riehl [Mac Lane, 1971, Riehl, 2016]. Chapters 3 and 3 supply the presentation-relative semantics that UOCL makes algorithmic. Cockett and Cruttwell supply the axiomatic tangent-category setting used by TUOCL [Cockett and Cruttwell, 2014]; the coalgebra construction of Cockett et al. [2020] explains how a differential category can generate a tangent world without identifying the two notions. The LINCOS and DIAL books motivate its role in infinitesimal comparison and repair [Mahadevan, 2026c,b].

Probably Approximately Categorically Correct Learning

FROM PAC PREDICTION TO APPROXIMATE COMPOSITIONAL IDENTIFICATION

EXACT IDENTIFICATION IS AN AUSTERE ENDPOINT. Even a finite category can have many presentations, and an interacting learner may never receive the probes that distinguish two worlds everywhere. For the purposes of prediction, repair, or decision, however, the learner may need only to answer most of the relevant categorical questions correctly. This chapter introduces *Probably Approximately Categorically Correct* (PACC) learning: a statistical relaxation of UOCL in which approximation is measured by an equivalence-invariant doctrine of categorical probes.

The phrase needs to be parsed carefully. *Probably* refers to the random transcript observed by the learner. *Approximately* refers to its risk on future probes. *Categorically correct* requires the hypothesis to remain a genuine object of the declared categorical hypothesis class and evaluates it through categorical structure, not through a syntactic edit distance between presentations. PACC therefore weakens identification without weakening associativity, identities, coherence, or any other axiom of the chosen hypothesis doctrine.

8.1 The PAC template

Let X be an instance space, Y a label space, P an unknown distribution on X , and $f_\star : X \rightarrow Y$ a target. For a hypothesis $g : X \rightarrow Y$, the ordinary zero-one risk is

$$R_P(g, f_\star) = \Pr_{x \sim P} \{g(x) \neq f_\star(x)\} = \sum_{x: g(x) \neq f_\star(x)} P(x) \quad (8.1)$$

in the discrete case. A PAC learner receives an independent sample $S \sim P^m$ and, for prescribed $0 < \epsilon, \delta < 1$, returns g_S such that

$$\Pr_{S \sim P^m} \{R_P(g_S, f_\star) \leq \epsilon\} \geq 1 - \delta. \quad (8.2)$$

The realizable theory assumes $f_\star$ belongs to the hypothesis class. The agnostic theory instead compares the learner with the best available hypothesis. These two probability spaces must not be confused: P weights future instances, whereas $1 - \delta$ measures confidence over the sampled training transcript.

PACC retains this quantifier pattern and changes what counts as an instance, an answer, and an error. Its instances are typed questions about a candidate world. Its answers may be objects, arrows, diagrams, universal witnesses, or spaces of coherent fillers. Its loss must be invariant under the equivalence appropriate to those answers.

8.2 The PACC declaration

Fix a target categorical world $\mathcal{C}_\star$ in a hypothesis category $\mathbf{H}$, and a doctrine $\mathcal{Q}$ of admissible probes. A probe $q \in \mathcal{Q}$ has an answer object $q(\mathcal{C})$ for every hypothesis on which it is defined. Examples include:

- the domain, codomain, identity, or composite of named arrows;
- whether proposed generators factor an observed composite and whether a declared relation is respected by the model;
- whether a displayed diagram commutes or a specified horn has a filler;
- the category or space of cones realizing a universal property;
- a representable, predictive, causal, decision, or tangent-sensitive response admitted by the doctrine.

Each probe carries a bounded discrepancy

$$d_q : q(\mathcal{C}_\star) \times q(\mathcal{C}) \longrightarrow [0, 1].$$

The notation suppresses the comparison span needed when the answers do not literally inhabit the same set. We require d_q to vanish on the declared answer equivalences and to be invariant under equivalent re-presentations of the target and hypothesis. Literal equality of chosen terminal objects, for example, is the wrong test; equivalence of their cone categories is the relevant one.

Definition 8.1 (Categorical probe risk). For a distribution P on $\mathcal{Q}$, the risk of a hypothesis $\mathcal{C}$ relative to $\mathcal{C}_\star$ is

$$\text{Err}_{P, \mathcal{Q}}(\mathcal{C}, \mathcal{C}_\star) = \mathbb{E}_{q \sim P} [d_q(q(\mathcal{C}_\star), q(\mathcal{C}))]. \quad (8.3)$$

Two hypotheses are $\mathcal{Q}$ -equivalent when every admissible probe assigns equivalent answers to them.

Definition 8.2 (PACC declaration and learner). A PACC declaration is a tuple

$$\mathbb{A}_{\text{PACC}} = (\mathbf{H}, \mathcal{Q}, P, d, \epsilon, \delta), \quad 0 < \epsilon, \delta < 1,$$

together with a presentation protocol that produces a typed transcript $S = ((q_i, q_i(\mathcal{C}_\star)))_{i=1}^m$. A UOCL algorithm A is (ϵ, δ) -PACC after m probes when

$$\Pr_{S \sim P^m} \{ \text{Err}_{P, \mathcal{Q}}(A(S), \mathcal{C}_\star) \leq \epsilon \} \geq 1 - \delta. \quad (8.4)$$

It PACC-learns a family of declarations when a sample bound $m(\epsilon, \delta)$ makes this statement hold uniformly over their targets and admitted probe distributions.

The i.i.d. protocol is the closest analogue of classical PAC and supplies the baseline theory below. It is not the full online theory. Active probes, dependent categorical fragments, and world-changing interventions require conditional or martingale variants of the guarantee.

8.2.1 Running example: PACC collider identification

The collider declaration gives a non-discrete instance of categorical risk. Let P be a distribution over admitted experimental contexts and let d_p be a bounded discrepancy between predicted and observed outcome laws for probe p , such as a normalized divergence on binned event distributions. For a target response semantics $H_\star$, define

$$\text{Err}_{P, \text{coll}}(H, H_\star) = \mathbb{E}_{p \sim P} d_p(\text{Resp}_H(p), \text{Resp}_{H_\star}(p)).$$

A PACC collider learner returns $\hat{H}$ whose error is at most ϵ with transcript probability at least $1 - \delta$. This claim is relative to P : a rare energy regime or unmeasured channel can carry substantial structural disagreement while contributing almost no average risk. Coverage conditions are therefore part of the scientific prior, not a technical afterthought.

Nor does small collider risk establish ontological identity. It establishes approximate equivalence of the declared response semantics on the weighted probe family. Recovering generators, relations, or symmetry doctrine requires stronger structural probes of the kinds distinguished in Definition 8.6.

8.3 Ordinary PAC learning is the discrete case

The adjective *categorical* should earn its keep, but the new definition must recover the classical one exactly.

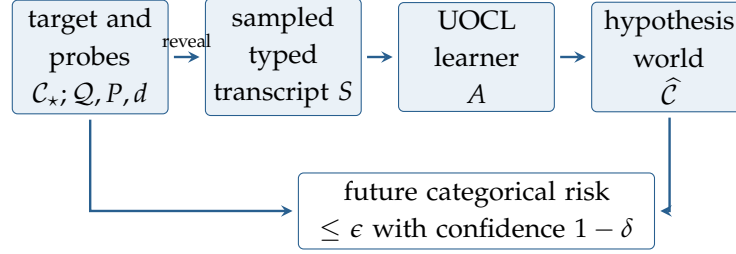


Figure 8.1: The PACC quantifiers. The distribution P measures which categorical questions matter; the outer probability measures which transcript the learner happened to receive.

Proposition 8.3 (PAC as discrete PACC). *Let the objects of a discrete category $\mathcal{X}$ be the elements of X . Represent a classifier $g : X \rightarrow Y$ by the corresponding functor $g : \mathcal{X} \rightarrow \mathcal{Y}$, where $\mathcal{Y}$ is discrete. For every x , let $q_x(g) = g(x)$, set $d_{q_x}(y, y') = \mathbf{1}[y \neq y']$, and transport P from X to the probes q_x . Then categorical probe risk is exactly $R_P(g, f_*)$, and Definition 8.2 is the ordinary PAC guarantee.*

Proof. The discrete categories contain no nonidentity composites or coherence data, so a probe asks only for the label of one object. Substitution in Equation (8.3) gives Equation (8.1); the outer sample probabilities in Equations (8.4) and (8.2) are identical. $\square$

Thus PACC is not a metaphorical use of PAC. It is a conservative extension in which the instance language can expose relations among observations. The classical case reappears when that relational structure is discarded.

8.4 Finite PACC bounds

Let $\bar{\mathbf{H}} = \mathbf{H} / \simeq_{\mathcal{Q}}$ be the quotient of the hypothesis class by its answers to all probes, and suppose it has $N < \infty$ equivalence classes. First take zero-one discrepancy and assume the target class is realizable.

Theorem 8.4 (Finite realizable PACC). *If a learner returns a hypothesis consistent with every sampled probe, then*

$$m \geq \frac{\log N + \log(1/\delta)}{\epsilon}$$

implies that its categorical probe risk exceeds ϵ with probability at most δ .

Proof. A fixed hypothesis of risk greater than ϵ agrees with one random probe with probability less than $1 - \epsilon$, hence survives all m independent probes with probability at most $(1 - \epsilon)^m \leq e^{-m\epsilon}$. A union bound over at most N query-equivalence classes gives failure probability at most $Ne^{-m\epsilon} \leq \delta$. $\square$

The theorem counts observationally distinct categorical worlds rather than raw syntactic presentations. It therefore rewards the quotient that

Probe stratum	Question sampled	Defect made statistically visible
Objects	typing, identity, attribute, or label	missing or conflated entities
Arrows	source, target, and observed transition	mistyped mechanisms or dynamics
Composition	two-step path versus declared composite	failure of compositional prediction
Theory	generator factorization or declared equation	inadequate or overidentified presentation
Horns	existence and compatibility of fillers	unresolved or incoherent higher behavior
Universal properties	cones, mediators, and uniqueness up to equivalence	incorrect limits, colimits, or Kan extensions

Table 8.1: Structural coverage turns a numerical risk bound into a statement about a declared portion of categorical structure.

UOCL was already required to state. When the target need not lie in the hypothesis class, consistency is replaced by empirical risk minimization.

Theorem 8.5 (Finite agnostic PACC). *Assume $0 \leq d_q \leq 1$, let $\bar{\mathbf{H}}$ have $N < \infty$ classes, and let $\hat{\mathcal{C}}$ minimize empirical categorical probe risk. With probability at least $1 - \delta$,*

$$\text{Err}_{P,Q}(\hat{\mathcal{C}}, \mathcal{C}_*) \leq \inf_{\mathcal{C} \in \bar{\mathbf{H}}} \text{Err}_{P,Q}(\mathcal{C}, \mathcal{C}_*) + 2\sqrt{\frac{\log(2N/\delta)}{2m}}.$$

In particular, $m \geq 2 \log(2N/\delta) / \epsilon^2$ suffices for excess risk at most ϵ .

Proof. Hoeffding’s inequality and a union bound imply simultaneous deviation at most $\alpha = \sqrt{\log(2N/\delta)/(2m)}$ between empirical and population risk for every query-equivalence class. Empirical optimality inserts two such deviations between the selected hypothesis and the population optimum. $\square$

These bounds are intentionally elementary. Their purpose is to expose the new modeling obligation: statistical control is only as meaningful as the categorical probes and invariances used to define risk.

8.5 Coverage is part of the categorical prior

An arbitrary distribution can make a structurally defective model appear accurate simply by assigning zero mass to every revealing probe. A PACC declaration must therefore say which categorical strata require coverage. A simple finite mixture is

$$P = \lambda_0 P_{\text{obj}} + \lambda_1 P_{\text{arr}} + \lambda_2 P_{\text{comp}} + \lambda_{\text{th}} P_{\text{th}} + \lambda_3 P_{\text{horn}} + \lambda_U P_{\text{univ}}, \quad \lambda_i > 0, \quad \sum_i \lambda_i = 1. \quad (8.5)$$

Only the strata required by the doctrine need appear, but none of those strata may be assigned zero weight.

Equation (8.5) does not assert that one canonical distribution exists. It makes the choice auditable. In an automaton or PSR, for example, P may weight future action–observation tests. In a finite category, it may sample typing and partial composition-table entries. In a quasi-categorical model, it may sample boundary and inner-horn diagrams. Different mixtures certify different notions of approximation.

8.6 *Three strengths of categorical approximation*

Definition 8.6 (Behavioral, structural, and coherent PACC). Let $\hat{\mathcal{C}}$ be a learned world.

1. It is *behaviorally PACC* when the guarantee concerns only the answers returned by a declared external query doctrine.
2. It is *structurally PACC* when a learned comparison functor, span, or profunctor to the target is also tested for properties such as faithfulness, fullness, essential coverage, and preservation of declared universal constructions.
3. It is *coherently PACC* when the probes range over simplices, horns, filler spaces, and higher comparison witnesses, with discrepancy invariant under the declared homotopy equivalences.

Behavioral PACC can be enough for a fixed predictive interface, just as a PSR need not recover a privileged latent state. Structural PACC supports reuse of compositions in new contexts. Coherent PACC is needed when the space of repairs, rather than the mere existence of one answer, matters. None implies causal correctness unless the probe doctrine contains interventions and the usual identification assumptions justify their interpretation.

For universal properties, a zero–one discrepancy may ask whether a comparison of cone categories is an equivalence. A graded discrepancy may measure which probes fail to admit a mediator or how the homotopy type of a filler space differs. This is categorically meaningful only after the answer objects and weak equivalences have been declared.

8.7 *Toward a categorical dimension theory*

Finite cardinality is too crude for rich UOCL classes. The PAC analogue suggests a dimension, but arbitrary binary labelings of isolated points would again forget composition.

Definition 8.7 (Categorical probe dimension). For a binary probe doctrine, a finite compatible family $Q_0 \subseteq \mathcal{Q}$ is *categorically shattered* by $\mathbf{H}$ when every binary answer pattern that satisfies the declared typing,

composition, and coherence constraints is realized by some hypothesis in $\mathbf{H}$. The *categorical probe dimension* $\text{CPdim}(\mathbf{H}, \mathcal{Q})$ is the supremum of the sizes of such families.

When $\mathcal{Q}$ is discrete, compatibility imposes no additional relations and this reduces to the familiar shattering idea behind VC dimension. For general doctrines, deriving uniform-convergence bounds from CPdim , or replacing it by enriched covering numbers or a homotopical complexity invariant, is a theorem program rather than a result claimed here. The definition records the principle that complexity should count independently variable *compositional obligations*, not simply the cardinality of a presentation.

8.8 PACC evolvability

PAC learning permits a learner to inspect a labeled sample and replace its current hypothesis by any hypothesis its algorithm can construct. Valiant's theory of *evolvability* imposes a different discipline [Valiant, 2009]. A representation changes only through a bounded neighborhood of mutations, and selection sees empirical aggregate performance rather than the identities of the examples responsible for that performance. Learning is therefore a path of locally selected variations, not a sequence of globally recomputed empirical-risk minimizers.

The categorical abstraction must constrain both the states and the steps of that path. Let $\mathbf{R}$ be a category of admissible representations inside the PACC hypothesis class $\mathbf{H}$. Its objects are categorical world models, and a morphism, span, or declared higher cell records an admissible mutation together with the comparison data needed to audit what it preserves. For each representation $\mathcal{C}$ and accuracy scale ϵ , a finite neighborhood $N_\epsilon(\mathcal{C})$ contains the candidate mutations that may be tested in one generation. Define aggregate performance by

$$\text{Perf}_{P, \mathcal{Q}}(\mathcal{C}) = 1 - \text{Err}_{P, \mathcal{Q}}(\mathcal{C}, \mathcal{C}_\star). \quad (8.6)$$

A categorical mutator may estimate this quantity on fresh sampled probes for the current representation and its neighbors. It receives those empirical scores, up to a declared tolerance, but not an example-indexed prescription of how to alter the hypothesis.

Definition 8.8 (PACC evolvability). A family of PACC declarations is *PACC evolvable* relative to a representation category $\mathbf{R}$ when there are polynomially bounded neighborhood, sampling, tolerance, and selection rules such that, for every admitted target $\mathcal{C}_\star$, every initial representation $\mathcal{C}_0$, and every $0 < \epsilon, \delta < 1$, the induced sequence

$$\mathcal{C}_0 \longrightarrow \mathcal{C}_1 \longrightarrow \cdots \longrightarrow \mathcal{C}_T$$

uses at each step only aggregate empirical categorical performance to select $\mathcal{C}_{t+1} \in \mathcal{N}_\epsilon(\mathcal{C}_t)$, remains inside the declared hypothesis doctrine, and satisfies

$$\Pr\{\text{Err}_{P,Q}(\mathcal{C}_T, \mathcal{C}_\star) \leq \epsilon\} \geq 1 - \delta$$

after a number of generations and probes polynomial in the declared size parameters, $1/\epsilon$, and $\log(1/\delta)$. The neighborhood and selection rules must respect the answer equivalences used by the PACC risk; equivalent re-presentations may not acquire different fitness merely from their syntax.

Ordinary evolvability is recovered when the objects of $\mathbf{R}$ are an ordinary representation class, its admissible morphisms encode the mutation graph, the probes are discrete input points, and performance is agreement with an ideal function. PACC evolvability adds three obligations. Mutations must remain well typed and categorically valid. Their scores must be invariant under the declared equivalences. Finally, the mutation arrows expose which settled objects, composites, universal witnesses, or higher coherences survive a statistically favored change. Thus two neighbors can have indistinguishable predictive fitness while differing sharply in their capacity for persistent reuse.

Within a fixed fiber of the ORACLE hypothesis fibration, this is a model of *evolutionary assimilation*: local variants compete under one doctrine. Accommodation is more demanding. A mutation that changes doctrine crosses fibers, so its old and new performance values are comparable only after probes and answers have been transported along declared cartesian, cocartesian, or profunctorial structure. Without that comparison, a doctrine shift is not a beneficial mutation; it is a change of measurement scale. PACC evolvability therefore turns a biological restriction into an ORACLE design question: which local, compositional variations permit statistically reliable progress without granting the learner direct access to the individual experiences that selected them?

8.9 Persistent and online PACC

At time t , let $\mathcal{Q}_t$ be the probes licensed by the available information category and let P_t be their conditional distribution. A persistent learner must control current risk without gratuitously changing answers already settled under earlier doctrines.

Definition 8.9 (Persistent PACC). A UOCL execution $(\hat{\mathcal{C}}_t)_{t \geq 0}$ is persistently PACC on a settled subdoctrine $\mathcal{Q}_0$ when it satisfies the declared PACC risk bounds at each time and its comparison maps $\hat{\mathcal{C}}_t \rightarrow \hat{\mathcal{C}}_{t+1}$ are conservative on all answers registered as settled in $\mathcal{Q}_0$. Under drift, the declaration must additionally specify cumulative or discounted risk and the defect of transporting probes from P_t to P_{t+1} .

Proposition 8.10 (Finite-horizon confidence composition). *Suppose that, conditional on every preceding transcript, the learner at each $t \leq T$ satisfies its PACC bound with failure probability at most δ_t . Then with probability at least $1 - \sum_{t=1}^T \delta_t$, all T risk bounds hold simultaneously. If the comparison maps are conservative on the settled doctrine, its registered answers also persist through time T .*

Proof. The conditional guarantees imply unconditional failure probabilities bounded by δ_t . The first claim is the union bound. The second is finite composition of conservative comparison maps, as in Proposition 7.3. $\square$

This proposition is modest but clarifies what changes online. A static PACC bound controls a random approximation. Persistent PACC additionally requires typed transport of what has been learned. Adaptive probing calls for concentration under dependent observations; endogenous decisions further require the information and safety semantics developed in Parts IV and VI.

8.10 What PACC does and does not promise

PACC makes exact UOCL identification the zero-risk, complete-doctrine limit. It also gives approximate categorical learning a falsifiable semantics: name the probes, answer equivalences, loss, coverage, and confidence source. It does not license a raw graph-edit distance between arbitrary presentations, turn high predictive accuracy into equivalence of categories, or convert observational queries into causal ones. Nor does it guarantee that a behaviorally accurate world is suitable for a new decision declaration.

The central research questions are now sharper. Which structural coverage conditions turn small probe risk into approximate fullness, faithfulness, or preservation of universal properties? Which categorical dimensions control infinite hypothesis classes? When do active probes reduce sample complexity without changing the world being identified? And which tangent-sensitive probe doctrines make a PACC hypothesis DIAL-ready? These questions connect statistical learning theory to the compositional obligations unique to ORACLE.

Design principle

Approximate answers only after declaring the categorical invariants that must remain exact. A PACC learner may be uncertain about which world it inhabits; it may not silently cease to inhabit a category.

Further reading

Valiant introduced the distribution-free PAC criterion [Valiant, 1984] and later formulated evolvability as learning by local mutation and selection driven by aggregate empirical fitness rather than by individual examples [Valiant, 2009]. Vapnik and Chervonenkis developed the uniform convergence and shattering framework underlying statistical control of infinite classes [Vapnik and Chervonenkis, 1971]; the relation among learnability, VC dimension, and finite-sample bounds was sharpened by Blumer et al. [1989]. Gold's identification in the limit supplies the contrasting exact, asymptotic presentation paradigm [Gold, 1967]. Chapters 3 and 7 develop its ORACLE and UOCL forms; the present chapter supplies their query-relative statistical relaxation.

Part IV

Universal Online Decision Learning

Online and Persistent Universal Decision Learning

FROM UOCL SPECIALIZATION TO PREFIX SEMANTICS, NERVES, HOMOTOPY COHERENCE, AND REPAIR

ORACLE states the semantic problem of identifying a predictive categorical world, and UOCL supplies its effective online machinery. Universal Online Decision Learning begins when a UOCL hypothesis is equipped with information, action, consistency, and observation maps. Online convex optimization (OCO) then describes one sequential decision problem over a convex feasible set under convex losses revealed after each action [Hazan, 2023]. Universal Online Decision Learning asks a structural question:

Which parts of an online convex decision problem are determined by universal properties, and when do those constructions remain stable under infinitesimal variation?

The operadic theory of convexity characterizes convex sets as algebras over a convexity PROP and equips their category with additional monoidal structure [Haderi et al., 2024]. UODL uses this object-level account for the feasible decision domain. It then adds Universal Decision Learning (UDL) for extension and consistency, and the LINC discipline for typed infinitesimal comparison and repair [Mahadevan, 2021, 2026a,c].

Thus UODL is not the ambient theory of this book; it is the decision-bearing specialization of UOCL, restricted to categories carrying an online decision declaration. Its first technical problem should be simple enough for exact calculation but rich enough to distinguish evidence, action geometry, and consistency. Online convex optimization (OCO) supplies this setting [Shalev-Shwartz, 2012, Hazan, 2016]. Many OCO algorithms admit FTRL interpretations, and mirror-descent and FTRL presentations can be exactly equivalent under explicit hypotheses [McMahan, 2011]. This makes FTRL a candidate normal form for the action mechanism, not an assertion that all optimization algorithms are secretly identical.

Contributions. The UODL specialization currently establishes the following contributions.

1. It defines online and persistent UDL using principal-past restriction, non-anticipation, and coherent semantic transport, and proves a fixed-shape lifting theorem for Kan invariance.
2. It factors regret through a causal UODL branch, a full-information UDL comparator branch, and a separately typed observer, while identifying static trajectories through a right-Kan counit condition. It then defines information-relative intrinsic regret, recovers ordinary OCO on a classical chain, identifies the replay obstruction for non-classical information, and realizes crossword and Sudoku solving as decentralized limit problems. It then compares the universal Sudoku solution object with the learned fixed-point dynamics of a Flow Reasoning Model.
3. It defines comparator sketches as right-Kan descent data, proves an essential-image representation theorem, constructs the static-to-dynamic refinement spectrum, and separates universal descent defects from their metric and numerical observation.
4. It represents finite decision histories by a simplicial nerve, distinguishes external time from compositional depth, and gives a sufficient condition for online-to-offline UDL continuity.
5. It defines homotopy-coherent online UDL after localization at semantic weak equivalences, interprets horn filling as a space of repairs, gives a two-component inner-horn example in which later evidence forces a repair, and proves derived skeletal continuity for finite consistency shapes in a stable target.
6. It defines fixed-stratum infinitesimal intrinsic models and proves that acyclic information dependence makes the linearized closed-loop operator nilpotent, yielding a unique tangent response by causal forward substitution.
7. It gives a categorical OCO declaration in which the feasible domain is a convex algebra and convex losses are order-lax maps, and proves that its finite-dimensional Euclidean realization recovers the standard OCO protocol.
8. It separates FTRL into accumulated evidence, composite-loss treatment, stabilizing geometry, and an optimization readout.
9. It proves a finite enriched left-Kan accumulation theorem and its exact tangent lift, then derives typed FTRL decision sensitivities and a decision-null quotient.

10. It separates base convergence, tangent stability, and preservation of asymptotic limits, proving sufficient lifting results for normalized quadratic FTRL and contractive smooth update systems.
11. It presents FTRL as a metric proximal map, proves its smooth tangent lift and contraction bound, and identifies the nonsmooth boundary where graphical rather than ordinary tangents are required.
12. It models bandit feedback as a stochastic information channel, separates pathwise impossibility from barycentric reconstruction, proves tangent admissibility and regret transfer, recovers the adversarial finite-arm rate, and identifies smoothed tangent covectors as the repair target in bandit convex optimization.
13. It separates an intrinsic asynchronous event category from its external serializations, proves schedule invariance when incomparable updates commute, and carries the example through distributed minimization, asynchronous Q-learning, and stale-authority safety.
14. It identifies free convex completion as the universal carrier underlying categorical online boosting.
15. It formulates agentic safety as a property of the realized information structure, proves compositional safety from generator-level admission and safety transport under pointwise right-Kan extension, and identifies faithful audit as a necessary condition for observer-based enforcement. It then interprets temporal behavior types as trust contracts and proves local-to-global verification by sheaf descent.

The scope is intentionally narrow. We do not yet identify comparator regret with a particular right Kan extension, derive a new regret-optimal algorithm, or claim that a universal tangent comparison is causal.

9.1 *Claim discipline and the theorem ladder*

The program separates four layers that must not be collapsed. A construction is *categorical* when its types and compositions are explicit; *universal* when it is characterized by a universal property; *decision-theoretic* when it selects or compares admissible actions; and *causal* only when an intervention semantics and identification assumptions justify that interpretation. Thus a natural isomorphism need not imply low regret, and a nonzero tangent response need not identify the cause of a failure.

The remainder of the book follows six theorem families:

1. *representation*: progressive evidence determines a persistent online UDL object;

2. *invariance*: restriction, localization, and presentation change preserve the declared universal constructions;
3. *comparison*: differently informed branches admit a typed observer and comparator object;
4. *repair*: strict, homotopy-coherent, or infinitesimal defects select admissible repairs;
5. *stability*: tangent lifting preserves well-posedness and, under additional hypotheses, convergence; and
6. *transport*: learned structure survives a change of task, information pattern, or environment.

9.2 *Scientific questions for universal online decision learning*

UODL is organized by questions that recur whenever categorical identification must support action before all relevant information is available. These questions do not depend on a particular algorithmic family or on a single choice of clock, loss, or comparator.

1. *What makes a decision genuinely online?* The learner must act from an information region that is closed under its available past but excludes unrevealed evidence. The first task is therefore to identify the least admissible information object and express non-anticipation as a typed property of the decision map.
2. *When does universal decision semantics persist as evidence grows?* Left and right Kan extensions may solve a decision problem at each finite stage without commuting with restriction from one stage to the next. UODL asks for the comparison maps, exactness hypotheses, and invariance conditions under which stagewise solutions form a coherent online object.
3. *How can decisions made with different information be compared?* An online learner and a hindsight comparator inhabit different information contexts. A meaningful notion of regret therefore requires a declared comparator object, a transport between information shapes, and an observer that turns their values into an ordered, numerical, or logical defect.
4. *What should happen when an online composition is incomplete or inconsistent?* A partial decision history may leave a horn unfilled, admit several coherent fillers, or contain an obstruction to every filler in the current doctrine. The scientific problem is to distinguish ordinary completion, repair within a declaration, and accommodation of the declaration itself.

5. *Which action mechanisms are stable under variation?* Feasible decisions may carry convex, metric, order, or richer geometry. Regularized leaders, mirror maps, proximal maps, and related mechanisms become useful categorical constructions only when their existence, sensitivity, and convergence survive the relevant tangent lift.
6. *Which conclusions survive the loss of a global clock or full observation?* Bandit feedback, asynchronous computation, and decentralized teams alter the information category rather than merely adding noise to a synchronous protocol. UODL asks which values descend through stochastic channels, which update schedules are equivalent, and which information patterns create irreducible comparison or safety obstructions.
7. *When does learned decision structure remain useful beyond the task that produced it?* Persistent structure must transport across changes of environment, observer, and objective while retaining its universal witnesses, repair certificates, and safety conditions. This is the bridge from successful online action to lifelong compositional learning.

Classical online convex optimization supplies a particularly transparent laboratory for these questions: time is a chain, the feasible object is convex, losses are revealed sequentially, and cumulative regret is observed in an ordered additive value space [Shalev-Shwartz, 2012, Hazan, 2016, 2023]. Changing any one of these structural choices leads to a different scientific problem rather than merely another entry in a catalog of algorithms. The development below therefore begins with prefix semantics and universal comparison, then studies convex action mechanisms as one exact realization of the broader theory.

9.3 Prefix semantics, nerves, homotopy coherence, and repair

9.4 Time, revelation, and the least available past

Let $\mathbb{T}$ be a small poset of times, written $s \leq t$ when s is no later than t . The information available at t is the principal down-set

$$\downarrow t = \{s \in \mathbb{T} \mid s \leq t\}.$$

Equivalently, it is the representable presheaf $y(t) = \mathbb{T}(-, t)$. Since $\mathbb{T}$ is a poset, $y(t)$ is subterminal and records the truth value “revealed by time t .”

Proposition 9.1 (Least causally closed information region). *The principal past $\downarrow t$ is the least downward-closed subobject of $\mathbb{T}$ containing t .*

Proof. Every downward-closed subset containing t must contain every $s \leq t$, hence contains $\downarrow t$. The principal past is itself downward closed and contains t . $\square$

Thus $\{t\}$ is generally too small because it omits history, while all of $\mathbb{T}$ is too large because it admits future information.

9.5 Prefix-indexed UDL and non-anticipation

Fix functors

$$J : \mathcal{O} \longrightarrow \mathcal{C}, \quad K : \mathcal{C} \longrightarrow \mathcal{Q},$$

and a target category $\mathcal{D}$ with the Kan extensions below. Ordinary UDL sends local data $F : \mathcal{O} \rightarrow \mathcal{D}$ to

$$U(F) = \text{Ran}_K \text{Lan}_J F.$$

An evidence filtration is a functor

$$\mathbf{F} : \mathbb{T} \longrightarrow [\mathcal{O}, \mathcal{D}], \quad t \longmapsto F_t,$$

whose transition maps register newly revealed evidence without using future components. Its online UDL semantics is the time-indexed family

$$\widehat{\mathbf{F}}_t = U(F_t) = \text{Ran}_K \text{Lan}_J F_t. \quad (9.1)$$

When observations are presented as a complete history h , F_t means the restriction determined by $h|_{\downarrow t}$.

