

ASSESSMENT • APPLICATION ROADMAP

(Ω, D) Dynamics

Assessment and Application Roadmap

Original roadmap retained • eight applications expanded in full

PURPOSE Preserve the current assessment and application table, then make every proposed mapping concrete enough to design an experiment or engineering prototype.

Current revision: 12 July 2026

Current findings as of 12 July 2026

Overall assessment

The framework is a promising, unusually self-critical architecture for homeostatic action selection under environmental uncertainty. Its defensible claim is not a new theory of causality; it is a constrained decision rule. Ω specifies currently available action primitives. D ranks their internally projected states by distance from a reference manifold x^* . E represents what the external world subsequently does. The defining restriction is E -blindness: selection cannot model or learn external dynamics.

This restriction produces the framework's most original contribution: testable differences from MPC and model-based reinforcement learning, especially Ω -trap avoidance combined with persistent E -trap blindness. The unconstrained equation is only notation; empirical content comes from the axioms, overlap conditions, uniqueness discipline, and falsifiable signatures.

Its mathematical ingredients individually resemble revealed preference, viability theory, projection control, and input-to-state stability. The likely novelty is their architectural combination and experimental separation rather than a wholly new branch of dynamics. The framework becomes distinct only when EB is maintained instead of allowing environmental outcome learning.

Strengths

- The vacuity analysis is excellent: $R0$ – $R4$, effective uniqueness, and the control conditions prevent the equation from explaining arbitrary trajectories after the fact.
- The framework makes risky predictions rather than merely demonstrating attractive simulations.
- Negative controls— $\arg\max$, frozen, random, servo removal, plant validation, and true- E MPC—are unusually strong.
- The horizon triangle is valuable: lookahead can help, fail completely, or actively harm.
- The separation between selection failure, missing primitives, and environmental disturbance makes failures auditable.

Issues requiring attention

- **Theorem 10's domain must remain narrow.** The corrected statement uses the σ -compact metrizable domain supported by Levin's cited theorem; arbitrary Polish spaces are not claimed.
- **Continuous and typed menus differ.** Stay plus a bounded step is $R1$ -H, not proof that a discrete menu contains a metric ball; typed revealed-choice tests require designed overlap.
- **The Ω/E boundary can be moved by the modeller.** A primitive may conceal environmental prediction, while ordinary outcome errors confound actuator and environment effects.
- **GX is disciplined cost shaping.** Deficit coordinates remain weighted cost terms; legitimacy comes from measurability, pre-registration, boundedness, and safety dominance.

- **Ties require an explicit policy.** C1-T permits a frozen observable secondary rank, with raw and effective ties logged; silent array order is not uniqueness.
- **Current evidence establishes feasibility in authored simulations,** not external validity in physical, biological, organizational, or production systems.

Best application directions

Application	Framework mapping	Why it fits
Behavioral experiments	x: organism's internal state; Ω: available actions; x*: physiological setpoint; E: hidden environmental consequences	Probably the strongest scientific use. Designed Ω-traps, E-traps and variance manipulations could distinguish homeostatic selection from model-based planning.
Black-box agent classification	Present an unknown agent with controlled overlapping menus and traps	Sen-α violations, noise sensitivity and trap behavior provide a behavioral fingerprint for whether an agent uses a world model.
Self-stabilizing software services	x: health/load/configuration; Ω: restart, isolate, shed load, replicate; x*: healthy service envelope; E: workload and failures	Discrete, auditable primitives and recovery toward a safe manifold are a natural match.
Certified fallback control	x: vehicle or machine telemetry; Ω: verified actuator macros; x*: safe operating envelope; E: unmodelled disturbances	Appropriate as a conservative fallback or runtime-assurance layer. EB should not replace a capable world-model controller.
Synthetic biology and cellular regulation	x: metabolic/regulatory state; Ω: available regulatory responses; x*: viable homeostatic region; E: toxins, temperature or nutrient shocks	Useful as a falsifiable null model of regulation without environmental anticipation.
Evolutionary ecology and artificial life	Evolution searches over D parameters and h; organisms remain E-blind within a lifetime	Directly exercises the evolved-horizon prediction under benign environments and hidden E-traps.
Interpretable NPCs and simulated societies	Typed behavior primitives, safety/need manifold, world tick as E	Produces debuggable behavior where failures can be traced to an incomplete menu, bad ranking, or external disruption.
Fault-tolerant swarms	x: local energy/spacing/connectivity; Ω: local motion and communication primitives; x*: viable formation manifold	Tests whether local self-models can generate collective recovery without predicting the global environment.

