



E N K I D U

— M I O V S K I G R O U P —

C O N C E P T U A L A R C H I T E C T U R E

Aegis Assure

A standards-aligned framework for assessing and accrediting data governance and ethics across the lifecycle of machine-learning, neural-network, and AI systems — delivered as a MIOVSKI GROUP service, grounded in the Origi trust discipline.

English · Complete edition

Version 0.1 · First Draft

4 July 2026

Prepared by MIOVSKI GROUP

Classification · Confidential · Draft for Client Distribution

Glossary

Key terms used throughout this document. Framework terms shared with the Enkidu stack (Tangram, Origi, Helios Grid, Aegis Defence) are carried forward and not redefined except where this service extends them.

Aegis Assure

The MIOVSKI GROUP AI data-governance-and-ethics assessment and accreditation service: a structured, evidence-based method for evaluating how an organization governs the data that trains and operates its AI systems, and for accrediting that governance against recognized standards.

Training data

Data used to develop, train, validate, and test AI models — historical datasets, labelled data, synthetic data, and data used for tuning and evaluation.

Inference data

Data used during operational deployment, when the system makes predictions, recommendations, or decisions in production — real-time inputs and operational data.

Data lifecycle

The complete journey of data from acquisition through disposition, assessed here across ten stages.

Criterion

A single assessable control with two faces: the evidence a practitioner must produce, and the guidance an assessor applies to score it. The framework defines sixty-three unique criteria.

Practitioner / Assessor

The two-role separation at the heart of the method: practitioners collect and document evidence; assessors independently evaluate it and assign scores. The roles are never held by the same person for the same criterion — a separation of duties directly parallel to the Origi trust model.

RAG score

The Red / Amber / Green (plus Grey — not yet evaluated) compliance rating applied to each criterion, weighted by criticality to produce a system-level result.

Criticality

The weight and gating power of a criterion: Critical, High, Medium, or Low. A failed Critical criterion blocks accreditation regardless of overall score.

Assurance case

Borrowed from Origi: a claim, the argument that decomposes it, and the concrete evidence that supports it. Each accreditation result is, in effect, an assurance case for the AI system's data governance.

Model card / Dataset card

Structured documentation describing, respectively, a model's characteristics, intended use, limitations, and performance; and a dataset's characteristics, collection method, and known limitations.

Data drift / Concept drift

Change over time in, respectively, the statistical properties of input data; and the relationship between inputs and outputs — both of which degrade model reliability and trigger reassessment.

High-impact / High-risk AI system

An AI system posing significant risk to health, safety, or fundamental rights. The threshold for mandatory assessment under this service and the focus of the recognized regulatory regimes.

Abstract

Every organization deploying AI now faces the same demand from regulators, boards, auditors, and customers: not "does the model work?" but "can you prove the data behind it was governed responsibly?" That question is difficult to answer because the evidence is scattered across a data lifecycle that few organizations have ever mapped — acquisition, ingestion, integration, classification, storage, maintenance, operation, sharing, archival, and disposition — and because the standards that define a good answer have shifted sharply in the last eighteen months.

Aegis Assure is MIOVSKI GROUP's service for answering that question rigorously. It is a standards-aligned assessment and accreditation framework of sixty-three criteria spanning the full data lifecycle, the distinct obligations of training versus inference data, and the organizational governance that sits above both. Each criterion pairs a defined evidence requirement for practitioners with structured guidance for independent assessors, scored on a Red-Amber-Green scale and weighted by criticality, so that the output is not an opinion but a defensible, reproducible accreditation dossier.

This document is a reformatted, revalidated edition of the underlying framework. It corrects the regulatory foundation to the mid-2026 landscape — a landscape materially changed since the framework's first drafting: Canada's AIDA has died on the order paper, the EU AI Act has moved from text to a phased compliance calendar, and the United States has reversed its federal posture twice. And it positions the framework where it belongs within the Enkidu stack: as the governance-and-ethics assurance layer that complements Origi, sharing its assurance-case structure, its separation of duties, and its insistence that a claim is worth only the evidence chain behind it. Aegis Assure is how MIOVSKI GROUP extends the trust discipline it builds into its own products outward, as a service, to any organization that must prove its AI is governed well.

Trust in an AI system is not a slogan to be asserted; it is a position to be argued and evidenced. Aegis Assure is how that position is built, criterion by criterion.

