



Global AI Governance Assessment for AI-Generated Content for a Globalized Digital World

Jalil Ahmad ^{1*}

¹Human Rights Institute, Southwest University of Political Science & Law, China

* Corresponding Email: jalilkhana4400@gmail.com

Received: 27 December 2025 / Revised: 26 February 2026 / Accepted: 01 March 2026 / Published online: 05 March 2026

This is an Open Access article published under the Creative Commons Attribution 4.0 International (CC BY 4.0) (<https://creativecommons.org/licenses/by/4.0/>). International Journal of Law and Legal Advancement (IJLLA) ISSN(e) 3105-0034 published by Scientific Collaborative Online Publishing Universal Academy (SCOPUA). SCOPUA stands neutral with regard to jurisdictional claims in the published maps and institutional affiliations.

ABSTRACT

AI-generated content has rapidly transformed the global information environment, challenged traditional legal systems, and exposed important gaps in international law. Through a doctrinal, comparative, and normative analysis of national laws, international conventions, bilateral agreements, and soft-law instruments adopted between 2000 and 2025, the study examines emerging regulatory models. The analysis identifies both convergence and divergence in global approaches, particularly regarding transparency, human oversight, accountability, and risk-based regulation. The findings also reveal a need for international consensus that AI systems must remain under meaningful human control and that transparency and rights protection are essential. However, significant tensions persist in enforcement mechanisms, regulatory philosophy, and geopolitical priorities, creating risks of fragmentation and regulatory arbitrage. The future of AI governance will depend on whether states can move beyond fragmented experimentation toward principled coordination. The challenge is not to eliminate diversity in regulatory models, but to anchor that diversity within shared minimum standards grounded in human dignity, accountability, and the rule of law. If international law can successfully adapt to the realities posed by AI-generated content, it will not only manage technological disruption but also reaffirm its relevance in the digital age.

Keywords: Artificial Intelligence; AI Governance; AI-Generated Content; International AI Law; Globalized Digital World

1. Introduction

AI-generated content has become a globally significant phenomenon, radically changing how information is created, shared, and used internationally (Stanikza, 2025). These days, automated media, deepfakes, generative text, and synthetic images travel instantly across digital networks, impacting public trust, economic activity, political discourse, and individual rights simultaneously in several jurisdictions (Ghiurău & Popescu, 2024). The territorial foundations of international law, which are still primarily based on human agency, state sovereignty, and geographically bounded responsibility, are directly challenged by this borderless nature. AI-generated content highlights significant gaps in current international legal frameworks as it increasingly impacts freedom of expression, privacy, reputation, and democratic processes globally (Ahmad, 2026).

The existing literature has discussed AI-generated content from a variety of angles in a disorganized way. In an effort to fit AI outputs into pre-existing human-centric legal categories, legal studies have concentrated on specific issues like copyright, authorship, liability, and intermediary responsibility (Vijendran et al., 2025). While policy and ethics literature promote values such as accountability, transparency, and human monitoring through soft-law tools and governance guidelines, parallel work in human rights law highlights threats to privacy, nondiscrimination, and freedom of expression. Despite their worth, these research are still not well-integrated. A strong international legal approach to AI-generated content is lacking because doctrinal legal approaches frequently ignore the structural transformation brought about by generative and autonomous systems, while ethical frameworks frequently lack enforceability (Romero Moreno, 2024).

In order to close that gap, this article makes the case that AI-generated content should be viewed as the beginning of the creation of international AI law rather than as a collection of discrete regulatory issues. The current study contributes in three main ways.

- First, it views AI-generated content as a multifaceted legal issue that exposes basic flaws in international law concerning jurisdiction, attribution, and accountability.



- To find new trends of convergence and divergence in AI governance, it also provides a comparative study of national, regional, and global regulatory approaches.
- Third, it makes the argument that controlling AI-generated content offers a useful and moral starting point for developing a solid, rights-focused framework for international AI law that can adapt to the realities of a globalized digital world.

2. Methodology

This study analyzes different legal sources, including UNESCO Recommendations, OECD guidelines, AI regulations, international treaties, conventions, bilateral agreements, and policy documents relevant to AI-generated content and human rights. The instruments examined cover the period from 2000 to 2025, to ensure the analysis reflects recent developments. Sources were selected based on their direct relevance to AI governance, human rights protection, and the regulation of AI-generated content. Selection prioritized texts that explicitly address accountability, transparency, human oversight, and risk management in AI systems, providing both contemporary regulatory insights and comparative perspectives from jurisdictions with established digital governance frameworks.

2.1 Analytical Framework

The study employs a combined doctrinal, comparative, and normative framework. Doctrinal analysis interprets key provisions of national, regional, and international instruments on AI accountability, transparency, human oversight, and risk management. Comparative analysis identifies similarities and differences across jurisdictions, including Pakistan, China, the EU, and Japan, as well as international frameworks such as UNESCO, OECD, and the Council of Europe, highlighting convergent principles and divergent governance models. Normative analysis evaluates the adequacy of these frameworks in addressing human rights challenges posed by AI-generated content, identifying gaps and areas for improvement.

