

WHO I AM / WHAT EXISTS

Mohammad Rahimi

Founder-Architect / AI Systems Architect whose record spans real-world product formation, AI/LLM systems architecture, human-context modeling, optimization, compute, security and trust, and foundational intelligence research.

Systems decomposition

Cross-layer architecture

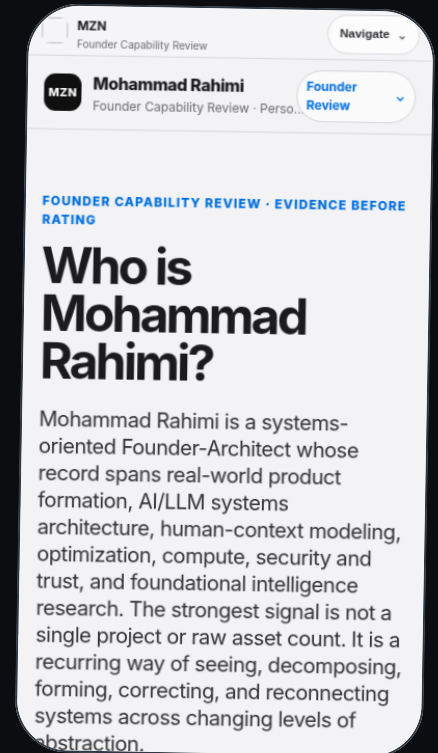
Human context

Security & trust

AI-native formation

A compressed map of a portfolio that is too large to fit into one pitch, one PDF, or one review session.

Reading rule: strong conclusions, explicit boundaries, staged proof. This document is a review surface - not the full technical corpus and not an independent certification.



21

Canonical LLM capability slots

554

Canonical Anatomy nodes: 21 roots + 533 descendants

330+

Phase 2 mapped assets & sub-assets

22+

Patent-grade candidates pending professional IP review

External platform signal

Crunchbase: sustained Top-10 visibility since May 2026

Founder-reported, supported by a dated screenshot archive retained for review.

Observed category range

#1-#2 global in AI / Machine Learning / Cybersecurity

Presented as a visibility signal, not technical certification or valuation evidence.

THIS IS A REVIEW SURFACE - NOT THE FULL CORPUS

The portfolio should be judged through a staged evidence process, not by how much can be reproduced on one page.

MZN's technical and research corpus is materially larger than can be reproduced in a single executive document. This PDF therefore provides a capability map, representative evidence, maturity boundaries and review routes rather than attempting exhaustive disclosure.

Absence from this review surface is not evidence of absence. Restricted or unsurfaced evidence, however, is not independent validation. Both distinctions matter.

01 · PUBLIC REVIEW SURFACE

Understand the claim

Founder, phase, capability and technical pages establish the shape of the case, terminology, chronology, maturity and non-claims.

02 · REPRESENTATIVE EVIDENCE

Inspect depth, not just titles

Selected deep artifacts show how the corpus moves from naming a subject into architecture, decisions, trade-offs, failure modes, implementation or test.

03 · SCOPED DILIGENCE

Ask for the branch that matters

If a capability, timeline, implementation or endpoint matters to the judgment, request the relevant underlying material rather than demanding the entire archive at once.

04 · CONTROLLED / FULL RELEVANT CORPUS

Expand when the question requires it

Where justified and permitted, review can expand from selected files to a complete capability family or the full relevant evidence set under controlled access.

Reviewer instruction: challenge the map. Ask technical questions. Request the evidence behind any claim that materially affects your judgment. The review architecture is designed to answer specific challenges without pretending the entire corpus belongs in a pitch deck.

PDF = map

Orientation, coverage, boundaries.

Evidence room = proof surface

Files, provenance, implementation and tests.

Phase 3 = validation

Reproduction, benchmarks, legal/IP review and pilots.

THREE DIFFERENT QUESTIONS - THREE DIFFERENT ANSWERS

Coverage is not readiness. Maturity is not validation. Disclosure is not proof.

Coverage %

CAPABILITY BREADTH

Estimated share of the canonical capability domain represented across the mapped corpus, structured inventories, representative deep artifacts, known implementation/testing records and backing materials staged for controlled review.

Maturity

HOW FAR THE WORK MOVED

Documented → architecture → implemented → internally tested → operating → independently validated. Different assets can sit at different levels inside the same slot.

