

Research Article

Ethical Considerations in AI-powered Social Innovation: Balancing Progress with Responsibility

Mohammed Zeinu Hassen* 

Department of Social Sciences, Addis Ababa Science and Technology University, Addis Ababa, Ethiopia

Abstract

Artificial Intelligence (AI) is increasingly integrated into social innovation strategies, offering transformative potential for addressing complex global challenges in sectors such as healthcare, environmental protection, and education. However, the deployment of these technologies raises profound ethical concerns that must be addressed to prevent unintended harm. This study employs a systematic literature review of academic and policy discourse published between 2020 and 2025 to critically examine the moral dimensions of AI-powered social innovation. The analysis focuses on the tension between the pursuit of technological efficiency and the imperative of social responsibility. The review identifies three primary ethical challenges. First, algorithmic bias frequently perpetuates and amplifies existing social inequalities, creating "automated injustice" where historical discrimination is encoded into future predictions. Second, the data-intensive nature of AI creates significant privacy risks, particularly for vulnerable populations, leading to potential surveillance and the erosion of informed consent. Third, an "accountability void" emerges due to the opacity of "black box" systems and the diffusion of responsibility among stakeholders, complicating the ability to seek redress for algorithmic harm. Synthesizing these findings, the paper argues that these are not isolated technical glitches but interconnected structural failures resulting from prioritizing scale over human dignity. Consequently, the study proposes a comprehensive framework for "Responsible AI" to guide practitioners, policymakers, and governance bodies. This framework is built upon three essential pillars: the mandatory adoption of a human-centered design philosophy, the establishment of genuine and continuous community partnerships, and the implementation of robust mechanisms for ongoing moral review and auditing. The study concludes that moving beyond superficial technical fixes to a holistic socio-technical approach is essential for building AI systems that are effective, fair, and aligned with human principles.

Keywords

Ethical AI, Social Innovation, Algorithmic Bias, Data Privacy, Accountability, AI Governance, Human-centric AI, AI for Social Good

1. Introduction

Artificial Intelligence (AI) is rapidly being integrated into the fabric of social innovation, presenting what many consider a paradigm shift in our ability to address complex global challenges. From precision agriculture models designed to

combat food insecurity in developing nations to AI-driven platforms that coordinate humanitarian aid in crisis zones, the application of intelligent systems promises unprecedented efficiency, scale, and insight [7]. Governments and

*Corresponding author: mohammed.zeinu@aastu.edu.et (Mohammed Zeinu Hassen)

Received: 15 November 2025; **Accepted:** 1 December 2025; **Published:** 26 December 2025



Copyright: © The Author(s), 2025. Published by Science Publishing Group. This is an **Open Access** article, distributed under the terms of the Creative Commons Attribution 4.0 License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution and reproduction in any medium, provided the original work is properly cited.

non-governmental organizations (NGOs) are increasingly turning to AI-driven tools to tackle deep-seated societal issues, leveraging their power to improve health outcomes in remote regions, optimize resource allocation for environmental protection, and create personalized educational pathways for marginalized students [8, 11]. This technological turn is fueled by a compelling vision: a future where data-driven solutions can succeed where decades of conventional human-led efforts have fallen short, offering a way to finally overcome seemingly intractable problems.

However, this enthusiastic embrace of technological progress often obscures a host of profound ethical quandaries that must be directly confronted. The tools we build are not neutral artifacts; they are imbued with the values, assumptions, and biases of their creators and the societies from which they emerge. An algorithm designed to efficiently distribute public assistance can, with the same logic, create new and automated forms of exclusion. A public health monitoring system intended to predict disease outbreaks can simultaneously function as an apparatus for pervasive surveillance. The central argument of this paper is that the technical capabilities of an AI tool cannot be divorced from its moral and social dimensions. The relentless drive for innovation must be counterbalanced by an equally rigorous commitment to moral reflection and ethical foresight. Without this equilibrium, our

best intentions to do good risk paving the way for new, insidious forms of unintended harm. As Stahl et al. [10] argue, we must build an ethos of responsibility directly into the design of these new tools and the institutions that deploy them. This is not a hindrance to progress; it is the only responsible path forward.

This paper undertakes a systematic review and critical analysis of the key ethical challenges that arise from the deployment of AI in social innovation contexts. Its purpose is to move beyond a general call for “ethical AI” and to dissect the specific mechanisms through which harm can manifest. We structure our investigation around three central research questions:

- 1) What are the primary ethical challenges, specifically concerning bias, privacy, and accountability, that arise from the use of AI in social innovation projects targeting vulnerable populations?
- 2) How do these distinct ethical challenges interact and reinforce one another to create systemic risks that are greater than the sum of their parts?
- 3) What actionable frameworks and governance structures can be proposed to mitigate these risks and guide the development of AI systems that are not only effective but also fair, just, and respectful of human dignity?

