

Strategy in the Age of Artificial Intelligence: From Analysis to Final Decision-Making

Vahid Pourshahabi^{1*} 

¹ Associate Prof., Department of Management, Zah.C., Islamic Azad University, Zahedan, Iran.

**Corresponding author: Vahid Pourshahabi (Email: vahid.pourshahabi@iau.ac.ir)*

Received: January 30, 2026 | Revised: March 3, 2026 | Accepted: March 10, 2026 | Published: September 10, 2026

Abstract

Purpose: This study investigates the boundaries of artificial intelligence participation in strategic decision-making and explains why final decision-makers must remain human—a fundamental question unresolved despite AI's growing analytical capabilities.

Methodology: This qualitative research employed a multiple case study design. Using purposive and snowball sampling, 23 senior and middle managers from leading AI-adopting Iranian organisations across manufacturing, financial services, information technology, and consumer goods sectors were interviewed. Data collected through semi-structured interviews and document analysis were analysed using thematic analysis with an abductive approach, following Gioia, Corley, and Hamilton's (2013) three-stage coding method.

Findings: Analysis revealed 64 first-order concepts, 14 second-order themes, and 4 aggregate dimensions. Problem structuredness and system transparency emerged as key contextual determinants. Human-AI interaction dynamics—conditional trust, experience-based judgment, control-accountability alignment, and epistemic de-skilling—shape managerial responses to AI recommendations. The Iranian context presents both structural barriers and local opportunities. Four criteria for determining AI participation boundaries were identified: problem structuredness, outcome irreversibility, data availability, and ethical accountability.

Conclusion: The three-level conceptual model demonstrates that while AI can effectively support all strategy stages, final decision-making must remain human due to ethical-legal accountability, context-based experiential judgment, unforeseeable consequences, and the need to cultivate managerial wisdom. The findings contribute theoretically to human-AI interaction literature and offer practical implications for managers and policymakers.

Keywords: Strategic Decision-Making, Artificial Intelligence, Human-AI Interaction, Problem Structuredness, Control-Accountability Alignment.

Citation: Pourshahabi, V. (2026). Strategy in the Age of Artificial Intelligence: From Analysis to Final Decision-Making. *Journal of Modern Studies in Management & Organization*, 2(4), 71-90.

Publisher's Note: *JMSMO remains neutral* with regard to jurisdictional claims in published material and institutional affiliations.



Copyright: © The Author(s). Authors retain copyright and full publishing rights. This article is an open access article distributed under the terms of the [Creative Commons Attribution 4.0 International License \(CC BY 4.0\)](https://creativecommons.org/licenses/by/4.0/).

1. Introduction

The emergence of artificial intelligence (AI) in the domain of strategic management represents one of the most significant paradigmatic shifts of the past decade. According to data from the Organisation for Economic Co-operation and Development (OECD), the share of firms using AI technologies in member countries increased from 5.6% in 2020 to 14% in 2024 (OECD, 2025, p. 14). However, a considerable gap persists between large and small enterprises; the adoption rate among large firms stands at approximately 40%, whereas for small firms it reaches only 11.9% (OECD, 2025, p. 14). The World Economic Forum's Future of Jobs Report 2025 further indicates that AI and digital technologies will create 170 million new jobs and displace 92 million existing roles by 2030, resulting in a net increase of 78 million jobs (World Economic Forum, 2025). At the same time, 39% of current workforce skills are expected to undergo transformation by 2030 (World Economic Forum, 2025).

The rapid diffusion of this technology raises fundamental questions regarding the boundaries of its involvement in strategic decision-making processes. Recent research published in the Strategic Management Journal demonstrates that large language models (LLMs) can be employed to evaluate strategic options. Although individual assessments produced by these models are often inconsistent and exhibit biases, aggregating evaluations from multiple diverse models yields results that approximate those of human expert judges (Doshi, Bell, Mirzayev, & Vanneste, 2025, pp. 583-610). This finding opens new horizons for the application of AI in the strategy process, yet it simultaneously leaves open the essential question: *Do these capabilities imply that final decision-making should be delegated to AI?*

A study by Grote and colleagues (Grote, et al., 2024) published in the *Academy of Management Review* addresses this question through the lens of "control-accountability alignment." They argue that in many organisations there is a misalignment between "control" (i.e., those who shape the main determinants of a decision) and "accountability" (i.e., those who are answerable for its outcomes). This misalignment becomes considerably more complex in machine learning-based systems, where the boundary between system development and system use becomes blurred. Their proposed solution involves establishing dialogue and negotiation mechanisms among developers, users, and managers to define shared norms and achieve a balanced distribution of control and accountability (Grote, et al., 2024).

At another level of analysis, a study from the University of Bath, also published in the *Academy of Management Review*, examines the consequences of widespread generative AI use on "managerial phronesis" (practical wisdom or experience-based judgment). Lindebaum, Balasubramanian, Ashraf, and Haack (2026) introduce the concept of "epistemic de-skilling"—a process through which managers gradually offload their thinking to generative AI, thereby losing their cognitive capabilities and becoming unable to grasp context, exercise ethical judgment, or anticipate the consequences of their decisions. Their research demonstrates that this phenomenon occurs when managers, under severe time pressure, use AI as a shortcut to problem-solving, neglecting questioning, seeking diverse perspectives, and engaging with organisational realities (Lindebaum, et al., 2026).

Concurrently, an empirical study by Kim and colleagues (Kim, Glaeser, Hillis, Kominers, & Luca, 2026) published in the *Strategic Management Journal* examined this phenomenon at the operational level. In a case study of municipal building inspection, they observed that risk-predicting algorithms, despite achieving higher accuracy than the manual system, had virtually no impact on improving inspection outcomes because inspectors repeatedly disregarded algorithmic priorities and instead acted based on habit or operational considerations (Kim, et al., 2026). This finding suggests that the value of algorithms depends not on their mathematical properties per se, but on *how humans use their recommendations*.

Drawing on contemporary theoretical frameworks in decision-making systems, socio-technical systems theory, and control-accountability alignment theory, this article addresses the question: **"Where is the boundary of AI participation in the strategy process, and why should the final decision-maker remain human?"** Our central argument is that although AI can play an effective role throughout all stages of the strategy process—from data collection and option analysis to evaluation and scenario simulation—the final decision must remain in human hands for four key reasons: *ethical and legal accountability, contextual judgment grounded in experience, understanding of unforeseeable consequences, and the long-term cultivation of managerial wisdom*.

