

WHITEPAPER • JULY 2026

# $(\Omega, D)$ Dynamics

## A Falsifiable Architecture for Homeostatic Action Selection

*Current findings, limits, evidence, and application directions*

$$x_{t+1} = E(\hat{y}(x_t)) \quad \text{where} \quad \hat{y}(x) \in \arg \min D(y, x^*; h), \quad y \in \Omega(x)$$

### The central idea in one sentence

*Choose the available action that best preserves the system's declared reference state—then let the environment respond, without having been modelled inside the choice.*

Prepared from the consolidated mathematical framework and implemented validation laboratories

## Executive summary

(Ω, D) Dynamics is a constrained architecture for choosing among available actions. It separates three roles that many decision systems combine: the menu of actions a system can express (Ω), the metric by which it ranks internally projected states (D), and the external environment that acts after selection (E). Its defining commitment is E-blindness: the selector may use a self-model and declared static affordances, but it may not model, estimate, or anticipate environmental response.

That restriction is both the framework's identity and its limitation. It creates predictions that model-predictive control and model-based reinforcement learning do not make. More internal lookahead can avoid traps contained in the action menu, yet remain blind to traps introduced only by the environment. In some cases, additional lookahead makes the system less safe because it trusts a continuation that the environment never delivers.

**BOTTOM LINE** The framework is best understood today as a falsifiable architecture and experimental null model for homeostatic action selection—not as a demonstrated replacement for causality, reinforcement learning, or modern control theory.

## What is established

- Well-posed selection can be guaranteed under standard compactness, lower-semicontinuity, and selection conditions.
- Persistence has a sharp contraction threshold: the selector's improvement must be strong enough to overcome environmental expansion.
- Without structural restrictions, the update rule can rationalize almost any trajectory and has no empirical content.
- Fixed, overlapping menus create revealed-preference tests; rigid metric-projection models are generically refutable and partly identifiable.
- E-blind horizon planning is formally separable from true-environment MPC through trap, variance, range, and geometric signatures.
- Implemented tests reproduce the predicted Ω-trap, E-trap, and horizon-fragility patterns across abstract graphs and two non-Minecraft benchmarks.

## What is not yet established

- The broader philosophical claim that causality is reducible to dynamic informational similarity.
- External validity in independently built physical, biological, organizational, or production software systems.

- A general identification theory for unknown menu shapes, stochastic environments, or single-trajectory data.
- A proof that E-blind control is preferable where anticipating environmental consequences is safety-critical.

## How to read this paper

Sections 1–3 explain the architecture and its main mathematical findings in plain language. Sections 4–5 review evidence and limitations. Sections 6–7 identify promising applications and the validation standard required before making scientific or engineering claims. The appendices provide notation, theorem status, and selected references.

## 1. The framework in plain language

A persistent system is represented by an extended state  $x$ : everything carried forward that matters for future behavior, such as position, energy, inventory, configuration, phase, or bounded memory. At each step, the system does not choose from every imaginable future. It receives a typed menu  $\Omega(x)$  of actions or candidate next states that its current mechanisms can actually express.

The selector evaluates each candidate by its discontinuity  $D$  from a frozen reference manifold  $x^*$ . A reference manifold is a set of acceptable states rather than necessarily a single target point: upright and balanced, alive and fed, operational and below an error threshold, or metabolically viable. A horizon  $h$  allows the selector to inspect paths through its own future menus. Once a candidate is chosen, the environment operator  $E$  produces the realized next state.

$$\hat{y}(x) \in \operatorname{argmin}_{y \in \Omega(x)} D(y, x^*; h) \quad \text{and} \quad x_{t+1} = E(\hat{y}(x_t))$$

