



Original Article

Cross-country artificial intelligence policy mapping: Human-centered governance and industrial adoption prospects in Indonesia

Dedy Permadi

Department of International Relations, Universitas Gadjah Mada, Yogyakarta, Indonesia

Corresponding author: dedypermadi@ugm.ac.id

Article Info	Abstract
Article history: Received Mar 29th, 2026 Revised Jun 3rd, 2026 Accepted Jun 19th, 2026 Published Jun 30th, 2026 Keywords: <i>Artificial intelligence; technology policy; comparative study; human-centered artificial intelligence</i>	This study analyzes Artificial Intelligence (AI) governance through a comparative study of regulatory frameworks in the United States (innovation-based), the European Union (risk-based), Spain (regulatory sandbox experimentation), and China (state-led sectoral). These four jurisdictions were selected to illustrate the operationalization of Human-Centered Artificial Intelligence (HCAI) principles, namely reliability, safety, and trust, through different governance pathways. Current literature on AI governance predominantly focuses on the discussion of the great power rivalry, while the gap between global standards and domestic readiness of middle-power nations like Indonesia remains underexamined. As a Global South country with complex development contexts, Indonesia cannot adopt a single governance model. This study argues that Indonesia's AI governance must transition from a risk-mitigation focus toward a hybrid approach based on readiness, incentives, and participation. A hybrid approach is more effective in integrating ethical guidelines, proportional regulation, policy experimentation, and sectoral rules in a phased manner. Integrating HCAI principles into Indonesia's digital transformation will enable equitable AI diffusion, strengthen human oversight, and align national agenda with global norm standards. This study emphasizes the urgency of strategic policies to ensure AI adoption in Indonesia fulfills socio-economic development needs, rather than serving as mere performative policy. How to Cite (APA Style): Permadi, D. (2026). Cross-country artificial intelligence policy mapping: Human-centered governance and industrial adoption prospects in Indonesia. <i>Jurnal Ilmu Sosial</i>, 25(1), 144-170. https://doi.org/10.14710/jis.25.1.2026.144-170

INTRODUCTION

The application of Artificial Intelligence (hereinafter referred to as "AI") has undergone a profound transformation over the last decade, shifting from specialized research laboratories into the mainstream of public discourse, governance, and everyday life. The acceleration of AI deployment is not confined to one dimension of technical progress but encompasses a wide range of fields: machine vision, natural language processing (NLP), speech recognition, robotics, and, most prominently, the emergence of generative and multimodal systems. Within this trajectory, generative AI (GenAI) is possibly the most notable development due to its accessibility and capacity to present information engagingly and conversationally reminiscent of human interaction.

As technology shifts, AI has substantial economic implications. A report indicates that the AI market surpassed US\$184 billion in 2024, reflecting an increase of approximately US\$50 billion from the previous year (Thormundsson, 2024). This upward trend is projected to continue, with estimates suggesting the market could reach US\$826 billion by 2030 (Thormundsson, 2024). Such figures reflect not only the commercial viability of AI but also its integration into strategic industries ranging from healthcare, energy, and finance to defense, agriculture, and education. The concentration of private investment in particular regions underscores the geopolitical dimensions of AI development. In this regard, the leading destinations for private investment in the AI sector continue to be the United States, China, and the United Kingdom. Unquestionably, numerous commercial prospects have been created due to digital- and internet-driven activities induced by the ongoing digital technology advancement, which is also relevant in the context of Indonesia (Afdhal et al., 2022; Solihati et al., 2024).

Indonesia's Coordinating Minister for Economic Affairs has stated that the country's digital sector attracted approximately US\$22 billion in investments in 2023 (Sutrisno, 2024). With the anticipated rapid growth of Indonesia's digital sector, the utilization of generative AI is increasingly demonstrating its potential as a vital driver of economic transformation and productivity across various industries (Sahin & Karayel, 2024; Salari et al., 2025). According to Boston Consulting Group's report in 2024, businesses stand to acquire considerable value by leveraging AI technologies, ranging from enhanced productivity and improved decision-making to reduced costs (Boston Consulting Group [BCG], 2024). For instance, AI has made a significant impact on research and development (R&D) within the biopharma, medtech, and automotive sectors. Additionally, in the realm of customer service, AI delivers substantial value to the banking and insurance industries (BCG, 2024). With the accelerating frequency, the intensity and breadth of AI policymaking have intensified. Policymaking in this area extends across multiple levels of governance, including local and regional administrations as well as multilateral institutions.

Public authorities, private sector firms, universities, and civil society organizations have all engaged in shaping AI-related policies, reflecting the cross-sectoral and transboundary nature of the technology. The current global landscape of AI development reflects a dynamic and productive trend toward establishing trustworthy and responsible governance policies. Over 70 countries worldwide are currently formulating their AI policies and initiatives (Organisation for Economic Cooperation and Development [OECD], 2024). Moreover, an increasing number of governments are adopting AI technologies to enhance public services, modernize the public sector, formulate data-driven policies, and create ethical frameworks for AI governance (Androniceanu, 2024).

Furthermore, advancements in AI across the globe have led to a range of technological innovations impacting various aspects of life, accompanied by both opportunities and challenges. A report by IBM underlines that many current AI solutions have effectively enhanced labor productivity and efficiency, making the adoption of AI increasingly appealing and relevant in everyday contexts (Kosinski & Scapicchio, 2024). In particular, AI possesses the capability to analyze vast amounts of data and keeps continuing to evolve autonomously through machine learning that reduces the need for human oversight. One significant advantage of AI lies in its ability to serve as a practical tool that eases monotonous or straightforward tasks, allowing individuals to concentrate on more meaningful work (Androniceanu, 2024). This capability is among the most prevalent applications of AI across various sectors, aimed at reducing the burden of everyday responsibilities. This surge of activity highlights the extent to which AI-driven applications are beginning to permeate nearly every dimension of human existence.

At the same time, it illustrates the persistent lag between technological advancement and policy response as governments and institutions struggle to keep pace with the rapid evolution of AI.

Policymakers around the world are, in effect, racing to establish frameworks that can both harness AI's transformative potential and contain its associated risks. The urgency of this task derives from the fact that, alongside its benefits, the negative consequences of AI have also become more pronounced. Without safeguards, over-reliance on AI can exacerbate wealth inequality and facilitate mass surveillance.

Regardless of the possibility, the impact of AI on the workforce is becoming increasingly evident across the globe. A recent report by Forbes (2024), for example, documented how AI-driven restructuring has led to a surge in layoffs within the US. In April 2024, it was reported that around 325,000 job losses occurred due to companies opting for AI solutions, particularly in the technology sector, such as Tesla. In March 2024, Tesla laid off approximately 14,000 employees, which represents about 10% of its workforce. This trend has raised significant concerns among workers, particularly in the US, as the implementation of AI has contributed to a decline in employment levels and resulted in the highest number of layoffs since 2016 (Forbes, 2024). Such developments illustrate the inherently double-edged character of AI: while it offers prospects for efficiency gains and innovation, it simultaneously generates vulnerabilities related to labor precarity, the concentration of power within a small number of global technology firms, and the emergence of new ethical and governance dilemmas.

Against this backdrop, a Human-Centered AI (HCAI) approach becomes crucial. HCAI emphasizes that AI should not be developed solely to maximize technical performance or market efficiency. Instead, it should be designed and governed to enhance human agency, safeguard rights, and promote social well-being. This orientation builds on sociotechnical systems thinking, which views technologies as embedded in networks of human actors, institutions, and cultural norms rather than as neutral tools. For policy, this perspective shifts attention away from purely technical regulation toward broader institutional arrangements that shape the interaction between AI systems and society.

Global governance frameworks increasingly reflect this orientation. The Organisation for Economic Co-operation and Development (OECD) AI Principles consolidate the trend by stressing inclusive growth, sustainability, and accountability as central to AI governance. The European Union's Ethics Guidelines for Trustworthy AI (2019), produced by its High-Level Expert Group on AI, articulated principles rooted in dignity, democracy, and equality, emphasizing that AI must be lawful, ethical, and robust. Spain has localized this supranational model by experimenting with regulatory sandboxes that allow for controlled testing of AI systems while maintaining oversight and citizen protection. The United States represents a contrasting model. Its AI ecosystem has historically been shaped by market-driven innovation and private-sector leadership. Regulatory interventions have been fragmented, though recent initiatives such as the Blueprint for an AI Bill of Rights signal a growing recognition of the need to protect citizens. China exemplifies a third trajectory, characterized by strong state-led governance. AI is embedded in long-term development strategies and framed as central to national modernization and digital sovereignty.

Existing scholarship on global AI governance has largely focused on the regulatory dynamics of established technological powers, particularly the United States, the European Union, and China. Much of this literature situates AI governance within broader geopolitical competition, including the intensifying US-China rivalry and the framing of AI development as a contemporary technological contest between competing political and economic systems (Hine & Floridi, 2021; Williams, 2025). In this landscape, the United States is often associated with a more decentralised and private-sector-led AI ecosystem, supported by major technology firms such as Microsoft, Google, and OpenAI, while China has advanced a more state-centred approach that connects AI development with national planning, data governance, and surveillance capacity (Nawaz et al., 2025). The European Union offers

another influential model through the AI Act and its broader regulatory architecture on data protection, cybersecurity, and digital markets, which shape AI development and deployment across sectors (Jørgensen, 2025).

While these studies are valuable in explaining how major powers define the global terms of AI governance, they provide limited guidance for countries such as Indonesia, where the central policy challenge is not only to respond to global AI standards, but also to translate them into governance mechanisms that fit domestic institutional capacity, industrial readiness, and inclusive development priorities. This research contributes to that gap by moving beyond a monolithic comparison of major AI powers. It examines the United States, the European Union, Spain, and China as four governance archetypes that offer different lessons for Indonesia. The United States illustrates the role of research capacity and innovation incentives; the European Union demonstrates how risk-based regulation can protect rights and build accountability; Spain provides a more localised example of regulatory experimentation, including the use of AI to support anticipatory governance in areas such as fire and drought prediction through large-scale data analysis (Ros-Medina et al., 2025); while China shows how AI governance can be integrated into national industrial strategy and state-led coordination.

The novelty of this article lies in using comparative AI governance as a basis for formulating a readiness-based HCAI governance pathway. This approach positions HCAI beyond ethical principle. It takes HCAI as a policy framework that connects reliability, safety, and trust with the practical conditions required for responsible AI adoption. For Indonesia, this is especially relevant because middle-power countries are increasingly expected to engage with AI as a tool for shaping norms, strengthening institutional capacity, and advancing cooperative approaches to security and development (Hanggarini, 2025). In this sense, Indonesia's AI governance pathway must address both risk mitigation and responsible adoption, ensuring that AI supports industrial innovation, public-sector transformation, and wider participation from startups, MSMEs, and other domestic actors.

This research proceeds by first discussing the conceptual development of AI and the relevance of HCAI as a governance principle, followed by an explanation of the research method and case selection. It then analyses further the AI governance approaches of the four selected jurisdictions before synthesising their relevance for Indonesia's policy readiness, industrial adoption, and inclusive digital transformation. A synthesis of key considerations in Global South discourse will also be added. The article concludes by highlighting the contribution of a readiness-based HCAI framework for Indonesia's future AI governance agenda.

AI development: Definition, scope, and landscape

Despite its ubiquity in public discourse, AI remains conceptually unsettled. There has yet to be a single generally accepted definition of the term, if not necessarily a need for one. As a novel innovation, AI might be better understood as a phenomenon where scholars and policymakers alike are still trying to gather a deeper understanding of it and its place in the world. To define it broadly, AI could be comprehended as an imitation of human intelligence by means of computing technologies, although it is not an entirely accurate definition either (Sheikh et al., 2023).

