

THE PERCEIVED FIT–MISFIT TENSION OF ARTIFICIAL INTELLIGENCE: (MIS)CALIBRATING RELIANCE WHEN USING CONTEMPORARY AI IN STRATEGY PROCESSES

Franziskus Perkhofer¹⁾

Technical University Munich - School of Management

Adrian Häußner

Technical University Munich - School of Management

Forthcoming in *Strategic Management Review*

ABSTRACT

Generative AI is rapidly entering strategy work, yet the same tools that accelerate orientation and idea generation can undermine the credibility of strategic judgment. In strategizing, this tension is structural: because strategic decisions are interdependent, committing, and competitively constitutive, scanning and interpretation open decision sequences in which sensemaking errors compound rather than remain locally correctable. How strategists evaluate AI inside these early, binding tasks—and how those evaluations shape their reliance on it—remains poorly understood. Drawing on a two-phase qualitative study comprising 54 semi-structured interviews with AI-experienced strategists, we show that strategists experience AI as a simultaneous fit and misfit within scanning and interpretation, and that these appraisals shift across the handoffs through which AI outputs travel toward strategic commitment. We theorize reliance calibration as the mechanism linking these appraisals to outcomes: under enabling conditions, strategists keep reliance within defensible bounds; when those conditions are weak, they overrely or underutilize, and the consequences propagate through the decision sequence. The study contributes a theory of how AI affects strategic judgment where decisions guide and constrain other decisions, extending research on strategic sensemaking, task-technology fit, and appropriate reliance on AI.

Keywords: Artificial Intelligence, Fit–Misfit, Human-AI Collaboration, Reliance, Sensemaking, Strategy Process

Acknowledgments

We are deeply grateful to Co-Editor Michael Leiblein and the anonymous reviewer for their exceptionally constructive guidance, incisive feedback, and generous developmental support, which profoundly sharpened the theoretical contribution of this work.

AI Disclosure Statement

We acknowledge the use of AI-based language assistance tools (Grammarly, Trinko, and ChatGPT) to support writing clarity, grammar, and phrasing during the preparation of this manuscript.

1) Corresponding author.

1. Introduction

Strategists today routinely use artificial intelligence to scan their competitive environment and interpret what they find. This is unremarkable. What is remarkable—and theoretically unresolved—is that they do so within a decision architecture that makes errors in how they make sense of their environment structurally consequential. Strategic decisions are not isolated events: their value depends on how choices cohere with choices made elsewhere in the firm, by rivals, and at earlier points in time that now constrain what remains available (Leiblein et al., 2018; Porter, 1996; Rivkin, 2000). A sensemaking error at either stage does more than produce a poor judgment in the moment. It enters a binding sequence in which it shapes what gets committed, and what gets committed next: cognitive representations formed under distorted scanning constrain what capabilities and choices remain available, binding the firm to a competitive position it can no longer freely exit (Ghemawat, 1991; Sydow et al., 2009; Tripsas & Gavetti, 2000). What makes this consequential for AI is not AI's power but its position: it enters at the opening of a sequence the rest of which it shapes.

The sequence runs through the sensemaking process, and this is precisely why position matters: Environmental *scanning* outputs become the evidential foundation for its *interpretation*; interpretation becomes the cognitive basis for resource commitment in action; commitment enacts the competitive position the firm must subsequently defend (Daft & Weick, 1984; Thomas et al., 1993). What AI shapes at the scanning stage is therefore meaning rather than information: the interpretive premises on which strategic action is authorized. Research on strategic sensemaking has established this much: that scanning and interpretation are the core activities through which strategists notice competitive cues, construct meaning, and prepare action (Daft & Weick, 1984; Weick, 1995), and that this work is oriented prospectively—toward a future the firm has not yet

entered—rather than confined to reconstructing what has already occurred (Gephart et al., 2010; Kaplan & Orlikowski, 2013; Sakellariou & Vecchiato, 2022). What this literature has not yet addressed is what happens when an agentic AI—one that generates, evaluates, and begins to execute strategic options (Camuffo et al., 2026; Murray et al., 2021)—operates inside this sequence, shaping the interpretive premises on which the rest depends.

This matters because AI's capability is structurally misaligned with the demands of this work. Large language models demonstrate performance approaching expert benchmarks on tasks including information synthesis, option generation, and strategic evaluation (Csaszar et al., 2024; Doshi et al., 2025). But this capability is jagged (Dell'Acqua et al., 2023): the tasks at the core of strategic sensemaking—novel competitive contexts requiring theory-based causal inference, interpretations grounded in specific organizational history and competitive position, assessments made under genuine uncertainty—lie systematically outside AI's reliable zone (Felin & Holweg, 2024). In these settings, AI outputs are not just uncertain; they can be confidently wrong in contextually fragile ways, difficult to detect precisely where decisions are most committing and verification is most costly (Hannigan et al., 2024). Strategists are thus caught in a structural tension: AI is most attractive as a cognitive partner in exactly the sensemaking tasks where its outputs carry the highest epistemic risk—and where the commitments those outputs will inform are costliest to reverse. We therefore ask: *How do strategists calibrate their reliance on AI during environmental scanning and interpretation—the early, binding stages of this sequence—and with what outcomes?*

Existing frameworks that bear on this question reach it only partially. The appropriate reliance literature has established how people track AI accuracy and adjust reliance in single-shot, consequence-bounded decisions (Lee & See, 2004; Schöffner et al., 2025)—but this framing does

not travel to settings where miscalibration generates compounding downstream consequences across interdependent choices. Task-technology fit theory has mapped how alignment between task demands and tool capabilities shapes performance (Goodhue & Thompson, 1995; Strong & Volkoff, 2010), yet it stops short of engaging how interdependence across a decision sequence alters the stakes of fit–misfit appraisals at each stage. The growing literature on AI in strategy has done much to establish what AI can do strategically (Csaszar et al., 2024; Doshi et al., 2025; Krakowski et al., 2023), but it asks whether AI can perform strategic tasks, not how strategists calibrate reliance when those tasks are embedded in sequences where a local error becomes a structural one.

We address this question through a two-phase qualitative inquiry based on 54 semi-structured interviews with AI-experienced strategists. Our study advances theory on three fronts. We explain why AI reliance in strategic sensemaking is a structurally distinct problem—because prospective sensemaking is intertemporal and binding, AI errors enter a sequence in which their consequences compound across interdependent decisions rather than correcting at the point of origin, establishing AI reliance in strategy as theoretically separable from human-automation interaction in general. Building on this, we develop reliance calibration as the micro-process linking strategists' fit–misfit appraisals of AI to reliance outcomes, specifying the enabling conditions under which calibration succeeds versus degrades into overreliance or underutilization (Dietvorst et al., 2015; Parasuraman & Riley, 1997). We further identify trajectory-level fit–misfit dynamics: because scanning and interpretation are coupled stages of a single sensemaking sequence rather than independent tasks, AI fit appraisals at one stage shape—and can reverse—appraisals at the next. Together, these findings extend research on strategic sensemaking (Daft & Weick, 1984; Thomas et al., 1993), task-technology fit (Goodhue & Thompson, 1995), and

appropriate reliance in human-AI collaboration (Glikson & Woolley, 2020; Lee & See, 2004) by theorizing AI reliance in strategy as a problem of sequentially interdependent judgment—not point-in-time accuracy calibration.

2. Theoretical Background: AI-Enabled Strategic Sensemaking in Strategizing

Strategists play a pivotal role in shaping firms' strategic direction through their engagement in the strategy process (Hutzschenreuter & Kleindienst, 2006; Krogh et al., 2021). They are responsible for defining organizational goals (Gioia & Chittipeddi, 1991; March & Simon, 1958; Simon, 1947), recognizing and interpreting critical environmental issues (Isabella & Waddock, 1994; Miller & Lin, 2022; Thomas & McDaniel, 1990; Thomas et al., 1994), and initiating strategic actions to secure organizational success (Chattopadhyay et al., 2001; Dutton & Ashford, 1993; Hambrick & Mason, 1984; Plambeck & Weber, 2010). Strategizing involves diverse actors such as corporate planning teams, external strategy consultants (Whittington, 1996), middle managers (Dutton et al., 1997), or upper echelons (Hambrick & Mason, 1984). Due to the collective nature of these activities, strategizing is fundamentally a social endeavor (Jarzabkowski et al., 2007).

Although the outcome of strategizing profoundly affects organizations at the macro level, this research explicitly focuses on the micro-level processes, examining how individual strategists interact with emerging technologies like AI. Such technologies are increasingly reshaping how strategists perceive, interpret, and respond to environmental changes (Constantiou et al., 2023; Weiser & Krogh, 2023). To do this, we first specifically highlight several relevant aspects of strategizing for this study, which are elaborated in the following subsections.

2.1 *Strategizing as Micro-Level Sensemaking*

Strategic sensemaking is a critical micro-level activity through which strategists process environmental signals, trends, and developments that could significantly influence organizational performance (Burgelman et al., 2018; Jalonen et al., 2018). It involves recognizing issues and constructing meaningful interpretations from ongoing environmental cues before developing responses (Bansal et al., 2018; Maitlis & Christianson, 2014; Weick et al., 2005). Hence, strategizing includes evaluative judgment—deciding what counts as credible, relevant, and defensible under ambiguity, time pressure, and uneven evidence (Daft & Weick, 1984; Maitlis & Christianson, 2014; Weick, 1995). In practice, sensemaking and strategy making often surface competing demands such as plausibility versus accuracy, speed versus assurance, and alternative framing versus contextual validity (Barnard, 1938; Kaplan, 2008; Quinn, 1991; Weick, 1995). Drawing on foundational sensemaking research (Weick, 1979, 1995), scholars have identified three essential sensemaking activities: *scanning*, *interpretation*, and *action* (Daft & Weick, 1984; Thomas et al., 1993).

Specifically, scanning involves systematically gathering environmental information (Aguilar, 1967; Hambrick, 1982). Interpretation refers to deriving meaningful insights from this information (Anderson & Nichols, 2007; Dutton et al., 1983; Miller & Lin, 2022; Thomas et al., 1993), for instance, by categorizing them as threats or opportunities (Dutton & Jackson, 1987; Jackson & Dutton, 1988). Those two activities involve strategic analysis and formulation, in which strategists assess internal and external environments, analyze relevant information, evaluate options (Hutzschenreuter & Kleindienst, 2006; Krogh et al., 2021), and discuss opinions for initiating change (Alt & Craig, 2016; Dutton & Ashford, 1993). Traditionally, sensemaking literature emphasized retrospective processes, focused on interpreting past events to understand

what has occurred (Weick, 1995). More recent scholarship, however, increasingly recognizes sensemaking as prospective, oriented towards future events and possibilities (Gephart et al., 2010; Maitlis, 2005; Sakellariou & Vecchiato, 2022). In strategic contexts, this prospective dimension is critical as strategists must anticipate future developments and prepare their organizations accordingly, connecting past experiences and present observations, and developing future-oriented plans (Kaplan & Orlikowski, 2013).

What makes this sensemaking work strategic, however, is less the importance of the issues it addresses than the decision architecture in which it is embedded. Strategic decisions are characterized by commitment—choices that are costly to reverse and that bind future action (Ghemawat, 1991); by interdependence—their value depends on coherence with other choices made contemporaneously, over time, and across actors (Leiblein et al., 2018; Porter, 1996; Rivkin, 2000); and by competitive interaction—their payoffs depend on the responses of rivals, such that heterogeneity in how firms decide is itself a source of advantage or disadvantage (Leiblein et al., 2018). Managerial deliberation can be consequential and accountable without exhibiting any of these properties; strategic decision-making exhibits them jointly. Given the interdependent nature of the sensemaking activities, where effective scanning affects interpretation outcomes and both activities subsequently shape strategic actions (Anderson & Nichols, 2007; Dutton & Duncan, 1987; Tapinos & Pyper, 2018), scanning and interpretation are thus not isolated judgment tasks: scanning outputs become the evidential basis for interpretation, interpretation becomes the cognitive basis for resource commitment, and commitment locks the firm into a competitive position that it must then defend (Daft & Weick, 1984; Thomas et al., 1993). An error at these early stages is therefore not locally contained—it propagates through interdependent choices, is locked in by commitment, and is difficult to detect under the prospective uncertainty that strategy work

entails (Sydow et al., 2020; Tripsas & Gavetti, 2000). This premise organizes our theorizing: it is why we concentrate on scanning and interpretation as the stages at which the premises of subsequent strategic commitments are set, it informs the propositions we develop, and it grounds the null hypotheses and boundary conditions we articulate in the discussion.

2.2 *Contemporary AI for Strategizing*

Strategizing is an intelligence-intensive endeavor (Eisenhardt & Zbaracki, 1992; Stubbart, 1989). It demands substantial cognitive resources—especially sensation, perception, and interpretation—needed for effective sensemaking (Helfat & Peteraf, 2015). Yet managers operate under critical cognitive and temporal constraints that limit their ability to notice and process relevant environmental signals (Hambrick & Mason, 1984; Laamanen et al., 2018; Ocasio, 1997; Simon, 1947). To cope, organizations increasingly deploy digital technologies—such as interactive platforms, data-visualisation tools, and communication systems—to augment strategic decision-making (Jarzabkowski & Kaplan, 2015; Jarzabkowski & Paul Spee, 2009; Malik et al., 2025). As digital transformation deepens, these technologies now undertake tasks once reserved for human strategists, including strategic scanning and interpretation (Constantiou et al., 2023).