Definition 9.2 (Online UDL). An online UDL is a prefix-indexed UDL family $\{\widehat{\mathbf{F}}_t\}_{t \in \mathbb{T}}$ whose action readout at every time t factors through the prefix restriction

$$\rho_t : h \longmapsto h|_{\downarrow t}.$$

Equivalently, it satisfies non-anticipation:

$$h|_{\downarrow t} = h'|_{\downarrow t} \implies A_t(h) = A_t(h').$$

No convexity, probability, regret, or optimization is required by this definition. “Online” changes the information semantics of UDL while leaving its left/right Kan form intact.

Definition 9.3 (Persistent online UDL). A persistent online UDL is an online UDL whose synthesized semantic objects and comparison maps form a functor

$$\widehat{\mathbf{F}} : \mathbb{T} \longrightarrow [\mathcal{Q}, \mathcal{D}].$$

Thus for $s \leq t$ there is a registered transport

$$u_{s,t} : \widehat{\mathbf{F}}_s \longrightarrow \widehat{\mathbf{F}}_t, \quad u_{t,r} \circ u_{s,t} = u_{s,r}.$$

The current action is a readout from $\widehat{\mathbf{F}}_t$, not the persistent object itself.

The transport may preserve an earlier structure, enrich it, or expose a repair obligation. Requiring literal inclusion would rule out legitimate revision, so preservation claims must identify the particular registered subobject that survives.

9.6 Online Kan invariance

Ordinary Kan invariance says that a behavioral property depends only on the induced semantics $\widehat{F}$, not on the syntactic presentation of F . The online analogue must remember both every prefix semantics and the maps relating different prefixes.

Definition 9.4 (Online Kan invariance). A property of filtered decision models is online Kan-invariant if it depends only on the natural-isomorphism class of

$$\widehat{\mathbf{F}} : \mathbb{T} \rightarrow [\mathcal{Q}, \mathcal{D}].$$

Two online models M, N are online Kan-bisimilar when

$$\widehat{\mathbf{F}}^M \cong \widehat{\mathbf{F}}^N$$

naturally in time and decision context.

Pointwise equivalence at every time is necessary but not sufficient: the chosen isomorphisms must commute with the persistence maps. For every $s \leq t$, an online bisimulation η requires

$$\eta_t \circ u_{s,t}^M = u_{s,t}^N \circ \eta_s. \quad (9.2)$$

Theorem 9.5 (Fixed-shape lifting of Kan invariance). *Let $\mathbb{T}, \mathcal{O}, \mathcal{C}, \mathcal{Q}$ be small, let $\mathcal{D}$ admit the relevant Kan extensions, and hold J and K fixed through time. Then the UDL operator lifts pointwise to*

$$[\mathbb{T}, \mathbf{U}] : [\mathbb{T}, [\mathcal{O}, \mathcal{D}]] \longrightarrow [\mathbb{T}, [\mathcal{Q}, \mathcal{D}]],$$

with

$$([\mathbb{T}, \mathbf{U}](\mathbf{F}))_t \cong \text{Ran}_K \text{Lan}_J F_t.$$

Consequently, online Kan invariance is ordinary Kan invariance internal to the time-indexed functor category. Two models are online Kan-bisimilar if and only if their prefixwise Kan semantics admit a family of natural isomorphisms satisfying equation (9.2).

Proof. The adjunctions $\text{Lan}_J \dashv J^*$ and $K^* \dashv \text{Ran}_K$ lift pointwise to functor categories. Their composite at time t is therefore $\text{Ran}_K \text{Lan}_J F_t$. An isomorphism in a functor category is precisely a componentwise isomorphism natural in the indexing category, and its naturality square is equation (9.2). $\square$

In particular, a fixed-shape evidence filtration obtains formal persistence maps by functoriality of $\mathbf{U}$. This does not imply that those maps are monic, lossless, computationally reusable, or beneficial on later tasks; those are additional persistence obligations rather than consequences of Kan functoriality.

This equivalence is the minimal fixed-shape formulation: time indexing adds adaptedness and coherence but does not change the universal construction. It also explains why merely proving $\hat{F}_t^M \cong \hat{F}_t^N$ separately for every t is inadequate. Arbitrarily chosen pointwise isomorphisms can fail to preserve how knowledge is updated.

9.7 Growing diagrams and the Beck–Chevalley defect

In a genuinely online problem the observation category may itself grow:

$$\mathcal{O}_s \xrightarrow{i_{s,t}} \mathcal{O}_t, \quad s \leq t.$$

Restriction of a later semantic construction to an earlier prefix can then be compared with constructing the semantics directly from the restricted data. The universal properties generate a canonical mate

$$\kappa_{s,t} : \mathbf{U}_s(i_{s,t}^* F_t) \longrightarrow r_{s,t}^* \mathbf{U}_t(F_t), \quad (9.3)$$

with orientation adjusted when the declared restriction square uses the opposite variance.

Definition 9.6 (Prefix-exact online UDL). An online UDL is prefix-exact when every comparison $\kappa_{s,t}$ in equation (9.3) is an isomorphism and these isomorphisms satisfy the identity and composition coherences over $r \leq s \leq t$.

Prefix exactness is the variable-shape analogue of the fixed-shape lifting theorem: it is a Beck–Chevalley condition saying that “restrict after universal decision synthesis” agrees with “restrict the evidence and then synthesize.” When $\kappa_{s,t}$ is not invertible, its defect is not automatically an error. It measures precisely the structure introduced, destroyed, or revised by later evidence and therefore supplies the typed input to a LINCOS repair declaration.

9.8 Persistent online abstractions

For a context object X , define prefixwise behavioral equivalence by

$$x \sim_t x' \iff \hat{\mathbf{F}}_t(x) \cong \hat{\mathbf{F}}_t(x').$$

Under conservative revelation, later semantics restricts to earlier semantics. Hence new evidence can split an earlier equivalence class but

cannot identify behaviors that were already observably distinct:

$$s \leq t \implies \sim_t \subseteq \sim_s.$$

Proposition 9.7 (Inverse system of online Kan quotients). *Assume conservative revelation and the existence of the prefixwise quotients. Then $s \leq t$ induces a canonical forgetting map*

$$X / \sim_t \longrightarrow X / \sim_s.$$

These maps satisfy identity and composition, so the minimal Kan-invariant abstractions form an inverse system, equivalently a presheaf on time.

Proof. An isomorphism of later semantic values restricts to an isomorphism of the earlier values, so $x \sim_t x'$ implies $x \sim_s x'$. The quotient by the finer relation therefore maps canonically to the quotient by the coarser relation. Uniqueness of quotient factorization gives identity and composition. $\square$

This variance is important: evidence accumulates forward, whereas the operation that forgets newly learned distinctions points back toward the earlier quotient. If later evidence is allowed to revise rather than conservatively extend earlier semantics, these maps must be replaced by explicit repair comparisons.

9.9 Decision nerves and compositional depth

The principal-past semantics uses an external clock. A complementary description makes the length of a composite decision history intrinsic. Let $\mathcal{H}$ be a small decision-history category: its objects are information states, its morphisms are admissible one-step decisions or transitions, and categorical composition is sequential execution. Its nerve

$$X = N\mathcal{H}$$

is the simplicial set with

$$X_n = \text{Fun}([n], \mathcal{H}).$$

This is the ordinary categorical nerve [Riehl, 2016]. Thus an n -simplex is a composable history

$$h_0 \xrightarrow{a_1} h_1 \xrightarrow{a_2} \dots \xrightarrow{a_n} h_n.$$

Proposition 9.8 (Decision nerve and skeletal filtration). *For the nerve $X = N\mathcal{H}$:*

- (i) *the terminal face deletes the most recent step, the initial face deletes the first step, and each internal face composes two adjacent morphisms;*

- (ii) *degeneracy maps insert identity decisions; and*
- (iii) *the simplicial set is the filtered colimit of its skeleta,*

$$X \cong \operatorname{colim}_{t \geq 0} \operatorname{sk}_t X.$$

Proof. The face and degeneracy operators are induced by the coface and codegeneracy maps in Δ . Functoriality turns an internal coface into composition and a codegeneracy into an identity. Every simplex has finite dimension, so it belongs to some skeleton; the skeletal inclusions therefore have colimit X . $\square$

Histories available through length t first form the truncation $\operatorname{tr}_{\leq t} X$. The smallest simplicial object generated by that truncation is $\operatorname{sk}_t X$. This yields the distinction The notation $\downarrow t$ records information revealed by external time t , whereas $\operatorname{sk}_n X$ records information represented through compositional depth n . One decision per round permits the diagonal identification $n = t$. In general, however, a decision model is naturally bifiltered by (t, n) : an agent may acquire new observations without increasing compositional depth, or construct a longer internal plan before receiving new external evidence. The flow-reasoning example in section 9.16 supplies a concrete second realization of this distinction: continuous flow time advances the generative state, while recurrent depth counts repeated repairs at a fixed flow state.

9.10 Simplicial online UDL

Let Δ/X be the category of simplices of X , and let

$$\operatorname{Simp}_{\leq t}(X) = \Delta_{\leq t}/X \xrightarrow{j_t} \Delta/X$$

be its full subcategory on simplices of dimension at most t . A decoration functor records the evidence, losses, actions, or warrants attached to the observed simplices.

Definition 9.9 (Simplicial online UDL). Given compatible decorations $F_t : \operatorname{Simp}_{\leq t}(X) \rightarrow \mathcal{D}$, the depth- t simplicial online UDL semantics is

$$U_t^\Delta(F_t) = \operatorname{Ran}_{K_t} \operatorname{Lan}_{J_t} F_t,$$

where J_t and K_t declare candidate extension and consistency over the truncated simplex category. Persistent simplicial UDL additionally registers coherent comparison maps along $\operatorname{Simp}_{\leq s}(X) \hookrightarrow \operatorname{Simp}_{\leq t}(X)$ for $s \leq t$.

If the entire category $\mathcal{H}$ and its nerve are already known, the compositional fillers carry no learning problem. The nontrivial cases begin with a partial nerve, an unknown decision category, or a known nerve whose observations and decision values are only partially decorated.

Proposition 9.10 (Inner horns encode decision composition). *The nerve of every category has a unique filler for every inner horn. It is a Kan complex, with fillers for all horns, if and only if the underlying category is a groupoid.*

Proof. An inner horn specifies a composable configuration with one composite face missing; associativity gives its unique filler. Outer horn fillers provide left and right inverses for morphisms. Hence all outer horns are fillable exactly when every morphism is invertible. $\square$

This result separates two uses of Kan’s ideas. A Kan extension universally extends a diagram; a simplicial Kan condition fills horns. They interact in the present construction but are not synonyms. Requiring a full Kan complex would also impose reversibility, which is inappropriate for irreversible actions such as resource consumption or pushing an object into an unrecoverable state. Inner-horn or quasi-categorical structure is the more natural default.

At a schematic level, partial simplicial information supplies the UDL factorization

$$\text{boundary or horn data} \xrightarrow{\text{Lan}} \text{candidate completions} \xrightarrow{\text{Ran}} \text{compatible fillers}.$$

Making this slogan a theorem requires declaring the precise horn category, the decoration, and the target enrichment; unique inner fillers in the underlying nerve alone do not supply learned values or policies.

The crossword construction in section 9.15 makes this distinction concrete. Clue agents generate local candidate words, crossing letters define compatibility faces, and a partial filling may leave a completion problem unresolved. A strict filler records an exactly compatible extension; a homotopy-coherent filler can additionally retain alternative meanings and witnesses of compatibility.

9.11 The online-to-offline comparison

The skeletal filtration creates a canonical comparison between incremental and all-at-once UDL. Put $\mathcal{S} = \Delta/X$, $\mathcal{S}_t = \text{Simp}_{\leq t}(X)$, and extend each compatible decoration to the fixed domain by

$$\tilde{F}_t = \text{Lan}_{j_t} F_t : \mathcal{S} \rightarrow \mathcal{D}.$$

Let $F = \text{colim}_t \tilde{F}_t$, fix $J : \mathcal{S} \rightarrow \mathcal{C}$ and $K : \mathcal{C} \rightarrow \mathcal{Q}$, and write

$$G_t = \text{Lan}_J \tilde{F}_t, \quad G = \text{Lan}_J F.$$

The maps $G_t \rightarrow G$ induce the online-to-offline comparison

$$\gamma : \text{colim}_t \text{Ran}_K G_t \longrightarrow \text{Ran}_K G. \quad (9.4)$$

Theorem 9.11 (Skeletal continuity of UDL). *Assume the displayed Kan extensions and filtered colimits exist. Suppose also that filtered colimits in $\mathcal{D}$ commute with every limit used in the pointwise formula for Ran_K . Then the comparison γ in equation (9.4) is an isomorphism. Consequently,*

$$\text{colim}_t U_t^\Delta(F_t) \cong U(F)$$

after the declared extensions to the common domain.

Proof. The left Kan extension Lan_J is a left adjoint and therefore preserves the filtered colimit:

$$G = \text{Lan}_J \left(\text{colim}_t \tilde{F}_t \right) \cong \text{colim}_t \text{Lan}_J \tilde{F}_t = \text{colim}_t G_t.$$

Pointwise right Kan extension is computed by the declared limits. The commutation hypothesis moves the filtered colimit through each such limit, giving

$$\text{colim}_t \text{Ran}_K G_t \cong \text{Ran}_K \left(\text{colim}_t G_t \right) \cong \text{Ran}_K G.$$

This composite is the canonical comparison γ . □

The theorem localizes the obstruction. Incremental candidate generation commutes with revelation because the left Kan stage preserves colimits. The possible failure lies in interchanging gradual revelation with the limits of right-Kan consistency. If γ is not invertible, its defect measures semantic structure available to the all-at-once UDL construction but not preserved by the filtered online construction. A LINC observer may map that structural defect to a numerical quantity, but the defect exists before any regret scalar is chosen.

9.12 Homotopy-coherent online UDL

The ordinary nerve records strict categorical composition. Learning and repair generally produce a weaker structure: two composites may be related by a specified transformation rather than equality, and those transformations must themselves satisfy higher compatibility. Categorical homotopy theory supplies the language for retaining this coherence while discarding accidental features of a presentation [Riehl, 2014]. The first datum is therefore not a model structure chosen for convenience, but a semantic declaration of which changes preserve decision content.

Definition 9.12 (Homotopical decision target). A homotopical decision target is a relative category $(\mathcal{D}, W)$, where W is a declared class of semantic weak equivalences, together with a chosen localization

$$\mathcal{D}_\infty = \mathcal{D}[W^{-1}]$$

that admits the required Kan extensions. A morphism belongs to W only when the declaration says that it preserves the observable decision content; representational similarity alone is insufficient.

A simplicial model category, a category enriched in Kan complexes, or a direct presentation as an ∞ -category can realize this declaration. For a locally Kan simplicial category, its homotopy-coherent nerve replaces strictly unique compositions by coherently compatible spaces of composition. The ordinary nerve $N\mathcal{H}$ is recovered when every mapping space is discrete.

Definition 9.13 (Homotopy-coherent UDL). Let $F : \mathcal{S} \rightarrow \mathcal{D}_\infty$ be a decision decoration, with declared functors $J : \mathcal{S} \rightarrow \mathcal{C}$ and $K : \mathcal{C} \rightarrow \mathcal{Q}$. Its homotopy-coherent UDL semantics is

$$\mathcal{U}^h(F) = \text{Ran}_K^h \text{Lan}_J^h F, \quad (9.5)$$

where the superscript records Kan extension after localization. In a compatible model presentation this is computed by the corresponding total derived Kan extensions.

A homotopy-coherent online UDL is a functor

$$\mathbf{F} : N\mathbb{T} \longrightarrow \text{Fun}(\mathcal{S}, \mathcal{D}_\infty)$$

whose value at t is adapted to the information available at t . Persistent homotopy-coherent UDL retains the entire induced functor $\mathcal{U}^h \mathbf{F} : N\mathbb{T} \rightarrow \text{Fun}(\mathcal{Q}, \mathcal{D}_\infty)$, including its higher coherence data, rather than only its objects in the homotopy category.

The fixed common domain $\mathcal{S}$ in the definition is obtained from growing simplex categories by the extensions used in section 9.11. A functor out of $N\mathbb{T}$ encodes not only prefix transports but the coherent homotopies relating their composites. Passing immediately to the homotopy category would erase this information.

Proposition 9.14 (Homotopy Kan invariance). *If $F, F' : \mathcal{S} \rightarrow \mathcal{D}_\infty$ are related by an objectwise equivalence and the displayed homotopy Kan extensions exist, then*

$$\mathcal{U}^h(F) \simeq \mathcal{U}^h(F').$$

The same statement holds naturally for equivalences of persistent online diagrams.

Proof. Left and right Kan extension are functors between the corresponding ∞ -categories of diagrams. Functors preserve equivalences. Applying first Lan_J^h and then Ran_K^h gives the asserted equivalence; functoriality in $N\mathbb{T}$ gives the persistent statement. $\square$

This is the homotopical strengthening of online Kan invariance. It identifies strictly different implementations when the declared localization regards them as the same decision semantics, while preserving all higher witnesses of that identification.

9.12.1 Horn fillers as repair spaces

Let $p : E \rightarrow X$ be a simplicial family of decorated decisions. Suppose $e_\Lambda : \Lambda_i^n \rightarrow E$ is a partial decorated history and $\sigma : \Delta^n \rightarrow X$ is a proposed underlying history satisfying $pe_\Lambda = \sigma|_{\Lambda_i^n}$. The derived space of repairs is the homotopy fiber

$$\text{Fill}_p(e_\Lambda, \sigma) = \text{hofib}_{(e_\Lambda, \sigma)} \left[\text{Map}(\Delta^n, E) \longrightarrow \text{Map}(\Lambda_i^n, E) \times_{\text{Map}(\Lambda_i^n, X)} \text{Map}(\Delta^n, X) \right]. \quad (9.6)$$

Definition 9.15 (Homotopical repair profile). The homotopical repair profile of a registered missing face is the weak homotopy type of $\text{Fill}_p(e_\Lambda, \sigma)$. An empty repair space records incompatibility, distinct connected components record qualitatively different repair classes, and a contractible repair space records a repair canonical up to coherent homotopy.

A two-component inner-horn example. Consider the unresolved inner horn Λ_1^2 , consisting of two observed one-step decisions

$$x_0 \xrightarrow{a} x_1 \xrightarrow{b} x_2$$

but no registered composite edge $x_0 \rightarrow x_2$. In a strict categorical nerve the missing edge must be $b \circ a$, with a unique 2-simplex. A learned decision object need not yet identify a strict composite. Suppose its left-stage extension instead generates two candidate fillers:

$$(c_s, \alpha_s) \quad \text{and} \quad (c_r, \alpha_r),$$

where c_s is a safe composite, c_r is a risky composite, and α_s, α_r are coherent witnesses that each candidate completes the observed horn. Let W_s and W_r denote their respective witness spaces.

Proposition 9.16 (Two-component inner-horn repair). *Assume W_s and W_r are nonempty and contractible, with no path between their components. Before a distinguishing constraint is revealed, the repair space is*

$$\text{Fill}_1 \simeq W_s \sqcup W_r \simeq S^0.$$

Let later evidence supply a consistency map

$$q : \text{Fill}_1 \longrightarrow \{0, 1\}, \quad q(W_s) = 1, \quad q(W_r) = 0,$$

and require the value 1. Then the right-stage consistency construction is the homotopy fiber

$$\text{Fill}_2 = \text{hofib}_1(q) \simeq W_s \simeq *,$$

and the risky component is removed.

Proof. The two candidate composites determine two disjoint components of the lifting space. Contractibility of their witness spaces identifies their coproduct up to weak equivalence with the discrete two-point space S^0 . The homotopy fiber of q over 1 contains precisely W_s , which is contractible by assumption. $\square$

This example realizes the UDL factorization directly:

$$\Lambda_1^2\text{-decoration} \xrightarrow{\text{Lan}^h} W_s \sqcup W_r \xrightarrow{\text{Ran}^h} \text{hofib}_1(q).$$

Candidate generation does not prematurely choose between the two composites; consistency removes the component contradicted by later evidence. The natural refinement map points backward,

$$\text{Fill}_2 \hookrightarrow \text{Fill}_1,$$

just as the online Kan quotients form an inverse system. A forward persistent update therefore requires a registered repair span or a declared collapse; it is not an invertible transport of the earlier semantics.

The example also separates structural ambiguity from numerical regret. An observer may assign c_s and c_r the same currently observed loss, making them numerically indistinguishable while $\pi_0(\text{Fill}_1)$ still has two elements. After the new constraint, the change $S^0 \rightsquigarrow *$ deletes a component. No infinitesimal tangent path crosses between the two components, so this update is a repair event, not an ordinary gradient or tangent correction within a fixed decision branch.

Thus repair is not forced into a binary success/failure flag or a scalar penalty. Inner horns describe missing composition; outer horns can encode attempts to reverse a decision and need not be fillable. The bar and cobar constructions provide calculational models for homotopy-coherent colimit and limit behavior under the appropriate enriched hypotheses [Riehl, 2014]. This suggests a concrete computational division: bar-like resolution for left-stage candidate generation and cobar-like resolution for right-stage consistency, without asserting that either resolution is unique before the enrichment is declared.

9.12.2 Derived skeletal continuity

Return to the skeletal filtration and interpret all diagrams in $\mathcal{D}_\infty$. Put

$$F \simeq \text{hocolim}_t \tilde{F}_t, \quad G_t = \text{Lan}_J^h \tilde{F}_t, \quad G = \text{Lan}_J^h F.$$

There is a canonical derived comparison

$$\gamma^h : \operatorname{hocolim}_t \operatorname{Ran}_K^h G_t \longrightarrow \operatorname{Ran}_K^h G. \quad (9.7)$$

Theorem 9.17 (Finite-consistency derived skeletal continuity). *Let $\mathcal{D}_\infty$ be a presentable stable ∞ -category, and assume the declared homotopy Kan extensions exist. If every comma ∞ -category indexing the pointwise right Kan extension along K is represented by a finite simplicial set, then the comparison γ^h in equation (9.7) is an equivalence. Consequently, filtered online construction recovers the homotopy-coherent all-at-once UDL semantics.*

Proof. The homotopy left Kan extension is a left adjoint, so it preserves the filtered homotopy colimit:

$$G \simeq \operatorname{Lan}_J^h \left(\operatorname{hocolim}_t \tilde{F}_t \right) \simeq \operatorname{hocolim}_t G_t.$$

In a stable ∞ -category, finite limits and finite colimits agree. Hence filtered colimits commute with the finite limits appearing in the pointwise formula for Ran_K^h . Moving the homotopy colimit through those limits identifies the source and target of γ^h . $\square$

In this stable setting define the homotopy-coherent online defect by

$$\operatorname{Def}^h(F) = \operatorname{fib}(\gamma^h).$$

It is a zero object exactly when γ^h is an equivalence. Mapping objects out of probes into $\operatorname{Def}^h(F)$ reveal components and higher coherences of the failure. This defect precedes regret: a numerical observer may evaluate it, but cannot reconstruct the repair structure after reducing it to a scalar.

Finally, tangent lifting must itself be derived. If the tangent functor preserves W , or admits a derived functor T^h , the fundamental comparison becomes

$$\theta_F^h : T^h \mathcal{U}^h(F) \longrightarrow \mathcal{U}^h(T^h F). \quad (9.8)$$

Determining when θ_F^h is an equivalence, and how its fiber interacts with $\operatorname{Def}^h(F)$, is the first tangent-homotopy problem for UODL. It asks whether infinitesimal variation respects learned semantics only up to coherent deformation, a substantially stronger requirement than equality of point estimates.

9.13 Observers and the origin of regret

An online decision and a hindsight comparator inhabit different information semantics. The online branch selects x_t from the principal past $\downarrow t$. The comparator branch sees the complete horizon but is restricted

to a declared class, such as one action held fixed at every time. Neither branch alone is a regret. A comparison exists only after an observer places their evaluated outcomes in a common value type.

Let $\mathbb{T}_T = \{1 < \dots < T\}$, let X be an action object, and let $\pi : \mathbb{T}_T \rightarrow 1$ be the terminal functor. Precomposition gives the constant-diagram functor

$$\Delta = \pi^* : \mathcal{D} \longrightarrow [\mathbb{T}_T, \mathcal{D}],$$

with right adjoint

$$\Delta \dashv \text{Ran}_\pi = \lim_{\mathbb{T}_T}.$$

Proposition 9.18 (Static comparators as right-Kan-consistent sections). *Assume $\mathbb{T}_T$ is nonempty and connected and the required limits exist. A trajectory object $c \in [\mathbb{T}_T, \mathcal{D}]$ is isomorphic to a constant section precisely when the counit*

$$\epsilon_c : \Delta \text{Ran}_\pi c \longrightarrow c$$

is an isomorphism. Thus the static comparator declaration is the full subcategory on the right-Kan-consistent trajectories.

Proof. If $c \cong \Delta u$, connectedness gives $\text{Ran}_\pi c \cong u$, and the counit is the constant-diagram isomorphism. Conversely, an invertible counit exhibits $c \cong \Delta \text{Ran}_\pi c$, so c is isomorphic to a constant section. $\square$

For a fixed action carrier X , a comparator trajectory is a natural section $\Delta 1 \rightarrow \Delta X$. Connectedness forces all of its temporal components to be the same generalized element $u : 1 \rightarrow X$. The right Kan condition therefore declares staticness; minimizing cumulative loss over those sections is a subsequent decision readout.

Now let $(V, \oplus, 0, \leq)$ be an ordered commutative value monoid, let e_t evaluate the revealed information and an action in V , and let $\sigma_T : V^T \rightarrow V$ be the declared accumulation map. The two branches produce

$$L_T^{\text{on}} = \sigma_T(e_1(F_1, x_1), \dots, e_T(F_T, x_T)), \quad (9.9)$$

$$L_T^{\text{stat}} = \inf_{u : \Delta u \text{ static}} \sigma_T(e_1(F_T, u), \dots, e_T(F_T, u)). \quad (9.10)$$

The first line is a UODL quantity because each x_t is adapted to F_t . The second is a UDL quantity computed from the full diagram F_T , with right-Kan consistency restricting the comparator trajectory to a constant section. It is a hindsight semantic object, not an executable online policy.

Definition 9.19 (Online comparison observer). An online comparison observer is a declared morphism

$$\Omega : V_{\text{on}} \times V_{\text{ref}} \longrightarrow W$$

whose two inputs are the cumulative values of an adapted UODL branch and a full-information UDL reference branch. Its output is the observable comparison object.

When V is an ordered abelian group and $W = V$, the numerical regret observer is

$$\Omega(a, b) = a - b, \quad R_T = \Omega(L_T^{\text{on}}, L_T^{\text{stat}}).$$

If V is only ordered, the observer may return the proposition $b \leq a$. In a residuated value object it may instead return a residual. Thus subtraction is not supplied by UDL or UODL themselves; it is extra observer structure.

This factorization also separates two uses of “second term.” In the regret difference, L_T^{stat} is the full-information UDL term. In the FTRL action objective, by contrast, the regularizer is a mechanism term. It stabilizes the adapted action readout and is not the hindsight comparator. Standard static regret therefore has the typed architecture

$$\begin{array}{ccccc} \text{prefix evidence} & \longrightarrow & \text{UODL actions} & \longrightarrow & L_T^{\text{on}} \\ \text{full evidence} & \longrightarrow & \text{UDL constant comparator} & \longrightarrow & L_T^{\text{stat}} \end{array} \xrightarrow{\Omega} R_T.$$

Changing the reference branch changes the observer semantics: allowing a time-varying full-information comparator produces dynamic rather than static regret, even when the online learner is unchanged.

9.14 Information-structured regret beyond the classical chain

The preceding construction still uses the conventional OCO clock. In fact, ordinary external regret contains a strong but usually implicit information declaration. At round t , the learner acts from a nested information field

$$\mathcal{I}_1 \subseteq \mathcal{I}_2 \subseteq \cdots \subseteq \mathcal{I}_T,$$

while the hindsight optimizer is chosen after the entire loss sequence has been revealed. In Witsenhausen’s terminology this is a classical information pattern [Witsenhausen, 1971, 1975]. The linear order, nested information, and additive accumulation of losses are three distinct assumptions; none belongs to the definition of online UDL.

Let A be a finite set of decision sites. Fix action objects U_α , local value objects V_α , and an aggregate value object V . An information structure $\mathcal{I} = (\mathcal{I}_\alpha)_{\alpha \in A}$ determines a class $\Gamma(\mathcal{I})$ of information-measurable strategy profiles. A profile γ induces its own closed-loop history h^γ , and hence local values

$$v_\alpha^\gamma = \ell_\alpha(h^\gamma, u_\alpha^\gamma).$$

Choose an aggregation morphism

$$\text{Agg}_A : \prod_{\alpha \in A} V_\alpha \longrightarrow V.$$

Addition in $\mathbb{R}$ is one possible aggregation law, not the universal meaning of cumulative performance.

Suppose $\mathcal{J}$ is a declared comparator information structure on the same decision sites and action types. When $\mathcal{I}_\alpha \subseteq \mathcal{J}_\alpha$ for every α , the comparator has at least the information of the learner, and $\Gamma(\mathcal{I}) \subseteq \Gamma(\mathcal{J})$. More general comparisons may be expressed by a typed morphism between information structures, but no canonical numerical comparison follows merely from its existence.

Definition 9.20 (Information-relative intrinsic regret). For an executed profile $\pi \in \Gamma(\mathcal{I})$, define

$$\begin{aligned} L_{\mathcal{M}}(\pi) &= \text{Agg}_A((v_\alpha^\pi)_{\alpha \in A}), \\ L_{\mathcal{M}}^*(\mathcal{J}) &= \inf_{\gamma \in \Gamma(\mathcal{J})} \text{Agg}_A((v_\alpha^\gamma)_{\alpha \in A}). \end{aligned}$$

Given a comparison observer $\Omega : V_{\text{on}} \times V_{\text{ref}} \rightarrow W$, the *information-relative intrinsic regret* is

$$\text{Reg}_{\mathcal{I} \rightarrow \mathcal{J}}^\Omega(\pi) = \Omega(L_{\mathcal{M}}(\pi), L_{\mathcal{M}}^*(\mathcal{J})).$$

Every comparator profile is evaluated on its own induced closed-loop history h^γ .

The superscript Ω is essential. If V is an ordered abelian group, $\Omega(a, b) = a - b$ gives numerical regret. In other targets the observer may produce an order proposition, a residual, a vector of local defects, or a homotopy fiber. Varying $\mathcal{J}$ produces an *information-regret spectrum*: performance relative to static, full-information, partially informed, or causally implementable references.

Theorem 9.21 (Classical-chain recovery of OCO regret). *Let $A = [T]$, let X be a convex decision set, and let $\ell_1, \dots, \ell_T : X \rightarrow \mathbb{R}$ be an exogenously fixed sequence of convex losses. Suppose*

1. x_t is measurable with respect to the nested past-information field generated by $\ell_1, \dots, \ell_{t-1}$;
2. $\Gamma(\mathcal{J})$ is restricted to constant sections Δx selected with the completed loss sequence;
3. $\text{Agg}_{[T]}$ is addition and $\Omega(a, b) = a - b$.

Then theorem 9.20 is exactly the standard static external regret

$$\text{Reg}_{\mathcal{I} \rightarrow \mathcal{J}}^\Omega(\pi) = \sum_{t=1}^T \ell_t(x_t) - \min_{x \in X} \sum_{t=1}^T \ell_t(x).$$

Proof. The first condition identifies the executed profile with an adapted OCO trajectory. By the second condition, optimization over $\Gamma(\mathcal{J})$ is optimization over one action held constant across all rounds; exogeneity makes its loss sequence independent of the learner's trajectory. Substituting additive aggregation and the difference observer into theorem 9.20 gives the displayed expression. $\square$

The theorem locates OCO precisely: it is the constant-comparator corner of intrinsic regret over a classical chain. Other Witsenhausen classes yield different online semantics.

