

People, Processes, Platforms: A Coding Framework and Comparative Benchmark for Global AI Governance

Shera Potka

University of Victoria
Department of Computer Science
Victoria, British Columbia, Canada

Jens Weber

University of Victoria
Department of Computer Science
Victoria, British Columbia, Canada

Abstract

What determines whether an AI regulatory framework functions as a coherent governance system or a collection of independent provisions? Existing comparative analyses catalogue cross-jurisdictional differences but offer limited structural explanation for why they arise. This paper introduces a 17-sub-dimension coding framework based on the People–Processes–Platforms (PPP) model and applies it to 12 regulatory instruments across 10 jurisdictions, producing a provision-level comparative matrix of global AI governance. All 170 classifications are validated against an independent AI-assisted coding pipeline (87.6% agreement, $n=170$), with reliability predicted by regulatory text legibility rather than document complexity. From this dataset, we identify two structural features that account for substantial variation in the observed data. First, AI system classification functions as a *cross-dimensional regulatory trigger*: jurisdictions with formal classification schemes (EU, China, Colorado, Canada) exhibit cascading obligation structures averaging approximately 12 mandatory provisions, while those without (UK, Singapore, Japan, OECD, UNESCO) average fewer than one. This difference is primarily driven by regulatory architecture rather than stated ambition. Second, we identify an *institutional creation threshold* that bounds global governance convergence: provisions implementable through existing institutional capacity are addressed in 76% of jurisdiction–sub-dimension pairs, while provisions requiring new governance infrastructure (certification, incident reporting, regulatory sandboxes) are addressed in only 27%. This pattern holds across all regulatory philosophies and explains why voluntary frameworks plateau at 65% coverage. Among the jurisdictions studied, these findings suggest that the most consequential barriers to harmonization may not be principled disagreements over values but differences in regulatory architecture and the capacity to create new governance institutions. We derive five testable propositions and maintain the coding framework as a living benchmark at <https://ppp-ai-governance.vercel.app>.

Keywords

AI regulation, AI governance, comparative regulatory analysis, People–Processes–Platforms, regulatory trigger, AI-assisted coding

1 Introduction

In 2024, the European Union enacted the most comprehensive binding AI legislation to date, addressing governance requirements ranging from human oversight to foundation model transparency. In the same year, the United States federal government addressed none of these dimensions with binding force, though state-level legislation began to emerge. Between these poles, eight other jurisdictions and international bodies occupy distinct positions, yet

no systematic instrument exists to compare their approaches at the provision level. This paper introduces such an instrument.

The absence of a common analytical framework for AI regulatory comparison creates three problems. First, **analytical fragmentation**: most scholarship examines individual instruments or pairs of jurisdictions in isolation, leaving researchers without a shared vocabulary for cross-regime comparison. Second, **comparison without structure**: cross-jurisdictional analyses proceed descriptively, cataloguing what each jurisdiction does, without enabling dimension-level analysis of *how* and *why* provisions differ. Third, **practical coordination barriers**: organizations deploying AI across borders cannot systematically identify areas of regulatory convergence and divergence.

This paper addresses these problems through four contributions. First, we introduce an **operationalized PPP coding framework** with 17 sub-dimensions, validated through an AI-assisted coding protocol achieving 87.6% agreement across 170 classifications ($n=170$). Second, we produce a **systematic provision-level comparative matrix** of global AI governance, covering 12 instruments across 10 jurisdictions as of early 2026. Third, we identify two **structural determinants of regulatory coherence**: (a) AI system classification functions as a *cross-dimensional regulatory trigger*: jurisdictions with classification average approximately 12 mandatory provisions while those without average fewer than one; and (b) global governance convergence is bounded by an *institutional creation threshold*: provisions implementable through existing capacity are addressed in 76% of cases, while provisions requiring new institutional infrastructure are addressed in only 27%, independent of regulatory philosophy. Fourth, we derive **five testable propositions** that reframe the global AI governance gap as a structural problem of regulatory architecture and institutional capacity rather than a problem of principled disagreement.

The survey is guided by five research questions:

- RQ1:** How do major AI regulatory frameworks distribute governance requirements across the People, Processes, and Platforms dimensions, and what configuration patterns emerge?
- RQ2:** What governance philosophies underpin different frameworks, and how do these philosophies shape the relative emphasis placed on People, Processes, and Platforms provisions?
- RQ3:** On which PPP sub-dimensions do global AI regulatory frameworks exhibit the greatest convergence, and on which do they diverge most significantly?
- RQ4:** How do interdependencies between People, Processes, and Platforms dimensions manifest within and across regulatory regimes?

RQ5: What regulatory gaps emerge from the comparative analysis, and what do they imply for prospects of international regulatory interoperability?

Our analysis covers regulatory instruments in force, officially adopted, or historically significant as of early 2026, including the EU AI Act, the U.S. Executive Order on Safe, Secure, and Trustworthy AI (2023; revoked January 2025, retained as a historical reference), the NIST AI Risk Management Framework (AI RMF 1.0), the Colorado Artificial Intelligence Act (2024), China’s Deep Synthesis Provisions (2022) and Generative AI Measures (2023), Canada’s Directive on Automated Decision-Making, the UK’s Pro-Innovation AI Regulation White Paper (2023), Singapore’s Model AI Governance Framework, Japan’s AI Promotion Act (2025), the OECD AI Principles, and the UNESCO Recommendation on the Ethics of Artificial Intelligence.

The remainder of this paper is organized as follows. Section 2 reviews the evolution of AI governance and positions this survey relative to existing literature. Section 3 presents the PPP analytical framework, its theoretical foundations, and our operationalization into 17 sub-dimensions. Section 4 describes the methodology, including the AI-assisted coding approach. Sections 5–7 present the comparative analysis organized by PPP dimension. Section 8 provides cross-dimensional analysis, including a regulatory philosophy typology, convergence and divergence patterns, and gap identification. Section 9 discusses implications for international harmonization and the utility of the PPP framework. Section 10 concludes.

2 Background and Related Work

2.1 The Evolution of AI Governance

AI governance has evolved through three broadly distinguishable phases. The first phase (2016–2019) was characterized by the proliferation of *ethical principles and voluntary guidelines*. Jobin, Ienca, and Vayena (2019) documented 84 AI ethics guidelines published during this period, revealing broad convergence on principles such as transparency, fairness, and accountability, but limited operationalization. Hagendorff (2020) argued that most guidelines lacked enforcement mechanisms and failed to address power asymmetries. The OECD AI Principles, adopted in May 2019 and initially adhered to by 42 countries, marked the culmination of this phase as the first intergovernmental standard on AI, structured around five values-based principles and five policy recommendations.

The second phase (2019–2023) saw a transition toward *binding regulatory instruments*. China became an early mover, enacting the Provisions on the Management of Deep Synthesis Internet Information Services (effective January 2023) and the Interim Measures for the Management of Generative Artificial Intelligence Services (effective August 2023), among the earliest binding regulations specifically targeting generative AI. Canada’s Directive on Automated Decision-Making, originally issued in 2019 and subsequently amended, established mandatory algorithmic impact assessments for federal government institutions. The UNESCO Recommendation on the Ethics of Artificial Intelligence, adopted in November 2021 by 193 member states, represented the broadest normative consensus achieved. As Cath (2018) and Stahl (2021) had argued, effective AI governance requires bridging ethical principles with enforceable

institutional mechanisms, a transition that the second phase began but did not complete. In the United States, the Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence (October 2023) directed federal agencies to manage AI risks; it was subsequently revoked in January 2025, but is included in this survey as a historical reference reflecting the federal government’s AI governance posture during the study period. The NIST AI Risk Management Framework (AI RMF 1.0, January 2023) provided a voluntary Govern–Map–Measure–Manage lifecycle and remains in effect.

The third and current phase (2024–present) is characterized by *comprehensive legislation and implementation*. The EU Artificial Intelligence Act (Regulation (EU) 2024/1689) entered into force on August 1, 2024, establishing a four-tier risk classification system with phased implementation: the Act is fully applicable from August 2, 2026, with certain product-safety high-risk system rules applying from August 2, 2027. The Colorado Artificial Intelligence Act (SB 24-205), signed in 2024, introduced mandatory impact assessments for high-risk AI systems at the state level, with an effective date subsequently extended to June 30, 2026 by SB25B-004. Japan enacted the AI Promotion Act in 2025, establishing an institutional framework for AI governance through the AI Strategy Headquarters. The UK, while maintaining its principles-based approach articulated in the 2023 White Paper, has signaled forthcoming comprehensive legislation.

2.2 Comparative AI Regulation: Existing Surveys

A growing body of literature examines AI regulation comparatively. Veale and Zuiderveen Borgesius (2021) and Madiega (2021) analyzed the EU AI Act’s risk-based architecture, while Bradford (2020) theorized the EU’s capacity to export regulatory standards globally through the “Brussels Effect.” Smuha (2021) framed the international landscape as a “race to AI regulation” driven by regulatory competition. Roberts et al. (2021) provided a detailed analysis of China’s distinctive approach, characterizing it as state-directed governance with content-safety priorities. Kerry et al. (2023) assessed the EU AI Act’s likely global impact, finding significant but limited regulatory diffusion. Mökander and Floridi (2023) explored how AI governance principles can be operationalized through ethics-based auditing.

However, the existing literature exhibits four gaps that this survey addresses:

- (1) **Jurisdiction-organized analysis:** Most comparative studies organize by country, producing descriptive accounts that limit systematic cross-jurisdictional comparison at the provision level. To our knowledge, no existing survey applies a consistent coding scheme across all major frameworks.
- (2) **Unoperationalized frameworks:** The People–Process–Technology triad has been discussed conceptually in AI governance (Coglianese, 2023) but has not been formalized into a coding instrument applied across a broad set of regulatory texts.
- (3) **Temporal coverage:** The regulatory landscape has changed substantially since most surveys were published. The EU

AI Act entered into force in 2024, Japan enacted the AI Promotion Act in 2025, and the Colorado AI Act takes effect in 2026.