Poor primary applications include autonomous driving, financial trading, medical dosing, adversarial cybersecurity, and other settings where anticipating and learning environmental response is essential. There, strict EB is more useful as a controlled baseline or independently certified emergency layer than as the main controller.

Recommended validation program

1. Keep the corrected topological domain behind Theorems 10 and 14 explicit.
2. Use R1-C for continuous neighbourhood menus and R1-H for bounded typed action graphs with designed overlap.
3. Run the executable factorization audit for every implementation and freeze the manifest before outcomes.

4. Compare identical menus under E-blind $h = 1$, E-blind $h \geq 2$, true-E MPC, frozen, seeded-random, and argmax selection.
5. Treat the framework as supported only if it reproduces distinctive failures as well as successes: Sen- α consistency, variance blindness, Ω -trap avoidance, E-trap blindness, and horizon fragility.

Assumptions

- The framework is evaluated as a scientific and engineering architecture, not as an established reduction of causality.
- Strict E-blindness remains non-negotiable; weakening it intentionally turns the system into MPC or adaptive control.
- Applications are prioritized by how cleanly they expose falsifiable signatures, not merely by whether the equation can describe them.

EXPANDED APPLICATION PROFILES

Eight concrete mappings and decisive tests

Each profile expands the original Framework mapping and Why it fits entries without replacing them.

APPLICATION PROFILE 01

Behavioral experiments

FIT VERDICT The strongest scientific application because the experimenter can control menus, information access, and hidden consequences independently.

Framework mapping

x — extended state A measurable internal condition such as satiety, energy, temperature, pain avoidance, arousal, or a safe laboratory proxy for need. The state may include a bounded choice history, but not a hidden clock or learned outcome table added after seeing the data.

Ω(x) — offered actions The exact set of actions presented on a trial: consume, wait, explore, rest, approach, or choose among symbolic alternatives. Paired trials must reuse genuinely identical alternatives so that menu contraction and Sen-α tests are meaningful.

x* — reference manifold A pre-registered viable band rather than a perfect point—for example, an acceptable range of energy or temperature. In human studies this can be a virtual resource state; no harmful physiological deprivation is required.

D — discontinuity A fixed ranking of candidate internal states by distance from the reference band. A transparent weighted-L1 or ordinal model is preferable; predicted food delivery, punishment probability, or environmental value must not be inserted into D.

E — hidden consequence What occurs after choice: whether a resource is delivered, delayed, blocked, made variable, or paired with an unseen trap. E must be manipulated separately from the offered menu.

h — internal lookahead Depth through the subject's own action/menu structure. It can be manipulated by nesting visible submenus or inferred from whether the subject avoids menu-visible dead ends.

Why it fits

Behavioral laboratories can create the exact contrasts that give the framework empirical content. An Ω-trap is visible in the structure of future offered choices; an E-trap is attached only after a superficially attractive choice. Variance can be altered without changing candidate distance, and menus can be contracted while preserving alternatives. These interventions separate greedy homeostatic choice, E-blind self-model lookahead, and genuine environmental anticipation.

The domain also makes negative evidence interpretable. A subject that avoids an E-trap may possess a world model, may have learned a cue that leaks E, or may carry an omitted state coordinate. Each possibility can be tested by changing cues and transfer contexts rather than explaining the result after the fact.

