



RESPONSIBLE
AI TRUST

NOVEMBER 2025

GLOBAL AI GOVERNANCE ALIGNMENT MAP

Navigating AI Regulation, Standards,
and What's Coming Next

PRESENTED AND LED BY

LEHAR GUPTA, RESPONSIBLE AI TRUST CEO

AUTHORS

Mike Wood, Research Associate, Responsible AI Trust

With guidance & co-authorship from:

Lehar Gupta, Research Director, Responsible AI Trust

REVIEWED BY:

Patrick Sullivan, VP of Strategy and Innovation, A-LIGN



Responsible AI Trust
Flagship Brief
#1.25

Powered by



MetricsLM

| Continuous AI assurance

TABLE OF CONTENTS



Foreword from the Founder

About this Brief	1
------------------	---

Leader's Dashboard	2
--------------------	---

Executive Takeaways	4
---------------------	---

Abstract & Introduction	5
-------------------------	---

Background and Context	7
------------------------	---

The Alignment Blueprint - From Fragmentation to Foresight	10
---	----

Global AI Governance Matrix - Framework Comparison	12
--	----

Crossroads Analysis - Overlaps, Obligations, and Omissions	14
--	----

Regional Governance Postures - US, EU, UK, China, Africa, Brazil	20
--	----

Global AI & Cybersecurity Timeline (1975-2040+)	24
---	----

Conclusion - Mapping Today, Preparing for Tomorrow	26
--	----

What This Means for You (Now vs Next)	35
---------------------------------------	----

Reflections	
-------------	--

Appendices - 1, 2,3,4	
-----------------------	--



FOREWORD FROM THE FOUNDER



At Responsible AI Trust, our purpose is simple yet vital: to strengthen global confidence in intelligent systems.

Across industries and borders, AI is redefining how decisions are made, risks are managed, and accountability is shared. With this transformation comes a collective responsibility to ensure that technology remains aligned with human values, transparency, and fairness.

Each brief we publish represents collaborative work from researchers, advisors, and practitioners who believe that trust must be earned through clarity, governance, and evidence. Together, we translate complex regulation into actionable insight, helping leaders navigate uncertainty with structure and foresight.

Responsible AI is not a trend, it is the foundation of sustainable innovation. As systems grow more capable, our frameworks for oversight must grow equally intelligent, adaptive, and globally connected.

Thank you for being part of this effort to turn principles into practice, and ideas into accountability.

Lehar Gupta

Founder & CEO, Responsible AI Trust

✉ Lehar@ResponsibleAITrust.com



ABOUT THIS BRIEF

Global AI Governance Alignment Map

The Global AI Governance Alignment Map is Responsible AI Trust's flagship research publication, designed to provide enterprises, investors, and policymakers with a clear, structured view of the world's emerging AI governance frameworks. It maps how regulations, standards, and voluntary codes align and diverge across jurisdictions, from the EU AI Act to ISO/IEC 42001 and the NIST AI RMF, and identifies the converging principles shaping global interoperability.

Purpose and Motivation:

AI governance has moved from aspiration to implementation. This brief was conceived to help leaders translate complex, fragmented regulations into actionable frameworks for responsible innovation. It seeks to clarify how risk management, transparency, and security-by-design can serve as universal anchors across regions and sectors.

Audience:

This publication is written for executives, policymakers, compliance officers, and investors navigating the evolving landscape of AI regulation. It supports decision-makers building trust infrastructure, those responsible for aligning innovation with governance, strategy, and assurance.

Methodology:

The brief combines primary and secondary research, cross-referencing over 30 global frameworks including binding laws, voluntary principles, and technical standards. Comparative matrices, literature reviews, and expert advisory inputs were synthesised into a single alignment model.

Version:

Public Release 1.0 - November 2025

Flagship Series: Responsible AI Trust Brief #1

Future updates will incorporate emerging standards, treaties, and evaluation benchmarks to maintain a live, global reference.

● [View the latest live version](#)

Disclaimer:

The "Global AI Governance Alignment Map" has been independently developed by Responsible AI Trust. The analyses, views, and recommendations expressed herein are solely those of the report authors and do not necessarily represent the opinions or endorsements of any organisations referenced. This publication is provided "as is," without warranty of any kind. Responsible AI Trust accepts no liability for actions taken based on its content. This document is for informational purposes only and should not be relied upon as legal or regulatory advice. All rights reserved. No part of this report may be reproduced, distributed, or transmitted in any form or by any means, including photocopying, recording, or other electronic or mechanical methods, without the prior written permission of Responsible AI Trust, except in the case of brief quotations used in reviews or scholarly works with proper citation. Any unauthorised sale, alteration, or redistribution of this report is strictly prohibited and should be reported to info@responsibleaitrust.com.

LEADER'S DASHBOARD



General Information

Brief Title : Global AI Governance Alignment Map - Responsible AI Trust

Authors : Mike Wood (Author) | Lehar Gupta (Co-Author & Research Director) | Patrick Sullivan (Advisor & Reviewer)

Date Prepared : November, 2025

Executive Summary

AI governance has shifted from principles to proof, risk management, transparency, and security-by-design are emerging as the universal backbone of global regulation.

Strategic Insight

The world is converging on three anchors:

- Risk Management: universal organising principle (EU AI Act, ISO 23894, NIST AI RMF).
- Transparency: mandatory disclosure (EU registries, NIST system cards, C2PA provenance).
- Security-by-Design: the new passport for interoperability (NIS2, Cyber Resilience Act, ENISA).

But divergence remains across:

- Assurance & Enforcement – penalties and proof vary widely.
- Clarity & Precision – principle-based vs prescriptive laws.
- Incident Reporting – fragmented thresholds and formats.

Executive Priorities - “Now vs Next”

Now (2025–2030)	Next (2030–2040)
Implement AI Bills of Materials (AI-BOMs) and model evaluation frameworks.	Prepare for International AI Safety Accord and compute governance thresholds.
Align ISO/IEC 42001 with NIS2 and EU AI Act.	Adopt universal provenance layers and AI ESG metrics.
Treat security-by-design as baseline, not optional.	Link AI governance to trade, procurement, and ESG reporting.

LEADER'S DASHBOARD



Maturity Model Snapshot - "TrustScore Readiness"

Level	Description	Analogy
1. Unverified	Ad-hoc policies, no formal evidence.	Pre-compliance
2. Documented	Policies mapped to frameworks.	Policy readiness
3. Auditable	Internal assurance, partial ISO alignment.	Emerging compliance
4. Certified	Third-party audit (ISO 42001, NIST alignment).	Market-ready
5. Trusted	Demonstrable, portable compliance.	AI trade passport

Key Takeaways for Leaders

- Proof beats promise. Regulators are moving from "intent" to "evidence."
- Security is global currency. Align cybersecurity and AI governance early.
- Transparency is reputation. Public trust is earned through explainability.
- Harmonisation is coming. The next decade favours organisations ready for interoperability.

Strategic Call to Action

Build a portable AI governance file:

- Risk register
- Model evaluations
- Incident log
- AI-BOM
- Provenance documentation

These will be tomorrow's licence to operate in the AI economy.

Presented & Led by

Lehar Gupta

Founder & Research Director, Responsible AI Trust
(Co-Author, Global AI Governance Alignment Map)

Lehar@ResponsibleAITrust.com



EXECUTIVE TAKEAWAYS

Global AI Governance Alignment Map

1. **Three converging principles:** Risk management, security-by-design, and transparency are the universal anchors of AI governance.
2. **Three main divergences:** Enforcement, clarity & precision, and incident reporting split sharply across jurisdictions, driving compliance fatigue and uncertainty.
3. **Beyond the West:** Brazil's PL 2338/2023, the African Union AI Strategy, and UAE AI Office initiatives show momentum outside the EU/US, shaping the global balance.
4. **The road ahead:** Expect AI Bills of Materials (AI-BOMs), provenance/authenticity systems, compute governance thresholds, and ESG-style AI reporting to become baseline requirements in the next decade.



ABSTRACT & INTRODUCTION

Artificial Intelligence governance has moved from the margins to the mainstream, now standing alongside economic stability, democratic resilience, and global trade as a systemic priority. The European Union's AI Act has entered into law, the Council of Europe has opened the world's first AI treaty, and a patchwork of national laws, standards, and voluntary codes is rapidly redefining how technology must be designed, deployed, and audited.

AI has left the lab. It now underpins healthcare, defence, finance, critical infrastructure, and democratic processes. With its rapid adoption, governments are racing to regulate while industry grapples with fragmented frameworks. This uncertainty creates both risk and opportunity:

For governments, enterprises, critical infrastructure suppliers, and compliance leaders, the challenge is immediate and strategic: how to operate across multiple, sometimes conflicting frameworks while maintaining speed, trust, and delivery.

This whitepaper delivers a Global AI Governance Alignment Map: a structured overview of binding laws, voluntary initiatives, and technical standards, showing where they converge and where they diverge. It compares leading frameworks (EU AI Act, NIST AI RMF, ISO/IEC 42001), cross-references cybersecurity overlays like NIS2 and the Cyber Resilience Act, and flags emerging priorities such as model evaluations, provenance, and compute governance. The goal is to give senior decision-makers an actionable roadmap for today, while also signalling where AI governance is heading over the next decade and beyond.

The risk: Compliance fatigue from multiple, unaligned frameworks and regulatory gaps where no framework yet applies.

The opportunity: Three converging principles - risk management, transparency, and security by design - suggest a shared foundation is emerging.

This paper seeks to map those answers and provide an actionable guide for decision-makers navigating today's patchwork while preparing for the governance frameworks of tomorrow.



ABSTRACT & INTRODUCTION

At the heart of this challenge lie several pressing questions:

1. Where do leading frameworks such as the EU AI Act, NIST AI RMF, and ISO/IEC 42001 align and diverge?
2. How do security overlays like NIS2, the Cyber Resilience Act, and ENISA/ETSI guidance interact with AI-specific requirements?
3. What do the various different regional approaches across the EU, US, Asia, Latin America, Africa, and the Gulf reveal about alignment opportunities and friction points?

