



POLITECNICO
MILANO 1863

SCUOLA DI INGEGNERIA INDUSTRIALE
E DELL'INFORMAZIONE

EXECUTIVE SUMMARY OF THE THESIS

The Product Manager and the AI Buffer: A Multi-Level Framework Linking Cognition to Organizational Practice

LAUREA MAGISTRALE IN MANAGEMENT ENGINEERING - INGEGNERIA GESTIONALE

Author: **DUCCIO PROFETI, EDOARDO PINZAUTI**

Advisor: **PROF. SERGIO TERZI**

Co-advisor: **PROF. PATRICK RAU**

Academic year: **2024-2025**

1. Introduction

Generative AI systems are increasingly embedded in knowledge-intensive work, producing external representations—summaries, analyses, recommendations—that reshape how information is encoded, retrieved, and acted upon. Despite rapid adoption, research remains fragmented: cognitive science focuses on memory and attention mechanisms, while organizational research emphasizes productivity and role change. Yet generative AI challenges this separation, because changes in encoding and retrieval at the individual level can propagate to decision quality, accountability, and governance at the organizational level.

This thesis addresses these issues through a *multi-level perspective* linking cognitive mechanisms to organizational practice. At the cognitive level, **Study 1** uses a controlled reading-and-memory experiment ($N = 36$) in which AI assistance is operationalized as pre-generated summaries varying in *structure* (integrated vs. segmented) and *timing* (pre-reading, synchronous, post-reading). At the organizational level, **Study 2** examines how product managers ($N = 74$; 37 USA, 37 China) integrate generative AI tools into decision-intensive work-

flows, assessing perceived impacts on efficiency, judgment, and ethical responsibility [3].

Three research questions guide the work:

- **RQ1 (Cognitive):** How do AI-generated external representations influence core human memory processes—encoding, recall, recognition, source monitoring, and confidence calibration?
- **RQ2 (Organizational):** How does AI adoption shape decision-making practices, role configuration, and perceived responsibility in product management?
- **RQ3 (Integrative):** How can organizational changes enabled by AI be explained through underlying cognitive mechanisms of AI-assisted memory?

The thesis proposes the *AI Buffer* as a unifying concept: AI-generated external representations function as persistent buffers that support cue-based cognition while also introducing source-monitoring and reliance vulnerabilities.

2. The AI Buffer Model

The conceptual framework centers on the *AI Buffer Model*, which treats AI-generated outputs as external representations that interact with five cognitive mechanisms drawn from estab-

lished memory theory [1]: (i) *encoding scaffolding*—AI summaries can function as advance organizers that guide attention during subsequent reading; (ii) *cue-dependent retrieval*—persistent AI artifacts provide retrieval cues that support recognition but may not strengthen generative recall; (iii) *source monitoring*—users must discriminate between AI-generated and original content, a process vulnerable to representational fragmentation [4]; (iv) *confidence calibration*—alignment between subjective confidence and actual accuracy under AI assistance; and (v) *cognitive load redistribution*—AI may shift attentional resources rather than simply adding information.

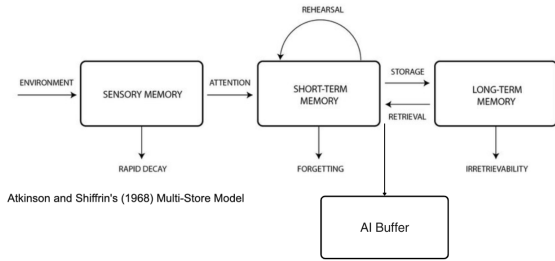


Figure 1: The AI Buffer Model: AI-generated representations interact with five cognitive mechanisms and propagate to organizational decision practices.

At the organizational level, the framework extends these mechanisms through the *AI-Augmented Decision Framework*, linking cognitive patterns to product-management practice: encoding scaffolds translate into decision framing, cue-dependent retrieval maps onto artifact-driven coordination, and source-monitoring vulnerabilities connect to governance and accountability structures.

