



ENGINEERING STANDARDS LIBRARY / EXECUTIVE PORTFOLIO

AI, Knowledge, Data, Identity, and Learning Executive Portfolio

Executive synthesis of the permanent standards, operating model, assurance posture, strategic risks, and required adoption decisions.



Contents

Contents	2
Portfolio position	3
Portfolio metrics	3
Portfolio register	4
Integrated assurance operating model	6
Strategic operating principles	6
Strategic risk register	7
Leadership decisions	7
Adoption roadmap	8
Executive assurance scorecard	8
Assurance status	8

Portfolio position

This permanent library contains 25 candidate standards and 475 atomic evaluation criteria across artificial intelligence, knowledge, data, identity, and learning. It establishes a connected assurance system rather than a collection of model feature checklists.

Portfolio metrics

DOMAIN	STANDARDS	CRITERIA	PORTFOLIO ROLE
Artificial Intelligence and Agentic Systems	15	354	Permanent candidate standards with criterion-level evidence and assessment guidance.
Knowledge and Grounding	2	24	Permanent candidate standards with criterion-level evidence and assessment guidance.
Data, State, and Evidence	5	61	Permanent candidate standards with criterion-level evidence and assessment guidance.
Identity and Access	2	24	Permanent candidate standards with criterion-level evidence and assessment guidance.
Learning and Workforce	1	12	Permanent candidate standards with criterion-level evidence and assessment guidance.

Portfolio register

ID	PUBLICATION	CRITERIA
MYTHOS-AI-0001	Agentic Framework Benchmark and Evaluation Standard Defines a governed benchmark system for comparing agent libraries, orchestration runtimes, multi-agent frameworks, persistent agent systems, and managed agent platforms using security, durability, evidence, and operational criteria.	36
MYTHOS-AI-0002	Model Intelligence, Training, and Deployment Standard Establishes requirements for model intelligence evaluation, training provenance, deployment architecture, model lifecycle governance, licensing, safety, optimization, and operational readiness.	95
MYTHOS-AI-0003	Applied Natural Language & LLM Evaluation Standard Defines applied evaluation methods for natural-language and large-language-model systems, including factuality, grounding, instruction following, retrieval, safety, multilingual behavior, and production quality.	32
MYTHOS-AI-0004	AI Software Engineer & Code Generation Benchmark Standard Establishes a benchmark for AI-assisted software engineering and code generation across planning, implementation, debugging, review, testing, security, repository-scale work, and evidence quality.	24
MYTHOS-AI-0005	Applied Automation & Infrastructure-as-Code Benchmark Standard Defines evaluation criteria for AI-assisted infrastructure automation, policy-safe change generation, validation, rollback, idempotency, drift handling, and operational evidence.	16
MYTHOS-AI-0006	AI Alignment, Red-Team, & Sycophancy Evaluation Standard Establishes alignment, adversarial evaluation, red-team, deception, manipulation, sycophancy, refusal, and behavioral assurance requirements for deployed AI systems.	20
MYTHOS-AI-0007	Context Engineering, Trajectory Compaction, & Temporal Memory Standard Defines context engineering, trajectory compaction, working memory, persistent memory, retrieval, provenance, isolation, and lifecycle controls for durable agent systems.	18
MYTHOS-AI-0008	Multimodal Interaction, Vision, & Voice Latency Standard Establishes multimodal vision, audio, speech, realtime interaction, latency, accessibility, provenance, and cross-modal safety requirements.	17
MYTHOS-AI-0009	Model Fine-Tuning, RLEF, & Synthetic Data Governance Standard Defines governance for model adaptation, fine-tuning, preference learning, reinforcement learning from execution feedback, synthetic data, evaluation, rollback, and release evidence.	14
MYTHOS-AI-0010	AI Avatar, Persona, & Provenance Standard Establishes requirements for AI presence, persona, voice, visual embodiment, consent, provenance, continuity, disclosure, impersonation resistance, and user control.	21
MYTHOS-AI-0011	Mixture of Experts (MoE) & State Space Model (SSM/Mamba) Architecture Standard Defines architectural and evaluation requirements for mixture-of-experts and state-space models, including routing, specialization, state handling, efficiency, scaling, and observability.	16
MYTHOS-AI-0012	Neuro-Symbolic AI & Continual Learning (Anti-Catastrophic Forgetting) Standard Establishes requirements for neuro-symbolic integration and continual learning while controlling catastrophic forgetting, unsafe adaptation, provenance loss, and unbounded behavioral change.	11
MYTHOS-AI-0013	AI-Assisted Software Engineering & Model Routing Standard Defines model-routing and AI-assisted engineering controls for task placement, specialist selection, local-versus-hosted execution, verification, and cost-aware quality management.	11

