

PUBLIC POLICY AND AI: A PRAGMATIC APPROACH

Lokesha M K

Associate Professor
GFGC KR Pura, Bangalore

ABSTRACT

Artificial Intelligence (AI) is rapidly reshaping the contours of governance, economic organisation, and social life. Policymakers across the world face the dual imperative of harnessing AI's transformative potential while mitigating its attendant risks — ranging from algorithmic discrimination and privacy erosion to systemic economic disruption and national security vulnerabilities. This article argues that neither uncritical permission nor reflexive precautionary prohibition constitutes an adequate policy posture. Instead, drawing on pragmatist political philosophy, comparative policy analysis, and emerging regulatory frameworks in the European Union, United States, United Kingdom, China, and Singapore, it develops a pragmatic framework for AI governance. That framework rests on five pillars: proportionate risk-based regulation, adaptive governance mechanisms, multi-stakeholder deliberation, rights-centred design, and international policy coordination. The article concludes that pragmatic AI governance is not merely an instrumental accommodation to technological complexity, but a principled commitment to evidence-informed, context-sensitive, and democratically accountable policymaking.

Keywords: Artificial Intelligence, public policy, AI governance, regulatory frameworks, pragmatism, adaptive regulation, risk-based approach, democratic accountability

1. INTRODUCTION

Few technological developments have posed as complex a challenge to public policy as Artificial Intelligence. AI is not a single, bounded technology but an evolving constellation of techniques — machine learning, large language models, computer vision, reinforcement learning, and beyond — each with distinct capabilities, failure modes, and societal implications. This complexity has confounded policymakers accustomed to regulating technologies with relatively stable characteristics and clearly demarcated boundaries.

The stakes are high. AI systems are being deployed in consequential domains: criminal justice risk assessment, medical diagnosis, credit allocation, hiring and admissions, border control, autonomous weapons, and critical infrastructure management. Errors or manipulations in these systems can cause grievous harm at scale, and that harm disproportionately falls on already-vulnerable populations. At the same time, the economic and social benefits of well-governed AI — in healthcare, education, climate science, and productivity — are enormous and, if foregone through over-regulation, represent a real cost to human welfare.

The central argument of this article is that the binary choice often posed in policy debates — between permissive innovation-first approaches and precautionary restrictions — is a false dilemma. A pragmatic approach, grounded in the philosophical tradition of American pragmatism and recent innovations in adaptive governance theory, offers a more productive path. Pragmatism, as articulated by John Dewey and extended by contemporary scholars such as Cass Sunstein and Elizabeth Anderson, insists that policy should be judged by its consequences rather than its adherence to ideological priors; that it should be experimental

and revisable in light of evidence; and that it should be embedded in robust processes of democratic deliberation.

Thesis: Effective AI governance does not require choosing between innovation and precaution — it requires building institutions and processes that can navigate that tension intelligently, adaptively, and accountably.

The article proceeds as follows. Section 2 surveys the policy landscape, reviewing major national and supranational AI regulatory frameworks. Section 3 develops the theoretical foundations of pragmatic AI governance. Section 4 articulates the five pillars of the proposed framework. Section 5 examines three cross-cutting challenges: international regulatory divergence, the governance of general-purpose AI, and democratic legitimacy. Section 6 offers conclusions and policy recommendations.

2. THE GLOBAL AI POLICY LANDSCAPE

2.1 A Fragmented Regulatory Terrain

The past five years have witnessed an explosion of AI governance initiatives. The Organisation for Economic Co-operation and Development (OECD) adopted its AI Principles in 2019, the first intergovernmental standard for trustworthy AI. The European Union's AI Act, which entered into force in 2024, represents the world's most comprehensive binding AI regulatory framework. The United States has pursued a more decentralised approach, combining executive orders, sector-specific agency guidance, and voluntary frameworks — most notably the NIST AI Risk Management Framework (2023). China has enacted a series of targeted regulations addressing algorithmic recommendation systems, deepfakes, and generative AI, embedded within a broader national AI strategy. The United Kingdom, post-Brexit, has sought a 'pro-innovation' regulatory posture through a principles-based, sector-led model. Singapore has developed the Model AI Governance Framework and accompanying AI Verify toolkit, positioning itself as a hub for responsible AI.