The High-Level Expert Group of the European Commission has advanced a more functional definition, explaining AI as “systems that display intelligent behavior by analysing their environment and taking actions—with some degree of autonomy—to achieve specific goals” (European Commission, 2018). While the framing has the advantage of flexibility, it remains debatable as an accurate interpretation, reflecting persistent ambiguity surrounding the phenomenon. This ambiguity is not incidental but reflects the unsettled stage of AI's development (Gabriel, 2020; Manyika, 2022). Historically, two competing visions emerged in the 1950s. The first, symbolic AI, sought to formalize knowledge and reasoning through logic, rules, and heuristic search, and for decades dominated the field through the rise of expert systems. The second, connectionist AI that

drew inspiration from the brain and pursued learning through networks of simple units (or perceptrons). Although initially marginalized, this approach resurged in the late 1980s and laid the groundwork for the deep learning and reinforcement learning systems that dominate today (Littman et al, 2022). Yet, despite these achievements, contemporary AI remains narrow in scope. Current systems excel in specific domains, such as language generation or image recognition, but lack the general reasoning, adaptability, and contextual awareness that characterize human intelligence.

The distinction between narrow AI and the still-speculative artificial general intelligence (AGI) is critical for governance debates. Narrow AI already permeates everyday life through recommendation systems, automated decision-making, fraud detection, and digital assistants, whereas AGI remains a matter of conjecture and contested timelines (Littman, et al. 2022). Governance frameworks must therefore attend to the tangible risks and impacts of narrow AI, while preparing for the speculative possibility of AGI. The limitations of current AI approaches are as important as their successes. Schwartz (2022) highlights persistent shortcomings: systems that produce biased outputs, lack transparency, or generate errors with serious social consequences. These concerns are amplified by structural issues, including a lack of diversity in AI research and development, which shapes both the design and the societal impact of AI systems. At a deeper level, unresolved “hard problems” remain, including causal reasoning, common-sense understanding, cross-domain generalization, and the embedding of AI within complex sociotechnical environments.

This ambiguity is compounded by the challenges of governance. International efforts to regulate AI have been slow, reflecting the difficulty of designing stable rules for what Ruschemeier (2023) calls an “opaque, complex, allegedly biased and rapidly changing” domain that interacts poorly with legal imperatives of certainty, transparency, explicability, and equal treatment (p. 363). One of the most prominent attempts at definitional clarity of AI has been the European Union’s AI Act. In the Act’s Annex I of Article 3, AI systems are defined into three different scopes of understanding, which are: (1) Machine learning approaches, (2) Logic and knowledge-based approaches, and (3) Statistical approaches (European Commission, 2021). These scopes were established in an effort to provide better clarity on the types of AI systems that have been most prominently featured in the discourse of AI evolution, despite arguments that it may lead to more legal uncertainty as to what can or cannot be included in regulating AI (Ruscheimer, 2023).

Furthermore, this effort to regulate and capture a better understanding of AI also stems from an observation of states’ behavior around the governance of AI, which, according to Papyshev and Yarime (2023), can be comprehended through three broad orientations of policy shared by several notable countries: (1) leading role in facilitating and initiating AI projects (e.g. China, South Korea, Japan); (2) prioritizing AI ethics and its regulation for national security purposes (e.g. countries in the EU, Norway, Mexico, and Uruguay); (3) indirect role in facilitating and promoting AI projects for other stakeholders to take over (e.g. the UK, the US, Singapore, Saudi Arabia, New Zealand, Lithuania, Ireland, India, and Australia). These themes were based on research regarding the governing sociotechnical systems that states with active regulatory mechanisms on AI have implemented, further illustrating the variety of approaches that have existed thus far. It also aims to emphasize further the experimental nature of where we are in understanding AI globally, highlighting a sense of awareness from international actors when dealing with AI, to opt for more long-term strategies in dealing with its fast-paced changes, and to be able to adjust better when said changes occur (Papyshev & Yarime, 2023).

Human-centeredness: A fundamental principle of AI

The seminal work *Human-Centered Artificial Intelligence* (2022) by Ben Shneiderman was one of the first attempts to systematically articulate what it means to conceptualize and operationalize the notion of human-centered in AI systems. Shneiderman (2022, p. 57) departs from a binary

understanding of automation versus human control by advancing a two-dimensional framework in which both automation and human agency can be maximized simultaneously. In this framing, intelligent systems are not designed to replace human decision-making but to augment it in ways that safeguard rights, enhance creativity, and promote accountability (Human-Centered Artificial Intelligence Institute, 2022). He argues that high-performance AI is not achieved through automation alone but requires a convergence of three mutually reinforcing elements: reliable, safe, and trustworthy (RST).

Reliability, as Shneiderman (2022) notes, emerges from rigorous technical practices such as maintaining audit trails to review failures, benchmarking for verification and validation, and continuously testing data quality and bias as contexts evolve. Safety, in turn, is cultivated through management strategies that create organizational cultures attentive to risk. Trust, however, extends beyond internal practices and requires independent oversight. The critical role of professional organizations, regulatory bodies, auditing firms, and certification agencies in creating accountability structures that reassure stakeholders.

Industry practice reinforces the feasibility of this approach. International Business Machines (IBM) Research (2021) has advanced a HCAI strategy that seeks to “amplify and augment rather than displace humans.” IBM operationalizes HCAI across three areas: human-AI collaboration, responsible and human-compatible AI, and interdisciplinary design for interaction. AutoAI, one of IBM’s flagship tools, automates model selection and parameter tuning but incorporates visualization techniques that allow users to compare performance and bias metrics, exemplifying the balance of automation with oversight. These practices illustrate how reliability can be engineered through auditability, how safety can be advanced by embedding fairness checks and explainability, and how trust can be cultivated by providing transparency to end-users. Industry practice thus demonstrates that RST is not merely normative but practically attainable at scale.

However, the practical feasibility of HCAI in industry does not eliminate the need for a more critical discussion of its limitations. Humans have a relevant and significant stake in technological development, which calls for a closer investigation of the human-centered perspective in AI governance (Schmager et al., 2025). In practice, HCAI design is often understood through three main aspects: technology, human factors, and ethics. The technological aspect emphasizes that AI systems should remain under human control to avoid threaten or replace human roles. The human factor aspect stresses the importance of preserving meaningful human decision-making in the design, deployment, and use of AI. Meanwhile, the ethical aspect focuses on ensuring justice, fairness, privacy, and accountability throughout the AI lifecycle (Xu et al., 2023). However, HCAI’s promise becomes limited when the language of human control and ethics is applied to AI systems that are still driven by market competition, surveillance incentives, and labour-saving pressures. Therefore, HCAI should not be treated merely as a design principle or ethical aspiration, but as a governance challenge that requires continuous negotiation between technological innovation, institutional accountability, and social protection.

The convergence between RST and international AI governance initiatives is noteworthy. A shared objective across the efforts is to ensure the adaptability of state actors in implementing AI and its rapid advancements while maintaining safeguards that align with ethical standards, transparency, and risk-mitigation values in order to foster trust and accountability in all of its various applications, which are also capable of maximizing its full potential. Such alignment is essential not only for fostering trust and accountability but also for enabling AI systems to achieve their full social and economic potential. Reflecting this consensus, the OECD reports that more than 70 countries have policies related to AI governance (OECD, 2024).

To guide these efforts, OECD (2024) asserts five principles of trustworthy and responsible AI governance: 1) Inclusive Growth, Sustainable Development, and Well-being. This principle refers to the obligation of stakeholders to be proactive in the responsible management of AI; thus, it can contribute to growth and prosperity for all parties. 2) Respect for the Rule of Law, Human Rights and Democratic Values, Including Fairness and Privacy. In this respect, AI should be consistently developed based on human-centered values, such as equality, non-discrimination, fairness, data protection and privacy, as well as rights governed by international law. 3) Transparency and Explainability. Transparency involves understanding how AI systems are developed, trained, and used. Explanation ensures that those affected by AI results can understand and test the rationale for those results. 4) Robust, Secure, and Safe. This principle emphasizes the application of a risk management approach in maintaining a robust, secure, and safe AI system. In this context, a risk management approach can assist in identifying, assessing, prioritizing, and mitigating potential risks that may arise from AI system outcomes. 5) Accountability. This principle emphasizes the responsibility of parties that the functioning of the AI systems they design, develop, manage, and operate has been in accordance with the applicable regulatory framework (OECD, 2024).

In addition to the policies issued by the OECD, other international initiatives have contributed substantially to shaping the emerging landscape of human-centered AI governance. Recommendation on the Ethics of Artificial Intelligence issued by UNESCO (2022), also sets out a comprehensive ethical and human-centered vision grounded in principles. These principles include: 1) Proportionality and Do No Harm; 2) Safety and Security; 3) Fairness and Non-Discrimination; 4) Sustainability; 5) Right to Privacy, and Data Protection; 6) Human Oversight and Determination; 7) Transparency and Explainability; 8) Responsibility and Accountability; 9) Awareness and Literacy; and 9) Multi-Stakeholder and Adaptive Governance and Collaboration (United Nations Educational, Scientific and Cultural Organization—UNESCO, 2022).

Aside from the international instruments, in the regional sector, the Association of Southeast Asian Nations (ASEAN) issued the “ASEAN Guide on AI Governance and Ethics” in 2024 to contextualize these debates within the priorities of Southeast Asia, outlining seven fundamental principles of AI governance: 1) Transparency and Explainability. Transparency refers to the openness of information on the use of AI systems, such as the type of data used and/or the purpose of its use. Explainability means providing a clear comprehension of the factors that influence AI results; 2) Fairness and Equity. AI-generated results should be free from bias, discrimination, or injustice. To prevent unfair and discriminatory results, AI developers shall take appropriate steps from data collection and pre-processing, training, and inference; 3) Security and Safety. AI systems must be safe and protected from dangerous threats such as data poisoning, dataset destruction, model inversion, byzantine attacks, and other cyberattacks that can endanger AI users; 4) Human Centricity. AI systems should prioritize human-centered values and aim to improve the well-being, health, and benefits of society. Therefore, this principle should be applied from the design, development, and implementation stages of AI. In addition, both AI developers and users need to ensure that the application of AI technology will not disrupt the labor market or existing job opportunities; 5) Privacy and Data Governance. Data protection needs to be made comprehensively, thus providing clarity on who has access to the data and when the data can be accessed. Moreover, the way data is collected, stored, processed, and deleted must be in accordance with applicable laws and regulations and ethical principles; 5) Accountability and Integrity. This principle emphasizes the responsibility of AI system managers in the event of failure or misuse that causes negative impacts, as well as the importance of mitigation measures to prevent similar events in the future; 6) Robustness and Reliability. In order to create a safe and reliable AI system, regular monitoring and effective mitigation of potential risks are necessary (ASEAN, 2024).

In essence, this convergence indicates that HCAI is no longer a peripheral concept but a central axis of global AI governance. The principles articulated by the OECD, UNESCO, and ASEAN underscore three interrelated priorities: the protection of fundamental human rights, the establishment of accountability mechanisms that distribute responsibility across developers, operators, and users, and the minimization of risks that could undermine trust in AI systems. This alignment reflects a growing recognition that the legitimacy of AI governance depends on fostering innovation and on ensuring that technological development is socially sustainable and ethically defensible. By foregrounding human dignity, fairness, and transparency, these frameworks collectively embody the essence of HCAI. They accentuate that AI should be developed and deployed as a means of augmenting rather than displacing human agency, while safeguarding societies against harms.

METHODS

This research uses a comparative study and qualitative analysis method. Wilson argues that a comparative study involves expanding national legal studies to identify potential challenges or barriers by critically assessing a country's national legal system (McConville & Chui, 2017). The comparative approach is chosen because AI governance is inherently transnational, involving legal, ethical, institutional, and technological issues that surpass national borders. By situating Indonesia's emerging AI governance within a comparative framework, this research identifies both convergences, namely areas of global normative alignment, and divergences, namely context-specific differences in regulatory design, institutional authority, and policy implementation.