Looking forward, the decade ahead will be defined less by drafting new flagship laws and more by deepening existing frameworks. The EU AI Act will be operationalised through standards, the GPAI Code of Practice and the EU Cyber Resilience Act will harden supply-chain security, and model evaluation will evolve from voluntary practice to a procurement prerequisite. Companies unable to produce AI governance reporting will increasingly be excluded from tendering. Beyond 2030, convergence on incident taxonomies, sectoral playbooks, certification ecosystems, and ESG-style AI governance metrics will accelerate global alignment. And beyond 2040, an International AI Safety Accord, compute governance thresholds, universal provenance layers, and secure-by-design trade requirements may define the global baseline.

This whitepaper seeks not just to describe today's patchwork, but to anticipate tomorrow's architecture. Guiding governments, enterprises, and multilateral bodies toward governance models that are adaptive, resilient, and globally interoperable.



BACKGROUND AND CONTEXT

Artificial Intelligence governance is no longer theoretical; it is an active and expanding regulatory landscape. The EU AI Act has become the world's first comprehensive binding law, setting obligations based on risk tiers and establishing conformity assessment routes. In parallel, the Council of Europe AI Convention opened for signature in 2024, creating the first global treaty linking AI directly to human rights, democracy, and the rule of law.

Elsewhere, national governments are moving quickly. Denmark and the Netherlands prioritising anti-deepfake laws. The United States has leaned on executive orders, NIST's AI Risk Management Framework, and OMB guidance to set a de facto baseline for federal use, while states like Colorado are introducing their own AI laws. These are all alongside the US fast tracking and adopting bills such as the TAKE IT DOWN act which became public law in May 2025. Canada's AIDA is in progress, Brazil has passed a Senate bill, and across Asia-Pacific, Singapore, Japan, and South Korea are publishing governance frameworks that mix voluntary guidelines with statutory obligations. China has already enforced binding rules on generative AI and synthetic media, further national standards come into effect in November 2025. While the African Union has endorsed a continental AI strategy.

Internationally, voluntary principles still matter and show willingness.

The OECD AI Principles, the G7 Hiroshima Process, and the integration of GPAI into the OECD show movement towards multilateral coordination, while safety commitments from the Bletchley Declaration and the Seoul Frontier AI Agreements highlight the role of industry self-commitments.

Standards bodies are also stepping up: ISO/IEC 42001 and ISO/IEC 23894 provide management system and risk guidance, while ETSI and ENISA are defining baseline cybersecurity requirements for AI systems.



BACKGROUND AND CONTEXT

Challenges

Despite the momentum, global alignment remains limited. The friction falls into three core divergences and two structural drivers that amplify them.

Assurance & enforcement (the proof-and-penalty gap)

What counts as “proven” compliance — and what happens if you fall short — varies widely. The EU and a few peers require formal assurance routes and attach real penalties; principle-based regimes don’t. Result: over-engineering in some markets, under-shooting in others, and little portability of evidence.

Clarity & precision (Innovation vs Control)

Some frameworks are prescriptive (risk tiers, filings, CE-style routes); others set broad principles. Precision creates certainty but can slow product cycles; flexibility enables iteration but opens grey zones and loopholes. Teams are left guessing what “good enough” looks like, per market.

Incident reporting (fragmented thresholds, formats, timelines)

Obligations, definitions of “serious” harm, and reporting clocks differ across regimes. In many they are completely absent. This forces bespoke playbooks by jurisdiction and weakens cross-border learning from near-misses.

Structural drivers that amplify the above:

Jurisdictional fragmentation & data sovereignty

Different definitions of AI, role duties, and cross-border rules (for data, models etc) complicate supply chains and procurement, even when the underlying controls are similar.

Pace of change & compliance fatigue

New laws, standards, and guidance land faster than governance programmes mature.



BACKGROUND AND CONTEXT

Overlapping audits, duplicative attestations, and shifting expectations drain capacity — especially for SMEs this delays delivery increasing cost and reducing innovation.

Existing solutions

Several tools exist to help organisations manage this complexity, but none are complete.

1. Voluntary frameworks like the NIST AI RMF that provide a common risk vocabulary and lifecycle structure, but without enforcement they rely on uptake and are open to interpretation.
2. Standards such as ISO/IEC 42001 help organisations operationalise governance through certification, but they are only just beginning to be adopted and are not yet widely recognised by regulators. They can also be costly to implement for many organisations that prioritise innovation over structure.
3. International principles (OECD, UNESCO, UN resolutions) establish high-level goals, but offer little practical advice, structure or detail on assurance or enforcement.
4. Corporate commitments at forums like Bletchley or Seoul show good will, but they remain pledges rather than obligations. Pledges are not proof.

The result is a governance patchwork that if fully adhered to would cover some areas too heavily with repeated over-regulation and have other blind-areas that are sparsely covered, open to interpretation and exploitation. Some areas (risk management, transparency, lifecycle controls) show early convergence, while others (assurance, penalties, sector carve-outs or Freemium services) remain fragmented.

This context sets the stage for why a global alignment map is urgently needed.



THE ALIGNMENT BLUEPRINT

From Fragmentation to Foresight

The challenges outlined above make one point clear: the world needs more than scattered laws and voluntary pledges. It needs a Global AI Governance Alignment Map, a blueprint that cuts through the maze of fragmented rules, overlapping audits, and shifting definitions. Instead of piecemeal compliance, leaders, governments, and enterprises require a single structured view that compares obligations, highlights convergence, and flags what's coming next.

The Blueprint is designed to deliver a framework that enables leaders to chart their compliance path across jurisdictions today, while preparing for tomorrow's standards.

Key Features and Components

Global AI Governance Framework Matrix:

A comprehensive grid with frameworks as rows (EU AI Act, NIST AI RMF, ISO/IEC 42001, national laws, voluntary codes) and obligations as columns (scope, risk tiers, lifecycle controls, transparency, cybersecurity, human oversight, reporting, enforcement, cross-border rules, sector carve-outs). The matrix is designed for filtering, expansion, and visualisation, highlighting gaps and overlaps visible at a glance.

Crossroads Analysis:

Targeted comparisons where frameworks meet and diverge. For example: EU AI Act vs. NIST AI RMF, EU AI Act vs. ISO/IEC 42001, and overlays between AI-specific rules and broader cybersecurity frameworks (NIS2, Cyber Resilience Act). These analyses reveal duplication, blind spots, and friction points distilled into simple visuals.

Global Snapshots:

Half-page profiles of leading jurisdictions (EU, US, UK, Canada, Brazil, Japan, Singapore, India, China, AU/NZ, African Union, Gulf States). Each outlines current obligations, upcoming measures, and alignment opportunities, grounding the global picture in regional realities.



THE ALIGNMENT BLUEPRINT

From Fragmentation to Foresight

Timeline & “What’s Next” Tracker:

A forward-looking roadmap from 1975 to 2040+. It shows when key laws, standards, and treaties were introduced, and maps what’s expected in the next decade, from model evaluation programmes and provenance standards to compute governance thresholds and AI Bills of Materials.

Together, these components transform the Alignment Blueprint from a static report into a navigation tool. By exposing where frameworks converge, where they clash, and what’s on the horizon, it equips leaders with clarity in a space often defined by confusion. What follows is a deep dive into the laws, standards, and initiatives shaping AI governance and the practical insights that will help you stay ahead of them.



GLOBAL AI GOVERNANCE MATRIX

Please see Appendix 4 for the Framework Comparison

The Global AI Governance Alignment Matrix was designed as a structured way to compare a highly fragmented field. Each framework, whether a binding law, voluntary code, or technical standard, has been broken down into a consistent set of dimensions: scope and definitions, risk categorisation, lifecycle controls, transparency and documentation requirements, security linkages, human oversight, incident reporting, assurance routes, enforcement mechanisms, cross-border rules, and sector carve-outs. By standardising the categories, the matrix makes very different instruments comparable at a glance. For example, it shows how a treaty like the Council of Europe AI Convention and a technical specification like C2PA can be evaluated against the same obligations, even though they originate from completely different contexts.

The process also allows us to highlight the relationship between frameworks: where obligations repeat, where they diverge, and where they leave blank spaces. This makes the matrix not just a catalogue, but a diagnostic tool. It can be updated as new laws, standards, and guidance emerge, allowing policymakers, enterprises, and critical-infrastructure providers to see the regulatory terrain as a living map rather than a static checklist.

Insights Emerging from the Matrix

A few key patterns stand out. First, risk management is the universal denominator: nearly every framework, from EU AI Act to ISO/IEC 23894 to NIST AI RMF, uses risk as the organising principle, though the level of enforcement varies. Second, security-by-design is becoming the global passport: EU instruments like NIS2 and the Cyber Resilience Act hard-wire it into law, while ETSI, ENISA, and Singapore's CSA push it as best practice. Third, transparency is now a reputational and regulatory lever: whether through EU registries, NIST system cards, or C2PA provenance tags, disclosure and explainability are converging as shared expectations.

At the same time, the gaps are just as telling. Assurance and auditability remain patchy, with OECD and many Asian frameworks offering principles but no enforcement. Incident reporting is fragmented, with serious-incident thresholds in the EU, vulnerability disclosures



GLOBAL AI GOVERNANCE MATRIX

Please see Appendix 4 for the Framework Comparison

in cybersecurity law, but voluntary or absent requirements elsewhere. And sector carve-outs, healthcare (BS 30440), elections/media (C2PA), finance (MAS FEAT) show how quickly vertical domains are building their own overlays.

Together, the alignment matrix shows both the backbone and the blind spots of global AI governance. For organisations, it clarifies where one governance file can travel across borders and where local adaptations are unavoidable. For policymakers, it reveals where convergence is already happening, and where international coordination is most urgently needed.



CROSSROADS ANALYSIS

Overlaps, Obligations and Omissions

01

EU AI Act ↔ ISO/IEC 42001

→ Strong fit on management systems and auditability; ISO provides the organisational backbone for EU compliance.

02

EU AI Act ↔ NIST AI RMF/OECD

→ Divergence: EU mandates conformity and penalties; NIST/OECD remain voluntary. Organisations need an “assurance bolt-on.”

03

NIST AI RMF ↔ OECD/Singapore/Japan

→ Shared voluntary, risk-proportional ethos; diverges from EU because there is no enforcement route.

