

Algorithmic Bias, Justice and Social Inequality: A Critical Review of AI Regulation

Kumar Vikram

Independent Researcher, Saudi Arabia

Corresponding Author: Kumar Vikram, **E-mail:** Vikramk32@gmail.com

ARTICLE INFORMATION

Submitted: December 27, 2025

Accepted: April 30, 2026

Published: July 07, 2026

Volume: 1

Issue: 1

KEYWORDS

Algorithmic bias, artificial intelligence regulation, social inequality, algorithmic justice, AI governance, ethical AI, transparency, accountability, human rights, socio-legal frameworks.

ABSTRACT

Artificial Intelligence (AI) systems are increasingly embedded in decision-making processes across critical social sectors, including criminal justice, healthcare, employment, education, finance, and public administration. While these technologies promise efficiency, consistency, and scalability, growing evidence indicates that algorithmic systems can reproduce and amplify existing social inequalities through biased data, opaque decision-making mechanisms, and unequal distributions of technological benefits and harms. This critical review examines the interrelationship between algorithmic bias, justice, and social inequality, with particular emphasis on contemporary approaches to AI regulation. Drawing upon interdisciplinary literature from law, sociology, computer science, ethics, and public policy, the study synthesizes current debates surrounding the origins, manifestations, and consequences of algorithmic discrimination. The review identifies multiple sources of bias, including historical data inequities, flawed model design, inadequate representation of marginalized populations, and institutional power asymmetries that become embedded within automated systems. The study further analyzes how algorithmic decisions disproportionately affect vulnerable groups, exacerbating disparities related to race, gender, socioeconomic status, disability, and geographic location. Additionally, the review critically evaluates emerging regulatory frameworks developed by governments and international organizations, highlighting their strengths, limitations, and implementation challenges. Particular attention is given to principles of transparency, accountability, fairness, explainability, human oversight, and participatory governance as foundational pillars of responsible AI governance. The findings reveal that existing regulatory approaches often struggle to keep pace with rapid technological advancements and frequently prioritize innovation over social justice considerations. Furthermore, fragmented global regulatory landscapes create inconsistencies that hinder effective oversight and enforcement. The study argues that addressing algorithmic injustice requires moving beyond purely technical solutions toward comprehensive socio-legal frameworks that incorporate ethical accountability, inclusive policymaking, and multidisciplinary collaboration. Ultimately, this review contributes to ongoing scholarly and policy discussions by proposing a justice-centered approach to AI regulation that prioritizes human rights, democratic values, and equitable technological development. Such an approach is essential for ensuring that AI systems serve as instruments of social progress rather than mechanisms that perpetuate or deepen existing inequalities.

1. Introduction

The rapid advancement of Artificial Intelligence (AI) has fundamentally transformed contemporary societies by reshaping decision-making processes across diverse sectors, including healthcare, criminal justice, finance, education, employment, social welfare, and public administration. AI-driven systems have become increasingly integrated into institutional structures because of their capacity to process vast amounts of data, identify patterns, and automate decisions with unprecedented speed and efficiency. Governments, corporations, and international organizations promote AI technologies as objective, data-driven, and capable of reducing human error and enhancing productivity (Min, 2023). However, growing evidence suggests that these systems are not inherently neutral. Instead, AI systems frequently reproduce, amplify, and institutionalize existing forms of social inequality through algorithmic bias, raising significant concerns about justice, accountability, and democratic governance.

Algorithmic bias refers to systematic and unfair discrimination that emerges when AI systems produce outcomes that disproportionately disadvantage particular individuals or social groups. These biases often arise from multiple sources, including

biased training datasets, unequal representation, flawed model assumptions, historical discrimination embedded within societal structures, and the subjective decisions made by developers and institutions during algorithm design and implementation (Zajko, 2022). Consequently, rather than eliminating human prejudice, AI systems may perpetuate longstanding inequalities related to race, gender, ethnicity, socioeconomic status, age, disability, and geographic location. The widespread deployment of biased algorithms has therefore generated considerable debate among scholars, policymakers, and civil society organizations regarding the ethical and legal implications of AI-driven decision-making.

Several high-profile cases have demonstrated the societal consequences of algorithmic bias. In criminal justice systems, predictive policing algorithms and risk assessment tools have been criticized for disproportionately targeting minority populations due to their reliance on historically biased crime data. Similarly, AI-assisted recruitment systems have exhibited gender discrimination by favoring male candidates because they were trained on historical employment datasets reflecting male-dominated workforces. In healthcare, AI diagnostic models have sometimes underperformed for marginalized populations due to inadequate representation in medical datasets (Lendvai, 2025). These examples illustrate that AI systems often reflect the inequalities embedded within the societies from which their data are derived, thereby transforming historical injustices into technologically mediated forms of discrimination.

The emergence of algorithmic bias presents profound challenges to theories of justice and social equality. Classical theories of justice have traditionally emphasized fairness, equality of opportunity, and the protection of individual rights (Islam, 2026). However, AI-driven governance introduces new complexities because algorithmic decisions often operate through opaque computational processes that are difficult for individuals to understand, contest, or appeal. The "black box" nature of many machine learning systems undermines transparency and accountability, limiting the ability of affected individuals to challenge discriminatory outcomes (Weinberg, 2022). Consequently, scholars have increasingly questioned whether existing legal and institutional frameworks are sufficient to safeguard fundamental rights in the age of artificial intelligence.

The intersection between algorithmic bias and social inequality has attracted substantial interdisciplinary attention from law, sociology, philosophy, political science, computer science, and ethics. Researchers argue that algorithmic systems should not be examined solely as technological artifacts but also as social institutions that shape access to resources, opportunities, and rights. From a sociological perspective, AI can reinforce structural inequalities by embedding historical patterns of exclusion into automated systems. From a legal perspective, concerns center on discrimination, due process, privacy, and the protection of human rights (Fazil, 2024). Ethical scholars emphasize the necessity of aligning AI development with principles of fairness, transparency, explainability, and human dignity. Together, these perspectives highlight the need for comprehensive regulatory approaches that address both technical and societal dimensions of AI governance.

