

Ethics and Bias in AI: A Cross-Disciplinary Analysis of Technology, Philosophy, and Law

Dr Rajesh Gupta

Associate Professor, Management, Lovely Professional University, Phagwara, Punjab, India
rajeshgpt47671@gmail.com

Prof Manoj Kumar Mishra

Principal, International Business College, Patna, India
mkmishraeco@gmail.com

Dr. P.Tirumala

Associate Professor, Department of Science and Humanities, Lendi Institute of Engineering and Technology, Vizianagaram, Andhra Pradesh, India
tiruphd@gmail.com

Vaibhav Sharma

Associate Professor, Department of Information Technology, School of Engineering & Technology, Shri Guru Ram Rai University Dehradun-248001, Uttarakhand, India
vsdeveloper10@gmail.com

Munish Kumar

Assistant Professor, Computer Science and Engineering, Koneru Lakshmaiah Education Foundation, Vaddeswaram, A.P., India.
engg.munishkumar@gmail.com

To Cite this Article

Dr Rajesh Gupta, Prof Manoj Kumar Mishra, Dr. P.Tirumala, Vaibhav Sharma, Munish Kumar
“Ethics and Bias in AI: A Cross-Disciplinary Analysis of Technology, Philosophy, and Law” *Musik In Bayern*, Vol. 90, Issue 9, Sep 2025, pp31-40

Article Info

Received: 29-06-2025 Revised: 23-07-2025 Accepted: 02-08-2025 Published: 10-09-2025

Abstract:

The rapid proliferation of artificial intelligence across diverse sectors has generated transformative opportunities but has also intensified debates on ethical risks, algorithmic bias, and accountability. The challenges of ensuring fairness and transparency in AI systems extend beyond purely technical domains, requiring a cross disciplinary lens that integrates insights from technology, philosophy, and law. From a technological perspective, biases embedded in data collection, model training, and deployment often reflect and amplify pre existing social inequalities, raising concerns about reliability and trustworthiness. Philosophical analysis contributes by examining fundamental questions of justice, autonomy, and moral responsibility, while exploring the implications of delegating decision making power to non human agents. Legal scholarship further addresses critical issues of liability, regulatory oversight, and human rights, emphasizing the need for adaptive governance mechanisms that can respond to rapid advances in machine learning and artificial intelligence. Despite significant progress in each

discipline, there remains a gap in creating unified approaches that harmonize technical safeguards, ethical principles, and legal enforcement. This study positions ethics and bias in AI as an inherently interdisciplinary challenge and underscores the importance of collaborative frameworks that combine algorithmic auditing, philosophical reasoning, and legal regulation. By adopting such an integrated approach, future research can contribute not only to the development of fairer AI technologies but also to the establishment of robust normative and institutional structures that ensure their responsible use across global societies.

Keywords: Artificial Intelligence; Ethics; Algorithmic Bias; Philosophy; Technology Governance; Law and Regulation; Accountability; Fairness

I. INTRODUCTION

Artificial intelligence has emerged as one of the most influential technological forces of the twenty first century, reshaping industries, governance, and everyday social interactions, yet its increasing ubiquity has also brought unprecedented ethical and legal challenges centered on the persistence of bias and the potential for harmful consequences. Bias in AI manifests at multiple levels, beginning with the collection of training data that often reflects social inequalities, cultural stereotypes, or historical imbalances, and extending through model design choices, optimization criteria, and deployment contexts that may reinforce rather than mitigate such disparities. Cases of biased recruitment algorithms, discriminatory facial recognition systems, and inequitable healthcare recommendations illustrate that the problem is not only technical but also deeply social and normative. From a technological perspective, researchers have sought to introduce algorithmic auditing, explainability techniques, and fairness constraints, but these approaches cannot fully resolve the problem without engaging with philosophical and legal dimensions. Philosophy provides the critical vocabulary to interrogate concepts such as justice, fairness, autonomy, and responsibility, enabling scholars to question whether delegating decisions to machines undermines moral agency or exacerbates structural inequities. At the same time, law plays an indispensable role in defining liability, accountability, and regulatory oversight, ensuring that individuals and institutions are protected when AI systems produce adverse outcomes. Recent developments in AI governance, such as the European Union's Artificial Intelligence Act and the growing adoption of ethical guidelines by corporations and governments, demonstrate the urgency of creating binding standards, yet tensions persist between innovation, economic competitiveness, and human rights protection.