04

ISO/IEC 42001 ↔ NIST + EU

→ Portable bridge across frameworks; diverges from China’s filings and localisation duties.

05

China GenAI Rules ↔ EU AI Act

→ Shared transparency signals, but diverge sharply from NIST/ISO with licensing, content controls, and localisation.



CROSSROADS ANALYSIS

Overlaps, Obligations and Omissions

06

South Korea AI Framework Act

→ Closer to EU (binding obligations) than Japan/Singapore (voluntary). An APAC enforcement anchor.

07

Security frameworks (NIS2, CRA, ETSI, ENISA, Singapore CSA)

→ Converge strongly; diverge from OECD/NIST where security is optional. Security-by-design is becoming the global passport.

08

Voluntary frontier pledges (G7, Bletchley, Seoul)

→ Align with OECD values and evaluation commitments but lack enforcement. Useful for reputation, not legal sufficiency.

09

Sector specificity matters

→ BS 30440 (healthcare) and C2PA (provenance) show that specialised frameworks can shortcut trust in critical domains.

10

Regional strategies (Brazil, Canada, AU)

→ Mirror EU's scaffolding but with weaker enforcement. They are direction-setting baselines rather than fully operational frameworks.



CROSSROADS ANALYSIS

Overlaps, Obligations and Omissions

The EU AI Act sits closest to ISO/IEC 42001 and ISO/IEC 23894, which both emphasise management systems, evidence gathering, and risk methodologies. These standards can serve as the organisational backbone for EU conformity. However, the Act diverges sharply from voluntary frameworks such as the NIST AI RMF or OECD AI Principles. Where the EU imposes conformity assessments and penalties, NIST and OECD stop at guidance. For organisations, this means mapping NIST or OECD into EU practice requires adding an assurance and audit layer to satisfy regulators.

ISO/IEC 42001 stands out as a bridge. It is structurally close to both the EU AI Act and the NIST RMF, covering evidence, roles, and continuous improvement, while aligning well with lifecycle risk management. Yet it diverges from China's Generative AI and Deep Synthesis rules, which require filings, localisation, and content-level obligations that ISO does not address. For global operators, 42001 offers portability, but must be bolted onto local filings and data residency duties.

The Council of Europe AI Convention connects neatly with the OECD Principles and, to some extent, the EU AI Act, in its rights-centric framing. But it diverges from China's content-moderation posture, where rights are not the organising principle. Its real value lies in providing governments with a baseline umbrella, aligning national laws with human rights guarantees.

China's Generative AI and Deep Synthesis Measures are closer to the EU AI Act in their emphasis on transparency signals, labelling, and risk controls. Yet they diverge from NIST and ISO by imposing administrative filings, licensing, and localisation duties. The significance for enterprises is clear: platforms must prepare for filings, moderation infrastructure, and provenance controls as non-negotiables.

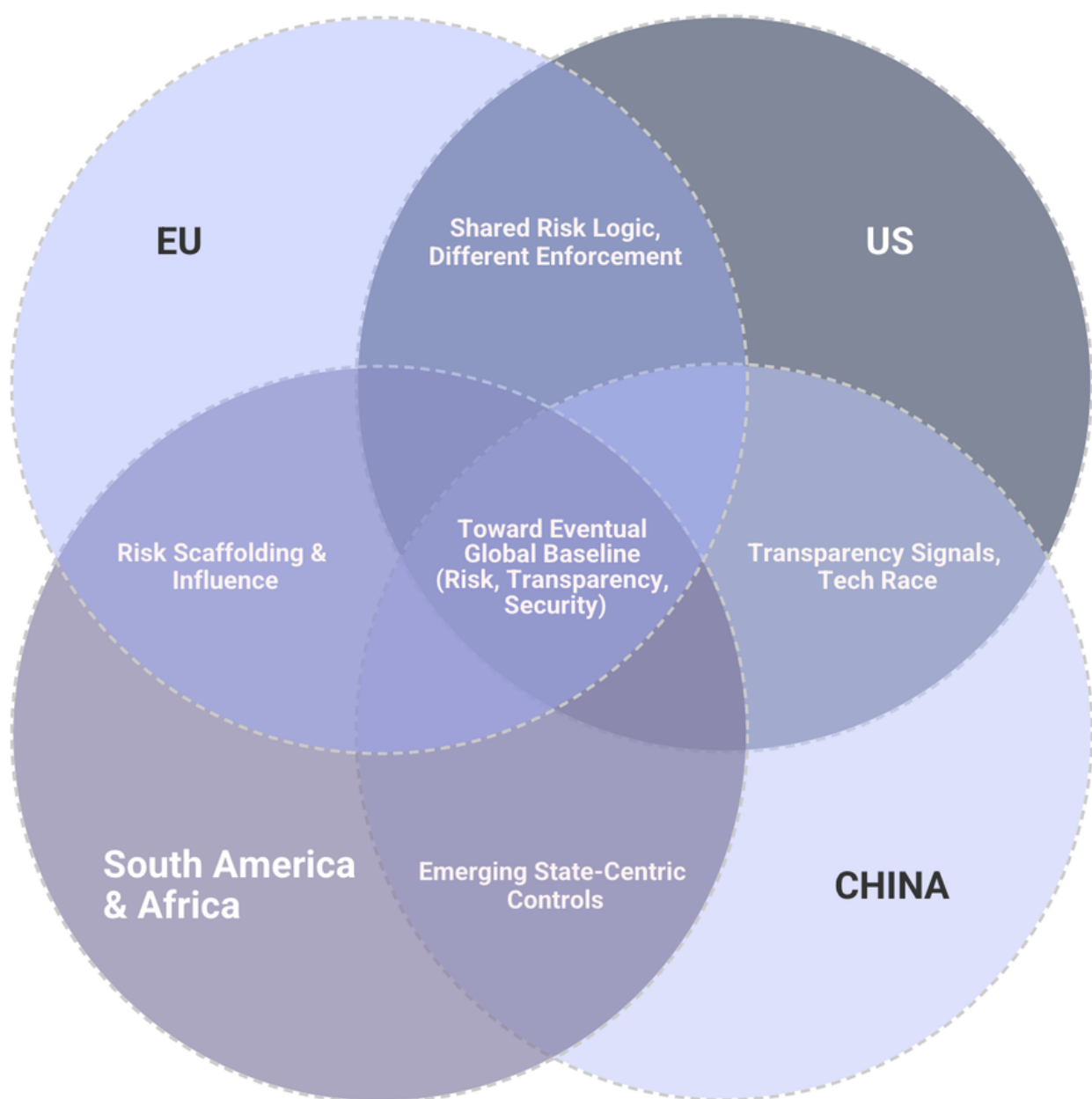
In Asia-Pacific, South Korea's AI Framework Act mirrors the EU's binding obligations and ISO/NIST lifecycle structures, making it the region's enforcement anchor. By contrast, Japan and Singapore remain voluntary, aligning more with OECD and NIST.



CROSSROADS ANALYSIS

Overlaps, Obligations and Omissions

This creates a divergence inside APAC itself: multinationals can design to NIST/ISO and then layer South Korean specifics, but cannot rely on voluntary codes alone.





CROSSROADS ANALYSIS

Overlaps, Obligations and Omissions

Singapore's Model AI Governance and CSA Securing AI guidelines sit closest to NIST RMF and ETSI/ENISA best practices for security, diverging mainly from the EU in their lack of penalties. These are excellent operational templates for multinationals, provided they add EU-style assurance when exporting to Europe.

Japan's AI Guidelines likewise align with NIST and OECD, focusing on proportionality and lifecycle. But they diverge from the EU's enforcement culture, offering no penalties or mandated assurance. These guidelines are useful for product and engineering teams, but will not by themselves deliver external trust. Pairing them with certifiable standards like ISO/IEC 42001 is the pragmatic route.

In the US, the Colorado AI Act (SB 24-205) is the first enforcement foothold, sitting close to the EU AI Act in its risk-based, anti-discrimination duties, while diverging from OECD/NIST in its enforcement posture. For multinationals, it foreshadows a trend of US states moving toward accountability with real penalties.

The NIST AI RMF itself is closest to OECD principles and regional voluntary codes such as Japan's AI Guidelines or Singapore's Model AI Governance. They share a focus on proportionality, lifecycle mapping, and values. But the RMF diverges from the EU model by avoiding CE-style assurance and financial penalties. The practical takeaway is that NIST is an excellent operating model for product teams, but proof of compliance will still need to be adapted to jurisdictions where enforcement is binding.

Canada's AIDA sits between EU and NIST. It borrows risk-tier language from the EU and risk methodology from ISO/IEC 23894, but diverges in its still-immature enforcement provisions. Enterprises should prepare for impact assessments and evidence requirements, even if penalties are still being defined.

Brazil's PL 2338/2023 mirrors the EU's risk-based scaffolding and OECD's principles, but diverges from Europe's hard security mandates. It signals LATAM's baseline direction, organisations should track liability allocation and sector carve-outs as they evolve.



CROSSROADS ANALYSIS

Overlaps, Obligations and Omissions

The NIS2 Directive and the Cyber Resilience Act are closest to ETSI and ENISA security controls, embedding security-by-design and vulnerability handling as mandatory. They diverge from OECD and NIST, where security is optional guidance. The message is simple: these are the EU's non-negotiable spine, and they will increasingly serve as references for AI.

ETSI SAI and ENISA guidance likewise align with NIS2 and Singapore CSA, though they diverge from OECD's advisory stance. Implementing them makes organisations effectively "security-ready" for most frameworks.

At the voluntary frontier, the G7 Hiroshima Principles, Bletchley Declaration, and Seoul Commitments align with OECD values and evaluation pledges but diverge from binding EU enforcement. These are soft-law norm setters: crucial for reputation and buyer trust, but not legally sufficient.

BS 30440 in healthcare is an outlier: it aligns with EU high-risk medical device standards and ISO medical norms, while diverging from general AI frameworks. For health deployments, however, it is the fastest route to regulatory trust.

The C2PA/Content Authenticity Initiative aligns with EU transparency ideals and Seoul commitments on provenance, but diverges from binding assurance, remaining market-driven. It is nonetheless the practical provenance rail for media, elections, and brand integrity.

