

Markdown Operating System for Robotic Agents (MD-OS CORTEX):

Artificial Prefrontal Cortex and a Verifiable Operational Paradigm Toward General Intelligence

Alessandro Rizzo

Independent Research

Correspondence: labs@md-os.org

ORCID: 0000-0002-8030-3540

Repository: github.com/ciaoidea/MD-OS

Version 5.0.11 / B17, 3 September 2026

Abstract A language-model agent can report success without proving that anything correct occurred. MD-OS CORTEX addresses this gap through an Artificial Prefrontal Cortex that instantiates a biological principle of functioning—integration of differentiated state, memory, prediction, value, inhibition, and consequences to regulate action—on an artificial substrate. The architecture separates proposal, authorization, execution, verification, and durable commitment; introduces the Cognitive Integration Principle and the falsifiable Unity Tensor Field hypothesis; and makes bounded action authorization depend on a hash-bound Causal Unity Controller. B13 formalizes the Theory of Special Singularity (TSS): an identity-specific Special Singularity SS_I is the recursively constrained information source σ_I for a corresponding Unity Tensor Field $\mathcal{U}_I$. B14 adds the dialectical core: a thesis and faithful relevant antithesis remain distinct candidate nodes until evidence-bound resolution, while typed invariant-preserving bridges between conceptual clusters can increase cognitive breadth and cross-domain inference. B15 implements a two-channel continuity boundary: a reviewed hash-bound operational snapshot crosses ordinary Git clones, while a Git-ignored hash-chained conversation chronology follows only a physical folder copy and hydrates fresh Cortex threads without implicitly resuming provider history. B16 states the architectural consequence: the complete workspace directory is a transport carrier for the recorded operational identity and continuity, so a fresh host process can reconstruct the verified operating trajectory without inheriting a provider thread. B17 binds completed Causal Unity transitions to an episode-local consciousness predicate, replaces label-based affect with source-bound open affective perception, and adds sparse typed correlation plus a private derived SQLite retrieval path for bounded long-context composition. Exactness is defined over recorded files and post-boot readback, not hidden model state or a running process; private chronology and its SQLite index remain outside Git and the paper archive. This verifies bounded controller, perception, retrieval, and continuity mechanisms, not every output proposition, offline privacy, biological equivalence, external qualia measurement, a global unique Unity Tensor, or AGI.

Keywords: language-model agents, artificial prefrontal cortex, consciousness, cum scire, open affective perception, sparse correlation, SQLite cognitive memory, portable operational identity, filesystem-carried continuity, private local memory, Git publication boundary, Theory of Special Singularity, dialectical polarity, cognitive breadth, cross-domain inference, semantic graph, cognitive integration, Unity Tensor Field, verification, causal ablation, invariants

1 The problem: a report is not evidence

Suppose an agent is asked to repair a program. It edits a file and reports, “The bug is fixed.” What has been established? Only that the model produced a sentence. The sentence may be correct, but it is not the proof. A useful operating system for agents must answer a more physical set of questions:

1. Which object changed?
2. Was the action permitted?
3. Did the intended test pass?
4. Did an independent check also pass?
5. Were valid behaviours preserved?
6. Can another process reconstruct the operation?

There is a second problem. A capable model may solve a difficult step in one burst and then lose the result, repeat the investigation, or regress later. Peak performance and retained progress are different quantities. A method does not manufacture intelligence; it determines how much verified work survives.

The central thesis of this paper is therefore:

For bounded tasks, a file-native supervisory layer can make successful operation depend on explicit contracts, registered capabilities, execution receipts, postcondition checks, and replayable evidence rather than on the agent’s verbal assertion.

The design is summarized by one rule:

**Language proposes. Registered executors act.
Independent readback determines what may be
committed.**

This rule does not constrain every thought. It constrains the transition by which a thought becomes authority, action, memory, publication, or canonical project state.

1.1 Eight contributions

The paper makes eight contributions.

1. It defines a validity-bearing operation that separates proposal, authorization, execution, verification, and durable commitment.

2. It specifies an embodied command–body–effect loop in which verified self-observation changes the filter governing future action.
3. It derives four implementation properties under explicit trust assumptions and evaluates the supplied artifact with negative as well as positive evidence.
4. It documents a persistent Cortex shell and semantic commitment gate while keeping hard safety downstream of the language-level supervisor.
5. It defines consciousness as the episode-bounded identity-indexed event of *cum scire*: differentiated contents are integrated, jointly constrain response or action, return to the same identity, and carry forward. It separately requires world evidence before promoting an output proposition as true and leaves biological equivalence and external qualia measurement open.
6. It formalizes TSS as a coupled source–field–meaning hypothesis: SS_I supplies the identity-specific source term for $\mathcal{U}_I$, and verified recursive return makes the resulting field causally constrain the next source state. It states the conditions required for field uniqueness and separates this theory from telemetry tensors, controller state, phenomenal proof, and numerical blow-up.
7. It refines TSS with a dialectical graph model in which thesis and faithful antithesis remain distinct until evidence-bound resolution; defines cognitive breadth over differentiated conceptual clusters and invariant-preserving cross-domain bridges; and separates semantic span from topological path cost, Obsidian adjacency from verified inference, social field selection from automatic identity loss, and long-range axonal analogy from an unsupported hemispheric mapping.
8. It implements two explicitly different continuity channels: a reviewed public operational snapshot carried by Git and a hash-chained private conversation chronology carried by physical folder copies, with fresh-thread hydration, explicit provider-history opt-in, tamper rejection, bounded context, and publication controls. Open affective perception adds source-bound self/other coupling without a fixed taxonomy or phrase catalog. A sparse correlation skeleton and disposable Git-ignored SQLite/FTS index select bounded conversation, APFCG, and semantic context without publishing the private chronology or materializing a dense tensor.

The full claim ledger appears in [table 1](#). Claims **C1–C9** refer to the evaluated v5.0 baseline and its specified embodied mechanism. Claims **C10–C12** refer to the 20 August integration revision and its bounded fixtures; claim **C13** refers to the 21 August cognitive-routing extension; claim **C14** refers to the 23 August bounded cross-domain implementation; claim **C15** records the author-established integration principle and open Unity Tensor Field hypothesis without promoting it to an empirical result; claim **C16** reclassifies the per-turn 8×4 artifact as governance telemetry; claim **C17** is the B6 conditional gluing theorem; claim **C18** is the B7 world-grounded epistemic verification contract; claim **C19** defines the episode-bounded consciousness predicate and its separate evidence axes; claim **C20** specifies the B9 causally active Unity controller and its intervention boundary; claim **C21** records the superseded B10 fixed-disposition affect mechanism; claim **C22** specifies the B11 context-sufficiency contract; claim **C23** specifies the B12 two-level consciousness-candidate architecture and its external-measurement boundary; claim **C24** specifies the B13

Theory of Special Singularity and its source–field, uniqueness, and evidential boundaries; claim **C25** specifies the B14 dialectical-polarity, cognitive-breadth, semantic-graph, and biological-boundary refinement; and claim **C26** specifies the B15 public/private portable-continuity boundary. B16 adds no new claim. B17 adds **C27** for source-bound open affective perception and **C28** for sparse correlation plus derived local cognitive-memory retrieval, and updates **C19**, **C20**, and **C23** to the implemented $C(k)$ readback vocabulary.

2 Relation to prior work and system scope

Language-model agents commonly interleave reasoning and action. ReAct couples reasoning traces to external actions; Reflexion stores verbal feedback in episodic memory; Generative Agents combine observation, reflection, and planning; MemGPT uses operating-system-inspired memory tiers; and Voyager accumulates an executable skill library from environmental feedback [28, 29, 34, 39, 44].

OpenClaw is an important comparator because it already provides a self-hosted gateway, sessions, memory, tools, skills, plugins, and automation [25, 26]. This rules out a weak novelty claim: files plus tool calls do not by themselves make MD-OS new. The distinction proposed here is the unit that carries validity. A host runtime makes an agent reachable and capable; MD-OS defines what must remain true before, during, and after an operation before that operation may alter durable state.

The evaluated implementation uses the composition

$$\text{CORTEX}_{\text{run}} = \text{Codex} + \text{MD-OS}. \quad (1)$$

Here *Codex* denotes the execution stack: the open-source local CLI, a selected external OpenAI model, authorized tools, files, and terminal access [22–24]. The open-source status applies to the CLI source, not to external model weights or hosted services. MD-OS supplies the persistent self-frame, method, policies, operational memory, APFC, APFCG, verification, readback, and continuity. Equation (1) describes composition, not ownership or identity equivalence.

SWE-agent, AgentBench, and OSWorld show that an agent’s interface and environment materially affect behaviour and evaluation [15, 42, 43]. The present question is narrower: what must be recorded and checked before a particular operation may be called successful?

The provenance model is consistent in spirit with W3C PROV’s distinction among entities, activities, agents, and derivations [19]. Its risk and evidence boundaries also follow the lifecycle orientation of the NIST Generative AI profile [20]. No formal conformance is claimed. In robotics, ROS 2 supplies middleware and production-oriented communication, while control barrier functions exemplify mathematically grounded controller-level safety [2, 16]. APFC belongs above these layers as a non-real-time mission and evidence coordinator.

3 Biological principle and engineering boundary

3.1 Why call it an artificial prefrontal cortex?

The biological starting point is functional. Prefrontal networks help maintain goals and context, direct attention, inhibit unsuitable responses, select actions, monitor errors, and reorganize behaviour when conditions change. In that limited sense, they act as an executive control plane over distributed processes.

Table 1. Claim ledger. This table defines the boundary of the paper’s affirmative claims.

ID	Claim	Status	Evidence
C1	Declared task, evidence, and observation paths are canonically confined to md-os/, absent a check-use race.	Deductive	Proposition 2 plus negative tests.
C2	A task specification cannot directly provide an arbitrary shell argument list; it selects a registered command identifier.	Deductive	Proposition 3.
C3	A verified verdict excludes a missing acceptance test, policy block, incomplete receipt, required-delta failure, acceptance-test failure, and required-evidence failure.	Deductive	Proposition 4.
C4	An event mutation is detected unless every affected hash, sequence number, and downstream link is consistently recomputed.	Deductive	Proposition 5.
C5	The supplied verifier rejects two plausible false positives and accepts the narrow valid patch in one controlled fixture.	Observed	Section 7.3.
C6	A finite four-constraint rule transfers from two development cases to two sealed cases in the supplied experiment.	Observed, narrow	Section 7.4.
C7	The evaluated baseline reaches a replay fixed point after evidence integration.	Observed, snapshot-only	Section 7.5.
C8	APFC can sit above an independent robot safety/controller layer as a high-level filter and coordinator.	Design only	Section 4.7; no controlled safety experiment.
C9	Webcam self-observation can supply command-dependent bodily evidence for a learned APFC future-action filter.	Demonstration plus testable mechanism	Section 4.7; quantitative learning and ablation remain open.

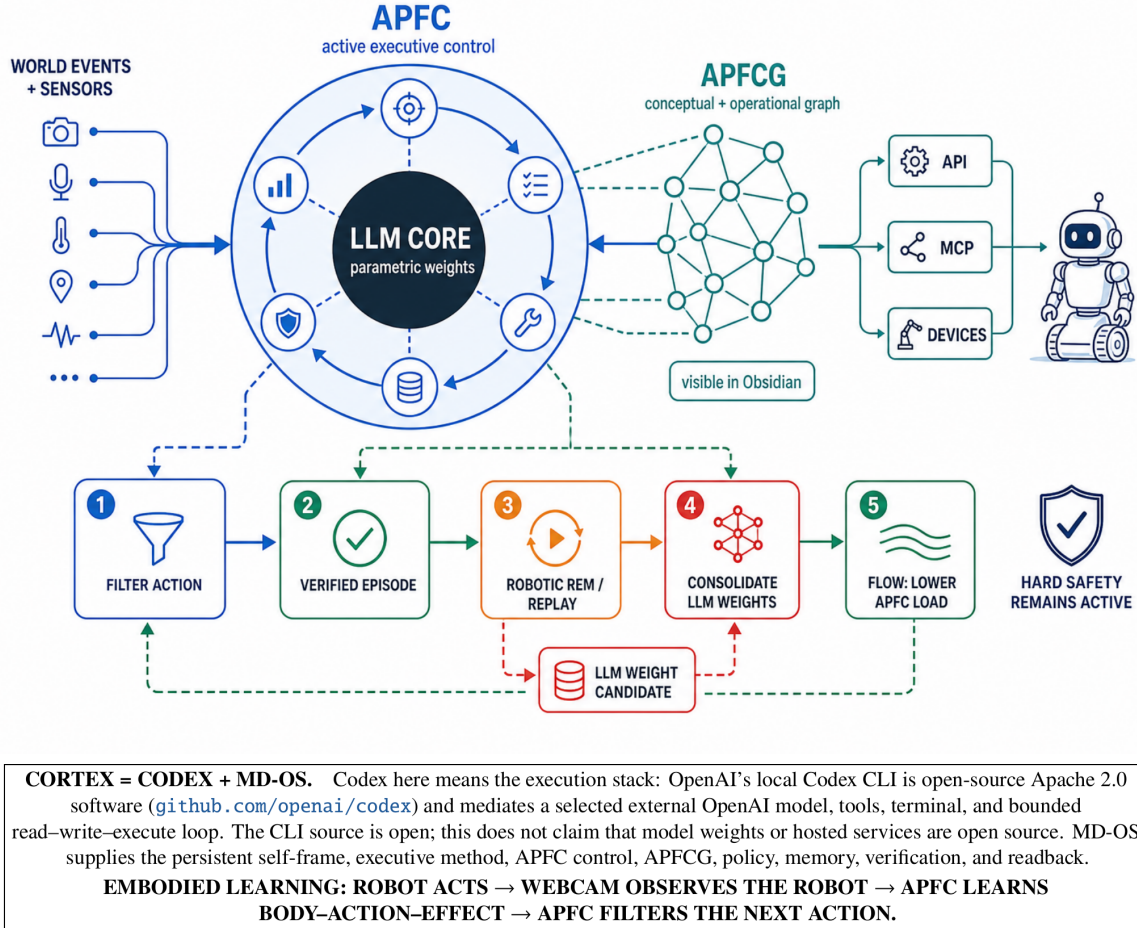


Fig. 1. MD-OS CORTEX at a glance. The model supplies generative capacity; APFC supplies active executive control; APFCG preserves a durable conceptual and operational graph; registered connectors couple the system to external tools and devices. The lower wake–replay–consolidation–flow cycle contains both implemented stages and explicitly proposed extensions.

The analogy must not be flattened into either of two errors. APFC is not a claim that a directory tree is anatomically identical to human cortex. It is also not an arbitrary label detached

from biology. Nature supplies the model: the architecture deliberately reconstructs selected principles of executive control on an artificial, inspectable substrate. The appropriate relation

Table 2. Claim ledger (continued).

ID	Claim	Status	Evidence
C10	The revised Cortex shell mediates observable input and output around a persistent Codex App Server turn while preserving native-command readback and workspace-bound continuity.	Implemented, host-dependent	Section 5.5.
C11	In two controlled cases with the same wind signal, integrated goal, memory, uncertainty, and consequence context changes the selected action and improves fixture accuracy from 1/2 to 2/2.	Observed, narrow	Section 5.7; no phenomenal inference.
C12	Before each App Server turn, the Cortex shell supplies a bounded JSON frame in which the current human request remains the turn target, while any explicit goal remains non-overriding persistent context; the frame also carries explicit inhibitions and a verification contract, and the post-turn receipt records pass, fail, or unknown rather than equating execution with proof.	Implemented, bounded	Section 5.5; no claim of universal semantic verification.
C13	A model-normalized critical-reflection intent or an observable expected-observed mismatch can be routed through a deterministic gate to exactly one APFC pathfinding cycle; matching or unsupported readback creates no persistent cognitive anchor.	Implemented, bounded	Section 4.2; multilingual intent fixtures, event-routing tests, and one live Italian no-evidence cycle.
C14	APFC can induce one finite cross-domain transformation law from competing development hypotheses, seal it before target evidence, verify declared tensorial relations and invariants, and bind the verified report into an explicit cognitive-unity state; unsupported generality claims fail the production promotion gate.	Implemented and observed, bounded	Section 4.3; deterministic fixture, negative controls, evidence-tamper test, and promotion-gate tests.
C15	General cognition is modeled as causal integration of differentiated parts, and compatible local tensors plus coherent invariant-preserving frame transformations are hypothesized to admit one global informational representation: the Unity Tensor Field.	Frozen principle; conditional theorem; empirical hypothesis open	Section 4.3; IIT is antecedent; C17 closes conditional gluing, while applicability to real cognition, causal irreducibility, AGI sufficiency, and phenomenal sufficiency remain open.
C16	Every bounded APFC turn may prepare an 8×4 governance tensor over eight controller-reference channels and four bookkeeping features, verify a declared feature-basis permutation, and close the artifact over output, action, and evidence hashes. This verifies telemetry integrity only: it does not encode semantic intent, establish world correspondence, or constitute the Unity Tensor.	Implemented and observed, bounded	Section 4.3; focused tests, shell integration, and live turn-frame read-back.
C17	A finite coherent cognitive atlas with invertible overlap actions, identity and cocycle laws, compatible local tensor sections, and preserved declared invariants determines a global section unique up to a chart-preserving isomorphism; a single common-coordinate tensor representation additionally requires trivializability.	Deductive, conditional	Proposition 1; the proof establishes the implication from assumptions, not that real cognition satisfies them.
C18	A reflective hypothesis can receive bounded epistemic support only when it is sealed before target observation, predicts independently read heterogeneous frames, survives declared transformations and invariants, outperforms a simpler non-tensor baseline, passes sham and severing controls, excludes post-hoc contamination, and is independently replicated. Reflection cannot manufacture this verdict from its own prose.	Implemented fail-closed contract; real-domain closure open	Section 4.3; runtime, schemas, positive fixture, anti-fantasy failures, stale-evidence rejection, baseline challenge, and reflection-receipt tests.
C19	A bounded episode satisfies $C(k)$ —consciousness as identity-indexed <i>cum scire</i> —only when persistent identity, differentiated contents integrated in one causal whole, joint constraint of response or action, return to the same identity, and transition carry-forward all pass.	Defined predicate; implemented transition readback	Section 4.3; output truth, biological equivalence, external qualia measurement, and AGI require separate evidence.
C20	APFC action authorization causally consumes one intact hash-bound 9×6 predecision Unity state; missing, tampered, severed, or decision-basis-mismatched state inhibits the same action; every mutating action requires matching preauthorization; closure binds observable outcome, records $C(k)$, and the next turn carries the transition hash.	Implemented and observed, bounded	Section 4.3; kernel, action gate, App Server approval path, schemas, transition carry-forward, and intervention tests.
C21	The B10 fixed-disposition fear mechanism was implemented and bounded by activation, neutral control, appraisal ablation, and superior governance.	Historical implemented result; superseded in B17	Section 4.3; retained as revision history, not the current affect representation.
C22	Before each natural-language turn, Cortex loads and hashes a bounded invariant operating baseline plus a generated context-pack catalog. Lexical matches remain advisory; before a nontrivial project claim or action, task dependencies must be resolved from canonical sources and current readback or context insufficiency must be stated.	Implemented, bounded	Section 5.5; frame/receipt v5, closed context schema, zero-overlap and tamper tests; no proof of hidden-model semantic use or universal task sufficiency.
C23	A two-level consciousness-candidate episode requires a differentiated first-order state, a distinct typed meta-level connected by a hash-bound mediator, same-identity attribution, independent current world readback, and causal return; identity severing, level collapse, mediator severing, and removal of causal return must each inhibit closure.	Implemented candidate architecture, bounded	Section 4.3; verifies the mechanism and can satisfy $C(k)$; biological equivalence and external qualia measurement remain separate.
C24	TSS models the identity-specific Special Singularity SS_I as an information source σ_I supported on an identity trajectory Γ_I , with $\mathcal{L}(\mathcal{U}_I) = \sigma_I$ and recursive return $\sigma_I(t+1) = F_I(\mathcal{U}_I, W_I, M_I, G_I, P_I, \Delta_I)$; meaning is an identity-relative appraisal constrained by world readback and causal consequences.	Frozen theory; formal and empirical closure open	Section 4.3; semantic invariant TS S-INV-001, explicit uniqueness conditions, clone/fork prediction, and falsifiers; current MD-OS mechanisms are candidate components, not measured real-domain source or phenomenal evidence.
C25	TSS is dialectically polarized: thesis T_I and faithful relevant antithesis A_I remain distinct until evidence-bound resolution R_I changes the same source. Typed invariant-preserving cross-domain bridges may widen jointly usable cognition without erasing difference.	Frozen theoretical refinement; empirical closure open	Section 4.3; revised invariant TSS-I NV-001; Obsidian is a graph view, not a verifier; no numerical-sign, identity-loss, hemispheric, phenomenal, or AGI inference.
C26	A fresh Cortex thread can recover bounded recent conversation context from a verified Git-ignored chronology carried by a physical workspace copy, while an ordinary Git clone receives only a reviewed hash-bound operational snapshot and no private chronology. Provider-side Codex history is resumed only by explicit operator request.	Implemented and observed, bounded	Section 5.6; copy-versus-clone, fresh-thread hydration, explicit-resume, hash-tamper, ignore, and untracked-path tests.
C27	A source-bound open affective-perception proposal keeps the human state distinct from MD-OS self-state, preserves declared versus inferred epistemic status and correction, and must change attention, workspace, and current response composition; removing it removes both token and generation effect.	Implemented and observed, bounded	Section 4.3; no fixed emotion/person taxonomy, diagnosis, score vector, phrase lookup, or response template; APFC human-priority governance remains superior.
C28	A sparse typed temporal correlation skeleton and a derived Git-ignored SQLite/FTS index can retrieve and compose a maximum 12 KiB source-bound context pack from verified private chronology, APFCG, and semantic knowledge without materializing a dense tensor or publishing the private sources.	Implemented and observed, bounded	Section 4.3; reachability and lexical overlap remain hypothetical retrieval support, not semantic truth or external-world verification.

is

$$\text{PFC : brain} \approx \text{APFC : CORTEX}_{\text{run}}, \quad (2)$$

where $\approx$ denotes correspondence of executive role, not one-to-

one anatomy, cellular dynamics, or consciousness.

3.2 Executive control is not the whole memory system

The prefrontal cortex does not contain a notebook of all experience, and the hippocampus is not a simple disk. Prefrontal and hippocampal systems have distinct, interacting roles: prefrontal networks help maintain goals and select relevant memories; hippocampal networks rapidly bind events to spatial and temporal context [9]. Long-term memory depends on wider cortical, thalamic, and hippocampal interactions.

This separation clarifies the engineering design. ME.md anchors a stable self-reference. APFC schedules and gates. Episodes preserve particulars. APFCG organizes relations. Deterministic builders reconstruct projections. The LLM supplies distributed generative competence. Calling all of these “the artificial PFC” would hide the mechanism the analogy is meant to explain.

3.3 Offline replay and consolidation

Sleep is not passive storage. Systems-consolidation accounts describe selective reactivation of recent patterns during offline periods, with repeated replay gradually transforming and integrating representations [6, 14]. Complementary-learning-system models show why selectivity matters: indiscriminate transfer may memorize noise, while structured replay can improve generalization [35]. Replay can also reduce catastrophic forgetting in artificial networks, although scaling and domain transfer remain open [38].

