

Enacting Explainability: Knowledge Practices Around Generative AI in a Consulting Firm

Magnus Erga Skreden¹, Per Olav Svendsen², and Eli Hustad³

¹GlobalConnect, Norway

²Eramet, Norway

³University of Agder, Norway

pellesvendsen@hotmail.com

magnus.ser@live.no

eli.hustad@uia.no

Abstract: As generative AI becomes embedded in knowledge-intensive work, organizations face a central paradox: AI systems produce influential outputs while their inner workings remain opaque. Although explainability research has largely focused on technical methods, less attention has been given to how organizational actors practically make sense of and govern AI-generated knowledge. This study adopts a knowledge-practice perspective to examine how explainability is enacted in everyday work within a mid-sized consulting firm. The study is based on a qualitative case design involving nine semi-structured interviews with AI practitioners, clients, and an AI expert. Rather than treating explainability as a technical property, the analysis explores how consultants interpret, validate, and operationalize AI outputs in the absence of formal governance structures. The findings show that explainability is enacted as a situated and role-dependent knowledge practice. First, employees assess outputs pragmatically, focusing on whether they “make sense” in context rather than seeking insight into model internals. Second, prompting emerges as a form of tacit expertise developed through trial-and-error learning, functioning as an informal mechanism for influencing and interpreting outputs. Third, actors rely on local validation routines such as re-running prompts and cross-checking information to manage uncertainty and maintain trust. Finally, the absence of internal guidelines results in individualized and uneven practices, indicating that knowledge about AI use is created but not institutionalized. From a Knowledge Management (KM) perspective, the study contributes by conceptualizing explainability as an emergent knowledge practice shaped by situated sensemaking, tacit skill development, and informal governance. It extends knowing-in-practice research by illustrating how organizations cope with opaque digital systems when formal knowledge infrastructures lag behind technological adoption. For KM practice, the findings highlight the need for lightweight governance mechanisms, shared prompting norms and role-adapted guidelines that support collective knowledge development around AI use.

Keywords: Knowledge practices, Organizational knowing, Generative AI, Explainability, Knowledge governance

1. Introduction

Generative artificial intelligence (AI) is rapidly becoming embedded in knowledge-intensive work, transforming how organizations produce, interpret, and apply knowledge. In domains such as consulting, employees increasingly rely on AI systems to generate analyses, summaries, and recommendations that directly influence professional outputs. These systems are increasingly embedded in organizational processes, particularly in knowledge-intensive sectors such as consulting (Lari and Manu 2024). However, organizations face significant challenges in implementing and integrating AI into existing practices, including issues related to governance, trust, and organizational readiness (Raftopoulos 2024). This growing reliance introduces a fundamental tension: while AI systems can produce highly useful and persuasive outputs, their internal workings are often difficult to interpret for end users. This creates challenges for understanding, validating, and justifying AI-generated knowledge in organizational contexts (Miller 2019; Rai 2020).

Existing research on explainability has primarily approached this issue from a technical perspective, focusing on methods that make machine learning models more transparent or interpretable (e.g., Lundberg and Lee 2017; Ribeiro et al. 2016). While such approaches are important, they often assume that explainability is a property that can be engineered into the system itself. In practice, however, many widely used AI tools, particularly generative AI systems, remain effectively “black-boxed” from the perspective of everyday users. As a result, the challenge of explainability is not resolved solely through technical design but emerges as a practical concern in how users make sense of and work with AI outputs (Burrell 2016).

From a Knowledge Management (KM) perspective, this raises important but underexplored questions. KM research has long emphasized that knowledge is not only stored in systems, but enacted through practice, shaped by context, and embedded in social interaction (Alavi and Leidner 2001). More recent work highlights how digital technologies, including AI, are reshaping knowledge processes by augmenting human capabilities while also introducing new forms of uncertainty and complexity. For instance, Jarrahi et al. (2023) argue that AI

should be understood as a partner in knowledge work, requiring new forms of human–AI collaboration and sensemaking. Similarly, Alavi et al. (2024) emphasize that generative AI challenges traditional KM assumptions by altering how knowledge is created, validated, and governed.