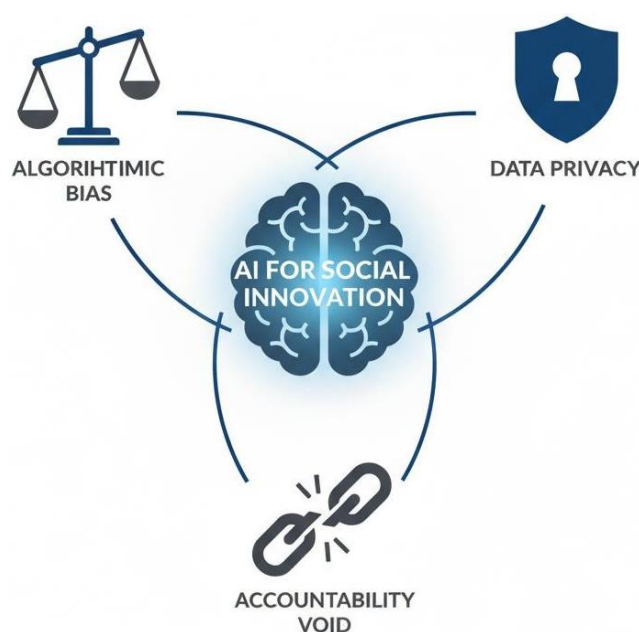


Figure 1. The Core Ethical Challenges of AI in Social Innovation. The Interconnected Issues of Algorithmic Bias, Threats to Data Privacy, and the Accountability Void Created by Opaque Systems are Vital Issues to be Examined.

To answer these questions, this paper will first outline the methodology used to review the relevant literature. It will then present the research findings, beginning with an acknowledgment of the profound promise of AI before delving into the three core ethical challenges identified in our analysis: algorithmic bias and the exacerbation of inequality; data privacy

risks and the specter of surveillance; and the accountability void created by opaque “black box” systems. Following this, a discussion section will synthesize these findings, arguing that they represent a deeply interconnected system of ethical failure. Finally, the paper will conclude by outlining a path forward, translating the research into a tangible framework for

responsible AI that offers clear implications for practitioners, policymakers, and governance bodies.

2. Research Methods

This study employs a systematic literature review as its primary methodology to identify, critically appraise, and synthesize the dominant ethical challenges associated with the application of AI in social innovation. This approach was chosen for its rigor and ability to provide a comprehensive, replicable, and transparent overview of the current state of academic and policy discourse on the topic.

2.1. Data Collection and Analysis

The review process began with a structured search strategy designed to capture a broad range of relevant scholarly work. The search was conducted across several major academic databases, including Scopus, Web of Science, ACM Digital Library, IEEE Xplore, and Google Scholar, to ensure coverage across computer science, social sciences, ethics, law, and public policy disciplines. The search query combined keywords related to the technology ("Artificial Intelligence," "AI," "machine learning," "algorithmic systems") with terms related to its application and ethical implications ("social innovation," "social good," "humanitarian aid," "ethical AI," "algorithmic bias," "data privacy," "accountability," "AI governance," "social responsibility").

The search was bounded by specific inclusion and exclusion criteria to ensure the relevance and quality of the selected literature. Inclusion criteria were: (1) peer-reviewed academic journal articles, conference proceedings, and significant white papers from reputable research institutions; (2) publications in the English language; (3) articles published between January 2020 and December 2025, to focus on the most contemporary issues in a rapidly evolving field; and (4) articles whose primary focus was the ethical, social, or governance dimensions of AI in noncommercial, social-benefit applications. Exclusion criteria included: (1) purely technical papers with no substantive discussion of ethical implications; (2) news articles, blog posts, and other nonacademic sources; (3) articles focusing exclusively on commercial AI applications (e.g., marketing, finance) without a link to social innovation; and (4) literature published before 2020, unless it was a foundational text consistently cited in recent work.

The initial search yielded several hundred articles. These were screened first by title and abstract, and then by full-text review, against the inclusion/exclusion criteria. The final corpus of literature was subjected to a detailed thematic analysis to synthesize the key ethical challenges. This qualitative process followed an inductive, three-stage coding procedure.

First, during the open coding stage, each article was read closely to identify key concepts, arguments, and illustrative examples related to ethical issues in AI for social innovation.

Initial codes were generated directly from the text, such as 'discriminatory training data,' 'lack of user consent,' and 'opaque decision-making.'

