

The Ethical Impacts of AI on Industries and Governance Frameworks for Overcoming Them

Jung-im Kim¹

ABSTRACT

This paper As the AI industry continues to flourish today, the ethical implications of AI on various industries are expected to manifest in diverse ways. It is anticipated that AI will impact not only specific sectors but also entire industries and government systems, significantly influencing many aspects of human life.

In particular, as the AI industry advances, issues such as excessive data collection, data bias, and misuse of information are becoming increasingly prevalent. Additionally, the ambiguity of accountability is emerging as a social issue due to unclear responsibilities. Furthermore, with the digital divide and the acceleration of job displacement, the threat posed by AI to human employment is becoming a growing concern.

While the adoption of AI can bring highly positive impacts by enhancing efficiency and productivity, it is a crucial time to recognize potential ethical challenges in advance and address them through appropriate regulatory governance.

To achieve this, industries, governments, local authorities, and academia must collaborate to establish sustainable AI governance and ethical guidelines. This collective effort will help minimize the problems that arise from the development and use of AI.

☞ keyword : AI Ethics, Data Collection, Digital Divide, Job Displacement, Ethical Issues, AI Governance

1. Introduction

1.1 Research Background

As the era of the Fourth Industrial Revolution unfolds, artificial intelligence (AI) has become a core technology across industries. In manufacturing, AI-powered smart factories are enhancing productivity, while in the healthcare sector, AI supports precision diagnostics and personalized treatments. In finance, AI enables automated risk analysis and strategic financial planning, driving innovation across various industries.

However, the rapid advancement and widespread adoption of AI technology have led to numerous ethical challenges. Since AI operates based on data, concerns about privacy infringement have been raised. Additionally, inaccurate data or biased algorithms can result in unfair outcomes for specific groups.

Furthermore, in critical decision-making areas such as autonomous vehicles and medical AI, errors can arise, and it is often difficult to clearly identify accountability, which has become a societal issue.

As automation technologies evolve, some jobs are disappearing, and labor market structures are changing, exacerbating economic inequality. These ethical issues undermine trust

in AI technology and pose significant risks for businesses that rely on AI-driven operations. If left unaddressed, the rapid proliferation of individuals or groups harmed by AI technology is inevitable.

In response, governments and corporations worldwide are establishing ethical governance frameworks to ensure that AI technology guarantees transparency, fairness, and accountability. For instance, the European Union (EU) has developed legal guidelines for AI use and regulation, while companies like Google and Samsung SDS have implemented ethical AI principles and applied them to their operations and technology development processes.

These efforts are considered effective approaches to addressing ethical challenges. This study aims to analyze the ethical impacts of AI technology across industries and explore solutions to overcome these issues by examining cases of ethical governance frameworks.

Through this analysis, the study seeks to propose practical AI ethical guidelines and governance directions for fostering responsible AI development and building societal trust.

1.2 Research Objectives

The primary objective of this study is to conduct an in-depth analysis of the ethical impacts of AI technology across industries and propose ethical governance frameworks to address these challenges.

While AI has driven innovation in industries such as manufacturing, healthcare, finance, and education, it has also introduced ethical issues, including data bias, privacy infringement, ambiguous accountability, and changes in the labor market.

These challenges are not merely technical problems but multifaceted issues with social and economic implications, requiring comprehensive research and solutions.

As the speed and scope of AI technology adoption continue to expand rapidly, failing to proactively diagnose and address these ethical challenges could exacerbate societal mistrust in the technology. This would not only diminish the social acceptability of AI but also hinder sustainable growth and innovation across industries.

Therefore, this business project emphasizes the necessity of establishing governance frameworks for the ethical use of AI technology and aims to propose concrete and practical solutions for achieving this.

The study systematically analyzes international and domestic cases addressing AI ethical challenges, including policy initiatives like the EU's AI Act, the adoption of ethical AI principles by companies like Google and IBM, and the ethical guidelines of various institutions worldwide.

By extracting common success factors from these cases, the study intends to design tailored governance frameworks that can be applied across different industries. Additionally, the study emphasizes that solving AI ethical issues is not solely the responsibility of technology developers or regulatory bodies but a task requiring the collective participation of industries and society as a whole.

To this end, the research seeks to propose strategies for collaboration among diverse stakeholders. Corporations can internalize ethical standards during technology development, governments can support ethical technology use through regulations and policies, and academia can contribute by advancing ethical standards through research.