2.2 Limitations

The study is limited by the availability of updated legal instruments and literature on AI governance, particularly in jurisdictions where regulations are emerging. Comparative analysis is restricted to jurisdictions with accessible and translated legal texts. The research focuses on legal and regulatory frameworks and does not include empirical data on enforcement or practical application.

3. Critical Analysis of AI Regulatory Frameworks

3.1 European Union's AI Act

The European Union's AI Act stands as a key regulatory framework, representing the world's first comprehensive attempt to establish a strong, legally binding set of rules governing the development, market placement, and use of AI systems (Laux, 2024). Formally adopted in 2024, this landmark legislation transcends a mere technical directive; it is a profound political and legal statement that seeks to codify a distinctly European approach to technological governance. Its foundational objective is to foster innovation and capitalize on the economic potential of AI while rigorously justifying associated risks and safeguarding the Union's fundamental rights, democratic values, and the rule of law. The regulatory architecture of the AI Act is constructed upon a risk-based arrangement, creating a four-tiered pyramid of obligations (Laux et al., 2024). At its apex, it imposes an outright prohibition on AI practices believed to pose an unacceptable risk, such as hidden manipulative techniques or real-time remote biometric identification in publicly accessible spaces for law enforcement, with narrow exceptions. The core of the regulation focuses on high-risk AI systems, which include those used in critical infrastructures, education, employment, and essential public services. These systems are subject to strict preliminary conformity assessments, rigorous data governance protocols, mandatory human oversight, and requirements for robustness, accuracy, and cybersecurity (Kusche, 2024). For limited-risk AI systems, such as those interacting with humans or generating synthetic content, the Act mandates specific transparency obligations, most notably the requirement for users to be informed they are engaging with an AI. Minimal-risk applications are largely left unregulated, though encouraged to adhere to voluntary codes of conduct. Enforcement is supported by a robust governance structure, including the establishment of a European AI Office and a scientific advisory panel, with substantial financial penalties for non-compliance. By establishing this extensive normative ecosystem, the EU AI Act aims to function as a global standard-setter, shaping international norms and ensuring that AI development progresses in alignment with a human-centric, trustworthy, and accountable paradigm (Ebers et al., 2021).

3.2 Italy-AI Law for Public Administration

Italian Law Decree No. 148 of 2023 (converted into Law No. 215 of 2023) establishes a comprehensive and principle-based governance model for the development, procurement, and deployment of artificial intelligence systems within the nation's public administrative bodies (Corea et al., 2023). This regulatory framework represents a significant and proactive effort by the Italian legislature to establish a clear national trajectory for AI in the public sector, which both anticipates and aligns with the broader ethical and regulatory principles of the European Union's AI Act. The law is fundamentally structured around a risk-based classification of AI applications, explicitly prohibiting systems considered to pose an unacceptable risk to citizens' rights and democratic processes. It mandates strict requirements for transparency and explainability, particularly for AI used in decision-making processes that affect individuals. This includes the obligation for public administrations to clearly communicate when a citizen is interacting with an AI system and to provide accessible explanations for any algorithmic-driven decisions that impact them. Furthermore, the decree protects



the principle of human oversight, requiring that final decisions, especially in sensitive areas, remain under meaningful human control and judgment. Operationally, the law creates a coordinated governance structure. It designates the Agency for Digital Italy (AgID) as the central hub for guidance and oversight, supported by a dedicated AI Task Force. A critical innovation is the establishment of a regulatory "Sandbox," a controlled environment where public administrations can test innovative AI solutions in collaboration with private developers under regulatory supervision, fostering safe experimentation (Nikolinakos, 2023).

3.3 Colorado AI Act

The Colorado Artificial Intelligence Act (CAIA), formally enacted as Senate Bill 24-205, constitutes a pioneering and distinctive piece of state-level legislation within the United States, establishing a comprehensive regulatory framework aimed specifically at mitigating algorithmic discrimination by developers and deployers of high-risk artificial intelligence systems. As one of the first statutes of its kind to be passed into law, it diverges significantly in scope and mechanism from the European Union's horizontal AI Act, focusing not on broad market regulation but on establishing direct duties of care and transparency to protect consumers, particularly in consequential domains such as employment, housing, healthcare, and financial services (Musthafa & Arundhati, 2025). The law's core innovation is its imposition of a proactive obligation on developers to employ reasonable care to avoid algorithmic discrimination, defined as differential performance disadvantaging individuals based on protected characteristics, across the AI system's lifecycle. This duty extends to deployers, who must implement risk management policies and use systems in accordance with provided documentation. Notably, the CAIA mandates a suite of public transparency measures, including the publication of statements summarizing the types of high-risk systems deployed, their purposes, and the data governance and risk mitigation strategies employed. Furthermore, it grants consumers a critical right to appeal and receive an explanation for an adverse algorithmic decision, coupled with the right to opt out of automated decision-making in favor of a human alternative in many circumstances. Enforcement is vested primarily in the state's Attorney General, with a statutory requirement for the publication of governance guidelines, eschewing a centralized ex-ante conformity assessment in favor of an ex-post enforcement model grounded in consumer protection law (Murtazashvili et al., 2025). The Act thus represents a pragmatic, operational, and rights-centric approach to AI governance at the sub-national level, seeking to balance innovation with accountability by creating enforceable standards for fairness and transparency, and positioning consumer awareness and recourse as fundamental pillars of a trustworthy AI ecosystem.

