



The Indian Journal for Research in Law and Management

Open Access Law Journal – Copyright © 2026

Editor-in-Chief – Dr. Muktai Deb Chavan; Publisher – Alden Vas; ISSN: 2583-9896

This is an Open Access article distributed under the terms of the Creative Commons Attribution-Non-Commercial-Share Alike 4.0 International (CC-BY-NC-SA 4.0) License, which permits unrestricted non-commercial use, distribution, and reproduction in any medium provided the original work is properly cited.

HARMONIZING GLOBAL AI GOVERNANCE: A COMPARATIVE STUDY OF CORPORATE COMPLIANCE FRAMEWORKS UNDER THE EU AI ACT V. EMERGING US AND ASIAN REGULATORY MECHANISMS

~ NANDANA.B.S.

ABSTRACT

Artificial intelligence does not respect borders, but the law that governs it does. A single foundation model trained in California, fine-tuned in Bengaluru, and deployed to a hospital in Frankfurt now traverses at least four incompatible regulatory philosophies before it ever touches a patient. This paper examines the resulting fragmentation through the lens of the multinational firm that must actually comply. It juxtaposes the European Union's rights-based, prescriptive architecture under Regulation (EU) 2024/1689, or the 'AI Act,' entered into force on 1 August 2024 and applies on a staggered schedule. Following the 2026 Digital Omnibus amendments, the transparency requirements under Article 50 apply from 2 August 2026, while the obligations for high-risk systems under Annex III apply from 2 December 2027 and those for high-risk AI embedded in regulated products under Annex I apply from 2 August 2028. The central argument is that "harmonization", the reflexive policy aspiration of every international AI declaration since 2023, is largely a fiction at the level of binding obligation, and that convergence, where it exists, is occurring not through treaty or comity but through the extraterritorial gravity of the Brussels effect and the private ordering of technical standards. The paper concludes that the realistic near-term future is not a harmonized global regime but a stratified compliance market in which the most demanding jurisdiction sets the operational floor. For India specifically, this counsels neither imitation of Brussels nor abdication to Washington, but a calibrated interoperability strategy anchored in its existing digital public infrastructure.

I. INTRODUCTION: THE PROBLEM OF GOVERNING A BORDERLESS TECHNOLOGY WITH BORDERED LAW

There is a peculiar mismatch at the heart of artificial intelligence regulation. The technology is fungible, replicable at near-zero marginal cost, and indifferent to geography. The law that seeks to constrain it is territorial, sovereign, and increasingly divergent. This mismatch is not new, as the data protection lawyers have lived with it since the Safe Harbour litigation, but AI intensifies it, because the harms in question are neither confined to the point of deployment nor traceable to a single actor in the value chain.

Consider the multinational enterprise. It cannot choose one regime. If it operates in the European Union, it inherits the world's most detailed prohibitions and high-risk obligations. If it sells in the United States, it now confronts a federal government actively hostile to regulation yet a patchwork of assertive state legislatures. If it builds in East Asia, it must reconcile South Korea's newly enforceable framework with Japan's studied light touch. Moreover, if it seeks India's vast market, it encounters a government that, as of late 2025, had declined to enact a dedicated, comprehensive AI statute. The firm's compliance function does not experience "global AI governance." It experiences four legal universes that happen to describe the same object.

This paper interrogates whether the widely invoked goal of "harmonizing" these regimes is achievable, or whether it is a rhetorical placeholder for something more modest and more real: interoperability. The distinction matters. Harmonization implies substantive convergence on shared rules. Interoperability implies only that compliance in one jurisdiction can be leveraged, without contradiction, toward compliance in another. The former is a legislative and diplomatic project that has, on the evidence, failed. The latter is a market and standards phenomenon that is quietly succeeding.

The analysis proceeds in five movements. Part II establishes why the fragmentation is a genuine legal problem deserving academic attention rather than a transient inconvenience. Part III dissects the EU AI Act as the regulatory benchmark against which every other regime is now measured. Part IV turns to the United States, where the December 2025 executive order marks a decisive pivot from the Biden-era safety posture. Part V examines the Asian mechanisms, South Korea, Japan, and, at length, India. Part VI synthesizes the comparison and advances

concrete reform proposals. The paper is written from an Indian vantage point, but its concern is global.