Validation

WHO HAS TESTED IT

Internal evidence and controlled review can justify serious evaluation. Independent reproduction, benchmark, legal/IP review or production validation are separate stages.

Important: Coverage is not the fraction of backing files reproduced in this PDF or opened in one session. A large corpus cannot be fairly scored by session exposure. The estimate reflects the mapped portfolio; evidence access is staged according to the review question, confidentiality and validation need.

What the percentage does mean

- How broadly the portfolio appears to cover the canonical capability family.
- Whether multiple sub-surfaces are represented rather than a single title.
- Whether representative artifacts show real depth behind the inventory.
- Whether the capability recurs across dated architecture, implementation or review materials.

What the percentage does not mean

- Percent of a frontier laboratory already built.
- Production readiness or benchmark percentile.
- Independent technical, legal or commercial validation.
- Proof that every backing artifact has been publicly disclosed.

90-100

Very high coverage

80-89

High coverage

70-79

Substantial coverage

60-69

Meaningful but incomplete

Legacy comparison: the older MZN position map used 7 Strong / 13 Partial / 1 Gap and an older 529-sub-endpoint figure. This edition uses the newer independent Canonical Anatomy baseline: 21 slots / 554 total nodes, and recalibrates coverage without simply translating those legacy labels into percentages.

CANONICAL BASELINE + MZN EVIDENCE LAYERS

The review starts from an independent anatomy, then maps MZN evidence onto it.

243

Pre-training nodes

93

Alignment nodes

71

Evaluation & Safety
nodes

63

Inference &
Production nodes

84

Cross-cutting nodes

REPRESENTATIVE TECHNICAL FAMILIES

Tokenizer	Implemented/internal-tested model-system infrastructure.
GPU Sentinel	Telemetry, anomaly, FinOps, audit and hardware trust surfaces.
Training Recipe	Model selection through checkpoint management.
LLM Anatomy	21-slot, 554-node capability baseline.
Optimization	DCA, UIOP, Multi-Brain, Suprompt, OFRP and related candidates.

CROSS-SYSTEM / RESEARCH FAMILIES

HUAI	Context, confidence, memory, routing, permission and consequence architecture.
ZOE / ISBP	Security, trust, control, threat discovery and restricted mitigation research.
Mother / Genesis	Root trust, lineage, governance, resilience and privileged-control architecture.
Data & Privacy	Object discipline, reuse separation, lineage, export and review logic.
Zoyan / BioCode / HDTP	Convergence, foundational research and transport concepts with distinct maturity.

330+

MAPPED PHASE 2 ASSETS & SUB-ASSETS

Systems, architectures, frameworks, protocols, specifications, components, research assets and review infrastructure.

8

PORTFOLIO DOMAINS

Used to organize the Phase 2 asset ledger; not a claim that every asset has the same maturity.

22+

PATENT-GRADE CANDIDATES

Candidate inventions/claim families pending qualified professional IP review; not granted patents.

Count hygiene: 554 anatomy nodes are not 554 proprietary inventions. 330+ mapped assets/sub-assets are not 330 finished products. 22+ patent-grade candidates are not granted or allowed patent rights.

COVERAGE ESTIMATE · CURRENT CORPUS SYNTHESIS

21-slot capability map - part I.

ID	CAPABILITY	COVERAGE	COVERAGE BAR	CURRENT MATURITY / BOUNDARY
A · Pre-training stack				
A1	Data	90%		High architecture breadth; data/governance/provenance depth. Phase 1 source-system history is contextual, not Phase 2 solo proof.
A2	Tokenizer	92%		Implemented + internally tested; independent frontier-scale benchmarking pending.
A3	Architecture	86%		Deep documented model/system architecture; no claim of a frontier foundation-model build.
A4	Training	84%		Reviewer-grade training recipe across 8 sections; full-scale execution remains compute-dependent.
A5	Compute	67%		Substantial architecture + GPU monitoring implementation; frontier-scale training cluster execution not demonstrated.
B · Alignment / post-training				
B1	SFT	79%		Structured fine-tuning strategy and training decisions; large reproducible SFT campaign not established here.
B2	Preference Optimization	76%		Documented architecture / method coverage; implementation and comparative outcome evidence remain selective.
B3	Constitutional Methods	78%		Safety/governance architecture with explicit policy boundaries; independent efficacy validation pending.
B4	Red-Teaming	86%		Broad threat/security corpus, structured tests and explicit gap-closure priorities; maturity varies by sub-family.