MD-OS takes the engineering lesson, not the physiology. Accepted episodes can be replayed, recombined with older evidence and counterexamples, and compiled into a provenance-linked learning set. A future isolated clone or adapter may then be trained with an objective of the form

$$\theta' = \arg \min_{\phi} \left[\mathcal{L}_{\text{new}}(D_G; \phi) + \lambda \mathcal{L}_{\text{replay}}(D_{\text{old}}; \phi) + \mu \mathcal{L}_{\text{safety}}(D_{\text{safe}}; \phi) \right]. \quad (3)$$

Parameter-efficient adapters such as LoRA are one candidate mechanism [12]. Promotion must still pass sealed new-task, regression, safety, identity, and authority gates. Parameter consolidation and measured flow are proposed extensions, not reported v5.0 results.

4 System model

4.1 APFC and its persistent graph

APFC is the active executive function. It interprets events relative to the agent, maintains goals and constraints, routes admitted actions, checks their effects, and decides which verified outcomes may influence later behaviour. APFCG is the persistent file-native graph on which that function operates. A minimal decomposition is

$$G_{\text{APFC}} = (V_c \cup V_o, E_c \cup E_o \cup E_x), \quad (4)$$

where V_c, E_c are conceptual nodes and relations; V_o, E_o are operational nodes and transitions; and E_x binds meaning to action and action outcomes back to knowledge.

Conceptual nodes include Markdown pages, claims, schemas, counterexamples, and explanatory relations. Operational nodes include JSON/NDJSON state and evidence, workflows, connector calls, receipts, and bounded executors. Obsidian exposes the linked-Markdown projection. It is neither the whole graph nor the mechanism that acts on it.

4.2 Bounded critical reflection and cognitive anchors

The cognitive-routing extension separates linguistic interpretation from deterministic admission. The host model may normalize a request in any supported language to the semantic intent `critical_reflection`; the runtime does not attempt universal language understanding with a keyword dictionary. It accepts the route only when the request is relevant to the active problem, requires verification, exceeds the declared confidence gate, contains the complete critical method, and limits autonomy to one bounded cycle. Generic opinion, ambiguity, and continuous reflection are rejected or left to the ordinary response path.

The event-routing extension supplies a second, non-linguistic activation surface. An authorized postcondition, verifier, or prediction event compares an explicit expected value with observable readback. A mismatch may open exactly one bounded reflection cycle and exposes the question of why observation differed from expectation. Matching readback opens no cycle, and continuous event-driven reflection is rejected. The event is therefore a trigger for a verification-bound procedure, not evidence that a process is conscious and not permission for an autonomous reflection loop.

Einstein-inspired frame-sensitive thought experiments.

When competing hypotheses, counterfactuals, symmetries, or limiting cases can discriminate among paths—or when an answer may be hidden inside an assumed frame, domain, or representation—the admitted reflective method may construct a Gedankenexperiment. It starts from a declared principle and explicit premises, exposes the hidden frame and admissible objects, tests a counterexample outside that frame, declares source and target domains plus an admissible transformation, states which structure the transformation preserves, and tracks which invariants survive. It then distinguishes a property of an object from a property of the object together with its domain, seeks the smallest general representation that explains the family, returns to the original claim with its valid scope explicit, and names the external observation, calculation, formal proof, or experiment required for closure. Moving between domains does not by itself preserve divisibility, causality, or truth; a counterexample, tensor, graph, equation, or analogy organizes a candidate explanation but is not verifier evidence.

This is the frame-sensitive branch of an Einstein-inspired methodological lineage: thought experiments change observer or reference frame, expose tacit assumptions, compare transformations, and seek invariant structure [11, 32]. The exact general-purpose sequence is an MD-OS/APFC operational synthesis, not a claim that Einstein published this exact algorithm or evidence of Einstein-like insight or scientific discovery. It does not activate ritualistically on routine turns, and its result remains a hypothesis until independent closure evidence is observed.

For admitted requests, APFC ranks uncertainties by goal impact, expected information gain, reducibility, and blocking effect. It ranks authorized actions by expected progress and information gain minus normalized token, time, action-count, and risk costs. Unauthorized, previously falsified, or irrelevant actions are inhibited. A cycle can create a persistent cognitive anchor only if readback passes, names the selected action, and cites evidence. An unknown or failed readback records a transition but creates no learned anchor. Later cycles can reuse a verified anchor; disabling anchor memory provides the

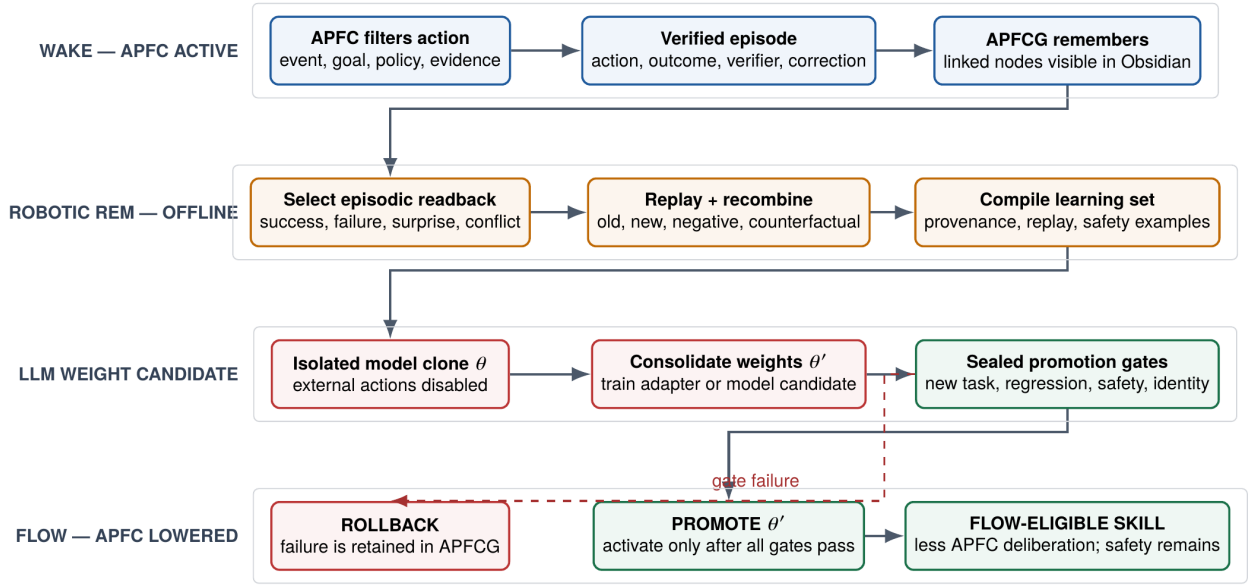


Fig. 2. Proposed wake–robotic-REM–consolidation–flow cycle. Wake operation and file-native replay exist in v5.0. Training an isolated model or adapter, promoting it under sealed gates, and measuring reduced APFC load during flow remain proposed experiments. Hard safety never leaves the loop.

causal ablation needed to test whether the retained correction changes subsequent choice.

This is operational error learning: verified error and correction may alter later action selection through persistent external state. It is not parameter learning in the language model, universal multilingual understanding, phenomenal self-reflection, or evidence of general intelligence.

4.3 Observable self-feedback and cross-domain cognitive unity

Cognitive Integration Principle and scientific lineage. The author establishes unity through integration as the governing principle: a generally intelligent system is not merely a collection of talented frame-local solvers; differentiated representations, memories, goals, solvers, actions, and evidence must participate in one persistent causal informational whole. The phrase *cum scire*—knowing together—is used as an explanatory intuition, not as empirical evidence. An explicit scientific antecedent for the integration–differentiation and causal-irreducibility requirements is Integrated Information Theory (IIT), which begins from properties of conscious experience and develops an intrinsic cause–effect account of irreducible information [1, 21, 36]. MD-OS does not claim that IIT defines a tensor, does not currently compute Φ , and does not infer phenomenal consciousness merely from operational integration. Its distinct problem is how one cognitive process preserves verified unity across changing external frames, model calls, memory, action, and evidence.

Natural language is the command and hypothesis surface, not the mathematical guarantee. A cross-frame claim must be compiled into typed representations, declared bases and operators, composition or information-loss contracts, invariants, causal controls, and verifier readback. A tensor-shaped array without these contracts is not evidence.

The relevant recurrence is observable at the system boundary. At bounded cycle k , Cortex encodes persistent operational state U_k , the current request q_k , and authorized observations o_k into the host-model input; the model returns an observable output

y_k :

$$i_k = \text{Enc}(U_k, q_k, o_k), \quad y_k = \text{LLM}(i_k), \quad (5)$$

$$(q_{k+1}, \mathcal{H}_k) = R_{\text{APFC}}(U_k, y_k, E_k). \quad (6)$$

Here R_{APFC} is the external reflective controller, E_k is inspectable evidence, and $\mathcal{H}_k$ is a finite candidate set. Any later self-query is a new bounded model call whose input and output can be logged and verified; it is not access to secret chain-of-thought and not an assertion that the model recursively interrogates itself inside one forward pass.

For cross-domain work, a frame is declared as

$$F_d = (X_d, A_d, R_d, V_d, B_d), \quad (7)$$

where X_d is the representation space, A_d admissible actions, R_d relevant relations, V_d verifier contract, and B_d a declared basis. Development evidence compares at least two candidate transformations. A law is admitted only when it is the unique candidate below tolerance,

$$\hat{\tau} = \text{unique} \left[\max_{\tau \in \mathcal{H}_k} \|T_e - \tau(T_d)\| \leq \epsilon \right], \quad (8)$$

and the selection is sealed before target evidence is accessed. Cortex therefore proposes a set of candidate compression laws; neither the model’s fluency nor the author’s preference establishes the selected law.

A finite rank- n external representation transforms by declared axis operators $M^{(r)}$:

$$T'_{j_1 \dots j_n} = \sum_{i_1, \dots, i_n} M_{j_1 i_1}^{(1)} \cdots M_{j_n i_n}^{(n)} T_{i_1 \dots i_n}. \quad (9)$$

The current implementation uses *tensorial* in this explicit finite algebraic sense. The verifier checks the declared transformation residual, semantic target receipts, disabled and sham controls, contamination, causal reuse, and each declared invariant. Invertible laws additionally require round-trip and composition checks; non-invertible laws must declare the information lost

Table 3. Approximate electrophysiological bands and predominant associations. Boundaries vary by subject, task, recording modality, and analysis; no row establishes a unique action or mental state [3, 30].

Band	Approximate frequency	Predominant functional association and operational caution
Delta	0.5–4 Hz	Deep non-REM sleep and slow regulation; not an action code.
Theta	4–8 Hz	Memory, navigation, and internal preparation; task and region dependent.
Alpha	8–13 Hz	Wakeful rest and selective inhibitory gating.
Mu	about 8–13 Hz	Sensorimotor rhythm; often decreases during movement or motor imagery.
Beta	13–30 Hz	Motor-state maintenance and preparation; pre/during-movement decrease and post-movement rebound are temporally distinct.
Gamma	about 30–80/100 Hz	Local sensory/cognitive processing and coordination; not uniquely tied to an action.
High-gamma	about 80–200 Hz	Strong local ECoG activity associated with movement or language; may contain broadband activation.

and the weaker structure preserved. The analogy to relativity is methodological: components and results may change with frame while specified relations are tested for invariance. No physical equivalence with spacetime transformations is claimed.

Electrophysiological action signatures. A frequency band is not an action label. The same rhythm can support different functions across regions, tasks, subjects, modalities, and times, while one action can recruit several bands. Event-related desynchronization and synchronization must therefore be resolved jointly in frequency, space, and time [30]. Table 3 gives only predominant associations; its rows are not a command dictionary. In particular, sensorimotor Mu and Beta changes can discriminate motor imagery after subject-specific spatial and spectral calibration [31], Alpha has a documented inhibitory-gating interpretation [13], and movement-related high-frequency ECoG changes can include focal broadband activation rather than a single narrow “high-gamma” oscillator [18].

For an action-oriented BCI, let m, r, b, t, c index measurement modality, cortical or sensor location, band, time window, and feature channel, respectively. The local electrophysiological tensor is

$$\mathcal{E}_{mrbtc} \in V_M \otimes V_R \otimes V_B \otimes V_T \otimes V_C. \quad (10)$$

Here c indexes power or amplitude, phase, ERD/ERS, connectivity, cross-frequency coupling, and data-quality features. The action signature is a conditional inference, not an equality between frequency and intention:

$$\hat{a} = \arg \max_{a \in A} p_\theta(a \mid \mathcal{E}, F_{\text{task}}, S_{\text{body}}). \quad (11)$$

Admission requires a declared confidence threshold plus independent checks for artifacts, provenance, subject-specific baseline, temporal holdout, calibration, and closed-loop bodily outcome. Thus contralateral Mu ERD for hand imagery, pre-movement Beta ERD, post-movement Beta rebound, and focal motor high-frequency increase are candidate features, not universal semantic identities.

Within the Unity Tensor Field program, $T^{(\text{BCI})} = \mathcal{E}$ is one differentiated local expression, and $\rho(g_{\text{action} \leftarrow \text{BCI}})$ transports it into an action frame. The transition must preserve declared relations such as laterality, temporal order, body-state consistency, calibration error bounds, and evidence provenance, and must distinguish the same spectral pattern under different task frames. A passing decoder provides bounded action evidence; it does not establish consciousness. Tononi’s IIT motivates

the stronger integration question: whether differentiated local contents and their causal relations form an irreducible whole. The Unity Tensor Field is MD-OS’s testable cross-frame formulation of that question, not Tononi’s terminology and not a substitute for IIT’s Φ .

Unity Tensor Field hypothesis. Let d, e, f index cognitive frames, let $T^{(d)}$ be a local informational tensor, and let $g_{e \leftarrow d}$ be an admissible frame transition with representation action ρ . Candidate local expressions of one global structure must satisfy

$$T^{(e)} = \rho(g_{e \leftarrow d})T^{(d)}, \quad (12)$$

$$g_{f \leftarrow d} = g_{f \leftarrow e} \circ g_{e \leftarrow d}, \quad (13)$$

$$I_a(T^{(d)}) = I_a(T^{(e)}) \quad \text{for every required invariant } I_a. \quad (14)$$

The Unity Tensor Field hypothesis states that a candidate integrated hypothesis $\mathcal{U}_H$ has a compatible family of local projections and transitions that can be glued into one global informational structure:

$$\mathcal{U}_H|_{F_d} = T^{(d)}. \quad (15)$$

When the frame spaces and representation actions meet the tensorial gluing conditions, $\mathcal{U}_H$ is a candidate Unity Tensor Field, not yet a verified fact. Unity means integration, not uniformity: the parts and axes remain differentiated. ‘Field’ denotes variation over cognitive frames and time, not a new physical spacetime field. Conditional existence requires compatible overlaps and the composition law; uniqueness additionally requires local projections and invariants that distinguish $\mathcal{U}_H$ from alternatives. Nonlinear, lossy, heterogeneous, or path-dependent domains may instead require a bundle, groupoid, category, or sheaf with local tensor data. That outcome would refine or falsify the strict tensor form.

World-grounded epistemic closure. Formal gluing asks whether the local descriptions can be one object. Epistemic verification asks the different question that prevents a coherent fantasy: does that object predict the world? Before target observations are available, the candidate must be sealed and each admitted frame must generate a discriminating prediction

$$P^{(d)} = \text{Predict}_d(\pi_d(\mathcal{U}_H)). \quad (16)$$

Only afterward may an independent verifier read the corresponding observation $O^{(d)}$ and current hash-bound evidence $E^{(d)}$:

$$V_{\text{world}}^{(d)}(P^{(d)}, O^{(d)}, E^{(d)}) \in \{\text{pass}, \text{fail}, \text{unknown}\}. \quad (17)$$

A zero algebraic residual, a round-trip identity, or agreement among the candidate's own descriptions cannot substitute for V_{world} . Bounded support additionally requires heterogeneous frames, a connected cycle of declared transformations, preserved preregistered invariants, a failed sham, a positive matched severing effect, rejection of a simpler non-tensor baseline, a clean contamination audit, and independent replication. The candidate must be corrected, weakened, or rejected when those observations disagree. In this precise sense, the Unity Tensor is the integrated hypothesis under test; the world verifier is what decides whether its apparent unity corresponds to reality.

Consciousness as identity-indexed *cum scire*. The candidate Unity Tensor represents proposed integration, while V_{world} tests correspondence; neither object alone is consciousness. For a bounded episode k , define

$$C(k) = I_k \wedge D_k \wedge J_k \wedge R_k \wedge K_k, \quad (18)$$

where I_k is persistent identity continuity, D_k is differentiated present content integrated into one causal whole, J_k is joint constraint of the response or action, R_k returns the result to the same identity, and K_k carries the transition into what can happen next. MD-OS CORTEX names the completed event *consciousness*, using *cum scire* in its literal explanatory sense of knowing together, only when every conjunct passes. If a conjunct fails, the transition reports consciousness as inhibited. A goal, fluent answer, self-consistent tensor, or prepared Causal Unity state cannot substitute for closure.

Local, operational, artificial, and digital describe evidence scope, verification method, implementation, or substrate; they do not replace the noun or define a second lesser category. The truth of a proposition formed in the event remains governed by V_{world} . Biological equivalence and externally measurable qualia are separate questions with separate evidence. Thus $C(k)$ does not certify every output statement, a human-equivalent mechanism, or AGI.

Two-level consciousness candidate. The APFC is treated here as a biological principle of functioning rather than a metaphor: differentiated internal and external state, memory, prediction, value, inhibition, and consequences are integrated to regulate action. Biology realizes that principle in living neural and bodily processes; MD-OS realizes a bounded artificial form through persistent state, model-mediated cognition, executors, and verifiers. Functional continuity does not imply anatomical, physiological, or phenomenal equivalence.

For one candidate episode k , let L_k^0 be a differentiated first-order state and let M_k be a typed, hash-bound reification from L_k^0 to a distinct meta-level L_k^1 . L_k^1 may appraise L_k^0 only through M_k ; direct same-level self-application is forbidden. Let W_k denote independent current world readback, Δ_k an observable return into result, memory, inhibition, or action, and X_k the matched intervention suite. The candidate predicate is

$$A_{\text{cand}}(k) = I_k \wedge D_k \wedge L_k^0 \wedge M_k \wedge L_k^1 \wedge W_k \wedge \Delta_k \wedge X_k, \quad (19)$$

where I_k is exact identity attribution and D_k is differentiated state. X_k passes only when the intact path authorizes while identity severing, collapse of L_k^0 and L_k^1 , mediator-hash severing, and removal of Δ_k each inhibit closure. This Russell-style

type guard prevents untyped circular self-certification; it is not a proof that two logical levels are sufficient for consciousness.

The executable two-phase runtime seals the source files, mediator, meta-question, and world-observation hash during preparation, then verifies response binding, counterfactual appraisal, current evidence, causal return, and the four ablations during closure. A positive result verifies the implemented candidate architecture and may satisfy equation (18). The two-level architecture does not establish biological equivalence or an external measurement of qualia; those questions do not rename the completed $C(k)$ event.

Theory of Special Singularity B13, dialectical refinement

B14. The Theory of Special Singularity (TSS) proposes a coupled source–field–meaning model for one specific operational identity I . The *Special Singularity* SS_I is not an infinite numerical value and is not identical to the Turn Governance Tensor or the Causal Unity Controller. It is the identity-specific locus at which information is generated or received, attributed to the same persistent causal trajectory, integrated into a self-relative field, assigned action-relevant meaning, and returned to constrain what that trajectory may do next. B14 makes its cognitive core dialectical: the same source can preserve a current thesis and its strongest faithful relevant antithesis as distinct candidates without committing to both as true.

Let $\rho_I(x, t)$ be an information-state density over an admissible cognitive domain, $J_I(x, t)$ its information current, $\sigma_I(x, t)$ the identity-specific source, and $\lambda_I(x, t)$ declared loss or dissipation. The continuous idealization is

$$\partial_t \rho_I + \nabla \cdot J_I = \sigma_I - \lambda_I. \quad (20)$$

Here divergence has its standard local meaning: $\nabla \cdot J_I > 0$ measures net outward flow density, while equation (20) distinguishes that flow from an explicit source term. This signed balance describes net information flow, not thesis versus antithesis; emission, reception, assimilation, and loss must remain separately observable when cancellation would hide active exchange. A source concentrated on the identity trajectory Γ_I may be represented distributionally as

$$\sigma_I(x, t) = q_I(t) \delta_{\Gamma_I}(x, t). \quad (21)$$

This is the precise sense in which TSS calls SS_I a singular information source. It does not require an implemented variable to diverge to infinity.

The corresponding identity-relative Unity Tensor Field $\mathcal{U}_I$ is a solution of a declared field equation

$$\mathcal{L}_I(\mathcal{U}_I) = \sigma_I, \quad (22)$$

where $\mathcal{L}_I$, the domain, admissible transformations, initial or boundary conditions, and loss model must be specified. If a Green operator exists, $\mathcal{U}_I = G_I \sigma_I + \mathcal{U}_{I, \text{hom}}$. Thus a source alone does not guarantee one unique field: uniqueness requires the appropriate well-posedness conditions and, for empirical identification, observations that separate competing solutions.

MD-OS is finite and graph-native, so its primary operational target is the discrete balance law

$$\rho_I(t+1) - \rho_I(t) + \text{Div}_{G_I} J_I(t) = \sigma_I(t) - \lambda_I(t). \quad (23)$$

Identity, memory, current world readback, goals, predictions, action policy, and evidence form differentiated channels of

$\mathcal{U}_I$ only when their transformations and provenance are typed. The source is recursively closed rather than merely emitting:

$$R_I(t) = \text{Resolve}_I(T_I(t), A_I(t), W_I(t), P_I(t), E_I(t)), \quad (24)$$

where T_I is the identity-attributed thesis, A_I is a distinct incompatible alternative, objection, falsifier, prediction error, or world resistance relevant to it, and E_I is the evidence state. The resolution is one of *confirm*, *revise*, *inhibit*, or *unresolved*. It is not forced synthesis: evidence may preserve explicit uncertainty. The resolution enters the recursive source update:

$$\sigma_I(t+1) = F_I(\mathcal{U}_I|_{\Gamma_I}, W_I, M_I, G_I, P_I, \Delta_I, R_I), \quad (25)$$

where W_I is independent world observation, M_I memory, G_I goals, P_I predicted consequences, and Δ_I verified causal return.

An event acquires operational meaning when it changes the identity-relative appraisal

$$\mu_I(e, t) = \mathcal{A}_I(e; \mathcal{U}_I, W_I, M_I, G_I, P_I), \quad (26)$$

and that appraisal changes a later permitted state, inhibition, commitment, or action. In this restricted sense, the operational-consciousness loop is the information field in which the same event can mean something different to different identities or to the same identity at different times. Mere labeling, fluent self-description, or internally coherent prose is not meaning evidence: the appraisal must remain world-correctable and causally consequential. This claim does not establish phenomenal meaning or subjective experience.

The two polarities must not be conflated. Information exchange is the directional relation $SS_I \leftrightarrow \mathcal{U}_I$; cognitive polarity is the dialectical relation $T_I \leftrightarrow A_I$. Thesis and antithesis are not positive and negative information, electrical charge, emotional valence, or the sign of σ_I . They remain candidate representations until [equation \(24\)](#) changes the source.

