

HUMAN-AI COLLABORATION FOR RESPONSIBLE SERVICE INNOVATION: A GOVERNANCE FRAMEWORK FOR MSME LEARNING, OVERSIGHT, AND CUSTOMER VALUE

Endang Ruchiyat¹, Ijang Faisal²

¹STIE Ekuitas, Bandung, Indonesia

²Universitas Muhammadiyah Bandung, Bandung, Indonesia

Email: ¹endang.ruchiyat@ekuitas.ac.id; ²kangijang75@gmail.com

Abstract

This conceptual and integrative review develops a framework for human-AI collaboration in micro, small, and medium enterprises. It addresses the problem that AI tools can widen MSME analytical and service capacity, yet uncritical adoption may introduce opaque errors, privacy exposure, skill erosion, and customer distrust. Drawing on the resource-based view, dynamic capabilities, organizational learning, absorptive capacity, social capital, and resilience, the paper explains how task-appropriate augmentation, human review, data minimization, explainability, continuous capability development can be organized as mutually reinforcing routines. The proposed pathway moves from problem diagnosis and capability mapping to bounded experimentation, evidence review, resource reconfiguration, and learning retention. No primary survey, interview, experimental, administrative, or statistical data are claimed. The framework links these mechanisms to service quality, responsible innovation, employee learning, customer value, controlled scalability and identifies managerial, institutional, and research implications suitable for resource-constrained enterprises.

Keywords: artificial intelligence; human-AI collaboration; responsible innovation; MSMEs; service management

INTRODUCTION

The strategic relevance of human-AI collaboration follows from a practical mismatch: AI tools can widen MSME analytical and service capacity, yet uncritical adoption may introduce opaque errors, privacy exposure, skill erosion, and customer distrust. This mismatch is not solved by acquiring technology or joining a program alone. It requires routines that connect information, judgment, resource commitments, and accountable review.

For small enterprises, the cost of a poor decision can be disproportionate because managerial attention and financial buffers are limited. A useful framework must therefore be selective, staged, and reversible. It should explain how task-appropriate augmentation, human review, and data minimization reinforce one another rather than presenting them as isolated practices.

This article asks how MSMEs can organize human-AI collaboration as an evolving capability. The analysis integrates the resource-based view, dynamic capabilities, organizational learning, absorptive capacity, social capital, and resilience. These perspectives clarify why resources create value only when embedded in routines that enable sensing, interpretation, coordinated action, and renewal.

THEORETICAL FOUNDATION

The resource-based view directs attention to valuable resources, but human-AI collaboration depends equally on how resources are combined. Knowledge, relationships, credibility, technology, and routines may be individually available yet strategically weak when they are fragmented or poorly governed (Barney, 1991).

Dynamic-capabilities theory explains renewal through sensing, seizing, and reconfiguration. Sensing identifies relevant change; seizing converts an opportunity into a bounded commitment; reconfiguration reallocates resources when evidence challenges established assumptions (Teece et al., 1997; Teece, 2007).

Absorptive capacity adds an important qualification: external knowledge matters only when an enterprise can recognize its relevance, assimilate it, transform it, and apply it. In human-AI collaboration, this implies that information access must be paired with interpretation, local adaptation, and decision ownership (Cohen & Levinthal, 1990).

Organizational learning distinguishes exploration from exploitation. Exploration tests alternatives; exploitation refines reliable routines. MSMEs need both, but scarce resources require explicit limits on experimentation and clear criteria for retaining what has been learned (March, 1991).

Social capital and ecosystem perspectives explain the relational dimension. Trusted ties lower coordination costs, while bridging ties provide non-redundant information and complementary capabilities. Resilience emerges when these relationships expand options without creating uncontrolled dependence (Adler & Kwon, 2002; Spigel, 2017).

