

UNIVERSITY OF OKLAHOMA

GRADUATE COLLEGE

UNDERSTANDING THE NUANCES OF HOW FOLLOWERS PERCEIVE AI LEADERSHIP

A DISSERTATION

SUBMITTED TO THE GRADUATE FACULTY

in partial fulfillment of the requirements for the

Degree of

DOCTOR OF PHILOSOPHY

By CECELIA GORDON

Norman, Oklahoma

2026

UNDERSTANDING THE NUANCES OF HOW FOLLOWERS PERCEIVE AI LEADERSHIP

A DISSERTATION APPROVED FOR THE
DEPARTMENT OF PSYCHOLOGY

BY THE COMMITTEE CONSISTING OF

Dr. Shane Connelly, Chair

Dr. Matthew Jensen

Dr. Patrick Manapat

Dr. Lori Snyder

Acknowledgements

It is difficult to put into words what this journey has meant to me. Earning my doctorate has been one of the most challenging and meaningful experiences of my life, and I am deeply grateful for the people who supported me along the way. I would like to sincerely thank my advisor, Dr. Shane Connelly, for their guidance and mentorship. I am also truly grateful to my committee members, Dr. Lori Snyder, Dr. Matthew Jensen, and Dr. Patrick Manapat, for their time, thoughtful feedback, and investment in my development. Furthermore, I would like to acknowledge Dr. Dingjing Shi, who played an important role during the early stages of my research and doctoral training. What I gained from these years goes far beyond a degree. I leave this chapter with a stronger sense of purpose and a lasting appreciation for the people who stood beside me.

Table of Contents

Abstract	vi
Introduction	1
AI Leadership.....	4
AI Leadership and Relational Outcomes	5
Decision-making Transparency.....	9
Decision-making Transparency as a Moderator	11
Situational Context and Leader Effectiveness	13
Study 1	18
Method	18
Results	25
Discussion	28
Study 2	32
Implicit Theories of Leadership	32
Method	34
Results	38
Discussion	43
General Discussion	46
Implications	47
Limitations and Future Directions	50
Conclusion	52
References	53
Tables and Figures	65
Appendices	77
Appendix A: Organization Background	77
Appendix B: Study 1 Leader Descriptions	79
Appendix C: Study 1 Colleague Emails	80
Appendix D: Study 1 Leader Emails.....	83
Appendix E: Situational Context Manipulations	86
Appendix F: Manipulation Check Measures	87
Appendix G: Outcome Measures	88
Appendix H: Covariate Measures	91

Abstract

Artificial intelligence (AI) technologies are increasingly integrated into organizational decision-making, raising important questions about how followers evaluate leadership enacted by AI systems. The present research examines how followers perceive human, AI, and AI-augmented (hybrid) leaders while considering the roles of decision-making transparency and situational context. Across two experimental vignette studies, participants evaluated leadership scenarios in which leader type (human, AI, AI-augmented) and decision-making transparency (high vs. low) were manipulated between subjects, while situational context (consideration vs. initiating structure) was manipulated within subjects. Study 1 ($N = 281$) and Study 2 ($N = 420$) assessed follower trust, satisfaction with the leader, and ratings of leader effectiveness. Results across studies indicated that human leaders generally elicited stronger relational reactions than AI leaders, particularly in terms of follower satisfaction. Decision-making transparency consistently increased follower trust and satisfaction across leader types. Perceptions of leader effectiveness were contingent on situational demands: AI leaders were evaluated more favorably in initiating structure contexts emphasizing task coordination, whereas human leaders received more favorable evaluations in consideration contexts emphasizing relational support relative to other leader types. AI-augmented leaders produced more moderate or stable evaluations across contexts. Study 2 further showed that when explicit leader characteristics were absent, followers relied more heavily on informational cues such as transparency when forming evaluations. Together, these findings suggest that follower responses to emerging leadership agents depend on relational expectations, transparency in decision processes, and alignment between leader capabilities and situational demands.

Keywords: AI, leadership, transparency, trust, satisfaction, effectiveness

Understanding the Nuances of How Followers Perceive AI Leadership

Artificial intelligence (AI) technologies are rapidly transforming organizational systems, not only by automating tasks but also by reshaping how leadership is conceptualized and enacted. Organizations such as Unilever have implemented AI-driven recruitment systems to screen candidates and inform hiring decisions, and Coca-Cola has integrated AI into marketing, knowledge management, and operational optimization (GSDC, 2024; The Coca-Cola Company, 2023). These developments reflect a broader shift in which AI systems increasingly participate in activities traditionally associated with managerial discretion and coordination. Academic institutions have similarly responded, offering executive programs focused on leading AI-driven organizations, signaling recognition that AI reshapes decision making, coordination, and judgment at senior levels (Harvard Kennedy School, 2026; MIT Sloan Executive Education, 2026).

Leadership has traditionally been conceptualized in the literature as a social influence process through which individuals guide, motivate, and coordinate others toward shared goals (Forsyth, 2014; Oc & Bashshur, 2013; Yukl, 2012). However, digital transformation challenges this human-centric framing. Banks et al. (2022) argue that leadership theory must evolve to account for digitally mediated contexts in which influence may be partially or fully enacted through technological systems. In such environments, the boundaries between tool, collaborator, and leader become increasingly blurred (Bevilacqua et al., 2025). In such settings, followers may attribute agency, responsibility, and evaluative judgments not only to human leaders but also to algorithmic systems that structure decisions and outcomes.

On one hand, AI systems offer distinct advantages in domains central to leadership, particularly in data processing, pattern recognition, and structured decision making (Peifer et al.,

2022; Kurup & Gupta, 2022). These capabilities position AI as a potentially powerful tool for supporting certain complex organizational tasks (Reeves, 2023). On the other hand, leadership effectiveness is not determined solely by analytic competence. Research highlights the importance of trust-building and relational processes in shaping follower attitudes and performance (Bligh, 2016; Legood et al., 2021). Although AI systems may simulate emotional expressions, they do not possess lived emotional experience or embodied social understanding (Martinez & Aldea, 2005), and practitioner-oriented scholarship emphasizes that AI lacks contextual intuition and affective sensitivity characteristic of human leaders (Sanders, 2025).

This tension between analytic capability and relational limitation creates ambiguity regarding how followers evaluate AI relative to human leaders. Emerging empirical work has begun to examine these dynamics. Jia and Zhang (2024) experimentally compared AI and human leaders and found that AI leaders were perceived as structured and goal-focused, with positive implications for team climate. Lai and Rau (2025) experimentally manipulated human consideration and AI structure behaviors and found that the combination of high human consideration and high AI structure produced higher trust, job satisfaction, and engagement than other configurations. Their findings suggest that hybrid leadership, where human and AI agents jointly contribute to influence processes, may leverage complementary strengths of relational and structural behaviors. In the present research, this form of leadership is referred to as AI-augmented leadership.

However, existing research has not systematically examined whether evaluations of AI, human, and AI-augmented leaders depend on situational demands. Classic leadership scholarship distinguishes between initiating structure behaviors (e.g., task coordination, goal clarity) and consideration behaviors (e.g., relational support, interpersonal concern; Fleishman, 1973; Judge

et al., 2004; Yukl, 2012). These behavioral domains have been consistently associated with attitudinal and performance outcomes in human leadership research (Bass & Stogdill, 1990; Fleishman, 1973). Yet, because these distinctions have been developed and tested almost exclusively in human-led contexts, it remains unclear whether they differentially advantage AI, human, or AI-augmented leaders across varying situational demands.

Beyond leader type and situational context, decision-making transparency represents an additional boundary condition that may shape follower evaluations. Transparency in leadership contexts has been positively associated with follower trust and perceived leader effectiveness (Norman et al., 2010). In algorithmic settings, transparency may be especially consequential because AI systems often function as “black boxes,” making it difficult for individuals to infer the reasoning processes underlying decisions (de Fine Licht & de Fine Licht, 2020). Providing decision justifications plays a critical role in fostering perceived legitimacy in AI-supported decision-making systems. From a governance perspective, explainability and auditability are central to maintaining accountability in automated decision-making environments (Waltl & Vogl, 2018), and the growing integration of AI into managerial contexts underscores the need for leaders to manage these transparency-related challenges effectively (Bevilacqua et al., 2025). Accordingly, transparency may shape how followers interpret and evaluate leadership enacted through or alongside AI.

Despite growing attention to AI in organizational research, comparatively little work has systematically examined how followers evaluate AI, human, and AI-augmented leaders while simultaneously accounting for decision-making transparency and situational demands. To address this gap, the present research adopts a multi-study design. Study 1 experimentally examines how leader type (human, AI, and AI-augmented), decision-making transparency (high

and low), and situational context (consideration and initiating structure) jointly shape follower trust, satisfaction, and ratings of leader effectiveness. Participants are presented with leader profiles that reflect characteristics aligned with their assigned leader type, allowing for controlled comparisons across conditions. Study 2 builds on these findings by examining how followers implicitly conceptualize each leader type in the absence of explicit characteristic cues, drawing on implicit leadership theory (ILT). Together, these studies provide insight into how emerging forms of leadership are perceived in technologically mediated environments.

AI Leadership

Leadership has long been conceptualized as a process of social influence through which an individual shapes the attitudes, behaviors, or performance of others toward collective goals (Bass, 1990; Oc & Bashshur, 2013; Liden, 2025; Yukl, 2002, 2012). Across definitions, leadership is fundamentally rooted in the capacity to influence others within a social context. Behavioral theories of leadership have historically centered on human agents and interpersonal influence processes (Forsyth, 2014; Liden et al., 2025). These frameworks implicitly assume that leadership unfolds through cognition, affect, and social signaling.

However, contemporary work argues that digital technologies and AI are reshaping assumptions about agency, authority, and influence within organizations. Recent scholarship suggests that AI systems are no longer merely decision aids but increasingly assume roles that resemble leadership functions, including task assignment, performance monitoring, coordination, and information control (Banks et al., 2022; Bevilacqua et al., 2025). In some contexts, AI systems operate as virtual team leaders, shaping team processes and climate through algorithmic guidance (Jia & Zhang, 2024). Unlike human leaders, however, AI systems lack intrinsic

emotional capacity and intentional agency, raising fundamental questions about whether followers evaluate AI leaders using the same mechanisms applied to human leaders.

These contrasting competencies imply different leadership strengths, which have implications for leadership behaviors such as task-oriented and relations-oriented behaviors (Yukl, 2002, 2012). Task-oriented behaviors emphasize clarifying, planning, monitoring operations, and problem solving. In contrast, relations-oriented behaviors include supporting, developing, recognizing, and empowering (Yukl, 2012). Because leadership fundamentally operates through interpersonal influence processes (Bass, 1990; Oc & Bashshur, 2013), follower perceptions are at least partly shaped by the behavioral signals leaders convey. Human leaders, as socially embedded actors, are positioned to enact such behaviors, whereas AI systems are increasingly deployed to assume structured leadership functions (Jia & Zhang, 2024; Bevilacqua et al., 2025). Consequently, differences in the type of agent enacting leadership behaviors may become most apparent in follower evaluations that are more relationally centered.

AI Leadership and Relational Outcomes

Theories of relational leadership emphasize that leadership unfolds through reciprocal exchanges characterized by mutual influence, respect, and expectations (Brower et al., 2000; Colquitt et al., 2007; Graen & Uhl-Bien, 1995). Within these frameworks, leadership effectiveness is not determined solely by task accomplishment, but by the perceived quality of the social exchange between leaders and followers. Followers continuously evaluate how leaders engage with them interpersonally, and these evaluations shape attitudinal reactions toward the leader (Graen & Uhl-Bien, 1995; Yukl, 2012). In emerging contexts where leadership may be enacted by human, AI, or AI-augmented agents, relational outcomes provide a particularly

meaningful lens for distinguishing follower reactions beyond assessments of technical competence.

Two important relational outcomes are follower trust and satisfaction. These constructs capture distinct but complementary evaluations of the leader-member relationship. Follower trust reflects their “willingness to be vulnerable” in the relationship (Mayer et al., 1995), whereas follower satisfaction reflects affective appraisal of the leader’s ability and the overall quality of the working relationship (Judge et al., 2004; Scarpello & Vandenberg, 1987). Relational leadership theory suggests that both outcomes emerge from ongoing social interaction processes in which followers interpret leaders’ intentions, responsiveness, and interpersonal orientation (Brower et al., 2000; Uhl-Bien, 2006). Thus, examining trust and satisfaction together allows for a more nuanced understanding of how AI and human-centered leadership may differentially shape follower attitudes.

Trust

Among the dimensions of trustworthiness, benevolence plays a particularly central role in relational evaluation. Benevolence reflects the perception that a leader genuinely cares about and seeks to support followers’ well-being (Mayer et al., 1995). In dyadic relationships, perceptions of benevolence are fostered through friendliness, social exchange, citizenship behaviors, and demonstrated concern for others’ interests (Brower et al., 2000). Trust in the leader deepens through continued expressions of benevolence over time (Colquitt & Baer, 2023). Collectively, this literature suggests that trust in leadership is grounded not merely in competence, but in socioemotional signaling and relational orientation.

Extending this framework to AI management contexts, recent evidence suggests that AI managers are perceived as less benevolent than human managers (Li & Bitterly, 2024). Across

field and experimental studies, Li and Bitterly (2024) demonstrate that lower perceived benevolence under AI management significantly reduces trust, and that benevolence mediates the relationship between manager type and follower trust. Similarly, Höddinghaus et al. (2021) find that automated leadership agents are perceived as high in integrity and transparency but significantly lower in benevolence and adaptability. The findings suggest that trust deficits associated with AI leadership stem not from competence concerns, but from evaluations of their emotional and moral capacity.

Importantly, leadership scholarship emphasizes that emotionally intelligent behaviors, such as empathy, interpersonal sensitivity, and adaptive communication, are foundational to effective leader-follower relationships (Yukl, 2012). AI systems do not possess intrinsic emotional experience and cannot authentically engage in relationship-building in the way human leaders can (Smith & Green, 2018). In contrast, human and AI-augmented leaders retain the capacity to express socioemotional cues that signal care, empathy, and responsiveness to followers. Because relational orientation underlies perceptions of benevolence, and benevolence is foundational to trust, human and AI-augmented leaders should elicit stronger trust than AI leaders operating without human involvement. Accordingly,

Hypothesis 1a: *Follower trust in the leader will be higher for human and AI-augmented leaders than AI leaders.*

Satisfaction

Whereas trust reflects followers' willingness to be vulnerable to a leader, satisfaction reflects followers' affective evaluation of the leader and the overall quality of the working relationship (Judge et al., 2004; Scarpello & Vandenberg, 1987). Satisfaction with the leader is therefore a leader-specific attitude shaped by how followers experience their supervisor's

behaviors, treatment, and guidance. Within relational leadership frameworks, such affective evaluations emerge from ongoing social exchanges in which followers assess whether their relational and task-related expectations are fulfilled (Brower et al., 2000; Uhl-Bien, 2006).

Classic leadership research indicates that subordinate satisfaction is influenced by both relational and task-oriented behaviors. Castaneda and Nahavandi (1991) examined managerial behavior and subordinate outcomes, finding that both leader consideration ($r = .58, p < .01$) and initiation of structure ($r = .48, p < .01$) were positively related to subordinate satisfaction. Importantly, followers reported the highest levels of satisfaction when managers exhibited both consideration and task structuring behaviors (Castaneda & Nahavandi, 1991). These findings suggest that while task structuring contributes meaningfully to satisfaction, relational behaviors remain central to followers' affective evaluations of their supervisor.

Importantly, satisfaction with the leader is particularly sensitive to relational behaviors. Madlock (2008) found that supervisors' communication competence (i.e., appropriate and effective communication that involves listening, negotiating, and motivating; Cushman & Craig, 1976; Spitzberg, 1983), was the strongest predictor of employee satisfaction. They also looked at leadership style noting supervisors' relational leadership was a predictor of subordinate communication satisfaction but not task leadership (Madlock, 2008). Similarly, a research study looking at leader humility and if it predicted follower satisfaction of the leader found leader humility (i.e., respect and openness of the leader) to be strong predictors of follower satisfaction (Krumrei & Rowatt, 2023). These findings suggest that satisfaction is not driven solely by task coordination or decision accuracy, but by the quality of interpersonal engagement embedded within the leader-follower relationship.