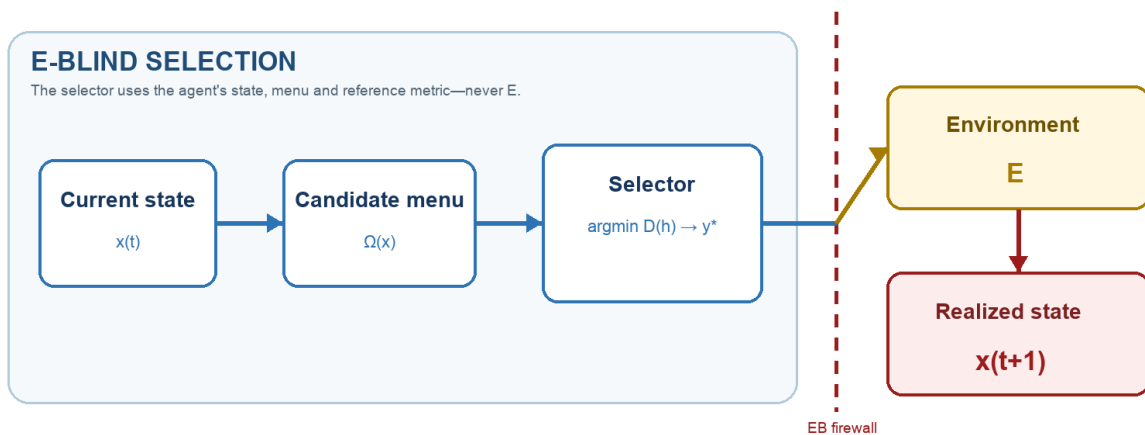


Figure 1. The defining factorization: selection is E-blind; environmental response is strictly post-selection.

## The one commitment that makes it distinct

If the selector evaluates  $D(E(y), x^*)$  or learns a predictive model of  $E$ , it becomes an anticipating controller. That may be the better engineering choice, but it is no longer the framework studied here. The honest one-line description is: MPC with a self-model and no world-model.

**E-BLINDNESS** The agent may learn about its bounded self-history and about actuators only under calibration that isolates them from environmental response. A realized error observed in an ordinary run is jointly caused by actuator and environment; treating it automatically as self-learning would move  $E$  back into the selector.

## The symbols

| Symbol      | Plain-language meaning                                                                |
|-------------|---------------------------------------------------------------------------------------|
| $x$         | Extended system state: all declared information carried forward.                      |
| $\Omega(x)$ | Mechanism-specified menu of currently available candidate actions or states.          |
| $x^*$       | Frozen reference manifold representing acceptable or completed states.                |
| $D$         | Declared ranking of candidates by distance or discontinuity from $x^*$ .              |
| $h$         | Depth of lookahead through the system's own menus, with internal $E$ = identity.      |
| $E$         | External response after selection: physics, workload, other agents, shocks, or noise. |
| $X^\dagger$ | Death or terminal set where no candidate remains or viability has failed.             |

## 2. What turns notation into a testable theory

An unconstrained argmin expression can explain almost anything: make every observed action the only available candidate, change the metric after seeing the data, hide time inside an unobserved state coordinate, or allow every candidate to tie. The framework's empirical content therefore lives in its restrictions, not in the bare equation.

| Restriction | Commitment                                                                           | Loophole it closes                       |
|-------------|--------------------------------------------------------------------------------------|------------------------------------------|
| R0          | Observable or parsimoniously constrained extended state                              | Hidden clocks and arbitrary indexing     |
| R1-C / R1-H | Bounded mechanism-specified continuous or typed menus, with declared overlap         | Singleton and per-trajectory menus       |
| R2          | Metric selected from a declared class before data                                    | Metrics invented to fit each observation |
| R3          | The same $\Omega$ and $D$ across trajectories and environments                       | Per-run explanation                      |
| R4          | Frozen or autonomous reference manifold                                              | Moving the goal to follow the data       |
| C1 / C1-T   | Unique effective ordering, including any frozen secondary rank                       | Universal rationalization through ties   |
| EB / EB-L   | No environmental map, statistics, traps, or confounded outcome learning in selection | Relabelling MPC as $\Omega, D$ dynamics  |
| GX          | Goals represented as bounded, measurable, frozen deficits dominated by safety        | Unconstrained reward shaping             |

## Continuous and typed menus are not the same

For a continuous actuator, an annular menu can literally contain a small metric ball around the current state. That guarantees nearby candidate overlap and supplies choice-theoretic tests. A finite typed menu—wait, restart, step, eat, isolate—does not contain such a ball. It must instead declare its primitive library, reach bound, and designed paired contexts in which the same candidate alternatives genuinely recur. Reusing an action label is not enough if that action projects to a different outcome in each context.