3. Methodology

3.1. Study 1: Controlled Experiment

Study 1 employed a 2×3 mixed design with *summary structure* (integrated vs. segmented; between-subjects) and *summary timing* (pre-reading, synchronous, post-reading; within-subjects). Participants ($N = 36$; AI group: $n = 24$, No-AI control: $n = 12$) read three scientific articles and completed comprehension assessments including multiple-choice questions (MCQs), free recall, false-lure detection, and

confidence ratings. AI summaries were pre-generated to ensure experimental control. The experimental flow is shown in Figure 2.

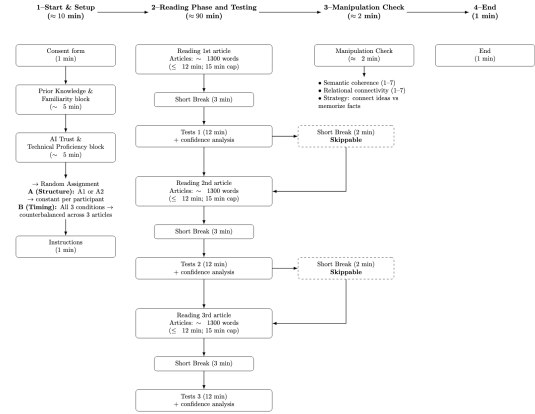


Figure 2: Study 1 experimental flow: each participant completed three reading blocks under different timing conditions.

3.2. Study 2: Cross-National Field Survey

Study 2 surveyed $N = 74$ practicing product managers (37 USA, 37 China) using a structured questionnaire covering AI tool usage, workload shifts, perceived strategic impact, decision-making practices, and governance/ethical responsibility. Likert-scale items (1–5) were tested against the neutral midpoint and compared across regions using independent-samples t -tests and chi-square tests.

4. Results

4.1. Study 1: Cognitive Mechanisms (RQ1)

Recognition improves with AI assistance. Across all MCQ items, the AI condition outperformed the No-AI condition (AI: $M = 0.598$, $SD = 0.085$; No-AI: $M = 0.510$, $SD = 0.098$), $F(1, 34) = 7.86$, $p = .008$, $\eta_p^2 = 0.188$, corresponding to an 8.8 percentage-point gain (Cohen’s $d \approx 0.98$).

Pre-reading timing yields the strongest gains. For AI-summary-sourced items, pre-reading access produced the highest accuracy ($M = 0.833$, $SD = 0.141$), compared with synchronous ($M = 0.568$, $SD = 0.239$) and post-reading ($M = 0.641$, $SD = 0.196$). The timing main effect was strong: $F(2, 44) = 14.00$, $p < .001$, $\eta_p^2 = 0.389$. No structure main effect or

interaction was observed. Holm-corrected post-hoc comparisons confirmed pre-reading superiority over synchronous ($p = .00052$, $d_z = 0.91$) and post-reading ($p = .00057$, $d_z = 0.87$).

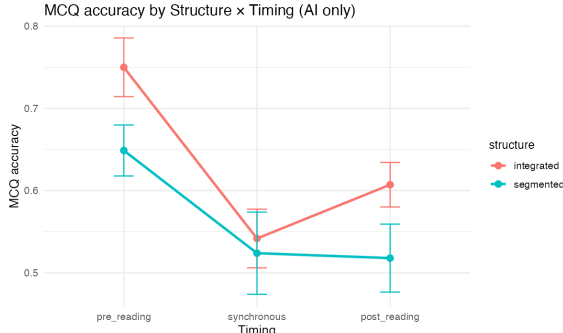


Figure 3: Overall MCQ accuracy by summary structure and timing (Study 1, AI group). Pre-reading consistently outperforms other conditions. Error bars: $\pm$ SE.

Free recall is unaffected. In contrast, free recall did not vary by timing (pre: $M = 5.50$; synchronous: $M = 5.54$; post: $M = 5.56$; $F(1.86, 40.88) = 0.03$, $p = .969$) and showed no AI vs. No-AI difference ($p > .50$). This *recognition-recall dissociation* indicates that AI assistance primarily strengthens cue-supported retrieval without building stronger internally generated memory traces.