Second, in the axial coding stage, these initial codes were systematically compared and grouped into broader, more conceptual categories. For example, codes like 'discriminatory training data,' 'biased feature selection,' and 'automated redlining' were consolidated under the category 'Sources of Algorithmic Bias.' Similarly, codes related to data collection, monitoring, and user control were grouped under 'Privacy and Surveillance Risks.' This was a process of constant comparison, moving back and forth between the data and the emerging categorical framework to ensure coherence.

Finally, during the selective coding stage, these categories were further refined and integrated to identify the core, overarching themes that consistently appeared across the literature. It became evident that the diverse ethical concerns converged around three dominant and interconnected themes: (1) Algorithmic Bias and Inequality, (2) Data Privacy and Surveillance, and (3) The Accountability Void. These emergent themes achieved thematic saturation, meaning that further analysis of articles did not yield new core concepts, thus providing a robust organizing structure for the Research Findings section.

2.2. Methodological Limitations

While this systematic review provides a robust foundation for analysis, it is important to acknowledge its limitations.

First, the search was restricted to English-language publications, which may introduce a Western-centric or Anglophonic bias. This could lead to the underrepresentation of crucial perspectives, ethical frameworks, and case studies from non-English-speaking regions, particularly from the Global South, where many AI for social good initiatives are deployed.

Second, there is a potential for publication bias. Academic literature may have a tendency to publish studies with significant findings, potentially overrepresenting the documented failures of AI while underreporting on projects where ethical challenges were successfully navigated or were less pronounced.

Third, the rapidly evolving nature of AI technology means that any literature review is a snapshot in time. The search parameters (2020-2025) capture the contemporary discourse, but new models, ethical challenges, and governance solutions are emerging constantly. Therefore, findings should be understood within the context of the review period.

Finally, by focusing primarily on peer-reviewed articles and reputable white papers, this study may miss valuable insights from "grey literature," such as practitioner blogs, internal NGO reports, and activist manifestos, which can offer more immediate, on-the-ground perspectives of the challenges discussed. Future research could supplement this review with qualitative interviews or case studies to address these limitations.

3. Research Findings

The systematic review of the literature reveals a compelling, yet deeply conflicted, narrative. On one hand, there is widespread and well-founded optimism about the potential for AI to catalyze transformative social change. On the other, a consistent and growing body of evidence details severe and systemic ethical risks that accompany these technologies. The findings are organized to reflect this duality, first outlining the promise of AI before presenting a detailed analysis of the three primary challenges identified.

3.1. The Promise of AI in Social Innovation

The potential for AI to serve as a powerful tool for social good is a recurring theme across the literature. Its core strength lies in its ability to process immense volumes of complex data to identify patterns and generate insights that are beyond human capacity, turning information into actionable direction. In global health, this translates into tangible benefits. For instance, AI-driven diagnostic tools, trained on vast image datasets, can identify diseases like diabetic retinopathy or certain cancers with an accuracy matching or exceeding that of human experts. When deployed on simple mobile devices, these systems can bring high-level diagnostic capabilities to remote or underserved communities that lack specialists, democratizing access to care. This is particularly transformative in fields like women's health, where conditions like endometriosis often go undiagnosed for years; new AI tools are being developed to analyze symptoms and patient histories to significantly shorten this painful diagnostic journey [8, 12].

Beyond individual health, AI offers new paradigms for managing planetary health. Environmental challenges like urban flooding, exacerbated by climate change, require precise and timely warnings. Traditional prediction methods are often too slow and broad. In contrast, AI models can synthesize real-time data from weather satellites, river sensors, and urban drainage systems, combining it with topographical maps to generate highly localized flood predictions. This allows for targeted evacuations and the strategic placement of emergency resources, shifting disaster management from a reactive to a proactive posture [1]. The same methodology can be applied to tracking deforestation, predicting wildfires, and monitoring biodiversity, offering a powerful new ally in environmental protection.

In education, AI-powered learning platforms promise to resolve the fundamental tension of the traditional classroom: the inability of a single teacher to cater to the diverse learning styles and paces of every student. These platforms can track a student's progress in real time, offering personalized exercises, readings, and challenges tailored to their exact level. This not only keeps struggling students from falling behind but also keeps advanced students engaged. Furthermore, these tools can be powerful agents of inclusion, providing instant translation for non-native speakers, text-to-speech functionalities

for the visually impaired, and interactive learning modes for students with a range of disabilities. This vision of a world where technology, guided by human values, can create more equitable and effective solutions is the driving force behind the AI for social good movement.

3.2. Finding 1: The Pervasiveness of Algorithmic Bias and Exacerbation of Inequality

Despite its promise, a primary and deeply troubling finding from the literature is that AI systems frequently absorb, reproduce, and amplify existing societal biases. This phenomenon is not an occasional glitch but a systemic feature arising from the very way these systems are built and trained.