- (4) **Missing interdependency analysis:** Few studies examine how provisions in one regulatory dimension condition provisions in another. For example, how AI classification schemes determine the intensity of human oversight requirements.

2.3 The People–Process–Technology Triad in Governance

The conceptual foundation for our analytical framework originates in Harold J. Leavitt’s (1965) model of organizational systems, which conceptualized organizations as composed of four interdependent components: people, structure, technology, and task. Leavitt’s central insight was that change in any one component produces compensatory changes in the others, a systems perspective that influenced decades of organizational research.

Over subsequent decades, Leavitt’s model was distilled into the People–Process–Technology (PPT) triad, gaining prominence in IT management and information security, as reflected in the work of Schneier (2000) and others who emphasized that security depends on the alignment of people, processes, and technology. Baxter and Sommerville (2011) extended the socio-technical systems perspective, arguing that effective technology governance requires integrated attention to human actors, organizational procedures, and technical systems. The triad’s core proposition has proven durable across domains.

In the AI governance context, Coglianese (2023) advanced a “people-and-processes approach,” analyzing Executive Order 14110 through a management-based regulatory lens. Coglianese argued that effective AI governance cannot rely on prescriptive technical standards alone because AI systems are heterogeneous and rapidly evolving. Instead, governance must focus on ensuring that organizations have the right *people* (with appropriate expertise, authority, and accountability) and the right *processes* (including impact assessments, auditing, and continuous monitoring). This argument provides direct scholarly precedent for applying the PPP framework to comparative regulatory analysis.

3 The PPP Analytical Framework

3.1 From Concept to Coding Scheme

This survey extends the PPP framework from a conceptual lens to an operationalized analytical instrument. We decompose each dimension into sub-dimensions defined with sufficient precision to enable consistent coding across regulatory texts of different legal forms and jurisdictional origins.

The framework comprises 17 sub-dimensions: 5 in the People dimension, 6 in the Processes dimension, and 6 in the Platforms dimension. The asymmetric distribution (5–6–6) reflects the different granularity required by each dimension rather than an artificial symmetry constraint. The People dimension naturally decomposes into five broad categories of human governance (oversight, accountability, institutional roles, workforce capacity, and individual rights). The Processes and Platforms dimensions each require six sub-dimensions because collapsing any would merge provisions

that are regulated independently across jurisdictions: for example, certification (PR4) and incident reporting (PR5) are institutionally distinct and show different adoption patterns, as do content labeling (PL5) and GPAI provisions (PL6). Table 1 presents the complete coding scheme.

3.2 Rating Scale

Each sub-dimension is coded per jurisdiction on a three-level regulatory intensity scale:

- **Mandatory (M):** The framework imposes a legally binding requirement with specified obligations. Non-compliance may result in enforcement action or penalties.
- **Recommended (R):** The framework officially encourages or expects this practice but does not impose a binding legal obligation.
- **Absent (–):** The framework does not address this sub-dimension.

This scale is deliberately coarse to enable comparison across instruments of different legal forms (binding legislation, executive directives, voluntary frameworks, intergovernmental standards). Finer-grained distinctions are captured in the textual evidence accompanying each coding decision.

3.3 Framework Justification and Limitations

The PPP framework is adopted for three reasons. First, *analytical coverage*: the three dimensions collectively address the principal concerns of AI regulation: who is responsible (People), what procedures govern development and deployment (Processes), and what technical requirements apply to the systems themselves (Platforms). This aligns with Floridi and Cowls’ (2019) observation that AI governance must simultaneously address ethical principles, institutional processes, and technical safeguards. Second, *scholarly precedent*: the framework builds on established organizational systems theory (Leavitt, 1965; Baxter & Sommerville, 2011) and has been specifically applied to AI governance by Coglianese (2023). Third, *comparative utility*: the tripartite structure produces a manageable dimensionality for cross-jurisdictional comparison while remaining fine-grained through sub-dimensions. We note that alternative frameworks exist, including the OECD’s five-principle structure and Baldwin, Cave, and Lodge’s (2012) regulatory instruments taxonomy, but these either lack the dimensional symmetry needed for matrix-based comparison or were not designed for provision-level coding across legal forms.

The framework has known limitations. The 17 sub-dimensions were derived inductively from provisions that appear in at least one instrument in the corpus; they do not claim to exhaust all possible governance concerns. Aspects such as liability regimes, environmental impact of AI training, competition law, and intellectual property are relevant to AI governance but fall outside the PPP framework as defined here. Claims of “complete coverage” (e.g., the EU’s 17M profile) refer to coverage of the 17 sub-dimensions in this coding scheme, not to all conceivable regulatory provisions. Additionally, PPP does not directly capture *regulatory philosophy* or *institutional architecture* (centralized vs. distributed enforcement,

Table 1: PPP Coding Scheme: Dimensions, Sub-dimensions, and Definitions. The framework decomposes AI regulatory provisions into 17 sub-dimensions across three interdependent dimensions. Each sub-dimension is defined with sufficient precision to enable consistent deductive coding across regulatory texts of different legal forms and jurisdictional origins. Codes P1–P5 address human actors, PR1–PR6 address governance procedures, and PL1–PL6 address technical systems.

Dim.	Code	Sub-dimension	Definition
People	P1	Human oversight	Mandated human review, intervention capability, or human-in-the-loop requirements
	P2	Accountability	Allocation of responsibility and liability among developers, deployers, users, and regulators
	P3	Regulator roles	Designated regulatory bodies, coordination mechanisms, and stakeholder consultation
	P4	Workforce training	Requirements for training, upskilling, or AI literacy
	P5	Affected persons’ rights	Rights of individuals subject to AI-assisted decisions
Processes	PR1	Risk assessment	Required assessment of AI system risks prior to or during deployment
	PR2	Auditing & compliance	Mechanisms for internal or external audit, monitoring, and compliance verification
	PR3	Transparency	Obligations for disclosure, explainability, or public reporting
	PR4	Certification	Conformity assessment, certification schemes, or standardized governance procedures
	PR5	Incident reporting	Requirements for reporting AI-related incidents or harms
	PR6	Regulatory sandboxes	Provisions for controlled testing environments
Platforms	PL1	AI classification	Risk-tiered categorization of AI systems determining regulatory treatment
	PL2	Infrastructure	Requirements for computing infrastructure or deployment standards
	PL3	Data governance	Standards for training data, input data quality, or data provenance
	PL4	Safety & robustness	Technical standards for reliability, robustness, and cybersecurity
	PL5	Content labeling	Requirements for marking AI-generated content
	PL6	GPAI provisions	Specific regulatory treatment of general-purpose AI or foundation models

binding vs. voluntary instruments). We address this by supplementing the PPP coding with a separate characterization of each jurisdiction’s regulatory approach (Section 8). The boundaries between dimensions are not always sharp; we code provisions to the dimension of their primary regulatory effect, with cross-references noted. Finally, PPP originates in organizational theory, not regulatory theory; we draw on regulatory governance scholarship (Baldwin, Cave & Lodge, 2012) to contextualize the analysis.

4 Methodology

4.1 Research Design

The survey employs a qualitative comparative documentary analysis structured in three phases: (1) document collection and corpus assembly, (2) structured coding against the PPP scheme, and (3) comparative analysis using a jurisdiction-by-sub-dimension matrix.

4.2 Document Corpus

Table 2 presents the regulatory instruments included in the analysis. Documents were included if they met four criteria: (1) officially adopted, in force, or with confirmed implementation dates as of early 2026; (2) binding legislation, executive directives, official policy frameworks, or intergovernmental standards; (3) available in English (original or official translation); (4) within the ten jurisdictions selected for regulatory significance, typological diversity, and geographic coverage. The EU, United States, and China serve as primary cases (detailed provision-level coding); remaining jurisdictions serve as secondary cases for comparative contrast.

The inclusion of the Colorado AI Act, the sole subnational instrument in the corpus, warrants brief justification. The U.S. federal instruments (Executive Order 14110, revoked in January 2025, and NIST AI RMF 1.0) are, respectively, a now-superseded executive

directive and a voluntary framework; neither imposed enforceable obligations on private-sector AI developers or deployers. The Colorado Artificial Intelligence Act (SB 24-205), by contrast, is the first U.S. state law to establish comprehensive, binding requirements for high-risk AI systems, including mandatory impact assessments and a deployer duty of care. Including it alongside the federal instruments captures the full spectrum of U.S. AI governance, from voluntary guidance to enforceable law, and provides a necessary binding counterpoint without which the United States would appear to lack any mandatory AI regulation. We acknowledge the inclusion of only one subnational jurisdiction as a scope limitation (Section 9).

4.3 Coding Procedure

Each regulatory document was coded against the 17 PPP sub-dimensions. For each sub-dimension, the coding captured: (1) presence or absence; (2) regulatory intensity (Mandatory, Recommended, or Absent); (3) specific provision references (articles, sections, clauses); and (4) a textual summary of the provision’s content and scope. Coding decisions are documented in full in Appendix A.