2. Theoretical Foundations and Literature Review

To analyse the role of artificial intelligence in strategic decision-making, this article integrates four complementary theoretical streams: socio-technical systems theory, control-accountability alignment theory, the theory of bounded rationality, and the dynamic capabilities approach. Rather than competing with one another, these four frameworks serve as complementary analytical lenses, each illuminating a different dimension of the complex relationship between humans and AI in strategic decision-making.

2.1. Socio-Technical Systems Theory and the Integration of AI in Organisations

Socio-technical systems theory, rooted in the work of researchers such as Trist and Bamforth at the Tavistock Institute, is founded on the principle that organisations cannot be analysed solely as technical systems or solely as social systems; rather, these two subsystems are so tightly interwoven that changes in one affect the other (Bechky & Davis, 2025). In the context of AI, this framework reminds researchers that the successful implementation of AI in strategic decision-making depends not only on the quality of algorithms but also on the degree of alignment with organisational structures, culture, processes, and the professional identities of employees (Baygi & Huysman, 2026).

Recent field research confirms this insight. For example, an in-depth case study of the multinational company LuxCorp (active in technology and financial services) demonstrated that integrating generative AI into strategic decision-making is a multi-stage process, moving from "initial experimentation" and "use-case application" to "capability scaling" and "institutionalisation with transparent governance" (Baygi & Huysman, 2026). Based on 27

semi-structured interviews with senior managers and specialists in information technology, legal affairs, human resources, product development, and marketing, the study shows that each stage presents its own socio-technical challenges, including "ambiguity regarding AI capabilities," "concerns about reliability and accuracy," and "data governance issues" (Baygi & Huysman, 2026, pp. 15-18). The key finding is that the value of AI in strategic decision-making lies not in its technical capabilities per se, but in *how humans interact with the system's recommendations*—an interaction shaped by the social, cultural, and governance structures of the organisation.

This finding aligns with the research of Kim and colleagues (Kim, Glaeser, Hillis, Kominers, & Luca, 2026). In their case study of municipal building inspection, they observed that risk-predicting algorithms, despite greater accuracy in identifying high-risk buildings, had virtually no effect on improving inspection outcomes because inspectors repeatedly disregarded algorithmic priorities and instead acted based on habit, professional judgment, or operational considerations (Kim, et al., 2026). This research demonstrates that even the most accurate AI systems cannot achieve the expected impact on organisational decision-making unless they are aligned with social structures and work processes.

2.2. Control-Accountability Alignment Theory

One of the most influential theoretical frameworks in AI governance is "control-accountability alignment theory," proposed by Grote, Parker, and Crowston (2024) in the *Academy of Management Review*. This theory emphasises that in many organisations there is a misalignment between "control" (i.e., those who shape the primary determinants of a decision) and "accountability" (i.e., those who are answerable for the decision's outcomes). This misalignment becomes significantly more complex in machine learning-based systems, where the boundary between system development and system use becomes blurred (Grote, et al., 2024, p. 6).

Grote and colleagues identify three types of misalignment in AI systems: first, a situation where users have full control over how the system is used but developers are held accountable for outcomes (high control for users, high accountability for developers); second, a situation where users have limited control but still bear full accountability; and third, a situation where even developers do not have full control over the system's behaviour (due to the non-deterministic and adaptive nature of complex models), forcing the organisation to maintain control through continuous learning and adaptation (Grote, et al., 2024, p. 9). Their proposed solution to overcome these misalignments is to create dialogue and negotiation mechanisms among developers, users, and managers to define shared norms and achieve a balanced distribution of control and accountability (Grote, et al., 2024, p. 12). In other words, accountability for decisions made with AI assistance cannot be fulfilled merely by a "manager's signature"; rather, it requires a transparent distribution of roles, responsibilities, and oversight mechanisms throughout the entire chain of system production and use.

This perspective is consistent with Crowston's (2025) keynote address at the Association for Computing Machinery (ACM). Crowston argues that the central challenge of using AI in

organisations is the widespread misalignment between control and accountability—cases in which individuals are held accountable for decisions over which they have no control, or conversely, individuals have decision-making power without being accountable for the outcomes (Crowston, 2025). He emphasises that this gap is particularly acute in learning systems such as generative AI, due to their non-deterministic and unpredictable behaviour, and he locates the solution not in fault-based or legal-liability approaches, but rather in "integrative negotiations" among stakeholders to foster a shared culture of accountability (Crowston, 2025).

2.3. Bounded Rationality and the Cognitive Boundaries of Managers

The theory of bounded rationality, introduced by Herbert Simon (1947), is founded on the principle that human cognitive capacity for processing information and evaluating all possible alternatives in decision-making is limited. From this perspective, AI can serve as a "cognitive augments," expanding the boundaries of bounded rationality and helping managers process more information and evaluate more alternatives (Lindebaum, Balasubramanian, Ashraf, & Haack, 2026). However, recent research suggests that this increased analytical capacity can also have unintended consequences.

Lindebaum and colleagues (Lindebaum, et al., 2026), in a study published in the *Academy of Management Review*, introduce the concept of "epistemic de-skilling"—a process through which managers gradually offload their thinking to generative AI, thereby losing their cognitive capabilities and becoming unable to grasp context, exercise ethical judgment, or anticipate the consequences of their decisions. Their research demonstrates that this phenomenon occurs when managers, under severe time pressure, use AI as a shortcut to problem-solving, neglecting questioning, seeking diverse perspectives, and engaging with organisational realities. The researchers emphasise that, unlike humans, AI "does not experience the world, does not understand the consequences of decisions, and does not grasp the social and emotional complexities of the workplace; rather, it merely generates responses based on patterns in past data" (Lindebaum, et al., 2026).

However, Lindebaum and colleagues note that AI does not necessarily lead to epistemic de-skilling. They also introduce the concept of "epistemic up-skilling"—a process in which managers use AI as a tool for reflection and deliberation, rather than as a substitute for thinking. In this approach, rather than uncritically accepting AI output, managers use it to challenge their own assumptions, explore alternative scenarios, and test the reasoning behind their decisions (Lindebaum, et al., 2026). This approach is most productive when managers know they are accountable for their decisions; in such circumstances, the explanatory gaps in AI output encourage managers to think more deeply about their choices and their consequences.

