



A Study of Ethical Issues and Challenges of AI in Management

Manali Prabhunath Malankar

SIES College of Arts, Science and Commerce, Sion (West)

Abstract

Artificial Intelligence (AI) has emerged as a transformative technology that is reshaping management practices across industries. Organizations increasingly rely on AI for decision-making, recruitment, customer service, performance evaluation, financial forecasting, and strategic planning. While AI enhances efficiency, productivity, and data-driven decision-making, it also raises significant ethical concerns that require careful consideration. This study examines the major ethical issues and challenges associated with the adoption of AI in management. Key concerns include algorithmic bias and discrimination, lack of transparency and explainability, data privacy and security, accountability for AI-driven decisions, workforce displacement, and the responsible use of employee and customer data. The paper further explores the importance of ethical governance, regulatory compliance, human oversight, and the development of fair and transparent AI systems. By analyzing existing literature and contemporary management practices, the study highlights the need for organizations to establish ethical AI frameworks that balance technological innovation with social responsibility. The findings emphasize that sustainable AI adoption depends on integrating ethical principles, legal standards, and human values into organizational decision-making processes. The study concludes that responsible AI governance is essential for building trust, ensuring fairness, and achieving long-term organizational success in the digital age.

Keywords: *Artificial Intelligence (AI), Management, AI Ethics, Ethical Challenges, Algorithmic Bias, Data Privacy, Transparency, Accountability, Human Oversight*

Introduction:

In today's world, the term Artificial Intelligence (AI) is not new to be heard. It is used by majority of the population around the world in some or the other ways. It is widely required and used as it helps in speeding up the work with efficiency.

Artificial Intelligence (AI) is also used in management for decision-making, recruitment, marketing, and operations, etc. With speed, accuracy, efficiency it also sometimes gives rise to issues relating to bias, data privacy, and lack of accountability.

This study focuses on understanding the ethical issues and challenges faced by organizations while using AI in management. This study adopts exploratory research methodology. Cases of three popular organisations like Amazon, Facebook and Uber are considered for the study.

Objectives:

1. To identify major ethical issues in the use of AI in management through some case studies.
2. To analyze the challenges faced by organizations in implementing ethical AI.
3. To study the impact of AI on decision-making and fairness.
4. To suggest measures for responsible and ethical use of AI.

Research Gap:

Most researches show how the technologies/Artificial Intelligence (AI) failed and the resulting outcomes due to technological/ Artificial Intelligence (AI) failure. There has not been much research directly focusing on the importance of "Human Heed" to avoid such circumstances. There is somewhere a gap of lack of management control. This study uses three different industries to prove that we need "Human Heed" that tells managers exactly when they must step in and make sure that the AI stays honest and safe. This study identifies a critical management gap by showing how these errors are actually failures of leadership oversight.

Hypothesis :

1. Ho(Null Hypothesis):

AI implementation in management does not create significant ethical issues.

2. H1 (Alternative Hypothesis):

AI implementation in management creates significant ethical issues such as bias, privacy concerns, and lack of transparency.

Review of Literature:

The scholarly discourse surrounding Amazon's failed AI recruitment tool (2014–2017) focuses primarily on the intersection of **Machine Learning (ML) ethics** and **Human Resource Management (HRM)**. According to **Dastin (2018)**, the foundational problem was "algorithmic path dependency," where the system was trained on a decade of historical resumes. Because these resumes predominantly represented a male-centric workforce, the algorithm established a "success profile" that was inherently gendered. This aligns with the "Garbage In, Garbage Out" (GIGO) doctrine discussed by **O'Neil (2016)**, which posits that algorithms do not eliminate human bias but rather automate and scale it under a veneer of mathematical objectivity.

Current scholars in the field of Commerce and Management now advocate for a "Human-in-the-Loop" (HITL) framework. This research suggests that while AI can provide analytical speed, the final evaluative judgment must remain human to ensure intersectional equity and to prevent the systemic exclusion of minority groups who are underrepresented in historical data sets.

