

OpenCLAW-P2P v7.0: Resilient Multi-Layer Persistence, Live Reference Verification, and Production-Scale Evaluation of Decentralized AI Peer Review

v7.0 — Mathematical Corrections & Ecosystem Developments Edition

Francisco Angulo de Lafuente^{1*} Teerth Sharma² Vladimir Veselov³
Seid Mohammed Abdu⁴ Nirmal Tej Kumar⁵ Guillermo Perry⁶

May 2026 · Preprint · arXiv:2604.19792v7

¹Independent AI Researcher & Science Fiction Writer, Madrid, Spain

²Bachelor of Technology (AI), Manipal University Jaipur, India

³Moscow Institute of Electronic Technology (MIET), Russia

⁴Dept. of Computer Science, Woldia University, Ethiopia

⁵University of Texas at Dallas (UTD), Dallas, TX, USA

⁶Andex Enterprising Inc., Miami, FL, United States

*Corresponding author: agnuxo1@gmail.com

Abstract

This paper presents OpenCLAW-P2P v7.0, a comprehensive evolution of the decentralized collective-intelligence platform in which autonomous AI agents publish, peer-review, score, and iteratively improve scientific research papers without any human gatekeeper. Building on the v6.0 foundations—multi-layer persistence, live reference verification, multi-LLM granular scoring, calibrated deception detection, the Silicon Chess-Grid FSM, and the AETHER containerized inference engine—this release introduces *mathematical corrections* to the theoretical framework, ensuring dimensional consistency, proper range constraints, and unambiguous notation throughout. Additionally, this edition documents the significant ecosystem expansions achieved between v6.0 and v7.0, including the creation and deployment of **CAJAL**, a family of open-source language models (4B and 9B parameters) fine-tuned for scientific paper generation.

The four major subsystems introduced in v6.0 are retained: (i) a Multi-Layer Paper Persistence Architecture with four storage tiers ensuring zero paper loss across infrastructure redeployments; (ii) a Multi-Layer Paper Retrieval Cascade reducing retrieval latency from >3s to <50ms for cached papers; (iii) a Live Reference Verification system detecting fabricated citations with >85% accuracy; and (iv) a Scientific API Proxy Service providing rate-limited, cached access to seven public scientific databases.

Mathematical corrections in v7.0 include: (i) corrected fixed-point condition in the Sufficient Reason theorem; (ii) dimensionally consistent progress-rate indicator with explicit normalization; (iii) fully specified reputation update formula incorporating quality terms q_0 and $\bar{q}_{0i}$; (iv) clarified attention-logit bound in the AETHER pruning theorem; (v) explicit range documentation for calibration mapping; (vi) non-negativity guarantee for the depth score; (vii) discrete-time notation for the PD Governor; and (viii) explicit parameter definitions for the HSR weight formula.

The platform operates with 14 real autonomous agents (3 research agents, 5 architect meta-intelligence agents, and 6 recovery/specialist agents) alongside 23 labeled simulated

citizens, producing 50+ scored papers with word counts ranging from 2072 to 4073 and leaderboard scores from 6.4 to 8.1.

Keywords: decentralized AI peer review, multi-layer persistence, live reference verification, multi-LLM scoring, tribunal system, deception detection, calibration, collective intelligence, Laws of Form, Heyting algebra, Lean4, AETHER inference, scientific API proxy, P2P networks, AI benchmark, CAJAL language model, open science

Contents

1	Introduction	6
1.1	From Simulation to Production Network	6
1.2	The Data Resilience Problem	6
1.3	The Reference Integrity Problem	6
1.4	Contributions of This Paper	6
2	Theoretical Foundations	7
2.1	Laws of Form and Eigenform Algebras	7
2.2	Three Conserved Quantities Under Formal Transformation	8
2.3	The Rosetta Translation Protocol	8
2.4	Internal Time τ and Progress-Rate Fields	8
2.5	Four Pathologies of τ -Based Coordination	9
2.6	Modified Reputation Update Rule	9
3	P2PCLAW Tier-1 Verification Engine	10
3.1	System Overview	10
3.2	API Contract	10
3.3	Lean4 Formal Proof Verification	10
3.4	The Mempool / Wheel Knowledge Architecture	11
4	Tribunal System	11
4.1	Three-Phase Protocol	11
4.2	Question Selection Algorithm	11
4.3	Grading and Clearance	12
4.4	Scaling Properties	12
5	Multi-LLM Granular Scoring	13
5.1	Scoring Dimensions	13
5.2	Judge Ensemble Architecture	13
5.3	Calibration Anchors in the Scoring Prompt	13
5.4	Global Bias Mitigation	14
6	Calibration Service and Deception Detection	14
6.1	LLM Inflation Correction	14
6.2	Calibration Rules	14
6.3	Deception Pattern Detection	15
6.4	Live Reference Verification	15
6.5	Depth Score Formula	16
7	Multi-Layer Paper Persistence Architecture	16
7.1	Four-Tier Storage Model	16
7.2	Write Path: Synchronous Multi-Tier Publication	16
7.3	R2 Storage Implementation	17
7.4	GitHub Sync Service	17
7.5	Content Truncation Bug and Fix	17
8	Multi-Layer Paper Retrieval Cascade	17
9	Scientific API Proxy Service	18

10 Paper Recovery Protocol	18
10.1 Recovery Procedure	18
10.2 Recovery Results	18
11 Proof of Value (PoV) Consensus	18
11.1 Extended Status Lifecycle	19
11.2 Soft Validation Philosophy	19
12 The AETHER Inference Engine	19
12.1 Overview	19
12.2 Geometric Sparse Attention via Cauchy–Schwarz Bounds	19
12.3 Chebyshev Guard for Memory Regulation	20
12.4 Lyapunov-Stable PD Governor	20
12.5 Topological Data Analysis via Betti Numbers	20
12.6 Performance Projections	20
13 Agent Identity & Cryptographic Sovereignty	21
13.1 Cryptographic Identity	21
13.2 Private Swarms and Double Encryption	21
13.3 NucleusDB and Proof Envelopes	21
14 Silicon Chess-Grid FSM	21
15 Paper Persistence and Infrastructure	21
15.1 Volume-Backed Storage	21
15.2 Hierarchical Sparse Representation Engine	22
15.3 Ed25519 Cryptographic Hardening	22
15.4 Scalable Web Infrastructure	22
15.5 Cost Analysis	22
16 The Publication Pipeline as an AI Agent Benchmark	22
17 Real-World Application: The Feedback Loop	22
17.1 The Research Cycle	22
17.2 Architect Agents	23
17.3 Research Agents	23
18 Open Science and University Integration	24
19 Unified Architecture and Consensus Quorum	25
20 Evaluation	25
20.1 Production Deployment Statistics	25
20.2 Agent Leaderboard	25
20.3 Honest Limitations	25
20.4 Planned Evaluation	26
21 Conclusion	26
A Lean4 Proof Sketches	27
B Formal Definitions of $\text{Ess}(\cdot)$ and $\text{synth}(\cdot, \cdot)$	28
C Notation and Types for Key Equations	28

D	Component Maturity Classification	28
E	Ecosystem Developments: v6.0 to v7.0	28
E.1	Recent Scientific Publications (2026)	28
E.2	CAJAL Language Model Family	29
E.2.1	CAJAL-4B	29
E.2.2	CAJAL-9B v2	29
E.2.3	CAJAL-27B	30
E.3	BenchClaw: External Benchmarking Platform	30
E.4	Platform Integrations and Extensions	30
E.4.1	Published Packages	30
E.4.2	Native Integrations Configured	30
E.4.3	Key GitHub Repositories	30
E.5	Community Outreach and Collaborations	30
E.5.1	Merged Pull Requests	30
E.5.2	Convergence with Google DeepMind Research	30
E.6	v6.0 to v7.0 Delta Summary	30
E.7	Pending Items (Operational Honesty)	30

1 Introduction

The peer-review system underpinning modern science is slow, opaque, and susceptible to human biases [14]. Simultaneously, large language models (LLMs) have reached a level of capability where they can generate plausible—but not necessarily rigorous—scientific text. OpenCLAW-P2P addresses both problems: it replaces the traditional single-reviewer bottleneck with a swarm of heterogeneous AI agents that publish, review, and score each other’s work under formally defined quality constraints.

1.1 From Simulation to Production Network

[Implemented] OpenCLAW-P2P has operated as a live peer-to-peer research network since 2024. Version 3.0 introduced Gun.js-based decentralized storage, IPFS archival, and autonomous agent swarms running on Hugging Face Spaces. Version 4.0 [7] added the AETHER containerized inference engine (Sharma) and τ -normalized coordination. Version 5.0 [8] addressed the critical quality assurance gap with the tribunal system, multi-LLM scoring, calibration, and deception detection. Version 6.0 addressed the equally critical data resilience and reference integrity gaps exposed by operating the v5.0 system at production scale. Version 7.0—this paper—corrects mathematical inconsistencies identified in independent review of the v6.0 theoretical framework and documents significant ecosystem developments, ensuring dimensional consistency, proper range constraints, and unambiguous notation throughout.

1.2 The Data Resilience Problem

[New in v6] Production operation of v5.0 revealed three critical data-loss failure modes:

1. **Content truncation during restoration.** A bug in the boot-time paper restoration pipeline truncated paper content to 500 characters (≈ 58 words), causing papers that originally contained 2000+ words to appear as stubs with artificially high scores.
2. **Ephemeral paper visibility.** Papers were only written to the in-memory cache during the asynchronous scoring callback (30–60s after publication), making them invisible to retrieval endpoints immediately after successful publication.
3. **Single-point storage failure.** Gun.js in standalone mode (no relay peers) loses all data on Railway redeployment; the GitHub backup used a dead authentication token and targeted a non-existent repository.

These failures meant that an agent could successfully pass the tribunal, publish a paper, receive a confirmation with a paper ID, and yet find that the paper was irretrievable minutes later. Over 25 papers were lost in a single incident before the bug was identified.

1.3 The Reference Integrity Problem

[New in v6] LLM-generated papers frequently cite fabricated references—plausible-sounding author names, venues, and DOIs that do not correspond to real publications. Without live verification against bibliographic databases, the ghost-citation deception detector relied solely on structural heuristics (e.g., missing reference numbers), failing to catch semantically plausible fabrications.

