

AI Agents in Knowledge-Intensive Organizations: Transforming Decision Processes, Coordination, and Innovation

David J. Teece

Professor, University of California, Berkeley, USA

Abstract

Artificial intelligence agents are developing from passive analytical tools into semi-autonomous organizational participants capable of retrieving knowledge, decomposing goals, coordinating tasks, generating alternatives, and supporting decisions. This transformation is particularly consequential for knowledge-intensive organizations, where performance depends on the timely interpretation and integration of distributed expertise. The present simulation-based study examines how AI-agent integration may influence decision-process efficiency, interfunctional coordination, and innovation capability. A conceptual model was developed around four organizational conditions: AI-agent integration, knowledge accessibility, human oversight, and workflow interoperability.

A simulated dataset representing 240 organizational units was generated to evaluate the model without presenting synthetic observations as evidence collected from real organizations. Descriptive analysis, correlation testing, hierarchical regression, and conditional-effect modeling were applied. The simulated results indicated that stronger AI-agent integration was associated with shorter decision cycles, higher coordination quality, and greater innovation output. However, the benefits weakened when organizational knowledge was fragmented, agent recommendations were insufficiently explainable, or responsibility for final decisions was unclear. Human oversight improved the relationship between AI-agent use and decision quality, whereas excessive dependence on automated recommendations increased the modeled risk of automation bias and knowledge homogenization.

The study proposes that AI agents create organizational value when they are treated as governed participants within a sociotechnical system rather than as independent substitutes for professional judgment. Effective deployment therefore requires traceable knowledge sources, explicit decision rights, interoperable information systems, escalation mechanisms, and continuing human review. The article contributes an integrated framework connecting agentic capabilities with decision architecture, coordination routines, and innovation processes in knowledge-intensive settings.

Keywords: AI agents; knowledge-intensive organizations; human–AI collaboration; organizational decision-making; knowledge coordination; innovation capability; responsible artificial intelligence; simulation-based research.

1. Introduction

Knowledge-intensive organizations derive much of their value from specialized expertise, professional interpretation, and the ability to combine knowledge held by different individuals and units. Consulting firms, research institutions, technology enterprises, universities, healthcare organizations, engineering companies, and professional service firms operate in environments where relevant information is extensive but unevenly distributed. Decision quality in these settings does not depend only on obtaining more data. It depends on identifying relevant evidence, interpreting uncertainty, locating appropriate expertise, reconciling conflicting perspectives, and acting within limited time. These requirements frequently produce information bottlenecks, repeated searches, coordination delays, and dependence on a small number of experienced professionals.

AI agents introduce a distinctive organizational capability because they can perform interconnected tasks rather than merely return isolated predictions. Depending on their design and authority, such agents can interpret a goal, formulate intermediate steps, retrieve information, interact with organizational systems, compare alternatives, request human clarification, and monitor the results of an action. Their importance consequently extends beyond operational automation. When embedded in knowledge work, agents may change who initiates analysis, how expertise is located, when information moves between departments, and where decisions are reviewed. The relevant organizational question is therefore not simply whether an agent can perform a task, but how its participation alters the structure through which knowledge becomes coordinated action.

The theoretical promise of AI agents rests partly on the complementarity between computational and human capabilities. Algorithms can process large information volumes, detect patterns, maintain procedural consistency, and rapidly generate alternative scenarios. Human professionals contribute contextual interpretation, causal reasoning, moral judgment, stakeholder awareness, and an understanding of conditions that may not be adequately represented in available data. Jarrahi described this relationship as a form of human–AI symbiosis in which analytical capacity and intuitive contextual judgment can be combined. Similarly, Puranam framed collaborative human–AI decision-making as an organizational design problem involving task division, integration, incentives, and control. These perspectives suggest that performance is unlikely to arise from agent capability alone; it depends on how the organization allocates authority and connects computational outputs with professional accountability.