4.4 AI-Assisted Coding and Validation

A distinctive methodological feature of this survey is the systematic validation of human coding against an independent AI-assisted pipeline. Recent work has demonstrated that large language models can match or exceed crowd-worker accuracy on text annotation tasks (Gilardi, Alizadeh, & Kubli, 2023) and hold potential to transform computational social science (Ziems et al., 2024). However, their application to regulatory document classification, where distinctions between binding and voluntary provisions require legal judgment, is nascent. We develop a replicable protocol that serves

Table 2: Regulatory Document Corpus. The 12 instruments analyzed in this survey, spanning 10 jurisdictions and international initiatives. Documents were selected for regulatory significance, typological diversity, and geographic coverage. The EU, United States, and China serve as primary cases (detailed provision-level coding); remaining jurisdictions serve as secondary cases. All instruments were officially adopted, in force, or included as historical references at the time of coding. The U.S. Executive Order was revoked in January 2025 and is retained as a historical reference reflecting federal AI governance intent during the study period.

Jurisdiction	Framework	Legal Form	Status
EU	Artificial Intelligence Act (Reg. (EU) 2024/1689)	Binding legislation	In force (Aug 2024)
USA (Federal)	Exec. Order on Safe, Secure, and Trustworthy AI (2023)	Executive directive	Issued Oct 2023; revoked Jan 2025
USA (Federal)	NIST AI Risk Management Framework (AI RMF 1.0)	Voluntary framework	Published Jan 2023
USA (State)	Colorado AI Act (SB 24-205, 2024)	State legislation	Effective Jun 2026
China	Deep Synthesis Provisions (2022)	Binding regulation	Effective Jan 2023
China	Generative AI Measures (2023)	Binding regulation	Effective Aug 2023
Canada	Directive on Automated Decision-Making	Federal directive	In force (amended)
UK	Pro-Innovation AI Regulation White Paper (2023)	Policy framework	Published Mar 2023
Singapore	Model AI Governance Framework (2nd ed. 2020)	Voluntary framework	Active
Japan	AI Promotion Act (2025)	Legislation	Enacted 2025
International	OECD AI Principles	Intergov. standard	Adopted 2019; updated 2024
International	UNESCO Recommendation on Ethics of AI	Global standard	Adopted Nov 2021

as both a practical augmentation and an empirical contribution to comparative regulatory methodology.

Protocol. All 12 regulatory documents were independently coded by a large language model (Claude, Anthropic) against all 17 PPP sub-dimensions. The model received structured prompts presenting (1) sub-dimension definitions from Table 1, (2) the three-level rating scale with definitions, and (3) the regulatory text. For each sub-dimension, the model provided a rating, cited specific provisions, and supplied a justification. No information about the manual coding was provided. The resulting 170 classifications (17 sub-dimensions $\times$ 10 jurisdictions) were compared against manual expert coding to assess inter-method agreement. Note that while the corpus contains 12 instruments, the United States (three instruments) and China (two instruments) are each coded at the jurisdiction level, producing 10 jurisdiction columns rather than 12 instrument columns. Where a jurisdiction has multiple instruments, the coding reflects the combined regulatory posture across all instruments for that jurisdiction.

Validation sequence. The experiment proceeded in two stages. In Stage 1, a pilot set of three frameworks (EU AI Act, NIST AI RMF, Canada DADM) was validated, achieving 84.3% agreement (43/51). In Stage 2, the remaining seven frameworks were coded and validated using the identical protocol. The full results (per-jurisdiction, per-dimension, and per-legal-form agreement rates) are reported in Section 9.

Scope. All final coding decisions in this paper are human-verified. The AI-assisted pipeline serves as a validation instrument and efficiency tool, not as a replacement for expert judgment. The full prompt templates and AI-generated outputs are available in the supplementary materials.

4.5 Comparative Analysis Procedure

The coded data were organized into a jurisdiction $\times$ PPP sub-dimension matrix (presented across Sections 5–7). Analysis proceeded in three stages:

- (1) **Within-jurisdiction profiling:** for each jurisdiction, we constructed a PPP profile summarizing how regulatory provisions distribute across dimensions and how dimensions interrelate (RQ1).
- (2) **Cross-jurisdictional comparison:** we compared profiles to identify convergences, divergences, and characteristic PPP configurations associated with different regulatory philosophies (RQ2, RQ3).
- (3) **Gap and interdependency analysis:** we identified sub-dimensions weakly addressed across most jurisdictions (RQ5) and analyzed how provisions in one dimension condition provisions in others (RQ4).

5 People Dimension: Comparative Analysis

The People dimension examines how regulatory frameworks address the human element of AI governance: who is responsible, who oversees AI systems, who is affected, and what rights and obligations apply to human actors throughout the AI lifecycle. Table 3 presents the comparative coding across all jurisdictions.

5.1 Human Oversight (P1)

Human oversight requirements represent one of the strongest areas of global convergence. All ten jurisdictions address this sub-dimension, with four imposing mandatory requirements and six recommending oversight measures.

The EU AI Act establishes the most detailed human oversight regime. Article 14 requires that high-risk AI systems be designed to allow effective human oversight, including the ability to fully understand system capabilities and limitations, to properly monitor operation, and to decide not to use the system or to override its output. The level of oversight must be commensurate with the risks posed.

China’s regulations require providers to bear responsibility for content generated by AI services and to implement human-based

Table 3: People Dimension: Cross-Jurisdictional Comparison. Each cell indicates the regulatory intensity for a given sub-dimension: **M** = legally binding requirement, **R** = officially encouraged but not binding, **–** = not addressed. Human oversight (P1) and accountability (P2) show the broadest convergence, while affected persons’ rights (P5) and workforce training (P4) show the greatest divergence.

	EU	USA (Fed.)	USA (CO)	China	Canada	UK	Singapore	Japan	OECD	UNESCO
P1: Oversight	M	R	M	M	M	R	R	R	R	R
P2: Accountability	M	R	M	M	M	R	R	R	R	R
P3: Regulator roles	M	R	M	M	M	R	R	M	R	R
P4: Workforce	M	R	–	R	R	–	R	R	R	R
P5: Rights	M	–	M	R	M	R	R	–	R	R

content moderation. Canada’s Directive on Automated Decision-Making mandates human-in-the-loop involvement proportionate to the assessed impact level, with higher-impact systems requiring greater human intervention capability. The Colorado AI Act imposes a deployer duty of care that includes human oversight of high-risk AI decision-making.

Among the voluntary frameworks, the NIST AI RMF addresses human oversight through its Govern function, emphasizing organizational policies for human-AI configuration. Singapore’s Model AI Governance Framework recommends three approaches to human involvement: human-in-the-loop (full control), human-over-the-loop (supervisory), and human-out-of-the-loop (autonomous), determined by a risk-based assessment of harm severity and probability.

5.2 Accountability Structures (P2)

Accountability provisions show a similar pattern, with four mandatory regimes and six voluntary ones. The EU AI Act allocates obligations to both providers (Articles 16–22) and deployers (Article 26), creating a dual accountability structure that spans the AI value chain. China places primary liability on service providers, with algorithm filing requirements with the Cyberspace Administration of China serving as an accountability mechanism. Canada holds federal institutions directly accountable for automated decisions, while Colorado assigns deployer-level accountability for high-risk AI decisions.

The distribution of accountability between developers and deployers remains a significant source of divergence. The EU’s dual-accountability model contrasts with China’s provider-first approach and with the UK’s delegation to existing sectoral regulators without AI-specific accountability provisions.

5.3 Regulator and Stakeholder Roles (P3)

Institutional architecture varies substantially. The EU has established a multi-layered governance structure comprising national competent authorities, the EU AI Office, and the European Artificial Intelligence Board. China centralizes regulatory authority in the Cyberspace Administration of China. Canada assigns oversight to the Treasury Board Secretariat. Japan’s AI Promotion Act established the AI Strategy Headquarters. In contrast, the UK delegates AI oversight to existing sectoral regulators, and Singapore relies on the

Infocomm Media Development Authority (IMDA) and the Personal Data Protection Commission (PDPC) without binding mandate.

5.4 Workforce Training and AI Literacy (P4)

Workforce provisions represent one of the weakest areas of global coverage. Only the EU imposes mandatory AI literacy obligations: Article 4 of the AI Act requires providers and deployers to ensure that their staff have sufficient AI literacy. All other jurisdictions either recommend training (US federal, China, Singapore, Japan, OECD, UNESCO) or omit the topic entirely (Colorado, UK).

5.5 Affected Persons’ Rights (P5)

Rights for individuals affected by AI-assisted decisions show the sharpest divergence. The EU provides a right to explanation for decisions based on high-risk AI systems that produce legal or similarly significant effects on individuals (Article 86), subject to applicable Union or national law. Canada requires that affected individuals receive plain-language explanations of automated decisions and provides a right of recourse. Colorado mandates consumer notification and opt-out rights. In contrast, the US federal framework has no AI-specific individual rights provisions, and Japan’s AI Promotion Act is silent on affected persons’ rights. China provides user complaint mechanisms but stops short of enforceable individual rights.

6 Processes Dimension: Comparative Analysis

The Processes dimension examines the governance procedures and compliance mechanisms that regulatory frameworks require or recommend. Table 4 presents the comparative coding.

6.1 Risk Assessment (PR1)

Risk assessment represents one of the strongest areas of convergence after human oversight. Four jurisdictions mandate risk assessment (EU, Colorado, China, Canada), five recommend it, and only Japan’s promotional framework is silent.

The EU AI Act requires a comprehensive risk management system for high-risk AI (Article 9), including identification and analysis of known and reasonably foreseeable risks, estimation and evaluation of risks arising from intended use and foreseeable misuse, and adoption of appropriate risk management measures. The NIST AI RMF structures risk assessment through its Map and Measure

Table 4: Processes Dimension: Cross-Jurisdictional Comparison. Risk assessment (PR1) and transparency (PR3) show near-universal adoption. Certification (PR4) is the most divergent sub-dimension across all 17: only the EU mandates conformity assessment. Incident reporting (PR5) and regulatory sandboxes (PR6) are similarly sparse, representing significant global regulatory gaps.