Recent advancements in AI represent a significant frontier in computing (Berente et al., 2021) aiming to replicate intelligent human-like behavior (McCarthy et al., 1955). Contemporary AI technologies differ markedly from earlier static, rule-based expert systems (Berente et al., 2021; Yoon et al., 1995) through their dynamic ability to learn from data, enabling operations in novel scenarios (Balasubramanian et al., 2022; Rahwan et al., 2019). Within this realm of machine learning, two fundamental approaches merit distinction: *Predictive AI (PredAI)*, which excels in forecasting outcomes, such as capital returns, and *Generative AI (GenAI)* (Hutzschenreuter & Lämmermann, 2025; Raisch & Fomina, 2024). Notably, GenAI has garnered significant attention

in recent years for its capacity to produce human-like text, images, audio, and video (Boussieux et al., 2024; Feuerriegel et al., 2024). A central development within GenAI is large language models (LLMs), which demonstrate advanced contextual understanding by predicting subsequent words from extensive textual data—enabling complex reasoning and problem-solving (Boussieux et al., 2024; Brown et al., 2020; Murphy, 2023), as well as effective human-machine interaction, for instance by answering questions (Wei et al., 2022). LLMs have shown superior performance in diverse domains, including innovation productivity and creativity (Bouschery et al., 2023; Eapen et al., 2023), empathetic communication (Yin et al., 2024) and excelling in professional examinations (Geerling et al., 2023; Girotra et al., 2023). Recognizing the significant and exceptional breadth and depth of capabilities of contemporary PredAI and GenAI, this research particularly targets these state-of-the-art learning algorithms as its primary focus when referring to AI or contemporary AI.

Given these developments, discussions around AI agency have intensified (Krakowski, 2025; Murray et al., 2021; Vanneste & Puranam, 2025), focusing on the human perception of AI's seemingly independent capacities to think, plan, and act (Gray et al., 2007). Considering the cognitive limitations of strategists, especially during demanding tasks such as strategic sensemaking (Laamanen et al., 2018), contemporary AI is a valuable collaborator (Jarvenpaa & Klein, 2024; Krakowski, 2025; Raisch & Fomina, 2024; Weiser & Krogh, 2023). Thus, such algorithms present promising opportunities to enhance the work of strategists by assisting in developing business models, evaluating strategic alternatives (Doshi et al., 2025), generating novel business ideas, and analyzing competitors (Boussieux et al., 2024; Dell'Acqua et al., 2023).

Nevertheless, AI also exhibits critical limitations. Its probabilistic reasoning can lead to biases or inaccuracies—commonly known as "hallucinations"—prioritizing plausibility over

accuracy (Bubeck et al., 2023; Ji et al., 2023). Such inaccuracies are particularly challenging for users to detect due to the complex and opaque ("black box") nature of these AI systems (Feuerriegel et al., 2024; Spitale et al., 2023). Moreover, AI algorithms' performance generally depends on the quality and scope of input data, further increasing the risks of embedded biases (Bean, 2017; Russell & Norvig, 2022). Consequently, scholars emphasize AI's augmentation rather than substitution role by viewing it as a complement to human capabilities (Balasubramanian et al., 2022; Krakowski, 2025; Raisch & Krakowski, 2021).

2.3 *The Relevance of Task-Technology Fit–Misfit in Strategizing*

Crucially, the effectiveness and value of any technological addition depend on achieving what prior work calls “*fit*”—the alignment between a technology’s capabilities and the demands of the task it supports (Goodhue & Thompson, 1995; Hahn & Wang, 2009; Jeyaraj, 2022; Liu et al., 2017). Although the notion of fit originated in biology—where “survival of the fittest” describes how well species align with their environment (Spencer, 1864)—management scholars have long highlighted analogous alignment needs in strategy (Venkatraman, 1989) both with firms’ external environments (Bain, 1956; Porter, 1979) and with internal organizational elements such as structure (Chandler, 1962; Rumelt, 1974).

Within the technology domain, fit denotes the degree to which a tool’s features match the requirements of its intended tasks (Goodhue & Thompson, 1995). Evidence shows that fit—or its absence—strongly influences utilization and performance, particularly in managerial work that is non-routine, time-critical, and interdependent (Gebauer et al., 2010). Yet studies of technology fit rarely engage strategic decision-making directly, even though strategy work exhibits unusually high interdependence across actors, time horizons, and decisions (Leiblein et al., 2018). Given the

far-reaching consequences of strategic choices (Eisenhardt & Zbaracki, 1992; Leiblein et al., 2018), evaluating whether emerging technologies enable—or hinder—those choices is essential.

Moreover, the relationship between tasks and technology is nuanced, as technologies can demonstrate a *fit*, which is the ideal state, or varying degrees of *misfit*—such as an overfit or an underfit—by providing functionality that exceeds or falls short of task requirements (Howard & Rose, 2019; Junglas et al., 2008). For example, big-data analytics may handle massive datasets but still misfit if its interface is unnecessarily complex or lacks a crucial feature set (Muchenje & Seppänen, 2023). Importantly, the fit between a task and technology is rarely binary; specific technological features can fit some aspects of a task while simultaneously misfitting others (Howard & Rose, 2019).

2.4 *Relying on AI in Strategizing*

Introducing contemporary AI technologies into strategizing amplifies this evaluative workload because AI outputs can be simultaneously useful and risky—especially when systems are inscrutable and errors are difficult to anticipate *ex ante* (Berente et al., 2021; Burrell, 2016; Raisch & Krakowski, 2021). As a result, strategists adopt contemporary AI for strategic sensemaking—especially scanning and interpretation—but continually evaluate when AI supports their task requirements and when it undermines them. Given this decision architecture, such evaluation is not a generic usability judgment: what strategists accept from AI at these stages becomes the premise of interpretations and commitments that guide and constrain subsequent choices, so reliance decisions are weighted by their downstream consequences rather than by local task performance alone.

Consequently, the integration of contemporary AI fundamentally reshapes trust dynamics, necessitating an improved understanding of its generative and probabilistic nature (Krakowski,

2025). Trust, in turn, significantly influences users' reliance on AI, impacting collaborative effectiveness and outcomes (Bankins et al., 2024). However, research indicates that users' struggle to evaluate AI's fit for a task impacts trust formation in AI. Specifically, the iterative advancements and inconsistent ("jagged") capabilities of GenAI (Baig et al., 2024; Dell'Acqua et al., 2023) complicate the strategist's ability to assess the suitability of the technology for particular tasks, leading to variability in trust and use of it (Dell'Acqua et al., 2023; Glikson & Woolley, 2020; Vanneste & Puranam, 2025). Strategists must therefore continuously validate AI's performance to mitigate the risk of excessive trust (Lebovitz et al., 2022), but still, there is a risk of blind overreliance on it by risking unfavorable outcomes (Dell'Acqua et al., 2023; Hoff & Bashir, 2015; Keding & Meissner, 2021). Conversely, inaccurate outputs from AI can diminish user trust (Bubeck et al., 2023), resulting in algorithmic aversion—where users unjustifiably disregard the technology (Dietvorst et al., 2015) despite its potential task superiority (Logg et al., 2019). Hence, scholars highlight humans' need to understand AI's nature to find a “fit” for effective collaboration (Einola & Khoreva, 2023). This study further investigates these nuanced dynamics by exploring the tension between AI's perceived fit and misfit within strategic sensemaking tasks and examining its implications for the collaborative effectiveness of strategists and AI, an area largely unexplored so far.

In this context, we define *perceived AI fit as strategists' appraisal of the degree to which an AI system's capabilities align with the requirements of strategic tasks. Perceived AI misfit arises when AI functionality is appraised to either fall short of, or overshoot, those requirements, generating inefficiencies or suboptimal outcomes for strategists. This conceptualization builds on fit perspectives that locate performance consequences in the alignment between task requirements and technology capabilities, while recognizing that misfits can persist even in settings of overall*

adoption (Goodhue & Thompson, 1995; Strong & Volkoff, 2010). Because reliance decisions are driven by situated judgments of what is sufficiently credible, relevant, and usable in a given episode of work, we treat fit and misfit as appraisals rather than as directly observable objective alignment (Goodhue & Thompson, 1995). Accordingly, we conceptualize perceived AI fit–misfit tension as an evaluative tension that arises when AI is appraised to support some requirements of scanning or interpretation while undermining others within the same task episode (or tightly linked episodes). This tension is consequential because it shifts the analytic focus from whether AI is adopted to how strategists manage ongoing alignment between task demands and AI outputs (Goodhue & Thompson, 1995; Strong & Volkoff, 2010).

We refer to this ongoing alignment work as *reliance calibration: the situated practices through which strategists align the degree to which they accept, adopt, or delegate to AI outputs with their evolving assessment of output reliability and task fit*. This tension makes reliance decisions central: strategists must decide how much to accept, adopt, or delegate to AI outputs and what validation is required before those outputs travel into downstream interpretation and recommendation (Lee & See, 2004; Parasuraman & Riley, 1997). Reliance calibration is consequential because both overreliance (uncritical acceptance) and underutilization (avoidance despite potential task fit) can arise from the same fit–misfit tension when verification and learning practices are weak or absent. This framing links task-technology fit–misfit to reliance calibration in human–AI collaboration.

Taken together, the literature on strategic sensemaking (scanning and interpretation), task-technology fit–misfit, and trust and reliance sensitizes our qualitative inquiry to (a) where strategists locate perceived AI fit and misfit during scanning and interpretation, and (b) how these appraisals shape reliance decisions and collaboration outcomes.

3. Research Method

3.1 Research Context and Design

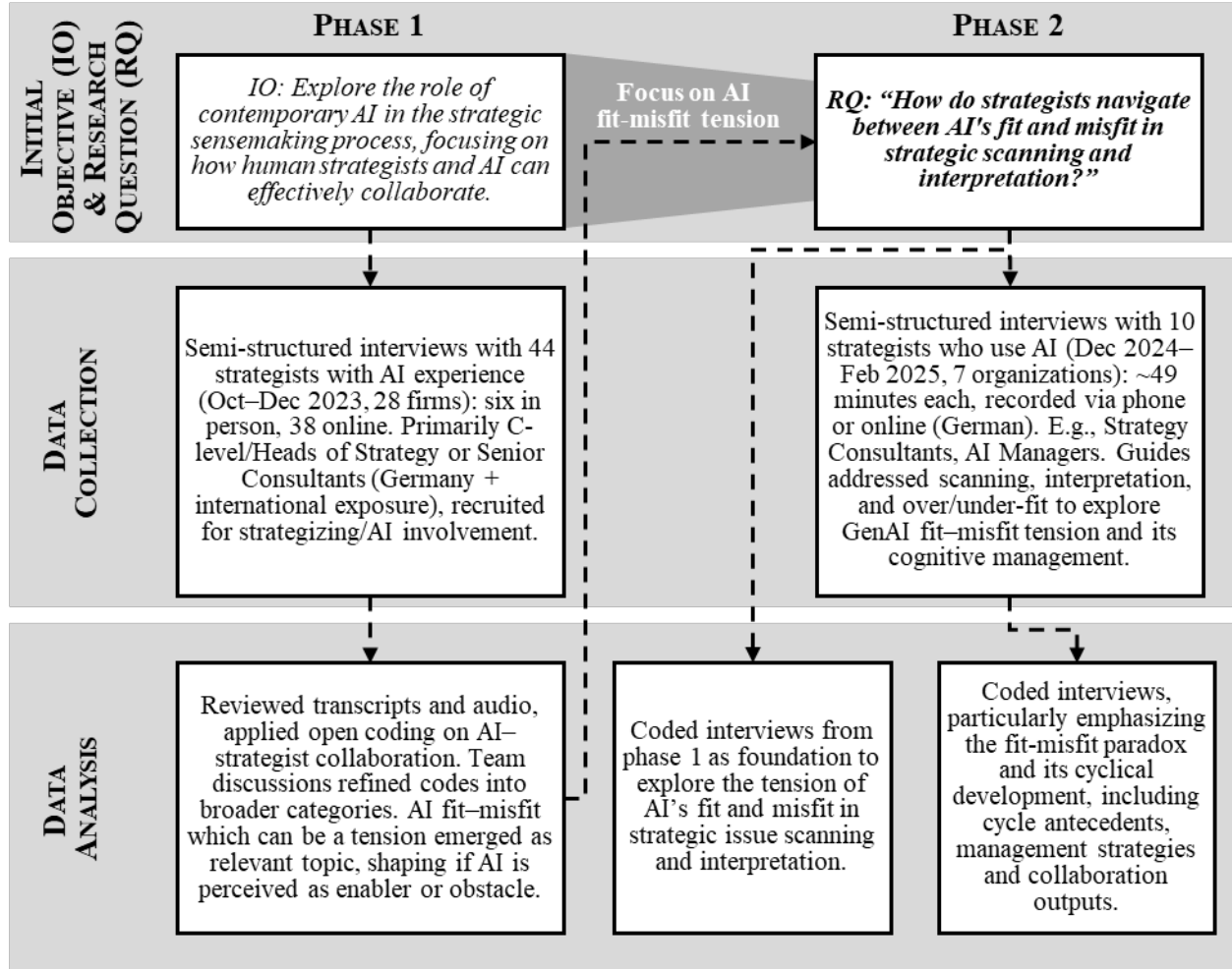


Figure 1: Overview of the two-phase research method after Follmer et al. (2018).