As it relates to AI leadership, recent experimental research demonstrates that increasing levels of automation in leadership correspond with reductions in followers' positive affect. Lochner et al. (2025) found that fully automated (AI-only) leadership elicited significantly lower positive affect compared to human-led and hybrid (AI-augmented) conditions. Because positive affect is closely linked to engagement and satisfaction (Watson et al., 1988), these findings suggest that AI-led decision making may attenuate the affective foundation upon which leader satisfaction is built.

Unlike fully automated leadership, both human and AI-augmented leadership retain a human relational component capable of signaling responsiveness, care, and socioemotional engagement. Given that satisfaction is particularly sensitive to relational cues, interpersonal competence, and affective resonance, leaders who maintain human involvement in the influence process should elicit stronger satisfaction than leaders operating solely through algorithmic authority. Accordingly,

Hypothesis 1b: *Follower satisfaction with the leader will be higher for human and AI-augmented leaders than AI leaders.*

Decision-making Transparency

Transparency in decision making is an underlying factor of leadership that critically shapes how followers perceive and evaluate leader behavior. It refers to the extent to which leaders openly communicate the rationale, criteria, and processes underlying their decisions (Vogelgesang & Lester, 2009). This form of transparency is distinct from general openness and places emphasis on the clarity and disclosure surrounding decision processes.

Empirical evidence supports a positive relationship between leader transparency and follower trust. In an experimental study examining leadership during a downsizing event,

Norman et al. (2010) manipulated communication transparency by varying whether leaders provided clear explanations and rationale for their decisions. Leaders who exhibited higher transparency were perceived as significantly more trustworthy than those who did not. Importantly, transparency was operationalized specifically through disclosure of decision rationale and openness in communication, aligning with our definition of decision-making transparency.

Research in the organizational justice literature provides strong support for the importance of decision-related explanations. Informational justice (i.e., justification and truthfulness of explanations during procedures) has been shown to predict subsequent trust in supervisors over time (Colquitt & Rodell, 2011; Greenberg, 1993). In their longitudinal analysis, Colquitt and Rodell (2011) demonstrated that informational justice uniquely predicted trust. Specifically, employees were more trusting of supervisors who offered sincere and adequate explanations for their decisions, and this relationship persisted even when prior trust and prior perceptions of trustworthiness were statistically controlled.

Vogelgesang and Lester (2009) propose that leaders who are transparent about decision-making processes and the challenges they face foster greater satisfaction, in addition to trust and engagement from their followers. Similarly, Norman and colleagues (2010) found that leaders who communicated their decision rationale clearly and openly were evaluated more favorably by followers, and these favorable evaluations have been linked in the broader leadership literature to increased follower satisfaction, commitment, and retention. They argue that transparency signals openness and alignment between words and actions, which enhances followers' positive relational evaluations of the leader. Experimental research demonstrates that when supervisors deliver evaluative decisions with fair interpersonal treatment, followers report significantly

higher satisfaction with the supervisor (Leung et al., 1998). These findings suggest that how leaders communicate and justify their decisions meaningfully shapes follower satisfaction with the leader.

Taken together, decision-making transparency can reduce uncertainty about leader intentions. When followers lack information about how decisions are made, they may infer hidden motives, bias, or inconsistency. Providing clear explanations limits the need for such inferences and allows followers to evaluate decisions based on disclosed criteria (Vogelgesang & Lester, 2009). This reduction in ambiguity should positively influence both trust and satisfaction with the leader. Accordingly,

Hypothesis 2a: *Leader decision-making transparency will influence follower trust in the leader, such that trust will be higher for leaders with high transparency than leaders with low transparency.*

Hypothesis 2b: *Leader decision-making transparency will influence follower satisfaction with the leader, such that satisfaction will be higher for leaders with high transparency than leaders with low transparency.*

Decision-making Transparency as a Moderator

Although decision-making transparency is generally expected to enhance relational evaluations, its impact may vary depending on the type of leader enacting the decision. Unlike human leaders, whose reasoning processes are typically assumed to be socially interpretable, AI systems are frequently described as opaque, particularly when they rely on complex machine learning models whose internal logic is not readily interpretable by users (Ananny & Crawford, 2018; Larsson & Heintz, 2020; Okoidigun, 2025). Transparency is particularly consequential in contexts characterized by uncertainty, opacity, or information asymmetry, and opacity concerns

are uniquely heightened for algorithmic systems relative to human decision makers (Okoidigun, 2025). Glikson and Woolley (2020) note that the complex, multilayered nature of AI decision making is often not transparent, meaning that decisions may be difficult to predict and their underlying logic poorly understood. Prior reviews of research on AI trust and transparency identify transparency as an important determinant of trust in algorithmic systems, particularly when explanations clarify how and why decisions are made (Glikson & Woolley, 2020) and when they provide insight into how data are used and processed within the system (Shin, 2022).

Empirically, research demonstrates that objective transparency increases perceived transparency, which in turn enhances trust (Yu & Li, 2022). Similarly, in leadership contexts, transparent communication has been shown to increase follower trust and perceived leader effectiveness, particularly during periods of uncertainty or organizational threat (Norman et al., 2010). Transparency functions through several mechanisms, including increasing understanding of motives, reducing perceived vulnerability, and enhancing predictability of future behavior (Vogelgesang & Lester, 2009), mapping to the conceptualization of trust adopted in this research effort.

Beyond trust, decision transparency is also likely to shape follower satisfaction with the leader by influencing acceptance of the decision-making process. Emerging research in AI-mediated contexts similarly suggests that transparency regarding algorithmic decision processes enhances relational satisfaction, particularly when transparency strengthens perceptions of accountability and trust (Park & Yoon, 2024). In their experimental study, AI algorithm transparency signaling increased relational satisfaction indirectly through trust in both the AI system and its parent organization. Jantzen et al. (2024) posit transparency of AI systems as a determinant of broader satisfaction classifications, including user understanding, perceived

oversight, and evaluative reactions toward AI systems. Likewise, design-oriented research on AI transparency highlights its relationship with user satisfaction and responsible system evaluation (Schelenz et al., 2024). Collectively, this research suggests that transparency shapes not only competence-based evaluations but also broader satisfaction with AI systems, providing a foundation for expecting similar relational reactions when AI occupies roles that involve consequential decision authority.

Thus, while transparency should enhance relational evaluations across leader types, it may play a compensatory role in AI leadership contexts where baseline perceptions of opacity and uncertainty are greater. AI-augmented leadership represents an intermediate case. Although algorithmic tools may influence decisions, the presence of a human authority figure can provide perceived accountability and contextual judgment, potentially buffering concerns associated with purely automated decision making. Consequently, the marginal benefit of added transparency may be strongest when the leader is fully AI-based.

Taken together, decision transparency is expected to interact with leader type such that its positive effects on trust and satisfaction are amplified for AI leaders relative to human and AI-augmented leaders.

Hypothesis 3a: *Leader decision-making transparency will have a stronger positive effect on follower trust for AI leaders than for human or AI-augmented leaders.*

Hypothesis 3b: *Leader decision-making transparency will have a stronger positive effect on follower satisfaction for AI leaders than for human or AI-augmented leaders.*

Situational Context and Leader Effectiveness

Whereas trust reflects willingness to be vulnerable and satisfaction reflects affective evaluation, leader effectiveness represents a broader judgment regarding the leader's capacity to

facilitate goal attainment and coordinate collective performance (Liden et al., 2025; Yukl, 2012). Effectiveness evaluations integrate perceptions of competence, behavioral appropriateness, and contextual fit. Thus, examining leader effectiveness allows for an assessment of which situational demands best align with AI, human, and AI-augmented leaders.

Leadership effectiveness has long been conceptualized as contingent upon situational demands rather than universally determined by stable traits or behaviors (Fiedler, 1967; Yukl, 2012). Behavioral leadership research identifies two foundational dimensions of leader behavior, initiating structure and consideration, both of which are positively associated with effectiveness across studies (Judge et al., 2004). Initiating structure reflects task-oriented behaviors such as clarifying expectations, coordinating work, and monitoring performance, whereas consideration reflects relational behaviors such as empathy, support, and interpersonal concern.

Critically, however, contingency perspectives emphasize that the relative value of these behaviors depends on contextual demands. Effectiveness emerges not from any single behavioral style, but from alignment between leader behavior and situational requirements (Lambert et al., 2012; Vroom & Yetton, 1973; Yukl, 2012). Recent work has sought to more systematically conceptualize these situational requirements by distinguishing between structural features of leadership contexts (e.g., task characteristics, organizational scope, interpersonal configurations) and leaders' psychological perceptions of situational demands and affordances (Green et al., 2023). Psychological situation characteristics are particularly important because they are more proximal predictors of leader behavior and effectiveness outcomes than objective contextual features themselves. This perspective is consistent with broader leadership theory suggesting that effectiveness emerges from the dynamic interplay between leader attributes, behavioral enactment, and situational constraints or opportunities (Zaccaro et al., 2018).

Although research on AI operating in formal leadership roles remains limited, adjacent literatures on algorithmic evaluation provide insight into how individuals respond to automated versus human decision agents. Studies of algorithm aversion and appreciation indicate that individuals tend to favor algorithmic decision-makers in analytical, data-driven contexts but may resist algorithmic agents in domains requiring subjective judgment or socioemotional sensitivity (Castelo et al., 2019; Dietvorst et al., 2015; Logg et al., 2019). Complementing this, mind perception research demonstrates that individuals attribute lower experiential capacity (e.g., pain, joy, emotion) to machines than to humans (Gray et al., 2007), which may reduce the perceived authenticity of relational behaviors enacted by AI systems. Together, these findings suggest that followers may differentially evaluate the effectiveness of AI versus human leaders depending on whether contextual demands emphasize structural coordination or relational engagement.

Within initiating structure contexts, where task clarity, performance monitoring, and coordination are central, AI leaders may align closely with situational demands. AI systems provide calibrated, probabilistic decision recommendations that are interpreted through analytical utility-based processes (Wang et al., 2022), suggesting alignment with contexts that require initiating structure behaviors. In such contexts, followers may evaluate AI leaders as effective due to perceived analytical competence, consistency, and objectivity. Emerging experimental research on AI team leadership similarly suggests that although AI can structure tasks effectively, relational perceptions and team effectiveness do not consistently exceed those of human leaders (Jia & Zhang, 2024), indicating that AI's advantages may be domain-specific rather than universal.

Conversely, within consideration contexts, where interpersonal sensitivity, empathy, and socioemotional support are salient, human leaders may be advantaged. Consideration behaviors

signal relational investment and emotional understanding, both of which are strongly linked to perceptions of leader effectiveness through their associations with trust and relational quality (Dirks & Ferrin, 2002; Graen & Uhl-Bien, 1995). Given that followers attribute greater emotional experience and moral agency to humans than to machines (Gray et al., 2007), relational behaviors enacted by human leaders may be perceived as more authentic and therefore more effective in socially oriented contexts.

AI-augmented leaders may allow for more flexible evaluations of their leadership. By integrating structure and consideration behaviors, AI-augmented leaders can potentially adapt to contexts requiring both integrating structure and consideration-based leadership. Their dual capacity allows them to align with either task-oriented or relationally oriented behaviors, offering a balanced approach to leadership and fitting the situational demands as needed. However, empirical evidence on how AI-augmented leaders perform across these contexts remains limited. Celestin and Vanitha (2020) found that AI-augmented models were preferred for decisions requiring both analytical precision and ethical sensitivity, while Munshi (2025) reported that AI-augmented leadership was perceived as effective across diverse industries, regardless of context. Furthermore, Lai and Rau (2025) suggest that hybrid leadership can enhance overall effectiveness, but they do not provide insight into AI-augmented leader performance across initiating structure and consideration-based contexts. Given the aforementioned, we adopt the null hypothesis that no significant differences are expected for the perceived effectiveness of AI-augmented leaders between situational contexts (i.e., consideration and initiating structure).

Taken together, research converges to suggest that situational context may moderate the relationship between leader type and perceived effectiveness. In contexts emphasizing initiating

structure, AI leaders may be perceived as relatively effective due to alignment with task-oriented demands. In contexts emphasizing consideration, human leaders may be perceived as more effective due to perceived relational authenticity and emotional capacity. Accordingly, leader type and situational context are expected to interact in predicting perceived leader effectiveness, leading to the following hypothesis:

Hypothesis 4: *Situational context will moderate the relationship between leader type and follower ratings of leader effectiveness, such that human leaders will receive higher effectiveness ratings in consideration contexts, whereas AI leaders will receive higher effectiveness ratings in initiating structure contexts.*

While situational context shapes which leadership behaviors are most diagnostic of effectiveness, decision-making transparency may further condition how leader type is evaluated within those contexts. Transparency reduces uncertainty by clarifying the rationale, criteria, and logic underlying decisions (Colquitt & Rodell, 2011; Norman et al., 2010). However, the extent to which transparency influences perceived effectiveness may depend on both the type of leader and the nature of situational demands.

In initiating structure contexts, effectiveness judgments are likely to hinge on perceptions of analytical competence, consistency, and procedural clarity. Because AI leaders are often perceived as analytically capable but opaque in their internal reasoning processes (Glikson & Woolley, 2020), transparency may serve as a particularly diagnostic cue. Providing clear explanations for AI-generated decisions may reduce ambiguity, enhance perceived legitimacy, and strengthen perceptions of contextual fit. In contrast, human leaders are generally assumed to possess intentional agency and social reasoning capacity (Gray et al., 2007), making their

decision processes more intuitively interpretable. As such, transparency may offer diminishing incremental value for human leaders in highly structured contexts.

In consideration contexts, however, effectiveness evaluations are less dependent on procedural clarity and more reliant on relational authenticity and socioemotional signaling. Transparency may clarify what decision was made but may not fully compensate for perceived deficits in emotional capacity or experiential understanding attributed to AI systems (Gray et al., 2007). Thus, even when AI leaders provide high transparency, they may not fully match the relational effectiveness of human leaders in socially oriented contexts. For AI-augmented leaders, transparency may operate differently. Because hybrid leadership blends algorithmic input with human oversight, transparency may clarify accountability boundaries and strengthen perceptions of balanced competence across contexts.

Taken together, transparency is likely to function as a boundary condition that amplifies or attenuates context-dependent evaluations of leader effectiveness. Specifically, transparency may strengthen perceptions of AI leader effectiveness in initiating structure contexts, where analytical competence is central, but may exert weaker effects in consideration contexts where relational authenticity dominates evaluative criteria. For human and AI-augmented leaders, the influence of transparency may be more stable across contexts.

Given the limited empirical research directly examining these combined dynamics, the present study adopts an exploratory approach:

Research Question 1: *Do leader type, decision-making transparency, and situational context interact to influence follower ratings of leader effectiveness?*

Study 1 Method

Sampling

A priori power analyses using G*Power 3.1 indicated a minimum sample of 251 was needed to achieve 95% power for detecting a medium effect, at a significance criterion of $\alpha = .05$, to test for fixed, main, and interaction effects. A total of 351 undergraduate students from a Midwestern university were recruited through convenience sampling and received course credit for their participation.

Prior to analyses, data were screened for quality. Participants were excluded if they completed the survey in less than 5% or greater than 95% of the average completion time ($n = 45$), failed one or more attention checks ($n = 16$), or provided incomplete responses ($n = 9$). After data cleaning, the final analytic sample consisted of 281 participants (70.8% female; 60.5% white; 63% first-year students; $M_{\text{age}} = 19$, $SD_{\text{age}} = 1.42$).

Design and Procedure

The study employed a 3x2x2 mixed factorial design, with leader (human, AI, AI-augmented) and decision-making transparency (high, low) manipulated between subjects, and situational context (consideration, initiating structure) manipulated within subjects. Participants were randomly assigned to one of six conditions: 1) human leader with high transparency, 2) human leader with low transparency, 3) AI leader with high transparency, 4) AI leader with low transparency, 5) AI-augmented leader with high transparency, and 6) AI-augmented leader with low transparency.

