



Legislative Gaps in India's AI Regulation: Need for A Dedicated AI Law

Dr. Keval Govardhan Ukey^{1*} & Mrs. Tanavi Prasad Naik²

^{1*}Associate Professor at School of Law, Sandip University, Nashik

²Ph.D. Scholar at School of Law, Sandip University, Nashik

Abstract

Artificial Intelligence (AI) is reshaping various sectors in India, including healthcare, finance, governance, and surveillance. However, the country currently lacks a dedicated legal framework to address the ethical, legal, and privacy challenges posed by AI. The rapid adoption of AI technologies has introduced concerns such as algorithmic bias, lack of transparency, privacy vulnerabilities, and accountability issues gaps that existing laws fail to fully cover.

This paper critically examines the inadequacies of India's current AI-related legal framework and assesses the limitations of existing legislation, such as the Information Technology Act, 2000, and the Digital Personal Data Protection Act, 2023. While these laws regulate certain digital activities, they do not specifically address AI-related risks, such as the explainability of AI-driven decisions, liability issues, and mechanisms to prevent automated bias and discrimination. The paper also explores global AI regulatory models, including the European Union's AI Act, the United States' Algorithmic Accountability Act, and China's AI governance framework, to draw insights from international best practices.

Additionally, the study highlights the urgent need for a specialized AI law in India to promote responsible AI governance. It suggests a legal framework that includes AI risk classification, mandatory AI impact assessments, transparency mandates, and accountability mechanisms. Furthermore, it underscores the importance of public awareness and ethical AI deployment standards to foster a fair and inclusive AI ecosystem. By addressing these gaps through a dedicated legal framework, India can strike a balance between AI innovation and the protection of fundamental rights, ensuring AI serves society in an ethical and equitable manner.

Keywords Artificial Intelligence (AI), AI Regulation in India, Legal and Ethical Challenges, AI governance, Algorithmic Bias, AI Transparency and Accountability, AI Liability Framework, Privacy and Data Protection, Digital Personal Data Protection Act (DPDP Act), 2023, Information Technology (IT) Act, 2000, AI Risk Classification, AI Ethics and Fairness, Surveillance and Fundamental Rights, Global AI Regulations, Policy Recommendations for AI Law

Received: 04/04/2024, Accepted: 14/05/2024

DOI: <https://doi.org/10.53555/AJBR.v27i2.8260>

© 2024 The Author(s).

This article has been published under the terms of Creative Commons Attribution-Noncommercial 4.0 International License (CC BY-NC 4.0), which permits noncommercial, unrestricted use, distribution, and reproduction in any medium, provided that the following statement is provided. "This article has been published in the African Journal of Biomedical Research"

1. Introduction

AI-driven technologies are reshaping the digital landscape in India, influencing diverse sectors such as healthcare, banking, education, agriculture, and governance.³ However, the widespread deployment of

AI introduces a host of legal and ethical concerns that remain inadequately addressed within India's existing legal framework. The increasing reliance on AI for decision-making, automation, and predictive analytics raises issues related to data privacy, algorithmic bias, transparency, and the accountability of AI-driven outcomes. Without a

¹ The author is Associate Professor at School of Law, Sandip University, Nashik

² The author is Ph.D. Scholar at School of Law, Sandip University, Nashik

³ NITI Aayog, 2018

comprehensive AI-specific legal framework, India faces challenges in regulating AI in a manner that ensures fairness, non-discrimination, and protection of fundamental rights (Chakravarthi, 2023).

Currently, India relies on a patchwork of sectoral regulations, general IT laws, and data protection provisions to manage AI-related risks. The Information Technology Act, 2000, primarily addresses cybersecurity and digital crimes but does not cover AI-specific concerns such as liability in automated decision-making or ethical AI use (Ministry of Law and Justice, 2000). The recently enacted Digital Personal Data Protection Act, 2023, focuses on personal data protection but lacks provisions on broader AI governance issues like algorithmic transparency, AI explainability, and fairness in automated processing.⁴

2. Existing Legal Framework for AI in India

2.1 Information Technology Act, 2000

The Information Technology (IT) Act, 2000 primarily governs cyber activities, electronic commerce, and data security. However, it does not explicitly regulate AI, automated decision-making, or algorithmic transparency. The Act was enacted at a time when AI was not as prevalent as it is today, and its provisions mainly focus on cybercrime, digital signatures, and data protection in electronic transactions.⁵ It lacks any specific clauses on AI governance, ethical AI development, or accountability mechanisms for AI-driven decisions.

One of the key limitations of the IT Act is that it does not address AI-based decision-making processes, which are becoming increasingly common in financial services, healthcare, and government operations. AI systems can make critical decisions affecting individuals, such as approving or denying loans, diagnosing diseases, or determining eligibility for welfare schemes. The absence of regulatory oversight on AI-driven decision-making raises concerns about fairness, bias, and discrimination.

Additionally, the IT Act does not mandate transparency in AI algorithms or require organizations to explain AI-generated decisions. The lack of an explicit framework for algorithmic transparency and accountability means that individuals affected by AI decisions often have no legal recourse or mechanism to challenge unfair outcomes. In contrast, regulatory models such as the European Union's General Data Protection Regulation (GDPR) emphasize the right to explanation and fairness in automated processing,⁶ highlighting a significant gap in India's IT regulations.

Furthermore, while the IT Act includes provisions for cybersecurity and data protection, it does not address AI-specific risks such as deepfakes, automated misinformation, or AI-generated fraud. These emerging challenges necessitate a more robust legal framework that explicitly regulates AI technologies and ensures that they are used ethically and responsibly.

2.2 Digital Personal Data Protection Act, 2023 (DPDP Act)

The DPDP Act focuses on data privacy and protection but does not address broader AI-related issues such as algorithmic bias,

AI ethics, and liability in case of AI-driven harm. While the Act introduces essential provisions for the collection, processing, and storage of personal data, it lacks specific guidelines on how AI systems should handle data responsibly.

