

The AISIF Detailed Crosswalk

ISO/IEC 27001:2022 · NIST AI Risk Management Framework · ISO/IEC 42001:2023

Excerpted and expanded from the AISIF Practitioner White Paper v1.1, Section 5 · Companion resource to the AISIF toolkit · Prepared August 2026

Purpose and How to Use This Document

This document is the complete version of the AISIF standards crosswalk summarized in the practitioner white paper and referenced throughout the AISIF toolkit. It maps all ten AISIF control domains — organized under the framework’s five layers — to specific controls in ISO/IEC 27001:2022 Annex A, functions and categories of the NIST AI Risk Management Framework, and clauses of ISO/IEC 42001:2023.

Use it in three ways. First, as a control-design reference: when building or auditing a control for a given AISIF domain, the crosswalk names the specific external-standard clauses that control should satisfy, so the same piece of evidence can support more than one audit. Second, as a gap-assessment tool: an organization already certified to one of the three standards can use the relevant column to see exactly which AISIF domains its existing certification 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 Information Security Management System (ISMS) can be extended into an AI Management System (AIMS) using this crosswalk as the domain-by-domain map between the two.

Section 1 lists the ten control domains with their full descriptions. Section 2 presents the complete crosswalk table. Section 3 explains how the three standards relate to one another and to AISIF, and closes with certification guidance. All content is drawn directly from the AISIF Practitioner White Paper (v1.1, Section 5); 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
AI Risk Governance & Lifecycle Management	Governance	A.5.1 Security policies A.5.4 Management responsibilities A.5.8 Project security	Full GOVERN (1.1–1.7) Full MAP (1.1–5.2) Full MANAGE (1.1–4.2)	Cl. 4–10 Full AIMS Cl. 6.1–6.3 Planning Annex A.1–A.13
Shared Responsibility & Ownership	Governance	A.5.2 Security roles A.5.31 Legal requirements A.5.19 Supplier relations	GOVERN 1.1 Policies GOVERN 6.1 Third-party MAP 1.6 Org context	Cl. 4.3 AIMS scope Cl. 5.1 Leadership Cl. 8.6 Supplier relations
Regulatory Compliance & Documentation	Governance	A.5.31 Legal requirements A.5.32 IP rights A.5.36 Policy compliance	GOVERN 1.7 Regulatory alignment GOVERN 4.1 Risk policies MAP 1.1	Cl. 4.2 Interested parties Cl. 6.1.1 Risks Cl. 10.1 Corrective action
Data Governance & Lineage	Data	A.5.12 Classification A.5.13 Labelling A.8.10 Deletion / A.5.33 Records	MAP 3.5 Provenance GOVERN 1.4 Data policy MEASURE 2.8 Bias eval	Cl. 8.3 Data management Cl. 6.1.3 Quality & lineage Annex A.6 / B.7.4
Training Data Integrity & Poisoning Defense	Data	A.8.9 Config management A.8.29 Security testing A.8.7 Anti-malware	MAP 3.5 Training provenance MEASURE 2.5 MANAGE 1.3	Cl. 8.3.2 Training controls Cl. 8.4 Impact assessment Annex A.6.2 / B.7.5
Threat Modeling & Attack Simulation	Model	A.5.7 Threat intelligence A.8.8 Vulnerability mgmt A.5.25 Event assessment	MAP 1.5 MAP 2.3 MEASURE 2.5 / GOVERN 4.2	Cl. 6.1.2 Risk assessment Cl. 8.4 Impact assessment Annex B.6.2
Model Access Control & Authentication	Model	A.5.15 Access control A.5.16 Identity mgmt A.5.17 Authentication / A.8.2 Privileged access	GOVERN 1.2 Accountability MANAGE 4.1 Incidents MEASURE 1.1 Metrics	Cl. 7.2 Competence & access Cl. 8.5.3 Access controls Annex B.6.5 / B.8.2
Secure Deployment & Supply Chain	Model	A.5.19 Supplier relations A.8.30 Outsourced dev A.8.25 Secure dev lifecycle	GOVERN 6.2 Third-party policy MAP 5.2 Deployment context MANAGE 3.2 Change mgmt	Cl. 8.5.4 Deployment controls Cl. 8.6.2 Third-party controls Annex B.8.3
Prompt Injection & Input Validation	Application	A.8.28 Secure coding A.8.29 Security testing A.8.27 System architecture	MANAGE 2.2 MEASURE 2.6 MAP 2.3	Cl. 8.5 Design controls Cl. 9.1 Monitoring Annex B.6.1
Output Monitoring & Anomaly Detection	Operations	A.8.16 Monitoring A.8.15 Logging A.5.25 Event assessment	MEASURE 2.7 Performance MEASURE 2.9 Fairness MANAGE 2.4 Change response	Cl. 9.1 Monitoring Cl. 9.2 Internal audit 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. How the Three 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.

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 Section 2’s crosswalk 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 this crosswalk 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 the crosswalk 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.

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.

Licensed CC BY 4.0 — free to use, share, and adapt with attribution. Part of the AISIF toolkit at aisif.ai.