



ISSN: 2583-7753

LAWFOYER INTERNATIONAL JOURNAL OF DOCTRINAL LEGAL RESEARCH

[ISSN: 2583-7753]



Volume IV | Special Issue I

2026

Special Issue on the International E-Conference on

“Law in the Age of Artificial Intelligence: A Fundamental Perspective”

Organised by KES’ Shri J. H. Patel College of Law, Mumbai, on 24 January 2026

Official Publication Partner: LawFoyer International Journal of Doctrinal Legal Research

DOI: <https://doi.org/10.70183/lijdlr.2026.v4si1.38>

© 2026 LawFoyer International Journal of Doctrinal Legal Research

Follow this and additional research works at: www.lijdlr.com

Under the Platform of LawFoyer – www.lawfoyer.in

After careful consideration, the editorial board of LawFoyer International Journal of Doctrinal Legal Research has decided to publish this submission as part of the publication.

In case of any suggestions or complaints, kindly contact (info.lijdlr@gmail.com)

To submit your Manuscript for Publication in the LawFoyer International Journal of Doctrinal Legal Research, To submit your Manuscript [Click here](#)

KES' SHRI. JAYANTILAL H. PATEL LAW COLLEGE

About the College



KES' Shri. Jayantilal H. Patel Law College is committed to nurturing and empowering students through quality legal education. The college provides an inclusive and academically enriching environment that encourages students to excel in their studies, participate in co-curricular activities, engage in research, and develop a strong sense of civic responsibility.

The college is affiliated to the University of Mumbai, approved by the Bar Council of India, NAAC accredited with 'B++' grade, and ISO 9001:2015 certified. It offers programmes including B.A. LL.B., LL.B., LL.M., and specialised diploma courses.

With qualified and experienced faculty, a strong academic environment, moot court training, research activities and modern infrastructure, the institution seeks to prepare responsible, skilled and socially conscious legal professionals equipped to meet the challenges of a rapidly changing world.

The college functions under the aegis of Kandivli Education Society (KES), an institution established in 1936 with a long-standing commitment to quality, inclusivity and value-based education.

Affiliation	University of Mumbai
Recognition	Approved by the Bar Council of India
Accreditation	NAAC accredited 'B++'
Certification	ISO 9001:2015 certified

KANDIVLI EDUCATION SOCIETY

Established 1936



Kandivli Education Society (KES), established in 1936, has grown from a small educational initiative into a dynamic educational institution serving thousands of students.

With a legacy of quality, inclusivity and value-based learning, KES has expanded its academic presence across various disciplines. Its educational institutions are committed to providing meaningful learning opportunities and fostering academic, professional and personal development among students.

KES' Shri. Jayantilal H. Patel Law College is one of its prominent institutions dedicated to legal education. Through academic excellence, research, practical learning and co-curricular engagement, the college seeks to empower students with the knowledge, skills and values necessary to contribute meaningfully to society.

EDUCATIONAL COMMITMENT

- Quality and inclusive education
- Value-based learning
- Academic and professional development
- Research and practical learning
- Holistic and future-ready education

LAW IN THE AGE OF ARTIFICIAL INTELLIGENCE

A Fundamental Perspective



ABOUT THE INTERNATIONAL CONFERENCE

The International Conference on “Law in the Age of Artificial Intelligence: A Fundamental Perspective” provided an academic platform to examine the transformative impact of Artificial Intelligence on law, legal systems, governance and society.

The conference explored emerging legal and ethical questions relating to Artificial Intelligence, including intellectual property, authorship and ownership, liability, data privacy, cybersecurity, criminal justice, evidence law, human rights, ethics and accountability. It addressed the implications of Artificial Intelligence for legal education, commercial law, regulatory frameworks, policy development, research and scholarly communication.

The conference brought together academicians, research scholars, legal professionals and students to present research, exchange ideas and engage in interdisciplinary discussions. It encouraged meaningful academic dialogue and contributed towards the development of effective, ethical and inclusive legal responses to the rapidly evolving Artificial Intelligence landscape.

Conference	International Conference
Theme	Law in the Age of Artificial Intelligence: A Fundamental Perspective
Date	24 January 2026
Time	10:00 AM onwards
Mode	Online
Organised by	IQAC, Student Research Society and Library Committee
Publication Partner	LawFoyer International Journal of Doctrinal Legal Research [ISSN: 2583-7753]

PRINCIPAL'S MESSAGE

KES' Shri. Jayantilal H. Patel Law College



It gives me immense pleasure to welcome all academicians, research scholars, legal professionals and students to the International Conference on “Law in the Age of Artificial Intelligence: A Fundamental Perspective.”

Artificial Intelligence is transforming the way societies function and is increasingly influencing legal systems, governance, professional practices and education. These developments bring significant opportunities as well as complex legal, ethical and social challenges.

This conference provided a valuable platform for intellectual exchange and interdisciplinary dialogue on emerging issues such as data privacy, intellectual property, liability, criminal justice, human rights, ethics and the governance of Artificial Intelligence. The scholarly contributions and discussions during the conference encouraged deeper understanding and promoted responsible approaches to the use of Artificial Intelligence.

I extend my best wishes to all the speakers, researchers, presenters, participants and members of the organising team for a meaningful and intellectually enriching conference.

DR. VIRAL DAVE

I/c. Principal

KES' Shri. Jayantilal H. Patel Law College

ORGANISING COMMITTEE

International Conference on Law in the Age of Artificial Intelligence



I/c. Principal

Dr. Viral Dave

Convenor

Mr. Mahavir Mandot

Co-Convenors

Dr. Shweta Vijendra Pathak

Ms. Priyanka Khakhadia

Ms. Deepanshi Chandarana

Organised by IQAC, Student Research Society and Library Committee

KES' Shri. Jayantilal H. Patel Law College, Mumbai

*Papers presented at the conference are published in the Special Issue of the LawFoyer International Journal of
Doctrinal Legal Research [Volume IV, Special Issue I, 2026]*

