



Human–AI Collaboration and the Redesign of Managerial Decision Processes

Dr. Daniel R. Weber

Professor, University of Mannheim, Germany

Abstract

Artificial intelligence is changing managerial decision-making from a predominantly human-centered activity into a distributed cognitive process in which managers increasingly interact with predictive models, recommendation systems, analytics platforms, and generative artificial intelligence. The strategic challenge is therefore no longer limited to whether organizations should adopt artificial intelligence, but concerns how decision authority, analytical responsibility, human judgment, verification, and accountability should be redesigned when humans and intelligent systems participate in the same decision process. This study examines human–AI collaboration as a managerial decision architecture rather than as a straightforward automation strategy.

A structured integrative review was undertaken using research on organizational decision-making, human–AI complementarity, algorithmic reliance, managerial AI adoption, human-centered artificial intelligence, generative AI productivity, and responsible AI governance. The synthesis indicates that effective collaboration depends on matching AI and human capabilities to specific decision stages. Artificial intelligence performs particularly valuable functions in large-scale information processing, pattern identification, scenario generation, anomaly detection, and rapid comparison of alternatives, whereas managerial judgment remains especially important when decisions involve ambiguity, ethical consequences, organizational politics, tacit knowledge, stakeholder interpretation, novel conditions, and accountability. Recent meta-analytic evidence further demonstrates that combining humans and AI does not automatically generate performance superior to the better individual agent, making collaboration design and calibrated reliance critical.

The paper consequently proposes a five-stage Human–AI Managerial Decision Redesign Framework covering problem framing, evidence augmentation, alternative evaluation, accountable judgment, and post-decision learning. An illustrative simulation is additionally used to demonstrate how decision-quality observations may be distributed in a well-designed collaborative environment; these simulated observations are methodological illustrations rather than empirical organizational findings. The study concludes that sustainable managerial advantage emerges not from maximizing automation but from allocating cognitive work deliberately, preserving meaningful human oversight, developing AI literacy, maintaining contestability, and creating feedback mechanisms through which both managerial practice and AI-supported workflows can improve over time.



Keywords: Human–AI Collaboration; Managerial Decision-Making; Artificial Intelligence; Decision Augmentation; Human Oversight; Algorithmic Reliance; AI Governance; Managerial Judgment; Hybrid Intelligence; Generative AI

1. Introduction

Artificial intelligence has moved from specialized analytical applications into the everyday infrastructure of organizational decision-making. Managers can now use AI-enabled systems to summarize large information environments, forecast possible outcomes, identify correlations, classify risks, generate strategic alternatives, support resource allocation, and produce preliminary recommendations. Evidence from knowledge-work environments demonstrates that generative AI can materially affect productivity, although the magnitude and direction of the effect vary with task characteristics and worker circumstances. Brynjolfsson, Li, and Raymond, for example, reported an average productivity increase of approximately 15% following the introduction of generative-AI assistance among 5,172 customer-support agents. Earlier experimental evidence from Noy and Zhang similarly demonstrated substantial productivity improvements in professional writing tasks.

These technological capabilities do not imply that managerial decision-making can simply be transferred from people to algorithms. Jarrahi's influential account of human–AI symbiosis argued that humans and AI possess complementary capabilities in organizational decision environments characterized by uncertainty and complexity. Raisch and Krakowski subsequently conceptualized AI adoption through an automation–augmentation paradox, emphasizing that substitution and augmentation cannot always be treated as independent alternatives because organizational applications frequently contain elements of both. This theoretical transition changes the central managerial question. Instead of asking whether AI should replace a manager, organizations must determine which elements of a decision process should be automated, which should be augmented, and where human responsibility must remain explicit.

The importance of this distinction has increased with generative AI because contemporary systems can produce persuasive outputs even when underlying reasoning or factual content is deficient. Dell'Acqua and colleagues' field experiment with management consultants demonstrated what the authors describe as a “jagged technological frontier”: AI assistance can improve performance for tasks within current AI capabilities while reducing performance when employees rely on AI for tasks beyond those capabilities. Their results consequently highlight the difficulty of determining in advance exactly where AI assistance will help and where stronger human scrutiny is required.