This study focuses on four jurisdictions that represent distinct pathways in global AI governance: the United States, the European Union, Spain, and China. These cases were selected through purposive case selection to provide to capture the diversity of governance models currently shaping global debates instead of drawing an exhaustive mapping of all influential AI governance actors. The United States represents a decentralized and innovation-driven model; the European Union represents a rights-based and risk-sensitive regulatory model; Spain illustrates a localized and experimental sandbox approach within the EU framework; while China represents a state-led and sector-specific model that connects AI governance with national industrial strategy, data governance, and sovereignty concerns. Other influential actors, such as Singapore and Canada, are not excluded because they are insignificant, but because this study prioritizes cases that provide clearer variation across regulatory archetypes.

In addition, this research employs qualitative analysis through descriptive analysis and desk research (Kusumastuti & Khoiron, 2019). Qualitative methodology is used because it allows the study to examine the depth and meaning of policy documents, regulatory instruments, and institutional arrangements. Only data relevant to the research problem are analysed (Muhaimin, 2020). Desk research involves the systematic collection and review of secondary sources, including legislation, policy documents, academic literature, and reports from international organizations. The analysis proceeds by mapping the selected AI governance instruments, identifying how each jurisdiction operationalizes Human-Centered AI principles, and comparing their implications for Indonesia's policy readiness.

The comparison, therefore, does not treat global AI governance as a single benchmark to be adopted by Indonesia. Rather, it uses the selected jurisdictions to examine the gap between global regulatory models and domestic policy readiness. While existing literature has paid considerable attention to AI governance among major powers, particularly the United States, China, and the European Union, this study focuses on how such global lessons can be translated into a readiness-based governance pathway for a middle power and Global South country such as Indonesia. The four jurisdictions are therefore analysed as governance archetypes in order to assess how Indonesia may

develop a hybrid pathway that reflects its institutional capacity, industrial needs, and human-centered development priorities.

RESULTS AND DISCUSSION

Development of AI regulations on a global level

Just as no nation can isolate itself from the transformative impact of AI, states remain central actors in shaping the norms and institutions that govern AI. The global race to govern AI illustrates the enduring importance of the nation-state. While private firms and international organizations contribute to AI norm-setting, the ultimate authority to translate principles into binding law or enforceable practice rests with states. National approaches to AI governance have begun to crystallize, yet states have historically approached emerging technologies through distinct governance logics.

Hence, to analyse this variation, this study adopts a critical political economy framework, conceptualizing AI governance models as direct reflections of underlying political systems and state-capital relations. Through this lens, liberal market democracies tend to generate decentralized, market-driven regimes that prioritize private-sector innovation, whereas state-directed political systems integrate AI governance into broader structural mechanisms of social control and national sovereignty. The results section therefore focuses on how AI governance can be understood in terms of three overlapping functions: (1) agenda-setting, (2) regulatory experimentation, and (3) international coordination.

The United States of America—the white house executive order on the safe, secure, and trustworthy development and use of artificial intelligence

As the country with the largest number of AI models in the world, the US has established and implemented the Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence (Perrault & Clark, 2024). The Executive Order, issued by the Biden Administration on October 30, 2023, serves as a legal framework for understanding the potential benefits and mitigating the risks associated with AI (Executive Order No. 14110, 2023). The US already has several frameworks related to AI, including the 2022 Blueprint for an AI Bill of Rights. This Executive Order can be categorized as a breakthrough in strengthening the US government's commitment to the legal provisions of AI due to its binding nature and applicability to the executive branch of government, including federal agencies and departments. This feature enables it to direct how laws and policies should be implemented by the previously mentioned entities (Congressional Research Service, 2021; Executive Order No. 14110, 2023).

Although EO 14110 was later rescinded in 2025, it remains analytically useful as an illustration of the US research-driven and inter-agency approach to AI governance (National Institute of Standards and Technology [NIST], 2025). The central purpose of the Executive Order is to ensure that AI technologies are developed and utilized safely, ethically, and in a manner that reflects the principles of HCAI. As such, it seeks to balance innovation with responsibility, integrating technological advancement with societal protection. The Executive Order has several sections that specifically regulate how legal provisions can govern the development of national AI systems. For example, there is a section on the safety and security of AI technology. To achieve this objective, the Executive Order requires industries to conduct testing and evaluation, including post-deployment performance monitoring. These steps help ensure that AI systems function ethically, safely, and in line with federal laws and policies (Executive Order No. 14110, 2023). The technical procedures that ensure the safety and security of AI systems include setting standards and requiring large-scale AI model developers to share AI model test results (red-teaming) with the federal government (Eredon & Westin, 2024). The Executive Order implements these provisions by directing the National Institute

of Standards and Technology (NIST) to publish guidelines, standards, and best practices on AI safety and security. In addition, NIST must request reports on any large-scale, dual-use AI model development activities from companies with significant computing clusters (Eredon & Westin, 2024).

By institutionalizing such oversight mechanisms, the Executive Order embodies a HCAI logic that prioritizes human safety, data security, and transparency in algorithmic decision-making. It also emphasizes the social dimensions of AI adoption, particularly its implications on the labor market. Beyond testing and evaluation, the Executive Order underscores the importance of innovation, fair competition, and support for industry workers as an integral component of a comprehensive AI governance framework to maximize the use of AI systems (Eredon & Westin, 2024). Regarding innovation and competition, some of the directives provided by the Executive Order include using national AI research resources to support domestic research, providing technical assistance to developers and small companies using AI, and simplifying immigration procedures to attract AI experts. On the other hand, support for industrial workers includes producing a report on the impact of AI on the labor market, setting out mitigation measures, and publishing principles and good practices for employers to optimize the potential of AI systems and ensure the welfare of workers. All of these actions demonstrate the US government's special attention to the impact of AI technology on business and industry development, particularly in empowering industrial workers and utilizing AI technology to improve the effectiveness of business operations.

However, several important issues may affect the effectiveness of this Executive Order. As previously mentioned, the US already has several federal laws, frameworks, and guidelines related to AI. For example, the National AI Initiative Act of 2020, which lays the foundation for the expansion of AI research and development, and the White House Blueprint for an AI Bill of Rights, which emphasizes guidelines around fair access and use of AI systems (Pouget & O'Shaughnessy, 2023). These laws and guidelines share core principles as the Executive Order, namely the ethical and safe development of AI systems through extensive research and AI system feasibility assessments (White & Case, 2024). Based on these similarities, it is in the interest of the US government to align the Executive Order with existing laws and guidelines in order to create consistency in upholding and implementing the established principles in regulating the development of AI technologies. Consequently, the Executive Order can strengthen and complement existing guidelines and legal foundations.

In addition, the effectiveness of future Executive Orders on AI is also a crucial issue, given their general limitations. For example, Executive Orders may have relatively weaker legal force than laws passed by Congress due to their unilateral issuance by the President (Reinstein, 2009).

In addition, there is a possibility that a subsequent administration can reverse an Executive Order if it is deemed to be an abuse of the rules and does not represent the interests of said administration (Reinstein, 2009). Therefore, the principles in the Executive Order on AI need to be bipartisan and legislative reinforcement to gain the full support of all factions of the political parties.

European union—European Union Artificial Intelligence Act

In accordance with the US' Executive Order, the European Union Artificial Intelligence Act (EU AI Act) is a comprehensive legal framework that establishes explicit rules for harmonizing the development, market positioning, and use of AI in the EU. Formulating regulations, aside from granting membership status, emulation, foreign aid, and reconstruction, are some of the approaches regularly exercised by the EU to perform its mission as a global normative agent, which, in this case, is specific to the context of AI usage (Nurmeiga & Nugrahani, 2020).

In general, the regulatory implementation of the EU AI Act adopts a risk-based approach. It operationalizes the humane principles by calibrating the level of regulatory oversight according to the potential risks posed by AI systems, ensuring that rules are proportionally tailored to the risk level

of systems held by each service provider and user in the EU (Kosinski & Scapicchio, 2024). The EU AI Act risk division has four different levels (European Union, 2024). *First*, there is the unacceptable-risk level of AI, such as social scoring systems and biometric categorization based on race and gender, which is prohibited (Guadamuz, 2024). *Second*, the high-risk AI tier includes any automated personal data processing systems that are subject to additional regulatory requirements. *Third*, the limited-risk AI tier is subject to lighter transparency obligations. Finally, the minimal-risk tier, such as AI applications found in video games, is not governed by the legal framework.

Generally, the EU AI Act emphasizes pre- and post-market surveillance mechanisms of products and services that incorporate AI systems, especially high-risk AI systems, in the market (European Union, 2024). This dual system ensures accountability throughout the AI lifecycle and reinforces the ethical responsibility of providers. Regarding the pre-market surveillance, products and services that utilize AI systems must undergo conformity assessment under the applicable provisions. The assessment is conducted to obtain certification from an authorized body, enabling the product to be declared as conforming and then marketed. AI system providers are required to repeat the conformity assessment if there are substantial changes in the life cycle of the AI system (European Commission, 2024). Meanwhile, once an AI product or service is on the market, the authorities will be responsible for market surveillance and ensuring that providers have a post-market monitoring system in place (Novelli et al., 2024). AI system providers and users are also required to report serious incidents and damage to AI systems.

A notable innovation within the EU AI Act is the introduction of regulatory sandboxes, designed to support the development of Micro, Small, and Medium-sized Enterprises (MSMEs). As the main driver of the EU's economic strength, MSMEs play an important role in creating employment opportunities and facilitating economic development through various innovations (European People's Party Group, 2024; Gyamfi et al., 2024; Thomson, 2024). With the central role played by MSMEs, AI adoption is inevitable for businesses to keep up with the rapid pace of technology in the global industry. However, MSMEs in the EU still experience several challenges in the adoption of AI technology compared to other large companies (Arroyabe et al., 2024). These challenges include financial constraints, a lack of expertise and human resource skills in using AI technology optimally, and the difficulty of integrating AI into existing business processes (Oldemeyer et al., 2024). Departing from these problems, the EU AI Act presents a legal framework that can accommodate the interests of MSMEs in growing national industries within EU member states through the adoption of AI technologies.

The regulatory sandbox initiative has extensive impacts on AI system providers and regulators. For AI system providers, the regulatory sandbox offers guidance, training sessions, and specialized communication channels related to AI technology development (European Parliament, 2022). Such guidance ensures that businesses' exploration and development of products and services remain compliant with laws and regulations, thereby reducing risks and unintended consequences. Regulators also benefit from the regulatory sandbox. By guiding providers, regulators can gain a better understanding of the products and services developed by artificial intelligence service providers, enabling adequate rule-making, supervision, and enforcement policies (European Parliament, 2022). In addition, the regulatory sandbox can facilitate communication between all parties involved.

Nevertheless, several parts of the regulatory sandbox can be challenging for the EU AI Act. One challenge related to the harmonization of regulatory sandbox implementation. This issue is crucial considering that the EU AI Act itself regulates the implementation of the regulatory sandbox by requiring each EU member state to establish at least one AI regulatory tool at the national level (European Union, 2024). In addition, the EU AI Act permits the creation of regulatory sandboxes at

the regional or local level within each member state through cooperation with relevant and competent authorities. Achieving regulatory coherence, therefore, becomes crucial. It aims to standardize AI regulatory tools, thus MSMEs can continue to protect the fundamental rights of customers, uphold security, and promote ethical principles in all of their AI products and services.

Spain—Royal Decree 817/2023

As one of the EU member states, Spain is required to comply with the legal framework of the EU AI Act. One form of such compliance is to create a regulatory sandbox in accordance with the provisions contained in the EU AI Act. Thus, Royal Decree 817/2023 was issued by the Spanish government as a pilot program to complement the EU AI Act. This decree aims to provide an isolated testing environment, where authorities and companies developing AI systems could collaborate to create a soft law that would guide the future implementation of the EU AI Act (BonelliErede et al., 2024).