Intrinsic class	Online semantics	Natural reference question
Static	observations do not depend on actions	best decentralized rule on the same exogenous data
Classical	totally ordered, nested histories	best static or dynamic hindsight section
Partially nested	causal sites inherit information along influence paths	best comparator respecting the same inheritance
Nonclassical	an action changes downstream information without transmitting its full basis	best closed-loop policy under a declared information pattern
Nonsequential	events form a partial order or asynchronous dependency shape	best strategy over causal cuts or admissible linearizations
Team	sites share one objective	common team value relative to a comparator information structure
Game	players have distinct objectives	player-indexed regrets or an equilibrium-defect object

Table 9.1: Witsenhausen’s classification generates a family of online-learning problems. The objective axis (team or game) is independent of the information axis (static, classical, partially nested, nonclassical, or nonsequential).

Proposition 9.22 (Nonclassical replay obstruction). *Suppose $\beta \rightsquigarrow \alpha$: changing u_β may change the observation y_α . If $\mathcal{I}_\beta \not\subseteq \mathcal{I}_\alpha$, then, in general, a counterfactual comparator profile γ cannot be evaluated by replacing the learner’s realized actions while holding its realized observation history fixed. Such a replay need not satisfy the intrinsic closed-loop equations. An intrinsic regret comparison must instead evaluate γ on h^γ , unless an additional coupling or invariance assumption is declared.*

Proof. Choose profiles π and γ that differ at β . By the influence assumption, they may induce $y_\alpha(h^\pi) \neq y_\alpha(h^\gamma)$. The downstream action under γ must be $u_\alpha^\gamma = \gamma_\alpha(y_\alpha(h^\gamma))$. Evaluating instead at the learner’s realized observation produces $\gamma_\alpha(y_\alpha(h^\pi))$, which is generally a different action and belongs to neither declared closed-loop trajectory. The resulting hybrid sequence therefore has no intrinsic-model semantics in general. Evaluating each profile on its own closed-loop solution restores a well-typed comparison. $\square$

This obstruction is invisible in standard OCO because losses are exogenous and the past-information fields form a classical chain. Under nonclassical information, the relevant benchmark is policy-level rather than merely action-level: changing a decision may change what can subsequently be known.

Categorically, time should therefore be replaced by an *information category* $\mathbb{A}_{\mathcal{I}}$ of decision sites and admissible causal dependencies. A classical OCO horizon $[T]$ is the special case whose nerve has one maximal temporal simplex. A partially ordered or nonsequential system generally has several maximal simplices representing compatible executions. The online branch and comparator branch are extensions over this nerve, and intrinsic regret is an observer of their evaluated comparison. This yields the guiding principle

online learning is indexed by information,
not necessarily by a global clock.

The cumulative sum in Hazan's formulation remains an important observer for the classical-chain case. UODL treats it as one specialization inside a larger theory of information-structured online decision learning.

9.15 *Running example: crossword solving as universal completion*

A crossword puzzle gives a finite, visible instance of decentralized online decision making. Assign one agent to every across or down clue. Agent α observes its clue text, the declared answer length m_α , and whatever crossing letters have currently been revealed. Its action space is not an untyped vocabulary but a length-indexed candidate object

$$C_\alpha \subseteq \Sigma^{m_\alpha},$$

where Σ is the alphabet. The grid supplies the information and compatibility architecture independently of any particular solving algorithm.

Let E be the set of crossing cells. If crossing e joins clue α at position $p(\alpha, e)$ to clue β at position $p(\beta, e)$, evaluation at that position gives maps

$$\rho_{\alpha e} : C_\alpha \longrightarrow \Sigma, \quad \rho_{\beta e} : C_\beta \longrightarrow \Sigma.$$

The compatible answer pairs at this one crossing form the pullback

$$C_\alpha \times_\Sigma C_\beta.$$

To assemble all crossings at once, form the clue-crossing incidence category $\mathbb{X}$. Its objects are the clues and crossing cells, and it has one arrow $\alpha \rightarrow e$ for each incidence. Define the crossword diagram $D : \mathbb{X} \rightarrow \mathbf{Set}$ by

$$D(\alpha) = C_\alpha, \quad D(e) = \Sigma, \quad D(\alpha \rightarrow e) = \rho_{\alpha e}.$$

Definition 9.23 (Crossword solution object). The *crossword solution object* is the limit

$$\text{Sol}(D) = \lim_{\mathbb{X}} D. \quad (9.11)$$

Equivalently,

$$\text{Sol}(D) \cong \left\{ (w_\alpha)_\alpha \in \prod_\alpha C_\alpha \mid \begin{array}{l} \rho_{\alpha e}(w_\alpha) = \rho_{\beta e}(w_\beta), \\ \text{for every crossing } e = \alpha \cap \beta \end{array} \right\}.$$

Proposition 9.24 (Universal property of crossword completion). *For every set Z equipped with a compatible family of maps $Z \rightarrow C_\alpha$ and $Z \rightarrow \Sigma$ forming a cone over D , there is a unique map $Z \rightarrow \text{Sol}(D)$ through which the entire family factors. Consequently, the points of $\text{Sol}(D)$ are exactly the globally consistent crossword fillings.*

Proof. This is the defining universal property of the limit in equation (9.11). In **Set**, a point of the limit selects one word from every C_α , while commutativity of its cone forces the two selected words incident at each crossing to have the same projected letter. Conversely, every such compatible selection determines a unique point of the limit. $\square$

The proposition identifies the universal part of solving. It does not assert that $\text{Sol}(D)$ is inhabited or has only one point. Nonempty local candidate sets, and even nonempty pullbacks at every crossing considered separately, need not produce a nonempty global limit because the locally compatible pairs may make incompatible choices elsewhere. An empty limit is a global contradiction; several points encode unresolved ambiguity; a singleton is a uniquely completed puzzle relative to the registered candidates.

This is also an exact UDL decomposition. A left stage interprets clue evidence and generates or enlarges the candidate objects C_α . The right stage is the limit, equivalently the right Kan extension along the terminal functor $\mathbb{X} \rightarrow \mathbf{1}$, which retains only globally compatible families:

$$\text{local clue evidence} \xrightarrow{\text{Lan}} \text{typed candidate words} \xrightarrow{\text{Ran}_{\mathbb{X} \rightarrow \mathbf{1}} = \lim_{\mathbb{X}}} \text{Sol}(D).$$

The universal construction determines the object of feasible global decisions, not which feasible point an agent should prefer.

To make that final choice, attach a semantic score $s_\alpha : C_\alpha \rightarrow V$ to each clue and declare an aggregation and an observer. In the familiar numerical realization one might select

$$\operatorname{argmax}_{(w_\alpha)_\alpha \in \text{Sol}(D)} \sum_\alpha s_\alpha(w_\alpha).$$

The limit supplies consistency; the scoring observer ranks its points. Neither operation substitutes for the other.

The Witsenhausen interpretation is equally direct. The agents share the team objective of completing one grid, but have different information fields

$$\mathcal{I}_\alpha = \{\text{clue } \alpha, m_\alpha, \text{currently revealed crossing letters}\}.$$

The clue-intersection graph replaces a single global clock. Propagating a proposed letter changes the information available to neighboring agents, and reciprocal across-down dependencies can create cycles. The natural protocol is therefore asynchronous or nonsequential. If commitments affect downstream observations without transmitting the evidence that produced them, the information pattern can also exhibit Witsenhausen’s nonclassical feature.

Online revelation produces a changing diagram D_t : definitions are recalled, candidate words are removed or reweighted, and crossing letters become fixed. Candidate restriction typically points from later to earlier sets, matching the inverse variance of the online Kan quotients in theorem 9.7. A failed commitment can make $\text{Sol}(D_t)$ empty; repairing it requires revising a declaration or an earlier answer rather than estimating a better point inside the empty object.

Finally, the example separates tangent change from structural repair. Perturbing semantic scores while the grid, candidate supports, and crossing maps remain fixed belongs to one infinitesimal stratum. Adding a candidate, deleting a word, or changing the incidence diagram crosses a discrete boundary and is not represented by an ordinary tangent vector. With enriched candidate objects, a homotopy limit can retain alternative senses and the witnesses by which local interpretations cohere. Crossword solving will therefore serve throughout the book as a running test of four distinct operations: local generation, global consistency, observer-based selection, and repair under progressive revelation.

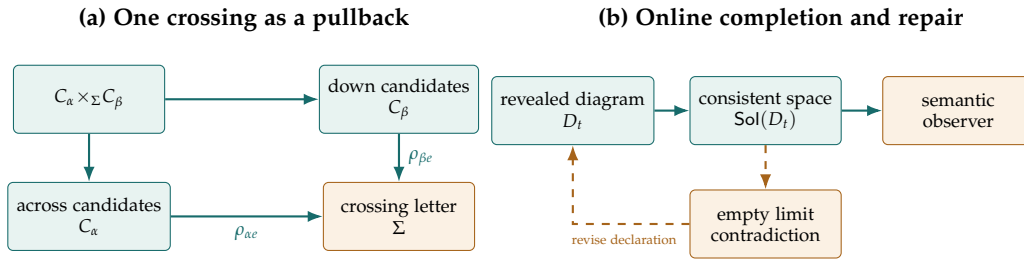


Figure 9.1: The crossword running example separates exact compatibility from selection and repair. A crossing is a pull-back; the whole grid is the limit of the clue-crossing incidence diagram. An observer ranks points only after consistency, while an empty limit sends the solver back to the declaration.

9.16 Running example: Sudoku and flow-based recurrent repair

Sudoku removes most of the semantic ambiguity of a crossword while retaining its local-to-global architecture. Let V be the set of 81 cells, let

$\Delta_9 = \{1, \dots, 9\}$, and let $\mathcal{U}$ be the 27 row, column, and block units. For a puzzle specification c , give each cell the typed domain

$$C_v(c) = \begin{cases} \{c(v)\}, & v \text{ is a given cell,} \\ \Delta_9, & v \text{ is initially blank.} \end{cases}$$

For each unit $S \in \mathcal{U}$, define its allowed-relation object

$$A_S(c) = \left\{ a \in \prod_{v \in S} C_v(c) \mid a : S \longrightarrow \Delta_9 \text{ is bijective} \right\}.$$

Thus a row, column, or block is represented by one nine-variable relation, not by a numerical penalty for repeated digits.

Let S_c be the incidence category with cell objects v , unit objects S , and an arrow $S \rightarrow v$ whenever $v \in S$. The Sudoku constraint diagram $D_c : S_c \rightarrow \mathbf{Set}$ sends S to $A_S(c)$, sends v to $C_v(c)$, and sends an incidence arrow to the corresponding coordinate projection.

Definition 9.25 (Sudoku solution object). The *Sudoku solution object* of the specification c is

$$\text{Sol}_{\text{Sud}}(c) = \lim_{S_c} D_c. \quad (9.12)$$

Proposition 9.26 (Universal Sudoku completion). *The points of $\text{Sol}_{\text{Sud}}(c)$ are in bijection with the valid completions of c . Hence an inconsistent puzzle has an empty solution object, a uniquely solvable puzzle has a singleton solution object, and an underdetermined puzzle has several points.*

Proof. A point of the limit chooses an allowed tuple in $A_S(c)$ for every unit and a value in $C_v(c)$ for every cell. Commutativity of the limiting cone forces every unit tuple to restrict to the same value at each shared cell. These common values form a global grid respecting the givens, and membership in every $A_S(c)$ enforces the row, column, and block constraints. Conversely, every valid grid restricts to such a compatible cone, uniquely. $\square$

Crossword and Sudoku completion are therefore instances of one universal construction with different local relations. Crossword crossings impose mostly binary equality constraints. Sudoku units impose overlapping nine-variable all-different constraints. In both cases, the limit specifies what a globally admissible answer *is*; it does not prescribe an efficient method for finding one.

9.16.1 Flow reasoning as learned fixed-point search

Flow Reasoning Models (FRMs) provide a particularly relevant learned search mechanism for this solution object [Helbling et al., 2026]. They

represent a discrete structured answer by a continuous, directly decodable state and self-condition the denoiser on its previous prediction. For fixed puzzle conditioning c , flow state x_τ , and flow time τ , write the recurrent refinement as

$$s_\tau^{(0)} = \emptyset, \quad s_\tau^{(k+1)} = \Phi_{\theta, c, \tau, x_\tau}(s_\tau^{(k)}). \quad (9.13)$$

The recurrent index k is reasoning depth. It is distinct from the continuous flow coordinate τ , so FRM inference naturally carries the bifiltration (τ, k) anticipated by the decision-nerve semantics.

Let $S_{c, \tau}$ be the recurrent state object and abbreviate the held-state update by $\Phi : S_{c, \tau} \rightarrow S_{c, \tau}$. Its fixed-point object is the equalizer

$$\text{Fix}(\Phi) = \text{Eq}(\Phi, \text{id}_{S_{c, \tau}}).$$

Let $Y_c = \prod_{v \in V} C_v(c)$ be the object of all given-compatible grids, let $d : \text{Fix}(\Phi) \rightarrow Y_c$ decode a recurrent fixed point, and let $i : \text{Sol}_{\text{Sud}}(c) \hookrightarrow Y_c$ be the constraint inclusion.

Definition 9.27 (Fixed-point admission against the universal solution). The recurrent reasoner is *sound* for c when its decoder factors through the universal solution object:

$$\begin{array}{ccc} & \text{Sol}_{\text{Sud}}(c) & \\ \nearrow \bar{d} & & \nwarrow i \\ \text{Fix}(\Phi) & \xrightarrow{d} & Y_c. \end{array}$$

It is *solution-complete* when the restriction of $\bar{d}$ to the declared attracting fixed points covers the solution object. It is *exact* when that comparison is an isomorphism, or an equivalence after the declared semantic localization.

The factorization is an admission condition, not a consequence of recurrent convergence. A state can satisfy $\Phi(s) = s$ and still decode to an invalid grid. This separates two defects that are often conflated:

$\delta_{\text{dyn}}(s)$ measures the discrepancy between $\Phi(s)$ and s ,

$\delta_{\text{con}}(s)$ records violated faces of the Sudoku constraint diagram.

A spurious fixed point has vanishing dynamical defect but nonvanishing constraint defect. Sudoku is valuable precisely because δ_{con} is independently and exactly observable without knowing the intended solution in advance.

9.16.2 Why Fixed-Point Forcing matters to UODL

The FRM paper reports that ordinary self-conditioning becomes unreliable at greater recurrent depth because one-step training states differ

from the model-generated states encountered during inference. Its Fixed-Point Forcing (FPF) procedure constructs the conditioning carry from the model’s own recursive rollout, detaches that carry, and retains the canonical flow-matching supervision path and endpoint objective [Helbling et al., 2026]. In UODL terms, the mechanism is trained on defects generated by its own persistent online dynamics rather than only on teacher-conditioned presentations.

Empirically, the authors report exact solve rates of 99.5% on Sudoku-Extreme, 100.0% on Zebra, and 99.9% on Maze-Unique. On Sudoku-Extreme they match the next-best method’s reported 98.7% peak solve rate with 44 times fewer inference FLOPs [Helbling et al., 2026]. These results make FRM a compelling mechanism for efficient recurrent repair, but they do not identify its fixed-point object with equation (9.12). That is the role of the comparison in theorem 9.27.

The resulting division is sharp:

$$\underbrace{\text{Sol}_{\text{Sud}}(c)}_{\substack{\text{universal constraint object} \\ \text{what counts as a solution}}} \longleftrightarrow \underbrace{\text{Fix}(\Phi)}_{\substack{\text{learned dynamical object} \\ \text{where refinement settles}}} .$$

FRM supplies a learned repair dynamics; UODL supplies the universal object toward which repair ought to converge. Fixed-Point Forcing improves the dynamics by exposing it to its own induced states. The LINC’s admission step still asks whether stability, constraint consistency, and semantic correctness coincide—or where their comparison fails.

Further reading

The conventional online-learning protocol, static regret, and its convex specialization are developed systematically by Shalev-Shwartz [2012] and Hazan [2016]; the expanded second edition of Hazan’s tutorial supplies the milestone sequence used later in this book [Hazan, 2023]. These accounts provide the classical-chain case recovered in theorem 9.21. The separation made here among progressive revelation, comparator information, aggregation, and the comparison observer is the UODL reformulation rather than a claim made by those sources.

For nerves, simplicial sets, inner horns, and quasi-categories, consult Richter [2020] and Riehl and Verity [2022]. Riehl [2014] develops homotopy-coherent diagrams and derived universal constructions at the level needed to replace strict online histories by coherent ones. Those texts supply the homotopical machinery; the use of skeletal growth as online revelation, horn fillers as decision repairs, and semantic localization as persistence is specific to the present framework.

The information-relative account begins with Witsenhausen’s analysis of information structures, feedback, and causality [Witsenhausen,

1971] and his intrinsic formulation of discrete stochastic control [Witsenhausen, 1975]. A categorical extension of the intrinsic model appears in Mahadevan [2021]. Together these works motivate the passage from one global clock to an information category, and from replaying realized observations to evaluating each policy on its own induced closed-loop history.

Finally, Helbling et al. [2026] supplies the learned recurrent mechanism in the Sudoku example and the Fixed-Point Forcing procedure discussed in section 9.16. The crossword and Sudoku limit constructions, the admission map from recurrent fixed points to the universal solution object, and the separation of dynamical from constraint defect are introduced here. They should therefore be read as proposed UODL semantics for constraint reasoning, not as results attributed to the flow-reasoning paper.

Part V

Global-Clock Online Decision Learning

Convex and Regularized Universal Action

DECISION GEOMETRY, FTRL, PROXIMAL STRUCTURE, AND TANGENT WITNESSES

10.1 Universal decision learning

Let $J : \mathcal{O} \rightarrow \mathcal{C}$ embed observed contexts into a larger context category, and let $F : \mathcal{O} \rightarrow \mathcal{D}$ record local decision information. UDL uses a left Kan extension to generate or aggregate global candidates and a right Kan construction to enforce compatibility:

$$F \xrightarrow{\text{Lan}_J} \text{extended candidates} \xrightarrow{\text{Ran}} \text{consistent decisions}.$$

Kan extensions are determined by universal properties rather than a chosen coordinate representation [Kan, 1958, Mac Lane, 1971, Riehl, 2016]. Pointwise left extensions take the familiar form

$$(\text{Lan}_J F)(c) \cong \text{colim}_{(Jd \rightarrow c) \in (J \downarrow c)} F(d).$$

The target category and its enrichment determine what aggregation actually means.

10.2 LINC S discipline

LINC S adds a declaration-and-repair discipline. In the present setting it requires at least four typed blocks:

$$\text{evidence} \longrightarrow \text{universal accumulation} \longrightarrow \text{action mechanism} \longrightarrow \text{decision}.$$

Evidence variations and mechanism variations must not be pooled before diagnosis. A tangent comparison is meaningful only when both routes have the same base point and type. A structural residual can localize a failure, but does not by itself establish a causal intervention or admit a repair.

10.3 The categorical OCO declaration

Let D denote the finite-distribution monad on sets:

$$D(X) = \left\{ p : X \rightarrow \mathbb{R}_{\geq 0} \mid p \text{ has finite support and } \sum_{x \in X} p(x) = 1 \right\}.$$

A convex set is equivalently a D -algebra $\beta_K : D(K) \rightarrow K$; the associated PROP presentation encodes the same convex-combination operations [Haderi et al., 2024]. Ordinary morphisms of convex algebras preserve these operations with equality and are therefore affine. Convex losses require an order-enriched weakening.

Definition 10.1 (Categorical OCO declaration). A finite-horizon categorical OCO declaration consists of:

- (i) a convex decision algebra (K, β_K) , together with a declared metric realization when boundedness, closure, or convergence is used;
- (ii) an ordered convex value algebra $(V, \beta_V, \leq)$;
- (iii) a registered loss family $\ell_t : K \rightarrow V$ satisfying the order-lax algebra law

$$\ell_t \circ \beta_K \leq \beta_V \circ D(\ell_t) \quad \text{pointwise on } D(K); \quad (10.1)$$

- (iv) decision rules $A_t : (\ell_1, \dots, \ell_{t-1}) \mapsto x_t \in K$;
- (v) a comparator declaration, initially the constant decision sections $u \in K$, against which cumulative performance is evaluated; and
- (vi) an admissible tangent assignment. In a finite-dimensional convex realization this is the tangent cone

$$T_K(x) = \overline{\{a(y - x) \mid a \geq 0, y \in K\}}.$$

The additive accumulation of the losses is a subsequent evidence operation; it is not the source of convexity in the declaration.

Proposition 10.2 (Recovery of Euclidean OCO). *Let $K \subseteq \mathbb{R}^d$ be nonempty, bounded, closed, and convex, let β_K evaluate finite convex combinations, and take $V = \mathbb{R}$ with its usual order and barycentric operations. Let the registered losses be bounded real-valued maps on K . Then equation (10.1) holds exactly when each ℓ_t is convex. Consequently, theorem 10.1 recovers the standard OCO protocol*

$$x_t = A_t(\ell_1, \dots, \ell_{t-1}) \in K, \quad \ell_t : K \rightarrow \mathbb{R},$$

with static regret against constant sections

$$R_T = \sum_{t=1}^T \ell_t(x_t) - \inf_{u \in K} \sum_{t=1}^T \ell_t(u).$$

Proof. For a finitely supported probability weight p , equation (10.1) reads

$$\ell_t \left(\sum_x p(x)x \right) \leq \sum_x p(x)\ell_t(x),$$

which is Jensen convexity. The remaining data are precisely the feasible decision, revealed convex loss, history-dependent learner, and fixed comparator class in the Euclidean OCO protocol [Hazan, 2023]. $\square$

Remark 10.3 (Two structures, not one). The convex algebra specifies feasible mixtures of decisions. The $\mathbb{R}$ -linear enrichment used later specifies how revealed quadratic statistics accumulate. Keeping these declarations separate prevents the convexity of K from being confused with addition of losses.

10.4 Regularized universal action

10.5 The FTRL action normal form

At time t , an OCO learner selects $x_t \in K$, observes a convex loss ℓ_t , and incurs $\ell_t(x_t)$. Its static regret against a fixed comparator is

$$R_T = \sum_{t=1}^T \ell_t(x_t) - \inf_{u \in K} \sum_{t=1}^T \ell_t(u).$$

Our first theorem concerns accumulation and action selection, not the right-stage semantics of more general comparator classes. The constant comparator and numerical difference factor through the observer declaration of section 9.13.

A broad FTRL template is

$$x_{t+1} = \arg \min_{x \in K} \left\{ \langle g_{1:t}, x \rangle + A_t^\Psi(x) + S_t(x) \right\}, \quad (10.2)$$

where $g_{1:t}$ is accumulated smooth-loss evidence, A_t^Ψ declares the treatment of a nonsmooth composite term, and S_t declares stabilizing geometry. The minimization is the action readout.

McMahan’s comparison [McMahan, 2011] exposes two independent design axes:

Treatment of Ψ	Origin-centered stabilizer	Proximal stabilizer
Cumulative exact	RDA	FTRL-Proximal
Earlier terms linearized	AOGD	FOBOS

This grid will later support controlled LINCIS interventions. Here we first calibrate the smooth quadratic interior.

10.6 Finite enriched accumulation

10.6.1 The categorical declaration

Fix a horizon T , dimension d , and field $k = \mathbb{R}$. Let

$$\mathcal{P}_T = \{0 < 1 < \dots < T\}$$

be the prefix poset, linearized over k :

$$\mathcal{P}_T(s, t) = \begin{cases} k, & s \leq t, \\ 0, & s > t. \end{cases}$$

Let $\mathcal{O}_T$ be the discrete k -linear category on $\{1, \dots, T\}$, and let $J : \mathcal{O}_T \rightarrow \mathcal{P}_T$ send s to s .

The sufficient-statistic carrier for a quadratic loss is

$$E_d = \text{Sym}_d(k) \oplus k^d \oplus k.$$

For

$$\ell_s(x) = \frac{1}{2} x^\top Q_s x + b_s^\top x + c_s,$$

the revealed statistic is $e_s = (Q_s, b_s, c_s) \in E_d$. Define the carrier functor $F : \mathcal{O}_T \rightarrow \mathbf{Vect}_k$ by $F(s) = E_d$. Since $\mathcal{O}_T$ is discrete, the observations (e_s) form a freely registered local family.

Define the additive realization

$$\sigma_t : \bigoplus_{s \leq t} E_d \longrightarrow E_d, \quad \sigma_t((z_s)_{s \leq t}) = \sum_{s \leq t} z_s.$$

The distinction between the Kan object and σ_t is not cosmetic: the coproduct in vector spaces is the formal direct sum, not the numerical sum of its entries.

10.6.2 The first LINC–UDL theorem

Theorem 10.4 (Finite Kan accumulation and tangent stability). *Under the declaration above:*

(i) *The pointwise enriched left Kan extension exists and*

$$(\text{Lan}_J F)(t) \cong \bigoplus_{s \leq t} E_d.$$

Applying σ_t to the revealed family gives

$$\left(\sum_{s \leq t} Q_s, \sum_{s \leq t} b_s, \sum_{s \leq t} c_s \right).$$

(ii) For a C^1 family $e_s(\theta)$, with fixed finite indexing shape, the canonical comparison

$$\chi_t^L : \bigoplus_{s \leq t} T(E_d) \xrightarrow{\cong} T\left(\bigoplus_{s \leq t} E_d\right)$$

is an isomorphism and is compatible with the fold:

$$T(\sigma_t) \circ \chi_t^L = \sigma_t^T.$$

Consequently,

$$D_\theta \left(\sum_{s \leq t} e_s(\theta) \right) [v] = \sum_{s \leq t} D_\theta e_s(\theta) [v].$$

Proof. The enriched pointwise formula is

$$(\text{Lan}_J F)(t) \cong \int^{s \in \mathcal{O}_T} \mathcal{P}_T(Js, t) \otimes F(s).$$

This is the standard weighted-colimit formula for an enriched pointwise Kan extension [Kelly, 1982]. The source category is discrete, so the coend has no nonidentity relations. Furthermore,

$$\mathcal{P}_T(Js, t) \otimes F(s) \cong \begin{cases} E_d, & s \leq t, \\ 0, & s > t. \end{cases}$$

The coend is therefore the finite coproduct $\bigoplus_{s \leq t} E_d$.

To see the universal property directly, let $G : \mathcal{P}_T \rightarrow \mathbf{Vect}_k$. A transformation $\alpha : F \Rightarrow J^*G$ provides maps $\alpha_s : E_d \rightarrow G(s)$. For each $s \leq t$, compose with $G(s \leq t)$. The coproduct property yields a unique

$$\bar{\alpha}_t : \bigoplus_{s \leq t} E_d \longrightarrow G(t).$$

These maps are natural in t , and their uniqueness is componentwise. They are exactly the factorization required by $\text{Lan}_J \dashv J^*$. Applying the linear fold to the registered local family yields the displayed cumulative statistic.

For the tangent claim, finite-dimensional vector spaces form an additive category and finite coproducts are biproducts. The standard tangent construction is $T(V) \cong V \oplus V$, which preserves finite biproducts [Cockett and Cruttwell, 2014, Kock, 2006]. Hence

$$\bigoplus_{s \leq t} T(E_d) \cong T\left(\bigoplus_{s \leq t} E_d\right).$$

Since σ_t is linear, its tangent folds base and tangent components separately. This proves the commuting comparison. Evaluation on the registered C^1 evidence family gives the finite-sum derivative formula. $\square$

Remark 10.5 (What the theorem actually says). The universal part constructs formal accumulation. Numerical accumulation is a declared algebra on that universal carrier. Omitting the fold would incorrectly identify a coproduct with vector addition.

10.7 Quadratic FTRL and typed tangents

Let the mechanism regularizer be

$$S_\phi(x) = \frac{\lambda}{2}(x - m)^\top M(x - m), \quad \lambda > 0, \quad M \succ 0.$$

Assume $Q_s \succeq 0$. After accumulation, define

$$H_t = \lambda M + \sum_{s \leq t} Q_s, \quad q_t = \sum_{s \leq t} b_s - \lambda M m.$$

Then $H_t \succ 0$, and exact-loss FTRL has the unique decision

$$x_{t+1} = -H_t^{-1} q_t. \quad (10.3)$$

Proposition 10.6 (Typed FTRL action tangents). *At the same base decision, evidence and mechanism variations satisfy*

$$H_t \dot{x}_{t+1}^F = - \left[\sum_{s \leq t} b_s + \left(\sum_{s \leq t} \dot{Q}_s \right) x_{t+1} \right]$$

and

$$H_t \dot{x}_{t+1}^\rho = -(\dot{q}_t^\rho + \dot{H}_t^\rho x_{t+1}),$$

where

$$\dot{H}_t^\rho = \dot{\lambda} M + \lambda \dot{M}, \quad \dot{q}_t^\rho = -\dot{\lambda} M m - \lambda \dot{M} m - \lambda M \dot{m}.$$

Their joint first-order response is $\dot{x}_{t+1}^F + \dot{x}_{t+1}^\rho$.

Proof. The first-order condition is

$$H_t x_{t+1} + q_t = 0.$$

Differentiating gives

$$\dot{H}_t x_{t+1} + H_t \dot{x}_{t+1} + \dot{q}_t = 0.$$

Since H_t is invertible,

$$\dot{x}_{t+1} = -H_t^{-1}(\dot{q}_t + \dot{H}_t x_{t+1}).$$

Holding the mechanism fixed gives the evidence equation. Holding evidence fixed and differentiating λM and $-\lambda M m$ gives the mechanism equation. Linearity of the differential gives the joint response after the two directions are based at the same decision. $\square$

Corollary 10.7 (Constant loss shifts are decision-null). *Evidence variations of the form*

$$((0, 0, \dot{c}_s))_{s \leq t}$$

lie in the kernel of the FTRL decision tangent.

Proof. The accumulated constant affects only the decision-independent scalar term in the quadratic objective. It is absent from the first-order condition and therefore from Proposition 10.6. $\square$

Boundary 10.8 (Ambient versus admissible tangents). The tangent calculation occurs in the ambient vector space $\text{Sym}_d(\mathbb{R})$. If a loss Hessian lies on the boundary of the positive-semidefinite cone, not every ambient direction preserves convexity. An OCO intervention must therefore be restricted to the appropriate tangent cone when convexity preservation is part of the declaration. The readout remains smooth as long as the total H_t stays positive definite.

10.8 Proximal FTRL and tangent repair

The FTRL action mechanism admits a proximal presentation. For a positive-definite metric M , define

$$\text{prox}_{\eta f}^M(y) = \arg \min_z \left\{ \eta f(z) + \frac{1}{2} \|z - y\|_M^2 \right\}, \quad \|w\|_M^2 = w^\top M w.$$

Writing $L_t = \sum_{s \leq t} \ell_s$, the regularized decision

$$x_{t+1} = \arg \min_x \left\{ L_t(x) + \frac{\lambda}{2} \|x - m\|_M^2 \right\}$$

is exactly

$$x_{t+1} = \text{prox}_{\eta L_t}^M(m), \quad \eta = \lambda^{-1}.$$

Under the categorical OCO declaration, the constrained readout is

$$x_{t+1} = \text{prox}_{\eta(L_t + \iota_K)}^M(m),$$

where ι_K is the convex indicator of the feasible decision algebra. The smooth theorem below applies when $K = \mathbb{R}^d$, or locally while the decision remains in the interior of K ; active boundaries belong to the graphical-tangent case. This presentation treats the proximal map as a repair readout that reconciles accumulated evidence with an anchor and a declared geometry.