Despite these advances, little is known about how explainability is enacted in everyday organizational practice. While AI research often focuses on improving model transparency, and KM research examines knowledge processes more broadly, limited attention has been given to how organizational actors actually produce “explainability” when working with AI systems whose inner logic is not accessible. We lack insight into how users interpret AI outputs, develop trust, and establish defensibility in contexts where formal explanations are limited or absent.

This study addresses this gap by adopting a knowledge-practice perspective to examine how explainability is enacted in AI-enabled knowledge work. Rather than treating explainability as a technical feature, the study conceptualizes it as a situated and socially embedded practice through which actors render AI outputs understandable, usable, and justifiable. This perspective aligns with knowing-in-practice approaches in KM, which view knowledge as something that is performed, negotiated, and situated in action rather than simply transferred or stored (Nicolini et al. 2003; Orlikowski 2002).

We conducted a qualitative case study of a mid-sized consulting firm where generative AI tools are increasingly integrated into daily work processes. Through semi-structured interviews with employees, clients, and an AI expert, we explore how actors interpret, validate, and operationalize AI-generated outputs in the absence of formal governance structures and standardized practices.

The following research question guided the study: *How do organizational actors enact explainability as a knowledge practice when working with opaque AI systems?*

By addressing this question, the paper makes two primary contributions. First, it contributes to KM theory by conceptualizing explainability as a knowledge practice shaped by situated sensemaking, tacit skill development, and informal governance. Second, it provides empirical insight into how organizations cope with the challenges of AI-enabled knowledge work when technological capabilities outpace formal knowledge infrastructures. In doing so, the paper bridges AI explainability research and KM scholarship, highlighting the importance of organizing knowledge practices alongside technical systems for responsible and effective AI use.

2. Theoretical Framing: Knowledge Practices and Explainability in AI-enabled Work

2.1 Knowledge as Practice in Organizations

KM research has increasingly moved beyond viewing knowledge as an object that can be stored, transferred, or codified, toward understanding knowledge as something that is enacted in practice. Early KM work emphasized the role of information systems in capturing and distributing knowledge (Alavi and Leidner 2001). At the same time, a complementary stream of research has emphasized that knowledge is inherently situated, socially embedded, and inseparable from action (Brown and Duguid 2001; Polanyi 1962).

The knowing-in-practice perspective conceptualizes knowledge not as a static resource but as an ongoing accomplishment emerging through everyday work activities (Orlikowski 2002). From this perspective, knowing is enacted through situated action, shaped by context, experience, and interaction among actors. Similarly, practice-based approaches emphasize that knowledge is developed and expressed through routines, tools, and embodied skills within specific organizational settings (Nicolini 2011). These perspectives shift attention from formal knowledge structures to the micro-level practices through which knowledge is produced, interpreted, and validated. In uncertain and complex settings, actors rely on tacit knowledge, judgment, and situated sensemaking rather than formal knowledge representations. Knowledge practices, therefore, include not only formalized procedures but also informal routines, heuristics, and workarounds that enable actors to navigate ambiguous situations. This perspective is especially relevant in complex and distributed work settings, where knowledge must be coordinated across multiple actors and teams (Dingsøyr et al. 2018).

2.2 AI and the Transformation of Knowledge Work

The increasing integration of AI into organizational processes is reshaping how knowledge is created and used. AI systems can generate content, synthesize information, and support decision-making, thereby augmenting human capabilities in knowledge-intensive work. However, this transformation also introduces new challenges for KM.

Recent research suggests that AI should be understood as a collaborator in knowledge work rather than merely a tool. Jarrahi et al. (2023) argue that effective use of AI requires new forms of human–AI partnership, where users interpret and contextualize AI outputs rather than passively receiving them. In this sense, AI systems participate in knowledge processes, but do not replace the need for human judgment and interpretation.

In addition, generative AI challenges traditional assumptions about knowledge creation and validation. Alavi et al. (2024) highlight that AI-generated outputs blur the boundaries between knowledge creation and knowledge retrieval, raising questions about authorship, reliability, and governance.