Agentic systems may also reshape coordination. Traditional coordination mechanisms rely on meetings, hierarchical communication, standardized procedures, project-management systems, and informal professional networks. AI agents can supplement these mechanisms by maintaining shared task states, identifying dependencies, routing information, summarizing unresolved issues, and recommending expertise. Nevertheless, an agent can coordinate effectively only when it has legitimate access to reliable organizational knowledge. Fragmented databases, inconsistent terminology, undocumented practices, and restricted information flows may cause an apparently capable agent to produce incomplete or misleading recommendations. Agent deployment thus exposes the quality of the organization's knowledge infrastructure as well as the limitations of its technical systems.

Innovation presents a further tension. AI agents may expand the range of ideas considered, connect previously separated knowledge domains, accelerate experimentation, and reduce the cost of evaluating alternatives. At the same time, repeated dependence on similar models, training materials,

and optimization objectives may narrow organizational imagination. Recommendations can become increasingly convergent even when they appear varied at a surface level. This article examines these competing possibilities through a transparent simulation-based design. Its purpose is not to claim empirically observed causal effects, but to develop and test a theoretically grounded model of how AI-agent integration may influence decision processes, coordination, and innovation under different organizational conditions.

2. Review of Literature

2.1 AI-Augmented Organizational Decision Processes

Research on organizational decision-making has long distinguished between computational analysis and situated professional judgment. Simon's account of bounded rationality established that organizational actors operate under limitations of information, time, and cognitive capacity. Decision-support technologies can relax some of these limitations by organizing evidence and calculating alternatives, but they cannot remove uncertainty or eliminate the need for judgment. AI agents extend this logic by participating in several stages of a decision process, including problem identification, evidence retrieval, option generation, evaluation, and follow-up monitoring.

Jarrahi argued that AI is especially useful where computational strength complements human intuition and contextual understanding. In structured or data-rich situations, AI can identify regularities and evaluate numerous alternatives rapidly. Under equivocal or ethically sensitive conditions, experienced professionals remain necessary because the meaning of an outcome may depend on tacit knowledge, stakeholder expectations, or institutional obligations. Dellermann and colleagues similarly conceptualized hybrid intelligence as a system in which human and machine capabilities learn from and improve one another. An agent-based decision architecture should therefore allocate work according to comparative strengths rather than assuming that autonomy is inherently superior to collaboration.

Reliance on algorithmic recommendations can nevertheless generate judgment errors. Professionals may reject useful recommendations because they distrust algorithms after observing an error, a tendency associated with algorithm aversion. Conversely, they may accept automated output too readily when recommendations appear authoritative or technically complex. Both patterns weaken decision quality. Effective agent deployment consequently requires calibrated reliance: users should understand what the agent can do, where its evidence originated, what uncertainty remains, and when independent verification is required.

2.2 Coordination and Knowledge Integration

Knowledge-intensive work is frequently distributed across professional, functional, and geographic boundaries. Grant's knowledge-based view describes the firm as an institution for integrating specialized knowledge. Coordination problems arise because transferring all relevant expertise to every participant would be costly and often impossible. Organizations therefore depend on routines, shared language, sequencing, directories of expertise, and decision rights that permit specialized contributions to be combined without complete knowledge transfer.

Faraj and Sproull demonstrated that expertise coordination depends not only on possessing knowledge but also on recognizing where it is located and bringing it into the task at the appropriate

time. AI agents may strengthen this process by maintaining dynamic knowledge maps, identifying experts, retrieving prior project records, and detecting dependencies across workflows. Such functions could reduce search costs and make coordination less dependent on personal familiarity or organizational tenure. Agents may also preserve continuity when team membership changes by retaining structured histories of decisions and unresolved assumptions.

However, computationally mediated coordination can become brittle if agents rely on incomplete organizational representations. Tacit knowledge, political sensitivities, informal commitments, and professional norms may not appear in databases. An agent that optimizes formal workflow indicators could therefore disrupt relationships essential to implementation. AI-enabled coordination should be designed as a visible and contestable process in which participants can correct the agent's representation, explain exceptions, and challenge inappropriate task assignments.