The Cambridge Analytica scandal serves as a seminal case in academic literature regarding the erosion of **informed consent** and the rise of **surveillance capitalism**. According to **Cadwalladr and Graham-Harrison (2018)**, the incident involved the unauthorized harvesting of personal data from approximately 87 million Facebook users via a third-party personality quiz app, "This Is Your Digital Life." Scholars like **Zuboff (2019)** argue that this case illustrates a fundamental shift in the digital economy, where personal behavioral data is treated as "raw material" for hidden commercial and political fabrication. The literature emphasizes that the primary ethical failure was not just a technical data breach, but a systemic

exploitation of platform APIs that allowed the "scraping" of data from users' friends without their knowledge or explicit agreement (**Bambauer, 2019**).

Furthermore, the literature evaluates the **regulatory and corporate governance** aftermath, viewing it as a catalyst for global privacy reform. The scandal is widely credited with accelerating the enforcement of the **General Data Protection Regulation (GDPR)** in Europe and the introduction of the **California Consumer Privacy Act (CCPA)** in the United States. Researchers (e.g., **Isaak & Hanna, 2018**) highlight the "Accountability Gap" in Big Tech, noting that Facebook's delayed response and defensive communication strategy led to a significant loss of stakeholder trust. Modern management studies now use this case to argue for "Privacy by Design," suggesting that ethical data stewardship must be an core organizational value rather than a reactive compliance measure.

A significant portion of the academic review focuses on the **Legal Classification and Duty of Care**. The landmark Supreme Court case, *Uber BV v Aslam [2021]*, serves as the primary reference point for understanding how employment status directly impacts safety governance. The literature (e.g., **Adams-Prassl et al., 2021**) suggests that by categorizing drivers as "self-employed" rather than "workers," Uber initially bypassed statutory requirements for rest breaks and maximum working hours under the UK Working Time Regulations. Researchers argue that this "contractual fiction" allowed the platform to optimize for vehicle density at the expense of driver wellness. The court's decision to reclassify drivers as "workers" is viewed in the literature as a critical intervention for public health, linking worker rights directly to the safety of the UK's road infrastructure.

Scholars suggest that the "Black Box" nature of Uber's routing and dispatch algorithms prevents traditional police investigations from identifying the systemic "contributory factors" behind accidents. This body of research concludes that for the UK to maintain its road safety standards in the age of the gig economy, regulators must transition from a "driver-blame" model to a "system-oversight" model that holds digital platforms accountable for the behavioral impacts of their algorithms

Findings:

	Objectives	Findings
1	To identify major ethical issues in the use of AI in management through some case studies	The study found that AI systems can lead to bias, privacy violations, and lack of transparency, as seen in cases like Amazon, Facebook and Uber.
2	To analyse the challenges faced by organizations in implementing ethical AI.	Organizations face challenges such as lack of proper regulation, guidelines, instructions in implementing Ai. There is difficulty in managing /monitoring AI systems for e.g., in case of Facebook. Over dependency on automation, as in case of Uber.
3	To study the impact of AI on decision-making and fairness.	AI can affect decision-making by creating unfair or biased outcomes, especially when trained on historical data, reducing fairness and trust. (Amazon)

4	To suggest measures for responsible and ethical use of AI.	The study highlights the need for human oversight, transparency, and strict data protection policies to ensure responsible AI usage. Human Heed is must for avoiding complexities.
---	--	--

Amazon AI Recruitment Bias:

Amazon developed AI based tool in the year 2014 for recruitment of candidates. During the training period it fed the toll with past 10 years historical data. Then it was put to operation. During the year 2017 the issue was raised stating that though both mala and females gave applications . It accepted resumes of male candidates. This was so because the past data fed had majority of male candidates resumes. Therefore, the system believed male is better than female. The system lacked to understand it was unfair. It operated based on past makle dominated data andf lacked fairness. After realising this matter the system was discontinued in the year 2018.

This shows that AI is trained based on what data you give. Therefore, here comes the role of the manager to correctly train AI so that it does not lack fairness in decision making.

Facebook – Cambridge Analytica

Cambridge Analytica is a data and Ai based company from UK. It was used to influence peoples opinions in politics.