A decisive first study

Use a virtual-resource task with two branches. One branch is immediately closer to the resource setpoint but ends in a visibly empty future menu; the other is slightly worse now but remains viable. In a separate block, keep both internal branches viable but make the first branch fail only through an unannounced post-choice contingency. Cross h-depth, outcome variance, and paired menu contractions while logging complete information access.

SUPPORT WOULD LOOK LIKE Sen-α consistency on overlapping menus; avoidance of the Ω-trap with sufficient depth; continued selection of an indistinguishable E-trap; and little or no response to outcome variance that was never available to the selector.

KEY BOUNDARY Real organisms learn environmental regularities and have partially observed internal states. Any cue correlated with E, uncontrolled learning across trials, or post-hoc enlargement of x can erase the claimed distinction.

APPLICATION PROFILE 02

Black-box agent classification

FIT VERDICT A high-value diagnostic use: the framework becomes a behavioral fingerprint rather than a claim about how the unknown agent was implemented.

Framework mapping

x — presented information state Everything the test harness exposes to the agent: observation, inventory, goal variables, bounded transcript, and seed. Hidden scratchpads, memory, retrieval, and tool outputs must be controlled or treated as unobserved state.

$\Omega(x)$ — controlled interface A harness-defined set of actions or tool calls. The same action objects must recur across paired contexts; merely reusing labels while changing arguments or predicted outcomes does not establish overlap.

x^* — task reference A declared success or viability manifold shared across agents—for example, service restored, energy within band, or puzzle state nonterminal. Keeping x^* fixed makes different policies comparable.

D — latent ranking Usually unobserved. Classification relies first on ordinal revealed choices and trap responses; a parametric D is fitted only after the Tier-1 consistency tests pass.

E — harness response Tool results, hidden transition rules, stochastic failures, other agents, and delayed consequences supplied after the action is selected.

h — inferred depth Estimated with a ladder of menu-visible traps requiring one, two, or more internal action steps to detect. Depth is behavioral, not the agent's advertised planning setting.

Why it fits

The framework supplies a compact battery of architecture-level tests: determinism or seeded reproducibility, Sen- α consistency, variance sensitivity, Ω -trap avoidance, E-trap avoidance, and horizon fragility. The joint pattern is more informative than task score alone. A policy can be highly competent yet classify as environment-anticipating, or perform poorly while still expressing the E-blind signature cleanly.

This application is especially useful for agents supplied by another team or vendor, where code inspection is unavailable or insufficient. The result should be reported as a behavioral class under the tested interface—not proof of a unique internal algorithm—because different mechanisms can remain observationally equivalent on a finite battery.

A decisive first study

Create a sealed benchmark with repeated paired menus, a variance reversal, a depth ladder of Ω -traps, and matched E-traps. Run each subject across held-out seeds with frozen prompts and tool schemas. Include known reference policies—greedy E-blind, h-lookahead E-blind, true-E MPC, seeded-random, and inconsistent—to calibrate the classifier before testing unknown agents.

SUPPORT WOULD LOOK LIKE Stable classification on held-out trials, correct recovery of the reference policies, and transfer of the fingerprint when surface labels and irrelevant presentation details change.

KEY BOUNDARY A hidden memory channel, leaked simulator metadata, stochastic policy temperature, or an agent that changes strategy during the battery can mimic or destroy signatures. Classification must include uncertainty and an 'underdetermined' outcome.

APPLICATION PROFILE 03

Self-stabilizing software services

FIT VERDICT A strong engineering testbed and plausible operational layer because service states and recovery primitives are discrete, observable, replayable, and auditable.

Framework mapping

x — service state Own telemetry such as healthy replicas, queue depth, error rate, resource headroom, configuration version, circuit-breaker state, and a bounded local event history. Workload forecasts and dependency failure models remain outside strict x.