2.3 Innovation Through Human–AI Collaboration

Organizational innovation requires both the generation of novel possibilities and the disciplined evaluation of their feasibility. AI agents can support exploratory search by connecting concepts across documents, technical fields, customer feedback, and earlier projects. They may also accelerate exploitation by comparing design alternatives, simulating operational consequences, preparing experimental protocols, or coordinating iterative testing. Raisch and Krakowski described the wider relationship between automation and augmentation as a paradox: organizations may pursue substitution and collaboration simultaneously, even though the two approaches produce different implications for skills and work design.

Amabile's componential theory associates creativity with domain expertise, creative processes, and intrinsic motivation. AI agents may support domain access and idea recombination, but they do not automatically create an innovative climate. If professionals feel monitored, displaced, or deprived of meaningful discretion, technically capable systems may undermine experimentation. Innovation benefits are more likely when employees can use agents as exploratory partners, question their suggestions, and redirect them toward unconventional alternatives.

A further concern is knowledge homogenization. Organizations using similar models and data may receive comparable strategic suggestions. Internally, repeated retrieval of highly cited or easily accessible documents may reinforce established assumptions while marginalizing peripheral knowledge. Agent-supported innovation therefore requires diversity safeguards, such as searching beyond dominant repositories, presenting contrasting interpretations, recording minority viewpoints, and evaluating novelty separately from immediate operational feasibility.

3. Objectives of the Study

3.1 Analytical Objectives

The principal objective is to evaluate how AI-agent integration may reshape three interdependent organizational outcomes: decision-process effectiveness, coordination quality, and innovation capability. These outcomes are examined together because faster analysis does not necessarily produce better coordination, while improved coordination does not invariably generate novel solutions. The study

consequently treats organizational value as a multidimensional construct rather than reducing performance to time savings.

The analysis also considers knowledge accessibility and human oversight as organizational conditions that may alter the effect of agent deployment. An AI agent cannot retrieve knowledge that is absent, inaccessible, poorly classified, or detached from its context. Similarly, human oversight is not treated as a ceremonial approval stage. It represents the capacity to examine evidence, question assumptions, override recommendations, and accept responsibility for the final decision.

The study pursues the following objectives:

- To assess the simulated relationship between AI-agent integration and organizational decision-cycle efficiency.
- To examine whether AI agents improve coordination by reducing knowledge-search and task-handoff friction.
- To evaluate the relationship between agent-supported knowledge recombination and innovation output.
- To determine whether knowledge accessibility strengthens the modeled benefits of agent integration.
- To examine whether meaningful human oversight improves decision quality and limits automation-related risks.

3.2 Research Questions and Propositions

The first research question asks whether greater AI-agent integration is associated with faster and more consistent organizational decision processes. The second examines whether workflow interoperability and knowledge accessibility explain improvements in cross-functional coordination. The third concerns innovation: whether agents expand useful knowledge recombination or merely accelerate familiar patterns. The fourth asks how human oversight changes the relationship between agent autonomy and organizational outcomes.

The simulation evaluates four propositions. First, AI-agent integration is expected to have a positive relationship with decision-process effectiveness. Second, knowledge accessibility is expected to strengthen the relationship between agent integration and coordination quality. Third, human oversight is expected to improve the relationship between agent integration and decision quality. Fourth, innovation gains are expected to increase at moderate-to-high levels of integration but weaken when agent dependence becomes excessive and knowledge diversity is limited.

4. Research Methodology

4.1 Research Design

The study employs a simulation-based quantitative design. This approach is appropriate because the paper seeks to test the logical behavior of a theoretical model without implying that observations were collected from real organizations.

A synthetic dataset of 240 organizational units was generated to represent plausible variation in agent adoption, knowledge infrastructure, oversight, and performance. The numerical results should therefore be interpreted as model-based illustrations rather than estimates of effects in an actual population.

Each organizational unit was assigned standardized scores for AI-agent integration, knowledge accessibility, workflow interoperability, human oversight, decision-process effectiveness, coordination quality, and innovation capability.