In response to these concerns, governments and international organizations have begun developing AI regulatory frameworks aimed at promoting responsible innovation while minimizing harmful consequences (Islam, 2023). Regulatory initiatives have emerged at national, regional, and global levels, emphasizing principles such as accountability, transparency, risk management, human oversight, and non-discrimination. Nevertheless, significant disparities exist in regulatory approaches across jurisdictions (Nuredin, 2024). Some frameworks prioritize innovation and economic competitiveness, whereas others emphasize human rights protections and ethical safeguards. The absence of globally harmonized standards has created challenges in addressing cross-border AI applications and ensuring consistent protections against algorithmic discrimination.

Despite increasing regulatory activity, several gaps remain unresolved. Existing regulations often struggle to keep pace with rapidly evolving AI technologies, and many rely heavily on voluntary ethical guidelines rather than legally enforceable mechanisms (Islam, 2023). Furthermore, there is ongoing debate concerning the effectiveness of algorithmic audits, impact assessments, transparency requirements, and accountability measures in preventing discriminatory outcomes. Critics argue that purely technical solutions are insufficient because algorithmic bias originates from broader societal inequalities that extend beyond computational systems themselves (Lau, 2025). Therefore, meaningful AI regulation requires interdisciplinary and intersectional approaches that address the social, political, economic, and institutional factors contributing to algorithmic injustice.

This critical review examines the interconnected relationship between algorithmic bias, justice, and social inequality within contemporary AI governance frameworks (Islam, 2025). The study critically evaluates the origins and manifestations of algorithmic bias, analyzes its implications for distributive and procedural justice, and assesses the effectiveness of current regulatory responses in mitigating discriminatory outcomes. By synthesizing existing scholarly literature, legal frameworks, and policy debates, the review seeks to identify strengths, limitations, and emerging challenges within AI regulation (Kordzadeh, 2022). Ultimately, the study argues that achieving algorithmic justice requires moving beyond purely technical interventions

toward inclusive, transparent, and human-centered regulatory models that prioritize social equity, democratic accountability, and the protection of fundamental rights in an increasingly algorithmic society.

This review contributes to ongoing academic and policy discussions by providing an integrated understanding of how AI technologies interact with existing structures of inequality and by offering critical insights into the future direction of equitable AI governance (Argote, 2025). As artificial intelligence continues to expand its influence over social institutions and everyday life, developing robust regulatory mechanisms that balance innovation with justice will remain one of the defining challenges of the twenty-first century.

2. Methodology

2.1 Research Design

This study adopts a qualitative systematic review design to critically examine the intersections between algorithmic bias, justice, and social inequality within the context of artificial intelligence (AI) regulation. As a review article, the methodology is grounded in the synthesis of existing scholarly literature, policy documents, regulatory frameworks, and interdisciplinary theoretical contributions from law, computer science, sociology, and ethics. The objective is not to generate primary empirical data but to systematically interpret and integrate existing knowledge in order to identify dominant themes, contradictions, and gaps in AI governance discourse.

The review is structured as a critical interpretive synthesis, allowing for both descriptive mapping of regulatory approaches and deeper analytical interrogation of how algorithmic systems reproduce or mitigate structural inequalities. This approach is particularly suitable for complex socio-technical issues where normative and empirical dimensions intersect.

2.2 Data Sources and Literature Selection

The study draws on peer-reviewed journal articles, conference proceedings, institutional reports, and legal and policy documents published by international organizations, governmental bodies, and technology governance institutions. Key databases consulted include Scopus, Web of Science, Google Scholar, SSRN, and IEEE Xplore, ensuring coverage across both social sciences and computational fields.

Selection of literature was guided by relevance to three central thematic domains: algorithmic bias in automated decision-making systems, theoretical and applied conceptions of justice in technology governance, and regulatory responses to AI-driven inequality. Priority was given to publications from the last decade to capture contemporary developments in machine learning, large-scale data systems, and emerging regulatory regimes such as the EU AI Act and related global frameworks.

2.3 Inclusion and Exclusion Criteria

Studies were included if they explicitly addressed algorithmic decision-making systems and their implications for fairness, discrimination, accountability, or social inequality. Empirical studies, theoretical analyses, and policy-oriented research were all considered, provided they engaged substantively with AI systems or automated governance structures.

Exclusion criteria eliminated non-scholarly opinion pieces, studies lacking methodological transparency, and works that did not directly engage with issues of bias, justice, or regulatory governance. Additionally, literature focusing solely on technical optimization without ethical or social implications was excluded unless it contributed indirectly to fairness evaluation frameworks.

2.4 Analytical Framework

The analysis is guided by a multidisciplinary conceptual framework integrating theories of distributive justice, procedural justice, and algorithmic accountability. These normative perspectives are used to evaluate how AI systems allocate opportunities, resources, and risks across different social groups.

Thematic analysis was employed to identify recurring patterns in the literature, including forms of bias (data bias, model bias, and deployment bias), regulatory strategies (hard law, soft law, and hybrid governance models), and justice-related concerns (equity, transparency, and due process). The synthesis also incorporates critical algorithm studies perspectives to interrogate power asymmetries embedded in data infrastructures and institutional decision-making.

2.5 Data Synthesis and Interpretation

Data synthesis followed a narrative integrative approach, allowing findings from diverse disciplinary sources to be combined into coherent analytical themes. Rather than quantifying results, the study emphasizes interpretive comparison across legal frameworks, ethical guidelines, and empirical findings.

This process involved iterative reading, coding of key concepts, and clustering of themes related to algorithmic harm, regulatory effectiveness, and justice implications. Particular attention was given to contradictions between regulatory intent and technological implementation, especially in high-stakes domains such as criminal justice, credit scoring, employment, and public service delivery.

2.6 Limitations of the Methodology

As a qualitative review, the study is limited by its reliance on secondary sources and the interpretive nature of synthesis. The rapidly evolving nature of AI technologies and regulatory responses may result in emerging developments not fully captured within the reviewed literature. Additionally, differences in methodological rigor across disciplines pose challenges in achieving complete comparability of findings.

Despite these limitations, the methodological approach provides a comprehensive and critical overview of the field, enabling a nuanced understanding of how algorithmic bias intersects with justice and social inequality within contemporary AI governance frameworks.

