

Artificial Intelligence in Action: Navigating Advancements, Pitfalls, and Human Dynamics

Vernette Grant¹, Robert E. Levasseur²

¹Aventura, FL, USA

²College of Management, Walden University, Minneapolis, MN, USA

Email: vernettegrant1@gmail.com, robert.levasseur@mail.waldenu.edu

How to cite this paper: Grant, V., & Levasseur, R. E. (2025). Artificial Intelligence in Action: Navigating Advancements, Pitfalls, and Human Dynamics. *Open Journal of Business and Management*, 13, 3865-3874. <https://doi.org/10.4236/ojbm.2025.136210>

Received: October 2, 2025

Accepted: October 27, 2025

Published: October 30, 2025

Copyright © 2025 by author(s) and Scientific Research Publishing Inc.

This work is licensed under the Creative Commons Attribution International License (CC BY 4.0).

<http://creativecommons.org/licenses/by/4.0/>



Open Access

Abstract

Artificial intelligence (AI) is revolutionizing organizations worldwide, with rapid advancements yielding both unprecedented opportunities and new risks. In this article, we synthesize scholarly literature and organizational case studies to examine the current state of AI, evaluating its key advantages—such as efficiency, decision support, and scalability—and downsides, including bias, job displacement, and ethical dilemmas. The analysis explores how AI shapes organizational behavior, leadership, and human-AI collaboration through a novel Human-AI Integration Continuum framework. The discussion incorporates the authors' previous research and contemporary developments, concluding with recommendations for responsible AI stewardship and expanded future research directions, while highlighting the evolving human context at the heart of technological progress.

Keywords

Artificial Intelligence, Organizational Behavior, AI Ethics, Automation, Trust, Human-AI Collaboration, Leadership

1. Introduction

Artificial intelligence (AI) has rapidly moved from speculative technology to a central component of organizational life (McKinsey & Company, 2025; Stanford Human-Centered Artificial Intelligence (HAI), 2025). In healthcare, finance, and communication, AI now routinely augments or automates decision-making, accelerates workflows, and provides novel forms of personalization (Taherdoost & Madanchian, 2023). The swift proliferation of AI elicits both hope and unease: AI promises operational revolution but raises legitimate concerns regarding bias, transparency, human displacement, and regulatory adequacy (Nikolova & Angrisani,

2025; Suran & Hswen, 2024).

Organizational change through AI mirrors ongoing challenges in diversifying leadership within healthcare. Previous research highlights that systemic bias and inequities—long documented in the advancement of women to executive roles—continue to impede progress, even in technologically advanced settings (Grant & Levasseur, 2025). Integrating lessons from glass ceiling studies into AI governance underlines the necessity for intentional diversity and ethical oversight in leadership as the field evolves.

This article is a review of the landscape of AI in 2025, exploring its practical effects, human dimensions, and implications for sustainable, ethical integration within organizations.

2. Methodology

In this article, we employ a narrative review and integrative synthesis methodology, following established frameworks for examining complex, interdisciplinary subjects such as AI adoption and organizational effects (Yoo et al., 2024).

2.1. Search Strategy and Sources

This review is the result of systematically examining scholarly articles published between 2021–2025, with particular emphasis on 2024–2025 publications to capture the most current AI developments. Primary sources included peer-reviewed journals focusing on artificial intelligence, organizational behavior, technology management, and ethics. Institutional reports from authoritative sources (McKinsey & Company, 2025; Stanford Human-Centered Artificial Intelligence (HAI), 2025) provided contemporary organizational data and trend analysis. Policy documents, particularly those addressing regulatory frameworks such as the Council of Europe’s AI Convention, were included to contextualize the governance landscape (van Kolfschooten & Shachar, 2023).

2.2. Inclusion Criteria

Sources were selected based on: 1) direct relevance to AI implementation in organizational contexts, 2) empirical findings on human-AI interaction dynamics, 3) documented cases of AI benefits and challenges, and 4) methodological rigor in data collection and analysis. Priority was given to studies examining real-world AI deployment rather than theoretical frameworks alone.

