

Artificial intelligence governance and social policy divergence: a comparative political economy perspective on global AI regulation

International
Journal of
Sociology and
Social Policy

Deepak Kumar

La Trobe University, Melbourne, Australia, and

Chinmaya Kumar Sahu

Faculty of Management Studies,

CMS Business School, Jain (Deemed-to-be University), Bengaluru, India

Received 10 February 2026

Revised 21 February 2026

Accepted 4 March 2026

Abstract

Purpose – To explain why artificial intelligence (AI) governance produces systematically divergent social policy outcomes across regions despite widespread convergence around ethical standards for AI. The study examines how political–economic institutions shape the allocation of social risks, regulatory authority, labour integration and ethics institutionalisation in AI governance, thereby driving these differences.

Design/methodology/approach – This study uses a comparative qualitative policy analysis based on 24 key AI policy documents published between 2018 and 2025 across the European Union, United States, China, and Indo-Pacific economies. Guided by theory, the documents are systematically coded across four institutional dimensions and converted into simple indices to compare governance approaches across the regions.

Findings – The findings show clear and systematic differences in how regions govern AI. Five distinct governance models emerge: rights-based (EU), market-driven (US), state-centric (China), hybrid (Australia–Japan–Singapore) and developmental (India). Although many regions use similar ethical language, substantial differences persist in risk allocation, regulatory enforcement, welfare integration and social protection. These differences reflect the historically embedded political–economic institutions shaping each regime.

Originality/value – The paper reframes AI governance as a form of social policy shaped by political and economic institutions. It develops a new, evidence-based typology of AI governance models and shows that differences across countries are driven by institutional structures and not by ethical principles alone.

Keywords Technology governance, Artificial intelligence governance, AI regulation, AI social policy, AI social risks, Comparative political economy

Paper type Research article

1. Introduction

Artificial intelligence (AI) is being widely used in economic production, public administration, and everyday social life. This has far-reaching implications for the labour markets, welfare systems, and public accountability (Wirtz *et al.*, 2019). AI systems are now common in processes such as hiring and workplace management, welfare eligibility assessment, content moderation, policing, and surveillance. This is reshaping how social risks are distributed and governed and raises pressing social policy questions concerning job displacement, algorithmic exclusion, due process, and the erosion or reconfiguration of social protection mechanisms (Acemoglu and Restrepo, 2020; Autor *et al.*, 2020).

Governments across the globe have rolled out policies to cater to this shift. However, the social policies governing AI differ markedly across regions. Over the past decade, international organisations have sought to promote shared ethical principles for trustworthy



© Deepak Kumar and Chinmaya Kumar Sahu. Published by Emerald Publishing Limited. This article is published under the Creative Commons Attribution (CC BY 4.0) license. Anyone may reproduce, distribute, translate and create derivative works of this article (for both commercial and non-commercial purposes), subject to full attribution to the original publication and authors. The full terms of this license may be seen at [Link to the terms of the CC BY 4.0 licence](https://creativecommons.org/licenses/by/4.0/).

International Journal of Sociology and
Social Policy
Emerald Publishing Limited
e-ISSN: 1758-6720
p-ISSN: 0144-333X
DOI 10.1108/IJSSP-02-2026-0129

and responsible AI with an emphasis on fairness, transparency, accountability, and human oversight (Floridi *et al.*, 2018; Jobin *et al.*, 2019; UNESCO, 2021). Despite these efforts, the responses to AI's social consequences remain largely divergent with Jurisdictions adopting different approaches to labour protection, welfare integration, regulatory enforcement, and ethical accountability.

Existing research on AI governance has largely focused on ethical frameworks, technical standards, and sector-specific regulation (Nikolinakos, 2023; Ferrell *et al.*, 2024; Kilian *et al.*, 2025; Currie *et al.*, 2025). The literature has focused on articulating normative principles and identifying governance gaps, with limited attention to the political-economic institutions through which such principles are translated into binding social policy. As a result, AI governance debates often assume that regulatory convergence is both desirable and attainable, framing policy divergence as a problem of regulatory lag, coordination failure, or insufficient ethical commitment (Mittelstadt, 2019; van Maanen, 2022).

Against this background, this study aims to explain why AI social governance policies diverge across regions and the consequences of this divergence for global AI governance. To address this question, the paper examines three related issues. First, it analyses how AI governance regimes differ in their allocation of social risk, regulatory authority, labour and welfare integration, and ethics institutionalisation. Second, it assesses whether distinct governance regimes emerge empirically from policy instruments rather than imposing *a priori*. Third, based on the findings of these two, it synthesises the divergence in AI social governance and its implications.

To address these questions, the study undertakes a systematic comparative analysis of AI social policy across four governance regimes: the European Union, the United States, China, and the Indo-Pacific group of countries (Australia, Japan, Singapore, and India). Drawing on comparative political economy (Baccaro and Pontusson, 2016; Hall and Soskice, 2001) and the concept of regulatory capitalism (Levi-Faur, 2005), the paper conceptualises AI governance as a form of social regulation embedded within nationally specific arrangements of state authority, market coordination, and social protection. From this perspective, AI does not disrupt social policy institutions as much as it rearticulates existing logics governing how societies allocate risk, regulate markets, and mediate technological change.

By answering these questions, the study makes three contributions to sociology and social policy scholarship. Firstly, it advances theory by explaining AI social policy divergence as institutionally embedded rather than normatively inconsistent through the lens of comparative political economy and its role in the governance of emerging technologies. Second, it provides an empirically grounded typology of AI social governance regimes based on systematic cross-regional policy comparison. Third, it offers critical insights into the limits of global AI harmonisation, demonstrating the insufficiencies of ethics-based coordination and soft-law frameworks to address uneven social protection and regulatory incompatibility. By reframing AI governance as a problem of social policy and institutional power rather than ethics or technology alone, the study contributes to broader debates on welfare states, risk governance, and the sociology of regulation.

2. Literature review and theoretical framework

2.1 AI governance and the limits of ethics-centric approaches

The rapid diffusion of AI has generated a substantial body of literature on its governance. But this discourse largely centres on ethical principles, accountability mechanisms, and technical safeguards. The early literary signs signalled the need for “trustworthy” or “human-centred” AI, articulating principles such as fairness, transparency, explainability, and human oversight (Floridi *et al.*, 2018; Jobin *et al.*, 2019). These principles were further consolidated through global soft law instruments issued by international organisations like OECD, UNESCO, and the World Economic Forum, aimed at fostering responsible and socially beneficial AI (UNESCO, 2021; World Economic Forum, 2022).

While this ethics-oriented literature and guidelines have been influential in shaping global discourse, scholars increasingly question their capacity to deliver substantive social protection. The argument that these ethical principles are frequently aspirational, weakly institutionalised, and poorly integrated into binding regulatory or social policy frameworks are not uncommon (Mittelstadt, 2019). Some scholars term this process as ethics washing, in which symbolic adherence to ethical norms substitutes for enforceable obligations or redistributive interventions (van Maanen, 2022). From a social policy perspective, ethics-centric governance often depoliticises AI-related harms by framing them as technical or moral challenges rather than as distributive and institutional problems.

The overemphasis on ethical convergence also obscures persistent cross-national variation in regulatory authority, labour protection, and welfare integration. While jurisdictions broadly endorse similar ethical language, they differ markedly in how and whether such principles are translated into enforceable social policy. This makes global AI governance exhibit normative alignment without institutional harmonisation, raising questions about the adequacy of ethics-based coordination as a response to AI's social consequences.

2.2 AI and social risk

Parallel to governance debates, a growing literature examines the social consequences of AI adoption in labour markets, inequality, and welfare systems (Feher and Veres, 2023). Research on automation and algorithmic management highlights the uneven distribution of benefits and risks, with AI contributing to job displacement, task polarisation, and intensified worker surveillance (Acemoglu and Restrepo, 2020; Autor *et al.*, 2020). In the public sector, AI systems shape access to welfare benefits, migration, decisions, and criminal justice outcomes. This raises concerns about bias, accountability, and procedural fairness of these systems (Wirtz *et al.*, 2019).

The recent labour sociology and digital capitalism research brings out why AI-related social risks cannot be treated as a downstream impact problem. Platform firms increasingly operate as governance actors that shape social and economic life through rules, rankings, and data infrastructures. This creates layered governance relationships among platforms, users, advertisers, and states rather than a simple regulator–regulated relationship (Gorwa, 2019). This platformisation is not just a technological shift but a reorganisation of capitalism in which financialisation, digitalisation, and privatisation converge into digital proprietary markets owned and regulated by transnational platform companies. This creates new discontinuities in social regulation as platforms compete with and partially displace public institutions (Törnberg, 2023).

Zhang *et al.* (2025) show that algorithmic management is an agentic boss and socio-technical process, with persistent tensions between efficiency and fairness, control and autonomy, and transparency and opacity. This also implies that welfare protections and labour institutions are central to risk allocation and redress. Extending this line of argument, Pepple and Muthuthantrige (2026) conceptualise AI as a socio-technical actor that can intensify displacement and structural precarity, eroding relational, moral, and emotional dimensions of work in ways that can resemble modern slavery dynamics. This underscores the need for inclusive policy frameworks and human-centred governance rather than narrow compliance or ex post mitigation.

