

Artificial Intelligence and Human Multiple Intelligences: Substitution, Complementarity, and the Future of Managerial Capability—A Scopus Evidence Map and Integrative Review

Nicole D'Silva^{1, 2}

¹PhD Scholar, Jamnalal Bajaj Institute of Management Studies, Mumbai, India

²Assistant Professor in Data Science and Technology, Mumbai, India

Keywords	Abstract
artificial intelligence, multiple intelligences, human-AI complementarity, task-level substitution, managerial capability, hybrid intelligence, open and distance learning, management education	Artificial Intelligence (AI) has revived the question of whether machines can rival human intelligence. The field expanded sharply after 2023, although the direct AI-multiple-intelligences intersection remained small. This paper argues that the comparison is better understood through multiple intelligences theory than through a single, general notion of intelligence. It combines a Scopus evidence map of 997 records with an integrative synthesis of 54 scholarly sources and four published meta-analyses to examine substitution, complementarity, and human distinctiveness. The synthesis indicates that AI can simulate or amplify selected linguistic, logical-mathematical, and spatial functions, but does not possess multiple intelligences in the human sense because human intelligence is embodied, developmental, socially situated, and morally accountable. AI substitutes most readily at the task level, complements people across workflows, and remains limited at the responsibility level. Published quantitative evidence also shows that AI assistance often improves performance relative to humans working alone, but does not reliably outperform the better standalone human or AI agent. For managers and ODL providers, the practical challenge is therefore to design hybrid systems that extend access, support, and analytical reach without outsourcing judgement or accountability.

Introduction

Artificial Intelligence (AI) is no longer a distant technological possibility or a laboratory curiosity. AI now drafts text, summarises reports, generates images and code, classifies data, and increasingly shapes the routines of managers, teachers, and open and distance learning (ODL) teams (Benbya et al., 2024; Brynjolfsson et al., 2025). As these tools enter managerial and learner-support workflows, the question is no longer simply whether AI performs well. It is what remains distinctively human and which capabilities institutions should continue to develop and govern. This study therefore examines, through multiple intelligences theory, which human capabilities AI can substitute for, which it can complement, and which remain dependent on human judgement and accountability at the task, workflow, and responsibility levels.

A simple human-versus-machine framing misses the texture of real capability. People work through language, analysis, social reading, self-control, timing, improvisation, and



contextual judgement. Multiple intelligences (MI) theory is useful here as an interpretive lens for keeping those differences visible. Gardner's (1983, 1999) domains are not treated as independent psychometric variables or as a basis for matching instruction to presumed learner types; those claims remain contested (Ferrero et al., 2021; Waterhouse, 2023). The analysis instead considers AI-enabled capability, human capability across MI domains, the relation of substitution or complementarity, and the level at which that relation occurs. Task structure, stakes, trust, context, and human oversight shape the outcome.

Literature Review

Management research has long resisted one-dimensional accounts of capability, viewing effective performance as a combination of technical, human, and conceptual skills (Boyatzis, 1982; Katz, 1955). Human-AI scholarship reaches a similar conclusion: the value of AI lies less in replacing people wholesale than in combining machine capability with human judgement, trust, and oversight (Dellermann et al., 2019; Jarrahi, 2018; Raisch & Krakowski, 2021). In parallel, a growing body of work has extended multiple-intelligences thinking into management and education. Davaei and Gunkel (2024) show that organisational performance draws on a broad mix of cognitive, emotional, social, and cultural resources; D'Silva and Pande (2025a, 2025b) connect multiple intelligences with managerial competency and sustainable leadership in management education; and D'Silva (2025) places AI support within a wider developmental framework that includes mentoring, well-being, and learner growth. AI-in-education research also highlights both the reach of technology-enabled support and the risks it creates for learner agency and accountability (Xia et al., 2025; Zawacki-Richter et al., 2019). However, these strands rarely speak directly to one another. The evidence map identified only 25 records at the direct intersection of AI and multiple intelligences, despite substantially larger literatures on human-AI collaboration, ODL, and organisational learning. The unresolved problem is therefore not a shortage of relevant research, but the absence of an integrating framework that distinguishes task execution, workflow design, and responsibility through a differentiated human-capability lens.

Objectives

Accordingly, the study focused on three objectives: (1) to map the evidence field and AI capabilities against domains of human multiple intelligences; (2) to determine where quantitative and conceptual evidence supports task-level substitution, workflow-level complementarity, or responsibility-level human distinctiveness; and (3) to derive implications for managerial capability, Learning for Development, technology-enabled learning and ODL, and organisational learning transformation.

Methods

Research Methodology

The study used a parallel two-strand review design. 'Strand A' mapped the Scopus evidence field and assessed whether the abstracts contained a defensible pool for a new meta-analysis. 'Strand B' used a structured integrative review to develop the conceptual framework, supported by a separate comparison of relevant published meta-analyses. Integrative-review guidance informed source selection and synthesis (Snyder, 2019; Whitemore & Knafl, 2005). PRISMA-ScR items were used only to report database identification, eligibility, deduplication, and record flow (Tricco et al., 2018). The 997 records were used for the evidence map, while the 54 sources

were selected separately for the integrative review. The 54 sources were not obtained by screening the 997 records down to 54.