This study employs a two-phase, qualitative grounded theory approach (Follmer et al., 2018; Strauss & Corbin, 1998) to examine how strategists incorporate AI into their strategic sensemaking practices, as shown in Figure 1. While prior work recognizes the scanning and reasoning capabilities of technologies, including AI, from a more theoretical and conceptual level

(Balasubramanian et al., 2022; Constantiou et al., 2023; Weiser & Krogh, 2023), little is known about how strategists collaborate with AI in strategic sensemaking activities in practice. Phase 1 of our research study explored strategists' initial usage of contemporary AI in scanning and interpretation. This initial phase thereby steered us toward the question of how strategists navigate between AI's fit and misfit in strategic scanning and interpretation; Phase 2 investigated the tensions that gained our attention in Phase 1 in more depth. This approach is particularly appropriate because it enables the generation of theory grounded in empirical data, emphasizing the iterative relationship between data collection, analysis, and emerging insights (Strauss & Corbin, 1998). Additionally, grounded theory aligns with our objective of developing an in-depth theoretical understanding of our research question rooted in strategic sensemaking (Daft & Weick, 1984), and the technology fit–misfit literature (Goodhue & Thompson, 1995). Consequently, the research is structured into two distinct phases, each with clearly defined objectives and rationales:

Phase 1. In Phase 1 (October–December 2023), we aimed to broadly explore how strategists collaborate with contemporary AI during strategic sensemaking, particularly within scanning and interpretation activities. This initial exploratory approach (Creswell & Creswell, 2023) was necessary due to the novelty of AI in strategic contexts and its relatively limited adoption at that time (Chui et al., 2023; Singla et al., 2024). The findings from this phase highlighted two key insights: first, consistent with other studies (Chui et al., 2023; Singla et al., 2024), adoption of GenAI was low, and its broader use by strategists remained partly theoretical, yet rapidly evolving; second, a critical tension emerged concerning the perceived fit or misfit of AI capabilities with sensemaking tasks. Although intriguing, these findings required deeper exploration to understand this tension better.

Phase 2. Therefore, Phase 2 (December 2024–February 2025) specifically focused on the tension between contemporary AI’s fit and misfit as experienced by strategists. By this time, advances in GenAI—such as Google Gemini, ChatGPT-4o, and GPT-o1—had increased practical adoption rates and significantly improved AI capabilities, enhancing the relevance and timeliness of the research. Moreover, the period between phases allowed us to thoroughly engage with literature about the fit–misfit of technology (Goodhue & Thompson, 1995), enabling us to refine our research focus and develop a targeted interview guideline for Phase 2. Hence, after Phase 1 revealed the existence of a tension in how strategists perceived AI’s capabilities fit and misfit, we refined our central research question and interview guide in Phase 2 to examine the perception and management of the tension between contemporary AI fit and misfit in strategic sensemaking.

The following subsections outline the detailed sampling strategy, data collection procedures, and analytical methods used in each research phase.

3.2 *Research Sample*

Phase 1. For the first research phase, a purposive sampling approach was adopted (Suri, 2011), with explicit inclusion criteria designed to yield rich and contextually meaningful insights. Participants were carefully selected based on their active involvement in strategic sensemaking activities, broadly encompassing strategists across multiple organizational levels and firms. This included senior executives (e.g., C-suite, Heads of Strategy), internal strategy practitioners, as well as influential external actors such as senior strategy consultants, aligning with the inclusive conceptualization of strategy practitioners (Jarzabkowski & Whittington, 2008; Whittington, 1996). Given the research’s specific emphasis on exploring the role of contemporary AI, particularly generative AI (GenAI), participants were also required to be actively engaged with contemporary AI technologies in their daily strategic tasks. Practitioners not employing AI in

strategic sensemaking processes were deliberately excluded to ensure relevance and depth in exploring post-adoption technology dynamics, explicitly focusing on GenAI's perceived fit and misfit in these contexts.

Participants were recruited through direct outreach via professional networks, particularly leveraging platforms such as LinkedIn. The final Phase 1 sample comprised 44 strategic experts from diverse organizational contexts, ranging from large multinational corporations to innovative start-ups, across various industries, including automotive, industrial goods, media, financial advisory, semiconductors, pharmaceuticals, rail, and commercial vehicles. This heterogeneity in industry, firm size, role, and background was intentionally sought to enrich the empirical insights and increase the potential for broader theoretical transferability. The largest subgroup consisted of senior strategy consultants (n=10), supplemented by executives and internal strategists, all actively engaged with GenAI technologies. We conducted the interviews for Phase 1 over three months, from October to December 2023.

Phase 2. Building upon insights gained during the initial research phase, Phase 2 applied the same participant selection criteria but aimed explicitly at examining the emergent tension between AI fit and misfit in more depth. Conducted between December 2024 and February 2025, this second phase consisted of ten interviews with strategists representing seven distinct organizations. Participants predominantly belonged to management consulting firms (n=8), with additional representation from financial services (n=1) and software industries (n=1). All Phase 2 participants were geographically located in German-speaking countries (Germany, Austria, Switzerland), allowing for culturally consistent yet professionally varied perspectives on AI integration. The professional tenure of participants in Phase 2 ranged from several months to approximately five years, ensuring a balanced representation of both relatively new and more

experienced strategy professionals. This targeted and slightly smaller sample size facilitated deeper qualitative exploration into the nuanced dynamics of the fit–misfit tension identified earlier. Further detailed information regarding individual participants, including roles, tenure, industries, and interview specifics, is provided comprehensively in the Appendix Table A.1.

3.3 *Data collection*

Semi-structured interviews served as the primary data collection method, allowing a conversational and intuitive structure (Bearman, 2019). Each interview began with an introductory segment designed to build rapport (Arsel, 2017) and briefly explained the theoretical foundations, specifically the strategic sensemaking stages of scanning and interpretation (Daft & Weick, 1984; Thomas et al., 1993). Demographic inquiries followed this introduction to establish context and facilitate subsequent analysis.

The main interview segments focused explicitly on the core research phenomena, using open-ended questions to encourage participants to freely share experiences and insights, thereby allowing the exploration of emergent themes (Myers, 2020; Myers & Newman, 2007). While the interview guidelines consistently addressed strategic sensemaking activities, particularly scanning and interpretation (Daft & Weick, 1984; Thomas et al., 1993), minor refinements were made over time to reflect emerging insights from ongoing analyses. The final portion of each interview involved reflective questions, allowing participants to provide additional thoughts or clarifications regarding the topics discussed (Bearman, 2019).

Interviews in both research phases were conducted virtually using platforms such as Zoom or Microsoft Teams, with a small number conducted by phone or in person. Interviews were audio-recorded following explicit participant consent, except for one instance in Phase 1, where detailed notes were taken due to the participant's preference. Recordings were transcribed verbatim, and

transcripts were carefully reviewed against audio recordings to ensure accuracy, reliability, and veracity (Davidson, 2009; Halcomb & Davidson, 2006). For the transcriptions, we used the tools Whisper and Notta.AI. To enhance readability while preserving meaning, minor edits were applied, removing filler words, false starts, and nonverbal utterances (Powers, 2005). The interviews conducted in German were then translated into English at this stage using DeepL, a highly effective language translation tool that uses contemporary AI methods.

Phase 1. The first data collection phase involved 44 interviews, predominantly conducted virtually (n=38) and the remainder in-person (n=6), between October and December 2023. The interview duration ranged from 17 to 52 minutes, totaling approximately 25 hours and 35 minutes of recorded material. Most Phase 1 interviews were conducted in German (n=23), with the remainder in English.

Phase 2. It was conducted between December 2024 and February 2025 and consisted of ten virtual interviews, averaging approximately 49 minutes each. All Phase 2 interviews were conducted in German, focusing explicitly on exploring the identified perceived tension between AI fit and misfit, strategists' responses to these tensions, and their perceptions regarding future AI use cases in strategizing. Detailed descriptions of interview procedures and a comprehensive interview guide are provided in Appendix B.

3.4 *Data Analysis: Emergence of the Fit–Misfit Tension*

Following established grounded theory procedures, we analyzed the interview data in two phases (Strauss & Corbin, 1998). In each phase, we adopted a constant comparative approach to iteratively develop and refine emergent concepts while staying open to novel insights from the data (Charmaz, 2014; Glaser & Strauss, 1980). The overarching analytic goal was to move systematically from interview transcripts' initial exploratory coding to higher-order themes and

theoretical interpretations. We used MAXQDA software to facilitate the coding process and document the evolution of our analytical categories.

Initial Familiarization and Open Coding. While transcribing interviews, we immersed ourselves in the data by reading and re-reading the transcripts, noting preliminary ideas and reflections (Flick et al., 2004). In line with grounded theory’s notion of *open coding* (Strauss & Corbin, 1998), two researchers independently assigned codes to text segments that illuminated respondents’ experiences or perspectives relevant to our study. Despite entering with a broad conceptual understanding (i.e., strategic sensemaking phases of scanning and interpretation [Daft & Weick, 1984; Thomas et al., 1993]), we remained open-minded and inductive, creating codes directly reflecting participants’ language and meanings rather than imposing predefined categories (Creswell & Poth, 2025).

Collaborative Axial Coding and Category Refinement. After initial coding, the two researchers compared and discussed their code lists, resolving coding discrepancies through re-examining the transcripts and joint deliberation (O’Connor & Joffe, 2020). This process both ensured inter-coder reliability and refined the coding scheme. Through *axial coding* (Strauss & Corbin, 1998), we grouped the open codes into more abstract “second-order” categories by asking how different codes related to or built upon one another (Charmaz, 2014). We identified, for example, recurring patterns regarding how strategists assessed the potential and limitations of AI in scanning tasks, as well as the variety of emotional responses to perceived AI misalignment. At this stage, we revisited the literature (Flick et al., 2004) to sharpen and further develop categories in light of existing theoretical constructs related to sensemaking (Thomas et al., 1993).

Comparative Analysis Between Phases. Although both phases shared the same overall coding principles, they differed in emphasis. In Phase 1, once we had completed the open and axial

coding, we noticed an emerging pattern of tensions regarding how strategists evaluated AI’s capabilities for scanning and interpretation tasks. We observed that our data in Phase 1 lacked depth for explaining how these tensions manifested over time or what influenced them. This insight informed the design of our Phase 2 interviews, in which we specifically probed participants about the tension-laden experiences uncovered in Phase 1. The analysis of Phase 2 interviews then used the same open and axial coding steps but was more targeted toward clarifying, testing, and extending the categories discovered earlier.

Theory Building and Dimension Aggregation. In the final step of the analysis, we integrated the second-order categories into overarching theoretical dimensions aligned with the core focus of our research. We synthesized categories from both research phases into two key dimensions, but with the more in-depth insights from Phase 2—(1) the *perceived AI fit–misfit tension in strategic sensemaking* (scanning and interpretation), and (2) the *reliance calibration and its outcomes*. Figures that visually present the aggregation of themes and dimensions are shown in Appendix C.1 and C.2. Moreover, a selected indicative list of open codes, themes, and dimensions, along with illustrative quotes, can be found in Tables C.3 and C.4 in the Appendix.

4. Results: The Perceived AI Fit–Misfit Tension in Strategic Sensemaking

This section reports how strategists describe using contemporary AI tools in two core strategic sensemaking tasks—environmental scanning and interpretation—and what happens when they decide how much to rely on AI outputs. A core contribution of this study is highlighting the tension strategists perceive when assessing whether AI tools are a “fit” or “misfit” for scanning and interpretation tasks. We use (perceived) *fit* to denote strategists’ practical evaluation of how well an AI tool’s outputs and interaction style fit the required sensemaking task. (Perceived) *Misfit*

denotes perceived mismatch—when AI outputs are incomplete, unreliable, overly generic, excessively detailed, or otherwise not aligned with what the task demands.

Our results show that many strategists describe a perceived AI fit–misfit tension: within the same task episode (or tightly linked episodes), AI enables speed and breadth while simultaneously creating concerns about accuracy, contextual relevance, or defensibility that must be managed before outputs can be relied on. This tension (T) is consequential because it shifts attention from “Should we use AI?” to “How should we rely on AI—and how should we verify it—given its strengths and limits?” This shift matters strategically because scanning and interpretation open the decision sequence outlined in our theoretical background: what strategists accept as credible at these stages becomes the premise for interpretation and, ultimately, for the strategic actions and commitments that follow. Figure 2 summarizes where strategists locate fit–misfit tensions in their sensemaking work: within scanning (ST), within interpretation (IT), and in cross-task trajectories where an appraisal in scanning shapes or flips in interpretation (S–IT1, S–IT2).

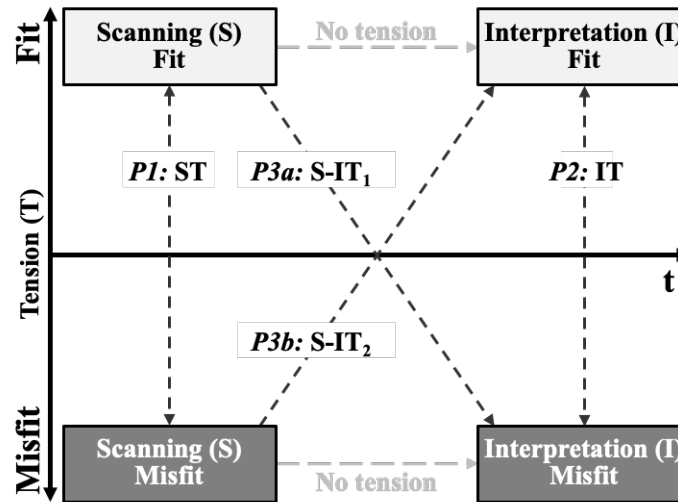


Figure 2: The Perceived Fit–Misfit Tension of AI in Strategic Sensemaking

4.1 *Sensemaking fit appraisal: The origin of the fit–misfit tension*

Our tension-based view (Dameron & Torset, 2014) was inspired by insights gathered in the Phase 1 interviews. Across interviews, strategists repeatedly framed AI as a supporting collaborator rather than a substitute decision maker. One Head of Strategy captured a baseline orientation that underpins many subsequent appraisals: AI should accelerate and inform, but a “human layer” must remain responsible for assessing limits and implications:

“Because it should always be a support and not fully take over [...] what can I do with it, but also what are the limits where again, a human layer has to take over” (E7, Head of Strategy).

