



International Journal of Artificial Intelligence and Machine Learning

Publisher's Home Page: <https://www.svedbergopen.com/>



Research Paper

Open Access

Human-Ai Collaboration in Management: Opportunities, Risks, and Best Practices

Dr. D. Paul Dhinakaran¹, Dr. S. Vijayalakshmi², Dr A Kalaivani³, Dr. V. David Raj⁴, Dr. Ramesh Kumar⁵, Dr Sundarapandiyan Natarajan⁶

¹Assistant Professor, Department of Commerce Jayagovind Harigopal Agarwal Agarsen College (Affiliated to University of Madras) Madhavaram, Chennai, Tamilnadu- 600060. E-mail: pauldhinakaranboss@gmail.com

²Assistant Professor (Aided) PG and Research Department of Commerce Arumugam Pillai Seethai Ammal College, Tirupattur, Sivagangai- 630211. E-Mail: vijibalasridhar@gmail.com

³Assistant Professor, Department of Computer Applications, Madanapalle Institute of Technology & Science (MITS), Deemed to be University, Madanapalle, Andrapradesh, India - 517325.

⁴Assistant Professor, Department of Commerce, St.Xavier's College (Autonomous). Palayamkottai Tirunelveli. E-mail: davidrajgam@gmail.com

⁵Associate Professor, Department of Commerce, PGDAV COLLEGE EVE, Delhi. E-mail: rameshdav2@gmail.com

⁶Professor and Head, Adithya School of Business Management Adithya Institute of Technology, Coimbatore. Orcid: 0000-0002-1303-2947, E-mail: nt_sundar@yahoo.com

Abstract

Artificial intelligence (AI) is increasingly embedded in managerial decision-making, from resource allocation and performance evaluation to strategic planning and customer relationship management. While AI promises efficiency, accuracy, and scale, its integration into management practice raises questions about trust, accountability, job redesign, and ethical risk. This study examines how perceived usefulness of AI, trust in AI systems, managerial AI literacy, perceived risk, and organizational support jointly shape human-AI collaborative performance. Using a cross-sectional survey of 410 practicing managers across five industry sectors, data were analyzed with SPSS 28 for descriptive and reliability statistics and SmartPLS 4 for partial least squares structural equation modeling (PLS-SEM). Results indicate that perceived usefulness ($\beta = 0.312$), trust in AI ($\beta = 0.284$), and AI literacy ($\beta = 0.198$) significantly and positively predict collaborative performance, while perceived risk exerts a significant negative effect ($\beta = -0.176$). Organizational support strongly predicts trust in AI ($\beta = 0.401$) and moderates the AI literacy-trust relationship. The model explains 47% of variance in collaborative performance. The study contributes an integrated framework linking technology acceptance and risk-perception literatures to human-AI teaming outcomes, and offers evidence-based best practices for managers and organizations seeking to design effective, responsible human-AI collaboration.

Keywords: Human-AI collaboration; artificial intelligence in management; trust in automation; AI literacy; PLS-SEM; responsible AI; organizational risk

1. Introduction

The rapid diffusion of artificial intelligence into organizational workflows has transformed AI from a back-office analytics tool into an active participant in managerial work. Generative AI assistants draft strategy documents, machine learning models forecast demand and flag fraud, and algorithmic systems increasingly inform hiring, promotion, and resource-allocation decisions once reserved exclusively for human judgment (Davenport & Ronanki, 2018; Wilson & Daugherty, 2018). This shift has given rise to a distinct organizational phenomenon: human-AI collaboration, in which managers and intelligent systems jointly perform tasks, share information, and co-produce decisions.

Human-AI collaboration is frequently framed as a source of competitive advantage, promising faster decision cycles, reduced cognitive load, and augmented analytical capability (Kolbjørnsrud, Amico, & Thomas, 2016). Yet the same literature documents significant risks: algorithmic opacity, over-reliance and automation complacency, erosion of managerial skill, bias amplification, and accountability gaps when AI-informed decisions produce adverse outcomes (Raisch & Krakowski, 2021). Despite growing interest, empirical evidence on the psychological and organizational antecedents of effective human-AI collaboration in management contexts remains fragmented.

This study addresses that gap by testing an integrated model of the drivers of human-AI collaborative performance among practicing managers. Specifically, the study examines how perceived usefulness, trust in AI, managerial AI literacy, perceived risk, and organizational support relate to collaborative performance outcomes, drawing on a sample of 410 managers and analyzed using advanced structural equation modeling techniques. The findings are intended to inform both theory and managerial practice by identifying which levers organizations can pull to capture the opportunities of AI while containing its risks.