The review asked how the evidence field has developed, how current AI capabilities align with the eight MI domains, where substitution or complementarity occurs, and what this means for managers and learning institutions.

Population and Sample

The searchable population comprised English-language Scopus-indexed journal articles and reviews published from January 1, 2000 to July 23, 2026. Four searches covered: AI and multiple intelligences; human-AI collaboration and managerial capability; MI in management and higher education; and AI in management education, technology-enabled learning, ODL, and organisational learning. The main term blocks and yields appear in Table 1. The searches identified 1,118 records. After removing seven duplicates and 114 records outside the year or document-type limits, 997 records remained: 899 Articles and 98 Reviews.

The 54-source integrative synthesis was a parallel purposive sample, not the final stage of the 997-record map. It comprised 10 direct Scopus records, 38 contemporary sources located through citation chaining or targeted checks of central authors and reviews, and six foundational pre-2000 works, including Parasuraman and Riley (1997). Sources were retained when they contributed directly to MI theory or criticism, managerial capability, human-AI collaboration and trust, management or ODL learning, or organisational transformation. Four published meta-analyses were examined separately.

Table 1: Scopus Search Streams, Core Term Blocks, Refinements, and Yields

Stream	Core Scopus Terms	Refinement	Raw / Eligible
S1	AI, GenAI, LLMs, ChatGPT, or agents AND multiple intelligence(s), MI theory, or Gardner	Excluded unrelated AMIS phrase	25 / 25
S2	Human-AI collaboration, interaction, hybrid intelligence, automation-augmentation, or trust	Relation terms in TITLE; management/workplace terms required	340 / 338
S3	Multiple intelligence(s) or MI theory	Management, higher/business education, distance education, or ODL; AMIS excluded	257 / 144
S4	AI, GenAI, LLMs, ChatGPT, or agents	Management education, TEL/ODL, Learning for Development, or organisational learning	496 / 496

Note: Searches used Scopus Advanced Search, 2000-2026, English, Article/Review, and Journal limits. Eligible stream totals overlap; the unique corpus contained 997 records. Exact syntax was archived with the search log.

Tools and Techniques: Reliability and Validity

Scopus exports supplied bibliographic data, abstracts, keywords, and citation counts. Records were profiled by year, source, and topic. Titles, abstracts, and author keywords were mapped through TF-IDF and nine-component non-negative matrix factorisation; the resulting labels described the field and did not decide inclusion. A charting form recorded evidence type, context, AI function, MI domain, capability level, benefits, risks, and contribution. Abstracts

were also checked for samples, experimental language, p -values, and effect indicators. Numerical flags were manually reviewed, but no new pooled estimate was calculated because the candidate studies did not share a sufficiently common population, comparison, outcome, and effect metric.

Reliability was strengthened through an archived search log, DOI/EID/title deduplication, explicit coding rules, repeated inspection of topic labels, and two-pass conceptual coding. Validity was supported by triangulation across management, information systems, psychology, education, and ODL and by keeping the evidence map, synthesis sources, published meta-analyses, and policy materials analytically separate. The review was conducted by one author, so no inter-rater reliability coefficient is claimed.

Procedure

The work proceeded through search and export, deduplication and technical eligibility, descriptive mapping, construction of the integrative sample, abstract-level poolability assessment, comparison of published meta-analyses, objective-wise coding, and interpretation against policy and practice. Figure 1 summarises the parallel review architecture.

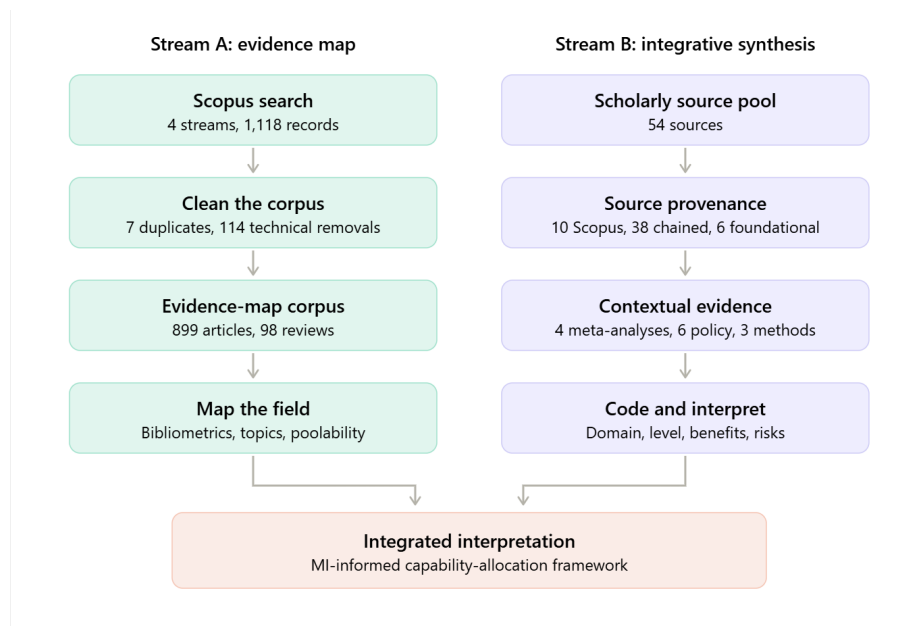


Figure 1: Parallel review architecture. (Source: Author design)