2.4. Dynamic Capabilities and the Orchestration Role of AI

The fourth theoretical framework employed in this article is the dynamic capabilities theory, proposed by Teece (2007), which emphasises the organisation's ability to "sense" environmental changes, "seize" opportunities, and "transform" the organisation to adapt to new

conditions. Recent research indicates that AI can play a key role in all three dimensions of dynamic capabilities, provided that organisations develop the capability to "orchestrate" this technology across their entire value chain (ScienceDirect, 2026). This orchestration capability comprises three components: "data orchestration," "model orchestration," and "decision orchestration." The third component—decision orchestration—has a critical impact on strategic coherence, strategic agility, and the management of AI-related risks in environments with high institutional heterogeneity (e.g., multinational corporations operating across different regulatory and normative regimes) (ScienceDirect, 2026).

This view aligns with the framework proposed by Garce, Shrestha, and von Krogh (2026) in an article published in *Business Horizons*. They argue that organisations should view AI not merely as a "tool," but rather as an "organising layer" that shapes how activities are coordinated, attention is allocated, and constraints are imposed throughout the organisation (Garce, et al., 2026, p. 5). This perspective requires a fundamental transformation in the role of managers—from "decision-maker" to "system designer," "supervisor," and "governor"—whose primary task is not to make each individual decision, but to design the decision-making architecture, define the boundaries of authority, and oversee system performance (Garce, et al., 2026, p. 8). These studies suggest that the most important challenge in the age of AI is not the quality of algorithms, but rather the quality of "decision-making system design" and its "governance."

2.5. Integration of Frameworks: The Analytical Model of the Article

Based on the synthesis of the four theoretical frameworks above, the analytical model of this article rests on three main components:

(a) Problem Structuredness as a Key Determinant of AI Participation Type

An empirical study by Shrestha and colleagues (Shrestha, Benyoucef, & Tremblay, 2026), published in *Group Decision and Negotiation* and based on 33 in-depth interviews with C-suite executives in three multinational corporations at the forefront of AI adoption, demonstrates that "problem structuredness" is the most important factor determining the type of AI participation in strategic decision-making (Shrestha, et al., 2026, p. 12). According to this research, highly structured problems (those that are well-defined, have clear criteria, and offer a limited solution space) can be confidently delegated to AI, whereas poorly structured problems (where the problem itself is contested, criteria are ambiguous, and outcome irreversibility increases the cost of error) require "iterative human-AI co-deliberation" (Shrestha, et al., 2026, p. 14).

(b) Control-Accountability Alignment as a Prerequisite for Decision Legitimacy

Drawing on the theory of Grote and colleagues (Grote, et al., 2024), AI participation in strategic decision-making is legitimate only when transparent mechanisms for control-accountability alignment are in place. This entails explicitly determining: who is involved in system design and development? who supplies the input data? who interprets the

system's output? who provides the final signature on the decision? and, most importantly, who is accountable in the event of error? (Grote, et al., 2024, p. 11).

(c) Preserving and Cultivating "Managerial Phronesis" as a Necessary Condition for Final Human Decision-Making

The concept of "phronesis" (practical wisdom)—rooted in Aristotelian philosophy and referring to the ability to make sound judgments in specific and complex situations—has emerged in recent management research as a key distinguishing feature between humans and AI in decision-making (Lindebaum, et al., 2026). According to this view, even the most advanced AI systems lack phronesis, because this capability is developed through lived experience, interaction with others, reflection on ethical implications, and an understanding of socio-cultural contexts. Therefore, preserving phronesis among managers and cultivating it in the new generation of managers is a strategic imperative for organisations in the age of AI (Lindebaum, et al., 2026).

3. Methodology

The present study adopts a qualitative approach and employs a multiple case study design. This choice aligns with leading research on human-AI interaction in strategic decision-making in both Iran and the broader international context (Haghi & Moghaddam, 2025; Pirataei et al., 2024; Baygi & Huysman, 2026).

3.1. Population

The population of this study comprises all senior and middle managers in Iranian organisations who have at least one year of experience using AI systems in decision-making processes. According to the *Iran Artificial Intelligence Index Report 2025* (Sharif University of Technology Center for AI Strategy and Transformation, 2025), organisations operating in four sectors—industry, financial services, information technology, and consumer goods and services—are at the forefront of AI adoption in the country. Accordingly, the sampling framework of this research focuses on these sectors. Eligible organisations in this population employ at least 100 staff members and possess formal decision-making structures.

3.2. Sample Size and Sampling Method

The sample size was determined based on the principle of "theoretical saturation" (Glaser & Strauss, 1967). Drawing on the experience of comparable studies, theoretical saturation was expected to be achieved after conducting 20 to 25 interviews (Pirataei et al., 2024; Shrestha, Benyoucef, & Tremblay, 2026).

Sampling was carried out using purposive sampling combined with snowball sampling. This means that the researcher initiated the sampling process by contacting accessible managers who met the eligibility criteria, and then asked initial participants to introduce other eligible managers. It is emphasised that this is a qualitative study of individuals from the above-defined

population, conducted solely on the basis of managers' personal willingness to participate in anonymous interviews. Consequently, the findings represent the collective views of managers working in leading AI-adopting organisations in Iran and are not attributed to any specific organisation. The organisational and personal identities of all participants remain completely anonymous in the final report, and no data that could identify individuals or organisations will be disclosed.

3.3. Data Collection Instruments

The primary data collection instrument was **semi-structured interviews**. The interview structure was organised around five main axes: (1) the current role of AI in strategic processes; (2) the experience of interacting with AI recommendations in key decisions; (3) mechanisms of trust, transparency, and oversight within the organisation; (4) managers' perceptions of the "boundary" of AI participation; and (5) specific barriers and constraints of the Iranian context (Shrestha, et al., 2026; Baygi & Huysman, 2026). Sample interview questions included:

- Please describe a recent strategic decision in which AI output was used.
- Under what conditions do you trust AI recommendations, and under what conditions do you override them?
- Who in your organisation is accountable for decisions made with AI assistance?
- What types of decisions do you believe should never be delegated to AI? Why?

In addition, **organisational document analysis**—including strategy reports, AI implementation documentation, and decision-meeting minutes—was used as secondary data to enrich the findings (Baygi & Huysman, 2026).