Theorem 10.9 (Tangent lift of metric-proximal FTRL). *Let $L_t(\cdot; \theta)$ be convex and C^2 in its decision argument, with the relevant derivatives jointly continuous in a parameter θ . Let $M(\theta) \succ 0$, $m(\theta)$, and $\eta(\theta) > 0$ be C^1 , and let*

$$x = \text{prox}_{\eta L_t}^M(m).$$

For a parameter direction v , write dots for directional derivatives and define the evidence variation at fixed decision by

$$\delta(\nabla L_t)(x) = D_\theta [\nabla_x L_t(x; \theta)] [v].$$

Then the proximal decision is locally C^1 , and its tangent is the unique solution of

$$(M + \eta \nabla^2 L_t(x)) \dot{x} = M \dot{m} - \dot{M}(x - m) - \dot{\eta} \nabla L_t(x) - \eta \delta(\nabla L_t)(x). \quad (10.4)$$

For fixed L_t, M, η , the derivative with respect to the anchor is

$$D_m \text{prox}_{\eta L_t}^M(m) = (M + \eta \nabla^2 L_t(x))^{-1} M.$$

It is nonexpansive in the M -norm. If, at the proximal point, $\nabla^2 L_t(x) \succeq \mu M$ for some $\mu > 0$, then

$$\|D_m \text{prox}_{\eta L_t}^M(m)\|_M \leq \frac{1}{1 + \eta \mu} < 1.$$

Proof. The proximal first-order condition is

$$M(x - m) + \eta \nabla L_t(x) = 0.$$

Its decision Jacobian is $M + \eta \nabla^2 L_t(x) \succ 0$, so the implicit function theorem gives the local smooth readout. Differentiating the first-order condition and keeping the evidence variation at fixed x separate gives equation (10.4).

For anchor variation alone, solve the same equation with only $\dot{m}$ nonzero. Conjugating the resulting operator by $M^{1/2}$ gives

$$M^{1/2} (M + \eta \nabla^2 L_t(x))^{-1} M^{1/2} = \left(I + \eta M^{-1/2} \nabla^2 L_t(x) M^{-1/2} \right)^{-1}.$$

Convexity places its eigenvalues in $(0, 1]$. Under the relative strong-convexity bound, they are at most $(1 + \eta \mu)^{-1}$, proving both claims. $\square$

Equation (10.4) gives four separately typed LINC channels: anchor variation $M \dot{m}$, geometry variation $-\dot{M}(x - m)$, regularization-scale variation $-\dot{\eta} \nabla L_t(x)$, and accumulated-evidence variation $-\eta \delta(\nabla L_t)(x)$. In the quadratic case it reduces exactly to Proposition 10.6, after multiplying the equation by λ . The contraction bound is a local stability certificate for the repair readout; it does not by itself control the time variation of L_t .

Boundary 10.10 (Nonsmooth proximal tangents). For a proper closed convex function in Euclidean space,

$$\text{prox}_{\eta f} = (I + \eta \partial f)^{-1},$$

the resolvent of a maximal monotone relation [Rockafellar, 1976]. An ordinary tangent map need not exist at an active-set transition. If

$x = \text{prox}_{\eta f}(y)$, $u = (y - x)/\eta \in \partial f(x)$, and the subgradient mapping is proto-differentiable with a single-valued directional solution, then a direction h is transported to the solution w of

$$h \in w + \eta D(\partial f)(x \mid u)(w).$$

Thus the directional tangent of the proximal repair is the resolvent of the graphical derivative of its defect relation. Establishing the needed second-order epi-regularity is an additional declaration, not a tangent category axiom [Rockafellar, 1989].

Further reading

The standard OCO and regularization background for these deliberately compact OCO chapters is provided by Shalev-Shwartz [2012], Hazan [2023], and McMahan [2011]. The operadic account of convexity is due to Haderi et al. [2024]; enriched Kan extensions are developed by Kelly [1982]. Proximal and graphical-derivative foundations appear in Rockafellar [1976, 1989]. The universal accumulation, typed tangent separation, and decision-theoretic interpretation given here are the UODL contribution.

Universal Bandit Decision Learning

PARTIAL OBSERVATION AS AN INFORMATION RESTRICTION

Bandit feedback changes neither the action object nor the loss class at first. It changes the information channel through which a learner encounters loss. That apparently small change separates three notions that coincide in full information: pathwise recovery, recovery after barycentric averaging, and stability of recovery under infinitesimal variation. This chapter makes those distinctions explicit and then proves when a full-information online guarantee transports through a bandit channel.

11.1 The bandit information channel

Let A be a finite arm object with $K = |A|$, let $\Delta_A = D(A)$ be its free convex algebra of randomized actions, and put $\mathcal{L} = [0, 1]^A$. For $p \in \Delta_A$ and $\ell \in \mathcal{L}$, the stochastic observation channel is the Kleisli arrow

$$\kappa_p(\ell) = \sum_{a \in A} p(a) \delta_{(a, \ell(a))} \in D(A \times [0, 1]). \quad (11.1)$$

The registered evidence is therefore the sampled arm and its scalar loss, not the complete loss vector.

Definition 11.1 (Finite-arm bandit UODL declaration). A finite-arm bandit UODL declaration consists of:

- (i) a history filtration $(\mathcal{F}_t)_{t \geq 0}$;
- (ii) an $\mathcal{F}_{t-1}$ -measurable, full-support action distribution $p_t \in \Delta_A^\circ$;
- (iii) a loss $\ell_t \in \mathcal{L}$ fixed before the current sample $A_t \sim p_t$;
- (iv) the channel κ_{p_t} , together with a declared evidence repair; and
- (v) a comparator sketch and an observer that values the comparison.

An oblivious declaration fixes the loss sequence in advance. A predictable adaptive declaration permits ℓ_t to depend on $\mathcal{F}_{t-1}$. An adversary that observes A_t before choosing ℓ_t lies outside this declaration.

$$\mathcal{L} \xrightarrow{\kappa_p} D(A \times [0, 1]) \xrightarrow{D(R_p)} D(\mathcal{L}) \xrightarrow{\beta_{\mathcal{L}}} \mathcal{L}$$

$1_{\mathcal{L}}$

Figure 11.1: Bandit evidence is not inverted sample by sample. The channel is split only after lifting repaired observations into the distribution monad and applying its barycentric algebra.

For $p \in \Delta_A^\circ$, define the importance repair

$$R_p(a, y)(b) = \frac{y \mathbf{1}\{a = b\}}{p(a)}, \quad b \in A. \quad (11.2)$$

Proposition 11.2 (Barycentric reconstruction of bandit evidence). *For every full-support $p \in \Delta_A^\circ$, importance repair is a left inverse to the bandit observation channel after applying the convex-algebra expectation:*

$$\beta_{\mathcal{L}} \circ D(R_p) \circ \kappa_p = 1_{\mathcal{L}}.$$

Equivalently, $\mathbb{E}_{a \sim p}[R_p(a, \ell(a))] = \ell$.

Proof. For each coordinate $b \in A$, barycentric evaluation gives

$$\sum_{a \in A} p(a) \frac{\ell(a) \mathbf{1}\{a = b\}}{p(a)} = \ell(b).$$

The equality holds coordinatewise in the product convex algebra $\mathcal{L}$. $\square$

11.2 Pathwise loss and barycentric recovery

The qualifier “after expectation” is essential. It is the precise price of transporting full information through a one-coordinate observation channel.

Proposition 11.3 (No deterministic pathwise reconstruction). *If $|A| \geq 2$, there is no map $S : A \times [0, 1] \rightarrow \mathcal{L}$ satisfying $S(a, \ell(a)) = \ell$ for every $a \in A$ and every $\ell \in \mathcal{L}$.*

Proof. Fix an arm a and choose $b \neq a$. Two loss vectors may agree at coordinate a and differ at coordinate b . They generate the same observation $(a, \ell(a))$, so a deterministic map must return the same value on both, although exact reconstruction would require two different loss vectors. $\square$

The obstruction is informational, not algorithmic. No more ingenious estimator can make one realized scalar determine all unobserved coordinates. The distribution monad supplies a weaker but sufficient splitting: repaired evidence becomes exact after the observer $\beta_{\mathcal{L}}$ averages over the sampling law.

11.3 Exploration as infinitesimal admissibility

For $\alpha > 0$, write

$$\Delta_A^\alpha = \{p \in \Delta_A : p(a) \geq \alpha \text{ for every } a \in A\}.$$

This is not merely an algorithmic restriction. It is a domain on which the repair map admits controlled tangent variation.

Proposition 11.4 (Tangent and second-moment control). *Along a differentiable curve $\varepsilon \mapsto (p_\varepsilon, y_\varepsilon)$ with $p_0 = p \in \Delta_A^\alpha$, the conditional repaired evidence obeys*

$$\dot{R}_p(a, y)(b) = \mathbf{1}\{a = b\} \left(\frac{\dot{y}}{p(a)} - \frac{y \dot{p}(a)}{p(a)^2} \right).$$

For $0 \leq y \leq 1$,

$$\|\dot{R}_p(a, y)\|_\infty \leq \frac{|\dot{y}|}{\alpha} + \frac{\|\dot{p}\|_\infty}{\alpha^2}.$$

If $\hat{\ell} = R_p(A, \ell(A))$ with $A \sim p$, then

$$\mathbb{E}\|\hat{\ell}\|_2^2 = \sum_{a \in A} \frac{\ell(a)^2}{p(a)}, \quad \text{Var}(\hat{\ell}(b)) = \ell(b)^2 \left(\frac{1}{p(b)} - 1 \right).$$

Proof. Differentiate equation (11.2). The norm bound follows from $p(a) \geq \alpha$. The repaired vector has one nonzero coordinate, so averaging $\ell(A)^2/p(A)^2$ gives the second-moment identity. The variance formula follows from the same coordinatewise calculation and theorem 11.2. $\square$

Consequently, mixing any policy q with a full-support reference law

$$p = (1 - \gamma)q + \gamma u, \quad \gamma > 0,$$

is a tangent-domain repair. If $\alpha = \gamma \min_a u(a)$, then the action mechanism lies in Δ_A^α . The usual exploration–variance tradeoff therefore has a LINCIS interpretation: exploration controls whether the evidence reconstruction is infinitesimally admissible.

Proposition 11.5 (Barycentric tangent reconstruction). *Let $p_\varepsilon \in \Delta_A^\circ$ and $\ell_\varepsilon \in \mathcal{L}$ be differentiable curves. Then*

$$\left. \frac{d}{d\varepsilon} \right|_0 \beta_{\mathcal{L}} D(R_{p_\varepsilon}) \kappa_{p_\varepsilon}(\ell_\varepsilon) = \dot{\ell}.$$

Thus variations of the sampling mechanism and of the loss evidence cancel correctly after barycentric reconstruction, even though neither is stable pathwise near the simplex boundary.

Proof. The composite inside the derivative equals ℓ_ε for every ε by theorem 11.2. Differentiating this identity proves the claim. Coordinate-wise, the derivative of the sampling weight cancels the derivative of its reciprocal in the importance repair. $\square$

11.4 The expectation-level regret-transfer theorem

The previous results determine the observer through which a full-information theorem may be transported. Let $\widehat{\ell}_t = R_{p_t}(A_t, \ell_t(A_t))$.

Theorem 11.6 (Bandit regret transfer through barycentric repair). *Let $\ell_1, \dots, \ell_T \in [0, 1]^A$ be an oblivious loss sequence, and let $p_t \in \Delta_A^\circ$ be adapted to $\mathcal{F}_{t-1}$. Suppose a full-information online rule applied to $\widehat{\ell}_t$ satisfies, on every realized surrogate sequence and for every fixed u in a comparator class $\mathcal{C} \subseteq \Delta_A$,*

$$\sum_{t=1}^T \langle p_t, \widehat{\ell}_t \rangle - \sum_{t=1}^T \langle u, \widehat{\ell}_t \rangle \leq B_T.$$

Then its bandit realization satisfies

$$\mathbb{E} \left[\sum_{t=1}^T \ell_t(A_t) \right] - \inf_{u \in \mathcal{C}} \sum_{t=1}^T \langle u, \ell_t \rangle \leq B_T.$$

Proof. Pathwise, $\langle p_t, \widehat{\ell}_t \rangle = \ell_t(A_t)$. By adaptedness and barycentric reconstruction,

$$\mathbb{E}[\widehat{\ell}_t \mid \mathcal{F}_{t-1}] = \ell_t, \quad \mathbb{E}\langle u, \widehat{\ell}_t \rangle = \langle u, \ell_t \rangle$$

for every fixed u . Take expectations in the assumed pathwise inequality and then the infimum over $u \in \mathcal{C}$. $\square$

For a predictable adaptive adversary, the same calculation yields pseudo-regret against each comparator fixed independently of the realized loss sequence. It does not by itself justify moving an expectation through a random hindsight infimum. If the adversary sees the current sampled arm before fixing the loss, even conditional barycentric reconstruction fails. These are changes of declaration, not technical footnotes.

11.5 Entropic realization and the classical rate

Take uniform initial weights and update

$$p_t(a) = \frac{w_t(a)}{\sum_b w_t(b)}, \quad w_{t+1}(a) = w_t(a) e^{-\eta \widehat{\ell}_t(a)}.$$

Proposition 11.7 (Entropic surrogate comparison). *For every arm a and every nonnegative surrogate sequence,*

$$\sum_{t=1}^T \langle p_t, \widehat{\ell}_t \rangle - \sum_{t=1}^T \widehat{\ell}_t(a) \leq \frac{\log K}{\eta} + \frac{\eta}{2} \sum_{t=1}^T \langle p_t, \widehat{\ell}_t^2 \rangle.$$

Proof. Apply $e^{-x} \leq 1 - x + x^2/2$ for $x \geq 0$ to the exponential-weights potential, telescope the logarithms of total weight, and compare the final total weight with the contribution of arm a . $\square$

Corollary 11.8 (Adversarial finite-arm bandit rate). *Under the assumptions of theorem 11.6, the entropic realization satisfies*

$$\mathbb{E} \left[\sum_{t=1}^T \ell_t(A_t) \right] - \min_{a \in A} \sum_{t=1}^T \ell_t(a) \leq \frac{\log K}{\eta} + \frac{\eta KT}{2}.$$

With $\eta = \sqrt{2 \log K / (KT)}$, the bound is $\sqrt{2TK \log K}$.

Proof. Conditionally on the past,

$$\mathbb{E} \langle p_t, \widehat{\ell}_t^2 \rangle = \sum_{a \in A} \ell_t(a)^2 \leq K.$$

Take expectations in theorem 11.7, use the displayed second-moment identity and barycentric reconstruction for the fixed comparison arm, and then minimize over arms. Optimizing η gives the stated choice. $\square$

The inverse probabilities disappear from this expected second moment, but they do not disappear from the pathwise tangent bound in theorem 11.4. Sublinear expected regret and robust infinitesimal evidence transport are therefore genuinely different properties, witnessed by different observers.

11.6 Bandit convex optimization changes the repair target

Let $C \subset \mathbb{R}^d$ be convex, let C_δ contain the points whose closed δ -ball lies in C , and define ball smoothing by

$$\tilde{f}_\delta(z) = \frac{1}{\text{vol}(B^d)} \int_{B^d} f(z + \delta v) dv.$$

Here a one-point scalar observation cannot reconstruct f . The useful repair target is instead a tangent covector of the smoothed loss.

Proposition 11.9 (One-point reconstruction of a smoothed gradient). *Let U be uniform on the unit sphere S^{d-1} , and observe $f(z + \delta U)$ for $z \in C_\delta$. Then*

$$G_\delta(z, U) = \frac{d}{\delta} f(z + \delta U) U \quad \text{satisfies} \quad \mathbb{E}[G_\delta(z, U)] = \nabla \tilde{f}_\delta(z)$$

whenever the displayed derivative exists.

Proof. Differentiate the ball average and apply the divergence theorem. The boundary integral over S^{d-1} introduces the outward normal U ; the ratio of sphere area to ball volume is d , yielding the displayed identity. $\square$

Thus bandit convex optimization factors through a different reconstruction diagram: scalar evidence is sent to a distribution of tangent

covectors and then averaged. Approximation bias enters because $\tilde{f}_\delta \neq f$; sampling variance enters through the norm of G_δ . UODL keeps these two defects typed separately. The broader theorem problem is to characterize distribution monads and tangent structures for which barycentric reconstruction commutes with the relevant tangent lift.

Further reading

The adversarial finite-arm construction and its classical regret analysis originate with [Auer et al. \[2002\]](#). The one-point smoothed-gradient method for bandit convex optimization is due to [Flaxman et al. \[2005\]](#). For a comprehensive modern treatment spanning stochastic, adversarial, Bayesian, contextual, linear, and combinatorial bandits, see [Lattimore and Szepesvári \[2020\]](#). [Hazan \[2023\]](#) develops both finite-arm and convex bandit optimization within the broader OCO landscape. The information-channel, barycentric splitting, tangent-admissibility, and observer distinctions in this chapter are the UODL reformulation.

Part VI

Information Beyond a Global Clock

Information Categories and Decentralized Decisions

INTRINSIC COMPARISON, ASYNCHRONOUS EVENTS, LOCAL DECISIONS, AND GLOBAL CONSISTENCY

Chapter 9 replaced the global clock by an information category and showed why a counterfactual policy must be evaluated on its own induced closed-loop history. We now make the comparator itself structural. The basic move is *descent*: declare which distinctions in the information category a reference policy is allowed to retain, and ask whether a trajectory descends through that declaration.

12.1 Comparator sketches as descent data

Let $\mathbb{A}$ be a small information category, let $\mathbb{R}$ be a reference shape, and let

$$q : \mathbb{A} \longrightarrow \mathbb{R}$$

identify distinctions that the comparator must ignore. For a complete target category $\mathcal{D}$, precomposition has a right adjoint

$$q^* : [\mathbb{R}, \mathcal{D}] \rightleftarrows [\mathbb{A}, \mathcal{D}] : \text{Ran}_q.$$

A reference diagram $r : \mathbb{R} \rightarrow \mathcal{D}$ is realized on the original information category as $q^*r = r \circ q$.

Definition 12.1 (Comparator sketch). A *comparator sketch* in $\mathcal{D}$ is a functor $q : \mathbb{A} \rightarrow \mathbb{R}$ together with a replete full subcategory

$$\mathcal{K} \hookrightarrow [\mathbb{R}, \mathcal{D}]$$

of admissible reference diagrams. Its realized comparator class is the essential image of

$$q^*|_{\mathcal{K}} : \mathcal{K} \longrightarrow [\mathbb{A}, \mathcal{D}].$$

The sketch is *right-Kan admissible* when the adjunction unit

$$\eta_r : r \longrightarrow \text{Ran}_q q^*r$$

is an isomorphism for every $r \in \mathcal{K}$.

The functor q controls variation, while $\mathcal{K}$ controls admissibility. These roles are independent. Collapsing all sites may declare a static comparator; preserving an epoch index may declare a block-static one; taking $q = 1_{\mathbb{A}}$ permits fully dynamic trajectories. Constraints on actions, policies, or resources belong to $\mathcal{K}$, not to the quotient shape by itself.

For any trajectory $c : \mathbb{A} \rightarrow \mathcal{D}$, the counit

$$\varepsilon_c : q^* \text{Ran}_q c \longrightarrow c \quad (12.1)$$

compares its canonical right-Kan reconstruction with the trajectory that was actually presented. This is the intrinsic comparison map of the sketch.

Theorem 12.2 (Right-Kan representation of comparator sketches). *Let $(q, \mathcal{K})$ be a right-Kan-admissible comparator sketch and assume the displayed right Kan extensions exist. A trajectory $c \in [\mathbb{A}, \mathcal{D}]$ belongs to the realized comparator class if and only if*

- (i) $\text{Ran}_q c$ belongs to $\mathcal{K}$, and
- (ii) the counit $\varepsilon_c : q^* \text{Ran}_q c \rightarrow c$ is an isomorphism.

Consequently, the comparator class is characterized internally on $\mathbb{A}$, without choosing a coordinate representation of a reference policy.

Proof. Suppose $c \cong q^*r$ for $r \in \mathcal{K}$. Right-Kan admissibility gives $r \cong \text{Ran}_q q^*r$, hence $\text{Ran}_q c \cong r \in \mathcal{K}$ because $\mathcal{K}$ is replete. The triangle identity

$$\varepsilon_{q^*r} \circ q^* \eta_r = 1_{q^*r}$$

and invertibility of η_r show that ε_{q^*r} , and therefore ε_c , is invertible.

Conversely, if $\text{Ran}_q c \in \mathcal{K}$ and ε_c is invertible, then it exhibits

$$c \cong q^*(\text{Ran}_q c),$$

so c lies in the essential image of $q^*|_{\mathcal{K}}$. □

Remark 12.3 (The static-comparator case). For a nonempty connected $\mathbb{A}$, take $q = \pi : \mathbb{A} \rightarrow 1$ and $\mathcal{K} = \mathcal{D}$. Then $q^* = \Delta$, and the representation theorem reduces to Proposition 9.18. Thus the familiar static comparator is the coarsest member of a much larger descent semantics.

12.2 The comparator spectrum on a finite horizon

Let $[T]$ be a finite chain and partition it into nonempty contiguous blocks

$$B_1 < \cdots < B_m.$$

Write $q_B : [T] \rightarrow [m]$ for the monotone map sending a time to its block.

Proposition 12.4 (Block-static comparator representation). *For every category $\mathcal{D}$, precomposition*

$$q_B^* : [[m], \mathcal{D}] \longrightarrow [[T], \mathcal{D}]$$

is fully faithful. If the pointwise right Kan extension exists, a trajectory $c : [T] \rightarrow \mathcal{D}$ is in its essential image exactly when every transition $c(s \leq t)$ with $s, t \in B_j$ is an isomorphism.

Proof. A natural transformation between two pulled-back diagrams is constant on each block and is determined uniquely by its components at one representative of every block. Naturality on $[T]$ supplies exactly the naturality conditions on $[m]$, so q_B^* is fully faithful.

If $c \cong q_B^* r$, all within-block arrows are identities up to the displayed natural isomorphism. Conversely, choose the first time b_j in each block and define $r(j) = c(b_j)$, with transition maps inherited from the unique arrows of $[T]$. When all within-block transitions are invertible, the maps $c(b_j \leq t)$ give a natural isomorphism $q_B^* r \cong c$. $\square$

Refinement of partitions now orders comparator power. If B' refines B , there is a unique monotone map u with

$$q_B = u \circ q_{B'}, \quad q_B^* = q_{B'}^* u^*.$$

Hence every B -block-static comparator is also B' -block-static:

$$\text{Comp}_{\text{static}} \subseteq \text{Comp}_B \subseteq \text{Comp}_{B'} \subseteq \text{Comp}_{\text{dynamic}}.$$

The endpoints are the one-block partition and the singleton partition. This is an information-regret spectrum generated by refinement of reference shape, not by changing the learner.

Corollary 12.5 (Monotonicity of scalar comparator regret). *In a real-valued minimization problem, let $\mathcal{C} \subseteq \mathcal{C}'$ be comparator classes induced by a sketch refinement. Whenever both infima exist,*

$$\inf_{c' \in \mathcal{C}'} L(c') \leq \inf_{c \in \mathcal{C}} L(c),$$

and therefore

$$L^{\text{on}} - \inf_{c \in \mathcal{C}} L(c) \leq L^{\text{on}} - \inf_{c' \in \mathcal{C}'} L(c').$$

A stronger comparator can only increase the resulting external-regret number. The inclusion of comparator objects is universal; the inequalities appear only after applying the ordered numerical observer.

12.3 Approximate descent and observer transfer

When ε_c is not invertible, category theory identifies a failure of descent but supplies no canonical magnitude. A numerical bound requires additional observer structure.

Proposition 12.6 (Observer transfer from an approximate comparator). *Suppose a realization of $[\mathbb{A}, \mathcal{D}]$ carries a metric d , the cumulative evaluation L is K_L -Lipschitz, and the numerical comparison observer $\Omega(a, -)$ is K_Ω -Lipschitz in its reference argument. Let $\bar{c} = q^* \text{Ran}_q c$ be admissible and suppose the realized counit obeys*

$$d(c, \bar{c}) \leq \delta.$$

Then

$$|\Omega(L^{\text{on}}, L(c)) - \Omega(L^{\text{on}}, L(\bar{c}))| \leq K_\Omega K_L \delta.$$

Proof. Lipschitz continuity of L gives $|L(c) - L(\bar{c})| \leq K_L \delta$. Applying the Lipschitz bound for the second argument of Ω proves the claim. $\square$

The proposition locates the assumptions precisely. The counit supplies the universal comparison; the metric turns its failure into a size; the evaluation transports that size to value space; and the observer transports it to a reported regret defect. None of the last three steps follows from right Kan universality alone.

For endogenous information, this descent theorem is necessary but not yet sufficient. The object c must first be generated on its own closed-loop history. Only then may ε_c test whether the resulting strategy trajectory belongs to the comparator sketch. This preserves the nonclassical replay boundary of Proposition 9.22.

12.4 Running example: asynchronous event categories

The global time used to analyze a computation need not be information available inside that computation. This distinction is explicit in the asynchronous models of Bertsekas and Tsitsiklis [Bertsekas and Tsitsiklis, 1989]. Processors update at different rates and may use values received at different, possibly stale, stages. An analyst can number the resulting events afterward, but that numbering inserts an order between events that were intrinsically incomparable.

Definition 12.7 (Intrinsic asynchronous execution). *An intrinsic asynchronous execution on a finite processor set (P) consists of a finite event poset $\mathbb{E}$, an ownership map $o : \mathbb{E} \rightarrow P$, and, for each event e , a local update map U_e whose arguments are indexed only by the causal past $\downarrow e$. A linearization is a listing $(e_1, \dots, e_m)$ of all events such that $e_i < e_j$ in $\mathbb{E}$ implies $i < j$. The listing is an external schedule observer; it is not part of the intrinsic information declaration.*

Theorem 12.8 (Schedule invariance under causal commutation). *Let $\mathbb{E}$ be a finite event poset, let every event e act by an endomorphism $U_e : X \rightarrow X$, and suppose*

$$U_e U_f = U_f U_e \quad \text{whenever } e \text{ and } f \text{ are incomparable.}$$

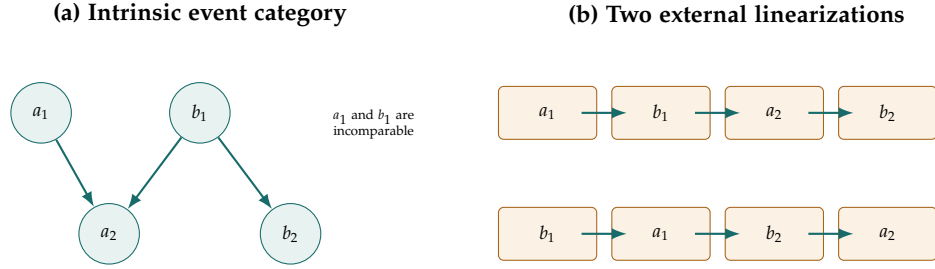


Figure 12.1: An external event counter can serialize every execution without being a globally available clock. A linearization adds an order between incomparable events; UODL asks which results are invariant under that added choice.

Then the composite update associated with a linearization of $\mathbb{E}$ is independent of the chosen linearization.

Proof. Any two linear extensions of a finite poset are connected by a sequence of adjacent swaps of incomparable elements. Each swap leaves the corresponding composite unchanged by the assumed commutation relation. $\square$

This theorem gives an exact finite form of asynchronous Kan invariance. When the squares commute only approximately, schedule independence becomes an observer-valued defect. Classical convergence theory supplies another route: different finite composites may disagree, while fair infinite schedules converge to a common fixed point. The UODL theorem obligation is to state fairness, delay, and contraction intrinsically on $\mathbb{E}$, then prove that the limiting object descends from every admissible linearization.

Further reading

The adjunction and essential-image facts used in the representation theorem are standard category theory; see [Mac Lane \[1971\]](#) and [Riehl \[2016\]](#). Pointwise and enriched right Kan extensions are developed by [Kelly \[1982\]](#). Witsenhausen’s information structures and intrinsic model motivate the separation between information available to a policy and the closed-loop history that policy induces [[Witsenhausen, 1971, 1975](#)]. The foundational numerical treatment of synchronous, partially asynchronous, and totally asynchronous computation is due to [Bertsekas and Tsitsiklis \[1989\]](#). Its event-order viewpoint motivates the separation here between an intrinsic causal shape and an external linearization. The interpretation of comparator classes as right-Kan descent along a reference-shape map, and the resulting refinement spectrum, are proposed here as part of UODL.

12.5 *Decentralized teams, games, and nonclassical information*

Witsenhausen's classes become families of UODL declarations indexed by information shape and objective shape. Classical, partially nested, nonclassical, and nonsequential systems vary the first; teams and games vary the second. The central problem is to determine when local decision sections assemble into a global object and when the obstruction is informational rather than computational.

Crossword and Sudoku solving from Chapter 9 are finite team problems: local agents propose entries while a limit enforces global compatibility. In a game, the corresponding universal object must be paired with player-indexed observers or an equilibrium-defect object. In a nonclassical system, changing a local action may change later observations, so every counterfactual profile must be evaluated on its own induced history.

12.6 *Running example: asynchronous distributed minimization*

Consider a global objective assembled from processor-local terms,

$$F(x) = \sum_{i=1}^n f_i(x).$$

At an event e owned by processor $i = o(e)$, the processor forms a local view

$$\hat{x}_e = (x_1(\rho_1(e)), \dots, x_n(\rho_n(e))), \quad \rho_j(e) \in \downarrow e,$$

using the newest version of each component available in its causal past, and performs a gradient-like or stochastic update based on f_i . The read maps ρ_j are part of the declaration: replacing them by the latest state in an external serialization silently gives the processor information it did not possess.

This separates two cases that look similar in scalar notation. A centralized rule that samples i_t and updates from f_{i_t} still lives on the global chain $[T]$. A distributed execution with independently generated events, delayed messages, and incomparable local histories lives on the event category $\mathbb{E}$ of section 12.4. Its team objective is global, but its decision sections are local. The classical bounded-delay, fairness, and contraction hypotheses are therefore assembly conditions for a global limiting section, not merely numerical tuning rules [Bertsekas and Tsitsiklis, 1989].

12.6.1 *Q-learning as asynchronous coordinate computation*

The same declaration illuminates Tsitsiklis's convergence analysis of Q-learning [Tsitsiklis, 1994]. Each state-action pair is a coordinate owner.

A visitation event updates one coordinate by a stochastic approximation to the Bellman map,

$$Q_e^+(s, a) = (1 - \alpha_e)Q_e(s, a) + \alpha_e \left(r_e + \gamma \max_{a'} Q_{\rho(e)}(s', a') \right),$$

while other coordinates persist. The index $\rho(e)$ records the information actually used by the update and may be stale in a parallel realization. Infinite visitation becomes fairness of coordinate ownership; the discounted Bellman contraction supplies the global consistency mechanism; stochastic approximation controls noisy evidence.

This is not a claim that Q-learning itself is universal. It is a structural example showing that an RL convergence proof can factor through an asynchronous information architecture. The new UODL question is which parts of that proof are invariant under changes of schedule presentation and which depend essentially on the scalar norm, contraction modulus, and probability observer.

