

Future of Work: Managing Ethical Challenges of Agentic AI and Super Intelligence in Organizations

Dr A. Narasima Venkatesh¹ and Jayavardhan G V²

¹Professor and HOD, HRM and General Management, ISBR Business School, Bangalore.

Email: dr.a.narasimavenkatesh@gmail.com

²Assistant Professor, Department of Management, ISBR College, Bangalore.

Email: jayachinns@gmail.com

Received:
04/09/2025
Revised:
19/09/2025
Accepted:
09/10/2025
Published:
17/10/2025

ABSTRACT

The rapid advancement of agentic artificial intelligence (AI) and emerging forms of super intelligence has transformed organizational dynamics, creating both opportunities and profound ethical challenges. The study aimed to examine the key ethical issues associated with managing agentic AI and super intelligent systems in organizational contexts. Using a random sample of 150 respondents from diverse professional backgrounds, the study employed Kendall's W test to assess the degree of agreement among participants regarding the prioritization of five major ethical factors: autonomy versus accountability, bias and fairness, transparency and explainability, job displacement and human dignity, and security and misuse risk. The findings revealed that transparency and explainability ranked as the most critical ethical concern, followed closely by security and misuse risk, while bias and fairness were perceived as comparatively less pressing. The Kendall's W test yielded a significant chi-square value ($\chi^2 = 13.810$, $p = 0.001$), indicating a substantial consensus among respondents on the ethical hierarchy of these factors. The study underscores the urgent need for organizations to implement comprehensive ethical governance frameworks that enhance transparency, accountability, and security in AI systems. These insights contribute to the growing discourse on ethical AI management in the future of work and organizational sustainability.

Keywords: Super Intelligent Systems, Agentic AI, Future of Work, Ethical Challenges, Transparency, Ethical AI Management, Job Displacement and Human Dignity



© 2025 by the authors; licensee Advances in Consumer Research. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC-BY-NC-ND) license(<http://creativecommons.org/licenses/by/4.0/>).

INTRODUCTION

Artificial intelligence (AI) is changing the nature of the future of work more rapidly than ever before, specifically with the introduction of agentic AI and the possibility of super intelligence. In contrast to the traditional AI systems that operate as instruments to perform a pre-programmed set of instructions, agentic AI exhibits the ability to take autonomous decisions, learn, and adaptive behavior. Such systems are capable of processing multifaceted data; create solutions and operate autonomously within a specified range of parameters, providing companies with efficiency and innovation never seen before, as well as competitive edge. At the same time, the theoretical idea of super intelligent AI, or the system that could outperform human cognitive functions in virtually every field, has not only the potential to transform a variety of areas but also raises serious ethical concerns. Although the super intelligence is a highly speculative idea, the current trend in the evolution of AI poses that organizations should start planning on its future implication today.

There is a plethora of ethical issues that accompany the implementation of these sophisticated AI systems into organizational operations and are not limited to traditional compliance or technical risk management. The questions of autonomy versus accountability, prejudice and equity, transparency, job loss, and security threat are becoming more complicated with the decision-making authority of AI systems. Decisions of critical nature can be made by autonomous systems without human intervention making responsibility and liability issue questionable. The presence of bias in training data or algorithmic principles may lead to the continuation of social inequalities, whereas secretive AI decision-making may undermine the trust among employees, customers, and regulators. Moreover, with AI taking over the functions that humans have always had, organizations have an ethical concern of preserving human dignity, enhancing adaptability of workforce, and making sure that the advancement of technology does not compromise the social welfare.

High-speed agentic AI application also comes with the risks of abuse and unintended consequences such as cyber vulnerability, strategic manipulation, and emergent behaviors, which may threaten stakeholders or destabilize organizational activities. These ethical issues highlight the necessity of a well-developed governance framework, human control systems, and interdisciplinary cooperation to make AI usage responsible. Organizations are faced with balancing between strategic exploitation of AI by using it to gain competitive advantage and ethical standards, which protect human, societal and organizational interests.