II. WHY FRAGMENTATION IS A LEGAL PROBLEM, NOT MERELY A BUSINESS INCONVENIENCE

It is tempting to treat divergent AI rules as an ordinary cost of doing business, the sort of friction that internationally active firms have always absorbed. That framing understates the problem in three respects.

First, AI regulation is *extraterritorial by design*. Article 2 of the EU AI Act extends its reach to providers and deployers located outside the Union where the output of the AI system is used within it.¹ A model developer in Hyderabad with no European establishment can therefore fall squarely within the Act's jurisdiction. This is not incidental spill over; it is the deliberate importation of the "market access" logic that made the GDPR globally consequential. When the most stringent regime claims the widest jurisdictional footprint, fragmentation ceases to be a neutral menu of options and becomes a hierarchy in which one regime dominates.

Second, the *value chain is fractured across actors* in ways that existing liability doctrine handles poorly. India's own AI Governance Guidelines identify the precise difficulty: when an AI diagnostic tool errs, allocating fault among the developer, the deploying hospital, and the supervising physician is genuinely intractable under conventional tort and product-liability frameworks.² Divergent regimes multiply this problem, because each jurisdiction assigns the risk differently, the EU loads obligations onto "providers," Korea onto "AI business operators," and India, for now, onto no one in particular pending statutory amendment.

Third, and most fundamentally, the fragmentation implicates *fundamental rights unevenly*. An algorithmic hiring tool that would be classified as high-risk and subjected to conformity assessment in Europe may be entirely unregulated where deployed. The result is regulatory arbitrage of a morally significant kind: firms can route their riskiest deployments through the most permissive jurisdictions, exporting the harm while retaining the profit. This is the AI

¹ Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 (Artificial Intelligence Act), art. 2, 2024 O.J. (L 1689).

² Ministry of Electronics & Information Technology, Government of India, India AI Governance Guidelines (Nov. 5, 2025) [hereinafter India AI Governance Guidelines].

analogue of pollution havens, and it is precisely the sort of race-to-the-bottom dynamic that justifies academic and policy attention.

The problem, then, is not that compliance is expensive. It is that fragmentation systematically disadvantages the jurisdictions least able to regulate, entrenches the extraterritorial power of those most able to, and leaves the individuals harmed by AI without a coherent forum for redress. That is a legal problem worthy of the name.

III. THE EUROPEAN UNION: THE PRESCRIPTIVE BENCHMARK

A. ARCHITECTURE AND THE RISK PYRAMID

Regulation (EU) 2024/1689, or the “AI Act”, came into force on 1 August 2024 and applies on a staggered schedule. Following the 2026 Digital Omnibus amendments, the obligations for high-risk systems under Annex III apply from 2 December 2027, while those for high-risk AI embedded in regulated products under Annex I apply from 2 August 2028. Its defining feature is a tiered, “risk-based taxonomy”. Practices posing an unacceptable risk of social scoring, certain biometric categorization, manipulative subliminal techniques, are prohibited outright, with those bans having taken effect on 2 February 2025. High-risk systems, enumerated principally in Annex III, are permitted but heavily conditioned. Limited-risk systems attract transparency duties, and minimal-risk systems are essentially free.

The genius, and the burden, of this design lies in the high-risk tier. A provider of a high-risk system must implement a risk-management system, satisfy data-governance requirements, maintain technical documentation, ensure logging and human oversight, and undergo conformity assessment before placing the system on the market. These are not aspirational principles; they are operational preconditions to market access, enforced through a CE-marking regime borrowed from EU product-safety law.

B. GENERAL-PURPOSE AI AND THE COMPLIANCE TIMELINE

The Act’s treatment of general-purpose AI (GPAI) models, which is the foundation-model layer, deserves particular attention, because it is here that the compliance burden falls hardest on the handful of firms capable of building frontier systems. GPAI obligations began applying on 2 August 2025, with providers of models already on the market before that date given until

2 August 2027 to achieve full compliance.³ The layered timeline is itself a compliance variable: a firm's obligations depend not only on what its model does but on when it was placed on the market.