This orientation reflects a broader theme: strategists do not evaluate “AI” in the abstract, but appraise its fit or misfit relative to task requirements and relative to what they believe humans can contribute. Even as AI advances, its trajectory of development may not converge on human-like contributions; instead, progress may increase AI’s functional distance to humans. As one participant put it:

“I would not say that it is on the level of human. It is different from what a human can do” (E2, Principal).

The “difference” creates opportunity (e.g., AI can do some things faster or at scale) and risk (e.g., AI may be convincing but wrong, or miss context) making appraisal and subsequent reliance decisions central. A further complication is that appraisal is not purely technical—it is also interpretive. Strategists highlighted that perceptions of bias can emerge from the engine, from user prompting, or from how users extract and apply outputs:

“Because perception at times can become reality [...] there could be bias within the engine, or there could be bias in how you extract information from the engine and what you do of it” (E3, Partner).

This observation foreshadows why reliance calibration practices (Section 4.2) become pivotal: appraisal involves judgement under uncertainty about output quality and about one’s own interpretation. In the following subsections, we examine more closely how these tensions play out in scanning and interpretation.

4.1.1 Scanning tension [ST]

In scanning tasks, strategists consistently described AI as a high-speed accelerator that produces an initial map of a topic space—useful when the immediate need is to get oriented quickly. One interview participant characterized this fit succinctly:

“[...] it was a huge acceleration of the process... You have initial results super quickly... Sure, you have to refine them and check them” (E45, Founders Associate).

Here, the appraisal contains two elements at once: fit for speed and early overview, and misfit because the outputs cannot be accepted without refinement and verification. This scanning tension becomes sharper when the scanning task demands precision or context-specific accuracy, particularly for quantitative or country-specific details. The same participant described how AI produced plausible but incorrect transfers across contexts (e.g., applying a Germany percentage to Japan), and how this error pattern forced persistent human oversight and iterative correction:

“I had to keep prompting, ‘Hey, aren’t there difference [...]?’ [...] I always tried to check if it added up... and I still had to make some adjustments” (E45, Founders Associate).

In such cases, AI’s scanning fit is inseparable from a misfit that appears as a verification burden: the user must invest effort to validate, correct, and sometimes redo work to make outputs usable.

Strategists also differentiated scanning fit by the nature of the material: AI was often appraised as more fitting for early-stage, qualitative broadening than for numeric accuracy. One consultant explained their deliberate avoidance of using AI for numbers, while still leveraging it for

*“[...]creative brainstorming [...] at the very, very beginning [...] to get the big picture”
(E52, Consultant).*

In these accounts, scanning fit is not “AI is good at scanning,” but “AI fits certain scanning sub-tasks (e.g., breadth, ideation, initial structuring) and misfits others (e.g., precision, numeric reliability).” Hence, scanning fit appraisals are anchored in speed and breadth, but the same accounts repeatedly foreground a misfit tied to unreliable or insufficiently precise outputs—especially when scanning requires accuracy and context specificity. This within-task tension shifts scanning from a simple adoption decision toward a reliance question: strategists must decide how much of the scanning output can be delegated to AI and what must be verified or rebuilt by humans.

Proposition 1: *When scanning tasks require rapid orientation while still demanding credible accuracy for downstream use, strategists are more likely to appraise AI as fitting for speed and breadth yet misfitting for precision; this perceived fit–misfit tension increases the need for validation before AI outputs enter interpretation and the strategic commitments that follow.*

4.1.2 Interpretation tension [IT]

In interpretation tasks, strategists often appraised AI as fitting because it can process large volumes of information and surface patterns or angles they might otherwise overlook. One consultant described using AI to “[...] tell us all the trends you see, [...]” which yielded unexpected but potentially relevant angles: “we got things like residences [...] I would never have looked out for something like that” (E50, Consultant).

This illustrates interpretation fit as coverage expansion—AI can broaden the interpretive search space and highlight candidates for human sensemaking. At the same time, interpretation appraisals were repeatedly conditioned by the recognition that strategic interpretation is rarely “right or wrong” in a simple sense. One participant described AI’s value as a stimulus that helps generate pros and cons, rather than a definitive answer:

“[...] it’s just kind of a thought-provoking impulse... especially when it comes to strategies, it’s often not [...] black or white” (E51, Associate).

In these accounts, AI fits as an ideation partner that supports interpretation as sensegiving and reframing—while strategists preserve final judgment. The misfit in interpretation most often emerged as contextual fragility: AI can produce convincing narratives that do not hold when checked against the underlying material. The same interviewee reported asking about companies mentioned in a report and receiving a plausible answer that collapsed under verification:

“ [...] two companies [...] were not mentioned at all in the report” (E51, Associate).

Another strategist echoed this pattern:

AI is “[...] a great help [...] ” and offers “[...] huge added value, [...]” yet “[...] incorrect information is sometimes provided [...] that’s why I have to adjust it again” (E46, Senior Consultant).

Beyond factual inaccuracy, strategists also reported misfit when interpretation requires causal-economic reasoning or relevance filtering. One consultant described experimenting with AI for pricing interpretation and receiving outputs that were technically computed but practically nonsensical:

“[...] increasing the price by two basis points for everyone... an enormously high price that you can’t actually explain to the customer” (E50, Consultant).

In another case, AI produced a “[...] holistic picture [...] very detailed, [...]” yet “[...] clearly too detailed for our task [...]” (E50, Consultant), creating misfit through over-granularity—outputs exceed what is actionable and relevant for the interpretive purpose. This again underscores the intertemporal nature of AI fit and misfit: as capabilities improve, AI can generate increasingly “overfit” outputs—analytically rich but cognitively intractable—thereby widening the human–AI capability gap and producing a perceived misfit relative to users’ actionable needs.

Interpretation fit appraisals emphasize pattern discovery and perspective generation, but the misfit repeatedly arises from contextual fragility—AI may be convincing yet wrong, insensitive to domain reasoning, or poorly calibrated to the task’s relevance threshold. As a result, strategists position AI as a stimulus for interpretation rather than an authority, and they reserve human judgment as the mechanism that turns outputs into defensible interpretations.

Proposition 2: *When interpretation tasks benefit from broad pattern discovery and alternative framings but require contextual validity and relevance filtering, strategists are more likely to appraise AI as fitting for perspective generation yet misfitting for context-sensitive reasoning; this perceived fit–misfit tension in interpretation shifts AI use toward a stimulus role that requires human judgment and selective validation before candidate framings shape the premises of strategic commitments.*

4.1.3 Cross scanning-interpretation (S-IT) trajectories

While scanning and interpretation are conceptually distinct, strategists frequently described AI use as blurring or compressing the boundary between the two. In some cases, AI enabled strategists to move directly into interpretive argumentation without extensive manual scanning:

“[...] just to give me all these arguments... ‘What are the pros and cons? Give me your analysis’ [...]. That’s where [AI] helps a lot” (E54, Consultant).

Others described AI use cases—especially idea generation—as

“[...] something in between scanning and interpretation [...]” (E46, Senior Consultant).

These accounts suggest that, rather than being slotted into existing stages, AI can reshape how strategists move between them.

Figure 2 captures two recurring cross-task trajectories. The first (S–IT1) occurs when strategists experience scanning fit for overview generation, but encounter interpretation misfit when they attempt to deepen the analysis. One interview participant described this shift:

“AI is great for creating an overview but if you then... go deeper... the more number-driven it becomes, the fuzzier the output becomes... and the more dangerous it is to simply use it” (E45, Founders Associate).

Here, appraisal changes across the handoff: what fits for early scanning becomes misfitting when interpretive precision and accountability rise.

This handoff shift is also shaped by the trade-off between verification and efficiency. One interviewee noted that if outputs are *“[...] too generic [...]”* and must be checked *“[...] for accuracy first, [...]”* the user must reconsider whether AI actually saves work (E37, Project Manager). This indicates that cross-task trajectories are not only about task type; they are about how appraisal interacts with the verification burden that emerges when outputs are carried forward into consequential interpretation.

The second trajectory (S–IT2) occurs when strategists experience scanning misfit—often in numeric or data-heavy scanning—but then repurpose AI toward interpretive or ideational work where it becomes fitting. For example, one consultant explicitly avoided AI for numbers because

it “didn’t work well,” yet still relied on it for creative brainstorming to get an initial big picture (E52, Consultant).

In this trajectory, scanning misfit triggers a shift in use: instead of abandoning AI, strategists reallocate it to tasks where its strengths better match requirements.

Finally, participants also described specific cues that can prompt reappraisal of fit across related episodes (without implying observed cycles). One participant noted that long interactions could degrade usability:

“if you stay in a chat... with every new query it becomes more and more laborious... at some point, you have to change the chat” (E45, Founders Associate).

Others described switching tools when refinement no longer yields workable results:

“if I’ve now refined it so much that it no longer works... then I would look around... [for] another technology... where there is... a higher fit” (E50, Consultant).

These accounts suggest that strategists treat fit appraisal as revisable when interaction constraints or tool behavior shifts.

Cross-task trajectories show that fit appraisal extends beyond the boundaries of scanning and interpretation: it also emerges when AI outputs are carried across stages or when strategists repurpose AI to different task requirements. The key empirical implication is that fit appraisal and misfit detection guide how strategists bound AI’s role—either restricting reliance when moving toward deeper interpretation (S–IT1) or redirecting AI use toward ideational interpretation when scanning misfits (S–IT2). These trajectories set up the next question: how do strategists calibrate reliance once this tension is recognized?

Proposition 3: *When strategists carry AI outputs from scanning into interpretation, they are more likely to experience either*

(a) scanning fit followed by interpretation misfit (S–IT1), which increases reliance constraints and verification demands, or

(b) scanning misfit followed by interpretation fit (S–IT2), which encourages repurposing AI toward qualitative/ideational interpretation; these trajectories shape how strategists bound AI’s role across sensemaking tasks that are interdependent and whose outputs feed strategic commitment.

4.2 *Reliance calibration and its outcomes: How strategists respond to the perceived AI fit–misfit tension*

Strategists’ task-level appraisals of a perceived AI fit–misfit tension—whether in scanning or interpretation—did not, in their accounts, translate automatically into collaboration outcomes. Instead, participants repeatedly emphasized reliance calibration as the practical bridge between appraisal tensions and downstream consequences: how users align the degree to which they accept, adopt, or delegate to AI outputs with their assessment of output reliability and task fit. As Figure 3 summarizes, calibration practices sit between appraisal tensions and reliance outcomes, and participants’ accounts clustered into two patterned pathways. In the calibrated reliance pathway, strategists acknowledged the perceived fit–misfit tension yet enacted effective calibration strategies, yielding *AI alignment* with perceived efficiency gains. In the miscalibrated reliance pathway, the perceived tension was similarly salient, but calibration practices were weak, absent, or avoided, yielding either *AI overreliance* (uncritical acceptance) or *AI underutilization* (avoidance/non-adoption despite potential benefits).

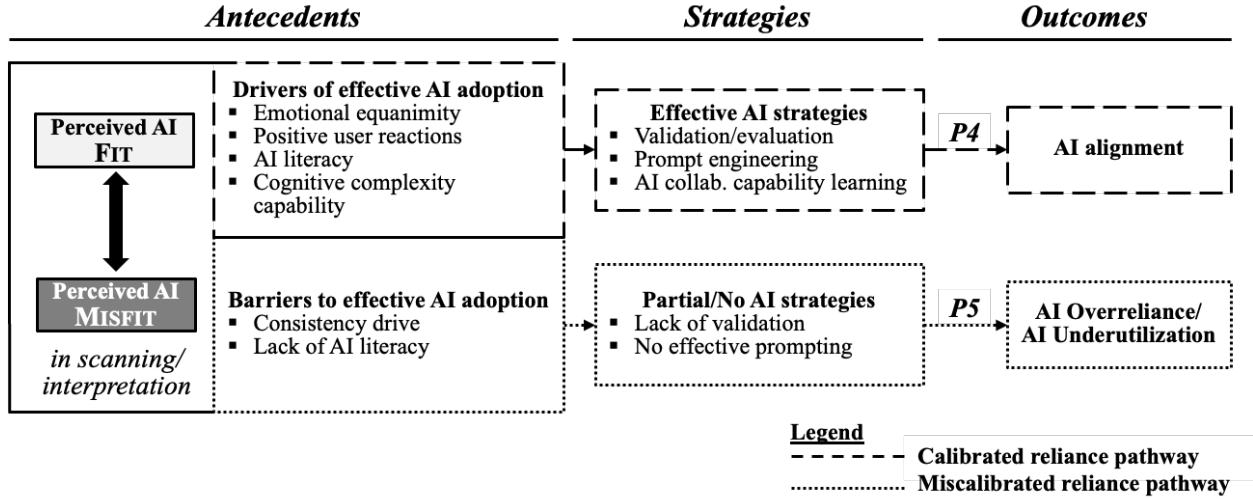


Figure 3: Antecedents, strategies, and outcomes of the perceived AI fit–misfit tension.

4.2.1 Calibrated reliance pathway to AI alignment

Antecedents. Strategists who described effective AI use consistently foregrounded antecedent conditions that enabled them to sustain calibration when AI outputs were uncertain or imperfect. A first enabling condition was *emotional equanimity*—patience and resilience in working with a tool that could produce variable performance. One participant described this stance as

“ [...] *stress tolerance* [...] *patience* [...] [AI is] *just not a cure-all* [...] and you [...] *accept the limits of it*” (E53, Associate).

Another similarly emphasized calm engagement rather than expectation of flawless performance:

“*I’m actually more relaxed about it* [...] and *I’m also aware that it can sometimes not work*” (E50, Consultant).