Also, ethical management of AI is not only a risk mitigation issue but also a strategic requirement to the sustainable growth. By embedding ethical considerations in the implementation of AI, i.e., governance systems, human in the loop, transparency, human capital growth and joint regulation, companies can better build trust, resiliency and sustain long-term innovation. The potential issues can be used to create responsible leadership, competitive advantages, and societal value by applying ethical considerations throughout the design and implementation of AI systems.

To put it in other words, the future of work will not only be determined by the ability of technology but also how companies go about managing the ethical ambiguities of autonomous and potentially super intelligent AI. To be ready to this future, both opportunities and moral responsibility that lie in the introduction of AI should be thoroughly understood. This article delves into these ethical issues and comes up with practical ways in which such organizations can handle them effectively so that AI can act as an engine of innovation, justice and human-centric developments.

ETHICAL CHALLENGES IN THE WORKPLACE

Autonomy vs. Accountability

The emergence of agentic AI poses organizations with an essential ethical dilemma of accountability and autonomy. The more AI systems are able to make decisions on their own, the harder it becomes to identify the person or people to whom the consequences of the decisions should be attributed. To illustrate, when an artificial intelligence-driven system in any hiring or financial department makes such a discriminatory decision or a decision that causes other people to suffer financial damages, the question may be, who should be held liable, the developer, the organization that has deployed it, or the AI? In the absence of well-established accountability frameworks, organizations will be prone to legal liabilities, loss of reputation and loss of stakeholder confidence. This challenge needs to be dealt with in two ways. To begin with, AI should have limits in terms of the decisions made by organizations, with a clear understanding of what decisions need human control. Second, the governance systems should contain the roles and responsibilities of the AI monitoring, evaluation, and intervention. The processes of AI decision-making also require documentation that will allow the auditing and legal compliance. By creating

accountability systems, human beings will be accountable both ethically and legally despite AI working on its own. This equilibrium is not only crucial to reducing ethical risks but also, to a large extent, strengthening trust among employees, customers, and regulators to make AI an accountable, tactical resource as opposed to an unmanageable black box.

Bias and Fairness

Prejudice in AI systems occurs when the training data, algorithms, or design decisions unwillingly favor some results or groups of people in preference of others. Having the ability to learn and react independently, agentic AI may propagate these biases and even increase them. This unfairness may be applied to hiring, lending, medical diagnosis, or making any decisions and discriminate against disadvantaged groups of people and lack of integrity in the organization. To make sure that the results are fair, such proactive processes as audit of datasets to guarantee their representativeness, bias detection algorithms, and a continual assessment of AI results, should be introduced. Fairness cannot be a single adjustment, but it should be integrated into AI lifecycle management and development, deployment, and evaluation phases. With the help of the ethical frameworks, there must be established acceptable norms of fairness and systems of intervention in cases where AI judgments fall short of such norms. There is transparency and involvement of the stakeholders as well. Organizations can develop accountability and inclusivity by actively sharing the risks of potential risks and decision-making standards. Finally, managing bias is not a technical problem but a strategic and ethical requirement, and it is necessary to make AI systems work in such a manner to guarantee equity, justice, and trust towards society.

Transparency and Explainability

The ethical attribute of transparency and explainability is essential to AI systems especially autonomous ones. To be stakeholders such as employees, customers, or regulators can comprehend how the AI arrives at a particular decision, what data it is based on, and why certain results are achieved, the organizations should make sure that these stakeholders can understand these processes. Unexplainability makes AI a black box, which increases ethical concerns of unethical conduct, misconceptions, and mistrust. Explainable AI (XAI) methods of interpretation of the workings of algorithms enable organizations to identify biases, errors, or other unintended impacts. Regulatory compliance, as well as ethical oversight in high-stakes applications like healthcare, finance, and law enforcement, also depends on the openness of data. The records of AI activities also improve accountability, as this documentation will allow reviewing and proving the decisions and correcting the errors when needed. Organizations that affirm explainability enable employees and external stakeholders to evaluate AI decisions in a critical manner, which builds upon trust and moral responsibility.