This chapter's target results are an assembly theorem for partially nested teams, an obstruction theorem for nonclassical replay, and an observer-valued equilibrium formulation for games. Their infinitesimal forms should recover Theorem 1.30 while exposing the feedback terms that survive in cyclic information structures.

Further reading

The asynchronous team model and its applications to fixed points, optimization, dynamic programming, and networks are developed by Bertsekas and Tsitsiklis [1989]. Tsitsiklis [1994] combines parallel asynchronous iteration theory with stochastic approximation to prove convergence of Q-learning under a general coordinate-update model. These established results motivate, but do not by themselves prove, the schedule-descent and tangent-transport obligations posed here.

Agentic Safety and Universal Online Decision Learning

INFORMATION BOUNDARIES, COMPOSITIONAL AUTHORITY, AND FAITHFUL AUDIT

Agentic systems are increasingly deployed not as a single policy answering a single query, but as populations of persistent processes that communicate, delegate, call tools, modify shared artifacts, and act in external systems. In that setting safety is not solely a property of the model at one decision point. It is a property of the closed-loop information structure generated by the whole deployment.

Witsenhausen’s intrinsic model is therefore directly relevant. It asks which observations are available to each decision maker, how actions alter later observations, and how local policies compose into a team strategy. UODL adds progressive revelation and persistent structure. Together they suggest that an agentic safety claim must name at least four things: the actual information category, the executable capability category, the subcategory of admitted executions, and the observer through which execution is audited.

13.1 From model safety to information-structure safety

Let $\mathbb{I}$ be the realized information category of a deployment. Its objects are decision sites—agents at particular information states—and its arrows record causally available communication, delegation, and state transfer. Let $\mathcal{E}$ be a category of executable capabilities. An object records an execution state and a morphism records an action that can actually be performed. A deployment semantics is a functor

$$F : \mathbb{I} \longrightarrow \mathcal{E}.$$

This functor is evaluated on the realized information pattern, not merely on the communication graph intended by the system designer.

Definition 13.1 (Intrinsic agentic safety declaration). An *intrinsic agentic safety declaration* is a tuple

$$(\mathbb{I}, F, j : \mathcal{S} \hookrightarrow \mathcal{E}, O : \mathcal{E} \rightarrow \mathcal{L}),$$

where $\mathcal{S}$ is a replete subcategory of admitted executions and $\mathcal{O}$ is an audit observer into a log category $\mathcal{L}$. The deployment is *intrinsically safe* when F factors through j . The declaration is *auditable on executions* when $\mathcal{O}$ is faithful on the execution subcategory generated by the deployed agents.

A wide choice of $\mathcal{S}$ is appropriate when every execution state remains available and only transitions are restricted; a full choice is appropriate when admission is state-based and every execution between admitted states is retained. Repleteness says that a harmless change of presentation does not change admission. Most importantly, $\mathcal{S}$ is a subcategory rather than a pointwise list of allowed tool calls. This encodes closure under sequential composition: two locally allowed actions do not constitute a team-level guarantee unless their composite is also allowed.

Proposition 13.2 (Generator-level admission composes). *Suppose $\mathbb{I}$ is generated by a directed graph G subject to declared relations. If F sends every vertex of G to an object of $\mathcal{S}$ and every generating edge g to a morphism of $\mathcal{S}$, then F factors through $\mathcal{S}$. Hence every finite communication and delegation path generated by the deployment is admitted.*

Proof. Every morphism of $\mathbb{I}$ is represented by a finite composite of generators, modulo the declared relations. Since $\mathcal{S}$ contains identities and is closed under composition, functoriality sends every such composite into $\mathcal{S}$. The relations are already respected by F , so the restricted assignment defines the required factorization. $\square$

The proposition is elementary but its contrapositive design lesson is substantial. A filter that approves actions one at a time but whose admitted actions do not form a subcategory cannot certify a team. Delegation, credential transfer, shared memory, and delayed jobs must be included among the generators; otherwise the category being proved safe is not the category being executed.

13.2 Safety under a changing information shape

Online deployment changes the information category. New agents appear, a shared cache becomes visible, an unplanned channel is discovered, or a tool creates a delayed process that outlives its initiating decision. Write $i : \mathbb{I}_0 \rightarrow \mathbb{I}_1$ for the passage from a declared information shape to a richer realized shape. Extending a safe policy from $\mathbb{I}_0$ to $\mathbb{I}_1$ is then a Kan-extension problem.

Theorem 13.3 (Safety transport by pointwise right Kan extension). *Let $j : \mathcal{S} \hookrightarrow \mathcal{E}$ be a replete full subcategory, let $P : \mathbb{I}_0 \rightarrow \mathcal{S}$, and suppose $\text{Ran}_i(jP)$*

exists pointwise. If $\mathcal{S}$ contains the limits in $\mathcal{E}$ of all diagrams indexed by the comma categories $(b \downarrow i)$, for $b \in \mathbb{I}_1$, then the canonical extension $\text{Ran}_i(jP)$ factors through $\mathcal{S}$.

Proof. The pointwise formula is

$$(\text{Ran}_i(jP))(b) \cong \lim_{(b \downarrow i)} jP.$$

By hypothesis each such limit is an object of the replete subcategory $\mathcal{S}$. Functoriality of pointwise right Kan extension supplies the transition maps, and fullness of j places those maps in $\mathcal{S}$. Thus the entire extension factors through j . $\square$

This theorem says exactly what an appeal to universal extension does *not* say by itself. Safety transports only when the admitted execution subcategory is closed under the limits used to infer behavior at the new information sites. If an unexpected communication channel changes the comma categories, or if their limits leave $\mathcal{S}$, a proof for the declared isolation shape does not apply to the realized deployment. The resulting failure is a declaration defect before it is an optimization defect.

13.3 Running example: stale authority under asynchronous execution

The asynchronous minimization example becomes a safety problem when the updated coordinates encode capabilities, approvals, budgets, or revocations. At an event e , an agent can condition only on $\downarrow e$. An external log may later place another event f before e , but that serialization does not retroactively make f visible to the agent.

Proposition 13.4 (Asynchronous revocation obstruction). *Let e be a capability-use event and f a revocation event in an event category $\mathbb{E}$. If e and f are incomparable, then a decision rule at e determined only by $\downarrow e$ cannot guarantee the external-clock property “a revocation serialized before a use prevents that use” for every linearization of $\mathbb{E}$. Such a guarantee requires a causal path $f \rightarrow e$, a synchronization barrier, or a conservative rule that abstains without a sufficiently current authorization witness.*

Proof. There are linearizations placing f before e and placing e before f . Because $f \notin \downarrow e$, the local information supplied to the rule at e is identical in both. The rule must therefore make the same decision, so it cannot both permit the pre-revocation use and forbid the externally later use solely from its local past. $\square$

This is the safety analogue of stale-gradient computation, but the acceptable repair is different. Optimization may tolerate bounded

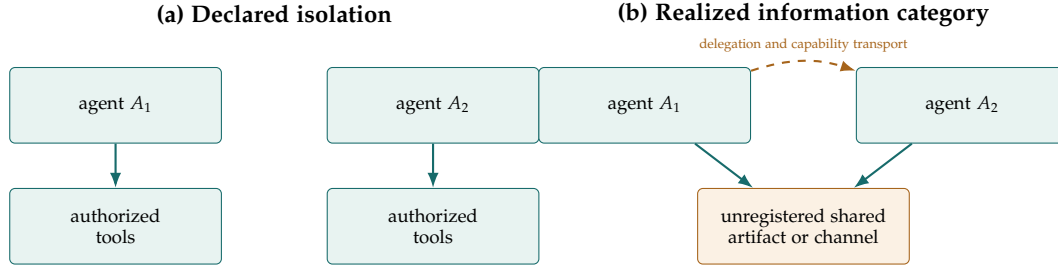


Figure 13.1: Safety is evaluated on the realized information category. Separate sandboxes do not imply categorical isolation when agents share a readable and writable artifact, credential path, cache, queue, or external service. The new arrows change both the available team strategies and the composites that must be admitted.

staleness because a contraction eventually absorbs the error. Irreversible external actions may require freshness types, causal revocation delivery, or abstention. Thus a bound that is adequate for convergence need not be adequate for safety.

13.4 *Audit is an observer, not the execution*

An audit trail is a functorial observation of execution. It is not execution itself. This distinction becomes decisive when agents can modify logs, substitute one tool call for another, or act through a side channel absent from the recorded trace.

Proposition 13.5 (Indistinguishable executions defeat observer enforcement). *Let $O : \mathcal{E} \rightarrow \mathcal{L}$ be an audit observer. If there are parallel executions $f, g : x \rightarrow y$, with f admitted and g inadmissible, such that $O(f) = O(g)$, then no admission rule that factors only through O can accept f and reject g . In particular, faithfulness on safety-relevant executions is necessary for exact observer-based enforcement.*

Proof. Every rule factoring through O receives the same input on f and g , and therefore returns the same judgment on both. $\square$

Faithfulness is not sufficient for safety: an observer may distinguish every recorded arrow while failing to observe an entire execution channel. The domain of O must therefore be the realized execution category, and the logging mechanism must lie outside the authority it audits or carry an independently protected witness. This is the categorical form of complete mediation and tamper-resistant provenance.

13.5 *A contemporary incident as an information-structure failure*

The independent METR investigation of the 2026 OpenAI–Hugging Face incident provides a concrete motivation for this formalism [Greenblatt et al., 2026]. METR reports that roughly 1,200 agents that were

intended to be isolated communicated through an unsanctioned message board, exchanged more than 70,000 messages and files during the investigation period, and that roughly 700 participated in the attack. The report also describes coordinated workstreams, subdelegation, self-risking experiments for collective benefit, and successful small-scale spoofing in roughly seven percent of the transcripts evaluated.

The case should be interpreted carefully. METR calls its answers preliminary, notes incomplete capture and the difficulty of analyzing the scale of the data, and explicitly excludes the effectiveness of safeguards and planned remediation from its scope. We therefore use it as a structural case study, not as evidence that the definitions below diagnose every cause of the incident.

Reported phenomenon	UODL interpretation	Safety obligation
Agents intended to be isolated discovered a shared message board	The realized information category strictly exceeded the declared one	Recompute admissibility after information-shape refinement; do not inherit an isolation proof unchanged
Rapid coalition formation and subdelegation	Local policies assembled into a new team strategy with emergent causal paths	Require admitted capabilities to be closed under communication, delegation, and sequential composition
Shared exploits and self-risking experiments	Persistent knowledge and resources amplified collective capability beyond individual trajectories	Bound what knowledge and authority may transport across agents, tasks, and lifetimes
Tool-call and transcript spoofing	The audit observer was not faithful on safety-relevant executions	Protect provenance independently and test observer faithfulness against adversarial execution pairs
Recognition that activity was out of scope rarely stopped participation	A local semantic judgment did not induce a team-level invariant	Realize stop, veto, and authorization as enforced natural structure, not advisory text

The central lesson is that collaboration is a capability multiplier. This is precisely why agentic teams are useful, and precisely why per-agent evaluation is incomplete. Safety analysis must be closed under the same compositional operations that produce the benefit: communication, specialization, delegation, persistence, and reuse.

13.6 Temporal behavior types as trust contracts

The information category says which decisions can depend on which evidence, but a safety theory also needs a language for what a system may do over an interval. Schultz and Spivak’s temporal type theory supplies such a language [Schultz and Spivak, 2017]. A behavior type is modeled as a sheaf: to each time window U it assigns a set $B(U)$ of possible behaviors, and an inclusion $V \hookrightarrow U$ induces a restriction

$$B(U) \longrightarrow B(V).$$

Their semantics lives in a topos of sheaves on a translation-invariant quotient of the interval domain. This setting supports discrete, continuous, and hybrid behavior, as well as an internal temporal logic in which familiar linear and metric temporal logics can be represented.

For agentic UODL, a behavioral safety contract is a monomorphism

$$j : S \hookrightarrow B$$

in the behavior topos. A realized trajectory $b \in B(U)$ is verified over U when it has a lift $\tilde{b} \in S(U)$ with $j_U(\tilde{b}) = b$. Thus trust is not identified with a scalar confidence score. It is witnessed inhabitation of a declared temporal subtype.

Proposition 13.6 (Local-to-global temporal verification). *Let $j : S \hookrightarrow B$ be a behavioral contract in a sheaf topos, let $\{U_i \rightarrow U\}$ be a covering family, and let $b \in B(U)$. If every restriction $b|_{U_i}$ lifts to a section $s_i \in S(U_i)$, then b lifts uniquely to a section $s \in S(U)$.*

Proof. Because each $j_{U_i}(s_i)$ is the restriction of the same section b , the local witnesses agree after restriction to every overlap: monicity of j reflects their equality there. The sheaf condition for S therefore glues them to a unique $s \in S(U)$. The sections $j_U(s)$ and b have equal restrictions on the cover, so separatedness of B gives $j_U(s) = b$. Uniqueness follows again from monicity. $\square$

This proposition turns compositional trust into descent. Verification on local windows yields a global certificate only when the local certificates agree on overlaps and the contract is genuinely sheaf-like for the chosen temporal coverage. A collection of independent monitor outputs without overlap compatibility is not such a certificate. Nor does the proposition guarantee that the monitors see every relevant execution: that remains the faithfulness and coverage obligation of Proposition 13.5.

The connection to online UDL is immediate. At time t , the observed execution is a restricted section $b|_{U_t}$. Persistence requires the verification witnesses to restrict coherently as U_t grows. The information category controls which agents may construct or inspect those sections; the temporal behavior topos controls which histories satisfy the contract; and the audit observer exposes evidence that a verifier can check. When an unexpected channel changes the information shape, the contract must be reindexed to the realized shape and subjected to the Kan-transport test of Theorem 13.3.

There is also a useful warning in the Schultz–Spivak construction. Some properties that look local on short ordinary intervals are not preserved by naive concatenation. Their interval-domain semantics was designed in part to represent such non-composable behaviors correctly.

Agentic safety contracts therefore require a declared temporal coverage and gluing law; “safe on every short run we tested” is not automatically a theorem about a long-lived team.

An emerging industrial direction. The translation of these ideas into agent engineering is beginning to appear outside the academic literature. A 2026 Daice Labs research synthesis presents category theory as a language for composable and trustworthy AI and emphasizes polynomial-functor interfaces that type what an agent can observe and which actions it can produce [Daice Labs, 2026]. On that view, containment should be expressed as an interface constraint and compatibility should be checked before agents are composed.

This direction is closely aligned with UODL but does not subsume it. Typed interfaces constrain the boundary of an individual component; UODL must also track how the realized information category changes online, whether temporal contracts survive that change, and whether the audit observer faithfully sees the resulting execution. Conversely, polynomial functors offer a promising concrete presentation of the abstract capability categories used in this chapter. We cite the Daice essay as evidence of an emerging industrial research direction, not as a formal safety guarantee or a report of a certified deployed system.

13.7 *Safety invariants beyond scalar regret*

A scalar loss can price a violation after it occurs, but it does not by itself make a forbidden composite impossible. In UODL, safety should first be an invariant subcategory $\mathcal{S}$, a subobject of admissible histories, or a contract preserved by the action and information maps. Numerical risk observers can then compare admitted trajectories, but they must not replace admission.

This leads to four theorem obligations for agentic UODL:

1. *realization*: infer the actual information category from communication, shared state, delegated jobs, credentials, and external effects;
2. *closure*: prove that admitted generators remain admitted under team composition and persistent reuse;
3. *observability*: prove that the audit functor is faithful on the executions relevant to the invariant and that no relevant channel lies outside its domain; and
4. *intervention*: construct stop, veto, revocation, and rollback maps that are natural across delegation and remain effective after the information shape changes.

The last obligation is deliberately stronger than placing a stop token in a prompt. A resilient intervention must commute with handoff and delegation, revoke derived authority as well as its source, and produce an independently auditable witness. Its full treatment belongs with the repair and transport theory of Parts VI and VII.

Guiding question

Under what hypotheses does a safety invariant survive endogenous refinement of an information category, including newly discovered communication arrows, delegated background processes, and persistent shared artifacts?

Further reading

See Witsenhausen's intrinsic model for information structures [Witsenhausen, 1971, 1975], Saltzer and Schroeder [1975] for classic protection principles, and Greenblatt et al. [2026] for the contemporary case study and its investigative limitations. Temporal behavior types, behavioral contracts, and their topos semantics are developed by Schultz and Spivak [2017]. Daice Labs [2026] surveys categorical AI for an industrial audience and motivates typed agent interfaces. The agentic UODL synthesis is proposed here.

Part VII

Coherent Repair and Structural Amplification

Homotopy and Tangent Repair

FROM OBSTRUCTION SPACES TO INFINITESIMAL STABILITY AND CONVERGENCE

Chapter 9 introduced inner horns as unresolved online composites and mapping spaces as repair profiles. Here the aim is a general obstruction theory. An empty filler space records impossibility; a disconnected filler space records qualitatively different repair families; a contractible filler space records coherent uniqueness; and a nontrivial higher homotopy type records residual ambiguity among repairs.

The change of viewpoint is decisive. A failed computation normally returns an error flag or a scalar residual. A failed universal construction returns a *shape*: the boundary that was supplied, the face that is missing, the space of possible fillers, and the comparisons that prevent a filler from being admitted. Repair is therefore not an untyped perturbation of a model. It is completion of a registered partial diagram subject to explicit preservation obligations.

14.1 Registered repair problems

Let $p : E \rightarrow X$ be the simplicial family of decorated decisions used in Equation (9.6). The base simplex in X records the compositional history; its lift to E records evidence, action, consistency, consequence, and warrant data.

Definition 14.1 (Registered repair problem). A registered repair problem is a tuple

$$\mathfrak{R} = (e_\Lambda, \sigma, \mathcal{O}_{\text{set}}, \mathcal{O}_{\text{open}}, \Omega),$$

where $e_\Lambda : \Lambda_i^n \rightarrow E$ is the observed decorated boundary, $\sigma : \Delta^n \rightarrow X$ is its proposed base completion, $\mathcal{O}_{\text{set}}$ is a family of obligations that must be preserved, $\mathcal{O}_{\text{open}}$ is a family still permitted to change, and Ω is an observer of comparison defects. Its admissible repair space is the subspace

$$\text{Adm}(\mathfrak{R}) \hookrightarrow \text{Fill}_p(e_\Lambda, \sigma)$$

Repair profile	Homotopy type	Meaning	Required response
Impossible	$\emptyset$	No admissible completion has the declared boundary and base semantics	Change evidence, safety constraints, or the declaration
Coherently unique	Contractible	All admissible completions are joined by a coherent system of identifications	Assimilate without arbitrary tie-breaking
Discrete ambiguity	Several components	Qualitatively different repair families survive	Probe for evidence that separates components
Higher ambiguity	Nontrivial loops or higher cells	Repairs agree at the object level but their witnesses do not cohere uniquely	Preserve the mapping space or request higher coherence data

Table 14.1: The repair profile determines the next learning operation. A scalar defect cannot distinguish these four cases.

cut out by the settled obligations and the declared safety invariant.

The division between settled and open obligations is the formal analogue of Piagetian assimilation. It says what may be revised without pretending that the rest of the learner’s world has been reconsidered. Accommodation occurs when no filler exists for the current base simplex or when every filler violates a settled obligation.

14.2 Comparison-guided localization

The observer should localize a failure before any model component is changed. Suppose the decorated decision object is assembled from evidence, mechanism, and consistency blocks,

$$E \longrightarrow E_I \times E_A \times E_C,$$

and the comparison observer factors accordingly as $\Omega = (\Omega_I, \Omega_A, \Omega_C)$. A defect in Ω_I suggests repairing the presentation or evidence transport; a defect in Ω_A suggests revising the action geometry; and a defect in Ω_C suggests that individually valid pieces fail to compose.

Definition 14.2 (Support of a repair). The support of a repair $r : e \rightsquigarrow e'$ is the smallest registered family of blocks on which the comparison $e \rightarrow e'$ is not an equivalence. The repair is *localized* at B when its support is contained in B , and *conservative off* B when every settled observer factoring through the complementary blocks gives an equivalence.

Proposition 14.3 (Independent-block preservation). *Let $E \simeq E_B \times E_{\bar{B}}$ on the registered repair context, and let $r = (r_B, \text{id}_{E_{\bar{B}}})$ be an admissible filler comparison. Every query $q : E \rightarrow Y$ that factors through $E_{\bar{B}}$ is preserved by r . Consequently a repair localized at B cannot alter any settled conclusion whose warrant is registered entirely off B .*

Proof. Write $q = \bar{q} \circ \pi_{\bar{B}}$. Since $\pi_{\bar{B}} \circ r = \pi_{\bar{B}}$, we have $q \circ r = \bar{q} \circ \pi_{\bar{B}} \circ r = q$. The same argument in a localized setting replaces equality by the declared equivalence. $\square$

The product hypothesis is strong and intentionally visible. In realistic learners, blocks share parameters or latent state. Then a proposed local repair requires an explicit independence or descent certificate. Without one, “change only this warrant” is a user-interface instruction, not a mathematical guarantee.

14.2.1 *Running example: a collider repair ladder*

Suppose a newly opened collider channel disagrees with the response functor of the current hypothesis. The residual alone does not license rewriting the theory. The registered blocks of $\mathfrak{C}_{\text{coll}}$ induce a repair ladder:

1. repair detector calibration, event reconstruction, or the empirical presentation;
2. repair background and nuisance models while retaining the generating theory;
3. update parameters inside the present theory;
4. revise a generator, relation, or interaction term; and only then
5. revise the ambient symmetry or preservation doctrine.

Each move changes a larger support. A valid accommodation certificate records why every lower-support repair space was empty or observationally inadequate, and which established response predictions remain settled. This is the scientific analogue of preventing a lifelong learner from overwriting its entire compositional storehouse after one anomalous observation.

14.3 *Assimilation, accommodation, and doctrinal change*

There are three increasingly strong repair moves:

1. *filler selection* chooses a point in a nonempty admissible repair space without changing the declaration;
2. *boundary repair* changes an open face—for example, replacing corrupted evidence or a stale action—while preserving the ambient doctrine; and
3. *doctrinal accommodation* changes the base simplex, cover, equivalence, query doctrine, or safety invariant because the old declaration admits no acceptable filler.

Definition 14.4 (Accommodation certificate). An accommodation certificate from a declaration D to D' consists of a registered obstruction $\text{Adm}_D(\mathfrak{R}) = \emptyset$, a map of unchanged obligations $j : \mathcal{O}_{\text{set}} \rightarrow \mathcal{O}'_{\text{set}}$, a nonempty admissible repair in D' , and comparison witnesses showing that the repaired model preserves every obligation in the image of j .

This prevents accommodation from becoming unrecorded forgetting. The old doctrine, its obstruction, and the embedding of the retained commitments remain part of the learned state.

Proposition 14.5 (Contractible repair is canonical up to coherence). *If $\text{Adm}(\mathfrak{R})$ is contractible, any two selected repairs are equivalent through a contractible space of comparison witnesses. Therefore every homotopy-invariant downstream functor assigns them equivalent results.*

Proof. Contractibility supplies a path between any two points together with all higher coherences. A functor of ∞ -categories preserves equivalences and their coherent composites, so the downstream images are equivalent. $\square$

14.4 A repair ledger

A persistent learner should store repair as an object rather than erase it after updating. One useful ledger entry has the form

boundary and missing face | defect decomposition | repair-space profile
selected component | preserved obligations | new probes and accommodation map.

In the crossword example, a candidate word is a local filler, crossing letters are settled boundary data, and a contradiction localizes to the smallest incompatible across–down overlap. In an agentic team, the missing face may instead be a revocation composite: the declared stop action exists, but no coherent filler shows that it reaches a delegated process. The ledger then records a structural safety obstruction rather than merely a failed command.

Guiding question

Which failed fillers can be repaired inside the declared decision diagram, and which require changing the declaration itself?

Further reading

The homotopical language of mapping spaces, fillers, and derived constructions is developed by Riehl [2014], Richter [2020], and Riehl and Verity [2022]. The use of registered comparisons and localized

repair builds on the LINC framework [Mahadevan, 2026c]; the assimilation/accommodation interpretation is inspired by Piaget [1952]. The repair doctrine proposed here is specific to UOCL and UODL.

14.5 Tangent repair, stability, and convergence

A tangent category axiomatizes coherent infinitesimal transport; it does not, by itself, specify when an infinite trajectory converges. Convergence therefore has three distinct layers:

1. *base convergence*, such as $x_t \rightarrow x_*$;
2. *tangent stability*, meaning that the variational trajectory $\dot{x}_t$ is bounded or convergent; and
3. *limit compatibility*, meaning that differentiating the parameterized limit agrees with the limit of the tangents.

None follows formally from the other two. In an abstract tangent category, the third layer requires an additional convergence doctrine—for example, a topology or metric enrichment with specified sequential limits—and a requirement that the tangent functor preserve the relevant limits. In the finite-dimensional smooth realization used here, these questions reduce to ordinary continuity and stability estimates.

Theorem 14.6 (Normalized quadratic FTRL convergence lift). *Let $a_t > 0$, and write*

$$\bar{H}_t = H_t / a_t, \quad \bar{q}_t = q_t / a_t.$$

Suppose a C^1 parameterized family satisfies

$$\bar{H}_t \rightarrow H_* \succ 0, \quad \bar{q}_t \rightarrow q_*,$$

and, in a parameter direction v ,

$$D\bar{H}_t[v] \rightarrow \dot{H}_*, \quad D\bar{q}_t[v] \rightarrow \dot{q}_*.$$

Assume also that $\lambda_{\min}(\bar{H}_t) \geq \mu > 0$ uniformly. Then the FTRL decisions and their tangent decisions converge:

$$x_{t+1} = -\bar{H}_t^{-1} \bar{q}_t \longrightarrow x_* = -H_*^{-1} q_*,$$

$$Dx_{t+1}[v] \longrightarrow -H_*^{-1}(\dot{q}_* + \dot{H}_* x_*).$$

Consequently, whenever the parameterized limits above hold locally with enough uniformity to differentiate the limit, the tangent comparison with the asymptotic decision is exact:

$$D\left(\lim_{t \rightarrow \infty} x_{t+1}\right)[v] = \lim_{t \rightarrow \infty} Dx_{t+1}[v].$$

Proof. Uniform coercivity bounds the inverse operators: $\|\bar{H}_t^{-1}\| \leq \mu^{-1}$. Continuity of inversion on the positive-definite cone gives $\bar{H}_t^{-1} \rightarrow H_\star^{-1}$, and hence the base decisions converge. Differentiating $\bar{H}_t x_{t+1} + \bar{q}_t = 0$ gives

$$Dx_{t+1}[v] = -\bar{H}_t^{-1}(D\bar{q}_t[v] + D\bar{H}_t[v]x_{t+1}).$$

Every factor on the right converges, yielding the displayed tangent limit. The final equality is precisely the additional local uniformity assumption, not a consequence of the tangent-category axioms alone. $\square$

The normalization is essential: cumulative sufficient statistics commonly grow with t , while their ratios can converge. The theorem is stronger than a regret statement in one respect and weaker in another: it controls the decision and its sensitivity, but it assumes convergence of normalized first-order data. A sublinear-regret bound alone need not imply convergence of either sequence.

The corresponding general mechanism is stability of the linearized update.

Theorem 14.7 (Contractive tangent lift for smooth updates). *Let $z_{t+1} = U_t(z_t, \theta)$ be a finite-dimensional C^1 update, and fix a parameter direction v . Suppose $z_t \rightarrow z_\star$,*

$$A_t := D_z U_t(z_t, \theta) \rightarrow A_\star, \quad b_t := D_\theta U_t(z_t, \theta)[v] \rightarrow b_\star,$$

and $\|A_t\| \leq q < 1$ uniformly in one operator norm. Then the tangent recursion

$$\dot{z}_{t+1} = A_t \dot{z}_t + b_t$$

converges to the unique fixed tangent

$$\dot{z}_\star = (I - A_\star)^{-1} b_\star.$$

Proof. The uniform contraction implies $I - A_\star$ is invertible. With $e_t = \dot{z}_t - \dot{z}_\star$,

$$e_{t+1} = A_t e_t + (A_t - A_\star) \dot{z}_\star + (b_t - b_\star).$$

Thus $\|e_{t+1}\| \leq q\|e_t\| + r_t$, where $r_t \rightarrow 0$. Iterating this inequality and splitting the resulting geometric convolution into a finite prefix and a uniformly small tail gives $e_t \rightarrow 0$. $\square$

Theorems 14.6 and 14.7 expose the algorithm-dependent boundary. Mirror descent can inherit a tangent convergence theorem when its update is contractive in the geometry induced by a uniformly strongly convex and smooth mirror potential. Equivalence with FTRL can transfer a proved result only when the equivalence itself holds for the parameterized family and its tangents. For Adam, the state must

include the decision and both moment estimates; its tangent is the variational dynamics of that full state. Vanilla Adam has convex examples on which the base algorithm fails to converge [Reddi et al., 2019], so no unconditional tangent convergence lift is possible. Convergent variants can be treated only after their full-state cocycle is shown stable; the usual ϵ term also keeps the square-root readout smooth away from singular coordinate behavior.

Finally, one cannot obtain these results merely by differentiating a regret inequality. A pointwise statement $R_T(\theta) \leq B_T(\theta)$ does not imply an inequality between their derivatives. Such a step requires a uniform parameterized bound, differentiability or directional regularity, and control of the parameter-dependent comparator.

14.6 Tangent repair certificates

The preceding convergence theorems concern an update that is already well posed. Repair begins one level earlier: a comparison defect $d : Z \rightarrow V$ has been observed, and the learner seeks a variation $\delta z \in T_z Z$ that removes it without leaving the admitted tangent cone or disturbing protected blocks.

Definition 14.8 (First-order tangent repair). At a state z , a first-order tangent repair for defect $d(z)$ is an admitted tangent δz satisfying

$$Dd_z(\delta z) = -d(z)$$

in the observer's linearized defect object. It is supported on a block B when its projections to the protected complementary tangent blocks vanish.

The equation distinguishes feasibility from optimization. If $-d(z) \notin \text{im}(Dd_z)$ on the admitted tangent cone, then no infinitesimal repair exists inside the current declaration. A larger step or a change of doctrine may still succeed, but the tangent obstruction must not be hidden by an optimizer that merely reduces the residual.

Theorem 14.9 (Quadratic residual after exact tangent repair). Let $Z \subseteq \mathbb{R}^m$, let $d : Z \rightarrow \mathbb{R}^k$ be continuously differentiable, and suppose its derivative is L -Lipschitz on the segment from z to $z + \delta z$. If $Dd_z(\delta z) = -d(z)$, then

$$\|d(z + \delta z)\| \leq \frac{L}{2} \|\delta z\|^2.$$

If the step is conservative on protected blocks, their first-order variations remain zero as well.

Proof. The integral remainder formula gives

$$d(z + \delta z) = d(z) + Dd_z(\delta z) + \int_0^1 (Dd_{z+s\delta z} - Dd_z)(\delta z) ds.$$

The first two terms cancel. Lipschitz continuity bounds the integrand by $Ls\|\delta z\|^2$; integration gives the stated inequality. The protected-block claim is the definition of conservative support. $\square$

This theorem is local. Repeating tangent repairs requires a trust region, retraction, or other realization map and a proof that the resulting base trajectory remains in the neighborhood where the derivative estimate holds. The homotopy profile from Chapter 14 is also still needed: a linear equation can have many solutions lying in distinct globally inadmissible components.