REGULATORY FRAMEWORK AND POLICY DEVELOPMENTS FOR ARTIFICIAL INTELLIGENCE

Veronica Gabriel Fernandes¹

I. ABSTRACT

Artificial Intelligence (AI) has become an integral part of modern society, influencing decision-making in sectors such as healthcare, banking, law enforcement, education, and governance. The integration of AI into these domains has led to improved efficiency, enhanced accuracy, and greater innovation. While AI offers efficiency and innovation, its unchecked use may result in serious legal and ethical challenges, including privacy violations, discrimination, lack of transparency, and absence of accountability. This research paper examines the regulatory frameworks and recent policy developments governing Artificial Intelligence at international and national levels. It analyses significant legal instruments such as the European Union's Artificial Intelligence Act, policy-based approaches adopted by the United States, regulatory measures in China, and the evolving Indian legal position. The paper also discusses key concerns relating to data protection, transparency, ethical AI, and liability. The paper highlights the need for balanced and adaptive regulation that safeguards fundamental rights while promoting innovation. The paper further explores key challenges associated with AI regulation, such as data protection and privacy safeguards, transparency and explainability of algorithms, ethical deployment of AI systems, and the determination of liability for harm caused by AI-driven decisions. By employing doctrinal legal research and drawing upon constitutional principles and judicial precedents, the study emphasises the necessity of a balanced, flexible, and human-centric regulatory framework. It ultimately advocates a harmonised, risk-based approach to AI governance that effectively safeguards fundamental rights while fostering responsible innovation and technological progress. The study concludes by recommending a risk-based, human-centric, and harmonized approach to AI governance.

¹ LLM, Second Year ,KES Shri. Jayantilal H. Patel Law College (India).Email: veronica1987ferns@gmail.com

II. KEYWORDS:

Artificial Intelligence Regulation, EU AI Act, Data Protection, Risk-Based Governance, Fundamental Rights

III. INTRODUCTION AND RESEARCH PROBLEM

One of the most important technological advances of the twenty-first century is Artificial Intelligence (AI), which is changing how people, corporations, and governments use technology. Artificial Intelligence (AI) is the term used to describe computer-based systems that are able to carry out tasks like data-driven learning, decision-making, problem-solving, language processing, and pattern recognition that have historically required human intelligence. AI has been deeply integrated in industries like healthcare, finance, education, law enforcement, public administration, and digital governance due to the rapid developments in machine learning, big data analytics, and automation. Credit scoring, recruitment procedures, medical diagnosis, predictive policing, facial recognition, and public service delivery are all increasingly using AI-driven systems, which has an impact on decisions that directly impact social structures and individual rights.

Artificial Intelligence is now extensively integrated into various aspects of everyday life across multiple sectors. In the banking sector, AI-driven systems are increasingly employed to assess creditworthiness and determine loan eligibility through predictive analytics and data-based risk assessment models. In the corporate sector, organizations utilize AI-based recruitment tools to screen and shortlist job candidates by analysing resumes and evaluating suitability based on predefined criteria. In the healthcare sector, AI technologies assist medical professionals in early disease detection, diagnosis, and treatment planning through advanced data analysis and pattern recognition. Similarly, governments are adopting AI-driven systems for crime prediction, public service delivery, and surveillance, thereby enhancing administrative efficiency and decision-making processes.

Even while Artificial Intelligence has a lot to offer in terms of productivity, creativity, and economic expansion, its broad application has raised major ethical, legal, and

regulatory issues. It might be challenging to comprehend how certain decisions are made because AI systems frequently operate through automated decision-making procedures that lack transparency and human control. AI systems that exploit big datasets create risk of discrimination, algorithmic bias, privacy concerns, and misuse of personal data. Predictive policing tools may disproportionately target vulnerable communities, AI-based recruitment tools may duplicate past bias in employment practices, and surveillance technologies may violate people's privacy and dignity. These advancements show that AI is a legal and constitutional issue that calls for cautious governmental action rather than just a technological advancement.

Globally, various jurisdictions throughout the world have taken diverse stances when it comes to regulating artificial intelligence. The Artificial Intelligence Act, which categorizes AI systems according to the degree of risk they pose to society and places stringent requirements on high-risk systems, is a comprehensive risk-based legislative framework created by the European Union. In order to promote innovation while controlling risks, the United States follows a sectoral, policy-driven strategy that depends on current legislation, executive orders, and voluntary regulatory guidelines to encourage innovation while managing risks. China has implemented a centralized regulatory structure that prioritizes national security, state control, and content restriction in AI systems. India, on the other hand is currently relying on existing laws like the Information Technology Act of 2000 and the Digital Personal Data Protection Act, 2023, while it is still in the process of creating a structured regulatory framework.

Despite increased global attempts to regulate Artificial intelligence, a clear and uniform legal framework is still lacking in many jurisdictions, especially in developing nations like India. The rapid pace of technological advancement has created a gap between existing legal frameworks and the emerging challenges posed by AI-driven decision-making. Issues like algorithmic accountability, automated liability, cross-border data governance, and the ethical use of AI systems were not intended to be addressed by current legislation. As a result, there is uncertainty regarding the protection of fundamental rights, determination of liability for AI-related harm, and the extent of regulatory oversight required to ensure responsible innovation.

A. Research Problem

The research problem of this study is the absence of a comprehensive and coherent legal framework governing Artificial Intelligence in India, particularly in relation to accountability, liability, data protection, and protection of fundamental rights. While the global models such as the European Union's risk-based approach and the United States' policy-driven framework offer useful regulatory directions, India continues to rely on fragmented legal provisions and non-binding policy initiatives. This creates regulatory uncertainty and leaves significant gaps in addressing issues such as algorithmic bias, transparency, liability for AI-related harm, and constitutional safeguards.

The study therefore aims to examine whether India needs a specific legislative framework for Artificial Intelligence and how international regulatory models can guide the development of an equitable and human-centric AI governance regime. It also aims to identify the legal gaps in the existing Indian framework and propose a regulatory approach that safeguards fundamental rights while promoting innovation and technological development.

B. Review of Literature

The growing influence of Artificial Intelligence on governance, economy, and society has attracted significant scholarly attention across legal, technological, and policy-oriented disciplines. Several scholars and international organizations have examined the ethical, legal, and regulatory implications of AI, particularly in relation to fundamental rights, data protection, and accountability. The existing literature highlights the need for a balanced regulatory framework that promotes innovation while safeguarding human rights and democratic values.

Virginia Eubanks, in her work *Automating Inequality* (2018), critically examines how automated decision-making systems in public administration can reinforce social and economic inequalities. She argues that algorithmic governance, especially in welfare distribution, policing, and public services, often reproduces existing structural biases and disproportionately affects marginalized communities. Her work emphasizes that AI systems must be regulated to prevent discrimination and ensure fairness in decision-making processes.