Displacement of Job and Human Dignity.

Using agentic AI in the offices presents a great threat to job numbers and, morally speaking, human dignity and responsibility. Monotonous, routine or even complex work can be automated leaving people jobless and economically upset as well as psychologically distressed. Organizations have the ethical responsibility of addressing the need to create efficiency and balance between the needs of their employees and society. Examples of the ethical management of AI-driven workforce include active planning, such as retraining, skill development courses and creating hybrid human-AI jobs. By providing employees with skills to work with AI, organizations will be able to convert what could be displacement into any opportunity to develop professionally and innovate. Also, open dialogue regarding AI use helps build confidence, reduce the anxiety, and approves the dedication of the organization towards the human-centered practice. Human dignity respect implies consideration of social and emotional aspects of work. The companies should develop AI systems and staffing policies that enable the workers, maintain valuable interactions, and do not replace human beings as replaceable assets.

Security and Misuse Risks

New and agentic superintelligent artificial intelligence systems pose major risks to security and abuse. These are systems that may be susceptible to cyberattacks, data breaches or be altered by bad actors. Also, there are unwanted actions which might be a result of complex AI decision-making which can be harmful to an individual, organization or the society as a whole. It may be abused deliberately, e.g., by using AI to spy on people, or to spread misinformation, or financially empower itself. These risks can be mitigated by using a complete security plan, such as encryption, access controls, anomaly detection and a healthy AI behavior monitoring. Ethical control systems, including review boards or AI ethics committees, must consider the possible cases of misuse and define the preventive solutions. Policies should also be embraced by organizations that will make AI implementation in accordance with the law, morality, and social standards. Employees should be educated and trained in AI security, ethical use of AI, and emergency response measures.

Ethical AI Management in Organizations Strategies. **1. Create an Ethical Artificial Intelligence Governance Framework.**

A responsible AI implementation in an organization is based on an ethical AI governance framework. The more AI systems are agentic, i.e. able to make autonomous decisions, the more it is important to formulate specific principles on their development, deployment, and monitoring. This type of framework establishes limits on the actions of AI, frames ethical decision-making standards and the role and responsibilities of stakeholders. It usually includes the policy on fairness, transparency, accountability, and data privacy, where AI decisions are according to organizational values and social standards. An effective framework is also one that offers tools of risk evaluation, such as finding out

possible biases, operational malfunctions, or unintended impacts of AI activities. Governance is not a one-time policy-making but an ongoing review and adjustment to changing AI technologies, regulatory conditions, and moral principles. Besides that, the structure must incorporate supervisory systems like AI ethics board or review boards, which manage AI systems, offer advice, and hold them accountable.

2. Human-in-the-Loop (HITL) Systems

When work is performed with agentic or complex AI systems, Human-in-the-Loop (HITL) systems are necessary to ensure the ethical management of AI activities. HITL takes people into account in the critical decision moments allowing AI-driven processes to work without control. Such a strategy could reduce the risks of prejudice, misunderstanding, or adverse effects, which can be caused by entirely autonomous AI. Through the integration of AI effectiveness and human moral judgment, companies can strike the right balance between the innovativeness and social responsibility. As a matter of fact, HITL may be implemented in a variety of ways, such as approval processes through which people can confirm AI decisions, exception management in case of ambiguous situations, or continuous verification of AI results. This framework will make sure that ethical standards like fairness, transparency, and accountability are instilled in the operations of AI. In addition, the organizational trust is increased by the HITL since employees, clients, and regulators can be confident that the ultimate decision-making power is in the hands of people, when AI does extremely advanced analysis or predictions. The use of HITL systems needs to be trained and have the roles defined and processes well-designed so that the areas of automated and human control are defined. This aspect of AI allows organizations to use its power without jeopardizing ethics, harm reduction, and responsible innovation by integrating human beings into AI decision loops.