1. Service Overview and Positioning

1.1 What Aegis Assure is

Aegis Assure is an assessment and accreditation service. An organization engages MIOVSKI GROUP to evaluate the data governance and ethics of one or more AI, machine-learning, or neural-network systems against a defined body of criteria, and to produce an accreditation result that stakeholders — regulators, boards, procurement authorities, insurers, customers — can rely on. The framework is the instrument; the service is the delivery of trustworthy findings through it.

It applies to any organization that designs, develops, deploys, or operates AI systems, with particular force for high-impact and high-risk applications — those affecting individuals' fundamental rights, safety, or well-being — where the recognized regulatory regimes concentrate their obligations.

1.2 Why it exists now

Three pressures converge. First, regulation has hardened: the EU AI Act's high-risk and general-purpose-AI obligations are now on a concrete compliance calendar rather than a distant horizon (Section 2), and organizations selling into or operating within those regimes need demonstrable governance. Second, the evidentiary bar has risen: "we take ethics seriously" is no longer an acceptable answer to a board or an auditor; what is required is documented, per-criterion, independently-assessed evidence. Third, the standards themselves have moved, and much existing governance documentation is now out of date — including, until this edition, the framework's own regulatory references.

1.3 Position in the Enkidu stack

Aegis Assure is the outward-facing member of the Enkidu family. The other frameworks govern how MIOVSKI GROUP builds AI systems: Tangram structures them, Origi makes their trustworthiness measurable and enforced, Helios Grid operates them at fleet scale, and Aegis Defence secures them cryptographically. Aegis Assure turns the same discipline outward — it assesses whether any organization's AI data governance meets the standard, using the same intellectual apparatus that Origi applies internally: the assurance case, the separation of assessment from action, and the rule that evidence, not assertion, is what a claim rests on.

- Framework — Question it answers — Direction
- Tangram — Where does work happen, and under what boundary conditions? — Internal — how we build
- Origi — May this action happen, and can we prove what happened? — Internal — how we build
- Helios Grid — Is the fabric healthy, current, and recoverable? — Internal — how we build
- Aegis Defence — Will every key and signature survive a quantum adversary? — Internal — how we build
- Aegis Assure — Is this organization's AI data governance responsible and defensible? — Outward — a service we deliver

The naming is deliberate: Aegis Defence protects the data in motion cryptographically; Aegis Assure protects the organization by proving its data was governed ethically. Both are shields — one technical, one evidentiary.

1.4 The relationship to Origi

Origis and Aegis Assure are the same philosophy at two scopes. Origi is a runtime trust harness that verifies each action of an AI system against policy as it runs; Aegis Assure is a lifecycle assurance framework that

verifies each stage of an AI system's data governance against criteria, periodically and at accreditation time. Both decompose a broad trust claim into checkable sub-claims tied to concrete evidence; both separate the party that gathers evidence from the party that judges it; both treat a passing result as valid only when the supporting evidence actually validates. An organization already assured under Aegis Assure is well-prepared to adopt Origi; an organization running Origi generates much of the evidence Aegis Assure requires as a by-product. They are designed to compose.

2. Regulatory Foundation

The framework is deliberately aligned to recognized regimes so that its accreditation results are externally defensible. That alignment must track a landscape that has changed materially since the framework's first drafting; this section states the position as of mid-2026 and flags what has moved. Specific articles, dates, and statutory status should be confirmed against official sources at each engagement, as this remains an actively shifting area.

2.1 European Union — the EU AI Act

The EU AI Act (Regulation (EU) 2024/1689) entered into force on 1 August 2024 and applies in phases rather than at once [1]. As of mid-2026 the operative milestones are: prohibited-practice bans and AI-literacy obligations since 2 February 2025; governance rules and general-purpose-AI (GPAI) model obligations since 2 August 2025; and the general application date — including most high-risk-system obligations and the transparency rules — at 2 August 2026, with high-risk AI embedded in regulated products extending to 2 August 2027 [1][2]. A November 2025 "Digital Omnibus" simplification proposal may further adjust the high-risk timeline by tying it to the availability of harmonized standards, which as of mid-2026 remain in development [1]. For this framework, Article 10 (data and data governance for high-risk systems) and Article 53 (GPAI training-data transparency and documentation) are the load-bearing provisions, and both are now within or approaching their application windows — which is precisely why demonstrable data governance has become urgent rather than aspirational.