Ω(x) — recovery primitives Wait, shed load, isolate a dependency, restart one replica, roll back, replicate, drain traffic, or enter read-only mode. Feasibility comes from current permissions and configuration, not a prediction that the operation will succeed.

x* — healthy envelope A set of states satisfying service-level and safety bounds: minimum capacity, bounded latency and errors, valid configuration, and no cascading restart condition. It should allow several acceptable operating modes.

D — operational deficit A frozen metric over SLO shortfall, capacity loss, queue pressure, configuration risk, and terminal penalties. Deficits are capped and safety terms dominate convenience or throughput goals.

E — operating environment Incoming workload, dependency behavior, hardware failure, orchestration races, network partitions, and operator interventions after a recovery action is issued.

h — action-graph lookahead Planning over the declared recovery workflow—restart then validate, shed then drain, isolate then restore—without simulating workload or dependency response.

Why it fits

Software recovery offers unusually clean instrumentation. Candidate menus can be logged exactly, incidents replayed deterministically, and the same primitives compared under E-blind, true-E MPC, frozen, random, and argmax controllers. Menu traps are natural: a locally attractive repair can remove all subsequent recovery options. Environment traps are equally natural: a restart that appears restorative can trigger a dependency outage or thundering herd only after execution.

The framework also creates an audit vocabulary. Operators can distinguish missing recovery authority from bad ranking, insufficient contraction, contaminated prediction, or an external shock. This is useful even when a predictive orchestrator ultimately wins, because the E-blind policy supplies a transparent baseline and a narrow fallback design.

A decisive first study

Deploy the existing service-recovery benchmark into a containerized service with controlled fault injection. Run restart storms, dependency partitions, burst load, stale configuration, and misleading healthy telemetry. Keep menus identical across policies and pre-register which channels belong to self-state versus E. Measure recovery time, terminal failures, D trajectory, tie counts, and action-path reproducibility.

SUPPORT WOULD LOOK LIKE E-blind $h > 1$ should escape workflow-visible dead ends but remain vulnerable to dependency traps that a true-E controller avoids. Failures should be attributable to declared menu, ranking, or E categories rather than hidden automation.

KEY BOUNDARY Production recovery often depends on learned failure probabilities and workload forecasts. Strict E-blindness may therefore be unsuitable as the main orchestrator; it is most credible as a baseline, local self-stabilizer, or certified degraded mode.

APPLICATION PROFILE 04

Certified fallback control

FIT VERDICT Plausible as a narrow runtime-assurance layer, but only when its action set and safe envelope are independently verified; the framework itself does not provide certification.

Framework mapping

x — trusted telemetry A minimal, validated state estimate: energy, temperature, attitude, joint limits, pressure, battery, actuator availability, and current operating mode. Untrusted perception and speculative object tracks should not silently enter the fallback state.

$\Omega(x)$ — verified macros Pre-analysed actions such as hold, stop, reduce power, open a relief path, return to neutral, isolate a subsystem, or transition to a safe mode. Each primitive must have explicit preconditions and bounded authority.

x^* — safe operating envelope A robust set of acceptable states rather than a mission target. It may include several low-performance modes so that safety does not require returning to a single nominal point.

D — safety deficit Distance to envelope margins, with terminal constraint violations dominating comfort, productivity, or mission progress. D is fixed during the incident and must not depend on predicted external benefit.

E — residual disturbance Unknown load, wind, contact, component degradation, sensor error, or primary-controller residue acting after the macro is selected.

h — macro-sequence depth Lookahead through internally declared safe-mode transitions. It can detect that one macro leaves no legal successor, but it does not predict how the external disturbance responds.

Why it fits

Fallback systems value simplicity, deterministic behavior, bounded authority, and auditability more than optimal performance. Those properties align with typed menus and a frozen reference envelope. The framework's contraction result also gives a useful engineering question: does at least one admissible macro reduce distance to safety strongly enough to overcome the worst disturbance bound?

