

The Enterprise AI Governance Crosswalk

Mapping the EU AI Act, NIST AI RMF and the ISO/IEC 42001 Family into One Operational System

A CXO Intelligence Governance Series Reference Companion guidance to Article 4 of the CXO Intelligence Governance Series, Part 4, "Three Frameworks, One System: The Convergence Playbook for AI Governance" Leadership in the Age of AI

This guide is a practical implementation reference. It does not reproduce the text of any standard. All clause and article references are pointers; consult the source documents for normative wording. Regulatory timelines and standard statuses reflect mid-2026 and should be re-verified before formal use, this field is moving quickly, and one central date changed materially in the months before publication.

PART 1, EXECUTIVE SUMMARY

Every enterprise deploying AI now faces the same structural problem: three major governance instruments have arrived at once, they overlap heavily, they use different vocabularies, and none of them alone is sufficient. Organizations respond by running three parallel programs, a legal team reading the EU AI Act, a risk team adopting the NIST AI Risk Management Framework, and a compliance team pursuing ISO/IEC 42001 certification, that duplicate effort, generate conflicting artifacts, and still leave gaps.

This is a waste, and it is avoidable. The three instruments are not competitors. They are three views of the same underlying reality, pitched at three different altitudes:

- **The EU AI Act** is law. It tells you *what you must do* to place AI on the European market without penalty. It is mandatory, risk-tiered, and enforceable with fines up to €35 million or 7% of global turnover.
- **ISO/IEC 42001** is a management system standard. It tells you *how to build the organizational machinery* that produces responsible AI repeatably. It is certifiable, auditable, and voluntary.
- **The NIST AI RMF** is a risk framework. It tells you *how to reason about and treat AI risk* in a structured, technically grounded way. It is voluntary, flexible, and outcome-oriented.

Put plainly: **NIST helps you think, ISO helps you operate, and the EU AI Act tells you what you are legally required to prove.** An enterprise that treats them as one integrated stack, a single set of controls generating a single body of evidence, mapped once to all three, spends a fraction of the effort and ends up with a stronger governance posture than one running three disconnected programs.

That integration is the purpose of this guide. The core deliverable is the **Crosswalk Matrix** in Part 5: a single table where each row is a governance capability your enterprise needs, and the columns show how each instrument addresses it, what evidence satisfies all three at once, who owns it, and how to prioritize it. Around that matrix, this

guide provides an operating model, a maturity model, a phased roadmap, a board dashboard, an original operating-stack framework, and a decision tool, the **AI Governance Navigator**, that tells any given enterprise which obligations actually apply to it.

One timing note that reframes everything for European deployment. As this guide went to press, the EU formally simplified the AI Act's timeline. The heaviest high-risk obligations, originally due 2 August 2026, were deferred. This does not reduce the work; it changes its sequencing, and it makes the case for building on the durable ISO and NIST machinery now, rather than racing a legal deadline, stronger than ever. Part 3 covers this in detail.

The enterprises that win the next two years will not be those that comply fastest with one instrument. They will be those that build governance machinery once, generate evidence once, and satisfy all three, and every instrument that follows, from the same operational foundation.

PART 2, WHY ENTERPRISES STRUGGLE WITH AI GOVERNANCE

Most AI governance programs underperform for reasons that have little to do with the technology and everything to do with operating design. Six failure patterns recur across enterprises of every size and sector.

2.1 The parallel-programs trap

The most common and most expensive mistake. Legal owns the EU AI Act, risk owns NIST, compliance owns ISO, and security owns the technical controls, and none of them share a data model. The same AI system is inventoried four times in four formats. The same risk is assessed under four methodologies. When an auditor, a regulator, or the board asks a question, four teams produce four partial, inconsistent answers. The fix is not more coordination meetings; it is a single control framework that maps to all instruments at once, which is what Part 5 provides.

2.2 No operating model

Enterprises write AI policies and then discover nobody owns their enforcement. A policy that says "high-risk AI must undergo human oversight review" is worthless if no named role is accountable for running that review, no artifact records it, and no trigger initiates it. Governance fails at the seam between policy and operation, the layer where a written commitment becomes a named owner, a defined trigger, a produced artifact, and a retained piece of evidence. Most programs have the policy layer and the technology layer, and nothing connecting them. Part 9 (Operating Model) and Part 13 (the Operating Stack) exist to fill this seam.

2.3 The missing inventory

You cannot govern what you have not inventoried, and most enterprises do not know how many AI systems they run. Shadow AI, models embedded in SaaS tools, personal-account usage, agents spun up by individual teams, third-party features that quietly became AI, means the real inventory is always larger than the official one. Every downstream governance activity (risk assessment, classification, monitoring, evidence) depends on an inventory that is complete and living. This is the single most common root cause of governance failure, and it is why every roadmap in this guide starts there.

2.4 Treating governance as documentation

A governance program that produces binders rather than behavior change is theatre. The test of a control is not whether a document exists but whether the document is generated as a byproduct of an actual operational activity, is current, has an owner, and could be produced on demand under audit. Evidence built as an after-the-fact compliance exercise is fragile, stale, and unconvincing. Evidence built into the operating process is durable. This distinction, between documentation and operational evidence, runs through the entire guide.

2.5 Static governance for dynamic systems

AI systems are not static software. Models drift, are retrained, are swapped for newer versions, and change behavior as their inputs change. A governance program that assesses a system once at deployment and never revisits it is governing a system that no longer exists. Effective governance is continuous: it defines triggers for reassessment (model change, scope expansion, performance degradation, incident) and treats monitoring as a first-class governance activity, not an operational afterthought.

2.6 Governance disconnected from the board

The board is accountable for AI risk, but most boards receive either nothing or an unreadable technical report. When the board cannot see AI risk in terms it can act on, exposure, incidents, compliance posture, resource adequacy, it cannot govern, and accountability breaks. A functioning program translates operational governance into a small set of executive metrics the board actually uses. Part 7 (Board Dashboard) addresses this directly.

The through-line: every one of these failures is an operating-model failure, not a knowledge failure. Enterprises do not struggle because they lack access to the standards. They struggle because they have not converted the standards into an operating system, owners, triggers, artifacts, evidence, and a rhythm of review. The rest of this guide is that conversion.

PART 3, THE GOVERNANCE LANDSCAPE

To integrate the three instruments, you first have to understand precisely what each one is, what it does well, and, critically, what it does not solve. Confusion between them is the source of most wasted effort.

3.1 The EU AI Act

What it is. The world's first comprehensive horizontal AI regulation (Regulation (EU) 2024/1689). It entered into force on 1 August 2024 and applies to any organization placing AI systems on the EU market or whose AI output is used in the EU, regardless of where the organization is based. It is binding law, enforced by national competent authorities and, for general-purpose AI models, the European AI Office.

How it works. The Act classifies AI systems into four risk tiers and attaches obligations to each:

- **Unacceptable risk**, a small set of prohibited practices (e.g. social scoring, certain manipulative or exploitative systems). Banned outright since February 2025.
- **High risk**, systems in sensitive domains (biometrics, critical infrastructure, education, employment, essential services, law enforcement, migration, justice) and AI that is a safety component of a regulated product. These

carry the heavy obligations: risk management, data governance, technical documentation, logging, transparency, human oversight, accuracy, robustness, and cybersecurity, plus conformity assessment and registration.

- **Limited risk**, transparency obligations (e.g. disclosing that content is AI-generated, that a user is interacting with a chatbot).
- **Minimal risk**, the majority of AI; no specific obligations.

General-purpose AI (GPAI) models sit on a parallel track with their own transparency, documentation, and, for models presenting systemic risk, additional obligations.

The timeline, and the change that matters. The Act phases in over several years. Two milestones have passed: prohibitions (February 2025) and GPAI obligations (August 2025). The pivotal milestone for most enterprises was the high-risk regime, originally set for **2 August 2026**.

That date moved. Through 2025 the implementation was visibly off track, national authorities were not designated in time and the harmonized standards providers need to demonstrate conformity were not ready. In response, the EU brought forward a simplification package (the "Digital Omnibus"). A political agreement was reached on 7 May 2026, and the Council gave final approval on 29 June 2026. Under the adopted package:

- High-risk obligations for **stand-alone Annex III systems** (recruitment, credit scoring, education, law enforcement, border control, and similar) now apply from **2 December 2027**.
- High-risk obligations for **AI embedded in regulated products under Annex I** (medical devices, machinery, and similar) apply from **2 August 2028**.
- **Article 50 transparency obligations** and the **Article 4 AI-literacy duty** remain on their original timeline and apply from **2 August 2026**, these were *not* deferred.
- New prohibitions (including AI-generated non-consensual intimate imagery and CSAM) were added, with certain provisions from **2 December 2026**.