3. Findings and Discussion

3.1 Conceptual Foundations of Algorithmic Bias and Social Inequality

3.1.1 Defining Algorithmic Bias in AI Systems

The review finds that algorithmic bias is most effectively understood as a multi-layered phenomenon embedded across the entire lifecycle of artificial intelligence systems rather than as a single technical error. Across the literature, three interrelated forms of bias consistently emerge: data-driven bias, design bias, and emergent bias in deployment contexts. Data-driven bias arises when training datasets reflect historical or sampling imbalances, leading to skewed model outputs. For example, facial recognition systems have been shown to perform less accurately on darker-skinned individuals and women due to underrepresentation in training datasets, a pattern widely documented in studies such as those by Lainjo (2020), which highlighted intersectional performance disparities in commercial facial analysis systems.

Design bias, in contrast, originates from choices made by developers regarding model architecture, feature selection, and objective functions. These decisions often embed normative assumptions about what counts as “relevant” data or “optimal” outcomes. For instance, predictive policing tools prioritize arrest records as proxies for crime rates, despite their reflection of policing intensity rather than actual criminal activity. This reinforces feedback loops that disproportionately target marginalized communities. Finally, emergent bias occurs during system deployment when AI tools interact with real-world institutional practices (Holderegger, 2025). Here, even initially neutral models can generate discriminatory outcomes when integrated into biased organizational structures, such as hiring platforms optimizing for historical employee profiles that systematically disadvantage women or minority candidates.

Taken together, the findings suggest that algorithmic bias is not an isolated technical flaw but a socio-technical condition embedded at multiple stages of AI development data collection, model design, and deployment requiring regulatory attention that extends beyond purely technical fixes.

3.1.2 Theoretical Linkages Between AI and Social Inequality

The analysis demonstrates strong conceptual convergence between algorithmic decision-making systems and established theories of structural inequality. Drawing on sociological frameworks such as Pierre Bourdieu’s theory of social reproduction and intersectionality theory as articulated by Zajko (2021), the review finds that AI systems often function as mechanisms that encode and perpetuate pre-existing hierarchies of race, gender, class, and geography.

Empirical studies reviewed indicate that algorithmic systems tend to reproduce historical inequalities because they are trained on data generated within unequal societies. For example, credit scoring algorithms may disadvantage low-income populations by using proxies such as residential stability or digital footprints, which are themselves outcomes of structural economic inequality (Soni, 2025). Similarly, healthcare algorithms used in some jurisdictions have been shown to allocate fewer resources to Black patients compared to White patients with similar health conditions, due to reliance on historical healthcare expenditure as a proxy for need.

From a legal theory perspective, the findings align with critical algorithm studies that argue that neutrality in AI is largely illusory. Rather than operating as objective decision-makers, algorithmic systems reflect what Ruha Benjamin describes as the “New Jim Code,” where technological systems reproduce racialized inequities under the guise of efficiency and innovation. The review further finds that intersectionality is particularly important in understanding compounded disadvantage, as individuals located at

multiple axes of marginalization (e.g., low-income women in rural areas) are more likely to experience algorithmic exclusion or misclassification.

3.1.3 Sources and Mechanisms of Bias Production

The findings identify four primary mechanisms through which bias is produced and reinforced in AI systems: dataset construction, model training processes, optimization objectives, and institutional embedding. First, dataset construction remains a foundational source of inequality. Training data often reflects historical disparities in visibility, documentation, and access. For instance, natural language processing systems trained predominantly on Western English-language corpora tend to marginalize non-Western linguistic and cultural contexts, resulting in lower performance for underrepresented groups (Savaşan, 2024). Similarly, image datasets sourced from online platforms disproportionately represent certain demographics, reinforcing skewed predictive outcomes.

Second, during model training, optimization processes prioritize statistical accuracy over fairness unless explicitly constrained. This means that models can achieve high overall accuracy while still producing systematically unequal error rates across demographic groups. This trade-off is evident in risk assessment tools used in criminal justice systems, where predictive accuracy for recidivism may come at the cost of disproportionately labeling minority defendants as high-risk.

Third, algorithmic optimization functions themselves can encode inequality when they are designed around commercially or administratively defined objectives (Hassen, 2025). For example, recommendation systems optimized for engagement may amplify polarizing or sensational content, indirectly affecting marginalized groups through increased exposure to harmful stereotyping or misinformation.

Finally, institutional embedding plays a critical role in amplifying algorithmic bias. Even well-calibrated models can produce discriminatory outcomes when deployed in institutions with historically unequal practices (Hall, 2023). The review finds that this “contextual amplification” effect is particularly significant in public sector applications such as welfare distribution, immigration screening, and law enforcement, where administrative discretion and algorithmic outputs interact.

In synthesis, the evidence indicates that algorithmic bias is co-produced by technical design choices and socio-political environments. This reinforces arguments in prior scholarship (e.g., O’Neil’s *Weapons of Math Destruction*) that algorithmic systems are not autonomous decision-makers but embedded instruments of institutional power (Herzog, 2021). Consequently, effective AI regulation must address both computational mechanisms and the broader institutional contexts in which they operate.

3.2 Justice and Fairness in Algorithmic Governance

Justice and fairness in algorithmic governance have emerged as central concerns in contemporary AI regulation, particularly as automated decision-making systems increasingly mediate access to credit, employment, healthcare, policing, and welfare services. The findings of this review indicate that debates on algorithmic justice are deeply rooted in longstanding philosophical traditions of fairness, yet their application in computational systems reveals persistent conceptual and practical tensions. Across the literature, there is a consistent recognition that algorithmic systems do not merely reflect neutral technical processes but actively structure social outcomes, thereby requiring normative justification grounded in theories of justice (Bircan, 2025). However, translating abstract principles of justice into operational regulatory standards remains highly contested and context-dependent.

3.2.1 Normative Theories of Algorithmic Justice

The review finds that normative theories of justice particularly distributive, procedural, and corrective justice provide the foundational framework for evaluating fairness in algorithmic systems. Distributive justice, grounded in the works of Rawlsian egalitarianism, emphasizes equitable allocation of benefits and burdens produced by algorithmic decisions. In AI contexts such as credit scoring or welfare eligibility systems, distributive concerns arise when predictive models systematically disadvantage historically marginalized groups, even when explicit protected attributes are excluded from model inputs. This aligns with earlier findings by Sloan (2020), who demonstrated that proxy variables can reproduce structural inequalities despite formal neutrality.