Royal Decree 817/2023 adopts EU's risk-based approach but offers more specific guidance regarding the high-risk AI systems. In line with the EU AI Act, the definition of high-risk AI systems includes systems that may pose risks to safety, health, and human rights (Arbona, 2023; Royal Decree 817/2023, 2023). This supervision applies to public administrations and certain private entities. Regarding the provisions of the surveillance procedure, Royal Decree 817/2023 also adopts similar measures to the EU AI Act, which include test development, compliance validation, monitoring, and incident mitigation (Royal Decree 817/2023, 2023). During the testing, development, and compliance validation stage, AI system providers are required to have a risk management system that serves as the basis for conducting product or service assessments, and subsequently create a declaration of compliance given to the authorized body. Meanwhile, the incident monitoring and mitigation stage requires artificial intelligence system providers to monitor and report incidents related to products and services that have been marketed.

Reflecting the EU AI Act, isolated testing mechanisms under Royal Decree 817/2023 are aimed at setting harmonized standards on AI to safeguard citizens' rights in their interactions with artificial intelligence, particularly for artificial intelligence systems developed and used by MSMEs. In addition to MSMEs, the testing facilitation is also aimed at providing implementation guidance for startups that develop and use high-risk artificial intelligence systems, general-purpose AI (GPAI), and foundational AI models. Royal Decree 817/2023 offers several benefits to the development of Spanish startups and MSMEs. One of the benefits is that businesses can access expert guidance, including insights and best practices for developing AI systems. This enables companies to optimize their AI technologies while ensuring compliance with existing regulatory requirements. These benefits can be realized because the Royal Decree regulates the provisions for the involvement of expert advisory groups in helping to maintain the compliance of participants of this regulation (Royal Decree 817/2023, 2023). In addition to expert guidance, the Royal Decree also provides specific feedback channels for cooperation between providers, such as developers, MSMEs, and technology companies, and users of AI systems—such as individuals, clients, and third parties (Tanna & Brown, 2023).

China—the interim measures for the management of generative artificial intelligence services and administrative provisions on deep synthesis

Over the past decade, China has emerged as one of the leaders in innovative research and studies related to generative AI (Omaar, 2024). This emergence has been supported by the development of intense research and academic institutions, such as Tsinghua University, which has produced most of China's top AI startups (Table, Briefings, 2024). To ensure ethical and responsible AI development, the Chinese government has enacted two key regulatory instruments: Interim Measures for the

Management of Generative Artificial Intelligence Services (2023) and the Administrative Provisions on Deep Synthesis (2023).

Interim Measures for the Management of Generative Artificial Intelligence Services was published on July 13, 2023. These regulations specifically address the use of generative AI (genAI) technologies in providing services to the public in China. The Interim Measures define genAI as models and technologies capable of generating content such as text, images, audio, or video (Cyberspace Administration of China, 2023). The term "public" in these regulations does not specifically refer only to consumers, although specific requirements seem to target consumer-facing services. Regarding technology development and governance, the Interim Measures encourage the application of genAI technologies in every field and industry. This is done by establishing adequate gen AI infrastructures that are supported by a qualified public data training resource platform (Cyberspace Administration of China, 2023). The implementation of said governance is carried out through several measures, such as requiring service providers to conduct security assessments in accordance with relevant regulations and submit certain information related to the use of algorithms to the Cyberspace Administration of China in accordance with the Algorithm Recommendation Provisions. It includes the name of the service provider, form of service, type of algorithm, and algorithm self-assessment report (Cyberspace Administration of China, 2023). In addition, the Interim Measures also require service providers to conduct pre-training, optimization training, and other activities that handle data training in accordance with the law. These measures ensure the development of AI technology governance is consistent with existing laws and regulations by emphasizing the truthfulness, accuracy, and objectivity of data.

A defining feature of the Chinese framework is the explicit requirement for service providers to apply core socialist values in the development of their AI technology. These values are reflected through prohibition against any content that incites subversion of national sovereignty or the overturning of the socialist system, endangers national security and interests, incites separatism, advocates terrorism or extremism, and promotes ethnic hatred and ethnic discrimination, violence and obscenity, as well as false and harmful information (Cyberspace Administration of China, 2023; Baker McKenzie, 2023). Applying these socialist values show the policy direction of the Interim Measures, which not only regulates the technical supervision of AI-enabled system products and services, but also the content contained in them. This ideological orientation marks a clear departure from liberal-democratic interpretations of humane principles, which typically emphasize rights-based ethics. Instead, the Chinese model articulates a form of state-embedded humanism, in which social harmony, collective welfare, and national security constitute the primary axes of human-centeredness.

Complementing the Interim Measures, Administrative Provisions on Deep Synthesis specifically regulates deep synthetic technology, commonly known as deepfakes. Deep synthetic service providers must label generated content to inform users whose information will be edited and obtain specific consent from those individuals (Sheehan, 2023; Zhang, 2023). This provision also requires deep synthetic technology providers to enhance their technical management and conduct regular reviews, evaluations, and validations of their generative or synthetic algorithm mechanisms and logic. If a deep synthesis technology provider offers tools that generate or edit faces, voices, or other biometric information, the provider is required to conduct a security assessment independently or through a professional institution (Zhang, 2023).

By embedding transparency and accountability requirements, these regulations aim to preserve public trust and mitigate the misuse of AI-generated content. These provisions are designed not only to enhance technical reliability but also to reinforce the legitimacy of state oversight in maintaining an orderly and "truthful" information environment. In this way, China's regulatory approach

institutionalizes a collectivist interpretation of human-centeredness, in which technological governance is intertwined with state-led moral frameworks and national developmental goals.

Human-centered AI governance in practice: comparative lessons from the united states, Spain, the European Union, and China

The comparative examination of AI regulatory frameworks in the US, the EU, Spain, and China reveals both convergence and divergence in their normative foundations, governance approaches, and institutional mechanisms. Central to this framework are three mutually reinforcing components, Reliability, Safety, and Trust (RST), which define how technical robustness, organizational responsibility, and independent oversight converge to ensure that AI remains both high-performing and ethically sound. Within this conceptual backdrop, the US, EU, Spain, and China have adopted regulatory frameworks that embody differing yet complementary interpretations of HCAI. Each seeks to operationalize the RST principle through distinct governance architectures, cultural orientations, and institutional mechanisms. These insights can serve as lessons learned for similar regulation efforts in other countries that are in the early stages of developing AI governance systems aligned with global standards, including Indonesia.

Table 1. Operationalization of Human-Centered AI (HCAI) in Major Jurisdictions

Dimensions	United States	European Union (EU)	Spain	China
Primary Instrument	Executive Order on Safe, Secure, and Trustworthy AI (2023)	EU Artificial Intelligence Act (2024)	Royal Decree 817/2023	Interim Measures on Generative AI (2023); Administrative Provisions on Deep Synthesis (2022)
Regulatory Approach	Innovation- and research-driven; decentralized agency coordination	Legally binding, risk-based framework harmonizing market and rights protection	Hybrid EU-aligned model with experimental sandbox governance	Model-specific, state-led governance emphasizing social stability
HCAI Orientation	Augmentation-focused; empowering marginalized communities and promoting fairness	Rights-based and participatory	Citizen-centric and collaborative	Collectivist and stability-oriented; prioritizing social harmony and safety
Reliability (R)	Technical validation and testing via NIST standards and audits	Conformity assessment and continuous monitoring	Iterative feedback and pilot testing through sandboxes	Government-led pre-deployment audit and content moderation
Safety (S)	Red-teaming, safety benchmarks, and voluntary testing	Legally mandated risk mitigation and human-in-the-loop control	Safe experimental spaces for developers and SMEs	Strict pre-launch reviews for societal risk prevention
Trust (T)	Public transparency and ethical research partnerships	Independent regulatory oversight	Multi-stakeholder engagement and policy co-creation	State trust through compliance with ideological norms
Institutional Oversight	NIST, OSTP, and sectoral agencies	European Commission and AI Office	Spanish Agency for the Supervision of Artificial Intelligence (AESIA)	Centralized state funding and supervision (National Governance Committee for the New Generation AI)
Normative Alignment	OECD & UNESCO voluntary guidelines	OECD & UNESCO human-centered principles	EU AI Act alignment emphasizing ethical innovation	Selective adherence emphasizing sovereignty and security

Source: Author's synthesis based on regulatory documents and secondary sources

Across the five selected policy instruments, AI governance is framed around a shared commitment to trustworthy and responsible development. Although each jurisdiction operationalizes these principles differently, they converge around human-centered values that prioritize safety, transparency, fairness, accountability, and responsible innovation. The US AI Executive Order translates HCAI into a research-driven and innovation-centric framework. It operationalizes these values by encouraging testing and evaluation to ensure AI technology can empower minority groups instead of attacking them (Executive Order No. 14110, 2023). The US model, therefore, advances HCAI through technical reliability and procedural safety, supported by institutions such as NIST and the OSTP that create voluntary but influential standards for trustworthy AI.

The same sentiment is echoed by the EU AI Act. It represents the most comprehensive legal codification of HCAI to date. The EU AI Act operationalizes safety and trust through a risk-based framework that classifies AI systems into four tiers of risk, ensuring that regulatory oversight scales with potential harm. High-risk systems must undergo human oversight, conformity assessments, and post-market monitoring, in which mechanisms that directly embody the OECD and UNESCO principles of transparency, proportionality, and human rights protection. The EU AI Act thus institutionalizes the trust dimension of HCAI by embedding independent enforcement through the AI Office and the European Commission, creating accountability structures that extend beyond technical compliance toward ethical responsibility. For instance, the EU AI Act prohibits the use of biometric categorizations that have the potential to incite bias against a particular race or gender (European Union, 2024).

Spain's Royal Decree 817/2023 mirrors this emphasis by embedding citizen protection and ethical oversight within its regulatory sandbox framework. Through its regulatory sandbox, Spain enables startups, MSMEs, and research entities to test AI systems in controlled environments, ensuring that ethical principles are not abstract ideals but lived practices. This approach reflects Shneiderman's vision of AI that augments human capability through iterative design and feedback, aligning with ASEAN's call for transparency, fairness, and human-centricity. Spain's model demonstrates how HCAI can be realized through co-regulation, where government guidance, civil oversight, and industry innovation reinforce each other.

By contrast, China's dual-track regulatory model, Interim Measures of Generative Artificial Intelligence Services (2023) and Administrative Provisions on Deep Synthesis (2022), interpreted HCAI within a state-centric governance framework. Here, human-centeredness is defined collectively rather than individually: the "human" in HCAI aligns with social stability, cultural integrity, and national security. The Chinese model enforces safety and reliability through stringent pre-launch audits and content moderation mechanisms, emphasizing risk containment and control over algorithmic transparency. While this diverges from liberal-democratic interpretations of HCAI, it nonetheless demonstrates a commitment to ensuring that AI operates within ethical and socially beneficial boundaries. Collectively, these instruments illustrate that HCAI has become a shared policy language. Thus, AI regulation must integrate normative principles of trust, responsibility, and inclusivity to balance innovation with human welfare.

In addition to the similarity in the principles and values contained in the regulations, there are notable similarities related to the implementation of assessment or testing for AI system providers, as part of the government's effort to standardize AI products and services and ensure they are safe and beneficial to customers. Each framework mandates pre- and post- deployment assessments for marketed AI products and services to evaluate safety, fairness, and reliability throughout the life cycle of AI technology. Moreover, all four jurisdictions provide guidelines for every AI system provider, especially business actors such as startups, businesses, and MSMEs, so that they can innovate AI

technology while still complying with existing rules and regulations. This shared orientation underscores a mutual policy goal: to cultivate an innovative, sustainable, and safe industrial ecosystem through two-way guidance between the government, authorized bodies, and AI system providers.