Results

Objective 1: Evidence Profile and AI-MI Domain Mapping

Profile of the Evidence Field

The four searches yielded 1,118 records. After deduplication and technical exclusions, 997 formed the evidence-map corpus. Of these, 795 (79.7%) were published from January 1, 2023 to July 23, 2026, including 305 in 2025 and 290 in the partial 2026 year. Figure 2 shows the sharp recent expansion. The records appeared across more than 600 source titles, confirming that the field is spread across management, education, information systems, human-computer interaction, and applied technology.

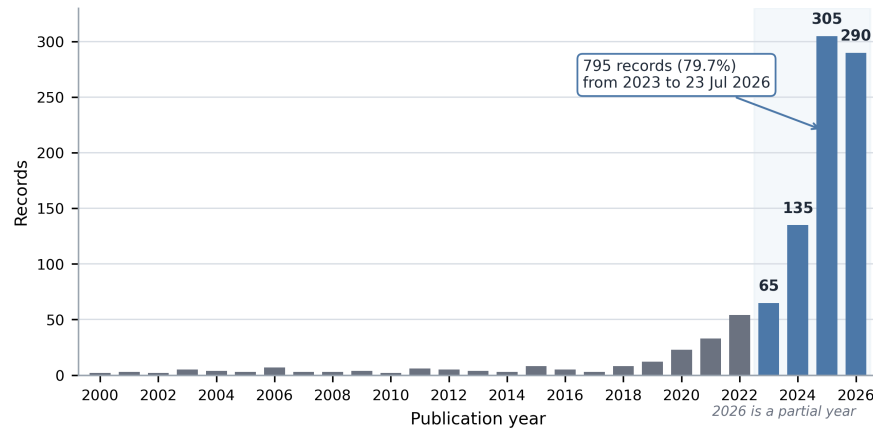


Figure 2: Annual publications in the 997-record Scopus evidence-map corpus, January 1, 2000-July 23, 2026. (Source: Author analysis of Scopus exports)

Only 25 eligible records came from the direct AI-MI search; adjacent streams on human-AI collaboration, management education, ODL, and organisational learning were much larger. Computational mapping reinforced this pattern. The largest topics concerned ODL and learner support (195 records), management education and generative AI (146), human-AI collaboration (146), organisational learning and AI capability (140), and MI and educational capability (111). Smaller clusters addressed trust, hybrid intelligence, leadership, and automation-augmentation. Figure 3 shows the distribution; these metadata topics guided the synthesis but were not treated as full-text findings.

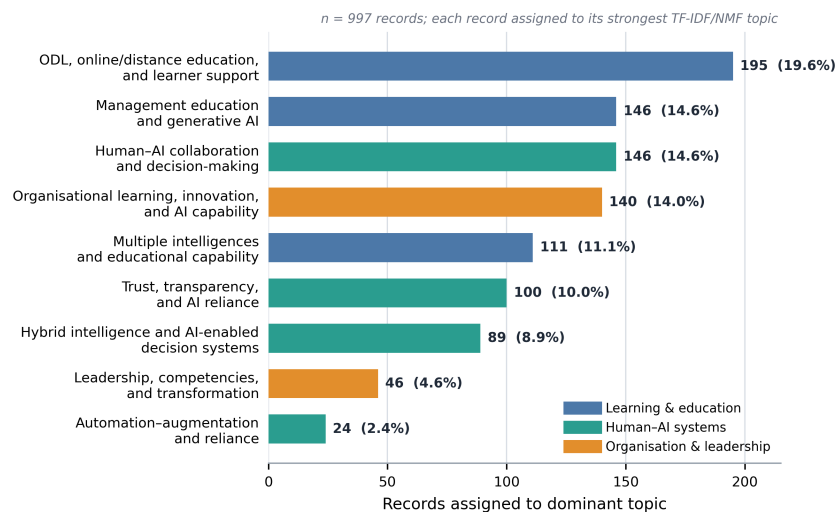


Figure 3: Dominant Scopus topics in the 997-record evidence map. (Source: Author analysis)

AI Capabilities across Multiple-Intelligence Domains

The domain mapping showed a consistent asymmetry (Table 2). AI was strongest where work could be represented through symbols, patterns, searchable data, or decomposable steps. It created the greatest substitution pressure in bounded linguistic and logical-mathematical tasks, while selected spatial tasks were mainly complementary. Interpersonal, intrapersonal, embodied, musical, and naturalistic capability retained a larger human component because they involve lived interpretation, tacit adaptation, cultural meaning, and responsibility.