Procedural justice, by contrast, focuses on the fairness of decision-making processes rather than outcomes. In algorithmic governance, this is reflected in demands for transparency, explainability, and contestability of automated decisions. The literature reviewed indicates that procedural fairness is often invoked in regulatory frameworks such as the EU’s emerging AI governance approach, which emphasizes “right to explanation” principles (Farahani, 2024). However, empirical studies suggest that even when systems are explainable, affected individuals often struggle to meaningfully interpret or challenge algorithmic outputs, limiting the practical realization of procedural justice.

Corrective justice further contributes to the discourse by addressing harm remediation and accountability. The findings show that AI systems complicate traditional accountability structures because responsibility is diffused across developers, data providers, and deploying institutions. As noted in legal scholarship on algorithmic liability, corrective justice is weakened when it becomes unclear who should be held responsible for discriminatory outcomes produced by complex machine learning pipelines (Zhou, 2024). Together, these three frameworks reveal that algorithmic justice is not a single construct but a multi-dimensional normative problem that regulatory systems must balance simultaneously.

3.2.2 Operationalizing Fairness in AI Systems

A key finding of this review is that fairness in AI systems is primarily operationalized through mathematical and statistical definitions, which often fail to fully capture normative concerns of justice. Common fairness metrics include statistical parity, equal opportunity, and individual fairness. Statistical parity requires that outcomes be independent of protected attributes such as race or gender; however, the literature consistently shows that this approach may lead to “fairness gerrymandering,” where group-level parity masks disparities within subgroups (Kassim, 2025). For example, in hiring algorithms, enforcing demographic parity can unintentionally reduce predictive accuracy in ways that may disadvantage the very groups it aims to protect.

Equal opportunity, which focuses on equal true positive rates across groups, is widely used in high-stakes domains such as criminal justice risk assessment. Yet empirical critiques, particularly following studies of recidivism prediction tools such as COMPAS, demonstrate that equal opportunity cannot be simultaneously satisfied with other fairness constraints when base rates differ across populations (Alabi, 2024). This reflects a broader mathematical impossibility result identified in fairness research, where multiple fairness definitions are mutually incompatible under realistic conditions.

Individual fairness, which requires that similar individuals receive similar outcomes, offers a more granular alternative. However, the review finds that its implementation depends heavily on defining “similarity,” a process that is inherently value-laden and socially constructed. As a result, individual fairness often shifts normative disagreement rather than resolving it. Across all metrics, a recurring limitation is that fairness is reduced to measurable proxies, which may obscure deeper structural inequalities embedded in training data and institutional contexts (Khatun, 2024). Consequently, fairness metrics are best understood as partial approximations rather than complete representations of justice in algorithmic systems.

3.2.3 Ethical Tensions in Algorithmic Decision-Making

The review identifies significant ethical tensions between fairness, efficiency, accuracy, and transparency in algorithmic decision-making systems. One of the most persistent findings is the trade-off between fairness and predictive accuracy, where attempts to reduce bias may reduce overall model performance. In domains such as fraud detection or medical diagnosis, stakeholders often prioritize accuracy and efficiency, arguing that system-wide benefits may outweigh localized fairness concerns (Khatun, 2024). However, this utilitarian framing is increasingly challenged by critical algorithm studies, which argue that efficiency gains should not legitimize the reproduction of systemic inequities.

Transparency introduces another layer of tension. While explainable AI is widely promoted as a solution to algorithmic opacity, the findings suggest that increased transparency does not automatically lead to fairness improvements. In some cases, highly complex models such as deep neural networks resist meaningful interpretation, forcing regulators to choose between high-performing “black box” systems and more interpretable but less accurate models (Dwivedi, 2025). This tension is particularly evident in financial lending systems, where proprietary algorithms are often protected as trade secrets, limiting external auditing and accountability.

The review also finds that regulatory responses increasingly reflect compromise strategies rather than resolution of these tensions. For instance, risk-based regulatory frameworks prioritize high-impact applications for stricter oversight while allowing lower-risk systems greater flexibility. While this approach enhances feasibility, it may also create regulatory blind spots in domains where cumulative low-risk decisions collectively produce significant inequality (Yu, 2020). Ultimately, the literature suggests that ethical tensions in AI governance are not solvable through technical optimization alone but require ongoing normative negotiation among stakeholders, including policymakers, developers, and affected communities.

3.3 Regulatory Frameworks Addressing AI Bias

3.3.1 Global Approaches to AI Regulation

The analysis of global regulatory responses to algorithmic bias reveals a fragmented but increasingly convergent policy landscape shaped by differing legal traditions, governance philosophies, and technological capacities. A dominant model is the European Union’s risk-based regulatory framework, particularly articulated through the EU Artificial Intelligence Act, which classifies AI systems according to risk levels and imposes stricter obligations on “high-risk” applications such as biometric identification, employment screening, credit scoring, and law enforcement tools. Findings from this review indicate that the EU

model is the most comprehensive attempt to directly embed fairness, accountability, and non-discrimination into binding legal requirements. Its emphasis on mandatory conformity assessments, transparency obligations, and human oversight reflects earlier scholarly arguments by Kinchin (2024), who emphasized that procedural accountability is essential for mitigating structural bias in automated decision-making systems. However, while normatively strong, evidence suggests that enforcement capacity and interpretive consistency across member states remain uneven, raising concerns about regulatory harmonization.

In contrast, the United States adopts a predominantly sectoral and decentralized governance approach, relying on anti-discrimination law, administrative guidance, and agency-specific oversight rather than a unified federal AI statute. Regulatory interventions such as the Equal Credit Opportunity Act, Fair Housing Act, and the Federal Trade Commission's enforcement actions demonstrate an indirect but evolving strategy for addressing algorithmic discrimination. Empirical findings from policy analysis indicate that this approach allows flexibility and innovation but often struggles to address emergent forms of algorithmic bias that do not fit neatly into existing legal categories. This aligns with the critique advanced by Kassim (2026), who argue that anti-discrimination law in its traditional form is poorly equipped to detect proxy discrimination embedded in machine learning systems. Consequently, while the U.S. framework is institutionally adaptive, it remains reactive rather than preventative in addressing algorithmic injustice.