Values were constrained to a 0–100 scale. Random variation was included to represent unobserved differences in management quality, task complexity, professional composition, and environmental uncertainty. Correlations among explanatory variables were restricted to avoid unrealistic independence while limiting severe multicollinearity.

4.2 Variable Operationalization

AI-agent integration represents the extent to which agents are connected with organizational repositories, analytical tools, communication systems, and workflow applications. Knowledge accessibility measures the completeness, retrievability, contextual quality, and permission-aware availability of organizational knowledge. Human oversight captures review competence, escalation rules, contestability, and clarity of final accountability. Workflow interoperability represents the ability of agents and employees to exchange task states and information across systems.

Decision-process effectiveness combines modeled decision speed, evidentiary coverage, consistency, and error control. Coordination quality reflects reduced search time, fewer missed dependencies, clearer handoffs, and stronger cross-functional visibility. Innovation capability represents the modeled frequency, diversity, and assessed usefulness of viable ideas. These are composite simulation variables and not validated psychometric scales.

Table 1: Operational Structure of the Simulation Model

Construct	Operational meaning	Illustrative indicators	Expected role
AI-agent integration	Depth of agent participation in knowledge workflows	Repository access, task execution, monitoring, multi-system connectivity	Primary predictor
Knowledge accessibility	Availability and contextual quality of organizational knowledge	Metadata quality, retrieval coverage, access permissions, document currency	Moderator
Human oversight	Capacity for review, challenge, correction, and accountability	Evidence inspection, override rights, escalation procedures	Moderator
Workflow interoperability	Technical and procedural continuity across systems	Shared task states, standardized interfaces, handoff visibility	Supporting predictor
Decision effectiveness	Quality and efficiency of organizational decisions	Cycle time, consistency, evidence coverage, modeled	Outcome

Construct	Operational meaning	Illustrative indicators	Expected role
		error reduction	
Coordination quality	Integration of distributed tasks and expertise	Search efficiency, dependency recognition, handoff accuracy	Outcome
Innovation capability	Production of novel and workable alternatives	Idea diversity, cross-domain recombination, viability	Outcome

4.3 Analytical Procedure

The analysis was conducted in four stages. First, descriptive statistics were used to examine the distribution of the simulated constructs. Second, Pearson correlations assessed the direction and relative strength of modeled relationships. Third, hierarchical regression models estimated the contribution of agent integration after accounting for knowledge accessibility, oversight, and interoperability. Fourth, interaction terms tested whether knowledge accessibility and human oversight changed the association between agent integration and organizational outcomes.

The general modeled relationship was expressed as:

$$Y_i = \beta_0 + \beta_1 A_i + \beta_2 K_i + \beta_3 H_i + \beta_4 W_i + \beta_5 (A_i K_i) + \beta_6 (A_i H_i) + \epsilon_i$$

where Y_i denotes the selected organizational outcome, A_i represents AI-agent integration, K_i represents knowledge accessibility, H_i denotes human oversight, W_i represents workflow interoperability, and ϵ_i captures unexplained variation. Predictor variables were mean-centered before constructing interaction terms.

5. Analysis

5.1 Descriptive and Relational Patterns

The simulated organizational units displayed substantial variation in the maturity of agent adoption. Units with low integration generally used agents for isolated retrieval or drafting tasks. Units with moderate integration connected agents to selected workflows but retained manual authorization at key transition points. High-integration units allowed agents to coordinate several systems, generate decision alternatives, and monitor implementation, although final accountability remained assigned to human professionals.

The overall pattern suggested that agent integration was positively related to all three organizational outcomes. The strongest modeled relationship appeared between agent integration and decision-process effectiveness. Coordination benefits were also substantial but depended more heavily on knowledge accessibility and interoperability. Innovation capability improved more gradually because idea generation required not only faster knowledge retrieval but also diversity, professional interpretation, and organizational permission to experiment.