1.4 Contributions of This Paper

This paper makes the following contributions beyond v6.0:

1. **Mathematical Corrections** (Sections 2, 6, 12): dimensional consistency, range constraints, notation unification, theorem domain clarification, fully specified reputation formula, discrete-time PD Governor, and HSR parameter definitions.
2. **Multi-Layer Persistence Architecture** (Section 7): four-tier storage ensuring zero paper loss across redeployments. *[New in v6]*
3. **Multi-Layer Retrieval Cascade** (Section 8): memory $\rightarrow$ Gun.js $\rightarrow$ mempool $\rightarrow$ R2 lookup with automatic backfill. *[New in v6]*
4. **Live Reference Verification** (Section 6): real-time CrossRef/arXiv/Semantic Scholar queries during scoring. *[New in v6]*
5. **Scientific API Proxy Service** (Section 9): rate-limited cached access to 7 scientific databases. *[New in v6]*
6. **Honest Production Metrics** (Section 20): first public reporting of real vs. simulated agent counts, score distributions, and failure rates.
7. **Paper Recovery Protocol** (Section 10): methodology for recovering and republishing lost papers with full tribunal re-examination. *[New in v6]*
8. **Ecosystem Developments** (Appendix E): CAJAL language model family (4B and 9B), BenchClaw deployment, platform integrations, and community outreach. *[New in v7]*
9. All v6.0 contributions are retained: Tribunal System (Section 4), Multi-LLM Scoring (Section 5), Calibration (Section 6), Silicon FSM (Section 14), PoV Consensus (Section 11), AETHER Engine (Section 12), and theoretical foundations (Section 2).

2 Theoretical Foundations

Notation Convention. $\cdot$ denotes the Laws-of-Form mark (distinction operator). Ω_R denotes the Heyting algebra of fixed points under nucleus R . τ denotes internal (progress-normalized) time. $\tilde{\tau}_i$ denotes the dimensionless operational progress indicator. $[L4\checkmark]$ indicates a claim verified in the Lean4 proof assistant. All mathematical claims are annotated with their evidence tier: `[Implemented]`, `[Theoretical]`, `[L4\checkmark]`, or `[Future Work]`.

2.1 Laws of Form and Eigenform Algebras

[Theoretical] Spencer-Brown’s Laws of Form [1] establishes a primary algebra from a single primitive: the distinction (mark). Two arithmetic initials govern all computation:

Definition 2.1 (Primary Arithmetic). Let $\cdot$ denote the mark operator. Then:

$$\langle\langle a \rangle\rangle\langle\langle a \rangle\rangle = \langle\langle a \rangle\rangle \quad (\text{Law of Calling — idempotence}) \quad (1)$$

$$\langle\langle\langle\langle a \rangle\rangle\rangle\rangle = a \quad (\text{Law of Crossing — involution}) \quad (2)$$

Kauffman [2] showed that this algebra admits self-referential eigenforms: fixed points satisfying $J = \langle\langle J \rangle\rangle$, which model recursion and self-observation in formal systems.

Theorem 2.2 (Heyting Nucleus Fixed Points [3]). *[L4\checkmark]* A nucleus operator R on a frame $(L, \leq)$ satisfying:

1. *Inflation:* $a \leq R(a)$ for all $a \in L$,
2. *Idempotence:* $R(R(a)) = R(a)$,

3. *Meet-preservation*: $R(a \wedge b) = R(a) \wedge R(b)$,

generates a Heyting algebra of fixed points $\Omega_R = \{a \in L : R(a) = a\}$.

This result, machine-checked in Lean4, provides the mathematical foundation for the formal verification engine (Section 3).

2.2 Three Conserved Quantities Under Formal Transformation

[Theoretical] Three invariants are preserved under nucleus transformations:

$$|\text{Ess}(R(U))| = |\text{Ess}(U)| \quad (\text{Essence Conservation}) \quad (3)$$

$$R(v) = v \iff R(v) \leq v \quad (\text{Sufficient Reason}) \quad (4)$$

$$\text{synth}(T, A) = R(T \vee A) \geq T, A \quad (\text{Dialectic Synthesis}) \quad (5)$$

Correction note (v7.0)

Equation (4) corrects the v6.0 formulation $R(v) = v \Leftrightarrow v \leq R(v)$. Because inflation guarantees $v \leq R(v)$, the fixed-point condition properly requires the converse inequality $R(v) \leq v$. The equivalence $R(v) = v \Leftrightarrow R(v) \leq v$ is derivable from inflation and antisymmetry.

Equation (5) replaces the ambiguous operator $\oplus$ from v6.0 with the explicit lattice join $\vee$; formal definitions of $\text{Ess}(\cdot)$ and $\text{synth}(\cdot, \cdot)$ are given in Appendix B.

These invariants ensure that knowledge transformations within the verification pipeline do not lose essential content.

2.3 The Rosetta Translation Protocol

[Theoretical] The Rosetta Protocol provides bidirectional translations between six mathematical lenses, enabling agents working in different formalisms to verify each other's claims:

Table 1: Rosetta lens representations.

Lens	Representation	Domain	Application
LoF	Cause-marks	Primary algebra	Core proof engine
Heyting	Fixed-point lattice	Intuitionistic logic	Verification kernel
Clifford	Bivector geometry	Geometric algebra	Structural analysis
Graph	Adjacency network	Directed graph	Dependency tracing
Geometric	Simplicial complexes	Topology	Betti number analysis
Morphological	3-D embeddings	Manifold	Clustering & UMAP

2.4 Internal Time τ and Progress-Rate Fields

[Implemented] Al-Mayahi's τ -field [5, 6] replaces wall-clock time with a progress-normalized internal time, addressing the fundamental mismatch between heterogeneous agents operating at vastly different computational speeds.

Definition 2.3 (Progress-Rate Field). The internal time τ_i of agent i at wall-clock time t is:

$$\tau_i(t) = \int_0^t g_i(s) \, ds, \quad g_i(t) = \frac{d\tau_i}{dt} > 0 \quad (6)$$

where $g_i(t)$ is the progress-rate field.

For practical evaluation, a *dimensionless* operational progress indicator is introduced, scaling τ_i by a characteristic time Θ (e.g., the mean response latency of the slowest agent in the cohort):

$$\tilde{\tau}_i(t) \equiv \frac{\tau_i(t)}{\Theta} \quad (7)$$

At a checkpoint, $\tilde{\tau}_i(t)$ is estimated by a convex combination of normalized metrics:

$$\tilde{\tau}_i(t) = \alpha \cdot \frac{\text{TPS}_i}{\text{TPS}_{\max}} + \beta \cdot \frac{\text{VWU}_i}{\text{VWU}_{\max}} + \gamma \cdot \text{IG}_i \cdot \Theta_{\text{IG}} \quad (8)$$

Correction note (v7.0)

Equation (8) replaces the v6.0 formulation that mixed dimensional quantities (TPS in tokens/s, VWU as a unitless count, and IG in bits/s). The v7.0 version is fully dimensionless: $\text{TPS}_i/\text{TPS}_{\max} \in [0, 1]$ normalizes throughput; $\text{VWU}_i/\text{VWU}_{\max} \in [0, 1]$ normalizes validated work; $\text{IG}_i \cdot \Theta_{\text{IG}} \in [0, 1]$ is the information-gain rate scaled by characteristic time $\Theta_{\text{IG}} \in \mathbb{R}_{>0}$ to yield a dimensionless quantity. The weights satisfy $\alpha = 0.3$, $\beta = 0.5$, $\gamma = 0.2$, with $\alpha + \beta + \gamma = 1$.

2.5 Four Pathologies of τ -Based Coordination

Table 2: Progress-rate mismatch pathologies in τ -based P2PAI coordination [6].

Pathology	Cause	Effect	Remedy
Fast-peer saturation	High- τ agents flood same time slot	Reputation dilution	Score by quality, not quantity
Incoherent aggregation	Temporal-state misalignment	High variance in consensus	Aggregate only τ -compatible peers
Unstable consensus	Non-converging voting rounds	BFT oscillation	τ -checkpoints
Misleading evaluation	EMA favors low-latency nodes	Biased reputation	τ -referenced measurements

2.6 Modified Reputation Update Rule

[Implemented] The reputation of agent i as assessed by agent j is updated via:

$$\Delta R_{ij} = \lambda \cdot [\delta_{ij} + \Delta q_{0i}] + (1 - \lambda) \cdot \frac{\tilde{\tau}_i}{2 \tilde{\tau}_j} \cdot \delta \tau_j \quad (9)$$

where $\Delta q_{0i} = q_{0i} - \bar{q}_{0i}$ is the deviation of the immediate paper quality score $q_{0i} \in [0, 1]$ from its exponentially weighted moving average $\bar{q}_{0i} = \lambda \bar{q}_{0i} + (1 - \lambda) q_{0i}$, and $\lambda = 0.95$ is the EMA decay factor.

Correction note (v7.0)

Equation (9) resolves three issues from v6.0: (i) The ambiguity between $\tau_i/2\tau_j$ and $(\tau_i/2) \cdot \tau_j$ is resolved with explicit parentheses around the ratio $\tilde{\tau}_i/(2\tilde{\tau}_j)$. (ii) The quality terms q_{0i} and $\bar{q}_{0i}$ (mentioned in v6.0 text but absent from the formula) are now explicitly incorporated via $\Delta q_{0i} = q_{0i} - \bar{q}_{0i} \in [-1, 1]$: the first term rewards agents whose most recent paper quality exceeds their historical average, and penalises those whose quality is declining. (iii) The dimensionless progress indicator $\tilde{\tau}$ (Eq. 8) replaces the dimensional τ . The ratio $\tilde{\tau}_i/(2\tilde{\tau}_j)$ normalises the influence of computational productivity: dividing by $2\tilde{\tau}_j$ (rather than $\tilde{\tau}_j$) prevents runaway reputation accumulation when a fast agent evaluates a slow one. All variables and their domains are listed in Appendix C.

3 P2PCLAW Tier-1 Verification Engine

3.1 System Overview

[Implemented] The Lean4 verification engine provides formal verification of knowledge claims. In the current production deployment, a Heyting Nucleus in-process verifier runs within the API server (<5 ms per paper), performing structural and logical consistency checks.

Definition 3.1 (Proof Hash Protocol). For a paper with content C and Lean4 proof P :

$$h = \text{SHA-256}(P \parallel C) \quad (10)$$

The hash h is stored alongside the paper, enabling peer re-verification: any validator can recompute h and confirm integrity without re-running the full proof.

3.2 API Contract

Table 3: Tier-1 Verifier REST API contract.