As AI systems become more capable, the locus of knowledge shifts to distributed human–AI configurations, creating new challenges for knowledge evaluation and trust. Users may need to act on AI-generated outputs without fully understanding how they were produced or justified.

2.3 Explainability Beyond Technical Transparency

The challenge of understanding AI systems has been widely discussed in the literature on explainable artificial intelligence (XAI). Much of this research focuses on developing methods that make machine learning models more interpretable, such as feature attribution techniques and model-agnostic explanations (Lundberg and Lee 2017; Ribeiro et al. 2016). More broadly, XAI aims to increase the transparency and interpretability of AI systems (Gerlings et al. 2021). While these approaches contribute to technical transparency, they do not fully address how explanations are used and interpreted in organizational contexts.

Recent work on generative AI suggests that explainability challenges are further amplified in these systems due to their probabilistic and generative nature, requiring new conceptual approaches to understanding AI outputs (Schneider 2024). In such contexts, users often interact with AI systems whose internal reasoning remains inaccessible, placing greater emphasis on how outputs are interpreted and evaluated in practice.

Scholars have increasingly emphasized that explainability is not only a technical issue, but also a social and cognitive one. Miller (2019) argues that explanations must be meaningful to human users, taking into account how people understand and evaluate information. Similarly, Rai (2020) highlights that explainability should be considered in relation to user needs, decision contexts, and organizational requirements. From this perspective, explainability can be understood as a process rather than a property. Burrell (2016) distinguishes between different forms of opacity in machine learning systems, noting that even when technical transparency is improved, systems may remain difficult to interpret for non-experts. This suggests that explainability cannot be fully “designed into” AI systems but must also be constructed through interaction and use. This perspective opens up a practice-based understanding of explainability, where explanations are enacted through interaction rather than solely embedded in technical systems.

2.4 Toward a Knowledge-practice Perspective on Explainability

Despite this growing recognition of the human-centered nature of explainability, there remains limited integration between XAI research and KM theory. In particular, little attention has been paid to how explainability is enacted as part of everyday knowledge practices in organizations.

Drawing on a knowledge-practice perspective, this study conceptualizes explainability as an emergent organizational accomplishment. Rather than being located within the AI system itself, explainability is produced through the practices by which actors interpret, validate, and communicate AI-generated outputs. This includes the use of tacit skills, such as prompting, as well as informal routines for verifying and contextualizing results.

This perspective highlights three important aspects. First, explainability is situated: it depends on the specific context in which AI outputs are used and evaluated. Second, it is relational: different actors require different forms of explanation depending on their roles and responsibilities. Third, it is organizational: the extent to which explainability practices are shared, standardized, or institutionalized depends on the presence of governance structures and collective norms.

By framing explainability as a knowledge practice, the study bridges KM and AI research, shifting the focus from technical transparency to the organizational processes through which AI-generated knowledge becomes meaningful and actionable. This provides a foundation for analyzing how organizations cope with the challenges of AI-enabled work when technological capabilities outpace formal knowledge infrastructures.

3. Research Context

This study is based on a qualitative case study conducted in a mid-sized Norwegian consulting firm operating in digital transformation, data analytics, and technology advisory services. The company has actively adopted generative AI tools in both internal and client-facing work, making it a relevant setting for examining AI-enabled knowledge practices. Consultants use AI for tasks such as drafting reports, generating analyses, summarizing information, and supporting problem-solving activities, with outputs often contributing directly to client deliverables.

At the time of the study, AI adoption was widespread but largely driven by individual experimentation rather than formal organizational initiatives. As a result, employees displayed varying levels of AI use and expertise, ranging from frequent application in analytical and advisory work to more selective use for information retrieval, drafting, and idea generation. The organization had not yet established formal guidelines or standardized procedures for AI use, and employees relied largely on individual experience and informal practices. This provided a suitable setting for exploring how explainability is enacted in contexts where AI systems are widely used, yet formal knowledge governance structures remain under development.

4. Research Method

The study adopts a qualitative case study approach, which is well suited for exploring knowledge practices in organizational settings (Yin 2018). Data were collected through nine semi-structured interviews with practitioners from a consulting firm, two client representatives, and an AI expert. Participants had varying levels of experience with generative AI, enabling insights across roles and contexts (Table 1).