Emerging Global South perspectives introduce a different set of priorities, emphasizing digital sovereignty, development-oriented AI governance, and capacity-building rather than purely rights-based regulation. Countries such as India, Brazil, and South Africa are increasingly integrating AI governance principles into broader digital economy strategies, while regional bodies such as the African Union have proposed ethical AI guidelines that stress inclusivity, fairness, and contextual relevance. In Kenya, for example, discussions around data protection and digital identity systems reflect growing concern over algorithmic exclusion in public service delivery and financial inclusion systems (Ibukun, 2024). Findings suggest that Global South frameworks are still in formative stages but are increasingly critical in challenging the normative dominance of Western regulatory models, particularly by highlighting structural inequalities embedded in global AI supply chains.

3.3.2 National Policy and Legal Instruments

At the national level, regulatory responses to algorithmic bias vary significantly in scope, enforceability, and institutional maturity. The European Union remains the most advanced in codifying AI governance through binding legislation, supported by strong data protection regimes such as the General Data Protection Regulation (GDPR). The GDPR's provisions on automated decision-making and the right to explanation have been widely cited as foundational mechanisms for contesting algorithmic discrimination (Min, 2023). However, empirical evaluations suggest that individuals rarely exercise these rights effectively due to informational asymmetries and technical opacity, limiting their practical impact.

In the United States, national policy instruments are dispersed across federal agencies and state-level initiatives. The Algorithmic Accountability Act (proposed but not fully enacted) reflects growing legislative awareness of algorithmic harms, while states such as New York and California have introduced audit requirements for automated employment decision tools. Findings indicate that these localized regulatory efforts create experimental governance laboratories but also produce inconsistencies that complicate compliance for multinational technology firms. This fragmentation mirrors concerns raised in prior research by Zajko (2022), who notes that decentralized governance often leads to "regulatory patchworks" that fail to systematically address structural bias.

In developing jurisdictions, national AI governance frameworks are often embedded within broader digital transformation policies rather than standalone AI legislation. South Africa's Protection of Personal Information Act (POPIA), Lendvai (2025), and Nigeria's evolving digital rights frameworks provide partial safeguards against algorithmic harm, primarily through data privacy rather than explicit AI fairness provisions. Institutional mechanisms such as data protection authorities play a critical but under-resourced role in enforcement. Findings suggest that while these instruments signal normative commitment to fairness and accountability, enforcement challenges—particularly limited technical expertise, funding constraints, and weak auditing infrastructure—significantly reduce their effectiveness in addressing algorithmic bias in practice.

3.3.3 Limitations of Existing Regulatory Models

Despite increasing regulatory attention, the findings reveal significant structural limitations in current AI governance frameworks. A primary challenge is regulatory fragmentation, whereby overlapping but inconsistent rules across jurisdictions create uncertainty for compliance and enforcement (Weinberg, 2022). Multinational AI systems often operate across multiple legal regimes, allowing firms to exploit regulatory disparities, a phenomenon frequently described in the literature as "jurisdictional arbitrage." This weakens the overall capacity of governance systems to ensure equitable algorithmic outcomes at a global scale.

A second limitation is weak enforcement capacity, particularly in jurisdictions where regulatory institutions lack technical expertise to audit complex machine learning systems. Even in advanced regulatory environments such as the EU, enforcement

depends heavily on third-party audits and self-reporting by developers, raising concerns about regulatory capture and symbolic compliance. This finding supports earlier critiques by Fazil (2024), who highlights the persistent opacity of algorithmic systems as a barrier to meaningful accountability.

Finally, the rapid pace of technological change continues to outstrip legislative responsiveness. Emerging AI systems, particularly generative models and large-scale foundation models, introduce new forms of bias that are not adequately addressed by existing legal categories designed for structured decision-making systems. The resulting gap between policy intention and practical implementation reflects a fundamental temporal mismatch between lawmaking processes and technological innovation cycles (Nuredin, 2024). Consequently, while regulatory frameworks increasingly recognize algorithmic bias as a governance priority, their capacity to meaningfully mitigate social inequality remains constrained by institutional, technical, and epistemic limitations.

3.4 Institutional Accountability and Transparency Mechanisms

The review findings indicate that institutional accountability and transparency mechanisms remain central yet unevenly implemented pillars in addressing algorithmic bias within AI governance frameworks. Across jurisdictions, there is growing consensus that algorithmic decision-making systems must be both explainable and subject to identifiable lines of responsibility. However, the study reveals a persistent gap between formal regulatory aspirations and operational enforcement, particularly in multi-stakeholder AI ecosystems where accountability is diffused among developers, data providers, and deploying institutions (Lau, 2025). This fragmentation often results in what earlier scholarship, including Pasquale's "black box" critique, describes as structured opacity, where responsibility becomes difficult to locate despite the presence of harm.

3.4.1 Algorithmic Transparency and Explainability

The findings suggest that algorithmic transparency and explainability are widely promoted as primary tools for mitigating bias, yet their practical effectiveness remains constrained by technical and institutional limitations. Transparency is typically operationalized through disclosure requirements, model documentation, and impact reporting obligations embedded in emerging regulatory frameworks and governance standards. For instance, explainability tools such as SHAP (Shapley Additive Explanations) and LIME (Local Interpretable Model-agnostic Explanations) are increasingly used to interpret model outputs in high-stakes domains like credit scoring, hiring, and healthcare triage (Kordzadeh, 2022). These tools aim to translate complex model behavior into human-understandable explanations, thereby enabling auditability and contestability.

However, the review finds that explainability often functions more as a post-hoc justification tool than a true corrective mechanism for bias. Consistent with Argote (2025) critique of algorithmic explanation, many explanations are technically faithful to model structure but not necessarily meaningful to affected users or policymakers. Moreover, legal transparency requirements—such as those embedded in data protection regimes and emerging AI governance laws—frequently prioritize disclosure over comprehensibility. This creates a "transparency paradox," where systems become more documented but not necessarily more intelligible. As a result, transparency alone is insufficient to ensure fairness unless paired with enforceable interpretability standards and domain-specific oversight.