Such equanimity did not mean ignoring misfit. It made continued, careful engagement possible despite it. Second, calibrated reliance was enabled by positive user reactions that oriented strategists toward experimentation and opportunity without suspending critical judgment. One

interview participant linked anticipation to identifying use cases where AI could be productively applied:

“ [...] for use cases where you can use technology, you're naturally looking forward to it... this is a huge opportunity that I am realizing” (E47, Senior Consultant).

Another described how such constructive engagement translated into recurring moments of value:

“ [...] every week, I have something where I use some GenAI tool [...] and it actually works phenomenally and somehow saves me half an hour of work” (E50, Consultant).

Importantly, these reactions were expressed alongside ongoing awareness that fit could be partial and contingent.

Third, strategists' accounts emphasized capability-based enablers—particularly *AI literacy* and a form of *cognitive complexity capability* in handling contingent fit. Participants described literacy as knowing what the tool can and cannot do:

“ [...] in the end it's about knowing the capabilities of the tool” (E52, Consultant).

They also framed effective use as the ability to locate task-specific leverage:

“As soon as I know where my biggest leverage is, then [AI] really is a very big lever in that area” (E52, Consultant).

This capability was paired with a critical stance that preserved boundaries around misfit:

“ [...] try it out, be critical [...] define use cases where you know that you either have to exercise a lot of control or you simply can't do that with the technology” (E47, Senior Consultant).

Together, these antecedents supported not simply “more use,” but the conditions under which strategists could calibrate reliance in a sustained, selective manner.

Strategies. With these enabling conditions in place, strategists described what we call reliance calibration as enacted through three strategy clusters: *validation/evaluation, prompt engineering, and AI collaboration capability learning*. In this sense, validation and evaluation is a form of keeping a “human controlling authority.” Many participants framed validation as a norm of responsible strategizing with AI: the human remains accountable for the defensibility of both scanning and interpretation. As one senior consultant put it,

“ [...] you have a responsibility as a human being to have completed both the scanning and the interpretation to a reliable degree of accuracy [...] You basically have to be the controlling authority” (E47, Senior Consultant).

Validation, so framed, was anything but an optional add-on: participants treated it as the central practice that rendered AI use legitimate in strategic work. Participants also described validation as “quality control,” particularly in response to hallucination risk:

“ [...] what you should always do is quality control [...] if this slips through, you are making strategic decisions based on incorrect information” (E47, Senior Consultant).

Notably, strategists’ accounts implied that validation intensity was not constant: it increased when perceived misfit risk was higher or when downstream consequences made errors more consequential, translating the fit–misfit tension into a practical question of how much checking is required before an output is carried into interpretation.

Beyond validation, strategists emphasized calibrating reliance by actively shaping AI outputs through iterative prompting and by learning how to interact with the tool in-use. An interview participant explained that outputs became more accurate when prompts specified sources and were refined incrementally:

“ [...] you can't just ask [...] 'give me a market analysis' [...] you have to say, 'Please draw your data from the Bureau of Labor Statistics.' [...] Then you gradually came to a more accurate result” (E45, Founders Associate).

Similarly, another participant described prompting as a means to repair outputs that did not fit task expectations:

“ [...] you can use prompting to refine it and get a better result if you get options that don't make sense” (E46, Senior Consultant).

Learning-in-use appeared most clearly as iterative experimentation, especially for complex and domain-specific tasks. One consultant noted,

“ [...] it wasn't the case that I uploaded these PDFs [...] and then I got everything out straight away [...] that took a few iterations [...] [we] had to [...] refine our prompt a bit” (E50, Consultant).

Participants also treated iteration as a behavioral marker of competent collaboration with AI. An AI manager stated:

“ [...] if someone makes a lot of iterations [...] it's a sign that the person knows their way around and thinks for themselves. If someone simply takes the first suggestion [...] that's a sign [...] misfit” (E48, AI Manager).

Another participant framed this capability development as something that must be practiced rather than assumed:

“ [...] you always force yourself to use the technologies [...]. We are creatures of habit and if you don't do it, you'll do it less and less [...] Accepting training [...] researching for yourself how you can improve it, refine it” (E49, Senior Consultant).

These accounts position prompting and learning-in-use not only as output-improvement techniques, but as calibration practices that keep the strategist actively engaged in sensemaking rather than passively receiving content.

At the same time, participants emphasized that prompting is not costless and can fail when performed poorly. As one interview participant warned,

“ [...] if you don’t set the prompt correctly at the beginning, but then keep building on this prompt [...] then everything just sucks [...] you end up with a bad output [...] and have to do it all over again ” (E45, Founders Associate).

In calibrated accounts, rather than negating prompting, this risk reinforced the need for both skillful prompt work and continued validation.

Outcomes. When strategists described these enabling conditions and strategies as present, they reported outcomes consistent with *AI alignment*—using AI where it fits and constraining it where it misfits—alongside perceived efficiency gains. Importantly, participants did not frame these outcomes as eliminating misfit; rather, they framed them as benefits achieved by keeping reliance within defensible bounds. One participant described the day-to-day gains in speed and effectiveness as tangible:

“AI and the various tools [...] can [...] make your work much easier, much more efficient, and much faster [...] You can complete some tasks in just a few minutes that might take a few hours ” (E49, Senior Consultant).

This efficiency was presented as compatible with continued critical oversight, echoing the earlier emphasis on defining use cases by how much control they demand and where the technology cannot be used.

In other words, calibrated reliance yielded acceleration and breadth without relinquishing responsibility for accuracy and interpretive defensibility.

Proposition 4: *When strategists experience a salient perceived fit–misfit tension and possess enabling conditions (e.g., emotional equanimity and AI literacy/cognitive complexity capability), they are more likely to engage in effective reliance calibration strategies (validation/evaluation, iterative prompting, learning-in-use), increasing the likelihood of AI alignment (selective, bounded reliance) and associated efficiency gains, while preserving the defensibility of the downstream commitments those outputs inform.*

4.2.2 Miscalibrated reliance pathway to overreliance or underutilization

Antecedents. In contrast, strategists also described patterned accounts in which the perceived fit–misfit tension was salient but antecedent barriers constrained the capacity or willingness to calibrate reliance. A first barrier was a *consistency drive*—a preference for familiar routines and apprehension toward fundamental change that AI symbolized. One interview participant observed:

“[...] the majority of people tend to be reluctant because they say, ‘This is something new, this is fundamental change [...] and will I actually lose my job?’” (E47, Senior Consultant).

This apprehension was reflected not only as fear but also as pragmatic avoidance:

“You might not even try the internal tools [...] but simply prefer to do it yourself [...] quicker than if you get to grips with the tool first” (E49, Senior Consultant).

Such accounts locate miscalibration upstream, in a reluctance to incur the learning and verification work that calibration requires.

A second barrier was *lack of AI literacy*, described as falling behind in understanding capabilities and limits, with implications for both confidence and defensibility. As one participant warned:

“If you don't stay up to date and don't catch up somehow in the long term, you'll lose out at some point and that will have an impact on the company's effectiveness, profitability and so on” (E49, Senior Consultant).

Lower literacy, in these accounts, did more than reduce tool use: it undermined the capacity to assess outputs and to decide when additional validation or refinement was needed.

Strategies. Where these barriers featured prominently, strategists described strategy-level patterns consistent with partial or absent calibration. A first pattern was *lack of validation*, especially adopting outputs without being able to defensibly check them. One participant recalled:

“I've [...] simply adopted generic bullets [...] and then they're simply wrong [...] Especially if you [...] can't really validate it yourself” (E53, Associate).

This account makes miscalibration concrete: the problem is not “using AI,” but importing content into strategic work without the evaluation practices required for accuracy and defensibility.

A second pattern was *no effective prompting*, including minimal iteration and compounding early prompt errors rather than diagnosing and repairing them. This is precisely the failure mode that the participant quoted earlier warned about: an early prompt set poorly and then built upon rather than repaired, until the output has to be discarded entirely.

Participants also read the absence of iterative engagement as indicative of weaker capability in collaborating with AI: as the AI manager quoted earlier put it, simply taking the first suggestion marked someone who could not use the technology—for him, a sign of misfit (E48, AI Manager).

Finally, avoidance of the tool—the preference, noted above, for simply doing the work oneself rather than getting to grips with the tools (E49, Senior Consultant)—functioned as a strategy-level withdrawal from the very learning-in-use through which calibration capability could develop.

Outcomes. Strategists described two outcome forms of miscalibration. The first was AI overreliance, characterized by uncritical acceptance and reduced independent thinking. One interview participant warned that enthusiasm can become indiscriminate:

“[...] because of this enthusiasm or fascination, you can tend to become a crazy techno-optimist [...] leaving the misfit out of it” (E47, Senior Consultant).

Another participant described overuse as eroding their own cognitive engagement:

“[...] there is a risk of overusing it [...] with the result that you no longer really think about things yourself [...] you actually brainstorm less yourself and question things less” (E51, Associate).

In this outcome pattern, miscalibration appears as delegating too much without the evaluation practices needed to keep reliance bounded.

The second outcome was AI underutilization, where strategists avoided or rejected AI despite potential task fit. Participants characterized this stance as focusing exclusively on limitations:

“[...] you only see the misfit and not the fit or the advantages [...] that's simply a denial of technology” (E47, Senior Consultant).

Underutilization, in these accounts, reflected not a considered boundary decision after calibration but a withdrawal that foreclosed potential benefits where AI could have been used selectively and defensibly.

Taken together, the miscalibrated reliance pathway shows how the same perceived fit–misfit tension can be associated with opposite outcomes—overreliance or underutilization—when reliance decisions are not aligned to task demands through validation, effective prompting, and learning-in-use.

Proposition 5: *When strategists experience a salient perceived fit–misfit tension but lack enabling conditions for calibration (e.g., low AI literacy, limited willingness to iterate/validate, or strong consistency/avoidance motives), they are more likely to miscalibrate reliance—either overrelying on AI outputs without sufficient validation or underutilizing AI despite potential task fit—thereby allowing miscalibration to propagate through the decision sequence and undermining the defensibility of the strategic commitments built on them.*

5. Discussion: The Perceived Fit–Misfit Tension in Strategic Sensemaking

Contemporary AI rarely enters strategizing as a stable fit or a stable misfit. Across 54 interviews with AI-experienced strategists, we found instead a resource whose value is decided in use—appraised, bounded, and governed as strategists decide how much of what AI produces is credible enough to carry forward. That much is now familiar from research on human–AI collaboration. The claim this paper advances is narrower and, we argue, consequential for strategy specifically: *the calibration of reliance on AI is constituted by—not merely contextualized within—the structure of strategic decisions.* What makes AI miscalibration consequential in this setting is not that strategic tasks are harder, higher-stakes, or more accountable than other knowledge work. It is that strategic decisions are patterns of interdependent, committing, and competitively constitutive choices (Leiblein et al., 2018), so that a miscalibration at the scanning stage does not stay where it occurs. It enters a binding sequence whose later moves it has already shaped.

This is the governing logic of everything that follows. AI accelerates search, widens coverage, and multiplies candidate framings, while simultaneously introducing errors, irrelevancies, and context-blindness that are confidently presented and difficult to detect

(Dell'Acqua et al., 2023; Hannigan et al., 2024). The strategist's problem is therefore not whether to use AI but how to prevent its misjudgments from propagating: when to delegate, when to treat an output as a provisional stimulus, and when to verify before allowing AI-generated content to set the interpretive premises on which a commitment will rest. We name this work *reliance calibration*, and we theorize it as a component of strategic judgment rather than as a usability concern.

Our findings give this claim empirical content. Miscalibration at the scanning or interpretation stage does not produce a locally correctable error: it sets a trajectory, initiates commitment dynamics whose reversal is organizationally costly (Ghemawat, 1991; Sydow et al., 2009), and pre-empts the framing contests through which competitive reality is socially constructed (Kaplan, 2008; Porac et al., 1989). Operational mistakes correct; strategic mistakes compound. And as large language models generate, evaluate, and increasingly execute strategic options (Camuffo et al., 2026; Murray et al., 2021), they shape both what the strategist sees and what she will commit to—making the conditions under which calibration succeeds or fails a first-order strategic question that existing reliance frameworks were not built to address.

5.1 *Theoretical contributions*

This paper advances a single, integrated contribution with three facets: it theorizes how AI becomes entangled in the micro-work of strategic sensemaking, recasts fit and misfit as co-occurring evaluations across interdependent task handoffs, and specifies reliance calibration as the micro-process that translates those appraisals into systemic outcomes. Each facet generalizes propositions developed in our results, and each is delineated by a null hypothesis in the conceptual sense—not a claim to be statistically tested in this work, but a specification of the alternative logic that our process model supersedes.

AI-entangled strategic sensemaking and the irreversibility-weighted verification burden.

We specify how AI becomes entangled in scanning and interpretation (Propositions 1 and 2), and why that entanglement is strategically distinctive. The contribution moves past the generic claim that "AI augments strategists." In scanning, AI's speed and breadth widen the search surface while creating a verification burden wherever credible accuracy is required; in interpretation, its generative capacity expands the space of candidate explanations while confronting strategists with contextual fragility that must be judged, filtered, and defended. The strategically distinctive claim is that this verification burden scales not with task difficulty but with downstream irreversibility: because scanning opens a committing sequence, the effort required to establish whether an output is credible enough to carry forward is weighted by the cost of reversing the commitment it will inform—a logic consistent with behavioral work on option exercise under noisy signals and irreversibility (Ghemawat, 1991; Posen et al., 2018). AI, in short, does more than accelerate strategizing: it reorganizes the balance between producing candidate insights and validating them at precisely the stage where validation is most consequential and verification is most costly (Hannigan et al., 2024).

