



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.v4si1.24>

© 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]*

HUMAN RIGHTS PROTECTION IN THE AGE OF ARTIFICIAL INTELLIGENCE: COMPARATIVE STUDY OF MODERN DEMOCRACIES

Dr. Sukdeo Ingale¹

I. ABSTRACT

Artificial Intelligence (AI) technology has the potential to transform governance, law enforcement, public administration and access to justice. While existing scholarship has largely emphasized risks such as privacy violations, algorithmic bias, mass surveillance, intellectual property concerns, automated profiling and exclusion, there is increasing recognition that AI can also be deployed as a tool for the protection and promotion of human rights. This paper undertakes a comparative legal study of selected modern democracies, namely the United Kingdom, the United States of America, Canada, Australia and India, to examine how AI-based systems are being used or regulated in relation to human rights enforcement. The study adopts a doctrinal and comparative methodology, relying on legal instruments, policy frameworks, judicial developments and scholarly literature to analyse AI applications in access to justice, equality, non-discrimination, welfare delivery, transparency and humanitarian response. The paper further examines challenges arising from AI deployment, including algorithmic bias, lack of transparency and explainability, fragmented regulation, digital divide and inadequate remedies for AI-generated harms. It recommends mandatory human rights impact assessments, explainable and accountable AI systems, effective grievance redressal mechanisms, inclusive public consultation, human supervision in consequential decisions, protection of vulnerable groups and stronger international cooperation on AI governance. These recommendations are directed towards ensuring that innovation does not weaken constitutional values, democratic accountability or substantive equality. The central argument of the paper is that AI has significant potential to advance human

¹ Asst. Professor, Department of Law, Vishwakarma University, Pune (India). Email: sukdeo.ingale@vupune.ac.in

rights, but its effectiveness depends upon a suitable rights-based legal framework, transparent governance norms, ethical design, institutional accountability and meaningful human oversight wherever rights and liberties are affected.

II. KEYWORDS

Artificial Intelligence, Human rights, AI Governance, Democracy, Law and Technology.

III. INTRODUCTION AND RESEARCH PROBLEM

The recent innovation in technology has opened Pandora's box of Artificial Intelligence (AI), which has become a core component of modern governance and administration. Governments, courts and public institutions are increasingly deploying AI to improve efficiency, accessibility, responsiveness and productivity within legal and administrative systems. Initially, in legal systems such as the European Union, Germany, the UK, France, Canada, the USA, Australia and Japan, greater attention was focused on the risks of AI, including privacy violations, systemic bias, mass biometric surveillance, data protection concerns, automated profiling, automated welfare decisions, predictive policing, ethical constraints and social or cultural risks.

Gradually, however, the potential of AI was also recognised, and its role expanded across several spheres of life, including the protection, promotion and advancement of human rights. The intersection of AI and human rights nevertheless raises critical normative questions concerning equality, dignity, accountability, transparency and human supervision. The research problem addressed in this paper lies in the tension between AI's rights-protective potential and its capacity to generate new forms of discrimination, exclusion and opaque decision-making.

There is also no settled comparative framework for assessing how modern democracies should regulate AI while preserving human rights and democratic accountability. Against this backdrop, the present paper studies selected legal systems to examine how AI may be used for protecting, promoting and advancing human rights.

A. Research Objectives

The present research paper aims to examine the role of Artificial Intelligence in the protection, promotion and advancement of human rights in modern democracies. The specific objectives of the study are:

1. To analyse the conceptual relationship between Artificial Intelligence and human rights in the context of democratic governance, public administration and access to justice.
2. To comparatively examine the AI governance models and human rights protection mechanisms adopted in the United Kingdom, the United States of America, Canada, Australia and India.
3. To identify the major risks associated with AI-based governance, including algorithmic bias, lack of transparency, automated discrimination, digital exclusion and inadequate remedies for AI-generated harms.
4. To evaluate the extent to which existing legal, judicial and policy frameworks ensure transparency, accountability, fairness and human supervision in the use of AI systems.
5. To formulate India-specific policy recommendations for developing a rights-based AI governance framework consistent with constitutional values, democratic accountability and international human rights principles.