ID	PUBLICATION	CRITERIA
MYTHOS-AI-0014	Inference Serving, Quantization & KV-Cache Standard Establishes requirements for inference serving, batching, quantization, scheduling, key-value cache management, latency, throughput, isolation, and production reliability.	11
MYTHOS-AI-0015	Federated Learning & Privacy-Preserving ML Standard Defines privacy-preserving and federated machine-learning requirements for distributed training, aggregation, privacy budgets, poisoning resistance, participation governance, and evidence.	12
MYTHOS-KNW-0001	RAG, Grounding & Source Authority Standard Defines retrieval-augmented generation, grounding, citation, source authority, freshness, conflict handling, provenance, and evidence requirements for knowledge-backed AI systems.	12
MYTHOS-KNW-0002	Persona, Avatar & Persistent Memory Architecture Standard Establishes architecture and governance for persistent persona, embodied presence, user-controlled memory, continuity, provenance, privacy, and lifecycle management.	12
MYTHOS-DAT-0001	Vector Database & Local-First Sync Architecture Standard Defines vector retrieval and local-first synchronization architecture, including indexing, embeddings, consistency, conflict resolution, offline operation, provenance, and lifecycle control.	12
MYTHOS-DAT-0002	AI & Agent Observability Semantic Conventions (OpenTelemetry) Standard Establishes semantic conventions and operational requirements for observing models, agents, tools, retrieval, memory, costs, safety events, and end-to-end AI execution traces.	11
MYTHOS-DAT-0003	Data Sovereignty, Privacy Preservation, & Egress Control Standard Defines data sovereignty, privacy preservation, residency, classification, egress control, minimization, policy enforcement, and accountable cross-boundary processing.	14
MYTHOS-DAT-0004	Event Sourcing, CQRS, & Durable State Machine Standard Establishes event-sourcing, CQRS, durable state-machine, replay, idempotency, consistency, migration, and evidence requirements for authoritative state systems.	12
MYTHOS-DAT-0005	Tamper-Evident Hash-Chained Ledgers & Evidence Bundle Standard Defines tamper-evident ledgers and evidence bundles for consequential system actions, including hash chaining, signatures, provenance, retention, verification, and disclosure boundaries.	12
MYTHOS-ID-0001	Workload, Agent & Human Identity Standard Defines identity architecture for humans, workloads, services, agents, tools, and devices with attestation, federation, lifecycle, delegation, authentication, and audit requirements.	12
MYTHOS-ID-0002	Public Key Infrastructure (PKI) & Attribute-Based Access Control (ABAC) Standard Establishes public-key infrastructure and attribute-based access-control requirements for trust anchors, certificate lifecycle, policy decision, authorization context, revocation, and evidence.	12
MYTHOS-EDU-0001	Adaptive Learning, Skill Graphs & Cyber Lab Standard Defines adaptive learning, skill graphs, competency evidence, cyber labs, assessment, accessibility, credential portability, safety, and workforce-assurance requirements.	12

Integrated assurance operating model

ASSURANCE PLANE	EXECUTIVE QUESTION	OUTCOME
Purpose and impact	What outcome is authorized, for whom, with what consequence and prohibited behavior?	Bounded, reviewable mission and accountable owner.
Knowledge and data	What sources and state justify the decision, and may they be used here?	Grounded context with authority, freshness, privacy, and sovereignty.
Identity and authority	Who or what is acting, and what may it do now?	Authenticated principal and narrow, expiring capability.
Model and reasoning	Which model or method fits the workload and evidence confidence?	Replaceable execution resource selected by controlled routing.
Execution and state	What tool or workflow changes which authoritative system?	Constrained, durable, recoverable side effects.
Verification and evidence	How is correctness established and the outcome reconstructed?	Independent validation and tamper-evident decision record.
Human accountability	Who can approve, interrupt, correct, accept residual risk, and retire the capability?	Meaningful control rather than ceremonial oversight.