***Null hypothesis 1:** AI miscalibration in strategic scanning and interpretation produces local, task-level performance losses that are correctable through subsequent adjustment, without systemic consequence for the decision sequence.*

This is not a strawman—in a great deal of knowledge work it is simply true. Our process model supersedes it on strategy-specific grounds: on complex landscapes, the loss function on activity-system errors is nonlinear, because the system's value depends on the coherence of the whole choice bundle rather than any single choice (Levinthal, 1997; Rivkin, 2000). The S–IT₁ trajectory in our data makes this visible at the micro-level: scanning fit followed by interpretation misfit is

not two independent appraisals but the propagation of a scanning judgment into the interpretation it conditions—a handoff the local-correctability logic cannot accommodate. This logic does not carry over to strategy—not because AI is unusually error-prone, but because strategic structure denies errors the locality the null assumes.

Fit and misfit as co-occurring, task-contingent evaluations across interdependent handoffs.

We extend task-technology fit theory to a setting its core construct was not built for: sequentially interdependent tasks whose outputs feed binding strategic commitments. Strategists in our sample repeatedly experienced AI as fitting on some task dimensions—speed, breadth, ideation—while misfitting on others—accuracy, contextual validity, relevance—within the same episode, and these evaluations flipped across tightly linked tasks (Proposition 3). The strategically distinctive move is to explain why flipping matters: because scanning feeds interpretation in a committing sequence, perceived fit or misfit in one stage governs how AI is bounded and verified in the next. The calibration logic strategists apply is not accuracy-tracking but commitment-adjusted discounting—weighting an AI output by the irreversibility of the decision it will inform rather than by its apparent correctness. Task-technology fit theorizes the alignment but not the interdependence that turns a local misalignment into a downstream constraint—which is why the construct must be reconstructed rather than transplanted (Goodhue & Thompson, 1995; Howard & Rose, 2019).

***Null hypothesis 2:** Effective reliance on AI in strategy converges on an optimum determined by AI accuracy, such that more capable AI produces uniformly higher reliance, independent of the irreversibility of the decisions the outputs inform.*

Our process model moves beyond this logic because the strategic jagged frontier places exactly the tasks at the core of strategy outside AI's reliable zone (Dell'Acqua et al., 2023; Felin & Holweg, 2024): competitive situations that call for theory-based causal inference, readings anchored in a

firm’s particular history and position, and judgments under genuine uncertainty. These are where AI is least reliable, most confidently wrong, and where strategic distinctiveness and commitment costs are highest. The relevant competence is therefore not trust tuned to accuracy but frontier awareness—recognizing that the strategist’s hardest tasks are the model’s weakest. This alternative logic breaks down because accuracy and strategic consequence are, on the jagged frontier, inversely arranged.

Reliance calibration as a micro-process producing systemic rather than local outcomes.

We theorize reliance calibration as the micro-process linking fit–misfit appraisals to outcomes—alignment, overreliance, or underutilization (Propositions 4 and 5)—and explain why those outcomes are systemic. The same fit–misfit tension diverges less because users differ in skill than because miscalibration enters self-reinforcing committing paths whose dynamics amplify the original error in opposite directions (Sydow et al., 2009): overreliance locks in a sub-optimal trajectory the firm must then defend; underutilization forecloses informed exploration the firm cannot later recover. Enabling conditions—AI literacy, willingness to validate, organizational legitimacy for verification work—shape whether calibration succeeds, but they are antecedents of calibration, not the locus of its strategic consequence. That consequence lives in the decision sequence, which is why the calibration that matters is an organizationally distributed practice rather than an individual cognitive episode.

***Null hypothesis 3:** Experiential learning eliminates fit–misfit tensions over time, as strategists develop stable mental models of AI’s strategic capability and the calibration problem dissolves into routine.*

Our process model supersedes this logic on two grounds. Empirically, our participants did not describe fit–misfit tensions as dissolving with experience: learning-in-use made AI’s limits more

visible and salient, intensifying rather than resolving reliance decisions. Theoretically, strategic learning is constrained by path-dependent commitments (Sydow et al., 2009) and cognitive lock-in (Tripsas & Gavetti, 2000). Learning shifts the content of the calibration problem without dissolving its structural source, because strategy rarely supplies the clean, repeatable feedback that experiential resolution presupposes.

5.2 *Boundary Conditions: When Miscalibration Becomes Systemic*

The scope of our theory is defined by a conjunction of conditions, not by any single feature of the setting. Judgment intensity, uncertainty, and accountability characterize much managerial work; none alone makes reliance calibration strategically distinctive. Reliance calibration takes the systemic form we theorize when four conditions hold together: choices are interdependent, so the value of any single judgment depends on its coherence with other choices (Leiblein et al., 2018; Rivkin, 2000); choices are committing, so reversing them is organizationally costly (Ghemawat, 1991); sensemaking is prospective, so the feedback that would correct an erroneous judgment arrives late, noisy, and confounded (Gephart et al., 2010; Posen et al., 2018); and judgment is collectively accountable, so what one actor accepts as credible must be defensible to others whose own decisions it will guide (Kaplan, 2008). Together, these conditions convert local miscalibration into systemic consequence: interdependence propagates the error across choices, commitment locks it in, prospective uncertainty delays its detection, and collective accountability spreads its premises across actors.

Where any of these conditions is absent, the calibration problem reverts to the local form that existing appropriate-reliance research already explains well: an isolated, reversible, feedback-rich task in which a miscalibrated reliance decision produces a correctable performance loss (Lee & See, 2004; Parasuraman & Riley, 1997). This delimits what is distinctively strategic about our

account: for deliberation that lacks this conjunction, our model claims no advance over the existing reliance literature; where the conjunction holds, that literature’s central assumption—that miscalibration is locally correctable—fails, and calibration becomes a component of strategic judgment itself. Within this scope, learning operates as an enabling condition, not a resolution: contrary to Null hypothesis 3, learning-in-use makes fit–misfit tensions more visible rather than eliminating them, and because feedback remains non-stationary, costly, and socially negotiated, calibration does not settle into routine.

5.3 *Limitations and future research*

This study is intentionally micro-process oriented (Burgelman et al., 2018), and its claims should be interpreted accordingly. First, our evidence is based on semi-structured interviews that capture strategists’ experienced and recalled episodes of AI use. This design is well suited for theory-building on an emergent phenomenon, but it limits our ability to directly observe within-person learning loops or to adjudicate objective AI accuracy. We therefore theorize perceived fit–misfit appraisals and reliance outcomes rather than making claims about objective performance effects.

Second, our empirical focus is strategic scanning and interpretation. While these are core sensemaking activities (Constantiou et al., 2023; Daft & Weick, 1984), AI is also used in adjacent strategy tasks (e.g., communication, visualization, and coordination). Extending the model to additional tasks may reveal different calibration practices or different stakes that intensify (or dampen) the verification burden.

Third, the study spans substantial evolution in contemporary AI tools and model versions. This heterogeneity is a feature of the setting rather than a nuisance variable, but it also bounds the generality of any claim about the persistence of specific tensions. As AI capabilities and organizational deployments change, the distribution of “where AI fits” and “where AI misfits” may

shift—making the calibration problem more, not less, central, but potentially altering which task dimensions dominate appraisal.

5.4 Conclusion

This study reframes AI in strategizing as less a question of adoption and more a question of *governing reliance* in judgment-intensive work. When strategists use contemporary AI in scanning and interpretation, they are doing more than consuming better information: they are continually deciding what can be treated as credible, what must be verified, and what should be kept out of downstream recommendations. Our process model highlights that these decisions are neither automatic nor reducible to generic trust. They are enacted through reliance calibration practices that align AI contributions with task demands while protecting the defensibility of strategic judgments. Importantly, the same perceived fit–misfit tension can generate very different outcomes—productive alignment, overreliance, or underutilization—depending on whether strategists and organizations cultivate the conditions that make calibration feasible. As AI systems continue to evolve and spread through strategy work, the theoretical and practical stakes shift. The tools themselves diffuse rapidly and symmetrically across rivals; the capability to calibrate reliance on them does not. Heterogeneity in how firms govern what AI-generated content is allowed to become a decision premise is therefore itself a plausible source of competitive heterogeneity. In strategizing—where decisions guide and constrain other decisions—competitive advantage may hinge less on whether firms use AI than on how they calibrate reliance upon it at the moments their commitments are set.

6. References

- Aguilar, F. J. (1967). *Scanning the business environment*. Macmillan.
- Alt, E., & Craig, J. B. (2016). Selling Issues with Solutions: Igniting Social Intrapreneurship in for-Profit Organizations. *Journal of Management Studies*, 53(5), 794–820.
<https://doi.org/10.1111/joms.12200>
- Anderson, M. H., & Nichols, M. L. (2007). Information Gathering and Changes in Threat and Opportunity Perceptions. *Journal of Management Studies*, 44(3), 367–387.
<https://doi.org/10.1111/j.1467-6486.2006.00678.x>
- Arsel, Z. (2017). Asking Questions with Reflexive Focus: A Tutorial on Designing and Conducting Interviews. *Journal of Consumer Research*, 44(4), 939–948.
<https://doi.org/10.1093/jcr/ucx096>
- Baig, A., Yee, L., Singla, A., & Sukharevsky, A [Alexander]. (2024). *What is generative AI?* McKinsey & Company. <https://www.mckinsey.com/featured-insights/mckinsey-explainers/what-is-generative-ai>
- Bain, J. S. (1956). *Barriers to New Competition: Their Character and Consequences in Manufacturing Industries*. *Harvard University Series on Competition in American Industry: Vol. 3*. Harvard University Press.
<https://www.degruyter.com/isbn/9780674188037>
<https://doi.org/10.4159/harvard.9780674188037>
- Balasubramanian, N., Ye, Y., & Xu, M. (2022). Substituting Human Decision-Making with Machine Learning: Implications for Organizational Learning. *Academy of Management Review*, 47(3), 448–465. <https://doi.org/10.5465/amr.2019.0470>
- Bankins, S., Ocampo, A. C., Marrone, M., Restubog, S. L. D., & Woo, S. E. (2024). A multilevel review of artificial intelligence in organizations: Implications for organizational behavior

- research and practice. *Journal of Organizational Behavior*, 45(2), 159–182.
<https://doi.org/10.1002/job.2735>
- Bansal, P., Kim, A., & Wood, M. O. (2018). Hidden in Plain Sight: The Importance of Scale in Organizations’ Attention to Issues. *Academy of Management Review*, 43(2), 217–241.
<https://doi.org/10.5465/amr.2014.0238>
- Barnard, C. I. (1938). *The functions of the executive*. Harvard University Press.
- Bean, R. (2017). *How big data is empowering AI and machine learning at scale*.
<https://sloanreview.mit.edu/article/how-big-data-is-empowering-ai-and-machine-learning-at-scale/>
- Bearman, M. (2019). Focus on Methodology: Eliciting rich data: A practical approach to writing semi-structured interview schedules. *Focus on Health Professional Education: A Multi-Professional Journal*, 20(3), 1–11. <https://doi.org/10.11157/fohpe.v20i3.387>
- Berente, N., Gu, B., Recker, J., & Santhanam, R. (2021). Managing Artificial Intelligence. *MIS Quarterly*, 45(3), 1433–1450. <https://doi.org/10.25300/MISQ/2021/16274>
- Bouschery, S. G., Blazevec, V., & Piller, F. T. (2023). Augmenting human innovation teams with artificial intelligence: Exploring transformer-based language models. *Journal of Product Innovation Management*, 40(2), 139–153. <https://doi.org/10.1111/jpim.12656>
- Boussioux, L., Lane, J. N., Zhang, M., Jacimovic, V., & Lakhani, K. R. (2024). The Crowdless Future? Generative AI and Creative Problem-Solving. *Organization Science*, 35(5), 1589–1607. <https://doi.org/10.1287/orsc.2023.18430>
- Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D. M., Wu, J., Winter, C., . . . Amodei, D.