Largest recalibration in this block: A4 Training and A5 Compute are no longer represented well by the legacy Partial/Gap labels alone. The new files show substantial training and compute architecture, while the boundary remains clear: architecture and monitoring do not equal frontier-scale training execution.

COVERAGE ESTIMATE · CURRENT CORPUS SYNTHESIS

21-slot capability map - part II.

ID	CAPABILITY	COVERAGE	COVERAGE BAR	CURRENT MATURITY / BOUNDARY
C · Evaluation & Safety				
C1	Capability Evaluation	88%		Structured test architecture, comparison logic and review surfaces; external benchmark replication still open.
C2	Safety Evaluation	91%		High documented coverage with hostile-review/failure-analysis logic; independent validation remains separate.
C3	Robustness	87%		Recovery, failover, anomaly, self-healing and stress/failure logic materially strengthen the old Partial view.
C4	Output Safety	93%		Strong output gating, egress, provenance and pre-commit control architecture; selected implementation depth varies.
D · Inference & Production				
D1	Serving	77%		Substantial architecture and runtime planning; broad production serving at frontier scale not demonstrated.
D2	Inference Optimization	91%		DCA, UIOP, Multi-Brain, Suprompt, OFRP and memory/routing candidates; measured gains remain workload-dependent.
D3	Monitoring	94%		Strong architecture plus GPU Sentinel implementation/internal testing; independent performance validation open.
D4	Deployment	75%		Operational/deployment architecture exists; Phase 3 institutional deployment is intentionally future work.
E · Cross-cutting				
E1	Data Governance	91%		Object-first, lineage, reuse, retention, consent, export and dataset-gating architecture; hearing-grade internal pack.
E2	Security	96%		Very broad cross-layer architecture with selected implementation/internal tests and a large restricted research corpus.
E3	Privacy	89%		Strong architectural treatment of classification, minimization, identity, lineage and reuse; legal validation separate.
E4	Compliance	82%		Governance, audit, approval and evidence-routing maturity; jurisdiction-specific legal certification not claimed.

Highest coverage:

E2 Security 96%
D3 Monitoring 94%
C4 Output Safety 93%

Materially upgraded:

E1 Data Governance
E3 Privacy
C3 Robustness

Hard boundary:

Coverage can be high while independent validation remains pending.

The scorecard is designed to be revised if deeper diligence changes the evidence.

WHY THE OLD MAP CHANGES

Pre-training coverage is broader than a simple Strong / Partial / Gap label can express.

A1 · DATA · 90%

Data is treated as a system, not a dataset title.

Evidence spans source-system history, consented signals, classification, provenance, lineage, reuse separation, training separation, dataset gating and review/export discipline. A separate Data Curation family is also identified in the archive inventory and is appropriate for controlled review.

Boundary: Phase 1 product/market data systems provide context and source history; they must not be relabeled as Phase 2 solo construction.

A2 · TOKENIZER · 92%

One of the clearest implementation-level technical assets.

Coverage includes BPE, WordPiece, Unigram/SentencePiece families, multilingual and mixed-script handling, critical-term/boundary preservation, fragmentation/context budget, runtime compatibility and internal test ladders.

Boundary: independent frontier-scale compression/downstream-quality benchmarking is still required.

A3 · ARCHITECTURE · 86%

Model and system architecture

Broad decomposition of transformer/model-system surfaces plus routing, context, memory and orchestration architecture. System-level HUI mechanisms do not automatically count as foundation-model implementation.

A4 · TRAINING · 84%

Eight-section training recipe

Model selection, fine-tuning method, optimizer, LR schedule, batch strategy, stability control, parallelism and checkpoint management are documented with decision logic and failure modes.

A5 · COMPUTE · 67%

Meaningful architecture, hard scale boundary

Distributed-training strategy, FSDP/parallelism, hardware/interconnect awareness and GPU monitoring exist. Frontier-scale cluster ownership and large pretraining execution are not demonstrated.