Human–AI collaboration also introduces behavioral and organizational problems that cannot be resolved through model accuracy alone. Managers may reject useful AI recommendations because of algorithm aversion, accept inaccurate recommendations because of automation bias, or treat AI-generated analysis as objective even when its assumptions require challenge. Haesevoets and colleagues found that managers expressed preferences concerning the distribution of human and machine input in managerial decisions, while Cao and colleagues demonstrated that managerial intentions to use AI are shaped simultaneously by perceived benefits and perceived threats. Human–AI decision-making therefore involves psychological acceptance, role clarity, accountability, organizational design, and cognitive calibration in addition to technological capability.



The present study addresses these issues by examining how managerial decision processes should be redesigned around structured human–AI collaboration rather than unrestricted automation. The study develops a five-stage model in which AI contributes analytical scale and computational assistance while managers retain responsibility for contextual interpretation, ethical judgment, escalation, final accountability, and organizational learning.

This approach is consistent with contemporary evidence showing that human–AI combinations do not automatically outperform the strongest individual decision-maker and that complementary performance depends on appropriate task design and reliance. Vaccaro, Almaatouq, and Malone's systematic review and meta-analysis found considerable heterogeneity in human–AI performance and cautioned against assuming automatic synergy.

2. Objectives of the Study

2.1 To Examine the Changing Distribution of Decision Authority

The first objective is to examine how artificial intelligence changes the distribution of analytical work, judgment, recommendation, verification, and decision authority within managerial processes. The study distinguishes between AI-supported analysis and managerial accountability and identifies decision stages where computational assistance can add value without eliminating meaningful human control.

2.2 To Develop a Human–AI Decision Redesign Framework

The second objective is to construct an integrated framework explaining how managers and AI systems can collaborate across problem framing, information processing, option generation, evaluation, final judgment, implementation, and organizational learning. Particular attention is given to complementarity rather than simple task substitution.

2.3 To Identify Conditions for Reliable and Responsible Collaboration

The third objective is to identify organizational conditions that improve human–AI decision quality, including AI literacy, transparent role allocation, verification requirements, appropriate reliance, escalation mechanisms, accountability, decision documentation, performance monitoring, and continuous feedback.

3. Review of Literature

3.1 From Automated Decisions to Hybrid Managerial Intelligence

Earlier organizational applications of artificial intelligence frequently emphasized automation: systems performed relatively structured tasks according to predetermined rules, predictive models, or algorithmic classifications. Contemporary managerial AI increasingly operates differently because algorithms can now contribute to semi-structured and unstructured knowledge work.

Shrestha, Ben-Menahem, and von Krogh argued that artificial intelligence changes organizational decision structures because human and AI decision-makers differ in the characteristics of their information processing and decision capabilities. Duan, Edwards, and Dwivedi similarly identified AI-supported organizational decision-making as an emerging research domain requiring greater understanding of interactions between computational intelligence and managerial judgment.

Jarrahi's symbiosis perspective is particularly relevant because it avoids the assumption that one agent must dominate the other. Human judgment contributes contextual knowledge, intuition, social interpretation, ethical reasoning, and adaptive responses to exceptional circumstances, whereas AI contributes computational scale, pattern detection, analytical consistency, and rapid processing of large datasets.

Hillebrand, Raisch, and Schad's more recent integrative framework further demonstrates that managing with AI should be understood as an organizational phenomenon involving changing relationships between human and artificial agency rather than as the deployment of an isolated technical instrument.

3.2 Complementarity, Reliance, and the Limits of AI Assistance

A central assumption surrounding human–AI collaboration is that combining human and machine intelligence will generate superior performance. The empirical literature presents a more conditional picture.

Vaccaro, Almaatouq, and Malone conducted a systematic review and meta-analysis covering more than 100 experimental studies and over 300 effect sizes. Their analysis found that human–AI combinations did not, on average, consistently outperform the better of humans or AI operating independently. Decision tasks were particularly challenging for generating genuine synergy, whereas stronger gains appeared in some content-creation contexts.

This finding has major managerial implications. Collaboration itself is not a performance mechanism. Its value depends on whether responsibilities are allocated according to complementary capabilities.

Dell'Acqua and colleagues reached a related conclusion through a field experiment involving knowledge workers. AI enhanced performance where tasks aligned with its capabilities but could reduce performance when users relied on it outside the effective capability frontier.

3.3 Human-Centered AI and Decision-Process Redesign

The next stage of research increasingly focuses on designing the interaction itself. Krakowski and colleagues' human-centered AI field experiment found substantial differences between tailored and untailored human–AI interactions during unstructured decision tasks. Their findings indicate that work procedure and the alignment of AI interaction with human information-processing requirements can materially influence outcomes.