The strategic reading: **the delay is real but narrow, and it is a deferral, not a dismantling**. The Act's architecture, risk tiers, conformity assessment, the GPAI track, the AI Office's authority, is fully intact. Transparency and literacy obligations are live in 2026 regardless. And the extra time should be spent building the durable governance machinery (via ISO and NIST) that makes high-risk compliance a byproduct rather than a scramble. Enterprises that treat the deferral as permission to stop have misread it; the backstop is fixed, the heaviest regime is coming, and the work is substantial.

Who should implement it. Mandatory for any enterprise with AI touching the EU market. Even outside the EU, it is becoming a de facto baseline because global enterprises standardize on the strictest applicable regime.

What it does NOT solve. The Act tells you *what* legal outcomes you must achieve, not *how* to build the organization that achieves them. It specifies obligations, not operating models. It does not tell you how to structure accountability, how to run a risk process day to day, or how to build repeatable machinery. It is a compliance target, not an implementation method, which is exactly the gap ISO and NIST fill.

3.2 ISO/IEC 42001

What it is. The first international management system standard for artificial intelligence, published in December 2023. It defines requirements for an **AI Management System (AIMS)**, the organizational structure, policies, processes, and controls through which an enterprise governs AI responsibly and repeatably. It is certifiable: an accredited body can audit an organization and issue a certificate, providing third-party assurance.

How it works. Like all modern ISO management standards, 42001 follows the common "Annex SL" high-level structure and runs on a Plan-Do-Check-Act cycle. Its requirements sit in **Clauses 4 through 10**:

- **Clause 4, Context.** Understand the organization, its internal and external issues, interested parties and their expectations, and define the AIMS scope.
- **Clause 5, Leadership.** Top management must demonstrate commitment, set an AI policy, and assign roles and responsibilities.
- **Clause 6, Planning.** Identify and address AI risks and opportunities, conduct AI risk assessments, run AI system impact assessments, and set objectives.
- **Clause 7, Support.** Provide resources, ensure competence, drive awareness, manage communication, and control documented information.
- **Clause 8, Operation.** Execute the operational controls, the AIMS in action across the AI lifecycle.
- **Clause 9, Performance evaluation.** Monitor, measure, analyze, audit internally, and conduct management reviews.
- **Clause 10, Improvement.** Handle nonconformities, take corrective action, and improve continually.

Alongside the clauses, **Annex A provides a reference set of 38 controls organized into 9 groups**, spanning AI policy, internal organization, resourcing, AI impact assessment, the AI system lifecycle, data for AI, information provided to stakeholders, responsible use of AI, and third-party and supplier relationships. Crucially, these controls are **not applied wholesale**: an organization performs its risk assessment, then selects applicable controls and documents its choices, inclusions and exclusions with justifications, in a **Statement of Applicability (SoA)**. The controls are deliberately high-level and principle-based rather than prescriptive technical steps, and Annex B offers informative implementation guidance that auditors nonetheless reference.

Who should implement it. Any enterprise that wants durable, auditable AI governance machinery and, optionally, a certificate that signals it to customers, regulators, and partners. It is especially valuable where AI is central to the product or where third-party assurance carries commercial weight. Because it reuses the management-system pattern of ISO 27001, organizations with an existing ISMS can extend rather than rebuild.

What it does NOT solve. ISO 42001 is deliberately non-prescriptive. It tells you to have a risk process, but not which specific technical tests to run. It tells you to manage the lifecycle, but not how to red-team an LLM or detect drift. It provides the management shell; it does not provide the technical depth (which NIST, OWASP, and MITRE supply) or the specific legal obligations (which the EU AI Act supplies). A certificate proves you have a functioning system, not that you meet any particular law.

3.3 The NIST AI Risk Management Framework

What it is. A voluntary framework published by the U.S. National Institute of Standards and Technology in January 2023, with a companion Generative AI Profile added in July 2024. It provides a structured, technically grounded way to identify, assess, and manage AI risk across the lifecycle. It is not certifiable and carries no legal force; its authority is intellectual and practical.

How it works. The framework is built on four functions that operate continuously and iteratively:

- **Govern**, the cross-cutting function: cultivate a culture of risk management, define accountability structures, policies, and processes. Govern underpins the other three.
- **Map**, establish context and frame risk: understand the system's purpose, its context of use, its stakeholders, and catalog what could go wrong. This is where the AI inventory and system framing live.

- **Measure**, analyze, assess, benchmark, and monitor risk using quantitative and qualitative methods: test for bias, robustness, security, explainability, and track performance over time.
- **Manage**, allocate resources to treat risks: prioritize, respond, recover, and communicate. This is where risk treatment decisions are made and acted on.

The framework is intentionally flexible, it adapts to any sector, organization size, and AI type, and it is outcome-oriented rather than checklist-driven. Its companion profiles (notably for generative AI) tailor the general framework to specific technology risks.

Who should implement it. Any enterprise that wants a rigorous, respected method for reasoning about AI risk, particularly useful for technical teams, for organizations without an EU nexus that still want strong governance, and as the analytical engine underneath an ISO management system or an EU AI Act compliance effort. It pairs naturally with both.

What it does NOT solve. NIST is a framework, not a management system and not a law. It will not certify you, and it will not, by itself, make you compliant with any regulation. It tells you how to think about risk but does not impose the organizational discipline of a management system (no mandatory documented information, no internal audit requirement, no management review cadence) or the enforceable obligations of a regulation. It is the reasoning layer; it needs ISO for operational discipline and the EU AI Act for legal grounding.

3.4 The complementary layer: Singapore's Model AI Governance Framework

The three core instruments are not the whole landscape. A fourth reference deserves a place in any serious enterprise program, not as a competitor to the EU AI Act, ISO, or NIST, but as a pragmatic, principle-level complement that many multinationals now use to harmonize governance across Asia-Pacific operations and to reason about generative and agentic AI specifically.

What it is. Singapore's Model AI Governance Framework, published by the Infocomm Media Development Authority (IMDA) and the AI Verify Foundation. It has evolved in three layers: the original Model AI Governance Framework for traditional AI (2019, revised 2020); the **Model AI Governance Framework for Generative AI** (30 May 2024); and, notably, the **Model AI Governance Framework for Agentic AI** (January 2026), described as the world's first governance framework aimed specifically at autonomous, goal-pursuing AI agents.

Why it matters for enterprises. It is voluntary, non-prescriptive, and deliberately practical. It carries no legal force and offers no certification, but it is influential, internationally referenced (including through the OECD), and unusually clear about generative and agentic risk. For a global enterprise, it functions as a principle-level input that maps cleanly onto the operating machinery ISO provides and the risk reasoning NIST provides.

The Generative AI framework's nine dimensions. The framework asks organizations to consider nine dimensions in totality: **accountability** (incentive structures across the AI value chain, including shared-responsibility models); **data** (quality and handling of contentious training data); **trusted development and deployment** (baseline safety and transparency in the build); **incident reporting** (structures for detection and response); **testing and assurance** (third-party validation); **security** (adapting security practices to AI); **content provenance** (watermarking and cryptographic provenance for AI-generated content); **safety and alignment** research and practice; and **AI for public good** (democratizing access, upskilling, sustainable development). These map onto the Crosswalk capabilities rather than replacing them, data governance, incident management, testing, security, and transparency all have direct homes in the matrix.

The Agentic AI framework's four dimensions. For autonomous agents, Singapore recommends four core dimensions: **assessing and bounding risks upfront**; **ensuring meaningful human accountability and oversight**; **implementing robust technical controls and processes across the agent lifecycle**; and **enabling end-user responsibility through transparency and training**. These align directly with the agentic controls this guide specifies (per-agent identity, least-privilege scoping, human-approval gates, containment) and reinforce that human oversight is an integral, non-delegable part of agentic governance.

How to use it. Treat Singapore's framework as a principle-level lens that sits alongside NIST in the reasoning layer, especially valuable for enterprises with Asia-Pacific exposure, for generative-AI-heavy portfolios, and for agentic deployments where it offers some of the clearest available guidance. It does not add a separate compliance burden; it enriches how you reason about risk and, for agentic systems, provides a governance vocabulary that the binding instruments have not yet fully articulated.

3.5 How they fit together

The core instruments are complementary by design, each operating at a different altitude, with Singapore's framework enriching the reasoning layer:

- **NIST AI RMF** is the **reasoning layer**, how you think about and treat risk.
- **ISO/IEC 42001** is the **operating layer**, how you build repeatable machinery and prove it works.
- **The EU AI Act** is the **obligation layer**, what you are legally required to achieve and evidence.

A mature enterprise uses NIST's Map and Measure functions, enriched by Singapore's generative and agentic dimensions where relevant, to run the risk and impact assessments that ISO Clause 6 requires, operates them through an ISO-structured management system, and configures the whole apparatus to produce exactly the evidence the EU AI Act's high-risk regime demands. One set of activities. One body of evidence. Multiple instruments satisfied. That is the integration this guide operationalizes, and the supporting frameworks (OECD principles, ISO 23894 and ISO 31000 for risk, Singapore's Model AI Governance Framework for principle-level and agentic guidance, OWASP and MITRE ATLAS for technical depth, and the major vendor frameworks) slot into this same structure as either principle-level inputs or technical-control detail.