This proliferation of frameworks is welcome evidence of policy attention, but it has also produced fragmentation. Firms operating across jurisdictions face a compliance patchwork that imposes significant costs and can create regulatory arbitrage incentives. More fundamentally, the absence of coherent international coordination means that powerful AI systems developed under permissive regimes can be deployed globally, undermining the protective purpose of stricter national frameworks.

2.2 Dominant Regulatory Paradigms

Beneath the diversity of national approaches, three broad regulatory paradigms can be identified. The first is a rights-based paradigm, exemplified by the EU AI Act and the Council of Europe Framework Convention on AI, which foregrounds the protection of fundamental rights and democratic values as the primary objective of regulation. The second is a risk-management paradigm, exemplified by the NIST AI RMF, which frames AI governance primarily as an enterprise risk management challenge, emphasising voluntary best practices and industry self-governance. The third is an industrial policy paradigm, most evident in China's approach, which instrumentalises AI governance to serve national technological leadership and social stability objectives.

Each paradigm captures important insights, and each has significant blind spots. The rights-based paradigm can become procedurally rigid and innovation-hostile if operationalised through categorical prohibitions untethered to empirical assessment of actual harm. The risk-management paradigm can become industry-captured and accountability-weak if it relies too

heavily on voluntary compliance. The industrial policy paradigm risks subordinating individual rights and global governance cooperation to state interests. A pragmatic framework must draw on the strengths of all three while remedying their respective deficiencies.

3. THEORETICAL FOUNDATIONS: THE CASE FOR PRAGMATISM

3.1 Pragmatism as a Policy Philosophy

Pragmatism, as a philosophical tradition, is frequently misunderstood as mere opportunism or unprincipled flexibility. In fact, classical pragmatism — associated with Charles Sanders Peirce, William James, and above all John Dewey — is a rigorous epistemological and ethical stance. For Dewey, inquiry is always purposive: we think in order to act, and the test of ideas is their consequences in experience. Applied to public policy, this means that regulatory choices should be evaluated by what they actually achieve — in terms of social welfare, rights protection, and democratic values — not by their formal consistency with ideological templates.

Several features of pragmatist thought are particularly salient for AI governance. First, pragmatism's anti-dogmatism counsels scepticism toward both the techno-optimist view that AI markets will self-correct without governance and the techno-pessimist view that AI poses existential threats requiring drastic precautionary intervention. Second, pragmatism's experimentalism suggests that AI governance should be designed for learning: policies should generate feedback, be open to revision, and be embedded in institutions capable of updating their approaches in light of evidence. Third, pragmatism's democratic commitments insist that governance processes must be inclusive and accountable, not merely technocratically competent.

3.2 Adaptive Governance Theory

Pragmatist philosophy finds institutional expression in the theory of adaptive governance, developed in the environmental governance literature and increasingly applied to emerging technologies. Adaptive governance recognises that complex sociotechnical systems cannot be governed through fixed rules derived in advance of experience. Instead, governance frameworks must incorporate monitoring, learning, and adjustment as constitutive features rather than afterthoughts.

For AI specifically, adaptive governance suggests that regulatory frameworks should be structured around ongoing assessment of actual impacts — including cumulative and systemic effects not visible in case-by-case compliance review — and should include mechanisms for updating requirements as AI capabilities and deployment contexts evolve. Sunset clauses, mandatory review cycles, regulatory sandboxes, and living guidelines are among the institutional tools through which adaptive governance can be operationalised.

3.3 The Limits of Precautionary and Permissive Extremes

The pragmatic case is reinforced by the demonstrated failures of both regulatory extremes. Strict precautionary approaches, which require proof of safety before deployment, face the fundamental problem that safety assessments for complex AI systems are context-dependent, dynamic, and subject to significant empirical uncertainty. Applied mechanically, precautionary regulation can suppress beneficial innovation without meaningfully reducing harm — particularly where developers shift to less regulated jurisdictions.

Permissive approaches, which defer regulation until harms have materialised and been judicially adjudicated, face equally serious problems. The scale and speed of AI deployment mean that harms can be widespread and entrenched before remedies are available. The power asymmetries between AI developers and affected individuals make market and litigation-based correction mechanisms structurally inadequate. And the systemic nature of certain AI harms — to democratic deliberation, to labour markets, to epistemic autonomy — may not be remediable through any ex post mechanism.