One major gap is the absence of provisions ensuring fairness and non-discrimination in AI-driven decision-making. AI models, particularly in sectors like finance, employment, and law enforcement, have been found to exhibit biases against certain groups, reinforcing existing inequalities. Without explicit requirements for bias detection, fairness audits, and explainability in AI decision-making, the DPDP Act does not provide adequate safeguards against AI-induced discrimination. Moreover, the DPDP Act does not outline accountability mechanisms for AI systems that process personal data. In cases where an AI system makes an erroneous decision leading to harm—such as wrongful denial of credit or incorrect medical diagnosis the law does not clarify whether liability falls on the AI developer, deployer, or data controller.

Additionally, the Act does not mandate transparency in AI decision-making. Many AI-driven processes operate as “black boxes,” making it difficult for individuals to understand how decisions impacting them are made. In contrast, regulations like the European Union's AI Act and the United States' Algorithmic Accountability Act emphasize transparency and explainability in high-risk AI applications,⁷ highlighting a key shortcoming in India's current data protection regime.

2.3 Sector-Specific Regulations

AI is transforming healthcare in India through diagnostics, robotic surgeries, and predictive analytics. Yet, the Clinical Establishments Act, 2010 does not cover AI-specific concerns such as accountability for AI-generated diagnoses or transparency in medical algorithms.⁸

The Clinical Establishments Act, 2010, was enacted to standardize healthcare services across the country, ensuring minimum standards for medical facilities. However, it does not address AI-related issues such as liability in case of incorrect AI-generated diagnoses, the explainability of AI decisions in medical treatment, and ethical concerns related to automated patient care.

The growing integration of AI in healthcare presents several regulatory challenges that need urgent attention. One major concern is liability and accountability when an AI-powered diagnostic tool leads to an incorrect diagnosis resulting in complications or even fatalities, it remains unclear whether the responsibility lies with the AI developer, the healthcare provider, or the hospital. Additionally, algorithmic transparency poses a significant issue, as many AI systems operate as “black boxes,” making it difficult for doctors and patients to comprehend the reasoning behind medical decisions. This opacity can undermine trust in AI-driven healthcare solutions. Furthermore, bias in AI healthcare systems is a pressing concern, as models trained on non-representative datasets may produce skewed results that adversely affect certain demographic groups. Patient data privacy is another critical area, as AI systems handle vast amounts of sensitive medical

⁴ Ministry of Electronics and Information Technology, 2023

⁵ Ministry of Law and Justice, 2000

⁶ European Parliament & Council, 2016

⁷ European Commission, 2021; U.S. Congress, 2022

⁸ Ministry of Health and Family Welfare, 2010; Chakravarthi, 2023

information, raising the risk of data breaches, unauthorized access, and non-compliance with privacy regulations. Given these concerns, there is a clear and pressing need for a comprehensive, AI-specific legal framework to govern its use in healthcare, ensuring ethical practices, patient safety, and accountability.

- **AI in Finance:** The financial sector is increasingly leveraging AI for fraud detection, credit risk assessment, algorithmic trading, and personalized financial services. AI-driven credit scoring models and automated loan approvals have streamlined financial processes, making services more efficient. However, the Reserve Bank of India (RBI), which regulates fintech and banking operations, does not have a specific framework addressing AI explainability, bias mitigation, and accountability.

Key regulatory gaps in the use of artificial intelligence within the financial sector include the prevalence of algorithmic bias in credit scoring systems, where AI models may reinforce historical discrimination against marginalized communities due to biased or incomplete datasets. Additionally, the lack of explainability in AI-driven decisions—such as loan approvals and fraud detection undermines transparency and makes it difficult for consumers to understand the basis of outcomes. The deployment of AI in high-frequency trading raises concerns about market manipulation, flash crashes, and systemic instability, which are not sufficiently addressed by existing financial regulations. Furthermore, current frameworks lack dedicated consumer protection and redressal mechanisms for harms caused by erroneous AI decisions, including wrongful denial of services. These challenges highlight the urgent need for a dedicated regulatory framework in India that incorporates AI explainability, bias audits, robust consumer protection, and comprehensive risk management in the financial sector⁹.

- **AI in Governance and Surveillance:** AI is increasingly being integrated into governance and surveillance systems in India. Technologies such as facial recognition, predictive policing, and biometric-based authentication (e.g., Aadhaar) are widely used for law enforcement, public service delivery, and national security. However, the legal framework governing these applications is inadequate in addressing AI-specific risks such as mass surveillance, privacy violations, and the potential misuse of AI-driven surveillance systems.

The Aadhaar Act, 2016, provides a legal basis for the collection and use of biometric data for identification and authentication; however, its integration with AI-powered surveillance technologies has sparked concerns regarding mass data collection, lack of oversight, and inadequate safeguards against misuse. India currently lacks a comprehensive legal framework to regulate the deployment of AI-driven surveillance tools, such

as facial recognition and predictive policing, which raises significant concerns about civil liberties and potential rights violations. These systems often function without public transparency or accountability, making it difficult to evaluate their fairness, accuracy, or ethical implications. Furthermore, global studies have highlighted the risk of discriminatory outcomes from such technologies, particularly their disproportionate impact on marginalized groups. The legal infrastructure also falls short in ensuring data protection, leaving sensitive personal data vulnerable to unauthorized access and misuse. These gaps underscore the urgent need for a dedicated regulatory framework that ensures ethical, transparent, and rights-respecting use of AI in governance and surveillance.¹⁰

3. Legislative Gaps in India's AI Regulation

3.1 Absence of AI-Specific Legal Definitions

One of the fundamental challenges in India's AI regulation is the absence of clear legal definitions for key AI-related terms. Indian law does not define terms such as “automated decision-making,” “algorithmic transparency,” “AI ethics,” or “AI accountability.” This regulatory ambiguity leads to inconsistent interpretations and enforcement challenges, making it difficult to develop uniform AI governance standards.

