

ARTIFICIAL INTELLIGENCE AND LAW: EMERGING CHALLENGES OF LIABILITY, PRIVACY, ACCOUNTABILITY AND JUSTICE IN THE DIGITAL ERA

Prof. Michael J. Anderson ¹, Prof. Elizabeth Carter ²

¹ School of Law and Technology Policy, California, USA.

² Institute for Artificial Intelligence, Law and Ethics, Washington, D.C., USA.

Received: 10 June 2026 **Revised:** 28 July 2026 **Accepted:** 28 August 2026 **Published:** 9 Sep 2026

ABSTRACT

Artificial Intelligence (AI) has emerged as one of the most transformative technologies of the twenty-first century, influencing healthcare, finance, education, employment, policing, judicial administration, commerce, national security and public governance. The rapid development of generative AI and increasingly autonomous computational systems has, however, created legal challenges that existing regulatory frameworks were not necessarily designed to address. Questions concerning privacy, data protection, algorithmic discrimination, intellectual property, liability, transparency, due process, consumer protection and human accountability have consequently become central to contemporary legal scholarship.

The relationship between artificial intelligence and law is particularly complex because AI systems are capable of processing enormous quantities of information, generating predictions and recommendations, producing content and, in certain contexts, taking or influencing decisions traditionally performed by human actors. The resulting legal problem is therefore not simply whether AI should be regulated, but how law should allocate responsibility among developers, providers, deployers, users and other actors throughout the AI lifecycle. The challenge becomes more significant where automated systems produce harmful or discriminatory outcomes and where affected individuals cannot understand or effectively challenge the reasoning behind an AI-assisted decision.

Keywords: Artificial Intelligence; AI and Law; Artificial Intelligence Regulation; Algorithmic Accountability; Data Protection.

© 2026. The Author(s). This is an open-access article distributed under the terms of the Creative Commons Attribution 4.0 International License (CC BY 4.0), which permits unrestricted use, distribution and reproduction in any medium, provided the original work is properly cited.

1. INTRODUCTION

Artificial Intelligence has moved from being primarily a subject of scientific research to becoming an important component of contemporary social, economic and governmental systems. AI technologies are now capable of performing tasks involving language generation, image recognition, prediction, classification, decision support, recommendation and increasingly complex forms of autonomous action. Generative AI has further expanded the practical reach of artificial intelligence by enabling systems to produce text, images, audio,

software code and other forms of content. These developments have generated substantial economic and social opportunities, but they have simultaneously exposed limitations in existing legal frameworks.

The central legal difficulty arises from the fact that AI systems can operate through complex computational processes that may be difficult for ordinary users, regulators and even developers to explain fully. Where an AI system contributes to a decision affecting employment, credit, healthcare, insurance, education, policing or access to public services, questions immediately arise concerning the legality of that decision. If the decision is discriminatory, who should be responsible? If an AI-generated output causes economic or physical harm, should liability rest with the developer, the provider, the deploying organisation or the individual user? If copyrighted material is used to train a generative AI system, what rights do authors and publishers possess? If personal data are processed by an AI system, what safeguards should protect individuals from unlawful surveillance or profiling?

These questions demonstrate that AI law cannot be treated as a single specialised branch of technology regulation. Rather, artificial intelligence interacts with multiple existing areas of law, including constitutional law, administrative law, contract law, tort law, criminal law, intellectual property law, consumer protection, competition law, privacy and data protection law. AI therefore challenges the boundaries between traditional legal disciplines.

The regulatory landscape is also rapidly developing. The European Union's Artificial Intelligence Act establishes a comprehensive legal framework based upon different categories of AI risk and seeks to promote human-centric and trustworthy AI while protecting fundamental rights, health, safety and democratic values. The current consolidated EU text records staged application, including application of the main framework from 2 August 2026, subject to specified exceptions.

The United States has pursued a different model. Rather than relying upon a single comprehensive federal AI statute, U.S. AI governance has increasingly developed through existing federal laws, agency authority, executive actions and sector-specific initiatives. Recent federal policy has emphasised maintaining U.S. leadership in AI, promoting innovation and addressing security risks associated with advanced AI systems.