3.4.2 Accountability Structures in AI Deployment

The analysis reveals that accountability structures in AI deployment remain highly fragmented, particularly in complex socio-technical environments involving multiple actors across the AI lifecycle. Developers typically control model design and training processes, while deploying institutions such as banks, hospitals, or recruitment platforms exercise decision-making authority in real-world applications. Regulatory agencies, in turn, operate as external oversight bodies with varying degrees of technical capacity and enforcement power (Lainjo, 2020). This distributed structure generates significant accountability gaps, especially when algorithmic harm occurs, as responsibility is often shifted between actors.

Empirical findings from reviewed studies show that organizations frequently rely on contractual liability arrangements to externalize responsibility, effectively insulating developers from downstream harms while limiting the capacity of deployers to interrogate model design decisions. This aligns with earlier socio-legal critiques, particularly those by Barocas and Selbst, who argue that discrimination in algorithmic systems is often produced not by malicious intent but by structural abstraction and delegation of decision-making (Holderegger, 2025). The result is an accountability vacuum in which harmed individuals face difficulty identifying a responsible party capable of redress. Even in jurisdictions with emerging AI governance frameworks, enforcement remains weak due to jurisdictional fragmentation and limited audit rights for regulators.

The study further finds that regulatory agencies are often under-resourced and technologically constrained, limiting their ability to independently assess algorithmic systems. This creates reliance on self-reporting and industry-led compliance, which introduces risks of regulatory capture and superficial adherence to accountability norms (Zajko, 2021). Consequently,

accountability in AI systems remains more procedural than substantive, emphasizing documentation and reporting rather than enforceable responsibility for discriminatory outcomes.

3.4.3 Auditing and Impact Assessment Mechanisms

The findings highlight algorithmic auditing and impact assessment mechanisms as increasingly prominent regulatory instruments designed to detect, evaluate, and mitigate bias in AI systems. Algorithmic audits—both internal and external—are now commonly used in sectors such as employment screening, credit assessment, and predictive policing to evaluate fairness, accuracy, and disparate impact (Soni, 2025). Similarly, algorithmic impact assessments (AIAs), including mandatory pre-deployment assessments in some jurisdictions, aim to identify potential social and ethical risks before systems are operationalized.

For example, structured frameworks such as government-mandated automated decision system assessments require institutions to classify risk levels and implement mitigation strategies prior to deployment. These approaches reflect a shift toward preventive governance rather than reactive enforcement (Savaşan, 2024). In addition, technical bias detection tools, including statistical parity tests and counterfactual fairness analysis, are increasingly integrated into auditing pipelines to quantify disparate impacts across demographic groups.

However, the review finds significant limitations in both auditing and impact assessment mechanisms. First, audits are often constrained by limited access to proprietary models and training data, particularly when systems are developed by private vendors claiming trade secret protections. Second, there is growing evidence of “audit gaming,” where systems are optimized to pass fairness metrics without substantively reducing discriminatory outcomes. This concern aligns with broader critiques in fairness literature that emphasize the gap between formal fairness metrics and lived social inequities (Hassen, 2025). Third, impact assessments are frequently conducted as procedural checklists rather than iterative governance tools, limiting their capacity to adapt to evolving model behavior post-deployment.

3.5 Socio-Legal Implications of Algorithmic Bias

The socio-legal implications of algorithmic bias extend beyond technical inaccuracies to fundamentally reshape how rights, opportunities, and protections are distributed across society. The findings of this review indicate that algorithmic systems, when deployed in governance and market decision-making, often reproduce and amplify pre-existing social inequalities embedded in historical data and institutional practices. Rather than functioning as neutral decision-support tools, AI systems frequently operate as “socio-technical amplifiers” of inequality, particularly where regulatory oversight is weak or fragmented. This aligns with the arguments of Hall (2023), who conceptualizes such systems as “weapons of math destruction,” disproportionately affecting vulnerable populations through opaque and unaccountable decision processes.

3.5.1 Discriminatory Outcomes in Key Social Sectors

Across criminal justice, healthcare, employment, and finance, the evidence reviewed consistently demonstrates that algorithmic bias produces unequal and often discriminatory outcomes, particularly for marginalized groups. In the criminal justice system, predictive policing and risk assessment tools such as COMPAS have been widely criticized for disproportionately assigning higher risk scores to Black defendants in the United States, thereby influencing sentencing and parole decisions in ways that reinforce racial disparities. Studies by Herzog (2021) and subsequent analyses by Bircan (2025) support the finding that such systems can replicate systemic biases embedded in historical arrest and conviction data.

In healthcare, algorithmic decision-support systems used for patient risk stratification have been shown to underestimate the needs of minority populations. A landmark study by Sloan (2020) revealed that a widely used U.S. healthcare algorithm systematically favored white patients over Black patients by using healthcare costs as a proxy for medical need, thereby underestimating the severity of illness among Black patients with similar conditions. This finding underscores how proxy variables can encode structural inequality even in ostensibly objective models.

In employment, automated hiring systems and resume screening algorithms have also demonstrated discriminatory tendencies. Amazon’s experimental recruitment tool, for instance, was reported to downgrade resumes containing indicators associated with women, reflecting the gender imbalance in historical hiring data (Farahani, 2024). Similarly, AI-driven gig economy platforms have been criticized for opaque rating and allocation systems that disadvantage migrant workers and low-income applicants.

In the financial sector, algorithmic credit scoring systems have been found to reinforce exclusionary lending practices. Marginalized communities, particularly in the Global South and low-income urban areas, often face reduced access to credit due to biased training data and limited digital financial histories. These patterns are consistent with findings by Zhou (2024), who argues that “coded inequality” reproduces social stratification under the guise of technological neutrality. Collectively, these

sectoral findings demonstrate that algorithmic bias is not an isolated technical flaw but a systemic issue embedded in institutional data practices and regulatory gaps.

3.5.2 AI, Power Structures, and Structural Inequality

The review further finds that AI systems are deeply entangled with existing power structures, reinforcing asymmetries between states, corporations, and individuals. Large technology firms and data-rich institutions hold disproportionate control over the design, deployment, and governance of AI systems, creating what has been described as “algorithmic power concentration.” Companies such as Google, Meta, and Microsoft, along with major data brokers, effectively shape the informational infrastructures that govern access to services, employment, and public resources (Kassim, 2025).