Kate Crawford, in *Atlas of AI* (2021), explores the social and political dimensions of Artificial Intelligence and highlights the hidden power structures behind AI systems. She argues that AI is not merely a technological tool but a system shaped by economic interests, data extraction practices, and institutional control. Crawford stresses the need for strong regulatory mechanisms to ensure transparency, accountability, and ethical governance of AI systems, particularly in democratic societies.

International organizations have also contributed significantly to the development of AI governance principles. The Organisation for Economic Co-operation and Development (OECD) adopted the OECD Principles on Artificial Intelligence (2019), which emphasize inclusive growth, human-centered values, transparency, robustness, and accountability in AI systems. These principles have influenced global regulatory discussions and have been adopted by several countries as guiding standards for responsible AI governance.

1. OECD – Key AI Reports and Publications:

- OECD Principles on Artificial Intelligence (2019)

The OECD AI Principles are the first intergovernmental standard promoting trustworthy, human-centric AI, transparency, accountability, and human rights protection.

They provide policy recommendations and ethical guidelines for governments and AI developers and have influenced global AI governance frameworks. Widely adopted by OECD member and partner countries as a foundation for responsible AI regulation.²

- State of Implementation of the OECD AI Principles (2021)

This report evaluates how countries are implementing OECD AI principles in national AI policies. It provides guidance for policymakers, classification frameworks for AI

² Organisation for Economic Co-operation and Development (OECD), *OECD Principles on Artificial Intelligence* (OECD 2019) <https://www.oecd.org/ai/principles/> accessed 10 January 2026.

systems, and implementation strategies. Focuses on practical governance mechanisms and policy coordination among states.³

- OECD AI Policy Observatory (OECD.AI)

A global platform providing data, research, policy tools, and best practices on AI governance. Supports governments in developing evidence-based AI regulations and policy frameworks. Helps coordinate international cooperation on AI governance.⁴

- OECD Due Diligence Guidance for Responsible AI (2024–2026)

Provides practical guidance for implementing responsible AI in public and private sectors. Focuses on risk management, transparency, and accountability in AI development.⁵

Similarly, the United Nations Educational, Scientific and Cultural Organization (UNESCO) issued the Recommendation on the Ethics of Artificial Intelligence (2021), which provides a comprehensive global framework for ethical AI governance. The UNESCO Recommendation stresses the importance of human rights, fairness, environmental sustainability, data protection, and accountability in AI systems. It encourages member states to develop legal and institutional frameworks that ensure AI is used in a manner consistent with democratic and ethical values.

The UNESCO framework encourages Member States to develop legal and institutional mechanisms to ensure ethical and socially responsible AI development. While the Recommendation was adopted by 193 Member States in November 2021, UNESCO's current references to 194 Member States relate to its present membership and should not be conflated with the historical adoption figure.⁵

2. UNESCO – Key AI Reports and Publications

- UNESCO Recommendation on the Ethics of Artificial Intelligence (2021)

³ OECD, *State of Implementation of the OECD AI Principles: Insights from National AI Policies* (OECD 2021).

⁴ OECD, *OECD AI Policy Observatory* (OECD) <https://oecd.ai> accessed 10 January 2026.

⁵ OECD, *OECD Due Diligence Guidance for Responsible AI* (OECD Publishing 2026) <https://doi.org/10.1787/41671712-en>.

First global normative framework adopted by 194 member states. Focuses on human rights, fairness, transparency, accountability, environmental sustainability and human oversight. Provides policy action areas for governments and institutions. Aims to ensure AI benefits humanity and democratic values.⁶

- UNESCO Global AI Ethics and Governance Observatory

Provides policy tools, research, and global best practices for ethical AI governance. Supports policymakers, researchers, and regulators in developing responsible AI frameworks. Encourages international cooperation and ethical implementation of AI.⁷

- UNESCO AI Ethics Policy Action Areas Framework

Provides implementation guidance on data governance, gender equality, environment, health, and education. Helps translate ethical principles into legal and institutional frameworks. Focuses on human rights-centered AI governance.⁸

Scholarly research on the European Union's Artificial Intelligence Act highlights it as one of the most comprehensive legal frameworks for AI regulation. Researchers have praised the EU's risk-based regulatory model for providing clear obligations for high-risk AI systems and banning harmful applications such as social scoring and manipulative AI practices. At the same time, some scholars argue that strict regulation may increase compliance costs and slow down innovation, indicating the need for a balanced regulatory approach.

Studies on the United States' AI governance model emphasize its flexible and innovation-oriented approach, which relies on sector-specific regulation and executive policy measures rather than a single comprehensive law. While this model promotes technological growth and private sector participation, scholars note that it lacks uniform legal accountability and may lead to fragmented regulatory oversight.

⁶ UNESCO, *Recommendation on the Ethics of Artificial Intelligence* (UNESCO 2021) <https://unesdoc.unesco.org> accessed 10 January 2026.

⁷ UNESCO, *Global AI Ethics and Governance Observatory* (UNESCO) <https://www.unesco.org/artificial-intelligence> accessed 10 January 2026.

⁸ UNESCO, *Ethics of Artificial Intelligence: Policy Action Areas* (UNESCO 2021).

In the Indian context, there is an absence of a dedicated statutory framework for Artificial Intelligence and the need to integrate constitutional principles such as the right to privacy, equality, and due process into AI governance. The recognition of privacy as a fundamental right by the Supreme Court has strengthened the argument for stronger data protection and AI regulation in India. Scholars have also emphasized the importance of developing a risk-based and human-centric regulatory model that aligns with international standards while addressing India's socio-economic and legal realities.

Overall, the existing literature suggests the need for responsible AI governance supported by legal accountability, ethical standards, and institutional oversight. However, there remains a significant research gap in analysing how international regulatory models can be effectively adapted to the Indian legal framework, particularly in relation to liability, constitutional safeguards, and enforcement mechanisms. This study seeks to address this gap by critically examining global AI regulatory approaches and proposing a structured legal framework suitable for India.

C. Objectives and Scope of the Study

The present research aims to examine the evolving regulatory framework governing Artificial Intelligence and to analyse the legal and policy developments at both international and national levels. The study focuses on understanding the need for AI regulation, evaluating comparative regulatory approaches, and identifying gaps in the Indian legal framework in order to propose a balanced and effective governance model.