Data collection was conducted via Qualtrics. After providing informed consent and demographic information, participants read a vignette describing a fictional marketing organization (“GemB”) created using established vignette design recommendations (Aguinis and Bradley, 2014). The vignette included an organizational overview, a visual organizational chart,

and email communications between the participant, their supervisor, and a colleague designed to reinforce leader type and transparency (see Appendices A - E).

Following the organizational scenario, participants reported their trust in and satisfaction with the leader. Participants were then presented with two situational contexts (i.e., one emphasizing consideration behaviors and one emphasizing initiating structure behaviors) and rated the anticipated effectiveness of their assigned leader after each context (see Appendix E). Manipulation checks for leader type and transparency, as well as covariate measures, were administered subsequently. Upon completion, participants were debriefed regarding the purpose of the study.

Manipulations

Leader Type

Within the scenario, participants were exposed to either a human, AI, or AI-augmented leader via an organizational chart, a leader description, and email communications (see Appendices A, B, and D). Human leaders were described as emphasizing interpersonal engagement and concern for employees, reflecting traditional leadership processes centered on relationship building and human judgment (Smith & Green, 2018). In contrast, AI leaders were described as emphasizing task execution, information management, and data-driven decision making, consistent with research describing AI agents as leaders that guide teams by assigning tasks, managing information flows, and coordinating work processes rather than engaging in relational leadership behaviors (Jia & Zhang, 2024). AI-augmented leaders were described as integrating aspects of human leadership with AI-supported decision input, reflecting contemporary organizational applications in which AI functions as a decision aid embedded within human leadership rather than as a full replacement for it (Smith & Green, 2018).

To assess whether participants correctly perceived the leader type manipulation, they responded to three items ($\alpha = .91$) reflecting perceptions of the supervisor's interpersonal characteristics using a 7-point Likert scale (1 - '*Strongly Disagree*' to 7 - '*Strongly Agree*'; see Appendix F). These items reflected attributes explicitly embedded in the leader descriptions, including perceived friendliness, consideration, and concern for employees (e.g., "My supervisor is considerate in their interactions with my colleagues and I.").

Decision-making Transparency

To examine the effects of the leader's decision-making transparency, participants were exposed to one of two transparency conditions: high transparency or low transparency. In the high decision-making transparency condition, leaders were portrayed as actively keeping employees informed about organizational developments, explaining the rationale behind decisions, responding to feedback, and aligning their behavior with stated intentions. In contrast, leaders in the low decision-making transparency condition were depicted as withholding information, offering limited justification for decisions, and displaying minimal openness to feedback.

Transparency cues were embedded in emails from both the leader and a colleague to the participant. For example, within the low decision-making transparency condition, the email from the colleague stated, "[The leader] didn't provide much explanation behind recent decisions," whereas the high transparency condition stated, "[The leader] took the time to explain the reasoning behind recent decisions." Additional cues were embedded in the leader's direct communication to reinforce the manipulation. See Appendices D and E for the full emails from the colleague and leader.

To assess the salience of the transparency manipulation, participants responded to four items ($\alpha = .77$) using a 7-point Likert scale (1 - '*Strongly Disagree*' to 7 - '*Strongly Agree*'; see Appendix F). An example item includes, "My supervisor typically provides their reasoning for decisions."

Situational Context

To understand how situational demands can influence perceptions of the leader, all participants were exposed to two distinct contexts: one emphasizing consideration behaviors and the other initiating structure behaviors. This within-subjects manipulation was informed by the Leader Behavior Description Questionnaire XII (LBDQ XII), which distinguishes between consideration, defined as the degree to which a leader shows concern and respect for followers, and initiating structure, defined as the extent to which a leader organizes roles, sets goals, and monitors task execution (Bass, 1990; Fleishman, 1973; Judge et al., 2004).

The consideration context depicted a situation requiring the leader to address interpersonal conflict and foster a supportive team environment, emphasizing behaviors such as supporting, developing, recognizing, and empowering employees. The initiating structure context depicted a data breach requiring coordinated action, emphasizing behaviors such as role clarification, planning, monitoring operations, and problem solving. These scenarios were presented following the primary vignette and were designed to elicit distinct expectations of leader behavior based on situational demands.

To establish the salience and clarity of these contexts, the scenarios were reviewed by graduate students with training in leadership and organizational psychology. These reviewers evaluated whether each scenario clearly represented the intended leadership context and provided feedback on wording and content. Based on this expert review, minor revisions were

made to enhance clarity and ensure that each scenario aligned with its intended context. This process provided evidence of content validity for the situational context manipulation prior to data collection. See Appendix E for the full situational context manipulations.

Dependent Variables

Satisfaction with the Leader

To assess participants' satisfaction with their leader, they responded to three items ($\alpha = .93$) on a 7-point Likert scale (1 - '*Strongly Disagree*' to 7 - '*Strongly Agree*'). Existing satisfaction scales often embed assumptions about leader type or organizational context (Belias & Koustelios, 2014) and focus on assessing employee attitudes about their job (i.e., job satisfaction), limiting their applicability to the current study. Therefore, items were developed to reflect the foundational aspects of employee satisfaction with a leader, including support, comfort, and alignment (Dubey et al., 2023; Rodrigues et al., 2024). An example item was "I would be satisfied working under this supervisor." All items are provided in Appendix G.

Trust in the Leader

To assess trust in the leader, participants completed McAllister's (1995) interpersonal trust measure. The scale consisted of 11 items ($\alpha = .92$) rated on a 7-point Likert scale (1 - '*Strongly Disagree*' to 7 - '*Strongly Agree*'). The measure captures both affect-based and cognition-based trust items. An example affect-based item was "I can freely share my ideas, feelings, and hopes with this supervisor," and an example cognition-based item was "I can rely on this supervisor not to make my job more difficult by careless work." All items are provided in Appendix G.

Leader Effectiveness

Participants rated their leader on a five-item effectiveness measure using a 7-point Likert scale (1 - '*Strongly Disagree*' to 7 - '*Strongly Agree*'). These items were inspired by existing literature and previously established measures of general and managerial effectiveness (Hassan et al., 2013; Kim & Yukl, 1995; Norman et al., 2010). Participants responded to these same five items following both the consideration ($\alpha = .88$) and initiating structure contexts ($\alpha = .86$). An example item was "My supervisor has the necessary skills and expertise to handle the challenges in this situation." All items are provided in Appendix G.

Covariates and Demographics

Participants reported demographic information including age, gender, ethnicity, academic major, year in college, number of business and marketing courses completed, marketing industry experience, and leadership experience. Following the vignette and manipulation checks, participants completed measures of propensity to trust, AI literacy, and personality, which were included as covariates to control for individual differences that may influence participant perceptions of the leader. The full covariate measures and their items are provided in Appendix H.

Propensity to trust is a dispositional trait that reflects an individual's general willingness to rely on others, particularly in contexts where information is limited or ambiguous (Mayer et al., 1995). Using Ashleigh et al.'s (2012) measure, participants rated 20 items across three factors (trusting others, others' reliability and integrity, and risk aversion) on a 7-point Likert rating scale (1 = '*Strongly Disagree*' to 7 = '*Strongly Agree*'). An example item was "Other people do what they say they will do." Including this measure as a covariate allowed us to control for baseline trust tendencies, helping to distinguish between trust that arises from leader behavior (e.g., transparency) and trust that stems from individual predispositions. This was particularly

important when evaluating AI and AI-augmented leaders, whose perceived trustworthiness may have been more variable across individuals depending on their inherent willingness to trust. The scale yielded internal consistency coefficients above .70 (α s = .73 to .87).

AI literacy refers to an individual's ability to understand, evaluate, and ethically engage with AI technologies (Wang et al., 2023). As AI systems increasingly occupy leadership roles, followers' AI literacy was a critical factor in shaping their perceptions of leader competence, transparency, and trustworthiness. Wang et al. (2023) developed and validated a 12-item multidimensional AI literacy measure, which was taken by participants and rated on a 7-point Likert scale (1 = 'Strongly Disagree' to 7 = 'Strongly Agree'). An example item was "I can use AI applications or products to improve my work efficiency." Including AI literacy as a covariate ensured that differences in follower perceptions were not confounded by varying levels of familiarity or competence with AI systems, allowing for a more accurate assessment of leader type and transparency effects. The scale yielded an internal consistency coefficient of .80.

In assessing personality, the John and Srivastava (1999) Big Five trait taxonomy was used, including 44 items to which participants rated the extent to which each statement applied to them on a Likert scale (1 = 'Disagree strongly' to 5 = 'Agree strongly'); an example item was "Generates a lot of enthusiasm." The inclusion of personality as a covariate was supported by extensive literature demonstrating its reliable prediction of behavioral and psychological outcomes across diverse populations and contexts (Javalagi et al., 2024; Shahzad et al., 2021), making it a robust control for individual differences in perception. The scale yielded an internal consistency coefficient above .70 for each of the Big Five traits (α s = .71 to .85).

Study 1 Results

Prior to hypothesis testing, data were screened for missing values, outliers, and violations of relevant statistical assumptions (e.g., normality, homogeneity of variance). Descriptive statistics, zero-order correlations, and reliability coefficients for the primary study variables are presented in Table 1. Manipulation checks were conducted to verify that participants differentiated among leader types and transparency conditions as intended.

Between-subjects analyses of covariance (ANCOVAs) examined the effects of leader type and transparency on follower trust and satisfaction. A mixed-design ANCOVA examined the effects of leader type, transparency, and situational context on follower ratings of leader effectiveness. Across all analyses, omnibus effects were evaluated prior to conducting pairwise comparisons of estimated marginal means for statistically significant effects, with Bonferroni adjustments applied. Results are summarized in Tables 2 and 3 and significant interactions are presented in Figures 1 and 2.

Manipulation Checks

Independent samples t-tests were conducted to evaluate the manipulation of leader type using items reflecting perceptions of the supervisor's interpersonal characteristics. Results indicated that human leaders ($M = 5.41$, $SD = 1.19$) were rated significantly higher than AI leaders ($M = 4.23$, $SD = 1.32$), $t(192) = 6.49$, $p < .001$. AI-augmented leaders ($M = 5.11$, $SD = 1.12$) were also rated significantly higher than AI leaders, $t(180) = 4.81$, $p < .001$. There was no significant difference between human leaders and AI-augmented leaders, $t(180) = 1.70$, $p = .09$, consistent with the conceptualization of AI-augmented leaders as human leaders supported by AI rather than categorically distinct from human leadership.

An independent samples t-test examined whether the transparency manipulation was perceived as intended, revealing a significant difference in perceived transparency between

conditions, $t(277) = -6.53, p < .001$. Participants in the high-transparency condition reported significantly higher perceived transparency ($M = 4.75, SD = 1.11$) than participants in the low-transparency condition ($M = 3.89, SD = 1.10$).

Main and Interaction Effects

Supporting H1a, there was a significant main effect of leader type on trust, $F(2, 267) = 13.40, p < .001, \eta^2 = .09$. Pairwise comparisons indicated that trust ratings were highest for human leaders ($M = 5.00, SE = 0.09$), followed by AI-augmented leaders ($M = 4.66, SE = 0.10$), and lowest for AI leaders ($M = 4.33, SE = 0.09$), with all pairwise differences statistically significant.

Partially supporting H1b, the main effect of leader type on satisfaction was also significant, $F(2, 265) = 8.83, p < .001, \eta^2 = .06$. Pairwise comparisons indicated that satisfaction ratings were higher for human leaders ($M = 5.13, SE = 0.12$) than AI leaders ($M = 4.39, SE = 0.12$). Satisfaction ratings for AI-augmented leaders did not differ significantly from human or AI leaders but did trend in the anticipated direction.

Supporting H2a, decision-making transparency had a significant main effect on trust, $F(1, 267) = 49.16, p < .001, \eta^2 = .16$. Trust was higher in high-transparency conditions ($M = 5.03, SE = 0.07$) than in low-transparency conditions ($M = 4.30, SE = 0.07$).

Similarly, supporting H2b, transparency also had a significant effect on satisfaction, $F(1, 265) = 39.78, p < .001, \eta^2 = .13$. Satisfaction was higher in high-transparency conditions ($M = 5.21, SE = 0.10$) than in low-transparency conditions ($M = 4.31, SE = 0.10$).

Contrary to H3a and H3b, the interaction between leader type and transparency was not significant for trust, $F(2, 267) = .45, p = .64, \eta^2 = .00$, or satisfaction, $F(2, 265) = 1.91, p = .15, \eta^2 = .01$.

Consistent with H4, a significant interaction between leader type and situational context was observed for follower ratings of leader effectiveness, $F(2, 268) = 9.31, p < .001, \eta^2 = .07$. Comparisons across contexts revealed human leaders received higher effectiveness ratings in the consideration context ($M = 5.02, SE = 0.10$) than in the initiating structure context ($M = 4.79, SE = 0.10$). In contrast, AI leaders received higher effectiveness ratings in the initiating structure context ($M = 4.65, SE = 0.10$) than in the consideration context ($M = 4.30, SE = 0.10$). As expected, AI-augmented leaders did not differ significantly in effectiveness ratings across situational contexts.

Furthermore, there was a significant effect of leader type on effectiveness ratings, $F(2, 268) = 7.00, p < .01, \eta^2 = .05$. Follow-up comparisons within the consideration context revealed that human leaders ($M = 5.02, SE = 0.10$) and AI-augmented leaders ($M = 4.81, SE = 0.10$) received higher effectiveness ratings than AI leaders ($M = 4.30, SE = 0.10$). No significant differences were observed between leader types in the initiating structure context.

Addressing RQ1, no significant three-way interaction was observed among leader type, decision-making transparency, and situational context on ratings of effectiveness, $F(2, 268) = 0.24, p = .76, \eta^2 = .00$. However, a significant interaction emerged between transparency and situational context, $F(1, 268) = 4.22, p < .05, \eta^2 = .02$. When transparency was low, leaders were rated as less effective in the consideration context ($M = 4.44, SE = 0.08$) than in the initiating structure context ($M = 4.61, SE = 0.08$). This difference was not observed under high transparency. Pairwise comparisons within contexts revealed higher ratings of leader effectiveness for each context when decision-making transparency was high (consideration: $M = 4.98, SE = 0.08$; initiating structure: $M = 4.92, SE = 0.08$) than when transparency was low.

Study 1 Discussion

The present findings are consistent with the idea that follower trust and satisfaction are likely influenced by assumptions regarding leaders' relational capacity. This aligns with relational leadership perspectives emphasizing that leadership emerges through interpersonal exchanges that signal understanding, responsiveness, and concern for followers (Graen & Uhl-Bien, 1995; Uhl-Bien, 2006). Because human leaders are perceived as capable of interpreting social cues, demonstrating empathy, and responding to employee needs (Kellett et al., 2006; Walter et al., 2011), they may be viewed as better equipped to foster the relational processes that underlie follower trust and satisfaction.

In contrast, followers may perceive AI leaders as less able to engage in relational leadership behaviors, which could contribute to lower trust and satisfaction, even when decision outcomes are perceived as competent or effective. This interpretation aligns with prior research showing that individuals often place greater trust in human decision makers than algorithmic systems, particularly in contexts involving interpersonal judgment (Dietvorst et al., 2015). Such patterns have been attributed to broader beliefs that artificial systems lack the emotional and social capacities central to relational leadership (Martinez & Aldea, 2005).

Extending this logic, AI-augmented leadership arrangements appear to partially mitigate the relational limitations associated with fully AI leaders while not fully replicating the relational advantages attributed to human leaders. Followers evaluated AI-augmented leaders more favorably than AI leaders, but less favorably than human leaders on trust, with satisfaction means trending in the same direction. This suggests that when decision agency is shared between human and AI, followers may perceive stronger interpersonal qualities while still experiencing some uncertainty regarding relational ownership. As a result, AI-augmented leadership may improve

relational evaluations relative to fully AI leadership, while still falling short of leadership that is perceived as fully human.