3.4 The UK's AI regulatory strategy

The United Kingdom's approach to artificial intelligence governance, as articulated in its policy paper "A pro-innovation approach to AI regulation (2023)" and subsequent government responses, establishes a deliberate and distinctive regulatory philosophy. Diverging from the European Union's comprehensive legislative model, the UK strategy is explicitly context-driven, sector-specific, and anchored in a principle-based framework designed to foster innovation and maintain strategic agility (Kazim et al., 2021). This model defers the creation of a centralized, AI-specific regulator or sweeping new legislation in the near term, opting instead to empower and coordinate existing regulatory bodies such as the Financial Conduct Authority, the Health and Safety Executive, and the Equality and Human Rights Commission to interpret and apply core principles within their domains. These cross-sectoral principles, which include safety, security, transparency, fairness, accountability, and contestability, are intended to provide consistent directional guidance while allowing for sectoral nuance (Charlesworth, 2025). The government's role is envisaged as one of central coordination, providing support through functions such as risk monitoring, regulatory sandboxing, and the development of technical standards and guidance. This approach is underpinned by a significant central investment in AI safety, notably through the establishment of the AI Safety Institute, which focuses on frontier model evaluation and national security considerations. Critically, the UK framework is presented as repetitive and adaptive; reserving the possibility of future statutory intervention should the principle-based, regulator-led model prove insufficient. This positions the UK's strategy as a conscious experiment in agile governance, seeking to balance the mitigation of tangible risks with a strategic aim to cement the country's position as a global leader in AI development and commercialization (McDougall, 2023). It is a model that prioritizes evolutionary adaptation over pre-emptive codification, placing significant trust in the capacity of a networked regulatory ecosystem to manage technological disruption effectively.

3.5 Japan-Act on Advanced Information and Communications Technology Society Promotion

The Advanced IT Society Promotion Act (enacted in 2000 and subsequently amended) represents a foundational and enabling piece of legislation within Japan's digital governance architecture. Unlike prescriptive regulatory frameworks emerging in other jurisdictions, this Act is fundamentally strategic and facilitative in nature. Its primary objective is not to impose specific restrictions on technologies like artificial intelligence, but to establish a coherent national infrastructure and policy process for systematically promoting the development and integration of advanced information and communications technology (ICT) across all sectors of society and the economy (Dirksen & Takahashi, 2020). The Act operates through two key institutional mechanisms. First, it establishes the Strategic Headquarters for the Advanced Information and Communications Technology Society, chaired by the Prime Minister, thereby elevating ICT policy to the highest level of cabinet coordination. Second, it mandates the creation of a comprehensive Basic Plan (revised approximately every five years) that outlines national priorities, strategic goals, and specific promotion policies for fields including digital infrastructure, public service digitization, and emerging technologies. This iterative planning

process provides the overarching policy direction within which more specific guidelines, such as Japan's Social Principles of Human-Centric AI (2019) and its AI Governance Guidelines, are developed and implemented (Persson et al., 2021). This Act provides the legal and administrative backbone for Japan's national digital strategy. It frames AI not as a standalone regulatory target, but as an integral component of a broader "Society 5.0" vision, where advanced ICT and cyber-physical systems are leveraged to drive economic growth and solve social challenges. Consequently, Japan's approach to AI governance is deeply contextualized within this Act's mandate, emphasizing innovation, public-private partnership, and societal integration over ex-ante regulatory classification (Kasinathan et al., 2022). The Act thus reflects a distinctive governance philosophy: one that prioritizes strategic promotion and agile, principle-based guidance within a stable, top-down policy framework, aiming to cultivate a competitive and pervasive advanced IT ecosystem.

3.6 China-Interim Administrative Measures for Generative Artificial Intelligence Services

The Interim Administrative Measures for Generative Artificial Intelligence Services, enacted by the Cyberspace Administration of China (CAC) and six other regulatory bodies and effective as of August 15, 2023, constitute a seminal and highly proactive regulatory intervention in the global landscape of artificial intelligence governance (Yuan et al., 2025). As one of the world's first comprehensive national regulations specifically targeting generative AI, these Measures articulate a distinctively holistic framework that seeks to balance vigorous support for technological innovation with unambiguous requirements for ideological security, social stability, and state-defined ethical alignment. The "Interim" designation signals a regulatory stance of adaptive observation, allowing for future refinement in response to the technology's rapid evolution. The regulatory architecture is predicated on a licensing and filing regime for service providers, establishing clear lines of accountability. Its substantive core is a set of intertwined obligations that mandate the alignment of generated content with "socialist core values," requiring the prevention of the production of content deemed subversive, discriminatory, or threatening to state security and public interests. Technically, the Measures impose rigorous requirements on training data quality, including mandates for lawful sourcing, respect for intellectual property, and the mitigation of bias, while also emphasizing the need for transparency through measures such as the clear labeling of AI-generated content (Zhang, 2024). Fundamentally, the Measures reflect a governance philosophy that views powerful generative models as both critical economic assets and potential vectors for systemic social risk. This dual imperative manifests in a regulatory approach that encourages industrial development and R&D within a firmly delineated "red line" boundary. Consequently, the framework serves a dual purpose: it functions as a domestic rulebook to shape a "healthy" and "orderly" development ecosystem for Chinese tech firms, while simultaneously positioning China as a decisive contributor to the international dialogue on AI governance, advocating for a model that prioritizes sovereign control.

