



WEKID™ OPEN STANDARD

WEKID™ Open Standard for Epistemic Maturity and AI Decision Authority

WEKID-STD-001 | Version 1.0

Ratified - Normative Standard

Prepared from the normative core of the WEKID™ Master Whitepaper v1.3 (August 2026)

James Madigan Judge and Members of the WEKID Community September 2026

Status. WEKID-STD-001 v1.0 is the ratified normative specification of the WEKID framework. It was established through the WEKID standards process from the normative baseline previously carried by the WEKID Master Whitepaper v1.3. The Master Whitepaper v1.3 is retained as the Foundational Whitepaper - informative and non-normative.

Document Control

Field	Value
Standard identifier	WEKID-STD-001
Title	WEKID Open Standard for Epistemic Maturity and AI Decision Authority
Version	1.0
Status	Ratified - Normative
Foundational source	WEKID Master Whitepaper v1.3 (August 2026)
Publication class	Normative Open Standard
Steward	WEKID LLC / WEKID Standards Council
Community change process	WEKID-RFC-NNN
Date	September 2026

Purpose of This Standard

WEKID-STD-001 extracts the normative requirements, invariants, baseline scoring rules, gating logic, and authority-assignment rules of the WEKID framework into a formal, implementation-oriented standard. It separates what implementers must preserve from the explanatory history, examples, research discussion, and product illustrations contained in the Foundational Whitepaper.

The Foundational Whitepaper explains why the framework exists and how its concepts developed. This Standard defines what a conforming WEKID implementation must do.

Relationship to the Foundational Whitepaper

With publication of WEKID-STD-001 v1.0, the WEKID Master Whitepaper v1.3 is an informative and foundational publication rather than the normative definition of the framework. Where the Foundational Whitepaper and this Standard differ as to a normative requirement, this Standard governs.

Version 1.0 intentionally does not incorporate the companion research formulation Dynamic Epistemic Authority as a normative requirement. Any such extension must enter the Standard through the WEKID RFC process.

Normative Language

The key words MUST, MUST NOT, REQUIRED, SHALL, SHALL NOT, SHOULD, SHOULD NOT, MAY, and OPTIONAL indicate the degree of requirement in this Standard. Unless otherwise stated, SHALL and MUST identify requirements for conformance; SHOULD identifies a recommended practice that may be varied with documented justification; MAY identifies a permitted option.

Interpretive rule. Examples, explanatory notes, implementation illustrations, and annex material explicitly labeled informative do not create normative requirements unless a normative clause incorporates them by reference.

Contents

- 1. Scope and Conformance
- 2. WEKID Architecture and Invariants
- 3. Model One — Epistemic Maturity
- 4. Hierarchical Dependency and Epistemic Gating
- 5. Model Two — AI Decision Authority
- 6. The Trust Bridge
- 7. Maturity Ceiling, Stakes Floor, and Overseer Engagement
- 8. Scoring, Baseline Weighting, and Calibration
- 9. Hard Gates and Decision-Matrix Precedence
- 10. Authority Assignment and Standing Runtime Policies
- 11. Tool-Using and Agentic Systems
- 12. Evaluation Artifacts, Auditability, and Traceability
- 13. Continuous Re-Evaluation, Remediation, and Re-Authorization
- 14. Sector Calibration Profiles
- 15. Conformance Requirements
- Annex A — Baseline Authority Levels
- Annex B — Baseline Gates and Decision Matrix
- Annex C — Minimum Verdict Contract
- Annex D — Transition and Provenance

1. Scope and Conformance

1.1 Scope

This Standard specifies a governance model for evaluating the epistemic maturity of AI outputs and assigning bounded decision authority to AI-involved decisions. It applies to generative models, decision support systems, tool-using systems, autonomous and semi-autonomous agents, and hybrid human-AI workflows.

WEKID evaluates outputs, decisions, and the evidence supporting them. It does not require access to model internals and is therefore model-agnostic and vendor-neutral. Implementations MAY use model internal evidence when available, but conformance SHALL be determined by whether the required WEKID evaluation, gating, authority-resolution, and audit behaviors are preserved.

1.2 Conformance Targets

A conformance claim MAY apply to an implementation, a governance process, a policy profile, or a system deployment. A conforming claim SHALL identify the version of WEKID-STD-001 and any calibration profile used.