Despite these emerging concerns, AI social policy is often treated as a downstream response to technological change rather than as a constitutive element of AI governance. The existing literature in the field mostly focuses on impact assessment and measuring effects on employment or inequality (Vicsek, 2021). There has been limited enquiry into how welfare regimes, labour institutions, and regulatory capacity shape governance choices. The questions of who bears AI-related social risks, how responsibility is allocated, and which institutions provide protection or redress remain under-theorised.

Comparative studies have documented regional differences in AI regulation, particularly between the European Union, the United States, and China. The EU is frequently characterised

as adopting a precautionary approach grounded in fundamental rights (Veale and Borgesius, 2021), US as prioritising innovation, relying on voluntary standards and sectoral oversight (Gasser and Almeida, 2017) and China associated with a model combining rapid technological development with strong central control (Cheng and Zeng, 2023). In the Indo-Pacific, the literature points to a mode of governance with selective elements from dominant models while responding to domestic institutional constraints (Keith, 2024). Much of this comparative literature remains descriptive and regulation-centric, focussing on legal instruments rather than the political–economic foundations of social policy divergence. These existing studies identify variation without fully explaining why such divergence persists despite increasing global coordination.

2.3 Comparative political economy and regulatory capitalism

To address this explanatory gap, this study draws on the theory of comparative political economy (CPE). The theory emphasises that economic governance is embedded in historically rooted institutional arrangements governing state–market relations, welfare provision, and labour regulation (Hall and Soskice, 2001; Baccaro and Pontusson, 2016). Hence, the policy responses to technological change are not neutral or technocratic but rather reflect enduring political compromises over risk distribution, social protection, and economic coordination.

CPE theory challenges the assumption that regulatory convergence is a natural or desirable endpoint. Instead, it suggests that divergence is a predictable outcome when global technologies interact with different welfare regimes and regulatory traditions. Societies differ systematically in how they handle social risks which plays a role in how emerging technologies are governed. Applied to AI, this implies that decisions about labour protection, welfare integration, and accountability are conditioned by pre-existing institutional logic and not just the technological characteristics.

While comparative political economy explains why institutional divergence persists, the concept of regulatory capitalism (RC) explains how governance is operationalised. RC is a mode of governance in which states rely on indirect regulation through standards, audits, certifications, and delegated authority instead of direct command-and-control intervention (Levi-Faur, 2005). This form of governance is particularly common in technologically complex and rapidly evolving domains such as artificial intelligence.

AI governance manifests RC in several respects including delegation of regulatory functions to private standard-setting bodies, professional associations, and technology firms. Compliance is enforced through mechanisms such as risk assessments, impact evaluations, and reporting requirements instead of substantive prohibitions. Ethical principles are operationalised through guidelines and best practices rather than embedded within redistributive social policy instruments.

This governance architecture can be used to explain why global AI initiatives produce normative convergence without regulatory harmonisation. Shared ethical language masks big differences in enforcement capacity, regulatory authority, and the role of welfare institutions across jurisdictions. Regulatory capitalism thus complements comparative political economy by highlighting the instruments through which institutional logics are enacted.

2.4 Integrated framework: explaining AI social policy divergence

Bringing these strands together, this study conceptualises AI governance as a form of social policy embedded within political–economic institutions and mediated through regulatory capitalism. AI social policy divergence is understood not as regulatory failure but as an institutional outcome reflecting how different societies govern social risk. The integrated framework focuses on four analytically distinct but interrelated dimensions:

- (1) **Social risk allocation (D1)**– whether AI-related risks are borne primarily by individuals, firms, or the state.

- (2) **Regulatory authority and enforcement (D2)**– whether governance relies on binding public law, delegated regulation, or voluntary standards.
- (3) **Labour and welfare integration (D3)**– the extent to which AI governance is embedded within labour law, welfare systems, and social protection mechanisms.
- (4) **Ethics institutionalisation (D4)**– whether ethical principles are aspirational, procedural, or legally enforceable.

Social risk allocation draws on welfare-state theory and the literature on social risk governance, which examines how risks are distributed between individuals, employers, and the state (Esping-Andersen, 1990). Regulatory authority and enforcement are informed by regulatory capitalism and hard–soft law scholarship, which differentiates binding legal instruments, institutional capacity, and coercive enforcement from voluntary or standards-based approaches (Abbott and Snidal, 2000; Levi-Faur, 2005). Labour and welfare integration reflects research on varieties of capitalism and labour market institutions, which emphasises the embedding of markets within social protection and employment regimes (Hall and Soskice, 2001). Ethics institutionalisation builds on sociological accounts of audit and procedural governance, distinguishing between aspirational norms and enforceable or proceduralized accountability mechanisms (Mittelstadt, 2019).

These dimensions translate abstract political–economic concepts into observable policy characteristics, providing the foundation for systematic comparison across regions. Using these dimensions, Figure 1 conceptualises AI social policy divergence as a multi-stage institutional process. Political–economic institutions (including welfare regimes, labour relations, and state–market configurations) shape the choice of regulatory instruments associated with regulatory capitalism (such as binding law, procedural oversight, or voluntary standards). These instruments, in turn, produce distinct AI social policy configurations that determine how social risks associated with AI, such as job displacement, algorithmic exclusion, and surveillance, are distributed across individuals, markets, and the state.

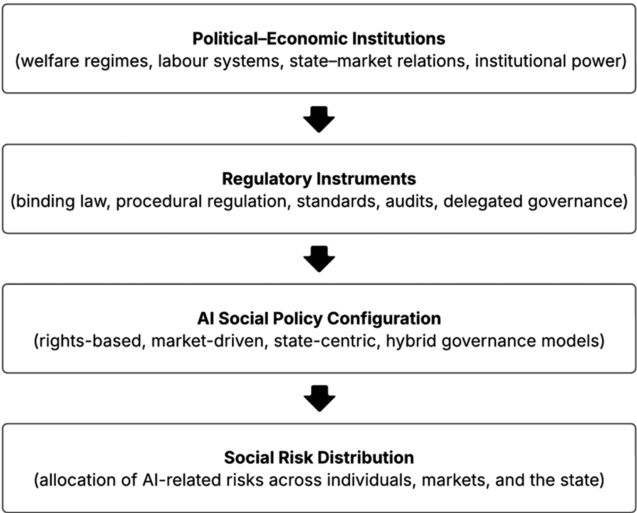


Figure 1. Political economy, regulatory capitalism, and ai social policy outcomes. **Source:** Authors’ own work

3. Research design and methodology

3.1 Research design and analytical orientation

This study adopts a comparative qualitative policy analysis to examine why AI social policies diverge across regions. Comparative policy analysis is particularly well-suited to sociological inquiry into governance variation, as it enables systematic comparison of institutional logics rather than isolated regulatory instruments (Dunlop and Radaelli, 2022). Rather than evaluating the effectiveness of individual AI regulations, the analysis focuses on governance design, asking how AI-related social risks are framed, allocated, and governed within different welfare and regulatory regimes.

3.2 Case Selection and scope

We focus on four governance regimes representing distinct configurations of political economy, welfare provision, and regulatory authority:

- (1) **European Union**, representing a rights-based and precautionary governance model
- (2) **United States**, representing a market-driven and innovation-oriented governance model
- (3) **China**, representing a state-centric and strategically coordinated governance model
- (4) **Indo-Pacific economies** (Australia, Japan, Singapore and India), representing adaptive and mixed regulatory arrangements

These cases are selected using theoretical sampling, not geographic convenience. Each regime reflects a distinct institutional configuration relevant to comparative social policy analysis, allowing meaningful comparison of how AI-related social risks are governed. The Indo-Pacific cases are grouped analytically to capture shared hybrid characteristics while retaining sensitivity to internal variation.

3.3 Data sources and corpus construction

The analysis draws on an original corpus of 24 authoritative policy documents (Appendix 2) issued between 2018 and 2025 across the four chosen governance regimes, during which AI governance became a formal policy priority. Documents were selected based on their relevance to AI governance and social policy, and include:

- (1) National and regional AI strategies
- (2) AI-specific legislation and regulatory proposals
- (3) Ethical guidelines and governance frameworks
- (4) Labour, welfare, and social protection policies explicitly addressing AI, automation, or algorithmic decision-making.

3.4 Coding framework and index construction

The integrated conceptual model described in Section 2.4 serves as the basis for the theory-driven policy coding framework used in this study. The framework translates four core analytical dimensions (social risk allocation, regulatory authority and enforcement, labour and welfare integration, and ethics institutionalisation) into 20 observable indicators. After the indicator set was first constructed using theory, a pilot coding phase was conducted to evaluate construct coverage. This identified two more governance tools that were not entirely represented by the initial indicators but were conceptually consistent with the analytical framework.

To increase measurement completeness, we developed and added two new indicators: Indicator_Monitoring_or_KPIs under regulatory authority and enforcement and Indicator_Contestability_or_Remedy under ethics institutionalisation. [Appendix 1](#) contains a list of indicators along with their dimensions and meanings. We finalised and locked the codebook prior to full coding to ensure consistency. This follows the standard practice in qualitative comparative policy analysis, where validity is improved without sacrificing replicability through limited refinement following pilot testing ([Fischer and Maggetti, 2017](#)).