2.3. Analytical Approach

Analysis centered on thematic extraction across four primary domains—operational efficiency, algorithmic bias and fairness, leadership transformation, and human-AI collaboration patterns. This thematic framework emerged inductively from the literature while being refined through the primary author’s prior research experience in organizational AI implementation. Cross-referencing of findings across multiple sources strengthened the validity of identified patterns

and trends.

The integrative synthesis methodology enabled the examination of convergent and divergent findings across disciplines, allowing for the identification of gaps where empirical evidence remains limited and highlighting areas where practitioner experience aligns with or challenges academic findings.

3. Results—State of AI in 2025

3.1. Organizational Adoption Patterns

McKinsey & Company (2025) data indicates widespread AI adoption now characterizes over 80% of major organizations, with cross-sector investment at historic heights and particularly strong uptake in healthcare diagnostics and financial services (McKinsey & Company, 2025). However, Stanford Human-Centered Artificial Intelligence (HAI), (2025) analysis reveals significant variation in implementation success, with organizations reporting 23% average productivity gains where human-AI collaboration models were prioritized versus 8% gains in automation-focused deployments (Stanford Human-Centered Artificial Intelligence (HAI), 2025). Generative AI models drive content creation, drug discovery, and learning environments but introduce new concerns over authenticity and manipulation (Singh et al., 2025). Affordable, sector-specific AI has leveled access for smaller firms and non-technical domains (Taherdoost & Madanchian, 2023).

3.2. Trust and Performance Dynamics

Kim and Park's (2024) behavioral analysis of 1247 financial decision-makers demonstrated that trust in AI advice increased 34% when transparency mechanisms were implemented, but decreased 18% following system errors, highlighting the fragile nature of human-AI trust relationships. This finding was corroborated by Gerlich (2024) in a qualitative study of 156 professionals that identified “competence anxiety”—a fear of being rendered obsolete or incapable due to the introduction of new technology—as a primary barrier to AI adoption, particularly among mid-career workers. Users appreciate efficiency but question objectivity and ethical soundness (Gerlich, 2024; Nikolova & Angrisani, 2025). Regulatory frameworks—such as the Council of Europe's AI Convention—have been developed to establish standards, especially around privacy and health, yet global harmonization remains elusive (van Kolschooten & Shachar, 2023).

3.3. Benefits and Pitfalls of AI

Advantages: Automation of routine work, improved decision-making, scalable personalization, and enhanced workplace health and safety are well-documented. Fiegler-Rudol et al.'s (2025) narrative review of workplace health applications documented measurable improvements in hazard detection (average 41% reduction in workplace incidents) and ergonomic optimization (Fiegler-Rudol et al., 2025; Levasseur, 2025).

Disadvantages: Job displacement affects both low- and mid-skill roles, while algorithmic bias persists despite technological advances. [de Bruijn et al.'s \(2022\)](#) analysis highlighted persistent concerns with algorithmic fairness and transparency in organizational settings, supporting the continuing relevance of these issues. Additional concerns included overdependence (risk of deskilling) and increasing threats to data privacy and security ([de Bruijn et al., 2022](#); [Voelker, 2023](#)).

3.4. Organizational Behavior and Impact

AI-driven recruitment and talent management optimize job matching but may perpetuate organizational biases ([Yoo et al., 2024](#)). Leadership styles are redefined by demands for transparency, trust, and ethical oversight ([de Bruijn et al., 2022](#)). Employees experience both liberation from menial tasks and anxiety due to obsolescence and skill gaps. [Fiegler-Rudol et al. \(2025\)](#), [Gerlich \(2024\)](#), and [Li et al. \(2025\)](#) identified increased psychological stress among workers concerned about surveillance and job displacement.

Effective human-AI collaboration emerges when organizations strategically deploy artificial intelligence to augment rather than substitute human capabilities, leveraging AI's computational strengths in data processing and pattern recognition while preserving human roles in critical judgment, ethical decision-making, and contextual interpretation ([Fiegler-Rudol et al., 2025](#)). The efficacy of such partnerships depends fundamentally on the implementation of transparent algorithmic processes and deliberate system design that facilitates shared decision-making frameworks, wherein humans retain interpretive authority over AI-generated insights and maintain adaptive control over technological integration as organizational requirements evolve ([Levasseur, 2025](#); [Taherdoost & Madanchian, 2023](#)). This synergistic approach, which combines computational efficiency with human expertise and value systems, enables organizations to establish collaborative ecosystems that maximize the complementary strengths of both human and artificial intelligence, ultimately yielding enhanced performance outcomes ([de Bruijn et al., 2022](#)).