India occupies an important position within this comparative landscape. The country is simultaneously seeking to encourage technological innovation, expand digital public infrastructure and address risks associated with data, automated decision-making and emerging AI systems. The continuing development of AI governance institutions within MeitY demonstrates that India's regulatory approach is evolving rather than static.

The fundamental challenge for law is therefore to establish a framework that does not unnecessarily restrict beneficial technological innovation while ensuring that technological power remains subject to legal accountability. AI systems should not become legal black boxes in which affected individuals are unable to determine why a decision was made, challenge an adverse outcome or obtain an effective remedy.

This paper examines these challenges through the concept of **responsible and accountable AI governance**. It argues that the central objective of AI law should be neither unrestricted technological freedom nor excessive regulation. Instead, the legal framework should establish a proportionate balance between innovation and fundamental rights, supported by transparency, human oversight, risk assessment, accountability and effective remedies.

2. RESEARCH OBJECTIVES

The primary objective of this research is to examine the emerging relationship between artificial intelligence and law and to identify the principal legal challenges created by the increasing deployment of AI systems in private and public sectors. The research seeks to

determine whether existing legal principles are sufficiently flexible to regulate AI-related risks or whether new legislative and institutional mechanisms are necessary.

The first objective is to examine the conceptual and legal evolution of artificial intelligence and understand why conventional legal categories may be insufficient to address systems capable of autonomous or semi-autonomous decision-making. Particular attention is given to the distinction between AI as a technological tool and AI as a system capable of influencing decisions that have direct consequences for individuals.

The second objective is to analyse the relationship between AI and **privacy and data protection**. AI systems frequently depend upon large datasets for training, operation and improvement. The study therefore examines questions concerning lawful data collection, purpose limitation, data minimisation, automated profiling, consent, surveillance and individual control over personal information.

The third objective is to examine **algorithmic accountability and discrimination**. AI systems may reproduce or amplify biases contained within training data or institutional practices. The research therefore considers whether existing equality and anti-discrimination principles can adequately address algorithmic decision-making and whether individuals should possess rights to explanation, review and human intervention in high-impact decisions.

The fourth objective is to examine **AI liability**. Traditional legal systems generally attribute responsibility to identifiable human or corporate actors. AI complicates this model where harm results from interactions among developers, providers, deployers, users and autonomous system behaviour. The research seeks to determine how liability should be distributed across the AI lifecycle.

The fifth objective is to investigate **generative AI and intellectual property**, particularly copyright-related questions concerning training data, AI-generated works, authorship, infringement and the relationship between human creativity and machine-generated content.

The sixth objective is to examine the use of AI in the **legal system and judicial administration**, including legal research, document analysis, predictive tools, evidence evaluation and judicial decision support. The study considers the need to preserve judicial independence, procedural fairness and human responsibility.

Finally, the research aims to conduct a comparative analysis of emerging approaches in India, the European Union and the United States and propose a rights-based, risk-sensitive and innovation-compatible framework for AI governance.

3. RESEARCH QUESTIONS

The study is guided by a series of interconnected research questions concerning the legal regulation and governance of artificial intelligence.

The first research question is whether existing legal frameworks are capable of adequately regulating AI systems or whether dedicated AI legislation is necessary. This question requires examination of whether existing principles of privacy, consumer protection, tort, contract, intellectual property and administrative law can be applied effectively to AI-related risks.

The second question concerns **algorithmic accountability**. When an AI system makes or significantly influences a decision affecting an individual, what level of transparency should the law require? Should individuals have a legal right to know that AI was used, understand the principal factors influencing a decision and obtain meaningful human review?

The third question concerns **privacy and data protection**. How should legal systems reconcile the large-scale data requirements of AI development with individual rights relating to personal information? The study considers whether conventional concepts such as consent and purpose limitation remain adequate when data may be reused in complex machine-learning systems.

The fourth question concerns **algorithmic discrimination**. How should law respond when an AI system produces systematically discriminatory outcomes without any explicit discriminatory intention on the part of its developer or user? The research examines the relationship between technical bias and legal discrimination.

