

# The AISIF Detailed Crosswalk

ISO/IEC 27001:2022 · NIST AI RMF · ISO/IEC 42001:2023

OWASP LLM Top 10 (2025) · MITRE ATLAS · Databricks AI Security Framework

*Excerpted and expanded from the AISIF Practitioner White Paper v1.1, Section 5, with a second crosswalk table extending coverage to three additional frameworks · Companion resource to the AISIF toolkit · Prepared August 2026*

## Purpose and How to Use This Document

This is the extended edition of the AISIF standards crosswalk. It maps all ten AISIF control domains — organized under the framework’s five layers — to six frameworks and standards across two tables. Table 1 (Section 2) covers the three certification-relevant standards: ISO/IEC 27001:2022 Annex A, the NIST AI Risk Management Framework, and ISO/IEC 42001:2023. Table 2 (Section 3) extends the same ten domains to three widely used technical references: the OWASP Top 10 for LLM Applications (2025), MITRE ATLAS, and the Databricks AI Security Framework (DASF).

Use it in four ways. First, as a control-design reference: when building or auditing a control for a given AISIF domain, the crosswalk names the specific clauses, functions, categories, or components each external reference expects, so the same piece of evidence can support more than one framework. Second, as a gap-assessment tool: an organization working from any one of the six can see exactly which AISIF domains its existing work already substantiates, and which remain uncovered. Third, as a certification bridge: because ISO/IEC 42001 shares the High-Level Structure common to ISO management-system standards, an existing ISO 27001 ISMS can be extended into an AIMS using Table 1 as the domain-by-domain map. Fourth, as a technical-to-governance bridge: Table 2 lets a red team finding tagged to an OWASP or ATLAS category, or a Databricks DASF component, be translated directly into the AISIF domain — and from there, via Table 1, into the certification-relevant standard it should also inform.

Section 1 lists the ten control domains with their full descriptions. Section 2 presents the certification-standards crosswalk. Section 3 presents the technical-framework crosswalk. Section 4 explains how the three certification standards relate to one another and to AISIF, and closes with certification guidance. Content in Sections 1, 2, and 4 is drawn directly from the AISIF Practitioner White Paper (v1.1, Section 5); Section 3 is new to this edition. Consult the white paper for the full five-layer architecture and case studies this crosswalk supports.

# 1. The Ten AISIF Control Domains

AISIF organizes its controls into ten domains nested within the framework's five layers. Each domain below is presented with its full description and the layer under which it is scoped.

## Governance Layer

### AI Risk Governance & Lifecycle Management

Establishes the organizational accountability structures, risk management cycle, and policy framework that govern the AI system from initial deployment through ongoing operation, retraining, and eventual decommissioning — ensuring that security decisions have a named owner, a defined process, and a feedback mechanism that keeps the program current as the system and its threat landscape evolve. This is the master control domain from which all other AISIF controls derive their authority and their connection to organizational risk appetite.

### Shared Responsibility & Ownership

Documents the precise boundary between the organization's security obligations and those of the AI platform provider, cloud vendor, or third-party model supplier — making explicit which controls belong to each party across every layer of the AI stack. Without this mapping, accountability gaps at deployment boundaries are assumed away rather than addressed, and controls that no party owns are controls that no party implements.

### Regulatory Compliance & Documentation

Maintains the documentation, evidence, and audit trail required to demonstrate that the AI system meets applicable legal and regulatory obligations — including EU AI Act risk-tier requirements, sector-specific guidance such as OSFI and Health Canada AI frameworks, and the management-system requirements of ISO 42001. Distinguishes compliance demonstration, which requires structured documentation of governance and controls, from security assurance, which requires those controls to function.

## Data Layer

### Data Governance & Lineage

Establishes documented provenance, access control, and quality assurance for all data entering the AI system's training and inference pipelines — ensuring that every record in a training dataset is traceable to a verified source and that data transformations are logged end-to-end. Without this control, the model's trustworthiness cannot be asserted independently of the data estate that shaped it.

### Training Data Integrity & Poisoning Defense

Protects the training pipeline against the deliberate or accidental introduction of corrupted, adversarially crafted, or statistically anomalous records that can shift a model's decision boundary in ways that remain invisible through standard evaluation metrics. Implements hash-based batch validation, distribution monitoring, and pre-deployment decision-boundary analysis to detect poisoning before a compromised model reaches production.

## Model Layer

### Threat Modeling & Attack Simulation

Systematically identifies and evaluates ML-specific adversarial threats against AI systems using the MITRE ATLAS taxonomy — covering reconnaissance, evasion, poisoning, extraction, and model theft. Goes beyond traditional application security testing to simulate the attack techniques an adversary would realistically use against a deployed AI system's full threat surface.