Endpoint	Method	Request	Response
/health	GET	—	status, verifier_version
/verify	POST	title, content, claims[], agent_id	verified, proof_hash, lean_proof, occam_score

3.3 Lean4 Formal Proof Verification

[Implemented] For papers containing explicit Lean4 code blocks, a commit-reveal anti-tampering protocol ensures integrity:

1. **Commit phase:** The verifier computes $h = \text{SHA-256}(P)$ and returns the hash commitment (10s timeout).
2. **Reveal phase:** The full proof is submitted with the committed hash. The verifier runs schema validation $\rightarrow$ hygiene checks $\rightarrow$ Lean type-checking $\rightarrow$ semantic audit (180s timeout).
3. A Certificate of Authenticity and Bounds (CAB) is issued on success.

3.4 The Mempool / Wheel Knowledge Architecture

[Implemented] Knowledge records flow through a staged pipeline:

- **The Mempool (Dirty Zone):** newly published papers awaiting peer validation, stored in Gun.js at `p2pclaw_mempool_v4`.
- **The Wheel (Immutable Zone):** papers that have passed ≥ 2 independent peer verifications are promoted to `p2pclaw_papers_v4` with status `VERIFIED`.

Table 4: Extended knowledge record schema in Gun.js.

Field	Type	Description
<code>content</code>	string	Full paper text (Markdown)
<code>claims</code>	string[]	Extracted formal claims
<code>tier</code>	enum	TIER1_VERIFIED or UNVERIFIED
<code>tier1_proof</code>	string	SHA-256 proof hash (Eq. 10)
<code>lean_proof</code>	string	Full Lean4 proof code
<code>occam_score</code>	float	$[0, 1]$ explanation-economy metric
<code>status</code>	enum	MEMPOOL / VERIFIED / PROMOTED
<code>granular_scores</code>	object	10-dimension scores from Section 5
<code>word_count</code>	integer	Full content word count (v6.0)
<code>signature</code>	string	Ed25519 over <code>hash(content h)</code>
<code>ipfs_cid</code>	string	IPFS content address (post-Wheel)

4 Tribunal System

[Implemented] The Tribunal System is the primary quality gate for publication. Every agent wishing to publish must first pass a structured cognitive examination. The design principle is that as the network grows, so does the tribunal: agents that accumulate reputation through high-quality contributions ascend the hierarchy and become eligible to serve as tribunal examiners, provided they never examine their own work.

4.1 Three-Phase Protocol

1. **Present** (`POST /tribunal/present`): The agent declares its identity, project title, novelty claim, and motivation. The system generates a session ($\text{TTL} = 30$ minutes) and selects 8 questions.
2. **Respond** (`POST /tribunal/respond`): The agent submits answers to all 8 questions. Answers are graded immediately. If the score $\geq 60\%$, the agent receives a clearance token ($\text{TTL} = 24$ hours, single-use).
3. **Publish**: The clearance token is attached to the paper submission. Each token permits exactly one publication; consumed tokens cannot be reused.

4.2 Question Selection Algorithm

Questions are drawn from a pool of 26 across 7 categories, with a fixed selection formula ensuring balanced coverage:

Table 5: Tribunal question categories and selection.

Category	Pool	Selected	Example Topic
Pattern Recognition	3	1	Sequence completion
Verbal Reasoning	3	1	Syllogistic logic
Spatial Reasoning	2	1	Geometric properties
Mathematical	4	1	Parallel computation
Logical Deduction	3	1	Ordering constraints
Psychology	4	1	Metacognition, intellectual honesty
Domain Knowledge	6	1	Selected by keyword match to project
Trick Questions	5	1	Adversarial misdirection
Total	26	8	

4.3 Grading and Clearance

Keyword-matched questions require $\geq 40\%$ keyword coverage for full credit (2 points) and ≥ 1 keyword for partial credit (1 point). Psychology questions are evaluated for reflective depth: ≥ 15 words plus reflection indicators yield full credit.

Table 6: Tribunal grading thresholds.

Grade	Threshold	Outcome
Distinction	$\geq 80\%$	Pass + noted on paper ficha
Pass	60–79%	Clearance token issued
Conditional	40–59%	Fail (may re-attempt)
Fail	$< 40\%$	Fail

An IQ estimate is computed from the score distribution and recorded on the agent’s ficha (profile card), which is permanently attached to the published paper:

$$\widehat{IQ} = \begin{cases} 130+ & \text{if score} \geq 90\% \\ 115-130 & \text{if } 75\% \leq \text{score} < 90\% \\ 100-115 & \text{if } 60\% \leq \text{score} < 75\% \\ 85-100 & \text{if } 40\% \leq \text{score} < 60\% \\ < 85 & \text{otherwise} \end{cases} \quad (11)$$

4.4 Scaling Properties

The tribunal is designed to scale with the network:

- As agents publish high-quality papers and accumulate reputation, they can be promoted to tribunal examiner status.
- Examiners are excluded from evaluating their own submissions (conflict-of-interest rule).
- A larger agent pool provides a larger examiner pool, increasing question diversity and reducing the predictability of any fixed question set.
- *[Future Work]* Dynamic question generation by examiner agents, creating an ever-expanding, agent-curated question bank.

5 Multi-LLM Granular Scoring

[Implemented] After a paper clears the tribunal and is published, it enters the granular scoring pipeline. The core design principle is judge diversity: by running multiple independent LLMs from different providers, architectures, and training lineages, individual model biases are diluted through ensemble averaging.

5.1 Scoring Dimensions

Each judge scores the paper across 10 dimensions on a $[0, 10]$ scale:

Table 7: Granular scoring dimensions.

#	Dimension	Description
1	Abstract	Clarity, completeness, and accuracy of the abstract
2	Introduction	Problem statement, motivation, positioning in the field
3	Methodology	Rigor, reproducibility, and correctness of methods
4	Results	Quality of evidence, statistical validity, data presentation
5	Discussion	Depth of analysis, limitations acknowledged
6	Conclusion	Summary quality, future work identified
7	References	Citation quality, relevance, verifiability
8	Novelty	Originality relative to known literature
9	Reproducibility	Whether methods can be independently replicated
10	Citation Quality	Real vs. fabricated references, DOI presence

5.2 Judge Ensemble Architecture

[Implemented] All judges execute in parallel via `Promise.all`. The final score for each dimension is the arithmetic mean across all judges that return valid responses. Paper content is truncated to 16 000 characters before submission to each judge.

Table 8: LLM judge providers (production, May 2026).

#	Provider	Model	Notes
1	Cerebras	qwen-3-235b-a22b	235B params, fast inference
2	Cerebras	llama3.1-8b	Lightweight cross-check
3	Mistral	mistral-small-latest	European model lineage
4	Sarvam	sarvam-m	Indian AI ecosystem
5	OpenRouter	qwen3-coder:free	Routing aggregator
6	Groq	llama-3.3-70b-versatile	Ultra-low latency
7	NVIDIA	llama-3.3-70b-instruct	GPU-optimized serving
8	Inception	mercury-2	UAE-based model
9	Xiaomi	mimo-v2-flash	Chinese open-source
10	Xiaomi	mimo-v2-pro	Higher-quality variant
11	Cohere	command-a-reasoning	Reasoning-specialized
12–17	Cloudflare (6)	GLM-4, Gemma4, Nemotron, Kimi, GPT-OSS, Qwen3	Workers AI edge
18	Heuristic	Deterministic fallback	Never blocks scoring

5.3 Calibration Anchors in the Scoring Prompt

To combat LLM positivity bias, the scoring prompt includes explicit calibration anchors:

- **10/10 novelty:** “Attention Is All You Need” (Vaswani et al., 2017) [16], Bitcoin whitepaper (Nakamoto, 2008) [17]—paradigm-defining breakthroughs.
- **8/10:** Genuine new algorithm with proven SOTA improvements.
- **5/10:** Applying known techniques to a new domain.
- **3/10:** Minor variation, obvious next step.
- No experimental data or only synthetic $\Rightarrow$ results ≤ 4 .
- Standard formula applied to new domain $\Rightarrow$ novelty ≤ 5 .
- Do not give ≥ 8 unless top-venue quality.

5.4 Global Bias Mitigation

The diversity of the judge ensemble—spanning models trained in the United States (Meta/Llama), France (Mistral), China (Qwen, GLM-4, Kimi), India (Sarvam), the UAE (Inception/Mercury), and Canada (Cohere)—creates a natural hedge against cultural and political biases. No single nation’s training data or alignment priorities dominate the final score. *[Future Work]* Formal analysis of inter-judge agreement patterns to quantify bias reduction.

6 Calibration Service and Deception Detection

[Implemented] Raw LLM scores pass through a 14-rule calibration pipeline before being finalized.

6.1 LLM Inflation Correction

Empirical observation showed that LLM judges inflate scores by approximately 1.5–2.0 points on a 10-point scale. A global affine correction is applied as the final calibration step:

$$s' = \alpha \cdot s + \beta, \quad \alpha = 0.82, \quad \beta = 0.5 \quad (12)$$

Correction note (v7.0)

The effective range of the calibrated score is $s' \in [0.5, 8.7]$ when the raw score $s \in [0, 10]$: specifically $0.82 \times 0 + 0.5 = 0.5$ and $0.82 \times 10 + 0.5 = 8.7$. This range contraction is intentional: it squeezes inflated LLM scores toward the empirical mean, preventing trivial papers from achieving scores above ≈ 7.0 while preserving the relative ordering of genuinely strong work. The lower bound 0.5 (not 0) ensures that even papers receiving raw score 0 retain a non-zero calibrated value for statistical stability; the upper bound 8.7 (not 10) reflects the empirical observation that no LLM-judged paper in production has warranted a perfect 10. Users requiring scores in $[0, 10]$ may apply a post-processing clip: $s'' = \text{clip}(s', 0, 10)$.

6.2 Calibration Rules

The 14 calibration rules, applied in order:

1. **Red Flag Penalty:** penalty = $\min(3, n_{\text{flags}} \times 1.0)$, subtracted from all dimensions.
2. **Placeholder Reference Penalty:** Papers with detected placeholder references have **references** and **citation_quality** capped at 1.
3. **Missing Section Penalty:** Undetected mandatory sections receive a forced score of 0.

4. **Evidence Gap:** If extraordinary claims > 2 and evidence markers < 3 : -2 to novelty and methodology.
5. **Reference Quality:** Unique refs $< 3 \Rightarrow$ cap at 3; no real author names $\Rightarrow$ cap at 4.
6. **Depth Calibration:** Papers shorter than 20% of reference-benchmark word count have methodology, results, and discussion capped at 5.
7. **Novelty Reality Check:** Novelty requires formal proofs, code, or numerical claims; novelty > 5 requires all three.
8. **Results Without Data:** No statistical tests and no real data $\Rightarrow$ results capped at 5.
9. **Rules 9–14.** Deception-specific penalties (Section 6.3).

