

ETHICS OF ARTIFICIAL INTELLIGENCE: WHO IS RESPONSIBLE FOR A MACHINE'S MISTAKE

Yakhyoev Azizjon

Bukhara State medical institute

Abstract: This article explores the ethical and legal challenges arising from the growing autonomy of artificial intelligence (AI) systems. It examines the problem of accountability when AI errors lead to harm or social consequences, highlighting the ambiguity of assigning responsibility between developers, users, and autonomous algorithms. Examples from healthcare, transportation, and digital platforms are analyzed to demonstrate the complexity of ethical evaluation. The author concludes that the key to responsible AI lies in transparency, shared accountability, and the preservation of human oversight in all critical decision-making processes.

Keywords: artificial intelligence, ethics, accountability, responsibility, autonomous systems, transparency, regulation, AI governance.

Introduction

The rapid advancement of artificial intelligence has transformed the way societies operate, bringing both unprecedented opportunities and serious ethical challenges. Unlike traditional software, modern AI systems can act autonomously, learn from experience, and make decisions that their creators cannot fully predict or control. This autonomy raises a crucial question: who is responsible when an intelligent system makes a mistake? Cases involving self-driving cars, medical diagnostics, and automated content moderation have already shown that AI errors can lead to significant harm. These examples highlight the urgent need for ethical frameworks and legal mechanisms that clearly define accountability in the age of intelligent machines.

Research Objective

The purpose of this study is to analyze ethical principles and current approaches to assigning responsibility for the actions and errors of AI systems.

The objectives include:

1. Identifying the specific nature of AI errors compared to human mistakes.
2. Reviewing existing models of accountability in AI governance.
3. Evaluating the need for new legal and ethical frameworks for autonomous decision-making.
4. Proposing directions for responsible AI development and oversight.

Main Body

1. The Nature of AI Errors

AI systems are not moral agents. Their “decisions” are the result of algorithmic processing rather than conscious intent. Errors often occur due to biased datasets, flawed design, or limitations in model training. Unlike human mistakes, they are systematic and reproducible. This distinction challenges traditional moral and legal categories such as guilt, negligence, or intent.

2. Models of Accountability

Several models have been proposed to address AI accountability:

- Developer responsibility: The creators of algorithms are responsible for ensuring safety, fairness, and transparency in design.
- User responsibility: Operators must understand the limitations of AI tools and avoid delegating critical moral or legal decisions entirely to machines.
- Shared accountability: Increasingly, experts argue that responsibility should be distributed among all actors in the AI lifecycle — from data collection to system deployment and oversight.

Some have suggested granting AI systems a form of “electronic personhood,” but this idea remains controversial, as machines lack consciousness or moral agency.

3. Ethical Principles and Governance

International institutions such as the OECD, UNESCO, and the European Union have developed frameworks for trustworthy AI, emphasizing the principles of transparency, accountability, human oversight, and non-discrimination. Ethical design requires that AI systems remain interpretable and auditable, ensuring that humans can understand and, when necessary, override automated decisions. The principle of “human-in-the-loop” remains central — a reminder that ultimate responsibility must always rest with human beings.

4. Case Studies

- Autonomous vehicles: Accidents involving self-driving cars reveal the complexity of shared accountability between manufacturers, software developers, and human drivers.
- Medical AI systems: Diagnostic algorithms may misclassify diseases due to data bias, raising questions about liability between the physician and the technology provider.
- Algorithmic decision-making: Financial and social media algorithms can unintentionally discriminate against certain groups, demonstrating the ethical importance of dataset transparency and diversity.

Conclusion

Ethics in artificial intelligence is not only about preventing harm but also about defining responsibility in a world where decision-making is increasingly shared between humans and machines. AI lacks intent, conscience, and empathy — therefore, moral and legal accountability must remain within the human domain. The future of ethical AI depends on transparent algorithms, human oversight, and collective responsibility across the entire technological ecosystem. Only through these principles can artificial intelligence be developed and deployed safely, ensuring that technological progress aligns with human values.

References

1. Floridi L., Cowls J. A Unified Framework of Five Principles for AI in Society. — Harvard Data Science Review, 2021.

2. Boddington P. Towards a Code of Ethics for Artificial Intelligence. — Springer, 2017.
3. Bryson J. The Artificial Intelligence of the Ethics of Artificial Intelligence: An Introductory Overview for Law and Regulation. — Law, Innovation and Technology, 2019.
4. European Commission. Ethics Guidelines for Trustworthy AI. — Brussels, 2019.
5. Jobin A., Ienca M., Vayena E. The Global Landscape of AI Ethics Guidelines. — Nature Machine Intelligence, 2019.
6. Russell S., Norvig P. Artificial Intelligence: A Modern Approach. — Pearson, 2021.