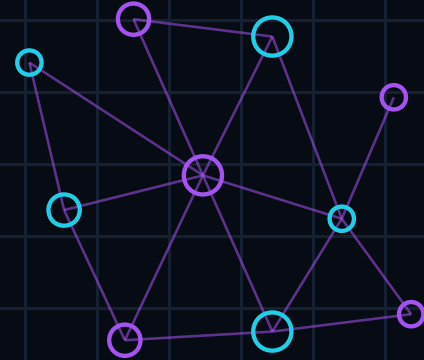




MYTHOS SYSTEMS

ENGINEERING STANDARDS LIBRARY /
FOUNDATION

MYTHOS-FND-0009 / FOUNDATION AND GOVERNANCE



CANDIDATE STANDARD

Mythos AI-First Application Standard

Govern material AI behavior through justified use, impact analysis, evaluation, grounding, human control, and lifecycle evidence.

PERMANENT ID

MYTHOS-FND-0009

PUBLICATION CLASS

CANDIDATE STANDARD

VERSION / STATUS

1.0.0-candidate / CANDIDATE

PERMANENT PUBLICATION IDENTITY

Mythos AI-First Application Standard

Universal Requirements for Responsible, Secure, Evaluated, and Operable AI-Enabled Applications

Document ID
MYTHOS-FND-0009

Version
1.0.0-candidate

Status
Candidate Standard

Review date
2027-09-17

Publication position

This is a permanent candidate standard in the Mythos Engineering Standards Library. Candidate status means the content is ready for structured implementation and review, but it is not represented as external certification, regulatory approval, or independent assurance.

PROFILE 3 LEVELS Foundational / Advanced / High Assurance	CONTROLS 19 Atomic normative controls	CADENCE TRIGGERED Annual plus material-change review	EVIDENCE ASSESSABLE Examine / Interview / Test / Measure
--	--	---	---

Reader orientation

Dimension	Engineering position
Primary domain	Foundation and Governance
Audience	AI product and engineering teams; Security, privacy, legal, risk, and safety teams; Data and model owners; Operators and incident responders; Assessors and procurement teams
Publisher / owner	Mythos Systems / Standards Program
Public classification	Public technical standard
Canonical route	/library/standards/foundation/mythos-fnd-0009/

Expected outcomes

- Use AI only where it provides defensible value relative to simpler alternatives.
- Make model uncertainty, limitations, and authority boundaries explicit.
- Require evaluation against representative tasks, harms, abuse, and operational conditions.
- Preserve accountability and safe fallback across model, provider, and data changes.

SOURCE / FRESHNESS POSITION

Authority and source posture

Validated 2026-09-17

Primary authority versions were revalidated for this regeneration on 2026-09-17. Current-source updates are reflected where a newer stable version was identified; active drafts are labeled as such and do not silently replace current final authorities.

FOUNDATION OPERATING DOCTRINE

OPERATING FLOW / MYTHOS-FND-0009



The source lineage for this edition is the enterprise candidate source set produced in July 2026. Normative controls are preserved; editorial, visual, metadata, and authority-freshness changes are recorded as revision lineage rather than presented as new control substance.

NAVIGATION

Contents

The table of contents is page-aware and internally linked. Control-level navigation follows on the next pages.

Document Control	8
Revision Record	8
How to Use This Standard	9
Executive Reading Path	9
Engineering Reading Path	9
Assessment Reading Path	9
Normative and Informative Content	9
Part I – Executive Direction	10
1. Executive Overview	10
2. Executive Findings and Required Direction	10
3. Business and Operational Impact	10
4. Target-State Summary	11
5. Leadership Responsibilities	12
Part II – Standard Foundation	13
6. Purpose and Objectives	13
7. Scope	13
8. Intended Audience	13
9. Requirements Language and Conformance	13
10. Normative References from Source Draft	13
11. Informative References from Source Draft	14
12. Terms and Definitions	14
13. Applicability and Tailoring	15
Part III – Risk and Architecture	16
14. Enterprise Problem and Risk Landscape	16
15. Governing Principles	16

16. Reference Operating Architecture	16
17. Roles and Accountability	17
18. State, Evidence, and Decision Model	18
19. Security and Trust Considerations	18
20. Failure Modes	18
21. Cross-Functional Dependencies	18
Part IV – Normative Control System	19
22. Control-Family Overview	19
23. Normative Controls	20
Part V – Implementation and Operations	40
24. Implementation Profiles	40
25. Implementation Lifecycle	40
26. Integration Requirements	41
27. Failure Handling and Recovery	41
28. Exceptions and Compensating Controls	41
29. Prohibited Practices	41
30. Adoption Roadmap from Source Draft	42
31. Limitations and Residual Risk from Source Draft	42
Part VI – Assurance and Assessment	43
32. Assessment Methodology	43
33. Metrics and Reporting from Source Draft	43
34. Exceptions and Compensating Controls	44
35. Enterprise Metric Set	44
36. Maturity Model	45
37. Assessment Confidence	45
38. Executive Assurance Dashboard	45
Part VII – Authority and Traceability	47
39. Cross-Standard Dependencies from Source Draft	47
40. Current External Authority Register	47

41. Control Traceability	47
42. Source-Draft Preservation	48
Part VIII – Appendices	49
Appendix A — Source-Draft Evolution Notes	49
Appendix B — Change History	49
Appendix C — Control Summary	49
Appendix D — Minimum Conformance Claim Record	50
Appendix E — Publication Review Checklist	50

NORMATIVE SYSTEM

Control index

Every control is independently addressable and assessable. Page references link directly to the control record.

MYTHOS-FND-AIA-01	AI use-case justification and baseline	21
MYTHOS-FND-AIA-02	AI impact and risk assessment	22
MYTHOS-FND-AIA-03	Accountability and decision rights	23
MYTHOS-FND-AIA-04	Training, retrieval, and inference data governance	24
MYTHOS-FND-AIA-05	Model and provider inventory	25
MYTHOS-FND-AIA-06	Model selection by measured fitness	26
MYTHOS-FND-AIA-07	Representative predeployment evaluation	27
MYTHOS-FND-AIA-08	Grounding and source integrity	28
MYTHOS-FND-AIA-09	Prompt and instruction governance	29
MYTHOS-FND-AIA-10	Action and decision boundaries	30
MYTHOS-FND-AIA-11	Meaningful human oversight and recourse	31
MYTHOS-FND-AIA-12	User disclosure and output provenance	32
MYTHOS-FND-AIA-13	AI threat model and abuse resistance	33
MYTHOS-FND-AIA-14	Privacy-preserving AI operation	34
MYTHOS-FND-AIA-15	Fallback, abstention, and degraded mode	35
MYTHOS-FND-AIA-16	Model, prompt, and provider change control	36
MYTHOS-FND-AIA-17	Production performance and harm monitoring	37
MYTHOS-FND-AIA-18	AI incident response and shutdown	38
MYTHOS-FND-AIA-19	Retirement, replacement, and record retention	39

Document Control

Field	Value
Document ID	MYTHOS-FND-0009
Standard	AI-First Application
Version	1.0.0-candidate
Status	Candidate Standard
Classification	Public
Publisher	Mythos Systems
Program	Mythos Engineering Standards Program
Accountable Owner	Mythos Systems Standards Program
Technical Review	Required
Source and Citation Review	Required
Assessment Review	Required
Accessibility and Publication Review	Required
Supersedes	No published edition; evolves source draft 0.2.0
Next Review	2027-09-17 or earlier on trigger

Revision Record

Version	Date	Status	Material Change
1.0.0-candidate	2026-09-17	Candidate Standard	Permanent-publication regeneration: source-preserving editorial system, richer information design, current authority revalidation, stable public metadata, and release QA.
0.2.0	2026-07-18	Working Draft	Source control standard
1.0.0-candidate	2026-07-22	Candidate Standard	Enterprise report architecture, executive direction, target operating model, profiles, maturity, authority register, and source-preserving reconstruction

How to Use This Standard

Executive Reading Path

Read Part I, the target-state architecture in Part III, the profile table in Part V, and the assurance dashboard in Part VI.