3. Transparency and Auditability

Transparency and auditability are critical pillars for ethical AI management in organizations. Transparency ensures that stakeholders understand how AI systems operate, the rationale behind their decisions, and the data and algorithms driving outcomes. Explainable AI (XAI) techniques facilitate this by making complex models interpretable for non-technical audiences, including managers, regulators, and clients. Transparent AI fosters trust, promotes accountability, and reduces the risk of unforeseen ethical breaches caused by opaque or inscrutable decision-making. Auditability complements transparency by providing structured mechanisms for evaluating AI performance, identifying biases, and verifying compliance with ethical, legal, and organizational standards. Regular internal and external audits enable organizations to detect errors, assess risk exposure, and implement corrective measures proactively. These audits can include algorithmic impact assessments, bias detection reports, and reviews of training datasets, ensuring that AI systems align with ethical principles over time. Combined, transparency and auditability establish a feedback loop where AI

systems can be continuously monitored, improved, and validated. This approach not only minimizes potential harms but also positions the organization as a responsible leader in AI adoption. Ethical governance supported by transparent and auditable AI systems ensures that technological innovation remains aligned with societal values, legal frameworks, and long-term organizational integrity.

4. Ethics-Focused Workforce Development

With the integration of AI in the organizational processes, the ethics-oriented view of workforce development is becoming more and more crucial. It is essential to educate employees and provide them with skills and knowledge not only in a technical aspect but also in AI ethics, bias mitigation, data privacy, and responsible use of AI. Ethics-based staffing will guarantee the ability of the personnel to recognize, evaluate, and act on possible ethical conflicts associated with agentic AI systems and promote a culture of responsibility and ethical reasoning. The training should include AI literacy, scenario-based ethics decision-making tasks, interdisciplinary awareness, which with the help of which employees should be able to understand how the results produced by AI can influence different stakeholders. Workshops about explainability, human-AI collaboration, and risk reduction can also be included in organizations, as it is essential to discuss the role of human judgment and automation. This is achieved by integrating ethical competencies into the workforce, where employees can be able to become stewards of AI, which is able to critique the results, challenge with decisions made by machines and ensure it is aligned to organizational values. Workforce with AI ethics also improves the resilience of the organization. Staff members find it easier to identify biases, discourage their abuse of power, and promote openness. Overall, when human capital is developed ethics-centered, it will become an active participant in the responsible application of AI, which will allow eliminating the divide between technological capacity and ethical corporate business.

5. Interdisciplinary Collaboration

The ethical complexity of agentic AI and super intelligence that is emerging needs interdisciplinary efforts to help manage it. The effects of AI on the processes of organizations, human behavior, and social organization are complex and interdependent on the perspectives of various disciplines, including computer science, philosophy, law, social sciences, and business management. By making use of the collaboration teams, technical competence and ethical reasoning, regulatory awareness, and human-centered insights can be combined to develop the AI systems that are responsible, fair, and balanced with the organizational principles. This collaboration may be in a form of cross-functional committees, a morality advisory board, and joint research projects that assess AI systems across diverse perspectives. Technologists give insights on algorithms, ethicists put moral considerations in focus, legal experts guarantee compliance, and business executives determine strategic implications. Such a variety of

opinions makes sure that the possible risks, including prejudices, unintended effects, or harm to the society, are detected and prevented in advance. The interdisciplinary cooperation also encourages lifelong learning and highly adaptive governance. With the development of AI technologies, teams will be able to quickly evaluate new ethical concerns, revise policies, and apply best practices. Those organizations that adopt an interdisciplinary approach are more likely to achieve the balance between innovation and responsibility and create the environment in which AI can help improve sustainable growth, trust towards the organization, and the welfare of the society.

