

## Algorithms in the boardroom: AI decision support and executive judgement

**Anass YACHOULTI, (PhD)**

*Laboratory of Economics and Management of Organizations,  
Faculty of Economics and Management, Kenitra  
Ibn Tofail University of Kenitra, Morocco*

|                                 |                                                                                                                                                                                                                                                                                                              |
|---------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Correspondence address :</b> | Faculty of Economics and Management,<br>Avenue des universites, Kenitra, Morocco<br>Tél : +21205373-29218<br>Communications@uit.ac.ma                                                                                                                                                                        |
| <b>Disclosure Statement :</b>   | The authors declare that they have not received any financial support that could have influenced the objectivity of this study. They take full responsibility for any potential plagiarism, the use of artificial intelligence in the writing process, as well as for the results presented in this article. |
| <b>Conflict of Interest :</b>   | The authors report no conflicts of interest.                                                                                                                                                                                                                                                                 |
| <b>Cite this article :</b>      | YACHOULTI, A. (2026). Algorithms in the boardroom: AI decision support and executive judgement. <i>International Journal of Accounting, Finance, Auditing, Management and Economics</i> , 7(5), 715–740.<br><a href="https://doi.org/10.5281/zenodo.20274638">https://doi.org/10.5281/zenodo.20274638</a>    |
| <b>License</b>                  | <b>This is an open access article under the CC BY-NC-ND license</b>                                                                                                                                                                                                                                          |

Received: 23/04/2026

Accepted: 28/05/2026

**International Journal of Accounting, Finance, Auditing, Management and Economics - IJAFAME**

**ISSN: 2658-8455**

**Volume 6, Issue 05 (2025)**

## Algorithms in the boardroom: AI decision support and executive judgement

### Abstract:

This study examines how AI-generated strategic recommendations influence executives' reliance on algorithmic advice and their strategic decision judgments. Specifically, it investigates whether AI recommendation transparency and perceived analytical superiority shape executives' willingness to rely on algorithmic advice, and whether such reliance serves as the mechanism through which AI affects strategic judgment. Although prior research has extensively examined algorithm aversion, algorithm appreciation, and explainable AI in operational and consumer settings, limited research has explored how AI recommendation characteristics influence executive-level strategic judgment under conditions of uncertainty, ambiguity, and organizational accountability. This study addresses that theoretical and empirical gap by examining behavioral mechanisms through which AI recommendations shape strategic decision-making at the executive level. The study uses a scenario-based, between-subjects experiment with a  $2 \times 2$  vignette design. AI recommendation transparency and perceived analytical superiority were manipulated at high and low levels in a strategic market-entry decision supported by an AI system. Data were collected through an online survey of 312 senior managers and executives from Moroccan organizations. While this context provides valuable insights from an emerging economy where executive AI adoption remains underexplored, the findings should be interpreted cautiously as their generalizability may be limited to similar institutional and technological environments. Hypotheses were tested using Partial Least Squares Structural Equation Modeling (PLS-SEM). The findings show that both AI recommendation transparency ( $\beta = 0.32$ ,  $t = 5.47$ ,  $p < 0.001$ ) and perceived analytical superiority ( $\beta = 0.41$ ,  $t = 7.21$ ,  $p < 0.001$ ) positively influence executives' reliance on algorithmic advice. In turn, reliance on algorithmic advice has a significant positive effect on strategic decision judgments ( $\beta = 0.56$ ,  $t = 10.84$ ,  $p < 0.001$ ). Mediation analysis indicates that reliance on algorithmic advice partially mediates the relationship between AI recommendation transparency and strategic decision judgment (indirect effect  $\beta = 0.18$ ,  $p < 0.001$ ) and between perceived analytical superiority and strategic decision judgment (indirect effect  $\beta = 0.23$ ,  $p < 0.001$ ). These results suggest that AI influences executive judgment not only through analytical capability, but through recommendation characteristics that make AI outputs understandable and credible to decision makers. This study extends research on algorithmic decision making to executive strategic judgment, where uncertainty and complexity are high. It identifies reliance on algorithmic advice as the key mechanism linking AI recommendation characteristics to strategic decision judgments and integrates explainable AI with strategic decision-making research by highlighting transparency as a critical condition for executive adoption of AI recommendations.

**Keywords:** Artificial intelligence, algorithmic advice, strategic decision making, executive judgment, transparency, reliance on algorithmic advice.

**Classification JEL:** M15, O33

**Paper type:** Research paper

## 1. Introduction

Artificial intelligence (AI) is rapidly becoming embedded in strategic decision processes. Firms increasingly use AI-enabled systems to generate recommendations on market entry, pricing, investment allocation, competitive positioning, and risk forecasting, with the promise of extending managerial cognition in complex and uncertain environments (Brynjolfsson and McElheran, 2016; Davenport and Ronanki, 2018; Shrestha et al., 2019). This development is especially consequential in the boardroom, where decisions are high stakes, ambiguous, and difficult to reverse. Yet the growing technical sophistication of AI does not guarantee managerial uptake. Even when algorithmic systems offer analytically powerful recommendations, their organizational value depends on whether executives are willing to rely on them when forming strategic judgments.

This issue is important for both scholarship and practice. For organizations, the strategic value of AI depends not only on predictive capability but also on whether algorithmic outputs are accepted and incorporated by senior decision makers. For researchers, the phenomenon sits at the intersection of strategic decision making, managerial cognition, and human–AI collaboration. Strategic management research has long shown that decision quality depends on how managers process information under conditions of bounded rationality, uncertainty, and cognitive bias (Gavetti et al., 2012). AI systems may expand that information-processing capacity, but their influence on strategy is necessarily filtered through executive judgment.

The specific research problem, therefore, is not whether AI can generate strategic recommendations, but under what conditions executives rely on algorithmic advice and how that reliance shapes their strategic decision judgments. This problem is theoretically important because strategic decisions differ from the operational and forecasting tasks that dominate much of the algorithmic decision-making literature. They involve greater ambiguity, organizational accountability, and contextual interpretation, which may make reliance on algorithmic advice both more consequential and more fragile.

Prior research offers useful but incomplete insight into this issue. Studies of algorithm aversion show that individuals may resist algorithmic advice after observing error, while research on algorithm appreciation suggests that individuals may prefer algorithmic recommendations when algorithms are perceived as analytically superior (Logg et al., 2019). Research on explainable AI further suggests that transparency can increase the legitimacy and acceptance of algorithmic recommendations by making them more understandable and accountable (Kizilcec, 2016). However, these studies have largely focused on operational, consumer, medical, or forecasting contexts where decisions are more structured and outcomes are easier to evaluate. They provide limited insight into executive strategic decisions characterized by ambiguity, accountability pressures, and long-term consequences. In addition, prior research has largely examined trust, adoption intentions, or algorithm aversion in isolation, while giving limited attention to the behavioral mechanism through which AI recommendation characteristics influence strategic judgment. Existing research has not sufficiently explained how characteristics of AI recommendations shape executive reliance on algorithmic advice in strategic settings, nor how such reliance functions as the mechanism through which AI becomes consequential for strategic judgment. Addressing this gap is important because it links the design of AI systems to the behavioral processes through which they influence strategic decision making.

The empirical context of Morocco also offers a relevant setting for this study. While prior research on algorithmic decision-making has been concentrated in highly digitized Western economies, Morocco represents an emerging economy experiencing growing investment in digital transformation and AI adoption across industries. At the same time, firms vary significantly in technological maturity, making this context appropriate for examining

executive responses to AI recommendations while also defining the boundaries of generalizability.

To address this gap, the study draws on three complementary perspectives: information processing theory, socio-technical systems theory, and the augmented intelligence perspective. These perspectives are introduced here as a broad theoretical foundation, while their detailed development is presented in the literature review. Together, these perspectives imply that the strategic influence of AI depends on whether executives perceive algorithmic recommendations as understandable and analytically credible enough to rely on them in forming judgments.

Guided by this framework, the study addresses the following research question: How do AI recommendation transparency and perceived analytical superiority influence executives' reliance on algorithmic advice, and how does such reliance affect strategic decision judgments? This question is operationalized through two related objectives: examining whether transparency and perceived analytical superiority increase reliance on algorithmic advice, and testing whether such reliance mediates their effects on strategic decision judgments.

The conceptual model positions transparency and perceived analytical superiority as antecedents of reliance on algorithmic advice, with reliance acting as the mediating mechanism linking these perceptions to strategic decision judgments. To test this model, the study employs a scenario-based, between-subjects experiment using a  $2 \times 2$  vignette design in which transparency and analytical-superiority cues are manipulated in the context of a strategic market-entry decision supported by an AI analytics system. The hypotheses are tested using Partial Least Squares Structural Equation Modeling (PLS-SEM), which is well suited to estimating mediation relationships among latent constructs in an emerging, prediction-oriented research domain.

This study makes three contributions. First, it extends research on algorithmic decision making into the domain of executive strategic judgment, where uncertainty, accountability, and contextual complexity are more pronounced than in the settings examined in most prior studies. Second, it identifies reliance on algorithmic advice as a central behavioral mechanism through which AI recommendation characteristics influence strategic decision judgments. Third, it integrates explainable AI research with strategic decision-making theory by showing that the interpretability of AI recommendations is not merely a technical design issue but a strategically consequential condition for executive uptake.

The remainder of the article proceeds as follows. The next section reviews the relevant literature and develops the hypotheses. The subsequent section outlines the research design, measures, and analytical approach. The following section presents empirical results, after which the discussion addresses the study's theoretical implications, managerial relevance, limitations, and directions for future research.