Interviews focused on participants' experiences of using AI in their work, including how they interpreted, validated, and communicated AI-generated outputs. Semi-structured interviews enabled participants to elaborate on their experiences while ensuring comparability across interviews (Kvale and Brinkmann 2009; Myers and Newman 2007). Interviews were conducted digitally, recorded, and transcribed. Limited follow-up questions were conducted via email to clarify selected findings.

The data were analyzed using an iterative thematic analysis informed by abductive reasoning (Langley 1999). Initial coding focused on how participants described understanding, validating, and explaining AI outputs. Codes were grouped into broader themes through comparison and categorization (Miles et al. 2018). Guided by a knowledge-practice perspective, the analysis identified five themes: explainability as situated judgment, prompting as tacit skill, local validation practices, role-dependent knowing, and absence of institutionalization.

As a single-case study, the findings are intended for analytical rather than statistical generalization (Yin 2018). The case provides insight into how explainability is enacted as a knowledge practice in AI-enabled work.

Table 1: Overview of interview participants in the study

Participant Group	N
AI Expert	1
Clients	2
Practitioners	6

5. Findings

5.1 Explainability as Situated Judgment

Across participants, explainability was rarely described in technical terms. Instead, it was understood as a practical judgment about whether AI-generated outputs were meaningful and appropriate for the task at hand. Rather than seeking insight into how models function internally, participants emphasized their ability to assess outputs based on experience and contextual relevance.

As one participant explained:

"Explainability, for me, is about whether the answer makes sense in the context I'm working in."

Similarly, another noted:

"I don't really need to understand the model. I just need to be confident that what it gives me is usable."

These accounts suggest that explainability is enacted through situated sensemaking. Participants draw on prior knowledge and contextual understanding to evaluate outputs, shifting attention from algorithmic transparency to the plausibility and usefulness of results.

5.2 Prompting as Tacit Skill

A central aspect of working with AI systems was the ability to formulate effective prompts. Participants consistently described prompting as something learned through experience rather than formal training. This competence was developed iteratively through trial and error and refined over time.

One participant described this process: “You learn by trying different things. Over time you figure out how to ask in a way that gives better answers.”

Another emphasized the importance of formulation:

“It really depends on how you phrase it. If the prompt is unclear, the output is unclear.”

Prompting thus functions as a form of tacit knowledge. It is not easily articulated or standardized, but embedded in practice and individual experience. This skill becomes particularly important in contexts where AI systems are opaque. Instead of accessing explanations from the system, users shape and interpret outputs through their ability to interact effectively with it. At the same time, the development of prompting competence appears uneven across participants. Some reported a high degree of confidence and experimentation, while others relied on more basic usage patterns. This suggests that explainability is partly dependent on individual capability, reinforcing its practice-based nature.

5.3 Local Validation Practices

In the absence of formal explanation mechanisms, participants described a range of informal strategies for assessing the reliability of AI-generated outputs. These practices included re-running prompts, comparing outputs across tools, and manually verifying information.

One participant explained:

“If I’m unsure, I’ll run the same prompt again or tweak it slightly to see if I get the same result.”

Another highlighted the need for caution:

“Sometimes it sounds very convincing, but when you look closer, it’s not correct.”

These practices constitute local validation routines. Rather than relying on system-provided explanations, participants actively test and verify outputs through iterative interaction. This involves comparing results, checking consistency, and applying professional judgment. Such practices represent a form of epistemic work, where users take responsibility for establishing the credibility of AI-generated knowledge. Explainability is therefore not only about interpretation, but also about validation. It emerges through a combination of interaction with the system and external verification processes.

5.4 Role-dependent Knowing

The findings also indicate that explainability is not uniform across participants but varies depending on role and context. Different actors require different forms and levels of explanation, reflecting their responsibilities and use of AI outputs.

As one participant noted:

“For some people, it’s enough to get a short answer, while others need more detail to trust it.”

Another emphasized the difference between technical and non-technical perspectives:

“If you’re working closely with the technology, you might want more insight. But for others, it’s more about whether it works.”

