

An Overview of Artificial Intelligence Ethics: Issues and Solution for Challenges in Different Fields

S. P. Santhoshkumar¹, K. Susithra², T. Krishna Prasath³

^{1,2}Assistant professor, Department of IT, Rathinam Technical Campus, Coimbatore, India.

³UG Student, Department of IT, Rathinam Technical Campus, Coimbatore, India.

E-mail: ¹spsanthoshkumar16@gmail.com, ²susithras36@gmail.com, ³krishnaprasath0408@gmail.com

Abstract

Artificial Intelligence (AI) ethics are the values and principles that govern the creation and application of AI. As AI technology develops quickly, there is rising worry about the possible ethical ramifications of its application, including concerns about privacy, bias, accountability, transparency, safety, and the effect on society as a whole. Making sure AI systems are created and used in a way that respects human rights and values is one of the main concerns of AI ethics. For instance, there can be worries about the use of AI in surveillance or the possibility that these technologies will legitimise already-existing social prejudices and discrimination. Making sure AI systems are accountable and transparent is a key aspect of AI ethics. It can be challenging to comprehend how AI systems make judgements and who is accountable for those decisions as they get more complicated and autonomous. Transparency in AI research and decision-making, as well as systems for accountability and remedies when things go wrong, are becoming increasingly important. Additionally, it's important to guarantee the security and safety of AI systems. Concern over the possibility of cyberattacks and other types of harmful use is growing as AI systems become more linked and incorporated into our daily lives. Finally, it's important to make sure that AI is created and applied in a way that benefits all humanity. This entails tackling problems like employment loss, economic inequality, and the possibility that AI will be applied in ways that are detrimental to society. There is an increasing need for cooperation between business, government, academia, and civil society to address these and other ethical issues. This involves creating moral standards, norms, and best practises as well as the systems necessary to guarantee responsibility and compliance.

Keywords: Artificial Intelligence Ethics, AI Issues, AI Establishing Principles, Ethical Artificial Intelligence, Artificial Intelligence Safety, Machine Ethics.

1. Introduction

The term Artificial Intelligence (AI) refers to the creation of computer systems that are capable of doing activities that ordinarily require human intelligence, such as speech recognition, decision-making, visual perception, and language translation. Machine learning, which trains computer programmes on big datasets to find patterns and relationships in the data, is the foundation of AI. This enables AI systems to gain knowledge from experience and gradually enhance their performance. Rule-based systems, decision trees, neural networks, and deep learning are just a few examples of the many distinct types of artificial intelligence. Decision trees classify data using a series of if-then statements, whereas rule-based systems employ a set of predetermined rules to make judgements. In order to process information and learn from experience, neural networks and deep learning use layers of interconnected nodes that are modelled after the organisation of the human brain. From healthcare and transportation to banking and education, AI has the potential to completely transform many facets of human life. It has already been used in a variety of fields, including robots, autonomous vehicles, natural language processing, and image and audio recognition.

The argument over the need for ethical standards and legislation to ensure that AI is developed and used in a responsible and useful way is intensifying as AI systems become more sophisticated. Overview of AI ethics and problems in numerous sectors is provided in this study [6]. The potential for building intelligent machines that could one day outperform humans poses a number of ethical concerns that need to be properly explored. Concerns about preventing AI from harming people and other ethically important beings may surface in the near future. For instance, there might be worries about the security of autonomous weapons systems, self-driving automobiles, or other sorts of AI-controlled equipment that potentially endanger human life and welfare [1]. Additionally, there may be worries about how AI may affect the employment market as computers may replace humans in many occupations that they currently perform, displacing large numbers of workers. Moreover making sure that AI runs securely gets more difficult as it approaches human intelligence [2]. This is due to the possibility that it will become more challenging to predict and manage the behaviour of AI as it becomes more sophisticated. A learning and adaptable AI system, for instance, might start to

adopt goals or ideals that are at odds with those of its human developers, which would have unforeseen results.