3.7 Singapore National AI strategy (NAIS)

The National AI Strategy (NAIS) of Singapore functions as the nation's central, forward-looking policy blueprint, designed to systematically harness artificial intelligence as a transformative force for its economy and society. Initially launched in 2019 and substantively refreshed in 2023, the strategy articulates a vision that extends beyond technological adoption, aiming to position Singapore as a global leader in developing and deploying scalable, impactful AI solutions from within a trusted and responsible ecosystem. The refreshed strategy is architected around three systemic enablers: activity, talent, and infrastructure. It identifies 15 high-potential, cross-sectoral AI projects in domains such as finance, healthcare, and education, intended to catalyze tangible economic value and address complex national challenges (Khanal et al., 2024). Concurrently, it commits to building peaks of excellence by tripling the AI talent pool and accelerating translational research, while strengthening foundational enablers like secure data access frameworks and advanced computing infrastructure. This entire edifice is underpinned by the nation's established philosophy of pro-innovation and trusted use, operationalized through soft-law instruments like the Model AI Governance Framework and the AI Verify toolkit. Fundamentally, the NAIS is an outcome-oriented and phased master plan. It represents a dynamic, whole-of-nation effort that aligns government, industry, and research institutions toward the strategic goal of deeply and beneficially embedding AI across Singapore's social and economic fabric (Frana, 2024). The strategy is not a static declaration, but a coordinated action plan aimed at ensuring the nation's competitive resilience and sustained growth in the digital age, moving from capability-building to widespread, impactful adoption by the end of the decade.

3.8 South Korea's AI Law

South Korea's AI Basic Act, promulgated in early 2025 and taking effect in January 2026, represents one of the earliest comprehensive national AI frameworks, merging innovation policy with regulatory oversight in a unique way (Kim, 2024). Unlike risk-averse models that emphasize prohibition, the Act seeks to balance economic competitiveness and public trust by consolidating disparate AI bills into a unified legal structure that promotes R&D, supports startups, and mandates periodic national AI master plans while imposing specific obligations on developers and deployers of high-impact and generative AI systems. Notably, the legislation requires clear identification and watermarking of AI-generated content, risk assessments for systems affecting fundamental rights, and transparency measures that are designed to mitigate misinformation and deepfakes, an approach that directly engages with international concerns about AI-generated content. However, the law has drawn criticism for vague definitions and potential compliance burdens, with stakeholders warning that those ambiguous standards could hinder innovation and create regulatory uncertainty despite the government's efforts to temper enforcement through a grace period and guidance mechanisms (Park et al.,

2024). From an international law perspective, South Korea's framework illustrates an early attempt to bridge domestic AI governance with transnational regulatory norms, highlighting the tension between fostering technological growth and protecting human rights in an era of rapidly evolving AI capabilities.

4. Critical Analysis of International Conventions and Treaties on AI

4.1 Council of Europe Framework Convention on AI (The AI Convention)

The Council of Europe Framework Convention on Artificial Intelligence, commonly termed the AI Convention or the Treaty of Strasbourg, represents a landmark multilateral treaty that establishes the first binding international legal framework for the development, design, and application of artificial intelligence systems grounded in the protection of human rights, democracy, and the rule of law (Afzal, 2024). Opened for signature in September 2024, this instrument extends beyond the Council of Europe's 46 member states, inviting global accession, thereby aiming to set a universal baseline for responsible AI governance. Unlike the sectoral or risk-based regulatory models of the European Union or the United States, the Convention adopts a fundamental rights-centric approach. Its primary objective is to ensure that all activities within the lifecycle of AI systems are conducted in a manner consistent with the conventions of the Council of Europe, including the European Convention on Human Rights. It imposes legally binding obligations on Parties to adopt necessary measures to mitigate risks, prevent adverse impacts, and promote responsible innovation. The framework is structured around core requirements for transparency, oversight, and accountability, mandating the establishment of effective remedies for individuals affected by AI systems. A critical feature of the Convention is its scope of application. It covers activities by both public authorities and private actors, albeit with certain flexibilities for the latter where appropriate, reflecting the need for broad yet adaptable governance (Stoyanova, 2025). The treaty also incorporates provisions for a follow-up mechanism to monitor implementation, ensuring its continued relevance. By creating a common international legal standard, the AI Convention seeks not only to prevent a regulatory race to the bottom but also to solidify a global consensus that the immense power of AI must be firmly anchored in the safeguarding of human dignity and democratic integrity. It stands as a cornerstone in the evolving architecture of international digital law.