1.1 Objectives of the Study

- To examine the influence of perceived usefulness, trust in AI, and managerial AI literacy on human-AI collaborative performance.
- To assess the effect of perceived risk on collaborative performance in AI-augmented management settings.
- To test the role of organizational support in building trust in AI and strengthening the AI literacy–trust relationship.
- To derive evidence-based best practices for managing human-AI collaboration responsibly.

2. Literature Review and Hypotheses Development

2.1 Opportunities of Human-AI Collaboration in Management

Prior research identifies three broad categories of opportunity created by AI in management: augmentation of analytical capacity, acceleration of decision cycles, and personalization at scale. AI systems can process volumes of structured and unstructured data beyond human capacity, surfacing patterns that inform pricing, staffing, and risk decisions (Brynjolfsson & McAfee, 2017). In collaborative arrangements, AI typically performs data-intensive, repetitive sub-tasks while managers retain judgment-intensive and interpersonal responsibilities, a division of labor associated with improved decision quality and efficiency (Wilson & Daugherty, 2018).

2.2 Risks and Challenges

Alongside these opportunities, scholars caution against over-automation and the erosion of human oversight. Raisch and Krakowski (2021) describe an 'automation-augmentation paradox' in which organizations that automate too aggressively lose the complementary human capabilities needed to supervise and correct AI systems. Additional risks include algorithmic bias that can systematically disadvantage protected groups, reduced transparency in decision rationale, skill atrophy among managers who defer excessively to algorithmic recommendations, and ambiguous accountability when AI-informed decisions cause harm (Shrestha, Ben-Menahem, & von Krogh, 2019).

2.3 Theoretical Framework

The research model integrates the Technology Acceptance Model (Davis, 1989), trust-in-automation theory (Lee & See, 2004), and organizational support theory (Eisenberger et al., 1986). Perceived usefulness and trust are theorized as proximal drivers of collaborative performance; AI literacy is modeled as a competence-based antecedent of both trust and performance; perceived risk is modeled as a countervailing force; and organizational support is positioned as a contextual resource that strengthens trust and amplifies the benefits of AI literacy.

2.4 Hypotheses

- **H1:** Perceived usefulness of AI positively influences human-AI collaborative performance.
- **H2:** Trust in AI systems positively influences human-AI collaborative performance.
- **H3:** Managerial AI literacy positively influences human-AI collaborative performance.
- **H4:** Perceived risk of AI use negatively influences human-AI collaborative performance.
- **H5:** Organizational support positively influences trust in AI systems.
- **H6:** Organizational support moderates the relationship between AI literacy and trust in AI, such that the relationship is stronger under high organizational support.

3. Research Methodology

3.1 Research Design

A quantitative, cross-sectional survey design was adopted to test the hypothesized relationships. This design is appropriate for theory testing across a large sample and enables the use of covariance- and variance-based structural equation modeling techniques to examine multiple interrelated hypotheses simultaneously.

3.2 Sample and Data Collection

Data were collected from 410 practicing managers (junior, middle, and senior levels) working in organizations that had implemented at least one AI-based tool in managerial workflows (e.g., predictive analytics dashboards, AI-assisted scheduling, generative AI writing assistants, or automated decision-support systems). Respondents were drawn from IT and software services, banking and financial services, manufacturing, healthcare, and retail/other services sectors using a combination of professional-network sampling and organizational gatekeeper referrals. Of 512 questionnaires distributed, 431 were returned (response rate 84.2%), and after removing incomplete and straight-lined responses, 410 usable responses remained, corresponding to a valid response rate of 80.1%. A post-hoc power analysis (G*Power) confirmed the sample was more than adequate to detect medium effect sizes at $\alpha = 0.05$ with power exceeding 0.99 for the structural model.

3.3 Measures

All constructs were measured with previously validated multi-item scales adapted to the AI-in-management context and rated on 5-point Likert scales (1 = strongly disagree, 5 = strongly agree). Perceived usefulness of AI (5 items) was adapted from Davis (1989); trust in AI systems (5 items) from Lee and See (2004); managerial AI literacy (4 items) was developed from recent AI-literacy scales; perceived risk of AI use (5 items) drew on perceived-risk research in technology adoption; organizational support (4 items) was adapted from Eisenberger et al. (1986); and human-AI collaborative performance (6 items) captured task efficiency, decision quality, and satisfaction with the human-AI division of labor. A pilot test with 40 managers preceded full data collection to refine item wording.

3.4 Analytical Tools

Data screening, descriptive statistics, and reliability analysis were conducted in IBM SPSS Statistics 28. The measurement and structural models were estimated using partial least squares structural equation modeling (PLS-SEM) in SmartPLS 4, selected for its suitability with prediction-oriented models, moderate sample sizes, and formative-reflective mixed constructs. Common method bias was assessed using Harman's single-factor test and the full collinearity VIF approach; all VIF values were below 3.3, indicating common method bias was not a serious concern. Bootstrapping with 5,000 resamples was used to estimate the significance of path coefficients.