Finally, the OECD AI Principles align closely with NIST, Japan, and Singapore's voluntary frameworks, but diverge from the binding enforcement models of the EU and China. They remain the lingua franca for dialogue, useful, but only when paired with certifiable or binding standards. Similarly, the African Union AI Strategy aligns with OECD-style principle frameworks but diverges from the EU's assurance detail. It is best read as a regional direction of travel, with member-state hardening still to come.



REGIONAL GOVERNANCE POSTURES

US, EU, UK, China, Africa, Brazil

Exec Summary:

The global map shows two anchors (EU and China), two bridges (US and UK), and two emerging baselines (Africa and Brazil), together defining the poles, pathways, and future foundations of AI governance.

United States

Internal posture: The US remains fragmented. At the federal level, the NIST AI Risk Management Framework, OMB guidance, and executive orders set a voluntary but widely adopted baseline for agencies and suppliers. States like Colorado have stepped in with binding rules targeting high-risk systems and algorithmic discrimination, hinting at a patchwork of state-led enforcement. The approach is pragmatic, risk-proportional, and deliberately light-touch to avoid stifling innovation.

Global position: Compared to the EU, the US lags on enforceable obligations but leads in operational frameworks (NIST AI RMF is the de facto lifecycle vocabulary). It aligns well with Japan and Singapore's voluntary approaches, but diverges from China and the EU on enforcement and penalties. Globally, the US posture is seen as values-driven but under-enforced, creating interoperability challenges for cross-border compliance.

European Union

Internal posture: The EU has taken the boldest step with the AI Act, the world's first comprehensive binding law. It defines risk tiers, mandates conformity assessments, creates enforcement routes via CE marking, and embeds linkages with NIS2 and the Cyber Resilience Act. The approach is heavy on assurance, accountability, and penalties, with harmonised standards under CEN/CENELEC now being drafted.

Global position: The EU is the reference point for binding obligations worldwide. Brazil and Canada mirror its risk-based scaffolding, South Korea echoes its statutory posture, and China overlaps on transparency but differs on administrative filings. Relative to the US and OECD, the EU is stricter; relative to China, more rights-centred. It is the hard anchor in the global matrix.



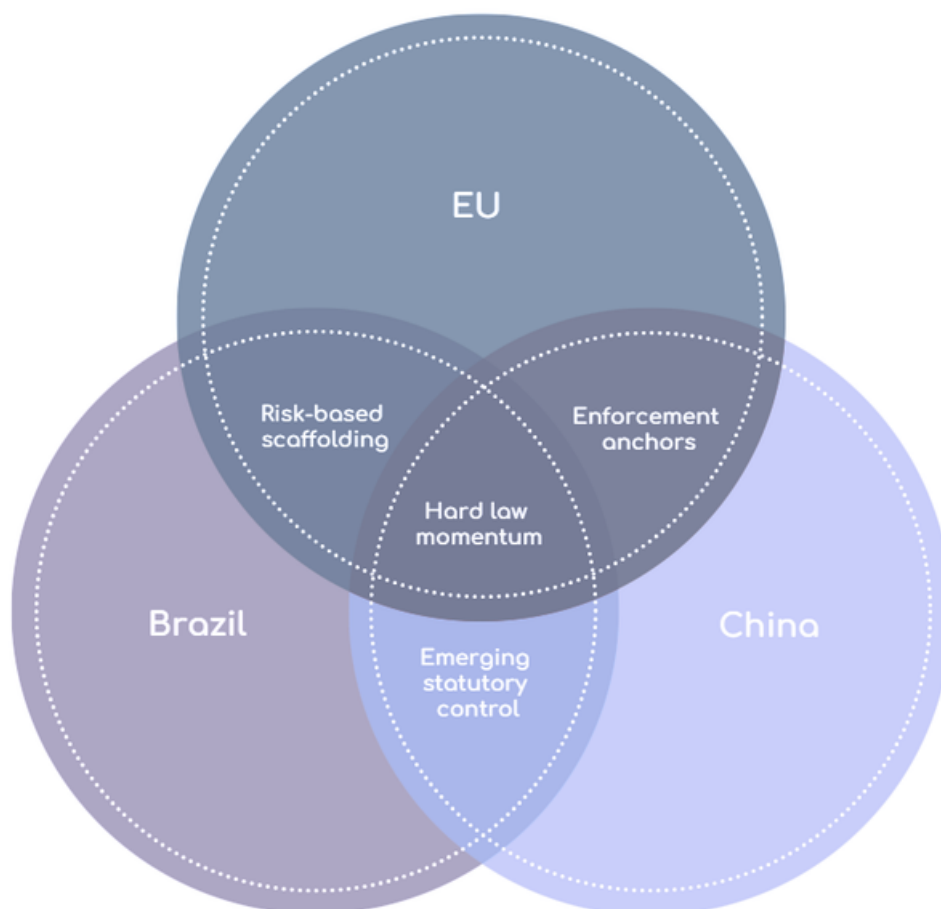
REGIONAL GOVERNANCE POSTURES

US, EU, UK, China, Africa, Brazil

United Kingdom

Internal posture: The UK has avoided a single binding AI law. Instead, it has leaned on regulator-led guidance, voluntary commitments (Bletchley Declaration, Seoul Frontier pledges), and sector-specific overlays. The Office for AI and the Digital Regulation Cooperation Forum encourage cross-sector consistency, but the posture is principle-first, law-later.

Global position: The UK situates itself as a convener and bridge-builder. It aligns closely with OECD, NIST, and G7 principles, diverging from the EU's statutory rigor and China's administrative control. Its global influence is outsized compared to its domestic enforcement, making it the diplomatic broker of AI governance.





REGIONAL GOVERNANCE POSTURES

US, EU, UK, China, Africa, Brazil

China

Internal posture: China has binding regulations already in force, the Generative AI Measures and Deep Synthesis Provisions. These emphasise provider accountability, mandatory filings, algorithm registration, data localisation, and content moderation aligned with “positive values.” The approach is administrative, platform-centric, and tightly integrated with national security objectives.

Global position: China shares overlap with the EU on transparency and risk controls, but diverges from OECD/NIST on enforcement style, localisation, and human-rights baselines. Its approach is more prescriptive than the US and Asia-Pacific peers, positioning China as the regulatory hardliner focused on sovereignty and control.

Africa (AU Continental Strategy)

Internal posture: The African Union’s Continental AI Strategy (endorsed in 2024) sets out principle-level goals: rights, inclusion, capacity building, and security development. Implementation is left to member states, which remain at very different levels of digital and governance maturity. It is a direction-setting rather than enforceable posture.

Global position: Africa aligns with OECD in principle-level commitments, and is beginning to mirror EU-style risk scaffolding in draft laws (e.g., Nigeria, Kenya). But enforcement, assurance, and incident reporting are largely absent, making Africa more aspirational at present. Globally, it is seen as the emerging voice of the Global South, emphasising inclusion and equity in governance debates.

South America (Brazil focus)

Internal posture: Brazil’s PL 2338/2023 sets out a risk-based approach inspired by the EU AI Act, with categories and obligations by role. It has cleared the Senate and is under review in the Chamber. The bill balances innovation with rights, embedding duties for providers and deployers while leaving enforcement structures still to be finalised.

Global position: Brazil positions itself as LATAM’s anchor, with a posture closer to the EU



REGIONAL GOVERNANCE POSTURES

US, EU, UK, China, Africa, Brazil

than to the US or OECD. It aligns on risk categories and transparency but lacks Europe's mature assurance and penalty systems. For now, Brazil is the regional bellwether, shaping a Latin American model that others are likely to follow.



GLOBAL AI & CYBERSECURITY

Timeline

Foundational (1980–2002)

- 1980 – OECD Privacy Guidelines (first international data principles).
- 1981 – Council of Europe Convention 108 (binding data protection treaty).
- 1992 (rev. 2002) – OECD Security Guidelines (culture of security).
- 1995 – EU Data Protection Directive (harmonised EU privacy rules).
- 2002 – US FISMA (federal information security programs).

Building Cybersecurity Baselines (2003–2019)

- 2003 – HIPAA Security Rule (US health data safeguards).
- 2016 – NIS Directive (EU’s first horizontal cybersecurity law).
- 2019 – EU Cybersecurity Act (ENISA mandate + certification framework).

Modern Convergence (2020–2025)

- 2022 – NIS2 Directive adopted (expanded scope, supply chain focus).
- 2023 – NIST AI RMF 1.0 (US voluntary lifecycle AI framework).
- 2023 – China’s AI rules: Deep Synthesis & Generative AI Measures.
- 2024 – EU AI Act adopted (first binding, comprehensive AI law).
- 2024 – EU Cyber Resilience Act (binding security-by-design).
- 2025 – South Korea AI Framework Act (effective 2026, APAC’s first binding AI law).

2025–2030: Implementation Wave

- EU AI Act standards
- GPAI Code of Practice
- CRA obligations
- Model evaluations shift from voluntary and become procurement-grade.



GLOBAL AI & CYBERSECURITY

Timeline



2030–2040: Sectoral Consolidation

- Convergence on incident taxonomies
- Mature certification ecosystems
- Sector playbooks
- Continuous monitoring
- AI Bill of Materials required

2040+: Global Harmonisation

- International AI Safety Accord
- Formalised compute governance
- Universal content authenticity
- Unified secure-by-design standards
- AI ESG reporting as standard.



CONCLUSION

Mapping Today, Preparing for Tomorrow

Artificial Intelligence governance is no longer a distant policy concern, it has become a defining operational challenge for governments, enterprises, and infrastructure suppliers. The Global AI Governance Framework Matrix highlights that while the landscape is fragmented, there are clear signals of both convergence and divergence that organisations must pay attention to.

Key Discoveries from the Alignment Matrix.

The matrix shows that risk management, security-by-design, and transparency are emerging as the universal backbone of AI governance, the areas where regulators, standards bodies, and industry are finding common ground. Yet alignment does not mean uniformity. Alongside these points of convergence, the matrix also exposes areas of sharp divergence, the gaps, inconsistencies, and conflicting obligations that create uncertainty for organisations operating across borders. These divergences are just as important to understand as the alignments, because they define where friction, compliance fatigue, and regulatory risk are most likely to arise.