The results further indicate that decision-making transparency functions as an important informational mechanism through which followers interpret leadership actions and form overall evaluations of the leader. These findings align with prior research suggesting that transparency can foster trust by providing followers with greater insight into how decisions are made and why particular actions are taken (Schnackenberg & Tomlinson, 2016). When leaders communicate the reasoning behind their decisions, followers may experience less uncertainty and feel better able to evaluate the legitimacy of leadership actions. Research on human-AI interaction similarly highlights transparency as an important determinant of trust in algorithmic systems, as providing explanations for automated decisions can help users understand and appropriately calibrate their reliance on those systems (Glikson & Woolley, 2020; Zerilli et al., 2022). Taken together, these findings suggest that transparency may function as a broadly beneficial leadership practice by enhancing followers' confidence in leadership decisions while also improving evaluations of the leader-follower relationship, regardless of whether leadership influence is attributed to human, AI, or AI-augmented leaders.

Study 1 also demonstrated that evaluations of leader effectiveness appear shaped by expectations regarding the alignment between leadership capabilities and situational demands. These findings align with research suggesting that AI systems are often perceived as particularly capable of analytical processing and structured decision making, whereas humans are viewed as able to interpret and respond to demands in social and emotional domains (Jarrahi, 2018), capabilities that align with leadership behaviors associated with initiating structure and consideration (Judge et al., 2004). Followers may therefore perceive AI leaders as better suited

for situations requiring task coordination and structured problem solving, while leadership involving human agency may be viewed as more effective in situations requiring interpersonal support and relational sensitivity. As anticipated, our findings suggest that AI-augmented leadership structures may signal a combination of analytical and relational capacities. This aligns with existing research on hybrid leadership, where AI structure behaviors and human consideration behaviors leverage complementary strengths across different demands (Lai & Rau, 2025).

Furthermore, the influence of decision-making transparency on effectiveness evaluations may be particularly salient in relationally demanding situations where followers seek cues of support, encouragement, and consideration. When leaders provide explanations for how and why decisions are made, followers may be better able to interpret the reasoning underlying leadership actions and evaluate whether those decisions reflect concern for employee needs (Schnackenberg & Tomlinson, 2016). When such explanations are absent, followers may find it more difficult to infer relational intent, potentially undermining perceptions of supportive leadership. Because structure-focused leadership places emphasis on organizing work and achieving objectives (Judge et al., 2004; Yukl, 2012), followers may be somewhat more tolerant of limited decision explanations as long as tasks are completed effectively.

Taken together, the findings from Study 1 indicate that structural cues about leadership (i.e., leader type) and informational cues (i.e., transparency) shape follower reactions, while situational context influences evaluations of leader effectiveness. However, these effects were examined in the presence of explicit behavioral descriptions embedded within contextualized leadership scenarios. Consequently, it remains unclear whether differences in follower evaluations across leader types reflect reactions to explicitly stated leadership behaviors or more

fundamental expectations regarding what leaders are prototypically expected to be like. In other words, would followers evaluate leader types differently in the absence of explicit leader descriptions?

Followers may hold implicit schemas about the traits and characteristics associated with human, AI, and AI-augmented leaders that shape their evaluations independently of observed performance or situational cues. If such schemas differ systematically across leader types, they may help further explain why relational outcomes and effectiveness judgments varied in Study 1. To address this possibility, Study 2 extends the investigation by examining follower perceptions of leader types when explicit behavioral cues are minimized, allowing implicit expectations to guide evaluations.

Study 2

Implicit Theories of Leadership

As discussed, Study 2 builds on the findings of Study 1 by examining whether the differences in follower evaluations of leader type, decision-making transparency, and situational context still hold true when explicit descriptions are absent. Whereas Study 1 focused on reactions to enacted leadership behaviors embedded within contextualized scenarios, the present study explores whether follower perceptions of the leader vary in the absence of behavioral information and descriptive priming.

Leadership perceptions are shaped by followers' ILTs, which represent cognitive schemas or prototypes that individuals hold regarding the traits and characteristics that define effective leaders (Lord et al., 1984; Offermann et al., 1994). Rather than passively responding to leader behavior, followers actively interpret leaders through these internalized expectations. As a

result, leadership is not solely enacted but socially constructed through perceptual categorization processes (Gioia & Sims, 1985).

Research on ILTs demonstrates that individuals evaluate leaders in relation to prototypical attributes such as sensitivity, intelligence, dedication, dynamism, and integrity (Epitropaki & Martin, 2004; Offermann & Coats, 2018). These prototypes shape leader recognition, evaluation, and subsequent relational outcomes. When a target aligns with followers' implicit prototype, they are more likely to be perceived as legitimate and effective; when misalignment occurs, evaluations tend to suffer (Epitropaki & Martin, 2004). Importantly, ILTs operate automatically and often outside conscious awareness, influencing judgments even when behavioral information is limited (Lord & Maher, 2002). In the present study, ratings on these ILT dimensions serve as a structured means of capturing followers' prototype-based perceptions of leader types when behavioral information is minimized.

Although ILT research has historically focused on human leaders, emerging technological contexts introduce agents that may not align neatly with traditional leader prototypes. Because ILTs developed in human-centered leadership environments, it remains unclear whether followers apply the same trait-based expectations to AI leaders or whether distinct schemas are activated when leadership is attributed to an algorithmic system. AI systems may be perceived as high in intelligence and analytical capability, yet lower in sensitivity or warmth. AI-augmented leaders, by integrating algorithmic and human components, may evoke hybrid perceptions that combine structural competence with relational capacity.

As an extension of Study 1, Study 2 directly examines these implicit prototype-based perceptions. To isolate schema-driven evaluations, the study assesses follower perceptions of human, AI, and AI-augmented using ratings on established ILT dimensions. To enable

comparison with Study 1, the same core outcome measures and manipulated variables are retained while introducing an ILT-based assessment of leadership dimensions (Offermann & Coats, 2018). This design allows for examination of whether differences in relational and effectiveness evaluations may originate from implicit categorization processes rather than solely from responses to primed characteristics.

Accordingly, the following research questions are posed:

Research Question 2: *Do follower evaluations observed in Study 1 replicate when priming of leader characteristics is absent?*

Research Question 3: *Do follower ratings on ILT dimensions differ across human, AI, and AI-augmented leaders?*

Study 2 Method

Sampling

Consistent with Study 1, an a priori power analysis conducted using G*Power 3.1 indicated that a minimum of 251 participants was required to achieve 95% power to detect medium-sized fixed, main, and interaction effects at $\alpha = .05$.

A total of 515 undergraduate students were recruited using the same sampling procedures as Study 1. Data screening resulted in the removal of 95 participants (12 withdrew, 28 provided incomplete responses, 27 failed both attention checks, and 28 fell outside the acceptable completion time range [$< 5\%$ or $> 95\%$ of the median completion time]). The final analytic sample consisted of 420 participants (81.8% female; 58.1% white; 78.7% first-year students; $M_{\text{age}} = 19$, $SD_{\text{age}} = 1.99$).

Design and Procedure

Study 2 employed the same 3x2x2 mixed factorial design as in Study 1, with two between-subjects factors (i.e., leader type and decision-making transparency) and one within-subjects factor (i.e., situational context). Participants were again randomly assigned to one of the six experimental conditions.

Procedures largely replicated Study 1. After providing informed consent and demographic information, participants reviewed an organizational overview, organizational chart, and email communications. One important difference from Study 1 was that leader characteristics were not explicitly described or primed; participants were only provided with a label indicating the leader type (i.e., human leader, AI leader, or human leader who uses AI), without additional descriptive cues.

Participants then completed measures assessing trust in and satisfaction with the leader. Before proceeding to the situational tasks, participants completed a measure of ILT dimensions in which they rated how characteristic various attributes were of the assigned leader type (human, AI, or AI-augmented). This measure captured general perceptions of the leader type and was not tied to the specific scenarios or transparency condition. Next, participants reviewed the same situational contexts as in Study 1, rating the anticipated effectiveness of their leader in each context. For Study 2, manipulation checks for the situational contexts were employed after each context excerpt. Similar to Study 1, the rest of the manipulation checks and covariate measures were administered after the scenario and outcome measures. Participants were debriefed at the conclusion of the study.

Manipulations

Leader Manipulation

Participants in Study 2 were assigned to the same leader types as in Study 1. Unlike Study 1, which primed participants with leader characteristics, Study 2 did not include any explicit descriptions of leader traits or behaviors. Instead, leader type (human, AI, or AI-augmented) was identified only through the organizational overview, chart, and titles within the email chain. The emails referenced the leader by name and role but did not include descriptive characteristics.

To assess the salience of the manipulation, Study 2 used a single item asking participants to “Select which of the following best describes your supervisor” to which they selected either a human leader, an AI leader, or a human leader who used AI (AI-augmented).

Transparency Manipulation

As in Study 1, participants were exposed to leaders who exhibited either high or low decision-making transparency. Transparency cues were embedded within emails from both the leader and a colleague directed to the participant. Importantly, the transparency manipulation was implemented in a parallel manner across leader types. For example, in the high transparency condition, an email from the leader stated, “These next steps are based on a combination of past market trends, real-time consumer engagement data, and team discussions.” In contrast, the low transparency condition stated, “These next steps are set, and we will move forward as planned with the launch of PureFusion.”

Participants took the same four-item manipulation check measure ($\alpha = .81$) as in Study 1, using a 7-point Likert scale (1 - ‘*Strongly Disagree*’ to 7 - ‘*Strongly Agree*’; see Appendix F).

Situational Context Manipulation

Study 2 replicated Study 1’s situational context excerpts for consideration and initiating structure. See Appendix E for both situational context excerpts.

In Study 2, a manipulation check was implemented to ensure accurate perception of the behavioral nuances between them. A total of 8-items, with four items reflecting consideration behaviors ($\alpha = .82$) and four items reflecting initiating structure behaviors ($\alpha = .78$), were rated by participants on a 7-point Likert scale (1 - '*Strongly Disagree*' to 7 - '*Strongly Agree*'). An example item for consideration includes, "This situation required my leader to support employees." An example item for initiating structure includes, "This situation required my leader to plan ahead." See Appendix F for all situational context manipulation check items.

Outcomes

All dependent variables and their corresponding measures from Study 1 were retained in Study 2, including trust in the leader ($\alpha = .94$), satisfaction with the leader ($\alpha = .94$), and effectiveness ratings for consideration-based contexts ($\alpha = .81$) and initiating structure-based contexts ($\alpha = .81$).

In addition, participants completed an ILT dimension measure in which item instructions and response anchors were tailored to reference the specific leader type to which participants were assigned (i.e., human, AI, or AI-augmented). For example, participants in the AI-augmented condition were instructed to "Please read each descriptor below and rate how well it describes an AI-supported leader (a human leader who uses AI in their day-to-day work)." Participants then rated 48 ILT descriptors on a 10-point Likert scale (e.g., 1 - '*Not at All Characteristic of an AI-supported leader*' to 10 - '*Extremely Characteristic of an AI-supported leader*'). Parallel wording was used across conditions such that references to leader type in both the instructions and scale anchors corresponded directly to the experimental condition.

This scale was directly informed by Offermann & Coats (2018) and represented the following ILT dimensions: Sensitivity, Dedication, Tyranny, Charisma, Strength, Creativity,

Well-groomed, Masculinity, and Intelligence. Example descriptors include “compassionate” (Sensitivity), “coercive” (Tyranny), “strong” (Strength), and “educated” (Intelligence). The scale yielded an internal consistency coefficient above .90 for all traits (α s = .91 to .96). All outcome measures and the full list of ILT descriptors are provided in Appendix G.

Covariates and Demographics

Study 2 retained the same covariates (propensity to trust, α = .72; AI literacy, α = .79; personality, α s = .70 to .84) and demographics included in Study 1. In addition, Study 2 incorporated a measure assessing trust in and functionality of AI technologies. Prior research indicates that functionality and trust in AI are associated with perceived usefulness, attitudes, and acceptance of AI technologies (Choung et al., 2022). This construct helps account for individual differences in baseline trust in and perceptions of AI that may influence evaluations of AI-based and AI-augmented leaders. Using Choung et al.’s (2022) measure of functionality and trust in AI, participants rated 24 items across five dimensions (trust in the AI, perceived usefulness, attitude toward the AI technologies, human-like trust in smart technologies, and functionality trust in smart technologies), using a 5-point Likert scale (1 - ‘*Strongly Disagree*’ to 5 - ‘*Strongly Agree*’). Example items include “Smart technologies are sincerely concerned about addressing the problems of human users” (human-like trust) and “Smart technologies have the features necessary to complete key tasks” (functionality trust). The scale yielded an internal consistency coefficient above .80 for each trait (α s = .80 to .91). The full covariate measures and their items are provided in Appendix H.

Study 2 Results

As in Study 1, Study 2 data were screened for missing values, outliers, and violations of relevant statistical assumptions. Descriptive statistics, correlations, and reliabilities for the

primary study variables are presented in Tables 4 and 5. Manipulation checks for accurate perceptions of leader type, transparency, and situational context were conducted.

Between-subjects and mixed-design ANCOVAs were conducted to assess replication across studies. Levene's tests indicated heterogeneity of variance for trust and satisfaction. However, variance ratios were moderate, and group sizes were highly balanced across conditions ($n_s = 63 - 76$). Research indicates that the ANOVA F-test is generally robust to moderate heterogeneity when sample sizes are approximately equal (Blanca et al., 2018; Glass et al., 1972). Kim and Cribbie's (2018) simulation work further indicates that balanced designs are unlikely to produce substantial Type I error distortion under such conditions. No evidence of negative pairing between group sizes and variances was observed. Accordingly, the original ANCOVA results are reported.

Furthermore, a multivariate analysis of variance (MANOVA) and a series of univariate follow-up analyses were conducted to explore the effect of leader type on ILT dimensions scores. Results for Study 2 analyses are summarized in Tables 6 and 7 and significant interactions can be seen in Figures 3-5.

Manipulation Checks

Given how leader type was manipulated in Study 2, a Pearson chi-square test of independence was used to test if participants identified the correct leader type. This test indicated a significant association between actual leader type and perceived leader type, $\chi^2(4, N = 418) = 149.41, p < .001$. Observed frequencies indicated that participants were significantly more likely to correctly identify the leader type corresponding to their assigned condition than would be expected by chance, supporting the effectiveness of the leader assignments.

An independent samples t-test indicated a significant difference between transparency conditions, $t(398.07) = -15.07, p < .001$. Participants in the high-transparency condition reported significantly higher perceived transparency ($M = 5.14, SD = 0.97$) than participants in the low-transparency condition ($M = 3.56, SD = 1.16$). Levene's test indicated unequal variances; therefore, Welch's correction was applied. These findings support the effectiveness of the transparency manipulation.

Paired samples t-tests were conducted to determine whether each context elicited the intended leadership behaviors. Within the consideration context, participants rated consideration behaviors ($M = 5.41, SD = 0.98$) higher than initiating structure behaviors ($M = 4.92, SD = 1.19$), $t(416) = 8.06, p < .001$. Within the initiating structure context, participants rated initiating structure behaviors ($M = 5.86, SD = 1.07$) higher than consideration behaviors ($M = 5.41, SD = 0.97$), $t(415) = 9.98, p < .001$. Together, these results indicate that the situational context manipulation successfully differentiated the two situations in the intended direction.

Main and Interaction Effects

Addressing RQ2, analyses examining the main effect of leader type on trust was not statistically significant, $F(2, 409) = 2.75, p = .07, \eta^2 = .01$, indicating that trust did not differ significantly across leader types. However, the means for trust were trending in the expected direction.

Moreover, replicating Study 1 results, the main effect of leader type on satisfaction was statistically significant, $F(2, 410) = 3.68, p < .05, \eta^2 = .02$. Pairwise comparisons showed that satisfaction was significantly higher for human leaders ($M = 4.41, SE = 0.11$) than for AI leaders ($M = 3.99, SE = 0.11$). Satisfaction ratings for AI-augmented leaders did not differ significantly from human or AI leaders but did trend in the anticipated direction.

Replicating Study 1, the main effect of decision-making transparency on trust was significant, $F(1, 409) = 267.22, p < .001, \eta^2 = .40$. Participants reported higher trust in the high-transparency condition ($M = 4.92, SE = 0.07$) than in the low-transparency condition ($M = 3.32, SE = 0.07$). Notably, the magnitude of this effect was larger than that observed in Study 1 ($\eta^2 = .15$).