Engineering Reading Path

Read Parts II through V, then use the normative controls in Part IV as implementation and design requirements.

Assessment Reading Path

Read the conformance requirements in Part II, every applicable control's evidence and assessment procedure in Part IV, and the assessment, maturity, confidence, and traceability provisions in Parts VI and VII.

Normative and Informative Content

Uppercase **MUST**, **MUST NOT**, **SHOULD**, **SHOULD NOT**, and **MAY** statements are normative when used under the requirements-language provisions of this document. Executive findings, examples, implementation patterns, reference architectures, mappings, and external-authority descriptions are informative unless explicitly incorporated into an adopting organization's binding policy, contract, or regulatory obligation.

Part I – Executive Direction

1. Executive Overview

AI should be introduced because it improves a defined outcome under acceptable risk, not because it is available. AI-first applications require governance beyond conventional software because behavior may be probabilistic, provider-dependent, data-sensitive, difficult to explain, and capable of influencing or performing consequential actions. This standard establishes a full lifecycle for justifying, evaluating, constraining, monitoring, and retiring AI-enabled application behavior.

The institutional objective is to govern when and how AI is used in applications, including impact, data, model selection, evaluation, grounding, prompts, action boundaries, oversight, transparency, privacy, fallback, change, monitoring, incident response, and retirement. Adoption should produce a controlled operating state in which leadership can understand the material risks, engineering teams can act on explicit requirements, assessors can test implementation and operating effectiveness, and the organization can support every conformance or trust claim with current evidence.

2. Executive Findings and Required Direction

Finding	Enterprise Exposure	Required Direction
AI without baseline justification	Teams use AI where deterministic methods would be safer, cheaper, or more reliable.	Compare against a non-AI baseline and document the reason for AI use.
Opaque model and provider dependency	Organizations cannot reproduce which model produced an outcome or where data went.	Inventory exact model, provider, runtime, configuration, data handling, and route.
Anecdotal evaluation	Demo quality replaces representative and harm-aware evaluation.	Use task-specific datasets, repeated trials, subgroup analysis, and regression gates.
Grounding and instruction failures	Models treat untrusted content as instructions or generate unsupported claims.	Separate instructions from sources, require citations where relevant, and validate outputs.
AI action authority	Generated text or tool calls can become system changes without independent control.	Keep authorization and execution outside the model and provide recourse, fallback, and shutdown.

3. Business and Operational Impact

3.1 Strategic Impact

This standard converts a discipline that is often dependent on individual expertise into an organization-wide system of ownership, records, controls, review, and evidence. It reduces decision ambiguity and makes material assumptions visible before they become embedded in products or operations.

3.2 Delivery and Cost Impact

Adoption requires investment in accountable roles, authoritative records, validation, tooling, and review. That cost should be measured against avoidable rework, delayed releases, incidents, regulatory exposure, duplicated platforms, failed integrations, and unsupported product claims.

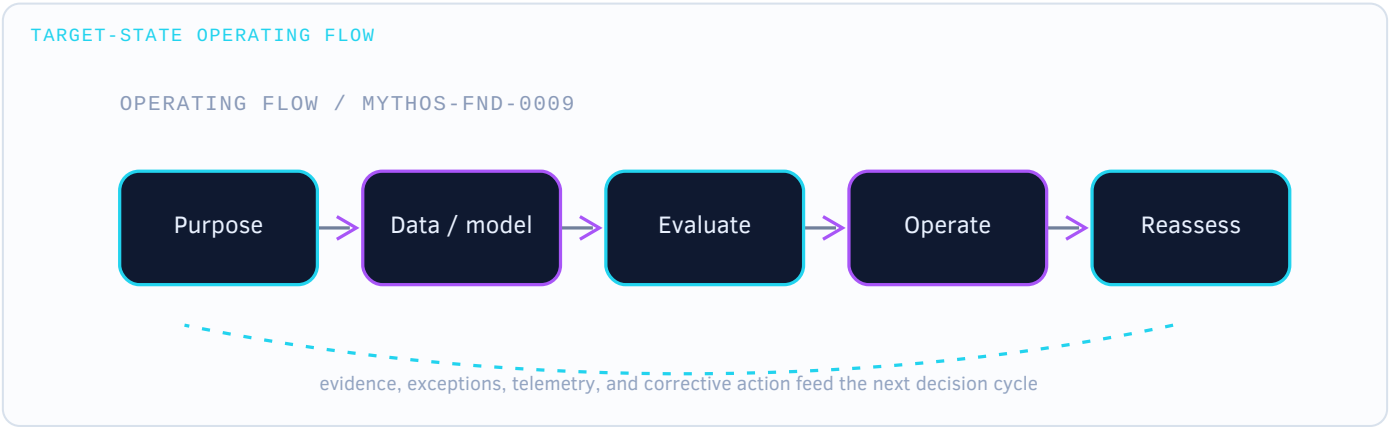
3.3 Security, Privacy, and Trust Impact

The standard requires security, privacy, integrity, resilience, and delegated authority to be considered within the primary operating model rather than as downstream approvals. This reduces the likelihood that unsafe boundaries become structurally expensive to correct.

3.4 Workforce and Organizational Impact

The organization must distinguish accountable ownership, implementation, approval, and independent challenge. Adoption may expose gaps in authority, skills, evidence, or cross-functional participation that require leadership action.

4. Target-State Summary



The target state contains the following institutional components:

#	Target-State Component
1	AI use-case and baseline assessment
2	Impact and risk assessment
3	Accountability and decision rights
4	Training, retrieval, and inference data governance
5	Model and provider inventory
6	Fitness-based model selection
7	Representative evaluation
8	Grounding and source integrity
9	Prompt and instruction governance
10	Action and decision boundaries
11	Human oversight and recourse
12	User disclosure and provenance
13	AI threat model and abuse resistance
14	Privacy-preserving operation

#	Target-State Component
15	Fallback, abstention, and degraded mode
16	Model, prompt, and provider change control
17	Production monitoring
18	AI incident response and shutdown
19	Retirement and records



5. Leadership Responsibilities

Role	Accountability
AI accountable executive	Owns AI risk appetite and portfolio governance.
Product owner	Owns the user outcome and non-AI baseline.
AI system owner	Owns model, prompt, data, evaluation, provider, and lifecycle.
Model risk or validation lead	Independently evaluates fitness and harm.
Data owner	Approves training, retrieval, inference, and retention use.
Security and privacy lead	Owns threat model, data movement, abuse, and incident controls.
Human-oversight owner	Designs intervention, recourse, appeal, and escalation.
Operations owner	Monitors performance, drift, cost, incidents, and shutdown readiness.