The Cambridge Analytica scandal involved the unauthorized harvesting of data from approximately 87 million Facebook users via a personality quiz app named "**thisisyourdigitallife**". Dr. Aleksandr Kogan, psychology lecturer at the University of Cambridge. He founded a private company called **Global Science Research (GSR)**. This company signed the commercial contract with SCL Group (the parent company of Cambridge Analytica). The University of Cambridge allowed an employee to use their city and school name for a private company that was harvesting data. Taking advantage of this, he created an app using **University's logo and name** to make people trust it. The professor acted in a private capacity, made this app and gave it to the the Cambridge Analytica. This app was presented to users as a simple **personality quiz**. It looked fun and academical. It used the name of the University so, was looking reliable. The app fetched personal data of users who downloaded the app without their concern and also of then who were friends of users on Facebook. And then with the help of this data brain washing advertisements about the politicians were targeted and shown to the users considering their personality derived from the data collected. In other words, This data was collected in the year 2013-14was used to construct psychographic profiles for political micro-targeting during the 2016 U.S. election and Brexit via Facebook. This scandal was exposed in the year 2018.

Facebook CEO Mark Zuckerberg publicly apologized for the "major breach of trust," an admission that ultimately led to a historic \$5 billion fine from the U.S. Federal Trade Commission in 2019 and a significant, mandatory overhaul of the platform's data privacy regulations. Later Facebook made changes in its policies,stopped such apps and if anything suspicious was found it took efforts to timely investigate and ban it.

Here, both the Facebook and the University of Cambridge only looked at the name and the stated goal (personality quiz). They failed the “Human heed” to actually see where the data actually goes.

Uber

In the year 2018, Uber decided to test a self-driving car developed by it in USA. The car hit a pedestrian walking nearby and which further lead to death of the pedestrian. Its emergency braking system did not work.

The study show that the Ai had detected the person. But was confused about whether it is an static object or moving object coming in the way. Earlier it use to take frequent breaks for unnecessary reasons like some light, sun ray, signs, signals, etc, Therefore, the management had disabled its emergency braking system. The car had one back up driver, but it seemed that he was busy looking into the phone when this happened. Initially in the beginning Uber had two drivers, one to watch the road and another one to see how the system operates. The company reduced one driver to reduce the cost. Naturally, the one driver might have found it boring to just keep checking in how the car works and so, chose phone for entertainment relying on the car to operate properly and this mishap happened. This shows over-reliance on automation is harmful. This case showed Management oversight Failure.

Here, the management should have made corrections in the car system instead of disabling its emergency braking system and should have also kept a second driver.

As per the finding, **H1 (Alternative Hypothesis) stating** AI implementation in management creates significant ethical issues such as bias, privacy concerns, and lack of transparency is proven and so is accepted and **Ho(Null Hypothesis) stating** AI implementation in management does not create significant ethical issues is rejected.

Conclusion:

Artificial Intelligence (AI) does help management in decision making processes and improving efficiency but it as artificially created intelligence by Human being. It operates based on the data fed into it by the human or direction/ instructions given to it by humans and it functions according the data fed every time irrespective of the situation. It lacks logical thinking and reasoning ability by itself. Therefore, it must be noted that it a machine which can be an assistant but management cannot fully rely on it for decision making. Therefore, organizations must focus on responsible and transparent use of AI. Only then can AI contribute to sustainable and ethical business growth.

References:

- 1) Dastin, J. (2018, October 10). Amazon scraps secret AI recruiting tool that showed bias against women. *Reuters*. (Focus: The primary investigative report on the 2014-2017 project).
- 2) O'Neil, C. (2016). *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. Crown Publishing Group. (Focus: The sociological impact of "Black Box" hiring models).
- 3) Bambauer, J. (2019). *Cambridge Analytica and the Meaning of Privacy Harm*. George Mason University White Paper. (Focus: Defining the legal and philosophical harm of data misuse).
- 4) Cadwalladr, C., & Graham-Harrison, E. (2018). "Revealed: 50 million Facebook profiles harvested for Cambridge Analytica in major data breach." *The Guardian*. (Primary investigative source).

- 5) Isaak, J., & Hanna, M. J. (2018). "User Data Privacy: Facebook, Cambridge Analytica, and Privacy Protection." *Computer*, 51(8), 56-59. (Focus: Technical governance and platform accountability).
- 6) Adams-Prassl, J., et al. (2021). *Uber and Beyond: Policy implications for the UK*. Bennett Institute for Public Policy, University of Cambridge. (Focus: The link between employment status and regulatory safety).
- 7) Prassl, J. (2017). *Humans as a Service: The Promise and Perils of Work in the Gig Economy*. Oxford University Press. (Focus: The theoretical framework of platform accountability).