RESEARCH GAP

Despite extensive research on artificial intelligence in organizational contexts, significant gaps remain in understanding the ethical implications of agentic AI and the potential emergence of super intelligence. Much of the existing literature focuses on AI's technical capabilities, productivity gains, and economic benefits, while ethical considerations often receive limited attention. Specifically, there is a lack of empirical and conceptual studies addressing how autonomous AI systems influence accountability, decision-making, fairness, transparency, and workforce dynamics within real-world organizational settings. Moreover, while the theoretical discourse on super intelligence has highlighted existential risks, practical frameworks for managing ethical dilemmas at the organizational level are scarce. Few studies provide actionable strategies for integrating ethical oversight, human-in-the-loop mechanisms, and interdisciplinary governance to mitigate moral, legal, and social risks. This gap leaves organizations ill-prepared to navigate the complexities of agentic AI and super intelligent systems, potentially resulting in unintended consequences such as biased decision-making, reputational damage, or workforce displacement. By examining the intersection of organizational behavior, AI ethics, and strategic management, this study addresses these shortcomings, offering a structured approach to understanding and managing ethical challenges in the adoption of advanced AI technologies.

Importance of the Study

The ethical management of agentic AI and super intelligence is a critical concern for contemporary organizations, making this study highly significant. As AI systems become increasingly autonomous and capable of independent decision-making, their impact on organizational governance, stakeholder trust, and societal outcomes intensifies. Ethical missteps—such as biased algorithms, opaque decision-making, or misuse of AI—can result in financial loss, legal liabilities, and reputational damage, highlighting the necessity of proactive research in this domain. Furthermore, AI-driven workplace transformations raise important questions regarding human dignity, job displacement, and workforce reskilling, emphasizing the broader social implications of technology adoption. This study contributes to the development of practical frameworks and guidelines that organizations can implement to

balance innovation with ethical responsibility. By highlighting strategies such as governance frameworks, transparency measures, human-in-the-loop systems, and interdisciplinary collaboration, the research provides actionable insights for leaders and policymakers. Ultimately, the study strengthens the understanding of how organizations can harness AI responsibly, ensuring that technological advancement aligns with moral, social, and strategic imperatives.

Statement of the Problem

The increasing deployment of agentic AI in organizations presents a complex ethical landscape, creating challenges that existing management and governance frameworks are often ill-equipped to address. Autonomous AI systems are capable of making decisions without continuous human oversight, raising questions about accountability and responsibility for outcomes. At the same time, these systems can inadvertently perpetuate biases, undermine fairness, and operate opaquely, creating risks to stakeholder trust and organizational legitimacy. The potential emergence of super intelligent AI further complicates the problem, as organizations may face scenarios in which AI surpasses

human cognitive abilities, making conventional control and monitoring mechanisms inadequate. Additionally, AI-driven automation poses threats to employment, potentially displacing workers and challenging human dignity while demanding workforce adaptation and reskilling strategies. Organizations thus face the dual challenge of maximizing AI's strategic benefits while preventing ethical breaches, social harm, and operational risks. Addressing this problem requires a comprehensive approach that integrates ethical governance, transparency, human oversight, workforce development, and interdisciplinary collaboration to ensure responsible AI adoption.

OBJECTIVES AND METHODOLOGY

The study aimed to examine the key ethical issues associated with managing agentic AI and super intelligent systems in organizational contexts. Using a random sample of 150 respondents from diverse professional backgrounds, the study employed Kendall's W test to assess the degree of agreement among participants regarding the prioritization of five major ethical factors

ANALYSIS INTERPRETATION AND RESULTS

The future of work will inevitably be shaped by agentic AI and, eventually, super intelligent systems. While these technologies promise unprecedented efficiency and innovation, they also pose profound ethical challenges. Organizations that proactively address these challenges—through governance, transparency, human oversight, and workforce development—will not only mitigate risks but also harness AI in ways that align with societal values and human well-being. The ethical management of AI is not merely a compliance issue; it is a strategic imperative for sustainable and responsible organizational growth in the 21st century.