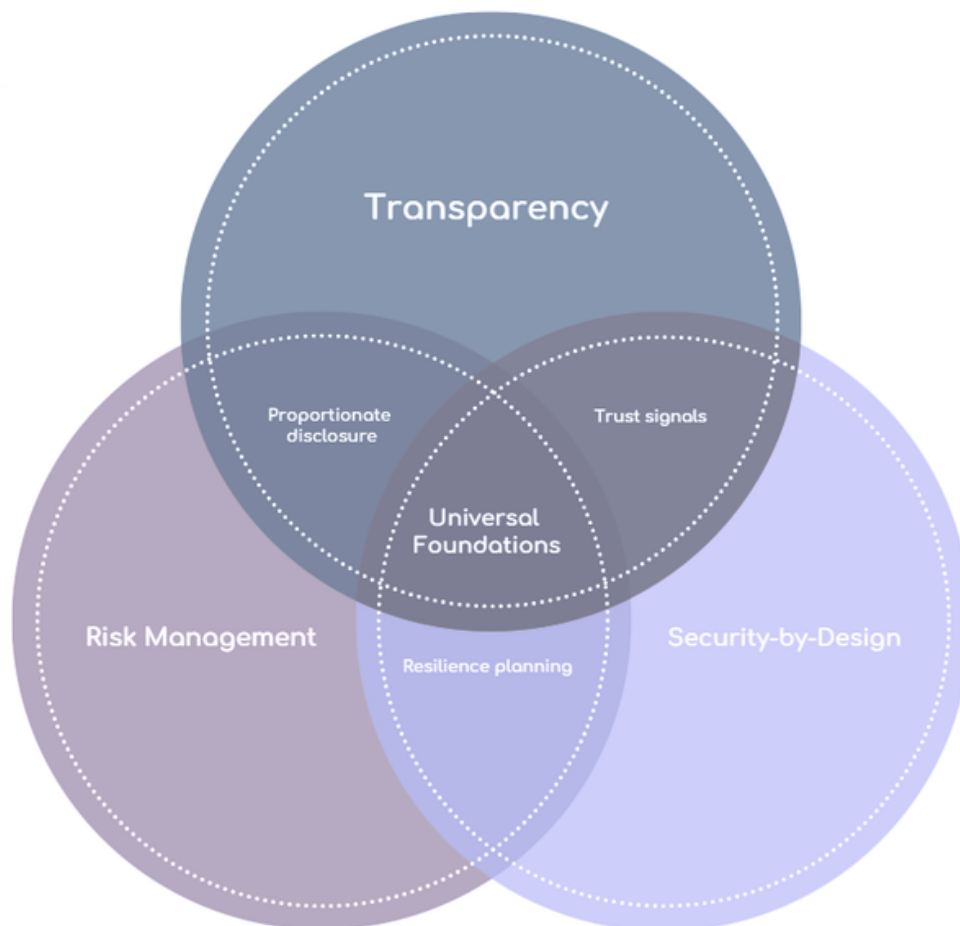
The three foundational principles:

Across binding laws, voluntary frameworks, and technical standards, three core principles consistently appear: risk management, transparency, and security-by-design. The terminology varies (risk tiers in the EU AI Act, lifecycle functions in the NIST AI RMF, management controls in ISO/IEC 42001) but the underlying logic remains the same. These common denominators provide the foundations for interoperability and a shared global baseline.



CONCLUSION

Mapping Today, Preparing for Tomorrow



It seems this convergence exists because these principles are universally recognisable, technically implementable, and politically defensible. Risk management gives regulators a structured way to scale obligations without stifling innovation. Transparency offers the public and policymakers reassurance that systems can be scrutinised, explained, and corrected. Security-by-design responds to widespread concern about AI misuse, cyber vulnerabilities, and resilience in critical infrastructure. Together, these common denominators form the most pragmatic foundation for interoperability and a shared global baseline. The areas where consensus is both easiest to achieve and hardest to ignore.



CONCLUSION

Mapping Today, Preparing for Tomorrow

1. Risk management as the organising principle

Whether framed as risk tiers in the EU AI Act, high-impact categories in Canada's AIDA, or lifecycle mapping functions in the NIST AI RMF, risk management is the organising principle of modern AI governance. ISO/IEC 23894 builds an entire methodology around risk identification, assessment, and treatment, while voluntary codes like the OECD AI Principles echo the expectation that AI should be deployed only in proportion to its risks.

This centrality exists because risk offers a scalable, adaptable way to govern a fast-moving technology. Unlike static rules, risk-based approaches can flex across contexts. from low-stakes consumer chatbots to high-stakes medical diagnostics. Regulators see it as a pragmatic way to balance innovation with safety, while enterprises value it as a structured framework that translates directly into compliance programs and audits. Risk management therefore functions as the common language across otherwise divergent frameworks, making it the natural starting point for global alignment.

Advice: Start with risk. By assessing risks early, organisations can understand value, prioritise controls, and make decisions before regulatory obligations harden into costly compliance gaps.

2. Security as a bridge between frameworks

Every major binding law now explicitly references cybersecurity obligations, often tied to wider digital product security frameworks such as NIS2 and the Cyber Resilience Act. China requires algorithm filings and security assessments, while South Korea's AI Framework Act mandates resilience and security-by-design. Even voluntary frameworks (OECD, G7, Bletchley, Seoul commitments) echo the importance of robustness, albeit without the penalty of enforcement.

This convergence arises because security is a non-negotiable foundation for trust. Failures in security create immediate, visible harms to individuals and sovereign entities alike. Data breaches, model manipulation, adversarial attacks are situations and potential catastrophes that regulators simply cannot ignore. By embedding security obligations into AI governance,



CONCLUSION

Mapping Today, Preparing for Tomorrow

lawmakers reassure both the public and industry that systems will be resilient against misuse. For enterprises, security also offers a common compliance denominator: if you are secure under one framework, you are better prepared to satisfy others. This is why security has become the anchor point for global convergence the place where alignment is happening fastest.

Advice: Double down on cybersecurity. A strong security baseline not only reduces exposure to direct threats but also acts as a passport for operating across multiple jurisdictions, where resilience and robustness are increasingly the shared measure of trust.

3. Transparency as a global expectation

From the EU's mandatory technical documentation and public registries, to the NIST RMF's system cards, to Singapore's disclosure guidelines, transparency has become a near-universal obligation. Even aspirational frameworks such as the OECD AI Principles elevate transparency as a core requirement, making it one of the few principles that cuts across binding law, voluntary guidance, and industry pledges.

The reason transparency travels so well is because it serves multiple audiences simultaneously. Policymakers see it as the antidote to opaque or biased decision-making. Regulators use it to enforce obligations without needing insider technical knowledge. Enterprises rely on it to build investor and consumer trust. And society demands it as the baseline for accountability in systems that affect lives and rights. This flexibility makes transparency the most politically defensible of the core principles, and therefore a safe bet for convergence across regions.

Advice: Invest in transparency. Clear documentation, model cards, and disclosure processes not only satisfy regulators but also differentiate organisations as trustworthy partners in a marketplace where credibility is as valuable as compliance.

Together, these three principles form the strongest points of alignment in today's governance landscape. They are the areas where policymakers, regulators, and enterprises



CONCLUSION

Mapping Today, Preparing for Tomorrow

are already speaking a common language, and where early convergence offers a path toward interoperability. However, the alignment Matrix also shows that beyond these shared foundations, significant gaps remain. To understand the real challenges of operating across jurisdictions, we must turn to the divergences. The points at which frameworks pull apart, creating friction, uncertainty, regulatory gaps and risk.

The three main divergences:

1. Assurance & enforcement (the proof-and-penalty gap)

Across jurisdictions, what counts as “proven” compliance and what happens if you fall short remain wildly inconsistent.

The EU AI Act and ISO/IEC 42001 define structured assurance routes (e.g., conformity assessment, third-party certification), while voluntary frameworks leave “sufficient proof” open to interpretation and even abuse. At the same time, penalties range from severe (EU fines, China’s administrative filings and sanctions, Colorado AG powers) to non-existent (Japan, Singapore, OECD/G7 soft-law). The result is a proof-and-penalty gap: multinationals are at risk of either over-investing in assurance to cover every possibility, or under-scoping and being out of bounds when operating in stricter markets. At the moment due to the politically charged nature of AI in the world today, this feels, region to region, on purpose.

Advice: Design for the strictest credible regime and let compliance “downshift” where lighter. Build a living risk register, documented testing logs, evaluation reports, audit trails, and decision logs that can slot into EU-style assurance or satisfy lighter frameworks without rework. Plan for the strictest frameworks. If your systems can withstand EU-level enforcement and penalties, they will almost always meet or exceed requirements elsewhere. If you can afford it, it will turn compliance from a burden into a competitive advantage.



CONCLUSION

Mapping Today, Preparing for Tomorrow

2 - Clarity and precision - Innovation Vs control

One of the sharpest divides in global AI governance is how clear or prescriptive obligations actually are.

A major fault line in AI governance lies in how precisely obligations are drafted. The EU AI Act, South Korea's Framework Act, and China's GenAI rules are highly prescriptive, specifying risk tiers, conformity routes, and filing duties with little scope for discretion. By contrast, the OECD Principles, Japan's guidelines, and Singapore's voluntary codes rely on broad statements of intent, leaving significant latitude for interpretation.

This divergence cuts both ways. Loose, principle-based frameworks encourage innovation by giving companies space to adapt practices without the drag of rigid requirements. But the same flexibility creates loopholes and uncertainty, as organisations second-guess what "good enough" compliance looks like. Conversely, high-precision frameworks provide certainty and comparability but risk slowing product cycles and raising entry barriers, especially for SMEs. The result is a compliance landscape where a system may be tightly constrained in Brussels, loosely guided in Singapore, and barely addressed in Nairobi.