PART 3B, THE ISO/IEC AI FAMILY: A CLAUSE-LEVEL MAPPING

ISO/IEC 42001 does not stand alone. It anchors a growing family of AI standards, and a serious enterprise program uses the family, not just the flagship, both because the companion standards fill specific gaps 42001 leaves open, and because the harmonized standards that will ultimately underpin EU AI Act conformity are drawn from this body of work.

A note on copyright and how to read this section. ISO/IEC standards are copyrighted commercial products. Nothing below reproduces any requirement, clause, annex, table, or control text. Every reference is by standard number, clause number, and a high-level statement of *intent* in original words. To implement against any of these standards, obtain a licensed copy and work from its normative text. This section is a navigational map, not a substitute for the standards.

3B.1 ISO/IEC 42001:2023, the management-system spine

PUBLISHED, certifiable

The clause structure, by intent:

- **Clause 4, Context.** Establish what the organization is, the internal and external issues bearing on its AI, the interested parties and their expectations, and the scope of the AI management system.
- **Clause 5, Leadership.** Top management demonstrates commitment, sets an AI policy, and assigns roles, responsibilities, and authorities.
- **Clause 6, Planning.** Identify and address AI-related risks and opportunities; conduct AI risk assessments and AI system impact assessments; set objectives and plan to meet them.
- **Clause 7, Support.** Provide the resources, competence, awareness, communication, and controlled documented information the system needs.
- **Clause 8, Operation.** Plan, implement, and control the operational processes, the management system in action across the AI lifecycle.
- **Clause 9, Performance evaluation.** Monitor, measure, analyze, and evaluate; conduct internal audit and management review.
- **Clause 10, Improvement.** Handle nonconformity, take corrective action, and improve the system continually.

Annex A, by intent, provides a reference set of 38 controls organized into nine control groups. The groups address, at a high level: AI policy; internal organization; resources for AI systems; assessing the impact of AI systems; the AI system lifecycle; data for AI systems; information provided to interested parties; responsible use of AI; and third-party and supplier relationships. Controls are **selected through risk treatment, not applied wholesale**, and the selection, inclusions and exclusions with justifications, is documented in a **Statement of Applicability**. Annex B provides informative implementation guidance that auditors nonetheless reference.

3B.2 The companion standards, mapped to the gaps they fill

Standard	Status	Intent (original words)	Gap it fills in a converged program	Where it maps
ISO/IEC 42005:2025	PUBLISHED, guidance	How to conduct an AI system impact assessment across the lifecycle	42001 Clause 6 requires impact assessment but does not detail the method; 42005 supplies it	Directly supports the EU AI Act fundamental-rights impact assessment (Art. 27) and Crosswalk rows 3, 14
ISO/IEC 42006:2025	PUBLISHED	Requirements for bodies that audit and certify AI management systems	Tells you what a certification body must do, hence what to prepare for	Audit & Certification layer; Crosswalk row 29
ISO/IEC 23894:2023	PUBLISHED, guidance	Guidance on managing AI-related risk, aligned to the ISO 31000 risk foundation	The risk-process detail beneath 42001 Clause 6	Pairs with NIST Map + Measure; Crosswalk row 2
ISO 31000:2018	PUBLISHED	General risk-management principles and process	The enterprise-risk foundation AI risk should integrate with	Underpins the Risk layer of the Operating Stack
ISO/IEC 5338:2023	PUBLISHED	AI system lifecycle processes	Defines the lifecycle stages governance gates attach to	Operations & Lifecycle domain; Crosswalk rows 8–11, 19–20
ISO/IEC 38507:2022	PUBLISHED	Governance implications of the use of AI for the governing body	The board-level accountability view above the management system	Governance domain; Crosswalk rows 23, 28
ISO/IEC 12792:2025	PUBLISHED	Transparency taxonomy for AI systems	Structures what transparency means and to whom	Transparency & Explainability domain; Crosswalk rows 7, 16
ISO/IEC TR 24028:2020	PUBLISHED, TR	Overview of trustworthiness in AI	Shared vocabulary for trustworthiness across the program	Glossary and framing
ISO/IEC TR 24368:2022	PUBLISHED, TR	Overview of ethical and societal concerns	Framing for ethically sensitive use review	Bias, Fairness & Ethics domain
ISO/IEC 5339:2024	PUBLISHED, guidance	Guidance for AI applications	Stakeholder and application-lifecycle framing	Complements Clause 4 context work
ISO/IEC 27001:2022	PUBLISHED, certifiable	Information security management system requirements	Access control, logging, incident response, supplier management, reuse, don't rebuild	Security, Logging, Vendor domains throughout
ISO/IEC 27701:2019	PUBLISHED	Privacy information management extension to 27001	Privacy controls and data-subject handling for AI data	Privacy; Crosswalk row 15
ISO/IEC 22989:2022	PUBLISHED	AI concepts and terminology	Canonical terms so the program speaks one language	Glossary

3B.3 The reuse principle, stated plainly

The single most valuable insight in this mapping is that **an enterprise already holding ISO/IEC 27001 is not starting AI governance from zero**. A large share of what AI governance requires, access management, logging and traceability, incident response, supplier due diligence, change control, business continuity, already exists in a mature

ISMS and can be *extended* to cover AI rather than rebuilt. The converged program treats 27001 as a foundation and 42001 as the AI-specific layer above it, mapping existing controls forward and building only the genuinely new AI-specific controls (impact assessment, bias, drift, explainability, agentic autonomy) from scratch. This is both faster and cheaper, and it is why the sequencing in Part 10 begins with mapping what you already have.

PART 4, ONE-PAGE COMPARISON

Dimension	EU AI Act	NIST AI RMF	ISO/IEC 42001
Purpose	Regulate AI to protect health, safety, fundamental rights	Provide a structured method to manage AI risk	Define requirements for an AI management system
Type	Binding law / regulation	Voluntary framework	Certifiable management system standard
Audience	Providers, deployers, importers touching the EU market	Any organization managing AI risk	Any organization governing AI operations
Mandatory?	Yes, if AI touches the EU market	No	No (but certification is a market signal)
Scope	AI systems and GPAI models on the EU market	All AI, all sectors, all lifecycle stages	The organization's AI management system
Risk focus	Risk-tiered obligations (unacceptable / high / limited / minimal)	Risk across the lifecycle via Govern-Map-Measure-Manage	AI-specific risks treated through a management system
Lifecycle	Design through post-market monitoring	Full lifecycle, continuous and iterative	Full lifecycle via Plan-Do-Check-Act
Controls	Legal obligations per risk tier	Outcome categories and subcategories	38 reference controls (Annex A), selected via risk
Evidence	Technical documentation, logs, conformity records	Documented risk decisions (organization-defined)	Documented information, records, Statement of Applicability
Certification	Conformity assessment + CE marking (high-risk)	None	Third-party accredited certification
Assessment	Conformity assessment (self or notified body)	Self-assessment against the framework	Internal audit + external certification audit
Governance	Mandated accountability, human oversight	Govern function; culture and accountability	Leadership, roles, management review (Clauses 5, 9)
Technical depth	Requirements stated, methods not prescribed	Strong (esp. with GenAI profile); method-rich	Management-level; technical method not prescribed
Operational discipline	Obligation-driven	Framework-driven, flexible	High, mandatory processes, audit, review
Enforcement	Fines up to €35M or 7% of global turnover	None	Loss of certification
Best suited for	Any enterprise with EU market exposure	Enterprises wanting rigorous risk reasoning	Enterprises wanting auditable, certifiable governance
Key limitation	Says what, not how	No certification, no legal force	No legal force, no prescribed technical depth

PART 5, THE CROSSWALK MATRIX

This is the core deliverable. Each row is a governance capability every enterprise needs. The columns show where each instrument addresses it, what evidence satisfies all three simultaneously, who typically owns it, the key artifacts, an implementation priority, and practical notes.

How to read the references. EU AI Act references cite article numbers by intent. NIST references cite the four functions (GV=Govern, MP=Map, MS=Measure, MG=Manage). ISO references cite clause numbers (4–10) and Annex A control groups (A.2–A.10) by intent, no control text is reproduced. **Priority** is P1 (foundational, do first), P2 (core), P3 (maturity).

A machine-readable CSV of this matrix accompanies the guide.