Nevertheless, the significant difference between the four regulations lies in the regulatory architecture and scope of enforcement. The EU AI Act has a risk-based approach that adjusts provisions based on the AI technologies' level of risk. Meanwhile, the approach taken by the Executive Order in the US is centered on innovation-based research and development, reflected in the involvement of research and education institutes in assisting authorized agencies to provide guidelines and oversee testing conducted by AI system providers. On the other hand, the two regulations issued by China follow a model-specific and state-centric approach that regulates explicitly different categories of AI models, namely genAI (the Interim Measures) and deep synthetic models (the Administrative Provisions on Deep Synthesis); hence, supervision and control can be carried out in a targeted and specific manner.

In addition, there are differences in specific provisions regarding the form of support for businesses that become AI providers and users in each country. Among the four regulations, the EU AI Act and Royal Decree 817/2023 have provisions that explicitly provide guidance, training sessions, and technical assistance to startups and MSMEs through the establishment of a regulatory sandbox. The regulatory sandbox provides a framework that allows businesses to test and evaluate products and services in a controlled environment, thereby encouraging ethical and safe AI technology innovation for national industries in each region.

Although the Executive Order on AI in the US and the two regulations issued by China also provide guidelines and technical assistance to AI system providers, the implementation of these guidelines applies to all AI service providers; thus, there is no single provision that regulates business entities specifically, like the regulatory sandbox of the EU AI Act and Royal Decree 817/2023. To elucidate how differing political and institutional contexts shape the governance of AI, the Table 1 provides a comparative synthesis of how HCAI principles, specifically RST, are embedded and operationalized within the regulatory frameworks of the US, the EU, Spain, and China.

Apart from the similarities and differences between each regulation, a significant point to be learned is the effort to balance the interests of innovation and the protection of user rights and welfare in the formulation of AI regulations. The balance of these two interests is highly crucial since AI technologies have become an inevitable phenomenon. Therefore, legal frameworks that are formulated and implemented by the government should provide room for AI providers and users to optimize the benefits offered by AI through ethical and responsible measures. This legal framework model has the potential to be applied by developing countries such as Indonesia in their efforts to utilize AI as a catalyst to accelerate national industry development. In addition, harmonization and alignment are crucial for ensuring consistency and support from cross-stakeholders, thereby enabling regulations to function as intended.

Beyond the major powers: academic studies of AI in the context of the Global South

While much of the global discourse on AI governance is shaped by the institutional architectures of the US, the EU, and China, it is crucial to recognize that the AI revolution is not geographically confined to the Global North. Countries in the Global South are increasingly engaging with the ethical, political, and economic implications of AI adoption, yet their perspectives remain marginalized in global standard-setting processes. According to Mokoena's research (2024) on the impact of AI in Africa, such perspectives are increasingly echoed across the Global South nations, particularly in Africa, where scholars and policymakers emphasize the need to shape AI development trajectories independently of Global North dominance and to address the longstanding underrepresentation of

these nations in multilateral AI governance dialogues. Considering that the existing AI policies and strategies have mainly originated from Global North nations, particularly in Europe and North America, it is increasingly crucial for Global South countries to ensure their inclusion in these discussions to prevent the formulation of AI-related policies that overlook their unique experiences and local needs.

This concern resonates with a broader call for decolonizing AI governance, which seeks to diversify the normative foundations of global policymaking. Current global initiatives, such as the OECD's Principles on AI (2024), UNESCO's Recommendation on the Ethics of AI (2022), and ASEAN's Guide on AI Governance and Ethics (2024), have contributed to the global convergence around HCAI. However, while these frameworks articulate universal values of fairness, transparency, and accountability, they often reflect Eurocentric conceptions of human-centeredness that emphasizes individual rights over collective or contextual values.

For the Global South, the task is to reinterpret these principles through the lenses of sovereignty, inclusivity, and socio-economic justice, ensuring that HCAI advances local agency rather than reinforcing dependency on foreign technology and capital. This socio-political imperative within Global South discourses is underscored by the need to assert their influence in shaping global narratives within international institutions, particularly regarding ethical considerations. Considering the EU's central role in global AI governance after the landmark establishment of the AIA, Mokoena (2024) highlights the risk of leaving the development of these policies solely to European or Global North entities. Therefore, it is indispensable for Global South countries to actively engage in multilateral policy-making processes, which have led to the creation of national strategies concerning AI of their own, to ensure that their indigenous ethical principles and values are duly considered in a multilateral setting.

Emerging examples from South Africa and India illustrate the Global South's growing normative and policy assertiveness. South Africa's Presidential Commission on the Fourth Industrial Revolution (2020) proposes an AI agenda that integrates ethics, regulation, and cultural values with developmental imperatives, signaling a vision of human-centeredness that is socially grounded and culturally situated. The text also included aspects that concern their country the most, such as regulation, ethics, and cultural aspects of the Internet, as well as technological or hardware developments in relation to AI, in addition to other aspects of broader digital transformation (Mokoena, 2024).

India, meanwhile, has adopted a policy stance that emphasizes data sovereignty and domestic innovation, linking AI development to its broader industrial strategy and economic self-reliance. As Joshi (2024) and Thormundsson (2024) note, this approach positions AI as a strategic sector contributing USD 254 billion to India's GDP in 2024, while framing data as a national asset rather than a transnational commodity. Together, these cases underscore how Global South nations are reclaiming the right to define human-centered AI in their own terms—balancing ethical imperatives with economic pragmatism. As for its stance on sovereignty, India believes in its capacity to produce its own domestic data-based technologies, not only as a potential incentive to the Indian economy as a whole but also to encourage the private sector to experiment with their “globalized data-based technology industry”, which relies on the datafication of their people and the commodification of their environment (Joshi, 2024). Both countries are merely two examples of Global South countries endeavoring to ensure that ethical and transparent use of these new technologies is undertaken, without compromising their own national policies and strategies. With the burgeoning growth of AI over the years, it is without a doubt that these countries will not let up with the pace that their Global North counterparts have established, with no exception.

In this sense, the Global South's engagement with AI governance represents not a passive reception of external norms but a transformative negotiation of global power. By embedding HCAI within developmental state frameworks, countries such as South Africa and India are articulating models that combine technological sovereignty with ethical responsibility, thereby expanding the conceptual scope of global AI governance. Their experiences highlight the necessity of viewing AI not solely as a technological artifact but as a site of political contestation, where issues of justice, representation, and epistemic equality are continuously renegotiated.

The potential of adopting ai incentive policies in Indonesia

Within this global landscape, Indonesia stands at a pivotal juncture. The country's AI trajectory, currently guided by the National Artificial Intelligence Strategy (*Stranas KA*) 2020–2024 and the Minister of Communications and Informatics Circular No. 9 of 2023 on AI Ethics, reflects an early-stage regulatory approach focused on risk prevention and ethical awareness. However, as AI adoption accelerates across industries, there is an urgent need to shift from reactive regulation to proactive facilitation. With the strategic role of government policy interventions in supporting the utilization and adoption of AI technology by industrial actors, Indonesia has the potential to implement a similar approach as an effort to develop its national economy. Moreover, the utilization of AI technology in Indonesia is expected to help the employment of 26.7 million people in the country (Kompas, 2023).

However, it is undeniable that the affirmative AI adoption policies for human-centeredness principles have not yet been established. Drawing lessons from the comparative cases, Indonesia can adapt a hybrid governance model grounded in HCAI principles. Although the countries studied in the comparative analysis are developed nations, it is crucial for developing countries to have AI adoption policies. Developing countries should not solely adopt the latest technology, but also implement policies that promote labor empowerment, which have been applied in high-income countries with adaptations to the needs of developing nations, including redesigning and rebuilding the technology (Korinek & Stiglitz, 2021). Thus, developing countries should not merely focus on adopting technology from developed nations but should direct the adoption to meet the specific needs of developing countries, one of which for Indonesia is to escape from the middle-income trap (Korinek & Stiglitz, 2021; Ministry of Communication and Digital, 2024).

Positioned as one of the largest digital economies in Southeast Asia, Indonesia occupies a strategically important role in shaping the regional discourse on HCAI. Yet, despite this potential, Indonesia's current AI governance remains primarily normative and precautionary, emphasizing ethical awareness over industrial facilitation. The National Artificial Intelligence Strategy (*Stranas KA*) 2020–2024 and Ministerial Circular No. 9 of 2023 on AI Ethics articulate the values of fairness, accountability, and transparency but lack clear mechanisms for operationalizing these principles into incentive-based policies.

Drawing lessons from the comparative experiences of the US, the EU, Spain, and China, Indonesia has the opportunity to develop a hybrid HCAI framework; one that balances regulatory oversight with developmental inclusivity. The core insight from these global regimes is that the institutionalization of HCAI is not merely a moral or technological question, but a political-economic choice that reflects differing state–market relations, innovation capacities, and social contracts. From the EU, Indonesia can draw lessons from a risk-based and rights-oriented approach that operationalizes safety and trust through proportional regulation. The EU AI Act demonstrates that human-centered governance can coexist with industrial competitiveness when oversight mechanisms, such as conformity assessments, post-market monitoring, and transparency obligations, are calibrated to the potential harm of AI applications.

This lesson is relevant for Indonesia because its current starting point of AI policy landscape has begun to articulate ethical and responsible AI principles through the National Artificial Intelligence Strategy and the Minister of Communication and Informatics Circular No. 9 of 2023 on AI Ethics. The Circular remains a form of soft regulation, intended as ethical guidance and a foundation for future higher-level regulation rather than a binding regulatory framework (Ministry of Communication and Digital Affairs, 2024). This shows that Indonesia has established an initial normative basis for AI governance. However, has not fully translated these principles into tiered regulatory interventions yet. For Indonesia, adapting the EU's lesson would mean moving from ethical guidance toward proportional regulation: stricter supervision for high-impact technologies, such as algorithmic decision-making in finance, labour, or public administration, and lighter, innovation-friendly regulation for lower-risk applications, such as agriculture or small-scale digital services.

Meanwhile, Spain's regulatory sandbox model offers a pragmatic blueprint for cultivating reliability through iterative testing and participatory design. For Indonesia, the relevance of this model lies in the fact that sandboxing is not entirely new. Indonesia has used sandbox approaches in sectors such as financial technology and digital health, indicating that regulators already have some institutional familiarity with controlled experimentation (Financial Services Authority, 2018; Ministry of Health, 2023). Nevertheless, it still remains unclear how the results of these sectoral sandboxes are systematically translated into broader policy learning, cross-sectoral regulatory standards, or AI-specific governance mechanisms. Spain's experience is therefore useful because it shows how a controlled testing environment can be designed as a structured mechanism for regulatory learning, stakeholder dialogue, and the operationalization of ethical principles. For Indonesia, where MSMEs contribute over 60 percent of GDP and remain central to national economic resilience, an AI-focused sandbox could serve as an innovation incubator and as a bridge between existing sectoral experimentation and a more coherent Human-Centered AI governance pathway (Purwanto, 2025).

Furthermore, the US' innovation-based and research-driven approach, exemplified by Executive Order 14110 (2023), highlights another dimension of HCAI relevant to Indonesia: the importance of inter-agency coordination and research infrastructure. The US model leverages universities, public research institutes, and private-sector laboratories to co-develop technical standards for fairness testing, red-teaming, and explainability. Indonesia is currently developing its national AI roadmap and white paper under the coordination of the Ministry of Communication and Digital Affairs, involving 39 ministries and agencies as well as academics, private-sector actors, and civil society (Ministry of Communication and Digital Affairs, 2025). This process shows that Indonesia has begun to build a multi-stakeholder foundation for AI governance. However, coordination alone is not sufficient unless it is translated into technical capacity for testing, auditing, and evaluating AI systems. By establishing a National AI Research and Testing Network, Indonesia could strengthen the reliability and auditability of domestic AI products while building human capital for long-term innovation. This would not only promote local R&D capacity, but also align Indonesia's national strategy with ASEAN's (2024) call for transparency, fairness, and multi-stakeholder collaboration.