4. A PRAGMATIC FRAMEWORK: FIVE PILLARS

The pragmatic approach to AI governance proposed here rests on five mutually reinforcing pillars, summarised in the framework table below.

Principle	Description	Policy Example
Incrementalism	Governance is built in iterative, evidence-informed steps rather than through sweeping legislative overhaul.	US NIST AI RMF (2023) — voluntary framework updated via iterative stakeholder consultation.
Proportionality	Regulatory burden is calibrated to actual risk levels; high-risk systems face stringent requirements, low-risk systems minimal intervention.	EU AI Act risk tiers: unacceptable, high, limited, and minimal risk categories.
Multi-Stakeholder Engagement	Policy legitimacy and effectiveness depend on inclusive consultation with industry, civil society, academia and affected communities.	OECD AI Policy Observatory — collaborative platform for global best-practice sharing.
Adaptive Governance	Regulatory frameworks include built-in review cycles and sunset clauses to respond to technological change.	Singapore Model AI Governance Framework — sector-specific guidance with annual review.
Rights-Centred Design	Human dignity, fairness, and non-discrimination are treated as non-negotiable constraints on permissible AI applications.	Council of Europe Framework Convention on AI (2024) — first binding multilateral AI treaty.

4.1 Proportionate Risk-Based Regulation

The foundational principle of pragmatic AI governance is proportionality: the intensity of regulatory intervention should be calibrated to the magnitude and nature of the risks presented by a given AI application in a given context. This principle, well-established in European regulatory law, resists both under-regulation (treating all AI as presumptively benign) and over-regulation (treating all AI as presumptively dangerous).

Risk-based regulation requires, first, a credible and continuously updated taxonomy of AI risks. The EU AI Act's tiered approach — distinguishing unacceptable, high, limited, and minimal risk categories — provides a useful starting point, but its specific categorical determinations should be treated as provisional and subject to empirical revision. Second, it requires robust mechanisms for risk assessment, including mandatory conformity assessments for high-risk systems, post-deployment monitoring requirements, and incident reporting obligations. Third, it requires enforcement capacity: the best-designed regulatory framework is worthless without adequately resourced, technically sophisticated regulatory bodies capable of meaningful oversight.

4.2 Adaptive Governance Mechanisms

Static regulatory frameworks are poorly suited to dynamic technology. AI capabilities are advancing rapidly; deployment contexts are shifting; the social and economic consequences of AI are still being understood. Pragmatic AI governance must therefore build adaptability into its institutional architecture.

Concretely, this means: mandatory periodic review of regulatory classifications and requirements, with clear procedural standards for triggering and conducting reviews; regulatory sandboxes that allow supervised testing of novel AI applications under real-world conditions, generating evidence to inform future regulatory decisions; living guidelines issued by regulatory agencies that can be updated on shorter timescales than primary legislation; and horizon-scanning functions within regulatory bodies to identify emerging AI capabilities and associated governance challenges before they become acute.

4.3 Multi-Stakeholder Deliberation

Pragmatic governance is necessarily deliberative governance. AI systems affect a vast range of stakeholders — users, workers displaced by automation, communities subject to algorithmic decision-making, civil society organisations, academic researchers — whose perspectives and interests must be incorporated into policy formation. Governance processes that exclude affected communities from deliberation may be formally legitimate but are likely to be substantively inadequate, failing to identify harms visible only to those who experience them.

Multi-stakeholder deliberation should be designed for genuine influence, not tokenistic consultation. This requires early and iterative engagement — before regulatory positions are entrenched — and meaningful capacity for less powerful stakeholders to participate effectively. It also requires transparency: the evidence base, assumptions, and reasoning behind regulatory decisions should be publicly accessible and contestable.

4.4 Rights-Centred Design

While pragmatism eschews rigid ideological frameworks, it is fully compatible with — and indeed requires — a robust commitment to fundamental rights as constraints on permissible policy choices. Rights are not merely the output of utilitarian calculation; they represent pre-commitments to protect individuals from being sacrificed for aggregate benefits, grounded in the equal dignity of persons.

In AI governance, rights-centred design means that certain applications should be categorically prohibited regardless of efficiency arguments — including social scoring systems that deny individuals opportunities based on algorithmic assessments of their character, biometric surveillance in public spaces for political monitoring, and AI systems that manipulate individuals through subliminal techniques. It also means that all AI systems deployed in consequential decisions affecting individuals must satisfy minimum standards of explainability, fairness, and due process — standards whose specific content should be developed through deliberative processes but whose normative basis is not negotiable.