	EU	USA (Fed.)	USA (CO)	China	Canada	UK	Singapore	Japan	OECD	UNESCO
PR1: Risk assess.	M	R	M	M	M	R	R	–	R	R
PR2: Auditing	M	R	M	M	M	R	R	–	–	R
PR3: Transparency	M	R	M	M	M	R	R	R	R	R
PR4: Certification	M	–	–	–	–	–	–	–	–	–
PR5: Incidents	M	R	–	M	–	–	–	–	–	–
PR6: Sandboxes	M	–	–	–	–	R	R	R	–	–

functions, providing a voluntary but detailed lifecycle approach (Govern–Map–Measure–Manage). Canada’s Directive requires an Algorithmic Impact Assessment (AIA) before any automated decision system is deployed in federal institutions, using a scoring system that determines the impact level (I through IV). China mandates security assessments before public deployment of generative AI services. Colorado requires impact assessments for all high-risk AI systems.

A notable difference lies in the *type* of assessment required. The EU emphasizes fundamental rights impact alongside safety risks. Canada’s AIA centers on impacts to individuals and communities. China’s security assessment prioritizes content safety and societal stability. This variation reflects different regulatory philosophies even where the procedural requirement converges.

6.2 Auditing and Compliance (PR2)

Auditing requirements show moderate convergence at the mandatory level (four jurisdictions: EU, Colorado, China, Canada) but notable absence in Japan and the OECD framework. The EU combines conformity assessment (Article 43) with post-market monitoring (Article 72). China requires algorithm filing and periodic review by the Cyberspace Administration. Canada mandates ongoing monitoring and periodic review of automated decision systems. Colorado requires annual review of impact assessments. Singapore offers the AI Verify testing toolkit as a voluntary self-assessment tool.

6.3 Transparency Requirements (PR3)

Transparency is the sub-dimension with the broadest coverage: all ten jurisdictions address it, with four imposing mandatory obligations (EU, Colorado, China, Canada). The EU AI Act imposes graduated transparency obligations by risk tier (Articles 13 and 50), including a requirement that AI-generated content be marked as such. China requires service providers to disclose the AI nature of their services. Canada mandates public disclosure of AIA results.

6.4 Certification and Governance Procedures (PR4)

Certification is the most divergent sub-dimension across all 17: only the EU mandates conformity assessment (Article 43), which may require third-party assessment for certain high-risk categories.

No other jurisdiction has established an AI-specific certification scheme. This represents a significant gap in the global regulatory architecture.

6.5 Incident Reporting (PR5)

Only the EU (Article 73, requiring serious incident reporting to market surveillance authorities) and China (illegal content reporting obligations) mandate AI-specific incident reporting. The US Executive Order (now revoked) directed safety reporting for dual-use foundation models as a recommended measure. The absence of incident reporting provisions in seven of ten jurisdictions represents a notable regulatory gap.

6.6 Regulatory Sandboxes (PR6)

Regulatory sandboxes are controlled environments established by regulators that allow developers to test AI systems under real-world conditions with regulatory oversight and temporary relaxation of certain requirements. The EU is the only jurisdiction that mandates their creation (Articles 57–60), requiring member states to establish AI regulatory sandboxes. The UK, Singapore, and Japan support sandbox environments on a voluntary or sector-specific basis. Most jurisdictions do not address this sub-dimension.

7 Platforms Dimension: Comparative Analysis

The Platforms dimension examines the technical and infrastructural requirements imposed on AI systems themselves. Table 5 presents the comparative coding.

7.1 AI System Classification (PL1)

AI system classification, i.e., how jurisdictions categorize AI systems for regulatory purposes, shows sharp divergence. Four jurisdictions employ formal classification: the EU uses a four-tier risk scheme (unacceptable, high, limited, minimal risk, as defined in Article 6 and Annex III); Colorado employs a binary classification (high-risk versus other); Canada uses an impact-level scoring system (levels I through IV); and China regulates by application category (deep synthesis, generative AI). The NIST AI RMF offers context-based risk framing without binding tiers. Six jurisdictions, including the UK, Singapore, Japan, OECD, and UNESCO, do not establish any formal AI classification scheme.

Table 5: Platforms Dimension: Cross-Jurisdictional Comparison. The Platforms dimension shows the sharpest divergence of any PPP dimension. Data governance (PL3) and safety (PL4) are broadly addressed, but GPAI provisions (PL6) and content labeling (PL5) are mandatory in only two jurisdictions (EU and China). AI classification (PL1), identified in this paper as a cross-dimensional regulatory trigger, is present in only four jurisdictions.

	EU	USA (Fed.)	USA (CO)	China	Canada	UK	Singapore	Japan	OECD	UNESCO
PL1: Classification	M	R	M	M	M	–	–	–	–	–
PL2: Infrastructure	M	R	–	M	–	–	–	R	R	R
PL3: Data governance	M	R	R	M	R	R	R	R	R	R
PL4: Safety	M	R	–	M	–	R	R	R	R	R
PL5: Content labeling	M	–	–	M	–	–	–	–	–	–
PL6: GPAI	M	–	–	M	–	–	–	–	–	–

The absence of a common global taxonomy for AI risk classification represents one of the most significant barriers to international regulatory interoperability.

7.2 Infrastructure and Deployment Standards (PL2)

Infrastructure requirements are addressed by only two mandatory regimes (EU and China) and four recommending jurisdictions. The EU AI Act mandates technical documentation requirements (Annex IV). China imposes technical standards for deep synthesis and generative AI infrastructure. Most jurisdictions leave infrastructure standards to voluntary adoption or sector-specific regulation.

7.3 Data Governance (PL3)

Data governance shows broad engagement: all ten jurisdictions address it, though only the EU and China impose mandatory requirements. The EU AI Act (Article 10) mandates data governance for training, validation, and testing datasets. China requires training data legality, quality, and provenance. The remaining jurisdictions recommend data quality and governance measures without binding force.

7.4 Safety and Robustness (PL4)

Safety and robustness requirements are broadly addressed: eight of ten jurisdictions cover this sub-dimension, with the EU (Article 15, requiring accuracy, robustness, and cybersecurity) and China imposing mandatory standards. Colorado and Canada do not address this sub-dimension. The NIST AI RMF addresses reliability and robustness through its trustworthiness characteristics. Singapore and Japan encourage robustness through guidelines.

7.5 Content Labeling and Provenance (PL5)

Content labeling, i.e., requiring AI-generated content to be marked as such, is mandatory only in the EU (Article 50) and China (deep synthesis provisions). This sub-dimension is absent from eight of ten frameworks, a notable gap given the growing significance of synthetic content.

7.6 GPAI and Foundation Model Provisions (PL6)

General-purpose AI regulation is the most sparsely addressed sub-dimension alongside certification: only the EU (Articles 51–56, establishing obligations for GPAI model providers including those posing systemic risk) and China (Generative AI Measures specifically regulating foundation model services) have enacted specific provisions. This gap is significant given the centrality of foundation models to current AI deployment.

8 Cross-Dimensional Analysis

8.1 Regulatory Philosophy Typology

Based on the observed PPP configuration patterns, we identify six distinct regulatory philosophies:

The typology reveals a governance spectrum from comprehensive mandatory regulation (EU) to normative guidance without enforcement (OECD, UNESCO). A striking quantitative pattern emerges: the EU’s 17 mandatory provisions represent complete PPP coverage, while the combined US federal and state approach yields only 8 mandatory provisions with 13 absent sub-dimensions, the most fragmented profile among jurisdictions with binding instruments. China’s 13 mandatory provisions place it closer to the EU than to any other jurisdiction, despite fundamental differences in regulatory philosophy and enforcement mechanisms.

8.2 Convergence and Divergence Patterns

Analysis of the comparative matrices reveals a hierarchy of global adoption:

Tier 1: Near-universal adoption (addressed by 9–10 jurisdictions): human oversight (P1), accountability (P2), regulator roles (P3), risk assessment (PR1), transparency (PR3), and data governance (PL3). Among the jurisdictions studied, these sub-dimensions constitute an emerging convergence on minimum AI governance requirements.

Tier 2: Common but not universal (addressed by 5–8 jurisdictions): workforce training (P4), affected persons’ rights (P5), auditing and compliance (PR2), safety and robustness (PL4), infrastructure standards (PL2), and AI classification (PL1). Coverage at this tier varies substantially in regulatory intensity.

Table 6: Typology of AI Regulatory Philosophies. Six regulatory philosophies identified from observed PPP configuration patterns. The PPP Profile column shows the count of Mandatory (M), Recommended (R), and Absent (A) sub-dimensions. The EU’s complete mandatory coverage (17M) contrasts sharply with the combined US approach (8M/13R/13A), the most fragmented profile among jurisdictions with binding instruments. China (13M) is structurally closer to the EU than to any other jurisdiction despite fundamental differences in governance philosophy.

Philosophy	Jurisdiction(s)	Characteristics	PPP Profile
Comprehensive risk-based	EU	Binding legislation; risk-tiered obligations; balanced coverage across all PPP dimensions	17M / 0R / 0A
State-directed	China	Binding regulations; application-specific; state-supervised; strong Platforms emphasis	13M / 2R / 2A
Public-sector focused	Canada	Mandatory for government use; strong People and Processes; limited Platforms	8M / 2R / 7A
Sector-based decentralized	USA (Federal + CO)	No central AI authority; voluntary federal + binding state; fragmented coverage	8M / 13R / 13A
Light-touch / non-comprehensive	UK, Singapore, Japan	Guidance-oriented or narrowly scoped legislation; sectoral implementation; Recommended-heavy profiles	1M / 29R / 21A
Normative-international	OECD, UNESCO	Values-based; no enforcement; broad principled coverage	0M / 21R / 13A

Tier 3: Sparse or emerging (mandatory in at most 2 jurisdictions): certification (PR4, 1/10 mandatory), incident reporting (PR5, 2/10), regulatory sandboxes (PR6, 1/10), content labeling (PL5, 2/10), and GPAI provisions (PL6, 2/10). These represent the frontier of AI regulation, where international divergence is greatest and harmonization most challenging.