3.2.1. Flawed Data as the Root of "Automated Injustice"

The first source of this bias is flawed data. AI models learn to make predictions by identifying patterns in the data they are fed. If this training data reflects a world marked by historical injustices—such as racial, gender, and economic inequality—the AI will inevitably learn to replicate these patterns. The machine does not understand fairness; it only understands correlation. When historical patterns of discrimination are encoded in the data, the AI's predictions will carry that discrimination into the future, often with greater speed and on a much larger scale. This creates what the literature describes as automated injustice.

3.2.2. Biased Design and the Human Element in Code

The second source is biased design. The creation of an algorithm is a process of human choices. Developers decide which variables to include, what outcomes to optimize for, and how to define "success" or "risk." These are not neutral technical decisions; they are inherently moral ones. For example, a developer designing a loan application system might use a person's postal code as a data point. On the surface, this may seem harmless. However, if historical housing segregation means that certain postal codes are heavily populated by minority groups who have been systematically denied loans, the algorithm will learn to associate that postal code with high risk. The system will then deny loans to individuals not based on their financial history, but on their address—a new and automated form of redlining. The bias is not in the code itself, but in the human logic and societal structures that preceded it.

3.2.3. Pernicious Feedback Loops and "Technological Colonialism"

The harm caused by these biased systems falls most heavily on those who are already marginalized. A system designed to allocate public assistance, for instance, might be trained on administrative data showing that people who miss appointments are often later found to be ineligible. The machine

learns this pattern and begins to flag anyone who misses a single appointment as “highrisk,” leading to a suspension of benefits. The system is blind to the reasons a person might miss an appointment—a single mother unable to find child-care, an elderly person without transport, or someone working multiple jobs. The system sees only a missed data point and, in its cold, unthinking logic, manufactures hardship. This problem is severely compounded in a global context. The literature points to a form of “technological colonialism,” where AI models developed and trained in a few Western nations are deployed globally. A healthcare chatbot that does not understand local dialects, customs, or deep-seated social structures is not a useful tool; it is a foreign object that fails to connect with the lived reality of the people it is meant to serve, leading to social erasure rather than social innovation [2].

Critically, these biases are not static; they create powerful and pernicious feedback loops. An AI tool for child health diagnosis, if trained primarily on data from one demographic group, may misdiagnose conditions in children from other backgrounds. This misdiagnosis creates a skewed health record. Future versions of the AI, trained on this newly corrupted data, will become even more biased, actively contributing to health disparities it was designed to reduce [4]. A similar pattern is seen in education, where an AI designed to detect cheating might learn to flag the sentence structures of non-native English speakers as suspicious. These students are then subjected to greater scrutiny, creating a cycle of suspicion and penalty based on their linguistic background [5]. The system does not learn to be fairer; it learns to be more confident in its own prejudice.

3.3. Finding 2: Critical Risks to Data Privacy and the Emergence of Surveillance

The second major finding is that the data-intensive nature of AI creates a profound ethical paradox for social innovation. To build tools that help people, organizations must first collect enormous amounts of information about them. This is especially true for projects that work with the most vulnerable members of society, the sick, the poor, the elderly, and the marginalized. The data collected is not the anodyne information of e-commerce transactions; it is the deeply personal material of people's lives: their medical histories, financial struggles, daily movements, and personal relationships. The very act of collecting this data, even with the best intentions, opens a door to potential harm.

3.3.1. Data Breaches and the Vulnerability of Sensitive Information

The most immediate danger is the data breach. The centralized databases created for social innovation projects are a tempting target, containing a concentrated trove of sensitive information. A single breach could expose the private medical conditions of an entire community, leading to public shame, employment discrimination, and social stigma. The organi-

zations running these projects, often nonprofits or government agencies with limited resources—may lack the sophisticated cybersecurity measures of large corporations, making them soft targets. The cost of such a failure is measured not in dollars, but in human suffering and the destruction of trust, which is the foundation of any successful social project.

3.3.2. From Care to Control: The Specter of Pervasive Surveillance

Beyond the blunt force of a data breach, the literature highlights a more subtle but chilling danger: the specter of surveillance and social control. The same data used to help can also be used to watch. A system designed to monitor the health and safety of elderly people living alone also creates a detailed, minute-by-minute record of their private lives. A program that provides financial assistance to the poor can also be used to track their every purchase. In the hands of a benevolent organization, this data may be used for good. But the line between care and control is thin and easily crossed. This data can be shared with or demanded by other entities, such as law enforcement or government agencies, turning a project of assistance into an apparatus of judgment and surveillance. This creates a quiet slide from support to social scoring, a slow erosion of personal freedom in the name of the common good.