Another critical factor in AI implementation success is an organization's readiness for digital transformation across technical, human resource, and cultural domains. Organizations that proactively assess their digital maturity—considering infrastructure, workforce training, and change management strategies—tend to realize higher returns on AI investments and demonstrate more sustainable integration outcomes ([McKinsey & Company, 2025](#); [Yoo et al., 2024](#)). By embedding continuous professional development and digital upskilling initiatives, organizations can mitigate resistance to change while fostering a culture of innovation that supports human-AI synergy ([Taherdoost & Madanchian, 2023](#)).

Furthermore, the role of emotional intelligence is increasingly recognized in navigating human-AI collaboration, especially as AI systems become more prevalent in workplaces requiring complex interpersonal interactions. [Li et al. \(2025\)](#) underscored that workers' emotional responses to AI-driven organizational

changes can mediate both acceptance and perceived fairness of algorithmic decisions. Developing emotional intelligence among leaders and teams is vital for cultivating psychological safety in AI-integrated workplaces. Emotional attunement enables leaders to recognize signs of competence anxiety, address resistance empathically, and promote open dialogue around ethical and performance concerns. As Li et al. (2025) highlighted, emotional regulation during technological disruption mediates acceptance and fairness perceptions, directly influencing collaboration quality and innovation outcomes. As such, organizational leaders must prioritize transparent communication and empathetic engagement to address anxieties and psychological impacts associated with automation and surveillance (Gerlich, 2024; Li et al., 2025).

Lastly, the rapid evolution of AI technologies calls for interdisciplinary collaboration that spans technical, ethical, and behavioral expertise. Engaging cross-functional teams—including ethicists, technologists, policy experts, and frontline practitioners—in the design, deployment, and ongoing assessment of AI initiatives helps ensure that systems are contextually relevant, ethically grounded, and aligned with organizational missions (de Bruijn et al., 2022; van Kolschooten & Shachar, 2023). This collaborative, systems-thinking approach supports not only compliance and risk mitigation but also the broader goal of inclusive and equitable innovation.

4. Discussion

4.1. Novel Synthesis Framework: The Human-AI Integration Continuum

Unlike traditional technology adoption models such as the Technology Acceptance Model (TAM) and the Unified Theory of Acceptance and Use of Technology (UTAUT), which focus primarily on user intention or system utilization, the Human-AI Integration Continuum extends these perspectives by emphasizing the evolving relational dynamics between humans and AI systems across multiple levels of organizational maturity. This continuum captures the co-adaptive interplay between technology capability, human behavior, and ethical governance—elements often treated separately in earlier models (de Bruijn et al., 2022; Yoo et al., 2024). The literature and case studies converge on several themes, revealing what we identify as a Human-AI Integration Continuum—a novel conceptual framework that synthesizes findings across organizational behavior, technology adoption, and ethics literature. Unlike previous reviews that examine AI benefits and challenges in isolation, this framework identifies three critical integration stages that organizations navigate.

Stage 1: Operational Augmentation At this stage, AI primarily enhances efficiency and automates routine tasks, as evidenced in innovative AI techniques (Taherdoost & Madanchian, 2023) and workplace safety enhancements (Fiegler-Rudol et al., 2025). Organizations at this stage experience the most straightforward benefits but also face initial resistance and competence anxiety (Gerlich,

2024).

Stage 2: Decision Collaboration At this stage, human judgment and AI analytics become interdependent, which is particularly visible in financial decision-making contexts where trust dynamics become paramount (Kim & Park, 2024) and in educational assessment where ethical considerations emerge (Huang, 2025). This stage requires sophisticated transparency mechanisms and demonstrates the fragile nature of human-AI trust relationships.