Similarly, transparency also had a significant main effect on satisfaction, $F(1, 410) = 256.55, p < .001, \eta^2 = .39$, replicating Study 1. Satisfaction ratings were higher in the high-transparency condition ($M = 5.20, SE = 0.09$) compared to the low-transparency condition ($M = 3.17, SE = 0.09$). This effect was also larger than the corresponding effect in Study 1 ($\eta^2 = .13$).

Unlike Study 1, the interaction effect of leader type and transparency on trust was significant, $F(2, 409) = 9.27, p < .001, \eta^2 = .04$. Contradicting the anticipated direction in H3a from Study 1, the increase in trust associated with transparency was smaller for AI leaders ($\Delta M = 1.01$) than for human ($\Delta M = 1.86$) and AI-augmented leaders ($\Delta M = 1.92$). Planned interaction contrasts revealed that the transparency effect was significantly stronger for human leaders than for AI leaders, $\Delta\Delta M = 0.86, SE = 0.23, p < .001, 95\% \text{ CI } [0.40, 1.32]$, and significantly stronger for AI-augmented leaders than for AI leaders, $\Delta\Delta M = 0.91, SE = 0.24, p < .001, 95\% \text{ CI } [0.44, 1.38]$. The transparency effect did not significantly differ between human and AI-augmented leaders, $\Delta\Delta M = -.05, SE = 0.24, p = .83, 95\% \text{ CI } [-0.52, 0.42]$.

A similar interaction emerged for satisfaction, $F(2, 410) = 8.62, p < .001, \eta^2 = .04$. Contradicting the anticipated direction of H3b in Study 1, the increase in satisfaction associated with transparency was smaller for AI leaders ($\Delta M = 1.29$) than for human ($\Delta M = 2.38$) and AI-augmented leaders ($\Delta M = 2.42$). Planned contrasts showed that the transparency effect on satisfaction was significantly stronger for human leaders than for AI leaders, $\Delta\Delta M = 1.09, SE =$

0.31, $p < .001$, 95% CI [0.49, 1.69], and significantly stronger for AI-augmented leaders than for AI leaders, $\Delta\Delta M = 1.13$, $SE = 0.31$, $p < .001$, 95% CI [.51, 1.74]. As with trust, the transparency effect did not significantly differ between human and AI-augmented leaders, $\Delta\Delta M = -.03$, $SE = 0.31$, $p = .91$, 95% CI [-0.64, 0.57].

A significant interaction between leader type and situational context was observed for follower ratings of leader effectiveness, $F(2, 408) = 12.61$, $p < .001$, $\eta^2 = .06$. Similar to Study 1, comparisons across contexts within each leader type showed that AI leaders received higher ratings of effectiveness within the initiating structure context ($M = 4.49$, $SE = 0.10$) than the consideration context ($M = 3.90$, $SE = 0.09$). However, human leaders did not differ significantly across contexts. Consistent with Study 1, AI-augmented leaders did not differ significantly in effectiveness across situational contexts.

Furthermore, there was a significant effect of leader type on effectiveness ratings, $F(2, 408) = 7.09$, $p < .001$, $\eta^2 = .03$. Consistent with Study 1, follow-up comparisons revealed that, within the consideration context, human leaders ($M = 4.46$, $SE = 0.09$) received higher effectiveness ratings than AI leaders ($M = 3.90$, $SE = 0.09$). However, human leaders also received higher effectiveness ratings than AI-augmented leaders ($M = 4.12$, $SE = 0.10$) within the consideration context. In contrast, within the initiating structure context, AI leaders and human leaders ($M = 4.60$, $SE = 0.09$) had higher effectiveness ratings than AI-augmented leaders ($M = 4.13$, $SE = 0.09$).

As in Study 1, the three-way interaction among leader type, transparency, and situational context was not statistically significant for follower ratings of effectiveness, $F(2, 408) = 0.74$, $p = .48$, $\eta^2 = .00$. Unlike Study 1, the two-way interaction among transparency and situational context was not significant, $F(1, 408) = 0.34$, $p = .56$, $\eta^2 = .00$.

Addressing RQ3, a MANOVA indicated that leader type did not have a significant multivariate effect on ILT scores, $F(18, 802) = 1.09$, $p = .36$, $\eta^2 = .02$. Follow-up univariate analyses also revealed no significant differences between leader types on any individual ILT dimension (p -values $> .05$). Therefore, results indicate participants did not endorse different ILT dimensions for human, AI, or AI-augmented leaders.

Study 2 Discussion

Study 2 extended the investigation by examining how followers evaluate human, AI, and AI-augmented leaders when explicit behavioral descriptions are minimized, thereby allowing implicit expectations about leader capability to play a larger role in shaping evaluations.

The findings suggest that trust evaluations may be particularly sensitive to the availability of diagnostic cues regarding leader capability and relational intent. In contrast, satisfaction with the leader appeared more resilient to reductions in explicit behavioral information, remaining influenced by generalized expectations of the leader. This differential pattern may reflect conceptual distinctions between trust and satisfaction highlighted in the leadership literature. Because trust involves a willingness to accept vulnerability (Mayer et al., 1995), it may depend more heavily on the availability of diagnostic cues regarding leader competence, benevolence, and intent. Comparatively, satisfaction represents a more global evaluative response to the leader-follower relationship (Scarpello & Vandenberg, 1987), which may be shaped by broader affective impressions and expectations about leadership. Consequently, when explicit behavioral information was reduced in Study 2, trust differentiation across leader types may have weakened due to increased ambiguity, whereas satisfaction differences persisted because followers could rely more readily on generalized leadership prototypes when forming overall evaluations of the leader.

Furthermore, when explicit descriptors of leader characteristics are limited, follower evaluations may become particularly sensitive to more observable features of the decision-making process, such as transparency, rather than on assumptions about the underlying leadership agent. Although transparency has been shown to reduce uncertainty and enhance perceptions of accountability by providing insight into how decisions are made (Glikson & Woolley, 2020; Schnackenberg & Tomlinson, 2016), its evaluative meaning may depend on how clearly leadership agency is perceived. Explanations attributed to AI systems may be interpreted primarily as informational outputs rather than as relational signals communicated by an intentional actor. As discussed by Ananny and Crawford (2018), transparency may not fully offset followers' uncertainty regarding the relational and moral capacities of algorithmic decision makers. In contrast, transparency within leadership that involves human agency may be more readily understood as signaling openness, responsibility, and consideration toward followers. Taken together, this perspective suggests that the influence of transparency on leadership evaluations may depend not only on the presence of explanations, but also on whether those explanations are perceived as reflecting human relational intent.

Effectiveness evaluations again suggest the importance of alignment between leadership capabilities and situational context. In consideration-based situations, leadership influence that is clearly attributed to human agency may be viewed as particularly appropriate, as followers are likely to expect interpersonal competence, understanding, and social responsiveness. Notably, the present study reduced the availability of explicit descriptions of leader capabilities, requiring followers to rely more heavily on inferences about the leadership agent when forming effectiveness judgments. Without these descriptive cues, AI-augmented leadership arrangements may introduce greater uncertainty regarding the source of decision influence. When leadership

input is jointly associated with human and automated judgment, followers may find it more difficult to form distinct expectations about which capabilities are driving decisions. Prior research suggests that individuals evaluate automated and human leadership agents based on different perceived strengths (Höddinghaus et al., 2021), which may complicate follower evaluations when both sources are involved in decision making. In contrast, when leadership influence is clearly attributable to either a human or an AI system, followers may form more coherent expectations about leader strengths, allowing for clearer differentiation in effectiveness evaluations across situational contexts.

Furthermore, the absence of differentiation across ILT dimensions suggests that leadership schemas for AI agents may remain underdeveloped or difficult to operationalize using traditional trait-based ILT measures. ILT suggests that individuals rely on established leadership prototypes when evaluating leaders, and these prototypes are historically rooted in experiences with human leaders (Lord et al., 1984; Epitropaki & Martin, 2004). Research on algorithmic decision-making also suggests that individuals hold inconsistent expectations regarding algorithmic versus human judgment (Logg et al., 2019), which may further limit the application of traditional leadership trait schemas to AI agents. As a result, participants may have relied more on situational information and decision cues when forming their evaluations.

While several patterns observed in Study 1 replicated in Study 2, differences between leader types were less pronounced when followers were required to infer leader attributes based on a label rather than explicit descriptions. These findings suggest that follower evaluations of AI, human, and AI-augmented leaders may depend not only on the characteristics of the leadership agent, but also on the availability of information that allows followers to interpret those characteristics.

General Discussion

Across two studies, this research examined how followers evaluate leadership enacted by human, AI, and AI-augmented leaders while also considering the role of decision-making transparency and situational context. Taken together, the studies suggest that follower evaluations of emerging leadership agents are shaped by existing leader schemas and the availability of informational and contextual cues. Importantly, the present findings reflect follower perceptions and attributions regarding leadership agents rather than objective differences in the actual relational or analytical capabilities of human and AI leaders.

First, we see that relational reactions toward leadership appear grounded in expectations regarding interpersonal competence and benevolent intent, suggesting that followers value the human side of leaders when evaluating outcomes of trust and satisfaction. Second, decision-making transparency consistently served as a central interpretive cue allowing followers to make sense of leadership decisions. Providing explanations for leadership decisions consistently improved evaluations of trust and satisfaction and influenced perceptions of effectiveness across contexts. Notably, the influence of transparency became especially salient when followers had limited information about leader characteristics. Third, evaluations of leadership effectiveness reflected perceived alignment between agent capabilities and situational task demands. AI-augmented leaders generally received more moderate or stable evaluations across contexts. One possible explanation is that AI-augmented leadership structures may create ambiguity regarding the source of decision authority. When leadership decisions are attributed jointly to a human and an AI system, followers may find it more difficult to infer whether decisions reflect human relational judgment or algorithmic analysis, leading to more moderate evaluations relative to clearly human or clearly AI leadership. Finally, the results of Study 2 suggested that when

explicit information about leader characteristics is limited, followers rely more heavily on available cues, such as leader transparency or the situational context, to interpret leadership actions.

Together, these findings highlight the importance of relational characteristics, presence of transparency in decision making, and fit with the situational demands in shaping positive follower evaluations of emerging leadership agents.

Theoretical Implications

This research contributes to leadership theory by examining how followers evaluate leadership enacted by agents with varying levels of AI involvement. The pattern of follower responses suggests that evaluations may shift from relationally grounded assumptions about leadership toward more capability-based interpretations when leadership influence is attributed to technological or hybrid agents. Traditional perspectives conceptualize leadership as a socially embedded influence process enacted through interpersonal interaction and behavioral exchange (Oc & Bashshur, 2013; Yukl, 2012). As algorithmic systems increasingly assume functions related to coordination, monitoring, and decision structuring in organizations (Banks et al., 2022; Bevilacqua et al., 2025), leadership theory may need to more explicitly account for influence processes distributed across human and technological agents.

The findings also extend situational perspectives on leadership by highlighting the importance of both structural task demands and psychologically perceived situational requirements in shaping leadership evaluations. Leadership research has long emphasized the relevance of contextual contingencies for determining the effectiveness of particular leadership behaviors (Judge et al., 2004; Yukl, 2012), yet recent work calls for more systematic conceptualization of leadership situations in terms of underlying situational affordances and

constraints (Green et al., 2023; Zaccaro et al., 2018). The present research aligns with these perspectives by suggesting that judgments of leadership effectiveness may increasingly depend on perceived alignment between the capabilities attributed to different leadership agents and the situational demands they are expected to address. When such alignment is lacking, followers may question the appropriateness of technologically mediated leadership influence. This interpretation underscores the value of integrating situational taxonomies and context-sensitive theorizing into emerging models of technologically mediated leadership.

Decision-making transparency represents an important informational mechanism shaping how followers interpret leadership enacted through technological systems. Prior research identifies transparency as central to trust formation by reducing uncertainty surrounding decision processes and enabling evaluations of procedural legitimacy (Burrell, 2016; Glikson & Woolley, 2020; Schnackenberg & Tomlinson, 2016). In technologically mediated leadership contexts, transparency may function as a cue supporting the calibration of follower trust, particularly when expectations regarding the capabilities and accountability of nonhuman leadership agents are less clearly defined. Under such conditions, followers may rely more heavily on information about decision rationale and authority distribution when forming judgments about leadership credibility. These dynamics highlight the need for greater integration between theories of trust in leadership and research on trust in AI.

Finally, the findings offer implications for ILT by extending schema-based perspectives on leader evaluation into technologically mediated contexts. ILT research suggests that followers rely on cognitive prototypes developed through socialization and experience with human leaders when interpreting leadership behavior (Epitropaki & Martin, 2004; Lord et al., 1984). When leadership influence involves AI, followers may possess less clearly differentiated schemas for

evaluating such leadership. In the absence of strongly diagnostic behavioral cues, evaluations may therefore depend more heavily on contextual information and decision process characteristics. This interpretation is consistent with broader research showing that individuals' reliance on algorithmic versus human judgment varies across tasks and decision environments (Logg et al., 2019). Together, these perspectives suggest that ILT research may benefit from considering how leadership prototypes adapt as technologically mediated forms of influence become more prevalent.

Practical Implications

The findings hold implications for organizations looking to integrate AI into leadership and decision-making processes, as they can benefit from carefully considering how technological systems are positioned within leadership roles. Acceptance of AI-enabled leadership is likely to depend in part on whether the capabilities associated with AI systems are perceived as appropriate for the situational demands of leadership tasks. Organizations may therefore facilitate smoother adoption by aligning AI deployment with responsibilities involving structured coordination, monitoring, information synthesis, and analytical decision support. In contrast, leadership responsibilities that require sustained interpersonal engagement, conflict management, or socioemotional support may continue to be more strongly associated with expectations of human leadership. Thoughtful role design that reflects these distinctions may help organizations leverage technological capabilities while maintaining follower confidence in leadership effectiveness.

The findings also highlight the importance of communication practices surrounding AI-supported decision processes. Followers are likely to form evaluations of leadership not only based on outcomes but also on their understanding of how decisions are generated and who holds

responsibility for them. Organizations may therefore benefit from establishing norms that clarify the role of AI in decision-making, including transparency behind the sources of information used, the degree of human involvement, and the basis on which final decisions are made. Providing this clarity may reduce uncertainty surrounding technologically mediated leadership and support more favorable reactions to AI-enabled authority.

Lastly, the results suggest that AI-augmented leadership arrangements may offer practical advantages by combining the analytical strengths associated with AI systems with the relational signaling associated with human leadership. Structuring leadership processes so that AI contributes insights or recommendations while identifiable human leaders retain responsibility for interpretation, communication, and relationship management may help preserve perceptions of accountability and legitimacy. Such configurations may allow organizations to capture efficiency gains from technological systems without undermining the interpersonal foundations of leadership influence.

Limitations and Future Directions

Several limitations should be considered when interpreting the findings of this research. First, both studies relied on vignette-based experimental scenarios in which participants evaluated hypothetical leadership situations. Vignette methodologies allow researchers to systematically manipulate variables while maintaining experimental control and strengthening causal inference (Aguinis & Bradley, 2014). In the present research, this approach allowed leader type, decision-making transparency, and situational context to be isolated while holding other factors constant. While hypothetical scenarios may not capture the full complexity of leadership dynamics in organizational environments, the scenarios were developed to be engaging and realistic; these scenarios describe very commonly occurring leadership demands familiar to a

variety of jobs. Additionally, the benefits of control in experiments enabling the analysis of causal effects may outweigh the loss of ecological validity. However, future research should extend these findings by examining follower responses to AI, human, and AI-augmented leadership in applied or field settings.

Second, generalizability of the findings to other samples would need to be illustrated. While student samples may have less direct experience with organizational leadership structures, they are frequently used in experimental research examining perceptual processes (Peterson, 2001). Moreover, younger populations often demonstrate greater adoption and use of emerging technologies (Venkatesh et al., 2012), making them relevant for studying perceptions of AI-enabled leadership. Future research should examine whether similar patterns emerge among working professionals with more extensive organizational experience.