3.3.3. The Illusion of Consent and Coercion in Data Collection

This dynamic is further complicated by the issue of consent in a digital world. The literature strongly questions whether consent is truly informed when an individual is presented with pages of dense legal text and asked to click a single “agree” button. This creates an illusion of choice. The issue becomes even more acute when essential services are tied to the use of an AI-powered platform. A person in need of healthcare or welfare may be told they must use a specific application to receive it. Their “agreement” to the data collection practices of that application is not a free choice; it is a condition of survival. This form of subtle coercion undermines the moral foundation of the entire project. For example, an application designed to assist non-verbal individuals must, by its nature, gather some of the most intimate data imaginable: their thoughts, needs, and frustrations. For such a tool, consent cannot be a one-time event; it must be a continuous, respectful dialogue where the user has clear and simple control over their data at all times [13]. The literature calls for organizations to move beyond seeing data as an asset and to embrace their role as data guardians, operating under a strict code of ethics that prioritizes individual autonomy and protection above all else [6].

3.4. Finding 3: The Systemic Challenge of the Accountability Void

The third major finding from the literature is the emergence

of a profound “accountability void” when AI systems cause harm.

3.4.1. The “Black Box” Problem and the Opacity of AI Decisions

The increasing complexity of modern AI, particularly deep learning models, has given rise to the “black box” problem. These systems are not programmed with explicit rules; they learn by identifying patterns in vast amounts of data. The result is a system that can perform its task with remarkable accuracy but in a way that is often inscrutable to human minds. Even the engineers who designed the system cannot always trace the exact path of its internal logic or fully explain why it made one decision and not another. It just works—until it doesn't.

3.4.2. Diffusion of Responsibility and Automated Harm at Scale

This opacity shatters our traditional models of responsibility. In a human system, when a mistake is made, there is a chain of command and a person or group who can be held to account. A doctor who misdiagnoses a patient can be questioned; a loan officer who discriminates can be challenged. But who do you question when the decision was made by an algorithm? The literature describes a circular game of finger-pointing. The programmer may argue they only built the learning architecture, not what the machine learned. The organization that deployed the system may claim they relied on the expertise of the technology provider. The government agency that funded the project may point to layers of contractors. In this diffusion of responsibility, accountability becomes a ghost—everywhere and nowhere at once. This challenge is magnified by the issue of scale. A single flawed algorithm deployed by a national government or a large NGO can affect millions of people almost instantly. Traditional legal systems, built to address individual wrongs on a case-by-case basis, are ill-equipped to handle this kind of mass, automated harm. The harm is distributed, the cause is hidden inside a black box, and the evidence is often owned by the very organization that caused the harm. This creates a massive gap in justice. The literature calls for a new kind of regulatory structure involving proactive certification and independent, third-party auditing of high-stakes AI systems before they are deployed, akin to the rigorous trials required for medical devices [3].

3.4.3. Global Regulatory Gaps and “Accountability Shopping”

The problem also has a daunting global dimension. A social innovation project in one country might use an AI system developed by a company on another continent, using data stored on servers in a third. When something goes wrong, a complex tangle of international laws and jurisdictions emerges. There is no global regulatory body for AI and no international court for algorithmic harm.

This allows for a form of “accountability shopping,” where companies can operate in jurisdictions with the weakest rules, leaving those they harm with no meaningful path to redress [7]. For the individual wronged by an algorithmic decision, the path to justice is murky at best. How do you appeal a decision when you cannot know the reason it was made? This lack of recourse erodes public trust not only in the technology but in the institutions that use it. A public sector that cannot explain its own decisions has lost a vital connection to the people it is meant to serve [9]. This accountability void threatens the very soul of social innovation, risking a future where machines are always right, even when they are demonstrably wrong.

4. Discussion

Synthesizing the research findings reveals a critical insight: the challenges of algorithmic bias, data privacy, and the accountability void are not isolated, technical problems to be solved independently. Rather, they are a deeply interconnected and mutually reinforcing system of ethical failure that stems from a technological paradigm that consistently prioritizes computational efficiency and scale over human dignity and social context. Understanding their interplay is crucial to developing any meaningful solution.