The fifth research question concerns **liability for AI-generated harm**. Where an autonomous or generative AI system produces harmful content, an incorrect recommendation or a damaging action, which participant in the AI ecosystem should bear legal responsibility? The study considers whether existing negligence, product-liability, consumer-protection and regulatory frameworks are sufficient.

The sixth question addresses **intellectual property**. Who should control or receive legal protection for AI-generated content? What rights do copyright holders possess when their works are used in training datasets, and how should the law distinguish between human creativity and machine-generated output?

The seventh question concerns the use of AI by courts, lawyers and public authorities. Can AI improve access to justice and administrative efficiency without undermining procedural fairness, confidentiality, judicial independence and human accountability?

Finally, the research asks what regulatory model India should adopt in light of the approaches developed in the European Union and the United States. Should India adopt a comprehensive AI statute, sector-specific regulation, a risk-based framework, or a hybrid model combining existing laws with dedicated AI governance mechanisms?

4. RESEARCH METHODOLOGY

This research adopts a **qualitative, doctrinal and comparative legal methodology**. The doctrinal component examines legislation, regulations, judicial decisions, governmental policies, regulatory documents and recognised principles of law concerning artificial intelligence, privacy, data protection, intellectual property, liability, human rights and digital governance.

The comparative dimension focuses principally upon **India, the European Union and the United States**. These jurisdictions provide useful contrasting models of AI governance. The European Union has adopted a comprehensive horizontal regulatory framework through the Artificial Intelligence Act, which establishes harmonised rules concerning the development, placing on the market and use of AI systems and seeks to protect fundamental rights while promoting trustworthy and human-centric AI.

The United States provides a contrasting regulatory environment characterised by existing federal legal authorities, sector-specific regulation, executive action and evolving national policy. Recent federal initiatives demonstrate an emphasis on AI innovation, cybersecurity, national security and a coordinated national policy framework.

India is examined as an emerging AI-governance jurisdiction. Rather than assuming that India currently possesses a single comprehensive AI statute, the research analyses the interaction between existing digital and data-protection laws, technology policy, sectoral regulation and emerging institutional arrangements. Current MeitY materials indicate continuing development of AI governance structures, including the establishment of an AI Governance and Economic Group in 2026.

The study also adopts a **rights-based normative approach**. AI regulation is evaluated against principles of privacy, equality, dignity, transparency, accountability, procedural fairness, human oversight and access to effective remedies. The research proceeds on the assumption that technological innovation should remain compatible with fundamental legal values.

A **risk-based analytical framework** is also employed. Rather than treating all AI applications as equally dangerous, the study distinguishes between applications according to their potential impact upon individuals and society. This approach is particularly relevant to the EU

model, where the AI Act establishes differentiated regulatory obligations according to the nature and level of risk.

The research relies upon secondary legal and academic materials and does not claim to present primary empirical findings. Its conclusions are based upon doctrinal interpretation, comparative analysis and normative evaluation.

5. Evolution of Artificial Intelligence and Its Legal Significance

Artificial Intelligence (AI) has evolved from a largely theoretical concept into a general-purpose technological infrastructure capable of influencing economic, social, governmental and legal decision-making. Early AI systems were primarily based on rule-based programming, in which computers performed tasks according to predetermined instructions. The subsequent development of machine learning enabled systems to identify patterns from large datasets rather than relying exclusively upon explicit human instructions. The emergence of deep learning, neural networks and generative AI has further transformed the technological landscape by enabling machines to generate text, images, audio, software code and other forms of content.

The legal significance of this transformation lies in the increasing movement of AI from **assistance to autonomy**. Traditional software generally performs a function predetermined by its developer. Contemporary AI systems, however, may generate outputs that were not specifically anticipated by their programmers. This creates an important legal question: **who should bear responsibility when an AI system produces a harmful, discriminatory, misleading or unlawful outcome?**

AI is now used in areas such as healthcare, financial services, education, employment, policing, transportation, advertising, public administration and judicial research. Consequently, decisions previously made exclusively by human beings may increasingly involve algorithmic recommendations or automated assessments. The resulting legal relationship is complex because several actors may participate in the AI lifecycle, including developers, data providers, model creators, deployers, users, intermediaries and regulators.