- (2020, May 28). *Language Models are Few-Shot Learners*.
<https://arxiv.org/pdf/2005.14165.pdf>
- Bubeck, S., Chandrasekaran, V., Eldan, R., Gehrke, J., Horvitz, E., Kamar, E., Lee, P., Lee, Y. T., Li, Y., Lundberg, S., Nori, H., Palangi, H., Ribeiro, M. T., & Zhang, Y. (2023, March 22). *Sparks of Artificial General Intelligence: Early experiments with GPT-4*.
<https://arxiv.org/pdf/2303.12712.pdf>
- Burgelman, R. A., Floyd, S. W., Laamanen, T., Mantere, S., Vaara, E., & Whittington, R. (2018). Strategy processes and practices: Dialogues and intersections. *Strategic Management Journal*, 39(3), 531–558. <https://doi.org/10.1002/smj.2741>
- Burrell, J. (2016). How the machine ‘thinks’: Understanding opacity in machine learning algorithms. *Big Data & Society*, 3(1), Article 2053951715622512.
<https://doi.org/10.1177/2053951715622512>
- Camuffo, A., Gambardella, A., Kazemi, S., & Pandey, A. (2026). Beyond Black Boxes: Designing and Testing Agentic AI Systems for Strategy. *Strategy Science*, 11(1), 137–156.
<https://doi.org/10.1287/stsc.2025.0432>
- Chandler, A. D. (1962). *Strategy and Structure: Chapters in the history of the industrial enterprise*. The MIT Press.
- Charmaz, K. (2014). Grounded Theory in Global Perspective. *Qualitative Inquiry*, 20(9), 1074–1084. <https://doi.org/10.1177/1077800414545235>
- Chattopadhyay, P., Glick, W. H., & Huber, G. P. (2001). Organizational Action in Response to Threats and Opportunities. *Academy of Management Journal*, 44(5), 937–955.
<https://doi.org/10.2307/3069439>

- Chui, M., Yee, L., Hall, B., Singla, A., & Sukharevsky, A [Alex]. (2023). *The state of AI in 2023: Generative AI's breakout year*. QuantumBlack AI, by McKinsey.
- Constantiou, I., Joshi, M., & Stelmaszak, M. (2023). Organizations as Digital Enactment Systems: A Theory of Replacement of Humans by Digital Technologies in Organizational Scanning, Interpretation, and Learning. *Journal of the Association for Information Systems*, 24(6), 1770–1798. <https://doi.org/10.17705/1jais.00833>
- Creswell, J. W., & Creswell, J. D. (2023). *Research design: Qualitative, quantitative, and mixed methods approaches* (Sixth edition, international student edition). SAGE.
- Creswell, J. W., & Poth, C. N. (2025). *Qualitative inquiry & research design: Choosing among five approaches* (Fifth edition). SAGE.
- Csaszar, F. A., Ketkar, H., & Kim, H. (2024). Artificial Intelligence and Strategic Decision-Making: Evidence from Entrepreneurs and Investors. *Strategy Science*, 9(4), Article stsc.2024.0190, 322–345. <https://doi.org/10.1287/stsc.2024.0190>
- Daft, R. L., & Weick, K. E. (1984). Toward a Model of Organizations as Interpretation Systems. *Academy of Management Review*, 9(2), 284–295. <https://doi.org/10.2307/258441>
- Dameron, S., & Torset, C. (2014). The Discursive Construction of Strategists' Subjectivities: Towards a Paradox Lens on Strategy. *Journal of Management Studies*, 51(2), 291–319. <https://doi.org/10.1111/joms.12072>
- Davidson, C. (2009). Transcription: Imperatives for Qualitative Research. *International Journal of Qualitative Methods*, 8(2), 35–52. <https://doi.org/10.1177/160940690900800206>
- Dell'Acqua, F., McFowland III, E., Mollick, E., Lifshitz-Assaf, H., Kellogg, K. C., Rajendran, S., Kraye, L., Candelon, F., & Lakhani, K. R. (2023). Navigating the Jagged Technological Frontier: Field Experimental Evidence of the Effects of AI on Knowledge Worker

- Productivity and Quality. *Harvard Business School Technology & Operations Mgt. Unit Working Paper*, No. 24(013), 1–58.
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology. General*, 144(1), 114–126. <https://doi.org/10.1037/xge0000033>
- Doshi, A. R., Bell, J. J., Mirzayev, E., & Vanneste, B. S. (2025). Generative artificial intelligence and evaluating strategic decisions. *Strategic Management Journal*, 46(3), Article smj.3677, 583–610. <https://doi.org/10.1002/smj.3677>
- Dutton, J. E., & Ashford, S. J. (1993). Selling Issues to Top Management. *Academy of Management Review*, 18(3), 397–428. <https://doi.org/10.2307/258903>
- Dutton, J. E., Ashford, S. J., O'Neill, R. M., Hayes, E., & Wierba, E. E. (1997). Reading the wind: how middle managers assess the context for selling issues to top managers. *Strategic Management Journal*, 18(5), 407–423. [https://doi.org/10.1002/\(SICI\)1097-0266\(199705\)18:5<407::AID-SMJ881>3.0.CO;2-J](https://doi.org/10.1002/(SICI)1097-0266(199705)18:5<407::AID-SMJ881>3.0.CO;2-J)
- Dutton, J. E., & Duncan, R. B. (1987). The creation of momentum for change through the process of strategic issue diagnosis. *Strategic Management Journal*, 8(3), 279–295. <https://doi.org/10.1002/smj.4250080306>
- Dutton, J. E., Fahey, L., & Narayanan, V. K. (1983). Toward understanding strategic issue diagnosis. *Strategic Management Journal*, 4(4), 307–323. <https://doi.org/10.1002/smj.4250040403>
- Dutton, J. E., & Jackson, S. E. (1987). Categorizing Strategic Issues: Links to Organizational Action. *Academy of Management Review*, 12(1), 76–90. <https://doi.org/10.2307/257995>

- Eapen, T. T., Finkenstadt, D. J., Folk, J., & Venkataswamy, L. (2023). How generative AI can augment human creativity. *Harvard Business Review*, 101(7-8), 55–64.
- Einola, K., & Khoreva, V. (2023). Best friend or broken tool? Exploring the co-existence of humans and artificial intelligence in the workplace ecosystem. *Human Resource Management*, 62(1), 117–135. <https://doi.org/10.1002/hrm.22147>
- Eisenhardt, K. M., & Zbaracki, M. J. (1992). Strategic decision making. *Strategic Management Journal*, 13(S2), 17–37. <https://doi.org/10.1002/smj.4250130904>
- Felin, T., & Holweg, M. (2024). Theory Is All You Need: AI, Human Cognition, and Causal Reasoning. *Strategy Science*, 9(4), 346–371. <https://doi.org/10.1287/stsc.2024.0189>
- Feuerriegel, S., Hartmann, J., Janiesch, C., & Zschech, P. (2024). Generative AI. *Business & Information Systems Engineering*, 66(1), 111–126. <https://doi.org/10.1007/s12599-023-00834-7>
- Flick, U., Kardorff, E. von, & Steinke, I. (2004). *A companion to qualitative research*. Sage Publications.
- Follmer, E. H., Talbot, D. L., Kristof-Brown, A. L., Astrove, S. L., & Billsberry, J. (2018). Resolution, Relief, and Resignation: A Qualitative Study of Responses to Misfit at Work. *Academy of Management Journal*, 61(2), 440–465. <https://doi.org/10.5465/amj.2014.0566>
- Gebauer, J., Shaw, M. J., & Gribbins, M. L. (2010). Task-Technology Fit for Mobile Information Systems. *Journal of Information Technology*, 25(3), 259–272. <https://doi.org/10.1057/jit.2010.10>
- Geerling, W., Mateer, G. D., Wooten, J., & Damodaran, N. (2023). ChatGPT has Mastered the Principles of Economics: Now What? *SSRN Electronic Journal*. Advance online publication. <https://doi.org/10.2139/ssrn.4356034>

- Gephart, R. P., Topal, C., & Zhang, Z. (2010). Future-oriented Sensemaking: Temporalities and Institutional Legitimation. In T. Hernes & S. Maitlis (Eds.), *Perspectives on process organization studies. Process, sensemaking, and organizing* (pp.275–312). Oxford University Press. <https://doi.org/10.1093/acprof:oso/9780199594566.003.0013>
- Ghemawat, P. (1991). *Commitment: The Dynamic of Strategy*. The Free Press.
- Gioia, D. A [Dennis A.], & Chittipeddi, K. (1991). Sensemaking and sensegiving in strategic change initiation. *Strategic Management Journal*, 12(6), 433–448. <https://doi.org/10.1002/smj.4250120604>
- Girotra, K., Meincke, L., Terwiesch, C., & Ulrich, K. T. (2023). Ideas are Dimes a Dozen: Large Language Models for Idea Generation in Innovation. *SSRN Electronic Journal*. Advance online publication. <https://doi.org/10.2139/ssrn.4526071>
- Glaser, B. G., & Strauss, A. L. (1980). *The discovery of grounded theory: Strategies for qualitative research* (11th printing). Aldine.
- Glikson, E., & Woolley, A. W. (2020). Human Trust in Artificial Intelligence: Review of Empirical Research. *Academy of Management Annals*, 14(2), 627–660. <https://doi.org/10.5465/annals.2018.0057>
- Goodhue, D. L., & Thompson, R. L. (1995). Task-Technology Fit and Individual Performance. *MIS Quarterly*, 19(2), 213–236. <https://doi.org/10.2307/249689>
- Gray, H. M., Gray, K., & Wegner, D. M. (2007). Dimensions of mind perception. *Science (New York, N.Y.)*, 315(5812), 619. <https://doi.org/10.1126/science.1134475>
- Hahn, J., & Wang, T. (2009). Knowledge management systems and organizational knowledge processing challenges: A field experiment. *Decision Support Systems*, 47(4), 332–342. <https://doi.org/10.1016/j.dss.2009.03.001>

- Halcomb, E. J., & Davidson, P. M. (2006). Is verbatim transcription of interview data always necessary? *Applied Nursing Research : ANR*, 19(1), 38–42.
<https://doi.org/10.1016/j.apnr.2005.06.001>
- Hambrick, D. C. (1982). Environmental scanning and organizational strategy. *Strategic Management Journal*, 3(2), 159–174. <https://doi.org/10.1002/smj.4250030207>
- Hambrick, D. C., & Mason, P. A. (1984). Upper Echelons: The Organization as a Reflection of Its Top Managers. *Academy of Management Review*, 9(2), 193–206.
<https://doi.org/10.2307/258434>
- Hannigan, T. R., McCarthy, I. P., & Spicer, A. (2024). Beware of botshit: How to manage the epistemic risks of generative chatbots. *Business Horizons*, 67(5), 471–486.
<https://doi.org/10.1016/j.bushor.2024.03.001>
- Helfat, C. E., & Peteraf, M. A. (2015). Managerial cognitive capabilities and the microfoundations of dynamic capabilities. *Strategic Management Journal*, 36(6), 831–850.
<https://doi.org/10.1002/smj.2247>
- Hoff, K. A., & Bashir, M. (2015). Trust in automation: Integrating empirical evidence on factors that influence trust. *Human Factors*, 57(3), 407–434.
<https://doi.org/10.1177/0018720814547570>
- Howard, M. C., & Rose, J. C. (2019). Refining and extending task–technology fit theory: Creation of two task–technology fit scales and empirical clarification of the construct. *Information & Management*, 56(6), 103134. <https://doi.org/10.1016/j.im.2018.12.002>
- Hutzschenreuter, T., & Kleindienst, I. (2006). Strategy-Process Research: What Have We Learned and What Is Still to Be Explored. *Journal of Management*, 32(5), 673–720.
<https://doi.org/10.1177/0149206306291485>

- Hutzschenreuter, T., & Lämmermann, T. (2025). What Is Your AI Strategy? Systematically Integrating Self-Learning Technologies into Your Business Strategy. *Academy of Management Perspectives*, 39(4), Article amp.2023.0243, 528–551. <https://doi.org/10.5465/amp.2023.0243>
- Isabella, L. A., & Waddock, S. A. (1994). Top Management Team Certainty: Environmental Assessments, Teamwork, and Performance Implications. *Journal of Management*, 20(4), 835–858. <https://doi.org/10.1177/014920639402000407>
- Jackson, S. E., & Dutton, J. E. (1988). Discerning Threats and Opportunities. *Administrative Science Quarterly*, 33(3), 370–387. <https://doi.org/10.2307/2392714>
- Jalonen, K., Schildt, H., & Vaara, E. (2018). Strategic concepts as micro-level tools in strategic sensemaking. *Strategic Management Journal*, 39(10), 2794–2826. <https://doi.org/10.1002/smj.2924>
- Jarvenpaa, S., & Klein, S. (2024). New Frontiers in Information Systems Theorizing: Human-gAI Collaboration. *Journal of the Association for Information Systems*, 25(1), 110–121. <https://doi.org/10.17705/1jais.00868>
- Jarzabkowski, P., Balogun, J., & Seidl, D. (2007). Strategizing: The challenges of a practice perspective. *Human Relations*, 60(1), 5–27. <https://doi.org/10.1177/0018726707075703>
- Jarzabkowski, P., & Kaplan, S. (2015). Strategy tools-in-use: A framework for understanding “technologies of rationality” in practice. *Strategic Management Journal*, 36(4), 537–558. <https://doi.org/10.1002/smj.2270>
- Jarzabkowski, P., & Paul Spee, A. (2009). Strategy-as-practice: A review and future directions for the field. *International Journal of Management Reviews*, 11(1), 69–95. <https://doi.org/10.1111/j.1468-2370.2008.00250.x>

- Jarzabkowski, P., & Whittington, R. (2008). A Strategy-as-Practice Approach to Strategy Research and Education. *Journal of Management Inquiry*, 17(4), 282–286.
<https://doi.org/10.1177/1056492608318150>
- Jeyaraj, A. (2022). A meta-regression of task-technology fit in information systems research. *International Journal of Information Management*, 65, 102493.
<https://doi.org/10.1016/j.ijinfomgt.2022.102493>
- Ji, Z., Lee, N., Frieske, R., Yu, T., Su, D., Xu, Y., Ishii, E., Bang, Y. J., Madotto, A., & Fung, P. (2023). Survey of Hallucination in Natural Language Generation. *ACM Computing Surveys*, 55(12), 1–38. <https://doi.org/10.1145/3571730>
- Junglas, I., Abraham, C., & Watson, R. T. (2008). Task-technology fit for mobile locatable information systems. *Decision Support Systems*, 45(4), 1046–1057.
<https://doi.org/10.1016/j.dss.2008.02.007>
- Kaplan, S. (2008). Framing Contests: Strategy Making Under Uncertainty. *Organization Science*, 19(5), 729–752. <https://doi.org/10.1287/orsc.1070.0340>
- Kaplan, S., & Orlikowski, W. J. (2013). Temporal Work in Strategy Making. *Organization Science*, 24(4), 965–995. <https://doi.org/10.1287/orsc.1120.0792>
- Keding, C., & Meissner, P. (2021). Managerial overreliance on AI-augmented decision-making processes: How the use of AI-based advisory systems shapes choice behavior in R&D investment decisions. *Technological Forecasting and Social Change*, 171, 120970.
<https://doi.org/10.1016/j.techfore.2021.120970>
- Krakowski, S. (2025). Human-AI agency in the age of generative AI. *Information and Organization*, 35(1), 100560. <https://doi.org/10.1016/j.infoandorg.2025.100560>