2.2 Canada — a changed picture

The framework's earlier reference to AIDA (the Artificial Intelligence and Data Act, within Bill C-27) as a "proposed framework" is no longer current and has been corrected. Bill C-27 died on the Order Paper when Parliament was prorogued in January 2025, and a subsequent federal election removed any near-term path to re-tabling it in its original form; Canada currently has no dedicated federal AI statute [3]. In its absence, Canadian AI governance is shaped by the existing federal privacy regime (PIPEDA, itself long overdue for reform), by federal Treasury Board instruments for government AI use, and by provincial action — notably Québec's Law 25 and Ontario's initiatives — with many organizations aligning to GDPR and the EU AI Act as the de-facto baseline. Aegis Assure therefore treats Canadian obligations as PIPEDA-plus-provincial today, while retaining AIDA's concepts — risk assessment, mitigation, transparency, accountability for high-impact systems — as durable good practice likely to reappear in any future statute. Federal public-sector engagements additionally consider the Directive on Automated Decision-Making and its Algorithmic Impact Assessment.

2.3 United States — a reversed posture

The framework's reference to Executive Order 14110 required correction. EO 14110 (the Biden "Safe, Secure, and Trustworthy AI" order) was rescinded in January 2025 and replaced by Executive Order 14179, "Removing Barriers to American Leadership in Artificial Intelligence," which reorients federal policy from risk-and-rights toward innovation and competitiveness, and by "America's AI Action Plan" (July 2025) and its accompanying orders [4]. Notably for this framework, the Action Plan directed a revision of the NIST AI RMF to remove certain references; the underlying anti-discrimination law (e.g., Title VII disparate-impact liability) is unchanged by the shift in federal guidance. The practical consequence for Aegis Assure: the U.S. federal floor has lowered, but the operative obligations for high-risk applications increasingly come from state law, sectoral regulators, and — for any organization touching the EU market — the EU AI Act. The framework does not rely on the rescinded order.

2.4 The standards that endure

Beneath the shifting statutes sit standards that have proven stable and that Aegis Assure treats as its technical backbone:

NIST AI Risk Management Framework (AI RMF 1.0) and its Generative AI Profile (NIST AI 600-1) — the Govern / Map / Measure / Manage functions and trustworthiness characteristics, mapped throughout this framework [5][6].

ISO/IEC 42001:2023 (AI management systems) and ISO/IEC 8183:2024 (AI data lifecycle framework) — the management-system and lifecycle backbones [7][8].

ISO/IEC 23894:2023 (AI risk management guidance) and ISO/IEC TS 5723:2022 (trustworthiness vocabulary) — shared with Origi's own alignment [9].

GDPR (Regulation (EU) 2016/679), PIPEDA, Québec Law 25, and CCPA/CPRA — the privacy regimes whose obligations the data-lifecycle criteria operationalize [10][11][12].

The deliberate consequence of this layering is durability: because the framework is anchored to standards and to the concepts of AI regulation rather than to any single statute's current status, a change of government or a lapsed bill shifts emphasis without invalidating the criteria.

3. Framework Structure

The framework organizes sixty-three unique criteria into three parts, reflecting a single insight: data governance obligations differ by lifecycle stage, by whether the data trains or operates the model, and by whether the control is technical or organizational.

- Part — Scope — Criteria
- Section A — Data Lifecycle — Ten stages from acquisition to disposition, assessed for both training and inference data (with the applicable phase noted per criterion). — 41
- Section B — Data Operation — Split into Model-Training operations and Model-Inference operations, with the same five criterion codes applied differently to each phase. — 10
- Section C — Governance and Ethics — Organizational-level criteria applying across all systems and data uses — not specific to any one project. — 12

3.1 The training / inference distinction

The framework's central structural discipline is that training data and inference data carry different obligations even when the criterion code is the same. Training data raises questions of provenance, consent scope, bias in historical datasets, and copyright of source material; inference data raises questions of real-time validation, adversarial and out-of-distribution inputs, production drift, and operational sharing. Criteria are therefore marked Training, Inference, or Both, and the Data-Operation criteria (Section B) are deliberately duplicated so that the same concern — quality monitoring, performance monitoring, human oversight, bias monitoring, real-time validation — is assessed once for the training phase and once for production.