In order to document the presence or absence of particular indicators, we begin coding the chosen policy documents at the indicator level using binary measures (0/1). NVivo qualitative data analysis software was used for coding in order to guarantee methodical text handling, procedure transparency, and reproducible outcomes ([O’Kane et al., 2019](#)). Semantic interpretation was used to manually code documents to predetermined indicator nodes after they were imported into NVivo. This ensured that coding decisions remained theory-led rather than software-driven and helped us prioritise substantive meaning over frequency counts.

Binary indicators were then summed within each dimension to construct four additive indices for every document. Each document therefore yields a four-dimensional governance profile (D1–D4), representing its configuration of AI social policy instruments. These document-level profiles constitute the atomic unit of analysis. For cross-regional comparison, document-level scores were aggregated to the regime level. Regime profiles are reported using mean values to capture central tendencies and minimum–maximum ranges to reflect internal heterogeneity. This aggregation strategy enables systematic comparison of governance orientation and intensity across regions while preserving within-regime variation.

3.5 Analytical strategy

The analysis proceeded in three stages. In order to determine the dominant AI social governance orientations within each regime across the four dimensions of social risk allocation, regulatory authority and enforcement, labour and welfare integration, and ethics institutionalisation, a within-regime analysis first looked at document-level index scores. Second, by combining document-level indices to regime-level profiles using mean values and minimum–maximum ranges, a cross-regime comparison evaluated systematic divergence. This phase identifies regional variations in governance intensity and orientation. Third, we evaluated the empirical recoverability of the suggested AI governance regimes from policy tools. In order to identify the institutional divergence, this required looking for patterns of similarity among documents and locating clustered configurations of governance tools.

3.6 Validity and reliability

The study uses several strategies to enhance analytical robustness. First, triangulation using multiple policy documents within each regime reduces reliance on single legislative instruments and mitigates document-level bias. Second, the coding framework and index-construction procedures theory-based with clear documentation, which ensures transparency and replicability. The use of NVivo further enhances reliability by maintaining a consistent coding structure and preserving all coded excerpts. Third, regime typologies are empirically validated through document-level similarity analysis. This reduces the risk of purely normative classification.

4. Findings

4.1 Regime-level AI social governance profiles

This section presents the regime-level results of AI social governance across the chosen regimes. Document-level indices were aggregated to the regime level to capture both central tendencies and internal heterogeneity within each governance model. [Table 1](#) reports the mean index scores and the minimum–maximum ranges for each regime. The results indicate

Table 1. Regime-level profiles of AI social governance

Dimensions/ Regimes	D1_ mean	D1_ min– max	D2_ mean	D2_ min– max	D3_ mean	D3_ min– max	D4_ mean	D4_ min– max
EU	4.4	3–5	3.4	0–6	3.2	2–5	4.2	2–6
US	2.6	1–3	3.6	2–5	1.8	1–3	4.6	3–5
China	2.4	2–3	5.6	4–6	1.4	1–3	5	2–6
AU–JP–SG	2.29	2–4	2	0–4	1.43	0–5	3.29	2–5
IN	4	4–4	2	1–3	2	2–2	2.5	2–3

Source(s): Authors’ own work

substantial cross-regional variation across all four dimensions, confirming that AI social governance is not organised around a single global model. At the same time, the reported ranges reveal meaningful within-regime variation, particularly where soft-law strategies coexist with binding regulatory instruments. This combination of central tendency and dispersion provides a robust empirical basis for comparing governance orientations while avoiding overgeneralisation.

The heat map in Figure 2 visualises regime-level AI social governance profiles across the four dimensions, highlighting systematic differences in governance orientation and intensity. The European Union consistently scores highly on social risk allocation (D1) and ethics institutionalisation (D4), reflecting a rights-based approach that emphasises collective responsibility, enforceable safeguards, and procedural accountability. Labour and welfare

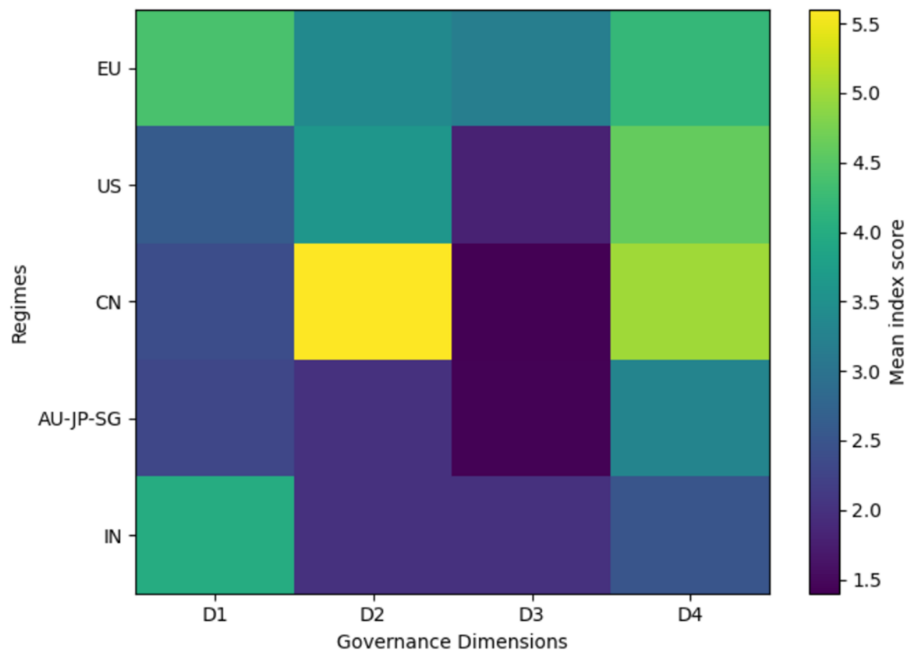


Figure 2. Regime-level AI social governance profiles across four dimensions. *Darker shading indicates a stronger institutional presence of the corresponding governance dimension. **Source:** Authors’ own work

integration (D3) is also comparatively strong, indicating closer embedding of AI governance within existing social protection frameworks.

The United States displays a contrasting profile, characterised by lower levels of social risk allocation and labour integration, alongside comparatively strong ethics-related signalling. This pattern reflects a market-driven governance model that prioritises innovation and organisational flexibility, relying more heavily on voluntary standards and ethical guidance than on welfare-linked regulation.

China's profile is distinguished by exceptionally high scores in regulatory authority and enforcement (D2), combined with relatively low labour and welfare integration. This configuration reflects a state-centric governance model in which AI-related social risks are managed primarily through centralised control and enforcement rather than redistribution or rights-based remedies.

The AU–JP–SG grouping occupies an intermediate position across most dimensions, combining moderate regulatory enforcement with limited labour integration and comparatively stronger ethical proceduralizing. This pattern is consistent with a pragmatic hybrid governance model that balances regulatory caution with policy flexibility. India, by contrast, exhibits relatively strong social risk allocation and labour-related scores, alongside weaker enforcement and institutionalisation of ethics, consistent with a developmental, capacity-building orientation focused on workforce adaptation rather than regulatory constraint.

These regime-level profiles demonstrate that AI social policy divergence is multidimensional and structured, rather than the result of isolated regulatory choices. These patterns provide the empirical foundation for examining the specific policy instruments driving divergence and for assessing whether the observed regimes emerge empirically from governance configurations.

4.2 Indicator-level drivers of regime divergence

We disaggregate the observed regime-level patterns by examining the prevalence of AI governance indicators across regions. Figure 3 highlights how distinct combinations of policy tools underpin divergent AI social governance regimes. Across all regimes, state responsibility is widely acknowledged as a core principle of AI governance, reflecting a shared baseline recognition of public authority in managing AI-related social risks. However, substantial divergence emerges in the use of welfare safeguards, which are nearly universal in the European Union and India but appear far less frequently in the United States and the AU–JP–SG group. This pattern indicates that while most jurisdictions accept a role for the state, only some institutionalise this responsibility through explicit social protection mechanisms, reinforcing cross-regional variation in how AI risks are redistributed.

The prevalence of binding law, sanctions or penalties, and risk classification reveals the clearest regulatory divergence. China shows consistently high prevalence across these enforcement-oriented instruments, underscoring a state-centric governance model grounded in hierarchical authority and compliance. The European Union also relies heavily on binding legal instruments, though with more variation across documents, reflecting a mixed ecosystem of hard law and soft guidelines. The United States and AU–JP–SG regimes display markedly lower reliance on binding law and sanctions, signalling a preference for standards-based or voluntary approaches over coercive enforcement.

There are pronounced differences in the integration of AI governance with labour institutions. Labour law references are common in the EU but largely absent in the United States and in the AU–JP–SG regimes, indicating that labour protection is not systematically embedded in AI governance in these contexts. Reskilling and training measures are prevalent across most regimes, particularly in India and the AU–JP–SG group, suggesting a shared emphasis on workforce adaptation. Importantly, this emphasis on skills development often substitutes for stronger labour protections, contributing to uneven social outcomes across regions.