Table 2: Mapping AI Capabilities against Domains of Human Multiple Intelligences

Intelligence Domain	AI-enabled Strengths	Human Capability that Remains Central	Dominant Relation	Implication
Linguistic	Drafting, summarising, translation, retrieval	Intent, tone, persuasion, authorship, accountability	Substitution + complementarity	Use for routine drafts; retain human ownership.
Logical-mathematical	Classification, prediction, optimisation, coding, scenarios	Problem framing, causality, assumptions, value trade-offs	Selective substitution; strong complementarity	Offload bounded analysis; require validation.
Spatial	Visualisation, image generation, pattern recognition, dashboards	Situated interpretation, physical fit, aesthetic purpose	Complementarity	Use for modelling; retain contextual judgement.
Interpersonal	Conversation, sentiment cues, scripts, first-line support	Empathy, trust, repair, political sensitivity, obligation	Weak substitution	Automate low-risk interaction; keep human escalation.
Intrapersonal	Reflection prompts, feedback synthesis, behaviour tracking	Self-knowledge, identity, moral struggle, self-correction	Assistive complement	Scaffold reflection; do not claim inner awareness.
Bodily-kinesthetic	Robotic repetition, movement sensing, procedural guidance	Tacit adaptation and embodied improvisation	Limited substitution	Useful in structured tasks with oversight.
Musical	Pattern imitation, composition support, variation	Expressive intention and cultural resonance	Complementarity	Treat as creative support, not accountable authorship.
Naturalistic	Classification, sensing, anomaly detection, monitoring	Field knowledge, ecological judgement, stewardship	Complementarity	Combine detection with local expertise.

Note: MI is used as an interpretive capability lens, not as a psychometric classification of individuals.

Linguistic and logical-mathematical capability illustrate the difference between output and responsibility. AI can draft, summarise, translate, classify, predict, compare scenarios, and detect anomalies, making routine communication and analysis faster (Benbya et al., 2024;

Brynjolfsson et al., 2025). Yet consequential intent, problem framing, causal interpretation, value conflict, and ownership of the decision remain human. Producing a plausible answer is not the same as deciding why it should be used or defending its consequences.

Spatial capability was largely complementary: AI can generate images, visualise alternatives, and support dashboards, while physical appropriateness and situated design judgement remain contextual. Agentic systems extend this reach by planning, calling tools, retaining memory, and coordinating parts of a workflow (Bandi et al., 2025; Sapkota et al., 2026). Even so, operational coordination is not interpersonal understanding or moral agency.

Interpersonal and intrapersonal intelligence were the least substitutable. Chatbots can sustain conversational form and offer first-line support, but cannot enter relationships as accountable participants or repair trust through lived commitment. AI can prompt reflection and organise feedback, but cannot possess identity or genuinely self-correct. In bodily, musical, and naturalistic domains, it can guide structured movement, imitate patterns, or classify signals, while embodied improvisation, expressive intention, ecological judgement, and stewardship remain human.

The unit of comparison therefore matters. AI may exceed people in speed, scale, and consistency while still depending on human framing, tacit knowledge, and accountability (Polanyi, 1966). Managerial and ODL work combine several intelligence domains in a single decision. The map is most useful not as a ranking of human and machine intelligence, but as a guide to what may be delegated, what should be augmented, and where work should return to a human decision-maker.

Objective 2: Published Quantitative Evidence and the Substitution-Complementarity-Responsibility Framework

Abstract-level Poolability Assessment and Published Meta-analytic Evidence

The abstract screen found sample information in 151 records, explicit p -values in 26, and readily extractable numerical effect indicators in 17. Manual review identified 13 primary studies with inferential effects, but they covered different populations, tasks, outcomes, and effect families. No group contained at least three independent studies with a common comparison and convertible effect metric. A new pooled estimate would therefore have been misleading.

Four directly relevant published meta-analyses were compared instead as independent quantitative syntheses (Table 3). They were not re-pooled because their baselines, tasks, study sets, and effect definitions differed.

Table 3: Published Meta-analytic Evidence Relevant to Human-AI Substitution and Complementarity

Review and Evidence Base	Contrast	Reported Result	Implication
Vaccaro et al. (2024): 106 experiments; 370 effects	Human-AI vs human alone	$g = 0.64$; 95% CI 0.53 to 0.74	Average augmentation
Vaccaro et al. (2024)	Human-AI vs better standalone agent	$g = -0.23$; 95% CI -0.39 to -0.07	No average synergy
Schemmer et al. (2022): 33 conditions; 5,083 participants	XAI-assisted vs human alone	SMD = 0.59; 95% CI 0.39 to 0.79	Assistance improved performance

Review and Evidence Base	Contrast	Reported Result	Implication
Schemmer et al. (2022): 11 conditions; 999 observations	XAI-assisted vs AI-assisted	SMD = 0.07; 95% CI - 0.15 to 0.30	No clear added benefit from explanations

Note: Estimates are reproduced from the cited published meta-analyses and presented as independent contrasts, not re-pooled effects. The main assistance and synergy estimates were heterogeneous.