#	Capability	EU AI Act (intent)	NIST RMF	ISO 42001 (clause/control intent)	Evidence that satisfies all three	Typical owner	Key artifacts	Priority
1	AI System Inventory	Registration of high-risk systems (Art. 49, 71)	MP	Clause 4 scope; A.2/A.6 organizational asset control	Living inventory with per-system risk tier, owner, status	AI Governance Lead	Inventory register, AIBOM	P1
2	Risk Assessment	Risk management system for high-risk (Art. 9)	MP, MS	Clause 6.1 risk; A.5 impact-related controls	Documented risk assessment per system, refreshed on triggers	Risk (CRO)	Risk register, assessment records	P1
3	AI Impact Assessment	Fundamental-rights impact assessment (Art. 27)	MP	Clause 6 planning; A.5 impact assessment controls	Impact assessment covering rights, safety, affected groups	Risk / Legal	Impact assessment report	P1
4	Risk Classification	Risk-tier determination (Art. 6, Annex III)	MP	Clause 6.1; risk-based control selection	Documented tier rationale per system	Legal + Risk	Classification record	P1
5	Data Governance	Data quality & governance for high-risk (Art. 10)	MP, MS	A.7 data-for-AI controls	Data lineage, quality checks, provenance, minimization records	CDO	Data governance policy, lineage records	P1
6	Human Oversight	Human oversight design & operation (Art. 14)	GV, MG	Clause 5; A.9 responsible-use controls	Oversight design per system; evidence gates are exercised	Business Owner	Oversight design + review logs	P1
7	Transparency to Users	Transparency obligations (Art. 13, 50)	MP, MG	A.8 information-to-stakeholders controls	Disclosures, AI-interaction notices, content labelling	Product	Transparency notices, labels	P2
8	Technical Documentation	Technical documentation for high-risk (Art. 11, Annex IV)	MP, MS	Clause 7.5 documented information	Complete, current technical file per high-risk system	Engineering	Technical documentation set	P2
9	Logging & Traceability	Automatic record-keeping / logs (Art. 12, 19)	MS, MG	Clause 8; A.6 operational controls	Immutable logs: user, model+version, I/O, tool calls, timestamps	Engineering / Security	Logging spec, retained logs	P1
10	Monitoring & Post-Market	Post-market monitoring (Art. 72)	MS, MG	Clause 9 performance evaluation	Continuous monitoring with defined metrics and thresholds	Engineering / Ops	Monitoring dashboards, reports	P2
11	Model Validation & Testing	Accuracy, robustness (Art. 15)	MS	A.6 lifecycle controls	Validation results, test coverage, acceptance criteria	Engineering / Data Science	Test reports, validation records	P2
12	Incident Management	Serious-incident reporting (Art. 73)	MG	Clause 10 improvement; A. 6 controls	AI-specific incident process, records, notifications	Security / Risk	IR playbooks, incident log	P1
13	Security (AI-specific)	Cybersecurity for high-risk (Art. 15)	MS, MG	A.6 lifecycle; general security controls	Adversarial testing, guardrail validation, security evidence	CISO	Security test results	P1

#	Capability	EU AI Act (intent)	NIST RMF	ISO 42001 (clause/control intent)	Evidence that satisfies all three	Typical owner	Key artifacts	Priority
14	Bias & Fairness	Non-discrimination via data & testing (Art. 10, 15)	MS	A.5 impact; A.6 lifecycle	Bias testing across groups, mitigation records	Data Science / Risk	Fairness assessment	P2
15	Privacy	Interaction with GDPR; data governance (Art. 10)	MP, MS	A.7 data controls; link to ISO 27701	DPIA where applicable, privacy controls, data-subject handling	DPO	DPIA, privacy records	P1
16	Explainability	Interpretability supporting oversight (Art. 13, 14)	MS	A.8, A.9 controls	Explainability approach appropriate to use and audience	Data Science	Explainability documentation	P3
17	Third-Party / Vendor AI	Provider-deployer obligations along the chain (Art. 25)	GV, MP	A.10 third-party & supplier controls	Vendor assessments, contractual controls, provenance	Procurement + Risk	Vendor assessment records	P2
18	Supplier Management	Responsibilities across the value chain (Art. 25)	GV	A.10 supplier controls	Supplier due diligence, ongoing monitoring	Procurement	Supplier register, contracts	P2
19	Model / System Retirement	Lifecycle end handling	MG	A.6 lifecycle controls	Decommissioning process, data handling, records	Engineering	Retirement checklist	P3
20	Drift Detection	Ongoing accuracy under monitoring (Art. 15, 72)	MS	Clause 9; A.6 controls	Drift metrics, alert thresholds, retraining triggers	Data Science / Ops	Drift monitoring reports	P2
21	Red Teaming	Robustness & security testing (Art. 15)	MS	A.6 lifecycle controls	Adversarial test campaigns, findings, remediation	Security	Red team reports	P2
22	Testing & QA	Accuracy & robustness (Art. 15)	MS	A.6 lifecycle controls	Test strategy, coverage, results, gates	Engineering / QA	Test plans and results	P2
23	Board Reporting	Governance & accountability (organizational)	GV	Clause 5 leadership; Clause 9 review	Regular board-level AI risk and compliance reporting	AI Governance Lead	Board dashboard, minutes	P2
24	Training & Awareness	AI literacy duty (Art. 4)	GV	Clause 7.2/7.3 competence & awareness	Role-based training records; literacy program	HR + AI Gov	Training records, curriculum	P1
25	Competence	AI literacy for relevant staff (Art. 4)	GV	Clause 7.2 competence	Competence definitions per role, evidence held	HR	Competence matrix	P2
26	Business Continuity	Robustness / resilience (Art. 15)	MG	A.6 controls; link to ISO 22301	Continuity plans covering AI system failure	Ops / Risk	BCP for AI systems	P3
27	Policy Framework	Governance foundation (organizational)	GV	Clause 5.2; A.2 AI policy controls	Approved AI policy set, enforced and reviewed	AI Governance Lead	Policy documents	P1

#	Capability	EU AI Act (intent)	NIST RMF	ISO 42001 (clause/control intent)	Evidence that satisfies all three	Typical owner	Key artifacts	Priority
28	Roles & Accountability	Clear responsibility allocation (Art. 26)	GV	Clause 5.3 roles & responsibilities	RACI with single accountability per activity	AI Governance Lead	RACI matrix	P1
29	Conformity Assessment	Conformity assessment for high-risk (Art. 43)	,	Clause 9; external certification analogue	Conformity records / CE marking evidence (high-risk)	Compliance	Conformity documentation	P2
30	Continuous Improvement	Post-market feedback loop (Art. 72)	MG	Clause 10 improvement	Nonconformity handling, corrective actions, review cycle	AI Governance Lead	Improvement log	P3
31	GPAI / Foundation Model Governance	GPAI obligations (Art. 51–55)	MP, MS	A.10 third-party; A.6 lifecycle	Model documentation, usage controls, provider evidence	Engineering + Legal	GPAI governance records	P2
32	Agentic AI Governance	Human oversight, robustness applied to agents (Art. 14, 15)	GV, MG	A.9 responsible use; A.6 lifecycle	Agent identity, least-privilege, kill-switch, oversight	CISO + Engineering	Agent governance records	P2

Using the matrix. Read each row as a single unit of work. The evidence column is deliberately written so one artifact satisfies all three instruments, build it once, map it three ways. Start with every P1 row; these are foundational and most other rows depend on them. The owner column is a starting default; adjust to your structure, but preserve the single-accountability principle from Part 9.

PART 6, ENTERPRISE IMPLEMENTATION CHECKLIST

A complete implementation checklist of 260+ items, grouped by the thirteen governance domains. Each item is designed to be tracked with a **Status** (Not Started / In Progress / Complete / N/A), an **Owner**, **Evidence** (the artifact that proves it), and a **Priority** (P1/P2/P3).

The full checklist accompanies this guide as a CSV for direct import into a GRC tool or spreadsheet. What follows is the complete item set in reference form.

Domain 1, Governance & Accountability

1. Establish an AI governance committee with cross-functional membership [P1]
2. Appoint a named AI Governance Lead with defined authority [P1]
3. Define board-level accountability for AI risk [P1]
4. Publish an approved enterprise AI policy [P1]

5. Define AI governance scope and boundaries [P1]
6. Build a RACI with single accountability per governance activity [P1]
7. Define escalation paths for AI risk decisions [P2]
8. Establish a model/system approval gate for high-risk AI [P1]
9. Define decision rights: who approves deployment [P1]
10. Integrate AI governance into enterprise risk governance [P2]
11. Establish a governance review cadence [P2]
12. Define exceptions and waiver process with documentation [P2]
13. Align AI policy with existing corporate and security policies [P2]
14. Establish conflict-resolution mechanism across functions [P3]
15. Define governance KPIs and report them [P2]
16. Document management commitment to the AIMS [P1]
17. Assign resources adequate to the governance program [P2]
18. Establish a management review process [P2]
19. Define organizational context and interested parties [P2]
20. Maintain a Statement of Applicability for selected controls [P2]