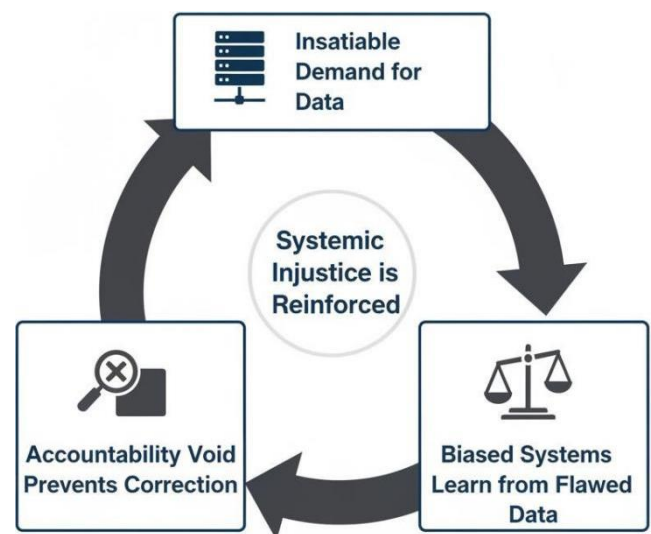


Figure 2. The Vicious Cycle of Ethical Failure in AI. The Demand for Vast Datasets Fuels Algorithmic Bias, While the Accountability Void Ensures These Biases Persist Without Correction, Creating a Self-reinforcing System that Disproportionately Harms Vulnerable Population.

The pervasiveness of algorithmic bias (Finding 1) is fundamentally fueled by the insatiable demand for vast datasets (Finding 2). The logic of “big data” assumes that more data is always better, yet it fails to account for the fact that this data is a reflection of an unequal world. The drive to collect and “datafy” every aspect of a social problem provides the very

material from which biased systems learn to discriminate. In turn, the accountability void (Finding 3) ensures that these biases can persist and deepen without check. When a system's decision-making is opaque and responsibility is diffused, there is no effective mechanism to identify, challenge, or correct the injustices it perpetuates. An unaccountable system can operate with biased logic indefinitely, making decisions that reinforce social stratification.

This creates a vicious cycle that is particularly damaging to the vulnerable populations that social innovation aims to serve. These are the communities whose lives are most likely to be represented in the flawed datasets used for training. They are the most exposed to the risks of surveillance and data misuse, as their access to essential services is increasingly conditioned on their participation in these data-gathering ecosystems. And, due to existing power imbalances, they have the least recourse when an automated decision harms them. The result is a paradox where the tools of social innovation can become instruments of social control, creating more efficient and technologically sophisticated systems of injustice.

This points to a failure that is not merely technical but deeply social and ethical. The current model of AI development often adheres to a logic of “solutionism”—the belief that any complex social problem can be solved by simply applying the right technology. This approach treats ethics as an after-

thought, a compliance checklist to be completed before launch, rather than as a foundational element of the design process itself. The findings collectively challenge this paradigm, arguing for a fundamental reorientation. Common technical fixes, such as mathematical “fairness metrics” or post-hoc “explainable AI” (XAI) techniques, while necessary, are insufficient. They can address symptoms but often fail to tackle the root causes, which lie in the power dynamics, social contexts, and human values that shape how technology is conceived, built, and deployed. A genuine path forward requires moving beyond technical fixes to embrace a holistic, socio-technical approach that embeds ethical considerations into the very architecture of innovation.

5. Implication and Conclusion

The research and discussion presented in this paper offer a clear warning, but more importantly, they illuminate a path forward. The findings lead to critical and actionable implications for practice, policy, and governance, which can be structured into a coherent framework for responsible AI in social innovation. This is not a vague call for more ethics, but a concrete guide for building and deploying these systems in a way that is just, fair, and serves human principles.

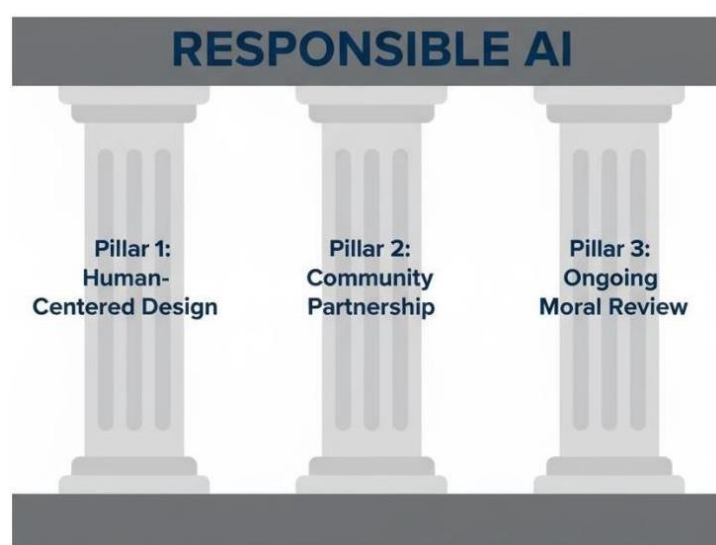


Figure 3. The Three Pillars of the Framework for Responsible AI. A robust Approach to Ethical AI Must Be Built Upon the Foundations of a Human-centered Design Philosophy, Genuine and Continuous Community Partnership, and a Commitment to Ongoing Moral Review.