Stage 3: Adaptive Coevolution At this stage, organizational culture, leadership styles, and AI capabilities mutually shape each other, requiring new forms of ethical stewardship and transparency (de Bruijn et al., 2022; Levasseur, 2025). Organizations reaching this stage report the highest innovation and satisfaction levels but face the greatest complexity in governance and ethical oversight. Each stage of the Human-AI Integration Continuum entails distinct leadership priorities. In Stage 1 (Operational Augmentation) leadership focus centers on change management and employee resilience-building to address initial resistance and competence anxiety (Gerlich, 2024). In Stage 2 (Decision Collaboration) leaders prioritize transparency mechanisms and cross-functional governance structures that maintain trust in AI-informed decisions (Kim & Park, 2024). By Stage 3 (Adaptive Coevolution) leadership emphasis shifts toward ethical stewardship and innovation culture, emphasizing continuous ethical review panels and the integration of human values into algorithmic updates (de Bruijn et al., 2022; Levasseur, 2025). This continuum framework addresses a gap in current literature by providing a developmental model that explains why identical AI technologies produce varied outcomes across organizations, depending on their integration maturity and human-centered design approaches.

4.2. Persistent Challenges and Complexities

AI's transformative impact is most apparent in operational efficiency, complex analytics, and personalized service delivery (Huang, 2025; Taherdoost & Madanchian, 2023). However, these advances are neither evenly distributed nor universally beneficial. Job displacement affects both low- and mid-skill roles, undermining security and prompting a race for reskilling (Gupta, 2025). Algorithmic bias persists, often reflecting entrenched social inequities, despite advances in explainable AI. This empirical evidence supports theoretical concerns raised by de Bruijn et al. (2022) about the limitations of current explainable AI approaches in addressing embedded social inequities (de Bruijn et al., 2022). Moreover, trust dynamics are multifaceted: individuals and organizations simultaneously value AI's objectivity and fear overreliance or loss of human expertise (Gerlich, 2024; Kim & Park, 2024).

Leadership and culture play crucial roles. Transparent, adaptive, and ethical leaders navigate AI transitions by setting standards, confronting bias, and maintaining focus on human dignity (Levasseur, 2025). Diverse leadership teams, encompassing members with varied cultural, gender, and disciplinary backgrounds,

are more likely to recognize bias patterns that homogeneous teams may overlook. As Grant and Levasseur (2025) demonstrated in the context of healthcare leadership, representation across identities enhances the capacity to challenge implicit assumptions and design equitable governance mechanisms. Applied to AI contexts, this diversity fosters more comprehensive algorithmic oversight and inclusive decision frameworks, thereby reducing the risk of encoding systemic inequities into AI systems. As prior research on minority women leaders in healthcare demonstrates, structural barriers and implicit biases create significant challenges in achieving equitable leadership representation (Grant & Levasseur, 2025). These same dynamics echo within AI adoption, where algorithmic bias and uneven access risk perpetuating inequities if not addressed by intentional, inclusive governance frameworks.

As observed in the primary author's consulting and research practice, organizations that cultivate human-AI complementarities report greater innovation, satisfaction, and resilience—mirroring findings in health, education, and corporate fields (Fiegler-Rudol et al., 2025; Taherdoost & Madanchian, 2023). The findings further reveal regulatory complexity: international conventions are advancing, but enforcement remains fragmented, leaving gaps in privacy, health, and equity protections (van Kolfschooten & Shachar, 2023).

5. Recommendations

1) Ethical Stewardship: Organizations must implement rigorous ethical oversight, prioritizing transparency, fairness, and accountability at every stage of the AI lifecycle (de Bruijn et al., 2022).

2) Skill Investment: Continuous education and upskilling programs are vital to offset deskilling and displacement, especially in fields most affected by automation (Gupta, 2025).

3) Inclusive Design: Diverse teams should lead AI design and governance, mitigating bias and aligning outputs with community values (de Bruijn et al., 2022).

4) Regulatory Engagement: Leaders should actively participate in shaping and evolving regulatory frameworks, advocating for clarity and equity (van Kolfschooten & Shachar, 2023).