A closer examination of task-appropriate augmentation shows that its value depends on operating discipline rather than formal adoption. In the context of human-AI collaboration, managers need to specify the decision being improved, the information required, the person responsible, and the condition that would trigger corrective action. This specification turns an abstract principle into a reviewable routine and makes its contribution to service quality visible.

The relationship between human review and data minimization is complementary but not automatic. The first mechanism creates an option, while the second determines whether that option can be applied consistently under resource constraints. Their alignment should therefore be reviewed through observable changes in work quality, exception handling, stakeholder response, and employee learning.

An important boundary condition concerns managerial attention. Even a technically sound approach to human-AI collaboration can fail when responsibilities are dispersed, review dates are unclear, or employees lack authority to act on emerging information. Proportionate documentation and short feedback meetings reduce this risk by linking assumptions to outcomes without creating excessive administrative burden.

CONCEPTUAL ARCHITECTURE

The framework begins with a focal decision rather than a broad aspiration. Managers define the problem, affected stakeholders, available evidence, resource ceiling, and review date. This discipline prevents human-AI collaboration from becoming a collection of disconnected activities.

A capability map then identifies existing strengths and gaps across task-appropriate augmentation, human review, data minimization, explainability, and continuous capability development. Mapping should distinguish resources already controlled by the enterprise from capabilities that depend on partners, platforms, advisers, or public institutions.

The third component is a bounded experiment. A useful experiment specifies the expected mechanism, observable intermediate outcome, responsibility, time horizon, and stop condition. It is designed to generate decision-relevant learning, not merely to demonstrate activity.

Evidence review connects action to service quality, responsible innovation, employee learning. Managers compare the result with the baseline and ask whether improvement can plausibly be attributed to the changed routine. Weak evidence triggers redesign or discontinuation rather than automatic expansion.

Learning retention completes the architecture. Assumptions, actions, exceptions, outcomes, and lessons are recorded in a format proportionate to enterprise size. This record becomes an organizational asset when it changes future choices and reduces repeated error.

A closer examination of task-appropriate augmentation shows that its value depends on operating discipline rather than formal adoption. In the context of human-AI collaboration, managers need to specify the decision being improved, the information required, the person responsible, and the condition that would trigger corrective action. This specification turns an abstract principle into a reviewable routine and makes its contribution to service quality visible.

The relationship between human review and data minimization is complementary but not automatic. The first mechanism creates an option, while the second determines whether that option can be applied consistently under resource constraints. Their alignment should therefore be reviewed through observable changes in work quality, exception handling, stakeholder response, and employee learning.

An important boundary condition concerns managerial attention. Even a technically sound approach to human-AI collaboration can fail when responsibilities are dispersed, review dates are unclear, or employees lack authority to act on emerging information. Proportionate documentation and short feedback meetings reduce this risk by linking assumptions to outcomes without creating excessive administrative burden.

Proposition 1. When task-appropriate augmentation is connected to explicit decision responsibility and a scheduled evidence review, improvements in service quality are more likely to be retained as an organizational routine than when the practice is introduced as a stand-alone activity.

Proposition 2. When human review is connected to explicit decision responsibility and a scheduled evidence review, improvements in responsible innovation are more likely to be retained as an organizational routine than when the practice is introduced as a stand-alone activity.

CAPABILITY DEVELOPMENT

Capability development should be anchored in recurring tasks associated with human-AI collaboration. Training has limited value when it is detached from decisions that employees and owners actually perform. Practice, feedback, and reflection should occur around a real workflow.

The first capability cluster concerns interpretation: employees must recognize relevant signals, assess information quality, and explain uncertainty. The second concerns coordination: responsibility for human review and data minimization must be visible across functional and organizational boundaries.

A third cluster concerns disciplined experimentation. Small tests protect scarce resources while exposing hidden operational constraints. A fourth concerns review: managers need simple indicators and predetermined decision rules for continuation, redesign, scaling, or exit.

Capability is cumulative when continuous capability development changes later practice. Checklists, short decision notes, after-action reviews, and versioned procedures can preserve learning without imposing excessive bureaucracy.