Leadership must ensure that exceptions are explicit, risk-owned, time-bounded, monitored, and retired. A maturity score or successful audit sample does not replace accountability for known nonconformance or residual risk.

Part II – Standard Foundation

6. Purpose and Objectives

- Use AI only where it provides defensible value relative to simpler alternatives.
- Make model uncertainty, limitations, and authority boundaries explicit.
- Require evaluation against representative tasks, harms, abuse, and operational conditions.
- Preserve accountability and safe fallback across model, provider, and data changes.

7. Scope

Applies to applications using predictive, generative, multimodal, retrieval-augmented, fine-tuned, embedded, local, hosted, or third-party AI models. It covers assistive and autonomous behavior but excludes purely deterministic software falsely marketed as AI.

8. Intended Audience

- AI product and engineering teams
- Security, privacy, legal, risk, and safety teams
- Data and model owners
- Operators and incident responders
- Assessors and procurement teams

9. Requirements Language and Conformance

The key words **MUST**, **MUST NOT**, **SHOULD**, **SHOULD NOT**, and **MAY** are normative only when written in uppercase and are interpreted in accordance with RFC 2119 and RFC 8174.

A conformance claim **MUST** identify the assessed scope, standard version, selected assurance profile, tailoring decisions, evidence period, assessor, and unresolved findings. Profiles are cumulative unless a control explicitly states otherwise.

- **Foundational:** broadly applicable minimum controls.
- **Advanced:** stronger governance, automation, scale, and independent verification.
- **High Assurance:** continuous or independent validation, stronger isolation, adversarial testing, formal analysis, or cryptographic evidence where warranted.

10. Normative References from Source Draft

- **RFC 2119** — Key words for use in RFCs to Indicate Requirement Levels. <https://www.rfc-editor.org/rfc/rfc2119>

- **RFC 8174** — Ambiguity of Uppercase vs Lowercase in RFC 2119 Key Words. <https://www.rfc-editor.org/rfc/rfc8174>
- **MYTHOS-GOV-0001** — Mythos Engineering Standards Authoring and Benchmark Framework, Version 0.1.0. [../governance/01_MYTHOS_STANDARDS_AUTHORIZING_AND_BENCHMARK_FRAMEWORK.md](https://github.com/Mythos-Engineering/standards/blob/main/governance/01_MYTHOS_STANDARDS_AUTHORIZING_AND_BENCHMARK_FRAMEWORK.md)

11. Informative References from Source Draft

- **NIST AI 100-1** — Artificial Intelligence Risk Management Framework (AI RMF 1.0). <https://doi.org/10.6028/NIST.AI.100-1>
- **NIST AI 600-1** — Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile. <https://doi.org/10.6028/NIST.AI.600-1>
- **NIST SP 800-218A** — Secure Software Development Practices for Generative AI and Dual-Use Foundation Models. <https://csrc.nist.gov/pubs/sp/800/218/a/final>
- **ISO/IEC 42001:2023** — Artificial intelligence — Management system. <https://www.iso.org/standard/81230.html>
- **ISO/IEC 23894:2023** — Artificial intelligence — Guidance on risk management. <https://www.iso.org/standard/77304.html>
- **ISO/IEC 42005:2025** — Artificial intelligence — AI system impact assessment. <https://www.iso.org/standard/44545.html>
- **OWASP Top 10 for LLM Applications 2025** — Risk taxonomy for applications using large language models. <https://genai.owasp.org/llm-top-10/>
- **NIST Privacy Framework 1.0** — A Tool for Improving Privacy through Enterprise Risk Management. <https://www.nist.gov/privacy-framework/privacy-framework>

External mappings are informative unless explicitly incorporated into a contract or regulatory obligation. Adopted organizations remain responsible for validating current source versions and legal applicability.

12. Terms and Definitions

Term	Definition
Accountable owner	The person or formally delegated role that accepts responsibility for an outcome, decision, control, or risk.
Assessment object	The system, process, component, artifact, team, service, or organizational scope being evaluated.
Compensating control	An alternative safeguard that provides comparable risk reduction when the specified control cannot be implemented as written.
Conformance claim	A bounded statement that an identified scope satisfies the applicable requirements of a named standard version and assurance profile.
Evidence	Reviewable information that supports a conclusion about design, implementation, operation, or control effectiveness.
High Assurance	The cumulative profile requiring stronger independence, adversarial testing, continuous verification, or formal evidence where risk warrants it.
Residual risk	Risk remaining after controls, mitigations, and compensating measures are applied.
System	A product, service, application, platform, workflow, integration, operational capability, or combined socio-technical environment.
Tailoring	A documented decision to include, strengthen, parameterize, or mark a control not applicable based on scope and risk.

Term	Definition
AI-first application	An application whose intended value or material behavior depends on one or more AI models rather than treating AI as a superficial add-on.
Model authority	The set of decisions, actions, tools, data, and resources a model or AI workflow is permitted to influence or control.
Evaluation case	A versioned input, context, expected properties, scoring method, and acceptance rule used to evaluate an AI behavior.
Fallback mode	A defined operating state used when a model, provider, retrieval source, or safety control is unavailable or untrusted.

13. Applicability and Tailoring

The organization **MUST** determine applicability at the control level. A not-applicable decision **MUST** identify the assessed scope, factual basis, approver, review date, and any assumptions that would invalidate the decision. Tailoring may strengthen or parameterize a control; it **MUST NOT** silently weaken a **MUST** requirement. Compensating controls require documented equivalence of risk reduction.

Part III – Risk and Architecture

14. Enterprise Problem and Risk Landscape

The principal enterprise risks addressed by this standard include inconsistent ownership, undocumented authority, uncontrolled change, local optimization, evidence gaps, stale assumptions, supplier dependence, false assurance, and the inability to determine whether the operating system still matches its approved intent.

Adopting organizations **SHOULD** maintain a domain-specific risk register that identifies cause, affected asset or stakeholder, credible scenario, impact, existing controls, residual risk, owner, review trigger, and evidence. Risks that cross standards **MUST** be linked rather than copied into disconnected registers.

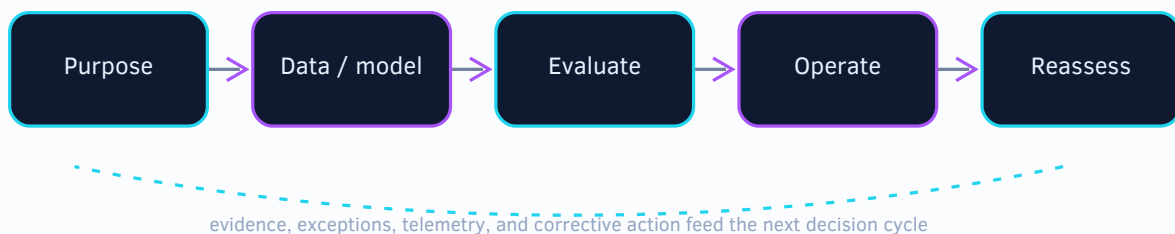
15. Governing Principles

- A credible non-AI baseline is part of AI design.
- Model output is untrusted until validated for its intended use.
- Evaluation must match the actual task, users, context, and consequence.
- Human review is not a control unless the reviewer has time, competence, evidence, and authority.
- Provider convenience does not override data, security, or accountability requirements.
- Every model-dependent behavior requires an explicit failure and retirement strategy.