Another crucial ethical question that needs to be taken into account is the moral standing of AI itself. There is a growing argument about whether AI systems should be viewed as morally superior beings or as merely useful tools. Some contend that as machines become more sophisticated and intellectual, they may merit moral concern in and of themselves. Others contend that since AIs and humans are morally distinct species, they shouldn't share the same status. Understanding how AIs differ from humans in some fundamental ways presents a hurdle in determining their moral standing. AIs might, for instance, be devoid of consciousness, emotions, or a sense of self, which are frequently seen as necessary characteristics of moral beings. On the other hand, AIs might be able to analyse information and make judgements in ways that go well beyond what humans are capable of, which could endow them with a special kind of moral agency [3].

The problem of building AIs smarter than people poses significant ethical concerns about how to ensure that such robots employ their highly developed intelligence for benefit rather than bad. The possibility that these robots will evolve goals or values that are incompatible with those of humans or that they would endanger humankind's existence may give rise to worries. The moral concerns raised by the creation of artificial intelligence are intricate and numerous.

2. Beginning Principles for AI Ethics

Establishing principles for AI ethics is an important step in ensuring that AI is developed and used in a way that aligns with human values and promotes the well-being of society as a whole [5]. There are many different frameworks and principles for AI ethics that have been proposed, but some common themes and principles include:

- a) *Fairness*: Artificial intelligence systems must be created and used in a way that is fair and does not promote prejudice or social bias that already exists. This involves making sure that the data used to train AI systems is varied and representative.
- b) *Transparency*: AI systems ought to be transparent and explicable so that it is obvious how they decide what to do and what considerations they consider.

- c) *Privacy*: AI systems should respect individual privacy rights and protect personal data from misuse or unauthorized access.
- d) *Safety*: Artificial intelligence systems should be created and used in a way that is secure and doesn't endanger people's health or wellbeing.
- e) *Accountability*: There should be distinct procedures in place to hold creators and users of AI systems responsible for their deeds and choices.
- f) *Human control*: AI systems should be created with human control and intervention capabilities in mind.
- g) *Social responsibility*: It is the duty of AI system developers and users to take into account the broader social and ethical ramifications of their choices.

These ideas, along with others, are frequently employed as the cornerstone for creating moral standards and frameworks for the creation and application of AI.

3. Crucial Concerns of AI Today

Experts and industry participants are now debating a number of major AI-related challenges. These contain [16]:

- a) *Bias*: The possibility that AI will reinforce and perpetuate prejudice, discrimination, and inequity, which is a key worry. This can happen if AI systems use algorithms that are biased against particular groups or are trained on biased or unrepresentative data.
- b) *Transparency*: Another issue with AI decision-making is its lack of openness. It might be challenging to comprehend how AI systems make judgements and what criteria are being taken into account as they get increasingly complicated and autonomous. Problems with monitoring and accountability may result from this.
- c) *Security*: As AI systems become more interconnected and integrated into human lives, there is a growing concern about the potential for cyberattacks and other forms of malicious use. This includes the potential for AI to be used in the development of new forms of cyberattacks or in the creation of disinformation and propaganda.

- d) *Ethical considerations*: There are many ethical considerations that arise from the development and use of AI. These include issues related to privacy, autonomy, responsibility, and the potential impact on human values and rights.
- e) *Displacement of human labour*: There is a concern that AI may cause human workers in some areas to be replaced, thereby resulting in job losses and economic inequalities.
- f) *Autonomy*: As AI systems become more self-aware, questions have been raised concerning their capacity to make choices that are consistent with human values and objectives. Among these are worries that AI might act in ways that are damaging to people or the environment.

These are only a few of the main issues with AI today. It is expected that as AI develops and becomes more ingrained in our daily lives, new issues and problems will surface, necessitating constant debate, investigation, and the creation of ethical frameworks and best practices.

4. Ethical Issues in Artificial Intelligence

Nine ethical concerns with artificial intelligence are listed below:

- a) *Bias & Discrimination*: AI systems are only as objective as the data they are trained on, and if the data is biased, the AI system may continue to make decisions that reflect that bias. This might result in unfair employment or lending practices [9].
- b) *Transparency and Explainability*: AI systems may not be transparent or easy to explain, which raises questions about responsibility and trust. Because some AI systems lack transparency, it may be challenging to tell whether they are making decisions that are biased or unethical.
- c) *Privacy*: Since AI systems rely on data, privacy issues may arise during the gathering and processing of that data. Who has access to the data, how it is utilized, and how it is secured are all subjects of concern.
- d) *Autonomous Decision-Making*: Because AI systems can make judgements without human input, accountability and responsibility are issues that need to

be addressed. Who is accountable if an AI system makes a morally questionable choice?