The primary objectives of this study are as follows:

1. To examine the concept and growing significance of Artificial Intelligence in modern governance and society.
2. To analyse the necessity of regulating Artificial Intelligence in order to protect fundamental rights, ensure transparency, and maintain accountability.

3. To study and compare international regulatory approaches to AI, particularly those adopted by the European Union, the United States, and China.
4. To evaluate the existing legal and policy framework governing Artificial Intelligence in India.
5. To identify legal lacunae in the Indian regulatory system, particularly in relation to data protection, liability, accountability, and ethical governance.
6. To propose a structured and human-centric regulatory framework for Artificial Intelligence in India based on international best practices.

This research primarily focuses on the legal and policy aspects of Artificial Intelligence regulation. The study examines statutory laws, policy documents, judicial decisions, and international regulatory instruments relating to AI governance. It includes a comparative analysis of global regulatory models, particularly those of the European Union, the United States, and China, to understand different approaches to AI regulation.

The research also focuses specifically on issues such as fundamental rights protection, data governance, ethical AI deployment, and liability for AI-related harm within the Indian legal framework. By concentrating on these aspects, the study seeks to provide practical and legally grounded recommendations for developing a comprehensive AI regulatory framework in India.

D. Hypothesis and Research Questions

1. Hypothesis

The present study is based on the hypothesis that the absence of a comprehensive statutory framework governing Artificial Intelligence in India creates significant legal and regulatory gaps in areas such as accountability, liability, data protection, and protection of fundamental rights, and that adopting a structured risk-based and human-centric regulatory framework inspired by international models, particularly the European Union's Artificial Intelligence Act, can ensure responsible innovation while safeguarding constitutional and human rights.

2. Research Questions

The research seeks to address the following questions:

1. Why is it necessary to regulate Artificial Intelligence in the modern legal and technological environment?
2. How have different jurisdictions, particularly the European Union, the United States, and China, approached the regulation of Artificial Intelligence?
3. What are the existing legal and policy frameworks governing Artificial Intelligence in India, and what gaps exist within them, and what legal mechanisms can be adopted in India to determine liability and accountability for harm caused by AI?
4. How can India develop a balanced regulatory framework that protects fundamental rights while promoting innovation and technological growth?

E. Research Methodology

The present study adopts a doctrinal and analytical research methodology to examine the regulatory framework and policy developments relating to Artificial Intelligence at both international and national levels. The doctrinal approach focuses on the systematic analysis of legal principles, statutory provisions, policy documents, judicial decisions, and international regulatory instruments governing Artificial Intelligence. The research primarily relies on secondary sources of data and involves a comparative analysis of different regulatory models to identify best practices and legal lacunae in the Indian framework.

1. Nature of Research

This research is descriptive and analytical in nature. It seeks to describe the existing legal and policy framework governing Artificial Intelligence and critically analyse its effectiveness in addressing issues such as privacy, algorithmic bias, accountability, and liability for AI-related harm. The study also evaluates international regulatory approaches in order to propose suitable reforms for India.

2. Sources of Data

The research is based on both primary and secondary sources:

- **Primary Sources:** Statutory laws such as the Information Technology Act, 2000, the Digital Personal Data Protection Act, 2023, international legal instruments such as the European Union Artificial Intelligence Act, policy frameworks issued by government authorities, and relevant judicial decision of court.
- **Secondary Sources:** Books, research articles, journal publications, government reports, international organization reports, and academic commentaries on Artificial Intelligence regulation.

3. Comparative Approach

The study adopts a comparative legal approach by analysing regulatory frameworks in the European Union, the United States, and China. This comparative analysis helps in understanding different regulatory models, including risk-based regulation, policy-based governance, and state-controlled frameworks. The comparison provides valuable insights for developing a structured and effective AI regulatory framework in India.

4. Scope of Legal Analysis

The research focuses on key legal issues such as protection of fundamental rights, data protection, transparency, ethical governance, and liability for harm caused by Artificial Intelligence systems. It also examines recent policy initiatives and regulatory developments in India to understand the evolving legal landscape.

5. Limitations of the Study

The study is limited to doctrinal legal research and does not include empirical or technical analysis of Artificial Intelligence systems. It relies on publicly available legal and policy documents and focuses primarily on regulatory aspects of AI governance. Rapid technological developments may also lead to future changes in regulatory frameworks, which are beyond the scope of the present research.

F. Research Analysis

Artificial Intelligence has significantly transformed governance, economic systems, and social structures across the world, creating both opportunities and regulatory challenges. The growing use of AI in decision-making processes has raised serious

legal and constitutional concerns relating to privacy, equality, accountability, transparency, and liability. Different countries have adopted varied regulatory approaches based on their legal traditions, political priorities, and economic goals. Some jurisdictions have developed comprehensive statutory frameworks, while others rely on policy-based or administrative regulation. India, however, is still in the process of developing a structured legal regime for Artificial Intelligence and currently relies on existing laws and policy initiatives. This section critically analyses the necessity of AI regulation, comparative international regulatory models, the evolving Indian framework, liability issues, and the challenges involved in regulating Artificial Intelligence.

1. Necessity of Artificial Intelligence Regulation

The regulation of Artificial Intelligence has become essential due to its increasing influence on fundamental rights, democratic governance, and social justice. AI systems are no longer limited to technical or industrial applications; they are now used in areas such as healthcare, banking, law enforcement, surveillance, education, and public administration. As a result, AI-driven decisions directly affect individuals' rights, opportunities, and freedoms. Without proper legal oversight, AI systems may create arbitrary, discriminatory, or unsafe outcomes, making regulatory intervention necessary.

One of the primary reasons for regulating AI is the protection of fundamental rights, particularly the right to privacy. AI systems depend heavily on large datasets, which often include personal and sensitive information. Technologies such as facial recognition, biometric identification, predictive analytics, and surveillance tools collect and process personal data on a massive scale. In the absence of strict legal safeguards, such systems may violate individual privacy and dignity. The recognition of privacy as a fundamental right by the Supreme Court in *Justice K.S. Puttaswamy v. Union of India* has strengthened the constitutional basis for regulating digital technologies and Artificial Intelligence.⁹ The judgment emphasized that any state

⁹ *Justice K.S. Puttaswamy v. Union of India* (2017) 10 SCC 1.

action involving data collection must be lawful, necessary, and proportionate, thereby establishing a constitutional framework for AI regulation in India.