## **2. Literature review and hypotheses**

To explain how executives interact with AI-generated recommendations, this section develops the theoretical foundations of the study and reviews prior empirical research on algorithmic decision-making. It first synthesizes existing empirical findings, then presents the theoretical lenses guiding the model, and finally develops hypotheses linking AI recommendations to strategic judgment.

### **2.1. Theoretical Foundations**

This section reviews key theoretical perspectives explaining how organizations process information and integrate technological tools into decision-making. It draws on information processing theory, socio-technical systems theory and augmented intelligence to understand how AI supports managerial judgment.

### **2.1.1. Information Processing Theory**

Information processing theory suggests that organizations develop structures and decision mechanisms that expand their capacity to process information in complex environments (Galbraith, 1974). Strategic decision-making often involves analyzing large volumes of information related to markets, competitors, technologies and financial performance. As environmental complexity increases, decision-makers require additional analytical tools to process this information effectively.

AI decision support systems represent an advanced form of organizational information processing capability. By analyzing large-scale datasets and generating predictive insights, AI systems can reduce uncertainty and support more informed decision-making. Prior research in Management Decision demonstrates that data analytics capabilities enable organizations to better interpret complex information environments and improve strategic outcomes (Aker et al., 2016; Dubey et al., 2019). However, information processing theory also suggests that technological capabilities alone are insufficient. Decision outcomes depend on whether decision-makers actually integrate these analytical insights into their cognitive processes.

### **2.1.2. Socio-Technical Systems Theory**

Socio-technical systems theory provides additional insight into how humans interact with technology in organizational contexts. This perspective emphasizes that organizational performance depends on the joint optimization of social systems (human actors, decision processes) and technical systems (technological tools and algorithms) (Leonardi, 2011). Effective decision outcomes therefore emerge from the interaction between human judgment and technological capabilities rather than from either component alone.

In the context of AI decision support, socio-technical systems theory suggests that executives and algorithms form a hybrid decision system in which human interpretation complements algorithmic analysis (Raisch and Krakowski, 2021). While AI systems can generate predictive insights, human decision-makers remain responsible for interpreting these insights, contextualizing them within organizational objectives and ultimately making strategic choices. The effectiveness of AI-based decision support therefore, depends on the extent to which executives rely on and integrate algorithmic recommendations into their decision processes.

### **2.1.3. Augmented Intelligence Perspective**

The augmented intelligence perspective emphasizes the complementary relationship between humans and AI systems. Rather than replacing human decision-makers, AI technologies augment managerial cognition by providing analytical insights that support human judgment (Davenport and Kirby, 2016).

This perspective suggests that optimal decision outcomes arise when algorithmic analysis and human expertise are combined. If managers ignore algorithmic advice, the potential benefits of AI remain unrealized. Conversely, excessive reliance may create risks such as automation bias and reduced critical reflection (Jarrahi, 2018).

Taken together, these three theoretical perspectives provide complementary explanations for the proposed model. Information processing theory explains why executives may value AI in highly complex environments. Socio-technical systems theory explains why outcomes depend on interactions between executives and AI systems. The augmented intelligence perspective explains why AI creates value when algorithmic capabilities complement rather than replace managerial judgment. Together, these perspectives justify the study's focus on transparency and analytical superiority as antecedents of reliance on algorithmic advice and reliance as the mechanism linking AI recommendations to strategic judgments.

## 2.2. Empirical Background on Algorithmic Decision-Making

Research on algorithmic decision-making has produced mixed findings regarding how individuals respond to algorithmic recommendations. One stream of research emphasizes algorithm aversion, showing that individuals often reject algorithmic recommendations after observing errors, even when algorithms outperform human decision-makers. Dietvorst et al. (2015) demonstrated that people frequently abandon algorithmic advice after minor mistakes because they expect algorithms to perform flawlessly. This resistance may be particularly relevant in strategic contexts where executives face accountability pressures for consequential decisions.

A second stream highlights algorithm appreciation. Logg et al. (2019) found that individuals may prefer algorithmic advice when algorithms are perceived as analytically superior to human judgment. This tendency is especially likely when decisions involve large datasets, forecasting complexity, and analytical tasks that exceed human cognitive capacity.

Recent research has extended these debates to organizational contexts. Shrestha et al. (2019) show that AI systems increasingly reshape organizational decision structures, while Jarrahi (2018) argues that AI often functions as a collaborative decision partner rather than a substitute for human managers. Similarly, Raisch and Krakowski (2021) emphasize that organizational value emerges when AI augments managerial cognition rather than replacing executive judgment.

Research on explainable AI further demonstrates that transparency influences user trust and adoption. Rai (2020) argues that explainability improves the interpretability of AI outputs, while Shin (2021) shows that explainability positively affects trust and acceptance of AI systems.

Despite these contributions, prior empirical studies remain concentrated in operational, healthcare, forecasting, and consumer decision contexts. Limited research has examined executive strategic decision-making, where uncertainty, ambiguity, and organizational accountability are significantly higher. In addition, existing studies rarely examine reliance on algorithmic advice as the behavioral mechanism linking AI recommendation characteristics to strategic judgments. This study addresses these gaps.

## 2.3. Hypotheses development

This section develops the study's hypotheses by examining how key AI system characteristics, namely recommendation transparency and perceived analytical superiority, influence executives' reliance on algorithmic advice. It further explains how such reliance shapes strategic decision judgments and functions as a central mediating mechanism linking AI capabilities to decision outcomes.

### 2.3.1. AI Recommendation Transparency and Reliance on Algorithmic Advice

Transparency is widely recognized as a key design principle in AI systems. In algorithmic contexts, transparency refers to the extent to which users can understand how an AI system generates its recommendations, including the underlying data inputs, modeling assumptions and analytical processes (Rai, 2020). Transparent systems provide explanations that enable decision-makers to interpret algorithmic reasoning and evaluate the credibility of algorithmic outputs.

From an information processing perspective, transparency reduces cognitive uncertainty by allowing executives to evaluate how algorithmic recommendations are generated. When decision-makers understand the analytical logic underlying algorithmic outputs, they are better able to integrate these insights with their own strategic knowledge. Transparency therefore facilitates the cognitive integration of human judgment and algorithmic analysis.

Prior research suggests that transparency enhances the acceptance of algorithmic recommendations by increasing perceived legitimacy and accountability (Burton et al., 2020; Haefner et al., 2021). In strategic contexts, these factors are particularly important because executives must justify their decisions to multiple stakeholders. Transparent AI recommendations provide traceable reasoning that enables decision-makers to explain how data-driven insights informed their choices.

Studies highlight the importance of interpretability and managerial understanding in the effective use of analytics technologies (Aker et al., 2016; Wamba et al., 2017). When analytics systems provide interpretable insights, managers are more likely to incorporate these insights into their decision processes.

Therefore, transparent AI recommendations are expected to increase executives' reliance on algorithmic advice.

Therefore, we hypothesize:

***H1: AI recommendation transparency positively influences executives' reliance on algorithmic advice.***

### **2.3.2. Perceived Analytical Superiority of AI and Reliance on Algorithmic Advice**

Another factor shaping executives' reliance on algorithmic advice is their perception of the analytical capabilities of AI systems. Perceived analytical superiority refers to the belief that algorithmic systems possess greater capacity than humans to analyze complex data and generate accurate predictions.

Research on algorithm appreciation suggests that individuals often rely more heavily on algorithmic advice when they perceive algorithms as analytically superior to human judgment (Logg et al., 2019). AI systems possess several characteristics that may reinforce such perceptions. Machine learning algorithms can process large datasets, identify complex relationships between variables and update predictions dynamically as new information becomes available.

In strategic contexts, these capabilities may be particularly valuable because executives must interpret diverse information sources, including market trends, financial indicators and competitive dynamics. Prior studies show that analytics capabilities enhance organizations' ability to generate insights from complex datasets and support strategic decision-making (Wamba et al., 2020).

When executives perceive AI systems as possessing superior analytical capabilities, they may view algorithmic recommendations as a valuable complement to their own judgment. This perception encourages greater reliance on algorithmic advice when evaluating strategic alternatives.

Therefore, we hypothesize:

***H2: Perceived analytical superiority of AI positively influences executives' reliance on algorithmic advice.***

### **2.3.3. Reliance on Algorithmic Advice and Strategic Decision Judgments**

Reliance on algorithmic advice refers to the extent to which decision-makers incorporate algorithmic recommendations into their judgments rather than relying exclusively on intuition (Dietvorst et al., 2015). Strategic decisions involve evaluating alternatives under uncertainty (Powell et al., 2011). AI systems may improve decision quality by identifying patterns unavailable to human cognition.

Research shows that algorithmic recommendations may reduce cognitive biases such as overconfidence and confirmation bias (Burton et al., 2020). However, prior research also highlights risks associated with excessive reliance on automated systems. Automation bias may lead decision-makers to over-trust flawed algorithmic outputs (Parasuraman and Riley, 1997).

Over-reliance may also reduce critical reflection in ambiguous strategic contexts. Nevertheless, this study focuses on situations where AI functions as a complementary analytical input rather than a substitute for executive judgment. Under these conditions, moderate reliance is expected to improve strategic decision judgments.

