

BLOG 02 — THE EIGHT TIERS, PART 0

Hundreds of AI frameworks distilled into eight

Sort by kind, not by count, and the landscape becomes a map

By Robin Beetge

Co-founder, laibrary

Part 0 of nine

Series opener · ungated

Ask how many AI frameworks exist and you will get an answer somewhere between three and a thousand, depending entirely on who you ask and what kind of week they are having.

Three, if you count only comprehensive laws in force. Around thirty-three, if you count binding legal instruments of any kind. Roughly ninety, if you count national governance frameworks. Approximately nine hundred, if you count formal technical standards. Several thousand, if you are willing to include every corporate ethics charter, every multilateral declaration and every set of principles a consultancy has ever published with a diagram attached.

3 comprehensive laws in force	33 binding legal instruments	90 national frameworks	900 technical standards	8 kinds of instrument
--	---	----------------------------------	-----------------------------------	---------------------------------

All of those numbers are defensible. None of them helps you.

I spent the better part of a year cataloguing this landscape properly, and the useful discovery was not a number. It was that the entire mess sorts cleanly into **eight kinds of instrument**, and that the kind tells you everything the count does not: what the thing demands of you, whether it can fine you, what evidence it wants, and when.

Sort by kind and the fog clears. Here is the map.

Tier 1 Comprehensive AI statutes

What it does to you: deadlines and fines.

These are laws that regulate AI horizontally, across every sector, in one instrument. Everyone talks about them. There are three in force worldwide.

The **EU AI Act** is the anchor, and the one every other regime gets measured against. Risk-tiered, extraterritorial, with penalties reaching €35 million or 7% of global turnover. **South Korea's AI Framework Act** came into force in January 2026, the first mandatory comprehensive AI law in Asia-Pacific, and it reaches foreign companies through user thresholds rather than establishment. **Japan's AI Promotion Act** is a statute with no penalties by deliberate design, which makes it a genuinely interesting hybrid: binding in form, promotional in substance.

Brazil's bill is close. Canada's died on the order paper in January 2025 and has not been revived.

That is the complete list. The most discussed category on the board is also the smallest.

Tier 2 International treaties

What it does to you: nothing directly, and it predicts almost everything.

There is exactly one: the **Council of Europe Framework Convention on Artificial Intelligence**, opened for signature in September 2024. The first legally binding international treaty on AI in history.

Signatories include the UK, the US, Canada, Japan, Israel, Norway, Switzerland, Ukraine and the EU, which became the first party to ratify in May 2026. Ratification is the slow part; entry into force requires five.

Almost nobody outside the profession watches this, and they should. A treaty binds states, not you, but ratification status is one of the cleanest leading indicators available for where binding domestic law arrives next — particularly in jurisdictions where treaties take direct effect.

Tier 3 Sub-national statutes

What it does to you: binds you today, probably without your knowledge.

If you sell into the United States, this is where your real obligations live while everyone else is reading about Brussels.

Texas brought its Responsible AI Governance Act into force in January 2026. California has four separate instruments running, covering frontier developers, training-data disclosure, automated decision-making and content provenance. Colorado repealed its landmark AI Act in May 2026 and replaced it with a narrower successor effective January 2027. Illinois amended its Human Rights Act. New York's frontier developer regime starts in January 2027. New York City has been running mandatory bias audits since 2023.

Over two thousand AI bills have been introduced across the states.

And here is the part that catches people. **Reach attaches through where your users are, not where your company is.** You cross into a state's jurisdiction silently, through ordinary growth, and you usually find out about it from a customer's security questionnaire rather than from a regulator.

Tier 4 Sectoral regulators

What it does to you: binds you today, through the regulator you already have.

This tier gets ignored in almost every AI governance conversation, which is remarkable given it is the one most likely to produce an examination finding in the next twelve months.

US banking supervisors replaced fifteen years of model risk management doctrine in April 2026 — and expressly placed generative and agentic AI *outside* the new framework's scope, leaving banks with a supervisory expectation and no supervisory instrument. More than twenty US states have adopted the insurance regulators' model bulletin, which requires a written AI programme enforced through market conduct examination. The FDA reversed its position on clinical decision support in March 2026. Certified health IT has carried structured AI transparency obligations since the start of 2025.

If you are in a regulated industry, Tier 4 binds you now, regardless of what any AI-specific law says or when it arrives.

Tier 5 Certifiable management standards

What it does to you: gives you something to be audited against, and a certificate at the end.

ISO/IEC 42001 is the first AI management system standard you can actually be certified to. It shares its structure with ISO 27001 and ISO 9001, so if you have run either, the shape is familiar: a management system, a control annex, a plan-do-check-act cycle.

The family matters as much as the flagship. ISO/IEC 23894 supplies the risk methodology. ISO/IEC 42005, published in 2025, covers impact assessment. And ISO/IEC 42006, also 2025, sets the requirements for the bodies doing the certifying — which is the unglamorous standard that made credible third-party certification possible at all.