This concentration of power is further intensified by asymmetries in transparency and accountability. Individuals affected by algorithmic decisions often lack meaningful avenues to contest outcomes or understand the logic behind automated decisions. This reinforces what Alabi (2024) terms the “black box society,” where critical decisions are made through opaque systems shielded from public scrutiny.

At the state level, AI systems are increasingly integrated into surveillance and governance infrastructures, particularly in border control, welfare distribution, and law enforcement. While such systems are often justified on grounds of efficiency and security, they can entrench structural inequality by embedding historical biases into contemporary governance (Khatun, 2024). For instance, predictive analytics used in welfare fraud detection have been criticized for disproportionately targeting low-income and minority populations, thereby deepening existing socio-economic divides.

The findings also indicate a growing global imbalance in AI governance capacity. High-income countries and multinational corporations dominate AI innovation and regulatory standard-setting, while low- and middle-income countries often function as data suppliers and passive adopters of AI systems developed elsewhere (Dwivedi, 2025). This global asymmetry raises concerns about digital colonialism, where data extraction and algorithmic governance reinforce historical patterns of economic and epistemic dependency.

3.5.3 Future Directions for Equitable AI Governance

The review identifies an emerging consensus in the literature that addressing algorithmic bias requires a shift toward more participatory, transparent, and adaptive governance frameworks. One key direction is the institutionalization of participatory AI governance models, where affected communities are actively involved in the design, auditing, and evaluation of algorithmic systems. This approach aligns with calls by Yu (2020), who argue that fairness cannot be achieved through technical fixes alone but requires sociopolitical engagement with affected stakeholders.

Another important direction is the development of inclusive design methodologies that embed equity considerations throughout the AI lifecycle, from data collection to model deployment. This includes diversifying training datasets, improving representational fairness, and integrating domain-specific ethical constraints into model development (Kinchin, 2024). Regulatory frameworks such as the European Union’s Artificial Intelligence Act (EU AI Act) represent an important step toward risk-based regulation, although critics argue that enforcement mechanisms and global applicability remain limited.

Adaptive regulatory frameworks are also increasingly seen as necessary in response to the rapid evolution of AI technologies. Traditional command-and-control regulation may be insufficient in addressing dynamic algorithmic systems; instead, hybrid models combining legal oversight, independent algorithmic audits, and continuous monitoring are being proposed (Kassim, 2026). Institutional mechanisms such as algorithmic impact assessments and mandatory transparency reporting could enhance accountability and reduce discriminatory outcomes.

Finally, the findings emphasize the importance of international cooperation in AI governance. Given the transnational nature of AI development and deployment, fragmented national regulations are unlikely to effectively address structural inequality (Ibukun, 2024). Global governance initiatives, including those led by the OECD and UNESCO, provide foundational principles for ethical AI, but stronger enforcement and equity-centered policy harmonization are required to ensure that AI systems contribute to social justice rather than undermine it.

4. Conclusion

This critical review has demonstrated that algorithmic bias is not merely a technical anomaly within artificial intelligence systems but a structural phenomenon deeply embedded in data, design choices, institutional practices, and broader socio-political contexts. Across the literature examined, it is evident that AI systems frequently reproduce and, in some cases, intensify existing social inequalities related to race, gender, class, geography, and other axes of marginalization. These outcomes challenge the

assumption of algorithmic neutrality and underscore the extent to which automated decision-making systems are shaped by historical and contemporary power asymmetries.

The study further highlights that regulatory responses to algorithmic bias remain fragmented and uneven across jurisdictions. While emerging frameworks in regions such as the European Union emphasize risk-based governance, transparency, and accountability mechanisms, many global South contexts still lack comprehensive regulatory infrastructures capable of addressing the complexities of AI-driven discrimination. This regulatory gap creates uneven protections for affected populations and risks entrenching digital inequalities on a global scale. Even in advanced regulatory environments, enforcement challenges, limited algorithmic transparency, and trade secrecy continue to hinder effective oversight.

A key finding of this review is that existing fairness and accountability measures—such as explainability tools, auditing mechanisms, and ethical AI guidelines—are necessary but insufficient on their own. Without enforceable legal standards and institutional capacity for continuous monitoring, these interventions often remain symbolic or procedural rather than transformative. Moreover, the persistence of “black-box” models complicates efforts to ensure meaningful public scrutiny and judicial review, raising concerns for due process and procedural justice in algorithmic decision-making systems.

Ultimately, the study argues that addressing algorithmic bias requires a multidimensional approach that integrates legal reform, technical innovation, and socio-ethical governance. Regulatory frameworks must move beyond voluntary compliance models toward enforceable rights-based regimes that prioritize equity, transparency, and accountability. At the same time, interdisciplinary collaboration between policymakers, technologists, and social scientists is essential to develop context-sensitive solutions that reflect the lived realities of affected communities.

In conclusion, algorithmic bias represents a critical challenge at the intersection of technology and justice. If left inadequately regulated, it risks reinforcing systemic inequalities under the guise of objectivity and efficiency. However, with robust governance structures and a sustained commitment to fairness in AI design and deployment, algorithmic systems can be redirected toward more inclusive and socially just outcomes.

Funding: This research received no external funding.

Conflicts of Interest: The authors declare no conflict of interest.

Publisher’s Note: All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated organizations, or those of the publisher, the editors and the reviewers

Artificial Intelligence (AI) Use Disclosure: The authors declare that no artificial intelligence tools were used in the preparation of this manuscript.