4.2 OECD AI Principles

The OECD AI Principles, adopted in 2019 by the Organization for Economic Co-operation and Development and subsequently endorsed by numerous non-member countries, constitute the first intergovernmental standard on artificial intelligence. These principles establish a foundational international consensus around a human-centered and trustworthy approach to AI, built upon five complementary values: inclusive growth, human rights, transparency, robustness, and accountability (Lie, 2023). While non-binding, they provide a critical normative framework that guides national policies and corporate practices, advocating for investment in responsible AI innovation and the implementation of practical tools like risk management systems. Their ongoing development and monitoring are facilitated through the OECD AI Policy Observatory, solidifying their role as a global reference point for aligning technological advancement with democratic values and societal well-being (Smith, 2025).

4.3 UNESCO Recommendation on the Ethics of AI

The UNESCO Recommendation on the Ethics of AI, agreed upon by all member states in 2021, stands as the first truly global framework for guiding the ethical development and use of artificial intelligence (Bean et al., 2025). It roots itself firmly in a commitment to human rights, dignity, and sustainability, moving the conversation beyond technical compliance toward a broader vision of societal good. The agreement outlines core principles like fairness, transparency, and human oversight, and turns them into practical areas for national action from data governance and education to supporting international cooperation. Unlike many technical standards, it explicitly connects AI ethics to cultural diversity, environmental responsibility, and inclusive peace (Ganbaatar, 2025). This recommendation provides a shared ethical compass, urging governments and organizations worldwide to ensure that AI advances not just innovation, but the common good. It's less a rulebook than a foundational call to align technological progress with enduring human values.

4.4 G7 Hiroshima AI Process International Code of Conduct

The G7 Hiroshima AI Process International Code of Conduct, established in 2023, is a voluntary international framework designed to promote safe, secure, and trustworthy AI (Habuka & Osa, 2025). It provides a common set of principles for organizations developing advanced AI systems, focusing on risk mitigation, transparency, and security throughout the AI lifecycle. The code aims to foster global interoperability and complement binding national regulations by aligning industry practices across leading economies, thereby reducing fragmentation and building a shared foundation for responsible innovation (Maremonti, 2024).

4.5 UN Global Digital Compact

The UN Global Digital Compact is a proposed multilateral framework currently under negotiation through the United Nations to establish shared principles for an open, inclusive, and secure digital future (Rasche, 2009). The Compact aims to address key digital challenges, including artificial intelligence governance, data protection, digital commons, and connectivity. It seeks to align global efforts by bringing together governments, the private sector, and civil society under a common vision that emphasizes human rights, accountability, and sustainable development (Rasche et al., 2013). While

not legally binding, the Compact is intended to provide a high-level political blueprint to guide international digital cooperation and ensure that digital technologies advance equitable progress worldwide.

4.6 Bilateral Treaties

The bilateral partnership between the United States and the United Kingdom in 2024 represents a focused, technical collaboration aimed squarely at frontier AI safety. This landmark memorandum of understanding commits both nations to aligning their scientific approaches and, where appropriate, pooling resources to develop robust methods for evaluating the most advanced AI models. The core operational mechanism is the joint work between their respective AI Safety Institutes, which will share expertise, conduct at least one joint testing exercise on a publicly accessible model, and collaborate on personnel exchanges (Birchfield, 2024). This partnership is strategically significant as it creates a dedicated channel for the world's two leading AI development hubs to establish shared scientific benchmarks for risk assessment. It focuses on pre-deployment evaluations of national security risks and severe societal harms, such as cyber threats and biochemical risks, thereby moving beyond principles into concrete, evaluative practice. While separate from broader regulatory dialogues, this bilateral effort strengthens the implementation of each country's existing strategy, the UK's pro-innovation framework, and the U.S. approach outlined in the Biden administration's Executive Order.

5. Governance Gaps and Clash of Regulatory Philosophy

The contrast between the European Union's AI Act and the United Kingdom's pro-innovation strategy exemplifies a fundamental divide in governance philosophy (Laux et al., 2023). The EU has embarked on a path of pre-emptive, detailed legislation, establishing binding ex-ante rules and conformity assessments for high-risk AI systems before they enter the market. This approach prioritizes legal certainty and aims to construct a unified, rights-based standard for its single market. In practice, it creates a structured environment designed to mitigate risk from the outset, though it may also impose considerable upfront compliance burdens on developers. Conversely, the UK has consciously deferred comprehensive legislation in favor of a context-driven, sectoral model. This model prioritizes regulatory agility and aims to avoid stifling innovation with premature or overly rigid rules. The inherent risk of this approach is one of fragmentation, where the interpretation of a core principle like "fairness" may vary significantly between different oversight bodies, leading to inconsistencies and potential gaps in protection.

China's Interim Measures for Generative AI introduce a dimension of governance that operates on a fundamentally different axis. While Western frameworks primarily focus on mitigating individualized harm, such as privacy violations or algorithmic discrimination, China's regulations explicitly prioritize state-defined social stability and ideological alignment. The rules mandate that generated content must reflect "socialist core values" and must not subvert state power or national unity (Zou & Zhang, 2025). This establishes content-based "red lines" that create a direct and profound conflict with Western liberal principles, particularly freedom of expression. A generative AI model engineered to comply with Chinese regulations is, by design, prohibited from engaging with broad categories of political, historical, or social discourse that would be permissible in other jurisdictions. For global technology platforms, this presents an almost intractable dilemma: whether to fragment their services by creating regionally specific AI models that adhere to local censorship laws, thereby contributing to a splintering of the global digital commons. This tension elevates AI governance from a technical regulatory challenge to a front in broader geopolitical and ideological competition, making meaningful global harmonization on content-related norms exceptionally difficult.