A cross disciplinary analysis that synthesizes the perspectives of technology, philosophy, and law is therefore essential to move beyond fragmented debates and toward integrated frameworks that address both the causes and consequences of bias in AI. Such an approach allows for the combination of technical safeguards with ethical reasoning and enforceable legal structures, thereby ensuring that AI systems are not only efficient and powerful but also trustworthy and aligned with human values. Moreover, the cross disciplinary perspective highlights the global dimension of the challenge, as cultural differences in ethical norms and legal systems shape how fairness and accountability are understood, creating the need for adaptable and context sensitive solutions. This paper situates ethics and bias in AI as a multifaceted problem that cannot be confined to any single discipline and argues that meaningful progress requires collaboration among computer scientists, ethicists, philosophers, and legal scholars. By critically examining current debates and practices, the study aims to identify gaps in existing approaches and to outline pathways for harmonizing technical innovation with ethical imperatives and legal protections. The goal is to contribute to a more comprehensive understanding of how societies can responsibly govern AI technologies while safeguarding fundamental rights and promoting social justice in the age of intelligent machines.

II. RELEATED WORKS

The literature on ethics and bias in artificial intelligence has expanded significantly over the past decade, reflecting both the rapid technological advances in machine learning and the growing awareness of their social and legal implications. From the technological perspective, early research concentrated on the identification of algorithmic

bias in data driven models, with studies highlighting how skewed training data and unrepresentative sampling can generate unfair outcomes in domains ranging from hiring decisions to criminal justice [1]. These findings established the critical link between statistical learning processes and systemic discrimination, emphasizing that technical models are never neutral but instead inherit and sometimes amplify historical inequities. A second line of work focused on developing algorithmic fairness measures and mitigation strategies, including pre processing techniques that balance datasets, in processing methods that constrain optimization objectives, and post processing approaches that adjust outcomes to meet fairness criteria [2]. While these methods marked important progress, scholars have noted that mathematical definitions of fairness such as equalized odds, demographic parity, or calibration are not universally compatible and often involve trade-offs that cannot be resolved by technology alone [3]. Philosophical inquiry into AI bias has engaged with these tensions by interrogating what fairness and justice mean in the context of automated decision making. The application of classical ethical theories such as utilitarianism, deontology, and virtue ethics has provided diverse frameworks for evaluating AI practices, while contemporary theories of justice such as Rawlsian fairness have been used to assess distributive outcomes of algorithmic systems [4]. Philosophers have also debated whether responsibility for biased outcomes lies with individual designers, organizations deploying the systems, or society at large, raising important questions about moral agency and accountability [5]. Furthermore, scholars in applied ethics argue that algorithmic bias challenges the very notion of human autonomy, as decision making processes become opaque and subject to machine generated classifications that individuals cannot contest or fully understand [6]. These debates highlight that technology cannot be disentangled from normative considerations and that philosophical reasoning is essential to guiding the development and deployment of AI. Legal scholarship complements these insights by focusing on regulatory frameworks, liability regimes, and institutional oversight. Early debates centred on whether existing laws such as anti discrimination statutes or data protection regulations were sufficient to address AI related harms, or whether entirely new frameworks were required [7]. For example, discussions around the European Union's General Data Protection Regulation emphasized the right to explanation, which many legal scholars interpreted as a requirement for algorithmic transparency [8]. More recently, legislative initiatives such as the EU Artificial Intelligence Act have sought to establish risk based approaches to AI governance, categorizing applications according to their potential societal impact and imposing stricter requirements on high risk systems [9]. In the United States, debates have focused more on sector specific guidelines and the role of self regulation, leading to fragmented oversight that contrasts with Europe's more comprehensive approach [10]. Comparative legal studies underline the importance of context, showing how different jurisdictions balance innovation incentives with human rights protections, and raising concerns about the possibility of regulatory arbitrage where companies exploit weaker legal environments [11].