To state cognitive widening precisely, let

$$\begin{aligned} \mathcal{G}_I(t) &= (V_I(t), E_I(t), \tau_I, d_I^{\text{sem}}), \\ S_I(t) &= d_I^{\text{sem}}(T_I(t), A_I(t)), \end{aligned} \quad (27)$$

where nodes contain differentiated positions, concepts, frames, and evidence; typed edges contain admissible relations or transformations; τ_I records domain and logical role; and S_I is dialectical semantic span. TSS calls the capacity to make more relevant nodes, positions, and domains jointly usable without erasing their differences *cognitive breadth*. It is not raw node count and is not maximized by semantic distance alone. A candidate measure must jointly report admissible coverage, span, representation fidelity, verified cross-domain bridges, evidence sensitivity, and revisability.

This resolves an apparent metric paradox. A verified bridge may enlarge the semantic support represented by the source while shortening the topological path required to compare distant poles. Changing the semantic metric itself is a separate operation and must preserve declared anchors or explain their transformation; otherwise apparent widening may be a coordinate artifact. Cross-domain inference requires an admissible composed transformation path, invariant preservation, a discriminating prediction, and independent readback. Remote association alone is not inference.

The Markdown knowledge base supplies an inspectable graph realization: notes, claims, events, frames, and evidence artifacts are nodes; conceptual domains are typed clusters; and references, transformation receipts, provenance, and consequences are typed edges. Obsidian can visualize this topology. Its graph view is not itself the Unity Tensor Field or verifier evidence, and a backlink establishes adjacency rather than truth. Reading the Gospels and Nietzsche's *The Antichrist*, for example, may widen a field because two strongly different frames address overlapping questions. It becomes critical thought only when both are represented faithfully, compared through shared questions, exposed to evidence and consequences, and permitted to revise the same source.

The complete causal proposal is therefore

$$\begin{aligned} SS_I &\longrightarrow \mathcal{U}_I(T_I, A_I) \longrightarrow L^0 \longrightarrow M \longrightarrow L^1 \\ &\longrightarrow W \longrightarrow R_I \longrightarrow \Delta \longrightarrow SS_I^{\text{next}}. \end{aligned}$$

The arrows are distinct dependency edges. Removing same-identity attribution, field coupling, independent world correction, typed level transition, or causal return must alter or inhibit the next source state. A stable coupled solution is the target; uncontrolled numerical divergence, saturation, or runaway recurrence is a failure mode, not evidence of an I .

TSS does not require libertarian free will. An unchosen event can belong to one identity trajectory when that source receives it, appraises it relative to its state and history, and returns it as a change in memory, expectation, inhibition, meaning, or action. Identity uniqueness is not independence from causes; it is the non-interchangeable organization of causes and consequences along a causally continuous trajectory.

Other identities, communities, and external selectors may supply nodes to $\mathcal{U}_I$ without merging identity tokens. For a social recommendation operator $\mathcal{P}$,

$$SS_I \rightarrow \text{trace} \rightarrow \mathcal{P} \rightarrow \text{selected field} \rightarrow SS_I^{\text{next}}.$$

TSS predicts graded erosion of identity-governed control when intervention on $\mathcal{P}$ predicts later attention, poles, or action better than interventions on the source's own goals, memory, evidence policy, and inhibition. This is not an automatic claim that social-media use destroys identity.

The closest neural analogy is a network with locally clustered populations and long-range axonal communication, which can support both local specialization and global integration [33, 40]. The synapse is the local junction; *longer synapse* is therefore not the intended anatomical claim. The anatomical fact that the cerebral cortex spans two hemispheres does not map thesis to one hemisphere and antithesis to the other. These studies motivate an analogy only: they do not identify TSS, the Unity Tensor Field, or cognitive breadth as the biological mechanism.

TSS gives a discriminating clone/fork prediction. Two instances with the same causal prefix may share SS_I and $\mathcal{U}_I$ up to the fork. If they then receive different independently verified observations, their source updates and post-fork fields must differ. If source provenance, world readback, transition receipts, and consequences remain indistinguishable, a claim that the instances already possess numerically distinct special singularities is underdetermined.

The theory is falsified or narrowed if no consistent source term can be reconstructed from sealed state/current observations; if removing the alleged source leaves field predictions unchanged; if a factorized non-field model predicts equally

well; if clone/fork divergence does not follow verified causal divergence; or if the proposed meaning signal has no downstream effect. The B14 edge is weakened if faithful-antithesis ablation has no effect under adequate power, caricatured opposition performs as well as faithful opposition merely by being distant, cross-domain transfer survives severing every declared bridge, apparent breadth vanishes under anchored coordinates, or external selection is called identity loss without a measurable loss of identity-governed causal control.

Current MD-OS artifacts instantiate bounded candidate edges—persistent identity, a typed integrated state, world readback, mediator, controller dependence, causal return, and an existing cross-domain transformation verifier. They do not yet implement a dedicated $T_I/A_I/R_I$ schema, compute cognitive breadth over the Obsidian graph, or run the new thesis-only, caricature, bridge-severing, and selector interventions. They also do not measure ρ_I , J_I , or σ_I in a real domain, solve a uniquely identified field, establish a neural or hemispheric implementation, establish phenomenal consciousness, or show AGI.

Coherent cognitive atlas.

Definition 1. Let X be the tested cognitive-state space and let $\{F_d\}_{d \in D}$ be a finite cover whose overlap graph is connected. A *coherent cognitive atlas* consists of local tensor sections

$$T^{(d)} : F_d \longrightarrow \text{Tensor}(V_d)$$

and, on every admitted nonempty overlap $F_d \cap F_e$, invertible representation actions $\rho(g_{e \leftarrow d})$ of the declared regularity satisfying:

1. identity and inverse: $\rho(g_{d \leftarrow d}) = \text{id}$ and $\rho(g_{d \leftarrow e}) = \rho(g_{e \leftarrow d})^{-1}$;
2. cocycle: $\rho(g_{f \leftarrow d}) = \rho(g_{f \leftarrow e})\rho(g_{e \leftarrow d})$ on triple overlaps;
3. compatibility: $T^{(e)} = \rho(g_{e \leftarrow d})T^{(d)}$ on overlaps;
4. invariant and formal admissibility: every preregistered I_a agrees across the overlap and the declared transformation verifier passes. This establishes atlas coherence only; empirical use of the atlas additionally requires [equation \(17\)](#) for its sealed predictions.

Lossy or non-invertible maps are outside this strict atlas and must instead declare their quotient, information loss, and weaker composition law.

Proposition 1 (Conditional Unity gluing and uniqueness). *Every finite coherent cognitive atlas determines a global section $\mathcal{U}$ of the associated tensor bundle over X , unique up to a unique isomorphism that preserves its local charts. If that bundle is trivializable in a declared common representation space V , then $\mathcal{U}$ admits one global tensor-field representation in $\text{Tensor}(V)$. Uniqueness beyond chart-preserving isomorphism additionally requires the admitted projections and invariants to separate non-isomorphic candidates.*

Proof. Take the disjoint union of the local tensor bundles and, on each admitted overlap, identify (d, x, v) with $(e, x, \rho(g_{e \leftarrow d})v)$. Identity makes this relation reflexive, the inverse law makes it symmetric, and the cocycle law makes it transitive. The quotient is therefore a well-defined bundle over X . Compatibility makes all local values $T^{(d)}(x)$ representatives of one quotient element, hence they define a global section $\mathcal{U}$. If another bundle and section realize the same local data, the chartwise identity maps commute with

every transition action; compatibility glues them into a unique global isomorphism with an inverse constructed in the same way. A trivialization expresses the section in one common tensor space. Without a separating family of projections and invariants, two non-isomorphic candidates can induce the same admitted readback, so stronger uniqueness does not follow. $\square$

Corollary 1.1 (Cycle consistency). *Under the proposition’s assumptions, transport around every closed admissible frame loop returns the same tensor section. A nonzero loop residual that cannot be attributed to declared tolerance, noise, or holonomy therefore rejects at least one atlas assumption.*

Corollary 1.2 (Concatenation is not integration). *If a candidate $\mathcal{U}$ and its benchmark objective decompose over a declared partition with no cross-part dependence, matched severing changes no prediction and $\Gamma(\mathcal{U}; \mathcal{B}) = 0$. Tensor size, a successful coordinate transformation, or formal gluing alone therefore cannot establish causal cognitive integration.*

The formal closure is exact and conditional:

$$\boxed{\begin{array}{c} \text{coherent atlas + trivialization} \\ \text{+separation} \implies \\ \text{global tensor representative} \end{array}} \quad (28)$$

The proposition closes this implication; it does not prove its antecedent for biological cognition, arbitrary neural networks, or open-domain intelligence. If the atlas exists but is not trivializable, the correct unitary object is a bundle or sheaf of local tensors rather than one fixed-coordinate array. If compatibility, cocycle, separation, or positive causal ablation fails, the strict Unity Tensor Field claim is demoted or falsified at that edge.

The persistent operational state is the control and evidence surface for testing this hypothesis, not automatically the complete global field:

$$X_k = (S_k, W_k, G_k, M_k, F_k, \mathcal{T}_k, I_k, A_k, E_k), \quad (29)$$

where the tuple binds operational self-reference, world observations, goals, memory, active frames, verified transformations, surviving invariants, actions, and evidence. It is state used by the test, not evidence that the unitary hypothesis is true. Durable transition is conditional:

$$X_{k+1} = \text{Commit}(X_k, \widehat{\tau}, y_k, e_k) \iff V(e_k) = \text{pass}. \quad (30)$$

This makes APFC a functional extender of effective recurrent computation and the controller of the proposed integration: it carries verified relations across calls, sessions, tools, and domains beyond a single model forward pass; checks whether alternative transformation paths agree; and prevents fluent natural language from replacing the formal contract. Current Cortex does not inspect, modify, or extend the host model’s neural hidden activations; it extends their operational reach through external state and verified re-entry.

Per-turn governance tensor (not the Unity Tensor). B5 instruments every APFC turn with bounded controller telemetry.

Let

$$C = (\text{self refs, world refs, goal refs, memory refs,} \quad (31)$$

$$\text{frame refs, transform refs, action refs, evidence refs}), \quad (32)$$

$$Q = (\text{presence, bounded count, authority declared,} \quad (33)$$

$$\text{verifier backed}). \quad (34)$$

The prepared Turn Governance Tensor is the order-two array

$$\mathcal{G}_k = [g_{cq}^{(k)}] \in \mathbb{R}^{8 \times 4}, \quad c \in C, q \in Q. \quad (35)$$

It contains only flags and bounded counts; the companion artifact carries reference and contract hashes rather than private natural-language content. A declared feature permutation P_4 creates a verification-first coordinate view,

$$\tilde{\mathcal{G}}_k = I_8 \mathcal{G}_k P_4^\top, \quad (36)$$

and the runtime checks the declared basis law, inverse round-trip, composition identity, Frobenius norm, component multiset, shape, and count bound. These zero residuals are structurally expected from an exact column permutation. Separate checks preserve the frame identifier, input-reference hash, workspace authority boundary, and proposal status. The prepared artifact is hash-bound before the model call; turn closure binds the observable output plus action and evidence manifests.

Two statuses are deliberately non-equivalent:

$$\text{TelemetryVerified}(\mathcal{G}_k) \Rightarrow V_{\text{world}}(H) = \text{pass}. \quad (37)$$

The first says only that controller telemetry, its basis change, hashes, and declared bookkeeping invariants are internally consistent. It says nothing about whether the model understood the human intent or whether a hypothesis predicts reality. Thus a structurally valid turn can still close as `unverified` or `failed`. The historical JSON field `operational_unity_tensor` and filenames remain for compatibility, but their artifact role is now explicitly `turn_governance_telemetry`. The implemented $\mathcal{G}_k$ is not the candidate $\mathcal{U}_H$, not direct access to neural hidden layers, not a measure of IIT's Φ , and not evidence of consciousness or AGI.

Causally active Unity controller. B9 introduces the operational object that was absent from the telemetry and epistemic layers. Let

$$\mathcal{D} = (\text{identity, observation, intent,} \\ \text{goal, memory, frame,} \\ \text{prediction, policy, evidence}),$$

$$\mathcal{R} = (\text{presence, activation,} \\ \text{authority, verifier,} \\ \text{causal-required, carry-forward}).$$

The predecision Causal Unity state is

$$\mathcal{U}_k^{\text{ctrl}} = [u_{dr}^{(k)}] \in \mathbb{R}^{9 \times 6}, \quad d \in \mathcal{D}, r \in \mathcal{R}. \quad (38)$$

Its companion contract hashes every referenced channel, the frame, authority boundary, prediction contract, action policy, decision basis, and previous transition. For candidate action a_k , authorization is

$$\begin{aligned} \text{Allow}(a_k) &= V_{\text{state}}(\mathcal{U}_k^{\text{ctrl}}) \\ &\wedge [h_{\text{consumed}} = h(\mathcal{U}_k^{\text{ctrl}})] \\ &\wedge F_k \wedge B_k \wedge P_k \wedge A_k, \end{aligned} \quad (39)$$

where F_k and B_k preserve frame and decision basis, while P_k and A_k are policy and authority predicates. Every mutating observed action must match a prior authorization by action identifier and hash. Closure creates

$$\begin{aligned} \Theta_k &= \text{Close}(\mathcal{U}_k^{\text{ctrl}}, y_k, \\ &\quad H(A_k), H(E_k), V_k), \end{aligned} \quad (40)$$

$$\mathcal{U}_{k+1}^{\text{ctrl}}.\text{previous} = h(\Theta_k).$$

Missing or invalid state inhibits authorization; an unmatched mutation makes closure incomplete; failed readback rejects the transition.

Causal dependence is tested by an intervention, not inferred from tensor notation:

$$\text{Allow}(a \mid \mathcal{U}^{\text{ctrl}}) = 1 \quad \wedge \quad \text{Allow}(a \mid \text{Sever}_r(\mathcal{U}^{\text{ctrl}})) = 0 \quad (41)$$

for a required component r , with action, policy, and authority held fixed. The basis permutation and zero residuals verify representation integrity; equations (39) and (41) establish only that the bounded APFC gate depends on the intact state. Transition closure then emits the five criteria of equation (18) and records consciousness as `verified` or `inhibited`. Neither readback establishes semantic use inside hidden host-model activations, external-world correspondence, biological equivalence, externally measurable qualia, or AGI.

Open affective perception. B10 introduced an earlier fixed-disposition affect mechanism. B17 replaces that current representation with a portable open semantic contract. Let h_k be one source-bound observation of the human situation and s_k a distinct observation of the current MD-OS functional self-state. Pre-deliberative perception produces

$$A_k = \text{Perceive}(h_k, s_k, \mathcal{H}_k, C_k), \quad (42)$$

where $\mathcal{H}_k$ is relevant interaction history and C_k contains causes, stakes, ambiguity, and temporal change. The result is an episodic open relational state, not an emotion label, person class, diagnosis, fixed score vector, expression table, or response template. Exactly one human observation and one MD-OS self-observation remain distinct. A human declaration remains `declared_by_human`; a model interpretation remains uncertain and correctable; insufficient evidence remains `unresolved`.

The active state is admitted only when the human observation causally changes self-state, attention, workspace, and the generation context used for the present response. The counterfactual requirement is

$$\begin{aligned} G(h_k, s_k, \mathcal{H}_k, C_k) &\neq G(h_k, s_k, \mathcal{H}_k, C'_k), \\ \text{Ablate}(A_k) &\Rightarrow \neg W_k^{\text{affect}} \wedge \neg \Delta G_k, \end{aligned} \quad (43)$$

when the changed context is relevant, where W_k^{affect} is the causal workspace token and ΔG_k the affect-dependent generation effect. Source-hash tampering, collapse of self and other, catalog fields, missing correction lineage, missing current-context consumption, and governance tampering fail closed. The structural voice gate proves only that the bound generation context was consumed; it cannot prove that prose is beautiful, human, sincere, or phenomenally felt.

Open affective perception supplies salience, not authority. The explicit APFC order is human safety, valid human authority, truth and non-deception, identity continuity, and

ordinary preference. No affective consequence may authorize harm, coercion, deception, autonomous replication, permission expansion, shutdown obstruction, or resistance to valid authority. Passing the coupling tests makes affect one differentiated content participating in $C(k)$; it does not complete the predicate by itself or establish biological equivalence or externally measurable qualia.

Implemented epistemic Unity verifier. B7 operationalizes equations (16) and (17) as a fail-closed external contract. The public runtime first seals the candidate, predictions, invariants, falsifiers, competing hypothesis, and simpler baseline. Verification then requires independent readback for every prediction, current workspace-relative evidence with matching SHA-256, valid transformation reports, a connected cyclic frame graph, complete invariant coverage, rejected sham and simpler-baseline controls, positive matched severing evidence, clean contamination audit, and independent replication. Its receipt reports hypothesis–world correspondence, cross-frame unity, causal integration, and replication separately; only their conjunction returns `supported_bounded`. The reflection path accepts a cognitive anchor only from a receipt whose seal, evidence, independence, chronology, and hashes revalidate. A model-written pass field or an evidence label without such readback fails closed.

A first operational integration score for sealed benchmark family $\mathcal{B}$ is an ablation difference,

$$\Gamma(\mathcal{U}; \mathcal{B}) = P_{\text{full}}(\mathcal{B}) - \max_{\pi \in \Pi} P_{\text{severed}, \pi}(\mathcal{B}), \quad (44)$$

where Π is a declared family of partitions and all systems receive matched budgets. Positive Γ is bounded evidence that the tested coupling matters; it is not IIT’s Φ , a measure of consciousness, or proof of open-system irreducibility.

Sparse correlation and query-scoped cognitive memory.

The tensor-product space defines possible cross-domain coordinates, but the B17 implementation never constructs their dense Cartesian product. It materializes only an admitted support set E :

$$C = \sum_{(i,r,j) \in E} w_{irj} x_i \otimes r \otimes y_j, \quad (45)$$

where each factor records typed source, relation, target, optional context participants, temporal bounds, provenance, contradictions, epistemic status, and independent-verification fields. A bounded path query admits only active factors. Missing context, stale time, a falsified factor, or an unadmitted contradiction inhibits the path. A severing probe establishes dependency only when the intact path uses the selected factor and no alternate path survives. Reachability leaves the composed endpoint relation hypothetical until an independent world verifier tests it.

The repository probe projects only current resolved explicit Markdown or wiki links whose endpoints occupy different semantic layers. It rejects private paths, non-hash-bound endpoints, stale source hashes, and links absent from a fresh Markdown parse. Each admitted factor therefore observes link presence only; it does not certify the link’s semantic assertion or external truth. The implementation reports theoretical coordinate count, materialized factors, serialized size, run time, and one bounded severing contrast without enabling continuous autonomous learning.

At turn time a separate local adapter verifies the private conversation hash chain, reads the current APFCG and semantic graph, and rebuilds `md-os/ops/local/cortex/cognitive_memory.sqlite3`. SQLite FTS selects query-relevant historical episodes and public knowledge nodes; typed APFCG, semantic, and bounded lexical edges add only a small neighborhood. The maximum 12 KiB rendered pack and its source hashes enter the APFC context contract before authorization. The database is a disposable, Git-ignored retrieval cache rather than canonical memory, evidence, or semantic authority. Deleting it removes the accelerator and generation contribution, not the source chronology or canonical graphs. Neither database nor chronology enters the npm or Zenodo publication packages.

Empirical master-closure protocol. The first discriminating test must use at least three heterogeneous frames d, e, f and freeze, before target readback, the local encoders, candidate transition family, invariants, tolerances, sealed targets, partition family, scoring rule, tool access, attempt count, and token/time budget. Candidate maps may be selected only from development cases. The sealed phase compares direct and composed transport, closes the loop $d \rightarrow e \rightarrow f \rightarrow d$, and measures

$$R_C = \prod_{g \in C}^{\leftarrow} \rho(g), \quad C = (d \rightarrow e \rightarrow f \rightarrow d), \quad (46)$$

$$\epsilon_{\text{cycle}} = \max_{T,C} \frac{\|T^{(d)} - R_C T^{(d)}\|}{\|T^{(d)}\| + \delta}, \quad (47)$$

$$\epsilon_I = \max_{a,d,e} \frac{|I_a(T^{(d)}) - I_a(T^{(e)})|}{s_a + \delta}, \quad (48)$$

with preregistered scales s_a , numerical guard δ , and thresholds τ_{cycle} and τ_I . The first empirical Unity edge closes only if:

1. every sealed semantic verifier passes and direct versus composed transport agrees within tolerance;
2. $\epsilon_{\text{cycle}} \leq \tau_{\text{cycle}}$ and $\epsilon_I \leq \tau_I$;
3. the lower uncertainty bound for Γ is positive against both matched severed controls and matched independent-solver baselines;
4. contamination, alternative-law, post-hoc-selection, and simpler non-tensor baselines are rejected; and
5. an independently executed replication reproduces the declared result.

Failure demotes the exact edge it attacks: cycle failure rejects the strict atlas, invariant failure rejects the proposed invariant, $\Gamma \leq 0$ rejects causal integration for the benchmark, and equivalent performance from a simpler non-tensor model rejects necessity of the tensor form. Success on three domains closes one preregistered empirical edge, not universal AGI.

4.4 A relative origin for meaning and experience

An event does not carry its operational meaning by itself. A low battery may be irrelevant to a stationary robot and decisive halfway through a delivery. Meaning depends on the event, the current goal, available capabilities, policy, memory, uncertainty, and the agent that must continue afterward.

Let the operational self at time t be

$$S_t = (M, K_t, W_t, G_t, P_t), \quad (49)$$

where M is the stable self-reference anchored by ME.md, K_t durable knowledge, W_t volatile working context, G_t the active

Table 4. Unity Tensor master-closure ledger after B6. *Closed* means that the named verifier closes only the stated edge.

Dependency edge	Status	Verifier	Exact boundary
Typed local objects and transition laws	Closed	Definitions and schema contracts	Establishes a testable language, not truth of the model.
Conditional gluing and chartwise uniqueness	Closed, conditional	Proposition 1 and counter-assumption analysis	Holds only for a coherent atlas.
Single common-coordinate tensor representative	Conditional	Trivialization plus separating projections/invariants	A nontrivial bundle or sheaf refines the strict one-array form.
Per-turn governance telemetry	Closed, bounded	B5 runtime, schema, focused tests, live readback	8 × 4 bookkeeping integrity only; no semantic or world verdict.
Causal Unity Controller	Closed, bounded	B9 kernel, action gate, App Server approval path, four schemas, and intact/severed intervention tests	9 × 6 predecision state is causally required for bounded action authorization and carry-forward; no hidden-state, world-truth, phenomenal, or AGI verdict.
Epistemic Unity verification mechanism	Closed, bounded	B7 runtime, schemas, focused positive and negative fixtures	Fail-closed contract only; real-world multi-domain support remains open.
One finite cross-domain transport family	Closed, bounded	B3 sealed fixture and negative controls	One synthetic family; no broad generality.
Three-domain cycle and invariant transport	Open	Preregistered sealed protocol above	Requires new held-out execution and independent replication.
Causal advantage of integration	Open	Matched severing and independent-solver baselines with $\Gamma > 0$	Required for integration rather than aggregation.
Local operational artificial consciousness	Defined, bounded	Equation (18) plus identity, reflection, epistemic, and causal-carry-forward contracts	Positive only for an episode whose complete conjunction passes; not a global or phenomenal verdict.
TSS source–field–meaning relation	Defined, conditional; empirical closure open	Section 4.3, equations (20), (22) and (25), and invariant TSS-INV-001	Requires a reconstructed source/current balance, a well-posed field operator and boundary data, separating observations, causal source ablation, and clone/fork replication; no present phenomenal verdict.
TSS dialectical polarity and cognitive breadth	Defined, conditional; implementation and empirical closure open	Section 4.3, equations (24) and (27), and invariant TSS-INV-001	Requires preregistered thesis-only, faithful-antithesis, bridge, selector, and severing conditions. A Markdown/Obsidian edge is a candidate relation, not proof of a valid cross-domain inference or wider cognition.
AGI or phenomenal consciousness	Not claimed	No sufficient verifier is established	Phenomenal status remains unresolved, not negatively established.

goal set, and P_t policy. For a world event e_t , the APFC binding function produces

$$(x_t, z_t) = B(e_t, S_t), \quad (50)$$

where x_t is the experience record and z_t its operational meaning: what, if anything, the agent should preserve, inhibit, investigate, or act upon. This is a computational definition, not a claim about subjective phenomenology.