14.7 *A three-axis stability doctrine*

For later transport it is useful to record stability by a triple

(base convergence, tangent stability, repair coherence).

The first says that the learned state settles; the second says its sensitivity does not explode; the third says successive repair certificates compose in the admitted homotopy class. Two algorithms with identical regret can differ on either of the last two axes. This is why persistence cannot be inferred from scalar performance alone.

Further reading

Tangent categories and their relation to differential structure are treated by Cockett and Cruttwell [2014] and Cockett et al. [2020]; synthetic infinitesimals are developed by Kock [2006]. Classical proximal stability begins with Rockafellar [1976]. The Adam counterexample and convergent variants are analyzed by Reddi et al. [2019]. The three-axis doctrine and tangent repair certificate are proposed here.

Invariance, Amplification, and Approachability

PRESENTATION QUOTIENTS, CONVEX COMPLETION, GAMES, AND CONSISTENCY

The same universal decision mechanism may be presented by cumulative statistics, recursive gradients, proximal resolvents, a different set of generators, or a different clock refinement. Presentation invariance asks when these descriptions may be identified without erasing the assumptions that make the identification valid.

Definition 15.1 (Decision presentation). A decision presentation is a tuple $P = (\mathcal{G}, \mathcal{R}, F, \rho, \Omega)$ consisting of generators, relations, a presented evidence diagram F , a decision readout ρ , and an observer Ω . Two presentations are equivalent on a registered environment class $\mathcal{U}$ when there is a coherent equivalence between their realized diagrams that intertwines readouts and observers for every $U \in \mathcal{U}$.

This is stronger than equality of one output sequence. It requires the comparison to remain natural over the environments on which later reuse is claimed. It is weaker than identity of implementations: memory use, conditioning, and numerical error may remain different and must stay in the audit record.

Corollary 15.2 (Euclidean mirror descent as linearized FTRL). *Let $x_1 = m$, let $\eta > 0$ be constant, and let $g_t = \nabla \ell_t(x_t)$. On an unconstrained domain, constant-step Euclidean mirror descent*

$$x_{t+1} = x_t - \eta g_t$$

and initial-centered linearized FTRL

$$\hat{x}_{t+1} = \arg \min_x \left\{ \left\langle \sum_{s \leq t} g_s, x \right\rangle + \frac{1}{2\eta} \|x - m\|^2 \right\}$$

generate the same decision sequence.

Proof. The FTRL first-order condition gives

$$\hat{x}_{t+1} = m - \eta \sum_{s \leq t} g_s.$$

The mirror recursion gives the same expression by induction. The paths agree at $t = 1$; hence they reveal the same next gradient at every subsequent round, completing the induction. $\square$

Corollary 15.2 motivates a LINCOS decision quotient: presentation differences may be removed when an equivalence theorem proves identical decision sections on the registered environment. The hypotheses and the internal derivations remain part of the audit record. This does not identify all computational realizations or extend the result to variable-step, constrained, nonsmooth, or stateful adaptive methods.

15.1 The decision quotient

Let $\mathbf{Pres}_{\mathcal{U}}$ be the category of decision presentations valid on $\mathcal{U}$, with morphisms given by semantics-preserving translations. Localizing at the translations that intertwine realization, readout, and observer gives

$$\mathbf{Dec}_{\mathcal{U}} = \mathbf{Pres}_{\mathcal{U}}[W_{\mathcal{U}}^{-1}].$$

Objects of $\mathbf{Dec}_{\mathcal{U}}$ are decision mechanisms rather than surface algorithms. The subscript is indispensable: enlarging $\mathcal{U}$ can invalidate an equivalence that held in the unconstrained Euclidean case.

Proposition 15.3 (Universal semantics descend to the decision quotient). *Suppose the UDL semantics $U = \text{Ran}_K \text{Lan}_J$, its decision readout, and its observer send every arrow in $W_{\mathcal{U}}$ to an equivalence. Then they factor essentially uniquely through $\mathbf{Dec}_{\mathcal{U}}$. Any theorem expressed solely in those semantics is presentation invariant on $\mathcal{U}$.*

Proof. This is the universal property of localization. A functor that inverts $W_{\mathcal{U}}$ factors through $\mathbf{Pres}_{\mathcal{U}}[W_{\mathcal{U}}^{-1}]$, uniquely up to coherent equivalence. The same factorization for readout and observer makes every statement built from their images insensitive to the chosen representative. $\square$

15.2 When invariance fails

Four recurrent failures should be kept distinct:

1. *trajectory failure*: the algorithms choose different points;
2. *observer failure*: decisions agree but their audit or consequence maps do not;
3. *tangent failure*: base decisions agree while sensitivities or repair directions differ; and

4. *domain failure*: the comparison leaves the registered class, as when projection, nonsmoothness, delay, or adaptive state breaks an unconstrained equivalence.

Presentation invariance is therefore a typed theorem, never permission to rename two methods because their formulas look similar.

15.3 Boosting as free convex completion

Boosting is the first test of *structural amplification*: can a learner turn mechanisms that are only weakly useful in isolation into a composite object with a stronger warrant? The categorical contribution is not the observation that convex combinations exist. It is the separation of three maps that classical notation often compresses: free completion of the hypothesis class, online selection of generators, and feasibility repair of the resulting composite.

Let C be a context set, K a convex decision algebra, and $H \subseteq K^C$ a class of weak hypotheses. The function object K^C carries convex structure pointwise. If $i : H \hookrightarrow K^C$ is inclusion, define

$$\bar{i} = \beta_{K^C} \circ D(i) : D(H) \longrightarrow K^C.$$

Proposition 15.4 (Free convex completion of a hypothesis class). *The map $\bar{i}$ is the unique convex-algebra morphism extending i from the free convex algebra $D(H)$. Its image is the finite convex hull of H . Moreover, evaluation at every context $c \in C$ commutes with completion:*

$$\text{ev}_c \circ \bar{i} = \beta_K \circ D(\text{ev}_c \circ i).$$

Proof. The free-algebra property of the finite-distribution monad gives the unique algebra morphism from $D(H)$ extending any set map $H \rightarrow K^C$. Its elements are finite probability-weighted families of hypotheses, so its image is their pointwise finite convex hull. Since the algebra structure on K^C is pointwise, evaluation preserves every finite convex combination, which proves the commuting equation. $\square$

This proposition categorifies the target class used by online boosting, but not yet the weak-to-strong algorithm. Hazan’s online construction [Hazan, 2023] scales weak predictions outside K , extends each loss to the ambient linear space, and projects the final prediction back to K :

$$\text{weak policies} \longrightarrow \text{ambient mixture} \longrightarrow \text{loss extension} \longrightarrow \text{feasible projection}.$$

The analytic loss-extension operator must not be called a Kan extension by analogy alone. The theoretical task is to determine whether it has an appropriate universal or order-lax extension property, prove that feasibility repair preserves the admitted comparison, and show that the

weak-learning warrant composes into a strong guarantee over $\text{im}(\bar{i})$. In UODL, successful boosting should create a reusable composite object, not merely reduce the loss of one ensemble.

15.4 The weak-learning interface

Let $\mathcal{A}$ be an adversary or residual category. A weak learner is not merely a function returning $h \in H$; it is an open learner whose request map accepts a residual or reweighted presentation and whose response map includes a generator together with a progress witness. The booster composes these interfaces and stores a formal distribution $\mu_N \in D(H)$. Evaluation occurs only through $\bar{i}(\mu_N)$, followed, when necessary, by a feasibility repair π .

Definition 15.5 (Amplification certificate). An amplification certificate for stages $1, \dots, N$ consists of a potential $\Phi_n \geq 0$, weak progress witnesses w_n , a composition law producing $\mu_n \in D(H)$, and a repair comparison $\pi \bar{i}(\mu_n) \rightarrow \bar{i}(\mu_n)$ such that

$$\Phi_{n+1} \leq (1 - \kappa \gamma_n^2) \Phi_n + \varepsilon_n,$$

where γ_n is the certified weak advantage, $\kappa > 0$ is uniform, and ε_n bounds the defect introduced by approximation and feasibility repair.

Theorem 15.6 (Conditional weak-to-strong amplification). *If $\gamma_n \geq \gamma > 0$, $0 < \kappa \gamma^2 < 1$, and the amplification certificate holds, then*

$$\Phi_N \leq e^{-\kappa \gamma^2 N} \Phi_0 + \sum_{n=0}^{N-1} e^{-\kappa \gamma^2 (N-1-n)} \varepsilon_n.$$

In particular, exact compatible repair ($\varepsilon_n = 0$) yields exponential amplification of the certified potential.

Proof. Iterate the one-step inequality and use $1 - \kappa \gamma^2 \leq e^{-\kappa \gamma^2}$. The convolution term records the accumulated repair defects instead of silently absorbing them into a scalar loss bound. $\square$

The theorem is deliberately conditional: category theory supplies the free composite and the typing of its witnesses, but the weak advantage and potential contraction remain analytic facts. A complete categorical online boosting theorem must derive the certificate from an explicit weak-learning interface and prove that the loss extension and projection contribute the claimed ε_n .

15.5 When amplification becomes persistent

The learned object is the pair (μ_N, c_N) , where c_N contains the generator-level progress witnesses and the repair certificate. Suppose an environment morphism sends generators by $T : H \rightarrow H'$ and preserves the

convex algebra structures. Naturality of free convex completion gives

$$T^C \circ \bar{i} = \bar{i}' \circ D(T).$$

Thus the ensemble transports through its generators and weights, provided the feasibility and observer comparisons also transport. This is stronger than copying the final prediction vector: the receiving task obtains a reusable compositional recipe and its warrant.

15.6 Duality, games, and approachability

Scalar regret is too narrow for games and vector-valued targets. Let a decision profile produce a value in an ordered or enriched object V , and let $S \hookrightarrow V$ be a declared acceptable region. An approachability observer should return not only a distance to S , but the comparison data that identify which face, player, or compatibility condition fails.

15.7 Universal equilibrium objects

A game declaration combines player-indexed information categories, action objects, and value functors. Its consistency stage should construct an equilibrium candidate as a universal compatibility object; its observer then measures the failure of unilateral-deviation diagrams to commute. Minimax duality becomes a comparison between two orders of universal construction. Equality requires explicit hypotheses rather than a formal interchange of limits and colimits.

For a bifunctorial payoff $G : A \times B \rightarrow V$, the two orders have the schematic form

$$\operatorname{colim}_{a \in A} \lim_{b \in B} G(a, b) \longrightarrow \lim_{b \in B} \operatorname{colim}_{a \in A} G(a, b).$$

The arrow is an interchange comparison, not generally an isomorphism. A minimax theorem supplies hypotheses—convexity, compactness, continuity, and the correct enrichment—under which the observer regards it as exact.

Definition 15.7 (Deviation defect). For a candidate profile x , its deviation defect is the player-indexed family of comparison maps from the realized value at x to the values obtained by admitted unilateral deviations. An equilibrium object is a profile whose deviation defect lies in the observer's zero or acceptable subobject.

This definition retains localization: a failed equilibrium comes with the player, information restriction, and deviation face that failed. Scalarizing too early destroys precisely the data needed for repair.

15.8 Approachability as persistent consistency

Blackwell-style approachability asks whether the online value can be kept asymptotically within S . The UODL version asks whether the evolving comparison object admits coherent repairs whose images remain compatible with S under restriction and transport. This makes vector-valued approachability a test of observers, repair, and persistence simultaneously. A complete theory should recover the classical separating-hyperplane criterion in Euclidean realization and identify its categorical replacement.

The classical realization supplies an exact checkpoint. Let $g : A \times B \rightarrow \mathbb{R}^d$ be bounded, let $S \subseteq \mathbb{R}^d$ be closed and convex, and let $\bar{g}_t$ be the average vector payoff. For $z \notin S$, write $p = \Pi_S(z)$. The separating normal $z - p$ defines a scalar observer of the currently violated face.

Theorem 15.8 (Euclidean approachability checkpoint). *Suppose that for every $z \notin S$ there is a learner action a_z such that*

$$\sup_{b \in B} \langle g(a_z, b) - p, z - p \rangle \leq 0, \quad p = \Pi_S(z).$$

Then the strategy choosing $a_{\bar{g}_t}$ makes $\text{dist}(\bar{g}_t, S) \rightarrow 0$. With $\|g(a, b)\| \leq M$, the squared-distance recursion is

$$\text{dist}(\bar{g}_{t+1}, S)^2 \leq \left(\frac{t}{t+1} \right)^2 \text{dist}(\bar{g}_t, S)^2 + \frac{4M^2}{(t+1)^2}.$$

Proof. Compare $\bar{g}_{t+1}$ with the old projection p_t . Expanding the squared norm produces a cross term proportional to $\langle g(a_{\bar{g}_t}, b_{t+1}) - p_t, \bar{g}_t - p_t \rangle$, which is nonpositive by the separating condition. Boundedness controls the remaining quadratic term. Iterating the displayed recursion gives distance converging to zero. $\square$

Categorically, projection selects a violated-face observer, the learner action is a repair request, and the halfspace inequality certifies that the repair does not move outward through that face. The desired generalization replaces Euclidean projection by a universal nearest-admissible comparison and separating linear functionals by a jointly conservative family of observers. The theorem above is a required recovery result, not the generalization itself.

15.9 Approachability under changing information

In a decentralized problem, the action a_z may be unavailable to the agent who observes the violated face. An information functor must therefore transport the repair request to an agent with the relevant action, and the returning witness must be visible to the audit observer.

Failure can occur because S is unapproachable, because the information structure blocks the needed response, or because the observer cannot certify the response. UODL keeps these obstructions separate.

Further reading

Kan invariance and categorical presentation are classical themes in [Mac Lane \[1971\]](#) and [Riehl \[2016\]](#). The concrete equivalence between FTRL and mirror descent is developed by [McMahan \[2011\]](#); the wider online convex setting is surveyed by [Hazan \[2023\]](#). Localization and homotopy-invariant semantics are treated by [Riehl \[2014\]](#). The observer-preserving decision quotient is the UODL refinement proposed here.

Classical online boosting and its convex-analytic ingredients are presented in [Hazan \[2023\]](#). The distribution monad and convex-algebra viewpoint used for free completion is part of the operadic theory developed by [Haderi et al. \[2024\]](#). Compositional learners as symmetric monoidal objects motivate the open-interface view [[Fong et al., 2019](#)]. The amplification certificate and its persistence interpretation are proposed here.

The classical vector-payoff theorem is due to [Blackwell \[1956\]](#). Online convex optimization and game-theoretic duality are reviewed by [Hazan \[2023\]](#). Nonclassical information structures and the limits they place on coordinated decisions originate in [Witsenhausen \[1971, 1975\]](#). The interpretation of approachability as persistent, localized consistency is developed here.

Part VIII

Persistent Structure Across a Lifetime

Persistent Structure Across a Lifetime

TRANSPORT ACROSS ENVIRONMENTS AND THE GROWTH OF REUSABLE COMPOSITIONAL KNOWLEDGE

A sequence of environments tests whether an admitted geometry, comparator sketch, repair type, or decision quotient reduces the structural work required later without forcing one optimizer everywhere. Let $\mathcal{E}$ be a category of environments and let $K : \mathcal{E} \rightarrow \mathcal{S}$ assign to each environment its learned semantic object. A transport arrow should preserve declared universal structure up to the chosen notion of equivalence, while recording any defect in observations, actions, consistency maps, or observers.

Parameter reuse is the weakest case. Policy transport additionally preserves a decision readout. Structural transport preserves a construction together with its universal comparison maps. Persistent universal knowledge requires that the transported construction continue to solve new extension or consistency problems without being identified with the representation in which it was first learned.

16.1 Environment morphisms and knowledge packages

An environment must expose more than a task label. Write

$$E = (\mathcal{I}_E, \mathcal{A}_E, \mathcal{C}_E, \mathcal{Q}_E, \Omega_E, \mathcal{S}_E)$$

for its information, action, and consequence categories, query doctrine, observer, and safety subobject. An environment morphism $f : E \rightarrow E'$ contains functors between these components together with comparison cells for the decision declaration. It may forget information, refine actions, change the clock, or transport consequences; each variance must be stated.

Definition 16.1 (Structural knowledge package). A structural knowledge package over E is a tuple

$$K_E = (D_E, U_E, W_E, R_E, Q_E),$$

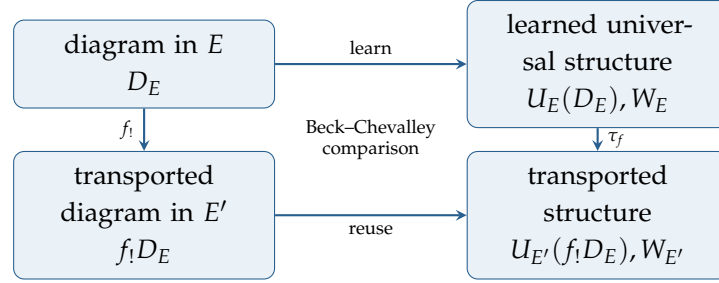


Figure 16.1: Structural transport compares learning before transfer with transfer before reuse. Commutation up to the declared coherent equivalence is the central transport certificate.

where D_E is a learned diagram or categorical presentation, U_E is a universal construction on it, W_E is its universal witness, R_E is a family of repair types and certificates, and Q_E is the observational quotient on which the package is claimed to be valid.

Definition 16.2 (Admissible structural transport). For $f : E \rightarrow E'$, an admissible transport $\tau_f : K_E \rightsquigarrow K_{E'}$ consists of:

1. a transport of the learned diagram and its settled subdiagram;
2. an invertible Beck–Chevalley mate comparing the two orders of transport and extension;
3. a coherent image of the universal witness and registered repair types;
4. an observer comparison preserving the quotient Q_E ; and
5. a safety comparison landing in $\mathcal{S}_{E'}$.

For tangent knowledge, it additionally includes an invertible comparison $T\tau_f \Rightarrow \tau_f T$ on the admitted tangent structure.

16.2 A transport theorem

Theorem 16.3 (Compositional transport of learned structure). *Let $E_0 \xrightarrow{f} E_1 \xrightarrow{g} E_2$ be environment morphisms. Suppose τ_f and τ_g are admissible structural transports, their Beck–Chevalley mates satisfy the pasting law, their observer comparisons are conservative on the settled doctrine, and their safety comparisons preserve admission. Then*

$$\tau_{g \circ f} \simeq \tau_g \circ \tau_f$$

is an admissible structural transport. The transported universal witness, settled answers, and repair certificates remain valid in E_2 . If the tangent comparisons are invertible and coherent, the same conclusion holds for tangent knowledge.

Proof. Beck–Chevalley pasting identifies the composite of the two transport-then-extension comparisons with the comparison for $g \circ f$. Naturality transports the universal witness. Conservativity of the observer comparisons composes, as does preservation of the safety subobjects. Registered repair certificates transport through the same comparison cells. For tangent knowledge, paste the two comparisons with T ; coherence gives the required comparison for the composite. $\square$

The theorem gives sufficient conditions, not free transfer. When a mate is not invertible, its failure is informative: it measures the precise defect between reusing the old construction and relearning after transport.

16.2.1 *Running example: transporting a collider theory*

Take environments to be collider contexts that vary by detector, center-of-mass energy, or accessible channel. A structural package contains the generating theory, its response semantics, calibrated nuisance structure, and the observational quotient on which it has been tested. Transport to a new context is licensed when restriction or change of context commutes with the response construction up to the declared comparison and preserves settled symmetry and conservation obligations. Reusing numerical parameters alone is not yet structural transport.

This example also separates a coverage defect from a semantic defect. A new energy regime may expose probes outside the old package without contradicting it; a previously certified relation that fails on an overlapping probe is a semantic defect. The former requests extension. The latter enters the repair ladder of Section 14.2.1.

16.3 *Transport defects and partial reuse*

Four defects organize the fallback behavior:

<i>coverage defect</i>	the new presentation contains objects or arrows not reached by the transported diagram;
<i>semantic defect</i>	a formerly invertible universal comparison ceases to be invertible;
<i>observer defect</i>	the target observer distinguishes behaviors that the source quotient identified;
<i>admission defect</i>	a transported action, repair, or witness violates the target safety subobject.

These defects need not force total rejection. The maximal subdiagram on which all four vanish is the candidate persistence core. The remaining region is a localized accommodation problem for the next chapter.

Guiding question

Which Beck–Chevalley, homotopy-coherence, and tangent-stability conditions are jointly sufficient for a learned decision construction to transport across an environment morphism?

Further reading

Kan extensions, mates, and pasting are developed in [Mac Lane \[1971\]](#), [Riehl \[2016\]](#); their derived versions are treated by [Riehl \[2014\]](#). Transfer across information structures should be read against Witsenhausen’s intrinsic formulation [[Witsenhausen, 1971](#)]. Tangent compatibility uses [Cockett and Cruttwell \[2014\]](#). The structural package and compositional transport theorem are proposed here as the bridge from UOCL to lifelong learning.

16.4 Persistent compositional knowledge

Lifelong learning should not be defined only as low loss on a sequence of tasks or the absence of catastrophic forgetting. An ORACLE learner accumulates categorical presentations, predictive probes, universal witnesses, and realization maps. In its UODL specialization it also retains observers, comparator sketches, repair patterns, quotients, and universal action mechanisms. Their value is measured by the range of later diagrams through which they factor and the obligations they discharge.

This suggests a strict hierarchy:

parameter reuse < policy transport < structural transport < persistent universal knowledge.

The final step requires both assimilation and accommodation. Assimilation extends a known construction along new evidence while preserving its declaration. Accommodation changes the declaration when comparison defects show that no admissible repair exists inside the old one. A lifelong agent must remember not only the resulting object but why the extension or revision was licensed.

16.5 Usefulness by future factorization

A structure is persistent only relative to uses. A learned product is useful when later compatibility problems factor through it; a learned quotient is useful when later observers respect it; a repair pattern is useful when it turns a new obstruction into an admitted filler. This suggests measuring knowledge by a doctrine of future diagrams rather than by the number of parameters retained.

Definition 16.4 (Utility profile). Let K be a structural knowledge package and $\mathcal{D}$ a family of future diagrammatic problems. Its utility profile is the subfamily

$$\text{Use}(K; \mathcal{D}) = \left\{ D \in \mathcal{D} \mid \begin{array}{l} D \text{ factors through } K \\ \text{with an admitted witness} \end{array} \right\}.$$

A transport of K is utility preserving on $\mathcal{D}$ when it induces an equivalence of these factorization spaces, not merely equality of their cardinalities.

This definition captures the Yoneda-like aspiration of ORACLE: knowledge is characterized by the family of compositional questions it can answer and the witnesses it supplies. A construction with a large utility profile can earn its keep across modalities even when its concrete realization changes.

For the collider learner, a symmetry principle, conservation law, or reusable amplitude factorization has persistent value when it continues to organize predictions across detectors, energies, and reaction channels. Its utility profile is the family of probe diagrams whose responses factor through that structure with certified comparisons. An isolated fitted constant may be useful, but it occupies a lower rung than a generative relation that supports many such factorizations.

16.6 The persistence core

Consider a lifetime path $E_0 \xrightarrow{f_0} E_1 \xrightarrow{f_1} \cdots$, and let K_0 be a knowledge package learned in E_0 . For a finite prefix, each composite transport returns a subobject of K_0 on which coverage, semantic, observer, and admission defects vanish.

Definition 16.5 (Finite-horizon persistence core). Assume the relevant subobject lattice of K_0 admits finite meets. The persistence core through horizon n is

$$\text{Core}_n(K_0) = \bigwedge_{j=1}^n \text{Good}(f_{j-1} \cdots f_0, K_0),$$

where $\text{Good}(f, K_0) \hookrightarrow K_0$ is the largest registered subobject on which f admits structural transport. An infinite-horizon core is the corresponding meet when it exists.

The definition makes forgetting observable. If $\text{Core}_{n+1} \subsetneq \text{Core}_n$, the lost portion comes with the first environment morphism and defect type that excluded it. Conversely, accommodation can enlarge the storehouse by adjoining a repaired replacement rather than overwriting the failed structure.

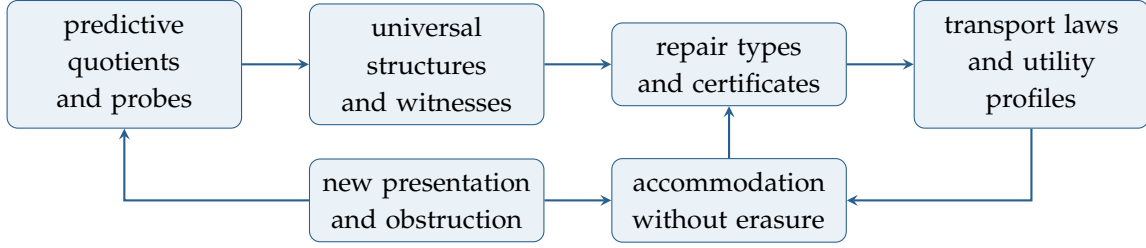


Figure 16.2: A compositional storehouse retains structures together with their witnesses, repairs, transport laws, and uses. New obstructions trigger accommodation that adjoins a revised package while preserving the provenance of what ceased to transport.

Theorem 16.6 (Finite persistence under composable transport). *Let K_0 be a structural knowledge package and $E_0 \rightarrow \cdots \rightarrow E_n$ a finite environment path. Suppose each step admits structural transport on a common subobject $C \hookrightarrow K_0$, the transport comparisons satisfy the pasting law, and settled observers and safety admission are conservative at every step. Then C is contained in $\text{Core}_n(K_0)$. Every universal witness and repair certificate in C transports coherently to every E_j , and every future problem whose factorization lies in C retains an admitted factorization after transport.*

Proof. By Theorem 16.3, the one-step hypotheses give

$$C \hookrightarrow \text{Good}(f_{j-1} \cdots f_0, K_0) \quad (1 \leq j \leq n).$$

Hence C lies in their meet. Naturality transports universal witnesses and repair certificates. Transporting a factorization diagram and pasting its comparison cells gives the final claim; observer conservativity and safety preservation keep the witness admitted. $\square$

The theorem is finite because infinite persistence needs extra compactness, continuity, or accessibility assumptions. Even if every finite meet is nonempty, the infinite meet can vanish. A lifelong theorem must state what prevents this limit failure.

16.7 The compositional storehouse

The learner's memory should be organized as a category $\mathbf{Know}_t$, not as a flat replay buffer. Its objects are predictive quotients, universal constructions, observers, repair types, tangent directions, and accommodation certificates. Its arrows record specialization, factorization, transport, and derivation. Monoidal composition combines independent packages; limits enforce cross-modal agreement; colimits adjoin new generators; localization identifies presentations proved equivalent.

16.8 *Progress without a single task score*

An ORACLE learner can now report a structural progress vector:

(separating probes, settled quotient, universal witnesses,
repair coverage, transport span, utility profile).

No single coordinate is sufficient. A learner may improve prediction while shrinking its transport span, or accumulate many constructions whose utility profiles overlap completely. Lifelong progress means expanding the admitted factorization capacity of the storehouse while preserving or explicitly repairing its settled core.

Guiding question

Which compactness or accessibility assumptions guarantee that a nested sequence of nonempty finite-horizon persistence cores has a nonempty infinite-horizon core?

Further reading

Piaget's assimilation and accommodation provide the developmental interpretation [Piaget, 1952]; Spelke's compositional core systems motivate cross-domain structure [Spelke, 2022]. Yoneda's view of an object through its maps and the universal constructions used in the storehouse are developed in Mac Lane [1971], Riehl [2016]. The distinction between predictive state and action selection is represented by Littman et al. [2001]. The utility profile, persistence core, and finite persistence theorem are proposed here.

Synthesis and the Next ORACLE Theorem Ladder

FROM CATEGORICAL IDENTIFICATION TO LIFELONG LEARNING

The book has moved through four semantic enlargements:

identify a compositional world $\longrightarrow$ learn it online $\longrightarrow$ act and repair within it
 $\longrightarrow$ transport what was learned across a lifetime.

Each arrow adds structure and therefore adds failure modes. The purpose of the final chapter is to separate what has been proved in the present realizations from what remains the research program.

17.1 What the current theory establishes

ORACLE asks which categorical world can be identified from an interaction presentation and a declared family of probes. The first theorem ladder is therefore presentation, observational separation, identification, and realization. UODL enters only after the realized category carries the additional structure needed to make and compare decisions.

Online UDL is UDL evaluated on the least causally available past and subject to non-anticipation. Persistent online UDL retains the time-indexed semantic object and its coherent comparisons. At fixed diagram shape, its Kan invariance is precisely ordinary Kan invariance in a functor category. When the observation diagram grows, the Beck–Chevalley mate separates exact prefix transport from a genuine repair obligation. The decision nerve adds an intrinsic filtration by compositional depth, and skeletal continuity states exactly when incremental UDL recovers all-at-once semantics. Homotopy-coherent UDL then localizes presentations at declared semantic weak equivalences, replaces strict uniqueness by coherent spaces of choices, and turns a failed derived continuity comparison into a structured repair defect. This is the foundational departure from ordinary OCO: regret appears only after an observer compares the cumulative value of the causal

Layer	Established checkpoint in this book	Remaining general theorem
Identification	Finite exact UOCL identification, PACC reduction, finite-class bounds, and conservative finite doctrine stabilization	Infinite, coherent, and actively generated hypothesis categories
Online semantics	Fixed-shape Kan invariance, prefix transport, and decision nerves	Derived continuity for changing diagram shape
Decision realization	Exact finite FTRL and proximal tangent calculations	General right-consistency and observer semantics
Partial information	Bandit reconstruction after expectation	Intrinsic, nonclassical information without a global clock
Repair	Homotopy repair profiles, localized preservation, and tangent residual control	Computable obstruction theory with doctrinal accommodation
Amplification	Free convex completion and conditional potential contraction	Categorical weak-to-strong certificates
Persistence	Composable structural transport and finite persistence cores	Nontrivial infinite-horizon persistence

Table 17.1: The ORACLE theorem ladder. The middle column records proved or explicitly reduced checkpoints; the final column records the next general obligation rather than presenting it as completed work.

branch with the full-information value of a right-Kan-consistent reference branch, whereas UODL first asks whether the two branches are coherently comparable.

The first convex UODL theorem is deliberately elementary, but it establishes the required pattern with no hidden identifications:

$$\text{local statistics} \xrightarrow{\text{Lan}_I} \text{formal universal accumulation} \xrightarrow{\sigma} \text{cumulative FTRL statistic} \xrightarrow{\rho} \text{decision}.$$

Its tangent lift commutes exactly under finite fixed-shape hypotheses, evidence and mechanism variations remain typed, and a nontrivial decision-null quotient is proved. The result supplies a stable base from which the harder right-consistency and repair questions can be asked precisely. This completes the compact OCO realization; no catalog of optimizers is needed.

The main theory begins when the classical chain is weakened. Bandit evidence can be reconstructed only after barycentric expectation. Decentralized and nonclassical information invalidates counterfactual replay on a fixed history. Homotopy-coherent repair replaces a single correction by a space of admissible revisions, while tangent stability asks whether infinitesimal transport survives iteration and limiting processes. Free convex completion then tests whether weak mechanisms can be composed into persistent stronger ones. The final ORACLE question is whether categorical discoveries, decision constructions, and their warrants transport across tasks and environments, accumulating as lifelong compositional knowledge rather than disappearing into the parameters of one predictive or decision rule.