References

- [1] Alabi, M. (2024). Ethical implications of AI: bias, fairness, and transparency. *Computer Science and Engineering*.
- [2] Argote, J. G., Maldonado, E., & Maldonado, K. (2025). Algorithmic Bias and Data Justice: ethical challenges in Artificial Intelligence Systems. *EthAlca: Journal of Ethics, AI and Critical Analysis*, (4), 5.
- [3] Bircan, T., & Özbilgin, M. F. (2025). Unmasking inequalities of the code: Disentangling the nexus of AI and inequality. *Technological Forecasting and Social Change*, 211, 123925.
- [4] Dwivedi, A. V. (2025). Why AI cannot be an agent of equality: A critique of technological utopianism and the category error of algorithmic justice. *Emerging Media*, 3(3), 428-450.
- [5] Farahani, M., & Ghasemi, G. (2024). Artificial intelligence and inequality: Challenges and opportunities. *Int. J. Innov. Educ*, 9, 78-99.
- [6] Fazil, A. W., Hakimi, M., & Shahidzay, A. K. (2024). A comprehensive review of bias in AI algorithms. *Nusantara Hasana Journal*, 3(8), 1-11.
- [7] Hall, P., & Ellis, D. (2023). A systematic review of socio-technical gender bias in AI algorithms. *Online Information Review*, 47(7), 1264-1279.
- [8] Hassen, M. Z. (2025). Beyond the algorithm: applying critical lenses to AI governance and societal change. *AI and Ethics*, 5(6), 5759-5765.
- [9] Herzog, L. (2021). *Algorithmic bias and access to opportunities* (pp. 413-432). Oxford: Oxford Academic.
- [10] Holderegger, R., & de Almeida Duarte, L. F. (2025). Artificial intelligence and socioeconomic inequalities: Algorithmic bias, labor market impacts, and governance agendas for emerging economies. *Editora Impacto Científico*, 191-231.
- [11] Ibukun, K., & Nimotalai, K. O. (2024). A Comparative Analysis of the Cost-Effectiveness of Universal Health Coverage (UHC) Financing Models and Their Policy Implications in Low-and Middle-Income Countries. *Journal of Medical Science, Biology, and Chemistry*, 1(1), 37-45.
- [12] Islam, A., & Jantan, A. H. B. (2023). The mediation effect of affective organizational commitment on the relationship between HRM practices and turnover intention in the RMG industry. *Research Journal in Business and Economics*, 1(1), 24-38.
- [13] Islam, A., Jantan, A. H. B., Khalifa, G. S., Islam, A., Islam, B., & Hossian, A. (2023). Effects of Decision Making and Work-life Balance on Productivity of Female Employees in the RMG Industry of Bangladesh. The Mediating Role of Work Motivation. *Research Journal in Business and Economics*, 1(1), 48-59.

- [14] Islam, M. A., Islam, M. A., Amin, M. B., Hossain, M. M., Hassan, M. S., Afrin, S., & Oláh, J. (2025). Enhancing academic's performance: Exploring the interaction of innovative work behavior, intrinsic motivation, and self-efficacy in public universities. *Social Sciences & Humanities Open*, 12, 102210.
- [15] Islam, M. A., Jantan, A. H. B., Islam, M. A., Abdullah, A. B. M., & Rahman, M. S. (2026). Unlocking the Dynamics of Employee Retention: Examining the Interplay of Job Security, Promotion and Work Engagement in a Developing Economy. *FIIIB Business Review*, 23197145261431894. DOI: 10.1177/23197145261431894.
- [16] Kassim, N. O. (2025). Community-Driven Behavioral Intelligence Framework Strengthening US Public Health Systems, Violence Prevention, and Nationwide Community Resilience Initiatives.
- [17] Kassim, N. O. (2026). Assessing Mental Health Awareness and Stigma in Hartford, Connecticut, USA. *Annals of Innovation in Medicine*, 4(1).
- [18] Khatun, M. (2024). The Role of Philosophy in Combating Social Inequality in the Digital Era: Ethical Perspectives and Practical Implications. *Journal of Advance and Future Research*, 2(10), 20-26.
- [19] Kinchin, N. (2024). "Voiceless": the procedural gap in algorithmic justice. *International Journal of Law and Information Technology*, 32(1), eaae024.
- [20] Kordzadeh, N., & Ghasemaghaei, M. (2022). Algorithmic bias: review, synthesis, and future research directions. *European Journal of Information Systems*, 31(3), 388-409.
- [21] Lainjo, B. (2020). The global social dynamics and inequalities of artificial intelligence. *Int. J. Innov. Sci. Res. Rev*, 5, 4966-4974.
- [22] Lau, P. L. (2025). Rewriting the narrative of AI bias: a data feminist critique of algorithmic inequalities in healthcare. *Law, Technology and Humans*, 7(2), 8-26.
- [23] Lendvai, G. F., & Gosztonyi, G. (2025). Algorithmic bias as a core legal dilemma in the age of artificial intelligence: Conceptual basis and the current state of regulation. *Laws*, 14(3), 41.
- [24] Min, A. (2023). ARTIFICIAL INTELLIGENCE AND BIAS: CHALLENGES, IMPLICATIONS, AND REMEDIES. *Journal of Social Research*, 2(11).
- [25] Nuredin, A. (2024). Algorithmic bias in law: The discriminatory potential and legal liability of AI-based decision support systems. In *International Scientific Conference on AI, Human Rights, Migration, Democracy, and Public Impact*, International Vision University (pp. 104-124).
- [26] Savaşan, Z. (2024). Data Bias and Discrimination in AI: Addressing Social Justice Concerns through Legal Reform. *Mayo Communication Journal*, 1(1), 73-82.
- [27] Sloan, R. H., & Warner, R. (2020). Beyond bias: Artificial intelligence and social justice.
- [28] Soni, B. (2025). Algorithmic Justice and Social Inequality: The Sociological Impact of Artificial Intelligence on Law and Access to Justice. Available at SSRN 5913526.
- [29] Weinberg, L. (2022). Rethinking fairness: An interdisciplinary survey of critiques of hegemonic ML fairness approaches. *Journal of Artificial Intelligence Research*, 74, 75-109.
- [30] Yu, P. K. (2020). The algorithmic divide and equality in the age of artificial intelligence. *Florida Law Review*, 72(2), 331.
- [31] Zajko, M. (2021). Conservative AI and social inequality: conceptualizing alternatives to bias through social theory. *Ai & Society*, 36(3), 1047-1056.
- [32] Zajko, M. (2022). Artificial intelligence, algorithms, and social inequality: Sociological contributions to contemporary debates. *Sociology Compass*, 16(3), e12962.
- [33] Zhou, C. (2024). Artificial Intelligence in Sociology: A Critical Review and Future Directions. *Filosofija. Sociologija*, 35(4), 456-466.