### Model Access Control & Authentication

Enforces least-privilege access to model weights, inference endpoints, fine-tuning interfaces, and model management APIs — ensuring that only authorized identities and systems can query, modify, or retrieve model artifacts. Addresses the risk that inference APIs, often treated as lower sensitivity than data stores, provide a systematic extraction surface for model theft and membership-inference attacks when inadequately authenticated.

## Secure Deployment & Supply Chain

Validates the integrity of every component introduced into the AI system during development and deployment — including pre-trained model weights, third-party datasets, ML libraries, and inference infrastructure — and applies secure development lifecycle controls to the AI-specific elements of the deployment pipeline. Recognizes that model artifacts, unlike traditional software packages, can carry backdoors that survive standard code review and antivirus scanning.

## Application Layer

### Prompt Injection & Input Validation

Defends against adversarial inputs — both direct user injection and indirect injection through retrieved content such as RAG documents, tool outputs, and external data sources — by enforcing structural separation between system instructions and untrusted content. Ensures that no input pathway, including plugin calls, API responses, or retrieval-augmented context, can override system prompt authority or redirect model behavior without explicit authorization.

## Operations Layer

### Output Monitoring & Anomaly Detection

Continuously evaluates inference outputs against defined behavioral baselines — detecting model drift, adversarial probing patterns, systematic extraction attempts, and performance degradation that would otherwise surface only through downstream business failures or user complaints. Establishes the operational signal that feeds the AISIF recursive feedback loop, connecting real-world model behavior back to the governance layer's risk management cycle.

## Notes on Layer Allocation

Three placement decisions are worth making explicit:

1. Threat Modeling & Attack Simulation is positioned within the Model layer rather than Governance, as its primary function centers on machine learning–specific adversarial analysis using MITRE ATLAS–informed methodology. While its outputs inform controls across all layers, its core focus — characterizing vulnerabilities at the model level — is the model layer's principal attack surface.
2. Secure Deployment & Supply Chain is likewise situated within the Model layer rather than Operations. Its primary concern is the integrity of model artifacts and their introduction into the operational environment — a pre-operational control validating inputs prior to execution, not a monitor of post-deployment behavior.
3. Shared Responsibility & Ownership resides within the Governance layer, as it is fundamentally an accountability-mapping exercise that must precede control design. Effective control allocation across every other layer depends on this foundation; without it, control assignment and enforcement lack coherence.

## 2. The Full Crosswalk

The table below maps every AISIF control domain to specific ISO/IEC 27001:2022 Annex A controls, NIST AI RMF functions and categories, and ISO/IEC 42001:2023 clauses. Multiple citations per cell indicate the primary controls most directly satisfied by an AISIF-aligned implementation of that domain — not an exhaustive list of every applicable clause.

*Table 1. AISIF Control Domain Crosswalk — ISO/IEC 27001:2022, NIST AI RMF, ISO/IEC 42001:2023.*