Domain 2, Risk Management

1. Define an AI risk assessment methodology [P1]
2. Maintain an AI risk register [P1]
3. Assess risk per AI system before deployment [P1]
4. Conduct AI system impact assessments [P1]
5. Conduct fundamental-rights impact assessment where applicable [P1]
6. Classify each system by EU AI Act risk tier [P1]
7. Document classification rationale per system [P1]
8. Define risk acceptance criteria and thresholds [P2]
9. Define risk treatment plans per identified risk [P2]
10. Define triggers for risk reassessment (model change, incident, drift) [P1]
11. Integrate AI risk into enterprise risk management [P2]
12. Assess risk of AI system interactions and dependencies [P2]
13. Assess systemic and aggregate risk across the AI portfolio [P3]
14. Document residual risk and formal acceptance [P2]
15. Link risks to owners and treatment status [P2]
16. Reassess risk on a defined periodic cycle [P2]
17. Assess third-party and supply-chain AI risk [P2]
18. Assess risk of foundation/GPAI model usage [P2]
19. Assess agentic AI-specific risks [P2]
20. Maintain risk assessment records as evidence [P1]

Domain 3, Security (AI-specific)

1. Extend the security program to cover AI systems [P1]
2. Conduct adversarial testing / red teaming of AI systems [P2]
3. Test for prompt injection (direct and indirect) [P2]
4. Validate input filtering and guardrails [P2]
5. Test output handling for downstream injection [P2]
6. Secure the model registry with least-privilege access [P2]
7. Verify model artifact integrity and signing [P2]
8. Scan models for unsafe serialization before load [P2]
9. Assess and pin AI dependencies; maintain an AIBOM [P1]
10. Check AI components against known CVEs [P2]
11. Secure vector stores and enforce access controls [P2]
12. Verify tenant isolation in shared AI systems [P1]
13. Enforce per-agent identity for agentic systems [P1]
14. Enforce least-privilege tool scoping for agents [P1]
15. Implement human-approval gates for sensitive agent actions [P1]
16. Implement and test kill-switch / containment for agents [P1]
17. Map AI threats using MITRE ATLAS [P2]
18. Assess against OWASP Top 10 for LLMs [P2]
19. Test guardrails for bypass [P2]
20. Deploy runtime monitoring for injection/exfiltration [P2]
21. Secure secrets and credentials from model/agent access [P1]
22. Assess supply-chain security of AI tooling [P2]
23. Define AI-specific security incident procedures [P1]
24. Validate fail-safe behavior of guardrails [P2]

Domain 4, Data Governance

1. Establish data governance for AI training and operation [P1]
2. Document data provenance and lineage [P1]
3. Assess and document data quality [P1]
4. Define data eligibility criteria for AI use [P1]
5. Implement data minimization for AI [P2]
6. Remove or protect PII in training data [P1]
7. Assess training data for bias [P2]
8. Version and control datasets [P2]
9. Document data sources and licensing [P2]
10. Enforce access controls on AI datasets [P1]
11. Define data retention and deletion for AI data [P2]
12. Validate data used in fine-tuning and RLHF [P2]
13. Assess synthetic data risks where used [P3]

14. Ensure retrieval respects document-level permissions (RAG) [P1]
15. Validate ingestion pipelines for poisoning [P2]
16. Maintain data governance records as evidence [P1]

Domain 5, Operations & Lifecycle

1. Define the AI system lifecycle with governance gates [P1]
2. Define development standards for AI systems [P2]
3. Define validation and acceptance criteria per system [P2]
4. Define deployment approval process [P1]
5. Implement change management for AI systems [P2]
6. Define model versioning and release control [P2]
7. Define rollback procedures [P2]
8. Implement drift detection with thresholds [P2]
9. Define retraining triggers and process [P2]
10. Define model/system retirement process [P3]
11. Handle data appropriately at decommissioning [P3]
12. Maintain lifecycle records per system [P2]
13. Define model-swap re-validation requirements [P2]
14. Manage configuration of AI systems [P2]
15. Define environment separation (dev/test/prod) [P2]
16. Define release gates tied to governance approval [P1]
17. Manage technical debt in AI systems [P3]
18. Define capacity and resource management [P3]
19. Document operational runbooks [P2]
20. Define business continuity for AI systems [P3]

Domain 6, Compliance & Legal

1. Determine EU AI Act applicability to the enterprise [P1]
2. Map each system to applicable obligations [P1]
3. Determine provider vs. deployer role per system [P1]
4. Track EU AI Act timeline and applicable deadlines [P1]
5. Prepare conformity assessment for high-risk systems [P2]
6. Register high-risk systems as required [P2]
7. Implement Article 50 transparency obligations [P1]
8. Implement the Article 4 AI-literacy duty [P1]
9. Assess GDPR interaction and run DPIAs where needed [P1]
10. Prepare technical documentation to Annex IV scope [P2]
11. Establish serious-incident reporting capability [P1]
12. Monitor regulatory change and update obligations [P2]
13. Determine applicability of sector regulation (GxP, BFSI, etc.) [P1]

14. Assess cross-border and multi-jurisdiction obligations [P2]
15. Maintain a compliance evidence repository [P1]
16. Prepare for ISO 42001 certification if pursued [P2]
17. Conduct internal audit of the AIMS [P2]
18. Track prohibited-practice compliance [P1]
19. Assess GPAI provider obligations if applicable [P2]
20. Document legal basis for AI processing [P2]

Domain 7, Monitoring & Measurement

1. Define monitoring metrics per AI system [P2]
2. Implement performance monitoring [P2]
3. Implement accuracy monitoring against thresholds [P2]
4. Implement bias/fairness monitoring [P2]
5. Implement drift monitoring [P2]
6. Implement security/anomaly monitoring [P2]
7. Define alerting and escalation from monitoring [P2]
8. Implement post-market monitoring for high-risk systems [P2]
9. Monitor cost and resource consumption [P3]
10. Monitor agent behavior sequences [P2]
11. Define monitoring review cadence [P2]
12. Retain monitoring evidence [P2]
13. Monitor third-party AI service performance [P3]
14. Define KRIs and track them [P2]
15. Feed monitoring insights into risk reassessment [P2]
16. Monitor user feedback and complaints [P3]

Domain 8, Human Oversight

1. Design human oversight per high-risk system [P1]
2. Define oversight roles and authority [P1]
3. Ensure overseers can understand system output [P2]
4. Ensure overseers can intervene or halt the system [P1]
5. Define when human approval is mandatory [P1]
6. Guard against automation bias in oversight design [P2]
7. Provide overseers adequate information and time [P2]
8. Train personnel performing oversight [P2]
9. Record oversight actions and decisions [P2]
10. Validate oversight is effective, not theatre [P2]
11. Define oversight for agentic systems specifically [P2]
12. Review and improve oversight mechanisms [P3]

Domain 9, Documentation & Evidence

1. Define required documentation per system [P2]
2. Maintain a central evidence repository [P1]
3. Create model cards for AI systems [P2]
4. Maintain technical documentation current [P2]
5. Document design and development decisions [P2]
6. Maintain the Statement of Applicability [P2]
7. Control documented information (versioning, access) [P2]
8. Map evidence to obligations across all instruments [P1]
9. Define evidence retention periods [P2]
10. Ensure evidence is audit-ready on demand [P1]
11. Document risk and impact assessments [P1]
12. Maintain conformity documentation [P2]
13. Document data governance decisions [P2]
14. Maintain an evidence index (artifact-to-control) [P1]
15. Ensure evidence is a byproduct of operation, not retrofit [P2]
16. Maintain records of management review [P2]

Domain 10, Vendor & Third-Party

1. Maintain a register of third-party AI systems and services [P1]
2. Assess vendor AI security and governance [P2]
3. Include AI-specific clauses in vendor contracts [P2]
4. Require audit rights and incident notification [P2]
5. Assess provider vs. deployer obligations per vendor [P1]
6. Verify vendor model provenance [P2]
7. Assess foundation-model provider evidence [P2]
8. Monitor vendor performance and compliance [P2]
9. Define vendor offboarding and data handling [P3]
10. Assess concentration and dependency risk [P3]
11. Verify vendor data-handling practices [P2]
12. Include right-to-audit and evidence access [P2]
13. Maintain vendor assessment records [P2]
14. Assess embedded AI in procured SaaS [P1]

Domain 11, Training & Competence

1. Define AI literacy requirements per role [P1]
2. Deliver AI literacy training to relevant staff [P1]
3. Define competence requirements for governance roles [P2]
4. Train developers on secure and responsible AI [P2]

5. Train oversight personnel [P2]
6. Train the board on AI governance [P2]
7. Maintain training and competence records [P1]
8. Refresh training on a defined cycle [P2]
9. Assess training effectiveness [P3]
10. Build role-specific curricula [P2]
11. Onboard new staff into AI governance [P2]
12. Maintain a competence matrix [P2]

Domain 12, Audit & Assurance

1. Define an internal audit program for the AIMS [P2]
2. Conduct internal audits on a cycle [P2]
3. Track and remediate audit findings [P2]
4. Prepare for external certification audit if pursued [P2]
5. Conduct management review of the AIMS [P2]
6. Verify control effectiveness, not just existence [P2]
7. Audit evidence completeness and currency [P2]
8. Conduct independent AI risk review [P3]
9. Verify third-party assurance where relied upon [P3]
10. Maintain audit records [P2]
11. Define nonconformity handling [P2]
12. Track corrective actions to closure [P2]
13. Report assurance results to the board [P2]
14. Benchmark governance maturity periodically [P3]