The absence of clear legal definitions for artificial intelligence in Indian law creates significant challenges across regulatory, ethical, and accountability domains. This vagueness leads to inconsistencies in how different sectors interpret AI-related responsibilities, resulting in fragmented compliance and enforcement mechanisms. Moreover, in the event of harm caused by AI-driven decisions such as biased outcomes or system failures it becomes difficult to determine liability due to the lack of established legal responsibility for developers, deployers, or users. The ambiguity also hampers the enforcement of ethical principles like fairness, transparency, and accountability, as these values remain aspirational without legal grounding. Therefore, the development of a comprehensive AI law in India is essential, one that clearly defines key AI-related terms and assigns legal obligations to all actors in the AI ecosystem. Such a framework would enhance regulatory clarity, promote responsible innovation, and safeguard fundamental rights.¹¹

3.2 Lack of AI Accountability Mechanisms

There are no clear legal provisions addressing accountability in AI systems. If an AI system makes an erroneous or biased decision, it is unclear whether liability falls on the developer, deployer, or the AI itself.

India's current legal framework lacks clarity on accountability for AI-driven decisions, leading to significant concerns across sectors. One of the primary issues is the absence of defined legal responsibility among AI developers, deployers, and users, making it difficult to determine who is liable when AI causes harm. Additionally, there are no dedicated redress mechanisms

⁹ Narayan, R., & Basu, A. (2020). *AI and financial inclusion in India: Opportunities and challenges*. Observer Research Foundation.

Sengupta, A., & Parsheera, S. (2019). *Regulating the Future: AI and the Role of the State in India*. Carnegie India.

¹⁰ Ramanathan, U. (2020). *Aadhaar: A Biometric History of India's 12-digit Revolution*.

Arun, C. (2021). *AI and the Rule of Law in India*. Digital Asia Hub.

¹¹ Sengupta, A., & Parsheera, S. (2019). *Regulating the Future: AI and the Role of the State in India*. Carnegie India.

for individuals adversely affected by AI decisions in sensitive domains like healthcare, employment, or financial services. The global debate on granting AI legal personhood or addressing its liability status has yet to gain traction in India, further complicating the legal landscape. Moreover, AI systems that exhibit algorithmic bias remain unchecked due to the absence of mandatory audit mechanisms and corrective obligations. To bridge these gaps, India must establish a comprehensive legal framework that delineates accountability, mandates transparency, and ensures developers and deployers are held responsible for ethical and lawful AI use.¹²

3.3 Weak AI Transparency and Explainability Standards

Unlike the EU's General Data Protection Regulation (GDPR), Indian law does not mandate AI systems to be explainable, which raises concerns about fairness and bias in AI-driven decisions. Explainability is crucial in ensuring that AI models are transparent, accountable, and do not operate as "black boxes" that make decisions without clear reasoning.

The absence of explainability requirements in AI systems presents critical challenges to transparency, accountability, and fairness. When AI models operate as "black boxes," users and regulators are often unable to understand the rationale behind decisions, which obstructs efforts to detect and correct bias or errors. This opacity severely limits the ability of individuals to seek recourse in situations where AI decisions affect their rights or opportunities, such as loan approvals, hiring processes, or legal outcomes. Moreover, regulatory bodies struggle to enforce principles of fairness and non-discrimination without access to interpretable decision-making processes. As a result, the lack of mandated explainability undermines both individual rights and institutional oversight, highlighting the need for legal standards that compel AI systems to be transparent and justifiable in their operations. To address these concerns, India needs to implement AI explainability requirements within its legal framework, ensuring AI systems provide interpretable and justifiable outcomes.

3.4 Ethical and Bias Concerns in AI Decision-Making

AI systems in India have shown biases, particularly in hiring, lending, and policing. Without strict regulations, AI models may reinforce societal inequalities and discrimination.

AI systems deployed across critical sectors such as hiring, lending, and policing have demonstrated significant potential for bias and discrimination. In recruitment, AI-driven tools have shown tendencies to favour specific demographics, thereby marginalizing qualified candidates from underrepresented groups. Similarly, credit scoring algorithms used in lending decisions can embed socioeconomic biases, disproportionately affecting marginalized communities and limiting their financial access. In law enforcement, predictive policing technologies have been found to reinforce systemic discrimination by disproportionately targeting certain populations. These biases are exacerbated by the lack of transparency and regulatory oversight, making it difficult for affected individuals to understand or contest AI-driven decisions. The ethical deployment of AI in these domains necessitates strong legal

frameworks that mandate fairness, explainability, and accountability to prevent the reinforcement of social inequities. To mitigate these risks, India needs regulations mandating bias audits, fairness assessments, and transparency requirements for AI models used in high-stakes decision-making.

3.5 Limited Regulation on AI in Surveillance and Law Enforcement

Given the growing deployment of AI-powered facial recognition and predictive policing in India, there is an urgent need for a comprehensive legal framework to mitigate associated privacy and civil liberty risks. Such a framework should begin with clear legal definitions that delineate the scope and permissible uses of AI surveillance technologies. To ensure accountability and protect human rights, an independent regulatory body must be established to oversee these deployments. Additionally, transparency and explainability mechanisms are crucial—law enforcement agencies should be mandated to disclose the functioning of AI systems and offer recourse to individuals wrongly identified or impacted. Robust data protection policies must govern the collection, retention, and deletion of biometric and personal data to prevent misuse. Finally, judicial and legislative safeguards are essential, with mandatory judicial oversight before AI surveillance tools are deployed in sensitive areas, thereby preventing potential overreach and upholding constitutional protections.