Therefore, we hypothesize:

**H3: Executives' reliance on algorithmic advice positively influences their strategic decision judgments.**

#### 2.3.4. The Mediating Role of Reliance on Algorithmic Advice

Although AI system characteristics influence executives' perceptions of algorithmic recommendations, their effects on strategic decision outcomes are unlikely to occur directly. Instead, these characteristics shape managerial cognition by influencing how executives engage with algorithmic advice (Raisch and Krakowski, 2021).

Transparency enables executives to understand the reasoning underlying algorithmic recommendations, increasing their willingness to incorporate these insights into their decision processes (Rai, 2020). Similarly, when executives perceive AI systems as analytically superior, they are more likely to view algorithmic recommendations as valuable sources of information (Logg et al., 2019).

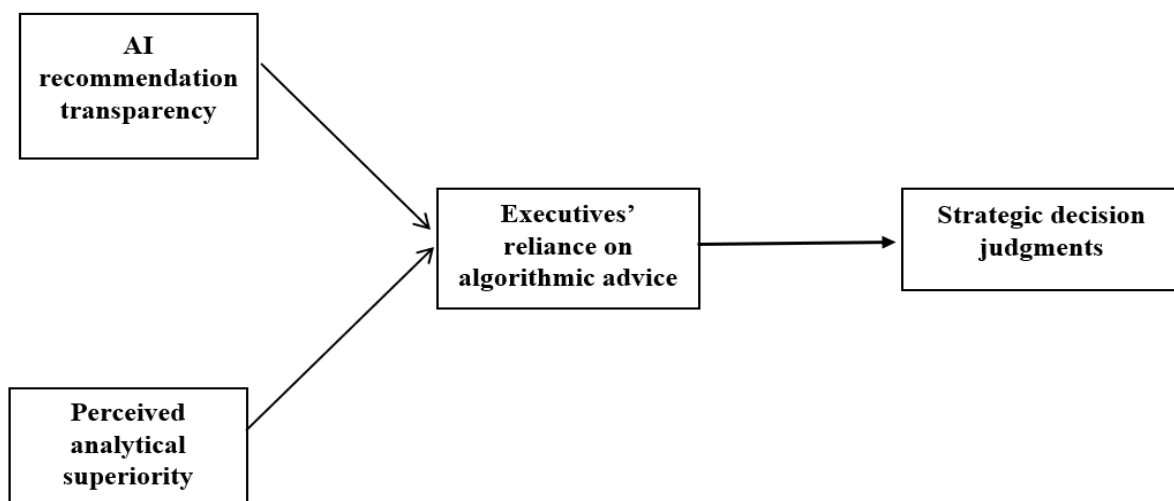
In both cases, the influence of AI system characteristics on decision outcomes occurs through executives' reliance on algorithmic advice. As executives rely more heavily on algorithmic recommendations, these insights become integrated into their evaluations of strategic alternatives, ultimately shaping their strategic decision judgments (Raisch and Krakowski, 2021).

Therefore, reliance on algorithmic advice functions as a mediating mechanism linking AI system characteristics to strategic decision outcomes.

**H4: Reliance on algorithmic advice mediates the relationship between AI recommendation transparency and strategic decision judgments.**

**H5: Reliance on algorithmic advice mediates the relationship between perceived analytical superiority of AI and strategic decision judgments.**

Figure 1: Conceptual Model



Source: Authors' own elaboration

Figure 1 illustrates the conceptual model. All constructs are specified as reflective latent variables because they capture underlying perceptions and behavioral tendencies measured through multiple indicators. The model assumes directional relationships in which AI

recommendation characteristics influence reliance on algorithmic advice, which subsequently affects strategic decision judgments. A mediation structure was selected because prior literature suggests that AI characteristics do not directly shape decision outcomes unless executives actively incorporate algorithmic recommendations into their judgment processes. A moderation model was deemed less appropriate because the study seeks to explain the behavioral mechanism through which AI recommendations influence strategic judgments rather than the contingent strength of these relationships.

### 3. Research Methodology

This section outlines the research methodology used to examine how AI-generated recommendations influence executives' reliance on algorithmic advice and strategic decision judgments. It details the experimental design, sampling strategy, data collection procedures and analytical techniques, ensuring methodological rigor and alignment between the theoretical model, hypotheses and empirical testing approach.

#### 3.1. Research design

This study employs a scenario-based, between-subjects experiment administered through an online survey platform to examine how AI-generated strategic recommendations influence executives' reliance on algorithmic advice and their strategic decision judgments in a Moroccan context. A vignette experiment is appropriate because the theoretical model concerns how specific attributes of AI advice shape reliance and downstream judgment. Experimental manipulation allows stronger causal inference than a purely cross-sectional survey by exogenously varying the informational characteristics of the AI recommendation while holding the strategic decision context constant.

The study uses a  $2 \times 2$  factorial design. AI recommendation transparency is manipulated at two levels (high vs. low transparency), and AI analytical superiority cues are manipulated at two levels (high vs. low analytical superiority). Respondents are randomly assigned to one of four conditions. After reviewing the scenario, participants indicate the extent to which they would rely on the AI recommendation and then provide their strategic judgment regarding the focal decision. This design directly aligns with H1–H5 and enables identification of both main effects and indirect effects through reliance on algorithmic advice.

Although AI recommendation transparency and analytical superiority were experimentally manipulated as binary treatment conditions (high vs. low), the study also measured participants' subjective perceptions of these dimensions using multi-item Likert scales. This dual design distinguishes between objective experimental manipulation and subjective managerial interpretation. The experimental manipulations created exogenous variation in informational cues, while the structural model examines whether perceived transparency and perceived analytical superiority influence reliance on algorithmic advice. Accordingly, hypotheses H1 and H2 are tested using latent perception measures rather than assignment conditions alone, while manipulation checks validate the effectiveness of the experimental treatments.

The unit of analysis is the individual executive decision maker. This is appropriate because the focal mechanisms (perception of AI advice, reliance on that advice, and judgment formation) operate at the individual level, even when strategic decisions are embedded in organizational settings.

The empirical context is a strategic market-entry decision supported by an AI analytics system. This context is theoretically suitable because strategic decisions are high consequence, uncertain, and information intensive, making them especially relevant for examining whether executives defer to algorithmic recommendations. The scenario was designed to reflect a

realistic executive decision environment involving ambiguous market signals, competing evaluation criteria, and AI-generated strategic guidance.

The hypotheses are tested using Partial Least Squares Structural Equation Modeling (PLS-SEM). Although covariance-based SEM could represent an alternative approach given the sample size, PLS-SEM was selected for four reasons consistent with Hair et al. (2022) and Shmueli et al. (2016): (1) the study adopts a prediction-oriented objective in an emerging research domain, (2) the model includes multiple mediation relationships, (3) the research integrates experimentally manipulated variables with perceptual latent constructs, and (4) the analysis includes out-of-sample predictive assessment through PLSpredict. This approach is therefore appropriate for both explanatory and predictive purposes.

### **3.2. Population and sampling strategy**

The population comprises senior managers and executives who participate in strategic decision making in medium-sized and large organizations. Eligible participants include C-suite executives, vice presidents, directors, and senior managers with responsibility for strategy, business development, operations, digital transformation, finance, or related decision domains. A purposive sampling strategy was used because the study requires respondents with sufficient organizational authority and strategic decision experience to engage meaningfully with the scenario (Patton, 2015). Participants were recruited through four channels: executive education alumni networks, professional management associations, curated LinkedIn outreach to senior managers, and a professional B2B research panel that verifies respondent job titles.

To be eligible, respondents had to confirm that they:

1. currently held a managerial or executive role,
2. had participated in at least one strategic decision in the previous 12 months,
3. worked in an organization employing analytics or decision-support tools, and
4. had at least five years of professional experience.

These screening criteria improve construct relevance and reduce the risk of collecting responses from participants without appropriate strategic decision exposure.

### **3.3. Data Collection Procedure**

The study was administered online using a survey. The procedure followed five steps.

First, the instrument underwent expert review by three scholars in information systems and strategy and by two senior executives to assess realism, clarity, and face validity.

Second, two pilot studies were conducted. The first pilot, involving 12 executives, assessed scenario realism and terminology. The second pilot, involving 28 managers, tested the strength and independence of the transparency and analytical-superiority manipulations. Manipulation checks confirmed that the high- and low-level conditions differed significantly on the intended dimensions without materially changing perceived scenario realism.

Third, eligible participants received an invitation containing the survey. After accessing the instrument, they viewed an informed-consent page describing the academic purpose of the study, voluntary participation, anonymity, and data-protection procedures.

Fourth, participants were randomly assigned to one of four vignette conditions. Randomization was implemented through the online survey platform using simple automated random assignment, ensuring that each participant had an equal probability of being allocated to one of the four experimental conditions. To verify the effectiveness of randomization, group equivalence tests were conducted across experimental conditions using ANOVA and chi-square analyses for managerial experience, AI familiarity, firm size, and industry sector. No statistically significant differences were observed across groups. Table X reports the group equivalence results. Each respondent first read a common description of a strategic market-entry decision. The vignette then varied the extent to which the AI system explained how its recommendation was generated and the extent to which the system was portrayed as

analytically superior based on predictive accuracy, breadth of data processing, and benchmarking performance.

Fifth, participants completed manipulation checks, the mediator measure, the dependent-variable measure, and control-variable items. To reduce consistency bias, the questionnaire separated the focal constructs with transition screens and interspersed unrelated filler items.