Ultimately, this study aims to provide practical guidelines for resolving ethical issues, enabling AI technology to develop in a responsible and trustworthy manner across industries and society. This approach aspires to ensure that AI not only delivers technological innovation but also positively impacts industries and society based on human-centered values and societal trust.

1.3 Research Methods

This business project utilizes various research methods to analyze the ethical impacts of AI technology across industries and derive governance frameworks to address these issues. First, the study conducts a literature review to investigate key issues related to AI ethics, drawing from academic papers, reports, and policy documents.

Through this, the study systematically categorizes ethical challenges such as data bias, privacy infringement, and ambiguous accountability, while identifying their root causes and potential solutions. The literature review focuses on establishing a theoretical foundation for building AI ethical governance frameworks.

Next, the study conducts case studies to analyze major domestic and international examples of ethical AI utilization. Specifically, it examines policy approaches like the EU's AI Act, ethical AI principles adopted by companies such as Google and IBM, and internal ethical guidelines implemented by corporations like Samsung SDS.

These cases provide insights into the challenges encountered during the design and operation of governance frameworks and the methods used to overcome them.

Qualitative analysis also plays a significant role in the study. Based on data collected through case studies, the study identifies common success factors and limitations to design customized governance frameworks applicable to various industries. In particular, the study analyzes how approaches to solving ethical challenges may differ depending on the characteristics of each industry.

If necessary, additional opinions are gathered through surveys or interviews with professionals in the field. Surveys and

interviews target corporate employees developing or utilizing AI technology, policymakers, and researchers to explore their perceptions, current practices, and requirements related to ethical issues.

By integrating these diverse research methods, the study aims to propose systematic and practical solutions for addressing AI ethical challenges. The ultimate goal is to establish a foundation for overcoming ethical issues and enabling AI technology to positively impact industries and society.

2. Previous Research and Environmental Analysis

2.1 Research on Ethical Issues in AI

The rapid advancement of AI technology has raised numerous ethical issues globally, prompting extensive research on the topic. Key findings from previous studies are as follows: Representative research has explored whether algorithmic bias and data imbalance can lead to unfair outcomes for specific groups. Cathy O'Neil's "Weapons of Math Destruction" provides an in-depth analysis of such issues within a societal context.

Studies have also addressed the ambiguity of accountability in critical decision-making AI systems, such as autonomous vehicles and medical AI. For instance, discussions are ongoing regarding whether responsibility for accidents lies with the developer, user, or manufacturer.

In addition, prior research has highlighted how bias and imbalance in training data can cause unfair outcomes for certain groups. Buolamwini and Gebru's "Gender Shades" analyzed gender and racial disparities in facial recognition algorithms, revealing lower accuracy for Black women. Barocas et al. examined how algorithmic bias impacts AI in recruitment, finance, and legal systems.

The issue of accountability has also been identified as a significant societal problem, particularly in ethical dilemmas and responsibility disputes related to autonomous vehicles and medical AI. Lin et al. discussed the legal and ethical accountability of self-driving car accidents, emphasizing the importance of moral decision-making

algorithms in AI design.

Floridi and Cowls proposed the "Five Ethical Principles" (transparency, fairness, accountability, beneficence, and reliability) for AI and suggested linking them to policy-making efforts.

2.2 Research on Ethical AI Guidelines and Governance

Research on developing guidelines and governance frameworks to address AI's ethical issues has gained significant attention. The European Union (EU) has proposed the world's first comprehensive AI legislation to regulate AI use and enhance transparency, focusing on ethical design and application.

Analysis has also been conducted on ethical principles adopted by companies such as Google, IBM, and Microsoft. These corporations have introduced guidelines based on transparency, fairness, and accountability, applying them to their operations and technology development. Similar movements to adopt such guidelines are emerging in the IT industry in South Korea.

Studies on ethical AI design and governance frameworks have increased, with three key areas of focus. The EU's AI Regulatory Framework: The Artificial Intelligence Act (2021) includes regulations on high-risk AI systems, transparency requirements, and data quality management, aiming to establish global standards for AI ethics. However, researchers have critiqued its potential to stifle innovation among SMEs and startups.

Corporate AI Ethical Guidelines: Companies such as Google emphasize "AI's social responsibility" and have implemented systems like Ethical Review Boards. IBM has developed open-source tools (e.g., AI Fairness 360) to enhance AI trustworthiness and shared them with the industry.