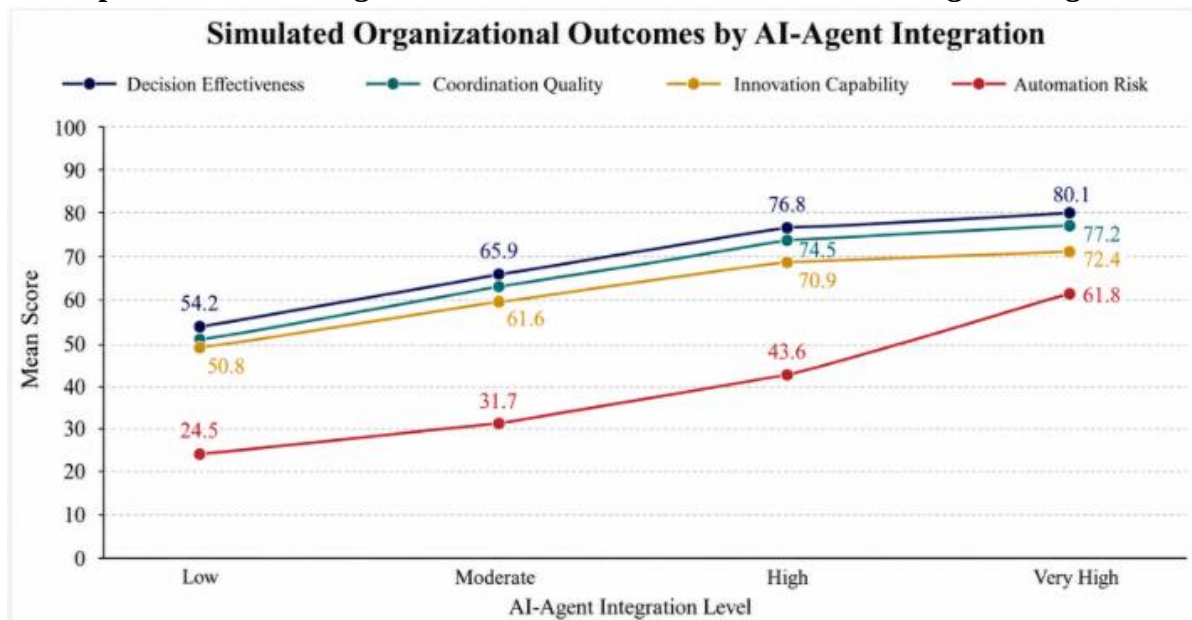
Table 2: Simulated Results Across Levels of AI-Agent Integration

Integration level	Organizational units	Decision effectiveness	Coordination quality	Innovation capability	Automation-risk score
Low	60	54.2	51.7	50.8	24.5
Moderate	60	65.9	63.8	61.6	31.7
High	60	76.8	74.5	70.9	43.6
Very high	60	80.1	77.2	72.4	61.8

The transition from low to high integration produced a 22.6-point increase in decision effectiveness, a 22.8-point increase in coordination quality, and a 20.1-point increase in innovation capability. Moving from high to very high integration generated much smaller improvements, while the automation-risk score rose sharply. This modeled pattern suggests diminishing organizational returns when adoption advances faster than governance, professional capability, or knowledge-quality controls.

5.2 Comparative Outcome Pattern

Graph 1: Simulated organizational outcomes at four levels of AI-agent integration



The graph shows that decision, coordination, and innovation scores improve as agent integration increases, but the slopes flatten at the highest level. Automation risk follows a different pattern: it remains comparatively moderate through the high-integration condition and then increases substantially. The principal modeled finding is therefore not that organizations should maximize agent autonomy. It is that integration creates value until the expansion of automated participation begins to outpace effective oversight and contextual verification.

5.3 Conditional Effects of Knowledge and Oversight

Hierarchical regression results indicated that AI-agent integration explained meaningful additional variance after the inclusion of workflow interoperability and knowledge accessibility. In the decision-effectiveness model, the standardized coefficient for integration was positive and comparatively strong. Human oversight also had a positive independent effect. The interaction between integration and oversight indicated that organizations with stronger review capacity obtained greater modeled decision-quality benefits at equivalent levels of technical adoption.