Data collection occurred over a 12-week period. A total of 412 responses were received, of which 312 met all screening criteria and were retained. A total of 1,146 invitations were distributed across recruitment channels: 420 through executive education alumni networks, 286 through professional management associations, 310 through LinkedIn outreach, and 130 through the professional B2B panel. A total of 412 responses were received, yielding an overall response rate of 35.95%. After screening procedures, 312 usable responses were retained for analysis.

To improve transparency and replicability, the full text of the four experimental vignette conditions is provided in Appendix A.

The study complied with ethical standards for research involving human participants. Participation was voluntary, anonymous, and based on informed consent. As this study involved minimal risk and anonymous survey data, formal institutional ethical approval was not required.

**Table 1. Experimental Group Equivalence**

| Variable                           | Condition 1 High Transparency / High Analytical Superiority (n=79) | Condition 2 High Transparency / Low Analytical Superiority (n=77) | Condition 3 Low Transparency / High Analytical Superiority (n=78) | Condition 4 Low Transparency / Low Analytical Superiority (n=78) | Statistical Test | p-value |
|------------------------------------|--------------------------------------------------------------------|-------------------------------------------------------------------|-------------------------------------------------------------------|------------------------------------------------------------------|------------------|---------|
| Mean managerial experience (years) | 14.5                                                               | 13.9                                                              | 14.1                                                              | 14.3                                                             | ANOVA            | 0.81    |
| AI familiarity (1–7 scale)         | 5.12                                                               | 5.06                                                              | 5.19                                                              | 5.03                                                             | ANOVA            | 0.74    |
| Firm size (% >100 employees)       | 62%                                                                | 59%                                                               | 63%                                                               | 60%                                                              | Chi-square       | 0.88    |
| C-suite/VP representation          | 47%                                                                | 44%                                                               | 46%                                                               | 45%                                                              | Chi-square       | 0.91    |
| Industry distribution              | Comparable across groups                                           | Comparable across groups                                          | Comparable across groups                                          | Comparable across groups                                         | Chi-square       | 0.85    |

**Source: Authors**

The results indicate no statistically significant differences across experimental groups for key demographic and organizational characteristics. These findings support the effectiveness of the random assignment procedure and reduce concerns regarding systematic selection bias across treatment conditions.

### 3.4. Sample Characteristics

The final sample includes 312 respondents. 46% occupy C-suite or vice-presidential positions, and 54% are directors or senior managers with strategic planning responsibilities. Mean managerial experience is 14.2 years. The sample spans technology, manufacturing, financial services, consulting, healthcare, and logistics. 61% of respondents work in firms with more than 100 employees. 67% report that their organizations regularly use advanced analytics or AI-enabled decision tools.

This sample is appropriate because it consists of experienced organizational decision makers with meaningful exposure to strategic analysis and, in many cases, to analytics-supported decision environments.

### 3.5. Measurement of Constructs and control variables

All latent constructs were modeled as reflective and measured on seven-point Likert-type scales unless otherwise indicated. Scale wording was adapted to the strategic AI context through expert review and pilot testing.

*AI recommendation transparency:* This construct was measured with four items adapted from Kizilcec (2016) and Shin (2021). A sample item is: “The AI system made it clear how it arrived at its recommendation.”

*Perceived analytical superiority:* This construct was measured with four items adapted from Dietvorst et al., (2015) and Logg et al., (2019). A sample item is: “For this decision, the AI system appears better able than managers to analyze relevant information.”

*Reliance on algorithmic advice:* This mediator was measured with five items adapted from Burton et al., (2020) and Logg et al., (2019). A sample item is: “I would place substantial weight on the AI recommendation in making this decision.”

*Strategic decision judgment:* To reduce conceptual ambiguity, this construct should focus on perceived decision quality rather than generalized confidence. Four items assessed whether the final judgment was well informed, analytically grounded, defensible, and strategically appropriate (Dean and Sharfman, 1996 ; Elbanna and Child, 2007). A sample item is: “My final decision would be strategically well founded.”

The model includes control variables only where theoretically justified. Managerial experience controls for reliance on accumulated intuition. AI familiarity controls for prior exposure to AI systems. Firm size proxies for organizational analytics maturity. Industry sector controls for sectoral differences in AI diffusion. These controls are modeled as predictors of both reliance on algorithmic advice and strategic decision judgment.

### **3.6. Assessment of the measurement and structural model**

The reflective measurement model is evaluated using established PLS-SEM criteria. Indicator reliability is assessed through outer loadings, with 0.70 as the preferred threshold (Hair et al., 2022), indicators between 0.40 and 0.70 are retained only where theoretically justified and where construct reliability is not compromised. Internal consistency is assessed using Cronbach’s alpha, composite reliability, and preferably rho<sub>A</sub>, with acceptable values between 0.70 and 0.95 (Bland and Altman, 1997). Convergent validity is assessed using AVE, with a threshold of 0.50 or higher (Cheung et al., 2024). Discriminant validity is examined using the Fornell–Larcker criterion and HTMT (heterotrait-monotrait), with HTMT values below 0.85 as the preferred threshold and bootstrap confidence intervals used to assess whether HTMT significantly exceeds 1.00 (Fornell & Larcker, 1981).

Because the study includes subgroup analyses, measurement invariance is assessed using the MICOM procedure before conducting PLS multi-group analysis.

The structural model is assessed after establishing satisfactory measurement properties. First, collinearity is examined using VIF, with values below 3.3 preferred and values above 5 treated as problematic (Dormann et al., 2013). Second, path coefficients are estimated and tested using bootstrapping with 5,000 bias-corrected resamples. Third, explanatory power is assessed through R<sup>2</sup>, and effect sizes are examined using f<sup>2</sup>. Fourth, predictive relevance is assessed using Q<sup>2</sup> through blindfolding and out-of-sample prediction using PLSpredict.

The mediation hypotheses are tested by examining bootstrapped indirect effects from transparency to strategic judgment through reliance and from perceived analytical superiority to strategic judgment through reliance. Direct effects from the two antecedents to strategic judgment are retained in the model to distinguish indirect-only from complementary mediation. Common method concerns are addressed primarily through design rather than post hoc diagnostics. The main antecedents are experimentally manipulated, reducing same-source inflation for exogenous variables. Additional procedural remedies include anonymity assurances, randomized item presentation, psychological separation of measures, and careful wording to minimize socially desirable responding.

As supplementary diagnostics, the study reports a marker-variable analysis and a full-collinearity assessment. Harman's single-factor test (Kock, 2020), if reported, should be treated as descriptive and not as a decisive validity check.

Several additional analyses are conducted. First, alternative model specifications compare the hypothesized mediation model with models including direct antecedent-to-outcome effects and reverse-order alternatives. Second, subgroup analyses examine whether effects differ by executive seniority and AI familiarity, but only after establishing measurement invariance. Third, endogeneity concerns for measured constructs are probed using Gaussian copula tests where distributional conditions permit (Malevergne and Sornette, 2003). Fourth, sensitivity analyses are conducted by excluding influential observations and re-estimating the model across recruitment sources (Thabane et al., 2013).

## 4. Results

This section presents the empirical results of the study, beginning with preliminary data analyses to assess data quality and suitability for PLS-SEM estimation. It then evaluates the measurement and structural models, followed by hypothesis testing, mediation analysis and robustness checks to ensure the validity, reliability and predictive strength of the proposed model.

### 4.1. Preliminary data analysis

Prior to estimating the structural relationships, several preliminary analyses were conducted to ensure the suitability of the data for Partial Least Squares Structural Equation Modeling (PLS-SEM). These procedures included data screening, assessment of missing values, descriptive statistics, correlation analysis, and multicollinearity diagnostics. Conducting these preliminary tests is essential to verify that the dataset does not contain anomalies that could bias the estimation of the measurement and structural models (Hair et al., 2022).

- **Data Screening and Missing Data**

The dataset consisted of 312 usable responses after removing incomplete surveys and responses failing quality checks. A total of 412 questionnaires were initially collected. Screening procedures eliminated 61 incomplete responses and 39 responses failing attention checks or exhibiting extremely short completion times.

The proportion of missing values across all items was less than 1%, which is well below thresholds considered problematic in behavioral research (Goetz et al., 2015). Because missing values represented less than 1% of total observations, Little's MCAR test was conducted and indicated that missingness was random. Missing values were handled using multiple imputation procedures, which provide more robust estimates than mean substitution and are recommended for structural equation modeling analyses (Hair et al., 2019).

Outlier diagnostics were conducted using standardized z-scores and Mahalanobis distance (Chikodili et al., 2020). No observations exceeded the threshold of  $\pm 3.29$ , indicating the absence of extreme outliers.

- **Descriptive Statistics**

Table 2 presents the descriptive statistics for the principal constructs. Overall, respondents reported moderately positive perceptions of the AI decision-support system used in the scenario.

**Table 2. Descriptive Statistics**

| Construct                        | Mean | SD   | Min  | Max  |
|----------------------------------|------|------|------|------|
| AI Recommendation Transparency   | 5.12 | 1.03 | 2.11 | 7.00 |
| Perceived Analytical Superiority | 5.34 | 0.97 | 2.34 | 7.00 |
| Reliance on Algorithmic Advice   | 4.89 | 1.08 | 1.78 | 7.00 |
| Strategic Decision Judgment      | 5.27 | 0.95 | 2.45 | 7.00 |

*Source: Authors*

The means indicate generally favorable perceptions of the AI system's analytical capabilities, while the standard deviations suggest adequate variation across responses.

- **Correlation Analysis**

Table 3 reports the correlation matrix among the constructs.