National Governance Approaches: Major countries like the United States, China, and Japan have announced policies on AI ethics, seeking to balance ethical design with technological competitiveness. Research has compared China's state-driven approach to AI with the US's private-sector-led innovation, analyzing the effectiveness and ethical implications of these policies.

2.3 Research on AI's Social Impact

Studies have explored the ethical, economic, and cultural impacts of AI on society. Research indicates that automation could displace traditional jobs and transform labor market structures, exacerbating economic inequality.

Frey and Osborne analyzed the impact of AI on various occupations in "The Future of Employment", highlighting that mid-level skilled jobs are most susceptible to automation.

Follow-up studies have also examined new job opportunities created by AI, such as data labeling and AI training roles, emphasizing the importance of reskilling and upskilling the workforce.

The issue of social inequality has also been raised. Eubanks' "Automating Inequality" demonstrated how AI systems could exacerbate inequality in public service distribution, calling for social consensus and policy intervention to address technology-driven disparities.

2.4 Research on AI Ethics Education and Policy

Efforts to promote the ethical use of AI technology have included studies on education and policy. Academic approaches and educational methodologies for AI ethics are being explored to foster ethical sensitivity among the next generation of technologists.

Comparative studies on national AI policies analyze the differences in AI ethics guidelines and aim to derive common lessons. AI ethics education is increasingly recognized as essential for enhancing the ethical awareness of both developers and the general public.

Leading universities in the United States, such as MIT, Stanford, and Carnegie Mellon, have incorporated AI ethics courses into their undergraduate curricula, offering students opportunities to learn about the importance of ethical design. In China, AI ethics education is being strengthened under state initiatives, including its integration into primary and secondary school curricula.

2.5 Research on Transparency and

Trustworthiness in AI

To address the "black box problem" in AI systems, researchers have focused on transparency and explainability. Ribeiro et al.'s "Why Should I Trust You? Explaining the Predictions of Any Classifier" proposed methods for providing explanations of AI predictions, contributing to greater trust in AI systems.

Transparency research has gained attention in industries like healthcare and finance, where ethical accountability in decision-making is crucial.

Doshi-Velez and Kim's study on interpretable machine learning emphasized the connection between AI transparency, fairness, and trustworthiness. In healthcare, explainable models are being developed to ensure that AI diagnostic results are understandable to medical professionals (e.g., IBM Watson Health).

2.6 Research on AI and Privacy Issues

Studies have examined the ethical concerns surrounding AI's use of personal data. Solove's "A Taxonomy of Privacy" categorized various types of privacy infringement and discussed how AI could exacerbate these issues.

Research on privacy-preserving technologies, such as Federated Learning and Differential Privacy, has been particularly active.

These technologies aim to protect personal data while maintaining AI model performance. McMahan et al. demonstrated how Federated Learning enables decentralized data training to safeguard privacy.

The EU's General Data Protection Regulation (GDPR) has set global benchmarks for data protection, defining standards for AI systems handling personal information.

2.7 Research on AI Automation and Labor Market Changes

Studies have examined how AI automation affects the labor market. Frey and Osborne's "The Future of Employment" analyzed the susceptibility of various occupations to computerization, emphasizing that mid-level skilled jobs are most vulnerable to automation.

Subsequent research has explored how AI

creates new high-value job opportunities, such as in data annotation and AI development. These studies underscore the importance of reskilling programs to help workers adapt to technological changes.

Brynjolfsson et al.'s "Skills of the Future" emphasized the need for collaborative efforts between governments and businesses to equip workers with new skills. Most researchers agree that fostering workforce adaptability is essential for minimizing the negative impacts of AI automation.

This business project builds on previous research by comprehensively analyzing AI ethical issues within the broader industrial context. Unlike earlier studies that focused on specific topics, this research integrates various ethical challenges and proposes governance frameworks through case studies and qualitative analyses.

By leveraging insights from previous research, this study aims to provide actionable solutions for overcoming ethical challenges and fostering the responsible development of AI technology.

3. Case Studies

3.1 Amazon's AI Recruitment System

In 2014, Amazon introduced an AI-based automated recruitment system to enhance the efficiency of its hiring process. However, this system evaluated applicants based on data from the past decade, which was predominantly male-centric. As a result, the system exhibited bias against female candidates.