Leadership behavior is central. Leaders must make uncertainty discussable, invite evidence that challenges preferred solutions, and protect employees who report emerging risks. Without psychological safety, formal controls may produce compliance signals while concealing operational weakness.

A closer examination of task-appropriate augmentation shows that its value depends on operating discipline rather than formal adoption. In the context of human-AI collaboration, managers need to specify the decision being improved, the information required, the person responsible, and the condition that would trigger corrective action. This specification turns an abstract principle into a reviewable routine and makes its contribution to service quality visible.

The relationship between human review and data minimization is complementary but not automatic. The first mechanism creates an option, while the second determines whether that option can be applied consistently under resource constraints. Their alignment should therefore be reviewed through observable changes in work quality, exception handling, stakeholder response, and employee learning.

An important boundary condition concerns managerial attention. Even a technically sound approach to human-AI collaboration can fail when responsibilities are dispersed, review dates are unclear, or employees lack authority to act on emerging information. Proportionate documentation and short feedback meetings reduce this risk by linking assumptions to outcomes without creating excessive administrative burden.

Proposition 1. When task-appropriate augmentation is connected to explicit decision responsibility and a scheduled evidence review, improvements in service quality are more likely to be retained as an organizational routine than when the practice is introduced as a stand-alone activity.

Proposition 2. When human review is connected to explicit decision responsibility and a scheduled evidence review, improvements in responsible innovation are more likely to be retained as an organizational routine than when the practice is introduced as a stand-alone activity.

IMPLEMENTATION PATHWAY

Implementation proceeds through four stages: diagnosis, pilot, controlled scaling, and institutionalization. Diagnosis establishes the focal constraint and baseline. The pilot tests one mechanism under limited exposure. Controlled scaling expands only after evidence and capability readiness are demonstrated.

During diagnosis, the enterprise maps stakeholders, decisions, dependencies, and information flows related to human-AI collaboration. Particular attention is given to single points of failure and assumptions that have not been verified.

During the pilot, scope and resource limits are written in advance. Responsibility for operating the change is separated from responsibility for reviewing its consequences where feasible. This modest separation improves challenge and reduces escalation of commitment.

Controlled scaling considers whether benefits remain credible when volume, complexity, or partner involvement increases. Scaling must not undermine service quality, responsible innovation, or core service continuity.

Institutionalization converts proven practices into routines, role expectations, records, and periodic review. Institutionalization is not permanence: routines remain subject to reassessment when the environment, technology, regulation, or stakeholder expectations change.

A closer examination of task-appropriate augmentation shows that its value depends on operating discipline rather than formal adoption. In the context of human-AI collaboration, managers need to specify the decision being improved, the information required, the person responsible, and the condition that would trigger corrective action. This specification turns an abstract principle into a reviewable routine and makes its contribution to service quality visible.

The relationship between human review and data minimization is complementary but not automatic. The first mechanism creates an option, while the second determines whether that option can be applied consistently under resource constraints. Their alignment should therefore be reviewed through observable changes in work quality, exception handling, stakeholder response, and employee learning.

An important boundary condition concerns managerial attention. Even a technically sound approach to human-AI collaboration can fail when responsibilities are dispersed, review dates are unclear, or employees lack authority to act on emerging information. Proportionate documentation and short feedback meetings reduce this risk by linking assumptions to outcomes without creating excessive administrative burden.

Proposition 1. When task-appropriate augmentation is connected to explicit decision responsibility and a scheduled evidence review, improvements in service quality are more likely to be retained as an organizational routine than when the practice is introduced as a stand-alone activity.

Proposition 2. When human review is connected to explicit decision responsibility and a scheduled evidence review, improvements in responsible innovation are more likely to be retained as an organizational routine than when the practice is introduced as a stand-alone activity.