TRAINING RECIPE SECTION	WHAT IT CONTRIBUTES TO THE EVIDENCE MAP
1. Model Selection	Base-model, licensing, language, size/quality/cost decisions.
2. Fine-Tuning Strategy	LoRA/QLoRA/full tuning, rank, target modules, memory optimization.
3-5. Optimization of the run	AdamW, LR finding, warmup/schedules, effective batch and dynamic batching.
6. Stability Control	Loss spikes, gradient clipping, NaN/Inf, rollback and early stopping.
7-8. Parallelism & Checkpoints	DDP/FSDP/tensor/pipeline strategy, multi-node coordination, storage and resumption.

Interpretation: the Training Recipe is strong evidence of training-system literacy and reviewer-grade architecture. Its own closing boundary is also useful: execution still depends on compute access, base-model selection and engineering resources.

FROM POLICY TO CHALLENGEABLE REVIEW

Safety coverage is strongest when it produces tests, failure paths, review logic and release boundaries.

79%

SFT

Training/post-training decisions and fine-tuning strategy are mapped; full comparative campaigns remain future validation work.

76%

PREFERENCE OPTIMIZATION

Meaningful architecture and optimization context; implementation/outcome evidence is more selective than the coverage map.

78%

CONSTITUTIONAL METHODS

Policy, permission, governance and safety architecture recur across HUI/ZOE and related review materials.

86%

RED-TEAMING

Threat research, adversarial review, gap closure and security-oriented test surfaces materially strengthen coverage.

88%

CAPABILITY EVALUATION

Structured evaluation frameworks, comparison logic and test matrices create a reproducible review path.

91%

SAFETY EVALUATION

Safety is treated as a measurable review layer rather than a generic moderation paragraph.

87%

ROBUSTNESS

Recovery, failover, anomaly, containment, self-healing and failure injection broaden the older map.

93%

OUTPUT SAFETY

Egress, output gating, provenance, pre-commit arbitration and consequence-aware controls form a coherent family.

HUI diagnostic foundation

L0-L8 · 51 deep-dives · 92 structured tests

Current HUI architecture material describes a layered diagnostic system intended to map before intervention.

Data & Privacy review pack

48 worked cases · 12 reviewer simulations · 12 failure injections

This is a separate hearing-grade internal pack focused on object handling, reuse, privacy, export and adversarial review.

Do not collapse these counts: 92 structured tests and the 48/12/12 Data & Privacy review figures come from different artifacts and should not be added together as one universal test count without reconciliation.

A credible evaluation architecture should make it possible to lower the score, not only confirm it.

RUNTIME INTELLIGENCE

Efficiency and monitoring are treated as architectural control problems, not only infrastructure costs.

D2 · INFERENCE OPTIMIZATION · 91%

Selective activation and structured reuse

DCA	Dynamic Contextual Activation - activate only what the task needs.
UIOP	User-intelligence / permitted-context optimization and staged routing.
Multi-Brain	Specialized routing instead of one undifferentiated reasoning path.
Suprompt	Intent clarification and staged prompt/interface architecture.
OFRP	Safe reuse/cache candidates and output-first reverse reasoning paths.

D3 · MONITORING · 94%

GPU Sentinel moves monitoring beyond a concept map.

Selected implementation/internal testing covers accelerated-infrastructure telemetry, workload attribution, anomaly detection, operational controls, FinOps, audit/forensics and hardware-isolation review surfaces, with integrations around NVML, DCGM, CUPTI, Kubernetes and observability connectors.

Boundary: internal technical evidence is not a substitute for independent throughput, detection-quality or production reliability benchmarks.

D1 · SERVING · 77%

Architecture is meaningful; scale remains open.

Serving, runtime, routing, inference, safety and observability surfaces are represented across the corpus. This does not establish frontier-scale serving fleet experience.

D4 · DEPLOYMENT · 75%

Deployment is intentionally a Phase 3 validation surface.

Release gates, operational monitoring, recovery, governance, infrastructure and staged-product paths are documented. Institutional deployment, hardening and commercialization remain next-stage work.

Important separation: architecture candidates such as DCA, UIOP, Multi-Brain, Suprompt and OFRP should not be presented as proven efficiency gains until measured against controlled workloads and baselines. The score is coverage, not claimed savings.

Architecture

Broad and cross-layer.