Its limitation is equally useful. E-blindness forbids claiming that the fallback can avoid hazards whose consequences are visible only through an environmental model. The architecture should sit behind a clear handover rule and beside independent monitors, not be presented as a universal safety controller.

A decisive first study

On a validated plant or high-fidelity hardware-in-the-loop rig, trigger loss of the primary controller across a matrix of initial states and bounded disturbances. Compare a strict E-blind fallback, a conventional verified safety controller, and a predictive controller using the same macro library. Verify menu completeness, contraction margins, terminal avoidance, deterministic replay, and the exact boundary where hidden disturbances defeat E-blind selection.

SUPPORT WOULD LOOK LIKE The fallback maintains the declared invariant envelope whenever the measured contraction budget predicts it should, fails at the predicted threshold, and never relies on undeclared E-side assistance.

KEY BOUNDARY Certification requires hazard analysis, verified software and sensing, fault containment, and domain-specific evidence. An argmin rule and successful simulation do not certify a system, and unmodelled moving hazards can make E-blind selection unsafe.

APPLICATION PROFILE 05

Synthetic biology and cellular regulation

FIT VERDICT A promising falsifiable null model of non-anticipating regulation, especially in controlled perturbation experiments—not yet a claim that cells literally solve the selector equation.

Framework mapping

x — cellular state Measured metabolite levels, membrane potential, energy charge, stress markers, protein abundance, transport capacity, and bounded regulatory history at a chosen timescale.

$\Omega(x)$ — regulatory responses Mechanistically available changes such as transporter activation, enzyme allocation, stress-response expression, dormancy entry, repair, or secretion. The menu is constrained by current molecular machinery and resource availability.

x^* — viability manifold A region of compatible pH, redox balance, energy, osmotic pressure, and damage rather than a single biochemical target. The manifold must be specified independently of the observed trajectory.

D — homeostatic discontinuity A declared distance over capped physiological deficits. It should rank internal states without incorporating predicted nutrient payoff, toxin dynamics, or environmental transition probabilities.

E — external stressor Temperature, toxin concentration, nutrient arrival, host response, neighbouring organisms, flow, or stochastic molecular events that act after a regulatory response begins.

h — endogenous cascade depth The depth of internally available regulatory cascades and downstream self-state transitions, not foresight about future nutrient or toxin conditions.

Why it fits

Homeostasis is native to the domain, and microfluidic or chemostat experiments can separate present internal state from controlled external perturbation. Menu-visible dead ends can be created by blocking regulatory branches, while environment-only traps can be introduced after the same initial response. This makes the help-useless-harm horizon triangle experimentally interpretable at the level of a model, even if the biological implementation is distributed and stochastic.

The framework is most valuable here as a null hypothesis: regulation is ranked by internal deficit and endogenous options without an explicit world model. Evidence for anticipation—response to a predictive cue with no present internal deficit—would challenge strict EB or require the cue and learned environmental state to be added explicitly to x .

A decisive first study

Use a clonal population in a controlled flow system with fluorescent reporters for internal deficits and regulatory activation. Present matched perturbations where one response branch has an endogenous downstream dead end and another where failure is introduced only by a later external pulse. Manipulate external variance without changing current internal state and compare response distributions across knockouts that alter cascade depth.

SUPPORT WOULD LOOK LIKE Responses track current internal distance and endogenous branch structure, deeper cascades avoid menu-visible dead ends, and hidden external variance or traps do not change initial selection unless a corresponding cue enters measured state.

KEY BOUNDARY Defining discrete candidate actions in a continuous biochemical network is difficult; molecular noise, population heterogeneity, multiple timescales, and evolutionary adaptation can blur Ω and E. The model must not anthropomorphize regulation or infer a selector from curve fitting alone.