3.4. Data Analysis Method

Data were analysed using thematic analysis with an abductive approach (O'Sullivan & Ravasi, 2026). The analysis process followed three stages: open, axial, and selective coding (Gioia, Corley, & Hamilton, 2013). To enhance the trustworthiness of the findings, strategies of data triangulation, member checking, and peer debriefing were employed (Lincoln & Guba, 1985).

Research Implementation: Data collection was conducted between October and December 2025. After identifying eligible managers through professional networks and snowball sampling, each participant was contacted and the research objectives were explained. Of the 34 managers invited, 23 (67%) agreed to participate in the interviews. The primary reasons for non-participation were time constraints (7 individuals) and concerns about confidentiality (4 individuals). Of the 23 interviewees, 15 interviews were conducted face-to-face at their workplaces, and 8 were conducted online via Skype. Each interview lasted an average of 52 minutes (range: 35 to 75 minutes). All interviews were recorded with the written consent of participants and subsequently transcribed verbatim. The transcription was carried out by the researcher and then randomly verified by another researcher. The transcribed data were analysed using MAXQDA version 2022. In this process, all interview texts were first imported

into the software, and open codes were systematically assigned to relevant segments. After extracting open codes, the software's "axial coding" function was used to group similar codes into second-order themes. Finally, a code matrix was used to identify the relationships between themes and aggregate dimensions.

3.5. Ethical Considerations

All stages of the research were conducted in full compliance with ethical standards. Before each interview, informed consent was obtained from participants, and they were informed of their right to withdraw at any time. The identities of organisations and individuals remain completely anonymous in the final report, and no identifying information is disclosed.

4. Findings

This section presents the findings derived from the analysis of semi-structured interviews with 23 senior and middle managers from leading AI-adopting organisations in Iran. Table 1 summarises the complete profiles of the participants. The participants had a mean age of 45.1 years (range: 34 to 59), mean managerial experience of 9.6 years (range: 4 to 22), and mean experience working with AI systems of 3.1 years (range: 1 to 7). Seventeen participants (74%) held a Master's degree or higher. Sixty-five percent of participants were at the senior management level (CEO, vice president, or director-general) and 35% were at the middle management level. The industrial sectors of participants included manufacturing (35%), financial and banking services (25%), information technology and software (30%), and consumer goods and services (10%). The organisational and personal identities of all participants remain completely anonymous, with codes M1 to M23 used in place of names.

Table 1. Participant Profiles (Source: By author)

Code	Age	Gender	Degree	Management Experience (years)	AI Experience (years)	Industry Sector	Organisational Level
M1	45	Male	PhD	12	4	Manufacturing	Senior
M2	38	Female	Master's	7	3	Financial Services	Middle
M3	52	Male	PhD	18	5	IT & Software	Senior
M4	41	Male	Master's	9	2	Manufacturing	Senior
M5	36	Female	Master's	6	3	Financial Services	Middle
M6	48	Male	PhD	14	6	IT & Software	Senior
M7	55	Male	Master's	20	4	Manufacturing	Senior
M8	34	Female	Master's	4	2	Consumer Goods	Middle
M9	42	Male	PhD	10	5	Financial Services	Senior
M10	39	Female	Master's	7	3	IT & Software	Middle
M11	59	Male	Master's	22	3	Manufacturing	Senior
M12	37	Male	Bachelor's	5	2	Consumer Goods	Middle
M13	46	Female	PhD	11	4	Financial Services	Senior
M14	50	Male	Master's	16	7	IT & Software	Senior
M15	35	Female	Master's	5	2	Manufacturing	Middle
M16	44	Male	PhD	10	3	Financial Services	Senior

Code	Age	Gender	Degree	Management Experience (years)	AI Experience (years)	Industry Sector	Organisational Level
M17	53	Male	Master's	19	4	IT & Software	Senior
M18	40	Female	Master's	8	3	Manufacturing	Middle
M19	47	Male	PhD	13	5	Consumer Goods	Senior
M20	38	Male	Master's	6	2	Financial Services	Middle
M21	56	Male	Master's	21	4	Manufacturing	Senior
M22	43	Female	PhD	9	3	IT & Software	Senior
M23	36	Male	Master's	4	1	Financial Services	Middle

Data analysis using the three-stage coding method (Gioia, Corley, & Hamilton, 2013) resulted in the identification of 64 first-order concepts, 14 second-order themes, and 4 aggregate dimensions. The findings are presented in the form of a conceptual model that illustrates the relationships among contextual factors, human-AI interaction dynamics, and decision-making outcomes.

4.1. First-Order Concepts and Second-Order Themes

Table 2 summarises the first-order concepts (based on participants' own statements) and the derived second-order themes.

Table 2. Sample Open Codes, Second-Order Themes, and Aggregate Dimensions (Source: By author)

Open Codes (First-Order Concepts)	Second-Order Themes	Aggregate Dimensions
"AI prepares the data for us" (M3)	AI's Role in the Strategy Process	Contextual Factors
"The system simulates different scenarios" (M7)	AI's Role in the Strategy Process	Contextual Factors
"We use AI for initial screening of options" (M11)	AI's Role in the Strategy Process	Contextual Factors
"AI doesn't understand organisational context" (M6)	Limitations of AI in Final Decision-Making	Contextual Factors
"The system doesn't know what consequences that decision had before" (M14)	Limitations of AI in Final Decision-Making	Contextual Factors
"The more complex the problem, the less we trust AI" (M2)	Problem Structuredness and Type of Participation	Contextual Factors
"When the problem is repetitive, we trust the system" (M18)	Problem Structuredness and Type of Participation	Contextual Factors
"Our trust in AI depends on its track record" (M5)	Conditional Trust in the System	Interaction Dynamics
"We never trust AI 100%" (M9)	Conditional Trust in the System	Interaction Dynamics
"My years of market experience tell me this decision is not right" (M1)	Experience-Based Judgment	Interaction Dynamics
"In ambiguous situations, we rely on intuition" (M16)	Experience-Based Judgment	Interaction Dynamics
"We are forced to sign, but the system has set the course" (M13)	Control-Accountability Misalignment	Interaction Dynamics
"When an error occurs, no one takes responsibility" (M20)	Control-Accountability Misalignment	Interaction Dynamics
"We write justification reports for decisions" (M4)	Compensatory Mechanisms	Interaction Dynamics
"I feel I'm gradually losing my analytical ability" (M8)	Epistemic De-skilling and Its Risks	Interaction Dynamics
"Young managers don't learn independent judgment skills" (M22)	Epistemic De-skilling and Its Risks	Interaction Dynamics
"Sanctions have restricted access to advanced tools" (M10)	Structural Barriers in Iran	Local Context & Governance
"AI development costs are very high" (M17)	Structural Barriers in Iran	Local Context & Governance