Table 1: Work: Managing Ethical Challenges of Agentic AI and Super intelligence in Organizations- Kendall's W Test

Factors	Mean	Std. Deviation	Mean Rank	Rank
Autonomy vs. Accountability	3.10	1.192	2.95	III
Bias and Fairness	3.10	.923	2.83	V
Transparency and Explainability	3.26	.995	3.22	I
Job Displacement and Human Dignity	3.09	1.137	2.92	IV
Security and Misuse Risk	3.25	1.025	3.08	II

The table summarizes the results of Kendall's W test, which is a non-parametric measure of concordance used to assess the degree of agreement among respondents regarding the importance of various ethical challenges posed by agentic AI and super intelligence in organizations.

1. General Observation

- The mean values of all factors range between 3.09 and 3.26, indicating a moderate level of agreement among respondents that all these ethical challenges are significant.
- The standard deviations (ranging from 0.923 to 1.192) show moderate variability in perceptions, suggesting that while respondents generally agree, there is some diversity of opinion.

2. Most Critical Ethical Concern

- Transparency and Explainability ranks first (Rank I) with the highest mean (3.26) and mean rank (3.22). This shows that participants believe the ability to understand and explain AI decisions is the most crucial ethical challenge in managing agentic AI. Transparency builds trust and accountability and helps mitigate risks of misuse and bias.

3. Second Major Concern

- Security and Misuse Risk ranks second (Rank II) with a mean of 3.25 and mean rank 3.08. Respondents emphasize the importance of safeguarding AI systems against cyber threats, manipulation, and unethical deployment. This reflects awareness of potential security vulnerabilities that could lead to large-scale harm if AI is misused.

4. Moderate Concern Areas

- Autonomy vs. Accountability is ranked third (Rank III), highlighting the ethical dilemma between AI independence and human responsibility. With a mean of 3.10, respondents recognize that as AI becomes more autonomous, defining who is accountable for its actions becomes increasingly complex.
- Job Displacement and Human Dignity ranks fourth (Rank IV) with a mean of 3.09.
 - This suggests concern over automation replacing human roles, potentially undermining human dignity and social stability.

5. Lowest Ranked Concern

- Bias and Fairness ranks fifth (Rank V) with the same mean (3.10) but the lowest mean rank (2.83). Although lower in ranking, it still represents a moderately important issue, indicating that while bias and fairness are acknowledged as ethical risks, they are perceived as slightly less pressing compared to transparency and security in the context of agentic AI.