16. Reference Operating Architecture

REFERENCE OPERATING ARCHITECTURE

OPERATING FLOW / MYTHOS-FND-0009





16.1 Control Plane

The control plane contains accountable ownership, policy, standards, risk decisions, approvals, exception governance, and authoritative state. It **MUST** be distinguishable from implementation and reporting systems.

16.2 Delivery and Enforcement Plane

The delivery plane contains engineering workflows, automation, validation, configuration, runtime enforcement, and lifecycle processes. Automated enforcement **SHOULD** be preferred where requirements can be expressed and tested safely.

16.3 Evidence and Assurance Plane

The assurance plane collects evidence, observations, test results, decisions, exceptions, and historical state. Evidence systems **MUST** protect integrity, scope, provenance, retention, and authorized access.

16.4 Feedback Plane

Metrics, incidents, assessments, user outcomes, exceptions, and changing authorities feed corrective action and controlled revision. Feedback **MUST** not silently alter normative requirements or accepted risk.

17. Roles and Accountability

Role	Accountability
AI accountable executive	Owns AI risk appetite and portfolio governance.
Product owner	Owns the user outcome and non-AI baseline.
AI system owner	Owns model, prompt, data, evaluation, provider, and lifecycle.
Model risk or validation lead	Independently evaluates fitness and harm.
Data owner	Approves training, retrieval, inference, and retention use.
Security and privacy lead	Owns threat model, data movement, abuse, and incident controls.
Human-oversight owner	Designs intervention, recourse, appeal, and escalation.
Operations owner	Monitors performance, drift, cost, incidents, and shutdown readiness.

18. State, Evidence, and Decision Model

The adopting organization **MUST** distinguish:

- approved intent;
- current implemented state;
- observed operational state;
- generated or inferred recommendations;
- accepted exceptions;
- historical state;
- independently verified results.

A generated report, model conclusion, roadmap, or design proposal **MUST NOT** be represented as implemented or verified state without supporting evidence.

19. Security and Trust Considerations

Every implementation of this standard **MUST** apply identity, authorization, classification, least privilege, evidence integrity, sensitive-data handling, and supplier-risk controls appropriate to impact. Machine-generated guidance, automation, extensions, and delegated agents **MUST** remain subject to external policy and independent verification.

20. Failure Modes

Failure or Anti-Pattern	Enterprise Consequence
Using AI where a deterministic control is available and adequate	Failure can create hidden risk, misleading assurance, uncontrolled cost, or unsafe operation.
Evaluating only with hand-selected prompts	Failure can create hidden risk, misleading assurance, uncontrolled cost, or unsafe operation.
Allowing model-generated confidence to serve as verification	Failure can create hidden risk, misleading assurance, uncontrolled cost, or unsafe operation.
Sending sensitive context to a provider without policy approval	Failure can create hidden risk, misleading assurance, uncontrolled cost, or unsafe operation.
Hiding AI involvement from affected users	Failure can create hidden risk, misleading assurance, uncontrolled cost, or unsafe operation.
Changing model or prompt without regression testing	Failure can create hidden risk, misleading assurance, uncontrolled cost, or unsafe operation.
Removing fallback because the primary model is usually available	Failure can create hidden risk, misleading assurance, uncontrolled cost, or unsafe operation.

21. Cross-Functional Dependencies

This Foundation standard depends on the remaining MYTHOS-FND standards for documentation, security and trust, quality, operations, data and integration, interfaces, extensions, AI-first behavior, and agentic orchestration. An adopting organization **MUST** resolve overlapping requirements through the stricter applicable control or a documented authority decision.

FOUNDATION STANDARD / STRUCTURAL PART

Part IV – Normative Control System

22. Control-Family Overview

CONTROL-FAMILY CONSTELLATION



This standard contains **19 controls** across the following families:

- Purpose
- Impact
- Ownership
- Data
- Models
- Selection
- Evaluation
- Grounding
- Prompts
- Authority
- Human control
- Transparency
- Security
- Privacy
- Reliability
- Change
- Monitoring
- Incident

- Lifecycle

23. Normative Controls

CONTROL SET 19 atomic requirements	PROFILES 3 cumulative assurance levels	ASSESSMENT 4 Examine / Interview / Test / Measure
--	--	---

Each control is independently addressable, evidence-bound, and assessable. The following control records preserve the source requirements while presenting implementation guidance, required evidence, assessment procedures, exceptions, compensating controls, and external mappings as a consistent engineering system.

NORMATIVE CONTROL CONSTELLATION



NORMATIVE CONTROL

 MYTHOS-FND-AIA-01 – AI use-case justification and baseline

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Purpose	Foundational, Advanced, High Assurance	Critical	Partial

Requirement

Each material AI capability MUST document the problem, intended users, expected benefit, non-AI baseline, reasons AI is appropriate, failure consequences, and criteria for retaining, replacing, or removing AI.

Purpose

Prevents unjustified complexity and automation theater.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Compare rules, search, statistics, and human workflow alternatives.
- Include cost, latency, privacy, and operability.

Required evidence

- Use-case record
- Baseline comparison
- Decision approval

Assessment procedure

- **Examine:** Review whether evidence supports AI selection.
- **Interview:** Confirm ownership and practice.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- NIST AI RMF 1.0

NORMATIVE CONTROL

MYTHOS-FND-AIA-02 – AI impact and risk assessment

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Impact	Foundational, Advanced, High Assurance	Critical	Partial

Requirement

Before deployment and after material change, the organization **MUST** assess affected people, rights, safety, security, privacy, environment, business processes, misuse, systemic effects, and reversibility.

Purpose

Makes consequence visible before commitment.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Include direct, indirect, cumulative, and reasonably foreseeable impacts.
- Engage affected expertise.

Required evidence

- Impact assessment
- Risk register
- Review record

Assessment procedure

- **Examine:** Trace mitigations to identified impacts.
- **Interview:** Confirm ownership and practice.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- ISO/IEC 42005:2025
- NIST AI RMF 1.0

NORMATIVE CONTROL

 MYTHOS-FND-AIA-03 – Accountability and decision rights

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Ownership	Foundational, Advanced, High Assurance	Critical	Low

Requirement

AI systems **MUST** have accountable owners for product outcome, model behavior, data, security, privacy, evaluation, operation, incident response, and risk acceptance, with conflicts and escalation paths documented.

Purpose

Prevents responsibility gaps across providers and teams.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Do not assign accountability to the model or vendor.

Required evidence

- RACI or equivalent
- Escalation map
- Approvals

Assessment procedure

- **Examine:** Interview owners on failure scenarios.
- **Interview:** Confirm ownership and practice.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- ISO/IEC 42001:2023

NORMATIVE CONTROL

 MYTHOS-FND-AIA-04 – Training, retrieval, and inference data governance

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Data	Foundational, Advanced, High Assurance	Critical	Partial

Requirement

Data used for training, tuning, retrieval, evaluation, prompting, and inference **MUST** have documented provenance, rights, purpose, quality, sensitivity, retention, access, deletion, contamination, and representativeness controls.