Without such a framework, the unchecked use of AI in policing and surveillance could lead to widespread violations of fundamental rights, reinforcing systemic biases and eroding democratic freedoms.

4. Comparative Analysis: AI Regulations in Other Countries

4.1 European Union: AI Act

The European Union's Artificial Intelligence Act (AI Act) is the world's first comprehensive AI regulation designed to ensure the safe and ethical use of AI technologies. Proposed in April 2021 by the European Commission, the AI Act adopts a risk-based approach, classifying AI systems into different categories based on their potential impact on human rights, safety, and public well-being. It imposes stricter regulations on high-risk AI applications while promoting innovation in low-risk AI systems.

Key Features of the EU AI Act

1. Risk-Based Classification of AI Systems

The European Union's Artificial Intelligence Act introduces a pioneering risk-based framework that classifies AI systems into four categories: unacceptable, high, limited, and minimal risk—based on their potential impact on fundamental rights and societal values. AI systems considered to pose an *unacceptable risk*, such as those used for social scoring, manipulative behavioural targeting, or real-time biometric surveillance in public spaces (except under strict law enforcement exceptions), are outright banned. This approach aims to safeguard democratic freedoms and human dignity by pre-emptively prohibiting technologies that can cause systemic harm.

High-risk AI systems, which operate in critical domains such as healthcare, education, employment, infrastructure, and law

¹² Arun, C. (2021). *AI and the Rule of Law in India*. Digital Asia Hub.

enforcement, are subject to stringent regulatory obligations. These include mandatory risk assessments, transparency in algorithmic operations, human oversight, and demonstrable model robustness. *Limited-risk AI* like virtual assistants and chatbots must comply with transparency duties, while *minimal-risk AI* systems such as spam filters and translation tools are largely exempt. This structured classification ensures regulatory resources are prioritized where the risks are highest, balancing innovation with ethical safeguards.¹³

2. Transparency and Accountability Requirements

The AI Act mandates that high-risk AI systems must:

- Provide clear documentation about how the AI system works.
- Ensure data quality to prevent algorithmic bias and discrimination.
- Allow external audits and human oversight in decision-making.

Companies developing AI models must disclose their training data sources and ensure that AI outputs are explainable to affected users.

3. Restrictions on Facial Recognition and Surveillance

The EU AI Act places strict limitations on real-time facial recognition in public spaces. Law enforcement agencies can only use facial recognition in specific cases, such as:

- Preventing terrorist attacks.
- Locating missing persons or criminals.
- Investigating serious crimes (e.g., human trafficking, organized crime).

Even in these cases, the use of real-time biometric surveillance requires prior judicial or regulatory approval to prevent misuse.

4. AI Regulatory Bodies and Compliance Mechanisms

The EU AI Act mandates the establishment of national AI regulators across each EU member state, responsible for monitoring compliance with its provisions. Companies involved in developing high-risk AI systems must submit regular compliance reports and undergo periodic audits to ensure adherence to regulatory standards. Non-compliance can result in significant penalties, including fines of up to €30 million or 6% of a company's global annual revenue for serious violations, and fines of €20 million or 4% of global revenue for failing to meet AI transparency requirements.¹⁴ These financial penalties are designed to enforce strict adherence to transparency, accountability, and ethical deployment standards within the AI ecosystem.

4.1. Implications for India

India, which currently lacks a dedicated AI law, can draw valuable lessons from the EU AI Act to build a robust AI regulatory framework. Key takeaways include the adoption of a risk-based classification system to regulate AI applications

based on their potential harm, ensuring that high-risk AI systems are subject to stringent compliance standards. Additionally, enforcing transparency requirements would help prevent biased and opaque AI decision-making, promoting accountability and fairness. Restricting AI-powered mass surveillance is crucial to safeguarding privacy rights and protecting individuals from potential overreach. Finally, establishing independent AI regulatory bodies would be essential to oversee compliance, monitor the deployment of AI technologies, and ensure accountability in their use across various sectors.

These steps would help India develop a balanced and effective regulatory approach to AI, addressing both innovation and ethical concerns. While India's AI ecosystem is still evolving, incorporating elements of the EU's AI Act could help create a robust legal framework that balances innovation with ethical AI governance.

4.2 United States: Overview of the Algorithmic Accountability Act

The United States does not yet have a comprehensive federal AI law, but several legislative efforts have been made to regulate AI, particularly focusing on fairness, bias mitigation, and consumer protection. One of the most notable proposed regulations is the Algorithmic Accountability Act (AAA), which aims to enhance transparency and accountability in automated decision-making systems.¹⁵

The Algorithmic Accountability Act was first introduced in 2019 and reintroduced in 2022 by U.S. lawmakers to address concerns about biased AI models, discriminatory decision-making, and the lack of oversight in AI-driven systems. The Act would require companies to assess, document, and mitigate risks associated with automated decision systems (ADS), particularly those affecting consumers in critical sectors like finance, healthcare, housing, and employment.

The Algorithmic Accountability Act introduces a comprehensive approach to regulating AI systems by requiring large companies to conduct risk assessments for their AI-driven tools, especially in high-risk areas like hiring, healthcare, and credit scoring. These assessments aim to evaluate the impact of AI on privacy, fairness, and bias, ensuring that AI models do not discriminate or cause harm. The Act also mandates companies to mitigate biases within their datasets and algorithms, addressing concerns of AI-driven discrimination, particularly in sectors like employment, lending, and law enforcement. Transparency is a key provision, with developers required to disclose how their algorithms make decisions, allowing consumers to challenge unfair automated outcomes. The Federal Trade Commission (FTC) is responsible for enforcing compliance, and companies failing to adhere to these requirements could face penalties.