8.3 Structural Determinants of Regulatory Coherence

Our analysis identifies two structural features that are associated with observed variation in regulatory coherence across the jurisdictions studied. We present these as the paper’s principal analytical contributions.

8.3.1 Classification as Cross-Dimensional Trigger. AI system classification (PL1) functions as a *regulatory trigger*, a Platforms-dimension provision that formally determines the regulatory intensity applied to People and Processes dimensions. The mechanism operates through three channels: *obligation activation* (classification creates discrete categories linked to specific obligation bundles), *proportionality anchoring* (classification provides the formal basis for calibrating governance intensity to risk level), and *enforcement gating* (classification determines which AI systems fall within mandatory regulatory scope).

This produces two distinct regulatory architectures. In the **cascading architecture** (EU, Colorado, Canada), classification triggers calibrated obligations across all PPP dimensions. The EU’s four-tier classification (Art. 6, Annex III) directly determines human oversight intensity (P1, Art. 14), accountability scope (P2, Arts. 16, 26), risk assessment depth (PR1, Art. 9), conformity assessment applicability (PR4, Art. 43), and data governance standards (PL3, Art. 10). Cross-dimensional coherence is high: provisions are not merely present but are *calibrated to* the classification level.

In the **flat architecture** (UK, Singapore, Japan, OECD, UNESCO), governance provisions operate independently, unanchored to system-level risk determinations. The UK’s five principles apply uniformly regardless of AI system type. Singapore recommends

risk-based human involvement but leaves the determination to organizational discretion. While organizational judgment plays a role in both architectures, in the flat model it is the primary determinant of governance intensity, whereas in the cascading model it operates within a regulatory structure that sets minimum obligations by risk tier.

The US federal approach reveals an instructive *hybrid*: the NIST AI RMF articulates cascading logic (its MAP function recommends context-based risk categorization), but because classification is voluntary, the cascade is never formally triggered. The result is a process-rich framework with zero mandatory provisions: the architecture of coherence without its structural foundation.

Quantitatively, among the jurisdictions studied, those with classification average 11.5 mandatory provisions; those without average 0.2. This difference is not explained by regulatory ambition: Canada’s narrowly-scoped DADM achieves higher coherence (8 mandatory) than the UK’s economy-wide White Paper (0 mandatory), because Canada’s instrument contains a classification scheme and the UK’s does not. The explanatory variable is architectural, not aspirational.

8.3.2 The Institutional Creation Threshold. A second structural pattern emerges when we partition the 17 sub-dimensions by the type of institutional capacity they require. Fourteen sub-dimensions, including human oversight, accountability, risk assessment, transparency, data governance, and AI classification itself, can be implemented by *adapting existing regulatory infrastructure*: extending mandates, adding requirements to existing compliance procedures, or issuing guidance within established frameworks. Three sub-dimensions, certification (PR4), incident reporting (PR5), and regulatory sandboxes (PR6), require *creating entirely new governance institutions*: certification bodies with AI-specific expertise, incident databases with standardized reporting protocols, and sandbox environments with dedicated legal frameworks.

The data reveal a stark divide. Adaptation provisions are addressed (at any intensity level) in 76% of jurisdiction-sub-dimension pairs (106 of 140). Creation provisions are addressed in only 27% of cases (8 of 30) and are mandatory in just 4 of 30 cases (13%),

with the EU accounting for three (PR4, PR5, PR6) and China for one (PR5). This divide holds across all regulatory philosophies: it is equally visible in binding legislation (China omits certification and sandboxes despite 13 mandatory provisions), narrowly-scoped directives (Canada omits all three despite 8 mandatory provisions), and voluntary frameworks.

The threshold explains the voluntary ceiling (Proposition 3) at the level of mechanism: light-touch frameworks plateau at approximately 65% coverage not because they lack governance ambition but because non-binding instruments are structurally incapable of mandating the creation of new institutions. It also explains why the residual global governance gap, the provisions that separate the EU from every other jurisdiction, consists precisely of creation provisions. The most consequential barrier to AI governance harmonization is not principled disagreement over values but the differential capacity and willingness of jurisdictions to invest in new governance infrastructure.

Together, these two structural features provide a parsimonious explanation for the majority of variation in our dataset. Jurisdictions with both classification and institutional creation (EU only) achieve complete PPP coverage. Jurisdictions with classification but without institutional creation (Colorado, Canada, China) achieve high but incomplete coherence. Jurisdictions with neither (UK, Singapore, Japan, OECD, UNESCO) achieve zero mandatory provisions regardless of stated governance ambition.

8.4 Regulatory Gaps

Five sub-dimensions are identified as significant global regulatory gaps:

- (1) **GPAI/Foundation model provisions (PL6):** Only the EU and China specifically regulate foundation models. Given that GPAI models are the underlying technology powering most current AI applications, this represents the most consequential gap.
- (2) **Certification (PR4):** Only the EU has established an AI-specific conformity assessment regime. The absence of certification frameworks elsewhere limits mutual recognition and international interoperability.
- (3) **Incident reporting (PR5):** Only the EU and China mandate AI-specific incident reporting. Without systematic incident reporting, jurisdictions lack the data needed to identify emerging risks and refine regulatory approaches.
- (4) **Content labeling (PL5):** Only the EU and China require AI-generated content labeling. As synthetic content becomes increasingly prevalent, this gap poses challenges for information integrity.
- (5) **Workforce training (P4):** Only the EU mandates AI literacy. The global underinvestment in AI workforce development limits the effectiveness of other governance provisions that depend on informed human actors.

8.5 Derived Propositions

From the comparative analysis, we derive five testable propositions that formalize the patterns observed and provide a foundation for future research.

Proposition 1 (Architectural Coherence). *In multi-dimensional technology regulation, the presence of a formal classification scheme is a stronger predictor of cross-dimensional regulatory coherence than regulatory philosophy, legal tradition, or the scope of the regulatory instrument.* The intuitive expectation is that comprehensive regulatory ambition produces coherent regulation. The data show otherwise: Canada's narrow directive (federal government only) is more coherent (8M) than the UK's economy-wide White Paper (0M), because Canada has classification and the UK does not. Among the jurisdictions studied, those with classification average 11.5 mandatory provisions regardless of scope; those without average 0.2 regardless of stated goals.

Proposition 2 (Institutional Creation Threshold). *Global convergence in AI governance is bounded by an institutional creation threshold: jurisdictions converge on provisions implementable through existing institutional capacity and diverge on provisions requiring new governance infrastructure, independent of regulatory philosophy or political system.* This offers an alternative to the common assumption that governance gaps primarily reflect different values. The UK and EU share nearly identical stated values (safety, transparency, fairness, accountability) yet differ on exactly the provisions requiring institutional creation. Adaptation provisions are addressed in 76% of cases; creation provisions in only 27%.

Proposition 3 (Voluntary Ceiling as Institutional Limit). *Voluntary AI governance frameworks exhibit a structural ceiling of 60–70% sub-dimension coverage because the provisions they systematically fail to address (certification, incident reporting, regulatory sandboxes) require institutional commitments that non-binding instruments cannot generate, regardless of the specificity or ambition of their guidance.* The deeper explanation is not that voluntary frameworks are “weaker” on what they address (Singapore's guidance on transparency is substantively rich) but that in practice they have not generated the institutional infrastructure these provisions require. In principle, certification, incident reporting, and sandboxes could be organized voluntarily through industry self-regulation or decentralized arrangements (e.g., company-led certification consortia or service-provider-managed sandbox programs). However, the empirical pattern in our data is that none of the voluntary frameworks in the corpus has produced such institutions: all four (UK, Singapore, OECD, UNESCO) uniformly omit PR4 and PR5, and while the UK, Singapore, and Japan recommend sandboxes (PR6), none mandates or has established them through voluntary coordination alone. Whether this reflects a structural limitation of non-binding instruments or a contingent outcome that could change as the field matures remains an open question for future research.

Proposition 4 (Philosophy–Architecture Decoupling). *Regulatory philosophy predicts which PPP dimension receives emphasis within a jurisdiction's covered provisions, but does not predict total coverage, which is determined by architectural features (classification and institutional creation) rather than governance goals.* Philosophy determines *shape*: rights-based regimes emphasize People (EU: 5/5 People mandatory; Canada: 4/5); state-directed regimes emphasize Platforms (China: 6/6 Platforms mandatory). Architecture determines *coverage*: China (state-directed, 13M) and the EU (rights-based, 17M) have opposing philosophies but similar coverage because both have classification. The UK and EU have similar

philosophies but radically different coverage (0M vs. 17M) because they differ on architecture.

Proposition 5 (Legibility Predicts AI-Coding Reliability). *The reliability of AI-assisted regulatory coding is determined by the formal legibility of the regulatory text, specifically the degree to which obligations specify identifiable duty-holders, enumerated actions, and explicit conditions, rather than by document complexity, subject matter, or jurisdiction. AI-assisted coding is therefore most viable precisely for the instruments most important to code correctly: binding legislation with explicit obligations.* The most complex document (EU AI Act, 113 articles) achieves 100% agreement; shorter voluntary frameworks achieve 76.5%. In our data, more complex binding legislation tends to be more explicit, and explicitness is associated with higher coding reliability. Binding legislation: 92.6%; directives/policy: 76.5%.