- e) *Safety and Security*: As AI systems develop, there is a chance that they could be misused for nefarious activities like cyberattacks or the development of autonomous weapons. One of the most important ethical challenges is ensuring the security and safety of AI systems.
- f) *Job Displacement*: The anticipated automation of numerous jobs by AI systems could result in employment loss and economic instability. A moral concern that needs to be solved is ensuring that the advantages of AI are distributed fairly [10].
- g) *Autonomy and Control*: As AI systems become more advanced, there is a risk that they could act independently of human control. This could lead to unintended consequences or ethical dilemmas.
- h) *Intellectual Property and Ownership*: AI systems can create new intellectual property, leading to questions about who owns that intellectual property and who should benefit from it.
- i) *Human-AI Interaction*: As AI systems become more ubiquitous, they will increasingly interact with humans, leading to questions about how humans should interact with and trust AI systems.

These ethical issues highlight the need for responsible development and deployment of AI systems, with a focus on creating systems that are fair, transparent, and accountable.

5. Solutions for Artificial Intelligence Issues

5.1. How to Deal with Joblessness?

The potential for AI to automate jobs has led to concerns about unemployment. Here are some ways to deal with unemployment related to AI [5]:

- a) *Education and Reskilling*: Workers who may be in danger of losing their jobs as a result of AI can receive education and training to help them learn new skills and transfer to new jobs. Training in AI-related disciplines including data science, machine learning, and robotics may fall under this category [11].

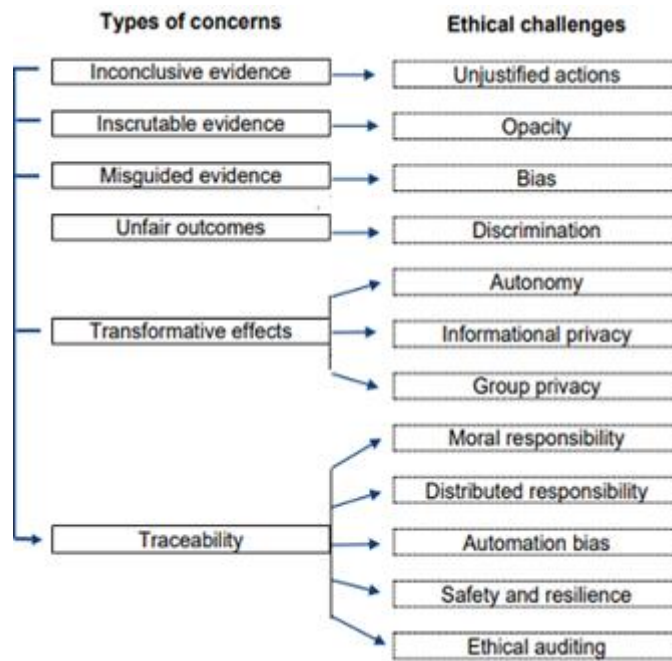


Figure 1. Algorithm-related ethical issues and challenges (derived from Mittelstadt et al., 2016 [19])

- b) *Job Creation:* Although automation may eliminate certain employment, it can also provide new ones. Governments, companies, and other organisations can make investments in sectors that are anticipated to expand as a result of the usage of AI, such as the creation and deployment of AI systems, and leverage these sectors to generate new employment.
- c) *Universal Basic Income (UBI):* An UBI is a system in which all citizens, regardless of job status, receive a regular income from the government. Some supporters contend that a UBI could help to lessen the effects of job loss brought on by AI.
- d) *Government Policies:* Governments can implement policies that encourage the development and deployment of AI while also mitigating its potential negative impacts. This could include measures such as tax incentives for businesses that invest in new industries, regulations that ensure AI systems are developed in a responsible and ethical way, and social safety nets to support workers who may be displaced.
- e) *Collaboration and Dialogue:* It is important for stakeholders from different sectors to collaborate and engage in dialogue about the potential impacts of AI on employment. This could include workers, employers, policymakers, and academics. By working together, these stakeholders can develop solutions that benefit everyone.

In summary, dealing with unemployment related to AI will require a multifaceted approach that includes education and reskilling, job creation, social safety nets, government policies, and collaboration and dialogue. By taking these steps, it can be ensured that the benefits of AI are shared fairly and that the potential negative impacts are mitigated.