The legal significance of AI can therefore be understood through four interconnected dimensions:

Dimension	Emerging Legal Concern
Individual	Privacy, autonomy, dignity and discrimination
Commercial	Intellectual property, contracts, consumer protection and liability
Governmental	Administrative decision-making, surveillance and public accountability
Societal	Employment, democracy, misinformation, inequality and fundamental rights

The emergence of AI therefore requires law to move beyond conventional technology regulation. The central challenge is no longer merely whether a technology is lawful, but **whether the manner in which that technology operates is consistent with constitutional values, human rights and principles of justice.**

6. AI Regulation and the Changing Nature of Law

The rapid development of AI has challenged traditional regulatory approaches. Conventional legislation generally assumes that the regulated activity can be clearly identified and that the responsible actor can be determined. AI systems complicate both assumptions. An AI application may simultaneously involve software developers, data processors, cloud-service providers, model developers and end-users located in different jurisdictions.

Modern AI regulation is therefore increasingly adopting a **risk-based and lifecycle-oriented approach**. Instead of treating every AI system identically, regulatory frameworks distinguish between systems according to the degree of risk they create. High-risk applications may

require stronger safeguards, documentation, transparency, human oversight and accountability mechanisms.

The European Union provides the most developed example of this approach through the **Artificial Intelligence Act (Regulation (EU) 2024/1689)**. The Act establishes a comprehensive framework for AI and uses differentiated regulatory obligations according to risk. Its principal application date is 2 August 2026, although certain provisions became applicable earlier and some obligations have later transition dates.

The significance of the EU approach extends beyond Europe because organisations operating internationally may need to adapt their AI governance practices to European requirements. It demonstrates a shift from regulating technology merely as a product toward regulating **the risks created by its deployment**.

The United States represents a substantially different regulatory philosophy. Recent federal policy has strongly emphasised innovation, national competitiveness, security and the development of a more uniform national framework. The White House's March 2026 AI legislative framework identifies innovation, intellectual property, free speech, child protection, workforce development and national competitiveness among its central objectives.

Similarly, the June 2026 Executive Order on advanced AI innovation and security emphasises AI innovation and cybersecurity while directing enforcement of existing federal criminal laws against unlawful computer access facilitated by AI.

India is developing an institutional governance model rather than relying exclusively upon a single comprehensive AI statute. The Ministry of Electronics and Information Technology (MeitY) has established structures concerned with AI and emerging technologies, while the **AI Governance and Economic Group (AIGEG)** was constituted in April 2026 to coordinate policy, examine regulatory gaps, promote responsible AI innovation and develop India's broader AI governance strategy.

These approaches demonstrate three broad regulatory philosophies:

European Union → Rights and risk-based regulation
United States → Innovation, security and national competitiveness
India → Responsible innovation, institutional coordination and adaptive governance

The comparative lesson is that no single regulatory model can adequately address every AI-related risk. Future AI governance will require a combination of innovation-friendly regulation, fundamental-rights protection, technical standards, institutional supervision and effective remedies.

7. Artificial Intelligence and the Right to Privacy

Privacy represents one of the most significant legal challenges created by artificial intelligence. AI systems depend heavily upon data, and increasingly sophisticated systems require enormous quantities of personal, behavioural, biometric and contextual information. The more information an AI system processes, the greater the possibility that individuals may lose meaningful control over information concerning themselves.

AI can affect privacy at multiple stages. Data may be collected from websites, mobile applications, social media platforms, public records, sensors, cameras and commercial databases. Machine-learning systems can then analyse this information to infer characteristics that an individual may never have directly disclosed.

This creates an important distinction between **data collection and data inference**. Traditional privacy law has generally focused on information that an individual provides or that an organisation collects. AI creates an additional problem: a system can generate new conclusions from apparently unrelated pieces of information.

For example, an AI system may analyse purchasing patterns, location information, online behaviour and demographic characteristics to predict an individual's preferences, financial

circumstances or behavioural tendencies. Even where the individual never expressly disclosed such information, the resulting inference may have significant consequences.