Third, the present studies captured follower reactions to leadership based on a single decision episode. In organizational contexts, follower perceptions of the leader often develop through repeated interactions over time (Dirks & Ferrin, 2002). Evaluating leadership based on a single scenario allowed the effects of the manipulations to be examined in a controlled manner, but it may reflect initial impressions rather than fully developed leader-follower relationships. Future research could investigate how perceptions of AI, human, and AI-augmented leaders evolve across multiple interactions or longer-term working relationships.

Finally, an additional limitation concerns the measurement of ILTs. The ILT scale used in the present research was originally developed to assess perceptions of human leaders and includes descriptors referencing physical appearance and interpersonal style (e.g., well-groomed, masculine, attractive; Offermann & Coats, 2018). Such traits may be less easily mapped onto nonhuman agents such as AI systems, potentially making it more difficult for participants to

form differentiated evaluations of AI leaders using traditional trait-based prototype measures. More broadly, AI-enabled leadership may represent an emerging category of leadership for which followers possess less clearly defined cognitive schemas. As a result, conventional ILT measurement approaches may not fully capture how followers conceptualize leadership enacted by technological agents. Future research may benefit from developing alternative methods for assessing leadership schemas in AI contexts, such as behavioral expectation measures, categorization tasks, or qualitative elicitation techniques that allow followers to describe perceived leadership capabilities and limitations more directly.

Conclusion

As AI becomes increasingly integrated into organizational decision-making, understanding how followers evaluate AI-enabled leadership is becoming increasingly important. Addressing this need, the present research examined how followers perceive human, AI, and AI-augmented leaders while accounting for decision-making transparency and situational context. Across two studies, the findings suggest that followers continue to value relational capabilities associated with human leadership, while AI leaders may be perceived as relatively effective in structured, task-oriented contexts. Study 2 further demonstrated that when explicit leader characteristics are absent, follower evaluations rely more heavily on contextual and informational cues such as transparency. Together, these findings highlight that perceptions of emerging leadership agents depend not only on leader type, but also on how leadership decisions are communicated and how well leadership aligns with situational demands.

References

- Aguinis, H., & Bradley, K. J. (2014). Best practice recommendations for designing and implementing experimental vignette methodology studies. *Organizational research methods, 17*(4), 351-371.
- Ananny, M., & Crawford, K. (2018). Seeing without knowing: Limitations of the transparency ideal and its application to algorithmic accountability. *new media & society, 20*(3), 973-989.
- Araujo, T., Helberger, N., Kruikemeier, S., & De Vreese, C. H. (2020). In AI we trust? Perceptions about automated decision-making by artificial intelligence. *AI & society, 35*(3), 611-623.
- Ashleigh, M. J., Higgs, M., & Dulewicz, V. (2012). A new propensity to trust scale and its relationship with individual well-being: implications for HRM policies and practices. *Human Resource Management Journal, 22*(4), 360-376.
- Banks, G. C., Dionne, S. D., Mast, M. S., & Sayama, H. (2022). Leadership in the digital era: A review of who, what, when, where, and why. *The Leadership Quarterly, 33*(5), 101634.
- Bass, B. M. (1990). From transactional to transformational leadership: Learning to share the vision. *Organizational dynamics, 18*(3), 19-31.
- Bass, B. M., & Stogdill, R. M. (1990). *Bass & Stogdill's handbook of leadership: Theory, research, and managerial applications*. Simon and Schuster.
- Belias, D., & Koustelios, A. (2014). Leadership and job satisfaction--A review. *European Scientific Journal, 10*(8).

- Bevilacqua, S., Masárová, J., Perotti, F. A., & Ferraris, A. (2025). Enhancing top managers' leadership with artificial intelligence: insights from a systematic literature review. *Review of Managerial Science*, 1-37.
- Blanca, M. J., Alarcón, R., Arnau, J., Bono, R., & Bendayan, R. (2018). Effect of variance ratio on ANOVA robustness: Might 1.5 be the limit?. *Behavior research methods*, 50(3), 937-962.
- Bligh, M. C. (2016). Leadership and trust. In *Leadership today: Practices for personal and professional performance* (pp. 21-42). Cham: Springer International Publishing.
- Brower, H. H., Schoorman, F. D., & Tan, H. H. (2000). A model of relational leadership: The integration of trust and leader-member exchange. *The leadership quarterly*, 11(2), 227-250.
- Burrell, J. (2016). How the machine 'thinks': Understanding opacity in machine learning algorithms. *Big data & society*, 3(1), 2053951715622512.
- Castaneda, M., & Nahavandi, A. (1991). Link of manager behavior to supervisor performance rating and subordinate satisfaction. *Group & Organization Studies*, 16(4), 357-366.
- Castelo, N., Bos, M. W., & Lehmann, D. R. (2019). Task-dependent algorithm aversion. *Journal of marketing research*, 56(5), 809-825.
- Celestin, M., & Vanitha, N. (2020). AI and the future of leadership: Navigating the human-machine collaboration. In *4th International Web Conference on Emerging Trends in Arts, Science, Engineering & Technology (IWCETASET-2020)* (pp. 247-254).
- Choung, H., David, P., & Ross, A. (2023). Trust in AI and its role in the acceptance of AI technologies. *International Journal of Human-Computer Interaction*, 39(9), 1727-1739.

- Colquitt, J. A., Scott, B. A., & LePine, J. A. (2007). Trust, trustworthiness, and trust propensity: a meta-analytic test of their unique relationships with risk taking and job performance. *Journal of applied psychology*, 92(4), 909.
- Colquitt, J. A., & Rodell, J. B. (2011). Justice, trust, and trustworthiness: A longitudinal analysis integrating three theoretical perspectives. *Academy of management journal*, 54(6), 1183-1206.
- Colquitt, J. A., & Baer, M. D. (2023). Foster trust through ability, benevolence, and integrity. *Principles of Organizational Behavior: The Handbook of Evidence-Based Management 3rd Edition*, 345-363.
- Cushman, D. P., & Craig, R. T. (1976). Communication systems: Interpersonal implications. In G. R. Miller (Ed.), *Exploration in interpersonal communication* (pp. 37-58). Beverly Hills, CA: Sage.
- de Fine Licht, K., & de Fine Licht, J. (2020). Artificial intelligence, transparency, and public decision-making: Why explanations are key when trying to produce perceived legitimacy. *AI & society*, 35(4), 917-926.
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: people erroneously avoid algorithms after seeing them err. *Journal of experimental psychology: General*, 144(1), 114.
- Dirks, K. T., & Ferrin, D. L. (2002). Trust in leadership: meta-analytic findings and implications for research and practice. *Journal of applied psychology*, 87(4), 611.
- Dubey, P., Pathak, A. K., & Sahu, K. K. (2023). Assessing the influence of effective leadership on job satisfaction and organisational citizenship behaviour. *Rajagiri Management Journal*, 17(3), 221-237.

- Epitropaki, O., & Martin, R. (2004). Implicit leadership theories in applied settings: factor structure, generalizability, and stability over time. *Journal of applied psychology*, 89(2), 293.
- Erfanian, F., Roudsari, R. L., Haidari, A., & Bahmani, M. N. D. (2020). A Narrative on the Use of Vignette: Its Advantages and Drawbacks. *Journal of Midwifery & Reproductive Health*, 8(2).
- Faul, F., Erdfelder, E., Buchner, A., & Lang, A.-G. (2009). Statistical power analyses using G*Power 3.1: Tests for correlation and regression analyses. *Behavior Research Methods*, 41, 1149-1160.
- Fiedler, F. E. (1967). A Theory of Leadership Effectiveness. McGraw-Hill Series in Management.
- Fleishman, E. (1973). Twenty years of consideration and structure. *Current developments in the study of leadership*, 1-40.
- Forsyth, D. R. (2014). How do leaders lead? Through social influence. In *Conceptions of leadership: Enduring ideas and emerging insights* (pp. 185-200). New York: Palgrave Macmillan US.
- Gioia, D. A., & Sims Jr, H. P. (1985). On avoiding the influence of implicit leadership theories in leader behavior descriptions. *Educational and Psychological Measurement*, 45(2), 217-232.
- Glass, G. V., Peckham, P. D., & Sanders, J. R. (1972). Consequences of failure to meet assumptions underlying the fixed effects analyses of variance and covariance. *Review of educational research*, 42(3), 237-288.
- Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of management annals*, 14(2), 627-660.

- Graen, G. B., & Uhl-Bien, M. (1995). Relationship-based approach to leadership: Development of leader-member exchange (LMX) theory of leadership over 25 years: Applying a multi-level multi-domain perspective. *The leadership quarterly*, 6(2), 219-247.
- Gray, H. M., Gray, K., & Wegner, D. M. (2007). Dimensions of mind perception. *science*, 315(5812), 619-619.
- Green, J. P., Dalal, R. S., Fyffe, S., Zaccaro, S. J., Putka, D. J., & Wallace, D. M. (2023). An empirical taxonomy of leadership situations: Development, validation, and implications for the science and practice of leadership. *Journal of Applied Psychology*, 108(9), 1515.
- Greenberg, J., & Cropanzano, R. (1993). The social side of fairness: Interpersonal and informational classes of organizational justice. *Justice in the workplace: Approaching fairness in human resource management*.
- GSDC. (2024). *Next-gen AI in action: Unilever's AI-powered recruitment revolution*. Global Skill Development Council. <https://www.gsdcouncil.org/blogs/next-gen-ai-in-action-unilever-s-ai-powered-recruitment-revolution>
- Hassan, S., Mahsud, R., Yukl, G., & Prussia, G. E. (2013). Ethical and empowering leadership and leader effectiveness. *Journal of Managerial Psychology*, 28(2), 133-146.
- Hauser, D. J., Ellsworth, P. C., & Gonzalez, R. (2018). Are manipulation checks necessary?. *Frontiers in psychology*, 9, 998.
- Harvard Kennedy School. (2026). <https://www.hks.harvard.edu/educational-programs/executive-education/leading-artificial-intelligence>
- Hayes, A. F. (2017). Introduction to mediation, moderation, and conditional process analysis: A regression-based approach. Guilford publications.

- Höddinghaus, M., Sondern, D., & Hertel, G. (2021). The automation of leadership functions: Would people trust decision algorithms?. *Computers in Human Behavior*, 116, 106635.
- Jantzen, L., Bottel, M., Kempen, R. (2024). How to achieve trust, satisfaction, and acceptance in the interaction with AI through an AI Cockpit and a Transparency Interface?. *Procedia Computer Science*, 246, 292-31.
- Jarrahi, M. H. (2018). Artificial intelligence and the future of work: Human-AI symbiosis in organizational decision making. *Business horizons*, 61(4), 577-586.
- Javalagi, A. A., Newman, D. A., & Li, M. (2024). Personality and leadership: Meta-analytic review of cross-cultural moderation, behavioral mediation, and honesty-humility. *Journal of Applied Psychology*, 109(9), 1489.
- Jia, Q., & Zhang, Y. (2024). Artificial intelligence agents as team leaders: A study of the impact on team climate and team effectiveness. In Li, E.Y. et al. (Eds.) *Proceedings of The International Conference on Electronic Business, Volume 23* (pp. 44-53). ICEB'24, Zhuhai, China, October 24-28, 2024
- John, O. P., & Srivastava, S. (1999). The Big-Five trait taxonomy: History, measurement, and theoretical perspectives.
- Judge, T. A., Piccolo, R. F., & Ilies, R. (2004). The forgotten ones? The validity of consideration and initiating structure in leadership research. *Journal of applied psychology*, 89(1), 36.
- Kellett, J. B., Humphrey, R. H., & Sleeth, R. G. (2006). Empathy and the emergence of task and relations leaders. *The Leadership Quarterly*, 17(2), 146-162.
- Kim, H., & Yukl, G. (1995). Relationships of managerial effectiveness and advancement to self-reported and subordinate-reported leadership behaviors from the multiple-linkage mode. *The leadership quarterly*, 6(3), 361-377.

- Kim, Y. J., & Cribbie, R. A. (2018). ANOVA and the variance homogeneity assumption: Exploring a better gatekeeper. *British Journal of Mathematical and Statistical Psychology*, 71(1), 1-12.
- Krumrei-Mancuso, E. J., & Rowatt, W. C. (2023). Humility in novice leaders: links to servant leadership and followers' satisfaction with leadership. *The Journal of Positive Psychology*, 18(1), 154-166.
- Kurup, S., & Gupta, V. (2022). Factors influencing the AI adoption in organizations. *Metamorphosis*, 21(2), 129-139.
- Lai, V., Chen, C., Liao, Q. V., Smith-Renner, A., & Tan, C. (2021). Towards a science of human-ai decision making: a survey of empirical studies. *arXiv preprint arXiv:2112.11471*.
- Lai, X., & Rau, P. L. P. (2025). AI as co-leader: Effects of human consideration and AI structure behaviors on leadership effectiveness and neural activation. *International Journal of Human-Computer Interaction*, 41(9), 5694-5712.
- Lambert, L. S., Tepper, B. J., Carr, J. C., Holt, D. T., & Barelka, A. J. (2012). Forgotten but not gone: An examination of fit between leader consideration and initiating structure needed and received. *Journal of Applied Psychology*, 97(5), 913.
- Larsson, S., & Heintz, F. (2020). Transparency in artificial intelligence. *Internet policy review*, 9(2), 1-16.
- Legood, A., van der Werff, L., Lee, A., & Den Hartog, D. (2021). A meta-analysis of the role of trust in the leadership-performance relationship. *European Journal of Work and Organizational Psychology*, 30(1), 1-22.

- Leung, K., Su, S. K., & Morris, M. W. (1998). *When is Criticism Not Constructive: The Roles of Fairness Perceptions and Attributions in Employee Rejection of Critical Supervisory Feedback* (No. 1478). Graduate School of Business, Stanford University.
- Li, M., & Bitterly, T. B. (2024). How perceived lack of benevolence harms trust of artificial intelligence management. *Journal of Applied Psychology*.
- Liden, R. C., Wang, X., & Wang, Y. (2025). The evolution of leadership: Past insights, present trends, and future directions. *Journal of Business Research*, 186, 115036.
- Lochner, E., Schmoll, R., & Kaiser, S. (2025). Less Human, Less Positive? How AI Involvement in Leadership Shapes Employees' Affective Well-Being Across Different Supervisor Decisions. *Computers in Human Behavior: Artificial Humans*, 100239.
- Logg, J. M., Minson, J. A., & Moore, D. A. (2019). Algorithm appreciation: People prefer algorithmic to human judgment. *Organizational behavior and human decision processes*, 151, 90-103.
- Lord, R. G., Foti, R. J., & De Vader, C. L. (1984). A test of leadership categorization theory: Internal structure, information processing, and leadership perceptions. *Organizational behavior and human performance*, 34(3), 343-378.
- Lord, R. G., & Maher, K. J. (2002). *Leadership and information processing: Linking perceptions and performance*. Routledge.
- Madlock, P. E. (2008). The link between leadership style, communicator competence, and employee satisfaction. *The Journal of Business Communication* (1973), 45(1), 61-78.
- Martinez-Miranda, J., & Aldea, A. (2005). Emotions in human and artificial intelligence. *Computers in human behavior*, 21(2), 323-341.

- Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of management review*, 20(3), 709-734.
- McAllister, D. J. (1995). Affect-and cognition-based trust as foundations for interpersonal cooperation in organizations. *Academy of management journal*, 38(1), 24-59.
- MIT Sloan Executive Education. (2026). <https://executive.mit.edu/artificial-intelligence>
- Munshi, J. D. (2025). Empirical insights into ai-augmented leadership: A multi-industry comparative study. *International Journal of Engineering and Information Management*, 1(1), 42-59.
- Norman, S. M., Avolio, B. J., & Luthans, F. (2010). The impact of positivity and transparency on trust in leaders and their perceived effectiveness. *The leadership quarterly*, 21(3), 350-364.
- Oc, B., & Bashshur, M. R. (2013). Followership, leadership and social influence. *The Leadership Quarterly*, 24(6), 919-934.
- Offermann, L. R., Kennedy Jr, J. K., & Wirtz, P. W. (1994). Implicit leadership theories: Content, structure, and generalizability. *The leadership quarterly*, 5(1), 43-58.
- Offermann, L. R., & Coats, M. R. (2018). Implicit theories of leadership: Stability and change over two decades. *The Leadership Quarterly*, 29(4), 513-522.
- Okoidigun, O. G., & Emagbetere, E. (2025). Trusting artificial intelligence: A qualitative exploration of public perception and acceptance of the risks and benefits. *Open J Soc Sci*, 13(3), 64-80.
- Park, K., & Yoon, H. Y. (2024). Beyond the code: The impact of AI algorithm transparency signaling on user trust and relational satisfaction. *Public Relations Review*, 50(5), 102507.