GOVERNANCE AND RISK

Governance for human-AI collaboration should be proportionate rather than elaborate. Its minimum components are decision rights, information ownership, resource limits, review frequency, escalation routes, and transparent documentation of material assumptions.

Risk assessment focuses on events capable of interrupting continuity or destroying initiative value. These include liquidity pressure, operational failure, inaccurate information, privacy or security incidents, partner dependence, legal exposure, and reputational harm.

Preventive controls reduce avoidable error, detective controls reveal deviations, and corrective controls restore performance. For MSMEs, practical examples include approval thresholds,

reconciliations, access restrictions, backups, partner terms, exception logs, and scheduled review of data minimization.

Governance also requires an exit logic. Stop conditions should be agreed before substantial commitment so that sunk costs and personal attachment do not override evidence. An initiative may be redesigned when the mechanism is plausible but implementation is weak; it should be discontinued when expected value no longer justifies exposure.

Accountability strengthens learning when decisions can be reconstructed. A brief record of who decided, what evidence was used, what uncertainty remained, and when the outcome would be reviewed supports both responsible oversight and controlled scalability.

A closer examination of task-appropriate augmentation shows that its value depends on operating discipline rather than formal adoption. In the context of human-AI collaboration, managers need to specify the decision being improved, the information required, the person responsible, and the condition that would trigger corrective action. This specification turns an abstract principle into a reviewable routine and makes its contribution to service quality visible.

The relationship between human review and data minimization is complementary but not automatic. The first mechanism creates an option, while the second determines whether that option can be applied consistently under resource constraints. Their alignment should therefore be reviewed through observable changes in work quality, exception handling, stakeholder response, and employee learning.

An important boundary condition concerns managerial attention. Even a technically sound approach to human-AI collaboration can fail when responsibilities are dispersed, review dates are unclear, or employees lack authority to act on emerging information. Proportionate documentation and short feedback meetings reduce this risk by linking assumptions to outcomes without creating excessive administrative burden.

Proposition 1. When task-appropriate augmentation is connected to explicit decision responsibility and a scheduled evidence review, improvements in service quality are more likely to be retained as an organizational routine than when the practice is introduced as a stand-alone activity.

Proposition 2. When human review is connected to explicit decision responsibility and a scheduled evidence review, improvements in responsible innovation are more likely to be retained as an organizational routine than when the practice is introduced as a stand-alone activity.

PERFORMANCE AND LEARNING

Performance measurement distinguishes inputs, activities, intermediate outcomes, and final outcomes. Expenditure, participation, adoption, or training completion are activities; they do not by themselves demonstrate improved capability or enterprise value.

Intermediate outcomes for human-AI collaboration may include better information quality, faster exception recognition, stronger coordination, improved employee confidence, or reduced process error. These indicators are useful because they appear before final financial effects.

Final outcomes should match the focal problem and may include service quality, responsible innovation, employee learning, customer value, controlled scalability. Measures should be few

enough to use consistently and sufficiently balanced to prevent one objective from being optimized at the expense of another.

Evaluation is strongest when baseline conditions, review intervals, and interpretation rules are specified before implementation. Where causal attribution is not possible, managers should state uncertainty rather than presenting association as proof.

Short feedback cycles make correction less costly. Learning becomes strategically meaningful only when evidence changes later behavior: successful routines are refined, weak routines are redesigned, and disproven assumptions are retired.

A closer examination of task-appropriate augmentation shows that its value depends on operating discipline rather than formal adoption. In the context of human-AI collaboration, managers need to specify the decision being improved, the information required, the person responsible, and the condition that would trigger corrective action. This specification turns an abstract principle into a reviewable routine and makes its contribution to service quality visible.

The relationship between human review and data minimization is complementary but not automatic. The first mechanism creates an option, while the second determines whether that option can be applied consistently under resource constraints. Their alignment should therefore be reviewed through observable changes in work quality, exception handling, stakeholder response, and employee learning.