China's dual-track approach, comprising the Interim Measures on Generative AI (2023) and Administrative Provisions on Deep Synthesis (2022), offers a reference point for Indonesia in figuring out its next policy step. It illustrates how AI governance can be translated into targeted rules for specific models, applications, and sectors. After establishing relatively general policy foundations and moving toward a national AI roadmap involving multiple ministries, agencies, and stakeholders across sectors, Indonesia's next policy challenge is to ensure that broad principles can be translated into sector-specific guidance that reflects different risks, institutional capacities, and development priorities. While Indonesia should avoid China's overly centralized and surveillance-oriented governance model, it can learn from its sector-specific regulatory logic. Differentiated guidelines for

AI deployment in key industries such as agriculture, energy, logistics, health, education, and public services would allow the state to maintain oversight while supporting domestic actors in scaling contextually relevant AI solutions.

Integrating these lessons, Indonesia's pathway toward HCAI should be conceived as a multi-level governance project involving three interrelated pillars: 1) *Regulatory Reliability* – Developing a risk-sensitive, transparent, and adaptive legal framework that clearly defines accountability mechanisms for AI developers, users, and regulators. This should include a national registry for high-risk AI systems, a code of conduct for algorithmic transparency, and sectoral compliance mechanisms harmonized with ASEAN's AI Governance and Ethics Guide (2024); 2) *Institutional Safety* – Establishing mechanisms that embed ethical and technical safeguards within organizations and public institutions. For example, AI auditing units within ministries or state-owned enterprises can ensure alignment with fairness, data protection, and human oversight standards, reinforcing organizational accountability; 3) *Social Trust and Inclusion* – Promoting equitable access to AI opportunities through public investment in education, reskilling, and digital literacy. Trust is not merely a product of technical transparency but of inclusive participation—ensuring that communities, workers, and SMEs understand, shape, and benefit from AI technologies.

By embedding these pillars into its economic and digital transformation agenda, Indonesia can create an HCAI ecosystem that unites industrial competitiveness with social justice. The policy objective should not be to replicate the regulatory sophistication of the EU or the innovation dynamism of the US, but to synthesize both within Indonesia's developmental context. This synthesis aligns with the vision of inclusive digital sovereignty, a governance paradigm in which AI serves national priorities while upholding universal ethical standards.

In operational terms, Indonesia could pursue several strategic actions: 1) Establish a National AI Governance Council to coordinate between ministries, academia, and private sectors, ensuring coherence across risk management, innovation incentives, and ethical oversight; 2) Implement regulatory sandboxes under the supervision of Ministry of Communication and Digital Affairs and National Research and Innovation Agency (Badan Riset dan Inovasi Nasional or BRIN) to pilot AI applications in high-potential sectors, such as agriculture, fisheries, and public health; 3) Provide fiscal and non-fiscal incentives (e.g., R&D tax breaks, innovation grants, or ethical certification rewards) for MSMEs that integrate HCAI principles in product design; 4) Develop regional AI competence centers, particularly in secondary cities, to decentralize innovation and build local capacity; 5) Institutionalize participatory mechanisms, such as ethics committees and stakeholder consultations, to ensure that public voices inform national AI strategy revisions.

The adoption policies for innovation and technology, including AI, needed by a developing country should focus on economic policies designed to manage, mitigate the impact of, and adapt to disruptions originating from innovation, ensuring that society can fully benefit from the adoption of these technologies without marginalizing any (Korinek & Stiglitz, 2021). For developing nations such as Indonesia, these contrasting models illuminate potential pathways for contextual adaptation. Thus, these policies should also determine the necessary forms of domestic investment.

Therefore, developing countries must develop a multi-pronged development strategy. If, previously, the export-based manufacturing strategy model was the most successful model for economic growth, the presence of automation technologies, including AI, requires industries to look beyond manufacturing activities and extend to secondary economic sectors, including encouraging digitization in primary sectors, like agriculture, or tertiary sectors, like service provision. When designing such strategies, public policy interventions become essential. As Greenwald & Stiglitz (2014) argue, development policies effectiveness depends on an integrated approach encompassing infrastructure developments, regulatory coherence, education investments, and taxation. Therefore,

the formulation of development strategies must consider these aspects, which also have internal (or national) impacts when determining the direction of technology adoption and technological development goals. Decisions made in determining development strategies, including technology adoption and diffusion in the national economy, will have long-term impacts (Korinek & Stiglitz, 2021).

These long-term impacts can manifest through technological learning and spillover effects. One example is from the perspective of technology that continues to evolve, thus enabling users to master technology from the early stages to advanced stages, creating spillover effects in the knowledge of the companies using the technology. On the other hand, companies that do not have access to this technology can understand the implications of using it (Korinek & Stiglitz, 2021).

However, these dynamics unfold within imperfect market structures that often constrain equitable participation in innovation, leading to a decline in workers' welfare. This condition has significant implications for developing countries, where similar structural constraints may amplify inequality. Hence, it is highly advised for these nations to combine industrial policies with labor market policies (Delli Gatti et al., 2012a; 2012b). Furthermore, a key issue in the development policy approach of developing countries that involves AI technology adoption is the condition where market prices do not adequately reflect social shadow values. While market prices can reflect the scarcity of resources and help determine efficiency-enhancing measures, they do not fully show the equitable distribution of resources or the necessary steps for economic activities. This limitation in distribution gives the government a crucial role in guiding innovation and fostering economic development to meet the social objectives needed (Korinek & Stiglitz, 2020).

AI itself can be a facilitator of this developmental governance role. When used strategically, AI enables governments and industries to identify new sources of productivity and inclusion. The shift in perspective, where technology is initially used to increase efficiency as a form of improved productivity, allows industry players to focus on increasing labor output through technology (Korinek & Stiglitz, 2021). This perspective can be applied, for example, in the agricultural sector, which can benefit from AI technology by helping agricultural industry players adjust their farming methods according to environmental conditions that optimize results. Moreover, AI and digital technologies can also help industry players reduce distribution routes in the agricultural supply chain (Korinek & Stiglitz, 2020).

The broader implication of these insights aligns with the principles of HCAI. The ultimate lesson from these comparisons is that technological governance must remain human-centered, ensuring that innovation strengthens social welfare, protects rights, and supports inclusive growth rather than deepening inequality. By synthesizing these international experiences with domestic developmental priorities, Indonesia and other Global South nations can cultivate AI ecosystems that are not only economically productive but also socially sustainable. In doing so, they align with the broader global movement toward HCAI governance, where technology serves as a means to augment human potential, safeguard collective welfare, and reinforce democratic accountability in the age of intelligent automation.

CONCLUSION

The development of AI technology is increasingly recognized as a transformative force with profound impacts in various areas, including politics, security, and economics. While public attention has largely focused on genAI, many countries have already begun issuing diverse policy instruments to regulate and harness AI. This comparative study demonstrates that the evolving landscape of AI governance reflects not only technical adaptation, but also broader ideological, institutional, and developmental contestations over how societies envision the role of technology in shaping human futures. Across the US, the EU, Spain, and China, regulatory frameworks increasingly converge around

the normative foundation of HCAI, anchored in RST, while diverging in their institutional logics, enforcement mechanisms, and interpretations of human-centeredness.

These differences reveal that AI governance is deeply intertwined with state capacity, developmental priorities, and the balance between innovation and accountability. The EU and Spain exemplify rights-based and participatory models that embed human oversight, ethical proportionality, and controlled experimentation through risk-based regulation and regulatory sandboxes. The US advances a more decentralized and innovation-driven approach, relying on research coordination, technical standards, and inter-agency mechanisms to safeguard fairness and transparency without constraining competitiveness. Meanwhile, China's state-centric model demonstrates how AI governance can also be framed around collective welfare, industrial strategy, and sector-specific control, although its centralized and surveillance-oriented elements should not be treated as a model to replicate.

The contribution of this study lies in showing that global AI governance cannot be treated as a uniform benchmark for countries with different levels of institutional and industrial readiness. Existing literature has widely discussed AI governance through the lens of major-power competition, particularly among the US, China, and the EU. However, the Indonesian case shows the importance of bridging the gap between global regulatory models and local governance realities. As a Global South country with a complex development context, Indonesia cannot be defined through a single governance characteristic or model. Its AI governance pathway should therefore be formulated not by replicating one dominant model, but by developing a hybrid and readiness-based approach that translates global lessons into domestic regulatory capacity, sectoral needs, and inclusive development priorities.

In the case of Indonesia, the government has begun to establish an initial foundation for AI governance through the National Artificial Intelligence Strategy, the Circular on AI Ethics, sectoral sandbox practices, and the ongoing development of a national AI roadmap. However, these instruments remain relatively general and have not yet been fully translated into operational, proportional, and sector-specific governance mechanisms. Indonesia's main policy challenge is therefore not only the absence of binding AI regulation, but also the need to connect existing ethical guidance and policy initiatives with concrete mechanisms for regulatory learning, technical testing, industrial adoption, and human oversight.

Accordingly, this study concludes that Indonesia must prioritize the formulation of AI-informed economic development policies that align with HCAI. This policy direction can be implemented through facilitation and incentives for industrial actors, including non-governmental actors, across the primary, secondary, and tertiary sectors, as part of diversifying economic development activities and implementing multi-pronged economic development strategies. Integrating HCAI principles into Indonesia's industrial and digital transformation agenda would enable a more equitable diffusion of AI, strengthen human oversight, and align national development with global ethical standards. By embedding RST within its economic policymaking, Indonesia can transform AI from a source of disruption into a catalyst for sustainable innovation and inclusive national progress, ensuring that AI adoption is not merely performative, but responsive to the country's economic and social development needs.

Funding Statement

This research received no external funding.

Conflict of Interest Statement

The authors declare no conflict of interest.

Data Availability Statement

The data supporting the findings of this study are available from the corresponding author upon reasonable request.