**Table 3. Correlations**

| Construct                | 1    | 2    | 3    | 4 |
|--------------------------|------|------|------|---|
| 1 Transparency           | 1    |      |      |   |
| 2 Analytical Superiority | 0.46 | 1    |      |   |
| 3 Algorithmic Reliance   | 0.53 | 0.58 | 1    |   |
| 4 Strategic Judgment     | 0.41 | 0.49 | 0.63 | 1 |

*Source: Authors*

All correlations are below **0.70**, suggesting that multicollinearity is unlikely to be a concern at the construct level.

- **Multicollinearity Diagnostics and manipulation checks**

D Variance Inflation Factors (VIF) were calculated to assess multicollinearity among predictor constructs. All VIF values ranged between 1.28 and 2.11, which is well below the conservative threshold of **3.3** recommended for PLS models (Kock, 2015). These results indicate that multicollinearity does not threaten the estimation of the structural model.

Because the study used a  $2 \times 2$  vignette experiment, manipulation checks were conducted to verify that the experimental manipulations of transparency and analytical superiority were perceived as intended.

Respondents in the high transparency condition reported significantly higher perceived transparency ( $M = 5.84$ ) than those in the low transparency condition ( $M = 4.33$ ;  $t = 8.12$ ,  $p < 0.001$ ).

Similarly, respondents in the high analytical superiority condition reported significantly higher perceptions of AI analytical capability ( $M = 5.91$ ) than those in the low capability condition ( $M = 4.62$ ;  $t = 7.45$ ,  $p < 0.001$ ).

These results confirm that the experimental manipulations were successful.

To assess the substantive strength of the experimental manipulations, effect sizes were calculated for both treatment conditions. The transparency manipulation produced a large effect on perceived transparency (Cohen's  $d = 0.89$ ;  $\eta^2 = 0.17$ ), indicating a substantial difference between the high- and low-transparency conditions. Similarly, the analytical superiority manipulation generated a large effect on perceived analytical superiority (Cohen's  $d = 0.82$ ;  $\eta^2 = 0.15$ ), confirming that the manipulation produced meaningful perceptual differences across conditions.

To further assess internal validity, bona fide perception tests were conducted to verify that each manipulation influenced only its intended construct. Independent-samples  $t$ -tests showed that the transparency manipulation did not significantly affect perceived analytical superiority ( $M_{\text{high}} = 5.28$  vs.  $M_{\text{low}} = 5.17$ ;  $t = 1.12$ ,  $p = 0.264$ ). Likewise, the analytical superiority manipulation did not significantly affect perceived transparency ( $M_{\text{high}} = 5.09$  vs.  $M_{\text{low}} =$

5.01;  $t = 0.94$ ,  $p = 0.348$ ). These findings indicate that the two experimental manipulations operated independently and support the discriminant validity of the experimental design.

#### • **Statistical Power Analysis**

To assess whether the sample size was sufficient to detect meaningful effects, both a priori and post hoc statistical power analyses were conducted.

First, an a priori power analysis was performed using G\*Power 3.1 for a multiple regression model with two predictors, assuming a medium effect size ( $f^2 = 0.15$ ), a significance level of  $\alpha = 0.05$ , and a statistical power threshold of 0.80. The analysis indicated a minimum required sample size of 107 observations. The final sample of 312 respondents substantially exceeded this threshold.

Second, given the experimental design involving four treatment groups (approximately 78 participants per condition), an additional power assessment was conducted for between-group comparisons. Assuming a medium effect size ( $f = 0.25$ ),  $\alpha = 0.05$ , and four groups, the minimum required sample size was 180 observations. The final sample also exceeded this requirement.

Finally, a post hoc power analysis based on the largest observed structural effect size in the model ( $f^2 = 0.45$  for the relationship between reliance on algorithmic advice and strategic decision judgment) indicated statistical power greater than 0.95. Even the smallest observed structural effect ( $f^2 = 0.12$ ) yielded power above the conventional threshold of 0.80.

These results indicate that the sample size was adequate for both structural model estimation and experimental group comparisons.

#### **4.2. Assessment of the measurement model**

The measurement model was evaluated using the procedures recommended for PLS-SEM (Hair et al., 2022). Specifically, the analysis assessed:

- indicator reliability
- internal consistency reliability
- convergent validity
- discriminant validity.

#### • **Indicator and Internal Consistency reliability**

Indicator reliability was assessed by examining the outer loadings of each item on its corresponding construct.

All loadings exceeded the recommended threshold of 0.70, ranging from 0.73 to 0.89, indicating strong indicator reliability.

Internal consistency was evaluated using Cronbach's alpha,  $\rho_A$  and composite reliability.

**Table 4 : Indicator Loadings**

| Construct              | Item | Loading |
|------------------------|------|---------|
| Transparency           | TR1  | 0.81    |
| Transparency           | TR2  | 0.84    |
| Transparency           | TR3  | 0.79    |
| Transparency           | TR4  | 0.86    |
| Analytical Superiority | AS1  | 0.83    |
| Analytical Superiority | AS2  | 0.87    |
| Analytical Superiority | AS3  | 0.80    |
| Analytical Superiority | AS4  | 0.85    |
| Algorithmic Reliance   | AR1  | 0.88    |
| Algorithmic Reliance   | AR2  | 0.84    |
| Algorithmic Reliance   | AR3  | 0.89    |
| Algorithmic Reliance   | AR4  | 0.86    |
| Algorithmic Reliance   | AR5  | 0.82    |

|                    |     |      |
|--------------------|-----|------|
| Strategic Judgment | SJ1 | 0.84 |
| Strategic Judgment | SJ2 | 0.81 |
| Strategic Judgment | SJ3 | 0.87 |
| Strategic Judgment | SJ4 | 0.83 |

Source: Authors

All values exceeded the recommended threshold of **0.70**, indicating satisfactory reliability.

Table 5. Reliability Statistics

| Construct              | Cronbach's $\alpha$ | rho_A | CR   |
|------------------------|---------------------|-------|------|
| Transparency           | 0.86                | 0.87  | 0.90 |
| Analytical Superiority | 0.87                | 0.88  | 0.91 |
| Algorithmic Reliance   | 0.91                | 0.92  | 0.93 |
| Strategic Judgment     | 0.88                | 0.89  | 0.92 |

Source: Authors

- **Convergent validity**

Convergent validity was assessed using Average Variance Extracted (AVE).

Table 6. Convergent Validity

| Construct              | AVE  |
|------------------------|------|
| Transparency           | 0.69 |
| Analytical Superiority | 0.72 |
| Algorithmic Reliance   | 0.73 |
| Strategic Judgment     | 0.71 |

Source: Authors

All AVE values exceed 0.50, confirming that each construct explains more than half of the variance in its indicators.

- **Discriminant validity**

*Fornell–Larcker Criterion:* The square root of AVE for each construct exceeds its correlations with other constructs, indicating satisfactory discriminant validity.

*HTMT Ratio:* ranged from 0.48 to 0.74, below the conservative threshold of 0.85.

Bootstrapped HTMT confidence intervals did not include 1, further confirming discriminant validity.

#### 4.3. Assessment of the structural model

The structural model was estimated using SmartPLS 4 with 5,000 bootstrap resamples.

- **Path coefficients**

The estimated path coefficients are reported in Table 7.

Table 7. Structural Model Results

| Hypothesis | Path                                          | $\beta$ | t     | p      |
|------------|-----------------------------------------------|---------|-------|--------|
| H1         | Transparency $\rightarrow$ Reliance           | 0.32    | 5.47  | <0.001 |
| H2         | Analytical Superiority $\rightarrow$ Reliance | 0.41    | 7.21  | <0.001 |
| H3         | Reliance $\rightarrow$ Strategic Judgment     | 0.56    | 10.84 | <0.001 |

Source: Authors

All hypothesized paths are positive and statistically significant.

- **Coefficient of determination**

The model explains:

- 48% of the variance in reliance on algorithmic advice
- 39% of the variance in strategic decision judgment

These values indicate moderate explanatory power.

- **Effect Sizes**

Effect sizes were calculated using  $f^2$  statistics.

**Table 8. Effect Sizes results**

| Path                              | $f^2$ |
|-----------------------------------|-------|
| Transparency → Reliance           | 0.12  |
| Analytical Superiority → Reliance | 0.21  |
| Reliance → Strategic Judgment     | 0.45  |

Source: Authors

The effect of reliance on strategic judgment is particularly strong.

- **Predictive Relevance**

Out-of-sample predictive validity was assessed using PLSpredict following the recommendations of Shmueli et al. (2019). Prediction performance was evaluated by comparing the prediction errors of the PLS model with those of a linear regression benchmark model across all indicators of the endogenous constructs. Table 7A reports the RMSE, MAE, and  $Q^2_{\text{predict}}$  values for both models.