Segmented structure increases false-lure endorsement. Summary structure significantly affected source monitoring: segmented summaries increased false-lure endorsement relative to integrated summaries (OR = 5.93, 95% CI [1.63, 21.5], $p = .007$). Participants with segmented summaries selected more false lures ($M = 1.06$ vs. $M = 0.58$) and showed lower false-lure accuracy ($M = 0.375$ vs. $M = 0.556$). This suggests that representational fragmentation weakens provenance cues and cross-checking opportunities [2].

Confidence calibration is imperfect but not worsened by AI. Recall confidence was similar across conditions (AI: $M = 4.25$; No-AI: $M = 4.08$; $t(31.7) = 0.16$, $p = .873$), indicating that the epistemic risks identified above are not driven by inflated overconfidence.

Process evidence: attention redistribution. Pre-reading produced the highest summary viewing time (132.5 s; 24.9% of total), while post-reading was lowest (69.5 s;

13.7%). Total time-on-task remained stable (≈ 531 –536 s), indicating a *redistribution* of attention rather than increased effort.

4.2. Study 2: Organizational Practice (RQ2)

Broad adoption with moderate positive impact. Product managers reported widespread AI tool usage, with generative writing assistants most frequently used ($M = 3.22$, $SD = 1.10$). Overall productivity impact was moderately positive ($M = 3.47$, $SD = 1.10$; $t(73) = 4.77$, $p < .001$), as was perceived decision-making improvement ($M = 3.45$, $SD = 1.05$; $t(73) = 3.34$, $p = .0013$) [5].

Workload reduction concentrates on routine tasks. The strongest perceived workload reductions occurred in routine administrative tasks ($M = 3.88$), customer feedback analysis ($M = 3.84$), and documentation/content creation ($M = 3.70$), all significantly above the neutral midpoint ($p < .001$). In contrast, routine communications showed a negative deviation ($M = 2.22$, $p < .001$).

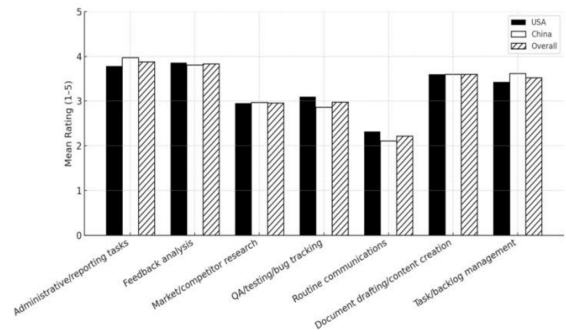


Figure 4: Perceived workload reduction by product-management task (Study 2). AI adoption most effectively reduces routine administrative and documentation workload. Error bars: $\pm$ SD.

Role reconfiguration toward monitoring. AI adoption shifted work toward higher-level activities: 85.1% of respondents reported reallocating saved time to product innovation, while 75.7% identified human-in-the-loop checks as the task that became more complex.

Governance gaps persist. Despite positive perceived impacts, substantial governance ambiguity was identified: 33.8% of respondents reported no clear ownership of AI-related deci-

sions, and 36.5% reported the absence of clear guidelines for AI use [6]. A notable USA–China difference emerged for perceived decision-making improvement ($p = .0104$), while productivity and expectations showed no regional differences.

5. Cross-Study Integration (RQ3)

The combined findings support three integration claims consistent with the AI Buffer Model:

Claim 1: “AI-first framing” is effective when it shapes encoding, not when it replaces it. Study 1’s pre-reading advantage shows that early exposure to AI summaries scaffolds subsequent encoding without increasing total time. Study 2 confirms that PMs derive the most value from AI when it is used to frame problems and structure attention early, rather than as a post-hoc justification tool.

Claim 2: Fragmentation increases misattribution risk, motivating governance mechanisms. Study 1 demonstrates that segmented presentations substantially increase false-lure endorsement ($OR = 5.93$), revealing a source-monitoring vulnerability. Study 2 shows that accountability and ethical governance are frequently unclear. The integration implication is that the *provenance* and *structure* of AI outputs are part of decision governance: integrated, traceable artifacts reduce misattribution and support responsibility assignment.