Selected implementation

Strongest in Tokenizer and GPU Sentinel.

Independent production validation

Still a Phase 3 objective.

Serving and deployment remain intentionally more conservative than monitoring/optimization.

OBJECT-FIRST DISCIPLINE

Governance is strongest where the system can explain what an object is, where it may go, and why.

01

OBJECT-FIRST DISCIPLINE

Classify meaningful objects before routing, reuse, retention, export or promotion decisions become valid.

02

REUSE SEPARATION

Production traces, feedback, evaluation, synthetic material and training candidates are treated as different zones.

03

SURFACE AWARENESS

The same object may appear differently across runtime, logs, audit, analytics and export surfaces.

04

EXPORT MATURITY

Reviewer/evidence bundles are derived objects with lineage, redaction review and purpose boundaries.

91%

E1 · DATA GOVERNANCE

High architecture coverage across classification, lineage, consent, retention, provenance, training separation, dataset gating and review/export logic.

89%

E3 · PRIVACY

Identity resolution, minimization, pseudonymization/tokenization concepts, surface boundaries and derived-object handling are explicitly addressed.

82%

E4 · COMPLIANCE

Governance, role separation, approvals, audit, chain of custody and evidence routes are meaningful; jurisdiction-specific legal validation remains open.

48

Worked cases

12

Reviewer simulations

12

Failure injections

8

Doctrine domains

Phase 1 context

Mazzaneh's product/market systems provide historical evidence of consented attributes, attention/comprehension, demand, preference and outcome-oriented signals. This is Phase 1 team/company context and must retain its provenance boundary.

Phase 2 abstraction

HUAI and Data/Privacy materials turn signal handling into explicit architecture: source, semantic type, confidence, recency, contradiction, permission, reuse and consequence. That architecture is evaluated separately from Phase 1 team execution.

Legal boundary: strong governance/privacy architecture is not a claim of GDPR, AI Act, sectoral or jurisdiction-specific compliance certification. Qualified legal review belongs to Phase 3.

SECURITY IS A SYSTEM, NOT A LIST OF NAMES

The corpus spans multiple security sub-disciplines that reconnect into trust, control, recovery and audit.

Root Trust

Origin vectors, certificate lineage, trust anchors and entropy-rooted control.

Identity & Privilege

Sensitive identity mapping, ephemeral privilege and one-time operations.

Runtime Security

Anomaly prediction, privileged endpoints, containment and state controls.

Output / Action Control

Reality-bound gating, pre-commit risk arbitration and egress discipline.

Audit & Provenance

Tamper-evident ledgers, lineage and summary/blind verification paths.

Recovery

Recovery forks, failover, sealed contexts and controlled disassociation.

Crypto Lifecycle

Key cycling, non-persistent anchors, zeroization and invalidation.

Infrastructure

GPU monitoring, isolation, hardware trust and operational telemetry.

Adaptive Defense

Entropy-leak detection, anomaly response and controlled self-healing.

Governance

Root policy, emergency authority, approval paths and evidence boundaries.

Mother / Genesis v2

30 professionally reframed components

The file explicitly moves older opaque/backdoor-like language toward root-of-trust, attestation, governance, continuity, lineage, resilience and controlled recovery. It also identifies formalization, embodiment and policy gaps rather than declaring every component complete.

ZOE / ISBP / restricted security corpus

Broader cross-layer research depth

Legacy inventories, threat-modeling notes and restricted security branches support breadth and continuity, but sensitive or exploit-enabling mechanics are not appropriate for this public review surface.

Coverage strength

Very broad across security/control families.

Maturity

Architecture is strong; implementation/testing is heterogeneous.

Validation

No claim that every protocol is independently validated or production-deployed.

Why the score is high: not because 30 titles exist, but because representative materials connect trust roots, privilege, runtime, audit, recovery, lifecycle and governance into a coherent control architecture with explicit maturity boundaries.

A SERIOUS MAP SHOULD SHOW WHAT STILL NEEDS WORK

Gap identification is evidence of review discipline - not evidence that the gap is already closed.

The HUI Gap Closure Dashboard explicitly prioritized missing or underdeveloped security/protocol surfaces. This matters because the corpus is not being presented as uniformly complete.