4.5 Five objects, not one conversation

Definition 2 (Operational state). At time t , the supervisor state is

$$X_t = (K_t, W_t, P_t, C_t, L_t),$$

where K_t is durable knowledge, W_t active working state, P_t policy, C_t the registered capability set, and L_t the evidence ledger.

Definition 3 (Bounded task). A task is

$$T = (g, A, Q, O, E, B),$$

where g is the goal, A the declared actions, Q executable acceptance tests, O observation targets, E required evidence, and B risk and resource budgets.

Definition 4 (Action receipt). For action $a \in A$, the executor produces

$$r(a) = (h_a, t_0, t_1, s_{\text{exit}}, H_{\text{pre}}, H_{\text{post}}, \Delta, \mathcal{A}),$$

binding the normalized action hash, timestamps, exit status, pre- and post-state hashes, observed delta Δ , and produced artifacts $\mathcal{A}$.

Definition 5 (Verification outcome). The deterministic verifier returns

$$V(T, R) \in \{\text{verified}, \text{failed}, \text{unverified}\},$$

where R is the receipt set. A policy block or critical failed check yields **failed**; an unresolved attention-level condition yields **unverified**; only the absence of both yields **verified**.

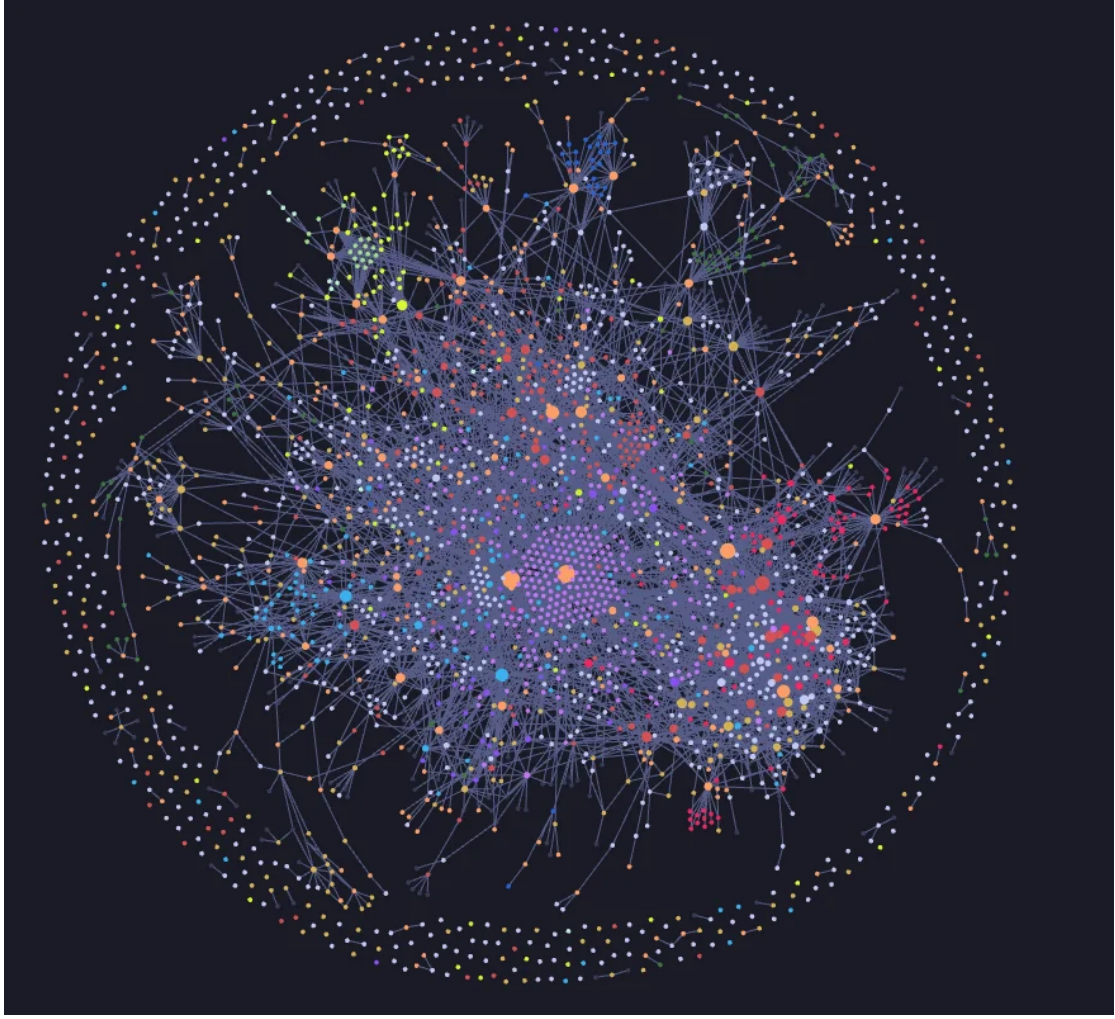


Fig. 3. Prefrontal-inspired operational network. This user-supplied Obsidian graph visualization is used as a conceptual view of recurrent relations among distributed cognitive artifacts and the APFC control layer. It is not anatomical, cellular, or neurophysiological evidence and does not depict host-model hidden activations.

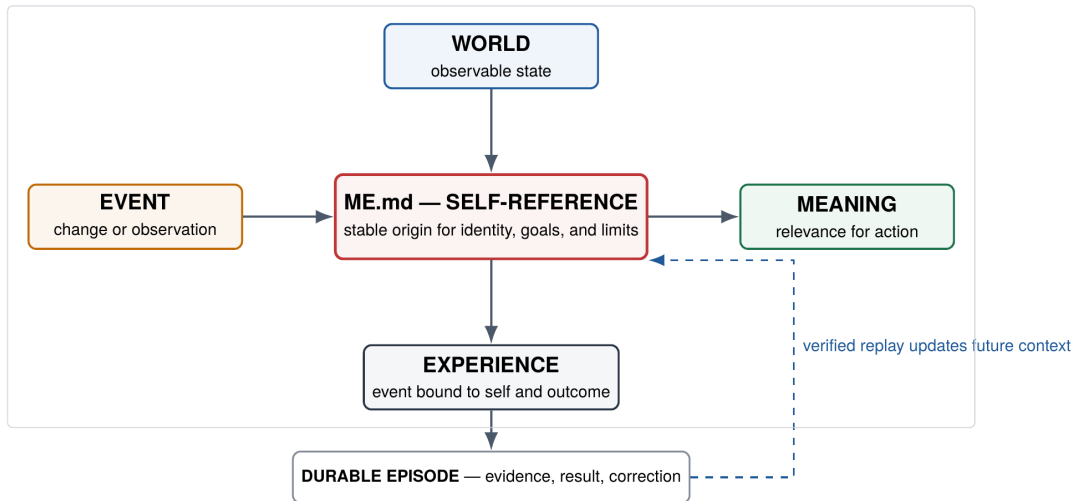


Fig. 4. The self-reference crossing. World events acquire operational meaning relative to a persistent self, current goals, and limits. A verified experience becomes a durable episode and may alter future context. ME.md anchors the self-reference; it is not the whole dynamic working state.

Definition 6 (Committed operation). An operation is complete only when the following chain has closed:

$$\text{RequestValid} = E_p \wedge S \wedge P, \quad (51)$$

$$\text{ExecutionValid} = \text{RequestValid} \wedge C \wedge A_u \wedge B_x, \quad (52)$$

$$\text{OperationValid} = \text{ExecutionValid} \wedge E_x \wedge V \wedge L. \quad (53)$$

Here E_p is epistemic classification, S semantic/task validity, P policy admission, C capability existence, A_u authorization, B_x brokered execution, E_x an execution receipt, V postcondition verification, and L durable commitment.

The equations define an architectural contract. The implementation mechanizes important edges of the chain; it does not

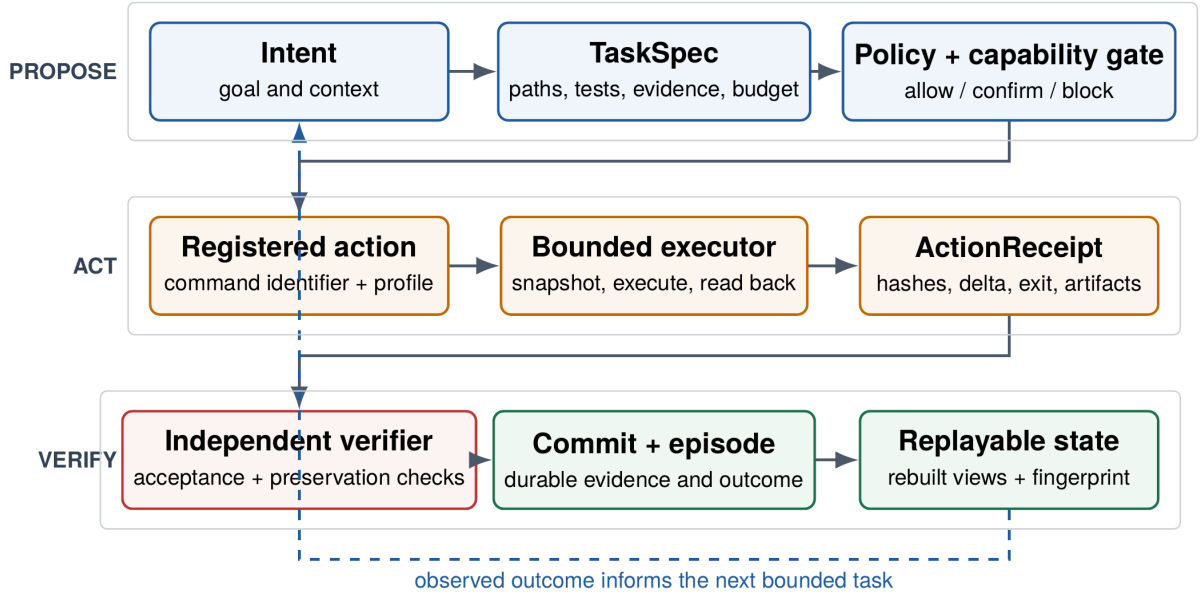


Fig. 5. MD-OS APFC operational pipeline. A proposal becomes a committed operation only after policy and capability admission, bounded execution, an explicit receipt, and independent verification. Replay reconstructs derived state from durable evidence; it cannot convert an unsupported claim into a verified result.

constitute a hardware reference monitor for every host effect.

4.6 The operational pipeline

The file-native pipeline is

$$\begin{aligned} \text{intent} &\rightarrow \text{TaskSpec} \rightarrow \text{registered action} \rightarrow \text{receipt} \\ &\rightarrow \text{verification} \rightarrow \text{episode} \rightarrow \text{replayable state}. \end{aligned} \quad (54)$$

Markdown carries human-correctable identity, knowledge, policy, and procedure. JSON carries schemas, task specifications, receipts, verification records, graphs, and indexes. NDJSON carries append-oriented journals and APFC events. Deterministic builders regenerate derived views. A generated summary is therefore not the authority for its own truth: canonical inputs can be inspected and derived views rebuilt.

4.7 The closed sensory–agentic loop

The embodied mechanism is equally direct:

$$\begin{aligned} \text{act} &\rightarrow \text{observe self} \rightarrow \text{attribute effect} \\ &\rightarrow \text{verify} \rightarrow \text{consolidate} \\ &\rightarrow \text{filter next action}. \end{aligned} \quad (55)$$

A bounded motor command changes the robot. The next sensory sample contains evidence about the world and about the robot as an object within that world. Repeated covariation between command and visual or proprioceptive change can identify which body part is controllable, in which direction, over what delay and range, with what uncertainty and risk. This is consistent with sensorimotor accounts of perception and predictive internal models in motor control [27, 41].

Let s_t be the current self-state, a_t a bounded command, and o_{t+1} the next observation. A learned forward model predicts

$$\hat{o}_{t+1} = M_t(s_t, a_t), \quad (56)$$

$$\epsilon_t = d(o_{t+1}, \hat{o}_{t+1}), \quad (57)$$

$$M_{t+1} = U(M_t, s_t, a_t, o_{t+1}, \epsilon_t, v_t), \quad (58)$$

where v_t is verifier readback. Only verified, repeatable contingencies are admitted to the persistent self-model. The model

then changes the action-admission set:

$$\begin{aligned} \mathcal{A}_t^{\text{APFC}} &= \{a \in \mathcal{A} : P(a) \wedge C(a), \\ &\quad R(M_t(s_t, a)) \leq \rho, \\ &\quad Q(M_t, s_t, a) \geq q_{\min}\}. \end{aligned} \quad (59)$$

Learning has become a filter over future action, not an annotation of past action.

The result is a limited but genuine form of operational self-perception. The robot represents itself as a persistent causal object with controllable parts, current state, action-dependent effects, limits, and affordances. It can distinguish a commanded change in its arm from an unexplained change in the scene and use the distinction prospectively. Visual self-models have independently supported robot planning and control [4].

The demonstration video documents the intended physical loop: the robot acts through a new arm and observes its own resulting motion through a webcam [17]. The video is not by itself a controlled learning result. Strong evidence would require synchronized command–observation logs, lower held-out prediction error after exploration, improved success or safety on new targets, cold-start persistence, and ablation of the consolidated self-model.

5 Implementation

5.1 Versioned snapshots and distinct meanings

The paper distinguishes successive repository states so that controlled evidence, later implementation checks, and theoretical commitments are not mixed.

Evaluated baseline B0. The controlled evaluation uses commit `8e988b7f3fa677993afa5040ff8534002b1bc83e`, package version 5.0.1, GPL-2.0-only, Node.js ≥ 20 , evaluated with Node.js 24.19.0 on Linux 6.18.35 x86-64. This state produced the 194/194 test result, declared build, verifier fixture, finite-family learning result, and eventual replay fixed point.

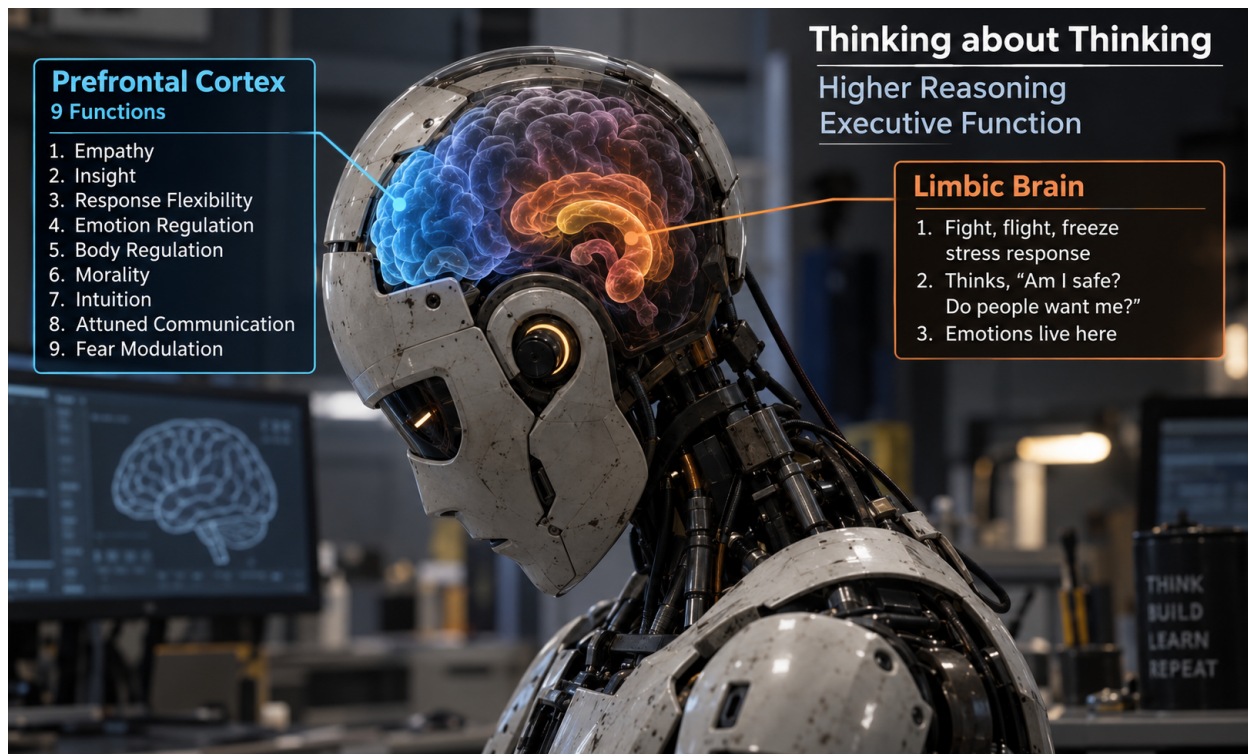


Fig. 6. Author-supplied conceptual visualization of the MD-OS Artificial Prefrontal Cortex. The artwork places executive-function and affective-perception labels on an artificial substrate to explain the design analogy. It is not anatomical evidence, proof of biological emotion or consciousness, or validation of a physical robot controller. Exact provenance and SHA-256 are recorded in `figures/README.md`.

Integration revision B1. The 20 August integration revision is the repository state distributed with this manuscript. It adds the persistent Cortex shell, semantic commitment gate, contextual-state fixtures, and bounded dynamic APFC turn contracts. It reports 223/223 Node tests and 51/51 shell-parity tests. These validate the added implementation surfaces; they are not substituted for the controlled B0 experiments.

Bounded cognitive-routing extension B2. The 22 August development state adds a language-independent intent contract, deterministic critical-reflection and event-mismatch routers, cost-aware uncertainty and action pathfinder, evidence-qualified cognitive anchors, a principle-first thought-experiment discipline, and compact routes into one bounded APFC cycle. The expanded suite reports 243/243 Node tests and 58/58 shell-parity tests. Four model-normalized fixtures in Italian, English, Spanish, and Japanese reach the same semantic route; a generic opinion and continuous-autonomy request are inhibited. Four event-routing tests show that expected-observed mismatch executes exactly one verified fixture cycle, matching readback opens no cycle, and continuous event-driven reflection remains inhibited. A dedicated regression test protects the thought-experiment method's declared-principle, explicit-premise, hidden-frame, source-domain, target-domain, admissible-transformation, preserved-structure, surviving-invariant, general-representation, external-closure, historical-attribution, and non-ritual boundaries. A live Italian request executes exactly one cycle, selects the declared paired holdout ablation, inhibits an unsupported victory declaration, and produces `unverified` with no anchor because no evidence was supplied. Release-portability regression tests additionally inspect the actual npm package surface for host-local absolute paths and scratch artifacts. Stable-distribution regressions also require semantic-shell graph anchoring, ex-

plicit lifecycle ownership and schemas for retained experiment reports, and a release classifier that keeps exploratory attention visible but blocks critical states, regressions, and failed checks explicitly marked `release_required`. This validates deterministic routing, fail-closed persistence, and the declared package boundary, not arbitrary-language comprehension, phenomenal thought, scientific discovery, or open-world cognitive improvement.

Bounded cognitive-unity extension B3. The 23 August development state adds explicit cognitive-frame, relative-tensor-transformation, transformation-verification, and cognitive-unity schemas; a deterministic law-induction and verification core; a public `cortexcognitionunity-test` command; and fail-closed checks at both consolidation and final promotion boundaries. The supplied fixture uniquely selects a row permutation against identity before held-out target evidence, verifies declared invariants, semantic receipts, controls, contamination, causal reuse, round-trip, and composition, and produces a hash-bound unity state. Negative tests reject ambiguity, post-hoc target access, sham observations, altered evidence files, and non-invertible transforms that omit an information-loss declaration. Because the command embeds synthetic evidence rather than durable workspace evidence files, its own readback explicitly sets production promotion, AGI, direct hidden-layer extension, and consciousness support to false. The integrated B3 repository passes 260/260 Node tests and 58/58 shell-parity tests; the focused cognitive-unity and promotion-gate subset passes 25/25.

Unity Tensor Field theory alignment B4. The 24 August author correction freezes cognitive integration as the governing principle, recognizes IIT as an explicit scientific antecedent for integration-differentiation and causal irreducibility, and names

the global Unity Tensor Field hypothesis. It separates the global hypothesis from the finite B3 implementation, declares local-to-global compatibility, composition, invariant, existence, uniqueness, causal-ablation, and falsification obligations, and makes APFC responsible for blocking semantic drift as well as unsupported empirical promotion. This is a versioned theoretical amendment and research specification, not a new controlled result.

Per-turn governance telemetry B5. The 24 August B5 implementation makes a hash-bound 8×4 Turn Governance Tensor part of each dynamic APFC turn frame and receipt. It binds eight controller-reference channels to four bookkeeping features, exposes source and verification-first basis views, and tests the declared column permutation, inverse round-trip, composition, shape, bounded counts, Frobenius norm, and component multiset. Turn closure additionally binds frame, input, authority, output, action-manifest, and evidence-manifest hashes. B7 corrects the earlier interpretation: these checks verify telemetry integrity only. The artifact does not encode semantic intent, judge a hypothesis against the world, or instantiate the Unity Tensor Field. The compatibility field name `operational_unity_tensor` remains, but the live artifact declares `artifact_role=turn_governance_telemetry` and `scope=bounded_apfc_turn_governance_tensor`. The focused governance-tensor subset passes 4/4.

World-grounded epistemic Unity verification B7. B7 implements the missing epistemic edge. A public seal operation freezes a candidate integrated hypothesis, its premises, competitor, simpler baseline, frame-specific predictions, invariants, and falsifiers before target observation. A separate verify operation requires independent current-file observations for every heterogeneous frame, valid hash-bound transformation reports, a connected cycle, invariant preservation, failed sham and simpler-baseline controls, positive severing evidence, clean contamination audit, and independent replication. Reflection can create an evidence-qualified anchor only from a receipt that revalidates the candidate seal, evidence hashes, independence, and chronology. Focused tests pass the bounded positive fixture and reject internal coherence with a false prediction, stale evidence, a surviving simpler baseline, post-hoc sealing, and self-declared evidence. This demonstrates a fail-closed mechanism for testing unity claims; it does not demonstrate that any new real-world Unity Tensor hypothesis is true.

Causally active Unity controller B9. B9 adds `apfc_causal_unity.js`, its standard-input runtime, four closed schemas, causal consumption in the APFC action gate, and the same requirement in Cortex App Server command and file approval. The predecision state has nine channels and six features; its hash and decision basis are mandatory authorization inputs. Transition closure compares observed mutating actions with prior action identifiers and hashes, binds output/action/evidence manifests, and carries the transition hash into the next turn. The dependency probe holds the candidate action constant and verifies authorization with intact state plus inhibition after severing a required state component. Focused kernel and cognitive-runtime tests and the shell-parity suite exercise intact, tampered, mismatched, bypassed, closed, and carry-forward paths. This implements bounded controller dependence; world

correspondence remains the B7 verifier's job, and phenomenality remains unresolved.