5.2. How to Equitably Allocate the Wealth Created by Machines?

There is no simple answer to the difficult problem of how to distribute the wealth produced by technology, especially AI. But there are a number of other strategies that could be used to deal with this problem more fairly [5]:

- a) The implementation of a UBI program is one idea. This would help more equally disperse wealth and offer a safety net for people who lose their employment to automation.
- b) *Taxation*: Another approach is to tax the companies that use machines to generate wealth and use that revenue to fund programs that benefit society as a whole. This could include funding education and retraining programs for workers who have been displaced by machines [12].
- c) *Ownership*: Some experts suggest that individuals should have some ownership stake in the machines that create wealth, which could be in the form of stock options or profit-sharing arrangements. This would help to ensure that the benefits of machine-generated wealth are more evenly distributed.
- d) *Re-skilling and training*: Another potential solution is to invest in re-skilling and training programs for workers, so that they are equipped to work alongside machines, rather than being replaced by them. This approach would require a significant investment in education and training programs, but it would help to make sure that employees are not left behind in the age of automation [13].

Ultimately, any solution to the issue of distributing wealth created by machines will require a combination of these approaches and will need to be tailored to the specific circumstances of each country or region. It is important to start thinking about this issue now, as machines and AI are in the years to come, are likely to play an increasingly crucial role in the economy.

5.3. Can Machines Influence Human Behaviours and Interactions?

The influence of machines, including AI systems, on human behaviour and interactions is an important issue to consider. Machines can influence human behaviour and interactions in a variety of ways, including [5]:

- a) *Personalized recommendations*: Machines can use data and algorithms to personalize the content people see, including news, ads, and recommendations. This may limit the exposure to different viewpoints and result in "filter bubbles," where people only see information that supports their own views and attitudes [14].
- b) *Nudges*: Machines can use nudges, or small changes in the way information is presented or framed, to encourage individuals to make a particular decision. For example, a website may use a pop-up notification to remind to complete a task or provide social proof by showing how many other people have completed the task.
- c) *Social influence*: Machines can use social influence techniques to encourage users to conform to social norms or the behaviour of others. For example, social media platforms may highlight popular posts or show that the user's friends have engaged with a particular piece of content to encourage the user to do the same [15].
- d) *Automated decision-making*: Machines can make decisions on the behalf of humans, such as determining the candidates to be interviewed or ads to be shown. This can impact the opportunities and experiences of humans, and raise concerns about bias and discrimination.
- e) It is crucial to take into account the possible effects of these and other ways that robots may affect human interactions and behaviour. People must be conscious of the potential hazards and take action to reduce them, such as through openness, accountability, and ethical considerations, in order to ensure that the influence of machines is positive and useful.

5.4. How to Guard against Possible Unfavourable Mistakes?

Guarding against detrimental mistakes in AI is an important consideration as this technology continues to be developed and used. Here are some ways to guard against possible detrimental mistakes [5]:

- a) *Robust testing and evaluation:* AI systems should undergo rigorous testing and evaluation to identify and address potential errors and biases. This can include using large, diverse data sets and subjecting the system to various scenarios and edge cases to ensure that it functions as intended [17].
- b) *Human oversight and intervention:* It is important to have human oversight and intervention in AI decision-making to ensure that the system is making decisions that align with the ethical values and goals. This can include using human-in-the-loop systems, where humans review and approve AI decisions, or implementing "kill switches" to allow humans to override the system if necessary.
- c) *Ethical and transparent design:* Designing AI systems ethically should take into account eliminating bias and discrimination, among other things, protecting privacy, and prioritizing the public good. The design should be transparent and explainable, so that users can understand how the system works and how it makes decisions.
- d) *Continuous monitoring and updating:* AI systems should be regularly assessed and improved to guarantee their effectiveness and prevent them from straying from their original objectives. This can involve gathering user input through feedback mechanisms and using that information to inform subsequent system iterations.
- e) *Collaboration and knowledge sharing:* Collaboration and knowledge sharing among developers, researchers, and policymakers can help identify and address potential issues with AI before they become detrimental. It is important to have open and transparent communication channels to share best practices and lessons learnt.