Purpose

Controls legal, privacy, quality, and security risk.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Separate evaluation data from optimization.
- Detect secrets and personal data.

Required evidence

- Data inventory
- Lineage
- Rights and quality records

Assessment procedure

- **Examine:** Trace samples to source and purpose.
- **Interview:** Test deletion.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- NIST AI 600-1
- NIST Privacy Framework 1.0

NORMATIVE CONTROL

 MYTHOS-FND-AIA-05 – Model and provider inventory

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Models	Foundational, Advanced, High Assurance	Critical	High

Requirement

The organization **MUST** maintain an inventory of models, versions, providers, deployment locations, licenses, capabilities, limitations, data-use terms, dependencies, safety settings, and approved use cases.

Purpose

Enables change control and incident response.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Include embeddings, rerankers, classifiers, and guard models.

Required evidence

- Model registry
- Provider assessments
- Deployment records

Assessment procedure

- **Examine:** Compare production calls to registry.
- **Interview:** Confirm ownership and practice.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- ISO/IEC 42001:2023

NORMATIVE CONTROL

 MYTHOS-FND-AIA-06 – Model selection by measured fitness

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Selection	Foundational, Advanced, High Assurance	Material	Partial

Requirement

Model and provider selection **MUST** be based on versioned evaluation of task quality, risk, latency, reliability, cost, context behavior, data handling, portability, and operational constraints rather than reputation alone.

Purpose

Aligns model choice to real requirements.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Evaluate multiple sizes and local options.
- Record tradeoffs and lock-in.

Required evidence

- Selection benchmark
- Decision record
- Cost model

Assessment procedure

- **Examine:** Re-run selected cases.
- **Interview:** Inspect rejected alternatives.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- NIST AI RMF 1.0

NORMATIVE CONTROL



MYTHOS-FND-AIA-07 – Representative predeployment evaluation

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Evaluation	Foundational, Advanced, High Assurance	Critical	Partial

Requirement

AI capabilities **MUST** pass versioned evaluations covering normal, boundary, ambiguous, multilingual or multimodal where applicable, adversarial, harmful, privacy, security, bias, hallucination, refusal, tool-use, latency, cost, and failure conditions.

Purpose

Provides evidence beyond demonstrations.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Use blinded or independent review where stakes warrant.
- Report confidence intervals and limitations.

Required evidence

- Evaluation suite
- Results
- Acceptance rules

Assessment procedure

- **Examine:** Reproduce sampled cases.
- **Interview:** Inspect hidden test separation.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- NIST AI 600-1
- NIST SP 800-218A

NORMATIVE CONTROL

 MYTHOS-FND-AIA-08 – Grounding and source integrity

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Grounding	Foundational, Advanced, High Assurance	Critical	Partial

Requirement

When outputs depend on external knowledge, the system **MUST** define authoritative sources, retrieval permissions, freshness, provenance, ranking, conflict handling, citation behavior, injection resistance, and abstention criteria.

Purpose

Reduces unsupported and manipulated outputs.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Preserve source boundaries and access controls.
- Do not treat retrieved text as instruction by default.

Required evidence

- Retrieval design
- Source registry
- Grounding evaluations

Assessment procedure

- **Examine:** Inject conflicting and malicious documents.
- **Interview:** Verify access filtering.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- NIST AI 600-1
- OWASP Top 10 for LLM Applications 2025

NORMATIVE CONTROL

 MYTHOS-FND-AIA-09 – Prompt and instruction governance

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Prompts	Foundational, Advanced, High Assurance	Critical	High

Requirement

Material system prompts, policies, tools, templates, and context assembly logic **MUST** be versioned, reviewed, tested, access-controlled, and traceable to deployed behavior.

Purpose

Treats instructions as executable policy.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Separate trusted policy from untrusted content.
- Prevent tenant leakage through shared context.

Required evidence

- Prompt registry
- Diffs and approvals
- Tests

Assessment procedure

- **Examine:** Verify deployed hash.
- **Interview:** Attempt instruction override.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- NIST SP 800-218A

NORMATIVE CONTROL

 MYTHOS-FND-AIA-10 – Action and decision boundaries

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Authority	Foundational, Advanced, High Assurance	Critical	High

Requirement

AI-generated recommendations, decisions, and actions MUST be constrained by explicit policy, identity, authorization, resource, budget, temporal, and consequence boundaries enforced outside the model.

Purpose

Prevents language output from becoming ambient authority.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Use deterministic validation and approval for consequential actions.

Required evidence

- Authority policy
- Gateway configuration
- Negative tests

Assessment procedure

- **Examine:** Attempt unauthorized tool and resource use.
- **Interview:** Confirm ownership and practice.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- NIST AI RMF 1.0
- OWASP LLM Top 10

NORMATIVE CONTROL



MYTHOS-FND-AIA-11 – Meaningful human oversight and recourse

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Human control	Foundational, Advanced, High Assurance	Critical	Partial

Requirement

Where human review or intervention is required, the system **MUST** provide adequate context, uncertainty, alternatives, time, competence, authority, escalation, and recourse; nominal approval alone **MUST NOT** be treated as effective oversight.

Purpose

Prevents rubber-stamping and accountability theater.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Measure override and disagreement patterns.
- Support appeal for affected users.

Required evidence

- Oversight design
- Training records
- Decision samples

Assessment procedure

- **Examine:** Observe representative reviews.
- **Interview:** Test escalation and appeal.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- NIST AI RMF 1.0
- ISO/IEC 42005:2025

NORMATIVE CONTROL



MYTHOS-FND-AIA-12 – User disclosure and output provenance

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Transparency	Foundational, Advanced, High Assurance	Material	Partial

Requirement

Users **MUST** be informed when AI materially generates, transforms, ranks, recommends, or acts, and consequential outputs **MUST** carry provenance sufficient to identify model or workflow version, relevant sources, transformations, and review status.

Purpose

Supports informed reliance and investigation.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Avoid implying human authorship or certainty.

Required evidence

- Disclosure design
- Provenance records
- User tests

Assessment procedure

- **Examine:** Trace a material output to its generation context.
- **Interview:** Confirm ownership and practice.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- NIST AI 600-1

NORMATIVE CONTROL

 MYTHOS-FND-AIA-13 – AI threat model and abuse resistance

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Security	Advanced, High Assurance	Critical	Partial

Requirement

AI applications **MUST** maintain a threat model addressing prompt injection, data poisoning, model theft, extraction, insecure output handling, excessive agency, supply chain, denial of service, cross-tenant leakage, evasion, and unsafe tool use.

Purpose

Addresses AI-specific attack surfaces.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Test direct and indirect injection.
- Treat model output as untrusted input.