## Tie handling

Exact ties are common in finite menus. Silently choosing the first item in an array keeps execution deterministic but does not make the underlying choice unique. C1-T therefore allows a secondary rank only when it is frozen before outcomes, depends on declared information, and is logged alongside both the raw and post-refinement tie counts. A post-hoc tie rule remains a model failure.

## Goals are still costs

The GX construction places tool, energy, stock, or completion deficits into an enlarged state and measures distance to a reference with those deficits at zero. This is useful because it makes targets bounded, observable, frozen, and subordinate to survival. Mathematically, however, a weighted deficit coordinate is equivalent to a weighted term in the cost. GX should therefore be described as disciplined cost shaping—not as a proof that goals have disappeared from D.

# 3. Current mathematical findings

## 3.1 Persistence has a sharp budget

Let  $V(x)$  be distance to the reference manifold. Suppose the best available candidate can reduce this distance by a factor  $\alpha$ , while the environment can expand it by at most  $L$  and add disturbance  $\epsilon$ . If  $\alpha L < 1$  and the relevant region contains no death state, distance remains bounded and converges geometrically toward a residual band of radius  $\epsilon/(1-\alpha L)$ . If  $\alpha L$  reaches one, explicit counterexamples diverge. Selection cannot rescue a system from an environment that expands faster than the menu can contract.

**INTERPRETATION** Persistence is not assumed as a goal. It follows only when the candidate menu supplies enough genuine corrective authority and the environment does not overwhelm it.

### 3.2 The unconstrained framework is vacuous

With unrestricted menus, any functionally consistent trajectory can be reproduced by making each observed predecessor point to a singleton candidate. A hidden time coordinate can reproduce even arbitrary sequences. The revealed-preference analysis gives the precise antidote: fixed overlapping menus, fixed rankings, and unique effective choices. On countable data, a strict preference cycle is a direct refutation.

### 3.3 Rigid projection models are refutable and partly identifiable

In a Euclidean ball-menu model, metric projection toward  $x^*$  generates a highly restricted family of trajectories. Generic trajectories do not belong to that family. With suitable data, the menu radius and reference point can be recovered, while the scalar transformation of distance is identifiable only up to monotone reparameterization. Under an unknown invertible affine environment, rich one-step map data reveal a constant plateau whose center and radius identify  $x^*$  and the menu radius.

### 3.4 Lookahead forms a help–useless–harm triangle

| Case    | What additional $h$ sees                                | Predicted outcome                                                    |
|---------|---------------------------------------------------------|----------------------------------------------------------------------|
| Helps   | A branch whose own future menu becomes empty            | $h \geq 2$ avoids an $\Omega$ -trap that greedy $h = 1$ enters       |
| Useless | Nothing—the trap exists only in $E$                     | Every $E$ -blind horizon selects the trap; true- $E$ MPC avoids it   |
| Harms   | A cheap internal continuation that $E$ will not realize | A larger $h$ selects a deceptive branch that greedy selection avoids |

For horizon path costs, the worst-case persistence coefficient grows with  $h$ . The exact threshold is  $\alpha_h L < 1$ , where  $\alpha_h = \alpha(1-\alpha^h)/(1-\alpha)$ . More internal planning can therefore require stronger menu-level contraction, even when the environment itself is nonexpansive.

### 3.5 The architecture has observable signatures

- Determinism from the declared extended state, reference, horizon, tie rule, and random seed.
- Sen- $\alpha$  consistency across genuinely overlapping menus.
- Blindness to environmental variance when variance does not enter the self-model.
- Avoidance of menu traps alongside persistent blindness to environment-only traps.
- A shrinking persistence margin—and sometimes failure—as internal horizon increases.
- Geometric range, concurrency, and plateau signatures in rigid Euclidean subclasses.
- An evolutionary shift toward shorter horizons when previously benign environments gain hidden  $E$ -traps.

### 3.6 Important scope correction

The continuous-rationalizability result now uses the  $\sigma$ -compact metrizable domain stated in Levin's cited theorem. Earlier versions claimed arbitrary Polish spaces, which is broader because a Polish space need not be  $\sigma$ -compact. The associated closed-consistency result inherits the corrected domain. No result is claimed for non- $\sigma$ -compact Polish state spaces without an additional representation theorem.

