

ARTIFICIAL INTELLIGENCE AND ETHICS: THE LINE BETWEEN BENEFIT AND RISK

*Djumakulova Shaxlo Davlyatovna,**Quyliyev Sardorbek Abdixoliq ugli**Computer science teachers of the Academic Lyceum of**Termez State University of Engineering and Agrotechnology*

Annotation: This text explores the ethical dimensions of artificial intelligence (AI), focusing on the delicate balance between its potential benefits and the risks it poses. It examines how AI impacts various sectors such as healthcare, education, security, and employment, while also highlighting concerns about algorithmic bias, lack of transparency, loss of privacy, social inequality, and the potential misuse of AI in warfare and disinformation. The discussion emphasizes the importance of human oversight, international cooperation, and ethical regulation to ensure that AI development aligns with human values and societal good.

Keywords: artificial intelligence, ethics, algorithmic bias, privacy, transparency, social impact, automation, employment, disinformation, AI in warfare, accountability, human rights, inequality, ethical regulation, technology risks, AI safety, artificial general intelligence, machine learning, data ethics, surveillance.

Artificial intelligence (AI) has developed rapidly in recent years, penetrating various aspects of our lives: healthcare, education, transportation, industry, finance, and even creative fields such as literature or music. These technologies make people's work easier, increase accuracy, and save time and resources. For example, with the help of AI, doctors are able to detect diseases at an early stage, or automated systems provide useful information to reduce traffic accidents. However, serious ethical questions arise against the background of these achievements. First, what basis do AIs rely on to make decisions? For example, facial recognition systems are known for incorrectly identifying certain racial groups. This leads to problems of discrimination and social injustice. The wrong decision-making of AI systems can have major consequences, especially in the judicial sphere.

The issue of responsibility assigned to AI also remains unclear. If an AI system makes a mistake, who is responsible: the programmer, the user, or the company that developed the system? The lack of clear answers to these questions is especially important for autonomous cars, robotic medical devices, or financial algorithms. Another ethical issue is privacy. AI-based systems collect, analyze, and use large amounts of personal data for profit. This leads to the violation of people's privacy and increased control and surveillance. In particular, the constant monitoring of citizens by the state or large corporations threatens civil liberties.

The loss of jobs is also not ignored. AI is taking over many ordinary jobs through automated systems. This is increasing social inequality and increasing unemployment. This problem is exacerbated by the fact that low-skilled workers are being replaced by AI. In addition, the development of AI poses the risk of transferring decisions that are of great importance in shaping the future of humanity to artificial systems. If humanity begins to

completely entrust its decisions to AI, a person may abandon his ability to think independently. This may lead to the loss of moral and cultural values. There is also no consensus on the ethical standards related to AI at the international level. While some countries are trying to establish strong control over these technologies, others have set the main goal of advancing in technological competition. This hinders the formation of global ethical standards and increases the risk of AI being used for wrong or dangerous purposes. In particular, weapons created on the basis of AI or disinformation distribution systems pose a real threat to human security.

Artificial intelligence (AI) is not only a technological revolution, but also a social phenomenon that has a fundamental impact on the life of society. Ethical issues related to AI are closely related to fundamental principles such as human dignity, rights, social equality, justice, and moral responsibility. How these principles should be preserved in technological progress has become a hot topic of discussion. The first important point is algorithmic transparency. Many AI systems operate on the principle of a "black box": that is, it is difficult to understand or explain how they came to a decision. This poses a major problem, especially in the fields of medicine, finance, and law, where important decisions are made. If an algorithm decides not to grant a loan or misdiagnoses a patient, a person cannot question this decision, because the internal workings of the system are not transparent. In this situation, ethical responsibility disappears, and human control is weakened.

The second problem is the prospects for artificial intelligence or general artificial intelligence (AGI). AGI is a type of AI that has the ability to think, learn, and make decisions like humans. Although this technology does not yet exist in practice, research is underway to create it. This raises the following questions: if a powerful artificial intelligence emerges from humanity, who will be responsible for its actions? Will it respect our interests or pursue its own interests? What will be the consequences if AGI gets out of control? The third main area is the problem of "deepfake" and false information. AI can be used to create images, videos, and audio signals that are very similar to reality, but are completely fake. These technologies lead to the spread of disinformation on social networks and in the media, political manipulation, and the incorrect formation of public opinion. Such technologies can become a dangerous weapon, especially in elections, wars, or crisis situations.

The fourth issue is global inequality. AI technologies are being developed mainly in developed countries, enhancing their economic and military superiority. On the other hand, developing countries do not have sufficient resources to either develop or effectively implement these technologies. As a result, global social and economic disparities are deepening through AI. From an ethical point of view, these technologies should serve the benefit of all humanity, but this principle is not yet fully implemented.

Also, the "learning" properties of AI systems sometimes reinforce existing social injustices by replicating human behavior. For example, if AI systems are trained on data based on pre-existing gender or racial discrimination, they may perpetuate the same discrimination. This will lead to the re-emergence of pre-existing problems in a modern form. In general, artificial intelligence opens up enormous opportunities for humanity. However, in order to use these opportunities correctly and fairly, it is necessary to resolve ethical issues related to AI, develop clear rules and control mechanisms. Finding a balance between benefits and risks is a key factor in ensuring that AI is beneficial to society. In this regard, the active participation of

not only technical experts, but also lawyers, philosophers, politicians and the general public is important.

References:

1. Russell, S., & Norvig, P. (2020). *Artificial Intelligence: A Modern Approach* (4th ed.). Pearson.
2. Bostrom, N. (2014). *Superintelligence: Paths, Dangers, Strategies*. Oxford University Press.
3. Floridi, L. (2013). *The Ethics of Information*. Oxford University Press.
4. Jobin, A., Ienca, M., & Vayena, E. (2019). *The global landscape of AI ethics guidelines*. Nature Machine Intelligence, 1(9), 389–399.
5. Moor, J. H. (2006). *The Nature, Importance, and Difficulty of Machine Ethics*. IEEE Intelligent Systems, 21(4), 18–21.