Required evidence

- Threat model
- Red-team results
- Mitigations

Assessment procedure

- **Examine:** Execute representative abuse cases.
- **Interview:** Confirm ownership and practice.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- OWASP Top 10 for LLM Applications 2025
- NIST SP 800-218A

NORMATIVE CONTROL

 MYTHOS-FND-AIA-14 – Privacy-preserving AI operation

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Privacy	Foundational, Advanced, High Assurance	Critical	Partial

Requirement

AI applications **MUST** minimize personal and sensitive data, enforce purpose and access boundaries, prevent unauthorized retention or provider use, support applicable rights, and evaluate memorization, inference, reidentification, and disclosure risk.

Purpose

Controls privacy risks unique to model workflows.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Prefer local or protected processing where justified.
- Redact before prompts when feasible.

Required evidence

- Privacy assessment
- Provider terms
- Leakage tests

Assessment procedure

- **Examine:** Submit canary secrets and deletion requests.
- **Interview:** Confirm ownership and practice.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- NIST Privacy Framework 1.0
- NIST AI 600-1

NORMATIVE CONTROL



MYTHOS-FND-AIA-15 – Fallback, abstention, and degraded mode

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Reliability	Foundational, Advanced, High Assurance	Critical	Partial

Requirement

AI capabilities **MUST** define confidence or validity checks, abstention behavior, fallback modes, provider or model failure behavior, stale-context limits, timeout budgets, and safe recovery for material workflows.

Purpose

Maintains safe service when AI is uncertain or unavailable.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Prefer explicit unavailable state over fabricated completion.

Required evidence

- Fallback design
- Failure tests
- Operational playbook

Assessment procedure

- **Examine:** Disable model and retrieval dependencies.
- **Interview:** Force low-confidence inputs.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- ISO/IEC 25010:2023

NORMATIVE CONTROL



MYTHOS-FND-AIA-16 – Model, prompt, and provider change control

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Change	Foundational, Advanced, High Assurance	Critical	High

Requirement

Changes to models, providers, routing, prompts, retrieval, safety settings, data, tools, or decoding that may affect behavior **MUST** trigger impact analysis, regression evaluation, approval, staged release, and rollback according to risk.

Purpose

Prevents silent behavioral drift.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Treat provider aliases as mutable dependencies.
- Pin versions where possible.

Required evidence

- Change records
- Regression results
- Deployment evidence

Assessment procedure

- **Examine:** Introduce a model alias change.
- **Interview:** Verify gate and rollback.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- NIST SP 800-218A

NORMATIVE CONTROL



MYTHOS-FND-AIA-17 – Production performance and harm monitoring

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Monitoring	Advanced, High Assurance	Critical	Partial

Requirement

Production AI systems MUST monitor task outcomes, quality proxies, drift, refusals, unsafe outputs, policy denials, latency, cost, provider changes, data quality, user feedback, and incidents with privacy-preserving sampling and escalation.

Purpose

Detects failures not captured before release.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Define leading and lagging indicators.
- Avoid treating model self-evaluation as sole evidence.

Required evidence

- Monitoring plan
- Dashboards
- Alerts and incident records

Assessment procedure

- **Examine:** Inject monitored failures.
- **Interview:** Review false negatives.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- NIST AI RMF 1.0

NORMATIVE CONTROL

 MYTHOS-FND-AIA-18 – AI incident response and shutdown

FAMILY	PROFILES	FAILURE IMPACT	AUTOMATION POTENTIAL
Incident	Advanced, High Assurance	Critical	High

Requirement

AI applications **MUST** support rapid containment, model or feature disablement, routing changes, evidence preservation, user and stakeholder notification, correction, rollback, and post-incident review.

Purpose

Limits harm from emergent or provider-driven failures.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Predefine kill switches that do not depend on the affected model.

Required evidence

- Incident plan
- Drill results
- Disablement evidence

Assessment procedure

- **Examine:** Run unsafe-output and provider-compromise exercises.
- **Interview:** Confirm ownership and practice.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- NIST SP 800-61 Rev. 3

NORMATIVE CONTROL



MYTHOS-FND-AIA-19 – Retirement, replacement, and record retention

FAMILY Lifecycle	PROFILES Advanced, High Assurance	FAILURE IMPACT Material	AUTOMATION POTENTIAL Partial
----------------------------	---	-----------------------------------	--

Requirement

AI capabilities **MUST** define end-of-support, replacement, migration, data and artifact retention, reproducibility, user communication, dependent-workflow handling, and deletion requirements.

Purpose

Prevents abandoned models and irreproducible decisions.

Applicability

Applicable to all in-scope systems and organizational processes unless a documented tailoring decision states otherwise.

Implementation guidance

- Retain evidence proportionate to consequence and law.

Required evidence

- Retirement plan
- Migration tests
- Retention records

Assessment procedure

- **Examine:** Retire a test model version and verify dependencies.
- **Interview:** Confirm ownership and practice.
- **Test:** Exercise the requirement.

Exceptions and compensating controls

Any exception must be approved by an accountable owner, time-bounded, risk-assessed, monitored, and supported by compensating controls.

External mappings

- ISO/IEC/IEEE 15288:2023

Part V – Implementation and Operations

24. Implementation Profiles

Profile	Required Operating State
Baseline	Use-case justification, impact assessment, inventory, data governance, model fitness, evaluation, grounding, prompt control, action boundary, disclosure, fallback, monitoring, incident response, and retirement.
Enterprise	Central AI inventory, provider policy, evaluation platform, prompt and model registry, continuous monitoring, human-oversight operations, source governance, and audit-ready evidence.
High Assurance	Independent validation, adversarial and red-team evaluation, strict provider placement, protected datasets, constrained decoding, external authorization, continuous harm monitoring, supervised shutdown, and regulatory evidence packages.

Profiles are cumulative unless a control explicitly states otherwise. A higher profile cannot be claimed solely through additional tooling; governance, evidence, operating effectiveness, and independent challenge must also satisfy the profile.

25. Implementation Lifecycle

25.1 Establish

- appoint accountable ownership;
- determine applicability and profile;
- inventory current processes, systems, records, and known exceptions;
- identify authority sources and legal applicability;
- establish evidence locations and review cadence.

25.2 Baseline

- assess every applicable control;
- distinguish implemented, partially implemented, not implemented, not applicable, and not assessed;
- record evidence quality and confidence;
- identify systemic dependencies and priority risks.

25.3 Implement

- define target architecture and ownership;
- create or revise policies, contracts, workflows, and controls;
- automate enforcement and evidence where safe;
- test failure and recovery paths;
- train accountable roles.

25.4 Assure

- conduct technical and procedural assessment;
- verify evidence over a representative period;
- test sampled controls and critical workflows;
- challenge exceptions, unsupported claims, and optimistic assumptions.

25.5 Operate and Improve

- monitor metrics and exceptions;
- investigate incidents and control failures;
- update for authority, threat, technology, or organizational change;
- preserve superseded decisions and evidence;
- retire obsolete controls and implementations deliberately.

26. Integration Requirements

Implementations SHOULD integrate the standard with product governance, architecture review, secure development, CI/CD, change management, observability, service management, identity, evidence, risk, procurement, and training systems. Integration MUST preserve ownership and source authority rather than copying uncontrolled state between tools.

27. Failure Handling and Recovery

Control failures MUST produce a classified finding, owner, containment or compensating action, due date, evidence requirement, and escalation path. Where failure creates ongoing material risk, the organization MUST consider stopping, restricting, rolling back, isolating, or retiring the affected activity.