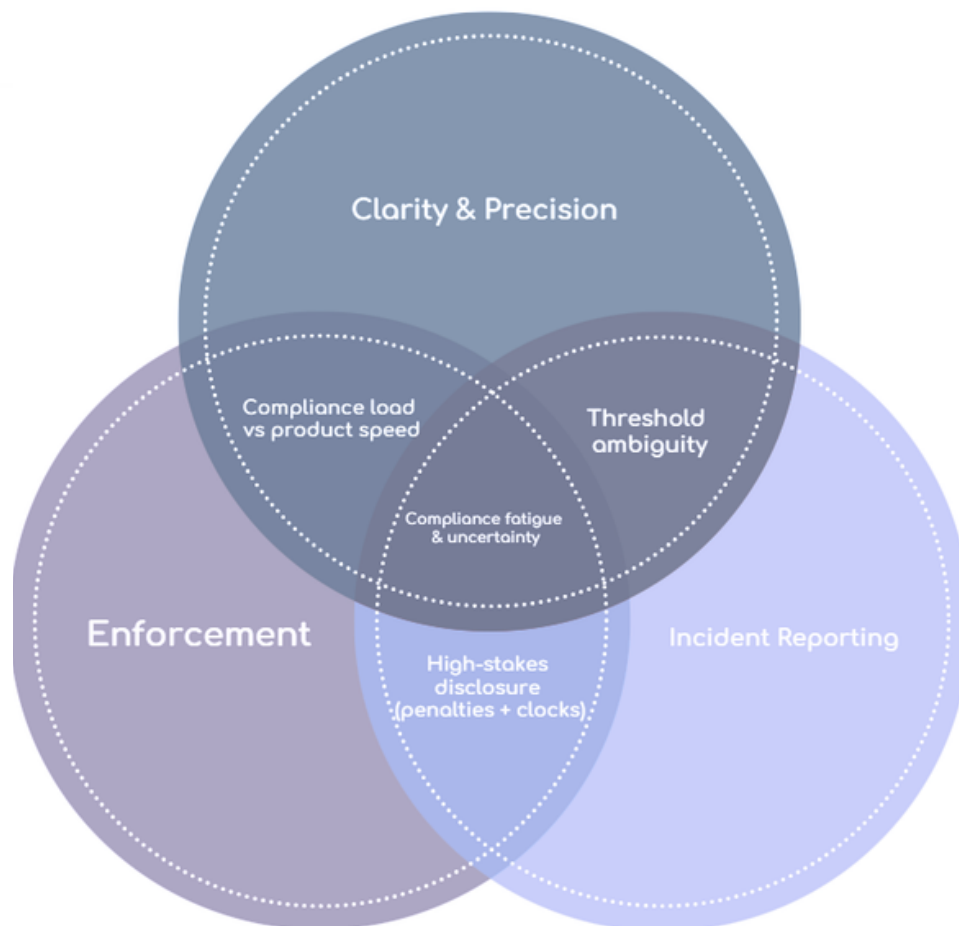
Advice: For innovators: Use flexibility to test and adapt, but document choices so they can be lifted into stricter regimes later.

For enterprises/operators: Be precise. Detail risks from the start and align to the rules of the markets you know you'll operate in.



CONCLUSION

Mapping Today, Preparing for Tomorrow



3. Incident reporting is fragmented and unclear

In Europe, mandatory reporting obligations apply under the EU AI Act, NIS2, and the Cyber Resilience Act. Elsewhere, incident reporting is encouraged (Singapore, OECD) or remains undefined (African Union, Brazil). Even within mandatory systems, the thresholds differ: what counts as a “serious incident” in Brussels may not even trigger reporting in Tokyo.

This divergence persists because no unified global consensus on harm definitions in AI exists.



CONCLUSION

Mapping Today, Preparing for Tomorrow

Advice: Adopt a “report more, not less” strategy. Build a centralised incident register that captures harms and near-misses consistently across your operations. This provides the flexibility to adapt reporting outputs to each jurisdiction without reinventing the process every time. International standards can also help here. ISO27035 (incident management) requires organisations to establish processes for incident detection, internal reporting, escalation, response, and lessons learned. ISO42001 (AI management systems) requires monitoring AI systems, documenting incidents and near-misses, and feeding these into continual improvement. These standards provide a solid baseline for consistent handling and record-keeping, but they are not legally binding. They are a good starting point but the regulators in each region still impose their own thresholds and obligations.

So, What happens next?

2025–2030: Implementation beats legislation:

The next decade will be less about drafting new flagship laws and more about operationalising and deepening existing frameworks. In Europe, the EU AI Act will be translated into practice through standards, with the GPAI Code of Practice and the EU Cyber Resilience Act driving security hardening across supply chains. Model evaluation programmes, today largely voluntary, are already on the path to becoming expected in government procurement, especially in critical infrastructure and frontier AI contexts. Companies that don’t produce the required AI governance reporting will be left out of procurement cycles. At the same time, the governance ecosystem will mature: convergence on incident taxonomies, certification schemes tied to ISO/IEC 42001 and sectoral standards, and the adoption of playbooks requiring safety cases, continuous monitoring, and an AI bill of materials (A-BOM) will become baseline expectations.



CONCLUSION

Mapping Today, Preparing for Tomorrow



2030–2040: Trade-linked obligations and global alignment:

Looking further ahead, the trajectory points towards global harmonisation and trade-linked obligations. An International AI Safety Accord is likely to emerge as the rights-and-safety counterpart to existing cyber treaties, while compute governance frameworks may formalise capacity thresholds and reporting duties for large-scale training runs. Provenance and authenticity systems, such as C2PA, will evolve into universal, tamper-evident layers, embedded at the camera, model, and platform level. Secure-by-design requirements for AI will become as ubiquitous and non-negotiable as today's CE or ISO marks, effectively serving as trade passports for digital products. Alongside these technical baselines, sustainability and governance pressures are converging: the World Economic Forum's call for AI governance and sustainability metrics foreshadows a world where AI ESG reporting becomes as standard as financial or carbon disclosure.



WHAT THIS MEANS FOR YOU

(Now vs Next)

For **leaders**, the immediate priority is to bake model evaluation and AI-BOM deliverables directly into procurement frameworks, while mapping product-security frameworks such as the CRA to AI-specific controls, beginning with common harm and incident taxonomies. They must also champion cross-border alignment toward an International AI Safety Accord.

Enterprises and suppliers, meanwhile, should focus on building a portable governance file that includes a risk register, evaluation results, incident logs, transparency packs, AI-BOMs, and provenance policies. Security-by-design and continuous monitoring must be treated as non-negotiable, with practices aligned to ISO/IEC 42001 and the NIST AI RMF. At the same time, piloting C2PA and authenticity measures across all synthetic outputs will ensure readiness for the procurement mandates that are rapidly emerging.

Multilateral bodies and **standards organisations** have a pivotal role to play in accelerating convergence: advancing incident taxonomy alignment, creating assurance templates, defining practical compute thresholds and evaluation benchmarks, and harmonising ESG-for-AI metrics so reporting becomes comparable, auditable, and procurement-ready.

The alignment map is only the beginning, what comes next is a steady march toward enforceable, interoperable, and globally recognised norms that treat AI governance as an essential condition for trust, trade, and long-term sustainability.



FINAL WORDS

It is clear that convergence is real, but far from complete. Risk management, transparency, and security-by-design are fast becoming the language of AI governance, offering organisations a pragmatic foundation for interoperability. Yet the divergences (assurance, enforcement, incident reporting and even the very definition of harm) are equally decisive, shaping the real-world friction points that governments and enterprises must navigate.

The task ahead is not just compliance. It is building adaptive, resilient governance models that thrive across jurisdictions as the regulatory map shifts. Organisations that treat governance as a strategic discipline rather than a compliance burden will be the ones best placed to operate, innovate, and lead in the global AI economy.



RESEARCHER REFLECTIONS



The Global AI Governance Alignment Map has been created to make sense of the chaos. We're in an era where every nation & company wants its own AI rulebook, but the technology itself doesn't recognise borders or follow orders. The challenge isn't just to keep up with the pace of innovation, it's to connect the dots between innovation, leadership, and policy in a way that actually works in practice. What stood out most through the research was that three ideas are doing the heavy lifting worldwide: risk management, transparency, and security-by-design. These are the shared languages

between developers, policymakers, and auditors. But the real story sits between the lines in the divergences. How assurance is enforced in Europe but voluntary elsewhere. How precision helps compliance but can slow creativity. And how "incident reporting" and the penalties that come with breaking the rules means different things depending on which side of a border you're on. This balance, between ambition, assurance, innovation and control, is what will define the next decade. We can't afford frameworks that smother progress, or innovation that ignores our social contract. Unfortunately this means that to get it right means treating governance as a living system, not just a checklist. And that is hard. It's hard because it requires building a governance principle that evolves through standards, shared taxonomies, and mutual trust. Ultimately, this isn't about regulation for its own sake. It's about global readiness and building the connective tissue that lets organisations innovate responsibly, lead confidently, and adapt to whatever comes next. Recommendations for policymakers and practitioners:

1. Prioritise interoperability: Anchor national frameworks in shared principles to avoid duplicative obligations.
2. Operationalise trust: Move from voluntary pledges to measurable standards and AIBOMs (AI Bills of Materials) that prove responsibility in practice.
3. Enable innovation through clarity: Draft governance that leadership teams can act on and communicate.
4. Unify incident reporting: Standardise harm definitions and timelines.
5. Plan for convergence: Treat ISO/IEC 42001, NIST AI RMF, and the EU AI Act as complementary not competing.

Mike Wood

Research Associate, Responsible AI Trust | Cohort 1:12

Head of Technology Operations & Compliance Lead - Hadean Supercomputing Ltd

Author, Global AI Governance Alignment Map - Responsible AI Trust



ADVISOR REFLECTIONS



This paper captures the state of global AI governance with clarity and discipline. It presents a world that is still fragmented but starting to align. Its focus on three shared anchors, risk management, transparency, and security by design, translates policy language into something leaders can act on.

ISO/IEC 42001 is the connective tissue in that landscape. The paper shows how it links the European Union's binding AI Act with the voluntary, lifecycle-based structure of the NIST AI RMF. This connection provides a path toward consistency and certification. Without it, compliance becomes reactive and repetitive. With it, governance becomes measurable, auditable, and adaptable across markets.

The paper also addresses what remains unresolved. Enforcement varies by jurisdiction. Rules differ in clarity and precision. Incident reporting still lacks a unified standard. These inconsistencies highlight where leadership must focus. Risk is not just a control function. It is the framework that aligns purpose with accountability and turns governance into an operational discipline.