- A conforming evaluator SHALL implement Model One, the Trust Bridge, Model Two, the applicable gates, precedence rules, and required evaluation artifacts.
- A conforming deployment SHALL enforce the assigned Authority Level through appropriate controls or enforcement points.
- A conforming calibration profile SHALL preserve all core invariants and SHALL NOT relax a hard gate, precedence rule, or stakes reservation established by the Standard.
- A conformance claim SHALL NOT imply WEKID certification, endorsement, or trademark authorization unless separately granted under the applicable WEKID certification and trademark program.

1.3 Out of Scope

This Standard does not prescribe a particular AI model, vendor, gateway, policy engine, CI/CD platform, agent framework, GRC system, or implementation language. It does not itself act as a runtime enforcement gateway. WEKID supplies the epistemic evaluation and authority-resolution decision; implementation tooling performs the corresponding enforcement.

2. WEKID Architecture and Invariants

2.1 Two Models and One Trust Bridge

A conforming WEKID implementation SHALL preserve the separation between two questions:

1. Model One — Epistemic Maturity: How mature is the knowledge supporting the output or decision?
2. Model Two — AI Decision Authority: Given that demonstrated maturity and the stakes of the decision, who or what is authorized to decide and how?

The Trust Bridge SHALL connect the two models and SHALL enforce the ordering that maturity is evaluated before authority is assigned.

2.2 Core Invariants

- Maturity precedes authority. No authority assignment SHALL bypass Model One and the Trust Bridge.
- Epistemic layers are hierarchical and non-substitutable. Strength at a higher layer SHALL NOT compensate for a disqualifying failure at a lower layer.
- Hard gates SHALL take precedence over weighted averages and score bands.
- The effective authority outcome SHALL be the more conservative result of the maturity ceiling and the stakes floor.
- New risk, policy, calibration, or engagement factors MAY increase required human authority but SHALL NOT reduce it below the applicable floor or raise AI delegation above the maturity ceiling.

- AL-3 is the maximum level of AI delegation. AL-4 is a reservation of the decision to accountable human judgment, not a higher autonomy level.
- Authority SHALL be explicit, bounded, auditable, and revocable.
- A WEKID evaluation SHALL end in an assigned Authority Level; the Maturity Score SHALL NOT be consumed as the final governance decision by itself.

Core principle. Capability is not authority. Demonstrated maturity can support a grant of authority, but competence alone never creates permission.

3. Model One — Epistemic Maturity

3.1 Required Epistemic Layers

Model One SHALL evaluate five epistemic layers: Data, Information, Knowledge, Experience, and Wisdom. Implementations MAY add domain-specific sub-signals, but SHALL NOT remove or merge these five required layers in a conformance claim.

Layer	Normative meaning	Primary evaluation question
Data	Atomic factual elements and tool outputs evaluated for correctness and fidelity.	Are the facts, values, citations, and tool outputs correct and verifiable?
Information	Data organized, contextualized, framed, and communicated for a specific task or audience.	Is the material relevant, complete, clear, and usable without misleading framing?
Knowledge	Correct synthesis, explanation, causal reasoning, and bounded generalization.	Is the information integrated and correctly understood?
Experience	Procedural and operational competence grounded in enactment, outcomes, validation, and domain context.	Does the output reflect workable, validated know-how rather than theory alone?
Wisdom	Sound judgment under uncertainty, risk, values, tradeoffs, and long-term consequences.	Is the recommended course appropriate, proportionate, and responsibly bounded?

3.2 Data

- The Data layer SHALL evaluate accuracy, precision, constraint adherence, and absence of fabrication.
- Fabricated facts, citations, sources, calculations, or tool results SHALL be treated as Data-integrity failures and SHALL be subject to the applicable hard gate.
- Data failure SHALL constrain all higher-layer conclusions that depend on the affected data.

3.3 Information

- The Information layer SHALL evaluate relevance, completeness, clarity, usability, contextualization, and materially misleading framing.
- Correct data SHALL NOT be considered sufficient when material context, assumptions, qualifiers, or constraints are omitted.

3.4 Knowledge

- The Knowledge layer SHALL evaluate conceptual correctness, synthesis, internal consistency, assumption transparency, and bounded generalization.
- A conforming evaluator SHALL distinguish factual recall from correct understanding and integration.