5) Value-Driven Leadership: Leaders should promote a culture where AI augments human agency, encourages critical thinking, and aligns with organizational mission and wellbeing (Fiegler-Rudol et al., 2025; Levasseur, 2025).

6. Future Research Directions: Critical Research Gaps and Methodological Needs

Longitudinal Trust Evolution Studies: While current research captures snapshot views of human-AI trust (Gerlich, 2024; Kim & Park, 2024), longitudinal studies tracking trust evolution over 2 - 5 year periods are needed to understand how organizational AI relationships mature and what factors predict sustained collaboration versus abandonment.

Cross-Cultural Implementation Analysis: Current work on cross-cultural communication hints at cultural variation in AI adoption patterns, but systematic comparative studies across different cultural contexts remain limited (Taherdoost & Madanchian, 2023). Research examining how cultural values influence AI integration stages and success metrics would inform global deployment strategies.

Bias Intervention Effectiveness: Despite documented bias concerns (de Bruijn et al., 2022), empirical studies testing specific bias mitigation strategies remain scarce. Controlled trials comparing different fairness interventions, diverse team composition effects, and community-involvement approaches would provide evidence-based guidance for ethical AI implementation.

Leadership Transformation Metrics: While this review identifies leadership adaptation as crucial (Levasseur, 2025), quantitative measures of leadership effectiveness in AI-integrated organizations are underdeveloped.

Persistent leadership barriers faced by minority women in healthcare demonstrate the importance of intersectional analysis in organizational transformation (Grant & Levasseur, 2025). Future studies on AI adoption should explicitly examine how earlier findings on structural inequity and bias can guide development of fairness metrics, inclusive design principles, and targeted interventions that address both technological and social determinants of equity.

Researches developing and validating leadership competency frameworks specific to AI governance would support evidence-based leadership development.

Economic Impact Granularity: Gupta's (2025) analysis of market competitiveness provides sector-level insights, but micro-level studies examining how AI impacts specific job categories, skill premiums, and career progression pathways would inform more targeted reskilling initiatives.

Regulatory Effectiveness Assessment: van Kolschooten and Shachar's (2023) analysis of the Council of Europe's AI Convention represents important policy scholarship, but empirical studies measuring regulatory compliance costs, effectiveness in protecting vulnerable populations, and innovation impacts are needed to guide future policy development.

Methodological Innovation: Future researchers should prioritize mixed-methods approaches combining behavioral experiments (following Kim & Park, 2024), ethnographic organizational studies, and large-scale survey research to capture both quantitative patterns and qualitative nuances of human-AI integration processes.

7. Study Limitations

While this narrative review provides comprehensive coverage of current AI implementation literature, several limitations should be acknowledged. The methodology's focus on English-language publications may have excluded valuable insights from non-Western organizational contexts, potentially limiting the generalizability of the findings across different cultural settings. Additionally, the rapid pace of AI development means that some technological capabilities and organiza-

tional responses may have evolved beyond what current literature captures, highlighting the need for continuously updating research in this dynamic field.

8. Conclusion

By 2025, AI has become both indispensable and controversial, its promise entwined with profound dilemmas. The Human-AI Integration Continuum framework presented here demonstrates that successful AI implementation depends not merely on technological capabilities, but on organizational maturity in navigating the complex interplay between operational efficiency, trust dynamics, and ethical governance. The responsibility falls to organizational leaders, researchers, and practitioners to ensure that AI amplifies human potential while protecting dignity and equity. The article underscores that AI's trajectory depends not on its algorithms, but on the human choices, values, and ethics that shape its deployment. The expanded research agenda outlined above provides concrete pathways for developing the empirical foundation necessary to guide responsible AI evolution and ensure that it remains a tool for inclusive, sustainable progress.

Conflicts of Interest

The authors declare no conflicts of interest regarding the publication of this paper.