B. Research Questions

The present research paper seeks to answer the following research questions:

1. How can Artificial Intelligence be used as a tool for the protection, promotion and advancement of human rights in modern democratic legal systems?
2. How do the AI governance frameworks of the United Kingdom, the United States of America, Canada, Australia and India differ in their approach to human rights safeguards, transparency, accountability and human supervision?

3. What are the common human rights risks arising from AI-based governance, particularly in relation to algorithmic bias, privacy, automated discrimination, lack of explainability and exclusion of vulnerable groups?
4. To what extent do existing legal, judicial and policy frameworks in the selected jurisdictions provide effective remedies against AI-generated harms?
5. What accounts for India's comparative regulatory gap in AI governance, and what legal and policy measures are required to develop a rights-based AI framework consistent with constitutional values and international human rights principles?

C. Research Methodology

The present study adopts a doctrinal and comparative legal research methodology. It is primarily based on the analysis of legal materials, policy documents and academic literature relating to Artificial Intelligence and human rights. The research relies on statutes, judicial decisions, government reports, policy frameworks, international instruments, scholarly articles and secondary literature to examine the relationship between AI governance and human rights protection.

The comparative method is used to study five democratic jurisdictions, namely the United Kingdom, the United States of America, Canada, Australia and India. These jurisdictions have been selected because they represent different models of AI governance, including data-protection-based regulation, decentralised civil-rights enforcement, public-sector automated decision-making controls, ethical AI principles and emerging digital governance frameworks. The study evaluates these models to identify common risks, institutional safeguards and policy lessons relevant to India.

The research is analytical in nature and does not involve empirical fieldwork. Its purpose is to examine existing legal and policy developments and to formulate recommendations for a rights-based AI governance framework that promotes transparency, accountability, non-discrimination, human dignity and effective human supervision.

IV. DEFINING HUMAN RIGHTS IN THE AI CONTEXT

The conceptual framework of human rights is reflected in foundational international instruments such as the Universal Declaration of Human Rights, 1948 (UDHR), the International Covenant on Civil and Political Rights, 1966 (ICCPR) and the International Covenant on Economic, Social and Cultural Rights, 1966 (ICESCR). The UDHR was adopted by the United Nations General Assembly on 10 December 1948, while both the ICCPR and the ICESCR were adopted and opened for signature, ratification and accession by General Assembly Resolution 2200A (XXI) on 16 December 1966. In the contemporary digital era, AI-based systems have the capacity to affect several human rights, including privacy, freedom of expression, equality, human dignity and participation in public life without fear or discrimination. The AI governance framework in recent times has to increasingly refer to human rights principles as integral part of ethical way of AI design, deployment and oversight.

The AI context with the human rights is as a tool for protection, promotion and advancement of human rights. AI can work as a tool in multiple areas including access to justice, detection of discrimination, promotion of transparency and accountability and also humanitarian response during governance. For access to justice, the AI platforms may assist individual persons in understanding nature scope and limitations of legal rights, process and procedure of filing complaints and for taking further follow ups. AI can detect discrimination through identification of patterns of unequal treatment given to some classes of categories of community during policing, employment, promoting corrective actions or in the implementation of welfare schemes. The AI analytics can make government processes more transparent and accountable. AI can also support for humanitarian responses during governance, especially during the crises and disaster management.

A. Human Rights Protection in the UK

The UK has developed a rights-based and sector-specific approach to AI governance, aligned with democratic norms, data protection principles and human rights obligations.

The Data Protection Act, 2018, read with the UK GDPR framework, remains central to the regulation of automated decision-making involving personal data. This framework requires fairness, transparency, accountability, lawful processing and safeguards where automated processing may significantly affect individuals. In the human rights context, these safeguards are important because AI systems may influence access to welfare benefits, policing decisions, immigration screening, public administration and access to services.