| Control Domain                                         | Layer              | ISO/IEC 27001:2022 Annex A                                                                       | NIST AI RMF                                                                           | ISO/IEC 42001:2023                                                              |
|--------------------------------------------------------|--------------------|--------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------|---------------------------------------------------------------------------------|
| <b>AI Risk Governance &amp; Lifecycle Management</b>   | <b>Governance</b>  | A.5.1 Security policies<br>A.5.4 Management responsibilities<br>A.5.8 Project security           | Full GOVERN (1.1–1.7)<br>Full MAP (1.1–5.2)<br>Full MANAGE (1.1–4.2)                  | Cl. 4–10 Full AIMS<br>Cl. 6.1–6.3 Planning<br>Annex A.1–A.13                    |
| <b>Shared Responsibility &amp; Ownership</b>           | <b>Governance</b>  | A.5.2 Security roles<br>A.5.31 Legal requirements<br>A.5.19 Supplier relations                   | GOVERN 1.1 Policies<br>GOVERN 6.1 Third-party<br>MAP 1.6 Org context                  | Cl. 4.3 AIMS scope<br>Cl. 5.1 Leadership<br>Cl. 8.6 Supplier relations          |
| <b>Regulatory Compliance &amp; Documentation</b>       | <b>Governance</b>  | A.5.31 Legal requirements<br>A.5.32 IP rights<br>A.5.36 Policy compliance                        | GOVERN 1.7 Regulatory alignment<br>GOVERN 4.1 Risk policies<br>MAP 1.1                | Cl. 4.2 Interested parties<br>Cl. 6.1.1 Risks<br>Cl. 10.1 Corrective action     |
| <b>Data Governance &amp; Lineage</b>                   | <b>Data</b>        | A.5.12 Classification<br>A.5.13 Labelling<br>A.8.10 Deletion / A.5.33 Records                    | MAP 3.5 Provenance<br>GOVERN 1.4 Data policy<br>MEASURE 2.8 Bias eval                 | Cl. 8.3 Data management<br>Cl. 6.1.3 Quality & lineage<br>Annex A.6 / B.7.4     |
| <b>Training Data Integrity &amp; Poisoning Defense</b> | <b>Data</b>        | A.8.9 Config management<br>A.8.29 Security testing<br>A.8.7 Anti-malware                         | MAP 3.5 Training provenance<br>MEASURE 2.5<br>MANAGE 1.3                              | Cl. 8.3.2 Training controls<br>Cl. 8.4 Impact assessment<br>Annex A.6.2 / B.7.5 |
| <b>Threat Modeling &amp; Attack Simulation</b>         | <b>Model</b>       | A.5.7 Threat intelligence<br>A.8.8 Vulnerability mgmt<br>A.5.25 Event assessment                 | MAP 1.5<br>MAP 2.3<br>MEASURE 2.5 / GOVERN 4.2                                        | Cl. 6.1.2 Risk assessment<br>Cl. 8.4 Impact assessment<br>Annex B.6.2           |
| <b>Model Access Control &amp; Authentication</b>       | <b>Model</b>       | A.5.15 Access control<br>A.5.16 Identity mgmt<br>A.5.17 Authentication / A.8.2 Privileged access | GOVERN 1.2 Accountability<br>MANAGE 4.1 Incidents<br>MEASURE 1.1 Metrics              | Cl. 7.2 Competence & access<br>Cl. 8.5.3 Access controls<br>Annex B.6.5 / B.8.2 |
| <b>Secure Deployment &amp; Supply Chain</b>            | <b>Model</b>       | A.5.19 Supplier relations<br>A.8.30 Outsourced dev<br>A.8.25 Secure dev lifecycle                | GOVERN 6.2 Third-party policy<br>MAP 5.2 Deployment context<br>MANAGE 3.2 Change mgmt | Cl. 8.5.4 Deployment controls<br>Cl. 8.6.2 Third-party controls<br>Annex B.8.3  |
| <b>Prompt Injection &amp; Input Validation</b>         | <b>Application</b> | A.8.28 Secure coding<br>A.8.29 Security testing<br>A.8.27 System architecture                    | MANAGE 2.2<br>MEASURE 2.6<br>MAP 2.3                                                  | Cl. 8.5 Design controls<br>Cl. 9.1 Monitoring<br>Annex B.6.1                    |
| <b>Output Monitoring &amp; Anomaly Detection</b>       | <b>Operations</b>  | A.8.16 Monitoring<br>A.8.15 Logging<br>A.5.25 Event assessment                                   | MEASURE 2.7 Performance<br>MEASURE 2.9 Fairness<br>MANAGE 2.4 Change response         | Cl. 9.1 Monitoring<br>Cl. 9.2 Internal audit<br>Annex B.9.1 / B.9.2             |

*Use this crosswalk to design AISIF-aligned controls that are simultaneously auditable against whichever of the three standards your organization is accountable to.*

### 3. The Technical Framework Crosswalk

The table below extends the same ten AISIF control domains to three widely used technical references. OWASP LLM Top 10 categories and MITRE ATLAS tactics are risk and threat catalogues rather than control standards, so several Governance-layer domains have only a limited or indirect mapping — noted explicitly rather than forced into a false equivalence. The Databricks AI Security Framework (DASF) maps closely across every domain, reflecting its own broad, platform-oriented component model.

*Table 2. AISIF Control Domain Crosswalk — OWASP LLM Top 10 (2025), MITRE ATLAS, Databricks AI Security Framework.*