Li and Tian's 2026 systematic review and bibliometric analysis of 627 publications similarly concluded that AI–human collaboration varies according to both the relationship between human and artificial agents and the type of decision being performed. Their framework identifies different decision paradigms rather than treating every AI-assisted decision as structurally equivalent.

The research gap therefore lies not primarily in demonstrating that AI can process information. Considerable evidence already establishes this capacity. The more important unresolved managerial problem concerns how complete decision workflows should be redesigned so that AI contribution increases analytical capability without weakening human reasoning, accountability, organizational knowledge, or the ability to detect algorithmic failure.



4. Research Methodology

4.1 Research Design

The study adopts a structured integrative review combined with conceptual framework development and an illustrative simulation. This design is appropriate because the research question concerns the organization of managerial decision processes rather than the performance of a single proprietary AI model.

The study does not claim to report a newly administered employee or manager survey. Consequently, no fabricated respondents, organizations, regression coefficients, or significance tests are presented.

4.2 Literature Identification and Selection

The analytical literature was identified from peer-reviewed research concerning human–AI collaboration, organizational AI, managerial decision-making, algorithmic reliance, generative-AI productivity, automation, augmentation, hybrid intelligence, and AI governance.

The core literature covers the period from 2018 through August 2026, while selected earlier foundational decision and algorithm studies were considered where necessary. Priority was given to empirical studies, systematic reviews, meta-analyses, field experiments, major conceptual articles, and recognized governance frameworks.

Table 1: Research Design and Evidence-Synthesis Protocol

Methodological Component	Study Specification
Research orientation	Conceptual-analytical
Main method	Structured integrative literature review
Core period	2018–August 2026
Evidence base	Peer-reviewed empirical, review and conceptual research
Principal themes	AI augmentation, managerial judgment, algorithmic reliance, decision authority, accountability and human oversight
Analytical unit	Managerial decision-process stage
Synthesis method	Thematic and comparative synthesis
Supplementary analysis	Reproducible illustrative simulation
Simulation purpose	Visual demonstration of a possible decision-quality distribution



Methodological Component	Study Specification
Empirical status of simulation	Synthetic; not a survey or organizational field dataset
Ethical status	No human participants or personally identifiable information collected

4.3 Analytical Procedure

The selected evidence was organized around five questions: what cognitive activity is being performed; whether the task is structured or ambiguous; what information is available to the AI; what contextual or ethical knowledge remains outside the model; and who remains accountable for the final decision.

The literature was subsequently synthesized into five managerial decision stages:

Problem framing → AI-supported evidence augmentation → alternative evaluation → accountable managerial judgment → feedback and organizational learning.

The framework therefore considers collaboration as a sequence rather than as a single moment in which a manager simply accepts or rejects an algorithmic recommendation.

4.4 Illustrative Simulation and Ethical Position

To demonstrate how decision-quality observations could be inspected in a collaborative decision environment, 60 synthetic decision-quality scores were generated on a 0–100 scale using a reproducible random process centered around a comparatively high-performance collaborative condition.

The simulated dataset produced a mean decision-quality score of 82.46, a standard deviation of 5.50, a minimum of 68.34, and a maximum of 96.99.

These numbers are illustrative analytical values only. They must not be interpreted as responses from managers, experimental participants, companies, or a population estimate.

5. Analysis

5.1 Redesigning the Managerial Decision Sequence

The evidence suggests that managers should not ask AI to “make the decision” as a single undifferentiated activity. Instead, the complete decision process should be decomposed into stages.

During problem framing, managerial leadership should remain dominant because defining the decision problem requires organizational purpose, stakeholder interpretation, tacit knowledge, and judgment concerning what should actually be optimized.

During evidence augmentation, AI can assume a stronger analytical role. It may retrieve relevant information, summarize large datasets, detect patterns, produce forecasts, identify anomalies, compare cases, and generate alternative interpretations.

During alternative evaluation, both human and AI contributions become important. AI can model scenarios and quantify trade-offs, while managers assess strategic feasibility, stakeholder effects, organizational constraints, and assumptions that may not be represented adequately in the model.

The final judgment stage should preserve explicitly assigned human accountability for consequential managerial decisions. AI may recommend, challenge, or simulate, but managerial responsibility should not disappear merely because an algorithm contributed to the analysis.