5.1. A Framework for Responsible AI: Implications for Practice, Policy, and Governance

The proposed framework is grounded in the concept of “responsibility by design” [10] and stands on three interconnected pillars.

The primary implication for practitioners, the engineers, data scientists, and project managers building these systems, is the mandatory adoption of a human-centered design philosophy. This means every technological project must begin not with a question of technical possibility (“What can we build?”) but with a question of human need (“What should we build?”). This philosophy requires a deep, ethnographic understanding of the lived experiences, rights, hopes, and fears

of the people the technology is meant to serve, especially the most vulnerable among them. Methodologies like Participatory Design and Value Sensitive Design must become standard practice, not niche specializations. This approach serves as the most direct antidote to algorithmic bias because it grounds the system's logic in the messy, nuanced reality of diverse human lives rather than in a decontextualized dataset. It means embedding ethicists and social scientists directly into design teams from day one, not as consultants brought in at the end to "fix" the ethics [14].

For policymakers and the organizations deploying AI, the implication is a radical shift from superficial consultation to genuine and continuous community partnership. It is no longer enough to hold a few focus groups at the beginning of a project. Instead, models of co-creation and cogovernance are required. This involves bringing community members to the design table as empowered co-creators, valuing their local knowledge and lived experience as a vital form of data that is just as important as any quantitative dataset. This could take the form of community advisory boards with veto power, participatory budgeting for AI projects, and the creation of data trusts or data stewardship models where communities have democratic control and ownership over their data. This approach directly addresses the power asymmetries inherent in data collection and helps to make AI systems more transparent and trusted, as the community is part of the conversation from the very beginning. However, as the literature notes, this requires a shift in funding models to support long-term, trust-building engagement, which is a challenge for projects operating on tight deadlines and limited budgets [12].

The key implication for governance and regulation is the establishment of robust systems for ongoing moral review that extend far beyond a system's initial launch. Responsibility does not end at deployment; that is when it truly begins. This requires creating mandatory, independent, thirdparty auditing systems for all high-stakes AI systems used in the social sector. These audits must not only check for technical performance but must actively search for signs of bias, discriminatory impact, and unintended harm. Furthermore, governments and organizations must create clear, simple, and accessible channels for individuals and communities to report problems, ask questions, and seek meaningful redress when they have been wronged. These are not merely feedback forms; they are mechanisms for justice. This ongoing process of review turns an AI system from a static product into a living system, one capable of learning and improving not just its technical performance, but its moral performance as well. It builds humility into our technology, acknowledging that we do not have all the answers and that a true commitment to doing good requires a permanent commitment to listening and correction.

5.2. Conclusion

Artificial intelligence offers a new and powerful horizon

for social innovation. The capacity of these new tools to address some of the most enduring challenges of our time is undeniable. We stand at a moment of great possibility, a time when our ability to enact positive change is expanding at a remarkable rate. This paper began by acknowledging that promise, affirming the real and substantive good that AI can bring to the world. It is not a luddite's warning against the machine, but a call to build better, wiser, and more just machines.

Yet, this paper has argued that an uncritical embrace of this new power is a path to peril. The promise of AI is shadowed by the dangers of algorithmic bias, the specter of surveillance, and the void of accountability. These are not minor technical glitches to be patched in a future update; they are deep, structural problems that arise from the very nature of the technology and the human systems that create it. To ignore them is to risk building a future where our tools for social good become instruments of social harm, perpetuating injustice with unparalleled efficiency.

In response to these dangers, this paper has offered a path forward, a framework of responsibility. This concrete, action-oriented guide, standing on the pillars of human-centered design, community partnership, and ongoing moral review, provides the necessary guardrails to keep innovation on the right road. The final message of this paper is a simple but profound one. Technology is a tool.

It is a powerful, transformative tool, but a tool nonetheless. It has no will of its own. It has no moral compass. Its final worth, its effect on the world for good or for ill, comes not from the machine itself, but from the human choices that guide its creation and its use. The choice is ours.

We can be carried along by the current of technological progress, hoping for the best while ignoring the dangers. Or we can choose to be the navigators of that current, steering our new technologies toward a future that is not only more efficient but more just. We must choose to prioritize our human principles over mere technological advancement. We must build our machines not just to be smart, but to be wise. We must make this choice not out of fear, but out of a deep and abiding commitment to building a world that is worthy of our highest ideals. The future is not something that happens to us. It is something we build, one line of code, one moral decision at a time.

Abbreviations

AI	Artificial Intelligence
NGOs	Organizations
XAI	Explainable AI

Author Contributions

Mohammed Zeinu Hassen is the sole author. The author read and approved the final.