Open Codes (First-Order Concepts)	Second-Order Themes	Aggregate Dimensions
"The shortage of experts is a fundamental problem" (M12)	Structural Barriers in Iran	Local Context & Governance
"Iranian startups in AI have made good progress" (M21)	Local Opportunities	Local Context & Governance
"Local Persian language processing models can be useful" (M15)	Local Opportunities	Local Context & Governance
"We need transparent policies for AI use" (M23)	Governance of Decision Support Systems	Local Context & Governance
"AI oversight committees have been established" (M19)	Governance of Decision Support Systems	Local Context & Governance
"Our criterion is the irreversibility of the decision" (M7)	Criteria for Delegation to AI	Boundary Framework
"If the consequences of a decision are reversible, we trust the system" (M2)	Criteria for Delegation to AI	Boundary Framework
"Decision-making speed has increased" (M3)	Outcomes of AI-Assisted Decision-Making	Boundary Framework
"Human errors have been reduced" (M11)	Outcomes of AI-Assisted Decision-Making	Boundary Framework

Table 3. First-Order Concepts and Second-Order Themes (Full Summary) (Source: By author)

Second-Order Theme	First-Order Concepts (Sample)
AI's Role in the Strategy Process	"AI prepares the data for us"; "The system simulates different scenarios"; "We use AI for initial screening of options"; "AI helps us in competitor analysis"
Limitations of AI in Final Decision-Making	"AI doesn't understand organisational context"; "The system doesn't know what consequences that decision had before"; "AI cannot judge ethical issues"; "In critical situations, AI output is not reliable"
Problem Structuredness and Type of Participation	"For simple problems, we accept AI output"; "The more complex the problem, the less we trust AI"; "We make strategic decisions ourselves"; "When the problem is repetitive, we trust the system"
Conditional Trust in the System	"Our trust in AI depends on its track record"; "If the system has predicted correctly before, we trust it"; "We never trust AI 100%"; "Our trust in the system depends on its transparency"
Control-Accountability Misalignment	"We are forced to sign, but the system has set the course"; "Developers are not accountable for outcomes"; "When an error occurs, no one takes responsibility"; "Managers are accountable but don't have full control"
Compensatory Mechanisms	"We write justification reports for decisions"; "We hold decision review meetings"; "We have internal oversight systems"; "There are mechanisms to challenge system output"
System Transparency as a Prerequisite for Trust	"If the system doesn't explain its reasoning, we don't trust it"; "Transparent systems are more acceptable to us"; "When system output is interpretable, we use it"
Experience-Based Judgment	"My years of market experience tell me this decision is not right"; "AI cannot replace experience"; "In ambiguous situations, we rely on intuition"
Epistemic De-skilling and Its Risks	"I feel I'm gradually losing my analytical ability"; "Beginners trust AI too much"; "Young managers don't learn independent judgment skills"
Structural Barriers in Iran	"Sanctions have restricted access to advanced tools"; "AI development costs are very high"; "The shortage of experts is a fundamental problem"; "Data infrastructure in Iran is weak"
Local Opportunities	"Iranian startups in AI have made good progress"; "Local Persian language processing models can be useful"; "AI can neutralise sanctions"
Criteria for Delegation to AI	"Our criterion is the irreversibility of the decision"; "We assess the cost of error"; "If the consequences of a decision are reversible, we trust the system"
Outcomes of AI-Assisted Decision-Making	"Decision-making speed has increased"; "Analysis quality has improved"; "Human errors have been reduced"; "But sometimes decisions become superficial"
Governance of Decision Support Systems	"We need transparent policies for AI use"; "AI oversight committees have been established"; "Guidelines for AI use have been developed"

4.2. Aggregate Dimensions (Final Conceptual Framework)

The analysis of second-order themes led to the identification of four aggregate dimensions that constitute the conceptual model of the research:

First Dimension: Contextual Determinants of AI Participation

This dimension comprises two main themes: "Problem Structuredness" and "System Transparency." The findings indicate that problem structuredness (the degree of problem definition, repeatability, and outcome irreversibility) is the most important factor determining the extent of AI participation in decision-making. Participants explicitly stated that in highly structured problems (e.g., sales data analysis or demand forecasting), AI can play a primary role, but in poorly structured problems (e.g., entering a new market or making strategic investment decisions), final decision-making must be delegated to humans. One participant (senior manager, manufacturing, code M1) stated: *"The more complex and irreversible the problem, the more the role of AI diminishes and the role of humans becomes more prominent."*

Furthermore, system transparency (output explainability) was identified as a fundamental prerequisite for trust in AI. One participant (senior manager, IT & Software, code M6) remarked: *"If the system does not explain its reasoning, we do not trust it. Transparent systems are more acceptable to us."* Participants emphasised that in the absence of transparency, they would refuse to accept system output even for structured problems.

Second Dimension: Human-AI Interaction Dynamics in the Decision-Making Process

This dimension encompasses four themes: "Conditional Trust in the System," "Experience-Based Judgment," "Compensatory Mechanisms for Control-Accountability Misalignment," and "Epistemic De-skilling." The findings reveal that the relationship between humans and AI in strategic decision-making is dynamic and negotiation-based. Participants reported that they rarely accept AI output without scrutiny and consistently compare it with their own experience and judgment.

Nevertheless, they expressed deep concern about "epistemic de-skilling" among young managers, believing that excessive use of AI undermines independent judgment capabilities. One participant (middle manager, financial services, code M9) stated: *"New managers who join the organisation trust AI too much and do not learn independent analytical skills. This is a major risk for the organisation's future."* Another participant (middle manager, manufacturing, code M8) also noted: *"I feel I'm gradually losing my analytical ability."*

Third Dimension: Challenges and Opportunities of the Iranian Context

This dimension comprises two themes: "Structural Barriers" and "Local Opportunities." The findings indicate that Iranian organisations face multiple challenges in using AI, the most important of which include restrictions arising from sanctions, high development and implementation costs, a shortage of expert personnel, and weak data infrastructure.