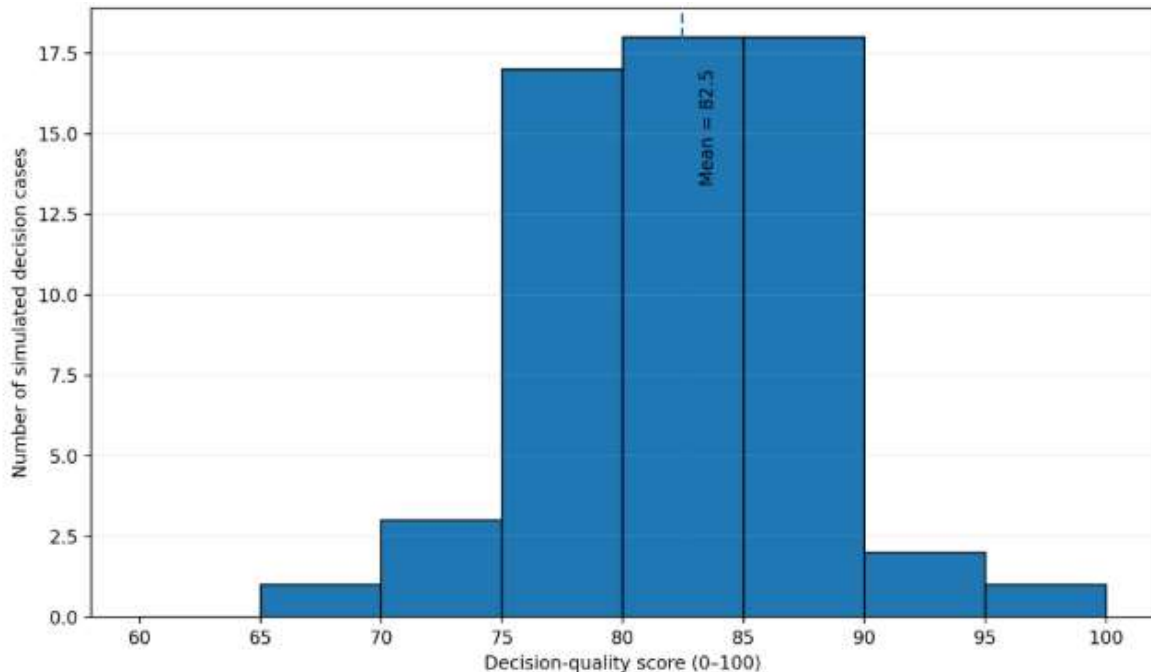
Finally, the learning stage should capture the decision, AI recommendation, human modifications, outcome, failure signals, and subsequent lessons. This creates organizational memory and enables future calibration.

Table 2: Human–AI Managerial Decision Redesign Matrix

Decision Stage	Primary Human Contribution	Primary AI Contribution	Required Control	Intended Outcome
Problem framing	Define objectives, boundaries and stakeholder priorities	Retrieve contextual evidence and identify related patterns	Human ownership of problem definition	Correct decision target
Evidence augmentation	Assess relevance and missing context	Process datasets, summarize evidence and detect patterns	Source and output verification	Improved information coverage
Alternative generation	Introduce strategic and experiential alternatives	Generate scenarios and alternative courses of action	Diversity and feasibility review	Expanded option set
Evaluation and challenge	Interpret trade-offs, ethics and organizational consequences	Quantify scenarios, compare alternatives and detect inconsistencies	Independent human challenge	Better calibrated judgment
Final decision	Exercise authority and accountability	Provide recommendation and uncertainty information	Named decision owner	Accountable choice

5.2 Histogram of Illustrative Collaborative Decision Quality

Figure 1: Illustrative Distribution of Human–AI Collaborative Decision Quality



The histogram represents 60 simulated decision cases, not real managers or organizations. The distribution is concentrated primarily within the upper-middle decision-quality range, with a simulated mean of approximately 82.46.

The important analytical observation is not the numerical mean itself. The distribution illustrates why managers should examine variation across decisions rather than relying exclusively on an average performance indicator. Even when average collaborative performance appears strong, individual cases may fall substantially below the central range.

This is particularly important in managerial AI because a system that performs well across routine cases may still fail in unusual strategic, ethical, or context-sensitive situations. Dell’Acqua and colleagues’ empirical evidence concerning the jagged technological frontier reinforces precisely this concern: similar-looking knowledge tasks may differ substantially in their suitability for AI assistance.