The UK's model is also shaped by its participation in the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, which was opened for signature on 5 September 2024 and is described by the Council of Europe as the first legally binding international treaty in this field. Its object is to ensure that activities within the AI lifecycle remain consistent with human rights, democracy and the rule of law. The UK Government confirmed its signature of the treaty on 5 September 2024, emphasising protection of human rights, democracy and the rule of law from potential AI-related harms. Therefore, the UK section should not merely state that AI is useful for human rights, but should also explain that the UK approach combines data protection safeguards, public-sector accountability, judicial caution and international treaty commitments

The UK has played an active role in the Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law. It was signed by the UK in 2024 under the Council of Europe's auspices (Reuters, 2024). This Convention is legally binding on the State parties and the States has to align use of AI with democratic values, human dignity and rights. The convention also ensures that the State Parties adhere to the necessary transparency and allow citizens to challenge decisions of AI. While interpreting this Convention, the UK Courts and judicial officials have also contributed for interaction of AI and human intelligence. In this regard, Sir Geoffrey Vos (2024), the Master of the Rolls, has advocated for a legal right to have cases decided by human intelligence rather than AI as there is intrinsic link between judgment and democratic values.

B. Human Rights Protection In The USA

The USA follows a decentralised and sectoral model of AI governance. In the absence of a single comprehensive federal AI statute, the protection of human rights in AI-related matters depends on constitutional principles, civil rights statutes, consumer protection law, privacy law, anti-discrimination norms and emerging State-level AI legislation. Existing statutes such as the Civil Rights Act, 1964, the Fair Housing Act, 1968 and equal employment opportunity laws remain relevant where AI systems produce discriminatory outcomes in employment, housing, credit, education or public services.

At the State level, Colorado has become an important example of the changing nature of AI regulation. Colorado originally enacted SB 24-205 to regulate high-risk AI systems, but this position has now been materially modified. SB 26-189, signed on 14 May 2026, replaced the earlier framework with a narrower law on automated decision-making technology, enforced through the Colorado Consumer Protection Act. The revised approach focuses on transparency, explanations for adverse outcomes, consumer recourse and meaningful human review in consequential decisions. This development shows that the American model is innovative but fragmented: it can respond quickly to algorithmic harms, yet it may also produce uneven protection across States. Accordingly, the USA section should analyse both the strengths and limits of a decentralised approach, especially where AI affects equality, due process, privacy and access to remedies.

However, some scholars argue that the existing AI driven mechanism shall be accompany by the traditional civil rights framework to recognize emergent algorithmic harms (Barocas & Selbst, 2016). Similarly, the theory of artificial immutability suggests that the groups generated through algorithms using AI do not align with the legally protected classes. It results in the discrimination for the traditional groups and hence in the context of AI, the traditional groups require protection (Wachter, 2022).

In the USA, some States like Colorado have adopted AI transparency and accountability laws regulating high-risky systems in employment, education and housing to protect human rights and prevent discrimination. Various Courts in the USA have also addressed AI related issues. In *Thaler v. Vidal* (2022) the US Court of Appeal held that

only human beings can be named as inventors and the Court specifically mentioned that AI artifacts shall not be shouldered with such legal rights. This type of judicial approach ensures that the 'legal agency' shall fundamentally remain human so that human rights can be promoted. The USA also uses AI in predictive policing and surveillance raising concerns about civil liberties and human rights.

C. Human Rights Protection in Canada

Canada has developed a structured public-law approach to automated decision-making, but the manuscript must correct the legal position regarding the proposed Artificial Intelligence and Data Act. AIDA was proposed as part of Bill C-27; however, it did not become enforceable law. Bill C-27 died on the Order Paper following prorogation of Parliament in January 2025. Therefore, the manuscript should not state that AIDA currently imposes mandatory obligations on organisations deploying high-impact AI systems.