Alongside the Algorithmic Accountability Act, the U.S. has introduced the AI Bill of Rights, a set of principles aimed at ensuring AI systems are safe, non-discriminatory, and respect data privacy. While the Bill of Rights is not a law, it provides

¹³ European Commission. (2021). *Proposal for a regulation laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)*

¹⁴ Regulation (EU) 2024/1684 of the European Parliament and of the Council of 13 June 2024 on Harmonised Rules on

Artificial Intelligence (Artificial Intelligence Act) and Amending Certain Union Legislative Acts, 2024 O.J. (L 168)

1.

¹⁵ Algorithmic Accountability Act of 2022, S.3572, 117th Cong. (2022).

essential guidelines for AI governance, particularly emphasizing human oversight and protections against algorithmic discrimination. Various U.S. agencies, such as the FDA for medical devices and the EEOC for employment-related AI, have implemented regulations specific to their sectors. However, AI regulation in the U.S. remains fragmented across different industries and states, with the absence of a federal law and ongoing challenges such as corporate resistance and a slow legislative process.¹⁶

India can take several lessons from these U.S. initiatives as it develops its own AI regulatory framework. Drawing on the principles of the Algorithmic Accountability Act, India could mandate AI impact assessments, ensure transparency in AI decision-making, and protect consumer rights by enabling individuals to challenge automated outcomes. Additionally, India should consider establishing a dedicated AI regulatory body to oversee AI ethics and compliance, much like the FTC's role in the U.S. By applying these lessons, India can create a legal framework that promotes fairness, accountability, and transparency in AI systems, particularly in critical sectors such as finance, healthcare, and employment, to prevent bias and discrimination.

4.3 China: AI Governance Regulations

China has taken a proactive approach to AI regulation, implementing AI-specific policies and legal frameworks to govern the development and deployment of AI technologies. Unlike Western nations that primarily focus on human rights, transparency, and fairness, China's AI governance emphasizes state control, security, and economic development while imposing strict regulations on deepfakes, facial recognition, and generative AI.

China's approach to AI regulation is marked by its proactive stance in addressing emerging technologies such as deepfake, facial recognition, and generative AI. In 2023, China implemented the "Provisions on the Administration of Deep Synthesis of Internet Information Services," which require AI developers to label AI-generated content clearly and obtain user consent for altering biometric data. This regulation aims to combat the spread of misinformation and safeguard privacy by ensuring that deepfake content is easily identifiable. In the domain of facial recognition, China has introduced various measures, including the Supreme Court ruling that businesses cannot force customers to use facial recognition unless absolutely necessary. The Personal Information Protection Law (PIPL) further ensures that biometric data collection is minimal and lawful, requiring companies to justify their necessity. Despite these regulations, there remain significant concerns about the use of AI for mass surveillance and potential human

rights violations, particularly with the continued use of facial recognition by state agencies.¹⁷

In addition to these sector-specific regulations, China's AI governance aligns with its broader national objectives, including economic growth, national security, and strict state control. The Chinese government has set a target for AI to become a key driver of economic growth by 2030, while also strengthening surveillance and cybersecurity measures. This state-centric approach to AI governance is reinforced by other laws, such as the Cybersecurity Law (2017), the Data Security Law (2021), and the AI Ethics Guidelines (2021), all of which support China's strategic goals. However, this approach has faced criticism due to its potential to infringe on individual freedoms, as well as the challenges it poses for businesses, with strict approval processes and regulatory burdens slowing AI deployment in certain sectors. Furthermore, the heavy censorship of AI-generated content to align with government narratives has raised concerns about the limits it places on free expression and innovation.¹⁸

For India, there are valuable lessons to be learned from China's regulatory model, especially in areas like deepfake prevention and facial recognition. However, India must adopt a more balanced approach to AI regulation, ensuring that AI development is aligned with ethical principles and human rights. Key areas for focus include the regulation of deepfake technology through mandatory content labelling and user consent, the establishment of clear privacy safeguards for biometric data, and the creation of ethical guidelines for AI content generation without stifling free speech and innovation. Additionally, India should consider implementing AI security reviews to prevent cyber threats and the spread of misinformation. While India can benefit from China's proactive stance in AI governance, it must ensure that the regulatory framework promotes innovation while protecting civil liberties and individual rights.

Given the increasing integration of AI across critical sectors in India, there is a clear need for a dedicated AI law. The current legal framework, including the Information Technology (IT) Act, 2000,¹⁹ and the Digital Personal Data Protection (DPDP) Act, 2023,²⁰ provides some level of oversight, but they fall short in addressing the unique challenges posed by AI. Issues such as algorithmic bias, AI accountability, and the ethical deployment of AI require specific legal provisions that are not adequately covered by existing laws. A dedicated AI law would establish clear legal definitions, accountability mechanisms, fairness principles, and safeguards tailored to AI technologies. This would foster responsible AI development while ensuring that AI systems operate in a way that respects the rights of individuals and promotes public trust in emerging technologies.

¹⁶ White House Office of Science & Tech. Policy, *Blueprint for an AI Bill of Rights: Making Automated Systems Work for the American People* (Oct. 2022), <https://www.whitehouse.gov/ostp/ai-bill-of-rights>.

¹⁷ Provisions on the Administration of Deep Synthesis of Internet Information Services (promulgated by Cyberspace Admin. of China, Dec. 25, 2022, effective Jan. 10, 2023) (China), https://www.cac.gov.cn/2022-12/11/c_1672292429950348.htm.

¹⁸ Personal Information Protection Law of the People's Republic of China (promulgated by the Standing Comm. Nat'l People's Cong., Aug. 20, 2021, effective Nov. 1, 2021) (China), translated in NPC Observer, <https://npcobserver.com/translated-laws/personal-information-protection-law/>.