3.2 Criteria distribution

- Applicability — Count — Where
- Training data only — 5 — Data Operation — Training (OP-01...05)
- Inference data only — 10 — Data Operation — Inference (5) + Data Sharing (5)
- Both training and inference — 36 — Acquisition through Disposition, excluding Operation
- Cross-cutting organizational — 12 — Governance, Documentation, Ethics, Training

3.3 The anatomy of a criterion

Every criterion is expressed in the two-role form that gives the framework its rigour: an evidence requirement the practitioner must satisfy, and assessor guidance applied to judge it. This is the same evidence-then-judgment separation Origi enforces at runtime, applied here to lifecycle assessment.

- Face — Content
- Evidence Required (Practitioner) — The concrete artifacts that demonstrate compliance — documented legal bases, provenance records, bias analyses, PIAs, encryption specifications, audit logs, model and dataset cards, and so on, specific to the criterion.
- Assessor Guidance — The structured questions the independent assessor applies to the evidence — adequacy, rigour, completeness, alignment to the applicable regime — producing a defensible RAG score rather than an impression.

The full per-criterion evidence-and-guidance tables are maintained in the operational assessment workbook (the delivery instrument of the service) and are summarized by stage in Section 4. Reproducing all sixty-three in full belongs to that workbook rather than to this conceptual architecture.

4. Section A — The Data Lifecycle Criteria

Ten lifecycle stages carry forty-one criteria. Each stage below states its scope, its applicable phase, and the concerns its criteria assess; the criterion codes are retained from the framework for traceability to the assessment workbook.

4.1 Data Acquisition — obtaining data with rights and provenance

Applies to: Training and Inference. Five criteria (AC-01...05) assess the legal basis for collection (consent, contract, legitimate interest, legal obligation, with GDPR Article 6 / PIPEDA / CCPA evidence); source documentation and provenance (chain of custody, methodology, versioning); bias in collection (selection bias, demographic representation, historical bias, mitigation); privacy and confidentiality (PIA/DPIA, minimization, anonymization, sensitive-data handling, cross-border transfer); and copyright and IP compliance (clearance, licensing, fair-use analysis, attribution, third-party rights). Acquisition is where the majority of downstream defensibility is won or lost.

4.2 Data Ingestion — importing with quality and security

Applies to: Training and Inference. Four criteria (IN-01...04) assess quality validation (completeness, accuracy, consistency, anomaly detection, thresholds); format standardization (schema, conversion, type, encoding); security during ingestion (encryption in transit, authentication, input validation, logging, rate limiting); and labeling and annotation quality (guidelines, inter-annotator agreement metrics such as Cohen's/Fleiss' kappa, annotator qualification, disagreement resolution).

4.3 Data Integration — combining without amplifying bias

Applies to: Training and Inference. Four criteria (IT-01...04) assess lineage tracking (end-to-end, transformation, source-to-destination mapping, impact analysis); integration quality assurance (testing, cross-source consistency, duplicate resolution, conflict rules); bias-amplification assessment (pre/post-integration analysis, representation shift, statistical fairness testing); and master data management (entity resolution, golden records, stewardship).

4.4 Data Classification — knowing what you hold

Applies to: Training and Inference. Four criteria (CL-01...04) assess sensitivity classification (schemes, PII/PHI/biometric identification, GDPR Article 9 special categories, access alignment); regulatory classification (regime identification, vulnerable-population categorization, jurisdictional and industry mapping); training-versus-inference labeling (purpose-based classification, retention differentiation, usage restrictions); and data-quality classification (tiering, scoring, fitness-for-purpose).

4.5 Data Storage — securing at rest

Applies to: Training and Inference. Five criteria (ST-01...05) assess security controls (encryption at rest — AES-256 or equivalent — key management, RBAC/ABAC, network isolation, physical security); retention and lifecycle management; backup and recovery (RTO/RPO, testing, offsite, DR plan); audit logging and monitoring; and storage-infrastructure resilience (redundancy, failover, capacity, IaC). Where these systems handle data whose sensitivity is long-lived, the encryption controls should be read together with Aegis Defence's quantum-resistant guidance — data harvested today under a classical cipher is at risk tomorrow.

4.6 Data Maintenance — keeping data fit over time

Applies to: Training and Inference. Five criteria (MA-01...05) assess currency management (refresh, staleness detection, outdated-data handling); cleansing and correction (error detection, correction workflows, audit trails); rebalancing and augmentation (class-imbalance monitoring, SMOTE and related strategies, synthetic-data quality); schema-evolution management (change control, backward compatibility, migration); and continuous improvement (quality KPIs, feedback loops, prioritization).