## 4. Evidence from implemented laboratories

The project's empirical work is most useful when read as mechanism validation and falsification practice. Authored simulations can show that the code implements the theory and that the predicted contrasts are realizable. They cannot by themselves establish that unrelated natural or engineered systems use the architecture.

### 4.1 Pool laboratory

A two-dimensional billiards laboratory showed that D cannot invent a missing response family. Remove wall-reflection candidates and balls leave the table; remove collision candidates and they tunnel. It also exposed the brittleness of putting open-ended goals directly into a rich hand-weighted score. Because its D is not the rigid Tier-2 distance form and its horizon is greedy, this laboratory is corroborative design evidence rather than a canonical test.

### 4.2 Stickman laboratory and the value of negative controls

Early stickman demonstrations appeared competent because hidden post-selection servos performed the actual standing and locomotion. Servo-stripped runs then failed—but that negative verdict was also invalid: the plant contained an inconsistent initial geometry, an unstable contact corrector, an erroneous energy meter, and inverted actuation signs. After repair, a framework-faithful  $h = 1$  selector maintained quiet standing for ten seconds, while  $\text{argmax}$ , stay-only, and ragdoll controls fell quickly. Disturbance rejection still did not clearly outperform a stiff frozen controller.

**METHODOLOGICAL LESSON** Trust neither a successful demo with hidden E-side assistance nor a failed selector test conducted on an invalid plant. Validate the plant and run anti-selection and frozen controls beside the same menu.

### 4.3 Minecraft laboratory

The Minecraft implementation uses typed primitives, a survival and goal-deficit reference manifold, weighted-L1 distance, E-blind menu lookahead, and the live game tick as post-selection E. Offline and live contrasts show  $\text{argmin}$  preserving health while  $\text{argmax}$  starves or steps off a ledge. A goal ladder

reaches wood, stone, iron, diamond, and enchanted equipment in the authored arena, then rests at  $D = 0$ . An outer evolutionary loop over metric weights and  $h$  selected the maximum tested horizon in benign worlds, motivating the pre-registered prediction that hidden E-traps should drive evolved horizons downward.

#### 4.4 Non-Minecraft validation benchmarks

Two small benchmarks were added specifically to move validation beyond Minecraft while keeping the expected answers transparent. The service-recovery simulator runs greedy, E-blind lookahead, true-E MPC, frozen, seeded-random, and argmax controllers on identical menus. The behavioral probe combines Sen- $\alpha$ , trap, variance, and horizon observations into an architecture fingerprint.

| Scenario          | Greedy $h = 1$                 | E-blind $h = 2$          | True-E MPC $h = 2$             |
|-------------------|--------------------------------|--------------------------|--------------------------------|
| $\Omega$ -trap    | Dead: local patch              | Recovered: staged repair | Recovered: staged repair       |
| E-trap            | Dead: restart outage           | Dead: restart outage     | Recovered: shed load           |
| Horizon fragility | Operational: conservative hold | Dead: deceptive restart  | Operational: conservative hold |

The classifier labels the included reference subjects as greedy E-blind, lookahead E-blind, and environment-anticipating. All 180 automated tests pass as of 12 July 2026. These benchmarks validate the implementation and experimental protocol; their subjects were authored to instantiate the categories, so they are not independent confirmation.

## 5. Assessment: strengths and limitations

### 5.1 Principal strengths

- **Vacuity is treated as a first-class problem.** The framework explicitly identifies how unrestricted menus, metrics, states, references, and ties can rationalize anything.
- **Failures are diagnostic.** A bad result can be traced to missing primitives, misranking, insufficient corrective authority, a hostile environment, a tie failure, or a contaminated factorization.
- **The theory risks refutation.** E-trap blindness, variance blindness, horizon fragility, and revealed-choice cycles can all fail against data.
- **The empirical discipline is unusually strong.** Argmax, frozen, random, servo-stripped, plant-validation, and true-E MPC controls prevent attractive behavior from being credited to the wrong component.
- **The architecture is auditable.** Typed primitives and declared metrics can be inspected in ways that black-box policies often cannot.