Domain 13, Incident & Continuity

1. Define AI-specific incident response procedures [P1]
2. Build playbooks for model poisoning, leakage, failure, rogue agent [P1]
3. Define incident severity classification [P2]
4. Define serious-incident reporting to authorities [P1]
5. Define incident detection and triage [P2]
6. Preserve forensic evidence (traces, logs) [P2]
7. Run tabletop exercises [P2]
8. Define kill-switch and rollback procedures [P1]
9. Define notification triggers and timelines [P2]
10. Conduct post-incident review [P2]
11. Feed incidents into risk and control improvement [P2]
12. Define business continuity for AI failure [P3]
13. Maintain an incident log as evidence [P1]
14. Test incident procedures periodically [P2]

Domain 14, Transparency & Explainability

1. Implement AI-interaction disclosure (chatbots) [P1]
2. Implement AI-generated content labelling [P1]
3. Define explainability approach per system [P3]
4. Provide appropriate information to affected persons [P2]
5. Document model limitations and intended use [P2]
6. Provide deployer information for downstream use [P2]
7. Ensure explanations match audience needs [P3]
8. Maintain transparency records [P2]

Domain 15, Bias, Fairness & Ethics

1. Define fairness objectives per system [P2]
2. Test for bias across relevant groups [P2]
3. Document bias mitigation measures [P2]
4. Monitor fairness in production [P2]
5. Define ethical principles for AI use [P2]
6. Assess societal and group impact [P2]
7. Establish a review process for ethically sensitive uses [P3]
8. Maintain fairness assessment records [P2]

Domain 16, Foundation Models & Agentic AI

1. Inventory foundation/GPAI models in use [P1]
2. Assess GPAI provider obligations if a provider [P2]
3. Govern foundation-model usage and access [P2]
4. Document foundation-model limitations and risks [P2]
5. Inventory agentic AI systems [P1]
6. Apply agent identity and authorization controls [P1]
7. Apply least-privilege and tool scoping to agents [P1]
8. Apply human oversight to agent actions [P1]
9. Implement agent containment and kill-switch [P1]
10. Assess multi-agent and composition risk [P2]
11. Govern memory and state in agentic systems [P2]
12. Monitor agent behavior in production [P2]

Domain 17, Program Foundations (cross-cutting)

1. Complete an initial AI system discovery sweep [P1]
2. Identify and register shadow AI [P1]
3. Establish the governance operating model [P1]

4. Define the target maturity level and roadmap [P2]
5. Secure executive sponsorship and budget [P1]
6. Establish a single control framework mapped to all instruments [P1]
7. Establish the evidence model early [P1]
8. Define success metrics for the program [P2]
9. Establish continuous-improvement rhythm [P2]
10. Communicate the program across the enterprise [P2]
11. Align governance with business strategy [P2]
12. Establish a horizon-scanning process for new regulation [P3]
13. Define build-vs-buy governance criteria [P2]
14. Review and refresh the program annually [P2]

PART 7, BOARD DASHBOARD

The board governs AI through a small set of metrics that translate operational governance into executive language. The dashboard below separates **KPIs** (are we governing well?), **KRIs** (what could hurt us?), and **executive metrics** (where do we stand strategically?).

7.1 Key Performance Indicators (governance health)

KPI	What it measures	Healthy signal
AI inventory coverage	% of known AI systems inventoried and classified	Approaching 100%, with shadow-AI sweep current
High-risk systems assessed	% of high-risk systems with current risk + impact assessments	100%
Evidence readiness	% of high-risk systems audit-ready on demand	High and rising
Human oversight coverage	% of high-risk systems with operating oversight	100%
Control implementation	% of applicable controls implemented	Rising toward target maturity
Training completion	% of relevant staff meeting AI-literacy requirements	High, refreshed on cycle
Governance gate adherence	% of deployments passing through the approval gate	100%

7.2 Key Risk Indicators (exposure)

KRI	What it signals	Watch condition
Unclassified AI systems	Ungoverned exposure	Any high-risk system unclassified
Overdue reassessments	Governance falling behind system change	Rising backlog
Open high-severity findings	Unremediated risk	Any past SLA
AI incidents (period)	Realized risk	Trend up, or any serious incident
Shadow AI discovered	Governance blind spots	Frequent new discoveries
Vendor AI without assessment	Unmanaged third-party risk	Any high-risk vendor unassessed
Drift/performance breaches	Systems degrading in production	Threshold breaches unresolved
Regulatory deadline exposure	Compliance timing risk	Approaching deadline with gaps

7.3 Executive Metrics (strategic posture)

Metric	Board question it answers
Compliance posture vs. applicable regime	Are we on track for the obligations that apply to us?
Governance maturity level (1–5)	How capable is our program, and where are we heading?
AI risk exposure summary	What is our aggregate AI risk, and is it within appetite?
Resource adequacy	Is the program adequately funded and staffed?
Portfolio value vs. risk	Is AI delivering value proportionate to the risk we carry?
Certification status	Where do we stand on ISO 42001 / conformity, if pursued?

Reporting cadence. KRIs and incidents: every board meeting. KPIs and maturity: quarterly. Full governance review: at least annually, or on material change (major incident, new regulation, significant new high-risk deployment). Keep the board view to one page; depth lives in the operating layers beneath.

PART 8, AI GOVERNANCE MATURITY MODEL

Five levels describe the journey from ad hoc to optimized. Each level defines what is true of the organization, and what it must do to advance.

Level 1, Initial (Ad Hoc). AI is used without systematic governance. No inventory, no defined ownership, no consistent risk assessment. Governance happens reactively, if at all. Shadow AI is unknown and unmanaged. *Characteristic:* the enterprise cannot answer "how many AI systems do we run, and what are their risks?" *To advance:* establish ownership, run an inventory sweep, secure executive sponsorship.

Level 2, Developing (Repeatable). Basic governance exists for some systems. An inventory has begun, a policy exists, and high-profile systems get risk assessments, but coverage is partial and processes are inconsistent.

Governance is owned but under-resourced. *Characteristic:* governance exists but is not yet comprehensive or reliable. *To advance:* build the single control framework, extend coverage to all high-risk systems, stand up the operating model.

Level 3, Defined (Established). Governance is systematic and documented. A complete inventory, consistent risk and impact assessments, a defined operating model with clear ownership, an evidence repository, and a functioning approval gate. The program maps to all three instruments through one framework. *Characteristic:* governance is comprehensive, consistent, and evidence-generating. *To advance:* mature monitoring and measurement, pursue certification if valuable, deepen technical controls.

Level 4, Managed (Measured). Governance is measured and continuously monitored. Metrics drive decisions, monitoring is continuous, drift and performance are tracked, and the board receives meaningful reporting. Controls are validated for effectiveness, not just existence. Certification may be achieved. *Characteristic:* the enterprise knows quantitatively how well it governs and where risk sits. *To advance:* optimize based on data, automate evidence, lead rather than follow regulatory change.

Level 5, Optimized (Leading). Governance is a strategic capability and a source of advantage. It is largely automated, continuously improving, and anticipates regulation rather than reacting to it. Evidence is generated automatically as a byproduct of operation. The enterprise contributes to standards and sets industry practice. *Characteristic:* governance enables faster, safer AI adoption than competitors can achieve. *Sustained by:* continuous improvement, horizon scanning, and cultural embedding.

Using the model. Assess honestly against the definitions. Most enterprises deploying AI at scale are between Level 1 and Level 2 today. Level 3 is the realistic near-term target and the point at which the program becomes genuinely defensible. Level 4–5 are differentiators. Advance one level at a time; skipping levels produces fragile governance that fails under audit or incident.

PART 9, OPERATING MODEL

Governance succeeds or fails on ownership. This operating model assigns every governance responsibility to a role, following one cardinal rule: **exactly one accountable owner per activity**. Others are consulted or informed, but accountability is never shared, shared accountability is no accountability.

Role	Governance responsibility
Board	Ultimate accountability for AI risk. Approves risk appetite, oversees the program, receives regular reporting, ensures adequate resourcing.
CEO	Owns the enterprise's AI posture. Ensures governance is resourced and empowered; accountable to the board for the program's effectiveness.
Chief AI Officer / AI Governance Lead	Owns the governance program day to day. The single accountable role for the AIMS: runs the committee, maintains the framework and evidence, drives the roadmap, reports to the board. <i>The role most enterprises have not yet created and most need.</i>
CTO / Engineering	Owns technical implementation of controls: logging, monitoring, testing, lifecycle, security controls, agent controls. Produces technical evidence.
CISO / Security	Owns AI-specific security: adversarial testing, guardrails, agent identity and containment, incident response, supply-chain security.
CRO / Risk	Owns the risk methodology, register, and treatment. Integrates AI risk into enterprise risk. Owns risk and impact assessments.
Legal / Compliance	Owns regulatory interpretation and obligation mapping: EU AI Act applicability, classification, conformity, transparency, incident reporting.
CDO / Data	Owns data governance for AI: provenance, quality, minimization, lineage, dataset controls.
DPO	Owns privacy: DPIAs, GDPR interaction, data-subject handling for AI.
Model Owners	Own individual AI systems end to end: assessment, documentation, monitoring, and lifecycle for their system. Accountable for their system's governance.
Business Owners	Own the business use and its human oversight. Accountable that the system is used within its intended purpose and that oversight operates.
Internal Audit	Provides independent assurance: audits the AIMS, verifies control effectiveness, reports findings. Independent of the program it audits.
Procurement	Owns third-party AI governance in the buying process: vendor assessment, contractual controls, provenance.
HR	Owns competence and AI literacy: training delivery, competence records, onboarding into governance.