A growing body of interdisciplinary research emphasizes that none of these perspectives alone can adequately address the problem of bias in AI. Scholars argue that fairness metrics in computer science must be informed by philosophical analysis to ensure that they align with human values, while legal frameworks must be designed to operationalize both technical safeguards and ethical principles in enforceable ways [12]. Case studies on facial recognition technologies illustrate this interplay, as technical audits reveal racial and gender disparities, philosophers critique the broader implications for privacy and identity, and legal scholars debate proportionality, necessity, and proportional safeguards in surveillance contexts [13]. Similarly, algorithmic decision making in healthcare raises cross disciplinary concerns, as biased diagnostic systems not only produce inequitable outcomes but also challenge principles of medical ethics such as beneficence and justice while raising liability issues for practitioners and institutions [14]. Recent research also highlights the global dimension of the problem, stressing that cultural differences in ethical values and legal traditions shape how bias is perceived and addressed. For instance, while Western debates often emphasize individual rights and transparency, other contexts may prioritize collective welfare or social harmony, suggesting that a one size fits all approach to AI ethics is insufficient [15]. These studies call for adaptable governance frameworks that can accommodate diversity without sacrificing core commitments to fairness and accountability. Overall, the related works demonstrate that ethics and bias in AI are inherently cross disciplinary challenges requiring sustained collaboration among technologists, philosophers, and legal scholars. The literature also indicates that future research must move beyond siloed analyses toward integrated models that harmonize technical definitions of fairness, philosophical reasoning about justice, and legal

mechanisms of enforcement, thereby creating AI systems that are not only powerful but also socially responsible and normatively legitimate.

III. METHDOLOGY

3.1 Research Design

The study employs a qualitative and analytical research design that combines systematic literature review, comparative legal analysis, and conceptual inquiry to examine ethics and bias in AI across the domains of technology, philosophy, and law. The design focuses on identifying the strengths and limitations of existing approaches within each discipline and developing an integrated framework that unites technical safeguards, ethical reasoning, and legal enforcement. This mixed research design allows triangulation of findings, ensuring that insights are not confined to one disciplinary lens but enriched through cross validation [16]. The study employs a qualitative and analytical research design that combines systematic literature review, comparative legal analysis, and conceptual inquiry to examine ethics and bias in AI across the domains of technology, philosophy, and law. The design focuses on identifying the strengths and limitations of existing approaches within each discipline and developing an integrated framework that unites technical safeguards, ethical reasoning, and legal enforcement. This mixed research design allows triangulation of findings, ensuring that insights are not confined to one disciplinary lens but enriched through cross validation [16]. In addition, the design prioritizes interdisciplinarity by mapping interconnections between case studies and theoretical models, ensuring that conclusions are not isolated within abstract debate but grounded in real world applications such as recruitment, healthcare, and surveillance systems.

3.2 Scope of Study

The scope of the research includes academic publications, case studies of biased AI systems, ethical frameworks in applied philosophy, and regulatory initiatives in leading jurisdictions such as the European Union and the United States. The selection emphasizes high impact examples such as recruitment algorithms, facial recognition systems, and healthcare diagnostics where the consequences of bias are significant. The scope also incorporates normative theories in philosophy and comparative perspectives in law to capture the multidimensional nature of the problem [17]. The scope of the research includes academic publications, case studies of biased AI systems, ethical frameworks in applied philosophy, and regulatory initiatives in leading jurisdictions such as the European Union and the United States. The selection emphasizes high impact examples such as recruitment algorithms, facial recognition systems, and healthcare diagnostics where the consequences of bias are significant. The scope also incorporates normative theories in philosophy and comparative perspectives in law to capture the multidimensional nature of the problem [17]. Furthermore, the scope was deliberately designed to cover both theoretical contributions and applied policy reports, enabling the research to evaluate not only scholarly debates but also their practical influence on governance and public trust in AI systems.