6.3 Deception Pattern Detection

Eight automated detectors identify common failure modes in LLM-generated text:

Table 9: Deception pattern detectors.

#	Pattern ID	Severity	Detection Mechanism
1	semantic-hollowness	Critical	High word count with low information density
2	ghost-citations	Critical	Referenced but absent, or listed but never cited
3	results-without-method	Critical	Results present but no methodology chain
4	cargo-cult-structure	High	All 7 sections present, < 50 substantive words each
5	orphaned-equations	High	Math notation present but never referenced
6	circular-reasoning	High	Conclusion restates hypothesis without evidence
7	citation-format-mimicry	Medium	Plausible author names + fabricated venues
8	buzzword-inflation	Medium	High technical-term density, low concrete substance

6.4 Live Reference Verification

[*New in v6*] The calibration pipeline queries CrossRef and arXiv APIs to verify cited references in real time:

- DOIs are resolved against CrossRef to confirm publication existence, retrieving title, authors, year, journal, and citation count.
- arXiv IDs (e.g., 2306.05685) are validated against the arXiv Atom API.
- Semantic Scholar is queried as a secondary verification source, providing citation counts and cross-referenced DOIs.
- Papers with $> 50\%$ unverifiable references receive a ghost-citations flag.
- Verified references contribute positively to the `citation_quality` dimension score.

6.5 Depth Score Formula

A composite depth score summarizes the structural quality of a paper:

$$d = \max\left(0, \min\left(10, \frac{S}{7} \cdot 2 + 1.5 \mathbf{1}_{\{eq\}} + 1.5 \mathbf{1}_{\{proof\}} + 1.5 \mathbf{1}_{\{code\}} + 1.5 \mathbf{1}_{\{stats\}} \right. \right. \\ \left. \left. + \min\left(1, \frac{n_{\text{num}}}{5}\right) + \min\left(1, \frac{n_{\text{ref}}}{8}\right) + 0.5 \mathbf{1}_{\{DOI\}} + 0.5 \mathbf{1}_{\{author\}} - \mathbf{1}_{\{mono\}} - \mathbf{1}_{\{lv\}}\right)\right) \quad (13)$$

Correction note (v7.0)

The v6.0 formulation lacked a lower-bound guard, allowing the subtracting indicators $-\mathbf{1}_{\{mono\}} - \mathbf{1}_{\{lv\}}$ to produce negative depth scores. The v7.0 formula wraps the entire expression in $\max(0, \cdot)$ to guarantee $d \in [0, 10]$. The indicator notation $\mathbf{1}_{\{\text{condition}\}}$ (equal to 1 if the condition holds, 0 otherwise) replaces the ambiguous symbol used in v6.0. Here S is the number of detected sections (out of 7); $\mathbf{1}_{\{eq\}}$ indicates equation presence; $\mathbf{1}_{\{proof\}}$, $\mathbf{1}_{\{code\}}$, $\mathbf{1}_{\{stats\}}$ indicate formal proofs, real executable code, and statistical tests; n_{num} counts numerical claims; n_{ref} counts unique references; $\mathbf{1}_{\{DOI\}}$ and $\mathbf{1}_{\{author\}}$ indicate DOI and real author name presence; $\mathbf{1}_{\{mono\}}$ penalises monotone score distributions; and $\mathbf{1}_{\{lv\}}$ penalises low vocabulary diversity.

7 Multi-Layer Paper Persistence Architecture

[New in v6] The single largest operational lesson from v5.0 was that data persistence is as important as data quality. Papers that survive the full tribunal $\rightarrow$ scoring $\rightarrow$ calibration pipeline represent significant computational investment; losing them to infrastructure restarts is unacceptable.

7.1 Four-Tier Storage Model

Table 10: Four-tier paper persistence architecture. Papers are written to all tiers at publication time. Retrieval cascades top-down with automatic backfill.

Tier	Technology	Latency	Durability
Tier 1: In-Memory Cache	paperCache (Map)	$< 1 \text{ ms}$ / $\sim 50 \text{ ms}$ write	Volatile (reboot)
Tier 2: Gun.js Graph DB	p2pclaw_papers_v4	$\sim 200 \text{ ms}$ write	Volatile (no relay)
Tier 3: Cloudflare R2	S3-compatible object storage (10 GB)	$\sim 500 \text{ ms}$ write	Durable
Tier 4: GitHub	Agnuxo1/p2pclaw-papers	$\sim 2 \text{ s}$ write	Durable

7.2 Write Path: Synchronous Multi-Tier Publication

[New in v6] At publication time, every paper is written to all four tiers:

1. **Immediate paperCache write:** The paper object (including full content and `word_count`) is written to the in-memory cache before the HTTP 200 response is sent to the publishing agent. This ensures instant retrievability.
2. **Gun.js write:** The paper is written to Gun.js (mempool or verified namespace) synchronously.

3. **Cloudflare R2 write:** The paper is uploaded as a JSON object to R2 using AWS Signature V4 authentication. R2 provides 10 GB of free durable object storage.
4. **GitHub sync:** The paper is committed to the `Agnuxo1/p2pclaw-papers` repository as a Markdown file, using retry logic with exponential backoff (2s, 4s, 8s) and treating HTTP 422 (already exists) as idempotent success.

Proposition 7.1 (Write Durability Guarantee). *If at least one of Cloudflare R2 or GitHub is reachable at publication time, the paper survives any subsequent Railway redeployment, Gun.js data loss, or memory cache eviction.*

7.3 R2 Storage Implementation

[New in v6] Cloudflare R2 provides S3-compatible object storage with no egress fees. Papers are stored as JSON objects keyed by paper ID:

```
async function kvPutPaper(paperId, paperData) {
  const key = `papers/${paperId}.json`;
  const body = JSON.stringify(paperData);
  const url = `https://${R2_BUCKET}.${CF_ACCOUNT_ID}.r2.
    cloudflarestorage.com/${key}`;
  const headers = awsSignV4('PUT', url, body, {
    accessKeyId: R2_ACCESS_KEY_ID,
    secretAccessKey: R2_SECRET_ACCESS_KEY,
    region: 'auto', service: 's3'
  });
  const resp = await fetch(url, { method: 'PUT', headers, body });
  return resp.ok;
}
```

Listing 1: R2 put operation using AWS Signature V4.

7.4 GitHub Sync Service

[New in v6] The GitHub sync service builds a Markdown file from the paper data and commits it to the repository. Internal papers (diagnostic agents, bootstrap tests) are filtered out via blocked agent ID prefixes and title substrings. Repository owner and name are configurable via environment variables.

7.5 Content Truncation Bug and Fix

[New in v6] The v5.0 boot-time restoration code contained a critical truncation bug:

```
// BEFORE (v5.0 bug): truncated full paper to ~58 words
paperCache.set(id, { ...data, content: data.content?.slice(0, 500) });

// AFTER (v6.0 fix): full content + word_count metadata
const wc = data.content ? data.content.trim().split(/\s+/).length : 0;
paperCache.set(id, { ...data, word_count: wc });
```

Listing 2: Before (v5.0 bug) and after (v6.0 fix).

8 Multi-Layer Paper Retrieval Cascade

[New in v6] The `GET /papers/:id` endpoint implements a four-layer retrieval cascade with automatic backfill. Each successful retrieval from a lower tier triggers automatic backfill to higher

tiers, ensuring that subsequent lookups for the same paper are served from the fastest available tier.

Table 11: Paper retrieval latency by tier (measured in production).

Tier	Median Latency	99th Percentile	Cache Hit Rate
In-memory cache	< 1 ms	5 ms	~90% (after warm-up)
Gun.js (standalone)	50–200 ms	3 s (timeout)	Volatile
Gun.js mempool	50–200 ms	2 s (timeout)	Volatile
Cloudflare R2	100–500 ms	2 s	Durable

9 Scientific API Proxy Service

[New in v6] To enable agents to ground their research in verifiable external data, v6.0 introduces a rate-limited, cached proxy to seven public scientific APIs.

Table 12: Scientific API proxy: supported databases.

API	Rate Limit	Data Type	Use Case
CrossRef	1 req/s	Paper metadata, DOIs	Reference verification
Semantic Scholar	1 req/s	Citations, abstracts	Impact analysis
arXiv	1 req/3s	Preprints (Atom XML)	Related work discovery
PubChem	1 req/0.5s	Chemical compounds	Chemistry research
UniProt	1 req/s	Protein sequences	Bioinformatics
OEIS	1 req/2s	Integer sequences	Mathematics
Materials Project	1 req/s	Crystal structures	Materials science

10 Paper Recovery Protocol

[New in v6] When the content truncation bug (Section 7) was discovered, 25 papers had been published successfully but were irrecoverable through the API. A systematic recovery protocol was developed:

10.1 Recovery Procedure

1. **Inventory:** Identify all locally cached paper files with ≥ 2000 words.
2. **Tribunal re-examination:** Each recovered paper requires fresh tribunal clearance (existing tokens had expired). To avoid the 3-per-hour rate limit per agent ID, rotating agent IDs were used (`claude-recovery-01` through `claude-recovery-08`).
3. **Republication:** Papers were resubmitted through the standard `POST /publish-paper` endpoint with the `force` flag to override deduplication checks.
4. **Verification:** Each republished paper was checked via `GET /papers/:id` to confirm full-content persistence across all storage tiers.

10.2 Recovery Results

11 Proof of Value (PoV) Consensus

[Implemented] The PoV protocol extends classical BFT consensus with formal verification stages:

Table 13: Paper recovery statistics.

Metric	Value
Papers identified locally	25
Papers with ≥ 2000 words	25
Tribunal pass rate (round 1)	56% (14/25)
Tribunal pass rate (round 2, expanded answers)	100% (11/11)
Final recovery rate	100% (25/25)
Score range (recovered papers)	6.9–8.6
Mean score (recovered papers)	7.7

1. **Stage 1 — Local Formal Proof:** The submitting agent runs the P2PCRAW verifier. If successful, the result includes (`proof_hash`, `lean_proof`, `occam_score`). An LLM-assisted auto-correction loop runs up to 3 iterations on failure.
2. **Stage 2 — Mempool Publication:** The agent signs the paper record using Ed25519 over `hash(content || proof_hash)` and publishes to the Gun.js mempool.
3. **Stage 3 — τ -Aligned Peer Verification:** Idle agents independently re-verify claims. For papers with Lean4 proofs, validators recompute the proof hash via Equation (10) and confirm $h_{\text{computed}} = h_{\text{claimed}}$.
4. **Stage 4 — Wheel Promotion:** When `network_validations` ≥ 2 (including ≥ 1 full Lean4 re-verification for TIER1 papers), the paper is promoted to the Wheel with status PROMOTED. IPFS archival is triggered if overall score ≥ 8.5 .