C. PENALTIES AND THE ENFORCEMENT TEETH

What gives the Act its gravitational pull is the sanction. Under Article 99, violations of the prohibited-practices provisions attract administrative fines of up to €35 million or 7% of total worldwide annual turnover, whichever is higher, a ceiling that deliberately exceeds even the GDPR's 4%.⁴ For a firm with global revenue in the tens of billions, this transforms compliance FROM A LEGAL NICETY INTO A BOARD-LEVEL FINANCIAL RISK.

D. THE BRUSSELS EFFECT AS DE FACTO HARMONIZATION

The critical insight for a comparative study is that the EU AI Act does not merely regulate Europe; it exports its standard. Because it is operationally cheaper for a multinational to build to a single, most-demanding specification than to maintain divergent product lines, the Act's high-risk requirements tend to become the global engineering default, the "Brussels effect" familiar from data protection and chemicals regulation. To the extent genuine convergence in corporate AI compliance exists today, it is not the product of any international agreement. It is the unilateral projection of European regulatory preference onto firms that would rather comply once than comply four times. This is harmonization by gravity, not by consent, and it is the single most important dynamic this paper identifies.

IV. THE UNITED STATES: DEREGULATION AND THE PREEMPTION TURN

If Europe is the prescriptive benchmark, the United States has become its deliberate antithesis. The trajectory is worth tracing, because it illustrates how quickly a jurisdiction's posture can invert.

A. FROM SAFETY TO INNOVATION

The Biden administration's approach, embodied in its 2023 executive order on the safe development of AI, leaned toward risk mitigation and reporting. That posture has been

³ Implementation Timeline, *supra* note 3; Regulation (EU) 2024/1689, arts. 53-55.

⁴ Regulation (EU) 2024/1689, art. 99.

decisively reversed. On 11 December 2025, President Trump signed an executive order titled “Ensuring a National Policy Framework for Artificial Intelligence,” the central purpose of which is not to regulate AI but to *prevent the states from doing so*.⁵

B. The Pre-emption Machinery

The order is remarkable for what it reveals about the American federal structure as applied to AI. It directs the Department of Justice to establish, within thirty days, a litigation task force dedicated to challenging state AI laws on constitutional (dormant Commerce Clause) and pre-emption grounds. It instructs the Commerce Department to identify “onerous” state laws, conditions federal broadband funding on state regulatory restraint, and tasks the Federal Communications Commission and Federal Trade Commission with crafting federal standards designed to displace conflicting state requirements.⁶ The animating concern, stated openly, is that state regulation is “overly burdensome, ideologically biased, or counterproductive to U.S. AI leadership.”

C. THE VACUUM AND ITS PARADOX

The consequence is a striking regulatory vacuum at the federal level combined with active hostility toward the state-level experimentation that had begun to fill it. Colorado enacted a comprehensive AI Act; California layered multiple statutes addressing workplace bias and whistle-blower protection; Utah pioneered AI consumer-protection legislation.⁷ Earlier in 2025, Congress had considered a ten-year moratorium on state enforcement, only for the Senate to strip it out by a near-unanimous 99-1 vote.⁸ The December order pursues by executive fiat what the legislature declined to grant.

For the multinational firm, the American position produces a paradox. There is minimal binding federal substance to comply *with*, yet enormous uncertainty about which state laws will survive federal challenge. Compliance planning must therefore hedge against litigation outcomes rather than statutory text, a fundamentally different exercise from the European task

⁵ Executive Order 14365, issued December 11, 2025, and published at 90 Fed. Reg. 58499.

⁶ Id.; Greenberg Traurig LLP, Trump Administration Issues Executive Order on National AI Policy Framework (Dec. 2025), <https://www.gtlaw.com/en/insights/2025/12/trump-administration-issues-executive-order-on-national-ai-policy-framework>.