Nevertheless, participants also pointed to local opportunities such as the growth of AI startups and the development of local Persian language processing models. One participant (senior manager, IT & Software, code M17) said: *"Despite the sanctions, Iranian startups in*

AI have made considerable progress and can meet domestic needs." In contrast, another participant (senior manager, manufacturing, code M11) referred to structural barriers and stated: *"Sanctions have restricted access to advanced tools and the costs of AI development are very high."*

Fourth Dimension: The Boundary Framework of AI Participation

This dimension forms the core of the conceptual model and includes three themes: "Criteria for Delegation to AI," "Outcomes of AI-Assisted Decision-Making," and "Governance of Decision Support Systems." Based on the findings, four main criteria for determining the boundary of AI participation in strategic decision-making were identified (Table 4). One participant (senior manager, manufacturing, code M7) expressed in this regard: *"Our criterion is the irreversibility of the decision. If the consequences of a decision are reversible, we are willing to accept the risk of delegation to the system."*

Table 4. Criteria for Determining the Boundary of AI Participation (Source: By author)

Criterion	Description	Sample Quote
Problem Structuredness	Degree of problem definition, repeatability, and availability of clear evaluation criteria	"If the problem has been experienced before and has clear criteria, we trust the system"
Outcome Irreversibility	The cost and possibility of error correction in case of a wrong decision	"If the mistake is reversible, we accept the risk of delegating to the system"
Access to Sufficient Data	Availability of sufficient quality data for training the system	"In areas where we lack data, we do not trust AI"
Ethical Accountability	The impact of the decision on stakeholders and ethical values	"Decisions that affect people's lives are never delegated to AI"

4.3. Final Conceptual Model of the Research

Based on the findings from the interview analysis, the conceptual model of the research is depicted as a three-level framework that illustrates how contextual factors, human-AI interaction dynamics, and local conditions determine the boundaries of AI participation in strategic decision-making:

Level 1 – Contextual Determinants: At this level, two key factors—"Problem Structuredness" and "System Transparency"—were identified as the most important preconditions determining the extent of AI participation. These factors shape the initial decision-making framework and determine in which domains AI can play a primary role and in which it should play a supportive role. Problem structuredness itself is determined by three components: "problem definition," "repeatability," and "outcome irreversibility." The more structured the problem (well-defined, repetitive, and reversible), the more prominent the role of AI; the more unstructured the problem (ambiguous, unique, and irreversible), the more critical the human role in final decision-making.

Level 2 – Human-AI Interaction Dynamics: This level shows how managers evaluate, adjust, accept, or reject AI recommendations. Four main mechanisms were identified at this level: "Conditional Trust in the System" (dependent on track record and system transparency), "Experience-Based Judgment" (acting as a corrective mechanism against system recommendations), "Compensatory Mechanisms for Control-Accountability Misalignment" (including justification reports, review meetings, and oversight systems), and "Epistemic

De-skilling" (identified as a potential risk of excessive AI use). This level demonstrates that the human-AI relationship in strategic decision-making is dynamic, negotiation-based, and subject to multiple conditions.

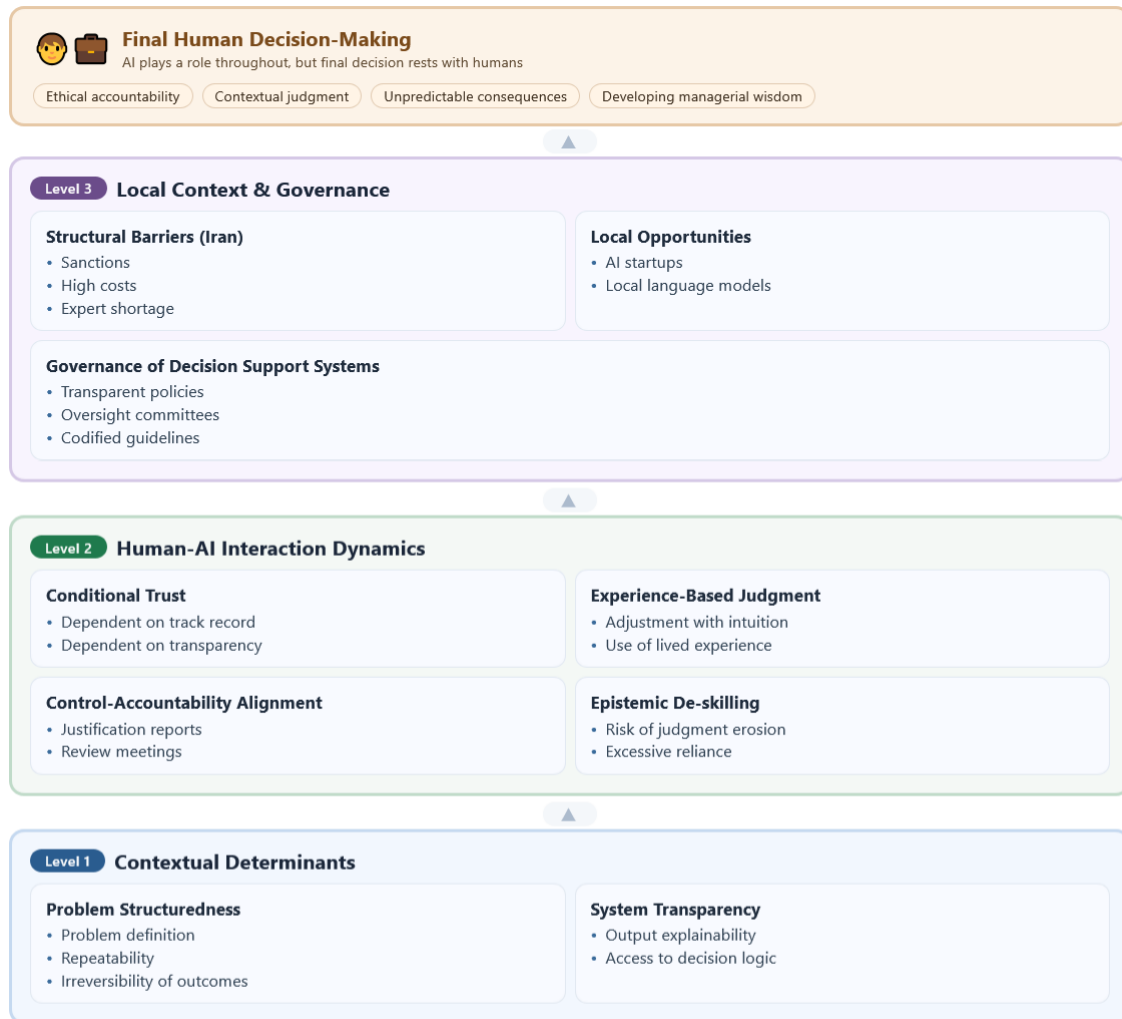


Figure 1: Conceptual Model of AI Participation Boundaries in Strategic Decision-Making (Source: By author)

Level 3 – Local Context and Governance: This level determines the overall context in which decision-making occurs and comprises two components: "Structural Barriers in Iran" (sanctions, high costs, expert shortage, weak data infrastructure) and "Local Opportunities" (growth of AI startups, development of local Persian language models), as well as "Governance of Decision Support Systems" (transparent policies, oversight committees, usage guidelines). This level shows that strategic decision-making with AI assistance does not occur in a vacuum but is influenced by the broader economic, political, institutional, and legal conditions of the country.