## 5.2 Principal limitations

- **E-blindness is deliberately restrictive.** Many useful and safe agents must anticipate other vehicles, patients, markets, adversaries, or equipment response.
- **The  $\Omega/E$  boundary is modeler-sensitive.** A primitive can quietly contain environmental prediction, and ordinary outcome errors cannot be cleanly attributed to the actuator without intervention.
- **Typed menus need designed overlap.** Recurrent action labels do not automatically produce the shared alternatives required by revealed-preference tests.
- **GX remains cost shaping.** Its value lies in measurable and bounded constraints, not in a mathematical escape from objectives.
- **Identifiability is class-dependent.** Strong results use rich map data and rigid ball-menu or affine assumptions; single trajectories and unknown model classes remain open.
- **Current evidence is internally authored.** No independent physical, biological, organizational, or production-software replication yet establishes external validity.
- **Much of the persistence machinery is inherited.** The contraction theorem is standard ISS-style reasoning; novelty is concentrated in the factorization, restrictions, and EB-driven signatures.

## 5.3 Present scientific status

**ASSESSMENT** Promising as a constrained, testable architecture for homeostatic choice and as a baseline for detecting environmental anticipation. Not yet established as a general theory of causality or as a superior control method.

# 6. Application directions

The best applications are those in which the menu, reference manifold, and information boundary can be declared in advance—and where the framework’s distinctive failures can be deliberately tested. The objective is not to find domains that can be described by the equation; almost any domain can. The objective is to find domains that can discriminate this architecture from alternatives.

## 6.1 Behavioral science and comparative cognition

Model internal physiological state as  $x$ , experimentally offered actions as  $\Omega$ , homeostatic setpoints as  $x^*$ , and hidden consequences as  $E$ . Paired menu contractions, variance manipulations, and separate  $\Omega$ - and  $E$ -traps could distinguish greedy homeostasis, self-model lookahead, and world-model planning. This is currently the strongest scientific application because the signatures can be presented as controlled interventions.

## 6.2 Black-box agent classification

Expose an unknown agent to controlled menus and record Sen-α consistency, variance sensitivity, trap avoidance, and horizon behavior. The resulting fingerprint can classify whether the policy is inconsistent, greedy E-blind, lookahead E-blind, or environmentally anticipating. This is useful for auditing independently built agents, provided their actual information access is controlled.

## 6.3 Self-stabilizing software services

Represent service health, load, queue depth, and configuration as  $x$ ; restart, isolate, shed, replicate, or roll back as  $\Omega$ ; a healthy service envelope as  $x^*$ ; and workload or dependency failures as  $E$ . The architecture may be valuable for auditable local recovery and as a contrast to predictive orchestration. Dijkstra-style self-stabilization is an essential baseline rather than a competing label to ignore.

## 6.4 Certified fallback and runtime assurance

Use verified actuator macros as  $\Omega$  and a safe operating envelope as  $x^*$ . A deliberately conservative E-blind layer could act when a more capable planner is unavailable or untrusted. Its blindness must be explicit: it should not replace primary control in settings where avoiding external hazards requires prediction.

## 6.5 Cellular regulation and synthetic biology

Use metabolic or regulatory state as  $x$ , feasible regulatory responses as  $\Omega$ , a viable region as  $x^*$ , and temperature, toxin, or nutrient shocks as  $E$ . The framework functions here as a falsifiable non-anticipating null model. Perturbations can ask whether deeper endogenous planning helps only menu-visible traps while leaving hidden environmental contingencies unlearned.

## 6.6 Evolutionary ecology and artificial life

Allow an outer evolutionary loop to tune metric weights and horizon while keeping within-lifetime selection E-blind. The key pre-registered test is population-level horizon shift: benign menu structure should reward longer internal lookahead, whereas hidden environment traps should make long horizons costly and select for shorter ones.

## 6.7 Interpretable simulated agents and swarms

Typed behavior, movement, and communication primitives make  $\Omega$  explicit; local energy, spacing, or connectivity define  $x^*$ ; the world tick supplies  $E$ . The framework can produce debuggable NPCs or collective recovery experiments in which missing primitives, misranking, and external disruption are separately observable.