The quantitative evidence supports a conditional, not universal, case for augmentation. Human-AI systems generally performed better than humans alone, but did not reliably outperform the better standalone human or AI agent (Vaccaro et al., 2024). Explainable-AI assistance improved performance relative to unaided humans, yet explanations did not show a clear additional benefit over ordinary AI assistance (Schemmer et al., 2022). Marketing and organisational meta-analyses likewise identified task fit, trust, and perceived risk as important conditions (Kumar et al., 2025; Ngo, 2025).

Task Substitution, Workflow Complementarity, and Human Responsibility

These findings support a three-level framework. At the task level, AI may substitute for bounded production such as drafting a standard response, summarising a document, checking code, classifying records, or generating options. The efficiency gain is real, but the task is only one part of a wider chain of interpretation, approval, negotiation, and consequence.

At the workflow level, complementarity is the more durable pattern. AI changes who performs the first pass, how many alternatives can be considered, and where human attention is concentrated. Hybrid value comes from combining machine scale and consistency with human flexibility, contextual judgement, and oversight (Dellermann et al., 2019; Jarrahi, 2018; Raisch & Krakowski, 2021). The manager increasingly becomes the framer, evaluator, integrator, and accountable owner of action.

At the responsibility level, human distinctiveness remains strongest. Organisations and learning institutions need decisions that can be justified to the affected people. Someone must determine what matters, whose interests are missing, what risk is acceptable, and how harm will be addressed. AI can inform these judgements but cannot bear institutional sanction, relational obligation, or moral responsibility.

Trust calibration connects all three levels. Reliance should correspond to demonstrated capability, context, and consequence rather than fluency or convenience (Fuegener et al., 2022; Lee & See, 2004). Verification, escalation, and stopping rules are therefore part of capability design. This matters because AI-supported work may improve immediate performance while weakening motivation or independent follow-through when tasks are poorly designed (Wu et al., 2025).

Four propositions summarise the framework. P1: substitution is most likely for bounded linguistic, logical-mathematical, and selected spatial outputs. P2: as AI use rises, human value shifts towards framing, interpretation, relational judgement, self-regulation, and accountability. P3: hybrid performance depends on task fit, relative human and AI capability, and calibrated delegation; collaboration is not automatically superior. P4: management education and ODL will underprepare learners if AI is treated only as a productivity tool rather than as support for agency, reflection, communication, and responsible action.

Figure 4 brings these relationships together. AI capability is filtered by task structure, stakes, context, trust, and human oversight. The institutional outcome depends not only on what the system can do, but on whether verification, escalation, and responsibility remain visible.

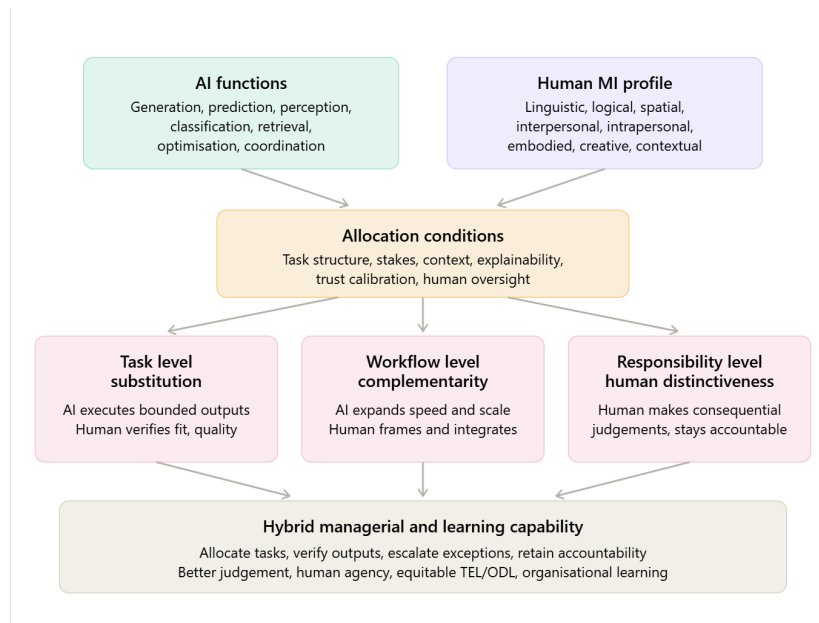


Figure 4: MI-informed framework for human-AI capability allocation. (Source: Author synthesis)

Objective 3: Implications for Managerial Capability and Learning Systems

The third objective shows that managerial capability must be rebalanced rather than simply digitised. Analytical and linguistic fluency remain necessary, but their value increasingly depends on problem formulation, assumption testing, contextual judgement, communication across interests, and the ability to slow or stop an automated process. AI literacy therefore includes understanding limitations, documenting reliance, validating outputs, and retaining decision ownership (Hossain et al., 2025).