All regimes exhibit a high prevalence of ethical principles, indicating global convergence in normative discourse. However, this convergence does not guarantee uniformity in

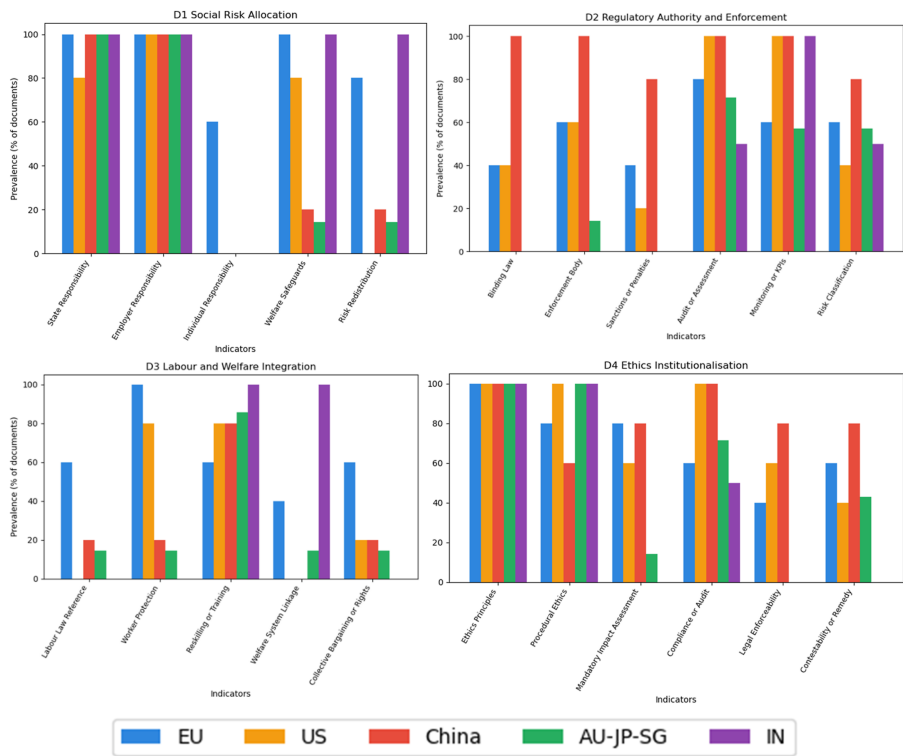


Figure 3. Prevalence of AI governance indicators across regimes. **Source:** Authors’ own work

approach. The European Union shows substantially higher prevalence of mandatory impact assessment and contestability or remedy, signalling a move from aspirational ethics to enforceable procedural mechanisms. In contrast, the United States and AU–JP–SG regimes rely more heavily on voluntary ethical principles, though they do not consistently embed them in legally enforceable processes. This divergence illustrates how similar ethical language can mask fundamentally different governance architectures.

These indicator-level patterns demonstrate that any single policy instrument does not drive regime divergence, but rather the systematic bundling of governance tools into distinct institutional configurations. While jurisdictions converge visibly around state responsibility and ethical principles, they diverge sharply in how these commitments are operationalised through binding law, welfare protections, labour integration, and enforcement capacity. In other words, what appears as shared normative language masks substantive differences in regulatory architecture and social risk allocation. To assess whether these patterned combinations translate into empirically coherent regimes rather than descriptive generalisations, the next section moves from instrument prevalence to document-level similarity analysis, examining whether distinct governance clusters emerge directly from the underlying indicator configurations.

4.3 Empirical validation of AI governance regimes

This section evaluates whether the AI social governance regimes can be empirically recovered from policy instruments. Using document-level indicator profiles, we compare policies based on their full configuration of governance tools to assess whether regimes emerge as coherent

clusters. We visualise these relationships using a document-level multidimensional scaling (MDS) similarity map. Multidimensional scaling is a dimensionality-reduction technique that represents pairwise similarities between documents as distances in a low-dimensional space, such that documents with more similar indicator profiles appear closer together.

The resulting map in Figure 4 shows distinct regime-based clustering. Policy documents associated with the European Union, the United States, China, the AU–JP–SG grouping, and India form largely distinct clusters, indicating that a coherent configuration of governance instruments characterises each regime. Importantly, this clustering emerges without imposing regime labels during the similarity calculation, demonstrating that regime distinctions are grounded in underlying policy features rather than post hoc classification. The observed structure suggests that AI social governance regimes are differentiated not by isolated regulatory choices, but by systematic combinations of social risk allocation mechanisms, enforcement tools, labour linkages, and ethics institutionalisation.

While regime-level clustering is pronounced, limited overlap between clusters is also visible. These boundary cases are most commonly associated with soft-law strategies, ethical guidelines, and coordination plans, which tend to occupy intermediate positions between more tightly clustered binding instruments. Such overlap does not indicate misclassification. Rather, it reflects governance hybridity at the document level, where shared ethical language or procedural commitments coexist with divergent enforcement capacities, labour protections, and welfare linkages. In particular, documents within the AU–JP–SG grouping frequently occupy bridging positions, consistent with a pragmatic governance orientation that combines elements of multiple models without fully converging on any single regime.

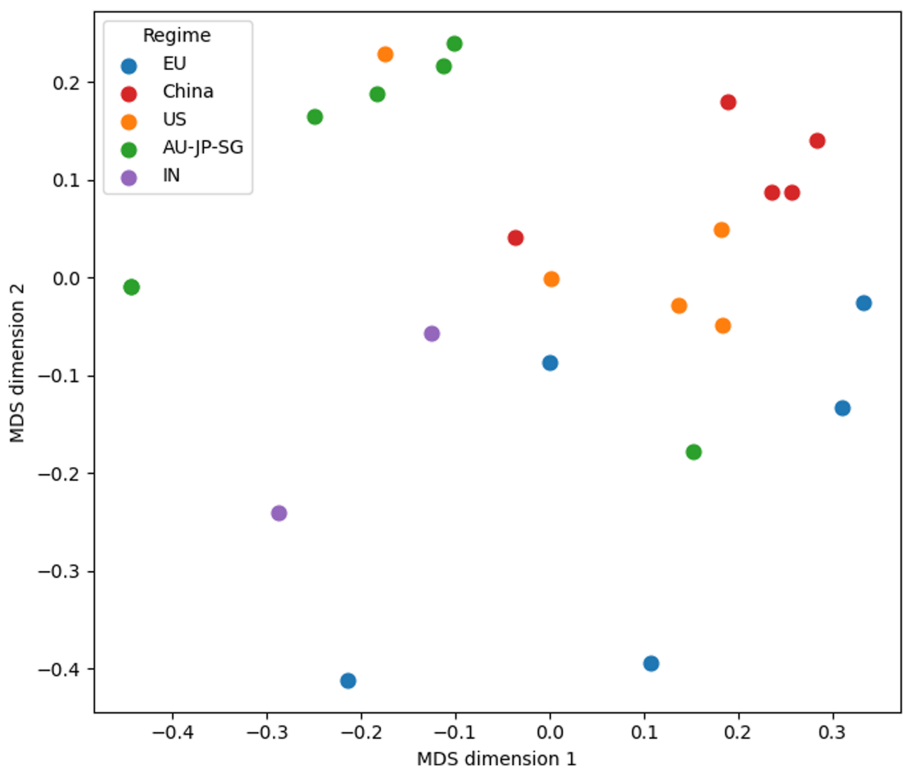


Figure 4. Document-level similarity of AI social governance policies. **Source:** Authors' own work

Notably, Indian policy documents do not collapse into the AU–JP–SG cluster despite partial overlap on capacity-building and reskilling instruments. This supports treating India as a distinct developmental regime rather than a subtype of hybrid governance. This document-level similarity analysis provides strong empirical support for the AI governance regime typology. Regimes differ systematically in their governance configurations, and these differences are instrument-driven rather than rhetorical. At the same time, the presence of boundary cases highlights that fragmentation in AI social governance is structured and patterned rather than absolute.

The empirical results show that AI social governance divergence is systematic, patterned, and institutionally grounded. The observed clustering of policy instruments across regimes demonstrates that divergence is not a by-product of administrative variation or draughting style. Rather, it reflects coherent governance architectures structured around distinct configurations of risk allocation, enforcement authority, labour integration, and ethics institutionalisation. The coexistence of global convergence in ethical language with deep divergence in binding regulation, welfare safeguards, and labour protections further suggests that shared normative discourse does not produce institutional uniformity. The question, therefore, is no longer whether AI governance differs across regions, but why do these patterned differences persist despite growing international coordination. The following section addresses this by situating the empirical configurations within the underlying political–economic institutional logics.

5. Discussions

5.1 *Why does AI social governance diverge?*

A central research question of this study is why AI social governance diverges systematically across regions despite growing international coordination and shared ethical discourse. The comparative results indicate that divergence is not primarily technological or administrative in origin. Instead, it reflects underlying institutional logics embedded within political–economic systems, which shape how societies allocate social risk, exercise regulatory authority, integrate labour and welfare considerations, and institutionalise ethics in the governance of artificial intelligence. AI social policy divergence thus emerges as an institutional outcome rather than a coordination failure.