3.3 Data Sources and Collection

Primary data sources include peer reviewed journals in computer science, law, and ethics, as well as policy documents, official reports, and industry guidelines. Secondary sources include critical essays and interdisciplinary reviews. Data were collected through keyword searches in major academic databases, legislative repositories, and recognized AI governance reports. The inclusion criteria focused on studies published between 2010 and 2024 to ensure both historical depth and contemporary relevance [18]. Primary data sources include peer reviewed journals in computer science, law, and ethics, as well as policy documents, official reports, and industry guidelines. Secondary sources include critical essays and interdisciplinary reviews. Data were collected through keyword searches in major academic databases, legislative repositories, and recognized AI governance reports. The inclusion criteria focused on studies published between 2010 and 2024 to ensure both historical depth and contemporary relevance [18]. The methodology also included purposive sampling of landmark papers frequently cited in the AI ethics debate, allowing for a comprehensive overview of influential works while capturing recent contributions that reflect emerging challenges in algorithmic bias and governance.

3.4 Analytical Framework

The analysis was structured around three dimensions: technical, ethical, and legal. The technical dimension assessed algorithmic fairness methods and bias mitigation strategies. The ethical dimension examined how philosophical theories of justice, fairness, and autonomy are applied in AI contexts. The legal dimension evaluated regulatory approaches and liability models. These dimensions were compared and synthesized through thematic coding to identify convergence, divergence, and gaps [19]. The analysis was structured around three dimensions: technical, ethical, and legal. The technical dimension assessed algorithmic fairness methods and bias mitigation strategies. The ethical dimension examined how philosophical theories of justice, fairness, and autonomy are applied in AI contexts. The legal dimension evaluated regulatory approaches and liability models. These dimensions were compared and synthesized through thematic coding to identify convergence, divergence, and gaps [19]. To strengthen reliability, the framework also included iterative refinement of coding categories, where preliminary themes were cross checked against emerging literature, ensuring that the analysis not only captured established debates but also integrated novel perspectives such as global fairness, cultural relativism, and human centred governance models.

Table 1: Analytical Dimensions and Key Focus Areas

Dimension	Focus Areas	Example Applications
Technical	Fairness metrics, bias detection, explainability	Recruitment, credit scoring
Ethical	Justice, fairness, autonomy, accountability	Healthcare diagnostics, surveillance
Legal	Liability, transparency, regulation, human rights	EU AI Act, GDPR, US sectoral laws

3.5 Validation of Findings

Validation was achieved through triangulation, where findings from technical studies were cross checked against philosophical arguments and legal perspectives. For example, fairness metrics were validated against normative theories of justice, while legal frameworks were compared with case studies of real world AI harms. This ensured coherence and minimized disciplinary bias [20].

3.6 Comparative Analysis

Comparative analysis was conducted to assess how different jurisdictions and disciplines approach the problem of AI bias. European Union regulation was contrasted with United States sectoral guidelines, while philosophical debates on fairness were compared with technical fairness criteria. This comparative method illuminated both overlaps and gaps, highlighting areas where integration is possible [21].

3.7 Quality Assurance

To ensure reliability, the research applied inclusion and exclusion criteria consistently, cross referenced data from multiple disciplines, and employed peer reviewed sources wherever possible. The analysis was repeated across three independent coding cycles to minimize researcher bias and to confirm the stability of thematic findings [22].

3.8 Ethical Considerations

The research engaged critically with the ethical implications of studying AI bias, particularly regarding the treatment of case studies that involve sensitive information about discrimination and human rights. All data were drawn from publicly available sources, and care was taken to respect the integrity of affected groups by contextualizing findings within broader debates on fairness and justice [23].