For executives, three priorities emerge:

- Build a complete AI governance file that contains risk registers, evaluation records, transparency documentation, and AI bills of materials.
- Treat security by design as a baseline requirement.
- Prepare for a market where proof, not intent, defines trust.

The closing message is straightforward. The organizations that succeed will be those that operationalize governance, not those that simply draft policies. AI governance is no longer about appearing responsible. It is about demonstrating it through evidence, structure, and continuous improvement.

Patrick Sullivan

Vice President, Innovation & Strategy | A-LIGN

ISO/IEC 42001 Lead Implementer | SC 42 Delegate | AI Governance Practitioner

Advisor & Reviewer, Global AI Governance Alignment Map - Responsible AI Trust



RESEARCH DIRECTOR REFLECTIONS



The Global AI Governance Alignment Map captures more than a comparison of frameworks, it reflects a global moment of alignment in how trust in intelligent systems is defined.

Mike Wood's research exposed the structural common ground forming across jurisdictions, while Patrick Sullivan's reflections reminded us that proof, not principle, is becoming the new currency of

responsibility. Together, it demonstrates that governance is no longer theoretical; it is the language of implementation.

Across every framework we studied, three constants emerged: risk management, transparency, and security-by-design. These are no longer optional ideals, they are the foundations of interoperability, and the earliest signals of a shared global standard. Our task at Responsible AI Trust is to operationalise these anchors so that enterprises, regulators, and innovators can navigate complexity with evidence and confidence.

This brief is part of a broader mission: transforming AI governance from compliance checklists into measurable trust systems. As we continue building that foundation, each collaboration, each associate, advisor, and partner, moves us closer to a world where responsible AI is not declared, but demonstrated.

This alignment map represents the first structured demonstration of measurable AI governance. It anchors Responsible AI Trust's mission to turn global compliance into proof-based trust: a direction I'm proud to lead alongside our global advisors and researchers.

Lehar Gupta

Founder & Research Director, Responsible AI Trust

Former AI Product and Governance Lead - Reuters, Altium, University of Cambridge, Gov.uk

Co-Author, Global AI Governance Alignment Map - Responsible AI Trust

✉ Lehar@ResponsibleAITrust.com



APPENDIX

Appendix 1 - Primary Legal and Regulatory Texts

EU AI Act (Regulation (EU) 2024/1689, Official Journal)

Council of Europe AI Convention (2024, opened for signature)

NIS2 Directive (Directive (EU) 2022/2555)

EU Cyber Resilience Act (Regulation (EU) 2024/2847)

National frameworks:

- United States: NIST AI RMF 1.0 (2023), CSF 2.0 (2024), Executive Order 14110 (2023)
- China: Generative AI Measures (2023), Deep Synthesis Provisions (2023)
- South Korea: AI Framework Act (2025, effective 2026)
- Colorado AI Act (SB 24-205) (2024, effective 2026)
- Canada: AIDA (Bill C-27, pending)
- Brazil: PL 2338/2023
- Japan: AI Business Guidelines (2024)
- Singapore: Model AI Governance Framework (2024 update), CSA Securing AI Systems (2024)
- African Union: Continental AI Strategy (2024)

International and Voluntary Frameworks OECD AI Principles (2019, updated 2024)

G7 Hiroshima Process and Code of Conduct (2023)

Bletchley Declaration (2023)

Seoul AI Safety Summit Commitments (2024)

GPAI integration with OECD (2024–25 work programme)

BRICS Declaration on AI Governance (2025)

Standards and Technical Guidance

ISO/IEC 42001 (AI Management System Standard, 2023)

ISO/IEC 23894 (AI Risk Management Guidance, 2023)

ETSI Securing AI (SAI) reports

ENISA Threat Landscape for AI (2023/2024)

BS 30440 (Healthcare AI validation)

C2PA / Content Authenticity Initiative documentation

National Cybersecurity & AI-Security Guidance

UK NCSC & US CISA: Guidelines for Secure AI System Development (2023)

Singapore CSA: Securing AI Systems (2024)

US OMB M-24-10 on federal AI use (2023)

OECD.AI Policy Observatory (global database of AI policies)



APPENDIX

Appendix 2 - Links and sources of information

Description	Link	Country / Region	Industry
Vendor-neutral explainer and hub covering AI governance concepts, major frameworks (EU AI Act, NIST, ISO/IEC 42001) and best practices; good “getting started” reference. (Modulos)	https://www.modulos.ai/guide-to-ai-governance/#ai-governance-frameworks-and-acts	Global	Cross-sector
Blog taxonomy that positions ISO/IEC 42001 within a broader governance stack; helpful for mapping policies → controls → certification. (Modulos)	https://www.modulos.ai/blog/ai-governance-taxonomy-iso-42001-and-beyond/	Global	Cross-sector
Two-page “cheat sheet” comparing EU AI Act, NIST AI RMF, and ISO/IEC 42001 (scope, objectives, who it targets); useful side-by-side for execs. (Trustible)	https://www.trustible.ai/post/cheat-sheet-comparing-eu-ai-act-nist-ai-rmf-and-iso-42001	Global (EU/US focus)	Cross-sector; Compliance
AI Seoul Summit (2024): list of Frontier AI Safety Commitments made by leading companies; voluntary but sets expectations around model evaluations, reporting, and risk management. (GOV.UK)	https://www.gov.uk/government/publications/frontier-ai-safety-commitments-ai-seoul-summit-2024/frontier-ai-safety-commitments-ai-seoul-summit-2024	Global (hosted by UK/KR)	Frontier AI developers; Platforms
Bletchley Declaration (2023): intergovernmental statement on AI safety, risk identification, and international cooperation - sets the tone for later G7/GPAI moves. (GOV.UK)	https://www.gov.uk/government/publications/ai-safety-summit-2023-the-bletchley-declaration/the-bletchley-declaration-by-countries-attending-the-ai-safety-summit-1-2-november-2023	Global	Cross-sector; Public policy
GPAI programme page (now under OECD stewardship): multi-stakeholder initiative supporting practical projects and policy bridges across jurisdictions. (OECD)	https://www.oecd.org/en/about/programmes/global-partnership-on-artificial-intelligence.html	Global (OECD)	Cross-sector



APPENDIX

Appendix 2 - Links and sources of information

Description	Link	Country / Region	Industry
Hiroshima AI Process (G7): official info hub (Japan); anchors the International Guiding Principles and Code of Conduct for organisations developing advanced AI. (Ministry of Foreign Affairs of Japan)	https://www.soumu.go.jp/hiroshimaaiprocess/en/index.html	G7 / Global	Frontier AI developers
OECD AI Principles: high-level, widely adopted policy baseline informing national frameworks and corporate governance. (OECD)	https://www.oecd.org/en/topics/sub-issues/ai-principles.html	Global (OECD)	Cross-sector
African Union – Continental AI Strategy: continental priorities and coordination for member states; signals regional direction of travel. (African Union)	https://au.int/en/documents/20240809/continental-artificial-intelligence-strategy	Africa (AU)	Cross-sector; Government
China regulatory tracker (White & Case): up-to-date overview of PRC AI measures (e.g., deep synthesis, generative AI) with links to primary rules. (White & Case)	https://www.whitecase.com/insight-our-thinking/ai-watch-global-regulatory-tracker-china	China	Cross-sector
Brazil AI Act explainer: status and content of PL 2338/2023 (Senate-approved; risk-based approach; penalties; authority). Includes link to Portuguese text. (Artificial Intelligence Act)	https://artificialintelligenceact.com/brazil-ai-act/	Brazil	Cross-sector
Canada AIDA – Companion Document: official explainer of the proposed Artificial Intelligence and Data Act (Bill C-27); scope, oversight, and risk-based approach. (Innovation Canada)	https://ised-isde.canada.ca/site/innovation-better-canada/en/artificial-intelligence-and-data-act-aida-companion-document	Canada	Cross-sector



APPENDIX

Appendix 2 - Links and sources of information

Description	Link	Country / Region	Industry
TAKE IT DOWN Act (S.146, 119th Congress): US bill targeting removal of intimate images/deepfakes; relevant to content authenticity and online harms. (Congress.gov)	https://www.congress.gov/bill/119th-congress/senate-bill/146	United States (Federal)	Online platforms; Safety
Colorado AI Act (overview by Skadden): first comprehensive US state-level AI law (SB 24-205) addressing high-risk AI and discrimination risk; effective 2026 with rulemaking. (Skadden)	https://www.skadden.com/insights/publications/2024/06/colorados-landmark-ai-act	United States (Colorado)	Cross-sector
Legal analysis of proposed deepfake laws in Denmark and Netherlands; highlights fit/misfit with copyright frameworks - useful for media/political content risks. (Legal Blogs)	https://legalblogs.wolterskluwer.com/copyright-blog/deepfake-bills-in-denmark-and-the-netherlands-right-idea-wrong-legal-framework/	Denmark; Netherlands	Media & platforms; Elections
Council of Europe – Framework Convention on AI: first global treaty tying AI to human rights, democracy, rule of law; open to non-members. (Portal)	https://www.coe.int/en/web/artificial-intelligence/the-framework-convention-on-artificial-intelligence	Council of Europe / Global	Cross-sector; Public sector
Council of the EU – AI policy hub: official page covering the EU AI Act journey, context, and related digital policy files. (Consilium)	https://www.consilium.europa.eu/en/policies/artificial-intelligence/#0	European Union	Cross-sector



APPENDIX

Appendix 3 - Cybersecurity Timeline





APPENDIX

Appendix 4 - Global AI Governance Alignment Matrix ([link](#))