Across the four dimensions analysed, such as social risk allocation, regulatory authority and enforcement, labour–welfare integration, and ethics institutionalisation, policy instruments cluster into coherent configurations rather than appearing randomly distributed. These configurations correspond to distinct institutional logics that predate AI itself. AI governance thus becomes layered onto existing welfare regimes and regulatory traditions, reproducing broader patterns of state–market relations. Divergence, therefore, emerges as a structural outcome of institutional diversity rather than a temporary coordination failure. To illustrate these patterns empirically, [Table 2](#) summarises regime-level orientations by translating additive index scores into low, medium, and high categories, which are tercile-based interpretations of regime-level mean index scores. Where scores lie close to adjacent tercile thresholds, intermediate labels (Low–Medium or Medium–High) are used to reflect hybrid positioning.

While these classifications condense quantitative patterns, they do not yet reveal the broader governance architectures that underpin them. To clarify how these empirical differences translate into distinct governance logics, [Figure 5](#) synthesises the comparative findings into a typological map of AI social policy models. This visualisation highlights that AI governance clusters into qualitatively different institutional configurations such as rights-based, market-driven, state-centric, hybrid, and developmental, each reflecting historically embedded approaches to allocating risk, responsibility, and authority.

[Figure 5](#) functions as an interpretive bridge between the quantitative profiles reported in [Table 2](#) and the institutional explanations developed below, guiding the reader from

Table 2. AI social governance profiles categories

Region	Social risk allocation	Regulatory authority and enforcement	Labour and welfare integration	Ethics institutionalisation
EU	High	Medium	High	High
CN	Low–Medium	High	Low	High
US	Low–Medium	Medium	Low	Medium–High
AU–JP–SG	Low–Medium	Low–Medium	Low	Medium
IN	High	Low–Medium	Medium–High	Low –Medium

Source(s): Authors’ own work

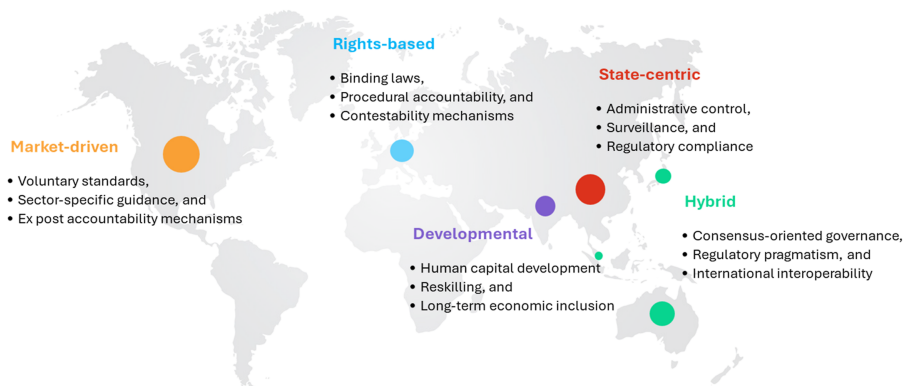


Figure 5. AI social policy typology and features. **Source:** Authors’ own work

measurement to mechanism. The following subsections, therefore, examine how each governance model emerges from its underlying political–economic foundations and how these foundations shape distinctive approaches to regulating the social consequences of artificial intelligence.

5.1.1 Rights-based regulatory capitalism (European Union). The European Union’s AI social governance profile reflects a model of regulatory capitalism rooted in strong welfare institutions, precautionary legal measures, and clear commitment to enforceable fundamental rights. Across the four dimensions examined, EU policy documents consistently frame AI not merely as an innovative technology but as a source of potential social risk requiring collective protection and public oversight. In social risk allocation, EU policy language repeatedly positions the state and public law as primary guarantors of protection. The AI Act explicitly recognises that AI “*may generate risks and cause harm to public interests and fundamental rights*” and therefore requires “*a consistent and high level of protection of public interests as regards health, safety and fundamental rights*” (European Parliament and Council, 2024).

Rather than individualising risk, responsibility is institutionalised through legal safeguards. This is operationalised through binding regulation and enforcement. The Regulation establishes “*harmonised rules applicable to the placing on the market, the putting into service and the use of high-risk AI systems*” and clarifies that existing protections relating to “*fundamental rights, employment, and protection of workers remain fully applicable*” (European Parliament and Council, 2024).

The EU AI governance also displays strong labour and welfare integration. AI systems used in employment contexts are treated as high-risk and subject to additional safeguards.

They need to undergo “*risk management, human oversight, and documentation requirements, with penalties for non-compliance*” (European Economic and Social Committee, 2025). The policy also calls for “*safe and fair working conditions*” and “*stronger involvement of workers and their representatives*” (European Economic and Social Committee, 2025).

This reflects moving beyond aspirational ethics toward institutionalised and enforceable norms. The Ethics Guidelines ground AI governance in “*fundamental rights enshrined in the EU Treaties, the EU Charter and international human rights law*”. It describes many of these rights as “*legally enforceable*” (High-Level Expert Group on AI, 2019). Ethics is thus framed not as soft guidance alone but as a complement to law. Procedural mechanisms such as impact assessments, oversight bodies, and accountability structures translate normative principles into practice.

These features illustrate coherent institutional logic. Not just relying on post hoc correction or market self-regulation, the EU emphasises ex ante precaution, legally binding safeguards, worker protection, and enforceable rights. AI is treated as a socio-technical system with distributive consequences that must be governed through collective legal institutions. This aligns closely with the tradition of European regulatory capitalism and welfare-state governance and explains the regime’s consistently high scores across all four dimensions and its classification as a rights-based model.

5.1.2 Market-driven, standards-based governance (United States). The United States (US) exhibits a AI governance logic that prioritises innovation, organisational flexibility, and decentralised responsibility. Its policies frame AI not as a domain requiring comprehensive social protection but as an area for risk management, voluntary coordination, and private-sector leadership. Federal policy instruments consistently rely on standards, frameworks, and guidance instead of binding statutory obligations. This is visible in the explicitly soft-law character of key documents. The Blueprint for an AI Bill of Rights states that it is “*non-binding and does not constitute U.S. government policy*” and “*does not constitute binding guidance for the public or Federal agencies and therefore does not require compliance*” (White House Office of Science and Technology Policy, 2022).

Similarly, the NIST AI Risk Management Framework is described as “*intended for voluntary use*” to support organisations in incorporating trustworthiness considerations (National Institute of Standards and Technology, 2023). The Executive Order repeatedly instructs agencies to develop “*guidance regarding the existing tools and practices*” and to solicit input from stakeholders (White House, 2023). This aims to establish and reinforce a consultative and standards-based approach. The ethical principles are prominent and widely articulated their policies, but their enforceability in law remains limited. This configuration reflects a market-driven, standards-based institutional logic in which AI governance is designed to minimise regulatory friction while preserving innovation capacity. In political-economic terms, responsibility for managing AI risks is decentralised to firms, professional communities, and individual adaptability rather than embedded within collective welfare institutions.

5.1.3 State-centric developmental control (China). Across China’s AI governance policies, artificial intelligence is treated as a strategic national technology requiring centralised oversight and strong enforcement capacity. In its policies, AI is framed as an object of state planning, security management, and administrative control. The State Council’s national strategy positions AI as central to national development and security, calling for the country to “*seize the major strategic opportunity for the development of artificial intelligence*” and to “*promote the overall elevation of national competitiveness and leapfrog development*” (State Council, 2017).

This strategic orientation is operationalised through binding regulatory authority and compliance mechanisms. The regulations emphasise direct state supervision and enforceable obligations. For example, the Algorithmic Recommendation Provisions require providers to “*strengthen algorithm security management*” and “*accept regulatory oversight by competent authorities*” (Cyberspace Administration of China, 2022). The Generative AI Measures state

that providers must ensure services are “safe, controllable, and in accordance with laws and regulations” (Cyberspace Administration of China, 2023). This reflects administrative command and regulatory compliance as key governance aims in its policies.

The labour and welfare protections are only indirectly embedded. Social risks are addressed primarily through macro-level planning and stability objectives. Individual rights or contestability mechanisms are not the primary concerns. The development plan links AI to “social governance” and “public security operations,” highlighting surveillance, coordination, and state capacity as core tools of risk (State Council, 2017). These features point to an institutional logic where AI is governed as a strategic infrastructure of state power, with risks managed through centralised planning, administrative enforcement, and regulatory control. Rather than embedding AI within welfare or rights-based safeguards, governance prioritises stability, security, and national development objectives. This aligns with a state-centric model of regulatory capitalism and explains the regime’s strong enforcement capacity combined with comparatively weaker labour–welfare and contestability provisions.

5.1.4 Pragmatic hybrid governance (Australia–Japan–Singapore). The Australia–Japan–Singapore grouping’s AI governance is characterised by regulatory flexibility, incrementalism, and soft institutional steering. AI governance is presented as a process of workforce adaptation, organisational accountability, and responsible innovation under all of these regimes. Human capital development, risk management procedures, and ethics guidelines are constantly emphasised in policy documents as the main instruments for controlling social risks associated with AI.

Japan’s approach ties societal inclusion and adaptation to AI governance. It says that new technologies should be implemented “in parallel with the reform of social systems” and that broad acceptance depends on enabling citizens to “realize concrete benefits through the introduction of AI” (Cabinet Office, 2022). It stresses the importance of “human resources development” and literacy to support AI diffusion, positioning skills formation as a central policy response to technological change rather than legal constraint (Cabinet Office, 2019). Governance is therefore embedded in coordination and capacity-building rather than ex ante statutory control.