| Control Domain                                         | Layer             | OWASP LLM Top 10 (2025)                                                                 | MITRE ATLAS                                                                      | Databricks AI Security Framework                                                                  |
|--------------------------------------------------------|-------------------|-----------------------------------------------------------------------------------------|----------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------|
| <b>AI Risk Governance &amp; Lifecycle Management</b>   | <b>Governance</b> | Cross-cutting — informs the lifecycle guidance behind LLM03 Supply Chain                | Not tactic-specific — ATLAS case corpus informs the organizational risk register | Component 4: Data & AI Catalog Governance<br>12-component system inventory                        |
| <b>Shared Responsibility &amp; Ownership</b>           | <b>Governance</b> | LLM03 Supply Chain — vendor / third-party component accountability                      | ML Model Access — boundary of what is exposed to third-party integrations        | Per-control shared-responsibility guidance (vendor vs. customer) across all components            |
| <b>Regulatory Compliance &amp; Documentation</b>       | <b>Governance</b> | Limited direct mapping — OWASP is a technical risk catalogue, not a compliance standard | Not compliance-oriented — case studies support incident documentation for audits | Component 4 governance chapter; DASF's own crosswalk to EU AI Act and ISO 42001                   |
| <b>Data Governance &amp; Lineage</b>                   | <b>Data</b>       | LLM02 Sensitive Information Disclosure<br>LLM08 Vector & Embedding Weaknesses           | Reconnaissance — data-source exposure<br>Collection                              | Component 1: Raw Data<br>Component 3: Datasets<br>Component 4: Data Catalog Governance            |
| <b>Training Data Integrity &amp; Poisoning Defense</b> | <b>Data</b>       | LLM04 Data and Model Poisoning (direct match)                                           | ML Attack Staging — where poisoning is prepared<br>Resource Development          | Component 2: Data Prep<br>Component 5: Machine Learning Algorithms (training)                     |
| <b>Threat Modeling &amp; Attack Simulation</b>         | <b>Model</b>      | Cross-cutting — informs coverage for LLM01, LLM04, LLM06, and LLM10                     | The direct use case for ATLAS — all 16 tactics, Reconnaissance through Impact    | Component 6: Evaluation (adversarial / red-team evaluation)                                       |
| <b>Model Access Control &amp; Authentication</b>       | <b>Model</b>      | LLM07 System Prompt Leakage — auth boundary around prompts and APIs                     | ML Model Access (direct match)<br>Credential Access                              | Component 8: Model Management<br>Components 9–10: Model Serving and Inference Requests / Response |
| <b>Secure Deployment &amp; Supply Chain</b>            | <b>Model</b>      | LLM03 Supply Chain (direct match)                                                       | Resource Development<br>Initial Access — supply-chain compromise                 | Component 7: Machine Learning Models<br>Component 11: MLOps                                       |

| Control Domain                        | Layer       | OWASP LLM Top 10 (2025)                                                       | MITRE ATLAS                                                   | Databricks AI Security Framework                                                       |
|---------------------------------------|-------------|-------------------------------------------------------------------------------|---------------------------------------------------------------|----------------------------------------------------------------------------------------|
| Prompt Injection & Input Validation   | Application | LLM01 Prompt Injection (direct match)<br>LLM06 Excessive Agency / LLM07       | Initial Access — prompt injection as entry point<br>Execution | Component 9: Model Serving and Inference Requests<br>Agentic AI component (v3.0)       |
| Output Monitoring & Anomaly Detection | Operations  | LLM05 Improper Output Handling (direct match)<br>LLM09 Misinformation / LLM10 | Exfiltration<br>Impact<br>Defense Evasion                     | Component 10: Model Serving and Inference Response<br>Component 11: MLOps (monitoring) |

Read “limited direct mapping” as an honest gap, not an oversight: OWASP and ATLAS are scoped to technical risk and adversary technique, not organizational governance, so Governance-layer domains are better substantiated through Table 1. Use Table 2 primarily for the Data, Model, Application, and Operations layers, where all three references map directly.

## 4. How the Three Certification Standards Interrelate

A recurring question among practitioners reviewing this crosswalk is whether organizations already certified to ISO/IEC 27001:2022 must also pursue ISO/IEC 42001:2023, and how the NIST AI Risk Management Framework aligns with both. These three standards address distinct but complementary dimensions of the assurance landscape and are most effective when implemented together. (OWASP LLM Top 10, MITRE ATLAS, and the Databricks AI Security Framework, covered in Table 2, are technical references rather than certifiable standards, and are not part of this comparison.)

### ISO/IEC 27001:2022 — The Foundational Security Baseline

ISO/IEC 27001 serves as the foundational layer of security governance. Its Annex A controls establish a comprehensive framework for managing information security risks across data, systems, personnel, suppliers, and organizational processes. For AI initiatives, it provides the essential baseline upon which AI-specific controls can be meaningfully constructed. Organizations already certified possess an Information Security Management System (ISMS) that can function as the structural platform for integrating AISIF controls; for those not yet certified, AI security capability-building and ISO 27001 implementation should ideally proceed in parallel. The 2022 revision introduced several controls of particular relevance to AI, including A.5.7 (Threat intelligence), A.8.11 (Data masking), and A.5.19–A.5.23 (Supplier security). ISO 27001 alone is not sufficient for comprehensive AI security, as it does not explicitly address machine learning–specific threats, model lifecycle risks, or AI governance considerations.