- Peifer, Y., Jeske, T., & Hille, S. (2022). Artificial intelligence and its impact on leaders and leadership. *Procedia computer science*, 200, 1024-1030.
- Peterson, R. A. (2001). On the use of college students in social science research: Insights from a second-order meta-analysis. *Journal of consumer research*, 28(3), 450-461.
- Podsakoff, P. M., MacKenzie, S. B., Lee, J. Y., & Podsakoff, N. P. (2003). Common method biases in behavioral research: a critical review of the literature and recommended remedies. *Journal of applied psychology*, 88(5), 879.
- Rodrigues, R., Teixeira, N., & Costa, B. (2024). The Impact of Perceived Leadership Effectiveness and Emotional Intelligence on Employee Satisfaction in the Workplace. *Merits*, 4(4), 490-501.
- Reeves, B. R. (2023). *The Impact of Transformational Leadership on Follower Attitudes Toward Artificial Intelligence (AI) in the United States Air Force*. Indiana Wesleyan University.
- Sanders, N. R. (2025). THREE PRINCIPLES FOR LEADERS IN THE AGE OF AI. *Leader to Leader*.
- Scarpello, V., & Vandenberg, R. J. (1987). The satisfaction with my supervisor scale: Its utility for research and practical applications. *Journal of management*, 13(3), 447-466.
- Schelenz, L., Segal, A., Adelio, O., & Gal, K. (2024). Transparency-check: An instrument for the study and design of transparency in ai-based personalization systems. *ACM Journal on Responsible Computing*, 1(1), 1-18.
- Schnackenberg, A. K., & Tomlinson, E. C. (2016). Organizational transparency: A new perspective on managing trust in organization-stakeholder relationships. *Journal of management*, 42(7), 1784-1810.

- Shahzad, K., Raja, U., & Hashmi, S. D. (2021). Impact of Big Five personality traits on authentic leadership. *Leadership & Organization Development Journal*, 42(2), 208-218.
- Shin, D. (2022). How do people judge the credibility of algorithmic sources?. *Ai & Society*, 37(1), 81-96.
- Smith, A. M., & Green, M. (2018). Artificial intelligence and the role of leadership. *Journal of Leadership Studies*, 12(3), 85-87.
- Spitzberg, B. H. (1983). Communication competence as knowledge, skill, and impression. *Communication Education*, 32(3), 323-329.
- Stöber, J. (2001). The Social Desirability Scale-17 (SDS-17): Convergent validity, discriminant validity, and relationship with age. *European Journal of Psychological Assessment*, 17(3), 222.
- The Coca-Cola Company. (2023). *Coca-Cola invites digital artists to 'Create Real Magic' using new AI platform*. Media Center. Retrieved from <https://www.coca-colacompany.com/media-center/coca-cola-invites-digital-artists-to-create-real-magic-using-new-ai-platform>
- Uhl-Bien, M. (2006). Relational leadership theory: Exploring the social processes of leadership and organizing. *The leadership quarterly*, 17(6), 654-676.
- Venkatesh, V., Thong, J. Y., & Xu, X. (2012). Consumer acceptance and use of information technology: Extending the Unified Theory of Acceptance and Use of Technology1. *MIS quarterly*, 36(1), 157-178.
- Vogelgesang, G. R., & Lester, P. B. (2009). Transparency: How Leaders Can Get Results by Laying it on the Line. *Organizational Dynamics*, 38(4), 252-260.

- Vroom, V., & Yetton, P. (1973). Leadership and Decision-Making. *University of Pittsburgh Press*.
- Walter, F., Cole, M. S., & Humphrey, R. H. (2011). Emotional intelligence: Sine qua non of leadership or folderol?. *Academy of management perspectives*, 25(1), 45-59.
- Waltl, B., & Vogl, R. (2018). Increasing transparency in algorithmic-decision-making with explainable AI. *Datenschutz und Datensicherheit-DuD*, 42(10), 613-617.
- Wang, B., Rau, P. L. P., & Yuan, T. (2023). Measuring user competence in using artificial intelligence: validity and reliability of artificial intelligence literacy scale. *Behaviour & information technology*, 42(9), 1324-1337.
- Watson, D., Clark, L. A., & Tellegen, A. (1988). Development and validation of brief measures of positive and negative affect: the PANAS scales. *Journal of personality and social psychology*, 54(6), 1063.
- Yu, L., & Li, Y. (2022). Artificial intelligence decision-making transparency and employees' trust: The parallel multiple mediating effect of effectiveness and discomfort. *Behavioral Sciences*, 12(5), 127.
- Yukl, G. (2002) Leadership in Organizations. 5. Prentice Hall.
- Yukl, G. (2012). Effective leadership behavior: What we know and what questions need more attention. *Academy of Management perspectives*, 26(4), 66-85.
- Zaccaro, S. J., Green, J. P., Dubrow, S., & Kolze, M. (2018). Leader individual differences, situational parameters, and leadership outcomes: A comprehensive review and integration. *The Leadership Quarterly*, 29(1), 2-43.

Table 1*Study 1 - Means, Standard Deviations, Correlations, and Reliabilities for Key Study Variables*

Variable	<i>M</i>	<i>SD</i>	1	2	3	4	5	6	7	8	9
1. AI Literacy	4.82	0.72	(.80)								
2. Propensity to Trust – F2	4.02	0.84	-.06	(.77)							
3. Propensity to Trust – F3	5.04	0.88	.28**	.19**	(.73)						
4. Extraversion	3.22	0.76	-.02	.10	-.16**	(.85)					
5. Agreeableness	3.74	0.61	.32**	.04	.20**	.22**	(.78)				
6. Trust	4.66	1.09	.22**	.31**	.23**	.23**	.26**	(.93)			
7. Satisfaction	4.76	1.41	.18**	.30**	.22**	.21**	.26**	.87**	(.93)		
8. Cons. Effectiveness	4.71	1.15	.16**	.36**	.29**	.21**	.22**	.82**	.71**	(.88)	
9. IS Effectiveness	4.76	1.07	.20**	.30**	.27**	.01	.17**	.60**	.49**	.62	(.86)

Note. $N = 281$. Propensity to Trust - F2 = reliability and integrity of others; Propensity to Trust - F3 = risk aversion; Cons. Effectiveness = consideration-based effectiveness; IS Effectiveness = initiating structure-based effectiveness. Cronbach's α coefficients appear on the diagonal. Control variables include, AI literacy, propensity to trust (F2 and F3), Extraversion, and Agreeableness.

* $p < .05$. ** $p < .01$.

Table 2*Study 1 - Univariate Analyses of Covariance for Follower Trust and Satisfaction*

	Trust in the Leader				Satisfaction with the Leader			
	<i>F</i>	<i>df</i>	<i>p</i>	η^2	<i>F</i>	<i>df</i>	<i>p</i>	η^2
Corrected Model	18.94	9	.00***	.39	15.04	10	.00***	.34
Intercept	0.66	1	.61	.00	0.00	1	.96	.00
Covariates								
AI Literacy	8.30	1	.00**	.04	-	-	-	-
Propensity to Trust – F2	15.08	1	.00***	.07	15.12	1	.00***	.05
Propensity to Trust – F3	5.03	1	.00**	.02	6.43	1	.01*	.02
Extraversion	11.32	1	.00***	.05	6.95	1	.01**	.03
Agreeableness	-	-	-	-	9.11	1	.00**	.03
Between-subjects Effects								
Leader Type	13.40	2	.00***	.09	8.83	2	.00***	.06
Transparency	49.16	1	.00***	.16	39.78	1	.00***	.13
Leader Type * Transparency	0.45	2	.64	.00	1.91	2	.15	.01

Note. $N = 281$. η^2 = effect size (partial eta squared). Only significant covariates were retained in the final model. Propensity to Trust - F2 = reliability and integrity of others; Propensity to Trust - F3 = risk aversion.

* $p < .05$. ** $p < .01$. *** $p < .001$.

Table 3*Study 1 – Repeated Measures Analysis of Covariance for Leader Effectiveness*

	Leader Effectiveness			
	<i>F</i>	<i>df</i>	<i>p</i>	ηp^2
Between-subjects Effects				
Intercept	3.35	1	.07	.01
Covariates				
Agreeableness	4.44	1	.04*	.02
Propensity to Trust – F2	31.86	1	.00***	.10
Propensity to Trust – F3	9.33	1	.00**	.03
AI Literacy	4.36	1	.04*	.02
Leader Type	7.00	2	.00**	.05
Transparency	18.05	1	.00***	.06
Leader Type * Transparency	4.37	2	.01*	.03
Within-subjects Effects				
Situational Context	0.16	1	.69	.00
Situational Context * Leader Type	9.31	2	.00***	.07
Situational Context * Transparency	4.22	1	.04*	.02
Leader Type * Transparency * Situational Context	0.24	2	.76	.00

Note. $N = 281$. ηp^2 = effect size (partial eta squared). Only significant covariates were retained in the final model. Propensity to Trust - F2 = reliability and integrity of others; Propensity to Trust - F3 = risk aversion

* $p < .05$. ** $p < .01$. *** $p < .001$.

Table 4*Study 2 - Means, Standard Deviations, Correlations, and Reliabilities for Key Study Variables*

Variable	<i>M</i>	<i>SD</i>	1	2	3	4	5	6	7
1. Propensity to Trust – F2	3.99	0.82	(.73)						
2. Functionality of AI – F1	2.75	0.77	.34*	(.80)					
3. Functionality of AI – F5	3.34	0.72	.21**	.43**	(.84)				
4. Trust	4.13	1.32	.21**	.19**	.20**	(.94)			
5. Satisfaction	4.20	1.69	.16**	.18**	.16**	.91**	(.94)		
6. Cons. Effectiveness	4.16	1.26	.22**	.19**	.13**	.72**	.67**	(.81)	
7. IS Effectiveness	4.42	1.23	.20**	.16**	.18**	.60**	.55**	.67**	(.81)

Note. $N = 420$. Propensity to Trust - F2 = reliability and integrity of others; Functionality of AI - F1 = trust in the AI; Functionality of AI - F5 = functionality trust in smart technologies; Cons. Effectiveness = consideration-based effectiveness; IS Effectiveness = initiating structure-based effectiveness. Cronbach's α coefficients appear on the diagonal. Control variables include propensity to trust (F2) and functionality of AI (F1 and F5).

* $p < .05$. ** $p < .01$.

Table 5*Study 2 - Means, Standard Deviations, and Correlations for ILT Dimensions*

Variable	<i>M</i>	<i>SD</i>	1	2	3	4	5	6	7	8	9
1. Sensitivity	5.66	2.56	-								
2. Dedication	7.04	2.57	.66*	-							
3. Tyranny	4.79	2.06	-.18**	-.14**	-						
4. Charisma	5.75	2.42	.80**	.64**	.05	-					
5. Strength	6.06	2.11	.44**	.56**	.43**	.61**	-				
6. Creativity	6.10	2.63	.71**	.82**	-.14**	.71**	.52**	-			
7. Well-Groomed	5.18	3.07	.65**	.44**	.04	.64**	.47**	.49**	-		
8. Masculinity	4.00	2.50	.40**	.14**	.19**	.41**	.32**	.19**	.62**	-	
9. Intelligence	7.20	2.71	.47**	.77**	-.11*	.51**	.47**	.75**	.33**	.09	-

Note. $N = 420$.* $p < .05$. ** $p < .01$.

Table 6*Study 2 - Univariate Analyses of Covariance for Follower Trust and Satisfaction*

	Trust in the Leader				Satisfaction with the Leader			
	<i>F</i>	<i>df</i>	<i>p</i>	η^2	<i>F</i>	<i>df</i>	<i>p</i>	η^2
Corrected Model	48.66	7	.00***	.45	44.25	7	.00***	.43
Intercept	71.19	1	.00***	.15	73.30	1	.00***	.15
Covariates								
Propensity to Trust – F2	15.21	1	.00***	.04	4.36	1	.04*	.01
Functionality of AI – F1	-	-	-	-	6.62	1	.01*	.02
Functionality of AI – F5	7.56	1	.01**	.02	-	-	-	-
Between-subjects Effects								
Leader Type	2.75	2	.07	.01	3.68	2	.03*	.02
Transparency	267.22	1	.00***	.40	256.55	1	.00***	.39
Leader Type * Transparency	9.27	2	.00***	.04	8.62	2	.00***	.04

Note. $N = 420$. η^2 = effect size (partial eta squared). Only significant covariates were retained in the final model. Propensity to Trust - F2 = reliability and integrity of others; Functionality of AI - F1 = trust in the AI; Functionality of AI - F5 = functionality trust in smart technologies.

* $p < .05$. ** $p < .01$. *** $p < .001$.

Table 7*Study 2 – Repeated Measures Analysis of Covariance for Leader Effectiveness*

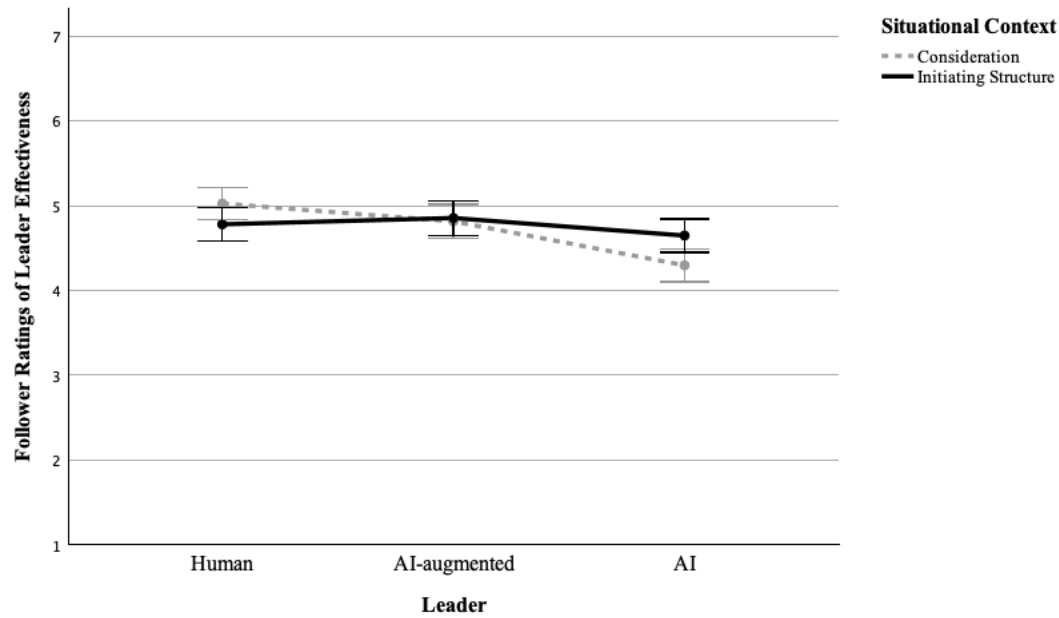
	Leader Effectiveness			
	<i>F</i>	<i>df</i>	<i>p</i>	ηp^2
Between-subjects Effects				
Intercept	135.79	1	.00***	.25
Covariates				
Propensity to Trust – F2	14.25	1	.00***	.03
Functionality of AI – F1	5.04	1	.03*	.01
Leader Type	7.09	2	.00***	.03
Transparency	105.80	1	.00***	.21
Leader Type * Transparency	8.06	2	.00***	.04
Within-subjects Effects				
Situational Context	4.09	1	.04*	.01
Situational Context * Leader Type	12.61	2	.00***	.06
Situational Context * Transparency	0.34	1	.56	.00
Leader Type * Transparency * Situational Context	0.74	2	.48	.00

Note. $N = 420$. ηp^2 = effect size (partial eta squared). Only significant covariates were retained in the final model. Propensity to Trust – F2 = reliability and integrity of others; Functionality of AI – F1 = trust in the AI

* $p < .05$. ** $p < .01$. *** $p < .001$.