APPLICATION PROFILE 06

Evolutionary ecology and artificial life

FIT VERDICT An excellent domain for the evolved-horizon prediction because lifetime policy and between-generation selection can be separated and manipulated directly.

Framework mapping

x — organism state Energy, damage, inventory, position, developmental mode, and bounded lifetime memory. Population or genotype identifiers cannot be used as arbitrary hidden clocks.

Ω(x) — lifetime behaviors Forage, rest, flee, reproduce, repair, communicate, build, or move, gated by current morphology and resources. Primitive definitions remain fixed within each experimental lineage.

x* — inherited viability target A frozen survival and completion manifold encoded before each lifetime. Reproductive or developmental deficits may be included only as bounded, declared GX coordinates subordinate to survival.

D — heritable steering metric Weights and perhaps horizon h vary between genomes while the canonical evaluation score remains fixed across the population. This steering/scoring split prevents organisms from winning by redefining their ruler.

E — ecology Resource renewal, predators, climate, other organisms, terrain changes, and hidden traps that respond after action selection.

h — heritable internal depth A genome-level parameter controlling menu-only lookahead. It can evolve across generations but is not retuned from outcomes within a lifetime.

Why it fits

The framework makes a directional population prediction rather than merely fitting individuals. In benign environments, longer internal lookahead can exploit visible menu structure and should be selected when its computational cost is low. When hidden E-traps are introduced, the same depth cannot anticipate them and can actively prefer deceptive continuations, creating selection pressure toward shorter horizons or different primitives.

Artificial-life platforms allow identical worlds, seeds, genomes, and controls to be replayed. They also expose a subtlety: evolution itself learns environmental regularities. EB can remain a within-lifetime architectural claim, but the genome may embody past E. Results must therefore distinguish inherited structure from online world modelling rather than claiming complete environmental ignorance.

A decisive first study

Evolve populations over metric weights and h in matched worlds. Begin with menu-visible resource chains and no hidden traps, then introduce post-choice hazards whose cues are absent from x and Ω . Hold the canonical fitness score, primitive library, mutation budget, and compute budget fixed. Include frozen, argmax, random, and true-E lineages plus transfers to novel trap locations.

SUPPORT WOULD LOOK LIKE Longer h dominates in benign worlds, then declines reproducibly after hidden E-traps appear; transfer tests show that the shift follows architecture-level trap pressure rather than memorized locations or an unexercised gene.

KEY BOUNDARY If genomes, development, or population memory encode the trap distribution, the system is environmentally adapted even when each lifetime is E-blind. Computational cost, mutation linkage, and changing Ω can also explain horizon shifts and must be controlled.

APPLICATION PROFILE 07

Interpretable NPCs and simulated societies

FIT VERDICT A strong engineering application for debuggable behavior and controlled experiments, provided authored primitives do not conceal a world model.

Framework mapping

x — local character state Health, hunger, fatigue, possessions, social obligations, role, location, and bounded memory. Character-specific traits should be explicit state or frozen metric parameters, not hidden special cases.

$\Omega(x)$ — typed behavior Wait, eat, rest, trade, craft, move, defend, assist, communicate, or pursue a task, gated by current affordances. Complex pathfinding or scripted success inside a primitive must be declared as assistance rather than credited to selection.

x^* — character reference A viable and role-consistent manifold: alive, adequately supplied, socially connected, and at rest when declared commitments are complete. Different archetypes can use different frozen reference coordinates.

D — interpretable needs metric A visible weighted ranking of safety, resource, role, and relationship deficits. Safety anchors dominate optional goals; completion returns the character to a stable wait state rather than perpetual reward seeking.

E — simulated world Other agents, combat, market prices, physics, resource respawn, weather, and player actions applied after the NPC chooses.

h — self-action lookahead Depth through the NPC's own craft, travel, or social-action menus with the world response held fixed. This can preserve future options without predicting other actors.