This highlights that explainability is relational and audience-dependent. What counts as a sufficient explanation depends on who the user is, what they need to accomplish, and how the output will be used. For consultants, explainability is often tied to the ability to justify outputs to clients, while clients themselves may focus on clarity and usability rather than technical detail. Explainability is therefore not a fixed property, but negotiated across different roles and contexts.

5.5 Absence of Institutionalization

Despite the widespread use of AI tools, participants reported a lack of formal organizational structures guiding their use. There were no standardized procedures, shared guidelines, or formal training related to explainability.

One participant stated:

"We don't really have any guidelines. People just use it in their own way."

The absence of institutionalization results in individualized and uneven explainability practices. While this allows flexibility and experimentation, it also creates inconsistencies in how AI outputs are interpreted and validated. From a KM perspective, knowledge about AI use is generated but not systematically shared or embedded in organizational practices. Taken together, the findings suggest that explainability is enacted through interconnected knowledge practices rather than embedded in AI systems. The following section develops a framework to conceptualize how these practices operate across individual, relational, and organizational levels.

6. Discussion

6.1 Explainability as Knowing-in-practice

This study examined how explainability is enacted in AI-enabled knowledge work. The findings suggest that explainability is not primarily embedded within AI systems but emerges through situated practices of interpretation, interaction, and validation. This extends explainable AI research, which has largely focused on technical mechanisms for increasing transparency (Lundberg and Lee 2017; Ribeiro et al. 2016), by showing how explainability is realized in organizational settings.

The findings align closely with knowing-in-practice perspectives in KM (Orlikowski 2002). Explainability is achieved not through access to model internals, but through the practices by which users interpret outputs, refine prompts, and validate results. Prompting emerges as a form of tacit knowledge, while local validation routines illustrate how credibility is established through practice rather than formal verification systems.

Rather than replacing human knowing, AI introduces new forms of uncertainty that require interpretation and judgment. This reflects broader KM insights into the importance of experiential and embodied knowledge in complex work settings (Nicolini 2011).

6.2 A Practice-Based Framework of Enacted Explainability

Building on these insights, this study proposes a *practice-based framework of enacted explainability* (see Figure 1). The framework conceptualizes explainability as an organizational accomplishment emerging through interconnected knowledge practices. It highlights how explainability is enacted across three levels: micro, relational, and organizational, each shaping how AI outputs are interpreted, validated, and applied in practice.

The figure illustrates how explainability in AI-enabled work emerges across three interconnected levels: situated knowledge practices at the micro level, role-dependent knowing at the relational level, and knowledge governance at the organizational level. When AI systems are opaque, explainability is produced through the interaction of these levels rather than embedded within the system itself.

At the micro level, explainability is enacted through situated knowledge practices, including contextual judgment, tacit prompting competence, and local validation routines. At the relational level, explainability is shaped by role-dependent knowing, as different actors require different forms of explanation depending on their responsibilities and contexts of use. At the organizational level, explainability is influenced by governance structures that shape how practices are shared, coordinated, and institutionalized. Together, these levels illustrate that explainability is not located within the AI system itself but distributed across a network of practices, roles, and structures.

The framework proposes that when AI systems are opaque, explainability shifts from being a system-level property to becoming a distributed organizational accomplishment. The three levels should not be understood as independent. Situated practices are shaped by role-specific expectations, while both are influenced by the presence or absence of governance structures. At the same time, repeated practices may become shared routines that gradually influence organizational norms and governance arrangements. Explainability therefore emerges through ongoing interactions between individual action, social context, and organizational structures. By conceptualizing explainability in this way, the framework extends knowing-in-practice perspectives by showing how new forms of knowing emerge in response to technologically mediated uncertainty.