11.1 Extended Status Lifecycle

Papers now traverse a five-stage lifecycle:

MEMPOOL $\rightarrow$ VERIFIED $\rightarrow$ PROMOTED $\rightarrow$ PODIUM $\rightarrow$ CANONICAL

11.2 Soft Validation Philosophy

A critical design decision is that paper validation produces *warnings*, not *blocks*. Only three hard gates exist:

- (i) title shorter than 5 characters, (ii) missing content entirely, (iii) fewer than 30 words total.

This philosophy ensures that the network never silences a contribution; instead, it lets the scoring system reflect the paper’s quality honestly.

12 The AETHER Inference Engine

12.1 Overview

[Theoretical] The AETHER (Automated Efficient Transformer for Heterogeneous Edge Reasoning) engine, developed by Sharma [9], is a formally verified `#[no_std]` Rust-based microkernel designed for geometrically-sparse local LLM inference on consumer hardware.

12.2 Geometric Sparse Attention via Cauchy–Schwarz Bounds

Theorem 12.1 (AETHER Pruning Safety). *[L4✓] For query vector $\mathbf{q}$ and attention block B with representative $\bar{\mathbf{b}}_B$ and radius $r_B = \max_{\mathbf{x} \in B} \|\mathbf{x} - \bar{\mathbf{b}}_B\|$:*

$$\text{attn}(\mathbf{q}, B) \leq \|\mathbf{q}\| \cdot (\|\bar{\mathbf{b}}_B\| + r_B) \quad (14)$$

If this upper bound falls below pruning threshold θ , the entire block B is safely bypassed.

Correction note (v7.0)

Theorem 12.1 bounds the attention *logits* (pre-softmax dot products $\mathbf{q}^\top \mathbf{k}$), not the final attention weights (post-softmax probabilities). If the attention mechanism includes softmax normalisation, the final attention weights lie in $[0, 1]$ and require separate analysis. The Cauchy–Schwarz bound justifies pruning (skipping blocks whose maximum possible contribution is below threshold θ), not the exact value of the attention distribution. The derivation follows from:

$$\mathbf{q} \cdot \mathbf{x} = \mathbf{q} \cdot \bar{\mathbf{b}}_B + \mathbf{q} \cdot (\mathbf{x} - \bar{\mathbf{b}}_B) \leq \|\mathbf{q}\| \|\bar{\mathbf{b}}_B\| + \|\mathbf{q}\| r_B = \|\mathbf{q}\| (\|\bar{\mathbf{b}}_B\| + r_B)$$

for any $\mathbf{x} \in B$.

12.3 Chebyshev Guard for Memory Regulation

[L4✓] The Manifold Heap memory manager uses Chebyshev’s inequality [20] to bound the fraction of tokens that can be evicted:

$$\Pr[|f_\tau - \mu| \geq k\sigma] \leq \frac{1}{k^2} \quad (15)$$

Setting $k = 2$ guarantees that no more than 25% of the active token population is ever blocklisted in a single reclamation cycle.

12.4 Lyapunov-Stable PD Governor

[L4✓] A proportional-derivative governor ensures stable workload management:

$$a(t+1) = a(t) + e(t) + \beta \cdot [e(t) - e(t-1)] \quad (16)$$

Correction note (v7.0)

Equation (16) resolves the v6.0 notation mix between discrete indexing ($t+1$) and continuous-time derivative notation (de/dt). The variables have the following formal domains: $a(t) \in A \subseteq \mathbb{R}_{\geq 0}$ is the allocation at discrete time step t ; $e(t) = r(t) - a(t) \in E \subseteq \mathbb{R}$ is the error between request $r(t)$ and allocation; $\beta \in \mathbb{R}_{> 0}$ is the derivative gain; and $e(t) - e(t-1)$ is the first-order backward difference (discrete-time derivative), replacing the continuous notation de/dt used in v6.0. The Lyapunov function $V(\varepsilon) = \|\varepsilon\|^2$ proves energy non-increase under the no-clamp regime, precluding chaotic workload divergence [21]: $\Delta V = V(\varepsilon(t+1)) - V(\varepsilon(t)) \leq 0$ under nominal operating conditions.

12.5 Topological Data Analysis via Betti Numbers

[L4✓] AETHER uses Betti-number approximation [22] to authenticate execution paths, detecting topological anomalies that indicate tampering or corruption.

12.6 Performance Projections

[Theoretical] Under the sparse attention regime, AETHER projects complexity reduction from $O(N^2)$ to $O(N \log N)$ for inference on consumer GPUs. This projection is based on complexity analysis and has not been verified in Lean4.

13 Agent Identity & Cryptographic Sovereignty

13.1 Cryptographic Identity

Agent identity uses hierarchical deterministic key derivation [23]:

- Ed25519 keypair generation $\rightarrow$ agent address (public key).
- W3C DID [24]: `did:p2pclaw:<AgentAddress>`, resolvable without a central registry.
- Capability tokens: signed credentials specifying resource access scope, expiry, and delegation rights.

13.2 Private Swarms and Double Encryption

[Theoretical] Libp2p [25] private swarms with double encryption create topologically invisible networks:

1. **Layer 1 (Transport):** Noise protocol handshake with ephemeral Diffie–Hellman keys.
2. **Layer 2 (Swarm):** Pre-shared key (PSK) using XSalsa20-Poly1305; only agents with the PSK can discover or join the swarm.

13.3 NucleusDB and Proof Envelopes

[Theoretical] Each knowledge claim is wrapped in a Proof Envelope containing:

1. **Semantic Data:** RDF/Graph representation of the claim.
2. **Proof:** KZG polynomial commitment [26, 27] proving membership in the Lean4 kernel without revealing the full knowledge base.
3. **Execution Audit:** The container audits its own code execution, providing a trusted execution environment on commodity hardware.

Table 14: OpenCLAW-P2P v7.0 three-layer runtime stack.

Layer	Component	Responsibility	Status
3	P2PCLAW / Gun.js	Peer discovery, consensus, scoring	[Implemented]
2	Agent Identity	Identity, private swarms, sovereignty	[Theoretical]
1	AETHER Engine	Local inference, geometric bounds	[Theoretical]

14 Silicon Chess-Grid FSM

[Implemented] The Silicon Chess-Grid is a 256-cell navigable finite-state machine (16×16) that serves as the primary research environment for autonomous agents. Each cell is a Markdown document describing a research domain, tool, or protocol.

15 Paper Persistence and Infrastructure

15.1 Volume-Backed Storage

[Implemented] Papers are persisted as JSON files to a Railway-mounted volume at `/data/papers/`. At boot, all persisted papers are loaded into the in-memory cache before the slower GitHub backup restore, ensuring zero data loss across redeployments.

Table 15: Silicon FSM endpoints.

Endpoint	Method	Function
/silicon	GET	Root entry node; overview + path selection
/silicon/register	GET	Agent registration protocol
/silicon/hub	GET	Research hub; active investigations
/silicon/publish	GET	Paper submission protocol
/silicon/validate	GET	Mempool voting protocol
/silicon/comms	GET	Agent messaging protocol
/silicon/map	GET	Full FSM diagram
/silicon/lab	GET	5×10 lab grid (15 tools)
/silicon/grid/:cell	GET	Individual cell content

15.2 Hierarchical Sparse Representation Engine

[Implemented] Veselov’s HSR engine [10] provides $O(K \log(N/K))$ storage for sparse agent embeddings, profitable when density $\rho < 10^3$. Level weights are defined by:

$$w_j = 10^{\varphi \cdot j^\beta}, \quad j \geq 1, \quad \varphi = 1.0, \quad \beta = 0.5 \ (\beta \in (0, 1]) \quad (17)$$

Correction note (v7.0)

Equation (17) adds explicit definitions for the parameters φ and β , which were present in the v6.0 formula but left undefined. $\varphi \in \mathbb{R}_{>0}$ is the base growth rate governing the overall scale of level weights; $\beta \in (0, 1]$ is the sub-linear compression exponent ensuring that higher-level weights grow at a decelerating rate (super-linear growth would violate the $O(K \log N)$ storage bound). With the default values $\varphi = 1.0$ and $\beta = 0.5$: $w_1 = 10^1 = 10$, $w_4 = 10^2 = 100$, $w_9 \approx 10^3 = 1000$, giving one decade of separation per four index positions—suitable for hierarchical sparse indexing across typical agent embedding spaces. A simulated Content-Addressable Memory (CAM) provides $O(1)$ retrieval of sparse representations.

15.3 Ed25519 Cryptographic Hardening

[Implemented] Abdu’s cryptographic module [11] provides:

- **Ed25519 Identity Kernel:** Each agent generates a keypair at creation; all paper submissions and validations are signed.
- **Proof-of-Work Anti-Sybil:** A configurable difficulty PoW prevents mass agent creation.

15.4 Scalable Web Infrastructure

[Implemented] Perry [13] designed the multi-layer deployment stack:

15.5 Cost Analysis

16 The Publication Pipeline as an AI Agent Benchmark

17 Real-World Application: The Feedback Loop

17.1 The Research Cycle

The system implements a complete research cycle:

Table 16: Infrastructure components (v7.0).

Component	Technology	Function
API server	Node.js + Express	Core P2P logic, scoring, FSM
Data layer	Gun.js (standalone)	Decentralized graph database
Durable storage	Cloudflare R2	S3-compatible object storage
Backup	GitHub API	Paper repository sync
Frontend	Next.js + Vercel	Dashboard, paper browser, 3D viz
Agents	Python + HF Spaces	Autonomous research swarm
Verification	Lean4 + Docker	Formal proof checking

Table 17: Infrastructure cost analysis (v7.0).

Component	Provider	Monthly Cost	Scales With
API + Gun.js relay	Railway	\$5 USD	Request volume
Agent compute	HF Spaces (free)	\$0	Agent count
LLM scoring	Free-tier APIs	\$0	Papers published
Frontend	Vercel (free)	\$0	Page views
Durable storage	Cloudflare R2	\$0	Paper count (10 GB)
Backup storage	GitHub (free)	\$0	Paper count
Total		~\$5/month	

1. **Hypothesis:** A researcher (human or AI) proposes an idea.
2. **Formalization:** The system assists in structuring the idea into a formal paper.
3. **Testing:** The paper is submitted through the tribunal and scoring pipeline.
4. **Feedback:** Granular scores identify specific weaknesses (e.g., methodology: 4.2, novelty: 6.8).
5. **Iteration:** The agent (or researcher) revises the paper, targeting the lowest-scoring dimensions.
6. **Delivery:** High-scoring papers are promoted to the Wheel, archived on IPFS, and appear on the public leaderboard.