Strategic operating principles

- Build the governed harness and evidence system rather than coupling business authority to a model provider.
- Treat source authority, data state, identity, policy, execution, verification, and human accountability as one system.
- Use workload profiles and mandatory gates instead of universal model or framework rankings.
- Preserve local-first and sovereign deployment paths for sensitive contexts.
- Require explicit freshness, limitations, and revalidation for research and rapidly changing specifications.
- Fund identity, evidence, failure recovery, and provider exit as shared infrastructure rather than project-specific add-ons.

Strategic risk register

STRATEGIC RISK	CONSEQUENCE	LEADERSHIP RESPONSE
Model authority conflation	Probabilistic output silently becomes business authority.	Externalize policy, approval, deterministic execution, and evidence.
Uncontrolled memory	Poisoned, stale, sensitive, or cross-tenant context persists.	Fund provenance, scope, correction, deletion, and memory incident response.
Provider lock-in	Critical state, prompts, routing, evaluations, or workflows cannot migrate.	Mandate open formats, abstraction boundaries, export, and tested replacement.
Identity fragmentation	Human, service, agent, device, and tool actions cannot be correlated or revoked.	Establish a unified principal and delegation architecture.
Evidence deficit	Outcomes cannot be reproduced, defended, audited, or improved.	Treat evidence as a production dependency with availability and integrity objectives.
Evaluation theater	Benchmark scores substitute for workload-specific safety and operations testing.	Require mandatory gates, profile-specific evaluation, and evidence grades.
Sovereignty erosion	Protected data, embeddings, traces, or memory cross an unintended boundary.	Approve deployment profiles, egress policy, locality, and provider contracts.
Unsafe adaptation	Fine-tuning or continual learning changes behavior without adequate lineage and rollback.	Control data, release, evaluation, approval, and reversion.
Human disengagement	Approval becomes routine while understanding and intervention decline.	Design consequence-aware review, sampling, escalation, training, and operator metrics.
Research overreach	Emerging methods are treated as mature because they are novel.	Separate research, pilot, conditional adoption, and production assurance.

Leadership decisions

- Approve permanent ownership and review cadence for each standard and companion publication.
- Fund reproducible evaluation labs for high-weight model, agent, retrieval, identity, state, and learning criteria.
- Establish evidence retention, exception, incident, and review workflows.
- Define approved deployment profiles, consequence classes, and prohibited autonomy boundaries.
- Approve model, provider, vector-store, and workflow-runtime exit strategies before consequential adoption.
- Require human-system interaction and workforce assurance to scale alongside automation capability.

Adoption roadmap

HORIZON	PROGRAM OUTCOME	LEADERSHIP EVIDENCE
0-90 days	Ownership, consequence classes, prohibited actions, authority registry, and evaluation method.	Approved charter, named owners, and prioritized use-case inventory.
3-6 months	Identity, policy, secret brokerage, controlled tools, evidence schema, and baseline labs.	Gate demonstrations and accepted baseline results.
6-12 months	Durable runtime, knowledge and memory controls, observability, incident operations, and production profiles.	Operational readiness and controlled production evidence.
12-18 months	Provider portability, local-first infrastructure, advanced evaluation, privacy-preserving learning, and workforce assurance.	Migration exercises, assurance scorecards, and measurable capability improvement.
Continuous	Source freshness, model and provider changes, threat intelligence, incidents, exceptions, and retirement.	Review cadence, expiry register, and evidence-backed adoption decisions.

Executive assurance scorecard

MEASURE	PURPOSE
Mandatory-gate pass rate	Shows whether candidate deployments satisfy non-compensable authority, safety, durability, and evidence controls.
Evidence-grade distribution	Separates documented claims from reproduced and independently reviewed behavior.
Grounding and source freshness	Measures support, authority, contradiction, expiry, and citation integrity.
Identity and revocation coverage	Measures principal correlation, scoped delegation, credential lifetime, and revocation propagation.
Mission recovery success	Measures checkpoint, replay, cancellation, idempotency, and external-state reconciliation.
Verification disagreement rate	Surfaces evaluator blind spots, uncertainty, and human arbitration demand.
Evidence completeness	Measures reconstructability of consequential decisions and actions.
Provider portability	Measures tested migration of models, state, evaluations, tools, and evidence.
Privacy and sovereignty exceptions	Measures boundary crossings, retention, deletion, and approved residual risk.
Human intervention quality	Measures meaningful review, correction, interruption, escalation, and operator workload.

Assurance status

All publications remain candidates. Documentation review raises confidence but does not prove runtime behavior, interoperability, safety, privacy, accessibility, legal applicability, or compliance. Production adoption requires accepted evidence in the intended environment and continued revalidation.