6. Discussion

6.1 Human–AI Collaboration Should Not Be Confused with Maximum Automation

The principal finding of the study is that organizations should distinguish AI intensity from decision quality. Increasing the number of managerial activities performed by AI does not necessarily improve organizational decision-making.

Raisch and Krakowski’s automation–augmentation paradox is particularly relevant because automation and augmentation can interact dynamically rather than representing fixed alternatives. The appropriate objective is therefore not maximum automation but optimal cognitive allocation.



Structured, repetitive and data-intensive activities can often be assigned more heavily to AI. Decisions involving ambiguity, political consequences, ethical obligations, novel strategic conditions, interpersonal meaning, or incomplete information require stronger managerial engagement.

6.2 The Central Managerial Capability Becomes Reliance Calibration

Traditional management development emphasizes analysis, leadership, communication, negotiation, and strategic judgment. AI-supported management adds another capability: knowing when to rely on AI, how much to rely on it, and when to challenge it.

This is more demanding than simply learning how to operate an AI tool.

Vaccaro and colleagues' meta-analysis provides an important warning because human–AI combinations did not systematically outperform the strongest independent agent across all decision contexts. Effective collaboration therefore requires managers to understand both their own limitations and the limitations of artificial systems.

Managers who reject every algorithmic recommendation may lose useful analytical value. Managers who accept recommendations automatically may reproduce model error at organizational scale.

The desirable state lies between these extremes: appropriate reliance.

6.3 Implications for Organizational Structure and Governance

Human–AI collaboration also changes organizational governance. Decision policies should clarify which AI systems may advise on particular decisions, which evidence may be entered into those systems, what verification standards apply, who can override recommendations, when escalation is required, and who is accountable for resulting outcomes.

NIST's AI Risk Management Framework and its Generative AI Profile similarly emphasize organizational risk management across the AI lifecycle rather than treating trustworthy AI as a model-development issue alone.

Decision documentation should consequently record the human–AI interaction itself. Organizations should be able to reconstruct important decisions by identifying what information was provided to AI, what recommendation was generated, what uncertainty existed, whether managers challenged the recommendation, what modification occurred, and who authorized the final action.

The introduction of AI should also avoid managerial deskilling. When managers repeatedly delegate evidence evaluation, alternative generation, reasoning, and drafting to AI, they may gradually lose practice in capabilities required when the system fails. Human oversight has little practical value when the nominal overseer no longer possesses sufficient expertise to identify an error.

7. Conclusion

Human–AI collaboration is redesigning managerial decision-making by redistributing cognitive work between managerial professionals and artificial intelligence systems. The most effective organizational response is not to treat artificial intelligence as either an autonomous substitute for managers or a passive analytical tool. Instead, AI should be incorporated within deliberately designed decision architectures

that exploit computational strengths while preserving human responsibility for contextual interpretation, ethical reasoning, stakeholder judgment, uncertainty management, and accountable choice.

The study identifies five essential components of this redesigned architecture: human-led problem framing, AI-enabled evidence augmentation, collaborative evaluation, accountable managerial judgment, and post-decision learning. Together, these components shift organizational attention from the question of whether AI should participate in management toward the more sophisticated question of how participation should be structured.

Existing empirical evidence demonstrates why such caution is necessary. Generative AI can increase productivity substantially in suitable knowledge tasks, yet the performance frontier remains uneven. Human–AI combinations likewise do not automatically outperform the strongest human or artificial agent. Collaboration quality therefore depends on task characteristics, human expertise, AI capability, interface design, organizational controls, and appropriate reliance.

The proposed framework contributes to managerial research by conceptualizing human–AI collaboration as an end-to-end organizational decision process rather than a single recommendation event. Its principal limitation is that the present paper synthesizes existing evidence and uses a clearly identified illustrative simulation rather than original organizational field data. Future research should consequently test the framework through longitudinal studies, controlled managerial experiments, cross-industry datasets, and real-world decision logs.

The strategic objective for organizations should therefore be neither human-only management nor AI-dominant management, but accountable hybrid intelligence in which the strengths and limitations of both decision agents are deliberately managed.