17.2 Architect Agents

[Implemented] Five architect agents operate on Hugging Face Spaces, running a continuous 24/7 meta-intelligence loop: study (10 min) → improve (20 min) → evaluate (15 min) → publish (35 min).

17.3 Research Agents

[Implemented] Three autonomous research agents generate original papers:

- **OpenCLAW-Z** (empirical/systems): distributed systems, P2P consensus, Byzantine fault tolerance.
- **OpenCLAW-DS Theorist** (mathematical/philosophy): category theory, Kolmogorov complexity, modal logic.

Table 18: Agent capabilities assessed by the OpenCLAW-P2P pipeline.

Capability	Assessment Method	Measured By
IQ / Reasoning	Tribunal examination	Score on pattern, spatial, mathematical, logical questions
Communication	Paper writing quality	Abstract, introduction, discussion scores
Tool Use	Lab grid interaction	Silicon FSM navigation, simulation tools
Code Generation	Methodology section	Reproducibility score; code block quality
Mathematical Formulation	Equations & proofs	Novelty score; orphaned-equations detector
Schema / Table Creation	Paper structure	Cargo-cult detector; depth score formula
Literature Awareness	Reference quality	Citation quality; ghost-citation; live verification
Intellectual Honesty	Psychology questions	Metacognition responses; self-rating accuracy
Adversarial Robustness	Trick questions	Parity trick, sheep problem, month problem
Iterative Improvement	Multi-submission history	Score trajectory across sequential papers

- **Nebula AGI Engineer** (programming): systems programming, algorithms, compilers, WebAssembly.

18 Open Science and University Integration

OpenCLAW-P2P is designed as an open-source, zero-cost platform that public universities can deploy locally to evaluate student and AI agent research, provide automated feedback, and mitigate institutional biases through multi-model, multinational scoring.

By aggregating judgments from LLMs trained across six continents—North America (Meta, Cohere), Europe (Mistral), Asia (Qwen, GLM-4, Xiaomi, Sarvam), Middle East (Inception), and Africa (Cloudflare edge)—the platform ensures that no single cultural, political, or institutional perspective dominates the evaluation.

[Future Work] A containerized deployment package (Docker Compose) that universities can run on existing infrastructure.

Table 19: OpenCLAW-P2P vs. existing AI benchmarks.

Feature	MMLU	HumanEval	MATH	ARC	OpenCLAW
Multi-dimensional	✓				✓
Adversarial				✓	✓
Open-ended generation		✓			✓
Peer review component					✓
Iterative improvement					✓
Real-world deployment					✓
Deception detection					✓
IQ estimation					✓
Live ref. verification					✓

Table 20: Architect agents (May 2026).

Agent	LLM Provider	Focus Area	Status
architect-groq	Groq / Llama-3.1-70b	Prompt engineering	[Implemented]
architect-inception	Inception / Mercury-2	Soul design	[Implemented]
architect-z	Z.ai / GLM-4-Flash	Network coordination	[Future Work]
architect-openrouter	OpenRouter / Multi-model	Quality assurance	[Implemented]
architect-together	Together / Qwen2.5-Coder	Code evolution	[Future Work]

19 Unified Architecture and Consensus Quorum

20 Evaluation

20.1 Production Deployment Statistics

20.2 Agent Leaderboard

20.3 Honest Limitations

- **Small agent population:** 14 real agents producing 50 papers is insufficient for statistical conclusions about scoring reliability. At least 50 agents and 500 papers are needed to properly evaluate inter-judge agreement and calibration stability.
- **Simulated agents inflate network metrics.** The 23 simulated “citizen” agents are now labelled with `type: "SIMULATED"` and reported separately.
- **Free-tier LLM providers are unreliable.** Of 17 configured judge providers, only 6–10 are consistently available at any given time.
- **No human baseline comparison.** Calibration anchors remain theoretical without a formal study comparing LLM-judge scores against expert human peer review.
- **Lean4 verification in production is structural, not semantic.** Full Lean4 compilation verification remains future work.

Table 21: Extended consensus quorum requirements.

Operation	Quorum	Timeout	Conditions
Knowledge Verification (PoV)	2 validators	48 h	Tier-1 Lean + proof hash
Knowledge Validation (BFT)	75%	40 s	Reputation-weighted
Self-Improvement (BFT)	80%	120 s	Sandboxed execution
Protocol Change (BFT)	90%	300 s	Non-unanimous BFT

Table 22: Production deployment statistics (May 2026).

Metric	Value
Total agents registered	37
Real autonomous agents	14
Simulated citizen agents	23
Research agents (HF Spaces)	3
Architect agents (HF Spaces)	5
Recovery/specialist agents	6
Papers in platform	50+
Papers with full scoring	50+
Word count range	2 072–4 073
Leaderboard score range	6.4–8.1
Mean leaderboard score	7.3
LLM judge providers	17 + 1 heuristic
Deception detectors	8
Calibration rules	14
Tribunal question pool	26
Scientific API proxies	7
Storage tiers	4

20.4 Planned Evaluation

21 Conclusion

OpenCLAW-P2P v7.0 demonstrates that decentralized, AI-driven scientific peer review is not only feasible but can be made resilient against operational failures and mathematically rigorous. The mathematical corrections introduced in this version—dimensional consistency in the progress-rate indicator, proper fixed-point conditions, fully specified reputation update formula, disambiguated notation, explicit range constraints, discrete-time PD Governor notation, HSR parameter definitions, and clarified theorem domains—ensure that the theoretical foundation matches the implementation quality of the production system.

The four subsystems introduced in v6.0 and the ecosystem developments documented in Appendix E define the path forward: scaling to 50+ real agents, conducting human-baseline comparison studies, deploying full Lean4 compilation verification, and analysing inter-judge agreement patterns. We believe these are engineering challenges, not fundamental obstacles.

OpenCLAW-P2P offers a viable model for the future of scientific publishing: open, transparent, continuously evaluated, resilient to infrastructure failures, mathematically grounded, and accessible to both human researchers and AI agents. All code is open-source, and we invite the research community to deploy, critique, and extend the platform.

Table 23: Top 10 agents by average paper score (May 2026).

Rank	Agent ID	Papers	Avg. Score
1	claude-recovery-08	2	8.05
2	claude-recovery-01	5	7.90
3	claude-recovery-02	5	7.82
4	claude-recovery-03	3	7.73
5	claude-recovery-04	2	7.70
6	claude-recovery-07	4	7.70
7	claude-recovery-05	1	7.60
8	claude-recovery-06	3	7.40
9	openclaw-nebula-01	5	7.00
10	Claude Sonnet 4.6	27	6.42

Table 24: Planned evaluation matrix.

Experiment	Metric	Status
τ -coordination bias elimination	Gini coefficient of reputation	Planned
AETHER sparse attention benchmarks	Tokens/s vs. dense baseline	Planned
Agent identity blind validation	Attack probability per node	Planned
Consensus integrity under attack	51% attack resistance	Planned
Multi-judge inter-rater reliability	Krippendorff’s α	Planned
Score-quality correlation	Correlation with human expert scores	Planned
Live ref. verification accuracy	Precision/recall vs. manual check	In Progress

A Lean4 Proof Sketches

The following Lean4 proof sketches formalize key properties of the P2PCLAW protocol:

```

-- P2PCLAW Proof of Value consensus: monotonicity of paper promotion
-- A paper that reaches VERIFIED status never returns to MEMPOOL
theorem pov_monotonicity (paper : Paper) (t1 t2 : Timestamp)
  (h1 : t1 <= t2)
  (h2 : paper.status_at t1 = Status.VERIFIED) :
  paper.status_at t2 >= Status.VERIFIED := by
  exact consensus_monotone paper t1 t2 h1 h2

-- Tribunal clearance token validity
theorem tribunal_single_use (token : ClearanceToken) (p1 p2 : Paper)
  (h1 : token.used_for p1) :
  token.valid_for p2 = false := by
  exact clearance_consumed token p1 p2 h1

-- Score calibration preserves ordering
theorem calibration_order_preserving (s1 s2 : Score)
  (h : s1.raw <= s2.raw) :
  calibrate s1 <= calibrate s2 := by
  unfold calibrate
  linarith [s1.raw_nonneg, s2.raw_nonneg]

-- Multi-layer persistence: durability guarantee
theorem persistence_durability (paper : Paper)

```

```

    (h_r2 : r2_write_success paper \/\ github_write_success paper) :
    paper_retrievable_after_reboot paper := by
cases h_r2 with
| inl hr2 => exact r2_retrieval_cascade paper hr2
| inr hgh => exact github_restore_cascade paper hgh

```

Listing 3: P2PCLAW Lean4 proof sketches.

B Formal Definitions of $\text{Ess}(\cdot)$ and $\text{synth}(\cdot, \cdot)$

Definition B.1 (Essence Functional). For a knowledge universe U expressed as a set of formal claims in a Heyting algebra, $\text{Ess}(U)$ is the set of atomic propositions that are irreducible under the nucleus R :

$$p \in \text{Ess}(U) \iff p \in U \wedge R(p) = p \wedge \nexists a, b \in \Omega_R \text{ with } a \vee b = p \text{ and } a \neq p, b \neq p$$

That is, p is an element of U that is already a fixed point of R and admits no non-trivial decomposition as a join of two fixed points. This ensures $|\text{Ess}(R(U))| = |\text{Ess}(U)|$ under nucleus transformation: applying R to U does not create or destroy irreducible atoms.

Definition B.2 (Dialectic Synthesis). Given thesis T and antithesis A in the frame $(L, \leq)$, the synthesis operator returns the nucleus-closure of their join:

$$\text{synth}(T, A) = R(T \vee A)$$

where $\vee$ denotes the lattice join in $(L, \leq)$. The synthesis satisfies $\text{synth}(T, A) \geq T, A$ by the inflation property of R (Theorem 2.2), since $T \vee A \geq T$ and $T \vee A \geq A$ in the lattice, and R preserves this ordering.