⁷ Colo. Rev. Stat. § 6-1-1701 et seq. (2024); Utah Code Ann. § 13-2-12 (2024); see Greenberg Traurig LLP, *supra* note 8.

⁸ Greenberg Traurig LLP, *supra* note 8.

of engineering to a fixed specification. The United States thus contributes to global fragmentation not by adding another set of rules, but by actively destabilizing the sub-national rules that existed, and by signalling that the world's largest AI economy will not participate in a rights-based harmonization project.

V. ASIAN REGULATORY MECHANISMS: THREE DIVERGENT PATHS

Asia does not present a single model. It presents at least three, and the distance between them is instructive.

A. SOUTH KOREA: THE FIRST ENFORCER

South Korea has taken the boldest legislative step in the region. Its Framework Act on the Development of Artificial Intelligence and Establishment of Trust Foundation, or the “AI Basic Act”, was enacted in January 2025, revised on 30 December 2025, and takes effect on 22 January 2026, making Korea the first country in the world to *enforce* a comprehensive, standalone AI statute.⁹

The Korean approach occupies a middle position. Like the EU, it classifies certain systems by risk, defining “*high-impact AI*” as systems capable of significantly affecting human life, physical safety, or fundamental rights, and imposing on their operators duties of risk management, human oversight, user protection, and ex-ante impact assessment.¹⁰ It mandates labelling, including watermarks of AI-generated content, and it introduces administrative fines for non-compliance, most notably for failing to disclose that a service is AI-based.¹¹ Yet its overriding statutory purpose remains developmental: to vault Korea into the global “top-three” AI powers. The revised Act’s renaming of its governance body as the National AI Strategy Committee underscores that industrial policy, not rights protection, is the animating logic.

B. JAPAN: THE STUDIED LIGHT TOUCH

Japan’s AI Promotion Act, enacted in May 2025, sits at the permissive end of the spectrum. It establishes no risk tiering, imposes *no penalties for non-compliance*, and functions primarily

⁹ Framework Act on the Development of Artificial Intelligence and Establishment of Trust Foundation, Act No. 20676 (S. Kor.) (enacted Jan. 21, 2025; eff. Jan. 22, 2026).

¹⁰ *Id.*; see Montreal AI Ethics Institute, *AI Policy Corner: AI Governance in East Asia — Comparing the AI Acts of South Korea and Japan* (2026), <https://montrealaiethics.ai/ai-policy-corner-ai-governance-in-east-asia-comparing-the-ai-acts-of-south-korea-and-japan/>.

¹¹ *Id.*

as an incentive-driven, promotional framework.¹² Japan's stated ambition is to become "the most AI-friendly country in the world," and its statute is engineered to remove friction rather than impose it. For the comparative analyst, Japan is the clearest illustration that a formal "AI law" can exist while imposing essentially no binding corporate compliance obligation, a caution against equating the presence of legislation with the presence of regulation.

C. INDIA: GOVERNANCE WITHOUT STATUTE

India merits the most sustained treatment, both because it is the paper's home vantage and because its chosen path is the most theoretically interesting.

On 5 November 2025, MeitY released the *India AI Governance Guidelines*, the culmination of a multi-year process that drew more than 2,500 stakeholder submissions and a drafting committee constituted in July 2025.¹³ The Guidelines' defining choice is what they decline to do: they do not propose a dedicated AI statute. Instead, the drafting committee concluded that "a substantial portion of the risks associated with artificial intelligence can be effectively managed within the framework of existing legislation," recommending targeted amendments only where analysis reveals genuine gaps.¹⁴

This is a deliberate wager on *adaptation over enactment*, and it rests on four existing pillars. The Information Technology Act, 2000, is to be amended to clarify whether generative AI systems qualify as intermediaries, to attribute liability across the value chain, and to determine the reach of safe-harbour protection for systems that autonomously generate or modify content.¹⁵ The Digital Personal Data Protection Act, 2023, is to be reconciled with the realities of model training, where data-principal rights of access, correction, and deletion collide with data embedded in trained models, and where purpose limitation strains against the repurposing endemic to AI.¹⁶ Copyright questions, crystallized in the pending Delhi High Court litigation between ANI Media and OpenAI over whether training on publicly available data constitutes fair use, await both judicial resolution and a DPIIT committee's deliberation.¹⁷ And content

¹² AI Promotion Act (Act on the Promotion of Research, Development and Utilization of AI-Related Technologies) (Japan) (enacted May 2025); see Montreal AI Ethics Institute, *supra* note 12.