Australia’s strategy exhibits a similar approach. It supports organisational and voluntary methods for governing AI. The framework encourages organisations to implement internal governance procedures and suggests that AI systems be implemented through “proportionate risk management” and “human oversight” (Department of Industry, Science and Resources, 2025). As a result, accountability is transferred to businesses and organisations, with the state serving mainly as a facilitator and standard-setter.

These policies demonstrate a hybrid institutional logic in which social risks are recognised but primarily addressed by means of reskilling programs, procedural safeguards, and ethics frameworks. This strategy falls somewhere between rights-based regulation and state-centric control, reflecting a practical hybrid governance.

5.1.5 Developmental capacity-building governance (India). The goals of India’s AI governance are inclusive economic transformation, labour absorption, and capacity building. AI is presented as an opportunity for job creation, skill enhancement, and general growth in all of its policy documents. The national strategy positions the state as a coordinator of social mobility and human capital formation by connecting AI to inclusive development goals.

This is visible in the emphasis on workforce transition and large-scale reskilling. India’s AI strategy notes that “Re-skilling of the current workforce will require integration with relevant existing skilling initiatives, building of new platforms that can enable improved learning, and novel methods of allowing large scale employment generation through promotion of AI” (NITI Aayog, 2018). Policy tools concentrate on growing training ecosystems, certification pathways, and industry partnerships rather than placing onerous compliance requirements on businesses.

Similarly, learning infrastructures are framed as the central governance lever. The document highlights the “Creation of open platforms for learning” that will “play an

instrumental role in large scale dissemination of requisite skills to some major sections of the employed workforce” (NITI Aayog, 2018). Institutional and financial incentives, such as “*co-funded models between the government and companies*” to reduce the opportunity cost of training, are suggested to encourage businesses and employees to engage in reskilling (NITI Aayog, 2018).

These directions demonstrate a governance framework that prioritises developmental coordination. AI-related social risks are mainly understood as employment and skill-related issues that need to be resolved through training, education, and labour market inclusion. AI governance thus operates through capacity-building, workforce preparedness, and long-term economic inclusion, distinguishing India’s model from both rights-based regulatory systems and market-driven approaches and explaining its classification as a developmental or capacity-building regime.

These regime-specific patterns demonstrate that AI social policy divergence is best understood as the product of institutional logics governing risk, responsibility, and redistribution and not as inconsistent or incomplete regulation. Each regime represents a coherent response to AI’s social implications, shaped by historically embedded political–economic arrangements. This institutional diversity provides the foundation for the structured differences observed in global AI governance.

5.2 Consequences of divergence in AI social governance

The institutional divergence identified in [section 5.1](#) has important consequences for how AI-related social risks are governed globally. The AI governance across the regions do not converge toward a unified regulatory model, but is characterised by structured differences, reflecting persistent differences in institutional capacity, welfare regimes, and regulatory philosophies. This shapes the regulatory design of AI governance including distribution of social protection, compliance burdens, and the prospects for international coordination.

5.2.1 Divergence as institutional diversity rather than governance failure. Disparities in AI social governance shouldn’t be taken as proof of incoherent policies or failed regulations. Rather, it represents the institutional diversity of how societies structure risk, redistribution, and responsibility. Each regime adopts a logical arrangement of governance tools in line with its underlying political-economic logic, as the empirical analysis shows. Therefore, the reason for the divergence is not a lack of coordination, but rather the fact that AI governance is a part of larger, heterogeneous systems of labour regulation, welfare provision, and state-market relations.

According to this interpretation, the divergence is not a transitional anomaly but rather a stable equilibrium. Attempts to enforce uniform governance models run the risk of ignoring the social contracts and institutional restraints that influence national policy decisions. AI social policy divergence thus mirrors long-standing patterns observed in areas such as labour law, social protection, and environmental regulation.

5.2.2 Regulatory interoperability and compliance challenges. In a globalised AI ecosystem, institutional diversity poses serious obstacles to regulatory interoperability, even though it may be internally coherent. Businesses that operate in multiple jurisdictions encounter disjointed compliance environments concerning risk assessment, labour protection, accountability systems, and ethical standards. Binding regulatory regimes that emphasise enforceability and contestability coexist with standards-based systems that prioritise flexibility and innovation, complicating cross-border alignment. These challenges are particularly acute for AI systems deployed transnationally, where regulatory obligations may conflict or be unevenly enforceable. This raises practical questions about mutual recognition, equivalence, and regulatory coordination, especially in the absence of supranational enforcement mechanisms.

5.2.3 Uneven social protection and distributive consequences. Divergence in AI social governance also produces uneven patterns of social protection across regimes. Where AI

governance is closely integrated with welfare systems and labour law, social risks associated with automation, surveillance, and algorithmic decision-making are more likely to be collectively managed. In contrast, regimes that rely on market-based adjustment or firm-level responsibility place a greater burden on individuals and workers to absorb AI-related risks. These differences have distributive implications, shaping who bears the costs of technological change and who benefits from AI-driven productivity gains. This divergence, therefore, has consequences not only for regulatory coherence but also for social inequality, labour precarity, and access to remedies in cases of harm.

5.2.4 Limits of ethics-based global coordination. The analysis also highlights the limits of ethics-based coordination as a mechanism for overcoming divergence in Social AI policies. Ethical principles such as fairness, transparency, and accountability are widely shared across regimes and travel easily through international forums, standards bodies, and multilateral initiatives. However, the institutionalisation of ethics varies substantially, particularly with respect to enforceability, remedies, and the linkage to welfare. As a result, convergence at the level of ethical rhetoric does not translate into convergence in governance outcomes. Ethics frameworks may facilitate dialogue and norm diffusion, but they cannot substitute for institutional arrangements governing enforcement, redistribution, and social protection.

These findings demonstrate that AI social policy divergence is structured, patterned, and institutionally grounded. Regimes diverge systematically in how they allocate social risk, exercise regulatory authority, integrate labour and welfare considerations, and institutionalise ethics. These differences are driven by enduring political-economic logics rather than by technological necessity or regulatory inertia. AI governance divergence should therefore be understood not as a transitional phase on the path toward harmonisation, but as a durable feature of global AI governance. Recognising this structured diversity is essential for developing realistic approaches to international coordination, regulatory interoperability, and social protection in the age of artificial intelligence.

6. Conclusion

This study set out to explain why AI social governance policies diverge across regions and what are the associated consequences. By conceptualising AI governance as a form of social policy embedded within political-economic institutions, the study makes three core contributions. First, it advances an integrated analytical framework that links AI governance to social risk allocation, regulatory authority and enforcement, labour and welfare integration, and ethics institutionalisation. This framework translates abstract political-economic concepts into observable policy characteristics, enabling systematic comparison across regions. Second, it provides empirical evidence of structured regime divergence, demonstrating that AI governance regimes are not rhetorical labels but empirically recoverable configurations of governance instruments. Third, it shows that global AI governance is characterised by durable and patterned divergence, rooted in institutional diversity rather than regulatory or coordination failure.

The findings carry important implications for international standard-setting bodies and international coordination efforts in AI governance. They suggest that expectations of comprehensive regulatory harmonisation are unlikely to be realised. Because AI governance is embedded within welfare systems, labour regimes, and state-market relations, convergence on technical standards or ethical principles does not translate into convergence on social protection, enforcement, or redistribution mechanisms. Rather than pursuing uniform global regulation, international coordination efforts may be more effective if they focus on regulatory interoperability. This includes mutual recognition of compliance mechanisms, equivalence assessments for risk management procedures, and alignment on minimum safeguards where feasible. Such approaches respect institutional diversity while reducing friction for cross-border AI deployment.

The analysis also highlights the limits of ethics-based coordination as a standalone governance strategy. While ethical principles provide a common normative language, they

remain weak substitutes for institutional arrangements governing enforcement, contestability, and welfare protection. Global AI governance initiatives that rely primarily on ethical convergence risk masking underlying disparities in social protection and accountability. The results underscore the importance of explicitly integrating labour and welfare considerations into international AI governance discussions. Differences in this domain have distributive consequences, shaping who bears the social risks of AI adoption. Without attention to labour transitions, reskilling, and social safety nets, global coordination efforts may reinforce existing inequalities rather than mitigate them.

The study has limitations typical of qualitative comparative policy research. It focused on governance design and institutional logic rather than policy implementation or effectiveness. While this constrains claims about real-world outcomes, it is appropriate given the study's explanatory aims and provides a foundation for future mixed-methods research examining implementation dynamics and social impacts. Future research could extend this framework through mixed methods designs that examine implementation practices, enforcement capacity, and social impacts across regimes. Further work could also expand the comparative scope beyond the regions examined here or apply the framework to sector-specific AI governance domains such as healthcare, finance, or public administration. Longitudinal analysis would be particularly valuable in assessing whether regime configurations evolve or remain institutionally stable.

AI social policy divergence is not a transient anomaly but a structural feature of global AI governance. Recognising this fragmentation as institutionally grounded rather than normatively deficient is essential for developing realistic, inclusive, and effective governance strategies. By situating AI governance within broader systems of social protection and political-economic organisation, this study contributes to a more grounded understanding of how societies govern the social implications of artificial intelligence and why they do so differently.