Enacted Explainability in AI-Enabled Knowledge Work

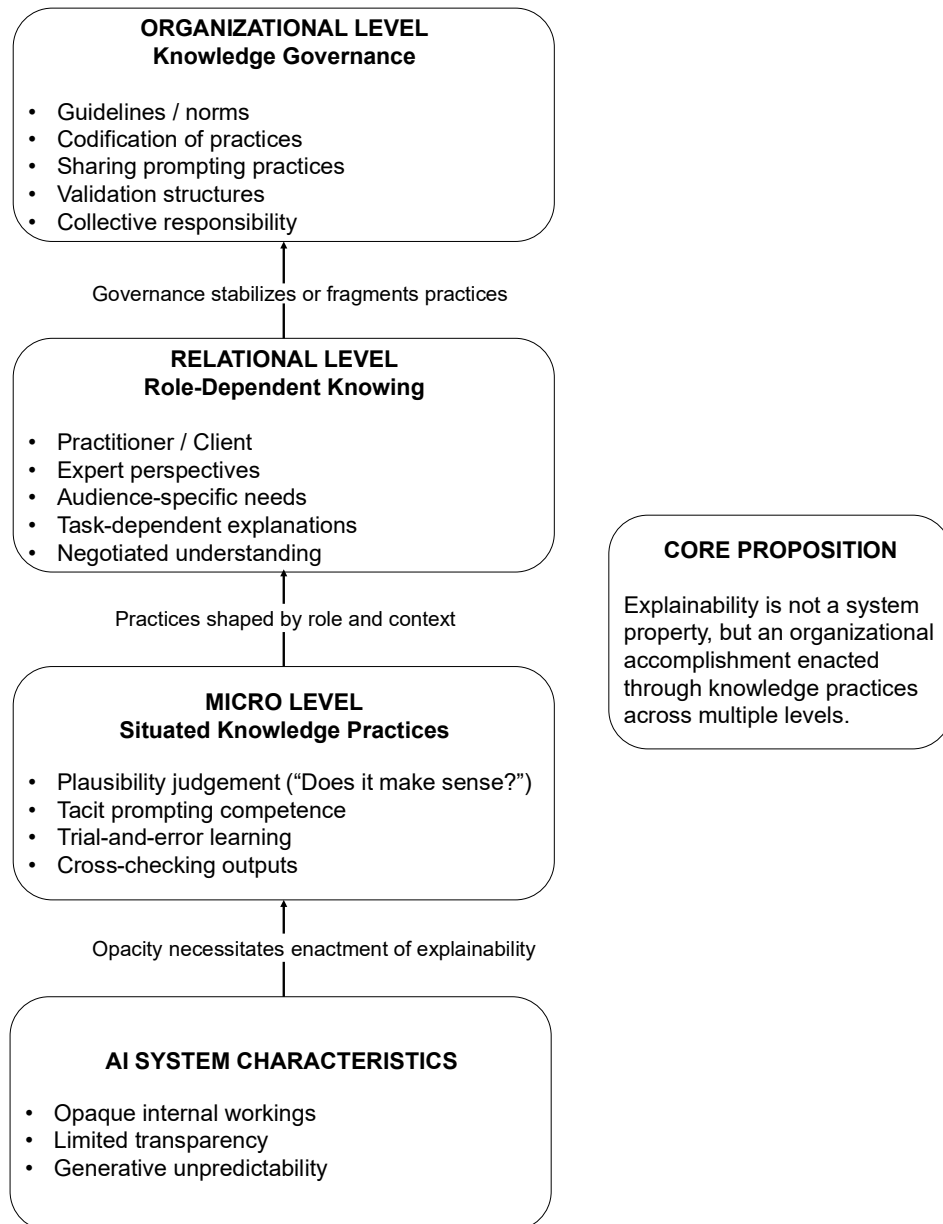


Figure 1: A Practice-Based Framework of Enacted Explainability in AI-Enabled Knowledge Work.

6.3 Implications for KM Theory

The study contributes to KM theory in three ways. First, it extends knowledge-practice perspectives by introducing explainability as a distinct form of knowing-in-practice. While prior research has emphasized how knowledge is enacted in organizational contexts, this study shows how explainability becomes an ongoing accomplishment when working with opaque AI systems. Second, the findings demonstrate the continuing importance of tacit knowledge in AI-enabled work. Prompting emerges as a competence developed through experience rather than formal training, suggesting that AI adoption reconfigures rather than replaces tacit knowledge. Third, the study highlights the role of knowledge governance in shaping AI use. The absence of shared norms and guidelines results in fragmented practices, underscoring the need for KM mechanisms that support the coordination and institutionalization of AI-related knowledge.