¹³ *India AI Governance Guidelines*, *supra* note 2; Khaitan & Co., *ERGO — India AI Governance Guidelines* (Nov. 11, 2025).

¹⁴ Khaitan & Co., *supra* note 15, at 1.

¹⁵ Information Technology Act, No. 21 of 2000, India Code; Khaitan & Co., *supra* note 15, at 1-2.

¹⁶ Digital Personal Data Protection Act, No. 22 of 2023, India Code; Khaitan & Co., *supra* note 15, at 2.

¹⁷ ANI Media Pvt. Ltd. v. OpenAI Inc., C.S. (Comm.) (Del. H.C.) (pending); Khaitan & Co., *supra* note 15, at 2.

authentication is being addressed through the October 2025 draft amendments to the Intermediary Guidelines Rules, which would require platforms enabling “synthetically generated information,” including deepfakes, to embed permanent metadata identifiers.¹⁸

Institutionally, the Guidelines propose a coordinating *AI Governance Group*, a supporting Technology and Policy Expert Committee, and an AI Safety Institute, while leaving front-line enforcement to sectoral regulators with MeitY as the nodal ministry.¹⁹ The instruments of choice are voluntary frameworks, regulatory sandboxes, techno-legal safeguards such as privacy-enhancing technologies and content-provenance tools, and a confidential national incident-reporting database designed to encourage candid disclosure without fear of penalty.²⁰

The Indian model is thus neither European prescription nor American abdication. It is a bet that a large developing economy can secure trustworthy AI through the flexible reapplication of existing law, leveraging its distinctive digital public infrastructure such as Aadhaar, UPI, the AIKosh data platform, and DEPA-style consent architectures, as governance instruments in their own right.²¹ Whether that bet pays off is the open question of the next decade.

D. THE ASIAN SPECTRUM IN SYNTHESIS

The three Asian regimes span nearly the entire regulatory range. Korea approaches European stringency in its high-impact tier while retaining a developmental soul; Japan legislates while regulating almost nothing; India governs deliberately without a dedicated statute at all. What unites them is a shared prioritization of innovation and national competitiveness over the rights-first orientation that underpins the EU Act. This is the deepest fault line in global AI governance, not East versus West, but development-first versus rights-first, and it is not a gap that technical harmonization can close.

VI. SYNTHESIS, REFORM, AND THE LIMITS OF HARMONIZATION

A. HARMONIZATION IS THE WRONG WORD

¹⁸ Information Technology (Intermediary Guidelines and Digital Media Ethics Code) (Amendment) Rules, 2025 (Draft) (Oct. 22, 2025); Khaitan & Co., *supra* note 15, at 2.

¹⁹ Khaitan & Co., *supra* note 15, at 3.

²⁰ *Id.* at 2-3.

²¹ *Id.* at 3.

The comparative picture yields an uncomfortable conclusion: substantive harmonization of AI governance is not occurring and, on present trajectories, will not occur. The EU builds an ever-more-detailed rights-based edifice. The US dismantles regulation and litigates against its own states. Japan declines to bind anyone. India routes AI governance through pre-existing statutes it has yet to amend. These are not positions on a spectrum converging toward a midpoint; they reflect irreconcilable first principles about what AI regulation is *for*.

What is occurring is *interoperability under duress*, driven by two forces. The first is the Brussels effect: the EU Act's extraterritorial reach and severe penalties make its high-risk specification the rational global engineering default. The second is "*private standardization*", the technical standards emerging from bodies such as ISO/IEC and the harmonized standards being developed to give the EU Act operational content increasingly function as the lingua franca that lets a firm's single compliance investment satisfy multiple regimes. Convergence, where it exists, is a market and standards phenomenon, not a diplomatic one.