8.6 Summary of Key Findings

Table 7 maps each research question to its principal findings.

9 Discussion

9.1 Implications for International Harmonization

Our comparative analysis identifies both opportunities and barriers for international AI regulatory harmonization. The Tier 1 sub-dimensions, where near-universal coverage exists, represent a potential foundation for mutual recognition agreements. All major jurisdictions agree, at minimum, that AI governance requires human oversight, accountability structures, risk assessment, transparency, and data governance. This convergence, while varying in regulatory intensity, provides a common vocabulary for international coordination.

The barriers are concentrated in Tier 3 sub-dimensions, where fundamental disagreements persist. The absence of a common AI classification taxonomy means that jurisdictions cannot straightforwardly map each other's risk categories. The EU's GPAI provisions have no counterpart in most other jurisdictions, creating compliance asymmetries for foundation model providers operating globally. The EU's unique certification regime lacks international equivalents, limiting mutual recognition.

9.2 AI-Assisted Coding: Full Experimental Results

We validated all 170 classifications (17 sub-dimensions $\times$ 10 jurisdictions) by comparing manual expert coding against independent AI-assisted coding. Table 8 summarizes the results.

The overall agreement rate of 87.6% across 170 classifications indicates substantial inter-method alignment. Three patterns emerge from the 21 disagreements:

Pattern 1: AI overrates coverage (9/21 disagreements). The LLM classified provisions as present where the manual coder judged them too tangential. For example, the LLM coded the UK White Paper as addressing risk assessment (PR1) based on the safety principle, which the manual coder assessed as too general to constitute a substantive risk assessment provision.

Pattern 2: AI underrates coverage (7/21 disagreements). The LLM coded provisions as Absent where the manual coder found them indirectly present. This concentrated in the infrastructure standards sub-dimension (PL2), where the LLM applied a strict interpretation requiring explicit infrastructure provisions, while the manual coder recognized indirect coverage through technical documentation or standards-development mandates.

Pattern 3: Intensity confusion (5/21 disagreements). The LLM disagreed on the Mandatory/Recommended boundary for provisions using strong normative language without explicit penalties.

These patterns support Proposition 5: agreement is highest for explicit, enumerated obligations (binding legislation: 92.6%; international standards with clear principle structure: 94.1%) and lowest for directive and policy instruments with mixed normative language (76.5%). Legal form, specifically the explicitness with which obligations are stated, predicts AI-coding reliability more strongly than document length, jurisdiction, or PPP dimension.

The protocol reduces initial classification time from approximately 4–6 hours per framework to under 30 minutes, with approximately 1 hour of human validation per framework. We recommend that future studies: (1) report per-jurisdiction and per-dimension agreement rates, (2) characterize disagreement patterns by regulatory text features, and (3) retain human validation as the final arbiter.

9.3 Framework Utility and Limitations

The PPP framework proved effective as a comparative instrument. Its primary strength is enabling dimension-level comparison across regulatory instruments of fundamentally different legal forms: the same coding scheme applies to binding EU legislation, voluntary US frameworks, and normative international standards. The 17 sub-dimensions provided sufficient granularity to identify meaningful variation without producing an unmanageable coding burden.

A principal limitation is the coarse three-level intensity scale, discussed further in the Limitations subsection below.

9.4 Limitations

Four limitations merit explicit acknowledgment. First, the *single-coder design* constrains inter-coder reliability. We mitigate this through documented coding decisions with provision-level evidence and through the AI-assisted validation (87.6% agreement across 170 classifications), but a multi-coder design would strengthen validity. Second, the *three-level coding scale* collapses substantial within-category variation: the EU's five-article GPAI regime and China's single-regulation approach both receive a Mandatory rating despite differing markedly in specificity. Future work could supplement the M/R/A scale with a provision-count or specificity metric. Third, the survey captures *regulatory design as of early 2026*, not implementation or enforcement. The gap between regulatory text and regulatory effect is a distinct research question requiring empirical methods beyond documentary analysis. Fourth, the *ten-jurisdiction scope*, while covering the principal regulatory models, excludes significant activity in Latin America, South Asia, and Africa; the propositions derived here should be tested against a broader set of jurisdictions before claiming generalizability.

Table 7: Summary of Key Findings by Research Question. Each research question is mapped to its principal empirical finding. Proposition numbers in parentheses indicate derived testable claims (see Section 8.5). RQ4 yields the paper’s central theoretical contribution: the identification of AI classification as a cross-dimensional regulatory trigger.

RQ	Principal Finding
RQ1	The EU is the only jurisdiction with mandatory coverage across all 17 sub-dimensions (17M). Among binding frameworks, mandatory coverage ranges from 17 (EU) to 1 (Japan). The three light-touch frameworks (UK, Singapore, Japan) converge on 0–1 mandatory and 8–11 recommended provisions, exhibiting a characteristic ceiling (Proposition 3).
RQ2	Six regulatory philosophies emerge: comprehensive risk-based (EU), state-directed (China), public-sector focused (Canada), sector-based decentralized (USA), light-touch/non-comprehensive (UK, Singapore, Japan), and normative-international (OECD, UNESCO). Philosophy predicts PPP emphasis: rights-based regimes prioritize People; process-oriented regimes prioritize Processes; technology-focused regimes prioritize Platforms.
RQ3	Highest convergence: transparency (PR3), risk assessment (PR1), human oversight (P1), and accountability (P2) are addressed by all 10 jurisdictions. Highest divergence: GPAI provisions (PL6), certification (PR4), and content labeling (PL5) are mandatory in at most 2 jurisdictions and absent from 6–8.
RQ4	AI classification (PL1) functions as a <i>regulatory trigger</i> : jurisdictions with formal classification (EU, China, Colorado, Canada) show strong cross-dimensional linkages where risk tier determines People and Processes intensity. Jurisdictions without classification show weaker interdependencies.
RQ5	Five critical gaps identified: GPAI regulation (PL6), certification (PR4), incident reporting (PR5), content labeling (PL5), and workforce training (P4). These Tier 3 sub-dimensions represent the highest barriers to international harmonization.

Table 8: AI-Assisted Coding: Agreement by Jurisdiction (left) and by Legal Form / PPP Dimension (right). Overall agreement is 87.6% across 170 classifications. Perfect agreement occurs for the EU AI Act, whose explicit legal obligations are unambiguous to both human and AI coders. Agreement is highest for binding legislation (92.6%) and international standards (94.1%), and lowest for policy directives with mixed normative language (76.5%), supporting Proposition 5.

(a) By Jurisdiction				(b) By Legal Form & PPP Dimension		
Jurisdiction	Agree	Disagree	Rate	Category	N	Rate
EU AI Act	17/17	0	100.0%	<i>By legal form</i>		
Japan AI Promotion Act	16/17	1	94.1%	Binding legislation	68	92.6%
OECD AI Principles	16/17	1	94.1%	Directives/policy	51	76.5%
UNESCO Recommendation	16/17	1	94.1%	Voluntary frameworks	17	88.2%
Colorado AI Act	15/17	2	88.2%	Intl. standards	34	94.1%
China (combined)	15/17	2	88.2%	<i>By PPP dimension</i>		
Singapore Framework	15/17	2	88.2%	People (P1–P5)	50	90.0%
NIST AI RMF	13/17	4	76.5%	Processes (PR1–PR6)	60	85.0%
Canada DADM	13/17	4	76.5%	Platforms (PL1–PL6)	60	88.3%
UK White Paper	13/17	4	76.5%			
Overall	149/170	21	87.6%			

9.5 Robustness Check: EU-Excluded Analysis

A potential concern is that the EU AI Act, as the only jurisdiction with complete mandatory coverage, disproportionately drives our structural findings. We therefore verify that the two principal patterns hold when the EU is excluded. *Classification trigger*: among the remaining nine jurisdictions, those with formal classification (Colorado, Canada, China) average 9.7 mandatory provisions; those without (US Federal, UK, Singapore, Japan, OECD, UNESCO) average 0.2. The pattern is not EU-dependent. *Institutional creation threshold*: among the nine non-EU jurisdictions, adaptation provisions are addressed in 73% of cases (92/126); creation provisions are addressed in 19% of cases (5/27) and are mandatory in 4% (1/27, Chinese). The gap between adaptation and creation provisions persists without the EU. *Voluntary ceiling*: the four voluntary frameworks (UK, Singapore, OECD, UNESCO) still converge on 10–11 of 17 sub-dimensions with uniform absence on PR4 and PR5, and none

mandating PR6. These checks confirm that our findings describe structural patterns in the global regulatory landscape, not artifacts of a single comprehensive outlier.

10 Conclusion

This paper has introduced a validated coding framework for comparative AI regulatory analysis and applied it to produce a systematic provision-level benchmark of global AI governance across 10 jurisdictions and 170 coded classifications. From this dataset, we have identified two structural features that together account for substantial variation in the observed regulatory landscape.

The first finding, *classification as cross-dimensional trigger*, reveals that the coherence of an AI regulatory system is determined by a single architectural feature. Jurisdictions with formal AI classification schemes (EU, China, Colorado, Canada) exhibit cascading obligation structures in which a taxonomic decision activates

calibrated requirements across all PPP dimensions (average: 11.5 mandatory provisions). Jurisdictions without classification (US Federal, UK, Singapore, Japan, OECD, UNESCO) exhibit flat structures in which provisions operate independently (average: 0.2). This difference is architectural, not aspirational: Canada's narrowly-scoped directive achieves higher coherence than the UK's economy-wide framework because Canada has classification and the UK does not.