Another important reason for regulating AI is the prevention of discrimination and algorithmic bias under Article 14 of the Constitution. AI systems often rely on historical data to make predictions and decisions. If the data used in training these systems contain bias or inequality, the AI system may produce discriminatory outcomes. For instance, AI-based recruitment tools may unfairly favour certain groups, credit scoring algorithms may discriminate against economically weaker sections, and predictive policing systems may disproportionately target specific communities. Such outcomes violate the principle of equality before the law and equal protection of laws guaranteed under Article 14.

The principle of non-arbitrariness under Article 14 requires that all state actions must be reasonable and non-discriminatory. In *State of West Bengal v. Anwar Ali Sarkar*, the Supreme Court held that arbitrary classification and unequal treatment violate the guarantee of equality under Article 14.¹⁰ This principle is highly relevant in the context of algorithmic decision-making, as AI systems must operate in a fair and transparent manner. If automated systems produce arbitrary or biased outcomes, they may be considered unconstitutional. Therefore, regulatory frameworks must ensure that AI systems follow fairness, transparency, and accountability standards in order to protect constitutional rights.

Transparency and accountability are also crucial for building public trust in Artificial Intelligence. Many AI systems function as “black box” technologies, where the decision-making process is not easily understandable. Individuals affected by automated decisions often do not know how or why a particular decision was made. This lack of explainability creates accountability gaps and undermines trust in technology. Legal regulation is therefore necessary to ensure that AI systems are explainable, auditable, and subject to oversight by regulatory authorities.

Public safety and national security concerns further justify the need for AI regulation. AI systems are increasingly used in autonomous vehicles, medical diagnosis, military

¹⁰ *State of West Bengal v. Anwar Ali Sarkar* AIR 1952 SC 75.

technology, and critical infrastructure management. Errors or failures in such systems can lead to serious harm, including loss of life or large-scale economic damage. Regulation ensures that safety standards, testing mechanisms, and risk assessments are in place before AI systems are deployed in sensitive sectors. Thus, the necessity of AI regulation arises from the need to protect fundamental rights, ensure equality, maintain transparency, and safeguard public safety.

2. International Approaches to AI Regulation

Different countries have adopted varied regulatory approaches to Artificial Intelligence based on their legal systems and policy priorities. The European Union, the United States, and China represent three distinct models of AI governance, while international organizations such as OECD and UNESCO provide guiding principles for ethical AI regulation.

The European Union has adopted one of the most comprehensive legal frameworks for Artificial Intelligence through the Artificial Intelligence Act, which follows a risk-based regulatory approach. The Act classifies AI systems into four categories: unacceptable risk, high risk, limited risk, and minimal risk. Unacceptable-risk AI systems, such as social scoring and manipulative technologies, are prohibited. High-risk systems used in sensitive areas such as healthcare, education, employment, and law enforcement are subject to strict safety, transparency, and governance obligations. Limited-risk systems are required to comply with disclosure obligations, while minimal-risk systems are subject to limited regulatory intervention.

The Act follows a phased implementation timeline: prohibited AI practices became applicable from February 2025, and obligations relating to general-purpose AI models became applicable from August 2025. Subsequently, under the EU's AI Omnibus simplification package, politically agreed on 7 May 2026 and entered into force on 27 July 2026, the application of rules for high-risk AI systems listed in Annex III was deferred from 2 August 2026 to 2 December 2027. This revised timeline reflects the EU's attempt to preserve legal certainty and fundamental-rights protection while allowing additional time for compliance with high-risk AI obligations.

The United States follows a flexible and policy-driven approach to AI regulation. Instead of a single comprehensive law, the US relies on existing legislation, federal

agencies, and voluntary guidelines to regulate AI. Institutions such as the National Institute of Standards and Technology (NIST), Federal Trade Commission (FTC), and Equal Employment Opportunity Commission (EEOC) play important roles in regulating AI-related activities. The US approach emphasizes innovation, private sector development, and market-driven growth. However, recent developments have shown a shift in regulatory priorities. The executive order on safe and trustworthy AI issued in 2023 was revoked in January 2025, reflecting a more deregulatory stance and a greater focus on technological competitiveness. This change highlights the dynamic nature of AI governance in the United States and the ongoing debate between regulation and innovation.

A major development in recent years was President Joe Biden's Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence issued in October 2023, which aimed to establish a coordinated federal framework for AI governance. The Executive Order required safety testing of advanced AI models, transparency obligations for developers, privacy protection measures, and federal oversight of high-risk AI systems, while also promoting responsible innovation and international cooperation.¹¹ This order was widely viewed as a significant step toward structured federal governance of AI in the United States, aligning regulatory oversight with national security, consumer protection, and civil rights concerns.

However, a significant shift occurred in January 2025 when President Donald Trump issued Executive Order 14179, which revoked the 2023 Executive Order on Safe, Secure, and Trustworthy AI.¹² The revocation signalled a major policy transition toward deregulation and technological competitiveness, emphasizing reduced federal oversight and greater reliance on private sector innovation. This development reflects a broader political and economic shift in US AI governance, prioritizing economic growth, national competitiveness, and innovation over centralized regulatory control. The repeal of the earlier executive order demonstrates the dynamic and evolving

¹¹ Executive Order 14110, *Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence* (30 October 2023), The White House.

¹² Executive Order 14179, *Revocation of Certain Artificial Intelligence Executive Actions* (23 January 2025), The White House.

nature of AI governance in the United States, where regulatory approaches are closely tied to changing political priorities and administrative policies.

This policy shift has important implications for global AI governance. It highlights the absence of a stable and unified regulatory framework in the United States and underscores the ongoing debate between innovation and regulation. While the US model promotes technological advancement and flexibility, the revocation of the 2023 Executive Order indicates potential regulatory uncertainty and fragmentation in AI governance. This contrasts with the European Union's comprehensive and binding AI Act and illustrates the divergence in international regulatory approaches. The evolving US stance therefore represents a dynamic and politically influenced regulatory environment, which significantly impacts comparative analysis of global AI governance frameworks and must be considered in contemporary legal scholarship.

China has adopted a centralized and state-controlled model of AI regulation. The Chinese government regulates AI through administrative measures, content control rules, and national security laws. Regulations such as the Interim Measures for the Management of Generative AI Services require companies to register algorithms, ensure content compliance, and follow government guidelines. The Chinese model prioritizes national security, social stability, and government control over technological development. While this approach allows rapid policy implementation and strict enforcement, it raises concerns about freedom of expression, privacy, and state surveillance.