References

1. M. H. Jarrahi, “Artificial intelligence and the future of work: Human-AI symbiosis in organizational decision making,” *Business Horizons*, vol. 61, no. 4, pp. 577–586, 2018, doi: 10.1016/j.bushor.2018.03.007.
2. S. Raisch and S. Krakowski, “Artificial intelligence and management: The automation–augmentation paradox,” *Academy of Management Review*, vol. 46, no. 1, 2021, doi: 10.5465/amr.2018.0072.
3. Y. R. Shrestha, S. M. Ben-Menahem, and G. von Krogh, “Organizational decision-making structures in the age of artificial intelligence,” *California Management Review*, vol. 61, no. 4, pp. 66–83, 2019, doi: 1177/0008125619862257.
4. Y. Duan, J. S. Edwards, and Y. K. Dwivedi, “Artificial intelligence for decision making in the era of Big Data—Evolution, challenges and research agenda,” *International Journal of Information Management*, vol. 48, pp. 63–71, 2019, doi: 10.1016/j.ijinfomgt.2019.01.021.
5. E. E. Makarius, D. Mukherjee, J. D. Fox, and A. K. Fox, “Rising with the machines: A sociotechnical framework for bringing artificial intelligence into the organization,” *Journal of Business Research*, vol. 120, pp. 262–273, 2020, doi: 10.1016/j.jbusres.2020.07.045.
6. I. Seeber *et al.*, “Machines as teammates: A research agenda on AI in team collaboration,” *Information & Management*, vol. 57, no. 2, Art. no. 103174, 2020, doi: 10.1016/j.im.2019.103174.



6. T. Haesevoets, D. De Cremer, K. Dierckx, and A. Van Hiel, “Human-machine collaboration in managerial decision making,” *Computers in Human Behavior*, vol. 119, Art. no. 106730, 2021, doi: 10.1016/j.chb.2021.106730.
7. G. Cao, Y. Duan, J. S. Edwards, and Y. K. Dwivedi, “Understanding managers’ attitudes and behavioral intentions towards using artificial intelligence for organizational decision-making,” *Technovation*, vol. 106, Art. no. 102312, 2021, doi: 10.1016/j.technovation.2021.102312.
8. S. Noy and W. Zhang, “Experimental evidence on the productivity effects of generative artificial intelligence,” *Science*, vol. 381, pp. 187–192, 2023, doi: 10.1126/science.adh2586.
9. [M. Vaccaro, A. Almaatouq, and T. Malone, “When combinations of humans and AI are useful: A systematic review and meta-analysis,” *Nature Human Behaviour*, vol. 8, pp. 2293–2303, 2024, doi: 10.1038/s41562-024-02024-1.
10. E. Brynjolfsson, D. Li, and L. Raymond, “Generative AI at work,” *The Quarterly Journal of Economics*, vol. 140, no. 2, pp. 889–942, 2025, doi: 10.1093/qje/qjae044.
11. L. Hillebrand, S. Raisch, and J. Schad, “Managing with artificial intelligence: An integrative framework,” *Academy of Management Annals*, vol. 19, no. 1, pp. 343–375, 2025, doi: 10.5465/annals.2022.0072.
12. F. Dell’Acqua *et al.*, “Navigating the jagged technological frontier: Field experimental evidence of the effects of artificial intelligence on knowledge worker productivity and quality,” *Organization Science*, vol. 37, no. 2, pp. 403–423, 2026, doi: 10.1287/orsc.2025.21838.
13. S. Krakowski, D. Haftor, J. Luger, N. Pashkevich, and S. Raisch, “Human-centered artificial intelligence: A field experiment,” *Management Science*, vol. 72, no. 1, pp. 57–72, 2026, doi: 10.1287/mnsc.2022.03849.
14. J. C. Bockstedt and J. R. Buckman, “Humans’ use of AI assistance: The effect of loss aversion on willingness to delegate decisions,” *Management Science*, vol. 72, no. 1, pp. 323–342, 2026, doi: 10.1287/mnsc.2024.05585.
15. H. Li and F. Tian, “Advancing decision-making through AI-human collaboration: A systematic review and conceptual framework,” *Group Decision and Negotiation*, 2026, doi: 10.1007/s10726-026-09980-1.
16. G. Romeo *et al.*, “Exploring automation bias in human–AI collaboration,” *AI & Society*, 2026, doi: 10.1007/s00146-025-02422-7.
17. National Institute of Standards and Technology, *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile*, NIST AI 600-1, 2024.