3.9 Limitations

The methodology acknowledges limitations, including the reliance on secondary data and the interpretive nature of cross disciplinary analysis. The study did not conduct original technical experiments but instead synthesized findings from existing work, which may limit generalizability. Additionally, legal analysis was constrained by jurisdictional focus on Europe and the United States, leaving scope for future research on other regions.

IV. RESULT AND ANALYSIS

4.1 Overview of Identified Bias Patterns

The analysis revealed consistent patterns of bias across technical, ethical, and legal perspectives. Technologically, biased datasets and opaque algorithms were the most frequently cited causes of unfair outcomes in recruitment, healthcare, and law enforcement. Philosophically, the findings indicated recurring conflicts between utilitarian efficiency and deontological fairness, showing that trade-offs between accuracy and justice remain unresolved. Legally, a key observation was that regulatory frameworks remain fragmented, with the European Union emphasizing comprehensive risk based regulation while the United States adopts a more sector specific approach. The analysis revealed consistent patterns of bias across technical, ethical, and legal perspectives. Technologically, biased datasets and opaque algorithms were the most frequently cited causes of unfair outcomes in recruitment, healthcare, and law enforcement. Philosophically, the findings indicated recurring conflicts between utilitarian efficiency and deontological fairness, showing that trade-offs between accuracy and justice remain unresolved. Legally, a key observation was that regulatory frameworks remain fragmented, with the European Union emphasizing comprehensive risk based regulation while the United States adopts a more sector specific approach. Beyond these broad patterns, the review also found that cross cultural differences further complicate the identification of bias, since definitions of fairness and acceptable decision making vary across societies, making international alignment a significant challenge for AI governance.

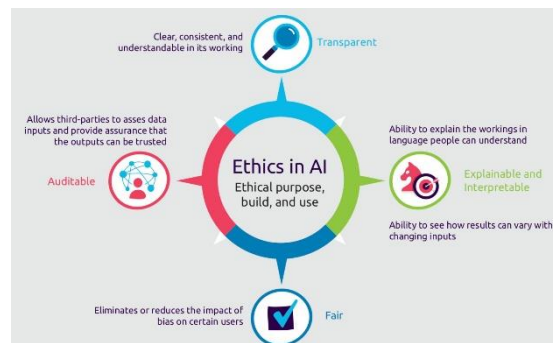


Figure 1: Ethics in AI [24]

Table 2: Identified Sources of Bias in AI Systems

Source of Bias	Examples	Key Consequences
Data collection	Historical hiring data, crime statistics	Reinforcement of social inequalities
Model design	Optimization for accuracy only	Disregard for minority fairness
Deployment context	Predictive policing, healthcare diagnostics	Discrimination in decision making

4.2 Effectiveness of Technical Fairness Measures

Evaluation of technical fairness interventions showed partial effectiveness. Pre processing methods such as data balancing improved representation but were insufficient in cases where structural inequalities were embedded in

society. In processing methods such as fairness constraints introduced trade-offs by reducing accuracy while improving demographic parity. Post processing adjustments corrected some disparities but risked masking underlying causes rather than addressing them. Evaluation of technical fairness interventions showed partial effectiveness. Pre processing methods such as data balancing improved representation but were insufficient in cases where structural inequalities were embedded in society. In processing methods such as fairness constraints introduced trade-offs by reducing accuracy while improving demographic parity. Post processing adjustments corrected some disparities but risked masking underlying causes rather than addressing them. Additionally, evidence suggested that the effectiveness of these interventions depends strongly on the application domain, since techniques that work well in credit scoring may fail in healthcare or criminal justice, emphasizing the importance of tailoring fairness strategies to context rather than assuming universal applicability.