International organizations have also contributed significantly to AI governance through soft law principles and ethical guidelines. The OECD Principles on Artificial Intelligence emphasize human-centered values, transparency, accountability, and inclusive growth. Similarly, the UNESCO Recommendation on the Ethics of Artificial Intelligence promotes fairness, environmental sustainability, human rights protection, and global cooperation. These international instruments provide guiding standards that influence national regulatory frameworks and encourage countries to adopt responsible AI governance practices.

The comparative analysis of international approaches demonstrates that there is no single universal model for AI regulation. The European Union focuses on strict legal regulation, the United States emphasizes innovation and flexibility, and China prioritizes state control and national security. These models provide valuable lessons for India in developing a balanced and effective regulatory framework.

3. Artificial Intelligence Regulation in India

India currently does not have a dedicated statute specifically governing Artificial Intelligence; however, several legislative provisions, policy initiatives, and regulatory advisories indirectly regulate AI-related activities. The Indian regulatory approach is gradually evolving and reflects a cautious and policy-driven model that emphasizes digital governance, data protection, innovation, and ethical AI development. Rather than adopting a comprehensive statutory framework, India has relied on sectoral laws, executive guidelines, and institutional initiatives to address emerging AI-related challenges.

The Information Technology Act, 2000 provides the foundational legal framework for digital transactions, cybersecurity, intermediary liability, and electronic governance in India. Although enacted prior to the emergence of Artificial Intelligence technologies, certain provisions relating to data protection, intermediary liability, and cyber offences may apply to AI systems operating on digital platforms. However, the Act is widely considered outdated in the context of modern technological developments, as it does not address critical issues such as algorithmic accountability, automated decision-making, liability for AI-generated harm, and governance of autonomous systems. This limitation highlights the need for legislative modernization to accommodate emerging technological risks.

A significant development in India's digital regulatory landscape is the enactment of the Digital Personal Data Protection Act, 2023, which establishes a comprehensive framework for the lawful processing of digital personal data, consent-based data usage, and accountability of data fiduciaries. Since Artificial Intelligence systems rely heavily on large-scale data processing, the Act plays a crucial role in indirectly regulating AI by requiring lawful, transparent, and purpose-specific handling of personal data. This framework strengthens privacy protection, enhances

accountability, and creates an essential regulatory foundation for responsible AI deployment in India.

The Digital Personal Data Protection Rules, 2025, notified by the Ministry of Electronics and Information Technology in November 2025, further strengthen this framework by operationalising the Digital Personal Data Protection Act, 2023. The Rules establish the practical compliance architecture for data fiduciaries and provide an 18-month phased implementation timeline extending up to May 2027. They require data fiduciaries to issue clear and standalone consent notices, prescribe protocols for personal data breach notification, strengthen the rights of data principals to access, correct, update, erase, and nominate persons in relation to their data, and impose enhanced obligations on Significant Data Fiduciaries, including independent audits, impact assessments, and stronger due diligence for deployed technologies.

The Rules also establish the Data Protection Board of India as a digital-first adjudicatory institution for complaints and enforcement. In the context of Artificial Intelligence, these Rules are particularly important because they convert the broad data-protection principles of the DPDP Act into operational obligations governing the collection, processing, storage, and use of personal data in AI-driven systems.

Another important regulatory step was taken by the Ministry of Electronics and Information Technology (MeitY) through its AI advisories issued in March 2024, which directed intermediaries and AI platforms to exercise due diligence and prevent misuse of AI technologies. The original advisory dated 1 March 2024 required platforms to obtain prior government approval before deploying certain under-tested or unreliable AI models in India. However, this requirement was withdrawn by the revised advisory dated 15 March 2024, which superseded the earlier position and removed the prior-approval requirement. The revised advisory continued to emphasise compliance with the Information Technology Rules, 2021, including due diligence, prevention of unlawful or harmful content, user awareness, and labelling or disclosure of under-tested, unreliable, or AI-generated content. This development reflects India's cautious but evolving approach toward platform accountability and responsible AI governance, while avoiding a mandatory pre-deployment approval regime for AI models.

India has also strengthened its policy framework through the launch of the IndiaAI Mission in March 2024, with a budget allocation of approximately ₹10,300 crores. The mission aims to build a robust AI ecosystem by developing computing infrastructure, supporting research and innovation, encouraging startups, and promoting ethical and inclusive AI development. The IndiaAI Mission focuses on enhancing national AI capacity through public-private partnerships, access to datasets, and institutional support, thereby positioning India as a global leader in responsible AI innovation.

Further reinforcing this approach, and building upon emerging sector-specific initiatives such as the RBI's FREE-AI framework, the Ministry of Electronics and Information Technology released the India AI Governance Guidelines in November 2025 under the IndiaAI Mission framework. These Guidelines are based on seven guiding principles, described as "Sutras": Trust is the Foundation, People First, Innovation over Restraint, Fairness & Equity, Accountability, Understandable by Design, and Safety, Resilience & Sustainability.

The framework also operates through six governance pillars intended to support safe, trusted, and innovation-oriented AI adoption in India. The Guidelines confirm India's preference for a principle-based, voluntary, and non-legislative governance model, emphasising ethical standards, institutional coordination, and responsible innovation rather than strict statutory control. This light-touch regulatory approach reflects India's attempt to balance innovation with accountability while avoiding overregulation of the technology sector. The Guidelines are particularly relevant in shaping India's long-term AI regulatory strategy and policy direction.

In addition to executive and policy initiatives, legislative interest in AI regulation is also emerging in India. The Artificial Intelligence (Ethics and Accountability) Bill, 2025, introduced as a Private Member's Bill in the Lok Sabha in December 2025, proposes statutory oversight mechanisms, bias audits, ethical review boards, and penalties for misuse of AI systems. Although the bill has not yet been enacted, it reflects growing parliamentary concern regarding AI governance and the need for a structured legal framework. The bill highlights issues such as accountability, transparency, and ethical regulation of AI, indicating a potential future direction for comprehensive AI legislation in India.

Overall, India's regulatory framework for Artificial Intelligence remains in a developmental stage and is characterized by a combination of existing laws, policy initiatives, executive advisories, and emerging legislative proposals. While significant progress has been made through the Digital Personal Data Protection Act, MeitY AI Advisory, IndiaAI Mission, and AI Governance Guidelines, the absence of a comprehensive statutory framework creates regulatory gaps in areas such as liability, algorithmic accountability, and enforcement mechanisms. This evolving policy landscape underscores the need for a dedicated and structured AI law that can ensure responsible innovation while protecting fundamental rights and promoting technological growth in India.