Table A1. Dimensions and indicators used for policy coding

Dimension	Indicator	Operational definition
D1 – Social Risk Allocation	State responsibility	Assigns primary responsibility to government/ public authorities for mitigating AI risks
	Employer obligations	Places compliance or risk-management duties on firms, deployers, or organisations
	Individual responsibility	Frames users or workers as bearing primary responsibility for managing risks
	Welfare safeguards	Provides social protection, safety nets, or compensatory support measures
	Risk redistribution	Redistributes AI-related risks collectively rather than concentrating them on individuals
D2 – Regulatory Authority and Enforcement	Binding law	Establishes legally binding rules or statutory obligations
	Enforcement body	Designates authorities with oversight or supervisory powers
	Sanctions/penalties	Provides fines, penalties, or corrective measures for non-compliance
	Audit/assessment	Requires audits, testing, conformity assessment, or evaluation procedures
	Monitoring/KPIs	Mandates ongoing reporting, monitoring, or performance indicators
	Risk classification	Categorises AI systems by risk level to determine regulatory treatment
D3 – Labour and Welfare Integration	Labour law reference	Links AI governance to employment or Labour legislation
	Worker protection	Protects workers from displacement, discrimination, or algorithmic harm
	Reskilling/training	Supports education, retraining, or workforce transition measures
	Welfare system linkage	Connects AI governance to social security or welfare institutions
	Collective rights	Recognises trade unions, social dialogue, or collective bargaining rights
D4 – Ethics Institutionalisation	Ethics principles	Articulates normative values (e.g. fairness, transparency, human oversight)
	Procedural ethics	Embeds ethics in organisational processes or governance structures
	Mandatory impact assessment	Requires formal AI or algorithmic impact assessments
	Compliance/audit	Establishes mechanisms to verify adherence to standards or rules
	Legal enforceability	Makes ethical or governance requirements legally binding
	Contestability/remedy	Provides rights or procedures to challenge automated decisions

Table A2. Policy documents details

Doc_ID	Region	Document_Title	Year	Official_URL	Notes
EU01	EU	Artificial Intelligence Act (Regulation (EU) 2024/1689)	2024	https://www.aiact-info.eu/full-text-and-pdf-download/	Core binding AI regulation
EU02	EU	Coordinated Plan on Artificial Intelligence (2021 Update)	2021	https://digital-strategy.ec.europa.eu/en/library/coordinated-plan-artificial-intelligence-2021-review	EU AI policy coordination
EU03	EU	General Data Protection Regulation (GDPR)	2016	https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:32016R0679	Algorithmic decision-making rights
EU04	EU	Pro-worker AI: levers for harnessing the potential and mitigating the risks of AI in connection with employment and labour market policies	2020	Opinion of the European Economic and Social Committee – Pro-worker AI: levers for harnessing the potential and mitigating the risks of AI in connection with employment and labour market policies (own-initiative opinion)	AI and worker protection
EU05	EU	Ethics Guidelines for Trustworthy AI	2019	Ethics Guidelines for AI	Core EU ethics framework
US01	US	Blueprint for an AI Bill of Rights	2022	https://marketingstoragerags.blob.core.windows.net/webfiles/Blueprint-for-an-AI-Bill-of-Rights.pdf	Rights-oriented AI principles
US02	US	National AI Initiative Act	2020	https://www.congress.gov/bill/116th-congress/house-bill/6216	National AI coordination
US03	US	Executive Order on Safe, Secure, and Trustworthy AI	2023	https://www.acc.com/sites/default/files/2024-08/2023.10.30-Executive-Order-on-the-Safe-Secure-and-Trustworthy-Development-and-Use-of-Artificial-Intelligence-_-The-White-House.pdf	Federal AI governance
US04	US	AI Risk Management Framework	2023	AI Risk Management Framework NIST	Procedural AI governance
US05	US	Guidance on AI and Algorithmic Accountability	2021	Using-Artificial-Intelligence-and-Algorithms-_-Federal-Trade-Commission.pdf	Consumer protection focus
CN01	China	New Generation Artificial Intelligence Development Plan	2017	https://digichina.stanford.edu/work/full-translation-chinas-new-generation-artificial-intelligence-development-plan-2017/	National AI strategy
CN02	China	Algorithmic Recommendation Management Provisions	2022	https://www.chinalawtranslate.com/en/algorithms/	Platform algorithm control
CN03	China	Data Security Law	2021	http://www.npc.gov.cn/englishnpc/c2759/c23934/202112/t20211209_385109.html	Data governance

(continued)

Table A2. Continued

Doc_ ID	Region	Document_Title	Year	Official_URL	Notes
CN04	China	Personal Information Protection Law	2021	https://digichina.stanford.edu/work/translation-personal-information-protection-law-of-the-peoples-republic-of-china-effective-nov-1-2021/	Privacy and algorithm use
CN05	China	Interim Measures for the Management of Generative AI Services	2023	https://cyrilla.org/api/files/1728288405530qs2ud3o27ys.pdf	Generative AI control
AU01	Indo-Pacific	Australia's Artificial Intelligence Ethics Framework	2019	https://www.industry.gov.au/publications/australias-artificial-intelligence-ethics-framework	Establishes voluntary ethical principles for AI development and deployment; no legal enforcement mechanisms
AU02	Indo-Pacific	National AI Plan	2025	https://www.industry.gov.au/publications/national-ai-plan	Focuses on coordination, industry enablement, and innovation support rather than social regulation
AU03	Indo-Pacific	Automation and the Future of Work	2020	The Future of Work	Addresses automation impacts through reskilling and workforce transition, not labour protection rights
JP01	Indo-Pacific	AI Strategy 2019/2022	2019/2022	AI Strategy 2022	Articulates state-led coordination for AI innovation with limited social policy intervention
JP02	Indo-Pacific	Social Principles of Human-Centric AI	2019	humancentricai.pdf	Sets high-level ethical principles; non-binding and non-enforceable
SG01	Indo-Pacific	Model AI Governance Framework	2019	https://www.pdpc.gov.sg/Help-and-Resources/2020/01/Model-AI-Governance-Framework	Voluntary governance guidance emphasising risk management and organisational accountability

(continued)

Table A2. Continued

Doc_ ID	Region	Document_Title	Year	Official_URL	Notes
SG02	Indo-Pacific	AI Verify Framework	2022	https://www.imda.gov.sg/programme-listing/ai-verify	Introduces voluntary AI auditing and assurance tools to build market trust
IN01	Indo-Pacific	National Strategy for Artificial Intelligence #AIForAll	2018	National Strategy for Artificial Intelligence	Frames AI as a developmental and capacity-building tool; limited social protection emphasis
IN02	Indo-Pacific	Responsible AI for All	2021	Approach Document for India: Part 1 - Principles for Responsible AI - OECD.AI	Articulates ethical aspirations and governance principles without binding obligations

References

- Abbott, K.W. and Snidal, D. (2000), "Hard and soft law in international governance", *International Organization*, Vol. 54 No. 3, pp. 421-456, doi: [10.1162/002081800551280](https://doi.org/10.1162/002081800551280).
- Acemoglu, D. and Restrepo, P. (2020), "The wrong kind of AI? Artificial intelligence and the future of labour demand", *Cambridge Journal of Regions, Economy and Society*, Vol. 13 No. 1, pp. 25-35, doi: [10.1093/cjres/rsz022](https://doi.org/10.1093/cjres/rsz022).
- Autor, D., Mindell, D. and Reynolds, E. (2020), "The work of the future: building better jobs in an age of intelligent machines", *MIT Task Force on the Work of the Future*, 2020-Final-Report4.pdf.
- Baccaro, L. and Pontusson, J. (2016), "Rethinking comparative political economy: the growth model perspective", *Politics and Society*, Vol. 44 No. 2, pp. 175-207, doi: [10.1177/0032329216638053](https://doi.org/10.1177/0032329216638053).
- Cabinet Office, Government of Japan (2019), "Social Principles of Human-Centric AI. humancentricai.pdf"
- Cabinet Office, Government of Japan (2022), *AI Strategy 2022*, Government of Japan. AI Strategy 2022, Tokyo.
- Cheng, J. and Zeng, J. (2023), "Shaping AI's future? China in global AI governance", *Journal of Contemporary China*, Vol. 32 No. 143, pp. 794-810, doi: [10.1080/10670564.2022.2107391](https://doi.org/10.1080/10670564.2022.2107391).
- Currie, W.L., Leimeister, J.M., Schlagwein, D. and Willcocks, L. (2025), "Rethinking technology regulation in the age of AI risks", *Journal of Information Technology*, Vol. 40 No. 3, pp. 236-245, doi: [10.1177/02683962251378815](https://doi.org/10.1177/02683962251378815).
- Cyberspace Administration of China (2022), *Provisions on the Administration of Algorithmic Recommendation Services*, available at: <https://www.chinalawtranslate.com/en/algorithms/>
- Cyberspace Administration of China (2023), "Interim measures for the management of generative artificial intelligence services", available at: <https://cyrilla.org/api/files/1728288405530qs2ud3o27ys.pdf>
- Department of Industry, Science and Resources (2025), "National AI plan", Australian Government, available at: <https://www.industry.gov.au/publications/national-ai-plan>
- Dunlop, C.A. and Radaelli, C.M. (2022), "Policy learning in comparative policy analysis", *Journal of Comparative Policy Analysis: Research and Practice*, Vol. 24 No. 1, pp. 51-72, doi: [10.1080/13876988.2020.1762077](https://doi.org/10.1080/13876988.2020.1762077).
- Esping-Andersen, G. (1990), *The Three Worlds of Welfare Capitalism*, Polity Press/Princeton University Press, Cambridge, Princeton, NJ.
- European Economic and Social Committee (2025), "Pro-worker AI: levers for harnessing the potential and mitigating the risks of AI in connection with employment and labour market policies (Own-initiative opinion C/2025/1185)", *Official Journal of the European Union*, C series, available at: <https://eur-lex.europa.eu>.
- European Parliament, and Council of the European Union (2024), "Regulation laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)", *Official Journal of the European Union*, available at: <https://www.aiact-info.eu/full-text-and-pdf-download/>
- Feher, K. and Veres, Z. (2023), "Trends, risks and potential cooperations in the AI development market: expectations of the Hungarian investors and developers in an international context", *International Journal of Sociology and Social Policy*, Vol. 43 Nos 1-2, pp. 107-125, doi: [10.1108/IJSSP-08-2021-0205](https://doi.org/10.1108/IJSSP-08-2021-0205).
- Ferrell, O.C., Harrison, D.E., Ferrell, L.K., Hair, J.F. and Hochstein, B.W. (2024), "A theoretical framework to guide AI ethical decision making", *AMS Review*, Vol. 14 No. 1, pp. 53-67, doi: [10.1007/s13162-024-00275-9](https://doi.org/10.1007/s13162-024-00275-9).
- Fischer, M. and Maggetti, M. (2017), "Qualitative comparative analysis and the study of policy processes", *Journal of Comparative Policy Analysis: Research and Practice*, Vol. 19 No. 4, pp. 345-361, doi: [10.1080/13876988.2016.1149281](https://doi.org/10.1080/13876988.2016.1149281).
- Floridi, L., Cows, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P. and Vayena, E. (2018), "AI4People—an ethical