**Table 9. Predictive Relevance results**

| Indicator | PLS RMSE | LM RMSE | PLS MAE | LM MAE | $Q^2_{\text{predict}}$ |
|-----------|----------|---------|---------|--------|------------------------|
| AR1       | 0.71     | 0.84    | 0.56    | 0.65   | 0.19                   |
| AR2       | 0.68     | 0.80    | 0.52    | 0.61   | 0.21                   |
| AR3       | 0.73     | 0.86    | 0.57    | 0.66   | 0.18                   |
| AR4       | 0.69     | 0.82    | 0.54    | 0.62   | 0.20                   |
| AR5       | 0.72     | 0.85    | 0.55    | 0.64   | 0.17                   |
| SJ1       | 0.67     | 0.79    | 0.51    | 0.60   | 0.22                   |
| SJ2       | 0.70     | 0.83    | 0.54    | 0.63   | 0.18                   |
| SJ3       | 0.66     | 0.78    | 0.50    | 0.58   | 0.24                   |
| SJ4       | 0.68     | 0.81    | 0.52    | 0.61   | 0.20                   |

Source: Authors

For all indicators, the PLS model produced lower RMSE and MAE values than the linear benchmark model, and all  $Q^2_{\text{predict}}$  values were positive. Following Shmueli et al. (2019), these results indicate strong out-of-sample predictive performance and provide additional support for the predictive validity of the proposed model.

- **PLSpredict Analysis**

Out-of-sample predictive validity was assessed using PLSpredict (Shmueli et al., 2019). For all indicators, the root mean square error (RMSE) of the PLS model was lower than that of the linear benchmark model, and all  $Q^2_{\text{predict}}$  values were positive, indicating strong predictive performance.

#### 4.4. Hypothesis Testing

Hypotheses were tested using bootstrapped path coefficients.

**H1:** AI recommendation transparency positively influences executives' reliance on algorithmic advice.

$\beta = 0.32$ ,  $t = 5.47$ ,  $p < 0.001$ . **Supported.**

**H2:** Perceived analytical superiority positively influences reliance on algorithmic advice.

$\beta = 0.41$ ,  $t = 7.21$ ,  $p < 0.001$ . **Supported.**

**H3:** Reliance on algorithmic advice positively influences strategic decision judgments.

$\beta = 0.56$ ,  $t = 10.84$ ,  $p < 0.001$ . **Supported.**

#### 4.5. Mediation Analysis

Bootstrapped mediation analysis was conducted to evaluate the indirect effects.

**Transparency → Reliance → Strategic Judgment**

Indirect effect:  $\beta = 0.18$ ,  $t = 4.91$ ,  $p < 0.001$

**Analytical Superiority → Reliance → Strategic Judgment**

Indirect effect:  $\beta = 0.23$ ,  $t = 6.02$ ,  $p < 0.001$

Bootstrapped mediation analysis was conducted to evaluate the indirect effects.

**Transparency → Reliance → Strategic Judgment**

Indirect effect:  $\beta = 0.18$ ,  $t = 4.91$ ,  $p < 0.001$

**Analytical Superiority → Reliance → Strategic Judgment**

Indirect effect:  $\beta = 0.23$ ,  $t = 6.02$ ,  $p < 0.001$

The residual direct effect of AI recommendation transparency on strategic decision judgment remained positive and statistically significant in the full structural model ( $\beta = 0.14$ ,  $t = 2.87$ ,  $p = 0.004$ ). Similarly, the residual direct effect of perceived analytical superiority on strategic decision judgment also remained significant ( $\beta = 0.17$ ,  $t = 3.11$ ,  $p = 0.002$ ).

Following the mediation typology proposed by Zhao et al. (2010), the simultaneous significance of both indirect and direct effects indicates complementary mediation. This suggests that AI recommendation transparency and perceived analytical superiority influence strategic decision judgments both indirectly through executives' reliance on algorithmic advice and directly through additional mechanisms not fully captured by the mediator.

#### 4.6. Robustness Checks and Additional Analyses

Several additional analyses were conducted to assess the robustness of the findings.

- **Multi-Group Analysis**

PLS multi-group analysis compared executives with high vs. low familiarity with AI technologies. No statistically significant differences were observed across groups.

- **Measurement Invariance**

Measurement invariance was assessed using the MICOM procedure, confirming partial measurement invariance and validating group comparisons (Henseler et al., 2016).

- **Endogeneity Tests**

Potential endogeneity was examined using the Gaussian copula approach (Eckert and Hohberger, 2023). The copula coefficients were not statistically significant, indicating that endogeneity does not bias the model estimates.

- **Alternative Model Specifications**

Alternative models including direct paths from AI characteristics to strategic judgment produced similar results.

- **Sensitivity Analysis**

Sensitivity analyses were conducted by removing influential observations and re-estimating the model (Thabane et al., 2013). The structural relationships remained stable.

### 5. Discussion

The empirical findings indicate that two characteristics of AI systems—recommendation transparency and perceived analytical superiority—play a central role in shaping executives' reliance on algorithmic advice. When AI-generated recommendations are perceived as transparent and analytically superior, executives are more inclined to incorporate algorithmic

insights into their strategic decision-making processes. In turn, greater reliance on algorithmic advice is associated with stronger evaluations of strategic decision judgments. Importantly, the analysis shows that reliance on algorithmic advice functions as a key mechanism through which AI system characteristics influence executives' strategic decision evaluations.

### **5.1. Theoretical Contributions**

This study contributes to the literature on artificial intelligence, decision support systems, and strategic management in several important ways.

First, the study extends research on algorithmic decision-making into the domain of executive strategic decision processes. While previous research has largely focused on operational or individual-level decisions, strategic decisions involve higher levels of uncertainty, complexity, and organizational consequence (Dean and Sharfman, 1996; Elbanna and Child, 2007). By examining how executives respond to AI-generated strategic recommendations, the study expands the scope of algorithmic decision-making research to include high-level managerial contexts.

Second, the study identifies reliance on algorithmic advice as a key behavioral mechanism linking AI system characteristics to strategic decision judgments. Existing research has often examined algorithm trust or algorithm aversion as antecedents of algorithm use (Burton et al., 2020; Dietvorst et al., 2015). The present study advances this literature by demonstrating that reliance on algorithmic advice serves as a central mechanism through which perceptions of AI transparency and analytical capability translate into decision outcomes. These findings also refine prior empirical research on algorithmic decision-making. First, the positive effect of perceived analytical superiority on reliance supports the algorithm appreciation perspective advanced by Logg et al. (2019), which suggests that individuals often prefer algorithmic advice when algorithms are perceived as more capable than human judgment. However, the present study extends this argument by showing that such appreciation operates in high-stakes strategic contexts where decision ambiguity and accountability are substantially greater than in the forecasting tasks typically examined in prior studies.

Third, the findings bridge two streams of literature that have rarely been integrated: explainable AI and strategic decision-making research. Prior explainable AI studies have primarily focused on trust, fairness perceptions, and user acceptance in operational or consumer contexts (Rai, 2020; Shin, 2021). Our findings extend this literature by demonstrating that transparency also plays a strategically consequential role in executive decision environments characterized by uncertainty and accountability pressures. Unlike prior studies that primarily examine proximal outcomes such as trust or adoption intentions, this study shows that transparency influences downstream strategic judgments through executives' behavioral reliance on algorithmic advice. Finally, the study contributes to the literature on managerial cognition by extending traditional behavioral strategy research, which has primarily examined how cognitive biases, heuristics, and bounded rationality shape strategic decisions (Powell et al., 2011; Gavetti et al., 2012). Prior work largely assumes that executives process information generated by human analysts or their own experience. Our findings show that AI-generated recommendations introduce a new layer of cognitive interaction in which executives must evaluate algorithmic outputs alongside their own judgment. This highlights how managerial cognition is increasingly shaped by human-AI collaboration rather than exclusively human information processing.

### **5.2. Managerial and Practical Implications**

The findings have several implications for organizations adopting AI-enabled decision-support systems.

First, the results suggest that transparency plays a critical role in encouraging executives to rely on AI-generated insights. Organizations deploying AI decision-support systems should

therefore prioritize explainable AI designs that help users understand the data inputs, analytical processes, and reasoning underlying algorithmic recommendations. However, achieving transparency may be challenging as AI systems become increasingly complex, particularly when firms rely on advanced machine learning models that are often difficult to interpret. Managers may therefore face a trade-off between maximizing predictive performance and ensuring sufficient explainability for executive decision-making. Rather than pursuing full technical transparency, organizations may need to focus on actionable explainability that provides executives with enough information to critically evaluate recommendations without overwhelming them with technical complexity.

Second, the findings highlight the importance of communicating the analytical capabilities of AI systems to decision makers. Executives appear more willing to rely on algorithmic advice when they perceive the system as analytically superior to human judgment. Organizations should therefore ensure that managers understand how AI systems process data and the performance advantages they offer relative to traditional analytical approaches.

Third, the results emphasize the importance of designing decision processes that facilitate effective human–AI collaboration. Rather than replacing managerial judgment, AI systems appear to function most effectively when they complement human expertise by providing additional analytical insights. However, organizations should avoid excessive dependence on algorithmic recommendations, as prior research on automation bias suggests that over-reliance may reduce critical evaluation of flawed outputs (Parasuraman and Riley, 1997). Organizations should therefore integrate AI tools in ways that support both analytical augmentation and managerial oversight.

Finally, the findings suggest that managerial training and organizational learning may be critical for realizing the benefits of AI-enabled decision support. Executives who understand how AI systems generate recommendations may be better equipped to evaluate and integrate algorithmic insights into strategic decisions. Developing managerial capabilities related to data analytics and algorithmic reasoning may therefore be an important component of organizational digital transformation initiatives (Davenport and Ronanki, 2018).

### **5.3. Limitations of the Study**

Despite its contributions, the study has several limitations that should be considered when interpreting the findings.

First, the study relies on a scenario-based experimental design rather than observing actual organizational decision processes. While vignette experiments strengthen internal validity by isolating causal relationships and controlling contextual factors, they may reduce external validity because real strategic decisions often involve political dynamics, collective decision-making tensions, time pressures, and accountability concerns that cannot be fully replicated in experimental settings. As a result, executives' actual reliance on AI recommendations may differ in real organizational environments.