Framework	Type	Region	Status	Scope/Definitions	Risk/Applicability	Lifecycle Controls	Transparency/Docs	Cybersecurity Linkage	Human Oversight	Incident Reporting	Assurance Route	Enforcement	Cross-Border/Data	Sector Cave-outs	Primary Source (URL)
CPA/ Content Authenticity Initiative	Technical standard/spec	Global	Active	Content provenance & signing	Applies to media supply chains	Provenance capture; manifests	User-facing provenance signals	Crypto signing; integrity	Governance of provenance policies	Not applicable (tech spec)	Conformance by ecosystem vendors	Market driven	Global interoperability	Media, elections relevance	https://cpa.ai/
UK - Blockchain & Social Commitments AIS	Voluntary + Program	UK / Global	Active (commitments); program active	Frontier AI safety; evaluations	Frontier capabilities focus	Pre-deployment mitigations	Safety reports; disclosure commitments	Model security/testing emphasis	Governance pledges	Voluntary reporting mechanisms	Evaluation programs (AIS)	Non-binding	Global cooperation focus	N/A	https://www.gov.uk/government/publications/blockchain-social-commitments-ais
Cyber Resilience Act (CRA)	Binding law (EU)	EU	Adopted (planning)	Products with digital elements	Risk classes; obligations	Secure-by-design; vuln handling; updates	CE elect. security info	Core product security law	Manufacturer responsibilities	Vulnerability/incident reporting	Conformity assessment & CE	Market surveillance; fines	EU internal market	Critical products lists	https://eur-lex.europa.eu/eli/reg/2022/2554/oj
OECD AI Principles	Voluntary principles	OECD / Global	Adopted by many states	High-level principles for trustworthy AI	Principle-based	Encouraged risk mgmt and monitoring	Transparency & explainability	Robustness & security principle	Human-centred values	Encouraged	None	Non-binding	Cross-border cooperation principle	None	https://oecd.ai/
Colorado AI Act (SB 24-209)	Binding law	USA (Colorado)	Enacted (future effective)	High-risk AI; algorithmic discrimination	Context risk-based; consumer impact focus	Risk assessments; governance program	Notices; documentation to AI upon request	Reasonable security; vendor mgmt	Internal accountability roles	Duty to notify AI for material impacts	Self-attestation; records	Attorney General enforcement	General privacy laws apply	Certain exclusions (e.g. small biz thresholds (likely))	https://leg.colorado.gov/bills/24-209
South Korea - AI Framework Act	Binding law	South Korea	Enacted (planning)	Trustworthy AI; duties across ecosystem	Risk-based; obligations scale by use	Governance; testing; monitoring	Labelling; user info (contextual)	Security-by-design; resilience	Oversight & accountability required	Report obligations TBD by right	Certification schemes anticipated	Regulator enforcement & penalties	Interop with PIPA; data transfer rules	National security exceptions	https://www.mlt.go.kr/eng
African Union - Continental AI Strategy	Regional strategy	Africa (AU)	Endorsed	Policy priorities; capacity building	Principle- and risk-based outline	Governance model recommendations	High-level commitments	Digital security capacity goals	Rights & inclusion emphasis	To be defined nationally	N/A	Non-binding strategy	Regional cooperation	N/A	https://au.int/en/press-releases/2022/07/2022-07-20-african-union-continental-artificial-intelligence-strategy
United States - OMB/NIST Federal Baseline	Policy + Framework	USA (Federal)	In effect (policy); framework voluntary	Federal agency AI use; AI tool scope	Contextual risk; critical uses flagged	Govern/Map/Measure/Manage	Impact assessments; inventories	NIST CSF/SSDF alignment	Agency oversight roles	Agency incident playbooks	Self-assurance; ATO processes	OMB oversight for agencies	Federal data rules apply	N/A (agency-specific)	https://www.nist.gov/ai
China - Generative AI & Deep Synthesis Measures	Binding regulations	China	In force	Generative services; synthetic media providers	Obligations on providers/platforms	Data quality; safety testing; tagging	Labelling; policy disclosures	Security assessment; algorithm filing	Content moderation responsibilities	Reporting to regulators	Filing/recordal frameworks	Administrative penalties; takedown	Data/localization & CAC rules	News/political content strict rules	https://www.cac.gov.cn/en/202304/20230401_6682620.htm
EU AI Act	Binding law	EU	In force (planning)	AI system; provider/supplier; high-risk; prohibited or used in EU	Tiered: Prohibited, High, Limited, Prohibited AI market obligations	Risk mgmt; data governance; testing/evals; PAM	Tech docs; instructions; model/system info; high-risk registry	Robustness; resilience; OAS/NSI interfaces	Required for high-risk	Submit incidents to national authorities/EU AI Office	Conformity assessment CE marking (many high-risk)	Administrative fines; market surveillance	Interacts with GDPR; no separate transfer framework	Military excluded; specific LE allowances	https://eur-lex.europa.eu/eli/reg/2024/426/oj
Brazil PL 2338/2023 (AI Bill)	Proposed law	Brazil	Legislative process	AI systems; rights & liability	Risk categories; duties by role	Risk mgmt; testing; monitoring	Disclosures; documentation	Security controls	Human-in-the-loop for critical uses	Serious incidents (death)	Conformity pathways TBD	Fines; authorities to be designated	LGPD Interplay	Public security exceptions possible	https://www2.camara.leg.br/legislacao/imprensa/2023/07/2023-07-20-proposta-de-lei-2338-2023
Canada AIDA (Bill C-27)	Proposed law	Canada	Legislative process	AI systems; high-impact AI	Risk-based; high-impact duties	Risk mgmt; data & monitoring controls	Public disclosures; record keeping	Security safeguards	Accountability; roles	Serious incident reporting	Ministerial reg; audits possible	Administrative/penalties	Interop with privacy reforms (COPPA)	TBD by regs	https://www242.intel.gc.ca/legislation/bills/2022/bill-27
ETSI SA1 (Securing AI)	Standard/Reports	Europe / Global	Ongoing publications	Security for AI & AI for security	Threat/mitigation focused	Secure design; testing; deployment	Security documentation	Core overlaps with NIS2/CIA	Operational controls & review	Security incident processes	Standards-based conformance (optional)	None (unless referenced)	Neutral	None	https://www.etsi.org/
Council of Europe AI Convention	Treaty (binding for parties)	Council of Europe / Global	Open for signature	AI aligned to HR, democracy, rule of law	Principle-based; party obligations	Safeguards; risk-based measures	Due diligence; transparency principles	Security and resilience as safeguards	Human rights-centred oversight	Encouraged; national implementation	Domestic mechanisms; cooperation	Treaty compliance via parties	Cross-border cooperation provisions	National security typically reserved	https://www.coe.int/t/e/treaties/ai-convention/
G7 - Hiroshima Principles & Code of Conduct	Voluntary principles/code	G7 / Global	Published	Advanced/Frontier AI developers	Risk-based expectations	Model safety testing; mitigations	Policies; objectives; records	Security testing & controls	Governance expectations	Encouraged	Voluntary adherence	Non-binding	International cooperation	N/A	https://www.g7summit.org/2023/07/13/hiroshima-principles-and-code-of-conduct-on-ai/
ISO/IEC 42001 (AIMS)	Standard (certifiable)	Global	Published	Core-level AI management system	Any org; scope defined for cert	Plan/Do/Check/Act for AI	Disclosures; documentation prompts	Integrates with ISO 27001/CS controls	Roles & responsibilities	Nonconformity & corrective action	3rd party certification available	Marked/contractual	Neutral, align with privacy standards	None	https://www.iso.org/standard/75401.html
ISO/IEC 23884 (AI Risk Management)	Standard (guidance)	Global	Published	Risk mgmt for AI across lifecycle	Contextual risk approach	Risk identification + treatment	Risk documentation	Security risks considered	Governance of risk decisions	Feedback & improvement	No certification (guidance)	None	Neutral	None	https://www.iso.org/standard/75401.html
Japan - AI Business Guidelines (MET/MIIC)	Guidance (non-binding)	Japan	Published	Business use of AI; ecosystem actors	Contextual risk approach	Governance; testing; monitoring	Disclosures; documentation prompts	Security & robustness expected	Roles/responsibilities	Recommended processes	Voluntary adoption	None (policy guidance)	Interop with Appi privacy framework	N/A	https://www.met.go.jp/en/press/2023/07/2023-07-20-010.html
Singapore - Model AI Governance & CSA Securing AI	Guidance (non-binding)	Singapore	Published	AI governance for industry; security for AI systems	Risk-based controls	Data testing; deployment; monitoring	Model/system info; disclosures	CSA security guidance; adversarial ML	Operator controls	Recommended practices	Voluntary, used in procurement	None (guidance)	Interop with PDPA; CBRP participation	Financial sector links (MAS FEAT)	https://www.mta.gov.sg/
BS 39440 (Healthcare AI)	Standard (sectoral)	UK / Global	Published	Healthcare AI validation	Health-sector risk view	Clinical validation; monitoring	Documentation & evidence	Safety & security expectations	Clinician oversight	Health incident frameworks	Conformance assessed by bodies	Via national regulator self-referenced	Health data governance context	Healthcare	https://www.bsigroup.com/
ENISA - AI/Cybersecurity Guidance	Guidance	EU	Published/updated	AI threat landscape; good practices	Risk-based recommendations	Secure lifecycle practices	Security evidence expectations	Aligns with NIS2	Operational responsibility	Under NIS2 frameworks	N/A	Indirect via ENISA/CSA references	EU context	Operators of essential services	https://www.enisa.europa.eu/content/cybersecurity
NIS2 Directive	Binding law (EU)	EU	Transposition phase	Essential/important entities	Risk mgmt; reporting duties	Security measures across lifecycle	Policies; evidence on request	Core cybersecurity directive	Mgmt accountability	Mandatory reporting	Audit/inspections by authorities	Fines & remedial action	Cooperation across MS	Sectoral specifics via annexes	https://eur-lex.europa.eu/eli/dir/2022/2555/oj



RESPONSIBLE
AI TRUST

NOVEMBER 2025

BUILDING TRUST IN INTELLIGENT SYSTEMS BEGINS WITH YOU.

At Responsible AI Trust, we believe alignment is not a debate, it's a design principle. Every framework, every audit, every standard leads to one shared goal: verifiable trust. Whether you're an enterprise **leader** shaping governance strategy, a **researcher** mapping global standards, or a **sponsor** advancing responsible innovation, your participation helps define how the world governs AI.

Join us in building the foundation for measurable, interoperable, and accountable AI.

✉ Lehar@ResponsibleAITrust.com

🌐 @ResponsibleAITrust

● [View the latest live version](#)



Join for updates on upcoming briefs, research insights, and governance resources.

 responsibleaitrust.substack.com

Powered by



MetricsLM

| Compliance Passport for AI