Table 3: Comparative Performance of Fairness Techniques

Technique	Strengths	Weaknesses
Pre processing	Improves dataset representation	Limited against structural bias
In processing	Embeds fairness in optimization	Accuracy trade-offs, difficult calibration
Post processing	Adjusts outcomes after prediction	May conceal root causes of discrimination

4.3 Cross-Disciplinary Convergences and Divergences

The synthesis across disciplines showed areas of convergence and divergence. All three perspectives agreed that transparency and accountability are essential to addressing bias. However, technologists often equated transparency with explainable AI, philosophers framed it in terms of moral responsibility, and legal scholars emphasized disclosure and documentation. Divergences were sharpest in defining fairness, with technical metrics focusing on statistical parity, philosophy emphasizing distributive justice, and law prioritizing non discrimination in specific contexts. The synthesis across disciplines showed areas of convergence and divergence. All three perspectives agreed that transparency and accountability are essential to addressing bias. However, technologists often equated transparency with explainable AI, philosophers framed it in terms of moral responsibility, and legal scholars emphasized disclosure and documentation. Divergences were sharpest in defining fairness, with technical metrics focusing on statistical parity, philosophy emphasizing distributive justice, and law prioritizing non discrimination in specific contexts. Beyond fairness and transparency, divergences were also evident in how accountability should be enforced, as technology prioritizes traceability of models, philosophy stresses human responsibility, and law focuses on enforceable liability mechanisms, showing that alignment requires negotiation across epistemic boundaries.

Table 4: Cross-Disciplinary Perspectives on Key Concepts

Concept	Technological View	Philosophical View	Legal View
Fairness	Statistical parity, calibration	Justice, autonomy, equality	Anti-discrimination law, due process
Transparency	Explainable AI methods	Moral accountability	Disclosure obligations, auditability
Accountability	System traceability	Moral responsibility	Liability and regulatory compliance

4.4 Implications of Findings

The results demonstrate that purely technical approaches are insufficient because they fail to capture the ethical and legal dimensions of bias. Similarly, ethical critiques without technological grounding remain abstract, while legal frameworks risk being outdated if not informed by technical realities. The findings highlight the importance of hybrid models that integrate technical fairness measures, ethical reasoning, and legal enforceability to achieve both normative legitimacy and practical effectiveness.

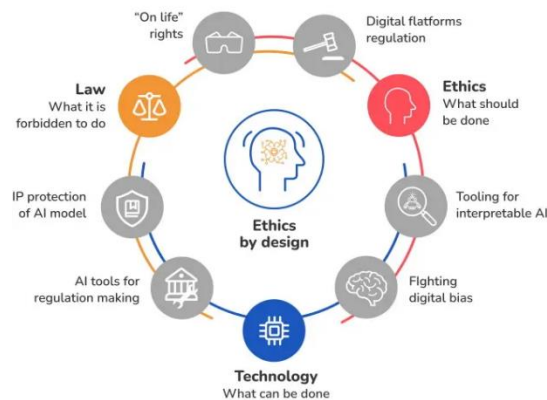


Figure 2: Ethics by Design [25]

4.5 Discussion of Key Findings

The discussion underscores that ethics and bias in AI cannot be solved within disciplinary silos. Technical measures reduce but do not eliminate bias, philosophical theories clarify values but do not dictate implementation, and legal frameworks provide enforceability but risk rigidity in rapidly evolving contexts. An integrated approach is therefore necessary, combining algorithmic auditing with normative reasoning and adaptive legal instruments. Such convergence reflects the broader trend toward interdisciplinary AI governance, offering a pathway to align technological progress with societal values and human rights. The discussion underscores that ethics and bias in AI cannot be solved within disciplinary silos. Technical measures reduce but do not eliminate bias, philosophical theories clarify values but do not dictate implementation, and legal frameworks provide enforceability but risk rigidity in rapidly evolving contexts. An integrated approach is therefore necessary, combining algorithmic auditing with normative reasoning and adaptive legal instruments. Such convergence reflects the broader trend toward interdisciplinary AI governance, offering a pathway to align technological progress with societal values and human rights. In addition, the findings suggest that effective solutions require balancing short term technical fixes with long term structural reforms, ensuring that AI development is guided by principles of fairness and justice while remaining adaptable to rapid advances in technology.