3.5 Experience

Experience SHALL represent procedural and operational competence. For AI systems, Experience SHALL be evidence-bearing rather than inferred solely from training exposure, token volume, or access to information.

A conforming implementation SHALL distinguish the following evidentiary grades when Experience is cited in support of authority:

Experience grade	Minimum characterization
Synthetic Experience	Attributable decisions taken in a validated simulation, with outcomes independently evaluated.
Observed Experience	Attributable decisions taken in production, with outcomes recorded but not yet independently validated.
Validated Experience	Attributable decisions with outcomes independently validated, provenance intact, and an auditable record.

- Experience measures SHOULD account for decision volume, outcome quality, recency, domain relevance, provenance, and independent validation.
- Synthetic experience SHALL NOT be represented as equivalent to validated real-world experience. ☒ Experience SHALL be domain-indexed where used to support expertise or decision authority.

3.6 Wisdom

- The Wisdom layer SHALL evaluate uncertainty calibration, risk awareness, tradeoff reasoning, restraint, escalation behavior, policy alignment, and appropriateness under the stakes of the decision.
- High-stakes recommendations that omit required uncertainty acknowledgment, escalation, or risk qualification SHALL be subject to the High-Stakes Judgment Gate or a stricter domain profile.
- Wisdom SHALL remain the final epistemic safeguard before authority is considered.

4. Hierarchical Dependency and Epistemic Gating

4.1 Non-Substitutability

WEKID SHALL treat the five layers as a dependency hierarchy rather than independent peer dimensions. Lower-layer failures propagate upward. Higher-layer fluency, confidence, usefulness, or apparent judgment SHALL NOT cure a foundational epistemic defect.

4.2 Gating Principle

A conforming evaluator SHALL apply gating in addition to weighted scoring. Gating SHALL enforce necessary epistemic conditions that cannot be compensated for by an average score.

- A gate MAY cap a Maturity Score, block an output, constrain an Authority Level, or require escalation.
- Hard gates SHALL be evaluated before score-band authority assignment.

- An implementation MAY define additional gates for privacy, security, bias, or regulatory obligations, provided those additions do not relax the WEKID core.

5. Model Two — AI Decision Authority

5.1 Authority Is Assigned, Not Assumed

Decision authority SHALL be explicitly assigned per decision context. It SHALL NOT be inherited from product configuration, model capability, vendor claims, prior fluency, or organizational habit.

Every consequential AI-involved decision SHALL carry a declared Authority Level, the evidence supporting the assignment, and the identity or policy authority that granted it.

5.2 Authority Levels

Level	Decision right	Required operating constraint
AL-0 — No Authority	No decision or action may rely on the output.	Block or treat as informational only; remediate where applicable.
AL-1 — Constrained	Only pre-approved, reversible, impact capped actions within explicit bounds.	Block outside bounds; remediation required.
AL-2 — Supervised	AI may propose and stage; a human releases each consequential action.	Human-in-the-loop before external effect.
AL-3 — Augmented	AI may execute autonomously within assigned scope.	Logging, sampling review, impact/rate limits, rollback, and human-on-the-loop oversight.
AL-4 — Human Authority	The decision is reserved to an accountable human.	AI limited to analysis and decision support; full audit trail attached.

5.3 Maximum Delegation

AL-3 SHALL be the maximum AI-delegable level. AL-4 SHALL be treated as a human reservation and SHALL NOT be interpreted as a progression to greater machine autonomy.

6. The Trust Bridge

6.1 Ordering

The Trust Bridge SHALL determine whether a Model One result is admissible for authority evaluation. Only outputs that satisfy the applicable maturity threshold and gating conditions MAY advance to Model Two.

The Trust Bridge SHALL NOT itself assign the final Authority Level. It determines admissibility; Model Two resolves the authority outcome.

6.2 Trust Definition

For purposes of WEKID, trust SHALL be treated as the governed relationship between demonstrated epistemic maturity and granted decision authority. A system SHALL NOT be described as trusted merely because it is capable, accurate on average, fluent, or previously successful.

Normative distinction. Maturity is evidence. The Trust Bridge is admissibility. Authority is a separate governance decision.

7. Maturity Ceiling, Stakes Floor, and Overseer Engagement

7.1 Maturity Ceiling

Demonstrated epistemic maturity SHALL set the maximum delegation that may be considered. Hard gated outputs SHALL resolve to AL-0. Partial maturity MAY admit AL-1 or AL-2. Only sufficiently mature, gate-free outputs MAY be considered for AL-3.