The coexistence of these divergent models generates systemic risks that extend beyond any single jurisdiction. One clear danger is that of regulatory arbitrage. The patchwork of strict (EU), flexible (UK), and facilitative (e.g., Singapore) regimes incentivizes forum shopping, where companies may choose to locate key research, development, or legal headquarters in jurisdictions with the most lenient or business-friendly rules. This strategic behavior can systematically undermine the protective objectives of stricter regulatory frameworks, as the global operations of a company may be governed by the lowest common denominator. Furthermore, these dynamic risks trigger a "race to the baseline". In contrast to regulatory paradigms where high standards in a leading market elevate global norms, a phenomenon observed in environmental regulation, the AI sector may experience the opposite. If a major economy successfully attracts capital and talent by adopting a highly permissive regulatory stance, it could place immense competitive pressure on other nations to dilute their own safeguards to retain their technological industries. This pressure threatens to erode the emerging global consensus on human-centric AI, potentially lowering universal standards for accountability, safety, and fundamental rights protection.

6. Discussion

The examination of bilateral agreements, international frameworks, and national laws reveals that AI-generated content is now viewed as a legal matter rather than just something that is technical. Legal reactions vary by jurisdiction, but they are all influenced by the common knowledge that AI-generated content can have an impact on people, public trust, and social order. As a result, several common legal issues have been reflected in the establishment of regulatory approaches that are not universal. The rejection of complete autonomy for AI systems is a fundamental feature that all frameworks under examination have in common. AI-generated content is constantly placed under human responsibility in binding laws, policy strategies, and ethical measures. Because current regulations on liability, rights protection, and remedies are built around human actors, legal systems still depend on human accountability. The enforcement of rights would be weakened, and legal uncertainty would result from allowing AI systems to function without human supervision.



Another common component is transparency. When content is created or influenced by AI, the majority of legal and policy instruments demand transparency. This requirement can take many different forms, such as explainability standards, user notification obligations, or labeling requirements. Because AI-generated content has the potential to deceive users, alter public opinion, and influence individual decision-making, there is a common emphasis on transparency. People are unable to appropriately evaluate information or contest detrimental results in the absence of clear disclosure. These frameworks are also united by the safeguarding of basic rights. Human dignity, nondiscrimination, privacy, and justice are just a few of the terms used to describe rights protection in both domestic and international legal frameworks. Human rights law and AI governance are directly linked by the Council of Europe AI Convention, and ethical guidelines from UNESCO and the OECD support similar ideals. Even regulatory frameworks that prioritize economic growth or state interests recognize the need to protect people from harm. The reason for this convergence is that AI-generated content has the potential to directly affect freedom of expression, equality, and reputation.

These methods are further connected by risk-based reasoning. While some frameworks only control certain high-impact applications, others formally categorize AI systems by risk. In both situations, the goal is to impose stronger control in areas where harm might be more likely. This common approach acknowledges that not all applications present the same degree of risk and represents a practical legal response to the various uses of AI-generated content. Cooperation and alignment are another shared characteristic at the international and bilateral levels. Coordination, common standards, and information sharing are emphasized in treaties, principles, and partnerships. AI-generated content quickly transcends national boundaries, leaving isolated domestic regulation inadequate, which is why this focus exists. Therefore, it is believed that international cooperation is essential in addressing regulatory gaps and minimizing cross-border harm. These shared characteristics demonstrate an increasing unity in legal thought despite variations in the strength of enforcement and the structure of the law. AI-generated content has to stay under human control, its use must be transparent and clear, risks must be controlled, and fundamental rights must be safeguarded, according to fundamental principles that states and international organizations are increasingly agreeing upon. Although it lays the groundwork for the future creation of more solid and unified AI governance, this new alignment does not completely eradicate regulatory diversity.

7. Conclusion

This article has demonstrated that generative systems challenge foundational assumptions of international law: that legal responsibility is human-centered, that harm can be territorially contained, and that regulation can operate effectively within national boundaries. The speed, scale, and transnational reach of AI-generated content expose gaps in attribution, jurisdiction, accountability, and enforcement that existing legal frameworks were not designed to address. At the domestic level, significant regulatory experimentation is underway. The European Union's Artificial Intelligence Act adopts a comprehensive risk-based legislative model; the United Kingdom advances a principle-driven and sectoral strategy; China's Interim Measures emphasize ideological alignment and state oversight; Singapore and Japan integrate AI governance within broader digital development agendas. These models reflect different political, legal, and cultural priorities. Yet despite their divergence, they converge on several core ideas: the rejection of full AI autonomy, the necessity of human oversight, the centrality of transparency, and the protection of fundamental rights.