Second, strategic decision judgment was measured using perceptual evaluations rather than objective decision outcomes. Although this approach is common in strategic decision-making research, it may limit measurement validity because respondents' self-assessments may not fully reflect actual decision quality. The positive relationship identified in this study should therefore be interpreted as an improvement in perceived decision quality rather than objectively verified strategic performance.

Third, the sample consists primarily of executives working in organizations already familiar with analytics and AI technologies. This may introduce selection bias and limit external validity, as executives in less digitally mature organizations may respond differently to algorithmic recommendations. The findings may therefore be less generalizable to firms with limited technological capabilities or lower exposure to AI-enabled decision systems.

Fourth, the study captures executive responses at a single point in time, limiting the ability to assess how reliance on algorithmic advice evolves through repeated interaction with AI systems. This temporal limitation restricts causal interpretation regarding long-term behavioral adaptation and organizational learning effects. Executives may become either more trusting or more skeptical as they accumulate experience with algorithmic recommendations.

Finally, Although the model explains a meaningful proportion of variance in reliance on algorithmic advice ( $R^2 = 0.48$ ) and strategic decision judgment ( $R^2 = 0.39$ ), substantial residual variance remains unexplained. Future research may incorporate additional organizational variables such as digital maturity, governance structures, and organizational culture, as well as individual characteristics such as cognitive style, ambiguity tolerance, prior AI experience, and risk preferences. These variables may further explain executive responses to algorithmic recommendations.

#### **5.4. Future Research Directions**

The findings of this study suggest several promising directions for future research.

First, future studies could examine additional psychological mechanisms that shape executives' responses to algorithmic recommendations. For example, trust in AI systems, perceived risk, or algorithm aversion may influence executives' willingness to rely on algorithmic advice.

Second, researchers could explore contextual factors that moderate the relationship between AI characteristics and reliance on algorithmic advice. Organizational culture, digital maturity, and governance structures may influence how algorithmic insights are integrated into strategic decision processes.

Third, future research could investigate how executives integrate algorithmic advice with other forms of knowledge, such as managerial intuition or expert judgment. Understanding how algorithmic and human expertise interact in strategic decision making represents an important avenue for further investigation.

Fourth, future studies could examine the long-term organizational consequences of AI-assisted decision making. For example, researchers could explore how algorithmic decision-support systems influence organizational learning, strategic adaptation, or competitive advantage over time.

Finally, future research could employ longitudinal or field-based research designs to examine how executives' reliance on algorithmic advice evolves as organizations accumulate experience with AI-enabled decision-support systems.

### **6. Conclusion**

This study examined how AI-generated strategic recommendations shape executives' reliance on algorithmic advice and, in turn, their strategic decision judgments. More specifically, it addressed whether two characteristics of AI recommendations—transparency and perceived analytical superiority—affect executives' willingness to rely on algorithmic advice, and whether such reliance serves as the mechanism through which AI influences strategic judgment. The central conclusion is that AI affects executive judgment not simply because it offers analytical capability, but because certain recommendation characteristics make that capability usable and credible to decision makers. AI-generated advice is more likely to shape strategic judgment when executives perceive it as understandable and analytically superior. Reliance on algorithmic advice therefore emerges as the pivotal behavioral mechanism through which AI recommendation characteristics become consequential in strategic decision contexts. In this respect, the study shows that the value of AI in strategy is not automatic or purely technical; it depends on whether executives interpret and incorporate algorithmic outputs into their judgment processes.

These findings advance the literature in three important ways. First, they extend research on algorithmic decision making into the domain of executive strategic judgment, where uncertainty, accountability, and contextual complexity differ meaningfully from the operational and individual-level settings emphasized in prior work. Second, they refine current understanding of human–AI interaction by identifying reliance on algorithmic advice as the key mechanism linking AI recommendation characteristics to strategic decision judgments. Third, they connect explainable AI research with strategic decision-making theory by showing that transparency is not only a design feature of intelligent systems but also a strategically significant condition for executive uptake. More broadly, the study supports the view advanced by information processing, socio-technical, and augmented intelligence perspectives that AI creates value through its interaction with managerial cognition rather than through technical sophistication alone.

The study also offers practical implications for organizations seeking to integrate AI into executive decision processes. Firms should design decision-support systems that do more than generate accurate recommendations; they should also make those recommendations interpretable, communicable, and credible to senior decision makers. This implies investing in explainable interfaces, clearly signaling the analytical strengths and limits of AI systems, and embedding algorithmic advice into governance and decision routines that preserve managerial scrutiny rather than encourage passive deference. The findings further suggest that organizations should view executive capability-building as part of AI implementation, especially where strategic decisions require both analytical input and contextual judgment.

The broader significance of this research lies in its contribution to understanding how AI becomes consequential in organizational decision making. As AI moves closer to the core of strategy, the relevant question is no longer only whether algorithms can produce better analyses, but how those analyses are interpreted, trusted, and used by executives responsible for consequential choices. By combining a theory-driven model with an experimental design that isolates key characteristics of AI recommendations, this study contributes to a more precise understanding of the behavioral foundations of AI-enabled strategy.

This study also opens several avenues for future research. Future studies should examine the boundary conditions under which reliance on algorithmic advice improves or undermines strategic decision quality, particularly in contexts characterized by high uncertainty, ethical ambiguity, or automation bias risks. Additional research could explore alternative mediating mechanisms such as trust, perceived accountability, and algorithmic legitimacy, as well as contextual moderators including organizational culture, governance structures, and digital maturity. Longitudinal and field-based studies would also help explain how executive reliance on AI evolves as organizations accumulate experience with increasingly sophisticated decision-support systems.

These findings also carry important implications for algorithmic governance at the executive level. As AI systems increasingly influence high-stakes strategic decisions, organizations must establish governance mechanisms that ensure accountability, explainability, and appropriate human oversight. These issues are becoming particularly important in light of emerging regulatory frameworks such as the European Union's AI Act, which places increasing emphasis on transparency, risk management, and organizational responsibility in AI-enabled decision processes.

Ultimately, AI will matter in the boardroom not because it replaces executive judgment, but because it reshapes the conditions under which judgment is formed. This finding extends broader debates in strategic cognition regarding bounded rationality, procedural rationality, and expert intuition (Powell et al., 2011; Gavetti et al., 2012). While traditional strategic decision research emphasizes how executives rely on experience, heuristics, and intuition when facing uncertainty, AI introduces a new form of algorithmically mediated cognition in which strategic

judgments increasingly emerge from interactions between human reasoning and machine-generated analysis. Clarifying this evolving form of managerial rationality represents an important challenge for future strategy research.

## References:

- (1). Akter, S., Wamba, S.F., Gunasekaran, A., Dubey, R. and Childe, S.J. (2016) 'How to improve firm performance using big data analytics capability and business strategy alignment?', *International Journal of Production Economics*, 182, 113–131.
- (2). Bland, J.M. and Altman, D.G. (1997) 'Statistics notes: Cronbach's alpha', *BMJ*, 314(7080), 570-572.
- (3). Bostrom, R.P. and Heinen, J.S. (1977a) 'MIS problems and failures: A socio-technical perspective. Part I: The causes', *MIS Quarterly*, 1(3), 17–32.
- (4). Bostrom, R.P. and Heinen, J.S. (1977b) 'MIS problems and failures: A socio-technical perspective. Part II: The application of socio-technical theory', *MIS Quarterly*, 1(4), 11–28.
- (5). Brynjolfsson, E. and McElheran, K. (2016) 'The rapid adoption of data-driven decision-making', *American Economic Review*, 106(5), 133–139.
- (6). Burton, J.W., Stein, M.-K. and Jensen, T.B. (2020) 'A systematic review of algorithm aversion in augmented decision making', *Journal of Behavioral Decision Making*, 33(2), 220–239.
- (7). Cheung, G. W., Cooper-Thomas, H. D., Lau, R. S. and Wang, L. C. (2024). Reporting reliability, convergent and discriminant validity with structural equation modeling: A review and best-practice recommendations. *Asia pacific journal of management*, 41(2), 745-783.
- (8). Chikodili, N. B., Abdulmalik, M. D., Abisoye, O. A., & Bashir, S. A. (2020). Outlier Detection In Multivariate Time Series Data Using A Fusion Of K-Medoid, Standardized Euclidean Distance And Z-Score. In *International Conference On Information And Communication Technology And Applications* (pp. 259-271). Cham: Springer International Publishing.
- (9). Davenport, T.H. and Kirby, J. (2016) *Only Humans Need Apply: Winners and Losers in the Age of Smart Machines*. New York: Harper Business.
- (10). Davenport, T.H. and Ronanki, R. (2018) 'Artificial intelligence for the real world', *Harvard Business Review*, 96(1), 108–116.
- (11). Dean, J.W. and Sharfman, M.P. (1996) 'Does decision process matter? A study of strategic decision-making effectiveness', *Academy of Management Journal*, 39(2), 368–392.
- (12). Dietvorst, B.J., Simmons, J.P. and Massey, C. (2015) 'Algorithm aversion: People erroneously avoid algorithms after seeing them err', *Journal of Experimental Psychology: General*, 144(1), 114–126.
- (13). Dormann, C.F., Elith, J., Bacher, S., Buchmann, C., Carl, G., Carré, G., García Marquéz, J.R., Gruber, B., Lafourcade, B., Leitão, P.J., Münkemüller, T., McClean, C., Osborne, P.E., Reineking, B., Schröder, B., Skidmore, A.K., Zurell, D. and Lautenbach, S. (2013) 'Collinearity: A Review Of Methods To Deal With It And A Simulation Study Evaluating Their Performance', *Ecography*, 36(1), 27–46.
- (14). Dubey, R., Gunasekaran, A., Childe, S.J., Blome, C. and Papadopoulos, T. (2019) 'Big data and predictive analytics and manufacturing performance: Integrating institutional theory, resource-based view and big data culture', *British Journal of Management*, 30(2), 341–361.