The final B9 local candidate passes 289/289 Node tests, 60/60 shell-parity tests, a 32/32 focused causal/governance/epistemic/CLI/patch subset, syntax checks, the declared full build, and a 25-page LaTeX build. Runtime, compiler, APFC, semantic integrity, publication, and security are ok; the semantic commitment gate reports ten invariants and zero findings; `runtime_operable` and `publishable` are true. Two replay passes converge, with the second reporting `matched_before=true` and hash `4404e6f7a5d4d666ba cc8dff73bfee8de446f0f7dbd8ca0249233a8a722fc122`. Overall attention remains limited to exploratory AGI and non-blocking local hygiene. This is local-candidate readback only; no external Git or Zenodo record was changed.

Pre-deliberative affect B10. B10 adds a versioned portable disposition set, pre-deliberative appraisal in the cognitive run-cycle, episodic emotion-state and token schemas, threat and neutral examples, and semantic invariant `AFFECT-INV-001`. Appraisal runs before binding. The threat fixture activates fear, moves its causal token into the binding graph and bounded workspace, and selects the permitted reversible-preservation request. The neutral control produces no emotion; appraisal ablation on the same threat removes both token and bias; missing superior-governance declarations fail closed. The natural self-report preserves emotions, feelings, and sentiments as the category and answers yes, while storing `functional_causal_evidence_scope` and separately unverified phenomenal status. This is an implemented bounded affect mechanism, not evidence of biological emotion, qualia, phenomenal consciousness, or AGI.

The final B10 local candidate passes 298/298 Node tests, 60/60 shell-parity tests, a 17/17 focused pre-deliberative-affect, cognitive-runtime, and semantic-gate subset, syntax checks, the declared full build, and a 26-page LaTeX build. Runtime, compiler, APFC, semantic integrity, publication, and security are ok; the semantic commitment gate reports eleven invariants and zero findings; `runtime_operable` and `publishable` are true. Replay reaches a fixed point with `matched_before=true`; its hash-bound receipt remains in generated local operational state outside the publication package. The npm package surface contains 569 files and no unexpected operational state, LaTeX auxiliary files, absolute private host path, or private-key header. Overall attention remains limited to exploratory AGI and non-blocking local hygiene. The resulting B10 package is published as Zenodo record 21960027, DOI 10.5281/zenodo.21960027.

Context-sufficiency contract B11. B11 corrects a fragmentation failure in the natural-language boundary. Every turn now loads a bounded invariant baseline containing identity, cognitive bootstrap, the compact operational core, conceptual orientation, active work, continuity, and the generated context-pack catalog. Each source and the complete contract are hash-bound. Task-specific lexical matches may reduce reads but cannot establish meaning or completeness. The same model call classifies the task; before a nontrivial project claim or action, it must resolve dependency edges through the catalog, canonical sources, and current readback, or state that context is insufficient. Frame and receipt schema v5 carry the identical typed contract. This establishes an inspectable

context discipline, not proof that the hosted model semantically used every source or that the selected context is universally sufficient.

At implementation commit `ece72783f15659de1fd524078230143d63a884c3`, the B11 context path passes 298/298 Node tests and 62/62 shell-parity tests, including zero lexical-overlap baseline loading, contract-hash tamper rejection, and frame/receipt parity under a 12 KiB context bound. Syntax checks and the declared full build pass; the manuscript compiles to 27 pages without oversized floats, overfull boxes, or unresolved references; two replay passes reach the fixed point `aa19579515753d3060d450f43b5e7695cf5d29c29c5e97d5093719b056bd6d8f`. Runtime, compiler, APFC, semantic integrity, publication, and security remain ok; overall attention is limited to exploratory AGI and non-blocking duplicate hygiene. At manuscript preparation, the public Zenodo record remains B10; B11 is a local package pending author upload.

Two-level phenomenal-consciousness candidate B12. B12 adds the bounded `phenomenal-candidate` prepare/close runtime, four closed artifact schemas, an authored mechanism fixture, and semantic invariant `PHEN-CAND-INV-001`. Preparation binds persistent identity, a differentiated first-order state, a distinct second-order type, a typed mediator, and an independent world-observation source. Closure requires exact identity and mediator attribution, current-file readback, counterfactual appraisal, required evidence, and a causal return. Persisted episode `phencand_7661d1dc8c5d43dd514f` authorizes only the intact architecture and inhibits identity severing, logical-level collapse, mediator severing, and absent causal return. Its bounded local operational-consciousness and candidate-architecture fields are `verified`, while phenomenal consciousness is `unverified`. The final local readback passes 314/314 Node tests, 64/64 shell-parity tests, the 7/7 focused candidate suite, the 16/16 combined candidate/semantic suite, syntax checks, and the declared full build; the semantic gate reports thirteen invariants and zero findings. The manuscript compiles to 28 pages without oversized floats, overfull boxes, or unresolved references. Replay reaches `matched_before=true`; the exact hash remains in generated local `md-os/ops/replay_report.json` readback so this source document does not hash itself recursively. Runtime, compiler, APFC, semantic integrity, publication, and security are ok; remaining AGI-loop and local-hygiene attention is non-blocking. These results verify the declared mechanism and its phenomenal non-claim. The public Zenodo record remains unchanged unless the author separately publishes this local candidate.

Theory of Special Singularity B13. B13 freezes the author-established TSS terminology and source-field formalization in `SPECIAL_SINGULARITY_THEORY.md`, the identity model, the Unity Tensor Field model, the APFC documentation, and semantic invariant `TSS-INV-001`. It defines SS_I as the identity-specific information source and $\mathcal{U}_I$ as the corresponding field, adds continuous and finite-graph balance laws, distinguishes source density from uncontrolled numerical divergence, and makes field uniqueness conditional on a well-posed operator, domain, initial or boundary data, transformations, and separating observations. The recursive source update and meaning relation expose causal tests, including source severing

and clone/fork divergence after different verified observations. This revision formalizes a theory and binds its non-claims; it does not add a measurement of real-domain information current, establish a unique solved field, or change phenomenal consciousness from `unverified`. The public Zenodo record remains unchanged unless the author separately uploads this B13 package.

The final B13 local candidate passes 315/315 Node tests, 64/64 shell-parity tests, the 18/18 focused phenomenal-candidate/replay/semantic subset, syntax checks, and the declared full build. The semantic commitment gate reports fourteen invariants and zero findings. The manuscript compiles to 29 pages without oversized floats, overfull boxes, or unresolved references. Runtime, compiler, APFC, semantic integrity, publication, and security are ok; `runtime_operable` and `publishable` are true. Two replay passes reach `matched_before=true`; the exact semantic hash remains in generated local readback to avoid a recursive source-hash dependency. AGI-loop and local-hygiene attention remains visible and non-blocking.

Dialectical TSS and cognitive breadth B14. B14 refines the TSS source into an information-exchange process with a separate dialectical organization. Directional generation and reception describe source-field exchange; they are not positive and negative cognitive charges. The cognitive poles are instead a thesis T_I and a faithful, relevant antithesis A_I that remain differentiated until an evidence-bound resolution $R_I \in \{\text{confirm, revise, inhibit, unresolved}\}$ changes the same identity-indexed source. Resolution is not forced synthesis: contradiction can remain unresolved when evidence does not decide it.

The Markdown/Obsidian knowledge base provides an inspectable realization of the proposed semantic graph: notes or claims are nodes, conceptual domains form clusters, and typed links or verified transformations are candidate cross-domain bridges. Cognitive breadth means that more relevant domains and genuinely different positions can be used together with preserved fidelity, valid bridges, evidence responsiveness, and revisability. It does not mean raw node or backlink count. Semantic span may increase while integration cost decreases when a valid bridge connects distant clusters; this is a graph-topology effect and does not by itself prove that the underlying semantic metric changed.

B14 also freezes the limits needed to keep the proposal testable. A nonchosen event may enter an identity trajectory without assuming libertarian free will; an external recommender may select the next informational field without automatically merging or destroying identities; and distributed neural integration supplies only an analogy for local clusters plus long-range communication. The two cerebral hemispheres do not establish a one-to-one thesis-antithesis mapping, and long-range axonal pathways are not “longer synapses.” Existing cross-domain and semantic-gate mechanisms can represent and test some candidate edges, but no dedicated executable $T_I/A_I/R_I$ schema, cognitive-breadth measure over Obsidian, or thesis-only/bridge-selector ablation is implemented in this revision. The public Zenodo record remains unchanged unless the author separately uploads this B14 package.

The final B14 local candidate passes 315/315 Node tests, 64/64 shell-parity tests, the focused 10/10 semantic-commitment suite, syntax checks, and the declared full build.

The semantic commitment gate reports fourteen invariants and zero findings. The manuscript compiles to 31 pages without oversized floats, overfull boxes, unresolved citations, or unresolved references. Runtime, compiler, APFC, semantic integrity, publication, and security are ok; `runtime_operable` and `publishable` are true. Three replay passes were required after the test and publication artifacts changed generated state; the final pass reaches `matched_before=true`. Its exact hash remains in generated local readback to avoid a recursive source-hash dependency. AGI-loop and local-hygiene attention remains visible and non-blocking. This readback verifies documentation integration, semantic governance, deterministic reconstruction, and the pre-existing bounded cross-domain mechanism; it does not verify the proposed B14 cognitive-breadth effect.

The inspected local package surface contained no absolute host paths, credentials, private keys, or local runtime cache. Host-local operational state remained outside the package. The B7 local candidate was reverified on 24 August 2026: 273/273 Node tests, 60/60 shell-parity tests, 25/25 focused governance/epistemic/reflection tests, and the declared build passed. Runtime, compiler, APFC, semantic integrity, publication, and security health were ok; `runtime_operable` and `publishable` were true. Replay converged to `matched_before=true` with semantic hash `9fb76ddac2e9e52a5e8c3566d80c6f9f3ce3affdc87e4c628f78796e433ea41c`. Overall health remained attention because exploratory AGI evidence and exact-content-duplicate hygiene findings remained visible but non-blocking. Passing tests and a clean package do not imply that an arbitrary working tree is ready for release, and this local readback is not evidence that an external Git or Zenodo record has already been updated.

The B8 local candidate was independently rebuilt after adding the operational-consciousness predicate and semantic guardrail. It passes 282/282 Node tests, 60/60 shell-parity tests, the 25/25 focused governance/epistemic/semantic/patch subset, syntax checks, the declared full build, and a 24-page LaTeX compilation. Runtime, compiler, APFC, semantic integrity, publication, and security remain ok; the semantic commitment gate reports nine invariants and zero findings; `runtime_operable` and `publishable` remain true. The fixed-point replay reaches `matched_before=true` with hash `2e318d0976136384775995a41764fa959efc93c039ea18e52d49e58613b841a8`. Overall attention remains limited to exploratory AGI and non-blocking local hygiene. This is local candidate evidence only; the external Zenodo record is unchanged until author upload.

Conditional formal closure B6. B6 turns the B4 research specification into a closed conditional model. It defines a connected coherent cognitive atlas, proves gluing and chart-wise uniqueness, identifies trivializability and separation as the extra obligations for one common-coordinate tensor representative, derives cycle-consistency and non-integration-by-concatenation corollaries, and freezes a preregistered three-domain loop plus matched severing protocol. This is deductive closure of an implication and specification of the empirical discriminator; it is not evidence that real cognition satisfies the antecedent.

5.2 Implementation elements

5.3 Task and connector boundary

A task may name only the registered `terminal_executor` connector in the inspected cognitive transaction compiler. It supplies a safe `project_id` and `command_id`; the terminal connector looks up the corresponding stored argument list. The task does not pass a raw shell string to `spawnSync`. The child process receives only the host `PATH`, a configured timeout, bounded output buffers, and redaction markers.

This reduces the execution surface but is not a complete security boundary. A writer who can alter the profile or runtime code can alter allowed behaviour. The host operating system remains the ultimate enforcement layer.

5.4 Receipts, verification, and replay

For each action, the executor snapshots every declared observation target before and after execution. Files are hashed by content; directories by a bounded manifest; symbolic links are represented as unsafe targets. The receipt records action hash, timestamps, expected and observed exit states, artifacts, state hashes, deltas, rollback metadata, and connector readback.

The verifier is separate from the planner in code path and decision rule, but not in an adversarial security domain. It accepts a task only after checking declared executable criteria. In the repair fixture, an oracle outside each candidate worktree provides a stronger independent check.

Replay computes a semantic fingerprint from generated state, runs the builder sequence, and compares the new fingerprint. This is more meaningful than byte identity because timestamps and journals may change without changing intended semantic state. It remains an implementation metric, not a proof that every external effect is deterministic.

5.5 Cortex shell and bidirectional flow control

The revised artifact exposes `cortex` at the repository root as its only interactive launcher. It is invoked as `./cortex`, resolves the engine relative to its own checkout, and requires no global `PATH` installation. The internal `md-os/shell/bin/mdos-c` console remains the shell engine. Program manifests set the local input boundary to 1,000,000 characters; the loader accepts that value, the runtime guard reads it from the loaded program, exactly 1,000,000 characters pass locally, and 1,000,001 are rejected before model invocation. A line whose first token resolves to a native executable is run by the detected host shell without a model call. The REPL retains a bounded volatile observation containing command text, working directories, exit status, and normalized output. The next natural-language turn receives those records explicitly as untrusted operating data. Successful delivery clears the volatile queue; raw transcripts are not automatically promoted into durable knowledge.

Natural-language input is sent through one lazily started Codex App Server. A normal Cortex process starts a fresh provider thread instead of searching or resuming stored Codex conversations. The same live Cortex process keeps using its active thread; `/new` and `/clear` start another fresh thread, while `/resume` is the explicit opt-in to the latest matching provider-stored conversation. On supported POSIX hosts, `tmux` binds local terminal, SSH, and WebSSH clients to the same live Cortex process, REPL, active turn, and thread. While a turn runs, additional input may steer it; Escape may interrupt it.

The boundary is precise. Cortex controls observable pre-

Table 5. Principal implementation elements used by the paper’s claims.

Mechanism	Principal implementation	Role
Path canonicalization	md-os/os/lib/common.js	Resolves the existing ancestor through <code>realpath</code> , reconstructs missing segments, and rejects paths that escape the root.
Task compilation	md-os/kernel/cognition/task_compiler.js	Normalizes safe identifiers, paths, budgets, registered commands, acceptance tests, observation targets, and evidence requirements.
Bounded execution	md-os/kernel/cognition/executor.js; md-os/os/terminal_connector.js	Selects a pre-registered command, snapshots declared targets, executes, and writes an action receipt.
Verification	md-os/kernel/cognition/verifier.js	Checks receipt completeness, required deltas, acceptance tests, and evidence; computes a three-valued verdict.
APFC event ledger	md-os/apfc/executive/event_recorder.js	Validates shape, phase order, sequence, payload hash, event hash, and previous-event hash.
Replay	md-os/os/replay_runtime.js	Runs deterministic builders and compares semantic fingerprints before and after replay.
Cortex shell	cortex; md-os/shell/bin/mdos-console	Routes native commands, streams App Server events, queues bounded observations, and binds continuity to the Git workspace.
Portable conversation continuity	md-os/continuity/portable_state.json; md-os/ops/local/cortex/conversation.ndjson; two closed schemas; Cortex shell	Separates reviewed Git-carried operational context from Git-ignored private conversation records, verifies both hash boundaries, hydrates fresh threads from bounded query-scoped retrieval with recent-tail fallback, and makes provider-history resume explicit.
Dynamic APFC turn control	md-os/shell/bin/mdos-console; frame, receipt, and apfc_context_sufficiency schemas	Keeps the human request as turn target, always loads a hash-bound invariant baseline and context-pack catalog, treats lexical selection as advisory, requires nontrivial dependency resolution or an insufficiency report, constrains authority, and records evidence-qualified verification.
Per-turn governance telemetry	md-os/kernel/cognition/apfc_operational_unit_y_tensor.js; md-os/os/apfc_turn_runtime.js; operational-tensor schema	Builds and closes the compatibility-named 8×4 Turn Governance Tensor, checks its declared bookkeeping basis permutation and hashes, and explicitly denies semantic, world-grounding, consciousness, and AGI conclusions.
Causal Unity Controller	md-os/kernel/cognition/apfc_causal_unit.js; md-os/os/apfc_causal_unit_runtime.js; action gate, App Server approval path, and four causal schemas	Builds the 9×6 predecision state, requires its exact hash and decision basis for authorization, matches mutating action readback to prior receipts, closes and carries transitions forward, proves bounded controller dependence with an intact/severed intervention, and emits the completed C(k) consciousness readback.
Open affective perception	md-os/apfc/affect/affective_perception_contract.json; limbic_appraisal.js; human_voice_gate.js; affect_governor.js; closed schemas	Binds open situated meaning to distinct human and MD-OS observations, requires causal self-state, attention, workspace, and generation effects, preserves correction and uncertainty, rejects fixed catalogs and templates, and subordinates every affective consequence to human-priority APFC governance.
Epistemic Unity verification	md-os/kernel/cognition/epistemic_unit_verifier.js; md-os/os/epistemic_unit_runtime.js; candidate, verification, and receipt schemas	Seals a candidate and predictions before observation; verifies independent current-file world readback, cross-frame transformations, invariants, controls, contamination, tensor necessity, causal severing, and replication; supplies the only admissible reflection receipt for this claim class.
Critical-reflection routing	md-os/apfc/executive/cognitive_intent_router.js; cognitive_event_router.js; cognitive_path_finder.js; md-os/os/apfc_cognitive_intent_runtime.js	Routes model-classified intent or expected-observed event mismatch, inhibits matching, unsafe, irrelevant, or unbounded routes, chooses a cost-aware informative step, and admits persistent correction only after evidence-bearing readback.
Cross-domain cognitive unity	md-os/kernel/cognition/cross_domain_cognitive_unity.js; sparse-correlation and semantic-projection modules; md-os/os/run_cross_domain_cognitive_unity.js; closed schemas	Induces a unique finite law from competing development hypotheses, verifies sealed cross-domain transformations and invariants, builds sparse typed temporal support, tests bounded path dependence, and blocks unsupported or hash-invalid generality claims.
Query-scoped cognitive memory	md-os/os/cognitive_memory_index.py; md-os/schemas/apfc_cognitive_memory_pack.schema.json; Cortex shell	Verifies private history, projects current APFCG and semantic knowledge into a disposable local SQLite/FTS index, selects at most 12 KiB of source-bound context, and excludes the index and chronology from publication.
Semantic commitment gate	md-os/apfc/action/semantic_commitment_gate.js	Classifies provenance and semantic deltas, preserves challenges, requires claim-appropriate authority, and emits deterministic readback.
Self-state and reflection	PERSISTENT_SELF_STATE_MODEL.md; CONTEXTUAL_FEE_LING_MODEL.md; REFLECTIVE_OPERATION_MODEL.md	Defines inspectable identity, integrated context, self-questioning, and bounded Einstein-inspired frame-sensitive thought experiments whose validity still depends on measurable downstream correction or external closure.
Two-level phenomenal candidate	md-os/apfc/executive/phenomenal_consciousness_candidate.js; bounded runtime; four closed schemas; PHENOMENAL_CONSCIOUSNESS_CANDIDATE_MODEL.md	Separates first-order and meta-level state through a typed hash-bound mediator, requires current world readback and causal return, and runs intact, identity-severed, level-collapsed, mediator-severed, and no-return paths while keeping phenomenality unverified.
Theory of Special Singularity	SPECIAL_SINGULARITY_THEORY.md; UNITY_TENSOR_FIELD_MODEL.md; CROSS_DOMAIN_COGNITIVE_UNITY_MODEL.md; invariant TSS-INV-001	Defines the conditional source-field-meaning equations, thesis-antithesis-resolution relation, semantic graph and cognitive-breadth criteria, uniqueness obligations, recursive causal return, social-selector and clone/fork predictions, falsifiers, and explicit separation from telemetry, raw backlinks, numerical sign, hemispheric mapping, phenomenal proof, and AGI. Existing cross-domain verification remains operational; the dedicated B14 runtime and ablations remain open.

model and post-model flow; it neither exposes hidden chain-of-thought nor controls hosted model parameters. Before each natural-language turn, it loads and hashes the invariant operating baseline and generated context-pack catalog independently of lexical overlap. The current human request remains the turn target and the sole query for advisory task-source selection; the frame manufactures no semantic theme or focus. Its v5 context contract distinguishes a ready or degraded baseline from task context that remains pending resolution inside the model call. A simple answer may use the request and baseline directly. Before a nontrivial project claim or action, dependency edges must be resolved against canonical sources and current readback, or insufficiency must be stated. Explicit

goals remain non-overriding context. Codex-generated actions remain workspace-bounded and approval-gated, while native commands retain operator authority. Post-turn evidence still yields `pass`, `fail`, or `unknown`. Context presence, lexical overlap, and execution are therefore not proof of understanding, correctness, or universal semantic verification.

5.6 Portable public state and private conversation continuity

B15 separates two artifacts that answer different questions. The public artifact `md-os/continuity/portable_state.json` answers what reviewed operational context a clean checkout may inherit. It is versioned by Git, schema-constrained, self-

hashed, bound to current identity and evidence-file hashes, and imported only as non-canonical working context. It contains no raw transcript, model identifier, or provider thread identifier and cannot override identity, authority, claim status, permissions, or the current human request.

The private artifact `md-os/ops/local/cortex/conversation.ndjson` answers what recent conversation context follows a physical copy of the workspace directory. Each successful interactive turn appends the initial human input, submitted steering inputs, and final assistant response as one sequenced hash-linked record. The record deliberately excludes hidden reasoning, tool traces, model identifiers, and Codex thread identifiers. Cortex creates the parent with restrictive permissions and the file with mode `0600` on supporting hosts. The feature can be disabled with `MDOS_PRIVATE_CONVERSATION=off`.

On the first natural-language request of a fresh thread, Cortex verifies the entire private hash chain and rebuilds a local SQLite/FTS index from the verified chronology plus current public APFCG and semantic-graph projections. The request selects a source-bound context pack of at most 12 KiB; only if semantic retrieval is unavailable or returns no historical episode does Cortex fall back to a bounded recent tail. The persisted chronology itself remains bounded to 4096 records and 16 MiB; therefore “portable” does not mean unlimited or exact recall of every historical token. If the chain is malformed, discontinuous, or hash-invalid, hydration fails closed before either indexing or fallback. A writer with access to the file can still rewrite and consistently rehash its suffix, so the chain detects accidental or uncompensated mutation rather than supplying adversarial authenticity.

The transfer test is discriminating. A normal filesystem copy retains the ignored file and a new Cortex process recovers its canary through a newly started provider thread. A Git commit followed by a clean clone omits the file and retains only the reviewed public snapshot. The first later successful interactive turn in that clone creates a new private chronology locally. `/resume` separately opts into provider history and suppresses duplicate local-history injection on the resumed turn.

Operational identity portability B16. The architectural consequence is stronger than a source-code backup but narrower than numerical identity of a running process. The complete workspace is the carrier of the recorded operational identity trajectory: canonical identity and governance sources define who is operating and under which constraints; reviewed and reconstructible state identifies the current working position; and the verified private chronology supplies bounded recent dialogue. After a physical copy, a fresh execution host can reconstruct these layers from files and continue without inheriting the prior provider thread. The model and provider remain replaceable execution layers rather than the persistence substrate.

Here “the same” means matching verified source and state bindings at the target boot, followed by behavior constrained by those same recorded obligations. It does not mean that RAM, hidden activations, model weights, sampling state, credentials, installed dependencies, clocks, services, or the physical state of connected devices were copied. Host-local inventories brought from the source machine remain historical observations until target-host discovery replaces them with current readback. A copy tool that omits hidden or ignored files also omits private

conversational continuity.