4.7 Data Sharing and Dissemination — controlling what leaves

Applies to: Inference only. Five criteria (SH-01...05) assess sharing governance (policies, approval workflows, data-use agreements, purpose limitation, activity logging); privacy protection in sharing (anonymization, differential privacy, re-identification risk, privacy-preserving methods); third-party processor management (due diligence, DPAs, sub-processor control, audit rights); training-data publication and transparency (EU AI Act Article 53 documentation, dataset and model cards, licensing transparency); and API and interface security (authentication, rate limiting, input validation, security testing, versioning). This stage maps directly onto the boundary-brokering that Origi's Trust Harness performs at runtime — every crossing checked against policy.

4.8 Data Archival — preserving with integrity

Applies to: Training and Inference. Four criteria (AR-01...04) assess archival policy and procedures (triggers, format standards, retrieval SLAs, cost); archival security and privacy (encryption, access control, long-term privacy, integrity verification); archival documentation (metadata preservation, cataloging, discovery); and long-term accessibility (format migration, obsolescence planning, media-degradation monitoring).

4.9 Data Disposition and Retirement — ending responsibly

Applies to: Training and Inference. Five criteria (DR-01...05) assess secure deletion (cryptographic erasure, multi-location deletion across backups and replicas, verification, hardware disposal); right-to-erasure compliance (request handling, identity verification, timelines, exceptions); model decommissioning (procedures, retention decisions, cleanup); disposition documentation (decision records, certificates of destruction, chain of custody); and legacy-system migration or disposal.

5. Section B — The Data Operation Criteria

Five criteria (OP-01...05) are assessed twice — once for the training phase, once for inference — because the same concern behaves differently in each.

- Criterion — In model training (Training data only) — In model inference (Inference data only)
- OP-01 Data-quality monitoring — Quality and distribution-drift monitoring across training runs, with root-cause analysis. — Real-time production input-quality monitoring, out-of-specification handling, trend reporting.
- OP-02 Performance monitoring — Train/validation/test tracking, overfitting/underfitting detection, experiment comparison. — Production accuracy, data- and concept-drift detection, A/B and challenger models, SLA tracking.
- OP-03 Human oversight — Review checkpoints in the training workflow, expert review, approval and escalation. — Human-in-the-loop decision points, override mechanisms, high-risk escalation, review sampling.
- OP-04 Bias and fairness monitoring — Per-iteration bias assessment, fairness metrics (demographic parity, equalized odds), trend analysis. — Continuous production fairness monitoring, disparate-impact analysis, incident response.
- OP-05 Real-time data validation — Pre-batch input validation, schema and constraint checking, corrupted-data handling. — Per-request validation, adversarial-input and out-of-distribution detection, sanitization.

The inference-side OP criteria are where Aegis Assure and Origi meet most directly: OP-03 (human oversight), OP-04 (fairness), and OP-05 (adversarial and OOD detection) describe, in assessment terms, exactly the checks the Origi Trust Harness performs as enforcement at inference time. An organization running Origi produces the evidence these criteria require as operational output.

6. Section C — Governance and Ethics Criteria

Twelve organizational criteria apply across every system and data use — they assess institutional maturity, not a single project.

6.1 Organizational Governance

Three criteria (GV-01...03): an AI governance structure (framework, roles such as an AI ethics board and data-governance committee, board-level oversight, escalation); a risk-management framework aligned to the NIST AI RMF (methodology, AI-specific risk register, risk appetite, integration with enterprise risk); and compliance management (regulatory-landscape monitoring, obligation mapping, testing, change management).

6.2 Documentation and Transparency

Three criteria (DC-01...03): technical documentation (standards, model cards for every deployed model, dataset cards for training data, architecture and version control); explainability and interpretability (XAI techniques, user-facing explanations, counterfactuals, validation); and algorithmic impact assessments (AIA methodology, fundamental-rights analysis, stakeholder consultation, periodic reassessment).

6.3 Ethical Principles and Accountability

Four criteria (ET-01...04): fairness and non-discrimination (policy, protected-characteristic identification, metrics and thresholds, remediation, auditing); transparency and accountability (disclosures, individual-rights procedures, contestation and appeal, accountability assignment); human-rights impact (assessment methodology, vulnerable-population protection, mitigation, stakeholder engagement); and environmental sustainability (carbon-footprint measurement, energy efficiency, sustainable infrastructure, green-AI practice).