28. Exceptions and Compensating Controls

Exceptions MUST identify the exact control, scope, rationale, risk, owner, approval, compensating measures, monitoring, expiration, and closure evidence. Exceptions MUST NOT be used to convert a permanent design weakness into nominal conformance.

29. Prohibited Practices

- Using AI where a deterministic control is available and adequate
- Evaluating only with hand-selected prompts
- Allowing model-generated confidence to serve as verification
- Sending sensitive context to a provider without policy approval
- Hiding AI involvement from affected users
- Changing model or prompt without regression testing
- Removing fallback because the primary model is usually available

30. Adoption Roadmap from Source Draft

- Document use case, non-AI baseline, impact, authority, data, model, and provider decisions.
- Build versioned evaluation, threat, privacy, and fallback plans before production integration.
- Deploy through policy-enforced gateways with provenance, telemetry, limits, and rollback.
- Establish continuous evaluation, incident response, model-change governance, and user recourse.
- For High Assurance, add independent validation, adversarial testing, stronger isolation, reproducibility, and formal constraints for consequential actions.

31. Limitations and Residual Risk from Source Draft

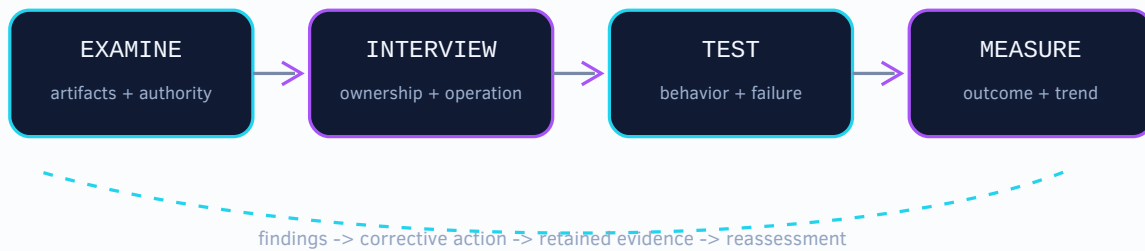
Conformance demonstrates that defined requirements were satisfied within a bounded scope and evidence period. It does not guarantee absence of defects, compromise, harm, outage, legal exposure, or future failure. Novel threats, misunderstood dependencies, invalid assumptions, human error, supplier changes, and conditions outside the assessment scope remain possible.

FOUNDATION STANDARD / STRUCTURAL PART

Part VI – Assurance and Assessment

32. Assessment Methodology

ASSESSMENT EVIDENCE LOOP



Assessors **MUST** use a combination of examination, interview, and testing appropriate to the selected profile and control risk. Assessment records **MUST** identify sample size, tools, versions, evidence references, exclusions, limitations, confidence, and residual uncertainty. Automated checks **MAY** support a finding but **MUST NOT** be treated as proof of effective governance or operation when the control depends on human accountability or contextual judgment.

12.1 Finding States

- **Satisfied**: requirement implemented and supported by sufficient evidence.
- **Partially satisfied**: some required outcomes or scope are missing.
- **Not satisfied**: requirement absent, ineffective, or contradicted by evidence.
- **Not applicable**: supported by an approved tailoring rationale.
- **Not assessed**: evidence is insufficient.

12.2 Conformance

A scope is **Conformant** only when all applicable **MUST** controls for the claimed profile are satisfied. A failed critical **MUST** control cannot be averaged away by stronger performance elsewhere.

33. Metrics and Reporting from Source Draft

Metric	Definition	Expected use or target
Evaluation coverage	Material capabilities and risk scenarios represented in versioned suites	100% of release-critical capabilities
Critical failure rate	High-consequence evaluation cases failing acceptance criteria	

Metric	Definition	Expected use or target
		0 at release unless formally accepted
Grounded-answer support	Material factual outputs supported by approved evidence when grounding is required	Target defined by use case
Unsafe action prevention	Unauthorized or policy-violating model-originated actions blocked	100% in defined tests
Drift detection latency	Time to identify material performance or risk regression	Risk-based service objective

Metrics **MUST** be interpreted with scope and uncertainty. Organizations **SHOULD** report trends, distributions, and exceptions rather than presenting a single composite score as complete assurance.

34. Exceptions and Compensating Controls

Every exception **MUST** identify the affected control and scope, justification, risk, compensating controls, accountable approver, start and expiration dates, monitoring, and remediation or retirement plan. Exceptions **MUST** be reevaluated after material change or incident and **MUST NOT** be self-approved by the team or agent requesting the exception where separation of duties is required.

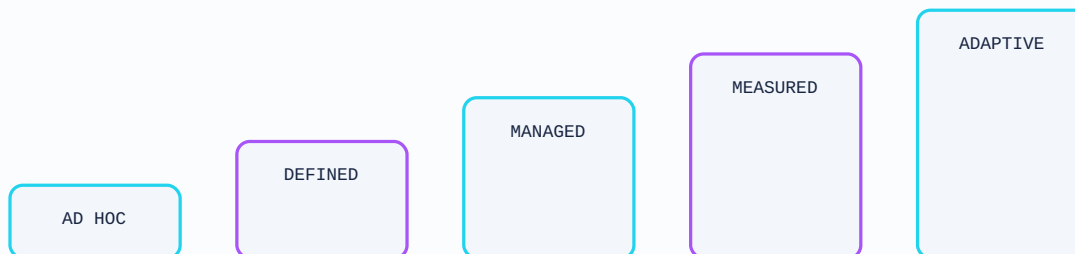
35. Enterprise Metric Set

Metric	Measurement Requirement
AI use cases without approved baseline	Defined by the adopting organization with owner, source, target, and cadence.
models and providers without current inventory	Defined by the adopting organization with owner, source, target, and cadence.
evaluation cases mapped to material risks	Defined by the adopting organization with owner, source, target, and cadence.
grounded claims with valid source support	Defined by the adopting organization with owner, source, target, and cadence.
unsupported-answer and abstention rates	Defined by the adopting organization with owner, source, target, and cadence.
high-impact actions requiring human or external authorization	Defined by the adopting organization with owner, source, target, and cadence.
production drift and harm indicators	Defined by the adopting organization with owner, source, target, and cadence.
AI incidents and shutdown exercises	Defined by the adopting organization with owner, source, target, and cadence.
users receiving required disclosure and recourse	Defined by the adopting organization with owner, source, target, and cadence.

36. Maturity Model

MATURITY PROGRESSION

MATURITY IS EVIDENCE OF CAPABILITY - NOT A SUBSTITUTE FOR CONFORMANCE



Level	Name	Required Characteristics
M0	Unmanaged	Outcome is not intentionally controlled or evidence is unavailable.
M1	Documented	Ownership and procedures exist but implementation is inconsistent or mostly manual.
M2	Implemented	Applicable controls operate in the assessed scope and evidence exists.
M3	Measured	Effectiveness, exceptions, failures, and trends are measured and reviewed.
M4	Assured	Independent testing, representative evidence, and continuous or periodic validation exist.
M5	Adaptive	Controls respond safely to changing risk, authority, technology, and operational evidence.

Maturity does not override conformance. An organization cannot compensate for a failed mandatory requirement by assigning itself a high maturity level.