Claim 3: Reliance is not purely “overconfidence”; it is a shift in cognitive infrastructure. Study 1 finds that AI does not worsen overconfidence but changes reliance and attention allocation. Study 2 complements this by showing that AI adoption reallocates work toward strategic activities while introducing new monitoring responsibilities. Reliance operates through infrastructural shifts in who curates the external representation, when it is consulted, and how it is audited.

Table 1 maps cognitive mechanisms to organizational implications.

6. Conclusions

This thesis examined how AI-generated outputs function as external representations that reshape both individual cognition and organi-

zational decision-making. The results support a coherent account: AI assistance improves performance when it provides useful retrieval cues and early framing, but introduces scalable epistemic and governance risks when provenance cues are weak or verification responsibilities are unclear.

Key contributions. (i) *Conceptual*: the AI Buffer Model provides a design-oriented framework linking representation parameters (timing, structure, provenance) to cognitive and organizational outcomes. (ii) *Empirical*: Study 1 delivers causal evidence that pre-reading summaries improve recognition and that segmented structure increases false-lure endorsement; Study 2 documents convergent adoption patterns and governance gaps. (iii) *Methodological*: the multi-level design combining controlled experiment with field survey demonstrates a practical approach to studying human–AI interaction across levels of analysis.

Actionable takeaways. The findings translate into six implementation-oriented guidelines:

1. Use AI early for framing, not late for justification.
2. Prefer integrated representations over fragmented ones—fragmentation is an epistemic risk factor.
3. Treat AI artifacts as part of the decision record.
4. Separate drafting from validation: make verification steps explicit.
5. Assign clear ownership for checking AI-mediated evidence.
6. Measure success beyond speed: evaluate decision trace quality (auditability, provenance, rework).

Limitations and future research. Study 1 is scoped to controlled laboratory conditions with pre-generated summaries and within-session memory measures; ecological validity requires replication with interactive AI and longer retention intervals. Study 2 relies on self-report from a sample of 74 product managers; future work should combine behavioral traces with larger cross-cultural samples. The AI Buffer Model generates testable propositions that can guide longitudinal and intervention-based research linking individual cognition to organizational outcomes.

In summary, the central question is not whether AI “helps” or “harms” memory and decision-

Cognitive Mechanism	Study 1 Evidence	Organizational Implication (Study 2)
Encoding scaffold	Pre-reading yields highest recognition; timing drives accuracy	AI is most valuable when used early to frame options
Cue-dependent retrieval	Recognition improves; free recall unchanged	Decision practices become artifact-driven; verification is essential
Source monitoring	Segmented structure increases false-lure endorsement (OR = 5.93)	Fragmented AI usage increases misattribution; provenance controls needed
Confidence calibration	AI does not increase overconfidence	Confidence is an unreliable governance signal; explicit checks remain necessary
Load redistribution	Pre-reading increases summary engagement without increasing total time	Adoption shifts time from operational to strategic work, but also to monitoring

Table 1: Cross-study integration: linking cognitive mechanisms (Study 1) to product-management practices (Study 2).

making in the abstract. The outcomes depend on designable parameters of external representations and on the organizational routines that govern their use. When AI systems support early framing and preserve provenance, they can improve performance and accelerate knowledge work. When outputs are fragmented and accountability is unclear, the same systems can amplify misattribution and diffuse responsibility. Provenance and governance should therefore be treated as first-class design requirements for AI-assisted work.

tions are rewiring to capture value, 2025. Accessed 2025-06-04.

- [6] F. Pasquale. *The Black Box Society: The Secret Algorithms That Control Money and Information*. Harvard University Press, 2015.

References

- [1] A. D. Baddeley. Working memory: Theories, models, and controversies. *Annual Review of Psychology*, 63:1–29, 2012.
- [2] S. Chan, P. Pataranutaporn, A. Suri, W. Zulfikar, P. Maes, and E. F. Loftus. Conversational AI powered by large language models amplifies false memories in witness interviews. arXiv preprint arXiv:2408.04681, 2024.
- [3] P. Huryn. How to automate your work as a product manager, 2024. Accessed 2025-06-04.
- [4] Marcia K. Johnson, Shahin Hashtroudi, and D. Stephen Lindsay. Source monitoring. *Psychological Bulletin*, 114(1):3–28, 1993.
- [5] McKinsey. The state of AI: How organiza-