Management education should not assess only polished products that AI can generate cheaply. AI-assisted analysis should be combined with oral defence, reflection, negotiation, scenario judgement, and collaborative problem solving. MI is useful here not as a learner-typing device, but as a reminder that managerial development involves analytical, linguistic, interpersonal, intrapersonal, spatial, and embodied modes.

For ODL, AI can improve first-line support, translation, resource discovery, feedback drafting, analytics triage, and access to practice. Evidence suggests benefits for engagement and self-assessment alongside hallucination, access, and context risks (Öncü et al., 2026; Xia et al., 2025). Low-risk enquiries can be automated, but high-stakes or ambiguous matters require visible human ownership and accessible escalation.

For organisational learning, AI can accelerate retrieval, pattern detection, documentation, and coordination. It creates value only when teams challenge outputs, connect them to tacit and local knowledge, and translate them into changed routines. Learning quality, transfer, motivation, adaptability, and accountability are therefore better indicators than time saved or content generated (Dote Pardo, 2026).

Discussion and Implications

Contribution in Relation to Prior Research

The paper's contribution is not the broad claim that humans and AI can complement one another; that is already well established. Its contribution is to show how that relationship changes across capability domains and levels of work. The MI lens predicts stronger substitution where performance can be externalised into symbolic products, and weaker substitution where capability is relational, self-reflective, embodied, or consequence bearing. The task-workflow-responsibility distinction then explains how the same technology can replace part of a job, improve the wider job, and increase the need for human judgement.

The evidence map also shows why this integration is needed. Research on AI, ODL, management education, trust, and organisational learning is expanding rapidly, but direct AI-MI work remains limited. MI is therefore used at its most defensible level: not as a settled psychometric taxonomy, but as a language for plural human capability. This complements managerial competency research, which has long resisted one-dimensional accounts of effectiveness.

The published meta-analyses add an important boundary condition. Assistance can outperform a human-only baseline, yet the combined system may still fall below the better standalone agent. Explanations, too, do not automatically improve decisions. The practical objective should therefore be capability amplification through task fit, division of labour, verification, and calibrated reliance rather than collaboration for its own sake. Agentic AI makes this more urgent because greater workflow autonomy enlarges the consequences of weak permissions, audit trails, and stopping rules (Bandi et al., 2025; Sapkota et al., 2026).

Implications for Learning for Development, TEL, and ODL

For 'Learning for Development', success should be judged by whether AI expands meaningful agency, inclusion, and institutional capability rather than simply increasing digital throughput. AI can widen multilingual or adaptive support, but it can also deepen dependence on opaque platforms or exclude learners through connectivity, language, disability, or data-cost barriers. Human-centred guidance therefore places dignity, capability, and contextual use at the centre of educational AI (Miao & Holmes, 2023; UNESCO, 2024; OECD, 2026).

TEL and ODL providers need a tiered service model. Routine navigation, reminders, translation, and formative practice can be automated. Feedback drafting, analytics alerts, and routine advising require human validation. Disability accommodation, distress, academic misconduct, progression, and appeals should remain human-led. Assessment should likewise combine AI-assisted artefacts with process evidence, explanation, reflection, applied judgement, and source verification so that learning remains visible.

Organisational Learning, Policy, and Practice

Organisational learning should be designed as a cycle of sensing, interpretation, action, reflection, and revision. AI is strong at sensing and retrieval; human teams remain central in connecting signals to strategy, values, tacit history, and stakeholder consequences. Every consequential process should therefore have a named owner, validation protocol, escalation route, and forum in which AI outputs can be questioned.

A practical governance test follows. Before adoption, institutions should specify the task and intended benefit, the consequences of error, the evidence needed to trust the output, the person who owns the decision, the appeal route, and the learning, equity, and motivation

indicators to be monitored. This moves policy beyond asking whether AI is allowed to ask when delegation is educationally, organisationally, and ethically defensible.

Limitations and Future Research

The review has important limits to acknowledge. It uses Scopus-indexed English-language articles and reviews, so local and non-indexed ODL evidence may be underrepresented. The evidence map and poolability screen are metadata/abstract based and cannot establish complete variance, dependence, risk of bias, or long-term effects. No homogeneous primary-study pool supported a new meta-analysis. The 54-source synthesis was transparent but purposive and single-author. Finally, MI was used as an interpretive lens, and 2026 is only a partial publication year.

Future research should test task-intelligence fit across managerial and educational roles; compare human-only, AI-only, and hybrid workflows at the task and responsibility levels; and examine whether repeated AI use strengthens or erodes independent judgement. ODL studies should evaluate tiered support across languages, disabilities, and low-connectivity settings. Management education should test assessment designs that combine AI-assisted production with oral defence, reflection, and collaborative judgement. Longitudinal organisational studies are also needed on motivation, deskilling, trust, knowledge retention, and accountability.