In the financial sector, an important sector-specific development was the Reserve Bank of India's Framework for Responsible and Ethical Enablement of Artificial Intelligence (FREE-AI) Committee Report, released in August 2025. The report provides a responsible AI framework for regulated financial entities and seeks to balance innovation with risk mitigation in the use of AI in banking and financial services. It is particularly relevant because it articulated seven guiding Sutras for responsible AI adoption and made recommendations across key areas such as infrastructure, capacity building, policy, governance, protection, and assurance. The FREE-AI framework may therefore be viewed as an important sectoral precursor to the broader India AI Governance Guidelines issued by MeitY in November 2025, which adopted a similar Sutra-based structure for responsible AI governance.¹³

4. Liability Framework for AI-Related Harm

One of the most critical issues in AI regulation is the determination of liability when AI systems cause harm. Unlike traditional technologies, AI systems may operate with a degree of autonomy, opacity, and adaptive decision-making that makes it difficult to identify a single responsible actor. The developer, manufacturer, deployer, service provider, data fiduciary, or end-user may all contribute to the functioning of an AI system, thereby complicating the attribution of legal responsibility. In India, liability

¹³ Reserve Bank of India, *FREE-AI Committee Report: Framework for Responsible and Ethical Enablement of Artificial Intelligence in the Financial Sector* (RBI, August 2025).

for AI-related harm is presently governed by general legal principles such as negligence, strict liability, product liability, and consumer protection law.

Under tort law, negligence may apply where harm results from lack of reasonable care in designing, training, testing, deploying, or monitoring an AI system. The Consumer Protection Act, 2019 may also be relevant where AI-based products or services are defective or deficient and cause injury or loss to consumers. However, these ordinary principles may be inadequate where harm results from highly autonomous, complex, or opaque AI systems and the injured person is unable to prove fault, causation, or defect with technical precision.

A more significant Indian doctrinal reference is the principle of absolute liability developed by the Supreme Court in *M.C. Mehta v. Union of India*, AIR 1987 SC 1086¹⁴, in the context of the Oleum Gas Leak case. The Court held that an enterprise engaged in a hazardous or inherently dangerous activity owes an absolute and non-delegable duty to the community to ensure that no harm results from such activity. Unlike the traditional rule of strict liability in *Rylands v. Fletcher*, the Indian doctrine of absolute liability does not permit conventional exceptions such as act of God, third-party act, or plaintiff's own fault. This doctrine is relevant to AI governance because certain high-risk AI systems, such as those used in autonomous transportation, healthcare diagnosis, critical infrastructure, biometric surveillance, policing, or financial decision-making, may create serious risks to life, liberty, privacy, equality, and public safety.

However, the application of absolute liability to AI-related harm must be approached with caution. Artificial Intelligence is not inherently hazardous in every form, and low-risk AI applications should not attract the same standard as high-risk autonomous systems. The doctrine may therefore serve as an analogue for high-risk AI deployments rather than as a blanket rule for all AI technologies. A future AI liability framework in India may adopt a graded approach by applying negligence-based liability to low-risk systems, product-liability principles to defective AI products or services, and an absolute or no-fault liability standard for high-risk AI systems deployed in sensitive sectors. Such a framework would ensure effective

¹⁴ *M.C. Mehta v. Union of India*, AIR 1987 SC 1086; (1987) 1 SCC 395.

compensation for victims while encouraging developers and deployers to adopt strong safety, audit, testing, and human-oversight mechanisms.

Therefore, the absence of a specific AI liability framework creates significant uncertainty and accountability gaps. A clear statutory mechanism is necessary to determine responsibility among developers, deployers, service providers, data controllers, and users, and to ensure compensation for victims of AI-related harm. Incorporating Indian doctrines such as absolute liability, along with tort and consumer protection principles, would strengthen accountability and provide a more constitutionally grounded response to AI-related risks.

5. Challenges in Regulating Artificial Intelligence

Regulating Artificial Intelligence presents several challenges due to the rapid pace of technological development and the global nature of AI systems. One major challenge is the speed of technological change, as laws often become outdated quickly. Another challenge is the lack of international consensus, which creates regulatory fragmentation.

Ethical concerns, data privacy risks, and accountability issues further complicate regulation. Balancing innovation with legal safeguards remains the biggest challenge for policymakers. Therefore, a flexible and adaptive regulatory framework is necessary to effectively govern Artificial Intelligence.

IV. CONCLUSION

The foregoing analysis demonstrates that Artificial Intelligence has become an integral part of modern governance and economic systems, making regulatory intervention both necessary and inevitable. As discussed in the previous sections, AI systems significantly influence decision-making in areas such as healthcare, banking, employment, law enforcement, and public administration, thereby directly affecting fundamental rights, public safety, and democratic governance. The absence of clear legal regulation creates risks of privacy violations, algorithmic discrimination, lack of transparency, and uncertainty in determining liability for harm caused by automated decision-making systems. In this context, the need for a structured and balanced

regulatory framework becomes essential to ensure that Artificial Intelligence serves human interests while maintaining constitutional and ethical safeguards.

The first research question examined the necessity of regulating Artificial Intelligence in the modern legal environment. The study concludes that AI regulation is essential for the protection of fundamental rights, particularly the right to privacy and equality, as well as for ensuring transparency, accountability, and public safety. AI-driven decisions, if left unregulated, may result in arbitrary and discriminatory outcomes, thereby violating constitutional principles and undermining public trust in technology. Regulation is therefore required to ensure that AI systems operate in a fair, transparent, and accountable manner.

The second research question analysed how different jurisdictions regulate Artificial Intelligence. The comparative analysis shows that the European Union has adopted a comprehensive risk-based legal framework through the Artificial Intelligence Act, ensuring strict obligations for high-risk systems and prohibiting harmful AI practices. The United States follows a flexible and policy-driven approach that prioritizes innovation and market development, while China has implemented a centralized and state-controlled regulatory framework focusing on national security and social stability. These diverse approaches demonstrate that there is no single universal model of AI regulation, but they collectively provide valuable lessons for developing countries.