The Canadian position is better described as a rights-based but incomplete regulatory framework. AI governance in Canada is presently shaped by the Canadian Charter of Rights and Freedoms, privacy law, public-sector automated decision-making policies and continuing legislative debate on algorithmic accountability. In human rights terms, the Canadian approach remains significant because it places emphasis on equality, privacy, non-discrimination, administrative fairness and proportionality in public decision-making. However, the absence of an enacted AIDA-type statute also reveals a regulatory gap. The Canadian model therefore demonstrates that strong constitutional values and public-sector norms can guide AI governance, but binding AI-specific duties are still necessary for consistent protection against algorithmic bias, opaque automated decisions and AI-generated harms.

The major focus of this mandate is on the issues of discrimination and privacy rights. This rights-based approach treats the human rights risk assessment as an integral part of AI governance (Government of Canada, 2022). AI regulation in Canada is based on the Canadian Charter of Rights and Freedoms. The Canadian Judiciary has engaged in constant internal debate about the use of AI in judicial processes. The main concerns in

this debate are about preserving democratic features of the governance and also the impartial human adjudication (The Conversation, 2024).

D. Human Rights Protection in Australia

Australia's experience is especially relevant because AI and automated systems have been used in welfare delivery, public administration and regulatory decision-making. The Australian approach has been influenced by voluntary AI Ethics Principles, which emphasise human, social and environmental wellbeing, human-centred values, fairness, privacy protection, reliability, transparency, contestability and accountability. These principles are useful for human rights protection because they recognise that automated systems should not operate beyond institutional responsibility or without avenues for review.

At the same time, Australia illustrates the risks of relying excessively on automation in welfare governance. Automated eligibility, debt recovery or compliance systems may adversely affect vulnerable persons where the data is incomplete, the model is inaccurate, or the affected person lacks the capacity to challenge the decision. For this reason, the Australian section should be expanded to examine AI not merely as a technological tool, but as an administrative power that must remain subject to legality, procedural fairness, human review and effective remedies. A rights-based Australian model requires that public authorities ensure transparency of automated decisions, preserve human accountability and maintain accessible complaint mechanisms for persons affected by AI-based governance.

The AI framework used by public agencies requires ethical risk assessments and human oversight. Australia has actively participated in international treaties like the Framework Convention on AI (2024) to strengthen its commitments to align use of AI based on the democratic norms.

E. Human Rights Protection in India

Use of AI is expanding in India in governance, legal services, translation, court administration, welfare delivery and digital public infrastructure. India does not yet have

a comprehensive AI-specific legislation, but its legal position has changed significantly because the Digital Personal Data Protection Act, 2023 and the Digital Personal Data Protection Rules, 2025 are directly relevant to AI systems that process personal data. The Ministry of Electronics and Information Technology has published the Digital Personal Data Protection Rules, 2025, including provisions relating to the establishment of the Data Protection Board of India. These developments should be incorporated into the manuscript because privacy, consent, data security, breach notification and individual control over digital personal data are central to rights-compliant AI governance.

In India, AI can strengthen human rights by improving access to justice through legal information platforms, translation tools, case management systems and assistive technologies for persons facing linguistic, economic or geographical barriers. However, the Indian context also raises specific concerns regarding digital exclusion, linguistic diversity, caste and gender bias in datasets, welfare exclusion, surveillance and the absence of clear statutory standards for high-risk AI systems. The Indian judiciary has shown cautious engagement with AI-assisted legal research, while also recognising that judicial reasoning and final adjudication must remain human functions.

Therefore, India's human rights approach to AI should be based on constitutional values, including equality, dignity, privacy, access to justice and due process. The manuscript should recommend that India move from a general "AI for All" policy approach towards a rights-based AI governance framework containing mandatory human rights impact assessments, explainability duties, grievance redressal, independent oversight and human review in consequential decisions.

The AI platforms like Nyaay AI and Adalat AI have been deployed for better case management system, legal research, transcription and translations enhancing access to justice (Oxford Institute, 2025). Although, Indian courts have experimented with use of AI for research in few cases like *Jaswinder Singh v. State of Punjab* (2023) and *Md. Zakir Hussain v. State of Manipur* (2024) using different AI tools, the Judiciary has expressed caution about AI's role in decision making. The Kerala High Court has explicitly

prohibited use of AI tools for judicial reasoning through a comprehensive policy issued (Times of India, 2025).