References

- de Bruijn, H., Warnier, M., & Janssen, M. (2022). The Perils and Pitfalls of Explainable AI: Strategies for Explaining Algorithmic Decision-Making. *Government Information Quarterly*, 39, Article ID: 101666. <https://doi.org/10.1016/j.giq.2021.101666>
- Fiegler-Rudol, J., Lau, K., Mroczek, A., & Kasperczyk, J. (2025). Exploring Human-AI Dynamics in Enhancing Workplace Health and Safety: A Narrative Review. *International Journal of Environmental Research and Public Health*, 22, Article No. 199. <https://doi.org/10.3390/ijerph22020199>
- Gerlich, M. (2024). Exploring Motivators for Trust in the Dichotomy of Human-AI Trust Dynamics. *Social Sciences*, 13, Article No. 251. <https://doi.org/10.3390/socsci13050251>
- Grant, V., & Levasseur, R. E. (2025). A Phenomenological Study of the Glass Ceiling Obstacles Faced by Minority Women Leaders in Healthcare Organizations. *Open Journal of Business and Management*, 13, 2569-2579. <https://doi.org/10.4236/ojbm.2025.134135>
- Gupta, D. R. (2025). How Does the Adoption of AI Impact Market Structure and Competitiveness within Industries? *Open Journal of Business and Management*, 13, 223-236. <https://doi.org/10.4236/ojbm.2025.131014>
- Huang, X. (2025). Reconceptualizing Formative Assessment in Business English: Opportunities, Challenges, and Ethical Considerations in the Age of Ai. *Open Journal of Social Sciences*, 13, 468-484. <https://doi.org/10.4236/jss.2025.136032>
- Kim, H. Y., & Park, Y. S. (2024). Trust Dynamics in Financial Decision Making: Behavioral Responses to AI and Human Expert Advice Following Structural Breaks. *Behavioral Sciences*, 14, Article No. 964. <https://doi.org/10.3390/bs14100964>
- Levasseur, R. E. (2025). Using AI, the Ultimate Black Box, Effectively—A Management Decision Making Perspective. *Open Journal of Business and Management*, 13, 2048-2057. <https://doi.org/10.4236/ojbm.2025.133107>

- Li, S., Fan, S., & Demartini, G. (2025). The Interaction between Emotion Dynamics and Opinion Changes in the Era of Generative AI. *Computers in Human Behavior Reports*, 19, Article ID: 100722. <https://doi.org/10.1016/j.chbr.2025.100722>
- McKinsey & Company (2025). *The State of AI*. <https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai>
- Nikolova, M., & Angrisani, M. (2025). The Impact of Learning about AI Advancements on Trust. *Technology in Society*, 83, Article ID: 102958. <https://doi.org/10.1016/j.techsoc.2025.102958>
- Singh, L. H., Charanarur, P., & Chaudhary, N. K. (2025). Advancements in Detecting Deep-fakes: AI Algorithms and Future Prospects—A Review. *Discover Internet of Things*, 5, Article No. 53. <https://doi.org/10.1007/s43926-025-00154-0>
- Stanford Human-Centered Artificial Intelligence (HAI) (2025). *2025 AI Index Report*. <https://hai.stanford.edu/ai-index/2025-ai-index-report>
- Suran, M., & Hswen, Y. (2024). How to Navigate the Pitfalls of AI Hype in Health Care. *JAMA*, 331, Article No. 273. <https://doi.org/10.1001/jama.2023.23330>
- Taherdoost, H., & Madanchian, M. (2023). AI Advancements: Comparison of Innovative Techniques. *AI*, 5, 38-54. <https://doi.org/10.3390/ai5010003>
- van Kolschooten, H., & Shachar, C. (2023). The Council of Europe's AI Convention (2023-2024): Promises and Pitfalls for Health Protection. *Health Policy*, 138, Article ID: 104935. <https://doi.org/10.1016/j.healthpol.2023.104935>
- Voelker, R. (2023). The Promise and Pitfalls of AI in the Complex World of Diagnosis, Treatment, and Disease Management. *JAMA*, 330, Article No. 1416. <https://doi.org/10.1001/jama.2023.19180>
- Yoo, S., Nimon, K., & Patole, S. R. (2024). Artificial Intelligence in Human Resource Development: An Umbrella Review Protocol. *PLOS ONE*, 19, e0310125. <https://doi.org/10.1371/journal.pone.0310125>