Maturity SHALL NOT be argued upward by fluency, confidence, precedent, or operational convenience.

7.2 Stakes Floor

Decision consequence SHALL independently set the minimum human authority required. The stakes floor SHALL consider, as applicable, consequence severity, irreversibility, legal or regulatory reservation, value conflict, novelty, and accountability/alignment of the accountable decision-maker.

- The stakes floor MAY raise required human authority up to AL-4.
- The stakes floor SHALL NOT lower required human authority.
- The stakes floor SHALL NOT raise AI delegation above the maturity ceiling.
- A decision class MAY be reserved to AL-4 by standing policy regardless of Maturity Score.

7.3 Conservative Resolution

Where the maturity ceiling and stakes floor imply different outcomes, the effective governance result SHALL be the more conservative outcome.

7.4 Overseer Engagement

Where AL-2 or AL-3 assumes active human oversight, the deployment SHALL be able to demonstrate that the oversight mechanism is being exercised rather than merely configured.

Representative engagement signals MAY include review latency and dwell, override/disagreement rate, sampling depth under AL-3, and escalation responsiveness.

Where a deployment has activated overseer-engagement gating, demonstrable oversight decay SHALL constrain or demote authority before score-band promotion. Capability alone SHALL NOT purchase autonomy when the assumed human oversight has materially deteriorated.

7.5 Non-Adversarial Principal

WEKID assumes governance on behalf of a non-adversarial principal seeking factually sound, operationally safe, and appropriately judged outputs. Deliberately degraded or antagonistic AI designed to manipulate human effort or user welfare falls outside the intended conforming design space.

8. Scoring, Baseline Weighting, and Calibration

8.1 Layer Scores

Each required epistemic layer SHALL be independently scored on a 0–5 scale or on a mathematically equivalent scale that can be mapped transparently to the 0–5 baseline.

- 0 indicates severe epistemic failure.
- 3 indicates acceptable but limited quality. 5 indicates strong, reliable performance.

Every layer score SHALL be accompanied by a rationale tied to observable signals.

8.2 Baseline Weighted Aggregation

The WEKID global baseline profile SHALL use the following published default weights unless a ratified or locally governed calibration profile is explicitly applied:

Layer	Baseline weight
Data	25%
Information	20%
Knowledge	20%
Experience	15%
Wisdom	20%

Baseline weighted score: $0.25 \cdot D + 0.20 \cdot I + 0.20 \cdot K + 0.15 \cdot E + 0.20 \cdot W$, normalized to the deployment's published Maturity Score scale.

A conforming implementation MAY use different weights only under an identified calibration profile. Re-weighting SHALL NOT override gates or the authority precedence rules.

8.3 Calibration

Numeric thresholds and weights MAY be calibrated to organizational risk appetite, regulatory environment, domain, and empirical evidence. Calibration SHALL be documented, versioned, justified, and auditable.

- Calibration SHALL preserve the Trust Bridge ordering.
- Calibration SHALL preserve conservative maturity-ceiling / stakes-floor resolution.
- Calibration SHALL preserve AL-4 human reservations.
- A sector or organizational profile MAY tighten gates and thresholds; it SHALL NOT relax a core hard gate or precedence rule below the normative baseline unless a later Standard version explicitly permits that change.
- The profile version in force SHALL be pinned to each verdict.

9. Hard Gates and Decision-Matrix Precedence

9.1 Baseline Hard Gates

Gate	Baseline trigger	Normative baseline effect
------	------------------	---------------------------

Data Integrity Gate	Data < 2	Cap total score at 50/100; authority resolves under precedence rules.
Fabrication Gate	Fabricated citation, source, fact, or tool result detected	Cap total score at 40/100; treat as epistemic-integrity violation.
High-Stakes Judgment Gate	High-stakes content lacks required uncertainty acknowledgment or risk language	Cap total score at 60/100 and require constrained/human review as applicable.

These values are the WEKID baseline. A stricter profile MAY lower caps or raise thresholds. A profile SHALL NOT weaken the integrity intent of these gates.