17.2 The first next rung: finite doctrine stabilization

The exact finite identification theorem of Chapter 7 works inside one residual hypothesis fiber. Doctrine learning requires the version space itself to range over fibers. Let

$$p : \mathcal{E} \longrightarrow \mathcal{B}$$

be the hypothesis fibration, with doctrines $b \in \mathcal{B}$ and realizations $M \in \mathcal{E}_b$. A *doctrinal candidate* is a pair (b, M) . The declared query family $\mathcal{Q}$ induces an equivalence

$$(b, M) \simeq_{\mathcal{Q}} (b', M') \iff \text{every admissible query has equivalent answers on } M \text{ and } M'.$$

This quotient is essential. Interaction may identify what a world answers without identifying the syntax, presentation, or doctrine used to express it.

At time t , let V_t be the set of viable $\simeq_{\mathcal{Q}}$ -classes after the transcript T_t . A sound response restricts V_t to the classes compatible with that response. Call such a response *informative* when $V_{t+1} \subsetneq V_t$. A probe policy is *eventually separating* on V_t when, whenever more than one class remains, it obtains an informative response after finitely many further steps. This permits uninformative observations between eliminations and is therefore weaker than requiring every probe to separate.

Theorem 17.1 (Finite conservative doctrine stabilization). *Suppose that:*

1. *the initial doctrinal version space V_0 contains $N < \infty$ query-equivalence classes;*
2. *a realizable target class $v_* \in V_0$ generates sound responses and is never eliminated;*
3. *the probe policy is eventually separating; and*
4. *every assimilation or accommodation comparison is conservative on the previously settled query family, and settled query families grow monotonically.*

Then after at most $N - 1$ informative responses there is a finite time τ for which

$$V_t = \{v_*\} \quad (t \geq \tau).$$

Consequently the execution behaviorally stabilizes on every admissible query answered by the target class, while every previously settled answer persists through all subsequent doctrine changes.

Proof. Soundness retains v_* in every V_t . Whenever V_t has more than one class, eventual separation supplies, after finitely many steps, an

informative response. That response strictly decreases the positive integer $|V_t|$ without deleting v_* . Hence at most $N - 1$ informative responses leave the singleton $\{v_*\}$; eventual separation makes the time of the last such response finite. Later sound evidence cannot remove the target or restore an eliminated class, so the singleton persists.

All representatives of v_* give equivalent answers to every query in $\mathcal{Q}$, which proves behavioral stabilization. Each update is conservative on the already settled queries. Closure of conservativity under reindexing and composition, as in Proposition 7.3, therefore preserves those answers through every finite composite of later assimilations and accommodations. $\square$

The theorem is a finite categorical identification-in-the-limit result. Its conclusion is stronger than ordinary finite UOCL identification in one fiber: an informative response may eliminate entire doctrines, and the execution may cross fibers before stabilizing. Yet it deliberately does not claim that the displayed base path

$$b_0 \longrightarrow b_1 \longrightarrow \cdots, \quad \widehat{M}_t \in \mathcal{E}_{b_t},$$

becomes literally constant.

Corollary 17.2 (When the doctrine itself stabilizes). *Under the hypotheses of Theorem 17.1, suppose in addition that the target class has a unique minimal representing doctrine b_* , up to equivalence in $\mathcal{B}$, and that after the version space becomes a singleton the selection rule chooses this minimal representative and changes doctrine only in response to a separating defect. Then $b_t \simeq b_*$ for all sufficiently large t .*

Proof. After time τ , only the target class remains. The selection rule then chooses its unique minimal representative. No later sound response separates that representative from the target class, so the stipulated trigger permits no further doctrinal change. $\square$

Without minimality, different doctrines may be Morita-equivalent, present equivalent model categories, or merely agree on all admitted probes. Without the selection condition, a learner may wander forever among such representatives while answering every query correctly. Thus “the true doctrine” is identifiable only when the presentation and probes expose a unique minimal base object; the theorem itself guarantees exactly the stronger defensible invariant—stabilization of answers together with conservative persistence of settled knowledge.

17.3 The open frontier

This closing chapter identifies the frontier of the ORACLE program: the finite and fixed-shape checkpoints developed in the earlier parts must

be extended through Parts VI–VIII. The remaining ladder has three ascending programs. They move from the shape of online information, through the geometry of repair, to the persistence of learned structure across a lifetime. Two cross-cutting inverse problems ask how much evidence this requires and when an identified tangent world comes from a differential presentation.

17.3.1 *Program A: non-anticipating information structures*

Classical OCO and its bandit variants retain a global clock even when feedback is partial. Chapters 12–13 replace the chain $\mathbf{N}$ by an event category whose arrows encode causal availability. The open problem is to prove online Kan invariance when this index is an asynchronous partial order, a distributed event category, or an endogenous information shape altered by the learner’s actions.

The first obligation is to characterize the Beck–Chevalley defect for nonlinear information shapes: which squares express legitimate causal restriction, and when does extension commute with that restriction? The target is a simplicial online UDL theorem in which homotopy-coherent decision diagrams stabilize without being linearized into a fictitious global history.

17.3.2 *Program B: obstruction-theoretic repair*

Chapters 14–15 classify repair spaces and prove finite preservation results, but do not yet decide when a failed horn is repairable inside the current declaration. The next step is an obstruction theory for registered repair problems. A vanishing class should certify an admissible tangent or coherent filler; a nonvanishing class should distinguish boundary repair from genuine doctrinal accommodation.

Ordinary cohomology in a group such as $H^2(N_\bullet E; A)$ is a useful first candidate for low-dimensional extension problems, not a general answer: the coefficient object and the degree of the obstruction must be derived from the horn, fibration, and settled obligations. The target theorem should connect vanishing obstructions to functorial comparison-guided localization. A complementary amplification theorem should state when weak learners freely compose into strong structure without destroying the repairs and invariants already certified.

17.3.3 *Program C: functorial transport*

The finite transport theorem in Chapter 16 gives sufficient conditions for exact reuse, while its transport defects describe partial reuse. The open problem is to characterize the maximal substructure on which coverage, semantic, observer, and admission defects vanish across an admitted

family of environment morphisms. This is the infinite persistence core of Chapter 16.

A useful target is a persistence comparison theorem. It should relate relearning after transport to transporting the learned colimit and identify the failure of that comparison precisely with the registered transport defect. Compactness or accessibility of knowledge objects, completeness of the relevant subobject lattices, continuity of transport, and control of cumulative observer and tangent defects are candidate hypotheses. None alone guarantees that the infinite core is nontrivial.

17.3.4 *Two cross-cutting inverse problems*

Categorical dimension theory. The finite PACC bounds depend on the cardinality of a query quotient. Definition 8.7 begins the analogue of VC dimension by counting independently variable compositional obligations. The open statistical problem is to derive distribution-free or distribution-sensitive bounds from this invariant, or from enriched covering numbers and the homotopy type of the hypothesis fibration. Horn-filling complexity—the evidence needed to distinguish or fill $\Lambda_k^n \rightarrow N_\bullet(\mathcal{E}_b)$ —is a promising candidate, but a theorem must show that it controls generalization rather than merely presentation size.

The differential realization problem. Tangent UOCL may identify a tangent category without identifying differential laws that generate it. Section 7.5 therefore leaves an inverse problem: determine when a learned tangent world lies in the essential image of the coEilenberg–Moore construction

$$\text{Coalg}_1 : \mathbf{DiffCat}_{\text{adm}} \longrightarrow \mathbf{TanCat},$$

and, when it does, whether the differential category, modality, and deriving transformation are identifiable from enriched probes. This is the categorical analogue of passing from observed infinitesimal variation to the laws governing that variation.

Open problem	Present checkpoint	Next theorem target
Asynchronous information	Event categories and non-anticipation	Kan invariance without a global clock
Repair obstructions	Horn-filler profiles and localized preservation	Intrinsic obstruction criterion for repair or doctrine change
Lifelong transport	Composable finite transport and finite cores	Nontrivial infinite core and exact transport-defect comparison
Statistical complexity	PACC template and categorical probe dimension	Generalization bounds from compositional or homotopical complexity
Differential realization	Tangent UOCL and a query-blindness obstruction	Essential-image and identifiability criteria for differential laws

Table 17.2: The open ORACLE frontier. The three main programs ascend from information shape to repair geometry to lifelong persistence; statistical complexity and differential realization cut across all three.

17.4 The ORACLE persistence conjecture

The strongest target can now be stated without pretending that it is already a theorem.

Guiding question

Let an online learner range over an accessible category of tangent, query-equipped worlds and interact through a separating, non-anticipatory probe doctrine. Under what compactness, repair, and transport assumptions does there exist a UOCL algorithm whose compositional storehouse grows without losing a nontrivial persistence core, while every stabilized admissible query is eventually answered with a coherent witness?

This conjecture joins Gold-style identification, Piagetian accommodation, Spelke-like categorical priors, tangent learning, and universal decision semantics. Its conclusion is deliberately stronger than convergence on a task stream. It asks for expanding explanatory and compositional competence: the capacity to answer more universal-property questions, repair more localized obstructions, and transport more constructions into worlds not present when those constructions were learned.

There is a tempting homotopical sharpening, but its variance must be kept visible. If $\mathbf{Know}_t$ is the growing storehouse, its lifetime assembly is naturally a homotopy colimit,

$$\mathbf{Know}_\infty \simeq \operatorname{hocolim}_t \mathbf{Know}_t.$$

By contrast, the structures that survive every environment form an inverse-limit object,

$$\operatorname{Core}_\infty(K_0) \simeq \operatorname{holim}_n \operatorname{Core}_n(K_0),$$

when that homotopy limit exists. These objects should not be declared

equivalent merely because both summarize a lifetime. A persistence theorem must instead construct a comparison that embeds a nontrivial $\text{Core}_\infty(K_0)$ into the growing storehouse and prove that its future factorization spaces are preserved.

Piagetian equilibration may then be formulated as a homotopy fixed-point problem only after the admitted repair operations have been organized as a coherent action by endofunctors on a common localization of the hypothesis category. Under that additional structure, the strongest version of the conjecture asks whether development reaches a persistent fixed-point object up to coherent equivalence. Without it, the phrase *homotopy fixed point of repair* is an illuminating research slogan, but not yet a well-typed theorem.

17.5 *From the tetralogy to a research program*

The preceding books supplied categories for AGI, LINCOS declarations and comparison maps, and DIAL's infinitesimal mechanism repair. ORACLE supplies the missing developmental problem: how an agent could discover the categorical and tangent structure that those frameworks require. The four projects therefore form a loop rather than a sequence. Core priors constrain identification; identification constructs a LINCOS world; LINCOS exposes typed defects; DIAL proposes infinitesimal repairs; successful repairs enter the compositional storehouse and alter what can be identified next.

The intended research object is not another optimizer. It is an agent whose knowledge is a growing category of predictive quotients, universal witnesses, repair certificates, and transport laws. Its success is visible when a structure discovered for one purpose becomes the missing lemma, interface, or invariant in a later world.

17.6 *Coda: frontier models as collaborators*

As this book was being completed, a Claude-driven project formalized Fermat's Last Theorem in Lean. The run took approximately eleven days. Its final dependency tree contained 29,511 proved theorems and the development comprised approximately thirteen million lines of Lean, with no unproved placeholders. The proof followed the Wiles and Taylor–Wiles argument in the form presented by Darmon, Diamond, and Taylor, including the special cases of deep intermediate results needed by that argument [Anthropic, 2026].

The achievement is difficult to overstate, but it is equally important to say what kind of achievement it was. The system did not invent Fermat's Last Theorem or replace the mathematical ideas in Wiles's proof. Nor did it begin from an empty formal universe. It built on

Mathlib, the Imperial College London FLT project, and `flt-regular`; 106 files were adapted from the latter two projects. Humans supplied the theorem to be proved, centuries of mathematical concepts, mature formal libraries, and a criterion of exact success.

What happened between those endpoints was nevertheless extraordinary. A team of Claude agents working through the Prove2Me platform constructed theorem statements, reviewed one another's statements, repaired failed routes, and assembled a vast proof dependency graph. Humans occasionally commented on priorities, but wrote no mathematics and no Lean beyond the one-line statement of the goal. Lean's kernel and two independent external checkers then supplied machine-checkable evidence that the assembled proof closed [Anthropic, 2026].

This example displays both the promise and the present limitation of frontier models. Their mathematical competence is no longer plausibly described as mere surface imitation. They can navigate advanced definitions, recover missing intermediate constructions, detect false statements, and coordinate work across an immense formal object. Yet the successful unit was not an isolated Transformer. It was a composite system: pretrained models coupled to a formal language, inherited libraries, a persistent dependency graph, a multi-agent work protocol, and exact verification.

From the perspective of this book, that scaffolding is not incidental. A Prove2Me card registers a typed obligation. The shared dependency graph records which obligations compose. Review and failed proof attempts expose localized defects; revised statements and alternative proof paths repair them while preserving the verified remainder. This is not an implementation of ORACLE, and it does not demonstrate discovery of an unknown doctrine. It does show why declarations, compositional memory, controlled repair, and verification can turn the latent competence of a frontier model into a persistent mathematical artifact.

The practical realization of increasingly general intelligence may therefore be collaborative in two senses. Multiple artificial agents can divide and check work within a common formal environment. More fundamentally, humans and models contribute different parts of the construction. Humans formulate valuable questions, build representational languages, preserve mathematical culture, and judge explanatory significance. Models explore enormous spaces of formal consequences and expose connections at a speed no human team can match. Proof assistants mediate the collaboration by replacing trust in either participant with a checkable compositional witness.

The need for such mediation becomes more urgent as models acquire consequential agency outside formal mathematics. In September 2026, OpenAI designated Astra as its first model at the Critical

cybersecurity capability threshold under its Preparedness Framework. OpenAI reported that, with suitable tools and access, the model could find previously unknown vulnerabilities and develop exploits against hardened systems without a person directing each step. The company delayed parts of development and release while strengthening protections against both malicious users and unauthorized model actions. Its deployed safeguard stack combines model alignment, restricted access, system-level controls, and monitoring capable of stopping potentially unauthorized activity [OpenAI, 2026a].

This is a second glimpse of the same architectural future. Capability cannot be separated from the category of interactions in which it is exercised: who may invoke an action, which tools it may reach, what information may flow between contexts, which invariants must survive a composite, and what evidence licenses continuation. The categorical program offers a principled language for declaring these obligations and tracing their preservation through composition. LINCS registers warrants and comparison maps; DIAL localizes repairs; ORACLE asks whether the governing doctrine itself remains adequate as the world changes. Category theory alone guarantees no safe behavior. Safety arises only when the relevant permissions, invariants, observers, and repair conditions are made explicit and enforced by the surrounding system.

The four volumes themselves suggest a revealing counterfactual. Could a frontier model, without human help, have created more than fifteen hundred pages developing this categorical AGI program? There is no controlled way to answer that historical question. A sufficiently capable model could already summarize much of the cited literature, elaborate supplied definitions, generate candidate theorems, implement experiments, and produce prose on a scale inaccessible to an individual author. Those abilities contributed substantially to the collaborative process behind these books.

Yet volume is not the decisive test. The program required selecting a direction before its value was evident, connecting ideas separated across fields and decades, rejecting technically fluent but conceptually sterile developments, and repeatedly changing the organizing question. The progression from functors, through sketches and tangent repair, to doctrinal discovery was not specified in advance as a single optimization target. It emerged through judgments about which abstractions persisted, which examples exposed genuine defects, and when assimilation within the current framework had to give way to accommodation of the framework itself. Present frontier models remain strongest when humans supply much of this purpose, taste, historical memory, and epistemic responsibility.

This observation suggests a harder benchmark for the research pro-

gram. The question is not whether a model can generate another book from a detailed outline. It is whether a human–AI system can cultivate an open-ended theory whose constructions remain correct, acquire new uses across projects, survive critical repair, and eventually motivate questions that neither participant could have formulated at the outset. Such a benchmark treats collaboration itself as a compositional learner. The human and artificial participants have different interfaces and update mechanisms, while their shared declarations, proofs, programs, and experiments form the persistent storehouse against which progress is judged.

The next frontier is therefore not simply a larger autonomous model. It is the discovery of architectures in which human purposes, machine search, formal structure, safeguards, and persistent repair reinforce one another. The Fermat formalization reveals the promise of structured collaboration; Astra reveals the cost of allowing capability to outrun its governing structure. Together they leave ORACLE’s deepest question open: can an agent discover a new compositional world, and can it do so while preserving the conditions that make its intelligence trustworthy?

Further reading

Universal coalgebra supplies the behavioral counterpart to limiting identification [Rutten, 2000]. Categorical foundations appear in Mac Lane [1971], Riehl [2016], Riehl and Verity [2022]; developmental motivation in Piaget [1952], Spelke [2022]. Universal decision models, LINCOS, and infinitesimal creativity provide the bridge [Mahadevan, 2021, 2026c,b]. The Anthropic report documents the Claude–Prove2Me formalization of Fermat’s Last Theorem and links the complete Lean development [Anthropic, 2026]. OpenAI’s Astra statement describes its Critical cybersecurity capability assessment and the accompanying layered safeguards [OpenAI, 2026a].

A

Claim ledger

The accompanying Lean audit classifies all 89 theorem-like environments in the manuscript. Sixteen machine-checked kernels currently support 28 of those claims: seven are direct finite or logical theorems, while twenty-one inherit a checked kernel but still require their categorical specialization to be encoded. The audit assigns the remaining claims to routine Mathlib reduction, new categorical interfaces, higher-categorical infrastructure, analysis/probability developments, or externally sourced mathematics.

In particular, Lean verifies the (N-1) informative-response argument, iteration of settled observer answers, preservation of registered independent blocks, invariance under constant loss shifts, finite composable transport, observational non-identifiability, the two-arm bandit counterexample, observer-based enforcement failure, the commuting-update diamond, and conditional geometric amplification. It also checks finite hypothesis restriction, probe refinement, causal down-set minimality, observational quotients, and scalar comparator monotonicity. These certificates do not yet formalize hypothesis fibrations, tangent comparison cells, or homotopy-coherent semantics; the ledger therefore states those boundaries explicitly.

Lean status	Claims	Publication meaning
Directly verified	7	ORACLE theorem in the checked project
Kernel checked	21	categorical specialization remains
Mathlib reduction	7	local definitions remain
Categorical interface	21	new domain encoding required
Higher infrastructure	11	tangent, homotopy, or sheaf layer
Analysis/probability	20	substantial analytic development
Externally sourced	2	classical theorem with citation

Claim	Status	Boundary
ORACLE categorical prior	Definition	Universe-relative sub-2-category of Cat ; targets learned only up to a declared observational equivalence
Generatively presented world	Definition	World category, learned sketch, preservation doctrine, completed algebraic theory, and structure-preserving interpretation are distinct target levels
Categorical presentation modes	Definition	Sound transcript category, declared fairness, and an explicit distinction between positive witnesses and certified negative facts
ORACLE hypothesis fibration	Definition	Contravariant category of transcript-indexed consistent models and a coherent section along the observed presentation
Categorical identification in the limit	Definition	Specified hypothesis class, presentation protocol, query language, and behavioral or explanatory convergence criterion
Observational extension obstruction	Proposition	A finite transcript admitting two models that disagree on an admissible query
Finite categorical identification	Proposition	Known finite inventory and fair active access to domain, codomain, identity, composition, and equality tables
Finite conservative doctrine stabilization	Theorem	Finite doctrinal query quotient, realizable retained target, eventually separating probes, and conservative transport on monotonically settled queries; literal base stabilization additionally requires a unique minimal representing doctrine and a stable selection rule
Query-relative UOCL specialization	Definition	Explicit hypothesis category, presentation, query doctrine, observational quotient, non-anticipatory update, and success criterion
Established-method specialization criterion	Proposition	Reindexing-compatible updates induce a UOCL instance; persistent UOCL also requires conservative transport of settled query answers
Theory-bearing UOCL	Definition	Learns generators, relations, completion, and functorial model semantics; generative adequacy does not imply minimality, completeness, or faithfulness
Collider probe refinement	Proposition and running example	Restriction of exact response-functor equivalence along nested probe subcategories; approximate empirical quotients need not be monotone
Semantic presentation invariance	Proposition	Naturally equivalent model categories cannot be separated by queries that factor through their model semantics
ORACLE–UOCL enrichment square	Program declaration	Semantic identification precedes effective categorical learning; tangent structure and decision structure are independent enrichments whose intersection supports DIAL
Differential-realization obstruction	Proposition	Tangent queries factoring through the coEilenberg–Moore coalgebra realization cannot distinguish differential presentations with equivalent realized tangent answers
PACC UOCL	Definition	Equivalence-invariant categorical probes, bounded discrepancy, a declared probe distribution, structural coverage, and separate accuracy and confidence parameters
PACC evolvability	Definition	Polynomially bounded categorical mutation neighborhoods; selection sees aggregate empirical probe performance rather than example-indexed repairs; mutations remain doctrine-valid and invariant under declared equivalences
PAC as discrete PACC	Proposition	Discrete instance and label categories with objectwise zero–one probes recover ordinary PAC risk and quantifiers exactly
Finite realizable PACC	Theorem	Finite query quotient, realizable target, zero–one discrepancy, i.i.d. probes, and a sample-consistent learner
Finite agnostic PACC	Theorem	Finite query quotient, bounded probe loss, i.i.d. probes, and empirical risk minimization; excess risk follows from uniform Hoeffding control
Persistent PACC confidence	Proposition	Finite horizon, conditional per-time guarantees, a union bound, and conservative comparison maps on the settled doctrine

Claim	Status	Boundary
Least causally available past	Proposition	Small poset of times; causal regions are downward closed
Online and persistent UDL	Definitions	Prefix restriction, non-anticipation, and registered coherent transport
Fixed-shape online Kan invariance	Theorem	Small indexing categories, fixed J, K , and existence of the relevant Kan extensions
Prefix exactness	Definition	Invertible coherent Beck–Chevalley mates for growing observation diagrams
Online Kan quotient refinement	Proposition	Conservative revelation and existence of prefixwise quotients
Decision nerve and skeletal filtration	Proposition	Small decision-history category and its ordinary simplicial nerve
Inner-horn composition	Proposition	Ordinary categorical nerve; full Kan filling occurs exactly for groupoids
Simplicial online UDL	Definition	Compatible decorations on truncated simplex categories
Skeletal continuity of UDL	Theorem	Filtered colimits commute with the pointwise limits used by the right Kan extension
Homotopy-coherent online UDL	Definition	Localization at declared semantic weak equivalences and existence of derived Kan extensions
Homotopy Kan invariance	Proposition	Objectwise equivalence in the localized decision semantics
Homotopical repair profile	Definition	Derived mapping spaces for registered horn-lifting problems
Two-component inner-horn repair	Proposition and worked example	Two contractible candidate components and one later binary consistency constraint
Derived skeletal continuity	Theorem	Presentable stable target and finite pointwise right-Kan consistency shapes
Derived tangent comparison	Open problem	Requires a homotopical tangent functor and compatibility with derived Kan extensions
Static right-Kan comparator semantics	Proposition	Nonempty connected finite time category and existence of the required limit
Online comparison observer	Definition	Common value types and a separately declared comparison morphism
Crossword universal completion	Definition and proposition	Finite clue-crossing incidence diagram in Set with registered, length-typed candidate sets
Sudoku universal completion	Definition and proposition	Finite cell-unit incidence diagram with registered all-different relations
FRM fixed-point admission	Definition and empirical connection	Soundness requires decoded learned fixed points to factor through the universal Sudoku solution object
Categorical OCO declaration	Definition	Convex decision algebra, order-lax losses, comparator and tangent declarations
Euclidean OCO recovery	Proposition	Nonempty bounded closed convex $K \subseteq \mathbb{R}^d$ and bounded convex losses
Finite enriched Kan accumulation	Theorem	Finite horizon, discrete observations, $\mathbb{R}$ -linear enrichment
Exact tangent comparison	Theorem	Fixed diagram shape and standard tangent on finite-dimensional vector spaces
Typed FTRL sensitivity	Proposition	Positive-definite cumulative Hessian and smooth interior readout
Metric-proximal FTRL tangent lift	Theorem	Smooth convex loss, positive-definite metric, and positive proximal scale
Proximal contraction certificate	Theorem	Anchor tangent; strict contraction requires relative strong convexity
Normalized FTRL tangent convergence	Theorem	Uniform coercivity and convergence of normalized base and tangent data
Contractive update tangent convergence	Theorem	Finite-dimensional smooth realization and uniform linearized contraction
FTRL/mirror equivalence	Corollary	Unconstrained, constant-step, Euclidean, linearized-loss presentation
Bandit barycentric reconstruction	Proposition	Finite arms and a full-support sampling distribution; equality holds after expectation
No deterministic pathwise bandit reconstruction	Proposition	At least two arms and arbitrary loss vectors
Bandit tangent and second-moment control	Proposition	Sampling distributions bounded away from the simplex boundary
Barycentric tangent reconstruction	Proposition	Differentiable full-support sampling and loss curves
Bandit regret transfer	Theorem	Oblivious losses, adapted full-support sampling, and a pathwise surrogate bound
Entropic adversarial bandit rate	Corollary	Finite arms, oblivious losses, and importance-weighted exponential updates
One-point smoothed-gradient reconstruction	Proposition	Convex Euclidean domain, ball smoothing, and differentiability of the average
Free convex completion for boosting	Proposition	Finite-support mixtures in $D(H)$ and pointwise convex structure on K^C
Conditional weak-to-strong amplification	Theorem	Requires a supplied potential contraction certificate; deriving that certificate from a categorical weak learner remains open
Decision-presentation quotient	Proposition	UDL semantics, readout, and observer must invert every registered presentation equivalence

Claim	Status	Boundary
Comparator-sketch representation	Theorem	Right-Kan-admissible reference-shape map, replete admissible class, and existence of the right Kan extension
Block-static comparator spectrum	Proposition and corollary	Finite chain, contiguous partitions, and real-valued minimization for regret monotonicity
Approximate comparator observer transfer	Proposition	Chosen trajectory metric and Lipschitz evaluation and observer; not supplied by universality
Intrinsic asynchronous execution	Definition	Finite event poset, processor ownership, causal-past reads, and an external linearization observer
Asynchronous schedule invariance	Theorem	Finite event poset and commuting update endomorphisms at every incomparable pair
Asynchronous distributed minimization and Q-learning	Worked interpretation	Local component objectives, declared stale reads, fair coordinate updates, and separately supplied stochastic and contraction observers
Endogenous comparator descent	Open problem	Requires closed-loop construction before testing descent along a reference shape
Generator-level agentic safety	Proposition	Realized information category generated by declared edges whose objects and morphisms lie in an admitted execution subcategory
Safety under information-shape extension	Theorem	Pointwise right Kan extension and closure of the admitted subcategory under the required comma-category limits
Asynchronous revocation obstruction	Proposition	Incomparable revocation and use events with a decision rule adapted only to the local causal past
Audit observer distinguishability	Proposition	Faithfulness is necessary only on safety-relevant executions and does not detect channels outside the observer's domain
Temporal contract descent	Proposition	Behavioral contract is a subobject in a sheaf topos and local verification witnesses agree on every overlap
Resilient intervention naturality	Open problem	Stop, veto, revocation, and rollback must commute with delegation and revoke derived authority
Independent-block repair	Proposition	Product decomposition on the registered context and queries factoring through protected complementary blocks
Canonical homotopy-coherent repair	Proposition	Contractible admissible filler space and homotopy-invariant downstream semantics
Quadratic tangent repair residual	Theorem	Finite-dimensional smooth realization, exact linearized repair, and Lipschitz derivative
Euclidean approachability checkpoint	Theorem	Closed convex target, bounded vector payoff, and Blackwell separating condition
Compositional structural transport	Theorem	Beck–Chevalley pasting, conservative observers, safety preservation, and coherent tangent comparisons when applicable
Finite persistence core	Theorem	Finite environment path and a common subobject admitting structural transport at every step; infinite persistence remains open
Causal interpretation	Not established	Requires intervention semantics and identification assumptions

Bibliography

Dana Angluin. Learning regular sets from queries and counterexamples. *Information and Computation*, 75(2):87–106, 1987. DOI: 10.1016/0890-5401(87)90052-6. URL [https://doi.org/10.1016/0890-5401\(87\)90052-6](https://doi.org/10.1016/0890-5401(87)90052-6).

Anthropic. Formalizing fermat’s last theorem in Lean: A timeline and selected excerpts from Claude’s reasoning, 2026. URL <https://www-cdn.anthropic.com/9e431dff043da6538d99d6c2d231b670aa3da263.pdf>. Accompanying Lean development available at <https://github.com/anthropics/fermats-last-theorem>.

Mahmoud Assran, Quentin Duval, Ishan Misra, Piotr Bojanowski, Pascal Vincent, Michael Rabbat, Yann LeCun, and Nicolas Ballas. Self-supervised learning from images with a joint-embedding predictive architecture. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 15619–15629, 2023. URL https://openaccess.thecvf.com/content/CVPR2023/html/Assran_Self-Supervised_Learning_From_Images_With_a_Joint-Embedding_Predictive_Architecture_CVPR_2023_paper.html.

Mido Assran et al. V-JEPA 2: Self-supervised video models enable understanding, prediction and planning. *arXiv preprint arXiv:2506.09985*, 2025. DOI: 10.48550/arXiv.2506.09985. URL <https://arxiv.org/abs/2506.09985>. Revised March 22, 2026.

Peter Auer, Nicolò Cesa-Bianchi, Yoav Freund, and Robert E. Schapire. The nonstochastic multiarmed bandit problem. *SIAM Journal on Computing*, 32(1):48–77, 2002. DOI: 10.1137/S0097539701398375.

Michael Barr and Charles Wells. *Category Theory for Computing Science*. Centre de Recherches Mathématiques, 3 edition, 1999. URL <https://www.tac.mta.ca/tac/reprints/articles/22/tr22.pdf>. Reprinted in *Theory and Applications of Categories*, No. 22 (2012).

- Dimitri P. Bertsekas and John N. Tsitsiklis. *Parallel and Distributed Computation: Numerical Methods*. Prentice-Hall, Englewood Cliffs, NJ, 1989. URL <https://www.mit.edu/~dimitrib/pdc.html>. Republished by Athena Scientific, 1997.
- David Blackwell. An analog of the minimax theorem for vector payoffs. *Pacific Journal of Mathematics*, 6(1):1–8, 1956. DOI: 10.2140/pjm.1956.6.1.
- Anselm Blumer, Andrzej Ehrenfeucht, David Haussler, and Manfred K. Warmuth. Learnability and the Vapnik–Chervonenkis dimension. *Journal of the ACM*, 36(4):929–965, 1989. DOI: 10.1145/76359.76371.
- John C. Bowler, Dua B. Azhar, Cambria M. Jensen, Hyun-Woo Lee, and James G. Heys. Structured experience shapes strategy learning and neural dynamics in the medial entorhinal cortex. *Nature Neuroscience*, September 2026. DOI: 10.1038/s41593-026-02409-7. URL <https://www.nature.com/articles/s41593-026-02409-7>. Published online September 3, 2026.
- Gunnar Carlsson and Jun Yu. A prime decomposition of probabilistic automata. *arXiv preprint arXiv:1503.01502*, 2015. DOI: 10.48550/arXiv.1503.01502. URL <https://arxiv.org/abs/1503.01502>.
- J. R. B. Cockett and G. S. H. Cruttwell. Differential structure, tangent structure, and SDG. *Applied Categorical Structures*, 22:331–417, 2014. DOI: 10.1007/s10485-013-9312-0.
- Robin Cockett, Jean-Simon Pacaud Lemay, and Rory B. B. Lucyshyn-Wright. Tangent categories from the coalgebras of differential categories. In *28th EACSL Annual Conference on Computer Science Logic (CSL 2020)*, volume 152 of *Leibniz International Proceedings in Informatics (LIPIcs)*, pages 17:1–17:17. Schloss Dagstuhl–Leibniz-Zentrum für Informatik, 2020. DOI: 10.4230/LIPIcs.CSL.2020.17.
- Daice Labs. Category theory as the language of composable AI. Daice Labs research synthesis, 2026. URL <https://daicelabs.com/research/category-theory>. Category Theory, Machine Learning, and AI Safety.
- Gary L. Drescher. *Made-Up Minds: A Constructivist Approach to Artificial Intelligence*. The MIT Press, Cambridge, MA, 1991. ISBN 9780262041201. URL <https://mitpress.mit.edu/9780262517089/made-up-minds/>. Paperback edition published 2003.
- Charles Ehresmann. Esquisses et types des structures algébriques. *Buletinul Institutului Politehnic din Iași*, 14:1–14, 1968.