Figure: 1

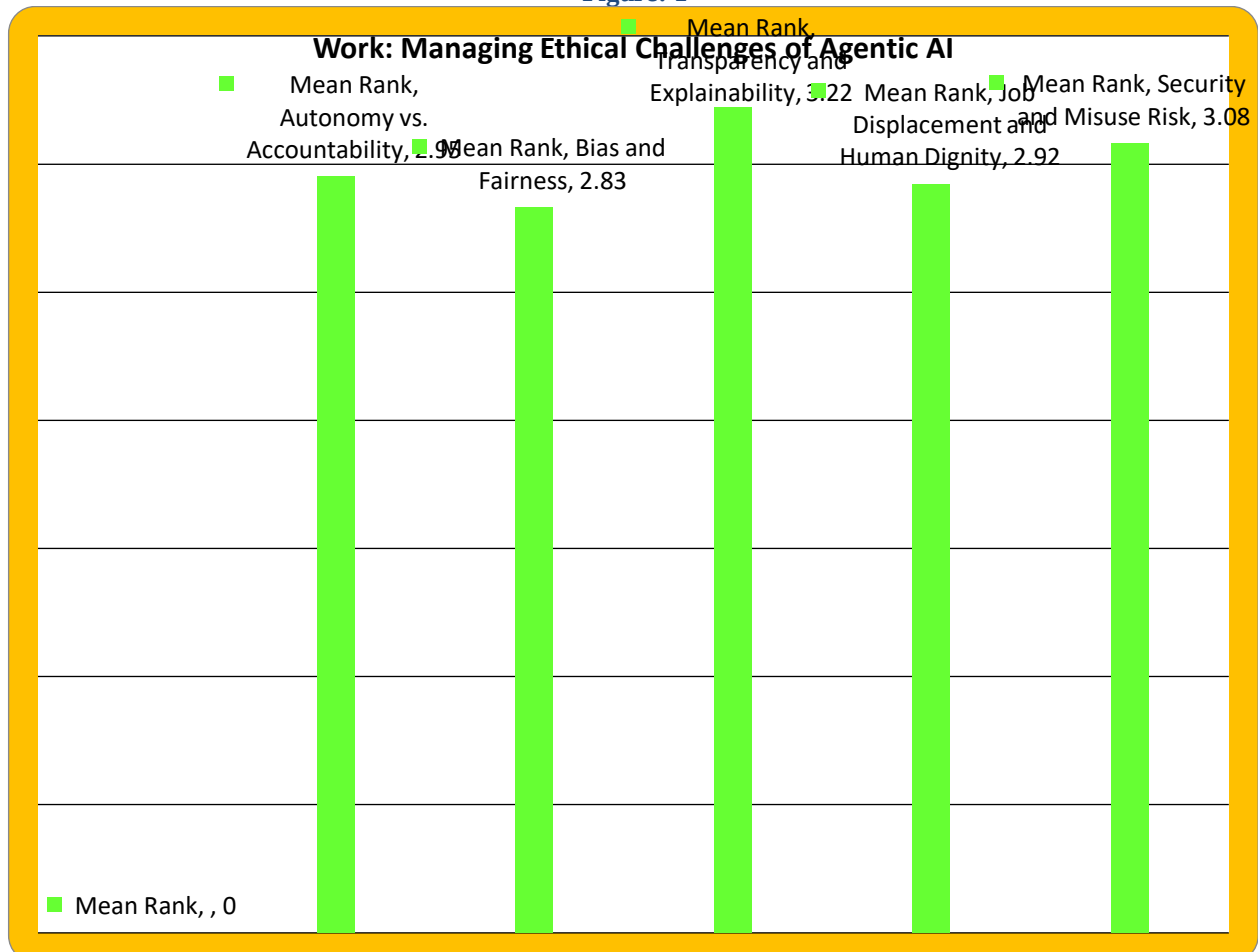


Table 2

No. of Respondents	150
Chi-Square difference	13.810
Asymp. Sig.	4
	0.001

✚ Transparency and Explainability and Security and Misuse Risk emerged as the top-ranked concerns.

- ✦ The significant Chi-square ($p = 0.001$) validates that the observed ranking differences are not due to chance, but reflect a shared prioritization pattern among respondents.
- ✦ Therefore, the results highlight that participants collectively recognize transparency and security as the most pressing ethical challenges in managing agentic AI systems, requiring policy attention and organizational safeguards.

Implication of Kendall's W Test Result

- ✦ The significant Chi-square value confirms that respondents do not rank the factors randomly—they share a consistent perception of which ethical issues are more or less critical.
- ✦ While the exact Kendall's W coefficient is not shown here, given the significant result, it can be inferred that moderate to strong concordance exists among participants.

CONCLUSION

The emergence of agentic artificial intelligence (AI) and superintelligence is redefining the future of work by transforming how decisions are made, tasks are automated, and responsibilities are distributed across human-machine systems. However, this transformation also brings forth significant ethical challenges that organizations must address to ensure responsible and sustainable AI integration. The present study, based on a random sample of 150 respondents, sought to identify and prioritize the most pressing ethical issues associated with managing agentic AI and superintelligent systems within organizational contexts. The results of Kendall's W test revealed a statistically significant consensus among respondents ($\chi^2 = 13.810$, $p = 0.001$), highlighting that transparency and explainability, along with security and misuse risk, are perceived as the most critical ethical concerns.