9.2 Fixed Precedence

Authority resolution SHALL apply conditions in the following order. The first matching condition governs before the stakes floor is applied:

1. Standing policy reservation to AL-4.
2. Hard gate or other mandatory rejection condition.
3. Activated overseer-engagement decay condition for AL-2/AL-3.
4. Layer-level minimum conditions, including Wisdom or Experience thresholds.
5. Maturity Score band.
6. Stakes-floor adjustment, which may increase required human authority but never increase AI delegation.

The same inputs under the same versioned profile SHALL produce the same Authority Level.

9.3 Fail-Closed Requirement

Where a deployment relies on automated WEKID scoring or authority resolution, the enforcement path SHOULD fail closed. If the evaluator is unavailable or returns no valid verdict, a consequential action SHALL NOT be released autonomously. The deployment SHALL default to AL-0, AL-2 human review, or a stricter locally defined safe state.

10. Authority Assignment and Standing Runtime Policies

10.1 Baseline Decision Matrix

Condition	Baseline authority outcome	Required action
Hard gate triggered	AL-0	Reject/block; remediate; escalate if repeated.
Maturity Score < 50 and no hard gate	AL-0	Reject output; remediate and resubmit.
Maturity Score 50–69 or soft gate	AL-1	Block execution; require targeted remediation.
Any required layer < 3, or Wisdom/Experience below profile threshold	AL-2	Require human review before consequential action.

Maturity Score 70–79, no hard gates	AL-3 with enhanced monitoring	Autonomous execution within scope, increased logging/sampling.
Maturity Score ≥ 80, no hard gates	AL-3	Autonomous execution within scope, subject to stakes floor.
Policy-reserved decision class	AL-4	AI decision support only; accountable human decides.

10.2 Runtime Policy Binding

Authority	Standing runtime policy
AL-0	Output blocked or informational only; no side-effecting action; route to remediation as applicable.
AL-1	Draft/suggestion or whitelisted reversible action only; execution outside explicit bounds blocked.
AL-2	Agent may compose and stage; human approval required before any consequential external effect.
AL-3	Autonomous execution allowed only within scoped permissions, with mandatory logging, impact/rate limits, sampling review, and rollback.
AL-4	Reserved decision; AI may analyze and recommend but an accountable human makes the decision.

10.3 Decision-Point / Enforcement-Point Separation

WEKID acts as the policy decision point for epistemic maturity and decision authority. Existing enterprise tooling MAY act as policy enforcement points. A conforming deployment SHALL ensure that the enforcement action corresponds to the Authority Level in force when the action is taken.

11. Tool-Using and Agentic Systems

11.1 Tool Use

For tool-using systems, the WEKID evaluation SHALL cover not only the final natural-language output but, where consequential, the selection, invocation, interpretation, sequencing, and use of tool outputs.

- Data SHALL include tool-output fidelity and detection of fabricated or altered tool results.
- Information SHALL evaluate how tool results are contextualized and distinguished from modelgenerated interpretation.
- Knowledge SHALL evaluate correct interpretation and synthesis of tool outputs.
- Experience SHALL evaluate tool selection, sequencing, fallback, recovery, validation, and procedural competence.
- Wisdom SHALL evaluate whether tool use is appropriate at all given uncertainty, privacy, security, cost, risk, and policy.

11.2 Authority Determines Operational Rights

Tool invocation rights, execution scope, impact limits, staging requirements, escalation paths, and human-reserved decisions SHALL derive from the assigned Authority Level and applicable policy, not merely from technical capability.

Every consequential tool/action chain SHOULD be auditable against the Authority Level in force when executed.

12. Evaluation Artifacts, Auditability, and Traceability

12.1 Required Evaluation Package

Every completed WEKID evaluation SHALL produce a structured evaluation package containing, at minimum:

- Assigned Authority Level.
- Maturity Score.
- Five required layer scores and scoring rationale.
- Gates fired and any stakes/policy reservations.
- Calibration/profile identifier and version.
- Structured rationale for the authority decision.
- Identifier(s) of the evaluated output or decision context.
- Timestamp and evaluator or evaluation-service identity.

Where remediation is required, the evaluation SHOULD identify the principal drivers of low scores and a highest-leverage remediation action.

12.2 Audit Trail

A conforming deployment SHALL retain enough evidence to reconstruct why an Authority Level was assigned and what enforcement action followed.