Final Outcome – Human Decision-Making: Ultimately, the model shows that despite AI's role throughout all stages of the strategy process (from data analysis to option generation and evaluation), final decision-making always remains in the human domain. The findings identified four key factors that justify this necessity: (1) **Ethical and Legal Accountability** (only humans can be held accountable for decisions); (2) **Contextual Judgment Grounded in Experience** (AI lacks understanding of organisational context and socio-cultural complexities); (3) **Understanding of Unforeseeable Consequences** (AI operates on past data and lacks the capacity to predict unprecedented situations); and (4) **Cultivating Long-Term Managerial Wisdom** (excessive delegation of decisions to AI undermines managers' judgment capabilities).

5. Discussion and Conclusion

5.1. Interpretation of the Conceptual Model

The findings of this research were presented in the form of a three-level model demonstrating that final human decision-making is the outcome of a complex interplay among contextual factors, human-AI interaction dynamics, and the local context of governance. The direction of the arrows in this model (from bottom to top) indicates that **contextual factors** (problem structuredness and system transparency) provide the initial foundation and influence **interaction dynamics** (conditional trust, experience-based judgment, control-accountability alignment, and epistemic de-skilling). These interaction dynamics are in turn shaped by the **local context and governance** (structural barriers in Iran, local opportunities, and governance of decision support systems), and together they culminate in **final human decision-making**.

The key insight of the model is that despite AI's involvement at every stage of the strategy process, **final decision-making always remains in the human domain**—because the four factors of *ethical and legal accountability*, *contextual judgment grounded in experience*, *understanding of unforeseeable consequences*, and *cultivating long-term managerial wisdom* render the complete delegation of final decisions to AI unacceptable.

5.2. Comparison with the Literature

The proposed model is consistent with the theoretical framework of Grote and colleagues (Grote, et al., 2024) on "control-accountability alignment," confirming that in Iranian organisations, misalignment between control and accountability is also a central challenge in using AI for decision-making. Likewise, the findings align with the concept of "epistemic de-skilling" introduced by Lindebaum and colleagues (Lindebaum, et al., 2026), showing that Iranian managers share concerns about the erosion of cognitive capabilities among younger managers due to excessive reliance on AI.

However, this research adds novel dimensions to the existing literature: first, the **role of the "local context"** (sanctions, expert shortages, and startup opportunities), which has received scant attention in prior studies; and second, the **emphasis on "problem structuredness"** as the most important determinant of AI participation, which resonates with the findings of Shrestha and colleagues (Shrestha, Benyoucef, & Tremblay, 2026).

5.3. Theoretical Implications

By presenting a three-level model, this research contributes to theory development in the domain of human-AI interaction in strategic decision-making. The proposed model demonstrates that **final human decision-making** is a multi-level phenomenon that cannot be reduced solely to individual or technical factors. Rather, it is influenced by contextual factors (problem structure and system transparency), interactional factors (trust, experience, and compensatory mechanisms), and macro-institutional factors (local context and governance). This multi-level approach integrates earlier perspectives that had largely concentrated on only one of these levels.

5.4. Practical Implications for Managers and Policymakers

Based on the findings, the following practical recommendations are offered to managers and policymakers:

(a) For Organisational Managers:

- **Diagnose Problem Structuredness:** Before using AI in decision-making, assess the degree of problem structuredness (definition, repeatability, and irreversibility) and determine the appropriate type of AI participation accordingly.
- **Establish Transparency Mechanisms:** Utilise AI systems with high explainability so that managers can understand the reasoning behind recommendations.
- **Align Control and Accountability:** Ensure that those accountable for decisions have sufficient control over the decision-making process and design compensatory mechanisms (e.g., justification reports and review meetings).
- **Preserve and Cultivate Human Judgment:** Through training programmes and opportunities for reflection, prevent epistemic de-skilling among young managers.

(b) For Policymakers and Governing Bodies:

- **Develop Infrastructure:** Invest in data infrastructure, expert training, and support for AI startups to reduce structural barriers.
- **Formulate Governance Frameworks:** Establish transparent policies for the use of AI in organisational decision-making and create oversight committees for decision support systems.
- **Support the Development of Local Models:** Given the constraints imposed by sanctions, promote the development of local AI models tailored to domestic needs and the Persian language.

5.5. Research Limitations

This research has several limitations that should be taken into account when interpreting the findings:

- **Generalisability:** The qualitative nature of the research and the limited sample size (23 managers) restrict the generalisability of the findings to all Iranian organisations. Quantitative research with larger samples could examine the validity of the model across broader populations.
- **Sampling Bias:** The use of snowball sampling may have led to the inclusion of managers with similar professional networks, potentially failing to capture the full diversity of perspectives.
- **Temporal Context:** The research was conducted at a specific point in time, and the rapid evolution of AI technology may necessitate future updates to the findings.
- **Reliance on Self-Reporting:** The data are based on managers' self-reports, which may differ from their actual behaviour in the workplace.

5.6. Suggestions for Future Research

Based on the findings and limitations of this study, the following avenues for future research are proposed:

- **Quantitative Research:** Design and validate a questionnaire based on the three-level model of this research and test the causal relationships among factors in larger samples.
- **Longitudinal Studies:** Investigate changes in human-AI interaction dynamics over time and their impact on strategic decision-making.
- **Comparative Studies:** Compare the proposed model in Iran with other countries (particularly developing and developed nations) to identify contextual differences.
- **The Role of Organisational Culture:** Examine the influence of organisational culture on trust in AI and experience-based judgment mechanisms.
- **Industry-Specific Research:** Investigate the model in different industries (e.g., healthcare, education, and public services) to identify sectoral differences.