The second finding, the *institutional creation threshold*, reveals that global AI governance convergence is bounded not by principled disagreement but by institutional capacity. Provisions that can be implemented by adapting existing regulatory infrastructure are addressed in 76% of jurisdiction-sub-dimension pairs. Provisions that require creating new governance institutions (certification bodies, incident databases, regulatory sandboxes) are addressed in only 27%, with mandatory adoption concentrated in the EU (three of three creation sub-dimensions) and China (incident reporting only). This threshold holds across all regulatory philosophies and explains both the voluntary ceiling (approximately 65% coverage) and the specific provisions that constitute the residual global governance gap.

Together, these findings suggest that the challenge of AI governance harmonization extends beyond value alignment. While most jurisdictions studied agree on transparency, accountability, and risk assessment, the observed gaps concentrate on architecture (whether a classification trigger exists) and institutional investment (whether jurisdictions create the new institutions that the most consequential governance provisions require). Our findings suggest that policy efforts focused solely on aligning principles may be insufficient to close these structural gaps.

The five propositions, the coding framework, and the companion website at <https://ppp-ai-governance.vercel.app> are maintained as a living resource for the research community.

References

- Leavitt, H. J. (1965). Applied Organizational Change in Industry: Structural, Technological, and Humanistic Approaches. In J. G. March (Ed.), *Handbook of Organizations* (pp. 1144–1170). Rand McNally.
- Baldwin, R., Cave, M., & Lodge, M. (2012). *Understanding Regulation: Theory, Strategy, and Practice* (2nd ed.). Oxford University Press.
- Coglianese, C. (2023). A People-and-Processes Approach to AI Governance. *Administrative and Regulatory Law News*, American Bar Association.
- Schneier, B. (2000). *Secrets and Lies: Digital Security in a Networked World*. John Wiley & Sons.
- Baxter, G., & Sommerville, I. (2011). Socio-technical systems: From design methods to systems engineering. *Interacting with Computers*, 23(1), 4–17.
- Floridi, L., & Cowls, J. (2019). A unified framework of five principles for AI in society. *Harvard Data Science Review*, 1(1).
- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(9), 389–399.
- Cath, C. (2018). Governing artificial intelligence: ethical, legal and technical opportunities and challenges. *Philosophical Transactions of the Royal Society A*, 376(2133).
- Hagendorff, T. (2020). The ethics of AI ethics: An evaluation of guidelines. *Minds and Machines*, 30(1), 99–120.
- Stahl, B. C. (2021). *Artificial Intelligence for a Better Future*. Springer.
- Bradford, A. (2020). *The Brussels Effect: How the European Union Rules the World*. Oxford University Press.
- Smuha, N. A. (2021). From a 'race to AI' to a 'race to AI regulation': regulatory competition for artificial intelligence. *Law, Innovation and Technology*, 13(1), 57–84.
- Roberts, H., Cowls, J., Morley, J., Taddeo, M., Wang, V., & Floridi, L. (2021). The Chinese approach to artificial intelligence: an analysis of policy, ethics, and regulation. *AI & Society*, 36(1), 59–77.
- Veale, M., & Zuiderveen Borgesius, F. (2021). Demystifying the Draft EU Artificial Intelligence Act. *Computer Law Review International*, 22(4), 97–112.
- Madiega, T. (2021). Artificial Intelligence Act. *European Parliamentary Research Service*, PE 698.792.
- Kerry, C. F., Meltzer, J. P., Renda, A., Engler, A., & Fanni, R. (2023). The EU AI Act will have global impact, but a limited Brussels effect. *Brookings Institution*.
- Mökander, J., & Floridi, L. (2023). Operationalising AI governance through ethics-based auditing: an industry case study. *AI and Ethics*, 3(2), 451–468.
- Ziems, C., Held, W., Shaikh, O., Chen, J., Zhang, Z., & Yang, D. (2024). Can large language models transform computational social science? *Computational Linguistics*, 50(1), 237–291.
- Gilardi, F., Alizadeh, M., & Kubli, M. (2023). ChatGPT outperforms crowd workers for text-annotation tasks. *Proceedings of the National Academy of Sciences*, 120(30).
- European Commission. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). *Official Journal of the European Union*.
- The White House. (2023). Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence. Revoked January 2025.
- National Institute of Standards and Technology. (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*. NIST AI 100-1.
- State of Colorado. (2024). Colorado Artificial Intelligence Act (SB 24-205).
- Cyberspace Administration of China. (2022). Provisions on the Management of Deep Synthesis Internet Information Services.
- Cyberspace Administration of China. (2023). Interim Measures for the Management of Generative Artificial Intelligence Services.
- Government of Canada. (2019, amended). Directive on Automated Decision-Making. Treasury Board of Canada Secretariat.
- UK Department for Science, Innovation and Technology. (2023). *A Pro-Innovation Approach to AI Regulation*. Cm 815.
- Infocomm Media Development Authority of Singapore. (2020). *Model AI Governance Framework* (2nd ed.).
- Government of Japan. (2025). Act on Promotion of Research and Development, and Utilization of Artificial Intelligence-related Technology (AI Promotion Act).
- OECD. (2019, updated 2024). *OECD Principles on Artificial Intelligence*. OECD Publishing.

- UNESCO. (2021). *Recommendation on the Ethics of Artificial Intelligence*.

A Supplementary Visualizations

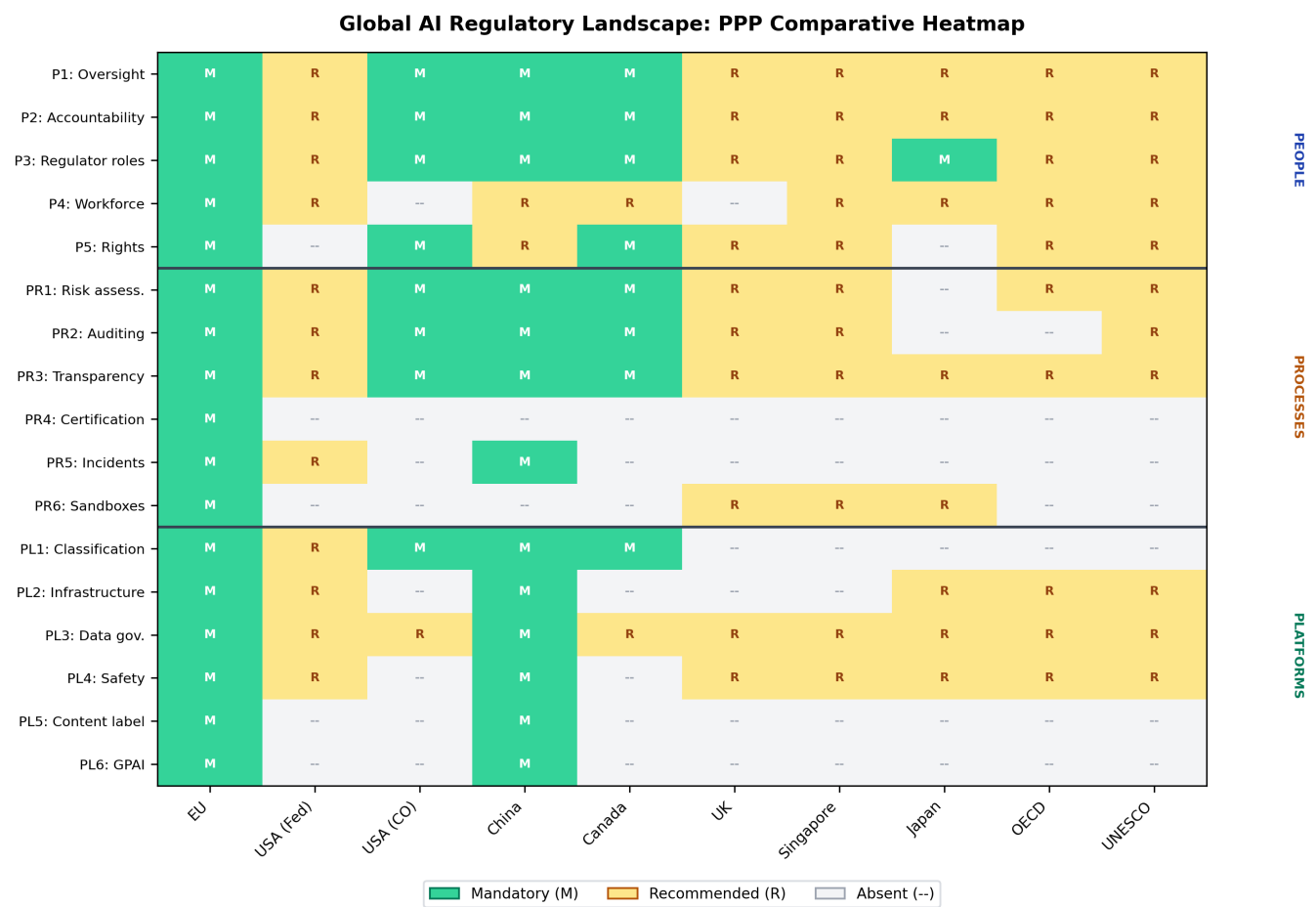


Figure 1: Full PPP Comparative Heatmap (17 sub-dimensions × 10 jurisdictions). Green cells indicate mandatory (binding) provisions, amber cells indicate recommended (voluntary) provisions, and gray cells indicate absent provisions. Three structural features are visible: (1) the EU’s complete mandatory column, (2) the “recommended ceiling” in voluntary frameworks (UK, Singapore, OECD, UNESCO) where coverage plateaus at 60–70% with systematic institutional gaps, and (3) the concentration of absent provisions in Tier 3 sub-dimensions (PR4–PR6, PL5–PL6). Sub-dimensions are ordered by PPP dimension, with horizontal dividers separating People, Processes, and Platforms.