The committee. These roles convene as an AI governance committee chaired by the AI Governance Lead, reporting to the board. The committee is where cross-functional decisions are made, classification disputes, high-risk approvals, exception requests, and where the single-accountability principle is enforced in practice.

The critical appointment. If one thing is missing in most enterprises, it is the AI Governance Lead, a single named owner for the whole program. Without it, governance work scatters across functions that each own a piece and none owns the whole, which is precisely the parallel-programs failure from Part 2. This role can be fractional in smaller enterprises, but it must exist and must be named.

PART 10, IMPLEMENTATION ROADMAP

A phased path from standing start to operating program. Each phase has a theme, concrete deliverables, and an exit criterion. Phases build on each other; do not start a phase before its predecessor's foundations are in place.

First 30 Days, Foundation and Visibility

Theme: know what you have and who owns it. - Appoint the AI Governance Lead and secure executive sponsorship. - Stand up the AI governance committee. - Run an initial AI system discovery sweep, including shadow-AI detection. - Build the first version of the AI inventory with owners assigned. - Draft the enterprise AI policy. - Determine EU AI Act applicability at the enterprise level. - **Exit criterion:** a named owner, a committee, and a first inventory exist. The enterprise can now answer "what AI do we run?"

31–60 Days, Framework and Classification

Theme: establish the single control framework and classify risk. - Adopt the single control framework mapped to all three instruments (the Crosswalk). - Classify each inventoried system by EU AI Act risk tier and record rationale. - Stand up the risk register and risk methodology. - Establish the evidence repository and evidence model. - Build the RACI with single accountability per activity. - Define the deployment approval gate. - **Exit criterion:** every system is classified, the control framework is adopted, and the approval gate is live.

61–90 Days, Core Controls for High-Risk Systems

Theme: implement P1 controls where risk is highest. - Run risk and impact assessments for all high-risk systems. - Implement human oversight for high-risk systems. - Implement logging and traceability to the required standard. - Establish data governance for AI. - Stand up AI-specific incident response and playbooks. - Implement Article 50 transparency and Article 4 AI-literacy obligations (live in 2026 regardless of the high-risk deferral). - **Exit criterion:** high-risk systems have current assessments, operating oversight, adequate logging, and the enterprise meets its 2026-applicable EU obligations.

91–180 Days, Depth and Assurance

Theme: extend coverage and validate effectiveness. - Implement monitoring and measurement (performance, drift, bias, security). - Conduct adversarial testing / red teaming of high-risk and agentic systems. - Complete vendor assessments for third-party AI. - Stand up training and competence programs. - Conduct the first internal audit of the AIMS. - Deliver the first full board dashboard. - Begin ISO 42001 preparation if certification is pursued. - **Exit criterion:** governance is comprehensive across the portfolio, monitored, and independently assured. The program is at Level 3 maturity.

12 Months, Optimization and Certification

Theme: measure, certify, and improve. - Complete ISO 42001 certification if pursued. - Prepare conformity assessment approach for high-risk systems ahead of the 2027/2028 deadlines. - Mature monitoring to continuous, metric-driven operation. - Automate evidence generation where possible. - Establish horizon scanning for regulatory change. - Run the annual governance review and refresh the roadmap. - **Exit criterion:** a measured, improving program (Level 4), certified if pursued, and positioned ahead of the high-risk deadlines rather than racing them.

On the EU timeline and sequencing. The high-risk deferral to December 2027 (Annex III) and August 2028 (Annex I) gives enterprises breathing room on the heaviest regime, but the 2026 transparency and literacy obligations are not deferred, and the durable machinery takes time to build. The right response is to use this roadmap to reach Level 3 on the ISO/NIST foundation now, so that high-risk conformity becomes a byproduct of an operating system rather than a last-minute project.

PART 11, TEMPLATES

Nine reusable templates translate the framework into working artifacts. Each is described here in structure; all accompany the guide as ready-to-use files. All are original and reproduce no standard's text.

11.1 AI Inventory Register

Fields: system ID; name; description; owner; business function; provider/deployer role; underlying models and versions; data sources; risk tier; deployment status; environments; dependencies; last assessment date; next review trigger; evidence location.

11.2 AI Risk Register

Fields: risk ID; linked system; risk description; category; likelihood; impact; inherent rating; controls; residual rating; owner; treatment plan; acceptance decision and approver; review date.

11.3 Model Card

Sections: model identity and version; intended purpose and out-of-scope uses; training data summary and provenance; performance and accuracy; limitations; fairness assessment; security testing summary; human oversight design; monitoring approach; owner and review cycle.

11.4 AI System Approval Form

Sections: system and owner; risk classification and rationale; risk and impact assessment reference; controls implemented; human oversight design; testing and validation results; residual risk and acceptance; approver and decision; conditions and review date.

11.5 Human Oversight Assessment

Sections: system and its decisions; oversight roles and authority; intervention and halt capability; information provided to overseers; automation-bias safeguards; mandatory-approval conditions; competence of overseers; oversight logging; effectiveness validation.

11.6 AI Incident Report

Sections: incident ID and severity; systems affected; detection and timeline; what happened; impact (data, individuals, business); containment and remediation; root cause; regulatory-notification decision; corrective actions; lessons fed back into risk and controls.

11.7 Vendor AI Assessment

Sections: vendor and service; AI components and models; provider/deployer obligation split; security posture; data handling; governance maturity; provenance; contractual controls (audit rights, incident notification); assessment outcome and conditions; review cycle.

11.8 Model Retirement Checklist

Items: decommissioning approval; dependency and downstream impact check; data handling and retention/deletion; access revocation; documentation archival; stakeholder notification; evidence retention; confirmation of clean termination.

11.9 AI Governance Audit Checklist

Sections: inventory completeness; classification accuracy; assessment currency; control implementation and effectiveness; evidence readiness; oversight operation; monitoring operation; incident-process readiness; training completion; nonconformity and corrective-action status.

PART 12, EXECUTIVE POSTER

An A1 landscape infographic accompanies this guide as a print-ready SVG. It renders the governance flow as a single memorable visual:

The integration flow (top band): EU AI Act (what you must prove) → ISO/IEC 42001 (how you operate) → NIST AI RMF (how you reason) → converging into **Enterprise Controls** → generating **Evidence** → supporting **Audit** → enabling **Certification & Conformity** → sustained by **Continuous Monitoring**, which loops back to controls.

The altitude model (left): the three instruments shown at their altitudes, Obligation (law), Operation (management system), Reasoning (risk framework), making visually clear that they are complementary layers, not competitors.

The operating stack (center): the ten-layer stack from Part 13, shown as a vertical spine.

The maturity ladder (right): Levels 1–5, as a rising path.

The poster is designed so a board member grasps the whole model in one glance and an implementer can use it as a wall reference. It is provided in the CXO Intelligence palette.

PART 13, THE ORIGINAL FRAMEWORK: THE AI GOVERNANCE OPERATING STACK

Standards tell you *what* good governance includes. They do not give you a single mental model for how the pieces stack into an operating system. The **AI Governance Operating Stack** is that model, an original, reusable framework that any enterprise can adopt. It arranges governance into ten layers, each depending on the ones below it, each with a clear question it answers, an owner, and an output that feeds the layer above.

The organizing insight: **governance is a stack, not a checklist**. Value flows up (each layer enables the next) and evidence flows down (each layer proves the intent of the one above). A weakness at any layer undermines everything above it, which is why programs that start at the control layer, skipping strategy and policy, produce controls nobody can trace to a purpose.

Layer 1, Strategy

Question: why are we using AI, and what is our risk appetite? The foundation. Defines the enterprise's AI ambition, the value it seeks, and the risk it will accept. Owned by the board and CEO. Without this layer, governance has no reference point for any decision. **Output:** AI strategy and risk appetite.

Layer 2, Policy

Question: what are our rules? Translates strategy into enforceable rules, the AI policy set covering acceptable use, approval, data, and third-party AI. Owned by the AI Governance Lead. **Output:** approved, enforceable policies.

Layer 3, Risk

Question: what could go wrong, and how much does it matter? The reasoning engine (this is where NIST's Map and Measure live). Identifies, assesses, and prioritizes AI risk per system and across the portfolio. Owned by the CRO. **Output:** risk register, assessments, treatment decisions.