4. Results

4.1 Sample Profile

Table 1 presents the demographic and organizational profile of respondents.

Table 1. Demographic Profile of Respondents (N = 410)

Characteristic	Category	Frequency (n)	Percent (%)
Gender	Male	228	55.6
	Female	182	44.4
Age	25–34 years	129	31.5
	35–44 years	158	38.5
	45–54 years	89	21.7
	55 years and above	34	8.3
Managerial level	Junior/first-line	103	25.1
	Middle management	201	49.0
	Senior/top management	106	25.9
Industry sector	IT & software services	112	27.3
	Banking & financial services	84	20.5

	Manufacturing	71	17.3
	Healthcare	58	14.1
	Retail & other services	85	20.7
Experience using AI-based tools	Less than 1 year	96	23.4
	1–3 years	187	45.6
	More than 3 years	127	31.0

4.2 Measurement Model Assessment

Internal consistency reliability, convergent validity, and discriminant validity were assessed prior to structural model estimation. As shown in Table 2, all constructs exceeded the recommended thresholds of 0.70 for Cronbach's alpha and composite reliability, and 0.50 for average variance extracted (AVE), supporting the reliability and convergent validity of the measurement model. Discriminant validity was confirmed using the Heterotrait-Monotrait (HTMT) ratio, with all values below the conservative 0.85 threshold.

Table 2. Reliability and Convergent Validity

Construct	Items	Cronbach's α	Composite Reliability	AVE
Perceived Usefulness of AI (PUA)	5	0.86	0.90	0.65
Trust in AI Systems (TAI)	5	0.88	0.91	0.68
Managerial AI Literacy (MAL)	4	0.83	0.88	0.64
Perceived Risk of AI Use (PRA)	5	0.85	0.89	0.62
Organizational Support (OSP)	4	0.82	0.87	0.63
Human–AI Collaborative Performance (HCP)	6	0.90	0.93	0.69

4.3 Model Fit and Explanatory Power

Table 3. Model Fit and Explanatory Power Indices

Fit Index	Obtained Value	Threshold
SRMR	0.052	< 0.08
NFI	0.91	> 0.90
R^2 (HCP)	0.47	Moderate–substantial
R^2 (TAI)	0.38	Moderate
Q^2 (predictive relevance, HCP)	0.31	> 0

The structural model explained 47% of the variance in human-AI collaborative performance ($R^2 = 0.47$) and 38% of the variance in trust in AI ($R^2 = 0.38$), both representing moderate-to-substantial explanatory power in behavioral research. The Q^2 statistic (0.31) confirmed the model's predictive relevance.

4.4 Hypothesis Testing

Table 3 summarizes the path coefficients, t-values, and hypothesis-testing outcomes obtained through bootstrapping.

Table 4. Structural Model Results and Hypothesis Testing

Hypothesis	Path	β	t-value	Decision
H1	PUA $\rightarrow$ HCP	0.312	6.84	Supported
H2	TAI $\rightarrow$ HCP	0.284	5.97	Supported
H3	MAL $\rightarrow$ HCP	0.198	4.21	Supported
H4	PRA $\rightarrow$ HCP	-0.176	3.65	Supported
H5	OSP $\rightarrow$ TAI	0.401	8.12	Supported
H6	MAL $\rightarrow$ TAI (moderation by OSP)	0.147	2.89	Supported

All six hypotheses were supported. Perceived usefulness emerged as the strongest positive predictor of collaborative performance ($\beta = 0.312$, $t = 6.84$, $p < 0.001$), followed by trust in AI ($\beta = 0.284$, $t = 5.97$, $p < 0.001$) and AI literacy ($\beta = 0.198$, $t = 4.21$, $p < 0.001$). Perceived risk significantly reduced collaborative performance ($\beta = -0.176$, $t = 3.65$, $p < 0.001$). Organizational support was a strong predictor of trust in AI ($\beta = 0.401$, $t = 8.12$, $p < 0.001$) and significantly moderated the AI literacy–trust relationship ($\beta = 0.147$, $t = 2.89$, $p < 0.01$), such that literacy translated into trust more strongly when organizational support was high.

5. Discussion

The findings affirm that human-AI collaboration in management is not a purely technical outcome but a socio-technical one, shaped jointly by perceptions of usefulness, trust, competence, and organizational context. Consistent with technology acceptance research, perceived usefulness remained the dominant driver of collaborative performance, suggesting that managers who see a clear line between AI outputs and better decisions are more willing to integrate AI meaningfully into their workflows rather than treating it as a peripheral tool.