PPP Regulatory Profiles: Radar Comparison

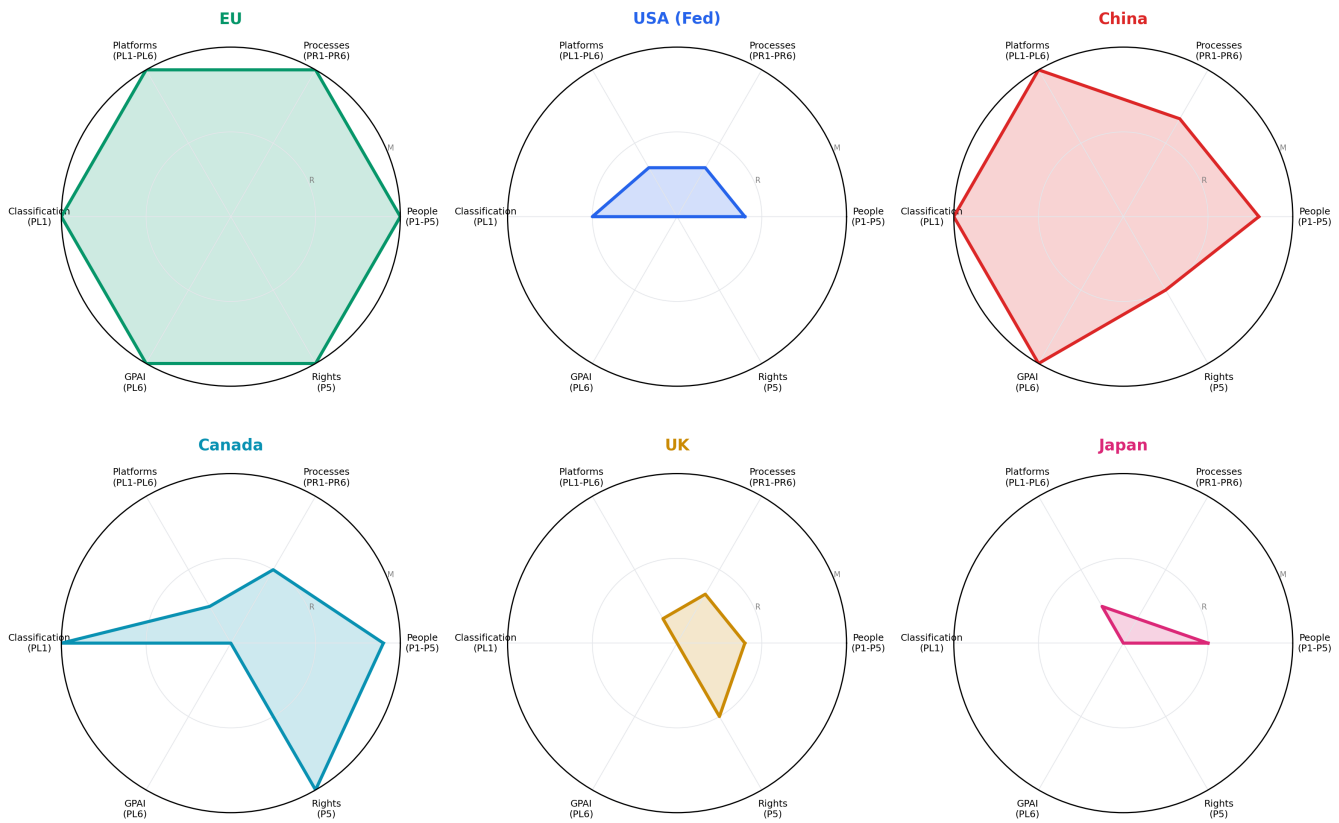


Figure 2: PPP Regulatory Profiles: Radar Comparison for Six Key Jurisdictions. Each radar chart shows a jurisdiction's regulatory intensity across six axes: aggregate People, Processes, and Platforms scores (averaged across sub-dimensions), plus three individual sub-dimensions of theoretical interest: Classification (PL1), GPAI provisions (PL6), and Affected Persons' Rights (P5). The EU presents a near-complete hexagon (full mandatory coverage). China shows strong Platforms coverage but weaker People provisions. The US federal profile is small and uniform (all voluntary). Canada shows an asymmetric shape reflecting strong People and Processes but weak Platforms. The UK and Japan profiles are small, confirming the voluntary ceiling (Proposition 3). The shape differences across jurisdictions visually encode the philosophy–dimension mapping identified in Proposition 4.

Classification as Cross-Dimensional Regulatory Trigger

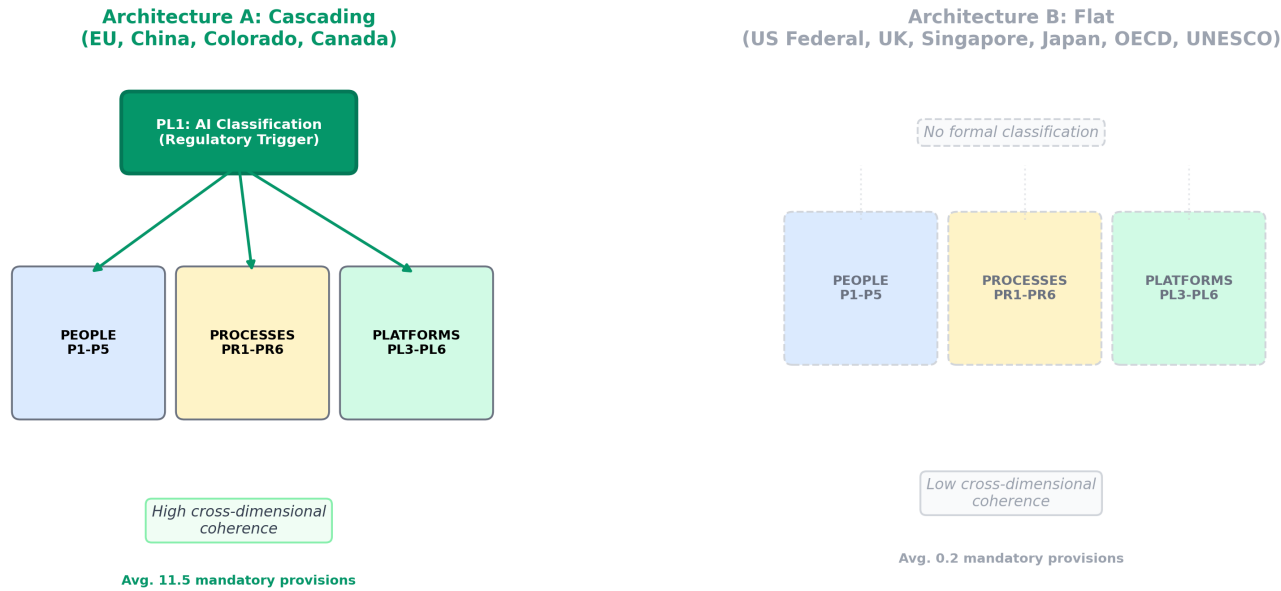


Figure 3: Classification as Cross-Dimensional Regulatory Trigger: Two Regulatory Architectures. Left (Architecture A, Cascading): In jurisdictions with formal AI classification schemes (EU, China, Colorado, Canada), classification (PL1) acts as a trigger that activates calibrated obligations across all three PPP dimensions. This produces high cross-dimensional coherence (avg. 11.5 mandatory provisions). Right (Architecture B, Flat): In jurisdictions without formal classification (US Federal, UK, Singapore, Japan, OECD, UNESCO), People, Processes, and Platforms provisions operate independently, producing low cross-dimensional coherence (avg. 0.2 mandatory provisions). This diagram formalizes the paper’s central theoretical contribution: the structural role of classification in determining regulatory architecture (Proposition 1).

Timeline of Global AI Regulatory Milestones (2019–2027)

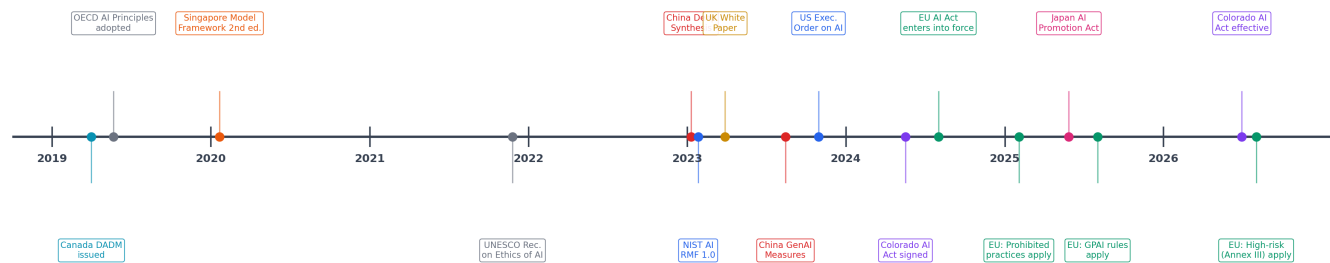


Figure 4: Timeline of Global AI Regulatory Milestones (2019–2027). Color coding indicates jurisdiction: green (EU), blue (USA), red (China), cyan (Canada), gold (UK), orange (Singapore), pink (Japan), gray (international). The timeline illustrates the three phases of AI governance evolution discussed in Section 2.1: the principles phase (2019–2021), the binding-instruments phase (2022–2023), and the comprehensive-legislation phase (2024–present). The EU AI Act’s phased implementation (2024–2027) is the most extended, with prohibited practices applying from February 2025, GPAI rules from August 2025, high-risk system obligations from August 2026, and certain product-safety high-risk rules from August 2027.