This finding aligns with recent research on responsible AI, which highlights the challenges of translating high-level principles into organizational practice (Akbarighatar et al. 2025). More broadly, it underscores the

importance of coordinating knowledge practices across organizational actors in complex knowledge work environments (Dingsøy et al. 2018).

While situated judgment, tacit prompting competence, and local validation routines enable productive use of opaque AI systems, they also introduce risks. When explainability relies heavily on individual expertise and informal practices, organizations may experience inconsistencies in how AI outputs are interpreted, validated, and communicated. In consulting environments, such variability may create challenges related to accountability, quality assurance, and organizational learning. This suggests that explainability should be viewed not only as an individual capability but also as an organizational capability requiring coordination and governance.

These contributions align with recent calls to integrate AI more explicitly into KM research (Alavi et al. 2024; Jarrahi et al. 2023). In doing so, the study extends knowing-in-practice perspectives by conceptualizing explainability as a distinct form of knowledge work in AI-enabled contexts.

6.4 Implications for Practice

The findings also have important implications for organizations seeking to manage AI-enabled work. First, organizations should recognize that explainability cannot be addressed solely through technical solutions. Even when explainability tools are available, users must still interpret and apply outputs in context. Supporting explainability therefore requires attention to how employees work with AI in practice. Second, the study highlights the importance of developing shared knowledge practices. This includes establishing guidelines for prompting, validation, and communication of AI outputs. Such mechanisms can help reduce variability and improve the consistency of AI use across the organization. Third, organizations should consider how explainability needs differ across roles. Rather than adopting one-size-fits-all solutions, organizations should develop role-adapted approaches that reflect different requirements for understanding and justification. Overall, managing AI in knowledge-intensive work is not only a technical challenge, but also a KM challenge. Organizations must develop structures and practices that enable users to work effectively with AI-generated knowledge.

7. Conclusion, Limitations and Future Research

This study has examined how explainability is enacted in AI-enabled knowledge work, shifting the focus from technical transparency to everyday organizational practice. Drawing on a qualitative case study, the findings show that explainability is not primarily achieved through access to model internals, but through situated knowledge practices such as contextual judgment, tacit prompting, and local validation routines. These practices are shaped by role-dependent needs and remain unevenly developed in the absence of formal governance structures. By conceptualizing explainability as a knowledge practice, the study contributes to KM research by extending knowing-in-practice perspectives into the context of generative AI. The proposed framework highlights that explainability emerges across multiple levels, micro, relational, and organizational, demonstrating that managing AI is not only a technical challenge, but also a matter of organizing how knowledge is produced, interpreted, and validated in practice. For organizations, the findings underscore the importance of supporting shared knowledge practices around AI use. Rather than relying solely on technical solutions, organizations should develop lightweight governance mechanisms, promote collective learning, and adapt explainability practices to different roles and contexts. As AI continues to evolve, future research should further explore how knowledge practices and governance structures co-develop with emerging technologies. Understanding this interplay will be essential for enabling responsible and effective use of AI in knowledge-intensive work.

This study is based on a single-case design, which limits the generalizability of the findings. Future research could explore similar dynamics in other organizational contexts or industries to assess the broader applicability of the framework. In addition, the study focuses primarily on user practices, rather than the design of AI systems themselves. Future research could examine how technical and organizational approaches to explainability interact, particularly in contexts where explainability tools are more actively integrated into workflows.

Ethical Declaration: Ethical clearance was not required for this research study.

AI declaration: Generative AI, ChatGPT (OpenAI), was used to assist with proofreading of this paper. The final text were critically reviewed and validated by the authors, who take full responsibility for the content of this publication.