V. CHALLENGES FACED BY AI BASED HUMAN RIGHTS PROTECTION MODELS

One of the most persistent challenges is algorithmic bias. AI models are based on historical database supplied to develop artificial intelligence. Hence, AI models often over-represent existing socio-economic inequalities, prejudices, and structural discrimination in the existing society. When such biased data is used without adequate safeguards, the AI models amplify the inequalities, prejudices and discrimination. For example, AI based predictive policing tools may disproportionately target people living in low income or minority localities creating a loop of over policing. This raises concerns of substantive equality, non-discrimination, faire treatment and human dignity. The AI based welfare eligibility systems used in the USA and Australia have been criticized due to this challenge for unfairly denying benefits to vulnerable groups in their societies.

Another major challenge concerns the opacity of AI systems and the lack of transparency, which significantly undermine accountability, procedural fairness and access to remedies. Many AI algorithms based on machine learning and deep neural networks operate in a way that their decisions are not easily explainable even to their developers. For example, in many countries, AI tools are used to assess risk profiles in immigration and asylum decision making. In such cases when any applicant receives adverse decisions, reasons behind decisions are not informed due to absence of explainability. Same is the situation when AI based credit scoring systems decide access to loans, housing and employment. This challenge of lack of transparency affects right to fair hearing, right to an effective remedy and right to information undermining the concerned human rights.

The absence of comprehensive regulatory frameworks also results in challenges as to the uniform treatment, consistent standards and even safeguards for protection of human rights. In the USA, AI governance increasingly relies on the existing civil rights and

consumer protection laws. These laws provide remedies without addressing AI specific risks of systematic bias, algorithmic explainability, etc. This regulatory gap operates without clear guidelines. Hence, in federal countries like India and Australia the State/Regional governments may respond AI risks differently and sometimes even contrary creating a patchwork approach resulting in unequal protection of rights in same country. This challenge results in uncertainty in protection of human rights, limited accountability of public authorities, and inadequate remedies for AI generated harms.

The challenge of digital divide is a structural challenge created by digital infrastructure technological literacy and institutional support available to different section of the society. When government increasingly promote use of AI based government schemes, healthcare access, grievance redressal mechanisms, etc without parallel offline mechanisms, it creates possibility that the vulnerable populations like elderly persons, differently-abled, socially-educationally or economically disadvantaged section in society, etc. may not have equal access to AI based infrastructure. For example, AI based legal assistance platforms and online dispute resolution system may exclude people lacking internet access, digital literacy or language support. This challenge affects right to equality, right to access public services, right to social inclusion, etc.

A further challenge is that AI regulation is developing unevenly across jurisdictions. The UK has treaty-based and data-protection-oriented safeguards; the USA relies on civil rights law and State-level experimentation; Canada has constitutional and policy-based safeguards but no enacted AIDA framework; Australia relies significantly on principles and administrative accountability; and India has a data protection framework but no comprehensive AI statute. This unevenness creates difficulty in developing uniform standards for human rights protection. It also affects cross-border AI systems, because the same AI tool may be subject to different levels of transparency, accountability and remedial control in different jurisdictions. From a human rights perspective, this fragmented position increases the risk that vulnerable persons may be left without effective remedies when AI systems cause exclusion, discrimination or denial of benefits.

VI. RECOMMENDATIONS FOR HARNESSING AI IN HUMAN RIGHTS PROTECTION

The effective use of AI as a tool for human rights protection requires a carefully drafted policy framework which can balance technological innovations with human right values, democratic accountability and social justice. For that purpose, the above-mentioned challenges can be tackled through the following policy recommendations.