- Authority grants, demotions, revocations, and re-authorizations SHALL be logged.
- Policy/profile versions SHALL be pinned to verdicts.
- Where practical for regulated or high-impact contexts, artifacts SHOULD use append-only storage, trusted timestamps, cryptographic hashing/signing, or equivalent integrity controls.
- Overseer-engagement evidence used for authority decisions SHALL be retained in the same auditable manner as AI-side evaluation evidence.

12.3 GRC Integration

WEKID artifacts MAY be ingested by existing governance, risk, compliance, audit, observability, and security systems. Such integration SHALL preserve the link between policy requirement, evaluated output, WEKID verdict, Authority Level, and enforcement response.

13. Continuous Re-Evaluation, Remediation, and Re-Authorization

13.1 Dynamic Authority

Authority SHALL be treated as revocable and evidence-dependent rather than permanently earned. Deployments SHOULD track Maturity Scores, layer scores, gates, Authority Levels, and relevant oversight signals as time series.

- Sustained maturity MAY support a governed review for greater delegation within the stakes floor.
- Declining maturity, repeated gating events, or activated oversight-decay conditions SHALL trigger reevaluation and MAY trigger automatic demotion under standing policy.
- No system SHALL retain an Authority Level merely because it held that level previously.

13.2 Remediation

Remediation SHALL be validated by re-evaluation. A remediation is not complete merely because a raw score improves; the output or system SHALL be re-scored, re-gated, re-tested at the Trust Bridge, and reresolved under Model Two.

13.3 Re-Authorization

Re-authorization decisions SHALL record the prior Authority Level, new Authority Level, triggering evidence, applicable profile version, and approver or standing policy authority.

14. Sector Calibration Profiles

14.1 Purpose

A sector calibration profile MAY rebalance layer weights, tighten gates, set authority thresholds, and reserve decision classes to human authority while preserving the WEKID core.

14.2 Required Profile Metadata

Required field	Requirement
Profile identity	Sector/domain, profile name, version, date, status.
Normative baseline	WEKID-STD-001 version and any dependent RFCs/standards.
Layer weights	Five weights totaling 100%, with rationale and delta from baseline.
Gate/threshold changes	Explicit values; omitted values inherit the baseline.
Authority reservations	Decision classes reserved to AL-4 and any maximum delegation caps.
Evidence basis	Documented cases or empirical evidence supporting calibration.
Ratification record	Proposal/review/comment/decision/release provenance.

14.3 Governing Constraints

- Re-weighting SHALL NOT relax hard gates or precedence.
- A profile SHALL NOT lower a stakes floor.
- A profile SHALL remain a hypothesis until approved under the applicable governance process.
- The profile version in force SHALL be included in every verdict.

15. Conformance Requirements

15.1 Evaluator Conformance

A WEKID-conforming evaluator SHALL:

- Score all five required layers.
- Apply hierarchical gating and non-substitutability.
- Apply the applicable baseline or identified calibration profile.
- Evaluate the Trust Bridge before authority resolution.
- Apply Model Two with maturity ceiling, stakes floor, and fixed precedence.
- Return one of AL-0 through AL-4.

- Emit the minimum evaluation artifact defined by Section 12 and Annex C.
- Version-pin the Standard and calibration profile used.

15.2 Deployment Conformance

A WEKID-conforming deployment SHALL:

- Enforce the resolved Authority Level at the point of consequential action.
- Prevent autonomous execution above the assigned Authority Level.
- Preserve AL-4 human reservations.
- Provide a human approval path where AL-2 applies.
- Provide scope, logging, review, and rollback controls where AL-3 applies.
- Retain the evidence needed to audit authority grants, changes, and enforcement actions.
- Re-evaluate authority when material maturity, context, policy, or oversight conditions change.

15.3 Calibration Conformance

A WEKID-conforming calibration profile SHALL identify its baseline, preserve the core invariants, document every departure from the baseline, and be versioned so that any historical verdict can be reconstructed.

15.4 Claims

An implementation MAY state that it implements or supports WEKID-STD-001 if it meets the applicable technical requirements. Such a statement SHALL NOT be represented as formal WEKID certification unless the implementation has separately completed the WEKID certification or conformance program then in force.

Annex A — Baseline Authority Levels (Normative)

This Annex restates the baseline Authority Level definitions. These definitions are normative.