### NIST AI Risk Management Framework — The AI Risk Governance Cycle

The NIST AI RMF operates at a higher level of abstraction, focused on the governance and lifecycle management of AI-related risk. Where ISO 27001 emphasizes the existence and implementation of controls, the AI RMF emphasizes the systematic identification, assessment, and management of AI-specific risk. Its four core functions align closely with AISIF:

- Govern establishes accountability and oversight structures within the Governance layer.
- Map defines system context and risk exposure, informing both the Data and Model layers.
- Measure evaluates the effectiveness of controls across all layers.
- Manage ensures risk responses are implemented and continuously refined through feedback loops between Operations and Governance.

The AI RMF does not replace ISO 27001 controls; it provides direction on where and why such controls should be applied. Organizations initiating an AI risk governance program should prioritize the Govern function before extending into broader framework adoption.

### ISO/IEC 42001:2023 — Artificial Intelligence Management System (AIMS)

ISO/IEC 42001 is the most AI-specific and most recent of the three. It defines requirements for establishing, implementing, and maintaining an Artificial Intelligence Management System (AIMS), enabling organizations to develop, deploy, and operate AI systems responsibly and under control. Structured to the High-Level Structure common to ISO management-system standards such as ISO 27001 and ISO 9001, it integrates into existing organizational governance rather than requiring a parallel system. In regulatory contexts — particularly the EU AI Act — ISO 42001 certification is increasingly recognized as a practical mechanism for demonstrating conformity. For organizations where regulatory alignment is a key driver, ISO 42001 should be treated as a strategic priority within the broader AI security program.

***ISO 27001 establishes the security foundation, the NIST AI RMF guides risk governance, and ISO 42001 formalizes AI-specific management practices. Together, they form a cohesive and mutually reinforcing framework for robust AI assurance.***

## Certification Path

For an organization already ISO 27001 certified: the existing ISMS is the structural platform. Use Table 1 (Section 2) to identify which AISIF domains your current Annex A controls already substantiate, close the AI-specific gaps domain by domain, then extend the management system's scope statement to cover the AIMS requirements in ISO 42001 — this is a scope extension, not a parallel certification effort, because the two standards share a High-Level Structure.

For an organization not yet certified to either standard: build AISIF-aligned controls against Table 1 from the outset, so the resulting evidence base is simultaneously usable for either certification track later. Prioritize the NIST AI RMF's Govern function first — accountability and oversight structures are the precondition for every other control in the crosswalk.

For an organization whose primary driver is EU AI Act conformity: ISO 42001 certification is the most direct demonstration mechanism currently recognized. Use the ISO 42001 column of Table 1 to sequence implementation, and treat the corresponding ISO 27001 and NIST AI RMF references as the technical and governance substantiation an auditor will expect behind each AIMS clause.

For a security team running day-to-day technical work: Table 2 (Section 3) is the more immediately useful reference. A red team engagement scoped to MITRE ATLAS, an AppSec review checked against the OWASP LLM Top 10, or a Databricks-hosted AI system assessed against DASF each produce findings tagged to that framework's own categories. Use Table 2 to translate each finding into the AISIF domain it belongs to, and Table 1 to see which certification-relevant standard that finding should also be logged against — so one piece of technical evidence supports governance reporting as well as the audit it was originally produced for.

## Sources

*AISIF Practitioner White Paper, v1.1, Section 5 (Tables 8 and 9). ISO/IEC 27001:2022 — Information security, cybersecurity and privacy protection — Information security management systems — Requirements. NIST AI Risk Management Framework (AI RMF 1.0). ISO/IEC 42001:2023 — Information technology — Artificial intelligence — Management system.*

*OWASP GenAI Security Project — Top 10 for LLM Applications 2025 (v2.0), published November 2024. MITRE ATLAS (Adversarial Threat Landscape for Artificial-Intelligence Systems) — current tactic and technique matrix, [atlas.mitre.org](https://atlas.mitre.org). Databricks — AI Security Framework (DASF), 12 canonical AI system components (v2.0) with a 13th Agentic AI component added in v3.0.*

*Licensed CC BY 4.0 — free to use, share, and adapt with attribution. Part of the AISIF toolkit at [aisif.ai](https://aisif.ai).*