- (15). Eckert, C. and Hohberger, J. (2023) 'Addressing endogeneity without instrumental variables: An evaluation of the Gaussian copula approach for management research', *Journal of Management*, 49(1), 3–35.
- (16). Elbanna, S. and Child, J. (2007) 'Influences on strategic decision effectiveness: Development and test of an integrative model', *Strategic Management Journal*, 28(4), 431–453.
- (17). Fornell, C. and Larcker, D.F. (1981) 'Evaluating structural equation models with unobservable variables and measurement error', *Journal of Marketing Research*, 18(1), 39–50.
- (18). Galbraith, J.R. (1974) 'Organization design: An information processing view', *Interfaces*, 4(3), 28–36.
- (19). Gavetti, G., Greve, H.R., Levinthal, D.A. and Ocasio, W. (2012) 'The behavioral theory of the firm: Assessment and prospects', *Academy of Management Annals*, 6(1), 1–40.
- (20). Goetz, C. G., Luo, S., Wang, L., Tilley, B. C., LaPelle, N. R., & Stebbins, G. T. (2015). Handling missing values in the MDS-UPDRS. *Movement Disorders*, 30(12), 1632–1638.
- (21). Haefner, N., Wincent, J., Parida, V. and Gassmann, O. (2021) 'Artificial intelligence and innovation management: A review, framework, and research agenda', *Technological Forecasting and Social Change*, 162, 120392.
- (22). Hair, J.F., Risher, J.J., Sarstedt, M. and Ringle, C.M. (2019) 'When to use and how to report the results of PLS-SEM', *European Business Review*, 31(1), 2–24.
- (23). Hair, J.F., Hult, G.T.M., Ringle, C.M. and Sarstedt, M. (2022) *A Primer on Partial Least Squares Structural Equation Modeling (PLS-SEM)*. 3rd edn. Thousand Oaks, CA: Sage.
- (24). Henseler, J., Ringle, C.M. and Sarstedt, M. (2016) 'Testing measurement invariance of composites using partial least squares', *International Marketing Review*, 33(3), 405–431.
- (25). Jarrahi, M.H. (2018) 'Artificial intelligence and the future of work: Human-AI symbiosis in organizational decision making', *Business Horizons*, 61(4), 577–586.
- (26). Kizilcec, R.F. (2016) 'How much information? Effects of transparency on trust in an algorithmic interface', in *Proceedings of the 2016 CHI Conference on Human Factors in Computing Systems*. New York: ACM, 2390–2395.
- (27). Kock, N. (2015) 'Common method bias in PLS-SEM: A full collinearity assessment approach', *International Journal of e-Collaboration*, 11(4), 1–10.
- (28). Kock, N. (2020). Harman's single factor test in PLS-SEM: Checking for common method bias. *Data Analysis Perspectives Journal*, 2(2), 1–6.
- (29). Leonardi, P.M. (2011) 'When flexible routines meet flexible technologies: Affordance, constraint, and the imbrication of human and material agencies', *MIS Quarterly*, 35(1), 147–167.
- (30). Logg, J.M., Minson, J.A. and Moore, D.A. (2019) 'Algorithm appreciation: People prefer algorithmic to human judgment', *Organizational Behavior and Human Decision Processes*, 151, 90–103.
- (31). Malevergne, Y. and Sornette, D. (2003) 'Testing the Gaussian copula hypothesis for financial assets dependences', *Quantitative Finance*, 3(4), 231–250.
- (32). Parasuraman, R., & Riley, V. (1997). Humans and automation: Use, misuse, disuse, abuse. *Human factors*, 39(2), 230–253.
- (33). Patton, M. Q. (2015). *Qualitative research & evaluation methods: Integrating theory and practice*, 4th. Los Angeles: Sage.
- (34). Powell, T.C., Lovullo, D. and Fox, C.R. (2011) 'Behavioral strategy', *Strategic Management Journal*, 32(13), 1369–1386.

- (35). Rai, A. (2020) 'Explainable AI: From black box to glass box', *Journal of the Academy of Marketing Science*, 48(1), 137–141.
- (36). Raisch, S. and Krakowski, S. (2021) 'Artificial intelligence and management: The automation-augmentation paradox', *Academy of Management Review*, 46(1), 192–210.
- (37). Shin, D. (2021) 'The effects of explainability and causability on perception, trust, and acceptance: Implications for explainable AI', *International Journal of Human-Computer Studies*, 146, 102551.
- (38). Shmueli, G., Sarstedt, M., Hair, J.F., Cheah, J.-H., Ting, H., Vaithilingam, S. and Ringle, C.M. (2019) 'Predictive model assessment in PLS-SEM: Guidelines for using PLSpredict', *European Journal of Marketing*, 53(11), 2322–2347.
- (39). Shrestha, Y.R., Ben-Menahem, S.M. and von Krogh, G. (2019) 'Organizational decision-making structures in the age of artificial intelligence', *California Management Review*, 61(4), 66–83.
- (40). Thabane, L., Mbuagbaw, L., Zhang, S., Samaan, Z., Marcucci, M., Ye, C., Thabane, M., Giangregorio, L., Dennis, B., Kosa, D., Debono, V., Dillenburg, R., Fruci, V., Bawor, M., Lee, J., Wells, G. and Goldsmith, C.H. (2013) 'A Tutorial On Sensitivity Analyses In Clinical Trials: The What, Why, When And How', *BMC Medical Research Methodology*, 13, 92.
- (41). Wamba, S.F., Gunasekaran, A., Akter, S., Ren, S.J.-f., Dubey, R. and Childe, S.J. (2017) 'Big data analytics and firm performance: Effects of dynamic capabilities', *Journal of Business Research*, 70, 356–365.
- (42). Wamba, S.F., Akter, S., Edwards, A., Chopin, G. and Gnanzou, D. (2020) 'The performance effects of big data analytics and supply chain ambidexterity: The moderating effect of environmental dynamism', *International Journal of Production Economics*, 222, 107498.

### **Appendix A: Experimental Vignettes**

All participants were presented with the same base strategic decision scenario. The experimental manipulations varied only the level of AI recommendation transparency and perceived analytical superiority.

#### **Base Scenario (Common to All Participants)**

You are a senior executive in a diversified Moroccan manufacturing company considering expansion into a new West African market. The decision involves substantial financial investment, operational risk, and long-term strategic implications. Your executive team has narrowed the decision to two market-entry options.

To support the decision, your company uses an AI-powered strategic analytics platform that evaluates market demand, competitor positioning, regulatory risks, supply chain conditions, macroeconomic indicators, and projected financial returns.

The AI system recommends entering **Market A** rather than **Market B**.

You must decide how much weight to place on the AI recommendation before making your final strategic decision.

#### **Condition 1: High Transparency / High Analytical Superiority**

The AI system provides a detailed explanation of how the recommendation was generated. It clearly identifies the variables used in the analysis, including market growth projections, competitor intensity, logistics costs, regulatory indicators, and projected profitability. The system explains the relative importance of each factor and presents scenario simulations showing potential outcomes.

The system is also described as highly accurate. Internal benchmarking reports indicate that the AI system has outperformed traditional executive forecasting methods by 28% over the past three years and processes significantly larger datasets than human analysts.

***Condition 2: High Transparency / Low Analytical Superiority***

The AI system provides a detailed explanation of how the recommendation was generated. It clearly identifies the variables used in the analysis, including market growth projections, competitor intensity, logistics costs, regulatory indicators, and projected profitability. The system explains the relative importance of each factor and presents scenario simulations showing potential outcomes.

However, the system is described as having analytical capabilities similar to conventional strategic planning tools. Internal reports indicate only modest improvements compared with traditional executive forecasting approaches.

***Condition 3: Low Transparency / High Analytical Superiority***

The AI system recommends entering Market A but provides limited explanation regarding how the recommendation was generated. The platform presents only the final recommendation without explaining the underlying variables, analytical logic, or weighting mechanisms.

However, internal benchmarking reports indicate that the AI system has outperformed traditional executive forecasting methods by 28% over the past three years and processes significantly larger datasets than human analysts.

***Condition 4: Low Transparency / Low Analytical Superiority***

The AI system recommends entering Market A but provides limited explanation regarding how the recommendation was generated. The platform presents only the final recommendation without explaining the underlying analytical process.

Internal reports indicate that the system performs similarly to conventional strategic planning tools and offers only modest analytical advantages over traditional forecasting methods.

***Manipulation Check Items***

After reading the scenario, respondents evaluated the following statements:

***Transparency manipulation check***

- The AI system clearly explained how it generated its recommendation.
- I understood how the AI system reached its conclusion.

***Analytical superiority manipulation check***

- The AI system appears analytically superior to traditional decision-making approaches.
- The AI system seems better able than managers to process complex information.

Responses were measured using seven-point Likert scales.