The significant negative effect of perceived risk reinforces automation-augmentation paradox theorizing (Raisch & Krakowski, 2021): even when managers find AI useful, concerns about bias, opacity, and accountability can suppress effective collaboration. The strong path from organizational support to trust, and its moderating role in the literacy–trust relationship, indicates that individual competence alone is insufficient. Organizations that invest in AI training without also building governance structures, escalation paths, and psychological safety for questioning AI outputs may see diminished returns on that training.

6. Key Risks in Human-AI Management Collaboration

- Automation complacency: excessive deference to AI recommendations, reducing critical scrutiny of edge cases and errors.
- Algorithmic bias: AI systems trained on historical data can reproduce or amplify discriminatory patterns in hiring, promotion, or performance evaluation.
- Accountability ambiguity: unclear ownership of decisions when outcomes are jointly produced by humans and algorithms.
- Skill erosion: prolonged reliance on AI-generated analysis may erode managers' independent analytical and judgment capabilities.
- Data privacy and security exposure: AI tools that ingest sensitive organizational or customer data increase the attack surface for breaches.
- Transparency and explainability gaps: 'black box' models can undermine managers' ability to justify decisions to stakeholders and regulators.

7. Best Practices for Managing Human-AI Collaboration

- Establish clear task allocation: assign AI to data-intensive, repetitive sub-tasks and reserve judgment-intensive, interpersonal, and ethically sensitive decisions for humans.
- Invest in structured AI-literacy training so managers understand both the capabilities and the limitations of the tools they use.
- Build organizational support structures — governance committees, escalation channels, and psychological safety — that make it acceptable to question or override AI recommendations.

- Require explainability and audit trails for any AI system used in consequential decisions (hiring, promotion, credit, resource allocation).
- Implement periodic bias audits and diverse review panels for AI systems influencing people-related decisions.
- Define clear accountability protocols that specify who is responsible for outcomes of human-AI co-produced decisions.
- Pilot new AI tools in low-stakes contexts before scaling to high-stakes managerial decisions, and monitor performance and risk indicators continuously.

8. Limitations and Future Research

This study relies on cross-sectional, self-reported data, which limits causal inference and may be subject to social desirability bias despite statistical checks for common method variance. The sample, while adequately powered, was drawn primarily from five industry sectors and may not generalize to public-sector or non-profit management contexts. Future research could employ longitudinal designs to capture how trust and collaborative performance evolve as AI tools mature, incorporate objective performance indicators alongside self-reports, and extend the model to examine team-level rather than purely individual-level human-AI collaboration.

9. Conclusion

Human-AI collaboration offers management significant opportunities for augmented decision-making and efficiency, but realizing these benefits depends on a configuration of individual and organizational factors: usefulness perceptions, trust, competence, and support structures that jointly counteract perceived risk. Drawing on survey data from 410 managers analyzed through PLS-SEM, this study offers empirical evidence that organizational investment in both AI capability-building and governance is essential to translating AI adoption into genuinely collaborative, high-performing human-AI work systems. Organizations that treat human-AI collaboration as a socio-technical design challenge — rather than a purely technical deployment — are best positioned to capture its opportunities while managing its risks responsibly.

References

1. Brynjolfsson, E., & McAfee, A. (2017). The business of artificial intelligence. *Harvard Business Review*, 95(4), 3–11.
2. Davenport, T. H., & Ronanki, R. (2018). Artificial intelligence for the real world. *Harvard Business Review*, 96(1), 108–116.
3. Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340.
4. Eisenberger, R., Huntington, R., Hutchison, S., & Sowa, D. (1986). Perceived organizational support. *Journal of Applied Psychology*, 71(3), 500–507.
5. Fornell, C., & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), 39–50.
6. Hair, J. F., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2022). A primer on partial least squares structural equation modeling (PLS-SEM) (3rd ed.). Sage.
7. Kolbjørnsrud, V., Amico, R., & Thomas, R. J. (2016). How artificial intelligence will redefine management. *Harvard Business Review*, 2, 1–6.
8. Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50–80.
9. Raisch, S., & Krakowski, S. (2021). Artificial intelligence and management: The automation-augmentation paradox. *Academy of Management Review*, 46(1), 192–210.
10. Shrestha, Y. R., Ben-Menahem, S. M., & von Krogh, G. (2019). Organizational decision-making structures in the age of artificial intelligence. *California Management Review*, 61(4), 66–83.
11. Wilson, H. J., & Daugherty, P. R. (2018). Collaborative intelligence: Humans and AI are joining forces. *Harvard Business Review*, 96(4), 114–123.