¹⁹ Information Technology Act, No. 21 of 2000, India Code (2000)

²⁰ Digital Personal Data Protection Act, No. 22 of 2023, India Code (2023)

5.1 Establish Clear Legal Definitions and Guidelines for AI Systems

One of the fundamental barriers to effective AI regulation in India is the absence of clear and standardized legal definitions related to artificial intelligence. Unlike jurisdictions such as the European Union, which offers a structured classification of AI systems based on risk levels under its AI Act, Indian law currently lacks precise terminology. Key concepts like "automated decision-making," "algorithmic transparency," "AI bias," and "explainability" remain undefined, leaving significant ambiguity in legal interpretation and enforcement.

This lack of clarity hampers the ability of regulators, courts, and stakeholders to assess accountability, monitor AI deployment, or establish industry-specific compliance standards. Therefore, a dedicated AI law in India must begin by establishing clear legal definitions and conceptual frameworks for AI systems and their components. It should also provide regulatory thresholds for high-risk applications, particularly in sensitive sectors such as healthcare, finance, and policing. Codifying these terms will serve as the foundational step toward building a robust, enforceable, and future-ready AI governance regime.

5.2 Ensure Accountability and Liability in AI Decision-Making

AI-driven decision-making often lacks transparency, and when harm occurs due to such decisions, the issue of accountability becomes complex and unresolved. In India's current legal framework, there is no clear delineation of responsibility in cases where AI systems produce biased, erroneous, or harmful outcomes. This raises critical concerns, especially in high-stakes domains like lending, recruitment, and healthcare, where the consequences of AI errors can severely impact individual rights and well-being.

Instances such as biased loan rejections, discriminatory hiring practices, and AI-induced medical misdiagnoses highlight the urgent need for accountability in AI governance. To address these concerns, a dedicated AI law should clearly establish liability frameworks that define whether developers, deployers, or users are responsible for specific harms. It must also mandate human oversight in critical sectors to prevent automated systems from making unchecked decisions. Additionally, individuals affected by AI errors should be provided accessible legal remedies to seek redress. Such measures will ensure ethical AI deployment while safeguarding public trust and individual rights.

5.3 Mandate Algorithmic Transparency and Fairness to Prevent Bias

AI models often lack transparency and have the potential to unintentionally reinforce existing societal biases, particularly when trained on historical or non-representative data. In India, there is no legal requirement for AI systems to be explainable, which makes it challenging to uncover and correct unfair or discriminatory outcomes. This lack of transparency undermines accountability and poses significant risks to fundamental rights,

especially when AI is used in critical areas like employment, finance, and law enforcement.

There have been notable instances where AI has produced biased outcomes in India such as recruitment tools that favor certain demographics over others, credit-scoring algorithms that disadvantage economically weaker applicants, and predictive policing systems that disproportionately target marginalized communities. To address these issues, a dedicated AI law should mandate algorithmic transparency, ensuring that the logic behind AI decisions is explainable and open to scrutiny. It should also require regular fairness assessments and enforce non-discrimination standards across sectors using AI. Such measures will not only enhance public trust in AI but also uphold principles of justice and equality in the digital age.

5.4 Provide Ethical AI Standards for Developers and Deployers

Currently, India lacks a formal AI ethics framework to guide the responsible development and deployment of artificial intelligence technologies. This regulatory vacuum creates risks of AI systems being designed or used in ways that violate fundamental rights, reinforce social biases, or operate without transparency. Ethical AI principles are critical to ensuring that AI remains human-centric and aligned with democratic values. Without such a framework, there is a heightened risk of irresponsible AI deployment causing harm, particularly in sensitive sectors like healthcare, policing, and governance.²¹

To address this gap, a dedicated AI law in India should introduce clear ethical guidelines rooted in principles such as privacy protection, fairness, accountability, and sustainability. This includes mandating that AI systems are trained on diverse datasets to mitigate bias, designed to produce explainable decisions, and deployed only after thorough impact assessments. Moreover, ethical AI governance should involve meaningful public participation and independent expert oversight to ensure that AI innovations serve the broader public interest. By embedding ethical standards into law, India can guide AI development in a way that balances innovation with societal well-being.

5.5 Regulate AI in Surveillance to Protect Fundamental Rights

India's increasing use of AI-powered facial recognition, predictive policing, and mass surveillance tools has outpaced the development of legal frameworks designed to regulate such technologies. These deployments, often carried out without adequate oversight or transparency, raise pressing concerns about privacy violations, misuse of biometric data—such as Aadhaar-linked facial recognition—and the erosion of democratic freedoms. Notable examples include the Delhi Police's facial recognition system, which has been criticized for enabling mass surveillance, and AI-based crowd monitoring at protests, which risks political profiling. Similarly, the use of AI in Aadhaar authentication has triggered alarms regarding data security and the lack of transparency in automated decision-making processes.

²¹ NITI Aayog, *Responsible AI for All: Part 1 – Principles for Responsible AI* (June 2020),

<https://www.niti.gov.in/sites/default/files/2020-07/Responsible-AI-22072020.pdf>.

To address these challenges, a dedicated AI law in India must introduce robust legal safeguards. This includes requiring judicial oversight before deploying AI surveillance tools, mandating transparency in the use of AI by law enforcement and public agencies, and protecting citizens from unchecked biometric data collection and surveillance. By ensuring that AI is used ethically and lawfully in surveillance, India can uphold constitutional values and protect civil liberties while still benefiting from technological advancements.

India urgently needs a standalone AI law that fills the gaps left by existing sectoral regulations such as the IT Act, DPDP Act, and RBI guidelines. This dedicated legislation should define AI-specific legal terms, establish accountability in AI-driven decision-making, mandate fairness and transparency in AI systems, and regulate AI surveillance to safeguard privacy and fundamental rights. A well-structured AI legal framework will not only promote responsible innovation but also ensure that AI development aligns with democratic principles and the public interest.

6. Recommendations for a Comprehensive AI Law in India

To effectively regulate artificial intelligence (AI) while fostering innovation, India requires a comprehensive AI legal framework that balances risk, accountability, fairness, and transparency. Based on global best practices and India's unique socio-legal landscape, the following key recommendations should be considered for a dedicated AI law in India.