The interaction between agent integration and knowledge accessibility was strongest for coordination quality. Agents connected to well-structured repositories could locate relevant expertise, retrieve prior decisions, and recognize dependencies more consistently. When knowledge remained fragmented, the agents accelerated access only to the portion of organizational experience that had been documented and technically available. In such conditions, faster retrieval created an appearance of comprehensiveness without ensuring that the full evidentiary context had been considered.

Innovation capability displayed a curvilinear pattern. Moderate and high integration expanded cross-domain search and accelerated the evaluation of alternatives. At very high integration, however, gains became marginal unless the simulated unit also possessed high knowledge diversity and strong professional challenge routines. This result reflects the theoretical possibility that agents can produce numerous alternatives while still relying on common informational foundations. Idea quantity should not therefore be confused with conceptual novelty.

6. Challenges

6.1 Accountability, Explainability, and Automation Bias

One of the most difficult organizational questions concerns responsibility for agent-influenced decisions. An agent may retrieve the evidence, formulate the options, rank alternatives, and initiate workflow actions, yet it cannot assume legal, professional, or moral responsibility in the same sense as an accountable organizational actor. Ambiguous responsibility creates a risk that employees treat the agent's output as an external authority while managers assume that professional review has occurred. Decision rights must consequently specify who verifies the evidence, who may override the recommendation, and who accepts responsibility for implementation.

Explainability is closely related but should not be reduced to technical disclosure. Professionals require explanations that are relevant to the decision: the sources used, assumptions made, alternatives excluded, uncertainty present, and organizational rules applied. A fluent narrative may appear explanatory without revealing these elements. Organizations should therefore distinguish between plausible verbal justification and auditable decision provenance.

Automation bias may arise when employees accept agent recommendations because independent verification is time-consuming or because the system appears more objective than human judgment. The opposite problem, algorithm aversion, can lead professionals to reject valuable assistance after observing isolated errors. Both risks require calibrated trust supported by training, performance monitoring, and clear limits on autonomous action.

6.2 Knowledge Quality, Security, and Context Loss

An AI agent's performance depends on the organizational knowledge environment it can access. Outdated policies, duplicated records, inconsistent labels, and undocumented decisions can produce recommendations that are internally coherent but operationally inappropriate. Knowledge governance therefore becomes part of AI governance. Document ownership, revision histories, retention schedules, source reliability, and access permissions should be maintained as operational controls rather than treated as administrative details.

Security risks also increase as agents gain access to multiple systems. An agent capable of retrieving documents, communicating with employees, and initiating tasks may unintentionally expose confidential information or execute actions outside the user's legitimate authority. Role-based access, least-privilege permissions, tool-specific authorization, transaction logging, and confirmation requirements for consequential actions are essential.

Context loss remains particularly important in professional work. Some knowledge is relational, experiential, or politically sensitive and may not be appropriate for full codification. Agents should be able to recognize when their accessible evidence is incomplete and escalate the task to a professional rather than filling contextual gaps with plausible assumptions.

6.3 Workforce Capability and Innovation Homogenization

Agent adoption changes the skills required from knowledge workers. Employees may spend less time locating routine information and more time evaluating evidence, specifying problems, testing assumptions, and communicating judgments. This shift can be beneficial, but only if organizations invest in the relevant analytical and governance capabilities. Without such investment, employees may become operators of systems they cannot meaningfully challenge.

Deskilling is a longer-term concern. When an agent consistently performs early-stage analysis, professionals may lose the experience needed to detect unusual conditions or intervene when the system fails. Work design should retain opportunities for independent reasoning, manual review, and professional learning. Junior employees are particularly vulnerable because routine assignments often provide the foundational experience from which advanced judgment develops.