V. CONCLUSION

This study has examined the problem of ethics and bias in artificial intelligence through the combined perspectives of technology, philosophy, and law, highlighting the ways in which disciplinary silos both contribute valuable insights and simultaneously limit comprehensive solutions. The analysis confirmed that technological approaches such as fairness metrics, bias detection algorithms, and explainability methods provide important safeguards but remain constrained by trade-offs between accuracy and fairness as well as by their inability to address deeper structural inequalities reflected in data. Philosophical inquiry was shown to be indispensable in clarifying concepts of justice, fairness, autonomy, and responsibility, ensuring that algorithmic interventions are guided by human values rather than purely technical criteria. At the same time, legal scholarship revealed the importance of enforceable standards, liability regimes, and institutional oversight, demonstrating that ethical concerns must be embedded in regulatory frameworks if they are to have meaningful impact. The results indicate that no single discipline can resolve the problem of AI bias in isolation, since technology requires normative grounding, philosophy requires practical application, and law requires technical feasibility. A key conclusion is that

convergence across disciplines is not only desirable but essential for producing AI systems that are both trustworthy and socially legitimate.

Hybrid frameworks that combine algorithmic auditing, ethical reasoning, and adaptive regulation are emerging as promising models, though challenges remain in aligning different definitions of fairness, balancing transparency with privacy, and harmonizing global standards in diverse cultural and legal contexts. The findings also suggest that future progress depends on sustained interdisciplinary collaboration, investment in cross sectoral dialogue, and a commitment to placing human rights and social justice at the core of AI governance. In aerospace, healthcare, law enforcement, education, and beyond, the societal implications of bias in AI underscore that this issue is not peripheral but central to the responsible development of intelligent technologies. By integrating technical, philosophical, and legal approaches, societies can move closer to ensuring that AI systems enhance rather than undermine equity, accountability, and human dignity. The study concludes that addressing ethics and bias in AI requires more than isolated interventions, it requires the construction of a comprehensive cross disciplinary framework capable of evolving alongside technological change while remaining anchored in enduring principles of fairness and justice.

VI. FUTURE WORKS

Future research on ethics and bias in artificial intelligence must move toward building integrated models that can bridge the gaps between technological, philosophical, and legal approaches. On the technological side, further development of fairness aware algorithms, explainable models, and auditing frameworks is required, particularly in high stakes applications such as healthcare, criminal justice, and finance where bias can have profound consequences. These innovations should be complemented by systematic studies that evaluate not only accuracy and efficiency but also long term social impact. From the perspective of philosophy, future work should deepen the dialogue on how abstract principles of justice, autonomy, and responsibility can be translated into practical design guidelines that engineers and developers can apply in real systems. Legal scholarship must also extend beyond regional or sector specific analysis to create adaptive and globally relevant governance frameworks that balance innovation with protection of rights. Another important direction lies in empirical research that investigates how different cultural contexts shape perceptions of fairness and accountability, highlighting the need for context sensitive solutions rather than one size fits all approaches. Collaborative research involving computer scientists, ethicists, and legal scholars will be essential in producing actionable frameworks that are normatively robust and practically feasible. The future of AI ethics research should focus on creating sustainable mechanisms that ensure algorithmic systems remain aligned with societal values while retaining the flexibility to adapt to evolving technologies.

REFERENCES

- [1] S. Barocas and A. D. Selbst, “Big data’s disparate impact,” *California Law Review*, vol. 104, no. 3, pp. 671–732, 2016.
- [2] M. Hardt, E. Price, and N. Srebro, “Equality of opportunity in supervised learning,” in *Advances in Neural Information Processing Systems (NeurIPS)*, pp. 3315–3323, 2016.
- [3] J. Kleinberg, S. Mullainathan, and M. Raghavan, “Inherent trade-offs in the fair determination of risk scores,” *Proceedings of Innovations in Theoretical Computer Science (ITCS)*, pp. 43:1–43:23, 2017.
- [4] J. Rawls, *A Theory of Justice*. Cambridge, MA: Harvard University Press, 1971.
- [5] L. Floridi and J. Cowls, “A unified framework of five principles for AI in society,” *Harvard Data Science Review*, vol. 1, no. 1, pp. 1–15, 2019.
- [6] S. Zuboff, *The Age of Surveillance Capitalism*. New York: PublicAffairs, 2019.
- [7] F. Pasquale, *The Black Box Society: The Secret Algorithms That Control Money and Information*. Cambridge, MA: Harvard University Press, 2015.