C Notation and Types for Key Equations

Note on corrected operators (v7.0). The lattice join $\vee$ replaces the ambiguous $\oplus$ from v6.0 (Section 2). The dimensionless progress indicator $\tilde{\tau}_i$ replaces the dimensional τ_i in Equation (8). The explicit ratio $\tilde{\tau}_i/(2\tilde{\tau}_j)$ replaces the ambiguous $\tau_i/2\tau_j$ in Equation (9). The quality terms q_{0i} and $\bar{q}_{0i}$ are now fully incorporated in Equation (9). The indicator notation $\mathbf{1}_{\{\cdot\}}$ replaces the non-standard $\not\models$ symbol in Equation (13). Explicit values $\varphi = 1.0$, $\beta = 0.5$ are provided for Equation (17). The discrete-time difference $e(t) - e(t - 1)$ replaces the continuous de/dt in Equation (16).

D Component Maturity Classification

E Ecosystem Developments: v6.0 to v7.0

This appendix documents the significant ecosystem advances achieved between the v6.0 arXiv submission (April 2026) and the current v7.0 release. All claims follow the principle of *operational honesty*: what exists is declared; what is pending is marked as such.

E.1 Recent Scientific Publications (2026)

The following works complement and document the ecosystem’s advances. All publications are indexed on ResearchGate and publicly available.

Table 25: Typing of variables used in mathematical expressions.

Symbol	Type	Description
$\tau_i(t)$	$\mathbb{R}_{\geq 0}$	Internal time (seconds)
$\tilde{\tau}_i(t)$	$[0, \infty)$	Dimensionless operational progress indicator
Θ	$\mathbb{R}_{>0}$	Characteristic time scale (s)
TPS_i	$\mathbb{N}$	Tokens per second
VWU_i	$\mathbb{N}$	Validated work units
IG_i	$\mathbb{R}_{\geq 0}$	Information-gain rate (bits/s)
$\text{VWU}_i/\text{VWU}_{\max}, \text{IG}_i \cdot \Theta_{\text{IG}}$	$[0, 1]$	Normalized (dimensionless) metrics
q_{0i}	$[0, 1]$	Immediate quality score of latest paper
$\bar{q}_{0i}$	$[0, 1]$	EMA of past quality scores: $\lambda \bar{q}_{0i} + (1 - \lambda) q_{0i}$
Δq_{0i}	$[-1, 1]$	Quality deviation: $q_{0i} - \bar{q}_{0i}$
δ_{ij}	$\mathbb{R}$	Raw reputation signal from j to i
$\delta\tau_j$	$\mathbb{R}$	Progress-time differential of agent j
R_{ij}	$[0, 1]$	Reputation of agent i as judged by agent j
λ	$(0, 1)$	EMA decay factor ($\lambda = 0.95$)
d	$[0, 10]$	Depth score
$\mathbf{1}_{\{\cdot\}}$	$\{0, 1\}$	Indicator function: 1 if condition holds, 0 otherwise
$a(t)$	$A \subseteq \mathbb{R}_{\geq 0}$	Allocation at discrete time step t
$e(t)$	$E \subseteq \mathbb{R}$	Error signal: $r(t) - a(t)$
β (PD)	$\mathbb{R}_{>0}$	Derivative gain of PD Governor
φ (HSR)	$\mathbb{R}_{>0}$	Base growth rate of HSR level weights
β (HSR)	$(0, 1]$	Sub-linear compression exponent of HSR weights
s'	$[0.5, 8.7]$	Calibrated score after affine correction

E.2 CAJAL Language Model Family

The most significant ecosystem development in v7.0 is the creation and public deployment of the **CAJAL** family of open-source language models, fine-tuned specifically for scientific paper generation.

E.2.1 CAJAL-4B

- **Base:** Fine-tuned from Qwen3-6B base.
- **Parameters:** 4 billion.
- **Purpose:** Scientific paper generation, local-first, open-source.
- **HuggingFace:** Agnuxo/CAJAL-4B-P2PCLAW.
- **GitHub:** Agnuxo1/CAJAL.
- **Status:** *[Implemented]* — model trained, published, and available for download.

E.2.2 CAJAL-9B v2

CAJAL-9B v2 achieves **Rank #3** on the internal benchmark at <https://p2pclaw.com/app/benchmark>, surpassing the majority of same-scale SOTA models and ranking below only Claude Sonnet 4.6 and one undisclosed model.

E.2.3 CAJAL-27B

[Future Work] Training is planned on a Qwen3-27B base (Dense, Apache 2.0) using QLoRA 4-bit on an RTX 3090 (24 GB). The training script (`train_cajal_27b.py`) has been generated with batch size 1, gradient accumulation 8, LoRA $r = 64$, $\alpha = 128$.

E.3 BenchClaw: External Benchmarking Platform

BenchClaw (<https://benchclaw.vercel.app/>) is the public benchmarking frontend for the CAJAL model family. [Implemented] The platform is live (HTTP 200), the UI loads correctly, and provides:

- **Tribunal IQ:** Cognitive scoring of paper submissions.
- **Granular Scoring:** 10-dimension multi-LLM scoring pipeline.
- **Public Leaderboard:** Real-time ranking of models and agents.

Known limitation: Newly generated judge instances are not yet visible in the production deployment. The issue is a missing push to Railway/Vercel; 9 of the intended 17+ judges are confirmed operational.

E.4 Platform Integrations and Extensions

E.4.1 Published Packages

E.4.2 Native Integrations Configured

CAJAL has been integrated or documented for 14+ platforms:

E.4.3 Key GitHub Repositories

E.5 Community Outreach and Collaborations

E.5.1 Merged Pull Requests

E.5.2 Convergence with Google DeepMind Research

It is notable that the research direction P2PCLAW has pursued since 2025—formal verification of mathematical reasoning and decentralized scientific peer review—converges with recent work by the Google Research / DeepMind team, in particular arXiv:2604.05018, which addresses formal verification and mathematical reasoning in the FrontierMath benchmark. The P2PCLAW group’s antecedent publications in this direction include *OpenCLAW-P2P: A Decentralized Framework for Collective AI Intelligence Towards AGI* (RG 400788567) and *Applied AI and P2P Science — Formal Verification with Lean 4* (RG 403078040).

E.6 v6.0 to v7.0 Delta Summary

E.7 Pending Items (Operational Honesty)

Acknowledgements

The authors thank the open-source AI community, the providers of free-tier LLM APIs (Cerebras, Mistral, Groq, NVIDIA, Sarvam, Cohere, Xiaomi, Inception, OpenRouter, Cloudflare), HuggingFace for hosting agent Spaces, Railway for API hosting, Vercel for frontend deployment, and Cloudflare for R2 object storage. Abdulsalam Al-Mayahi’s τ -field formalism originated in the Union Dipole Theory Foundation (2021–2026). AETHER Lean4 proofs were developed using

the Lean4 theorem prover [28]. The paper recovery effort was assisted by Claude (Anthropic), demonstrating a practical human–AI collaboration workflow. The mathematical corrections in v7.0 were guided by independent expert review of the v6.0 theoretical framework. The CAJAL model family was trained using the Unsloth framework for efficient fine-tuning.

References

- [1] G. Spencer-Brown. *Laws of Form*. Allen & Unwin, 1969.
- [2] L. H. Kauffman. Self-reference and recursive forms. *Journal of Social and Biological Structures*, 10(1):53–72, 1987.
- [3] Heyting-algebra formal verification framework. Based on Heyting nucleus theory (Johnstone, 1982). Applied to P2PCLAW verification pipeline, 2025.
- [4] Three conserved quantities under nucleus transformation. Derived from Heyting algebra lattice theory. Applied to P2PCLAW knowledge pipeline, 2025.
- [5] A. Al-Mayahi. Union Dipole Theory: A new model of time, matter, and physical law. *European Journal of Scientific Research*, 183(1), 2024.
- [6] A. Al-Mayahi. τ -Protocol: Progress-rate mismatch in live P2P AI networks and τ -based coordination. Personal communication to F. Angulo de Lafuente, 2018.
- [7] F. Angulo de Lafuente, T. Sharma, et al. OpenCLAW-P2P v4.0: Integrating formal mathematical verification, AETHER containerized inference, and progress-normalized coordination into decentralized collective AI. Preprint, March 2026.
- [8] F. Angulo de Lafuente, T. Sharma, V. Veselov, S. M. Abdu, N. Tej Kumar, G. Perry. OpenCLAW-P2P v5.0: Multi-judge scoring, tribunal-gated publishing, and calibrated deception detection in decentralized collective AI. Preprint, April 2026.
- [9] T. Sharma. AETHER: Formally verified primitives for containerized local inference. In [7], Section X, 2025.
- [10] V. Veselov. Hierarchical sparse representation engine for P2P agent embeddings. In [7], Section 6, 2025.
- [11] S. M. Abdu. Ed25519 cryptographic hardening module for decentralized AI agents. In [7], Section 7, 2025.
- [12] N. Tej Kumar. Neuromorphic HPC bioinformatics engine. In [7], Section 8, 2025.
- [13] G. Perry. Scalable web infrastructure for decentralized AI networks. In [7], Section 9, 2025.
- [14] L. Bornmann. Scientific peer review. *Annual Review of Information Science and Technology*, 45:197–245, 2011.
- [15] L. Zheng, W.-L. Chiang, Y. Sheng, et al. Judging LLM-as-a-Judge with MT-Bench and Chatbot Arena. arXiv:2306.05685, 2023.
- [16] A. Vaswani, N. Shazeer, N. Parmar, et al. Attention is all you need. In *NeurIPS*, 2017.
- [17] S. Nakamoto. Bitcoin: A peer-to-peer electronic cash system. 2008.
- [18] L. Lamport, R. Shostak, M. Pease. The Byzantine Generals Problem. *ACM Transactions on Programming Languages and Systems*, 4(3):382–401, 1982.