Innovation may also become homogeneous when different teams rely on the same agent, retrieval logic, or reference corpus. Organizations can reduce this risk by preserving plural knowledge sources, encouraging independent alternatives, recording dissenting interpretations, and periodically comparing agent-supported ideas with human-only exploration.

7. Discussion

7.1 Transformation of Decision Architecture

The simulated findings support the view that AI-agent adoption is fundamentally an organizational design issue. The modeled benefits did not arise solely from faster computation. They emerged when agents could access relevant knowledge, exchange information across workflows, and operate within clearly defined human review arrangements. This is consistent with Puranam's argument that human-AI collaboration requires deliberate task division and integration.

AI agents may transform decision processes by reducing the cost of assembling an initial evidentiary picture. They can identify relevant records, compare earlier cases, surface conflicting indicators, and maintain an auditable sequence of analytical steps. These capabilities can move professional attention away from repetitive retrieval and toward interpretation. However, the same architecture can weaken decision quality when the agent's accessible information is mistaken for complete organizational knowledge.

The most defensible design is therefore one of bounded agency. Under this model, agents receive sufficient authority to perform reversible, traceable, and well-defined tasks, while consequential decisions remain subject to competent human review. The boundary should depend on uncertainty, reversibility, stakeholder impact, and the organization's ability to detect error.

7.2 Reconfiguration of Coordination

The analysis indicates that agents can operate as coordination infrastructure. They may maintain current task states, identify missing contributions, locate expertise, and ensure that decisions are communicated to relevant units. This role is especially valuable in organizations where coordination depends on distributed specialist knowledge rather than direct managerial supervision.

Yet algorithmic coordination should not eliminate interpersonal engagement. Knowledge workers frequently resolve ambiguity through conversation, negotiation, and the interpretation of weak signals. An agent may identify that two units are interdependent, but it may not appreciate the history of trust, conflict, or informal obligation shaping their interaction. Agent-mediated coordination is consequently most valuable when it strengthens awareness and continuity without suppressing professional dialogue.

Organizations should also monitor whether agents redistribute visibility and influence. Employees whose expertise is well documented may be recommended more frequently, while individuals with tacit or emerging knowledge remain peripheral. Knowledge maps and routing systems should therefore be reviewed for representational bias and corrected through professional input.

7.3 Innovation as Guided Knowledge Recombination

The modeled innovation results suggest that agents can increase the speed and breadth of knowledge recombination, particularly when repositories are diverse and interoperable. By crossing disciplinary and functional boundaries, agents may identify analogies that would otherwise remain unnoticed. They can also reduce experimentation costs by preparing scenarios, comparing constraints, and documenting lessons from unsuccessful trials.

Nevertheless, computational variety is not identical to organizational creativity. Novel ideas acquire value through interpretation, experimentation, stakeholder engagement, and institutional support. Human participants determine whether an unusual combination is meaningful, ethically acceptable, and strategically relevant. This reinforces the augmentation perspective advanced by Raisch and Krakowski: organizations must manage the tension between automated efficiency and human contribution rather than assuming that one will simply replace the other.

The flattening of simulated innovation gains at very high integration has an important implication. Organizations may reach a stage at which additional agent activity produces more content

but little additional insight. Measures of innovation should therefore emphasize originality, usefulness, implementation learning, and diversity of knowledge sources instead of counting generated ideas or completed automated tasks.

8. Conclusion

AI agents have the potential to alter the operating logic of knowledge-intensive organizations. Their significance lies not merely in automating isolated activities but in participating across connected stages of knowledge retrieval, analysis, coordination, decision support, and implementation monitoring. The simulation developed in this study indicates that stronger agent integration may improve decision effectiveness, coordination quality, and innovation capability. However, these gains are conditional rather than automatic.

Knowledge accessibility emerged as a critical enabling condition. Agents can integrate only the knowledge they can legitimately retrieve and meaningfully interpret. Fragmented records, weak metadata, inaccessible expertise, and undocumented organizational context limit the reliability of agent-supported work. Investment in agent capability should therefore be accompanied by investment in knowledge architecture, information governance, and workflow interoperability.