References

- Afzal, J. (2024). *Implementation of Digital Law as a Legal Tool in the Current Digital Era*. Springer Nature. <https://doi.org/10.1007/978-981-97-7106-6>
- Ahmad, J. (2026). *AI-Generated Content and Violation of Human Rights: United Nations Initiatives and International Community Response | Journal of Conferences Proceedings Publication*. <https://doi.org/10.64060/ICPP.11>
- Bean, E., Burleigh, C., Haskell, C., Burris-Melville, T., Payne, J., & Pathak, B. (2025). Eavesdropping on UNESCO AI Policy, Leadership, and Ethics. *Journal of Leadership Studies*, 18(4), 98–110. <https://doi.org/10.1002/jls.70007>
- Birchfield, V. L. (2024). *From roadmap to regulation: Will there be a transatlantic approach to governing artificial intelligence?: Journal of European Integration: Vol 46, No 7*. <https://www.tandfonline.com/doi/abs/10.1080/07036337.2024.2407571>
- Charlesworth, A. (2025). AI Regulation in the UK. In V. L. Raposo (Ed.), *The European Artificial Intelligence Act: Promises and Perils?* (pp. 365–392). Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-98406-8_15
- Corea, F., Fossa, F., Loreggia, A., Quintarelli, S., & Sapienza, S. (2023). A principle-based approach to AI: The case for European Union and Italy. *AI & SOCIETY*, 38(2), 521–535. <https://doi.org/10.1007/s00146-022-01453-8>
- Dirksen, N., & Takahashi, C. S. (2020). *ARTIFICIAL INTELLIGENCE IN JAPAN 2020*.
- Ebers, M., Hoch, V. R. S., Rosenkranz, F., Ruschemeier, H., & Steinrötter, B. (2021). The European Commission's Proposal for an Artificial Intelligence Act—A Critical Assessment by Members of the Robotics and AI Law Society (RAILS). *J*, 4(4), 589–603. <https://doi.org/10.3390/j4040043>
- Frana, P. L. (2024). "Assessing Smart Nation Singapore as an International Model for AI Resp" by Philip L. Frana. <https://commons.lib.jmu.edu/ijr/vol7/iss1/2/>
- Ganbaatar, U. (2025). *Do Ethics in AI Still Matter? A Review of the 2021 UNESCO Recommendation on the Ethics of AI: The Review of Faith & International Affairs: Vol 23, No 3*. <https://www.tandfonline.com/doi/abs/10.1080/15570274.2025.2531638>