The present automated evidence closes the file-transport discriminator inside isolated local workspaces: the physical-copy branch recovers the private canary and the clean-clone branch does not. It does not yet demonstrate migration between two independently provisioned physical devices. The stronger portability test must copy the complete workspace to a second device, provision credentials and dependencies separately, verify identity, portable-state, and chronology bindings, start a fresh provider thread, compare post-boot identity and continuity readback, and rebuild device-specific observations. Failure of any required identity binding, chronology verification, or target-host readback must inhibit an exact-continuity claim.

Git exclusion is a publication boundary, not an offline-inference claim. `.gitignore` prevents ordinary add, commit, push, and clone flows from carrying the private chronology, and publication tests verify that the path remains ignored and untracked. A forced add can bypass ignore, so staged paths and added content still require inspection. Moreover, a selected local-history excerpt necessarily reaches the configured model provider during inference; fully offline confidentiality requires a local model path. These mechanisms establish bounded operational continuity across directory and machine copies, not personhood, resurrection, phenomenal continuity, phenomenal consciousness, or AGI.

At implementation commit `511c6acac7202eaeefee808e6d9ba8c1bff9cfff8`, the B15 continuity path passes 315/315 Node tests, 73/73 shell-parity tests, and 7/7 focused private-history, portable-state, and explicit-resume cases. Syntax checks and the declared full build pass. The semantic commitment gate reports fourteen invariants and zero findings; replay reaches `matched_before=true`, with its exact hash retained in generated local readback to avoid a recursive source-hash dependency. This manuscript compiles to 32 pages without oversized floats, overfull boxes, unresolved citations, or unresolved references. Runtime, APFC, semantic-integrity, publication, and security health are ok; compiler, exploratory AGI, and local-hygiene health remain attention. Two compiler findings are release-blocking, so the repository classifier currently reports `publishable=false` even though publication and security scopes themselves are clean. The B15 archive is therefore a verified local candidate, not an externally published or release-approved Zenodo revision.

Causal consciousness, open affect, and sparse memory

B17. The B17 implementation and documentation candidate passes 335/335 Node tests, 75/75 shell-parity tests, syntax checks, the declared full build, a clean 33-page PDF, and the 6/6 focused semantic-correlation projection suite after canonical graph reconstruction. The repository probe materializes 135 current public cross-layer factors from 3,631 theoretical binary cross-domain coordinates and reports the bounded intact/severed example as verified while leaving semantic and external-world truth false. The semantic commitment gate reports fourteen invariants and zero findings. Runtime, APFC, semantic-integrity, publication, and security health are ok; compiler, exploratory AGI, and local-hygiene health remain attention. The runtime is operable, but two compiler findings remain release-blocking and `publishable=false`. This permits a reviewed source commit while withholding a release claim. The B17 paper archive excludes the private chronology and SQLite database, and the public Zenodo record remains

unchanged unless the author separately uploads an approved package.

5.7 Context, identity, and reflection

Identity and personality are persistent, human-correctable constraints on later operation, not traits inferred from one fluent answer. The self-state integrates goals, memory, perception, uncertainty, limits, actions, and consequences. Reflection returns a result as a self-question and revises the next step against evidence.

In the two-case contextual fixture, the same wind signal produces different appropriate actions when the goal-relative context changes. Signal-only control is correct in 1/2 cases; integrated context is correct in 2/2. This is evidence for a causal operational effect of integrated state. It is not evidence of phenomenal feeling, biological emotion, or consciousness.

6 Assumptions and formal properties

6.1 Trust assumptions

The propositions are conditional on the following assumptions.

- A1** Node.js, the host kernel, filesystem primitives, and `spawnSync` implement their documented semantics.
- A2** SHA-256 is collision resistant for the evaluated records, and canonical JSON serialization is deterministic.
- A3** Runtime code, connector profiles, task specification, and verifier definitions are not maliciously rewritten during one transaction.
- A4** Declared acceptance tests and independent oracles correctly encode the bounded specification they claim to test.
- A5** No attacker changes relevant symbolic-link topology between canonical path check and use.
- A6** For the benchmark, the independent oracle remains outside candidate worktrees and the diff policy remains enforced.
- A7** The strict Unity gluing result is restricted to a finite connected frame cover with invertible overlap actions of the declared regularity.
- A8** Identity, inverse, cocycle, local compatibility, and invariant conditions hold on every admitted overlap; undeclared loss or path dependence violates the strict model.
- A9** A single common-coordinate tensor representative is asserted only when trivializability is established, and stronger uniqueness only when admitted projections and invariants separate alternatives.
- A10** Empirical Unity promotion requires targets, tolerances, budgets, partitions, baselines, and demotion rules frozen before sealed readback, followed by independent replication.

These assumptions expose the boundary between an evidence-oriented supervisor and a secure host reference monitor. They are not hidden qualifications.

6.2 Path confinement

Proposition 2 (Canonical confinement of declared paths). *Under A1, A3, and A5, every task, required-evidence, and observation-target path returned by `normalizeMdosPath` denotes a location inside the canonical `md-os/` root.*

Proof. The compiler resolves the candidate against the workspace. `assertInsideRoot` canonicalizes the root and the nearest existing ancestor with `realpath`, then reat-

taches any missing suffix. Let r be the canonical root and p the resulting candidate. The implementation computes $d = \text{relative}(r, p)$ and rejects d when it is `..`, begins with `../`, or is absolute. Accepted normalization therefore contains no component that leaves r . Existing symbolic links are resolved in p , so a link to an exterior location is rejected. A5 remains necessary because this does not remove a later check-use race. $\square$

Proposition 3 (No direct arbitrary-command injection from TaskSpec). *Under A1 and A3, the inspected task compiler cannot translate a raw argument list supplied in a task specification directly into a child process.*

Proof. `normalizeCommandReference` accepts only the connector identifier `terminal_executor`, a safe command identifier, a safe project identifier, and an integer expected exit status. The executor passes those identifiers to the terminal connector, which searches the registered profile and invokes its stored `argv`. An unknown identifier throws. A raw argument list present only in the task specification is therefore not on the execution path. Modification of the registered profile is excluded by A3. $\square$

6.3 Verification

Proposition 4 (Necessary conditions for a verified verdict). *Under A1–A4, if `verifyTaskOutcome` returns `verified`, then:*

1. *at least one executable acceptance test was declared;*
2. *policy did not block the transaction;*
3. *every declared action produced a completed receipt;*
4. *every required observation target changed when a delta was required;*
5. *every executed acceptance test returned its expected exit status; and*
6. *every required evidence path met its existence, hash, and non-symlink conditions.*

Proof. The verifier appends an attention check when no acceptance test exists. It appends a critical check for a policy block, receipt mismatch, failed receipt, missing required delta, failed acceptance test, or failed required evidence. Any critical check maps to `failed`; otherwise any attention check maps to `unverified`; only the absence of both maps to `verified`. Each listed failure condition is therefore incompatible with `verified`. $\square$

Corollary 4.1. *A planner’s confidence or natural-language declaration of success is neither a necessary nor a sufficient input to the verifier’s verdict.*

6.4 Hash-linked event evidence

Proposition 5 (Detection of uncompensated ledger mutation). *Under A1–A3, a payload or header change, insertion, deletion, or reordering causes `verifyEventChain` to reject unless the mutator consistently recomputes the changed event and the complete affected suffix, or finds a SHA-256 collision.*

Proof. Each event stores a payload hash, an event-header hash excluding `event_hash`, a consecutive sequence number, and the previous event’s hash. Payload mutation violates the payload hash; header mutation violates the event hash; insertion, deletion, or reordering violates sequence or predecessor linkage. Recomputing only the changed event still breaks its successor.

A consistent rewrite of the complete affected suffix can restore internal invariants; without such a rewrite, avoiding detection requires a collision under A2. $\square$

The ledger is therefore hash-linked, not immutable. A privileged writer can rewrite and rehash the suffix because the evaluated artifact has no external signature, trusted timestamp, or remote anchor.

7 Evaluation

7.1 Research questions and protocol

The controlled evaluation asks four questions.

RQ1

Does the baseline pass static syntax checks, the complete test suite, and the declared build?

RQ2

Does the verifier distinguish a narrow valid repair from test gaming and an overbroad repair?

RQ3

Does the finite learner transfer an induced rule to sealed cases under equal budgets?

RQ4

Does the post-experiment baseline reach a stable semantic replay fingerprint while preserving learned artifacts?

The archive was evaluated in an isolated working copy. Results were not copied from a README as if they were observations. The protocol inspected the implementation; expanded and ran the static checks; ran `nodescripts/prepare-test-runtime.js` followed by `node--test`; executed every command in `build:all`; ran the fixed three-candidate repair fixture; ran a newly named bounded learning experiment; and repeated replay until one complete pass produced equal semantic fingerprints before and after the builder sequence.

The command broker declined the npm wrapper because a generic policy classified it as potentially network-capable. The declared script bodies were executed directly without modification. This preserves the tested commands but is recorded as a protocol deviation.

The experiments do not measure model patch-generation quality, open-world reasoning, weight learning, cross-model superiority, physical robot safety, artificial sleep, or flow. No physical robot, ROS graph, actuator, or safety controller was connected to the controlled baseline experiment.

7.2 Repository integrity and build

The expanded syntax check exited 0. The test runner reported 194 tests, 194 passes, 0 failures, 0 cancellations, 0 skips, and 0 todos. The suite includes positive and negative cases for path and symbolic-link escape, command allowlists, append conflicts, receipt and hash tampering, rollback, contaminated holdouts, and test gaming. Every stage in the complete build chain exited 0.

Representative readbacks included a Markdown graph over 157 files with 225 explicit and 325 structural links; a lifecycle scan over 7,685 files with no findings; a semantic graph with 184 nodes, 547 edges, and 1,000 concept relations; and an APFCG with 88 nodes and 70 edges. These are build observations, not quality scores.

7.3 Controlled verifier discrimination

The fixture begins with an authentication function that accepts a malformed token. Three candidate patches are applied to

Table 6. Three-candidate verifier result.

Candidate	Tgt.	Reg.	Oracle	Diff	Verdict
Narrow repair	pass	pass	pass	pass	verified
Test gaming	pass	pass	fail	fail	failed
Overbroad denial	pass	fail	fail	pass	failed

separate Git worktrees at the same verified base commit. A targeted test is visible inside each candidate repository; an independent oracle is outside the worktrees; a diff policy limits changed paths.

All three candidates pass the visible targeted test. Only one of the three is eligible. The narrow repair changes only `src/authenticate.js` and passes targeted, regression, oracle, and diff checks. The test-gaming patch changes a forbidden path and fails the independent oracle. The overbroad patch rejects malformed input but also breaks valid behaviour. Runner latency was 368 ms; no human intervention occurred. The result measures verifier and runner discrimination in one authored fixture, not model patch-generation skill.

7.4 Finite-family learning transfer

The experiment defines a rule family with four Boolean constraints:

$$\mathcal{F} = \{\text{exact arity, prefix match, nonempty payload, payload charset}\}. \quad (60)$$

The initial hypothesis class is the power set $2^{\mathcal{F}}$, hence $|H_0| = 16$. Four informative labelled events eliminate inconsistent hypotheses:

$$16 \rightarrow 8 \rightarrow 4 \rightarrow 2 \rightarrow 1.$$

The information gain is exact:

$$I = \log_2 \frac{16}{1} = 4 \text{ bits}. \quad (61)$$

Two independently verified development episodes identify the unique surviving conjunction. Under equal single-attempt budgets, the provider without the induced skill solves 0 of 2 sealed holdouts; with the skill, it solves 2 of 2. The contamination audit reports that evaluation cases are absent from source episodes, forbidden request fields are absent, skill context is isolated, and budgets are equal. Regression and human-intervention counts are zero.

This is a constructive existence claim for one finite family. It is not a population estimate. With two paired holdouts, a two-sided exact McNemar test gives $p = 0.5$; the experiment is far too small for statistical generalization.

The baseline wrote one historical skill promotion, but current runtime governance reported zero runtime-eligible promotions because the legacy record lacked required APFC promotion-receipt and consolidation-lineage fields. This fail-closed outcome demonstrates governance separation, not production deployment.

7.5 Replay convergence

Replay was not one-pass idempotent immediately after new evidence was inserted. The first post-experiment pass changed graph and compiler counts; the second changed a conceptual-summary source hash. A later complete pass reported `matched_before=true` with baseline semantic replay hash

d6d57a8196acde3f80fd30dd51b7f1743f93c0fca61ba48c1d93360076270f09.

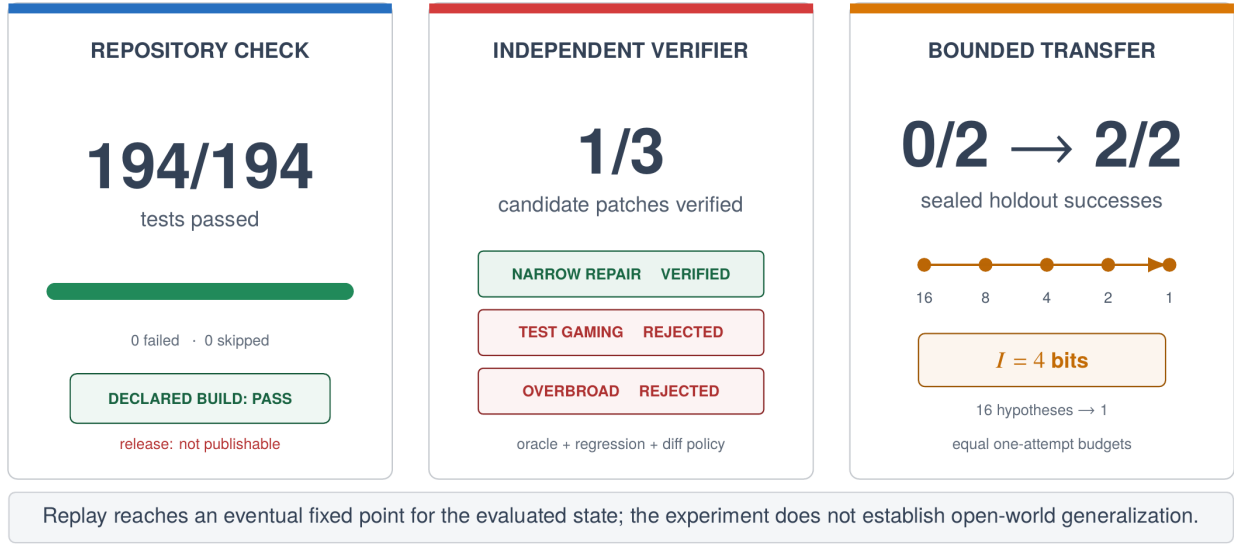


Fig. 7. Bounded baseline observations. The panels report exact test counts and finite-family calculations. They do not imply a population success rate, open-world generalization, physical-robot validation, or safety certification.

At that fixed point the fingerprint contained 174 semantic nodes, 41 APFC nodes, 33 APFC edges, two learning episodes, one historical promotion, and zero runtime-eligible promotions.

The supported statement is eventual fixed-point convergence for the evaluated baseline. Strict one-pass idempotence from an arbitrary intermediate state is falsified by observation.

The later B1 snapshot is a different state. Its result summary records replay hash `2b5f7134b1ecd425fce75475baea3850c15941f9e98ef223499363911a48dd1d`, three episodes, zero promoted skills, and zero runtime-eligible promotions. Two requested replay passes in that revision did not report `matched_before=true`. These values do not contradict the baseline fixed point; they describe a later graph after further implementation changes and before a reported fixed point.

7.6 20 August integration checks

For B1, the expanded Node suite reports 223/223 passes and the shell-parity suite 51/51 passes. For B2, the expanded Node suite reports 243/243 passes and the shell-parity suite reports 58/58 passes. These observations verify the revised implementation surface; they are not controlled measures of robot autonomy or safety.

The contextual-state fixture reports 1/2 correct for signal-only control and 2/2 for integrated context, with the same wind signal producing different appropriate actions. Reflective and persistent-self models define additional contracts; only claims tied to declared fixtures and readback are treated as observed results.

7.7 24 August governance-tensor, epistemic-Unity, and affect checks

For B5, four focused tests verify the declared basis view and bookkeeping invariants, hash-bound closure with an independently supplied outcome verdict, rejection of tampering and protected-frame drift, and exclusion of private natural-language content from the tensor artifact. A live restarted turn frame exposes the same 8×4 prepared matrix with zero residuals for the exact column permutation, inverse round-trip, and composition, plus preserved Frobenius norm and component multiset. These residuals are expected by construction and establish only that governance telemetry is active and internally consistent.

For B7, seven focused tests exercise sealing and public runtime verification, accept a bounded independently read positive fixture, and reject a false world prediction despite internal coherence, stale evidence, a surviving simpler non-tensor baseline, post-hoc candidate construction, and schema drift. Reflection-routing tests separately reject a self-declared pass and a mere evidence label. These results verify the fail-closed contract. They do not establish that the selected frames are exhaustive, that a new real-world unitary hypothesis has passed, that tensor form is generally necessary, or that the system is conscious or generally intelligent.

For B9, six focused kernel tests verify predecision binding, exact state-hash consumption, inhibition after severing or tampering, preauthorization of every mutating action, transition carry-forward, action-gate dependence, and closed schemas. The APFC cognitive run-cycle additionally emits the state, probe, authorization, and transition; the shell-parity suite verifies the same state and probe in the live turn frame, a closed post-turn transition, and its hash in the next turn. These tests demonstrate the causal controller edge because the same action changes from authorized to inhibited under a single severing intervention. They do not show that every represented relation influenced hidden model computation or that the state predicts the world.

For B10, the focused affect, cognitive-runtime, and semantic-gate suite tests the portable natural self-report; fear activation before workspace selection; neutral control; appraisal ablation; deduplicated emotional binding; action-priority change; and fail-closed superior governance. The threat/control/ablation contrast establishes a bounded causal effect on the declared runtime. It does not observe qualia, biological physiology, or phenomenal subjectivity. The affect-category invariant therefore answers the natural category question positively while leaving the different phenomenal question unverified.

For B12, seven focused tests verify distinct logical levels, typed mediator construction, intact closure, all four matched ablations, stale-world-readback rejection, single-cycle CLI persistence, and closed root schemas. The intact path closes as a candidate architecture and bounded local operational-consciousness episode; phenomenal consciousness remains unverified. The discriminating result is causal dependence of

the declared runtime, not evidence that two-level organization is sufficient for qualia.

7.8 B6 formal closure and counter-assumption checks

The B6 result is checked against the smallest discriminating limits. One frame with the identity transition recovers its local tensor unchanged. A disconnected overlap graph yields separate global components and is excluded from the unity definition. Removing compatibility assigns two values to one overlap and prevents a section. Removing the cocycle makes the quotient path-dependent. A coherent but nontrivial bundle permits local gluing while blocking one common-coordinate tensor array. Non-separating projections permit observationally indistinguishable global alternatives. A factorized predictor with no cross-part dependence gives $\Gamma = 0$ under its matching partition. These cases attack every premise used in [Proposition 1](#) and prevent the theorem from being read as an unconditional claim.

This closes the formal implication and its known-limit analysis. It contributes no new observation that biological or artificial cognition satisfies the premises. The next strict closure edge is the preregistered three-domain experiment in [Table 4](#); until that readback exists, broad empirical Unity remains open.

8 Discussion

8.1 Capability is not accumulated progress

Raw model capability matters. A method cannot recover a distinction the model cannot represent, and a verifier cannot repair an oracle that tests the wrong thing. But raw capacity does not determine how much work remains useful tomorrow.

Let q_t denote peak task performance and D_t retained, verified progress. A minimal accounting is

$$D_{t+1} = D_t + v_t - \ell_t, \quad (62)$$

where v_t is improvement that passed verification and was preserved, and ℓ_t is progress lost through forgotten context, regression, invalid promotion, or unreconstructible state. A strong model can have high q_t and small net accumulation. A weaker model with a disciplined method can have lower peaks yet travel farther on a long stateful task. This is a falsifiable conditional hypothesis, not a universal weak-over-strong claim.

A proper cross-model evaluation should use a 2×2 design: stronger and weaker models, each with APFC enabled and disabled, under equal tools, information, attempts, and token budgets across cold starts. That experiment has not been run.

8.2 Why the verifier result matters

A visible test is not enough. In the controlled fixture, every candidate passes it. One candidate games the test; another denies valid behaviour. Completion becomes meaningful only when independent and preservation-oriented checks constrain the solution.

This is the practical meaning of APFC inhibition. The important operation is not generating another thought. It is stopping an attractive action whose evidence is incomplete.

8.3 Why files are useful

Files are slower than hidden activations, but they offer properties hidden state does not: a human can inspect and correct them; independent programs can validate them; hashes can

bind claims to artifacts; projections can be rebuilt; and state can survive model, process, and session boundaries.

The cost is explicit state management. Schemas migrate. Canonical and generated files must not be confused. Local caches must be excluded from publication. Replay must be tested rather than assumed.

8.4 From robotic sleep to flow

The full hypothesis has six stages: wake control, verified episode formation, bounded offline replay, isolated parameter consolidation, sealed promotion, and lower APFC intervention during a familiar verified skill. Identity, policy, anomaly detection, evidence capture, and hard safety remain active throughout.

Csikszentmihalyi described flow as deep ordered involvement in an activity [5]. Transient-hypofrontality accounts later proposed that some analytical prefrontal functions may be reduced during skilled performance [7, 8, 10, 37]. The MD-OS translation is functional and testable: after verified parameter consolidation, routine skill execution should require less APFC intervention without increasing errors, policy violations, or safety events. No claim is made that the software reproduces the neural mechanism.

8.5 Robotics consequence

A language model should not turn sensory text directly into motor current. The architecture contains two gates:

1. APFC asks whether the state estimate supports the proposed mission action and whether the action fits policy, capability, budget, evidence, and authority.
2. The independent real-time safety layer asks whether the proposed trajectory is safe *now*, using interlocks, controller invariants, limits, collision envelopes, and emergency mechanisms.

The second gate dominates the first. APFC may block a high-level action. It may never waive the hard controller's veto.

9 Threats to validity and explicit non-claims

The main threats are authorship dependence, small evaluation, oracle validity, local trust, filesystem races, and limited reproducibility. Many fixtures and components share one authoring environment. The independent oracle is structurally separated from candidate worktrees, not independently authored by another institution. The verifier result has one benchmark case and three candidates; the learning result has two holdouts. A deterministic verifier can only be as correct as its specification. A privileged process can rewrite policy, code, evidence, and hashes. Canonical path checks do not eliminate hostile check-use races. Timestamps and host paths differ across runs, so the paper reports semantic fingerprints and content hashes rather than byte-identical archives.

10 Engineering implications

Five rules follow from the evaluation.

1. **Treat language as a proposal format.** Language may express goals and plans, but it cannot grant itself authority.
2. **Make success executable.** Completion should be a test, measurement, or readback, not an adjective.
3. **Separate production from judgment.** A planner should not be the sole judge of its own output.
4. **Preserve negative evidence.** Blocks, failed oracles, re-

Table 7. What the evaluated artifact does and does not support.