1. While developing any AI model, the policy makers shall carry on human rights impact assessment tests to identify, evaluate and mitigate potential human rights risks before deploying any AI systems.
2. The AI systems shall be developed which will explain logic of its decisions and allow affected persons to challenge decisions based on AI algorithms to enforce accountability.
3. A comprehensive rights-based AI framework has to be developed which priorities the remedies for harms caused by AI based actions, and decisions.
4. The public consultation before using AI systems for welfare schemes or policies can ensure social inclusivity and participatory policymaking by understanding capabilities, success/ failure probabilities and possibility of harm/injury due to use of AI systems.
5. International cooperation on AI governance for overcoming divergent national regulations, regulatory arbitrage can answer many challenges posed by integration of AI in human rights protection.
6. The recommendations should therefore be strengthened by adopting a layered rights-based governance model. First, high-risk AI systems used in welfare, policing, employment, housing, education, health care, immigration and judicial administration should be subjected to mandatory human rights impact assessment before deployment. Secondly, affected persons should have a legal right to receive meaningful explanations and to challenge adverse AI-assisted decisions before a human authority. Thirdly, AI systems used by public bodies should be regularly audited for bias, accuracy, proportionality and disparate impact on vulnerable

groups. Fourthly, India should develop sector-specific AI standards under a broader rights-based framework, ensuring consistency with constitutional guarantees and data protection law. Finally, international cooperation should be strengthened so that democratic jurisdictions can develop common standards on transparency, accountability, safety and human oversight without weakening innovation.

VII. CONCLUSION

Use of AI is neither inherently harmful nor inherently beneficial for the protection, promotion and development of human rights. Its impact depends on various factors like systems of governance, legal framework, judicial approach, and ethical design of a specific country and the people using AI in such a country. The comparative study of five leading democracies demonstrates that modern democracies are inclined to use AI to strengthen human rights. It is much more successful in cases of establishing equality, transparency, accountability, human dignity, and preventing discrimination, unfair treatments, and misuse of powers.

However, the challenges of algorithmic bias, systematic discrimination, lack of transparency and explainability, absence of a comprehensive regulatory framework, digital divide and exclusion needs immediate attention. These challenges can be cured by mandatory human rights impact assessment, ensuring transparency and explainability through necessary algorithmic development, guaranteeing uniform remedy-based approach, promoting inclusivity through strong infrastructure and intrastructure, and strengthening international co-operation on use of AI governance.

VIII. REFERENCES

1. Barocas, S., & Selbst, A. D. (2016). Big data's disparate impact. *California Law Review*, 104(3), 671–732.
2. Floridi, L., Cowls, J., Beltrametti, M., et al. (2018). AI4People - An ethical framework for a good AI society. *Minds and Machines*, 28(4), 689–707.
3. Government of Canada. (2022). *Artificial Intelligence and Data Act (AIDA)*.

4. Jaswinder Singh v. State of Punjab, CRM-M No. 31118 of 2022 (Punjab H. C. February, 06, 2023).
5. Md. Zakir Hussain v. State of Manipur & Ors., WP(C) No. 70 of 2023 (Manipur H.C. May 23, 2024).
6. NITI Aayog. (2018). *National strategy for artificial intelligence: #AIForAll*. Government of India.
7. Oxford Institute. (2025).
<https://www.techandjustice.bsg.ox.ac.uk/research/india?utm>
8. Reuters. (2024). *Council of Europe adopts landmark treaty on artificial intelligence and human rights*.
9. Thaler v. Vidal (2022) 43 F.4th 1207 (2022)
10. Times of India. (2025). Kerala High Courts new guidelines to district judiciary, 2025,
<https://timesofindia.indiatimes.com/technology/artificial-intelligence/kerala-high-courts-new-guidelines-to-district-judiciary-ai-tools-shall-not-be-used-to-/articleshow/122795589.cms?utm>
11. Veale, M., & Edwards, L. (2018). Clarity, surprises, and further questions in the GDPR. *Computer Law & Security Review*, 34, 398–404.
12. Vos, G. (2021). *Artificial intelligence and the future of justice*. Judiciary of England and Wales.
13. Wachter, S. (2022). The Theory of Artificial Immutability: Protecting Algorithmic Groups Under Anti-Discrimination Law, 97 Tul. L. Review. XX (2022-2023).