6.4 Training and Awareness

Two criteria (TR-01...02): staff training and competency (ethics and governance programs, role-specific requirements, effectiveness assessment); and organizational culture (ethical-AI values, speak-up and whistleblower mechanisms, leadership modeling, cross-functional collaboration).

7. Assessment Methodology

7.1 The two-role system

The method's integrity rests on a separation of duties: practitioners collect and document evidence for each applicable criterion; assessors independently evaluate that evidence and assign scores. No individual both produces and judges the evidence for the same criterion. This is the Origi separation-of-duties principle applied to assessment — the party that gathers cannot be the party that grants, so that the most dangerous failure mode (an organization marking its own homework) is removed by construction.

7.2 RAG scoring

Each criterion receives one of four ratings, which drive the accreditation result:

- Rating — Meaning
- Grey — Not yet evaluated — The default state; the criterion has not been assessed.
- Green — Compliant — Controls are effective and evidence supports the claim; no issues.
- Amber — Partial / minor gaps — Controls exist but carry minor design or operating deficiencies.
- Red — Non-compliant — Controls are missing, ineffective, or failing significantly.

7.3 Criticality weighting and gating

Criteria are weighted so that the result reflects what matters most, and so that certain failures gate accreditation outright:

- Criticality — Weight — Gating rule
- Critical (C) — 40% — A sub-threshold Critical criterion blocks accreditation until remediated.
- High (H) — 30% — Requires a remediation plan before accreditation.
- Medium (M) — 20% — Triggers an improvement roadmap.
- Low (L) — 10% — Continuous-improvement focus.

The gating rule is the point: a strong average cannot buy past a failed Critical control. This mirrors Origi's fail-closed posture — a single decisive failure blocks, rather than being averaged away.

7.4 The accreditation dossier

The output of an engagement is not a score but a dossier: per criterion, its applicability, its evidence, its assessor rationale, its RAG rating, and its remediation status — assembled into a system-level accreditation with a defensible audit trail. In Origi's language, this dossier is the assurance case for the system's data governance: a claim ("this system's data is governed responsibly") decomposed into sub-claims (the criteria) and tied to concrete, independently-assessed evidence. It is designed to be shown to a regulator, an auditor, a board, or a customer.

8. Implementation and Delivery

8.1 Practitioner workflow

Scoping (identify the system and its applicable training/inference criteria) -> planning (evidence-collection plan and timeline) -> evidence gathering (documentation, stakeholder interviews, system testing) -> documentation (complete the evidence templates) -> gap identification -> quality self-review -> submission to assessors.

8.2 Assessor workflow

Evidence review (completeness and quality) -> verification (sampling, testing, interviews) -> scoring (RAG against the guidance) -> documentation (findings, rationale, recommendations) -> calibration (peer review for consistency) -> reporting (the accreditation dossier) -> follow-up (track remediation).

8.3 Reassessment triggers

Accreditation is a position that must be maintained, not a certificate that is filed. Reassessment is triggered by: annual review for all deployed high-risk systems; material system change (new data sources, algorithm changes, use-case expansion); regulatory change affecting the system; significant incidents or bias discoveries; and transition from pilot to production. Each trigger is analogous to an Origi reassessment event — trust is earned continuously, not granted once.

8.4 How MIOVSKI GROUP delivers it

Aegis Assure is offered in three engagement shapes: a readiness assessment (a rapid, RAG-only gap scan to locate the biggest exposures before a formal effort); a full accreditation (the complete sixty-three-criterion evidence-and-assessment cycle producing the dossier); and continuous assurance (recurring reassessment against the trigger schedule, optionally integrated with an Origi deployment so that runtime evidence feeds the lifecycle assessment). The service draws on MIOVSKI GROUP's combined lineage in enterprise architecture, cybersecurity, and applied AI — and on the same trust discipline the firm builds into its own products.

9. Conclusion

The market has moved from asking whether an AI system works to demanding proof that the data behind it was governed responsibly — and the standards defining a good answer have shifted sharply enough that much existing governance documentation, including this framework's own earlier regulatory references, needed correction. This edition makes those corrections: AIDA is recorded as dead rather than pending, the EU AI Act as a live phased calendar rather than a distant text, and the U.S. posture as reversed rather than as EO 14110.