4.5 International Policy Coordination

AI governance cannot be purely national. The global character of AI development and deployment creates powerful incentives for regulatory arbitrage; governance gaps in one jurisdiction become vulnerabilities for all. Moreover, the most consequential AI risks —

autonomous weapons, AI-enabled cyberattacks, large-scale manipulation of democratic processes — are inherently transnational in character and cannot be addressed through unilateral national action.

International coordination does not require global regulatory uniformity, which is neither achievable nor necessarily desirable given legitimate differences in values, institutional contexts, and risk assessments. What it does require is: mutual recognition arrangements that reduce compliance fragmentation without race-to-the-bottom dynamics; common minimum standards for the most high-risk AI applications; shared incident reporting and information-sharing mechanisms; and coordinated approaches to frontier AI development, including international governance institutions with meaningful authority and technical capacity.

5. CROSS-CUTTING GOVERNANCE CHALLENGES

5.1 Governing General-Purpose AI

The proliferation of general-purpose AI systems — large language models and foundation models capable of performing a vast range of tasks without task-specific training — poses a distinctive governance challenge. These systems are developed by a small number of highly resourced actors, deployed by a much larger ecosystem of businesses and developers, and ultimately affect billions of users. The risk-based frameworks developed for specific AI applications are poorly suited to systems whose downstream uses cannot be fully anticipated at the point of development.

Pragmatic governance of general-purpose AI requires a two-layer approach: upstream regulation of foundation model developers, imposing requirements around safety evaluation, transparency, and capability disclosure; and downstream regulation of deployers and specific applications, maintaining the risk-based calibration appropriate to particular contexts. The EU AI Act's provisions on general-purpose AI models, introduced in the final stages of its negotiation, represent a first attempt at this approach, though their implementation will require further refinement.

5.2 Democratic Legitimacy and Technocratic Capture

A persistent tension in AI governance concerns the distribution of regulatory authority between technical experts and democratic institutions. Effective AI oversight requires deep technical knowledge that elected legislators and generalist civil servants typically lack, creating pressure to delegate regulatory authority to specialised agencies or technical bodies. Yet delegation carries risks: technocratic institutions may be captured by the industries they regulate, may lack democratic accountability, and may systematically discount the perspectives and values of those most affected by AI.

Pragmatic governance must navigate this tension without resolving it in favour of either pure technocracy or pure democracy. Expert regulatory bodies are necessary and legitimate, but their mandates should be clearly defined by democratic institutions, their deliberations should be transparent and open to challenge, their leadership should be publicly accountable, and their decisions should be subject to meaningful judicial review. Civil society organisations, academic researchers, and affected communities should have formal roles in regulatory processes, not merely advisory ones.

5.3 The Economic Dimension: Competition and Concentration

Discussions of AI governance often focus on safety and rights concerns, but economic governance dimensions are equally important. The AI sector exhibits powerful network

effects and scale economies that tend toward concentration: a small number of firms with access to massive compute, proprietary training data, and elite research talent increasingly dominate the development of frontier AI. This concentration poses risks that are simultaneously economic (reduced competition and innovation diversity), political (undue corporate influence over governance processes), and epistemic (homogenisation of the values and assumptions embedded in widely deployed AI systems).

Pragmatic AI governance must address these structural dynamics directly, through competition policy, data access regulation, and public investment in AI research and infrastructure. Regulatory frameworks that impose compliance costs primarily on smaller actors, or that create intellectual property regimes that entrench incumbent advantages, may produce the opposite of their intended effects — concentrating AI development in fewer, less accountable hands.

CONCLUSIONS AND POLICY RECOMMENDATIONS

The governance of Artificial Intelligence is among the defining policy challenges of the early twenty-first century. The analysis presented in this article supports four principal conclusions.

First, neither permissive nor precautionary approaches, applied in their pure forms, constitute adequate responses to the challenge. AI governance requires a calibrated, evidence-informed approach that takes seriously both the social costs of harmful AI and the social costs of foregone beneficial AI.

Second, pragmatist political philosophy and adaptive governance theory provide more useful foundations for AI policy than either uncritical techno-optimism or categorical techno-precautionism. Governance frameworks must be experimental, revisable, and democratically accountable.