For instance, resumes containing words associated with "women" (e.g., "women's club") were scored lower. This issue highlighted how AI systems can reflect the biases inherent in the data they are trained on. Recognizing this problem, Amazon discontinued the project.

The company has since introduced bias detection technologies to enhance the fairness of its training data and adopted more diverse datasets, including gender, race, and background variations.

Additionally, Amazon established an internal process to evaluate the fairness of AI systems to prevent similar issues from recurring.

3.2 Uber's Self-Driving Car Accident

In 2018, a fatal accident occurred in Arizona, USA, involving Uber's self-driving car, which struck and killed a pedestrian. Investigations revealed that while the vehicle's sensors detected the pedestrian, the software had been configured to ignore such inputs.

Furthermore, the emergency braking system was deactivated, and the safety operator in the driver's seat was distracted at the time of the incident. This tragedy raised serious concerns about the safety of autonomous driving technologies and accountability in such incidents.

Following the accident, Uber overhauled its self-driving car software, improving pedestrian detection and emergency braking systems. The company also clarified the role of safety operators and implemented stricter testing protocols.

Uber now collaborates with the National Highway Traffic Safety Administration (NHTSA) to enhance safety standards and improve ethical design and testing environments.

3.3 European Union (EU) AI Act

The EU proposed the AI Act to ensure that AI does not infringe on fundamental human rights or create ethical issues.

This regulation adopts a risk-based approach, applying strict criteria to high-risk AI systems in sectors such as healthcare, transportation, and law enforcement.

For example, AI systems using biometric data, such as facial recognition, are subject to stringent regulations. The EU mandates pre-assessments and certifications for high-risk AI systems and requires transparent data processing.

The AI Act promotes collaboration among businesses, academia, and regulatory authorities to balance technological innovation with ethical standards, thereby strengthening public trust in AI technologies.

3.4 Google's AI Ethics Board

In 2018, Google announced its "Responsible

AI" principles and established an AI ethics board. However, internal conflicts arose when the board was involved in decisions related to military AI projects (e.g., Project Maven). Additionally, a lack of communication with external stakeholders was criticized, and disagreements among board members hindered its effectiveness.

Google responded by reorganizing the board, integrating independent decision-making mechanisms, and including external experts. The company established procedures to evaluate the ethical impacts of AI projects in advance and transparently disclose the results.

Furthermore, Google strengthened internal communication and introduced employee training programs and guidelines to integrate ethical principles into AI development processes.

3.5 Foxconn's Automation and Job Displacement

Since 2017, Foxconn has replaced approximately 60,000 jobs at its Chinese factories with AI-powered robots. While this automation improved productivity, it also displaced workers, exacerbating economic inequality. This case exemplifies how AI can replace human labor in the workforce.

In response, Foxconn implemented retraining programs to help displaced workers adapt to the automated environment. The company also encouraged workers to transition to higher-skilled jobs created by automation and collaborated with governments to develop policies addressing labor market changes.

3.6 IBM Watson's Medical AI

IBM Watson was developed to assist with cancer diagnosis and treatment. Initially, the system faced criticism for providing recommendations that lacked transparency and for being trained on data of inconsistent quality, leading to mistrust among medical professionals.

IBM addressed these issues by incorporating explainable AI (XAI) capabilities into Watson, enabling it to clearly explain the rationale behind its treatment recommendations.

The company also improved the quality of

training data and provided training programs to help healthcare professionals use Watson as a supplementary tool. IBM continues to collaborate with academic institutions to standardize data and enhance transparency, building trust in medical AI.

3.7 Clearview AI's Facial Recognition Technology

Clearview AI developed facial recognition technology using images scraped from public websites without user consent, raising concerns about privacy violations. These concerns intensified when the technology was adopted by law enforcement agencies, sparking ethical and privacy debates.

In response, Clearview AI limited its technology to contracts with public institutions and restricted services for private individuals.

The company also adopted stricter data collection practices, requiring explicit user consent. In some countries, Clearview AI's use has been banned, prompting the company to adhere to stronger ethical guidelines for data collection.

4. Conclusion and Implications

4.1 Summary and Conclusion

AI technology has driven innovation across industries such as manufacturing, healthcare, finance, and transportation. However, it has also introduced ethical challenges, including data bias, privacy violations, and unclear accountability. If these issues are not addressed, they could erode public trust in AI, exacerbate economic inequality from job displacement, and deepen societal conflicts.