- framework for a good AI society”, *Minds and Machines*, Vol. 28 No. 4, pp. 689-707, doi: [10.1007/s11023-018-9482-5](https://doi.org/10.1007/s11023-018-9482-5).
- Gasser, U. and Almeida, V.A.F. (2017), “A layered model for AI governance”, *IEEE Internet Computing*, Vol. 21 No. 6, pp. 58-62, doi: [10.1109/MIC.2017.4180835](https://doi.org/10.1109/MIC.2017.4180835).
- Gorwa, R. (2019), “What is platform governance?”, *Information, Communication and Society*, Vol. 22 No. 6, pp. 854-871, doi: [10.1080/1369118X.2019.1573914](https://doi.org/10.1080/1369118X.2019.1573914).
- Hall, P.A. and Soskice, D. (Eds) (2001), *Varieties of Capitalism: The Institutional Foundations of Comparative Advantage*, Oxford University Press, doi: [10.1093/0199247757.001.0001](https://doi.org/10.1093/0199247757.001.0001).
- High-Level Expert Group on Artificial Intelligence (2019), *Ethics Guidelines for Trustworthy AI, Ethics Guidelines for AI*, European Commission, Brussels.
- Jobin, A., Ienca, M. and Vayena, E. (2019), “The global landscape of AI ethics guidelines”, *Nature Machine Intelligence*, Vol. 1 No. 9, pp. 389-399, doi: [10.1038/s42256-019-0088-2](https://doi.org/10.1038/s42256-019-0088-2).
- Keith, A.J. (2024), “Governance of artificial intelligence in Southeast Asia”, *Global Policy*, Vol. 15 No. 5, pp. 937-954, doi: [10.1111/1758-5899.13458](https://doi.org/10.1111/1758-5899.13458).
- Kilian, R., Jäck, L. and Ebel, D. (2025), “European AI standards – technical standardisation and implementation challenges under the EU AI act”, *European Journal of Risk Regulation*, Vol. 16 No. 3, pp. 1038-1062, doi: [10.1017/err.2025.10032](https://doi.org/10.1017/err.2025.10032).
- Levi-Faur, D. (2005), “The global diffusion of regulatory capitalism”, *The Annals of the American Academy of Political and Social Science*, Vol. 598 No. 1, pp. 12-32, doi: [10.1177/0002716204272371](https://doi.org/10.1177/0002716204272371).
- Mittelstadt, B. (2019), “Principles alone cannot guarantee ethical AI”, *Nature Machine Intelligence*, Vol. 1 No. 11, pp. 501-507, doi: [10.1038/s42256-019-0114-4](https://doi.org/10.1038/s42256-019-0114-4).
- National Institute of Standards and Technology (2023), *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*, U.S. Department of Commerce. AI Risk Management Framework | NIST, Gaithersburg, MD.
- Nikolinakos, N.T. (2023), “Ethical principles for trustworthy AI”, in *EU Policy and Legal Framework for Artificial Intelligence, Robotics and Related Technologies—The AI Act (Law, Governance and Technology Series*, Springer, Vol. 53, pp. 101-166, doi: [10.1007/978-3-031-27953-9_3](https://doi.org/10.1007/978-3-031-27953-9_3).
- NITI Aayog (2018), *National Strategy for Artificial Intelligence: #AIforAll*. Government of India, National Strategy for Artificial Intelligence, New Delhi.
- O’Kane, P., Smith, A. and Lerman, M.P. (2019), “Building transparency and trustworthiness in inductive research through computer-aided qualitative data analysis software”, *Organizational Research Methods*, Vol. 24 No. 1, pp. 104-139, (Original work published 2021), doi: [10.1177/1094428119865016](https://doi.org/10.1177/1094428119865016).
- Pepple, D. and Muthuthantrige, N. (2026), “Artificial intelligence, innovation and the new architecture of exploitation: towards reconfiguring humanness in the age of algorithmic labour”, *Journal of Innovation and Knowledge*, Vol. 11, 100878, doi: [10.1016/j.jik.2025.100878](https://doi.org/10.1016/j.jik.2025.100878).
- State Council (2017), “New generation artificial intelligence development plan”, available at: <https://digichina.stanford.edu/work/full-translation-chinas-new-generation-artificial-intelligence-development-plan-2017/>
- Törnberg, P. (2023), “How platforms govern: social regulation in digital capitalism”, *Big Data and Society*, Vol. 10 No. 1, doi: [10.1177/20539517231153808](https://doi.org/10.1177/20539517231153808).
- UNESCO (2021), *Recommendation on the Ethics of Artificial Intelligence*, United Nations Educational, Scientific and Cultural Organization, available at: <https://unesdoc.unesco.org/ark:/48223/pf0000381137>
- van Maanen, G. (2022), “AI ethics, ethics washing, and the need to politicize data ethics”, *Digital Society*, Vol. 1 No. 2, 9, Article 9 doi: [10.1007/s44206-022-00013-3](https://doi.org/10.1007/s44206-022-00013-3).
- Veale, M. and Borgesius, F.Z. (2021), “Demystifying the draft EU artificial intelligence act”, *Computer Law Review International*, Vol. 22 No. 4, pp. 97-112, doi: [10.9785/crl-2021-220402](https://doi.org/10.9785/crl-2021-220402).

-
- Vicsek, L. (2021), "Artificial intelligence and the future of work – lessons from the sociology of expectations", *International Journal of Sociology and Social Policy*, Vol. 41 Nos 7-8, pp. 842-861, doi: [10.1108/IJSSP-05-2020-0174](https://doi.org/10.1108/IJSSP-05-2020-0174).
- White, H. (2023), "Executive order on the safe, secure, and trustworthy development and use of artificial intelligence", available at: https://www.acc.com/sites/default/files/2024-08/2023.10.30-Executive-Order-on-the-Safe-Secure-and-Trustworthy-Development-and-Use-of-Artificial-Intelligence_-The-White-House.pdf
- White House Office of Science and Technology Policy (2022), "Blueprint for an AI Bill of Rights: making automated systems work for the American people", available at: <https://marketingstoragerags.blob.core.windows.net/webfiles/Blueprint-for-an-AI-Bill-of-Rights.pdf>
- Wirtz, B.W., Weyerer, J.C. and Geyer, C. (2019), "Artificial intelligence and the public sector—applications and challenges", *International Journal of Public Administration*, Vol. 42 No. 7, pp. 596-615, doi: [10.1080/01900692.2018.1498103](https://doi.org/10.1080/01900692.2018.1498103).
- World Economic Forum (2022), *Global AI Governance: A Multistakeholder Approach*, World Economic Forum, available at: https://www3.weforum.org/docs/WEF_Global_Technology_Governance.pdf
- Zhang, M.M., Cooke, F.L., Ahlstrom, D. and McNeil, N. (2025), "The rise of algorithmic management and implications for work and organisations", *New Technology, Work and Employment*, Vol. 40 No. 3, pp. 659-671, doi: [10.1111/ntwe.12343](https://doi.org/10.1111/ntwe.12343).

Corresponding author

Deepak Kumar can be contacted at: D.Kumar@latrobe.edu.au