Third, the five pillars identified — proportionate risk-based regulation, adaptive governance mechanisms, multi-stakeholder deliberation, rights-centred design, and international coordination — are mutually reinforcing and should be pursued together rather than traded off against each other.

Fourth, certain cross-cutting challenges — general-purpose AI governance, democratic legitimacy, and economic concentration — require specific institutional responses that existing frameworks have not yet adequately addressed.

ON THE BASIS OF THIS ANALYSIS, THE FOLLOWING POLICY RECOMMENDATIONS ARE ADVANCED:

- **Adopt risk-tiered regulatory frameworks** calibrated to actual harms and regularly updated in light of deployment evidence, with adequately funded and technically capable enforcement bodies.
- **Build adaptive mechanisms** into all AI governance frameworks, including mandatory review cycles, regulatory sandboxes, and horizon-scanning functions within regulatory agencies.
- **Establish inclusive deliberative processes** that give affected communities, civil society organisations, and independent researchers meaningful influence in AI governance decisions.

- **Entrench rights-based constraints** as non-negotiable floors on AI deployment, prohibiting applications that are incompatible with human dignity, democratic governance, and due process.
- **Pursue international coordination** through mutual recognition arrangements, common minimum standards for high-risk AI, and the development of multilateral governance institutions with meaningful authority.
- **Develop a coherent two-layer governance approach** for general-purpose AI, combining upstream regulation of foundation model developers with downstream risk-based regulation of specific applications and deployers.
- **Address AI market concentration** through competition policy, data access regulation, and public investment in AI research and infrastructure to ensure that AI governance serves broad public interests rather than incumbent corporate interests.

The pragmatic approach to AI governance is, ultimately, an expression of confidence in democratic societies' capacity to govern their own technological futures wisely, if they invest in the institutions and processes needed to do so. That confidence is not naive; it is grounded in the recognition that the alternative — governance by default, driven by technological momentum and corporate power — is far more dangerous than the challenges of deliberate, evidence-informed, accountable policymaking.

REFERENCES

1. Anderson, E. (2010). *The imperative of integration*. Princeton University Press.
2. Calo, R. (2017). Artificial intelligence policy: A primer and roadmap. *UC Davis Law Review*, 51(2), 399–435.
3. Council of Europe. (2024). *Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law*. CETS No. 225.
4. Dafoe, A. (2018). *AI governance: A research agenda*. Future of Humanity Institute, University of Oxford.
5. Delvaux, M. (2017). Report with recommendations to the Commission on Civil Law Rules on Robotics. European Parliament Committee on Legal Affairs.
6. Dewey, J. (1927). *The public and its problems*. Henry Holt.
7. European Parliament and Council. (2024). Regulation on a European approach for artificial intelligence (EU AI Act). OJ L.
8. Floridi, L., et al. (2018). AI4People — An ethical framework for a good AI society. *Minds and Machines*, 28(4), 689–707.
9. Hadfield, G. K., & Clark, J. (2021). Regulatory markets: The future of AI governance. *arXiv:2102.12874*.
10. Marchetti, V. (2023). Adaptive governance and AI: Rethinking regulatory theory for emerging technologies. *European Journal of Risk Regulation*, 14(1), 1–22.
11. NIST. (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*. National Institute of Standards and Technology.
12. OECD. (2019). *Recommendation of the Council on Artificial Intelligence*. OECD/LEGAL/0449.
13. Risse, M. (2023). *Political philosophy of artificial intelligence*. Oxford University Press.
14. Roberts, H., Cows, J., Morley, J., Taddeo, M., Wang, V., & Floridi, L. (2021). The Chinese approach to artificial intelligence: An analysis of policy, ethics, and regulation. *AI & Society*, 36(1), 59–77.
15. Sunstein, C. R. (2005). *Laws of fear: Beyond the precautionary principle*. Cambridge University Press.

16. Sunstein, C. R. (2021). *Sludge: What stops us from getting things done and what to do about it*. MIT Press.
17. Veale, M., & Borgesius, F. Z. (2021). Demystifying the Draft EU Artificial Intelligence Act. *Computer Law Review International*, 22(4), 97–112.
18. Winfield, A. F. T., & Jirotko, M. (2018). Ethical governance is essential to building trust in robotics and artificial intelligence. *Philosophical Transactions of the Royal Society A*, 376(2133).