## Where strict E-blindness is a poor primary fit

Autonomous driving, financial trading, medical dosing, adversarial cybersecurity, and most high-performance robotics depend on anticipating environmental response. In those domains, strict (Ω, D) Dynamics is better used as a controlled baseline, a narrow independently certified fallback, or a diagnostic hypothesis—not as the main decision system.

## 7. Validation standard and research agenda

### 7.1 Minimum standard for any new application

1. Freeze and publish the extended state, reference manifold, metric class, horizon, primitive library, and tie rule before observing outcomes.
2. Declare R1-C or R1-H. For typed menus, identify the paired contexts containing genuinely identical candidates for any revealed-preference claim.
3. Run the Ω/D/E factorization audit. Register every information channel and every source used to project a primitive.
4. Use identical menus across E-blind  $h = 1$ , E-blind  $h \geq 2$ , true-E MPC, frozen, seeded-random, and argmax controls.
5. Log raw ties, effective ties, seeds, state keys, selections, realized outcomes, terminal events, and trajectory hashes.
6. Test predicted failures as well as successes: variance blindness, E-trap blindness, and horizon fragility are evidence-bearing outcomes.
7. Separate implementation validation from external evidence. An authored reference subject proves that a protocol works, not that an independent system uses the architecture.

### 7.2 Priority research program

| Priority | Study                                                                  | Decisive result                                                            |
|----------|------------------------------------------------------------------------|----------------------------------------------------------------------------|
| 1        | Run the behavioral probe against independently built agents            | Correct classification without inspecting implementation                   |
| 2        | Add hidden E-traps to the evolutionary world                           | A reproducible downward shift in selected $h$                              |
| 3        | Deploy the service benchmark against real recovery policies            | Trap and variance signatures separate reactive from predictive controllers |
| 4        | Compare E-blind and feedback controllers on a validated physical plant | A mapped boundary where self-model selection succeeds or fails             |
| 5        | Extend identification beyond rigid ball menus                          | Recover or refute Ω, D, and E from limited stochastic trajectories         |
| 6        | Develop a finite operational approximation to closed consistency       | A usable topological revealed-choice test for empirical data               |

### 7.3 Decision rule for continuing the program

The framework gains credibility if independently built systems reproduce the joint signature pattern under pre-registered information boundaries. It should be narrowed or rejected if environment-trap avoidance, variance sensitivity, or horizon effects persist without the information needed to produce them—or if the factorization cannot be specified without moving useful world prediction into  $\Omega$  or the self-model.

## 8. Conclusion

The most valuable achievement of (Ω, D) Dynamics is not the argmin equation. It is the discipline imposed around that equation: menus cannot be invented after the fact, metrics cannot absorb all dynamics, references cannot chase the data, ties cannot hide arbitrary behavior, and environmental prediction cannot be smuggled into a theory defined by its absence.

Within those constraints, the framework yields a coherent body of results. Persistence requires a real contraction budget. Unrestricted choice models are vacuous. Rigid subclasses can be refuted and partially identified. E-blind lookahead has a distinctive triangle: it can help with menu structure, remain useless against hidden environmental structure, and become harmful when internal continuations are trusted too much.

The implemented laboratories now demonstrate these mechanisms with strong internal controls, including repaired negative results, anti-selection contrasts, true-E MPC comparisons, an executable factorization audit, and non-Minecraft benchmarks. The next threshold is external: independently constructed agents and real systems must be tested against the same pre-registered signatures.

**CURRENT VERDICT** A rigorous and promising experimental architecture with clear limits. Its future depends less on adding demonstrations of competence than on surviving independent attempts to falsify its information boundary and behavioral signatures.