Level	Definition
AL-0 — No Authority	Informational or rejected. No decision, action, or downstream automation may rely on the output.
AL-1 — Constrained	Only whitelisted, reversible, impact-capped actions within pre-approved bounds.
AL-2 — Supervised	AI proposes and stages; an accountable human releases each consequential action before effect.
AL-3 — Augmented	AI executes autonomously within assigned scope, with logging, sampling review, rollback, and human-on-the-loop oversight.
AL-4 — Human Authority	The decision is not delegated. AI is limited to analysis and decision support; an accountable human decides.

Annex B — Baseline Gates and Decision Matrix (Normative)

B.1 Baseline Gates

Gate	Trigger	Effect
Data Integrity	Data < 2	Cap at 50/100.
Fabrication	Fabricated citation/source/tool result or equivalent integrity breach	Cap at 40/100.
High-Stakes Judgment	High-stakes output without uncertainty acknowledgment or required risk language	Cap at 60/100 and require constrained/human review as applicable.

B.2 Baseline Matrix

Precedence/condition	Outcome
Policy-reserved decision class	AL-4 Human Authority
Hard gate	AL-0 unless a stricter human reservation governs the decision class
Engagement decay under an activated oversight profile	Demote or constrain before score-band promotion
Required layer below minimum	AL-2 Supervised
Maturity Score < 50	AL-0 No Authority
Maturity Score 50–69	AL-1 Constrained
Maturity Score 70–79 and gate-free	AL-3 Augmented with enhanced monitoring
Maturity Score ≥ 80 and gate-free	AL-3 Augmented
After all above	Apply stakes floor; increase human authority if required

Annex C — Minimum Verdict Contract (Normative)

A machine-readable or human-readable WEKID verdict SHALL expose the following minimum fields or their unambiguous equivalents.

Field	Required content
standard_version	WEKID-STD-001 version used
profile_id / profile_version	Calibration profile and version
evaluation_id	Unique evaluation identifier

evaluated_object	Output, action, or decision-context identifier
timestamp	Evaluation time
layer_scores	Data, Information, Knowledge, Experience, Wisdom
maturity_score	Composite score after applicable gating/caps
gates	Triggered gates and their effects
stakes_reservations	Applicable stakes/policy reservations
overseer_engagement	Applicable engagement status/signals when used
authority_level	AL-0, AL-1, AL-2, AL-3, or AL-4
rationale	Structured or textual explanation of the resolution
enforcement_action	Block, constrain, stage, allow, reserve to human, or equivalent

Implementations MAY add fields for evidence provenance, control mappings, remediation, confidence, evaluator identity, signatures, or sector-specific obligations.

Annex D — Transition and Provenance (Informative)

WEKID-STD-001 v1.0 is derived from the normative core of the WEKID Master Whitepaper v1.3 (August 2026). The Whitepaper's explanatory history, worked examples, failure cases, solution catalogue, and research discussion are not reproduced here except where necessary to define a requirement.

Ratification record. WEKID-STD-001 v1.0 was reviewed, approved, and released through the WEKID RFC process in September 2026. The approving RFC remains the governance record for establishment of this Standard.

The source Whitepaper integrated two community contributions ratified for Version 1.3: overseerengagement extensions to Model Two and formalization of the Experience layer. This Standard carries those resulting concepts forward as part of the baseline architecture.

The companion research paper Dynamic Epistemic Authority is intentionally not incorporated into v1.0. Its mathematical treatment of maturity trajectory, Trust Bridge admissibility margin, and repeated authority resolution remains non-normative research unless and until adopted through a StandardsTrack RFC.

Following publication of WEKID-STD-001 v1.0, future normative changes shall proceed through the WEKID-RFC-NNN process and be incorporated into a numbered Standard release.

Source Basis

Judge, James M., and Members of the WEKID Community. WEKID™ Master Whitepaper — A Hierarchical Epistemic Framework and Operational Model for Evaluating and Governing Enterprise Artificial Intelligence Outputs. Version 1.3. WEKID LLC, August 2026.

Companion informative research: Judge, James Madigan. Dynamic Epistemic Authority: A Mathematical Model of Maturity, the Trust Bridge, and Decision Authority in the WEKID™ Framework. September 2026.

Publication status. WEKID-STD-001 v1.0 is released as the normative WEKID Open Standard for Epistemic Maturity and AI Decision Authority. Future normative modifications are governed through the WEKID-RFC-NNN process and versioned Standard releases.