An important boundary condition concerns managerial attention. Even a technically sound approach to human-AI collaboration can fail when responsibilities are dispersed, review dates are unclear, or employees lack authority to act on emerging information. Proportionate documentation and short feedback meetings reduce this risk by linking assumptions to outcomes without creating excessive administrative burden.

Proposition 1. When task-appropriate augmentation is connected to explicit decision responsibility and a scheduled evidence review, improvements in service quality are more likely to be retained as an organizational routine than when the practice is introduced as a stand-alone activity.

Proposition 2. When human review is connected to explicit decision responsibility and a scheduled evidence review, improvements in responsible innovation are more likely to be retained as an organizational routine than when the practice is introduced as a stand-alone activity.

MANAGERIAL AND POLICY IMPLICATIONS

Owner-managers should begin with one high-consequence decision related to human-AI collaboration, establish a baseline, and test the smallest viable change. This approach concentrates limited attention and produces clearer evidence than broad transformation programs.

Support institutions should move beyond generic workshops toward problem-based assistance, follow-up coaching, and shared learning infrastructure. Universities can contribute diagnostic methods and evaluation capability; associations can aggregate experience; financial and public institutions can reduce access barriers.

Technology providers and ecosystem coordinators should make rules, risks, responsibilities, and exit options understandable to small enterprises. Design quality should be judged partly by whether it enables employee learning and preserves user agency.

Policymakers should favor interoperable support, transparent eligibility, proportionate compliance, and mechanisms that allow small firms to learn without taking irreversible risk. Program evaluation should examine capability and business outcomes rather than registration or participation alone.

The framework also cautions against uniform prescriptions. Sector, enterprise maturity, regional infrastructure, owner capabilities, and stakeholder exposure affect feasibility. Adaptation should preserve the underlying mechanism while allowing the operating form to vary.

A closer examination of task-appropriate augmentation shows that its value depends on operating discipline rather than formal adoption. In the context of human-AI collaboration, managers need to specify the decision being improved, the information required, the person responsible, and the condition that would trigger corrective action. This specification turns an abstract principle into a reviewable routine and makes its contribution to service quality visible.

The relationship between human review and data minimization is complementary but not automatic. The first mechanism creates an option, while the second determines whether that option can be applied consistently under resource constraints. Their alignment should therefore be reviewed through observable changes in work quality, exception handling, stakeholder response, and employee learning.

An important boundary condition concerns managerial attention. Even a technically sound approach to human-AI collaboration can fail when responsibilities are dispersed, review dates are unclear, or employees lack authority to act on emerging information. Proportionate documentation and short feedback meetings reduce this risk by linking assumptions to outcomes without creating excessive administrative burden.

RESEARCH AGENDA AND CONCLUSION

Future research should test the proposed mechanism chain in different sectors and regions. Longitudinal designs are especially valuable because human-AI collaboration develops through repeated decisions, feedback, and resource reconfiguration rather than a single intervention.

Researchers should distinguish capability presence from capability use and examine boundary conditions such as enterprise size, digital maturity, environmental turbulence, owner concentration, and ecosystem quality. Mixed methods can clarify why apparently similar practices generate different outcomes.

Measurement development should prioritize transparent constructs and observable routines. Comparative case studies can identify configurations of mechanisms, while surveys and administrative evidence can later test associations without inventing causal certainty.

In conclusion, human-AI collaboration is best understood as an organizational capability built through connected routines. The proposed framework links task-appropriate augmentation, human review, data minimization, explainability, continuous capability development to service quality, responsible innovation, employee learning, customer value, controlled scalability through diagnosis, bounded action, evidence review, and retained learning.

The central implication is practical: resilience and responsible development do not require maximal complexity. They require clear problems, proportionate commitments, credible evidence, accountable relationships, and the discipline to change course when learning demands it.

A closer examination of task-appropriate augmentation shows that its value depends on operating discipline rather than formal adoption. In the context of human-AI collaboration, managers need to specify the decision being improved, the information required, the person responsible, and the condition that would trigger corrective action. This specification turns an abstract principle into a reviewable routine and makes its contribution to service quality visible.