REFERENCES

- Afdhal, M., Prawiro, R., & Fenia, S. Z. (2022). Sosialisasi Penggunaan Media Sosial Tiktok Untuk Meningkatkan Penjualan di Kampung Akrilik Padang. *Jurnal Pengabdian Masyarakat Dharma Andalas*, 1(1), 102–107. <https://doi.org/10.47233/jpmda.v1i1.544>
- Aghion, P., Antonin, C., & Bunel, S. (2019). Artificial intelligence, growth, and employment: The role of policy. *Economie et Statistique*, 2019(510–512), 149–164. <https://doi.org/10.24187/ecostat.2019.510t.1994>
- Androniceanu, A. (2024). Generative artificial intelligence, present and perspectives in public administration. *Administrative Science and Management Public*, 43, 105–119. <https://doi.org/10.24818/amp/2024.43-06>
- Arbona, L. (2023). Artificial Intelligence Regulation: Between Innovation and Legislation, Spain's Pioneering Role. <https://veridas.com/en/artificial-intelligence-regulation/>
- Arroyabe, M. F., Arranz, C. F. A., De Arroyabe, I. F., & de Arroyabe, J. C. F. (2024). Analyzing AI adoption in European SMEs: A study of digital capabilities, innovation, and external environment. *Technology in Society*, 79. <https://doi.org/10.1016/j.techsoc.2024.102733>
- Association of Southeast Asian Nations (ASEAN). (2024). *ASEAN Guide on AI Governance and Ethics Contents*. https://asean.org/wp-content/uploads/2024/02/ASEAN-Guide-on-AI-Governance-and-Ethics_beautified_201223_v2.pdf
- Autor, D. H. (2015). Why are there still so many jobs? The history and future of workplace automation. *Journal of Economic Perspectives*, 29(3), 3–30. <https://doi.org/10.1257/jep.29.3.3>
- Bahrammirzaee, A., Ghatari, A. R., Ahmadi, P., & Madani, K. (2011). A hybrid credit ranking intelligent system using expert systems and artificial neural networks. *Applied Intelligence*, 34(1), 28–46. <https://doi.org/10.1007/s10489-009-0177-8>
- Baker McKenzie. (2023). *China: New interim measures to regulate generative AI*. https://insightplus.bakermckenzie.com/bm/attachment_dw.action?attkey=FRbANEucS95NMLRN47z%2BeeOgEFct8EGQJsWJiCH2WAWuU9AaVDeFglGa5oQkOMGI&nav=FRbANEucS95NMLRN47z%2BeeOgEFct8EGQbuwypnpZjc4%3D&attdocparam=pB7HEsg%2FZ312Bk8OIuOIH1c%2BY4beLEazirm3%2BK7wMU%3D&fromContentView=1
- BonelliErede, B., Bredin Prat, B. P., De Brauw, D. B., Hengeler Mueller, H. M., Slaughter and May, S., and M., & Uría Menéndez, U. M. (2024). *G7 Competition Authorities Adopt Joint Statement on the Importance of Competition in the Digital Sector*. <https://www.uria.com/documentos/galerias/7389/documento/13583/issue16.pdf?id=13583>
- Boston Consulting Group (BCG). (2024). AI Adoption in 2024: 74% of Companies Struggle to Achieve and Scale Value. <https://www.bcg.com/press/24october2024-ai-adoption-in-2024-74-of-companies-struggle-to-achieve-and-scale-value>
- Chen, H. J., Huang, S. Y., & Kuo, C. L. (2009). Using the artificial neural network to predict fraud litigation: Some empirical evidence from emerging markets. *Expert Systems with Applications*, 36(2), 1478–1484. <https://doi.org/10.1016/j.eswa.2007.11.030>
- Congressional Research Service. (2021). *Executive Orders: An Introduction*. <https://www.congress.gov/crs-product/R46738>
- Cyberspace Administration of China. Interim Measures for the Management of Generative Artificial Intelligence Services (2023). https://www.cac.gov.cn/2023-07/13/c_1690898327029107.htm
- Da Mota, M., & Scharf, S. A. Move Deliberately and Build Things: A Framework for Trustworthy Adoption of Artificial Intelligence in Canadian Public Services, Canadian Public Policy Vol. 52 No. S1 (2026). University of Toronto Press. <https://utppublishing.com/doi/epdf/10.3138/cpp.2025-013>
- Dong, Z., & Miao, Y. (2022). High-tech enterprise recognition policy and enterprise performance: Also on the moderating role of high-tech zone construction. *Macroeconomics*, 9, 141–160.
- Eloksari, E. A. (2020). AI to bring in \$366b to Indonesia's GDP by 2030. <https://www.thejakartapost.com/news/2020/10/09/ai-to-bring-in-366b-to-indonesias-gdp-by-2030.html>
- Eredon, B., & Westin, D. (2024). Overview of 'The Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence'. <https://www.pwc.com/jp/en/knowledge/column/generative-ai-regulation09.html>
- Ernst, E., Merola, R., & Samaan, D. (2019). Economics of Artificial Intelligence: Implications for the Future of Work. *IZA Journal of Labor Policy*, 9(1). <https://doi.org/10.2478/izajolp-2019-0004>
- European Commission High-Level Expert Group on Artificial Intelligence (2018). A definition of AI: Main capabilities and scientific disciplines. https://ec.europa.eu/futurium/en/system/files/ged/ai_hleg_definition_of_ai_18_december_1.pdf

- European Commission. (2024). Shaping Europe's digital future: AI Act. [https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai#:~:text=The%20AI%20Act%20\(Regulation%20\(EU,what%20AI%20has%20to%20offer](https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai#:~:text=The%20AI%20Act%20(Regulation%20(EU,what%20AI%20has%20to%20offer)
- European Commission. Annexes to the Proposal for a Regulation of the European Parliament and of the Council: Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act) and Amending Certain Union Legislative Acts (2021). Brussels.
- European People's Party Group. (2024). Empowering Entrepreneurship and SMEs: Fuelling European Prosperity. <https://www.eppgroup.eu/newsroom/empowering-entrepreneurship-and-smes-fuelling-eu-prosperity#:~:text=Europe's%20drive%20for%20economic%20prosperity,every%20corner%20of%20our%20continent>
- European Union. Regulation (EU) 2024/1689, Pub. L. No. 1689, Official Journal (OJ) of the European Union (2024). European Union. <https://artificialintelligenceact.eu/the-act/>.
- Executive Order No. 14110, Pub. L. No. 14110, The Executive Office of the President (2023). United States of America. <https://www.federalregister.gov/documents/2023/11/01/2023-24283/safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence>
- Financial Services Authority. (2018). *Peraturan Otoritas Jasa Keuangan Nomor 13/POJK.02/2018 tentang inovasi keuangan digital di sektor jasa keuangan* [Financial Services Authority Regulation Number 13/POJK.02/2018 concerning digital financial innovation in the financial services sector]. Otoritas Jasa Keuangan.
- Frana, P. L. Assessing Smart Nation Singapore as an International Model for AI Responsibility, *International Journal on Responsibility* Vol. 7 No. 1 (2024), 1-31. James Madison University. <https://doi.org/10.62365/2576-0955.1109>
- Frank, M. R., Autor, D., Bessen, J. E., Brynjolfsson, E., Cebrian, M., Deming, D. J., ... Rahwan, I. (2019, April 2). Toward understanding the impact of artificial intelligence on labor. *Proceedings of the National Academy of Sciences of the United States of America*. National Academy of Sciences. <https://doi.org/10.1073/pnas.1900949116>
- Frey, C. B., & Osborne, M. A. (2017). The future of employment: How susceptible are jobs to computerisation? *Technological Forecasting and Social Change*, 114, 254–280. <https://doi.org/10.1016/j.techfore.2016.08.019>
- Gabriel, I. (2020). Artificial intelligence, values, and alignment. *Minds and Machines*, 30(3), 411–437. <https://doi.org/10.1007/s11023-020-09548-2>
- Gasteiger, E., & Prettnner, K. (2017). *On the possibility of automation-induced stagnation* (7). <https://hdl.handle.net/10419/155784>
- Greenwald, B., & Stiglitz, J. E. (2014). Industrial Policies, the Creation of a Learning Society, and Economic Development. In J. E. Stiglitz & J. Y. Lin (Eds.), *The Industrial Policy Revolution I* (pp. 43–71). London: Palgrave Macmillan. https://doi.org/https://doi.org/10.1057/9781137335173_4
- Guadamuz, A. (2024). The EU's Artificial Intelligence Act and copyright. *The Journal of World Intellectual Property*, 28, 213–219. <https://doi.org/10.1111/jwip.12330>
- Gyamfi, S., Gerstlberger, W., Prokop, V., & Stejskal, J. (2024). A new perspective for European SMEs' innovative support analysis: Does non-financial support matter? *Heliyon*, 10(1). <https://doi.org/10.1016/j.heliyon.2023.e23796>
- Hanggarini, P. (2025). Defense Diplomacy of Middle Powers in Digital International Relations. *Andalas Journal of International Studies (AJIS)*, 14(1), 91-105. <https://dx.doi.org/10.25077/ajis.14.1.91-105.2025>
- Hine, E., & Floridi, L. (2024). Artificial intelligence with American values and Chinese characteristics: a comparative analysis of American and Chinese governmental AI policies. *Ai & Society*, 39(1), 257-278. <https://doi.org/10.1007/s00146-022-01499-8>
- Huang, X. (2022). Background of “machine for man” and the impact of government subsidy policy. *Rev. Ind. Econ*, 4, 81–101.
- Huin, S. F., Luong, L. H. S., & Abhary, K. (2003). Knowledge-based tool for planning of enterprise resources in ASEAN SMEs. *Robotics and Computer Integrated Manufacturing*, 19(5), 409–414. [https://doi.org/https://doi.org/10.1016/S0736-5845\(02\)00033-9](https://doi.org/https://doi.org/10.1016/S0736-5845(02)00033-9)
- Human-Centered Artificial Intelligence Institute. (2022). Human-centered artificial intelligence (Chapter 2). In *Artificial Intelligence Index Report 2022* (pp. 65–92). Stanford University. https://aiindex.stanford.edu/wp-content/uploads/2022/03/2022-AI-Index-Report_Master.pdf
- Issa, H., Jabbouri, R., & Palmer, M. (2022). An artificial intelligence (AI)-readiness and adoption framework for AgriTech firms. *Technological Forecasting and Social Change*, 182. <https://doi.org/10.1016/j.techfore.2022.121874>
- Jørgensen, B. N., & Ma, Z. G. (2025). Regulating AI in the energy sector: A scoping review of EU laws, challenges, and global perspectives. *Energies*, 18(9), 2359. <https://doi.org/10.3390/en18092359>
- Joshi, D. (2024). AI governance in India—law, policy, and political economy. *Communication Research and Practice*, 10(3), 328–339. <https://doi.org/10.1080/22041451.2024.2346428>
- Korinek, A., & Stiglitz, J. E. (2020). *Steering Technological Progress*.