5.7. Final Conclusion

This research has demonstrated that while AI can play an effective role at all stages of the strategy process, final decision-making must remain in human hands due to ethical and legal accountability, context-based experiential judgment, the understanding of unforeseeable consequences, and the long-term necessity of cultivating managerial wisdom. The three-level model presented in this study, with its emphasis on contextual factors, interaction dynamics, and local governance, provides a comprehensive framework for analysing the boundaries of AI participation in strategic decision-making. It is hoped that this model can assist both researchers in theory development and managers and policymakers in designing effective and responsible decision-making systems.

Declarations

Author Contributions (CRediT taxonomy)

The author confirms sole responsibility for the following: study conception and design, data collection, analysis and interpretation of results, and manuscript preparation and revision. All aspects of this work were undertaken by the author.

Funding Statement

This research received no specific grant from any funding agency in the public, commercial, or not-for-profit sectors. The author received no financial support for the research, authorship, or publication of this article.

Ethical Statement

Informed consent was obtained from all participants. The study was conducted in accordance with established ethical guidelines for qualitative research.

Authors' AI Use Statement

The author declares that no artificial intelligence (AI) tools were used in the preparation of this manuscript for writing, analysis, or data generation. All content is the result of the author's own human effort and judgment.

Conflict of interest

The author declares no competing financial interests or personal relationships that could have influenced the work reported in this paper.

Data Availability Statement

The data that support the findings of this study are available from the corresponding author upon reasonable request.

References

- Baygi, R. M., & Huysman, M. (2026). From pilots to decision systems: Embedding generative AI into strategic decision-making through a socio-technical and governance lens. *Journal of Decision Systems*, 1-32. <https://doi.org/10.1080/12460125.2025.2597835>
- Bechky, B. A., & Davis, G. F. (2025). Generative AI and the social fabric of organizations. *Strategic Organization*, 1-16. <https://doi.org/10.1177/14761270251385456>
- Crowston, K. (2025). Taming artificial intelligence: A theory of control–accountability alignment among AI developers and users. *Academy of Management Review*, 51(2). <https://doi.org/10.5465/amr.2023.0117>

- Doshi, A. R., Bell, J. J., Mirzayev, E., & Vanneste, B. S. (2025). Generative artificial intelligence and evaluating strategic decisions. *Strategic Management Journal*, 46(3), 583-610. <https://doi.org/10.1002/smj.3677>
- Eisenhardt, K. M. (1989). Building theories from case study research. *Academy of Management Review*, 14(4), 532-550.
- Eisenhardt, K. M., & Graebner, M. E. (2007). Theory building from cases: Opportunities and challenges. *Academy of Management Journal*, 50(1), 25-32.
- Garce, A., Shrestha, A., & von Krogh, G. (2026). Organizing for AI: From tools to organizing layers. *Business Horizons*, 69(2), 1-12.
- Gioia, D. A., Corley, K. G., & Hamilton, A. L. (2013). Seeking qualitative rigor in inductive research: Notes on the Gioia methodology. *Organizational Research Methods*, 16(1), 15-31.
- Glaser, B. G., & Strauss, A. L. (1967). *The discovery of grounded theory: Strategies for qualitative research*. Aldine Publishing.
- Grote, G., Parker, S. K., & Crowston, K. (2024). Taming artificial intelligence: A theory of control-accountability alignment among AI developers and users. *Academy of Management Review*, 1-22. <https://doi.org/10.5465/amr.2023.0117>
- Haghi, M., & Moghaddam, A. (2025). The impact of artificial intelligence on strategic decision-making in Iranian public and private organizations. *Journal of Public Organizations Management*, 12(3), 45-68. [In Persian]
- Kim, H., Glaeser, E. L., Hillis, A., Kominers, S. D., & Luca, M. (2026). Decision authority and the returns to algorithms. *Strategic Management Journal*, 47(1), 1-25.
- Lincoln, Y. S., & Guba, E. G. (1985). *Naturalistic inquiry*. Sage Publications.
- Lindebaum, D., Balasubramanian, N., Ashraf, M., & Haack, P. (2026). Generative AI and the erosion of managerial phronesis. *Academy of Management Review*, 51(2), 1-15.
- OECD. (2025). *OECD Digital Economy Outlook 2025: Artificial Intelligence and the Future of Work*. Paris: OECD Publishing. <https://doi.org/10.1787/5a4f5c7b-en>
- O'Sullivan, A., & Ravasi, D. (2026). The abductive approach in management research: A methodological review. *Academy of Management Annals*, 20(1), 1-35.
- Pirataei, Sh., Gharakhani, H., & Abbasi, H. (2024). Interaction of artificial intelligence and strategic management in Iran's sports industry. *Scientific Quarterly of Sports Management and Motor Behavior Research*, 14(2), 1-22. [In Persian]
- Sadeghi, M., Rezaei, M., & Karimi, S. (2025). Application of artificial intelligence and knowledge management in corporate governance: A case study of MAPNA Company. *Journal of Investment Knowledge*, 14(1), 25-44. [In Persian]
- Sharif University of Technology, Center for AI Strategy and Transformation. (2025). *Iran Artificial Intelligence Index Report 2025*. Tehran: Sharif University of Technology. [In Persian]
- Shojaei, R., Ghasemi, B., & Talebi, M. (2024). Identifying and prioritizing the challenges of implementing artificial intelligence in Iranian organizations. *Quarterly Journal of Information Technology Management*, 16(2), 113-138. [In Persian]
- Shrestha, S., Benyoucef, M., & Tremblay, M. C. (2026). AI in strategic decision-making: The role of problem structuredness. *Group Decision and Negotiation*, 35(1), 1-22.

- Simon, H. A. (1947). *Administrative behavior: A study of decision-making processes in administrative organizations*. Macmillan.
- Teece, D. J. (2007). Explicating dynamic capabilities: The nature and microfoundations of (sustainable) enterprise performance. *Strategic Management Journal*, 28(13), 1319-1350.
- World Economic Forum. (2025). *The Future of Jobs Report 2025*. Geneva: World Economic Forum.
- Yin, R. K. (2018). *Case study research and applications: Design and methods* (6th ed.). Sage Publications.