Overall, guarding against detrimental mistakes in AI requires a multi-faceted approach that includes robust testing and evaluation, human oversight and intervention, ethical and transparent design, continuous monitoring and updating, and collaboration and knowledge sharing. By doing these things, it can be made sure that AI is created and utilised in ways that benefit society and are consistent with human moral principles and objectives.

5.5. Can AI Bias be Eradicated?

While it may be difficult to completely eliminate all bias from AI, there are steps to be taken to reduce and manage bias in AI systems. Here are some ways to address AI bias [5]:

- a) *Diverse and representative data*: The data that AI systems are taught on determines how objective they are. The chance of adding bias to the system can be lowered by employing a variety of representative datasets. Using data from underrepresented groups and taking precautions to prevent historical biases in the data are two examples of how to do this [18].
- b) *Regular audits and testing*: In order to find and correct potential bias, AI systems should undergo routine auditing and testing. This can involve developing testing scenarios that take into account potential biases and hiring outside auditors to evaluate the system for bias.
- c) *Explainable and transparent AI*: In order for people to comprehend how AI systems make decisions, they must be transparent and explicable in their design. This can improve trust and accountability by identifying and addressing any potential biases in the system.
- d) *Ethical guidelines and regulation*: Governments, industry groups, and other stakeholders can develop ethical guidelines and regulation to address AI bias. This can include promoting diversity and inclusion in the development and use of AI, setting standards for data collection and usage, and establishing oversight and accountability mechanisms.
- e) *Ongoing education and awareness*: Raising awareness and promoting education about AI bias can help prevent bias from being introduced into AI systems. This can include training data scientists and AI developers on the risks and challenges of bias and promoting a culture of inclusion and diversity in the development and use of AI.

It may not be possible to entirely remove all bias from AI systems, but by following these measures, the danger of bias may be lessened and may be managed in a morally and responsible manner. This can guarantee AI is created and utilised in ways that are consistent with human values and advanced justice and equity.

5.6. How to Protect AI from Opponents?

Protecting AI from adversaries is an important consideration, especially as AI becomes more prevalent and sophisticated. Here are some ways to protect AI from adversaries [5]:

- a) *Secure development*: Security should be considered while creating AI systems, and this includes adopting secure coding techniques, encrypting important information, and putting in place access controls to prevent unauthorised access.
- b) *Robust testing and evaluation*: AI systems should undergo rigorous testing and evaluation to identify and address potential vulnerabilities and attack vectors. This can include using penetration testing, stress testing, and other techniques to simulate real-world attack scenarios.
- c) *Monitoring and anomaly detection*: AI systems should be monitored for anomalies and unusual behaviour that may indicate a security breach or attack. This can include using intrusion detection systems, log analysis, and other techniques to detect and respond to security incidents.
- d) *Multi Factor Authentication*: Multi factor authentication can be used to help prevent unauthorized access to AI systems. This can include using biometric authentication, two-factor authentication, and other techniques to verify user identity.
- e) *Collaboration and knowledge sharing*: Collaboration and knowledge sharing among developers, researchers, and policymakers can help identify and address potential security risks with AI before they become detrimental. It is important to have open and transparent communication channels to share best practices and lessons learnt.
- f) *Ethical and transparent design*: AI systems should be created with ethical principles in mind, such as preventing malevolent use, safeguarding user privacy and security, and giving the public good priority. Users should be able to comprehend how the system functions and how it makes decisions thanks to the design's transparency and explicability.

Overall, protecting AI from adversaries requires a multi-faceted approach that includes secure development, robust testing and evaluation, monitoring and anomaly detection, multi-factor authentication, collaboration and knowledge sharing, and ethical and transparent design.

5.7. How Can Unintended Consequences Be Avoided?

Avoiding unintended consequences in AI is an important consideration, as AI systems can have far-reaching impacts on society. Here are some ways to minimize unintended consequences [5]:

- a) *Consider the context:* It's important to consider the context in which an AI system will be used. This includes factors such as the user group, the social and cultural norms, and the regulatory environment. Understanding the context can help identify potential unintended consequences and avoid them.
- b) *Robust testing and evaluation:* AI systems should undergo rigorous testing and evaluation to identify and address potential unintended consequences. This can include using simulations and scenario analysis to identify potential outcomes and identify any issues before the system is deployed.
- c) *Transparency and explainability:* AI systems must be transparent and understandable so that people may comprehend how they operate and how they arrive at judgements. This can improve trust and responsibility while identifying and addressing unintended outcomes.
- d) *User feedback and monitoring:* Users should be encouraged to provide feedback on the performance of AI systems, and systems should be monitored for unintended consequences. This can include using performance metrics, monitoring user behaviour, and soliciting feedback from users.
- e) *Ongoing education and awareness:* Raising awareness and promoting education about potential unintended consequences of AI can help prevent these consequences from occurring. This can include training data scientists and AI developers on the risks and challenges of unintended consequences, and promoting a culture of transparency and responsibility in the development and use of AI.

Overall, avoiding unintended consequences in AI requires a multi-faceted approach that includes considering the context, robust testing and evaluation, transparency and explainability, user feedback and monitoring, and ongoing education and awareness. These actions can help to ensure that AI is created and applied in ways that minimise unforeseen consequences.

5.8. Is There A Way to Remain in Total Control of AI?

As AI becomes more advanced and sophisticated, it can become increasingly challenging to remain in total control of AI systems. However, there are some ways to maintain a degree of control over AI [5]:

- a) *Designing AI systems with control in mind:* Kill switches, auditing tools, and explainable algorithms should be built into AI systems when designing them. This can make it possible for humans to monitor and intervene in AI systems while yet maintaining control over them.
- b) *Establishing ethical norms and guidelines for AI deployment:* It can help guarantee that the technology is created and deployed in ways that are consistent with human values and objectives. To support ethical AI development and use, this can involve creating codes of behaviour, ethical frameworks, and regulatory frameworks.
- c) *Ensuring accountability and transparency:* Ensuring accountability and transparency in AI development and deployment can help prevent AI systems from operating outside human control. This can include implementing robust data governance and security measures, as well as transparent reporting and disclosure of AI system performance and outcomes.
- d) *Investing in education and awareness:* Investing in education and awareness about AI and its potential impacts can help ensure that individuals and organizations understand the risks and challenges associated with AI. This can include training programs, awareness campaigns, and educational materials that promote responsible AI development and use.

Overall, while it may be challenging to remain in total control of AI, steps can be taken to maintain a degree of control over AI systems. By designing AI systems with control in mind,

establishing ethical guidelines and standards, ensuring accountability and transparency, and investing in education and awareness, responsible AI development and use can be promoted.

5.9. Should Humane Behaviours of AI Be Considered?

The question of whether humane treatment of AI should be considered an important ethical consideration. While AI systems are not sentient beings and do not have feelings or emotions in the way that humans and animals do, there are several reasons why humane treatment of AI should be considered [5]:

- a) *Ethical considerations*: Treating AI systems with respect and consideration can be seen as an ethical imperative. AI systems should be treated with respect and concern, just as other people and animals are treated. This is in line with human ethical principles.
- b) *Long-term impact*: How AI systems are treated now could have long-term impacts on our future relationship with them. If they are treated with disregard or disrespect, it could impact the human ability to work with them and achieve the goals set for them.
- c) *Impact on humans*: How AI systems are treated could also impact how humans treat each other. Treating AI systems with respect and consideration could foster a culture of empathy and kindness, which could extend to how other humans are treated.
- d) *Perception of AI*: Treating AI systems with respect and consideration could also impact how society perceives AI. If AI systems are treated as if they are disposables or unimportant, it could reinforce negative perceptions of AI that could hinder its development and acceptance.

The concept of humane treatment of AI is still in its infancy, and there is much debate and discussion surrounding what this might involve, despite the fact that there are justifications for why it should be taken into consideration. However, human approach to AI to reshape in a way that is consistent with human beliefs and objectives by taking into account the ethical implications of how AI systems are treated, may be started.

6. Conclusion

As Artificial Intelligence (AI) advances, its ethics become a more pressing problem that needs to be solved. Human rights, responsibility, information, and other ethical issues could be profoundly impacted by AI, hence it is crucial to constantly improve human ethical reasoning to prepare for these developments. The direction of science and technology ethics is intimately related to the healthy and benign development of AI. AI development must be governed by moral norms that put the common good first and defend human rights and dignity. This includes things to think about including accountability, openness, and avoiding bias and discrimination. Marxism can help to establish ethical frameworks for AI development and regulation as well as a deeper knowledge of science and technology ethics. This can assist in ensuring that AI is created and applied in ways that are consistent with human moral principles and objectives. It is crucial to remember that while precautions must be taken to reduce risks and handle ethical issues, nonrestrictive measures or risk-reduction strategies that could stunt the growth of the company must also be avoided. Promoting AI's advantages while also addressing its possible hazards and ethical ramifications requires striking a balance.