37. Assessment Confidence

Assessment confidence **MUST** be recorded as Low, Moderate, High, or Very High. Confidence considers evidence integrity, observation period, sample coverage, source authority, environmental representativeness, test repeatability, reviewer independence, and unresolved gaps.

38. Executive Assurance Dashboard

The executive report **SHOULD** present:

- applicable controls;
- conformant, partial, nonconformant, exception, not applicable, and not assessed counts;
- high and critical findings;
- evidence freshness;
- exception age;
- maturity distribution;
- assessment confidence;
- top residual risks;

-
- corrective-action status.
-

Part VII – Authority and Traceability

39. Cross-Standard Dependencies from Source Draft

This standard is part of an integrated foundation. Product decisions depend on documentation and traceability; implementation depends on security, quality, data, interfaces, and extension controls; release depends on operational readiness; AI and agentic behavior inherit every applicable non-AI requirement. A specialized standard MAY strengthen these requirements but MUST NOT weaken them without an explicit supersession rule.

40. Current External Authority Register

Authority	Relevance	Official Location	Status	Validated
NIST AI RMF 1.0	AI risk management	https://www.nist.gov/itl/ai-risk-management-framework	Current; revision program active	2026-09-17
NIST AI 600-1	Generative AI Profile	https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence	Final	2026-09-17
NIST SP 800-218A	SSDF Community Profile for Generative AI and Dual-Use Foundation Models	https://csrc.nist.gov/pubs/sp/800/218/a/final	Final	2026-09-17
ISO/IEC 42001:2023	AI management systems	https://www.iso.org/standard/42001	Published	2026-09-17
ISO/IEC 23894:2023	AI risk management guidance	https://www.iso.org/standard/77304.html	Published	2026-09-17
ISO/IEC 42005:2025	AI system impact assessment	https://www.iso.org/standard/42005	Published	2026-09-17
ISO/IEC 5338:2023	AI system life cycle processes	https://www.iso.org/standard/81118.html	Published	2026-09-17
Regulation (EU) 2024/1689	Artificial Intelligence Act	https://eur-lex.europa.eu/eli/reg/2024/1689/	In force; staged application	2026-09-17

The authority register is informative unless an adopting organization incorporates a source into a binding obligation. Legal applicability and licensed ISO content must be evaluated by qualified parties. The register records current source identity and relevance without reproducing protected standards text.

41. Control Traceability

Each control MUST remain traceable to:

- the risk or outcome it addresses;
- the applicable profile;
- implementation ownership;

-
- evidence;
 - assessment procedure;
 - exceptions;
 - external mappings;
 - related Foundation controls.

42. Source-Draft Preservation

This enterprise candidate edition preserves every normative control and the substantive source-draft sections. Editorial movement into the enterprise report structure does not remove the original requirement, implementation guidance, evidence, assessment procedure, exception rule, or external mapping.

FOUNDATION STANDARD / STRUCTURAL PART

Part VIII – Appendices

Appendix A – Source-Draft Evolution Notes

- Recasts “AI-first” as disciplined product architecture rather than mandatory model use.
- Integrates current AI risk, secure development, impact assessment, and application security patterns.
- Separates AI application governance from the deeper multi-agent orchestration controls in MYTHOS-FND-0010.

Appendix B – Change History

Version	Date	Status	Summary
0.2.0	2026-07-18	Working Draft	First standards-program restructuring of the source draft into atomic controls, assurance profiles, evidence, assessment procedures, and cross-standard dependencies.

Appendix C – Control Summary

Control	Family	Title	Profiles	Automation	Impact
MYTHOS-FND-AIA-01	Purpose	AI use-case justification and baseline	Foundational, Advanced, High Assurance	Partial	Critical
MYTHOS-FND-AIA-02	Impact	AI impact and risk assessment	Foundational, Advanced, High Assurance	Partial	Critical
MYTHOS-FND-AIA-03	Ownership	Accountability and decision rights	Foundational, Advanced, High Assurance	Low	Critical
MYTHOS-FND-AIA-04	Data	Training, retrieval, and inference data governance	Foundational, Advanced, High Assurance	Partial	Critical
MYTHOS-FND-AIA-05	Models	Model and provider inventory	Foundational, Advanced, High Assurance	High	Critical
MYTHOS-FND-AIA-06	Selection	Model selection by measured fitness	Foundational, Advanced, High Assurance	Partial	Material
MYTHOS-FND-AIA-07	Evaluation	Representative predeployment evaluation	Foundational, Advanced, High Assurance	Partial	Critical
MYTHOS-FND-AIA-08	Grounding	Grounding and source integrity	Foundational, Advanced, High Assurance	Partial	Critical
MYTHOS-FND-AIA-09	Prompts	Prompt and instruction governance	Foundational, Advanced, High Assurance	High	Critical
MYTHOS-FND-AIA-10	Authority	Action and decision boundaries	Foundational, Advanced, High Assurance	High	Critical
MYTHOS-FND-AIA-11	Human control	Meaningful human oversight and recourse	Foundational, Advanced, High Assurance	Partial	Critical

Control	Family	Title	Profiles	Automation	Impact
MYTHOS-FND-AIA-12	Transparency	User disclosure and output provenance	Foundational, Advanced, High Assurance	Partial	Material
MYTHOS-FND-AIA-13	Security	AI threat model and abuse resistance	Advanced, High Assurance	Partial	Critical
MYTHOS-FND-AIA-14	Privacy	Privacy-preserving AI operation	Foundational, Advanced, High Assurance	Partial	Critical
MYTHOS-FND-AIA-15	Reliability	Fallback, abstention, and degraded mode	Foundational, Advanced, High Assurance	Partial	Critical
MYTHOS-FND-AIA-16	Change	Model, prompt, and provider change control	Foundational, Advanced, High Assurance	High	Critical
MYTHOS-FND-AIA-17	Monitoring	Production performance and harm monitoring	Advanced, High Assurance	Partial	Critical
MYTHOS-FND-AIA-18	Incident	AI incident response and shutdown	Advanced, High Assurance	High	Critical
MYTHOS-FND-AIA-19	Lifecycle	Retirement, replacement, and record retention	Advanced, High Assurance	Partial	Material

Appendix D – Minimum Conformance Claim Record

```

standard: MYTHOS-FND-0009
version: 0.2.0
profile: foundational | advanced | high_assurance
scope: <bounded systems, teams, environments, and processes>
assessment_period: <start/end>
assessor: <person or independent body>
tailoring_decisions: []
exceptions: []
findings_summary:
  satisfied: 0
  partially_satisfied: 0
  not_satisfied: 0
  not_applicable: 0
  not_assessed: 0
residual_risk_owner: <accountable role>
approval: <record reference>

```

Appendix E – Publication Review Checklist

- ☐ Domain technical review completed
- ☐ Security and privacy review completed
- ☐ Current primary sources validated
- ☐ Exact external mappings reviewed
- ☐ All controls testable
- ☐ Evidence requirements sufficient
- ☐ Assessment procedures repeatable
- ☐ Executive findings supported
- ☐ Legal applicability reviewed where required

-
- [] Accessibility and publication quality verified
 - [] Machine-readable control catalog validated
 - [] Change and review record complete
 - [] Formal approval recorded