First, India should adopt a risk-based classification system that categorizes AI applications into unacceptable, high-risk, limited-risk, and minimal-risk categories. This approach ensures that AI systems are regulated based on their potential to cause harm, enabling the government to impose stringent requirements on high-risk systems such as those used in hiring, credit scoring, and predictive policing while encouraging innovation in low-risk applications.

Second, the law must mandate AI transparency, explainability, and accountability. Users should have the right to receive understandable explanations of AI-driven decisions that impact their lives. Legal obligations must be placed on developers and deployers to conduct pre-deployment audits, periodic reviews, and bias assessments, especially for systems affecting public services and marginalized communities. These measures will build trust, prevent AI-induced discrimination, and safeguard fundamental rights in the digital era.²²

6.1 AI Risk Classification: Adopting a Risk-Based Regulatory Approach

India's AI governance strategy must be proactive, nuanced, and aligned with global best practices. One of the most effective models to achieve this is through a risk-based classification system, inspired by the European Union's AI Act. This system categorizes AI systems based on their potential impact on human rights, safety, and societal well-being, thereby allowing for tailored regulatory responses instead of a one-size-fits-all approach.

AI systems posing "unacceptable risk" such as social scoring mechanisms, biometric surveillance for mass control, manipulative political campaigning, and AI-driven racial profiling should be outrightly banned. These systems threaten fundamental rights, including privacy, dignity, equality, and democratic participation. Prohibiting such technologies ensures that AI is not misused to undermine civil liberties or facilitate state overreach.

High-risk AI applications, including AI tools used in hiring, credit scoring, predictive policing, healthcare diagnostics, and education, should be subjected to strict regulatory oversight. These systems directly influence people's lives and livelihoods and thus must comply with mandatory requirements such as impact assessments, bias audits, human oversight mechanisms, and compliance reporting. Developers and deployers should be held accountable for ensuring accuracy, fairness, and reliability to avoid discrimination or harm.

On the other hand, limited-risk AI systems, like chatbots, recommendation engines, or customer service automation tools, should be governed through transparency obligations, such as disclosing that users are interacting with AI. While these tools are less likely to cause serious harm, there is still a risk of misinformation or manipulation, making disclosure and explainability essential.

Lastly, minimal-risk AI applications such as AI-powered grammar checkers, smart filters for emails, and entertainment-based AI (e.g., music recommendations or gaming NPCs) can be subject to light-touch regulation. These technologies pose negligible risk to users and should be encouraged, as they foster digital innovation and productivity.

By adopting this tiered regulatory approach, India can effectively prioritize its oversight resources, focusing on monitoring and regulating high-risk AI while encouraging the safe development and deployment of low-risk applications. This strategy ensures the protection of fundamental rights without stifling technological progress, helping India become both a global AI leader and a defender of ethical AI practices.

6.2 AI Impact Assessments: Mandating Audits for High-Risk AI Applications

High-risk AI applications, such as those used in hiring, healthcare, and policing, should undergo mandatory AI impact assessments (AIAs) to address potential risks associated with bias, fairness, human rights, and the reliability of AI-driven decisions. These assessments would ensure that AI systems deployed in these sensitive areas do not reinforce discrimination or produce harmful outcomes. Pre-deployment audits would evaluate the readiness and ethical implications of AI systems, while periodic reviews would monitor their performance over time to ensure ongoing fairness and accuracy.²³

In addition, independent oversight from regulatory bodies is essential to ensure that AI systems are continually assessed by unbiased external entities. This would prevent companies from self-regulating and help ensure transparency in AI deployment. By implementing these AI audit requirements, India can

²² NITI Aayog, *Responsible AI for All: Part 2 – Operationalizing Principles for Responsible AI* (Feb. 2021), https://www.niti.gov.in/sites/default/files/2021-02/Operationalizing-ResponsibleAI_27Feb.pdf.

²³ OECD, *OECD Framework for the Classification of AI Systems* (Feb. 2022), <https://www.oecd.org/publications/oecd-framework-for-the-classification-of-ai-systems-cb6d9eca-en.htm>.

mitigate the risks of biased decision-making in key sectors, hold AI developers accountable for their creations, and build public trust in AI systems, ultimately ensuring a fairer, more ethical approach to AI technology.

6.3 Algorithmic Fairness Standards: Bias Detection and Mitigation Regulations

AI models often reinforce systemic biases, particularly in sensitive areas like hiring, credit approvals, and predictive policing. To address this, India's AI law should mandate the implementation of bias detection mechanisms in AI systems before deployment, ensuring that these technologies do not inadvertently amplify existing societal inequalities. Additionally, non-discrimination requirements must be enforced to prevent AI from perpetuating unfair advantages or disadvantages based on gender, region, socioeconomic status, or other factors. Transparency obligations should be introduced, compelling companies to disclose how their AI models make decisions and ensure that these decisions are free from biases.

In India, examples of AI bias are already emerging. AI-driven hiring tools have been found to favour certain genders, regions, or socioeconomic backgrounds, while lending algorithms have been criticized for denying loans to marginalized groups due to biased data. Predictive policing AI has raised concerns over disproportionately targeting specific communities, leading to over-policing and reinforcing existing systemic issues in law enforcement.

By addressing these concerns, India can prevent AI discrimination and ensure that AI-driven decisions are fair, transparent, and ethical. This will promote equal access to opportunities in areas such as finance, employment, and public services, ensuring that AI technologies contribute positively to society and are aligned with the principles of equity and justice.