B. REFORM PROPOSALS

The realistic reform agenda should therefore aim not at the mirage of a harmonized global code but at making fragmentation less costly and less unjust.

First, mutual recognition of conformity assessments. Rather than pursuing substantive harmonization, jurisdictions should negotiate agreements recognizing that a conformity assessment or impact assessment conducted under one regime satisfies the procedural requirements of another. This is the proven template of trade law as it does not require agreement on substance, only on process, and it would materially reduce duplicative compliance for high-risk systems.

Second, a shared value-chain liability taxonomy. India's Guidelines have correctly diagnosed that liability attribution across developers, deployers, and users is the central doctrinal gap.²² A common vocabulary defining these roles, even absent common liability rules, would let firms map obligations across borders and would give courts a shared analytical framework. This is a modest, achievable form of coordination.

²² India AI Governance Guidelines, supra note 2; Khaitan & Co., supra note 15, at 1-2.

Third, interoperable content-provenance standards. The one area of genuine cross-jurisdictional agreement is the need to authenticate AI-generated content. Korea mandates watermarking; India's draft rules require embedded metadata; the EU Act imposes transparency obligations under Article 50. A single technical standard for provenance would deliver real harmonization in the narrow domain where consensus already exists, the sensible place to start.

Fourth, for India specifically, calibrated interoperability rather than imitation. India should neither transplant the EU Act, which would burden a developing innovation ecosystem, nor drift toward the American vacuum, which would leave its citizens under-protected. Its comparative advantage lies precisely in its digital public infrastructure. By operationalizing DEPA-style consent, content-provenance tooling, and its incident-reporting database as *interoperable* governance primitives, India can position its compliance outputs to be recognized abroad while retaining a pro-innovation domestic posture. The near-term priority should be to convert the Guidelines' recommended amendments to the IT Act and DPDPA from aspiration into enacted law; because voluntary frameworks unsupported by clear, statutory liability will not withstand the first serious AI-caused harm.

C. BROADER INTELLECTUAL IMPLICATIONS

Step back, and a larger pattern emerges. The story of AI governance is a case study in a recurring twenty-first-century problem: the collision between technologies that scale globally and legal orders that remain stubbornly sovereign. The instinctive response, which is to harmonize the law, assumes a degree of shared political and moral consensus that manifestly does not exist across Brussels, Washington, Seoul, Tokyo, and New Delhi. What the comparative evidence suggests instead is that global governance of a borderless technology may proceed less through the public law of treaties and statutes than through the private law of standards, contracts, and market-driven engineering defaults. The most consequential AI regulator of the coming decade may not be any parliament, but the procurement clause of a multinational's vendor contract insisting that a supplier meet the strictest applicable standard.

That is a sobering prospect for anyone who believes that the governance of a technology this consequential should rest on democratic legitimacy rather than commercial convenience. The task for scholars and legislators is not to lament the failure of harmonization but to ensure that the interoperability filling its place is anchored, wherever possible, in public values,

transparency, accountability, and the protection of the individuals on whom these systems ultimately act. Harmonizing global AI governance, in the end, may be less about making the laws the same than about making the fragments fit together without leaving anyone in the gaps.

VII. CONCLUSION

Global AI governance is not converging toward a single legal grammar; it is settling into a hierarchy of overlapping obligations shaped by jurisdiction, market power, and technical standards. The comparison undertaken here demonstrates that the EU, the United States, South Korea, Japan, and India do not merely differ in regulatory intensity. They juxtapose different conceptions of what AI governance is meant to achieve: rights protection, national competitiveness, innovation, or institutional adaptability. For multinational firms, that divergence makes full substantive harmonization an improbable near-term project. The more credible objective is interoperability; mutual recognition of compliance processes, a shared vocabulary for value-chain responsibility, and interoperable standards for provenance and risk management. India should underpin that strategy with clear statutory liability while retaining the flexibility of its digital public infrastructure. The central regulatory challenge is therefore not to make every jurisdiction speak with one voice, but to ensure that their different legal voices do not leave consequential gaps in accountability.