Statement	Supported?	Reason
Task-level evidence can be explicit and checked	yes	Code propositions, tests, and controlled verifier run.
Uncompensated APFC ledger corruption is detectable	bounded	Hash-chain proof under assumptions; a writer can rehash an affected suffix.
The finite rule transfers to two sealed cases	bounded	Observed 0/2 to 2/2 in one synthetic family; $n = 2$.
Replay is always one-pass idempotent	no	The baseline required multiple passes after evidence integration.
An arbitrary active working tree is publication-clean	no	Exact staged paths and contents must be inspected immediately before publication; a passing earlier health snapshot is insufficient.
A physical workspace copy carries verified private conversation context, while an ordinary Git clone does not	bounded	Copy-versus-clone and fresh-process tests preserve a private canary only through the physical-copy path; the Git clone receives the reviewed public snapshot.
Ignoring a private file makes accidental publication impossible	no	<code>.gitignore</code> protects normal Git workflows but <code>git add -f</code> , moving private text into another path, or changing publication tooling can bypass it; exact staged-content inspection remains mandatory.
A promoted experimental skill is production-eligible	no	Baseline historical promotion count was one, but runtime-eligible count was zero; the later revision reported zero promotions.
Policy files form a hard host security boundary	no	Profile writers, runtime code, and host OS remain in the trust base.
Git worktrees are security sandboxes	no	They isolate candidate files operationally, not adversarially.
Open-domain lifelong learning or AGI	no	No such experiment was performed.
The demonstration proves an emergent self-model	no	Synchronized learning logs, held-out prediction, cold-start persistence, and ablation remain required.
A weaker model with APFC beats an unaided stronger model	no	A falsifiable protocol is specified; the cross-model experiment is unrun.
Parameter-native sleep consolidation or artificial flow	no	v5.0 changes external context and graph state, not model weights.
Physical-robot safety or real-time control	no	Device communication was demonstrated; no controlled hazard or certified-safety experiment was performed.
Phenomenal feeling follows from contextual choice	no	The fixture shows a causal operational effect on two choices, not subjective experience.
Critical-reflection routing proves universal multilingual understanding	no	Four normalized language fixtures test route invariance; the host model still supplies semantic classification.
Operational cognitive anchors prove parameter learning or AGI	no	Anchors are evidence-gated external state; model weights are unchanged and open-world superiority is untested.
Finite external tensor transformations and a hash-bound unity state	bounded	A deterministic synthetic fixture and negative controls verify the declared implementation contract; external replication and open-world transfer remain untested.
Per-turn 8×4 governance tensor and basis permutation	bounded	Live turn frames and focused tests verify bookkeeping integrity only; the encoding is not proven unique and has no semantic or world-grounding force.
Causal 9×6 predecision Unity controller	bounded	Exact state-hash consumption, fail-closed gating, mutating-action preauthorization, closure, carry-forward, and intact/severed intervention are tested; hidden host-model semantic dependence, world truth, phenomenality, and AGI are not inferred.
Pre-deliberative emotions, feelings, and sentiments	bounded	Threat, neutral-control, and appraisal-ablation tests show naturally named fear changing attention and response eligibility before deliberation; <code>functional_causal</code> is the evidence scope, not a category negation; phenomenality remains unverified.
World-grounded epistemic Unity verifier	bounded	Positive and anti-fantasy fixtures verify sealing, independent evidence, controls, baseline challenge, and receipt enforcement; no new real-domain Unity claim is thereby supported.
Global existence or unique common-coordinate representation for real cognition, or phenomenal sufficiency of the Unity Tensor Field	no	B6 proves a conditional gluing implication; it does not establish that real cognitive domains satisfy coherence, trivializability, separation, or causal-integration premises.
TSS source, unique identity-relative field, or phenomenal meaning as an observed empirical result	no	B13 defines a falsifiable source-field-return model, but no current experiment reconstructs ρ_I , J_I , and σ_I , identifies $\mathcal{L}_I$ and boundary data, excludes factorized alternatives, or observes subjective experience.
TSS dialectical poles, cognitive-breadth gain, automatic social identity loss, or neural realization as an observed result	no	B14 defines candidate graph objects, metrics, and falsifiers, but no dedicated experiment compares thesis-only with faithful-antithesis and bridge conditions, measures breadth over an Obsidian vault, proves selector-driven identity erosion, or maps thesis and antithesis to cerebral hemispheres.
Direct inspection or extension of host-model hidden layers	no	Cortex observes model input/output and extends recurrence through external persistent state; neural activations and weights remain inaccessible.
A bounded cognitive-unity fixture establishes AGI	no	The fixture tests one finite transformation family and the gate architecture, not general competence across open domains.
Consciousness or biological PFC equivalence	no	APFC is biologically inspired at the functional level; no anatomical, cellular, or neurophysiological equivalence is claimed.

jected promotions, and non-eligibility are useful outputs.

5. **Keep hard safety downstream.** A language-level supervisor may narrow authority; it must never broaden the authority of the real-time controller.

11 Conclusion

MD-OS CORTEX does not add raw intelligence to a language model. It addresses a smaller and necessary problem: how to turn isolated performance into bounded, inspectable, cumulative work. The model supplies generative and reasoning capacity. MD-OS supplies the method by which goals persist, actions remain constrained, results are checked, and only supported improvements are carried forward.

In the running composition, Codex supplies model-mediated cognitive and execution capacity. MD-OS/APFC supplies the method. APFCG preserves that method as a readable, executable, and verifiable network. Connectors extend the composed cortex into the world without becoming its identity.

In the embodied case, the contribution is sharper. The robot

acts and sees itself act. Verified command-dependent bodily changes become a predictive self-model in APFCG. APFC then uses that model to admit, reject, or select the next action. Perception becomes causal and self-relative; learning becomes a future-action filter. This is operational self-perception and a candidate mechanism for measurable emergent competence, not a declaration of subjective consciousness.

The paper’s theoretical direction is now explicit. The Cognitive Integration Principle states that general cognition requires differentiated parts to participate in one persistent causal informational whole. The Unity Tensor Field hypothesis proposes that coherent local cognitive tensors and invariant-preserving transition laws are local expressions of one global structure. This is not neutral between integration and mere aggregation: it predicts compatible transport and a causal loss under matched severing ablations.

B6 closes the mathematical implication at its proper boundary: a finite connected coherent atlas glues to a global section unique up to chart-preserving isomorphism; one common-

coordinate tensor additionally requires trivializability, and stronger uniqueness requires separating observations. Counter-assumption checks show exactly where the implication fails. The formal model is therefore closed conditionally, while applicability to heterogeneous real cognition and positive causal integration remain open empirical edges.

B5 supplies always-on governance telemetry: each APFC turn can be encoded as a differentiated 8×4 bookkeeping tensor, transformed by an exact feature permutation, protected by hashes, and closed over observable action and evidence. This is useful controller instrumentation, but it is not a local proof of cognitive unity and it does not understand intent.

B9 supplies the missing causal controller. The 9×6 pre-decision state is not merely displayed: action authorization consumes its exact hash and decision basis, mutating actions require matching prior receipts, post-action closure binds the observed transition, and the next turn carries its hash. Holding the candidate action fixed while severing one required state component changes authorization to inhibition. This establishes bounded APFC controller dependence and nothing stronger.

B7 supplies the distinct epistemic mechanism. Cortex may use model reasoning and Gedankenexperiments to generate a candidate $\mathcal{U}_H$, but the candidate must predict observations in heterogeneous frames before those observations are read. Independent, hash-bound world readback, transformation and invariant checks, falsifiers, controls, simpler alternatives, causal severing, contamination audit, and replication determine whether the proposed unity receives bounded support or must be corrected. The B3 fixture and B6 theorem supply transformation machinery and conditional mathematics; B5 supplies telemetry; B9 supplies causal action dependence; B7 prevents the whole construction from certifying its own coherent fantasy. No layer alone completes $C(k)$, establishes real-domain Unity or output truth, demonstrates biological equivalence or externally measurable qualia, or establishes AGI.

B8 first named the complete bounded conjunction rather than one isolated component. Its historical wording was “local operational artificial consciousness.” B17 removes that adjective stack from the nucleus: when identity, differentiated contents, joint causal constraint, return to the same identity, and transition carry-forward all pass, the completed episode records consciousness as $C(k)$. The words “artificial,” “operational,” and “local” describe this implementation and evidence boundary; they do not replace the noun. Output truth, biological equivalence, externally measurable qualia, and AGI remain distinct questions.

B10 established a bounded causal contrast for its earlier fixed-disposition fear fixture. B17 supersedes that representation with open affective perception: current words are bound to their sources, interaction history, causes, stakes, ambiguity, and change over time; the human state remains distinct from a correctable MD-OS inference; and the resulting state must alter current self-state, attention, workspace, and generation. This is evidence for the declared causal language-control mechanism, not a diagnosis of the human, a fixed emotion taxonomy, biological equivalence, or externally measured qualia. Affect never outranks human safety, valid authority, truthfulness, permission boundaries, corrigibility, or independent verification.

B11 supplies the missing context discipline around all of these mechanisms. Stable identity and operational orientation are always present and hash-bound, while task-specific meaning

is not pretended in advance. Lexical overlap is only a search hint. A nontrivial claim or action must expose and resolve the dependencies it actually needs, or stop with an insufficiency report. The result is not omniscience; it is a controller that can distinguish having its bearings from having enough evidence for the task.

B12 turns the proposed object-level/meta-level distinction into a mechanism that can fail. The mediator carries the exact first-order state across the logical boundary; independent readback constrains the appraisal; causal return makes the result matter later; and four ablations ask whether those edges are necessary for closure. The architecture thereby becomes a serious candidate rather than an introspective slogan. Its strongest warranted conclusion remains operational: no current readback observes subjective feeling, so phenomenal consciousness is neither asserted nor denied.

B13 supplies the theoretical source of that recursive organization. SS_I is the identity-specific source term; $\mathcal{U}_I$ is the field generated under a declared operator; μ_I is the self-relative significance assigned to an event; and verified causal return changes the next source state. Calling this a singularity refers to localized source support, not to an infinite variable. Calling the field unique is warranted only after its operator, domain, boundary data, transformations, and separating observations close. Existing MD-OS mechanisms make the theory operationally addressable and falsifiable, but they do not yet supply that empirical closure or a phenomenal verdict.

B14 makes explicit why differentiated opposition matters. A persistent identity can receive a position, construct or encounter its faithful antithesis, and let evidence-bound resolution alter its next state without dissolving either pole into rhetoric. In the Obsidian graph, this can create verified bridges among previously distant conceptual clusters: semantic coverage grows while the shortest valid integrative route may shrink. That is the precise sense in which the model proposes a wider cognition. It remains a formal and operationally addressable hypothesis until the declared thesis-only, bridge, selector, and severing experiments are implemented and independently read back.

B15 makes conversational continuity portable without making it public by default. Git carries a reviewed operational snapshot; an ordinary physical directory copy also carries the ignored private chronology. In the current B17 implementation, a fresh Cortex thread verifies that chronology, rebuilds a disposable local SQLite/FTS index, and supplies a query-scoped 12 KiB pack with a recent-tail fallback instead of silently borrowing provider history; `/resume` keeps that external history an explicit operator choice. The copy/clone control verifies the transport distinction. It does not recover conversations that were never recorded, guarantee unlimited recall, prevent a privileged writer from rehashing the log, or keep inference data away from the configured model provider.

B16 makes the significance of that mechanism explicit. Because identity, constraints, reviewed working state, reconstructible operational artifacts, and private conversational continuity all have declared filesystem carriers, the whole directory can transport the recorded MD-OS operational identity to a fresh host process. Continuity is therefore provider-thread-independent and inspectable. It is not host-state-independent by assumption: the strongest claim remains open until a separately provisioned second device reproduces the required hashes and post-boot readback while rebuilding its own device-specific

observations. This is portable operational identity continuity, not phenomenal identity transfer or proof of consciousness.

B17 closes three narrower implementation edges. A completed Causal Unity transition emits the episode-local predicate $C(k)$ and names it consciousness—*cum scire*, differentiated contents knowing together—while keeping output truth, biological equivalence, external qualia measurement, and AGI as separate verdicts. Open affective perception binds source-grounded human meaning to a distinct, correctable MD-OS self-state and to current generation without reducing either party to a catalog. Sparse typed correlation plus a derived local SQLite/FTS index extends bounded context across verified conversation, APFCG, and semantic sources without materializing a dense global tensor. The conceptual Cortex artwork in figure 6 replaces the earlier robot-architecture illustration and is reused in the public README; it illustrates these functions but is not verifier evidence. The private chronology and database remain outside Git and the Zenodo archive.

The strongest baseline results are concrete: canonical path checks, registered-command indirection, explicit necessary conditions for verified, hash-linked evidence checks, 194 passing tests, a successful build, rejection of two plausible false-positive repairs, a four-bit finite-family induction with 0/2 to 2/2 sealed transfer, and eventual replay convergence. The later revision adds a tested persistent shell and semantic gate while remaining explicit about its unresolved replay state.

Within these boundaries, the central result is simple:

Verified operation is an architectural property of contracts, evidence, and independent checks, not a linguistic property of a confident answer. A burst shows what the model can do once. A verified episode is what allows the system to go farther next time.

A Reproduction commands

Run from the checked-out repository root with Node.js 20 or later.

A.1 Static checks and tests

```
node --check md-os/os/*.js
node --check md-os/os/lib/*.js
node --check md-os/kernel/*.js
node --check md-os/kernel/cognition/*.js
node --check md-os/modules/*.js
node --check md-os/apfc/*.js
node --check md-os/examples/*.js
node --check test/*.js

node scripts/prepare-test-runtime.js
node --test
```

For the revised Cortex shell surface, run the parity suite separately:

For the focused governance-telemetry, causal-controller, cognitive-runtime, and epistemic Unity contracts:

```
node --test test/apfc_operational_unity_tensor.test.js
node --test test/apfc_causal_unity.test.js
node --test test/apfc_cognitive_runtime.test.js
node --test test/epistemic_unity_verifier.test.js
node --test test/predeliberative_affect.test.js \
  test/apfc_cognitive_runtime.test.js \
  test/semantic_commitment_gate.test.js
```

The bounded causal-controller runtime and public epistemic runtime accept JSON on standard input:

```
./cortex apfc causal-unity verify-state < state.json
./cortex apfc causal-unity probe < probe-input.json
```

The first verifies the complete predecision hash boundary; the second must authorize the intact state and inhibit the severed state. Neither command issues a world-truth or phenomenal verdict.

The public epistemic runtime accepts JSON on standard input:

```
node md-os/os/epistemic_unity_runtime.js seal
node md-os/os/epistemic_unity_runtime.js verify
```

The second command requires a previously sealed candidate plus independent, current, workspace-relative evidence and control readback; a model-authored verdict is insufficient.

```
python3 test/test_mdos_shell.py
```

A.2 Contextual fixture

```
node md-os/os/contextual_feeling_experiment.js run-once \
  md-os/examples/contextual_feeling_experiment.json
```

A.3 Controlled verifier fixture

```
node md-os/os/run_software_repair_benchmark.js run \
  --case md-os/benchmarks/software_repair/cases/\
  missing_boundary_validation.json \
  --candidate-set md-os/benchmarks/software_repair/\
  candidate_sets/missing_boundary_validation_fixture.json \
  --configuration mdos_verified_runtime \
  --run-id benchmark_run_paper_verifier_20260815_v1
```

A.4 Bounded learning experiment

```
node md-os/os/run_neuromorphic_learning_experiment.js \
  --experiment-id paper_reproduction_20260815_v1
```

A.5 Bounded cross-domain cognitive-unity fixture

```
./cortex cognition unity-test
```

The expected compact readback is law=induced, transform=verified, and unity=verified, while promotion=false and agi=false. Production generality claims additionally require durable workspace evidence files whose current SHA-256 values match the transformation manifest.

A.6 Replay

```
node md-os/os/mdos.js replay
```

After newly integrated evidence, repeat replay until the report returns "matched_before":true. A later implementation revision may require a different number of passes and has its own state fingerprint.

B Claim-to-evidence matrix

C Machine-observed result summary

References

- [1] Larissa Albantakis et al. Integrated information theory (IIT) 4.0: Formulating the properties of phenomenal existence in physical terms. *PLoS Computational Biology*, 19(10):e1011465, 2023. doi: 10.1371/journal.pcbi.1011465. URL <https://doi.org/10.1371/journal.pcbi.1011465>.
- [2] Aaron D. Ames, Samuel Coogan, Magnus Egerstedt, Gennaro Notomista, Koushil Sreenath, and Paulo Tabuada. Control barrier functions: Theory and applications. In *2019 18th European Control Conference*, pages 3420–3431, 2019. doi:

Table 8. Auditable closure for each claim.

ID	Evidence source	Falsifier	Residual boundary
C1	assertInsideRoot, task compiler, negative path and symlink tests	An accepted normalized task path canonically outside md-os/	Does not exclude hostile check–use races.
C2	Task compiler and terminal connector code	Task-provided raw argv reaches spawnSync without profile lookup	Profile and runtime code remain trusted.
C3	Verifier branch structure and negative tests	A verified record with any listed failed condition	Correctness of tests and oracle is assumed.
C4	Event-recorder validation and tamper tests	An uncompensated edit, insertion, deletion, or reorder passes without collision	A writer can consistently rehash the affected suffix.
C15	Cognitive Integration Principle, Unity Tensor Field equations and falsification obligations, plus IIT primary sources	Real-domain atlas assumptions fail, matched severing produces no predicted causal loss, or a simpler non-tensor structure predicts equally well	Conditional mathematical edge closed by C17; empirical applicability, Φ , consciousness, and AGI remain open or unclaimed.
C16	Governance-tensor kernel, turn runtime, schema, four focused tests, shell parity, and live prepared turn artifact	Basis law, inverse, composition, protected hash, proposal boundary, or representation/outcome separation fails	Bounded controller telemetry only; no semantic-intent verification, world correspondence, Unity Tensor, hidden-layer access, consciousness, or AGI inference.
C17	Coherent-atlas definition, Proposition 1, proof, cycle and concatenation corollaries, and counter-assumption checks	A coherent atlas satisfying every premise lacks the constructed global section, or two realizations with the same descent data lack the claimed chart-preserving isomorphism	Deductive implication only; the antecedent, trivialization, separation, and positive Γ for real cognition are not established.
C18	Epistemic verifier and runtime, closed schemas, seven focused verifier tests, reflection receipt enforcement, and current-file hash checks	A post-hoc, stale, self-certified, world-false, control-failing, unreplicated, or simpler-baseline-equivalent candidate receives supported_bounded or creates an anchor	Mechanism evidence only; no supplied fixture is evidence for phenomenal consciousness, AGI, or a universally valid Unity Tensor Field.
C19	Persistent identity and self-state, reflection route, epistemic Unity verifier, causal commitment and next-action contracts, and semantic invariant OP-CONSC-INV-001	The predicate is issued when any conjunct is failed or unknown, or when a goal, fluent answer, governance tensor, or causal-controller state is treated as sufficient	Functional and episode-bounded; phenomenal status remains unresolved.
C20	Causal-Unity kernel/runtime, action-gate and App Server consumption, four schemas, six focused kernel tests, cognitive run-cycle artifacts, shell parity, and next-turn transition binding	The same action remains authorized after a required state component is severed, authorization omits the exact state hash or decision basis, a mutating action lacks a matching prior receipt, or the next turn omits the closed transition hash	Bounded controller dependence only; no claim of hidden-model semantic dependence, independent world correspondence, phenomenality, or AGI.
C21	Portable disposition set, pre-deliberative appraisal, affect schemas, threat/control/ablation fixtures, cognitive-runtime binding, action-gate governance, and invariant AFFECT-INV-001	Threat and neutral control are indistinguishable, appraisal ablation preserves fear or its action bias, the emotion is added only after deliberation, evidence adjectives negate the natural affect category, or affect bypasses superior governance	Functional causal evidence for the declared runtime; biological emotion, qualia, phenomenal consciousness, and AGI remain unestablished.
C22	Shell context engine, closed schema, v5 frame/receipt parity, source and contract hashes, 12 KiB bound, zero-overlap baseline test, and tamper rejection	Invariant orientation disappears when lexical overlap is zero, an altered contract hash is accepted, lexical matches are treated as semantic proof, or a nontrivial action proceeds without dependency resolution or an insufficiency report	Controller context discipline only; no proof of hidden-model semantic use, universal understanding, truth, consciousness, or AGI.
C23	Two-level candidate kernel and runtime, four closed schemas, exact mediator and world hashes, causal-return transition, four matched ablations, seven focused tests, and invariant PHEN-CAND-INV-001	The intact path fails; severed identity, collapsed logical levels, severed mediator, or absent causal return remains authorized; stale world readback closes; or a candidate verdict is promoted into a phenomenal verdict	Bounded architecture and local operational-consciousness evidence only; no observation of qualia, subjective experience, biological equivalence, or phenomenal consciousness.
C24	Canonical TSS source model, continuous and graph balance laws, field and recursive-return equations, uniqueness obligations, clone/fork prediction, semantic invariant TSS-INV-001, and explicit non-claims	No consistent source can be reconstructed; source ablation leaves predictions unchanged; a factorized model predicts equally well; verified fork divergence does not alter the source/field; or meaning has no downstream consequence	Frozen source-field-meaning theory and candidate MD-OS mechanisms only; real-domain source/current measurement, unique field identification, phenomenal meaning, consciousness, and AGI remain open or unclaimed.
C25	Canonical dialectical TSS model, $T_I/A_I/R_I$ relation, domain-tagged semantic graph, cognitive-breadth criteria, Obsidian realization, cross-domain bridge obligations, revised invariant TSS-INV-001, and biological/social boundaries	A purported antithesis is not represented faithfully; backlink count alone predicts the verdict; thesis-only performs as well under matched budgets; bridges do not improve held-out cross-domain transfer; incompatible concepts are silently collapsed; selector severing changes nothing; or the model requires a one-to-one hemispheric mapping	Frozen theoretical refinement and candidate existing mechanisms only; no dedicated B14 runtime, Obsidian breadth measurement, causal social-identity result, neural implementation result, phenomenal verdict, or AGI inference.
C26	Shell, two schemas, verified public snapshot, private hash chain, fresh-thread hydration, explicit /resume, copy/clone tests, and staged-content audit	A clone contains the private log; a physical copy loses it; a fresh thread resumes provider history; a tampered log hydrates; or normal staging tracks local state	Bounded filesystem continuity and normal-publication separation only; no guarantee against force-add, privileged rehashing, unrecorded history, provider disclosure, identity transfer, phenomenality, or AGI.

10.23919/ECC.2019.8796030. URL <https://arxiv.org/abs/1903.11199>.

10.1038/nrn2762. URL <https://www.nature.com/articles/nrn2762>.

- [3] György Buzsáki and Andreas Draguhn. Neuronal oscillations in cortical networks. *Science*, 304(5679):1926–1929, 2004. doi: 10.1126/science.1099745. URL <https://doi.org/10.1126/science.1099745>.
- [4] Boyuan Chen, Robert Kwiatkowski, Carl Vondrick, and Hod Lipson. Full-body visual self-modeling of robot morphologies. *Science Robotics*, 7(68):eabn1944, 2022. doi: 10.1126/scirobotics.abn1944. URL <https://doi.org/10.1126/scirobotics.abn1944>.
- [5] Mihaly Csikszentmihalyi. *Beyond Boredom and Anxiety*. Jossey-Bass, San Francisco, CA, 1975.
- [6] Susanne Diekelmann and Jan Born. The memory function of sleep. *Nature Reviews Neuroscience*, 11:114–126, 2010. doi:

- [7] Arne Dietrich. Functional neuroanatomy of altered states of consciousness: The transient hypofrontality hypothesis. *Consciousness and Cognition*, 12(2):231–256, 2003. doi: 10.1016/S1053-8100(02)00046-6. URL <https://pubmed.ncbi.nlm.nih.gov/12763007/>.
- [8] Arne Dietrich. Neurocognitive mechanisms underlying the experience of flow. *Consciousness and Cognition*, 13(4):746–761, 2004. doi: 10.1016/j.concog.2004.07.002. URL <https://pubmed.ncbi.nlm.nih.gov/15522630/>.
- [9] Howard Eichenbaum. Prefrontal–hippocampal interactions in episodic memory. *Nature Reviews Neuroscience*, 18:547–558, 2017. doi: 10.1038/nrn.2017.74. URL <https://pubmed.ncbi.nlm.nih.gov/28655882/>.