Conclusion

This Scopus evidence map and integrative review brings AI into dialogue with multiple intelligences theory without treating either as a single, undifferentiated capability. The evidence field is large and growing, but the direct AI-MI intersection remains small and the available abstracts do not support one new pooled effect. Published syntheses instead show a conditional pattern: AI often improves human-only performance, yet hybrid systems do not reliably surpass the better standalone agent.

The more useful distinction is therefore among task substitution, workflow complementarity, and responsibility-level human distinctiveness. Managers increasingly need to design and govern hybrid work; management education and ODL should use AI to extend access, practice, feedback, and analytical reach while preserving learner agency, human escalation, and reflective judgement. Organisational learning requires faster information processing to become better at sensemaking and changed practice, rather than simply produce more generated content.

The strongest response is not to ask whether AI is smarter than people. It is to ask what kinds of intelligence good organisations, management programmes, and ODL systems still require from human beings, and how AI can extend those capabilities without displacing responsibility.

Declarations

Conflict of interest: The author declares no conflict of interest.

AI use statement: Generative AI tools were used for language refinement, structural editing, and conceptual-figure drafting. Python supported deduplication, descriptive mapping, and the abstract poolability screen. The author verified the search logic, numerical results, citations, and final text and accepts full responsibility.

References

- Bandi, A., Kongari, B., Naguru, R., Pasnoor, S., & Vilipala, S.V. (2025). The rise of agentic AI: A review of definitions, frameworks, architectures, applications, evaluation metrics, and challenges. *Future Internet*, 17(9), Article 404. <https://doi.org/10.3390/fi17090404>
- Benbya, H., Strich, F., & Tamm, T. (2024). Navigating generative artificial intelligence promises and perils for knowledge and creative work. *Journal of the Association for Information Systems*, 25(1), 23-36. <https://doi.org/10.17705/1jais.00861>
- Boyatzis, R.E. (1982). *The competent manager: A model for effective performance*. Wiley.
- Brynolfsson, E., Li, D., & Raymond, L. (2025). Generative AI at work. *The Quarterly Journal of Economics*, 140(2), 889-942. <https://doi.org/10.1093/qje/qjae044>
- Davaei, M., & Gunkel, M. (2024). *The role of intelligences in teams: A systematic literature review. Review of Managerial Science*, 18(1), 259-297. <https://doi.org/10.1007/s11846-023-00672-7>
- 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>
- D'Silva, M.N. (2025). Building resilience and academic success: An integrated framework of multiple intelligences, mentoring, psychological well-being, and AI support in higher education. *Journal of Marketing & Social Research*, 2(2), 685-689.
- D'Silva, M.N., & Pande, A. (2025a). Mapping the nexus between multiple intelligences and managerial competency: A systematic review. *Journal of Informatics Education and Research*, 5(3), 1450-1466. <https://doi.org/10.52783/jier.v5i3.3404>
- D'Silva, M.N., & Pande, A. (2025b). Multiple intelligences in management education: A path analysis for sustainable leadership. *Northern Economic Review*, 16(1), 70-85. <https://doi.org/10.10399/NER.2025592102>
- Dote Pardo, J.S. (2026). Artificial intelligence in organizational learning: A systematic review of applications, opportunities, and challenges. *Development and Learning in Organizations: An International Journal*, 40(2), 4-9. <https://doi.org/10.1108/DLO-06-2025-0228>
- Ferrero, M., Vadillo, M.A., & León, S.P. (2021). A valid evaluation of the theory of multiple intelligences is not yet possible: Problems of methodological quality for intervention studies. *Intelligence*, 88, Article 101566. <https://doi.org/10.1016/j.intell.2021.101566>
- Fuegener, A., Grahl, J., Gupta, A., & Ketter, W. (2022). Cognitive challenges in human-artificial intelligence collaboration: Investigating the path toward productive delegation. *Information Systems Research*, 33(2), 678-696. <https://doi.org/10.1287/isre.2021.1079>
- Gardner, H. (1999). *Intelligence reframed: Multiple intelligences for the 21st century*. Basic Books.
- Gardner, H. (1983). *Frames of mind: The theory of multiple intelligences*. Basic Books.
- Hossain, S., Fernando, M., & Akter, S. (2025). Digital leadership: Towards a dynamic managerial capability perspective of artificial intelligence-driven leader capabilities. *Journal of Leadership & Organizational Studies*, 32(2), 189-208. <https://doi.org/10.1177/15480518251319624>
- 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>
- Katz, R.L. (1955). Skills of an effective administrator. *Harvard Business Review*, 33(1), 33-42.
- Kumar, N., Kumar, R.R., & Raj, A. (2025). Establishing antecedents and outcomes of human-AI collaboration: Meta-analysis. *Journal of Computer Information Systems*. Advance online publication. <https://doi.org/10.1080/08874417.2025.2492885>
- Lee, J.D., & See, K.A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50-80. https://doi.org/10.1518/hfes.46.1.50_30392
- Miao, F., & Holmes, W. (2023). *Guidance for generative AI in education and research*. UNESCO.
- Ngo, V.M. (2025). Human-AI collaboration for marketing capabilities: A meta-analysis. *Marketing Letters*, 36, 839-855. <https://doi.org/10.1007/s11002-025-09783-5>