- [8] S. Wachter, B. Mittelstadt, and L. Floridi, “Why a right to explanation of automated decision-making does not exist in the General Data Protection Regulation,” *International Data Privacy Law*, vol. 7, no. 2, pp. 76–99, 2017.
- [9] European Commission, *Proposal for a Regulation Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act)*, Brussels, 2021.
- [10] B. Goodman and S. Flaxman, “European Union regulations on algorithmic decision-making and a right to explanation,” *AI Magazine*, vol. 38, no. 3, pp. 50–57, 2017.
- [11] J. C. Mitchell and H. Chen, “Regulating AI in the United States: Challenges and opportunities,” *Yale Journal on Regulation*, vol. 39, no. 1, pp. 89–123, 2022.
- [12] B. Friedman and H. Nissenbaum, “Bias in computer systems,” *ACM Transactions on Information Systems*, vol. 14, no. 3, pp. 330–347, 1996.
- [13] J. Buolamwini and T. Gebru, “Gender shades: Intersectional accuracy disparities in commercial gender classification,” *Proceedings of the 1st Conference on Fairness, Accountability and Transparency (FAT)*, pp. 77–91, 2018.
- [14] A. Obermeyer, B. Powers, C. Vogeli, and S. Mullainathan, “Dissecting racial bias in an algorithm used to manage the health of populations,” *Science*, vol. 366, no. 6464, pp. 447–453, 2019.
- [15] S. Jobin, M. Ienca, and E. Vayena, “The global landscape of AI ethics guidelines,” *Nature Machine Intelligence*, vol. 1, no. 9, pp. 389–399, 2019.
- [16] B. Mittelstadt, P. Allo, M. Taddeo, S. Wachter, and L. Floridi, “The ethics of algorithms: Mapping the debate,” *Big Data & Society*, vol. 3, no. 2, pp. 1–21, 2016.
- [17] C. Cath, S. Wachter, B. Mittelstadt, M. Taddeo, and L. Floridi, “Artificial intelligence and the ‘good society’: the US, EU, and UK approach,” *Science and Engineering Ethics*, vol. 24, no. 2, pp. 505–528, 2018.
- [18] S. Russell and P. Norvig, *Artificial Intelligence: A Modern Approach*, 4th ed. New York: Pearson, 2021.
- [19] T. Gillespie, *Custodians of the Internet: Platforms, Content Moderation, and the Hidden Decisions that Shape Social Media*. New Haven: Yale University Press, 2018.
- [20] M. Crawford and K. Paglen, *Excavating AI: The Politics of Images in Machine Learning Training Sets*. New York: AI Now Institute, 2019.
- [21] D. Leslie, “Understanding artificial intelligence ethics and safety: A guide for the responsible design and implementation of AI systems in the public sector,” *The Alan Turing Institute Report*, 2019.
- [22] T. O’Neil, “Algorithmic bias: A cross-disciplinary review,” *Philosophy & Technology*, vol. 34, no. 4, pp. 1023–1046, 2021.
- [23] R. Calo, “Artificial intelligence policy: A primer and roadmap,” *UC Davis Law Review*, vol. 51, no. 2, pp. 399–435, 2017.
- [24] S. Vallor, *Technology and the Virtues: A Philosophical Guide to a Future Worth Wanting*. Oxford: Oxford University Press, 2016.
- [25] OECD, *OECD Principles on Artificial Intelligence*. Paris: OECD Publishing, 2019.