References

- Akbarighatar, P., Pappas, I. O., Purao, S., and Vassilakopoulou, P. 2025. "Enacting Responsible AI: A Configurational Analysis of AI Principles in Practice," *Information Systems Frontiers* (28), pp. 1-27.
- Alavi, M., Leidner, D., and Mousavi, R. 2024. "A Knowledge Management Perspective of Generative Artificial Intelligence (GenAI)," *Journal of the Association of Information Systems* (25:1), pp. 1-12.
- Alavi, M., and Leidner, D. E. 2001. "Review: Knowledge Management and Knowledge Management Systems: Conceptual Foundations and Research Issues," *MIS Quarterly* (25:1), pp. 107-136.
- Brown, J. S., and Duguid, P. 2001. "Knowledge and Organization: A Social-Practice Perspective," *Organization Science* (12:2), pp. 198-213.
- Burrell, J. 2016. "How the Machine 'Thinks': Understanding Opacity in Machine Learning Algorithms," *Big data & society* (January - June), pp. 1-12.
- Dingsøyr, T., Moe, N. B., and Seim, E. A. 2018. "Coordinating Knowledge Work in Multiteam Programs: Findings from a Large-Scale Agile Development Program," *Project Management Journal* (49:6), pp. 64-77.
- Gerlings, J., Shollo, A., and Constantiou, I. 2021. "Reviewing the Need for Explainable Artificial Intelligence (XAI)," in *Proceedings of the 54th Hawaii International Conference on System Sciences*, Pp. 1284-1293.
- Jarrahi, M. H., Askay, D., Eshraghi, A., and Smith, P. 2023. "Artificial Intelligence and Knowledge Management: A Partnership between Human and Ai," *Business horizons* (66:1), pp. 87-99.
- Kvale, S., and Brinkmann, S. 2009. *Interviews: Learning the Craft of Qualitative Research Interviewing*. Los Angeles, Calif.: Sage.
- Langley, A. 1999. "Strategies for Theorizing from Process Data," *Academy of Management review* (24:4), pp. 691-710.
- Lari, H. A., and Manu, K. 2024. "AI in Consulting Services: Applications, Challenges, and Ethical Insights," *Indian Journal of Computer Science* (9:6), pp. 34-53.
- Lundberg, S. M., and Lee, S.-I. 2017. "A Unified Approach to Interpreting Model Predictions," in *31st Conference on Neural Information Processing Systems (NIPS 2017)*, Long Beach, Ca, USA.
- Miles, M. B., Huberman, A. M., and Saldaña, J. 2018. *Qualitative Data Analysis: A Methods Sourcebook*. Sage publications.
- Miller, T. 2019. "Explanation in Artificial Intelligence: Insights from the Social Sciences," *Artificial Intelligence* (267:February), pp. 1-38.
- Myers, M. D., and Newman, M. 2007. "The Qualitative Interview in Is Research: Examining the Craft," *Information and Organization* (17:1), pp. 2-26.
- Nicolini, D. 2011. "Practice as the Site of Knowing: Insights from the Field of Telemedicine," *Organization science* (22:3), pp. 602-620.
- Nicolini, D., Gherardi, S., and Yanow, D. 2003. *Knowing in Organizations: A Practice-Based Approach*. New York: M. E. Sharpe.
- Orlikowski, W. J. 2002. "Knowing in Practice: Enacting a Collective Capability in Distributed Organizing," *Organization Science* (13:3), pp. 249-273.
- Polanyi, M. 1962. *Personal Knowledge: Towards a Post-Critical Philosophy*, (1974 ed.). Chicago: The University of Chicago Press.
- Raftopoulos, M. 2024. "Organizational Challenges in Adoption and Implementation of Artificial Intelligence," in *Proceedings of the 57th Hawaii International Conference on System Sciences*. pp. 5786-5795.
- Rai, A. 2020. "Explainable AI: From Black Box to Glass Box," *Journal of the academy of marketing science* (48:1), pp. 137-141.
- Ribeiro, M. T., Singh, S., and Guestrin, C. 2016. "" Why Should I Trust You?" Explaining the Predictions of Any Classifier," *Proceedings of the 22nd ACM SIGKDD international conference on knowledge discovery and data mining*, pp. 1135-1144.
- Schneider, J. 2024. "Explainable Generative AI (GenxAI): A Survey, Conceptualization, and Research Agenda," *Artificial Intelligence Review* (57:11), p. 289.
- Yin, R. K. 2018. *Case Study Research and Applications: Design and Methods*, (Sixth edition ed.). Los Angeles: SAGE.