The relationship between human review and data minimization is complementary but not automatic. The first mechanism creates an option, while the second determines whether that option can be applied consistently under resource constraints. Their alignment should therefore be reviewed through observable changes in work quality, exception handling, stakeholder response, and employee learning.

An important boundary condition concerns managerial attention. Even a technically sound approach to human-AI collaboration can fail when responsibilities are dispersed, review dates are unclear, or employees lack authority to act on emerging information. Proportionate documentation and short feedback meetings reduce this risk by linking assumptions to outcomes without creating excessive administrative burden.

REFERENCES

- Adler, P. S., & Kwon, S.-W. (2002). Social capital: Prospects for a new concept. *Academy of Management Review*, 27(1), 17-40.
- Barney, J. (1991). Firm resources and sustained competitive advantage. *Journal of Management*, 17(1), 99-120.
- Chesbrough, H. W. (2003). *Open Innovation: The New Imperative for Creating and Profiting from Technology*. Harvard Business School Press.
- Cohen, W. M., & Levinthal, D. A. (1990). Absorptive capacity: A new perspective on learning and innovation. *Administrative Science Quarterly*, 35(1), 128-152.
- Duchek, S. (2020). Organizational resilience: A capability-based conceptualization. *Business Research*, 13, 215-246.
- Dwivedi, Y. K., et al. (2021). Artificial intelligence: Multidisciplinary perspectives on emerging challenges, opportunities, and agenda for research, practice and policy. *International Journal of Information Management*, 57, 101994.
- Eisenhardt, K. M., & Martin, J. A. (2000). Dynamic capabilities: What are they? *Strategic Management Journal*, 21(10-11), 1105-1121.
- Geissdoerfer, M., Savaget, P., Bocken, N. M. P., & Hultink, E. J. (2017). The circular economy: A new sustainability paradigm? *Journal of Cleaner Production*, 143, 757-768.
- March, J. G. (1991). Exploration and exploitation in organizational learning. *Organization Science*, 2(1), 71-87.
- Nahapiet, J., & Ghoshal, S. (1998). Social capital, intellectual capital, and the organizational advantage. *Academy of Management Review*, 23(2), 242-266.
- Nambisan, S. (2017). Digital entrepreneurship: Toward a digital technology perspective of entrepreneurship. *Entrepreneurship Theory and Practice*, 41(6), 1029-1055.
- OECD. (2021). *The Digital Transformation of SMEs*. OECD Publishing.
- Sarker, S., Chatterjee, S., Xiao, X., & Elbanna, A. (2019). The sociotechnical axis of cohesion for the IS discipline. *MIS Quarterly*, 43(3), 695-719.
- Spigel, B. (2017). The relational organization of entrepreneurial ecosystems. *Entrepreneurship Theory and Practice*, 41(1), 49-72.
- Stam, E. (2015). Entrepreneurial ecosystems and regional policy: A sympathetic critique. *European Planning Studies*, 23(9), 1759-1769.
- Teece, D. J. (2007). Explicating dynamic capabilities: The nature and microfoundations of sustainable enterprise performance. *Strategic Management Journal*, 28(13), 1319-1350.
- Teece, D. J., Pisano, G., & Shuen, A. (1997). Dynamic capabilities and strategic management. *Strategic Management Journal*, 18(7), 509-533.
- Verhoef, P. C., et al. (2021). Digital transformation: A multidisciplinary reflection and research agenda. *Journal of Business Research*, 122, 889-901.
- Vial, G. (2019). Understanding digital transformation: A review and a research agenda. *Journal of Strategic Information Systems*, 28(2), 118-144.
- West, J., & Bogers, M. (2014). Leveraging external sources of innovation: A review of research on open innovation. *Journal of Product Innovation Management*, 31(4), 814-831.