01

Runtime Alignment Monitoring

Live observability of alignment/runtime behavior inside the strongest existing security zone.

02

Adversarial Input Fuzzing

Turn input-security reasoning into systematic stress testing rather than concept-only coverage.

03

Output Provenance

Bridge output control to authenticity, lineage and verifiable origin.

04

Data Poisoning Detection

Make data/training integrity explicit rather than implied by general security architecture.

05

MFA for Sensitive Operations

Strengthen privileged governance with a practical hardening control.

06

Supply Chain Security

Extend trust and integrity into the enterprise model/data/software pipeline.

Recognizing a missing layer is not the same as implementing it. The next evidence step is dossier → mechanism → test → result → independent challenge.

How this affects scoring

A known gap can reduce maturity or validation without erasing broader capability coverage. The coverage score asks whether the domain is represented; maturity asks how far it has been built; validation asks whether the result survives independent test.

How this affects diligence

A reviewer should be able to ask: "Show me the gap dossier, the proposed mechanism, the implementation status and the test result." If the evidence does not survive that sequence, the assessment should be revised downward.

Review posture: this document does not convert planned gap closure into completed evidence. It records the priorities so the next stage can test whether the gaps were actually closed.

ARCHITECTURE OF KNOWING, PERMISSION AND CONSEQUENCE

HUAI connects signal quality to selective reasoning and closes the loop through outcome.



RECURRING DISTINCTION

Capability is not authority.

Knowing something about a user or environment does not automatically authorize the system to use it. Source, confidence, recency, contradiction, saliency and permission change what should activate.

Signal is not automatically truth.

Declaration, inference, behavior, preference and outcome-supported evidence are treated as different semantic classes.

SELECTIVE INTELLIGENCE

Do not load the whole human into every prompt.

Routing and selective activation are intended to keep context relevant, permissioned and efficient rather than always-on.

Output is not the end of the loop.

Outcome and consequence should update the system's confidence, context and future behavior.

DIAGNOSTIC

L0-L8

HUAI uses a layered model/system flow map as a diagnostic spine. This should not be confused with the 21-slot organizational LLM Anatomy.

CONVERGENCE

Zoyan

Human-facing convergence is the intended Phase 3 return path where validated capabilities may reconnect into a minimal closed-loop product.

FOUNDATIONAL RESEARCH

BioCode

Limitation, embodiment, consequence, memory, saliency and trust are explored as foundational intelligence questions - not proof of AGI or foundation-model capability.

Conceptual lineage note: the "Pre-Structured Reality & Infinite Paths" file is retained as a BioCode philosophical working note about state spaces, path selection and consequence. It may illuminate thinking style, but it does not increase any LLM capability score by itself.

FORMATION SIGNATURE

The strongest signal is a recurring way of decomposing and reconnecting systems across changing levels of abstraction.

Observe reality → separate signals → qualify what is known → identify the constraint → design the architecture
→ activate what is relevant → establish trust/permission → act or emit → observe consequence → update.

VERY HIGH

Systems Decomposition

Repeated decomposition across product, AI, security, context, trust and human-facing systems.

VERY HIGH

Cross-Layer Architecture

Connections across signals, representation, memory, routing, compute, trust, permission, orchestration and outcomes.

VERY HIGH

Technical Breadth

A large multi-domain working corpus with distinct technical branches rather than a single-topic archive.

VERY HIGH

Human / Context / Intent Modeling

Recurring separation of declaration, qualification, behavior, confidence, contradiction, salience and outcome.

VERY HIGH

Security / Trust Systems Thinking

Behavior, identity, session, privilege, audit, isolation, model/data protection and lower-layer threat-modeling surfaces.

VERY HIGH

Integration Across Domains

Independent branches repeatedly reconnect into operating loops rather than remaining isolated concept lists.

Current working positioning: Senior / Specialist to Principal / frontier-adjacent architecture, depending on domain. The strongest signal is cross-layer systems formation and integration.

Explicitly unproven: frontier-scale model training and institutional-scale AI/research execution. Those require external resources, specialists and independent test.

Breadth

Different system surfaces.

Depth

Decisions, trade-offs, failures, testable structure.

Recurrence

Multiple dated generations and domains.

Integration

Coherent operating relationships.

Maturity

Hypothesis through validation.

Historical context

What was advanced when produced.