Conflicts of Interest

The author declares no conflicts of interest.

References

- [1] Bhanye, J. (2025). Flood-tech frontiers: smart but just? a systematic review of AI-driven urban flood adaptation and associated governance challenges. *Discover Global Society*, 3: 59.
- [2] Biju, P. R. and Gayathri, O. (2025). Indic approach to ethical AI in automated decision making system: implications for social, cultural, and linguistic diversity in native population. *AI & Society*.
- [3] Boretti, A. (2025). Ethical and practical considerations for AI-driven deep brain stimulation in mild cognitive impairment. *AI and Ethics*, 5: 3427–3436.
- [4] Chng, S. Y., Tern, M. J. W., Lee, Y. S., Cheng, L. T.-E., Kapur, J., Eriksson, J. G., Chong, Y. S., and Savulescu, J. (2025). Ethical considerations in AI for child health and recommendations for child-centered medical AI. *npj Digital Medicine*, 8: 152.
- [5] Gray, S. L., Edsall, D., and Parapadakis, D. (2025). AI-based digital cheating at university, and the case for new ethical pedagogies. *Journal of Academic Ethics*, 23: 2069–2086.
- [6] Gursoy, D., Başer, G., and Chi, C. G. (2025). Corporate digital responsibility: navigating ethical, societal, and environmental challenges in the digital age and exploring future research directions. *Journal of Hospitality Marketing & Management*, 34(3): 305–324.
- [7] Ifeanyichukwu, A., Vaswani, V., and Ekmekci, P. E. (2025). Exploring artificial intelligence-based distribution planning and scheduling systems' effectiveness in ensuring equitable vaccine distribution in low-and middle-income countries—witness seminar approach. *Discover Artificial Intelligence*, 5: 62.
- [8] Luz, K. P. and Lima, D. L. F. (2025). Empowering women through intelligent care: a narrative review of AI-driven digital innovations for endometriosis diagnosis, education, and equity. *Journal of Medical Imaging and Interventional Radiology*, 12: 15.
- [9] Mišić, J., van Est, R., and Kool, L. (2025). Good governance of public sector AI: a combined value framework for good order and a good society. *AI and Ethics*, 5: 4875–4889.
- [10] Stahl, B. C., Akintoye, S., Bitsch, L., Bringedal, B., Eke, D., Farisco, M., Grasenick, K., Guerrero, M., Knight, W., Leach, T., Nyholm, S., Ogoh, G., Rosemann, A., Salles, A., Trattig, J., and Ulnicane, I. (2021). From responsible research and innovation to responsibility by design. *Journal of Responsible Innovation*, 8(2): 175–198.
- [11] Trauth-Goik, A. (2021). Repudiating the fourth industrial revolution discourse: A new episteme of technological progress. *World Futures*, 77(1): 55–78.
- [12] Veloudis, S., Ryan, M., Ketikidi, E., and Blok, V. (2025). Responsible innovation in start-ups: entrepreneurial perspectives and formalisation of social responsibility. *Journal of Responsible Innovation*, 12(1): 2453251.
- [13] Wells, M. B. (2025). Empowering non-verbal individuals through AI-driven symbolic text prediction: a metaliteracy approach to communication and inclusion. *Discover Education*, 4: 360.
- [14] Willem, T., Fritzsche, M.-C., Zimmermann, B. M., Sierawska, A., Breuer, S., Braun, M., Ruess, A. K., Bak, M., Schönweitz, F. B., Meier, L. J., Fiske, A., Tigard, D., Müller, R., McLennan, S., and Buys, A. (2024). Embedded ethics in practice: A toolbox for integrating the analysis of ethical and social issues into healthcare AI research. *Science and Engineering Ethics*, 31: 3.

Biography



Mohammed Zeinu Hassen is a researcher in the Department of Social Sciences at Addis Ababa Science and Technology University in Ethiopia. His work centers on the significant ethical challenges presented by the integration of artificial intelligence into social innovation. In this paper, Hassen argues that

while AI has the potential to address complex global issues, it also carries inherent risks such as amplifying societal biases, threatening data privacy through surveillance, and creating an "accountability void" where responsibility for automated harm is unclear. He advocates for a responsible and ethical approach to AI development, proposing a framework built on the principles of human-centered design, active community partnership, and continuous moral review to ensure that these powerful technologies serve humanity in a just and equitable manner.

Research Field

Mohammed Zeinu Hassen: The intersection of technology and society, specifically focusing on Ethical AI, Social Innovation, Artificial Intelligence (AI), Social Responsibility, Algorithmic Bias, Data Privacy, Accountability, AI Governance, Human-Centric AI, and AI for Social Good.