This is the closest thing we have to a globally portable AI governance credential. It is also, usefully, the thing a procurement team can understand without reading a regulation.

Tier 6 Risk management frameworks

What it does to you: technically nothing. In practice, quite a lot.

The **NIST AI Risk Management Framework** is entirely voluntary. Four functions — Govern, Map, Measure, Manage — across 19 categories and 72 subcategories, with a companion profile for generative AI carrying twelve GenAI-specific risk categories.

And yet: Texas made substantial compliance with it a statutory enforcement safe harbour. US federal agencies work to it under OMB direction. It is the most common baseline in American enterprise vendor questionnaires. It crosswalks more cleanly than anything else to ISO/IEC 42001, the EU AI Act and the state statutes.

A voluntary framework that became load-bearing more or less by accident. If you build to one thing, build to this and ISO/IEC 42001 together.

Tier 7 Security and threat taxonomies

What it does to you: turns governance into test cases.

This is where AI governance stops being paperwork and starts being engineering.

OWASP's Top 10 for LLM Applications, refreshed in August 2026, was the first edition built on incident evidence rather than opinion alone — nearly eight thousand real-world incidents weighted against a community vote. Excessive Agency climbed to number three. A new entry, Hidden Context Exposure, widened the attack surface from system prompts to retrieved documents, memory, application state and tool responses. There is a companion Top 10 for agentic applications with its own ten-item taxonomy.

MITRE ATLAS is the ATT&CK-style knowledge base of real adversary techniques against AI: 16 tactics, 84 techniques, with case studies. The **Cloud Security Alliance's AI Controls Matrix** does the same job for cloud deployments. **ISO/IEC 27090** covers AI-specific security threats and is at final publication stage.

If your governance programme has no Tier 7 content in it, you have written a policy, not a control.

Tier 8 Principles and soft law

What it does to you: supplies the vocabulary that the binding tiers are written in.

The **OECD AI Principles**, **UNESCO's Recommendation on the Ethics of AI**, the **G7 Hiroshima Process**, and the UN's new Scientific Panel and Global Dialogue. Non-binding, routinely dismissed as talking shops, and quietly the most influential tier of all.

Here is the proof. When the EU needed a definition of "AI system" for a law carrying fines of up to 7% of global turnover, it did not invent one. It reached for the OECD's, updated in May 2024 to accommodate generative AI. That definition now determines who is in scope.

Principles do not bind you. They are where the binding language gets drafted, three to five years early, in public, by people who will tell you exactly what they are doing.

The map, compressed

TIER	KIND	ENFORCEABLE?	WHAT IT WANTS FROM YOU
1	Comprehensive statutes	Yes, with fines	Conformity, documentation, registration
2	International treaties	On states	Nothing yet. Watch it
3	Sub-national statutes	Yes, with fines	Notices, audits, assessments
4	Sectoral instruments	Yes, via your regulator	Programme evidence at examination
5	Certifiable standards	Voluntary, auditable	A working management system
6	Risk frameworks	Voluntary, referenced in law	A risk methodology you can show
7	Threat taxonomies	Voluntary, operational	Test cases and controls
8	Principles	No	Definitions, and early warning

The thing the table doesn't show

The tiers are not independent. That is the insight the whole exercise was worth.

Tiers 5, 6 and 7 are the operational vocabulary through which Tier 1 to 4 obligations actually get satisfied and evidenced. When a notified body asks a provider to demonstrate a risk management system under the EU AI Act, what the provider hands over is ISO/IEC 42001 and 23894 artefacts. When a bank examiner asks about AI model validation, the answer is built from NIST AI RMF language. When a regulator asks how you tested for adversarial robustness, the credible answer is OWASP and ATLAS.

Which leads somewhere genuinely useful. Strip the jurisdictional packaging away from every instrument in all eight tiers and the same twelve requirements come back over and over:

01 Know what you have.	02 Classify it.	03 Manage the risk.	04 Assess the impact.
05 Govern the data.	06 Document it.	07 Log it.	08 Disclose it.
09 Keep a human in the loop.	10 Make it accurate and secure.	11 Watch it after launch.	12 Govern your suppliers.

Twelve. Not hundreds. The fragmentation is real at the instrument layer and mostly illusory at the requirement layer.

So you do not build a programme per framework. You build one programme against twelve requirements, produce evidence once, and map it many times. A single remediation should close findings in four places at once, because those four places are asking for the same thing in four dialects.

That is the whole strategy, and it only becomes visible once you stop counting frameworks and start sorting them.

NEXT IN THE SERIES

This is part 0 of a nine-part series. Next: Tier 1, the three comprehensive AI statutes in force worldwide, and why the list is so much shorter than the noise suggests.