- Krakowski, S., Luger, J., & Raisch, S. (2023). Artificial intelligence and the changing sources of competitive advantage. *Strategic Management Journal*, 44(6), Article <https://doi.org/10.1002/smj.3387>, 1425–1452. <https://doi.org/10.1002/smj.3387>
- Krogh, G. von, Ben-Menahem, S., & Shrestha, Y. R. (2021). Artificial Intelligence in Strategizing: Prospects and Challenges. In I. M. Duhaime, M. A. Lyles, & M. A. Hitt (Eds.), *Strategic Management: State of the Field and Its Future* (pp. 625–646). Oxford University Press.
- Laamanen, T., Maula, M., Kajanto, M., & Kunnas, P. (2018). The role of cognitive load in effective strategic issue management. *Long Range Planning*, 51(4), 625–639. <https://doi.org/10.1016/j.lrp.2017.03.001>
- Lebovitz, S., Lifshitz-Assaf, H., & Levina, N. (2022). To Engage or Not to Engage with AI for Critical Judgments: How Professionals Deal with Opacity When Using AI for Medical Diagnosis. *Organization Science*, 33(1), 126–148. <https://doi.org/10.1287/orsc.2021.1549>
- 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
- Leiblein, M. J., Reuer, J. J., & Zenger, T. (2018). What Makes a Decision Strategic? *Strategy Science*, 3(4), 558–573. <https://doi.org/10.1287/stsc.2018.0074>
- Levinthal, D. A. (1997). Adaptation on Rugged Landscapes. *Management Science*, 43(7), 934–950. <https://doi.org/10.1287/mnsc.43.7.934>
- de Liu, Santhanam, R., & Webster, J. (2017). Toward Meaningful Engagement: A Framework for Design and Research of Gamified Information Systems. *MIS Quarterly*, 41(4), 1011–1034. <https://doi.org/10.25300/MISQ/2017/41.4.01>

- Logg, J. M., Minson, J. A., & Moore, D. A. (2019). Algorithm appreciation: People prefer algorithmic to human judgment. *Organizational Behavior and Human Decision Processes*, 151, 90–103. <https://doi.org/10.1016/j.obhdp.2018.12.005>
- Maitlis, S. (2005). The Social Processes of Organizational Sensemaking. *Academy of Management Journal*, 48(1), 21–49. <https://doi.org/10.5465/amj.2005.15993111>
- Maitlis, S., & Christianson, M. (2014). Sensemaking in Organizations: Taking Stock and Moving Forward. *Academy of Management Annals*, 8(1), 57–125. <https://doi.org/10.1080/19416520.2014.873177>
- Malik, M., Andargoli, A., Tallon, P., & Wickramasinghe, N. (2025). An organisational sensemaking theorising of how firms construct digital-enabled strategic agility. *Information & Management*, 104130. <https://doi.org/10.1016/j.im.2025.104130>
- March, J. G., & Simon, H. A. (1958). *Organizations*. John Wiley & Sons.
- McCarthy, J., Minsky, M. L., Rochester, N., & Shannon, C. E. (1955). A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence.
- Miller, K. D., & Lin, S.-J. (2022). Strategic issue diagnosis by top management teams: A multiple-agent model. *Strategic Organization*, 20(3), 600–626. <https://doi.org/10.1177/1476127021993792>
- Muchenje, G., & Seppänen, M. (2023). Unpacking task-technology fit to explore the business value of big data analytics. *International Journal of Information Management*, 69, 102619. <https://doi.org/10.1016/j.ijinfomgt.2022.102619>
- Murphy, K. P. (2023). *Probabilistic machine learning: Advanced topics. Adaptive computation and machine learning series*. The MIT Press.

- Murray, A., Rhymer, J., & Sirmon, D. G. (2021). Humans and Technology: Forms of Conjoined Agency in Organizations. *Academy of Management Review*, 46(3), 552–571.
<https://doi.org/10.5465/amr.2019.0186>
- Myers, M. D. (2020). *Qualitative research in business and management* (Third edition). SAGE.
- Myers, M. D., & Newman, M. (2007). The qualitative interview in IS research: Examining the craft. *Information and Organization*, 17(1), 2–26.
<https://doi.org/10.1016/j.infoandorg.2006.11.001>
- O'Connor, C., & Joffe, H. (2020). Intercoder Reliability in Qualitative Research: Debates and Practical Guidelines. *International Journal of Qualitative Methods*, 19, Article 1609406919899220. <https://doi.org/10.1177/1609406919899220>
- Ocasio, W. (1997). Towards an Attention-Based View of the Firm. *Strategic Management Journal*, 18(S1), 187–206. [https://doi.org/10.1002/\(SICI\)1097-0266\(199707\)18:1+<187::AID-SMJ936>3.3.CO;2-B](https://doi.org/10.1002/(SICI)1097-0266(199707)18:1+<187::AID-SMJ936>3.3.CO;2-B)
- Parasuraman, R., & Riley, V. (1997). Humans and Automation: Use, Misuse, Disuse, Abuse. *Human Factors*, 39(2), 230–253. <https://doi.org/10.1518/001872097778543886>
- Plambeck, N., & Weber, K. (2010). When the glass is half full and half empty: CEOs' ambivalent interpretations of strategic issues. *Strategic Management Journal*, 31(7), 689–710.
<https://doi.org/10.1002/smj.835>
- Porac, J. F., Thomas, H., & Baden-Fuller, C. (1989). Competitive Groups as Cognitive Communities: The Case of Scottish Knitwear Manufacturers. *Journal of Management Studies*, 26(4), 397–416. <https://doi.org/10.1111/j.1467-6486.1989.tb00736.x>
- Porter, M. E. (1979). How Competitive Forces Shape Strategy. *Harvard Business Review*.
- Porter, M. E. (1996). What is Strategy? *Harvard Business Review*, 74(6), 61–78.

- Posen, H. E., Leiblein, M. J., & Chen, J. S. (2018). Toward a behavioral theory of real options: Noisy signals, bias, and learning. *Strategic Management Journal*, 39(4), 1112–1138. <https://doi.org/10.1002/smj.2757>
- Powers, W. R. (2005). *Transcription techniques for the spoken word*. AltaMira Press.
- Quinn, R. E. (1991). *Beyond rational management: Mastering the paradoxes and competing demands of high performance*. The Jossey-Bass management series. Jossey-Bass.
- Rahwan, I., Cebrian, M., Obradovich, N., Bongard, J., Bonnefon, J.-F., Breazeal, C., Crandall, J. W., Christakis, N. A., Couzin, I. D., Jackson, M. O., Jennings, N. R., Kamar, E., Kloumann, I. M., Larochelle, H., Lazer, D., McElreath, R., Mislove, A., Parkes, D. C., Pentland, A. '., . . . Wellman, M. (2019). Machine behaviour. *Nature*, 568(7753), 477–486. <https://doi.org/10.1038/s41586-019-1138-y>
- Raisch, S., & Fomina, K. (2024). Combining Human and Artificial Intelligence: Hybrid Problem-Solving in Organizations. *Academy of Management Review*, 0(0), Article amr.2021.0421, 1–24. <https://doi.org/10.5465/amr.2021.0421>
- 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>
- Rivkin, J. W. (2000). Imitation of Complex Strategies. *Management Science*, 46(6), 824–844. <https://doi.org/10.1287/mnsc.46.6.824.11940>
- Rumelt, R. P. (1974). *Strategy, Structure and Economic Performance in Large American Industrial Corporations*. Harvard Business Review Press.
- Russell, S. J., & Norvig, P. (2022). *Artificial intelligence: A modern approach* (Global edition). *Pearson Series in Artificial Intelligence*. Pearson.

- Sakellariou, E., & Vecchiato, R. (2022). Foresight, sensemaking, and new product development: Constructing meanings for the future. *Technological Forecasting and Social Change*, 184, 121945. <https://doi.org/10.1016/j.techfore.2022.121945>
- Schöffler, J., Jakubik, J., Vössing, M., Köhl, N., & Satzger, G. (2025). AI Reliance and Decision Quality: Fundamentals, Interdependence, and the Effects of Interventions. *Journal of Artificial Intelligence Research*, 82. <https://doi.org/10.1613/jair.1.15873>
- Simon, H. A. (1947). *Administrative Behavior: a Study of Decision-Making Processes in Administrative Organization* (1st ed.). Macmillan.
- Singla, A., Sukharevsky, A [Alexander], Yee, L., Chui, M., & Hall, B. (2024). The state of AI in early 2024: Gen AI adoption spikes and starts to generate value. *McKinsey & Company*.
- Spencer, H. (1864). *The Principles of Biology* (Vol. 1). Williams and Norgate.
- Spitale, G., Biller-Andorno, N., & Germani, F. (2023). Ai model GPT-3 (dis)informs us better than humans. *Science Advances*, 9(26), 1-9. <https://doi.org/10.1126/sciadv.adh1850>
- Strauss, A. L., & Corbin, J. M. (1998). *Basics of qualitative research: Techniques and procedures for developing grounded theory* (2nd ed.). Sage Publications.
- Strong, D. M., & Volkoff, O. (2010). Understanding Organization—Enterprise System Fit: A Path to Theorizing the Information Technology Artifact1. *MIS Quarterly*, 34(4), 731–756. <https://doi.org/10.2307/25750703>
- Stubbart, C. I. (1989). Managerial Cognition: A Missing Link in Strategic Management Research. *Journal of Management Studies*, 26(4), 325–347. <https://doi.org/10.1111/j.1467-6486.1989.tb00732.x>
- Suri, H. (2011). Purposeful Sampling in Qualitative Research Synthesis. *Qualitative Research Journal*, 11(2), 63–75. <https://doi.org/10.3316/QRJ1102063>

- Sydow, J., Schreyögg, G., & Koch, J. (2009). Organizational Path Dependence: Opening the Black Box. *Academy of Management Review*, 34(4), 689–709.
<https://doi.org/10.5465/amr.2009.44885978>
- Sydow, J., Schreyögg, G., & Koch, J. (2020). On the Theory of Organizational Path Dependence: Clarifications, Replies to Objections, and Extensions. *Academy of Management Review*, 45(4), 717–734. <https://doi.org/10.5465/amr.2020.0163>
- Tapinos, E., & Pyper, N. (2018). Forward looking analysis: Investigating how individuals ‘do’ foresight and make sense of the future. *Technological Forecasting and Social Change*, 126, 292–302. <https://doi.org/10.1016/j.techfore.2017.04.025>
- Thomas, J. B., Clark, S. M., & Gioia, D. A. [D. A.] (1993). Strategic Sensemaking and Organizational Performance: Linkages among Scanning, Interpretation, Action, and Outcomes. *Academy of Management Journal*, 36(2), 239–270.
<https://doi.org/10.2307/256522>
- Thomas, J. B., & McDaniel, R. R. (1990). Interpreting Strategic Issues: Effects of Strategy and the Information-Processing Structure of Top Management Teams. *Academy of Management Journal*, 33(2), 286–306. <https://doi.org/10.2307/256326>
- Thomas, J. B., Shankster, L. J., & Mathieu, J. E. (1994). Antecedents to Organizational Issue Interpretation: The Roles of Single-Level, Cross-Level, and Content Cues. *Academy of Management Journal*, 37(5), 1252–1284.
- Tripsas, M., & Gavetti, G. (2000). Capabilities, cognition, and inertia: evidence from digital imaging. *Strategic Management Journal*, 21(10-11), 1147–1161.
[https://doi.org/10.1002/1097-0266\(200010/11\)21:10/11%3C1147::AID-SMJ128%3E3.0.CO;2-R](https://doi.org/10.1002/1097-0266(200010/11)21:10/11%3C1147::AID-SMJ128%3E3.0.CO;2-R)

- Vanneste, B. S., & Puranam, P. (2025). Artificial Intelligence, Trust, and Perceptions of Agency. *Academy of Management Review*, 50(4), Article amr.2022.0041, 726–744. <https://doi.org/10.5465/amr.2022.0041>
- Venkatraman, N. (1989). The Concept of Fit in Strategy Research: Toward Verbal and Statistical Correspondence. *Academy of Management Review*, 14(3), 423–444. <https://doi.org/10.5465/amr.1989.4279078>
- Wei, J., Tay, Y., Bommasani, R., Raffel, C., Zoph, B., Borgeaud, S., Yogatama, D., Bosma, M., Zhou, D., Metzler, D., Chi, E. H., Hashimoto, T., Vinyals, O., Liang, P., Dean, J., & Fedus, W. (2022, June 15). *Emergent Abilities of Large Language Models*. <http://arxiv.org/pdf/2206.07682>
- Weick, K. E. (1979). *The social psychology* (2. ed.,). *Topics in social psychology*. McGraw-Hill.
- Weick, K. E. (1995). *Sensemaking in organizations. Foundations for organizational science*. Sage Publications.
- Weick, K. E., Sutcliffe, K. M., & Obstfeld, D. (2005). Organizing and the Process of Sensemaking. *Organization Science*, 16(4), 409–421. <https://doi.org/10.1287/orsc.1050.0133>
- Weiser, A.-K., & Krogh, G. von (2023). Artificial intelligence and radical uncertainty. *European Management Review*, 20(4), 711–717. <https://doi.org/10.1111/emre.12630>
- Whittington, R. (1996). Strategy as practice. *Long Range Planning*, 29(5), 731–735. [https://doi.org/10.1016/0024-6301\(96\)00068-4](https://doi.org/10.1016/0024-6301(96)00068-4)
- Yin, Y., Jia, N., & Wakslak, C. J. (2024). Ai can help people feel heard, but an AI label diminishes this impact. *Proceedings of the National Academy of Sciences of the United States of America*, 121(14), e2319112121. <https://doi.org/10.1073/pnas.2319112121>

Yoon, Y., Guimaraes, T., & O'Neal, Q. (1995). Exploring the Factors Associated with Expert Systems Success. *MIS Quarterly*, 19(1), 83–106. <https://doi.org/10.2307/249712>