The third research question focused on the existing legal and policy framework governing Artificial Intelligence in India and the gaps within it. The study finds that India currently relies on fragmented laws such as the Information Technology Act, 2000, and the Digital Personal Data Protection Act, 2023, along with policy initiatives such as the IndiaAI Mission and governance guidelines. While these measures represent important steps, they do not provide a comprehensive legal framework for regulating AI, particularly in areas such as algorithmic accountability, liability, and ethical governance. This creates regulatory uncertainty and highlights the need for a dedicated statutory framework. Additionally, the third research question addressed the issue of liability for harm caused by Artificial Intelligence. The study concludes that existing legal principles such as tort law and consumer protection provisions are

insufficient to address the complexities of AI-related harm. A clear statutory liability framework is therefore necessary to determine responsibility and ensure compensation for affected individuals.

The fourth research question examined how India can balance innovation and protection of rights through effective regulation. The analysis suggests that a risk-based and human-centric regulatory framework inspired by international best practices can help achieve this balance. Such a framework would allow innovation and technological growth while ensuring constitutional safeguards and accountability mechanisms.

In conclusion, the hypothesis of the study stands validated. India requires a comprehensive and structured Artificial Intelligence regulatory framework that integrates constitutional principles, international best practices, and ethical governance standards. A balanced and adaptive legal regime will ensure that Artificial Intelligence promotes innovation, protects fundamental rights, and contributes to sustainable and responsible technological development.

V. SUGGESTIONS

1. Enactment of a Comprehensive Artificial Intelligence Law in India

India should enact a dedicated and comprehensive Artificial Intelligence law to address the regulatory gaps identified in the existing legal framework. Current laws such as the Information Technology Act, 2000, the Digital Personal Data Protection Act, 2023, and the Digital Personal Data Protection Rules, 2025 provide important but largely indirect regulation and do not adequately address issues such as algorithmic accountability, transparency, liability, risk classification, and sector-specific oversight. In this context, the Artificial Intelligence (Ethics and Accountability) Bill, 2025, introduced as a Private Member's Bill in the Lok Sabha, represents an important legislative starting point because it proposes mechanisms such as an AI Ethics Committee, bias audits, ethical oversight, and penalties for misuse of AI systems. However, if the Bill is to be considered for enactment, it should be strengthened by incorporating a clearer risk-based classification of AI systems, detailed liability rules for AI-related harm, mandatory human oversight for high-risk AI systems,

independent audit requirements, grievance-redressal mechanisms, and coordination with the DPDP Act and sectoral regulators. A separate AI law would provide legal clarity, establish uniform regulatory standards, and ensure responsible AI deployment across sectors while protecting fundamental rights and promoting innovation.

2. Adoption of a Risk-Based Regulatory Framework

India should adopt a risk-based regulatory approach similar to international best practices in order to balance innovation with safety. AI systems that pose higher risks to fundamental rights, public safety, or economic stability should be subject to stricter regulation, while low-risk systems may be allowed greater flexibility. This approach ensures proportional regulation and prevents unnecessary restrictions on technological growth. A risk-based model would also provide legal certainty and encourage responsible innovation by clearly defining compliance requirements for different categories of AI systems.

3. Strengthening Data Protection and Privacy Safeguards

Strong data protection and privacy safeguards are essential for effective AI regulation. Since AI systems rely heavily on personal data, robust privacy laws and enforcement mechanisms must be implemented to prevent misuse and unauthorized access to data. The Digital Personal Data Protection framework should be effectively enforced, and additional safeguards such as algorithmic transparency and consent-based data usage should be incorporated. Strengthening privacy protection will help ensure that AI systems operate in a manner consistent with constitutional principles and human dignity.

4. Establishment of a Dedicated AI Regulatory Authority

India should establish a specialized regulatory authority to oversee Artificial Intelligence governance and ensure compliance with legal standards. A dedicated regulatory body would monitor AI deployment, conduct audits, enforce safety standards, and address grievances related to AI-related harm. Institutional oversight is necessary to ensure accountability and effective implementation of regulatory

policies. Such an authority would also facilitate coordination between government agencies, private companies, and research institutions.

5. Development of a Clear Liability and Accountability Framework

A clear liability framework is essential to address harm caused by Artificial Intelligence systems. The law should specify the responsibilities of developers, manufacturers, service providers, and users in order to determine accountability in case of failure or damage. A statutory framework incorporating elements of tort law, product liability, and consumer protection principles would ensure compensation for victims and promote responsible innovation. Clear liability rules would also increase public confidence in AI technologies and encourage ethical development.

6. Promotion of Ethical and Transparent AI Systems

Ethical principles such as fairness, transparency, accountability, and non-discrimination should be legally enforced in AI governance. AI systems should be designed in a manner that allows explainability and independent auditing to prevent bias and unfair outcomes. Legal requirements for transparency and ethical compliance will ensure that AI systems respect fundamental rights and operate in a socially responsible manner. Promoting ethical AI will strengthen public trust and align technological development with democratic values.

7. Capacity Building and Institutional Development

Effective AI regulation requires capacity building among policymakers, regulators, judges, and legal professionals. Training programs, research initiatives, and academic collaborations should be encouraged to enhance understanding of Artificial Intelligence and its legal implications. Institutional development will help ensure that regulatory authorities and judicial bodies are equipped to handle complex AI-related disputes and policy challenges. Building institutional capacity will strengthen the overall governance framework and ensure effective implementation of AI laws.

VI. REFERENCES:

1. Justice K.S. Puttaswamy v. Union of India, (2017) 10 SCC 1
2. State of West Bengal v. Anwar Ali Sarkar, AIR 1952 SC 75

3. Digital Personal Data Protection Act 2023 (India); Digital Personal Data Protection Rules 2025, notified by the Ministry of Electronics and Information Technology, November 2025.
4. Ministry of Electronics and Information Technology, *Advisory on Artificial Intelligence and Intermediary Compliance* (March 2024).
5. Government of India, *IndiaAI Mission* (Ministry of Electronics and Information Technology, March 2024).
6. Ministry of Electronics and Information Technology, *India AI Governance Guidelines* (November 2025).
7. Artificial Intelligence (Ethics and Accountability) Bill 2025, Bill No. 59 of 2025, introduced in the Lok Sabha as a Private Member's Bill.
8. OECD, *Principles on Artificial Intelligence*, 2019
9. UNESCO, *Recommendation on the Ethics of Artificial Intelligence*, 2021.