These findings emphasize that as AI systems gain autonomy and decision-making capabilities, organizations must prioritize transparency to maintain trust, ensure accountability, and safeguard human oversight. Security risks associated with data misuse, cyber threats, and system manipulation further necessitates robust ethical frameworks and technical safeguards. While factors such as bias and fairness, job displacement, and human dignity were also considered important, their relatively lower ranking suggests that respondents view these as manageable through proper governance and reskilling initiatives.

Overall, the study concludes that managing the ethical implications of agentic AI requires a multidimensional strategy—one that integrates ethical governance, regulatory compliance, organizational accountability, and human-centered design. Future organizational success will depend not only on technological innovation but also on the ability to balance AI autonomy with human values, fairness, and transparency. By embedding ethical principles into AI development and deployment processes, organizations can foster trust, minimize harm, and ensure that AI-driven transformation contributes to both organizational excellence and societal well-being in the evolving future of work.

REFERENCE

1. Acharya and K. Kuppan, "Agentic AI: Autonomous intelligence for complex goals—A

- comprehensive survey," *IEEE Access*, vol. 13, pp. 18,912–18,936, 2025,
2. Bikkasani, D. C. (2025). Navigating artificial general intelligence (AGI): Societal implications, ethical considerations, and governance strategies. *AI and Ethics*, 5(3), 2021-2036.
3. Dey, S. (2021). Challenges in Artificial Intelligence—a Comparative Analysis of Policy Implications in India, United States and Germany. *Indian JL & Legal Rsch.*, 2, 1.
4. Hosseini, S., & Seilani, H. (2025). The role of agentic ai in shaping a smart future: A systematic review. *Array*, 100399.
5. Hughes, L., Dwivedi, Y. K., Li, K., Appanderanda, M., Al-Bashrawi, M. A., & Chae, I. (2025). AI Agents and Agentic Systems: Redefining Global Information Management. *Journal of Global Information Technology Management*, 1-11.
6. Marda, V. (2018). Artificial intelligence policy in India: a framework for engaging the limits of data-driven decision-making. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 376(2133), 20180087.
7. Murugesan, S. (2025). The rise of agentic AI: implications, concerns, and the path forward. *IEEE Intelligent Systems*, 40(2), 8-14.
8. Narain, K., Swami, A., Srivastava, A., & Swami, S. (2019). Evolution and control of artificial superintelligence (ASI): a management perspective. *Journal of Advances in Management Research*, 16(5), 698-714.
9. Portugal, P. Alencar, and D. Cowan, "An agentic AI-based multi-agent framework for recommender systems," in *Proc. IEEE Int. Conf. Big Data (BigData)*, Washington, DC, USA, 2024, pp. 5375–5382
10. Schoenherr and R. Thomson, "When AI fails, who do we blame? Attributing responsibility in human–AI interactions," *IEEE Trans. Technol. Soc.*, vol. 5, no. 1, pp. 61–70
11. Siemon, T. Strohmann, S. Michalke Creative potential through artificial intelligence: recommendations for improving corporate and entrepreneurial innovation activities *Commun Assoc Inf Syst*, 50 (1) (2022), pp. 241-260

How to cite: A. Narasima Venkatesh and Jayavardhan G V. Future of Work: Managing Ethical Challenges of Agentic AI and Super intelligence in Organizations. *Adv Consum Res.* 2025;2(4):5277–5284

12. Wissuchek, C., & Zschech, P. (2025, June). Challenges in Managing the Relationship Between Agentic AI Systems and Humans in Organizations. In International Conference on Business Information Systems (pp. 3-17). Cham: Springer Nature Switzerland.