AI therefore challenges several established principles of privacy law:

1. **Purpose limitation** — information collected for one purpose may subsequently be used to train or operate another AI system.
2. **Data minimisation** — AI development often creates incentives to collect increasingly large datasets.
3. **Transparency** — individuals may not know how their information is processed by complex models.
4. **Accuracy** — incorrect data or inaccurate algorithmic inferences may produce harmful outcomes.
5. **Control and consent** — meaningful consent becomes difficult where individuals cannot understand future AI uses of their data.
6. **Right to correction or deletion** — removing information from complex trained models may be technically and legally difficult.

The problem becomes even more significant with facial recognition, biometric identification and predictive analytics. AI-enabled surveillance can enable governments and private organisations to identify, classify and monitor individuals on a scale that was previously difficult to achieve.

Privacy, therefore, should not be understood merely as secrecy. It is closely connected with **autonomy, dignity, informational self-determination and freedom from unjustified surveillance**. An effective AI governance framework must consequently examine not only whether data processing is technically secure but also whether the purposes, methods and consequences of processing are legally legitimate.

From a constitutional perspective, AI strengthens the argument that privacy protection must extend beyond conventional personal information. Individuals require protection against both **unauthorised collection of data and unjustified algorithmic inference from lawfully obtained data**.

8. Data Protection and AI Governance

Data protection constitutes the foundation upon which responsible AI governance must be constructed. Since modern AI systems depend heavily upon data, weaknesses in data governance can directly translate into weaknesses in AI decision-making.

A fundamental problem is that AI systems may reproduce the characteristics of the datasets upon which they are trained. If a dataset contains historical discrimination, incomplete information or systematic social inequalities, an AI system may reproduce or amplify those patterns. Consequently, data protection and AI governance cannot be treated as entirely separate legal disciplines.

Effective AI governance should address the complete **data-to-decision lifecycle**:

Data Collection → Data Storage → Data Processing → Model Training → Model Deployment → AI Output → Human Decision → Review/Remedy

At each stage, legal safeguards should be incorporated.

AI Lifecycle Stage	Principal Legal Safeguard
Data collection	Lawfulness, necessity and transparency
Data storage	Security and access controls
Data processing	Purpose limitation and proportionality
Model training	Data quality and bias assessment
Deployment	Risk assessment and human oversight

AI Lifecycle Stage	Principal Legal Safeguard
AI output	Explainability and accuracy
Decision-making	Accountability and procedural fairness
Post-decision	Appeal, correction and effective remedy

India's AI governance environment is developing through a combination of digital regulation, data governance, institutional initiatives and technology-policy mechanisms. MeitY's AI and Emerging Technologies Group identifies AI, machine learning, deep learning, robotics and related technologies as areas requiring policy and strategic development.

The establishment of AIGEG in 2026 is particularly significant because its terms of reference include coordination across ministries and regulators, examination of emerging AI risks and regulatory gaps, promotion of responsible AI innovation, and development of India's position and strategy on AI governance.

The future of data protection in the AI era should therefore move beyond a narrow model of **"protecting data"** toward a broader principle of **"protecting individuals throughout the data lifecycle."**

This requires meaningful transparency, strong security standards, responsible data acquisition, mechanisms for correcting inaccurate information, safeguards against discriminatory profiling, human oversight and accessible remedies.

The central principle emerging from this analysis is that **AI governance should be human-centred rather than technology-centred**. Innovation should remain an important legal objective, but technological progress cannot by itself justify interference with privacy, dignity, equality or procedural fairness.

Declarations

Conflict of interest: The authors declare that they have no conflict of interest.

Funding: This review received no specific grant from any funding agency in the public, commercial or not-for-profit sectors.

Ethical approval: Not applicable. This article is a review of published literature and did not involve human participants or animal subjects.

Data availability: No new data were generated or analyzed in support of this review. All sources relied upon are cited in the References.

Use of AI tools: No generative artificial intelligence tool was used in the conception, drafting or analysis of this article. The authors take full responsibility for the content.

Author contributions: Author contributed to the conception of the review, the survey of the literature, and the drafting and revision of the manuscript. Author read and approved the final version.