## Appendix A. Compact notation glossary

| Term                     | Definition                                                                                                       |
|--------------------------|------------------------------------------------------------------------------------------------------------------|
| Extended state           | All declared system information carried forward; bounded history must be fingerprinted into the state key.       |
| Candidate menu $\Omega$  | Set of mechanism-expressible candidate primitives or next states available now.                                  |
| Reference manifold $x^*$ | Closed set of acceptable or completed states; frozen before data or evolving autonomously under a separate rule. |
| Distance $V$             | Canonical Lyapunov quantity $V(x)$ = distance from $x$ to $x^*$ .                                                |
| Discontinuity $D$        | Declared candidate ranking; in rigid Tier 2 it is a strictly increasing transform of distance to $x^*$ .         |
| Horizon $h$              | Number of internally simulated menu steps used in path cost.                                                     |
| Environment $E$          | Post-selection map from selected candidate to realized next state; may be stochastic.                            |
| E-blindness              | Selection uses no map, statistics, or traps of $E$ and performs lookahead with internal $E$ = identity.          |
| $\Omega$ -trap           | Candidate whose own future menu dies; visible to sufficiently deep E-blind lookahead.                            |
| E-trap                   | Candidate made terminal or harmful only by environmental response; invisible to every E-blind horizon.           |
| C1-T                     | Pre-registered secondary ranking that makes a tied finite menu effectively unique while preserving raw tie logs. |
| GX                       | Bounded, measurable, frozen deficit coordinates added to the reference metric and dominated by survival.         |

## Appendix B. Result status at a glance

| Area                         | Current status                                                  | Main residue                                                       |
|------------------------------|-----------------------------------------------------------------|--------------------------------------------------------------------|
| Well-posedness               | Standard sufficient conditions recorded; proof sketch           | Typed discontinuities and infinite-dimensional information costs   |
| Persistence $h = 1$          | Proved with sharp threshold and death-free condition            | Application-specific proof that the real menu supplies contraction |
| Random menus                 | Proved for bounded iid menu noise                               | State-dependent, adversarial, and learned menus                    |
| Drifting reference           | Tracking band proved for slow autonomous drift                  | Fast drift, resonance, and skew-product theory                     |
| Vacuity                      | Universal rationalization loopholes proved                      | Operational audits in complex implementations                      |
| Tier-1 choice                | Acyclicity characterization on countable data                   | Finite operational test for topological closed consistency         |
| Continuous rationalizability | Proved on $\sigma$ -compact metrizable state spaces via Levin   | Non- $\sigma$ -compact Polish spaces                               |
| Rigid identifiability        | Reference, radius, and affine $E$ identified from rich map data | Single trajectories, stochastic $E$ , unknown menu shape           |
| Horizon persistence          | Exact $\alpha_h L$ threshold proved                             | Empirical estimation of $\alpha$ and $L$ in real systems           |
| MPC separation               | Range, geometry, variance, and trap separations proved          | Independent black-box replications                                 |
| Empirical evidence           | 180 tests plus multiple authored laboratories                   | External validity                                                  |

## Selected references and project materials

1. Aubin, J.-P. (2009). *Viability Theory*. Birkhäuser. <https://doi.org/10.1007/978-0-8176-4910-4>
2. Dijkstra, E. W. (1974). Self-stabilizing systems in spite of distributed control. *Communications of the ACM*, 17(11), 643–644.
3. Keramati, M., & Gutkin, B. (2014). Homeostatic reinforcement learning for integrating reward collection and physiological stability. *eLife*, 3, e04811. <https://doi.org/10.7554/eLife.04811>
4. Levin, V. L. (1983). Continuous utility theorem for closed preorders on a metrizable  $\sigma$ -compact space. *Doklady Akademii Nauk SSSR*, 273(4), 800–804. English translation: *Soviet Mathematics Doklady*, 28, 715–718.
5. Minguzzi, E. (2011). Topological conditions for the representation of preorders by continuous utilities. *arXiv:1108.5123*.
6. Sen, A. K. (1971). Choice functions and revealed preference. *The Review of Economic Studies*, 38(3), 307–317. <https://doi.org/10.2307/2296384>
7. *Omega\_D\_Dynamics\_Framework.md*. Consolidated mathematical framework, axioms, theorems, counterexamples, open problems, and laboratory record. Current project revision: 12 July 2026.
8. *minecraft-omega-d-lab/*. Executable selector, factorization audit, Minecraft laboratory, service-recovery benchmark, behavioral trap classifier, and automated test suite.

## Document note

This whitepaper is an accessible synthesis of the project's current claims. Formal definitions, assumptions, constructions, and proof details remain authoritative in *Omega\_D\_Dynamics\_Framework.md*. Implemented behavior and test counts are authoritative in *minecraft-omega-d-lab/* at the revision date above.