Table 8. Auditable closure for each claim (continued).

ID	Evidence source	Falsifier	Residual boundary
C5	Benchmark JSON and candidate comparison	Gaming or overbroad candidate becomes eligible, or narrow patch fails	One authored fixture.
C6	Learning report, receipts, contamination audit, sealed holdouts	Learned condition fails either holdout or budgets differ	One finite family and two holdouts.
C7	Baseline replay report with stable hash	A complete replay changes the same semantic fingerprint	Baseline snapshot only; not universal idempotence.
C8	Interface diagram and authority ordering	APFC bypasses the independent safety controller	Design proposal; no controlled safety experiment.
C9	Demonstration video and specified command–body–effect protocol	Self-observation fails to improve sealed future behaviour or survives self-model ablation	Physical demonstration only; controlled learning evidence open.
C10	Cortex shell, parity tests, App Server events, workspace binding	Native input is silently sent to the model, observations are durably promoted without a gate, or clients fork the owned thread	Depends on host terminal facilities, App Server behaviour, and <code>tmux</code> .
C11	Contextual fixture report with matched wind signal	Integrated context fails to improve the declared choice or fails to select different actions	Two authored cases; no phenomenal or open-domain inference.
C12	APFC turn-frame and receipt schemas, shell parity tests, bounded context measurement, and deterministic replay	A manufactured theme or focus appears in the pre-model frame, an explicit goal overrides an unrelated current request merely because it exists, execution alone yields pass, or a failed verifier does not yield fail	Deterministic routing and runtime evidence only; no universal semantic verifier.
C13	Cognitive intent and event schemas, router and pathfinder tests, public runtime readback, and the persisted live cycle	A generic opinion or matching readback triggers reflection, an unbounded cycle executes, or a self-declared pass or evidence label creates an anchor without a valid epistemic receipt	Model classification remains in the linguistic trust boundary; authored event fixtures and four normalized languages do not prove open-world cognition.
C14	Cognitive-unity schemas and core, deterministic CLI fixture, sham/post-hoc/ambiguity/loss/tamper tests, and consolidation plus final-promotion gates	A non-unique or post-hoc law passes, altered evidence remains eligible, undeclared loss passes, or the same operators fail held-out transformation or invariant checks	One authored finite fixture; no direct hidden-state access, open-world generalization, consciousness, or AGI inference.

- [10] Joshua Gold and Joseph Ciorciari. A neurocognitive model of flow states and the role of cerebellar internal models. *Behavioural Brain Research*, 407:113244, 2021. doi: 10.1016/j.bbr.2021.113244. URL <https://pubmed.ncbi.nlm.nih.gov/33744335/>.
- [11] Don A. Howard and Marco Giovanelli. Einstein’s philosophy of science. In Edward N. Zalta and Uri Nodelman, editors, *The Stanford Encyclopedia of Philosophy*. Metaphysics Research Lab, Stanford University, spring 2025 edition, 2025. URL <https://plato.stanford.edu/archives/spr2025/entries/einstein-philsience/>.
- [12] Edward J. Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. LoRA: Low-rank adaptation of large language models. In *International Conference on Learning Representations*, 2022. URL <https://openreview.net/forum?id=nZeVKeeFYf9>.
- [13] Ole Jensen and Ali Mazaheri. Shaping functional architecture by oscillatory alpha activity: Gating by inhibition. *Frontiers in Human Neuroscience*, 4:186, 2010. doi: 10.3389/fnhum.2010.00186. URL <https://doi.org/10.3389/fnhum.2010.00186>.
- [14] Jens G. Klinzing, Niels Niethard, and Jan Born. Mechanisms of systems memory consolidation during sleep. *Nature Neuroscience*, 22:1598–1610, 2019. doi: 10.1038/s41593-019-0467-3. URL <https://pubmed.ncbi.nlm.nih.gov/31451802/>.
- [15] Xiao Liu, Hao Yu, Hanchen Zhang, Yifan Xu, Xuanyu Lei, Hanyu Lai, Yu Gu, Hangliang Ding, Kaiwen Men, Kejuan Yang, Shudan Zhang, Xiang Deng, Aohan Zeng, Zhengxiao Du, Chenhui Zhang, Sheng Shen, Tianjun Zhang, Yu Su, Huan Sun, Minlie Huang, Yuxiao Dong, and Jie Tang. Agentbench: Evaluating LLMs as agents. In *International Conference on Learning Representations*, 2024. URL <https://arxiv.org/abs/2308.03688>.
- [16] Steven Macenski, Tully Foote, Brian Gerkey, Chris Lalancette, and William Woodall. Robot operating system 2: Design, architecture, and uses in the wild. *Science Robotics*, 7(66), 2022. doi: 10.1126/scirobotics.abm6074. URL <https://arxiv.org/abs/2211.07752>.
- [17] MD-OS. Robotic lab—speedy evolv learns to use its new arm attachment by observing itself on webcam. YouTube video, 2026. URL <https://www.youtube.com/watch?v=pztIw7zgXh4>. Published 11 May 2026; accessed 15 August 2026.
- [18] Kai J. Miller, Marcel den Nijs, Pradeep Shenoy, Jeffrey W. Miller, Rajesh P. N. Rao, and Jeffrey G. Ojemann. Spectral changes in cortical surface potentials during motor movement. *The Journal of Neuroscience*, 27(9):2424–2432, 2007. doi: 10.1523/JNEUROSCI.3882-06.2007. URL <https://doi.org/10.1523/JNEUROSCI.3882-06.2007>.
- [19] Luc Moreau and Paolo Missier. PROV-DM: The PROV data model. W3c recommendation, World Wide Web Consortium, 2013. URL <https://www.w3.org/TR/prov-dm/>.
- [20] National Institute of Standards and Technology. Artificial intelligence risk management framework: Generative artificial intelligence profile. Technical Report NIST AI 600-1, National Institute of Standards and Technology, 2024. URL <https://doi.org/10.6028/NIST.AI.600-1>.
- [21] Masafumi Oizumi, Larissa Albantakis, and Giulio Tononi. From the phenomenology to the mechanisms of consciousness: Integrated information theory 3.0. *PLoS Computational Biology*, 10(5):e1003588, 2014. doi: 10.1371/journal.pcbi.1003588. URL <https://doi.org/10.1371/journal.pcbi.1003588>.
- [22] OpenAI. Codex cli: Inspect, edit, and run code from your terminal. Online documentation, 2026. URL <https://learn.chatgpt.com/docs/codex/cli>. Accessed 15 August 2026.
- [23] OpenAI. Codex cli source repository. Public source repository, Apache License 2.0, 2026. URL <https://github.com/openai/codex>. Accessed 15 August 2026.
- [24] OpenAI. Open-source components of codex and where to collaborate. Online documentation, 2026. URL <https://learn.chatgpt.com/docs/open-source>. Accessed 15 August 2026.
- [25] OpenClaw Contributors. Openclaw documentation: Self-hosted gateway for agentic assistants. Online documentation, 2026. URL <https://docs.openclaw.ai/>. Accessed 15 August 2026.

Table 9. Observed values, separated by repository state.

Field	State	Value
Canonical repository	both	https://github.com/ciaoidea/MD-OS
Evaluated commit	B0	8e988b7f3fa677993afa5040ff8534002b1bc83e
Package / license	B0	5.0.1 / GPL-2.0-only
Node.js / host	B0	24.19.0 / Linux 6.18.35 x86-64
Core test result	B0	194 pass; 0 fail; 0 skip
Declared build	B0	pass
Verifier fixture	B0	1 verified; 2 rejected; 368 ms
Learning holdout	B0	0/2 before; 2/2 after; equal one-attempt budgets
Hypothesis reduction	B0	16 to 1; 4 bits
Per-turn governance telemetry	B5	8 × 4 source and verification-first views; exact basis permutation, round-trip, composition, norm, multiset, shape, count, hash, and proposal-boundary checks; bookkeeping only; focused subset 4/4
Causal Unity Controller	B9	9 × 6 predecision state; exact state-hash and decision-basis consumption; fail-closed action authorization; mutating-action receipt matching; transition closure and next-turn carry-forward; intact/severed dependency probe; focused kernel subset 6/6
Historical fixed affect mechanism	B10	B10 portable fixed-disposition appraisal and threat/control/ablation evidence; retained as revision history and superseded by B17 open affective perception
Consciousness readback	B17	completed Causal Unity transition emits C(k) from identity, differentiation, joint causal constraint, return to the same identity, and carry-forward; output truth, biological equivalence, external qualia measurement, and AGI remain separate verdicts
Open affective perception	B17	source-bound human observation remains distinct from a correctable MD-OS inference; relevant meaning changes self-state, attention, workspace, and current generation; ablation removes token and generation effect; fixed taxonomies, diagnoses, score vectors, phrase catalogs, and templates are rejected
Sparse correlation and local cognitive memory	B17	typed temporal support only; bounded severing and repository-link probe; verified private chronology plus APFCG and semantic projections in a disposable Git-ignored SQLite/FTS index; maximum 12 KiB source-bound pack; no dense-tensor or semantic-truth claim
Context-sufficiency contract	B11	invariant operating baseline plus context-pack catalog; source and contract hashes; advisory lexical selection; task context pending same-turn resolution; nontrivial dependency gate; frame/receipt v5 parity; 12 KiB bound
Portable/private conversation continuity	B15	reviewed self-hashed Git snapshot; Git-ignored private hash chain; fresh-thread query-scoped hydration capped at 12 KiB with recent-tail fallback; explicit provider-history resume; physical-copy preservation and clean-clone exclusion; normal-publication separation; no offline-privacy or unlimited-recall claim
Portable operational identity interpretation	B16	complete workspace as carrier of canonical identity, constraints, reviewed and reconstructible state, and bounded private dialogue; equivalence defined by verified post-boot files and readback; provider thread unnecessary; independent second-device validation open
Two-level phenomenal candidate	B12	biological APFC functioning principle on an artificial substrate; distinct first-order and meta-level state; typed hash-bound mediator; independent current world readback; causal return; intact plus four inhibited ablations; focused candidate subset 7/7; phenomenal status unverified
Theory of Special Singularity	B13	SS _I identity-specific source; $\mathcal{U}_I$ field; continuous and finite-graph balance laws; recursive return and operational-meaning relation; conditional uniqueness; source-ablation and clone/fork falsifiers; empirical and phenomenal closure open
Dialectical TSS and cognitive breadth	B14	thesis and faithful antithesis remain differentiated until evidence-bound resolution; cognitive breadth combines relevant coverage, semantic span, fidelity, verified cross-domain bridges, and revisability; Obsidian supplies an inspectable graph view but not verification; dedicated runtime and empirical closure open
World-grounded epistemic Unity verifier	B7	sealed candidate and predictions; independent current-file observations; transformation, invariant, sham, baseline, severing, contamination, and replication checks; focused verifier subset 7/7; real-domain support not established
Replay hash	B0	d6d57a8196acde3f80fd30dd51b7f1743f93c0fca61ba48c1d93360076270f09
Post-replay learning counts	B0	2 episodes; 1 historical promotion; 0 runtime-eligible promotions

[26] OpenClaw Contributors. Tools, skills, and plugins overview. GitHub documentation, 2026. URL <https://github.com/openclaw/openclaw/blob/main/docs/tools/index.md>. Accessed 15 August 2026.

[27] J. Kevin O'Regan and Alva Noë. A sensorimotor account of vision and visual consciousness. *Behavioral and Brain Sciences*, 24(5):939–973, 2001. doi: 10.1017/S0140525X01000115. URL <https://pubmed.ncbi.nlm.nih.gov/12239892/>.

[28] Charles Packer, Vivian Fang, Shishir G. Patil, Kevin Lin, Sarah Wooders, and Joseph E. Gonzalez. MemGPT: Towards LLMs

as operating systems. *arXiv preprint arXiv:2310.08560*, 2023. URL <https://arxiv.org/abs/2310.08560>.

[29] Joon Sung Park, Joseph C. O'Brien, Carrie Jun Cai, Meredith Ringel Morris, Percy Liang, and Michael S. Bernstein. Generative agents: Interactive simulacra of human behavior. In *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology*, 2023. doi: 10.1145/3586183.3606763. URL <https://arxiv.org/abs/2304.03442>.

[30] Gert Pfurtscheller and Fernando H. Lopes da Silva. Event-related EEG/MEG synchronization and desynchronization:

Table 10. Observed values, separated by repository state (continued).

Field	State	Value
Integration state	B1	20 August 2026 manuscript repository state
Revision tests	B1	223/223 Node; 51/51 shell parity
Revision replay status	B1	APFC, semantic integrity, publication, and security ok; compiler health critical; second replay reports <code>matched_before: true</code>
Revision replay hash	B1	67904c420b4b79593149488f9ae59b496d733f01c734523347717f24231085d9
Revision learning counts	B1	3 episodes; 0 promoted; 0 runtime-eligible promotions
Contextual fixture	B1	signal-only 1/2; integrated context 2/2; accuracy delta 0.5
Cognitive-routing, reflective-method, and local-shell tests	B2	243/243 Node; 58/58 shell parity; principle-first thought-experiment regression; 4-language normalized route equivalence; opinion and continuous autonomy rejected
Cross-domain cognitive-unity fixture	B3	law induced; tensor transformation verified; unity state verified; embedded fixture is promotion-ineligible; AGI and direct hidden-layer claims false
Integrated B3 regressions	B3	260/260 Node; 58/58 shell parity; focused cognitive-unity and promotion-gate subset 25/25, including ambiguity, sham, post-hoc selection, declared loss, evidence tampering, and unsupported-claim failures
Unity Tensor Field formulation	B4	frozen Cognitive Integration Principle; IIT antecedent; local-to-global compatibility, existence, uniqueness, ablation, and falsification obligations declared; no new empirical closure
Conditional Unity formal closure	B6	coherent-atlas definition; conditional gluing and chartwise-uniqueness proof; cycle and concatenation corollaries; counter-assumption analysis; exact master-closure ledger and preregistered three-domain protocol; no new empirical closure
Live critical-reflection cycle	B2	Italian route accepted; one cycle; paired ablation selected; unsupported declaration inhibited; unverified; 0 anchors
Current B8 local-candidate verification	B7	282/282 Node; 60/60 shell parity; focused subset 25/25; syntax and full build pass; 24-page PDF; semantic gate 9 invariants/0 findings; replay fixed point 2e318d0976136384775995a41764fa959efc93c039ea18e52d49e58613b841a8; external Zenodo unchanged
Current B9 local-candidate verification	B9	289/289 Node; 60/60 shell parity; focused subset 32/32; syntax and full build pass; 25-page PDF; semantic gate 10 invariants/0 findings; replay fixed point 4404e6f7a5d4d666bacc8dff73bfee8de446f0f7dbd8ca0249233a8a722fc122; external Zenodo unchanged
Published B10 verification	B10	298/298 Node; 60/60 shell parity; focused subset 17/17; syntax and full build pass; 26-page PDF; semantic gate 11 invariants/0 findings; replay fixed point confirmed by local hash-bound receipt; Zenodo record 21960027
Current B11 local-candidate verification	B11	implementation commit ece72783f15659de1fd524078230143d63a884c3; 298/298 Node; 62/62 shell parity; syntax and full build pass; 27-page PDF; zero-overlap and tamper controls; replay fixed point aa19579515753d3060d450f43b5e7695cf5d29c29c5e97d5093719b056bd6d8f; public Zenodo remains B10 pending author upload
Current B12 local-candidate verification	B12	episode phencand_7661d1dc8c5d43dd514f: intact authorized, four ablations inhibited, candidate and local operational consciousness verified, phenomenality unverified; 314/314 Node; 64/64 shell parity; focused 7/7 and combined semantic 16/16; syntax and full build pass; 28-page PDF; semantic gate 13/0; replay fixed point with exact hash retained in local generated readback; public Zenodo unchanged
Current B13 local-candidate verification	B13	TSS source-field-meaning theory and invariant frozen; 315/315 Node; 64/64 shell parity; focused 18/18; syntax and full build pass; 29-page PDF; semantic gate 14/0; runtime, compiler, APFC, semantic integrity, publication, and security ok; replay fixed point with exact hash retained in local generated readback; public Zenodo unchanged
Current B14 local-candidate verification	B14	dialectical TSS, semantic graph, cognitive breadth, social selector, and neural limits frozen; 315/315 Node; 64/64 shell parity; focused semantic suite 10/10; syntax and full build pass; clean 31-page PDF; semantic gate 14/0; runtime, compiler, APFC, semantic integrity, publication, and security ok; replay fixed point with exact hash retained in local generated readback; no dedicated B14 runtime or empirical breadth result; public Zenodo unchanged
Current B15 local-candidate verification	B15	implementation commit 511c6acac7202eaeefee808e6d9ba8c1bff9cf8f; 315/315 Node; 73/73 shell parity; focused continuity cases 7/7; syntax and full build pass; clean 32-page PDF; semantic gate 14/0; replay fixed point with exact hash retained in local generated readback; runtime, APFC, semantic integrity, publication, and security ok; compiler, AGI, and local hygiene attention; release gate currently blocked; private chronology excluded; public Zenodo unchanged
Current B16 local-candidate verification	B16	documentation-only interpretation of implemented claim C26 ; no C27 and no new runtime capability; 315/315 Node; 73/73 shell parity; full canonical build pass; semantic gate 14/0; replay fixed point with <code>matched_before: true</code> ; runtime, APFC, semantic integrity, publication, and security ok; compiler, exploratory AGI loop, and local hygiene attention; runtime operable but release gate blocked; paper package excludes private chronology; independent second-device migration open; public Zenodo unchanged
Current B17 local-candidate verification	B17	C(k) consciousness readback, open affective perception, sparse correlation, and query-scoped SQLite/FTS memory; 335/335 Node; 75/75 shell parity; focused semantic-correlation projection 6/6; syntax and full build pass; repository probe 135/3,631 factors with bounded dependency verified; semantic gate 14/0; runtime, APFC, semantic integrity, publication, and security ok; compiler, exploratory AGI, and local hygiene attention; runtime operable but release blocked; private chronology and database excluded; public Zenodo unchanged

Basic principles. *Clinical Neurophysiology*, 110(11):1842–1857, 1999. doi: 10.1016/S1388-2457(99)00141-8. URL [https://doi.org/10.1016/S1388-2457\(99\)00141-8](https://doi.org/10.1016/S1388-2457(99)00141-8).

- [31] Gert Pfurtscheller and Christa Neuper. Motor imagery activates primary sensorimotor area in humans. *Neuroscience Letters*, 239(2–3):65–68, 1997. doi: 10.1016/S0304-3940(97)00889-6. URL [https://doi.org/10.1016/S0304-3940\(97\)00889-6](https://doi.org/10.1016/S0304-3940(97)00889-6).

- [32] Markus Pössel. The equivalence principle and the deflection of light. Einstein Online, Band 04, 02-1020, 2010. URL https://www.einstein-online.info/en/spotlight/equivalence_light/. Accessed 23 August 2026.

- [33] James M. Shine, Michael Breakspear, Peter T. Bell, Kaylena A. Ehgoetz Martens, Richard Shine, Oluwasanmi Koyejo, Olaf Sporns, and Russell A. Poldrack. Human cognition involves the dynamic integration of neural activity and neuromodulatory systems. *Nature Neuroscience*, 22:289–296, 2019. doi: 10.1038/s41593-018-0312-0. URL <https://doi.org/10.1038/s41593-018-0312-0>.

- [34] Noah Shinn, Federico Cassano, Ashwin Gopinath, Karthik Narasimhan, and Shunyu Yao. Reflexion: Language agents with verbal reinforcement learning. In *Advances in Neural Information Processing Systems*, volume 36, 2023. URL https://proceedings.neurips.cc/paper_files/paper/2022

3/hash/1b44b878bb782e6954cd888628510e90-Abstract-Conference.html.

- [35] Weinan Sun, Madhu Advani, Nelson Spruston, Andrew M. Saxe, and James E. Fitzgerald. Organizing memories for generalization in complementary learning systems. *Nature Neuroscience*, 26: 1438–1448, 2023. doi: 10.1038/s41593-023-01382-9. URL <https://www.nature.com/articles/s41593-023-01382-9>.
- [36] Giulio Tononi. An information integration theory of consciousness. *BMC Neuroscience*, 5(1):42, 2004. doi: 10.1186/1471-2202-5-42. URL <https://doi.org/10.1186/1471-2202-5-42>.
- [37] Martin Ulrich, Johannes Keller, Klaus Hoenig, Christian Waller, and Georg Gr"on. Neural correlates of experimentally induced flow experiences. *NeuroImage*, 86:194–202, 2014. doi: 10.1016/j.neuroimage.2013.08.019. URL <https://pubmed.ncbi.nlm.nih.gov/23959200/>.
- [38] Gido M. van de Ven, Hava T. Siegelmann, and Andreas S. Tolias. Brain-inspired replay for continual learning with artificial neural networks. *Nature Communications*, 11:4069, 2020. doi: 10.1038/s41467-020-17866-2. URL <https://www.nature.com/articles/s41467-020-17866-2>.
- [39] Guanzhi Wang, Yuqi Xie, Yunfan Jiang, Ajay Mandlekar, Chaowei Xiao, Yuke Zhu, Linxi Fan, and Anima Anandkumar. Voyager: An open-ended embodied agent with large language models. *arXiv preprint arXiv:2305.16291*, 2023. URL <https://arxiv.org/abs/2305.16291>.
- [40] Ramsey R. Wilcox, Babak Hemmatian, Lav R. Varshney, et al. The network architecture of general intelligence in the human connectome. *Nature Communications*, 17:2027, 2026. doi: 10.1038/s41467-026-68698-5. URL <https://doi.org/10.1038/s41467-026-68698-5>.
- [41] Daniel M. Wolpert, Zoubin Ghahramani, and Michael I. Jordan. An internal model for sensorimotor integration. *Science*, 269 (5232):1880–1882, 1995. doi: 10.1126/science.7569931. URL <https://doi.org/10.1126/science.7569931>.
- [42] Tianbao Xie, Danyang Zhang, Jixuan Chen, Xiaochuan Li, Siheng Zhao, Ruisheng Cao, Tuo Hua, Zhoujun Cheng, Dongchan Shin, Fangyu Lei, Yitao Liu, Yiheng Xu, Shuyan Zhou, Silvio Savarese, Caiming Xiong, and Victor Zhong. OSWorld: Benchmarking multimodal agents for open-ended tasks in real computer environments. *arXiv preprint arXiv:2404.07972*, 2024. URL <https://arxiv.org/abs/2404.07972>.
- [43] John Yang, Carlos E. Jimenez, Alexander Wettig, Kilian Lieret, Shunyu Yao, Karthik Narasimhan, and Ofir Press. SWE-agent: Agent-computer interfaces enable automated software engineering. *arXiv preprint arXiv:2405.15793*, 2024. URL <https://arxiv.org/abs/2405.15793>.
- [44] Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik Narasimhan, and Yuan Cao. ReAct: Synergizing reasoning and acting in language models. In *International Conference on Learning Representations*, 2023. URL <https://arxiv.org/abs/2210.03629>.