Layer 4, Controls

Question: what do we do about it? Converts risk decisions into specific controls, the governance capabilities in the Crosswalk. This is where ISO Annex A control selection and the EU AI Act obligations become concrete measures. Owned jointly, coordinated by the AI Governance Lead. **Output:** an implemented, mapped control set.

Layer 5, Technology

Question: how do we implement controls in the systems themselves? The technical realization: logging, monitoring, testing, guardrails, agent identity, lifecycle tooling. Owned by Engineering and Security. **Output:** technical controls operating in production.

Layer 6, Evidence

Question: how do we prove it? The layer most programs neglect and auditors care about most. Evidence must be generated as a byproduct of operation, mapped once to all instruments, and audit-ready on demand. Owned by the AI Governance Lead. **Output:** a living evidence repository with an artifact-to-control index.

Layer 7, Monitoring

Question: is it still working? Makes governance continuous. Tracks performance, drift, bias, security, and behavior against thresholds, feeding breaches back into the risk layer. Owned by Engineering and Operations. **Output:** continuous monitoring with alerting and feedback.

Layer 8, Assurance

Question: can we independently trust it? Independent verification that controls are not just present but effective. Internal audit and management review. Owned by Internal Audit, independent of the program. **Output:** assurance findings and validated effectiveness.

Layer 9, Audit & Certification

Question: can a third party confirm it? External validation: ISO 42001 certification, EU conformity assessment, regulatory examination. Owned by Compliance. **Output:** certification and conformity records.

Layer 10, Improvement

Question: how do we get better? Closes the loop. Nonconformities, incidents, monitoring insights, and audit findings drive continual improvement, feeding back to strategy. Owned by the AI Governance Lead. **Output:** an improvement cycle that raises maturity over time.

Why the stack works. It gives every stakeholder a shared map. The board engages at Layers 1 and 10 (strategy in, improvement out). Engineers work at Layers 5 and 7. Auditors work at Layers 8 and 9. Everyone sees how their work connects to the whole, and every control traces down to a risk and up to a purpose. Adopt it as the enterprise's canonical governance model, and the three instruments stop being three programs, they become inputs to one operating system: NIST informs Layer 3, ISO structures Layers 2–10, and the EU AI Act sets requirements that shape Layers 4 and 9.

THE AI GOVERNANCE NAVIGATOR

The differentiator. A decision tool that tells any enterprise which obligations actually apply to it, which controls become mandatory, what evidence to collect, and who owns it.

Most governance content describes the landscape. The Navigator routes you through it. It is a structured decision tool: answer a sequence of questions about your AI use, and it outputs your applicable obligations, mandatory controls, required evidence, responsible owners, and remaining gaps. A full interactive version accompanies this guide; the logic is below so it can be applied directly or built into a tool.

How the Navigator works

The Navigator asks a sequence of gating questions. Each answer activates obligations, controls, and owners. At the end, it assembles your profile.

Q1, Are you building or buying AI?

- **Building** → you are likely a *provider*. Activates: full lifecycle controls, technical documentation, model validation, provider obligations. Owner: Engineering + Legal.
- **Buying** → you are likely a *deployer*. Activates: vendor assessment, deployer obligations, human oversight, usage controls. Owner: Procurement + Business Owner.
- **Both** → both sets apply per system; classify each system individually.

Q2, Are you using foundation / general-purpose AI models?

- **Yes** → activates GPAI/foundation-model governance: model documentation, usage controls, provider-evidence review; if you are a GPAI *provider*, the GPAI obligation track (Art. 51–55). Owner: Engineering + Legal.

- **No** → skip GPAI track; standard system governance applies.

Q3, Do any of your AI systems fall in a high-risk category under the EU AI Act?

(biometrics, critical infrastructure, education, employment, essential services, law enforcement, migration, justice, or AI as a safety component of a regulated product) - **Yes (Annex III, stand-alone)** → the full high-risk regime applies: risk management, data governance, technical documentation, logging, human oversight, accuracy/robustness/security, conformity assessment, registration. **Deadline: 2 December 2027.** Owner: cross-functional, led by AI Governance Lead. - **Yes (Annex I, embedded in regulated product)** → high-risk regime plus sector product rules. **Deadline: 2 August 2028.** Owner: cross-functional + product compliance. - **No, but limited-risk** → transparency obligations (Art. 50) and AI-literacy duty (Art. 4). **Deadline: 2 August 2026, not deferred.** Owner: Product + HR. - **No, minimal risk** → no specific EU obligations, but baseline governance still recommended.

Q4, Do you have EU market exposure at all?

- **Yes** → the EU AI Act applies; proceed with Q3's outcome.
- **No** → the EU AI Act does not directly bind you, but it is a strong baseline; NIST + ISO still fully apply and are recommended as your primary framework.

Q5, Are you in a regulated sector (GxP / life sciences, BFSI, healthcare, critical infrastructure)?

- **Yes** → sector regulation stacks on top of AI governance. Activates: sector-specific validation, documentation, and audit obligations; likely higher assurance and stricter human oversight. Owner: Compliance + sector SME.
- **No** → AI-specific obligations only.

Q6, Do you want ISO/IEC 42001 certification?

- **Yes** → activates the full AIMS: Clauses 4–10, Statement of Applicability, internal audit, management review, and external certification audit. Owner: AI Governance Lead + Compliance.
- **No** → ISO 42001 still valuable as an operating blueprint even without certification; adopt the machinery without the audit.

Q7, Are you deploying agentic AI (autonomous, tool-using systems)?

- **Yes** → activates agentic controls: per-agent identity, least-privilege tool scoping, human-approval gates, containment/kill-switch, memory-integrity governance, multi-agent composition risk. Owner: CISO + Engineering. *This is the fastest-moving risk area; treat as high priority.*
- **No** → standard system governance applies.

What the Navigator outputs

After the questions, the Navigator assembles a profile:

1. **Which regulations apply**, EU AI Act (with your specific tier and deadline), sector regulation, GPAI track, as activated.
2. **Which controls are mandatory**, the specific Crosswalk rows activated by your answers, separated into mandatory (legally required for your profile) and recommended (good practice).

3. **Which evidence to collect**, the evidence artifacts from the activated Crosswalk rows.
4. **Which owners are responsible**, the roles from the operating model mapped to your activated controls.
5. **Which artifacts to create**, the templates from Part 11 relevant to your profile.
6. **Which gaps remain**, activated controls not yet implemented, prioritized P1 → P3.

Worked example

A European bank building an AI system to assist credit decisions, using a foundation model, wanting ISO certification, not yet using agents:

- Q1 Building → provider obligations. Q2 Yes → foundation-model governance. Q3 Yes, Annex III (credit scoring) → **full high-risk regime, deadline 2 December 2027**. Q4 Yes → EU AI Act binds. Q5 Yes, BFSI → sector regulation stacks (model risk management expectations, stricter oversight). Q6 Yes → full AIMS + certification. Q7 No → standard controls.

Output: high-risk EU obligations + BFSI sector rules + GPAI governance + full ISO AIMS. Mandatory controls include (Crosswalk rows) risk management, data governance, technical documentation, logging, human oversight, accuracy/robustness/security, conformity assessment, plus bias/fairness given credit context. Evidence: the full high-risk technical file, impact and fundamental-rights assessments, oversight design and logs, validation and bias-testing records, conformity documentation, and the ISO Statement of Applicability. Owners: cross-functional, AI Governance Lead accountable, with model risk (CRO) and Legal heavily engaged. Priority gaps surface as P1 first. Target: Level 3 maturity well ahead of the December 2027 deadline.

Why the Navigator matters. It converts a generic framework into a specific, actionable roadmap for one enterprise's exact situation, which regulations, which controls, which evidence, which owners, which gaps. That specificity is what makes it usable rather than merely informative, and it is what makes the difference between a reference that sits on a shelf and a tool a practitioner returns to.

CLOSING NOTE

The three instruments, the EU AI Act, ISO/IEC 42001, and the NIST AI RMF, are not three programs to run in parallel. They are three views of one governance reality: what you must prove, how you operate, and how you reason. Build the machinery once, generate evidence once, map it three ways, and you satisfy all three, and position yourself for whatever comes next.

The high-risk deadline moving to December 2027 changed the timing, not the destination. The enterprises that use this window to build durable governance on the ISO and NIST foundation, rather than treating the deferral as a reprieve, will meet the deadline as a formality and will govern AI as a capability rather than a compliance burden.

That is the difference between compliance and governance. Compliance is passing an inspection. Governance is being the kind of organization that would pass any inspection, because the machinery is real, the evidence is a byproduct of operation, and someone is genuinely accountable.

The Enterprise AI Governance Crosswalk · A CXO Intelligence Governance Series Reference · Leadership in the Age of AI

This guide reproduces no copyrighted standard text. All clause and article references are pointers to source documents, which should be consulted directly for normative wording. Regulatory timelines and standard statuses reflect mid-2026; verify against current sources before formal use.