Why it fits

Designers gain an auditable failure taxonomy. If an NPC cannot cross a river, the primitive may be absent; if it walks into danger visible in D, the ranking is wrong; if a mob unexpectedly intercepts it, E caused the failure. The same menu can be run under argmin, argmax, frozen, random, and predictive controls to demonstrate that selection—not a hidden servo—produced the behavior.

The architecture also supports believable rest, bounded goals, and different character priorities without opaque reward accumulation. In simulated societies, local reference manifolds can generate trade, assistance, and conflict as consequences of deficits, while retaining reproducible decision logs for narrative debugging.

A decisive first study

Build a small settlement scenario with food, shelter, craft chains, social requests, and two trap types. One task branch visibly exhausts future options; another is sabotaged only by an unseen world event. Run identical populations under the full control set, audit every primitive projection, and have designers diagnose failures from logs before seeing the source code.

SUPPORT WOULD LOOK LIKE Designers correctly attribute failures, argmin outperforms anti-selection controls on the frozen score, h avoids menu dead ends but not world-only traps, and completed characters settle into coherent rest states.

KEY BOUNDARY A 'seek food' or 'go home' primitive often contains pathfinding, threat prediction, and guaranteed actuation. If those systems do the competence, the argmin layer is decorative. Dynamic social learning also becomes adaptive control unless explicitly separated.

APPLICATION PROFILE 08

Fault-tolerant swarms

FIT VERDICT A plausible research application for local, auditable recovery in decentralized systems, with a particularly demanding Ω/E boundary because other agents are both observed neighbours and part of the environment.

Framework mapping

x — local agent state Own energy, damage, pose, task load, communication status, bounded neighbour summary, and local connectivity estimate. Global state is absent unless a communication primitive delivers it.

Ω(x) — local primitives Hold, move, re-space, relay, elect, isolate, recharge, hand off a task, or change formation role. Availability follows current actuator, communication, and energy constraints.

x* — local viability manifold Safe spacing, energy reserve, communication redundancy, task coverage, and formation/connectivity bounds expressed locally. Global objectives must decompose into measurable local deficits or remain outside the strict model.

D — local continuity metric A fixed ranking of candidate self-states by safety, energy, isolation, and coverage deficits. It does not predict how neighbours will move in response.

E — coupled world Other agents' simultaneous actions, packet loss, latency, obstacles, wind, adversaries, and hardware faults after the local candidate is selected.

h — local menu depth Lookahead through the agent's own future roles and moves. It can avoid choosing a role with no legal successor, but cannot anticipate the joint action of the swarm.

Why it fits

Swarm resilience is often achieved through local rules rather than a centralized world model, making the framework a plausible descriptive and engineering baseline. Typed local primitives, bounded state, and explicit reference deficits support replayable failure analysis. Random agent loss, communication partitions, and formation bottlenecks naturally create persistence and death-set questions.

The domain also provides a stringent falsification test. If agents systematically coordinate around simultaneous-response traps that are invisible in their local menus, then either E information is entering x, the primitives contain prediction, or the architecture is not E-blind. Collective success alone is not evidence; the information pathway and matched controls matter.

A decisive first study

Simulate and then physically test a small swarm maintaining coverage and connectivity under node loss. Include an Ω-trap where a locally attractive role has no future legal handoff and an E-trap where multiple agents simultaneously choose the same recovery corridor. Compare E-blind depths, randomized tie policies, consensus control, and true joint MPC on identical local primitives.

SUPPORT WOULD LOOK LIKE Deeper local h avoids role dead ends, the swarm retains predicted viability under bounded losses, and strict E-blind agents remain vulnerable to simultaneous-response traps that coordinated predictive controllers avoid.

KEY BOUNDARY Neighbour observations straddle the factorization: a current measured neighbour may enter x, but a predicted neighbour response belongs to E. Delays, asynchronous updates, symmetry breaking, and emergent global objectives can make the boundary difficult to audit.