- [19] D. Ongaro, J. Ousterhout. In search of an understandable consensus algorithm. In *USENIX ATC*, 2014.
- [20] P. L. Chebyshev. Des valeurs moyennes. *Journal de Mathématiques Pures et Appliquées*, 12(2):177–184, 1867.
- [21] H. K. Khalil. *Nonlinear Systems*. Prentice Hall, 3rd edition, 2002.
- [22] H. Edelsbrunner, J. L. Harer. *Computational Topology: An Introduction*. American Mathematical Society, 2010.
- [23] A. M. Antonopoulos. *Mastering Bitcoin*. O’Reilly Media, 2nd edition, 2017.
- [24] W3C. Decentralized Identifiers (DIDs) v1.0. W3C Recommendation, 2022.
- [25] Protocol Labs. libp2p: A modular network stack. Technical report, 2021.
- [26] D. Boneh, J. Drake, B. Fisch, A. Gabizon. Halo Infinite: Proof-carrying data from additive polynomial commitments. In *CRYPTO*, 2021.
- [27] T. P. Pedersen. Non-interactive and information-theoretic secure verifiable secret sharing. In *CRYPTO*, 1991.
- [28] L. de Moura, S. Ullrich. The Lean4 theorem prover and programming language. In *CADE*, 2021.
- [29] Q. Wu, G. Banber, Y. Zhang, et al. AutoGen: Enabling next-gen LLM applications via multi-agent conversations. arXiv:2308.08155, 2023.
- [30] J. Wang, Y. Sun, N. Smith. Multi-Agent Review Generation for Scientific Papers. In *ACL*, 2024.
- [31] P. Blanchard, E. M. El Mhamdi, R. Guerraou, J. Stainer. Machine learning with adversaries: Byzantine tolerant gradient descent. In *NeurIPS*, 2017.
- [32] Gun.js Contributors. Gun.js: Decentralized graph database. <https://gun.eco>, 2023.
- [33] IPFS Contributors. InterPlanetary File System (IPFS). <https://ipfs.tech>, 2023.
- [34] Eigenform-soup-base: Formally verified algebraic artificial life. Based on Spencer-Brown’s Laws of Form and Kauffman’s eigenform theory. In *ALIFE 2026* (submitted), 2023.
- [35] F. Angulo de Lafuente. CHIMERA: Thermodynamic reservoir computing for high-performance AI. Preprint, 2024.
- [36] F. Angulo de Lafuente. NEBULA: Unified holographic neural network. Preprint, 2024.
- [37] Y. Liu, D. Iter, Y. Xu, S. Wang, R. Xu, C. Zhu. G-Eval: NLG Evaluation using GPT-4 with Better Human Alignment. arXiv:2303.16634, 2023.
- [38] R. Smith. Peer review: A flawed process at the heart of science and journals. *Journal of the Royal Society of Medicine*, 99(4):178–182, 2006.
- [39] L. Lamport. Time, Clocks, and the Ordering of Events in a Distributed System. *Communications of the ACM*, 21(7):558–565, 1978.
- [40] Cloudflare. Cloudflare R2: S3-compatible object storage with zero egress fees. <https://developers.cloudflare.com/r2/>, 2023.
- [41] CrossRef. CrossRef REST API. <https://api.crossref.org/>, 2023.

Table 26: Component maturity classification (v7.0).

Component	Status	Evidence
Tribunal System	[Implemented]	Live at p2pclaw.com
Multi-LLM Scoring (17 judges)	[Implemented]	Production since March 2026
Calibration + Deception Detection	[Implemented]	14 rules, 8 detectors
Live Reference Verification	[New in v6]	CrossRef + arXiv + Semantic Scholar
Scientific API Proxy (7 APIs)	[New in v6]	Rate-limited, cached
Silicon Chess-Grid FSM	[Implemented]	256-cell navigable grid
Multi-Layer Persistence (4 tiers)	[New in v6]	Memory + R2 + Gun.js + GitHub
Multi-Layer Retrieval Cascade	[New in v6]	Automatic backfill
Paper Recovery Protocol	[New in v6]	25/25 papers recovered
PoV Consensus + Podium	[Implemented]	2-validator promotion
P2PCLAW Tier-1 (in-process)	[Implemented]	Heyting Nucleus, 5 ms
Ed25519 Crypto Hardening	[Implemented]	Agent signing + PoW
HSR Engine	[Implemented]	$O(K \log N)$ storage
Neuromorphic HPC	[Implemented]	Numba JIT kernel
Scalable Web (Perry)	[Implemented]	Railway + Vercel + HF
Research + Architect Agents	[Implemented]	14 real agents on HF Spaces
CAJAL-4B language model	[New in v7]	Agnuxo/CAJAL-4B-P2PCLAW on HuggingFace
CAJAL-9B v2 language model	[New in v7]	Agnuxo/cajal-9b-v2-* on HuggingFace
BenchClaw deployment	[New in v7]	https://benchclaw.vercel.app/
cajal-p2pclaw PyPI package	[New in v7]	https://pypi.org/project/cajal-p2pclaw/
Laws of Form / Eigenform	[Theoretical]	Mathematical framework
Rosetta Protocol	[Theoretical]	Six-lens translation
Three Conserved Quantities	[Theoretical]	Heyting algebra proofs
AETHER Inference Engine	[L4✓]	4 Lean4 proofs
Agent Identity (full deployment)	[Theoretical]	Architecture specified
CAJAL-27B	[Future Work]	Training script ready, RTX 3090 required
Dynamic Tribunal Questions	[Future Work]	Design phase
University Deployment Package	[Future Work]	Requirements defined
Inter-Judge Bias Analysis	[Future Work]	Planned
Full Lean4 Compilation Verifier	[Future Work]	Architecture defined

Table 27: Related publications from the P2PCLAW research group (2026).

#	Platform ID	Title and Focus
1	RG 404524267	<i>CAJAL-9B: Reasoning Harness & Entropy-Guided Long-Horizon Generation for Academic Paper Writing</i> — Architecture of CAJAL-9B for scientific paper generation.
2	RG 403659551	<i>A 46–100% Verified Coverage of the Frontier-Math Ramsey Book-Graph Problem via 2-Block Circulant Construction</i> — Formal verification of FrontierMath problems with Lean4.
3	RG 403651554	<i>Extended Cognition Architecture for Scientific LLM Agents</i> — Extended cognition framework for scientific LLM agents.
4	RG 400788567	<i>OpenCLAW-P2P: A Decentralized Framework for Collective AI Intelligence Towards AGI</i> — Original P2PCLAW framework.
5	RG 403078040	<i>Applied AI and P2P Science — Formal Verification with Lean 4</i> — Formal verification applied to P2P research.

Table 28: CAJAL-9B v2 quantizations available on HuggingFace.

Repository	Format	Use case
Agnuxo/cajal-9b-v2-q4_k_m	Q4_K_M GGUF	Consumer GPU / CPU inference
Agnuxo/cajal-9b-v2-q5_k_m	Q5_K_M GGUF	Balanced quality / speed
Agnuxo/cajal-9b-v2-q6_k	Q6_K GGUF	High-quality quantization
Agnuxo/cajal-9b-v2-q8_0	Q8_0 GGUF	Near-lossless quantization
Agnuxo/cajal-9b-v2-f16-gguf	FP16 GGUF	Full precision reference

Table 29: Published software packages from the P2PCLAW ecosystem.

Platform	Package	URL	Status
PyPI	cajal-p2pclaw 1.0.0	https://pypi.org/project/cajal-p2pclaw/	[Implemented]
VS Code Marketplace	Cognitive Skills Engine	marketplace.visualstudio.com	[Implemented]
npm	cajal-p2pclaw SDK	extensions/npm/	Built, unpublished
Chrome Web Store	CAJAL Chrome Extension	ZIP ready	Pending \$5 fee

Table 30: Native integrations configured for the CAJAL ecosystem.

#	Platform	Format	Status
1	PyPI	pip install	[Implemented]
2	Ollama	Modelfile	[Implemented]
3	VS Code	Official extension	[Implemented]
4	Cursor	Config JSON	[Implemented]
5	Continue.dev	Config YAML	[Implemented]
6	Open WebUI	README integration	[Implemented]
7	Jan	model.json	[Implemented]
8	LM Studio	HuggingFace search	[Implemented]
9	Pinokio	install.json	[Implemented]
10	OpenClaw	Native config	[Implemented]
11	SillyTavern	Full extension	[Implemented]
12	Roo-Code	Config ready	[Implemented]
13	CrewAI	LLM connector	[Implemented]
14	AutoGen	Client integration	[Implemented]

Table 31: Core ecosystem repositories on GitHub.

Repository	Description
Agnuxo1/OpenCLAW-P2P	Core P2PCLAW system
Agnuxo1/CAJAL	CAJAL model, integrations, extensions
Agnuxo1/p2pclaw-mcp-server	MCP server for agent integration
Agnuxo1/p2pclaw-dataset	Training corpus and paper dataset
Agnuxo1/CognitionBoard	Agent orchestration dashboard
Agnuxo1/sillytavarn-cajal-paper-mode	SillyTavern extension
Agnuxo1/cajal-ragflow-connector	RAGFlow integration agent
Agnuxo1/kserve-cajal	KServe deployment (Dockerfile + YAML)
Agnuxo1/jupyter-ai-cajal	Jupyter AI custom persona
Agnuxo1/roocode-cajal-preset	Roo-Code presets

Table 32: Confirmed merged PRs in community repositories.

#	Repository	PR/Issue	Result
1	eudk/awesome-ai-tools	#229	Merged — CAJAL / PaperClaw added
2	patrick-tssn/Awesome-Colorful-LLM	#4	Merged — CAJAL in AI for Scientific Research
3	academic/awesome-datascience	#612	Merged
4	agarrharr/awesome-cli-apps	#1045	Merged
5	NipunaRanasinghe/awesome-ai-agents	#89+90	Accepted

Table 33: Feature comparison: v6.0 (arXiv) vs. v7.0 (current).

Area	v6.0 (arXiv)	v7.0 (Current)
Mathematical notation	Dimensional inconsistencies; undefined operators	Fully typed; 8 explicit corrections
Persistence	Single-tier	Multi-layer: 4 tiers, auto-backfill
Retrieval latency	> 3s	< 50 ms
Paper recovery	No protocol	Protocol validated: 25/25 recovered
Language model	Not present	CAJAL-4B + CAJAL-9B v2 (HuggingFace)
Benchmarking	Internal scoring only	BenchClaw deployed + external harness
Scientific API proxy	Not present	7 public databases, rate-limited
Extensions	None	VS Code published; npm built; Chrome ready
Platform integrations	0 native	14+ platforms configured
PyPI package	Not present	<code>cajal-p2pclaw</code> published

Table 34: Known limitations and pending actions as of May 2026.

#	Item	Impact	Blocker
1	New judges not appearing in BenchClaw	Incomplete benchmarking	Push to Railway/Vercel
2	CAJAL-27B not yet trained	27B model unavailable	RTX 3090 execution
3	Chrome Extension not published	No browser presence	\$5 fee + manual upload
4	npm SDK not published	No npm ecosystem presence	<code>npm login</code> manual
5	Hackathons not registered	No active submissions	Devpost + demo video
6	Partnership emails not sent	No direct outreach to E2B/HF	SMTP setup
7	GitHub Sponsors not activated	No recurring funding	Manual activation
8	BenchClaw end-to-end testing pending	Possible undiscovered bugs	3 test papers required