- Korinek, A., & Stiglitz, J. E. (2021). *Artificial Intelligence, Globalization, and Strategies for Economic Development* (NBER Working Paper No. 28453).
- Kosinski, M., & Scapicchio, M. (2024). What is the Artificial Intelligence Act of the European Union (EU AI Act)? <https://www.ibm.com/think/topics/eu-ai-act>
- Kusumastuti, A., & Khoiron, A. M. (2019). *Metode Penelitian Kualitatif*. (F. Annisya & S. Sukarno, Eds.). Semarang: Lembaga Pendidikan Sukarno Pressindo (LPSP).
- Littman, M. L., Ajunwa, I., Berger, G., Boutilier, C., Currie, M., Doshi-Velez, F., Hadfield, G., Horowitz, M. C., Isbell, C., Kitano, H., Levy, K., Lyons, T., Mitchell, M., Shah, J., Sloman, S., Vallor, S., & Walsh, T. (2022). Gathering strength, gathering storms: The One Hundred Year Study on Artificial Intelligence (AI100) 2021 Study Panel Report (arXiv preprint arXiv:2210.15767)
- Liu, G., & Liu, C. (2020). Research on the technology diffusion mechanism and the implementation strategy of the artificial intelligence technology industry. *Econ. Rev. J*, 9, 109–119.
- Manyika, J. (2022). Getting AI Right: Introductory Notes on AI & Society. *Daedalus*, 151(2), 5–27. <https://www.jstor.org/stable/48662023>
- McConville, M., & Chui, W. H. (Eds.). (2017). *Research Methods for Law* (2nd ed.). Edinburgh: Edinburgh University Press.
- Ministry of Communication and Digital Affairs. (2024, January 20). *Siaran Pers No. 48/HM/KOMINFO/01/2024 tentang Kembangkan ekosistem, Wamenkominfo: SE tata kelola AI jadi soft regulation* [Press release]. <https://www.komdigi.go.id/berita/siaran-pers/detail/siaran-pers-no-48-hm-kominfo-01-2024-tentang-kembangkan-ekosistem-wamenkominfo-se-tata-kelola-ai-jadi-soft-regulation>
- Ministry of Communication and Digital Affairs. (2025, June 26). *Indonesia susun roadmap AI nasional, target jadi pemimpin digital Asia* [Press release]. <https://www.komdigi.go.id/berita/siaran-pers/detail/indonesia-susun-roadmap-ai-nasional-target-jadi-pemimpin-digital-asia>
- Ministry of Health. (2023). *Keputusan Menteri Kesehatan Nomor HK.01.07/MENKES/1280/2023 tentang pengembangan ekosistem inovasi digital kesehatan melalui regulatory sandbox* [Minister of Health Decree Number HK.01.07/MENKES/1280/2023 concerning the development of a digital health innovation ecosystem through regulatory sandbox]. Ministry of Health of the Republic of Indonesia.
- Mokoena, K. K. (2024). A holistic Ubuntu artificial intelligence ethics approach in South Africa. *Verbum et Ecclesia*, 45(1). <https://doi.org/10.4102/ve.v45i1.3100>
- Muhaimin. (2020). *Metode Penelitian Hukum*. Mataram: Mataram University Press.
- National Institute of Standards and Technology. (2025). *Executive order on safe, secure, and trustworthy artificial intelligence*. <https://www.nist.gov/artificial-intelligence/executive-order-safe-secure-and-trustworthy-artificial-intelligence>
- Nawaz, A., Ali, S. R., Bakht, N., & Noor, Z. (2025). Artificial Intelligence and the New World Order: A Comparative Study of US-China Technological Rivalry. *South Asian Studies*, 40(1), 23. pp. 23 – 41. https://pu.edu.pk/images/journal/csas/PDF/2_40_1_25.pdf
- Ningsih, S. (2023). Keterkaitan Kepentingan Indonesia dalam Pertumbuhan Ekonomi Pengetahuan dan Transformasi Identitas UMKM oleh Hyperlocal Tokopedia. *Jurnal Ilmiah Hubungan Internasional*, 19(2), 126–143. <https://doi.org/10.26593/jihi.v19i2.5354.126-143>
- Novelli, C., Hacker, P., Morley, J., Trondal, J., & Floridi, L. (2024). A Robust Governance for the AI Act: AI Office, AI Board, Scientific Panel, and National Authorities. *European Journal of Risk Regulation*, 1–25. <https://doi.org/10.1017/err.2024.57>
- Oldemeyer, L., Jede, A., & Teuteberg, F. (2024). Investigation of artificial intelligence in SMEs: a systematic review of the state of the art and the main implementation challenges. *Management Review Quarterly*. <https://doi.org/10.1007/s11301-024-00405-4>
- Omaar, H. (2024). How Innovative is China in AI? <https://itif.org/publications/2024/08/26/how-innovative-is-china-in-ai/>
- Organisation for Economic Cooperation and Development (OECD). (2024). AI Principles. <https://www.oecd.org/en/topics/ai-principles.html>
- Papyshev, G., & Yarime, M. (2023). The state's role in governing artificial intelligence: development, control, and promotion through national strategies. *Policy Design and Practice*, 6(1), 79–102. <https://doi.org/10.1080/25741292.2022.2162252>
- Png, M.-T. (2022). At the tensions of South and North: Critical roles of Global South stakeholders in AI governance. In *FAccT '22: Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency* (pp. 1434–1445). Association for Computing Machinery. <https://doi.org/10.1145/3531146.3533200>
- Portulans Institute. (2024). *Network Readiness Index 2024*. <https://download.networkreadinessindex.org/reports/data/2024/nri-2024.pdf>
- Pouget, H., & O'Shaughnessy, M. (2023). Reconciling the US Approach to AI. <https://carnegieendowment.org/research/2023/05/reconciling-the-us-approach-to-ai?lang=en>

- Qin, Y., Xu, Z., Wang, X., & Skare, M. (2024). Artificial Intelligence and Economic Development: An Evolutionary Investigation and Systematic Review. *Journal of the Knowledge Economy*, 15(1), 1736–1770. <https://doi.org/10.1007/s13132-023-01183-2>
- Reinstein, R. J. (2009). The Limits of Executive Power. *American University Law Review*, 59(2), 259–337. <http://digitalcommons.wcl.american.edu/aulr>
- Ricaurte, P. (2022). Ethics for the majority world: AI and the question of violence at scale. *Media, Culture & Society*, 44(4), 726–745. <https://doi.org/10.1177/01634437211073547>
- Ros-Medina, J. L., Mayor-Balsas, J. M., Ferreira-Dias, T., & de Menezes-Neto, E. J. (2025). Spain and artificial intelligence: installed capacities, challenges, and opportunities. *Retos*, 225–240. Retrived from <https://doi.org/10.17163/ret.n30.2025.02>
- Royal Decree 817/2023. (2023). Royal Decree 817/2023, of November 8, establishing a controlled testing environment for compliance testing with the proposed Regulation of the European Parliament and of the Council establishing harmonized standards on artificial intelligence. https://www-boe-es.translate.goog/diario_boe/txt.php?id=BOE-A-2023-22767&x_tr_sl=es&x_tr_tl=en&x_tr_hl=en&x_tr_pto=sc
- Ruscheimer, H. (2023). AI as a challenge for legal regulation – the scope of application of the Artificial Intelligence Act proposal. *ERA Forum*, 23(3), 361–376. <https://doi.org/10.1007/s12027-022-00725-6>
- Şahin, O., & Karayel, D. (2024). Generative Artificial Intelligence (GenAI) in Business: A Systematic Review on the Threshold of Transformation. *Journal of Smart Systems Research*, 5(2), 156–175. <https://doi.org/10.58769/joinsr.1597110>
- Salari, N., Beiromvand, M., Hosseini-Far, A., Habibi, J., Babajani, F., & Mohammadi, M. (2025). Impacts of generative artificial intelligence on the future of the labor market: A systematic review. *Computers in Human Behavior Reports*, 18, 100652. <https://doi.org/https://doi.org/10.1016/j.chbr.2025.100652>
- Schmager, S., Pappas, I. O., & Vassilakopoulou, P. (2025). Understanding Human-Centred AI: a review of its defining elements and a research agenda. *Behaviour & Information Technology*, 44(15), 3771–3810. <https://doi.org/10.1080/0144929X.2024.2448719>
- Schwartz, R., Vassilev, A., Greene, K., Perine, L., Burt, A., & Hall, P. (2022). Towards a standard for identifying and managing bias in artificial intelligence (NIST Special Publication 1270). National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.SP.1270>
- Sheehan, M. (2023). China's AI Regulations and How They Get Made. <https://carnegieendowment.org/research/2023/07/chinas-ai-regulations-and-how-they-get-made?lang=en>
- Sheikh, H., Prins, C., & Schrijvers, E. (2023). Artificial Intelligence: Definition and Background. In C. Prins & F. Brom (Eds.), *Mission AI: The New System Technology* (pp. 15–41). Springer. https://doi.org/https://doi.org/10.1007/978-3-031-21448-6_2
- Shneiderman, B. (2022). Human-centered AI: Reliable, safe & trustworthy. Oxford University Press.
- Solihati, K. D., Wati, N. K., Herawati, A. R., & Sari, N. H. (2024). Post-TikTok Shop Closure: Will MSMEs Really Rise in Indonesia? *Jurnal Ilmu Sosial*, 23(2), 1–13. <https://doi.org/10.14710/jis.23.2.2024.1-13>
- Sutrisno, E. (2024). Indonesia Salip Vietnam dan Malaysia, Jadi Tujuan Investasi Digital Terbesar Kedua di ASEAN. <https://indonesia.go.id/kategori/editorial/8490/indonesia-salip-vietnam-dan-malaysia-jadi-tujuan-investasi-digital-terbesar-kedua-di-asean?lang=1>
- Table.Briefings. (2024). Tsinghua camp: backbone of China's AI startup scene. <https://table.media/en/china/sinolytics-radar/the-tsinghua-camp-the-backbone-of-chinas-ai-start-up-scene/>
- Tanna, M., & Brown, A. (2023). AI View. <https://www.simmons-simmons.com/en/publications/clptmxzfn0020u2ew5ig6152z/ai-view-28-november-2023>
- Thomson, E. (2024). These charts show which businesses are driving the EU economy. <https://www.weforum.org/stories/2024/01/chart-drive-eu-economy-small-business-sme/>
- Thormundsson, B. (2024). Artificial intelligence (AI) market size worldwide from 2020 to 2030. <https://www.statista.com/forecasts/1474143/global-ai-market-size#:~:text=AI%20market%20size%20worldwide%20from%202020%2D2030&text=The%20market%20for%20artificial%20intelligence,billion%20U.S.%20dollars%20in%202030>
- Timcke, S., & Schültken, T. (2023, October). *Towards a more comprehensive AI ethics: How global south perspectives can enrich current approaches to AI governance*. Research ICT Africa. https://researchictafrica.net/wp/wp-content/uploads/2023/12/AI-ethics-in-the-global-south_Oct23.pdf
- United Nations Educational, Scientific and Cultural Organization. (2024). *Readiness assessment methodology: Artificial intelligence - Indonesia*. <https://www.unesco.org/ethics-ai/en/indonesia>
- United Nations Educational, Scientific, and Cultural Organization (UNESCO). (2022). *Recommendation on the Ethics of Artificial Intelligence*. https://unesdoc.unesco.org/in/documentViewer.xhtml?v=2.1.196&id=p::usmarcdef_0000381137&file=/in/rest/ann

- otationSVC/DownloadWatermarkedAttachment/attach_import_e86c4b5d-5af9-4e15-be60-82f1a09956fd%3F_%3D381137eng.pdf&locale=en&multi=true&ark=/ark:/48223/pf0000381137/PDF/381137eng.pdf#1517_21_EN_SHS_int.indd%3A.8918%3A4
- Vermeulen, B., Kesselhut, J., Pyka, A., & Saviotti, P. P. (2018). The impact of automation on employment: Just the usual structural change? *Sustainability*, 10(5), 1661. <https://doi.org/10.3390/su10051661>
- Vijayakumar, A. (2024). *AI ethics from the Global South: Addressing the social costs of artificial intelligence* (Discussion Paper No. 296). Research and Information System for Developing Countries (RIS). <https://www.ris.org.in/sites/default/files/Publication/DP-296-Anupama-Vijayakumar.pdf>
- Vivarelli, M. (2014, March). Innovation, employment, and skills in advanced and developing countries: A survey of economic literature. *Journal of Economic Issues*. <https://doi.org/10.2753/JEI0021-3624480106>
- von Garrel, J., & Jahn, C. (2023). Design Framework for the Implementation of AI-based (Service) Business Models for Small and Medium-sized Manufacturing Enterprises. *Journal of the Knowledge Economy*, 14, 3551–3569. <https://doi.org/10.1007/s13132-022-01003-z>
- White & Case. (2024). AI Watch: Global Regulatory Tracker - United States <https://www.whitecase.com/insight-our-thinking/ai-watch-global-regulatory-tracker-united-states>
- Williams, B. K. (2025). Winning the defining contest: The us-china artificial intelligence race. *The Washington Quarterly*, 48(2), 153–171. <https://doi.org/10.1080/0163660X.2025.2517501>
- Wolff, J. G. (2014). Big data and the SP theory of intelligence. *IEEE Access*, 2, 301–315. <https://doi.org/10.1109/ACCESS.2014.2315297>
- Xu, W., Dainoff, M. J., Ge, L., & Gao, Z. (2023). Transitioning to human interaction with AI systems: New challenges and opportunities for HCI professionals to enable human-centered AI. *International Journal of Human–Computer Interaction*, 39(3), 494–518. <https://doi.org/10.1080/10447318.2022.2041900>
- Yang, C. H. (2022). How Artificial Intelligence Technology Affects Productivity and Employment: Firm-level Evidence from Taiwan. *Research Policy*, 51(6). <https://doi.org/10.1016/j.respol.2022.104536>
- Zhang, H., & Wu, Z. (2022). How high-tech enterprise recognition affects enterprise innovation behavior: empirical evidence from fuzzy regression discontinuity. *Bus. Manag. J.*, 44, 63–81.
- Zhang, L. (2023). China: Provisions on Deep Synthesis Technology Enter into Effect. <https://www.loc.gov/item/global-legal-monitor/2023-04-25/china-provisions-on-deep-synthesis-technology-enter-into-effect/>