- OECD. (2026). *OECD digital education outlook 2026: Exploring effective uses of generative AI in education*. OECD Publishing. <https://doi.org/10.1787/062a7394-en>
- Öncü, S.E., Gevher, M., Erdoğan, E., & Koçdar, S. (2026). Exploring the potential of generative AI for academic support in open and distance learning: A case study of learner experiences. *The International Review of Research in Open and Distributed Learning*, 27(2), 67-82. <https://doi.org/10.19173/irrodl.v27i2.9289>
- Parasuraman, R., & Riley, V. (1997). Humans and automation: Use, misuse, disuse, abuse. *Human Factors*, 39(2), 230–253. <https://doi.org/10.1518/001872097778543886>
- Polanyi, M. (1966). *The tacit dimension*. Doubleday.
- 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>
- Sapkota, R., Roumeliotis, K.I., & Karkee, M. (2026). AI agents vs. agentic AI: A conceptual taxonomy, applications and challenges. *Information Fusion*, 126, Article 103599. <https://doi.org/10.1016/j.inffus.2025.103599>
- Schemmer, M., Hemmer, P., Nitsche, M., Köhl, N., & Vössing, M. (2022). A meta-analysis of the utility of explainable artificial intelligence in human-AI decision-making. In *Proceedings of the 2022 AAAI/ACM Conference on AI, Ethics, and Society* (pp. 617-626). ACM. <https://doi.org/10.1145/3514094.3534128>
- Snyder, H. (2019). Literature review as a research methodology: An overview and guidelines. *Journal of Business Research*, 104, 333-339. <https://doi.org/10.1016/j.jbusres.2019.07.039>
- Tricco, A.C., Lillie, E., Zarin, W., O'Brien, K.K., Colquhoun, H., Levac, D., Moher, D., Peters, M.D.J., Horsley, T., Weeks, L., Hempel, S., ... Straus, S.E. (2018). PRISMA extension for scoping reviews (PRISMA-ScR): Checklist and explanation. *Annals of Internal Medicine*, 169(7), 467-473. <https://doi.org/10.7326/M18-0850>
- UNESCO. (2024). *AI competency framework for teachers*. <https://doi.org/10.54675/ZJTE2084>
- Vaccaro, M., Almaatouq, A., & Malone, T.W. (2024). When combinations of humans and AI are useful: A systematic review and meta-analysis. *Nature Human Behaviour*, 8, 2293-2303. <https://doi.org/10.1038/s41562-024-02024-1>
- Waterhouse, L. (2023). Why multiple intelligences theory is a neuromyth. *Frontiers in Psychology*, 14, Article 1217288. <https://doi.org/10.3389/fpsyg.2023.1217288>
- Whittemore, R., & Knafl, K. (2005). The integrative review: Updated methodology. *Journal of Advanced Nursing*, 52(5), 546-553. <https://doi.org/10.1111/j.1365-2648.2005.03621.x>
- Wu, S., Liu, Y., Ruan, M., Chen, S., & Xie, X-Y. (2025). Human-generative AI collaboration enhances task performance but undermines human's intrinsic motivation. *Scientific Reports*, 15, Article 15105. <https://doi.org/10.1038/s41598-025-98385-2>
- Xia, Q., Li, W., Yang, Y., Weng, X., & Chiu, T.K.F. (2025). A systematic review and meta-analysis of the effectiveness of generative artificial intelligence (GenAI) on students' motivation and engagement. *Computers and Education: Artificial Intelligence*, 9, Article 100455. <https://doi.org/10.1016/j.caeai.2025.100455>
- Zawacki-Richter, O., Marín, V.I., Bond, M., & Gouverneur, F. (2019). Systematic review of research on artificial intelligence applications in higher education—where are the educators? *International Journal of Educational Technology in Higher Education*, 16, Article 39. <https://doi.org/10.1186/s41239-019-0171-0>

Author Notes

Nicole D'Silva is an Assistant Professor of Data Science and Technology in Mumbai, India. She is also a PhD scholar in Management at Jamnalal Bajaj Institute of Management Studies, one of

Mumbai's leading business schools. She has two Master's degrees—an MSc in Statistics and an MMS in Finance. Her interdisciplinary research spans AI, human-AI interaction, data science, financial analytics, management education, and responsible technology adoption that connects analytical methods with contemporary organisational and societal challenges. Email: ndsilva321@gmail.com (<https://orcid.org/0009-0004-9239-7478>)

Cite as: D'Silva, N. (2026). Artificial intelligence and human multiple intelligences: Substitution, complementarity, and the future of managerial capability—A Scopus evidence map and integrative review. *Journal of Learning for Development*, 13(3), 441-454.

DOI: <https://doi.org/10.56059/06bjmb80>