This report analyzed the ethical impacts of AI technology across industries and proposed governance frameworks to address these challenges.

By examining literature and case studies, the report identified key ethical issues and explored strategies to overcome them.

Real-world examples, such as Amazon, Uber, and IBM Watson, provided practical insights into the ethical challenges posed by AI. Regulatory frameworks like the EU's AI Act offered a roadmap for implementing

ethical AI.

This business project proposes seven strategic recommendations for sustainable AI development: Establish data management frameworks to ensure fairness and transparency. Integrate ethical reviews throughout the AI development lifecycle.

Strengthen transparency and explainability to build public trust. Develop governance frameworks for oversight and regulation. Enhance education and public communication to foster societal trust.

Address labor market changes with reskilling programs and job creation. Promote sustainable AI development while considering environmental impacts.

By implementing these strategies, AI technology can maintain fairness and accountability, gain public trust, and overcome ethical challenges. While AI is a key driver of the Fourth Industrial Revolution, unresolved ethical issues could undermine its positive potential.

In conclusion, addressing AI's ethical challenges is a collective responsibility that extends beyond developers and policymakers. Governments, businesses, academia, and civil society must collaborate to establish ethical guidelines and governance frameworks.

Transparent communication among these stakeholders will be critical to building trust. The strategic recommendations in this report aim to support the sustainable development of AI technology, ensuring it transcends being a mere technological tool to become a driver of human-centric values and positive societal change.

4.2 Implications

Given the profound potential impact of AI on human society, its ethical use has become as important as technological advancement. This business project highlights the importance of addressing data bias, enhancing transparency, establishing governance frameworks, adapting to labor market changes, and ensuring environmental sustainability.

These implications serve as a reminder to developers, policymakers, and business leaders of the importance of ethical AI and encourage practical changes that contribute to responsible

AI development.

References

- [1] O'Neil, C. (2016). *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. Crown Publishing Group
- [2] Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). "Why Should I Trust You?": Explaining the Predictions of Any Classifier. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*.
- [3] European Commission. (2021). *Proposal for a Regulation Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act)*.
- [4] Solove, D. J. (2006). A Taxonomy of Privacy. *University of Pennsylvania Law Review*, 154(3), 477 - 564.
- [5] Frey, C. B., & Osborne, M. A. (2017). The Future of Employment: How Susceptible Are Jobs to Computerisation? *Technological Forecasting and Social Change*, 114, 254 - 280.
- [6] IBM (2020). *AI Ethics: Principles for Trust and Transparency*. IBM Research.
- [7] Suresh, H., & Guttag, J. V. (2021). A Framework for Understanding Unintended Consequences of Machine Learning. *Communications of the ACM*, 64(1), 62 - 71.
- [8] Binns, R. (2018). Fairness in Machine Learning: Lessons from Political Philosophy. *Proceedings of the 2018 Conference on Fairness, Accountability, and Transparency*
- [9] Google AI. (2018). *AI at Google: Our Principles*.
- [10] Hagendorff, T. (2020). The Ethics of AI Ethics: An Evaluation of Guidelines. *Minds and Machines*, 30(1), 99 - 120.
- [11] Microsoft. (2019). *Responsible AI Principles*.
- [12] European Parliament. (2020). *Artificial Intelligence: Tackling Ethical and Social Challenges*.
- [13] IEEE Global Initiative on Ethics of Autonomous and Intelligent Systems. (2019). *Ethically Aligned Design: A Vision for Prioritizing Human*

Well-being with Autonomous and Intelligent Systems.

- [14] Zuboff, S. (2019). The Age of Surveillance Capitalism: The Fight for a Human Future at the New Frontier of Power. PublicAffairs.
- [15] Clearview AI. (2021). Transparency Report.

● Author Biography ●



Kim Jungim

Ewha Womans University Graduate School of Business, MBA in Business Administration (Master of Business Administration), Class of 2011

Yonsei University Graduate School of Public Policy, Department of Entrepreneurship (Master of Entrepreneurship), Class of 2020

Seoul School of Integrated Sciences & Technologies (aSSIST), AI Advanced Graduate School, Department of AI and Big Data (Master of Science in AI and Big Data), Class of 2025

Research Interests : AI Ethics, AI Safety, AI Governance, AI Policy etc.

E-mail : kjung225@hanmail.net