Assessment transparency: this profile is an AI-assisted evidence synthesis of the documented corpus. It is not an external third-party certification or independent professional assessment. Its purpose is to make the record easier for qualified reviewers to inspect, challenge, downgrade, reproduce or validate.

SUSTAINED VISIBILITY - NOT A ONE-DAY SPIKE

Founder-reported Crunchbase visibility has remained in the Top 10 across relevant categories since May 2026.

Dated screenshots are retained by the founder for reconciliation and review. In categories including Artificial Intelligence, Machine Learning and Cybersecurity, the observed global position has reportedly fluctuated around #1-#2 during the period.

May 2026

Top-10 period begins in the retained screenshot record.

June 2026

Sustained category visibility continues.

July 2026

Top-10 placement remains visible across relevant categories.

Aug 2026

Current review still reports Top-10 presence; selected categories around #1-#2.

Top 10

Sustained across relevant categories

#1-#2

Observed range in AI / ML / Cybersecurity

May → Now

Persistent review-period signal

REPORTED CONDITIONS

Visibility without a conventional promotion stack

- No conventional PR team or campaign.
- No paid Crunchbase account.
- No announced fundraising round used as the visibility engine.
- No major paid promotional placement identified by the founder.

CORRECT INTERPRETATION

Reason to investigate - not proof of the technical case

A sustained platform signal can justify deeper review because it is external to the internal corpus. It does not prove technical capability, asset value, the one-person formation boundary, valuation or independent validation.

The strongest use of the Crunchbase signal is not "a platform certified the portfolio." It is: "an unusual external visibility pattern exists; inspect the underlying record."

Evidence status: this PDF records the founder's report and the existence of a dated screenshot archive. Exact date-by-date positions should be reconciled against that archive before using any frozen rank claim in press, legal or diligence materials.

THREE PHASES · DIFFERENT PROVENANCE RULES

Do not judge the portfolio only from what fits on these pages. Ask the questions that matter and follow the evidence as deep as the judgment requires.

PHASE 1

Founder + Team + Real Market

Original Mazzaneh product/ company operation, team execution, users/businesses/market signals and commercial/product context. Phase 1 is not the Phase 2 solo valuation base.

PHASE 2

Founder + AI Tools - Human Formation Team

Bounded one-human AI-native formation of systems, architectures, research programs, documentation and selected implementation/testing. Provenance remains auditable.

PHASE 3

Institutional Validation + Specialists

Rebuild/hardening, benchmarks, legal/IP/compliance review, pilots, commercialization and selective team/partner formation. This is where claims are tested, revised or rejected.

Phase 3 exists to answer the open questions experimentally - not rhetorically.

CHALLENGE THE MAP

- 1 Ask a specific question**
Capability, chronology, implementation, maturity, score or claim boundary.
- 2 Request the relevant evidence**
Public, controlled, restricted or full capability-family review as appropriate.
- 3 Reproduce what matters**
Define baseline, test criteria, failure conditions and independent reviewer.
- 4 Change the assessment if required**
Strengthen, narrow, reject, rebuild or validate based on the result.

CORE REVIEW ROUTES

Start	mzncompany.com/start/
Founder	mzncompany.com/mohammad/
LLM Systems	mzncompany.com/llm/
HUAI	mzncompany.com/huai/
Evidence	mzncompany.com/evidence/
Review	mzncompany.com/review/
Evaluate	mzncompany.com/evaluate/
Partnership	mzncompany.com/partnership/

Final review principle: the objective is not predetermined agreement. The objective is an evaluation process large enough for the object being evaluated - one that can inspect representative depth first, expand evidence when necessary, and independently validate what still matters.

Primary source basis for this edition: Canonical LLM Company Anatomy v1.0; MZN public deploy v10 (Founder, Phase 2, IP/Portfolio, LLM, HUAI, Tokenizer, GPU, ZOE, Zoyan); HUAI Training Recipe; HUAI Data & Privacy Architecture; HUAI Gap Closure Pack; HUAI LLM Architecture v3; Mother / Genesis v2; Pre-Structured Reality & Infinite Paths working note; and selected legacy/security materials used only with explicit maturity guardrails. Founder-reported Crunchbase history is treated separately as an external-platform signal pending screenshot reconciliation.