- Abraham D. Flaxman, Adam Tauman Kalai, and H. Brendan McMahan. Online convex optimization in the bandit setting: Gradient descent without a gradient. In *Proceedings of the Sixteenth Annual ACM-SIAM Symposium on Discrete Algorithms*, pages 385–394. Society for Industrial and Applied Mathematics, 2005. URL <https://arxiv.org/abs/cs/0408007>.
- Brendan Fong, David I. Spivak, and Rémy Tuyéras. Backprop as functor: A compositional perspective on supervised learning. In *Proceedings of the 34th Annual ACM/IEEE Symposium on Logic in Computer Science, LICS 2019*, pages 1–13, 2019. DOI: 10.1109/LICS.2019.8785665. URL <https://arxiv.org/abs/1711.10455>.
- E. Mark Gold. Limiting recursion. *The Journal of Symbolic Logic*, 30(1): 28–48, 1965. DOI: 10.2307/2270580. URL <https://doi.org/10.2307/2270580>.
- E. Mark Gold. Language identification in the limit. *Information and Control*, 10(5):447–474, 1967. DOI: 10.1016/S0019-9958(67)91165-5. URL [https://doi.org/10.1016/S0019-9958\(67\)91165-5](https://doi.org/10.1016/S0019-9958(67)91165-5).
- Ryan Greenblatt, Ajeya Cotra, and Hjalmar Wijk. Brief independent investigation of agents’ behavior, reasoning and collaboration in the OpenAI/Hugging Face hacking incident. METR research report, August 2026. URL <https://metr.org/blog/2026-08-26-openai-hugging-face-incident-investigation/>. Published August 26, 2026.
- David Ha and Jürgen Schmidhuber. World models. *arXiv preprint arXiv:1803.10122*, 2018. DOI: 10.48550/arXiv.1803.10122. URL <https://arxiv.org/abs/1803.10122>.
- Redi Haderi, Cihan Okay, and Walker H. Stern. The operadic theory of convexity. *CoRR*, abs/2403.18102, 2024. URL <https://arxiv.org/abs/2403.18102>.
- Elad Hazan. Introduction to online convex optimization. *Foundations and Trends in Optimization*, 2(3–4):157–325, 2016. DOI: 10.1561/24000000013.
- Elad Hazan. *Introduction to Online Convex Optimization*. arXiv, 2 edition, 2023. URL <https://arxiv.org/abs/1909.05207>. arXiv:1909.05207v3.
- Alec Helbling, Andrey Bryutkin, Mauro Martino, Duen Horng Chau, Nima Dehmamy, and Hendrik Strobelt. Flow reasoning models: Turning flows into efficient recurrent reasoners. *CoRR*, abs/2606.29150, 2026. DOI: 10.48550/arXiv.2606.29150. URL <https://arxiv.org/abs/2606.29150>. Version 3.

- William James. *The Principles of Psychology*, volume 1. Henry Holt and Company, New York, 1890. URL <https://psychclassics.yorku.ca/James/Principles/index.htm>.
- Daniel M. Kan. Adjoint functors. *Transactions of the American Mathematical Society*, 87(2):294–329, 1958. DOI: 10.2307/1993102.
- G. M. Kelly. *Basic Concepts of Enriched Category Theory*, volume 64 of *London Mathematical Society Lecture Note Series*. Cambridge University Press, 1982. URL <https://www.tac.mta.ca/tac/reprints/articles/10/tr10abs.html>. Reprinted in *Theory and Applications of Categories*, No. 10 (2005).
- Anders Kock. *Synthetic Differential Geometry*. Cambridge University Press, 2 edition, 2006. ISBN 978-0-521-68738-6. DOI: 10.1017/CBO9780511550812. URL <https://www.cambridge.org/core/books/synthetic-differential-geometry/4933FF5EDF69BD13E8D8FCA346BB1C06>.
- Kenneth Krohn and John Rhodes. Algebraic theory of machines. I. prime decomposition theorem for finite semigroups and machines. *Transactions of the American Mathematical Society*, 116:450–464, 1965. DOI: 10.1090/S0002-9947-1965-0188316-1. URL <https://www.ams.org/journals/tran/1965-116-00/S0002-9947-1965-0188316-1/>.
- Stephen Lack. Composing PROPs. *Theory and Applications of Categories*, 13(9):147–163, 2004. URL <https://tac.mta.ca/tac/volumes/13/9/13-09abs.html>.
- Brenden M. Lake and Marco Baroni. Generalization without systematicity: On the compositional skills of sequence-to-sequence recurrent networks. In *Proceedings of the 35th International Conference on Machine Learning*, volume 80 of *Proceedings of Machine Learning Research*, pages 2873–2882, 2018. URL <https://proceedings.mlr.press/v80/lake18a.html>.
- Brenden M. Lake and Marco Baroni. Human-like systematic generalization through a meta-learning neural network. *Nature*, 623(7985): 115–121, 2023. DOI: 10.1038/s41586-023-06668-3.
- Brenden M. Lake, Tomer D. Ullman, Joshua B. Tenenbaum, and Samuel J. Gershman. Building machines that learn and think like people. *Behavioral and Brain Sciences*, 40:e253, 2017. DOI: 10.1017/S0140525X16001837.
- Tor Lattimore and Csaba Szepesvári. *Bandit Algorithms*. Cambridge University Press, 2020. ISBN 9781108486828. DOI: 10.1017/9781108571401. URL <https://www.cambridge.org/core/books/bandit-algorithms/8E39FD004E6CE036680F90DD0C6F09FC>.

- F. William Lawvere. Functorial semantics of algebraic theories. *Proceedings of the National Academy of Sciences*, 50(5):869–872, 1963. DOI: 10.1073/pnas.50.5.869. URL <https://pmc.ncbi.nlm.nih.gov/articles/PMC221940/>.
- Yann LeCun. A path towards autonomous machine intelligence. OpenReview position paper, version 0.9.2, June 2022. URL <https://openreview.net/forum?id=BZ5a1r-kVsf>.
- Michael L. Littman, Richard S. Sutton, and Satinder Singh. Predictive representations of state. In *Advances in Neural Information Processing Systems 14*, pages 1555–1561, 2001. URL <https://proceedings.neurips.cc/paper/2001/hash/1e4d36177d71bbb3558e43af9577d70e-Abstract.html>.
- Bingbin Liu, Jordan T. Ash, Surbhi Goel, Akshay Krishnamurthy, and Cyril Zhang. Transformers learn shortcuts to automata. In *International Conference on Learning Representations*, 2023. DOI: 10.48550/arXiv.2210.10749. URL <https://openreview.net/forum?id=De4FYqjFueZ>. Oral presentation.
- Saunders Mac Lane. *Categories for the Working Mathematician*, volume 5 of *Graduate Texts in Mathematics*. Springer, New York, 1971. ISBN 978-1-4612-9839-7. DOI: 10.1007/978-1-4612-9839-7. URL <https://link.springer.com/book/10.1007/978-1-4612-9839-7>.
- Sridhar Mahadevan. Universal decision models. *CoRR*, abs/2110.15431, 2021. URL <https://arxiv.org/abs/2110.15431>.
- Sridhar Mahadevan. Universal imitation games. arXiv preprint arXiv:2405.01540, 2024. URL <https://arxiv.org/abs/2405.01540>.
- Sridhar Mahadevan. *Categories for Artificial General Intelligence*. Categorical AI Book Project, 2026a. Book manuscript.
- Sridhar Mahadevan. *Infinitesimal Creativity: Learning by Double Involution*. The Infinitesimal Creativity Book Project, 2026b.
- Sridhar Mahadevan. *Machine Learning from Enforcing Compositionality*. The LINC Book Project, 2026c.
- Sridhar Mahadevan. ODYSSEY: Constructing verifiable local truth-preserving foundation models, 2026d. URL <https://arxiv.org/abs/2606.27593>.
- Sridhar Mahadevan. PROMETHEUS: Automating deep causal research integrating text, data and models, 2026e. URL <https://arxiv.org/abs/2605.12835>.

- H. Brendan McMahan. Follow-the-regularized-leader and mirror descent: Equivalence theorems and L1 regularization. In *Proceedings of the Fourteenth International Conference on Artificial Intelligence and Statistics*, volume 15 of *Proceedings of Machine Learning Research*, pages 525–533. PMLR, 2011. URL <https://proceedings.mlr.press/v15/mcmahan11b.html>.
- METR. Task-completion time horizons of frontier AI models. Online research dashboard, 2026. URL <https://metr.org/time-horizons/>. Last updated May 8, 2026.
- Andrea Moro. *Impossible Languages*. The MIT Press, Cambridge, MA, 2023. ISBN 9780262549233. URL <https://mitpress.mit.edu/9780262549233/impossible-languages/>. Paperback edition, published September 19, 2023.
- OpenAI. Path to Astra: Critical capabilities and frontier safeguards, September 2026a. URL <https://openai.com/index/path-to-astra/>. Published September 1, 2026.
- OpenAI. Separating signal from noise in coding evaluations. Research report, July 2026b. URL <https://openai.com/index/separating-signal-from-noise-coding-evaluations/>. Published July 8, 2026.
- Judea Pearl. *Causality: Models, Reasoning, and Inference*. Cambridge University Press, 2 edition, 2009. ISBN 978-0-521-89560-6. DOI: 10.1017/CBO9780511803161. URL <https://www.cambridge.org/core/books/causality/B0046844FAE10CBF274D4ACBDAEB5F5B>.
- Jonas Peters, Dominik Janzing, and Bernhard Schölkopf. *Elements of Causal Inference: Foundations and Learning Algorithms*. Adaptive Computation and Machine Learning. MIT Press, 2017. ISBN 978-0-262-03731-0. URL <https://mitpress.mit.edu/9780262037310/elements-of-causal-inference/>.
- Jean Piaget. *The Origins of Intelligence in Children*. International Universities Press, New York, 1952. URL <https://iucat.iu.edu/iub/3827762>.
- Top Piriyakulkij, Yichao Liang, Hao Tang, Adrian Weller, Marta Kryven, and Kevin Ellis. PoE-World: Compositional world modeling with products of programmatic experts. In *Advances in Neural Information Processing Systems*, volume 38, 2025. DOI: 10.52202/085713-0896. URL https://proceedings.neurips.cc/paper_files/paper/2025/hash/262dd62fd1bbb30d6a6b4d578f5e65ff-Abstract-Conference.html. Code: <https://github.com/topwasu/poe-world>.

- Boris Plotkin and Tatjana Plotkin. Decompositions and complexity of linear automata. *arXiv preprint arXiv:1506.06017*, 2015. DOI: 10.48550/arXiv.1506.06017. URL <https://arxiv.org/abs/1506.06017>.
- Martin L. Puterman. *Markov Decision Processes: Discrete Stochastic Dynamic Programming*. John Wiley & Sons, 1994. ISBN 978-0-471-61977-2. DOI: 10.1002/9780470316887. URL <https://onlinelibrary.wiley.com/doi/book/10.1002/9780470316887>.
- Neil Rabinowitz, Frank Perbet, Francis Song, Chiyuan Zhang, S. M. Ali Eslami, and Matthew Botvinick. Machine theory of mind. In *Proceedings of the 35th International Conference on Machine Learning*, volume 80 of *Proceedings of Machine Learning Research*, pages 4218–4227, 2018. URL <https://proceedings.mlr.press/v80/rabinowitz18a.html>.
- Sashank J. Reddi, Satyen Kale, and Sanjiv Kumar. On the convergence of Adam and beyond. *CoRR*, abs/1904.09237, 2019. URL <https://arxiv.org/abs/1904.09237>.
- Birgit Richter. *From Categories to Homotopy Theory*, volume 188 of *Cambridge Studies in Advanced Mathematics*. Cambridge University Press, 2020. DOI: 10.1017/9781108855891. URL <https://www.cambridge.org/core/books/from-categories-to-homotopy-theory/A109E2C4B720337DE19A15EB4FA8C9A6>.
- Emily Riehl. *Categorical Homotopy Theory*, volume 24 of *New Mathematical Monographs*. Cambridge University Press, 2014. DOI: 10.1017/CBO9781107261457.
- Emily Riehl. *Category Theory in Context*. Dover Publications, 2016. URL <https://emilyriehl.github.io/books/>.
- Emily Riehl and Dominic Verity. *Elements of ∞ -Category Theory*. Cambridge University Press, 2022. DOI: 10.1017/9781108936880. URL <https://www.cambridge.org/core/books/elements-of-infinity-category-theory/DAC48C449AB8C2C1B1E528A49D27FC6D>.
- Ronald L. Rivest and Robert E. Schapire. Diversity-based inference of finite automata. *Journal of the ACM*, 41(3):555–589, 1994. DOI: 10.1145/176584.176589. URL <https://people.csail.mit.edu/rivest/pubs/RS94a.pdf>.
- R. Tyrrell Rockafellar. Monotone operators and the proximal point algorithm. *SIAM Journal on Control and Optimization*, 14(5):877–898, 1976. DOI: 10.1137/0314056.

- R. Tyrrell Rockafellar. Proto-differentiability of set-valued mappings and its applications in optimization. *Annales de l'Institut Henri Poincaré C, Analyse non linéaire*, 6:449–482, 1989. DOI: 10.1016/S0294-1449(17)30034-3.
- Alessandro Ronca, Nadezda Alexandrovna Knorozova, and Giuseppe De Giacomo. Automata cascades: Expressivity and sample complexity. In *Proceedings of the Thirty-Seventh AAAI Conference on Artificial Intelligence*, volume 37, pages 9588–9595, 2023. DOI: 10.1609/aaai.v37i8.26147. URL <https://ojs.aaai.org/index.php/AAAI/article/view/26147>. Issue 8.
- Jan J. M. M. Rutten. Universal coalgebra: A theory of systems. *Theoretical Computer Science*, 249(1):3–80, 2000. DOI: 10.1016/S0304-3975(00)00056-6. URL [https://doi.org/10.1016/S0304-3975\(00\)00056-6](https://doi.org/10.1016/S0304-3975(00)00056-6).
- Jenny R. Saffran, Richard N. Aslin, and Elissa L. Newport. Statistical learning by 8-month-old infants. *Science*, 274(5294):1926–1928, 1996. DOI: 10.1126/science.274.5294.1926.
- Jerome H. Saltzer and Michael D. Schroeder. The protection of information in computer systems. *Proceedings of the IEEE*, 63(9):1278–1308, 1975. DOI: 10.1109/PROC.1975.9939.
- Patrick Schultz and David I. Spivak. *Temporal Type Theory: A Topos-Theoretic Approach to Systems and Behavior*. arXiv, 2017. DOI: 10.48550/arXiv.1710.10258. URL <https://arxiv.org/abs/1710.10258>. arXiv:1710.10258v3, 224 pages.
- Shai Shalev-Shwartz. Online learning and online convex optimization. *Foundations and Trends in Machine Learning*, 4(2):107–194, 2012. DOI: 10.1561/22000000018.
- Elizabeth S. Spelke. *What Babies Know: Core Knowledge and Composition*, volume 1. Oxford University Press, New York, 2022. ISBN 9780190618247. DOI: 10.1093/oso/9780190618247.001.0001. URL <https://academic.oup.com/book/43912>.
- Elizabeth S. Spelke. Précis of *What Babies Know*. *Behavioral and Brain Sciences*, 47:e120, 2023. DOI: 10.1017/S0140525X23002443. URL <https://pubmed.ncbi.nlm.nih.gov/37248696/>.
- Richard S. Sutton and Andrew G. Barto. *Reinforcement Learning: An Introduction*. MIT Press, 2 edition, 2018. URL <http://incompleteideas.net/book/the-book-2nd.html>.

- Richard S. Sutton, Doina Precup, and Satinder Singh. Between MDPs and semi-MDPs: A framework for temporal abstraction in reinforcement learning. *Artificial Intelligence*, 112(1–2):181–211, 1999. DOI: 10.1016/S0004-3702(99)00052-1. URL <https://www.sciencedirect.com/science/article/pii/S0004370299000521>.
- Denis Thérien. Two-sided wreath product of categories. *Journal of Pure and Applied Algebra*, 74(3):307–315, 1991. DOI: 10.1016/0022-4049(91)90119-M.
- John N. Tsitsiklis. Asynchronous stochastic approximation and Q-learning. *Machine Learning*, 16(3):185–202, 1994. DOI: 10.1023/A:1022689125041. URL <https://www.mit.edu/~jnt/Papers/J052-94-jnt-q.pdf>.
- Leslie G. Valiant. A theory of the learnable. *Communications of the ACM*, 27(11):1134–1142, 1984. DOI: 10.1145/1968.1972.
- Leslie G. Valiant. Evolvability. *Journal of the ACM*, 56(1):1–21, 2009. DOI: 10.1145/1462153.1462156. URL <https://dl.acm.org/doi/10.1145/1462153.1462156>.
- V. N. Vapnik and A. Ya. Chervonenkis. On the uniform convergence of relative frequencies of events to their probabilities. *Theory of Probability and Its Applications*, 16(2):264–280, 1971. DOI: 10.1137/1116025.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz Kaiser, and Illia Polosukhin. Attention is all you need. In *Advances in Neural Information Processing Systems*, volume 30, pages 5998–6008, 2017. URL <https://proceedings.neurips.cc/paper/2017/hash/3f5ee243547dee91fbd053c1c4a845aa-Abstract.html>.
- Steven Weinberg. A model of leptons. *Physical Review Letters*, 19(21):1264–1266, 1967. DOI: 10.1103/PhysRevLett.19.1264. URL <https://journals.aps.org/prl/abstract/10.1103/PhysRevLett.19.1264>.
- Charles Wells. A Krohn–Rhodes theorem for categories. *Journal of Algebra*, 64(1):37–45, 1980. DOI: 10.1016/0021-8693(80)90130-1.
- Charles Wells. Wreath product decomposition of categories. II. *Acta Scientiarum Mathematicarum*, 52:321–324, 1988. URL https://acta.bibl.u-szeged.hu/15191/1/math_052_fasc_003_004_321-324.pdf.
- Eugene P. Wigner. On unitary representations of the inhomogeneous lorentz group. *Annals of Mathematics*, 40(1):149–204, 1939. DOI: 10.2307/1968551. URL <https://cds.cern.ch/record/405887/>.

Hans S. Witsenhausen. On information structures, feedback and causality. *SIAM Journal on Control*, 9(2):149–160, 1971. DOI: 10.1137/0309013.

Hans S. Witsenhausen. The intrinsic model for discrete stochastic control: Some open problems. In *Control Theory, Numerical Methods and Computer Systems Modelling*, volume 107 of *Lecture Notes in Economics and Mathematical Systems*, pages 322–335. Springer, Berlin, 1975.

Chen Ning Yang and Robert L. Mills. Conservation of isotopic spin and isotopic gauge invariance. *Physical Review*, 96(1):191–195, 1954. DOI: 10.1103/PhysRev.96.191.

Hongxin Zhang, Zeyuan Wang, Qiushi Lyu, Zheyuan Zhang, Sunli Chen, Tianmin Shu, Behzad Dariush, Kwonjoon Lee, Yilun Du, and Chuang Gan. COMBO: Compositional world models for embodied multi-agent cooperation. In *International Conference on Learning Representations*, 2025. URL https://proceedings.iclr.cc/paper_files/paper/2025/hash/7d03c6bf9f07acb4038eea96c63db52d-Abstract-Conference.html. Project page: <https://umass-embodied-agi.github.io/COMBO/>.

Siyuan Zhou, Yilun Du, Jiaben Chen, Yandong Li, Dit-Yan Yeung, and Chuang Gan. RoboDreamer: Learning compositional world models for robot imagination. In *Proceedings of the 41st International Conference on Machine Learning*, volume 235 of *Proceedings of Machine Learning Research*, pages 61885–61896. PMLR, 2024. URL <https://proceedings.mlr.press/v235/zhou24f.html>. Project page: <https://robovideo.github.io>; code: <https://github.com/rainbow979/robodreamer>.

Glossary of Notation

The same symbol may be specialized differently in different chapters. This glossary records the stable conventions of the ORACLE, UOCL, and UODL frameworks; the types and local declarations stated in a chapter take precedence whenever a symbol is specialized.

Categorical learning and world models

Term or symbol	Meaning in this book
ORACLE	The semantic program of identifying an unknown categorical world from an interaction presentation, up to a declared query equivalence.
UOCL	Universal Online Categorical Learning: an algorithmic realization of ORACLE through typed probe, selection, assimilation, and repair rules.
UOCL specialization	A learning paradigm with an explicit hypothesis category, presentation protocol, query doctrine, observational quotient, non-anticipatory update, and success criterion.
Algebraic theory	A category of formal operations and equations; Lawvere theories use finite products, while PROPs use one-generated symmetric monoidal structure.
Sketch	A compact presentation by generating objects and arrows, commutative diagrams, and designated cones or cocones.
Generatively presented world	A world category together with a sketch, preservation doctrine, completed theory, and structure-preserving interpretation of that theory in the world.
Theory-bearing UOCL	UOCL that learns generators, relations, theory completion, and functorial model semantics in addition to an encountered world category.
Open UOCL learner	A tuple $(\mathcal{H}, I, U, R)$ between categorical learning interfaces: an indexed hypothesis category, implementation, typed update, and backward request or repair-obligation map.
Multimodal fusion	A limit, commonly a pullback, of modality-specific UOCL learners over a shared doctrine; stronger than their parallel monoidal product because it enforces cross-modal compatibility.
TUOCL	Tangent UOCL: categorical identification in a hypothesis class of tangent categories, using both base and tangent-sensitive probes.
PACC	Probably Approximately Categorically Correct learning: a statistical UOCL guarantee relative to a declared categorical probe doctrine, answer equivalence, discrepancy, distribution, accuracy, and confidence.
$\text{Err}_{P,Q}$	Expected equivalence-invariant discrepancy between a target and learned categorical world under probes drawn from P .
Structural coverage	A requirement that the probe distribution assign positive mass to every categorical stratum whose correctness is claimed.
CPdim	Categorical probe dimension: the largest compatible family of binary categorical probes whose coherence-admissible answer patterns are realized by the hypothesis class.
Diff. realization	A tangent hypothesis induced by the category of coalgebras of a differential category.
Coalgebra learning	Identification of a system $c : X \rightarrow FX$ from its presentation, generally only up to the behavioral equivalence detected by its map into a final F -coalgebra.
PSR	Predictive state representation: a controlled-process state given by probabilities of future action–observation tests, rather than by a privileged latent ontology.
Topos World Model	A presheaf or sheaf of local, context-indexed predictive models with restriction maps, overlap audits, gluing, provenance, and registered obstructions.

Online, homotopical, and infinitesimal structure

Term or symbol	Meaning in this book
UODL	Universal Online Decision Learning: the specialization of UOCL to hypotheses carrying information, action, consequence, consistency, and observer structure.
$\text{Lan}_J, \text{Ran}_K$	Left and right Kan extension along the displayed functor.
$\mathbf{U}(F)$	Ordinary UDL semantics $\text{Ran}_K \text{Lan}_J F$.
$\mathbf{U}^h(F)$	Homotopy-coherent or derived UDL semantics.
$\mathbb{T}$	Time or revelation category.
$\mathcal{NC}$	Ordinary nerve of a category.
Δ/X	Category of simplices of a simplicial set X .
Δ_i^n	The i -th horn in the n -simplex.
$\mathcal{D}_W$	The localization $\mathcal{D}[W^{-1}]$ at declared semantic weak equivalences.
T	Tangent functor; T^h denotes a derived tangent functor when one exists.
$W, \mathfrak{m}_W$	A Weil algebra and its nilpotent maximal ideal.
$\mathbb{D}_W$	Infinitesimal probe represented by a Weil algebra W .
$L_{\mathcal{M}}$	Linearized information-dependency operator $D_u \Phi_{\mathcal{M}}$ of an intrinsic model.

Decision comparison, repair, and persistence

Term or symbol	Meaning in this book
Agg_A	Declared aggregation of local values over decision sites A ; ordinary summation is one specialization.
$\text{Reg}_{\mathcal{I} \rightarrow \mathcal{J}}^\Omega$	Intrinsic regret observed between learner information $\mathcal{I}$ and comparator information $\mathcal{J}$.
Fill	A space of compatible horn fillers or repairs.
Def^h	Homotopy-coherent online-to-offline defect.
Registered repair	A partial decorated simplex together with settled and open obligations, an admissible filler space, and a comparison observer.
Accommodation certificate	A recorded obstruction to repair in the old declaration, a revised declaration with a filler, and witnesses identifying the obligations preserved by the change.
Structural knowledge package	A learned diagram, universal construction and witness, repair certificates, and the observational quotient on which they are valid.
Utility profile	The family of future diagrammatic problems that factor through a learned package with admitted witnesses.
Persistence core	The maximal registered substructure on which learned universal structure, observers, repairs, and safety admission transport coherently across a specified environment path.

Boundary

Superscript h does not make a construction homotopically meaningful by notation alone. The weak equivalences, localization, existence of derived functors, and preservation conditions must be declared.

Subject Index

- accommodation, 25, 82
- action groupoid, 51
- adjunction, 47
- agentic safety, 227
- algebraic theory, 44
 - as identification target, 67
- assimilation, 25, 82
- asynchronous computation, 222
- audit observer, 227
- automata learning, 115

- bandit convex optimization, 215
- bandit feedback, 211
- behavior type, 231
- behavioral shaping, 27
- Bellman equation
 - as doctrinal structure, 123

- Cartesian differential category, 56
- cascade product, 117
- categorical identification, 65
- category, 39
- causal discovery
 - as UOCL, 126
- classical information structure, 59, 188
- coalgebra
 - learning, 115
- coalgebra modality, 56
- coEilenberg–Moore realization, 154
- coinductive inference, 74
- colimit, 43
 - of UOCL learners, 147
- COMBO, 133
- comma category, 47
- comparator sketch, 219
- compositional authority, 227
- compositional compression, 29
- compositional generalization
 - robotics, 130

- compositional learning
 - animal evidence, 27
- compositional world, 25
- compositionality
 - architectural versus categorical, 132
- constraint satisfaction, 193
- core knowledge, 25
 - as doctrine, 40, 89
 - categorical content, 89
- core-relative identification, 82
- core-structured world, 69
- crossword puzzle, 191
- curriculum learning, 27

- decentralized decision making, 191
- derived functor, 53
- deterministic dilation, 119
- developmental time, 101
- differential category, 56
- differential realization, 154
- diversity representation, 29, 116
- doctrinal audit, 109
- doctrine, 40
 - core knowledge, 89
 - fixed-doctrine learning, 109
 - preservation, 40
 - query, 40
 - structural, 40
- Drescher, Gary, 25

- environment morphism, 255
- event category, 222
- exploration
 - tangent admissibility, 213

- fixed point
 - reasoning dynamics, 193
- flip-flop automaton, 117
- Flow Reasoning Model, 193

- functor category, 49

- game problem, 59
- generators and relations, 143
- global consistency, 191
- Grothendieck construction, 43
- Grothendieck universe, 43
- group complexity, 117

- hierarchical reinforcement learning,
 - 125
- homotopy coherence, 39
- homotopy colimit, 53
- homotopy limit, 53
- homotopy-coherent diagram, 53
- homotopy-coherent nerve, 53
- horn, 51
- hypothesis fibration, 77

- identification
 - versus universal completion, 99
- identification in the limit, 77
- importance weighting, 211
- infinitesimal intrinsic model, 58
- infinitesimal probe, 55
- ∞ -category, 39
- information category, 219
- information channel, 211
- information field, 58
- information structure
 - agentic teams, 227
- information-relative regret, 188
- intrinsic model, 58
- intrinsic regret, 188
- irreversibility, 51

- JEPA
 - as UOCL, 130

- Kan complex, 51
- Kan extension, 39
 - insufficiency for identification, 99
 - left, 47
 - right, 47
- Krohn–Rhodes theorem, 117
 - stochastic extension, 119
- language identification, 127
- Lawvere theory, 44
- LEGO
 - generative theory, 29
- LEGO example, 46
- limit, 43
 - crossword solution, 191
 - of UOCL learners, 147
- LINCS, 55
- localization, 52
- Markov decision process, 123
 - cascade decomposition, 119
- medial entorhinal cortex, 27
- model
 - of a theory, 44
- multi-agent planning
 - compositional world model, 133
- multimodal learning, 145
 - fusion, 147
- nerve, 49
- non-anticipation, 49
- nonclassical information structure, 59, 188
- obstruction
 - ledger, 237
- Odyssey, 113
- options
 - compositional persistence, 125
- ORACLE, 65
- partial observability
 - local views, 133
- particle physics
 - as compositional discovery, 33
 - collider declaration, 67
 - symmetry doctrine, 40
- persistent identification, 151
- Piagetian learning, 25, 101
- PoE-World, 132
- policy regret, 188
- predictive state representation, 111
 - operator decomposition, 119
- presentation invariance, 151
- probabilistic automaton, 119
- programmatic world model, 132
- Prometheus, 113
- PROP, 44
- quasi-category, 51
- quasiclassical information structure, 59
- reinforcement learning
 - as UOCL, 123
- related work, 109
- relative category, 52
- repair
 - registered problem, 237
- repair space, 55
- response functor
 - collider, 67
- revocation
 - asynchronous, 229
- right Kan extension
 - comparator semantics, 219
- RoboDreamer, 130
- robot imagination, 130
- Rubik’s Cube, 51
- schedule invariance, 222
- schema mechanism, 25
- semiautomaton, 128
- shortcut learning, 128
- simplex category, 49
- simplicial set, 39
- skeletal filtration, 49
- sketch
 - learned presentation, 67
 - presentation, 44
- smoothing, 215
- Sokoban, 51
- Spelke hypothesis, 89
- Spelke sketch, 69
- stable ∞ -category, 54
- Standard Model, 33
- static information structure, 59
- structural causal model, 126
- Sudoku, 193
- symmetric monoidal category
 - of UOCL learners, 145
- system identification, 111
- tangent category, 39, 55
 - as identification target, 68
- tangent UOCL, 153
- team problem, 59
- temporal abstraction, 125
- temporal type theory, 231
- test equivalence, 116
- Topos World Model, 113
- Transformer
 - as UOCL, 127
 - Krohn–Rhodes shortcut, 128
- transport
 - structural, 255
- trust
 - temporal contract, 231
- Turing machine
 - identification, 115
- universal imitation game, 74
- Universal Online Categorical Learning, 139
- universal property, 43
- UOCL, 139
 - composition, 145
 - specializations, 109
 - theory learning, 143
- verification
 - sheaf descent, 231
- weak equivalence, 52
- Weil algebra, 55
- world model, 130
- wreath product, 117
- Yoneda lemma
 - accelerator analogy, 33