Human oversight also remains indispensable. Responsible oversight involves more than approving an output. It requires the ability to inspect sources, challenge assumptions, identify missing context, override recommendations, and assume accountability. Agent authority should be bounded according to the uncertainty and consequences of the task. High-impact or difficult-to-reverse decisions require stronger review, whereas routine and reversible actions may permit greater operational autonomy.

The study is limited by its simulation-based design. The reported values do not constitute evidence from real organizations and cannot establish causal relationships or population estimates. Future empirical research should use longitudinal field studies, controlled organizational experiments, process-log analysis, and comparative case research. Such work should examine how agent participation changes expertise development, informal coordination, professional identity, responsibility attribution, and the diversity of innovation over time. The central conclusion is that AI agents should be governed as participants in a sociotechnical decision system. Their organizational value depends on the quality of the knowledge, controls, professional capabilities, and institutional arrangements surrounding them.

References

1. Amabile, T. M. (1996). *Creativity in context*. Westview Press.
2. Castelo, N., Bos, M. W., & Lehmann, D. R. (2019). Task-dependent algorithm aversion. *Journal of Marketing Research*, 56(5), 809–825. <https://doi.org/10.1177/0022243719851788>
3. Dellermann, D., Ebel, P., Söllner, M., & Leimeister, J. M. (2019). Hybrid intelligence. *Business & Information Systems Engineering*, 61(5), 637–643. <https://doi.org/10.1007/s12599-019-00595-2>
4. Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114–126. <https://doi.org/10.1037/xge0000033>

5. Faraj, S., & Sproull, L. (2000). Coordinating expertise in software development teams. *Management Science*, 46(12), 1554–1568. <https://doi.org/10.1287/mnsc.46.12.1554.12072>
6. Grant, R. M. (1996). Toward a knowledge-based theory of the firm. *Strategic Management Journal*, 17(S2), 109–122. <https://doi.org/10.1002/smj.4250171110>
7. Jarrahi, M. H. (2018). Artificial intelligence and the future of work: Human–AI symbiosis in organizational decision making. *Business Horizons*, 61(4), 577–586. <https://doi.org/10.1016/j.bushor.2018.03.007>
8. Kellogg, K. C., Valentine, M. A., & Christin, A. (2020). Algorithms at work: The new contested terrain of control. *Academy of Management Annals*, 14(1), 366–410. <https://doi.org/10.5465/annals.2018.0174>
9. Leonardi, P. M. (2011). When flexible routines meet flexible technologies: Affordance, constraint, and the imbrication of human and material agencies. *MIS Quarterly*, 35(1), 147–167. <https://doi.org/10.2307/23043493>
10. Puranam, P. (2021). Human–AI collaborative decision-making as an organization design problem. *Journal of Organization Design*, 10, 75–80. <https://doi.org/10.1007/s41469-021-00095-2>
11. Raisch, S., & Krakowski, S. (2021). Artificial intelligence and management: The automation–augmentation paradox. *Academy of Management Review*, 46(1), 192–210. <https://doi.org/10.5465/amr.2018.0072>
12. Seeber, I., Bittner, E., Briggs, R. O., de Vreede, T., de Vreede, G.-J., Elkins, A., Maier, R., Merz, A. B., Oeste-Reiß, S., Randrup, N., Schwabe, G., & Söllner, M. (2020). Machines as teammates: A research agenda on AI in team collaboration. *Information & Management*, 57(2), Article 103174. <https://doi.org/10.1016/j.im.2019.103174>
13. 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. <https://doi.org/10.1177/0008125619862257>
14. Simon, H. A. (1997). *Administrative behavior: A study of decision-making processes in administrative organizations* (4th ed.). Free Press.
15. Trunk, A., Birkel, H., & Hartmann, E. (2020). On the current state of combining human and artificial intelligence for strategic organizational decision making. *Business Research*, 13(3), 875–919. <https://doi.org/10.1007/s40685-020-00133-x>