6.4 AI Liability Framework: Clearly Defining Legal Responsibility for AI Failures

Currently, the lack of clarity regarding who is responsible for erroneous or harmful decisions made by AI systems creates significant legal uncertainty. A dedicated AI law should address this by clearly defining AI liability, determining whether responsibility falls on the developer, employer, or user. It should also introduce AI product liability laws, ensuring accountability when an AI system, such as a medical diagnostic tool, causes harm or provides incorrect results. Additionally, the law must establish legal remedies for victims of AI harm, granting individuals the right to challenge AI-driven decisions and seek redress.²⁴

This framework will prevent AI developers from evading responsibility for the outcomes of their systems, ensuring that accountability is maintained at all stages of AI deployment and use. It will also protect individuals' rights by providing them with legal recourse when AI decisions cause harm. By defining clear liability and legal protections, India can encourage safer and more ethical AI innovation, as companies will be held

accountable for their AI systems' failures, promoting the development of more reliable and responsible technologies.

6.5 AI Ethics and Transparency Mandates: Establishing AI Explainability Requirements

India should mandate that AI systems be explainable and accountable, ensuring individuals can understand the rationale behind AI-driven decisions. To achieve this, the proposed AI transparency requirements should include clear explainability obligations, where AI systems must provide understandable reasons for their decisions in simple, accessible language. This will empower users to comprehend how and why certain decisions, such as loan approvals, job rejections, or medical treatment recommendations, are made by AI systems. Additionally, individuals should have the legal right to demand an explanation when an AI system impacts their personal life, ensuring accountability and fairness in AI-driven decision-making.

Another critical aspect of transparency is the requirement for AI labeling. AI-generated content, such as deepfakes or synthetic media, should be clearly labeled as such to prevent the spread of misinformation. This will help individuals distinguish between real and AI-generated content, promoting trust and credibility in media consumption. By making AI-generated content transparent, users can better understand the source and authenticity of the information they encounter.²⁵

These transparency measures are essential to prevent AI systems from becoming "black boxes" that make unaccountable decisions. By ensuring AI explainability, especially in high-risk applications like healthcare, finance, and law enforcement, India can create a system where individuals are not only protected but also empowered to challenge AI decisions when necessary. This approach will foster trust in AI technologies and promote ethical, responsible AI development in India.

6.6 Public Awareness and AI Education: Promoting AI Literacy and Legal Awareness

A dedicated AI law in India should prioritize promoting AI literacy and ensuring public awareness about the potential risks and rights associated with AI technologies. Many individuals and businesses are unaware of the implications AI systems can have on their lives, from automated decisions in hiring and lending to privacy violations. To address this, the law should include provisions for widespread AI literacy initiatives, such as public awareness campaigns that educate citizens on the benefits and risks of AI. Additionally, AI ethics training should be mandated for developers and policymakers to ensure that ethical considerations are at the forefront of AI development and regulation.

Introducing education on AI rights is crucial to empowering individuals with the knowledge of their legal protections against discrimination, bias, and other harmful AI-driven outcomes. This education will equip people to understand their rights when interacting with AI systems, enabling them to challenge unfair decisions and demand transparency. By ensuring that citizens

²⁴ European Commission, *Proposal for a Directive on Liability for Artificial Intelligence*, COM (2022) 496 final (Sept. 28, 2022), <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:52022PC0496>.

²⁵ UNESCO, *Guidance for Regulating Digital Platforms: A Multistakeholder Approach to Safeguarding Freedom of Expression and Access to Information* 48–50 (2023), <https://unesdoc.unesco.org/ark:/48223/pf0000385654>.

are well-informed about AI and their rights, India can foster a more engaged and responsible public who can actively participate in shaping AI governance.

Strengthening AI literacy is essential for fostering a balanced and fair relationship between citizens, developers, and AI technologies. Informed individuals will be better equipped to make decisions when engaging with AI systems, and developers will be encouraged to adhere to ethical guidelines that prioritize fairness, transparency, and accountability. Ultimately, this will create a more inclusive and ethical AI ecosystem, while also empowering the public to take part in the broader conversation about the future of AI in India.

7. Conclusion

India's current legal framework is inadequate to tackle the multifaceted challenges posed by artificial intelligence (AI). As AI technologies are increasingly integrated into critical sectors like healthcare, finance, governance, and surveillance, concerns about privacy, algorithmic bias, transparency, and accountability are growing. While laws like the Information Technology (IT) Act, 2000, and the Digital Personal Data Protection (DPDP) Act, 2023, provide some regulatory oversight, they fail to address AI-specific issues such as decision-making, liability, and bias. This gap in regulation, coupled with a lack of clear AI-related definitions and inconsistencies across sectors, underscores the urgent need for a dedicated AI law in India.

A comprehensive AI legal framework for India should include clear definitions and principles specific to AI, establish accountability mechanisms for developers, deployers, and users, and enforce mandates for algorithmic transparency and fairness to reduce bias. Additionally, ethical guidelines for AI development and deployment, as well as legal safeguards against AI-driven mass surveillance and violations of fundamental rights, are essential. By drawing on global regulatory models such as the EU AI Act, the U.S. Algorithmic Accountability Act, and China's AI governance policies, India can create a framework that aligns with its unique digital landscape. The goal would be to strike a balance between fostering technological innovation and ensuring the protection of privacy, fairness, and human rights.²⁶

Implementing a dedicated AI law would address existing regulatory gaps, promote ethical AI development, and enable India to harness AI's potential for economic growth while protecting its citizens from the risks associated with AI-driven technologies. This approach will ensure that India remains competitive in the global AI arena while safeguarding fundamental rights and advancing responsible AI practices.

²⁶ Vidushi Marda, *Artificial Intelligence Policy in India: A Framework for Engaging the Limits of Data-Driven Decision-*
643

Making, 10 *Philos. & Tech.* 1, 4-8 (2021),
<https://doi.org/10.1007/s13347-021-00460-6>.

Afr. J. Biomed. Res. Vol. 27, No. 2 (May) 2024

Dr. Keval Govardhan Ukey et al.