What does not change is the framework's structure and its value: sixty-three criteria across the full data lifecycle, the training-versus-inference discipline, organizational governance above both, and a two-role, evidence-then-judgment method that produces a defensible accreditation dossier rather than an opinion. Positioned within the Enkidu stack, Aegis Assure is the outward face of the same trust philosophy that Origi enforces inside MIOVSKI GROUP's own products — the assurance case, the separation of duties, the primacy of evidence — offered as a service to any organization that must now prove, and not merely assert, that its AI is governed well.

Trust in an AI system is not a slogan to be asserted; it is a position to be argued and evidenced. Aegis Assure is how that position is built, criterion by criterion.

References

Regulatory status in this area changes frequently; specific clauses, dates, and legislative status should be confirmed against the official publications before formal use.

[1] European Union. Regulation (EU) 2024/1689 (Artificial Intelligence Act); European Commission, AI Act — application timeline and Navigating the AI Act (phased dates: prohibitions 2 Feb 2025; GPAI and governance 2 Aug 2025; general application and high-risk 2 Aug 2026; embedded high-risk 2 Aug 2027; Nov 2025 Digital Omnibus simplification proposal). digital-strategy.ec.europa.eu, 2024–2026.

[2] EU AI Act Article 10 (data and data governance, high-risk systems) and Article 53 (GPAI transparency and training-data documentation). artificialintelligenceact.eu, 2024.

[3] Government of Canada. Bill C-27, Digital Charter Implementation Act, 2022, including the proposed Artificial Intelligence and Data Act (AIDA) — died on the Order Paper upon prorogation of Parliament, January 2025; not enacted. See Schwartz Reisman Institute, "What's Next After AIDA?" (2025–2026) and IAPP / DLA Piper analyses (Jan 2025). Canada currently has no dedicated federal AI statute; PIPEDA, the TBS Directive on Automated Decision-Making, and provincial laws (Québec Law 25; Ontario) apply.

[4] United States. Executive Order 14110 (Oct 2023, Safe, Secure, and Trustworthy AI) — rescinded January 2025; replaced by Executive Order 14179, "Removing Barriers to American Leadership in Artificial Intelligence" (23 Jan 2025) and "America's AI Action Plan" (Jul 2025), which directed a revision of the NIST AI RMF. [federalregister.gov](https://www.federalregister.gov); [whitehouse.gov](https://www.whitehouse.gov); NIST, 2025.

[5] NIST. AI Risk Management Framework (AI RMF 1.0), NIST AI 100-1, 2023. (Govern / Map / Measure / Manage; trustworthiness characteristics.)

[6] NIST. Artificial Intelligence Risk Management Framework: Generative AI Profile, NIST AI 600-1, 2024.

[7] ISO/IEC. ISO/IEC 42001:2023 — Information technology — Artificial intelligence — Management system.

[8] ISO/IEC. ISO/IEC 8183:2024 — Information technology — Artificial intelligence — Data lifecycle framework.

[9] ISO/IEC. ISO/IEC 23894:2023 (AI risk management guidance) and ISO/IEC TS 5723:2022 (trustworthiness vocabulary).

[10] European Commission. General Data Protection Regulation (GDPR) — Regulation (EU) 2016/679.

[11] Office of the Privacy Commissioner of Canada. Personal Information Protection and Electronic Documents Act (PIPEDA); Government of Québec, Law 25 (An Act to modernize legislative provisions as regards the protection of personal information).

[12] State of California. California Consumer Privacy Act (CCPA) and California Privacy Rights Act (CPRA), 2018 / 2020.

[13] Enkidu Enterprises / MIOVSKI GROUP (2026). Trust Assurance Overview (Origi), v1.0; Tangram Overview, v1.0; Helios Grid and Aegis Defence conceptual architectures, v0.1. (Assurance-case structure, separation of duties, chain of evidence — the internal frameworks Aegis Assure mirrors outward.)

Document version history — v1.0: initial comprehensive framework. v2.0: updated to the official criteria mapping with explicit Training / Inference / Organizational distinctions. v3.0 (this edition): reformatted to the Enkidu conceptual-architecture house style; regulatory foundation revalidated to the mid-2026 landscape (AIDA lapsed, EU AI Act phased dates, EO 14110 replaced); positioned within the Enkidu stack as the outward-facing Aegis Assure service, aligned to Origi.