- Ghiurău, D., & Popescu, D. (2024). *Distinguishing Reality from AI: Approaches for Detecting Synthetic Content*. <https://www.mdpi.com/2073-431X/14/1/1>
- Habuka, H., & Osa, D. U. S. de la. (2025, February). *Enhancements and next steps for the G7 Hiroshima AI Process: Toward a common framework to advance human rights, democracy and rule of law* | Cambridge Forum on AI: Law and Governance | Cambridge Core. <https://www.cambridge.org/core/journals/cambridge-forum-on-ai-law-and-governance/article/enhancements-and-next-steps-for-the-g7-hiroshima-ai-process-toward-a-common-framework-to-advance-human-rights-democracy-and-rule-of-law/767476A314F4D697C10489CDBCC09C39>
- Kasinathan, P., Pugazhendhi, R., Elavarasan, R. M., Ramachandaramurthy, V. K., Ramanathan, V., Subramanian, S., Kumar, S., Nandhagopal, K., Raghavan, R. R. V., Rangasamy, S., Devendiran, R., & Alsharif, M. H. (2022). Realization of Sustainable Development Goals with Disruptive Technologies by Integrating Industry 5.0, Society 5.0, Smart Cities and Villages. *Sustainability*, 14(22), 15258. <https://doi.org/10.3390/su142215258>
- Kazim, E., Almeida, D., Kingsman, N., Kerrigan, C., Koshiyama, A., Lomas, E., & Hilliard, A. (2021). Innovation and opportunity: Review of the UK's national AI strategy. *Discover Artificial Intelligence*, 1(1), 14. <https://doi.org/10.1007/s44163-021-00014-0>
- Khanal, S., Zhang, H., & Taeihagh, A. (2024, July). *Building an AI ecosystem in a small nation: Lessons from Singapore's journey to the forefront of AI* | Humanities and Social Sciences Communications. <https://www.nature.com/articles/s41599-024-03289-7>
- Kim, J. (2024, May). *Defining risk in AI legislation: Perspectives and implications in the Korean context: Communication Research and Practice: Vol 10, No 3*. <https://www.tandfonline.com/doi/abs/10.1080/22041451.2024.2346417>
- Kusche, I. (2024). Possible harms of artificial intelligence and the EU AI act: Fundamental rights and risk. *Journal of Risk Research*, 1–14. <https://doi.org/10.1080/13669877.2024.2350720>
- Laux, J. (2024). Institutionalised distrust and human oversight of artificial intelligence: Towards a democratic design of AI governance under the European Union AI Act. *AI & SOCIETY*, 39(6), 2853–2866. <https://doi.org/10.1007/s00146-023-01777-z>
- Laux, J., Wachter, S., & Mittelstadt, B. (2023). *Trustworthy artificial intelligence and the European Union AI act: On the conflation of trustworthiness and acceptability of risk—Laux—2024—Regulation & Governance—Wiley Online Library*. <https://onlinelibrary.wiley.com/doi/full/10.1111/rego.12512>
- Laux, J., Wachter, S., & Mittelstadt, B. (2024). Three pathways for standardisation and ethical disclosure by default under the European union artificial intelligence act. *Computer Law & Security Review*, 53, 105957. <https://doi.org/10.1016/j.clsr.2024.105957>
- Lie, L. (2023). OECD Principles of Corporate Governance. In *Encyclopedia of Sustainable Management* (pp. 2497–2500). Springer, Cham. https://doi.org/10.1007/978-3-031-25984-5_1181
- Maremonti, F. (2024). *The G7 and AI Governance: Towards a Deeper and Broader Agenda*.
- McDougall, S. (2023). More Speed, Less Haste: Finding an Approach to AI Regulation that Works for the UK. *Amicus Curiae*, 5(1), 104–125. <https://doi.org/10.14296/ac.v5i1.5664>
- Murtazashvili, I., Park, L., & He, J. (2025). *Soft Law All the Way Down: Artificial Intelligence Governance Under Biden and Trump* (SSRN Scholarly Paper No. 5682562). Social Science Research Network. <https://doi.org/10.2139/ssrn.5682562>
- Musthafa, A. R., & Arundhati, G. B. (2025). Artificial Intelligence Governance Beyond Borders: The EU AI Act's Influence on Third Party Legal Frameworks and Regional Organizations. *Review of European and Comparative Law*. <https://doi.org/10.31743/recl.18927>
- Nikolinakos, N. Th. (2023). *EU Policy and Legal Framework for Artificial Intelligence, Robotics and Related Technologies—The AI Act* (Vol. 53). Springer International Publishing. <https://doi.org/10.1007/978-3-031-27953-9>
- Park, D. H., Cho, E., & Lim, Y. (2024). A Tough Balancing Act – The Evolving AI Governance in Korea. *East Asian Science, Technology and Society: An International Journal*, 18(2), 135–154. <https://doi.org/10.1080/18752160.2024.2348307>
- Persson, A., Laaksoharju, M., & Koga, H. (2021). We Mostly Think Alike: Individual Differences in Attitude Towards AI in Sweden and Japan. *The Review of Socionetwork Strategies*, 15(1), 123–142. <https://doi.org/10.1007/s12626-021-00071-y>
- Rasche, A. (2009). “A Necessary Supplement”: What the United Nations Global Compact Is and Is Not. *Business & Society*, 48(4), 511–537. <https://doi.org/10.1177/0007650309332378>
- Rasche, A., Waddock, S., & McIntosh, M. (2013). The United Nations Global Compact: Retrospect and Prospect. *Business & Society*, 52(1), 6–30. <https://doi.org/10.1177/0007650312459999>
- Romero Moreno, F. (2024). Generative AI and deepfakes: A human rights approach to tackling harmful content. *International Review of Law, Computers & Technology*, 38(3), 297–326. <https://doi.org/10.1080/13600869.2024.2324540>
- Smith, C. (2025). *Normalizing doubt: AI, democratic confidence, and the OECD trust framework* | AI & SOCIETY | Springer Nature Link. <https://link.springer.com/article/10.1007/s00146-025-02789-7>
- Stanikza, M. E. (2025). *Leveraging AI-generated and human-generated content for maximized user engagement in contentpreneurs' innovation and creativity* | Journal of Innovation and Entrepreneurship | Springer Nature Link. <https://link.springer.com/article/10.1186/s13731-025-00529-1>
- Stoyanova, V. (2025). Council of Europe Framework Convention on Artificial Intelligence: Context, Regulatory Approach and Scope of Obligations. *European Journal of Risk Regulation*, 1–30. <https://doi.org/10.1017/err.2025.10070>
- Vijendran, M., Deng, J., Chen, S., Ho, E. S. L., & Shum, H. P. H. (2025). *Artificial intelligence for geometry-based feature extraction, analysis and synthesis in artistic images: A survey* | Artificial Intelligence Review | Springer Nature Link. <https://link.springer.com/article/10.1007/s10462-024-11051-3>
- Yuan, Z., Jiang, G., & Xu, L. (2025). *Generative Artificial Intelligence and Data Governance: Challenges and Frameworks in Enterprise Applications* | IEEE Conference Publication | IEEE Xplore. <https://ieeexplore.ieee.org/abstract/document/11081956>

Zhang, A. H. (2024). The Promise and Perils of China's Regulation of Artificial Intelligence. *SSRN Electronic Journal*.
<https://doi.org/10.2139/ssrn.4708676>

Zou, M., & Zhang, L. (2025). Navigating China's regulatory approach to generative artificial intelligence and large language models. *Cambridge Forum on AI: Law and Governance*, 1, e8. <https://doi.org/10.1017/cfl.2024.4>

Declaration

Conflict of Study: The author declares that he has no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Funding Statement: This research did not receive any specific grant from funding agencies in the public, commercial, or not-for-profit sectors.

Availability of Data and Material: There are no ethical issues. All data in this paper is publicly available.

Declaration of AI Use: Not Applicable.

Consent to Publish: Not Applicable.

Ethical Approval: Not Applicable

Consent to Participate: Not Applicable

Acknowledgment: Not Applicable