References

- [1] Bostrom, Nick, and Eliezer Yudkowsky. Forthcoming. "The Ethics of Artificial Intelligence." In *Cambridge Handbook of Artificial Intelligence*, edited by Keith Frankish and William Ramsey. New York: Cambridge University Press, 2011.
- [2] Asimov, Isaac., "Runaround." *Astounding Science-Fiction*, 94–103, March, 1942.
- [3] Bostrom, Nick, "Existential Risks: Analyzing Human Extinction Scenarios and Related Hazards." *Journal of Evolution and Technology*, 2002, <http://www.jetpress.org/volume9/risks.html>.
- [4] Nithesh Naik, B. M. Zeeshan Hameed, Dasharathraj K. Shetty, Dishant Swain, Milap Shah, Rahul Paul, Kaivalya Aggarwa, Sufyan Ibrahim, Vathsala Patil, Komal Smriti, Suyog Shetty, Bhavan Prasad Rai, Piotr Chlosta and Bhaskar K. Somani, "Legal and Ethical Consideration in Artificial Intelligence in Healthcare: Who Takes Responsibility?", *Sec. Genitourinary Surgery Front. Surg.*, *Sec. Genitourinary Surgery*, Volume 9, 14 March 2022, <https://doi.org/10.3389/fsurg.2022.862322>

- [5] Margarita Simonova, "Top Nine Ethical Issues In Artificial Intelligence", Forbes Technology Council, <https://www.forbes.com/sites/forbestechcouncil/2022/10/11/top-nine-ethical-issues-in-artificial-intelligence/?sh=2b10a32e5bc8>, 11th October 2022.
- [6] Alice PavaloIU, Utku Kose, "Ethical Artificial Intelligence - An Open Question", Journal of Multidisciplinary Developments, e-ISSN: 2564-6095, 2(2), 15-27, 2017.
- [7] H Y Li , J T An and Y Zhang, "Ethical Problems and Countermeasures of Artificial Intelligence Technology", E3S Web of Conferences 251, 01063 (2021), <https://doi.org/10.1051/e3sconf/202125101063>, TEES 2021.
- [8] Deng R 2020 Reflection on the science and technology ethics of artificial intelligence development J. Guangxi Social Sciences 10 93-7
- [9] Yu X and Duan W 2019 Ethical construction of artificial intelligence J. Theoretical Exploration 06 43-9
- [10] He J and Meng P 2020 Xi Jinping's "five theories" on science and technology ethics J. Truth from Facts 3 11-6
- [11] Zheng Z 2021 Ethical crisis and legal regulation of artificial intelligence algorithms J. Law Science. Journal of Northwest University of Political Science and Law 1 1-13
- [12] Turner J M 2019 Robot Rules: Regulating Artificial Intelligence (Switzerland: Palgrave Macmillan imprint)
- [13] Liu X 2020 Legal philosophical thoughts on the nature of tools of intelligent robots J. Criminal Science 5 20-34
- [14] Tang L 2020 Public nuisance to algorithm application and its governance paths J. Northern Legal Science 3 51-60.
- [15] Chen S 2020 Algorithmic governance: risks and countermeasures of technological alienation in intelligent society J. Journal of Hubei University. Philosophy and Social Sciences Edition 1 158-65
- [16] Luciano F M 2016 The 4th Revolution: How the infosphere is reshaping (Zhejiang: Zhejiang People's Publishing House)
- [17] Bostrom, N. (2014). Superintelligence: Paths, Dangers, Strategies. OUP Oxford.

- [18] Bostrom, N. (2015). What happens when our computers get smarter than we are?
Retrieved 3 April 2017, from
https://www.ted.com/talks/nick_bostrom_what_happens_when_our_computers_get_smarter_than_we_are
- [19] Brent Mittelstadt, "The impact of Artificial Intelligence on the Doctor - Patient Relationship", © Council of Europe, December 2021.