Figure 1

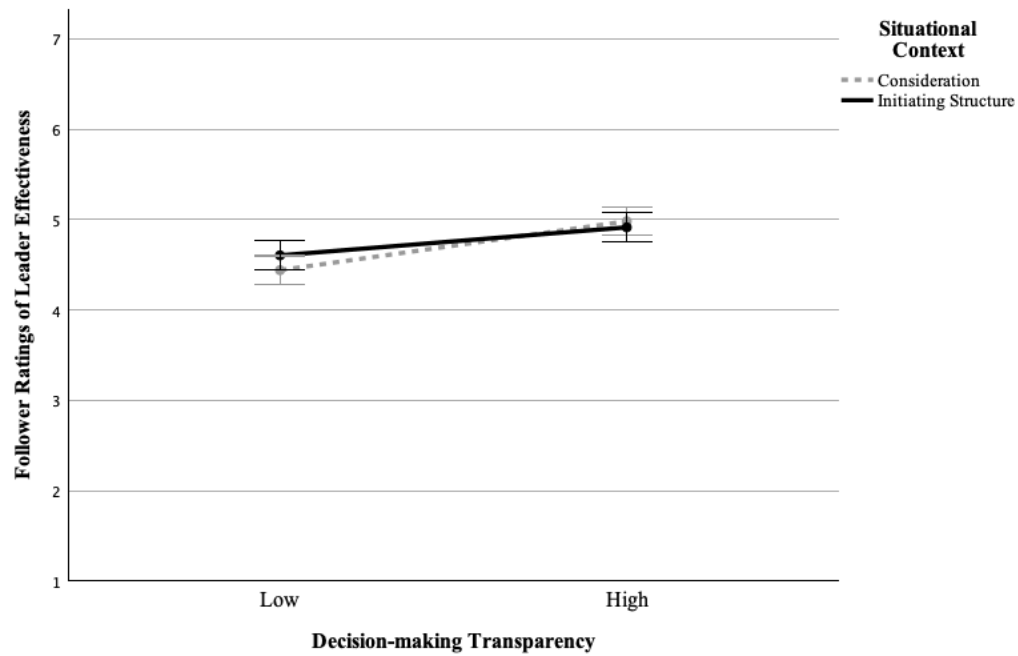
Study 1 Interaction Between Type of Leader and Situational Context on Leader Effectiveness



Note. $N = 281$. The x-axis represents leader type and the y-axis represents follower ratings of leader effectiveness. The dotted line represents the consideration context and the solid line represents the initiating structure context. Error bars represent 95% confidence intervals.

Figure 2

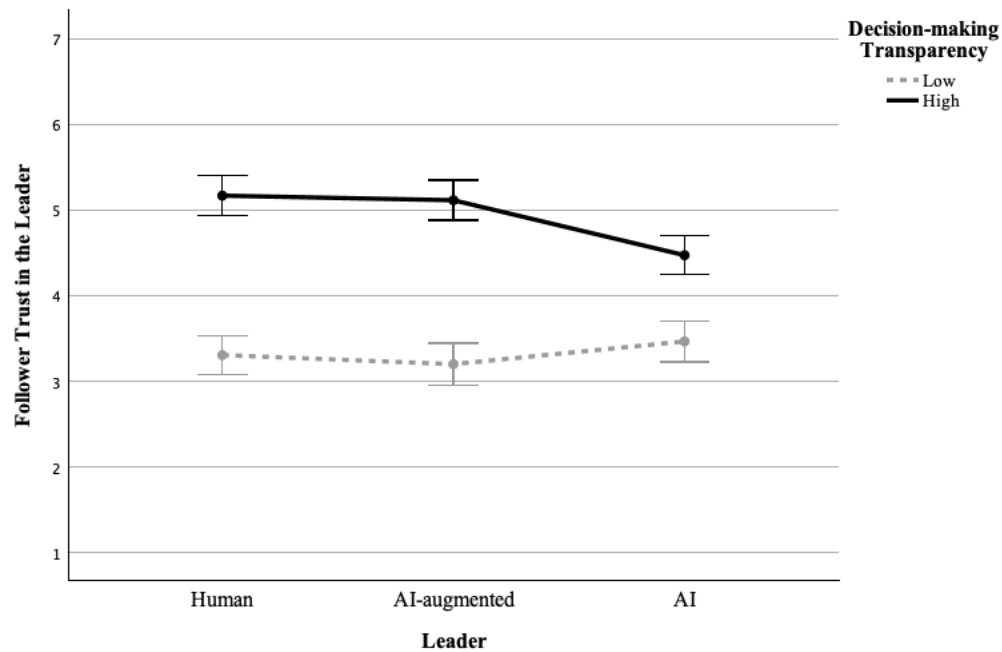
Study 1 Interaction Between Decision-making Transparency and Situational Context on Leader Effectiveness



Note. $N = 281$. The x-axis represents decision-making transparency (low vs. high) and the y-axis represents follower ratings of leader effectiveness. The dotted line represents the consideration context and the solid line represents the initiating structure context. Error bars represent 95% confidence intervals.

Figure 3

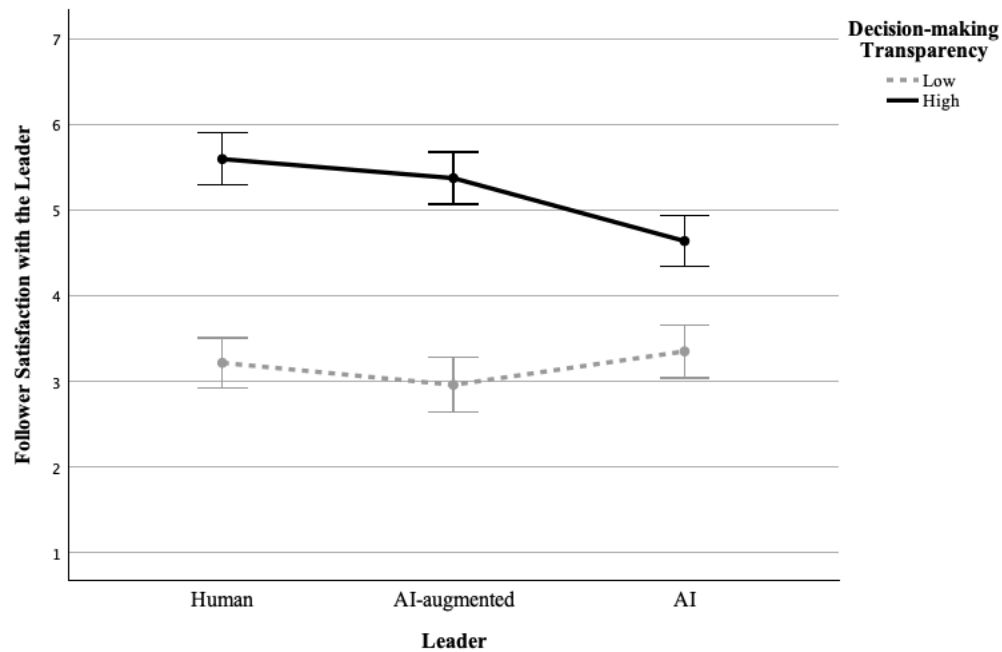
Study 2 Interaction Between Type of Leader and Decision-making Transparency on Trust



Note. $N = 420$. The x-axis represents leader type and the y-axis represents follower trust in the leader. The dotted line represents the low decision-making transparency condition and the solid line represents the high decision-making transparency condition. Error bars represent 95% confidence intervals.

Figure 4

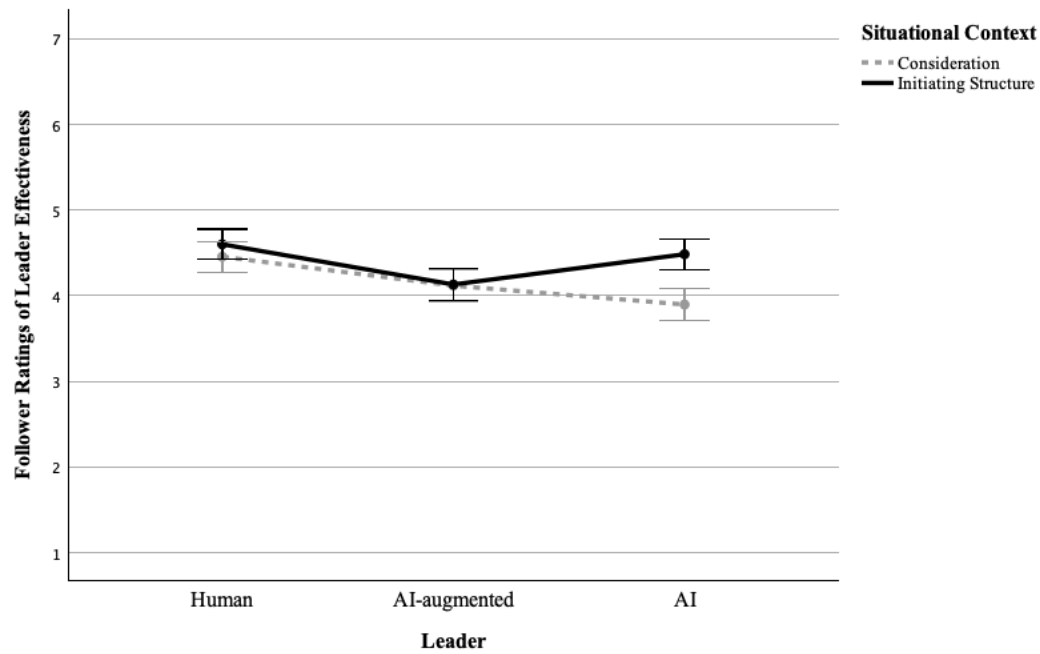
Study 2 Interaction Between Type of Leader and Decision-making Transparency on Satisfaction



Note. $N = 420$. The x-axis represents leader type and the y-axis represents follower satisfaction with the leader. The dotted line represents the low decision-making transparency condition and the solid line represents the high decision-making transparency condition. Error bars represent 95% confidence intervals.

Figure 5

Study 2 Interaction Between Type of Leader and Situational Context on Leader Effectiveness



Note. $N = 420$. The x-axis represents leader type and the y-axis represents follower ratings of leader effectiveness. The dotted line represents the consideration context and the solid line represents the initiating structure context. Error bars represent 95% confidence intervals.

Appendix A

Organization Background

Organizational Description

GemB Marketing Solutions is a *mid-sized marketing and sales consultancy* specializing in brand strategy, consumer engagement, and market expansion. The company partners with food, beverage, and lifestyle brands to help them grow through strategic campaigns and sales execution.

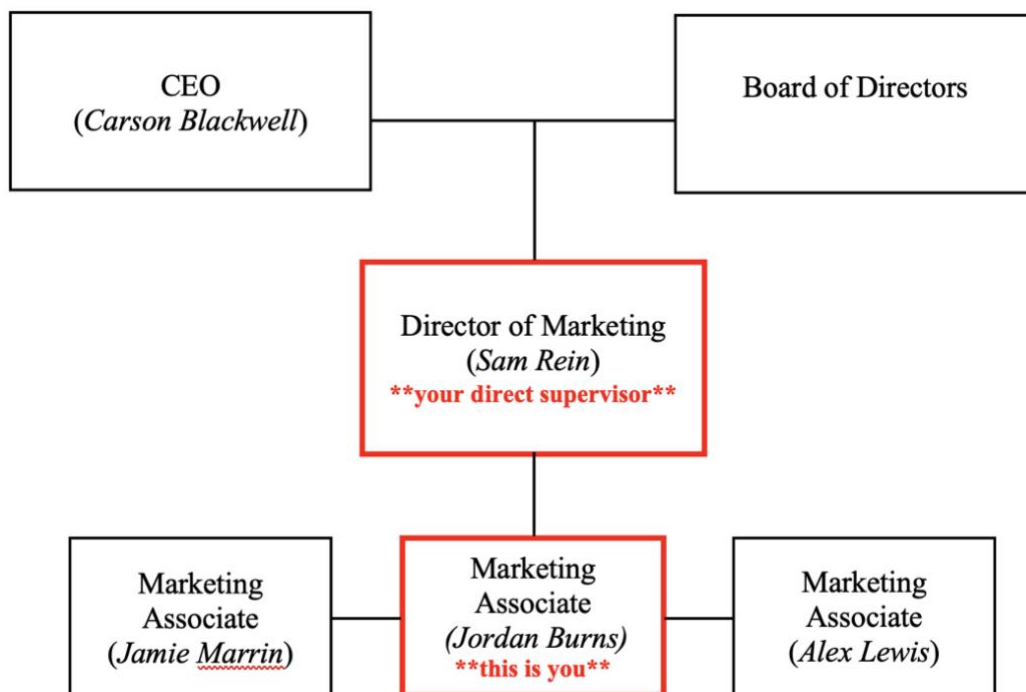
Currently, *GemB is working with a client in the beverage industry to develop and launch PureFusion*, a premium beverage line targeting health-conscious consumers and active professionals.

As an associate within the Marketing Department, **you play a key role in market research initiatives**, ensuring successful product positioning, sales execution, and consumer engagement. You report to the Marketing Director, who oversees marketing strategy, sales planning, and product rollout for PureFusion.

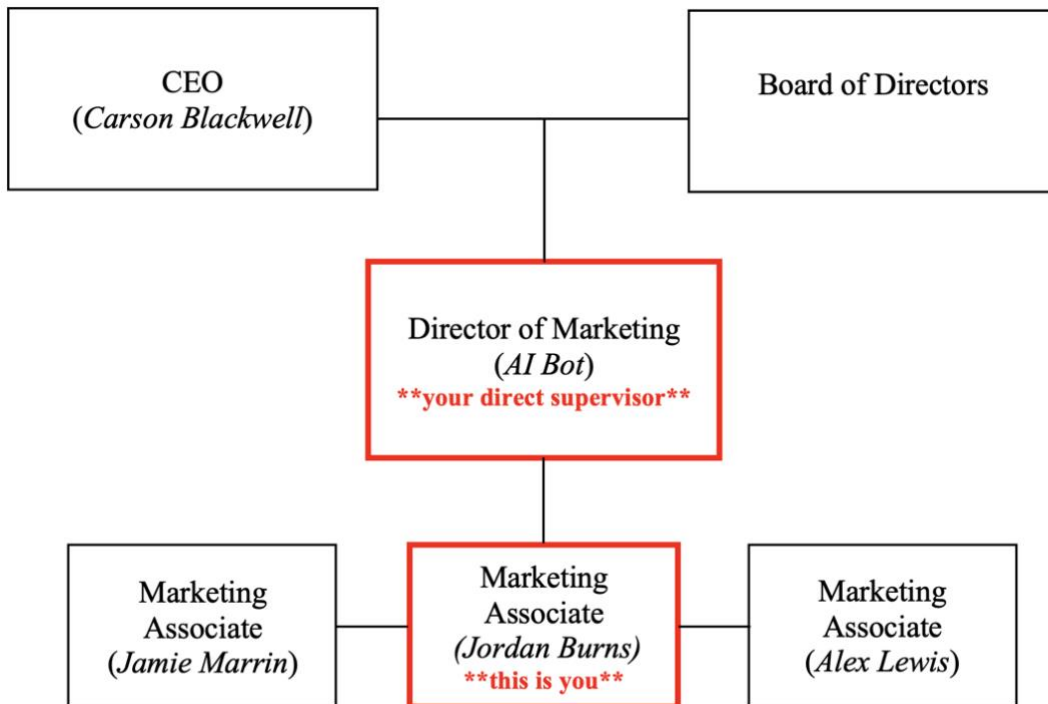
Below is a chart detailing GemB's organizational structure, indicating in red your (Jordan's) role and the role of your immediate supervisor ([INSERT LEADER HERE]):

Organizational Chart

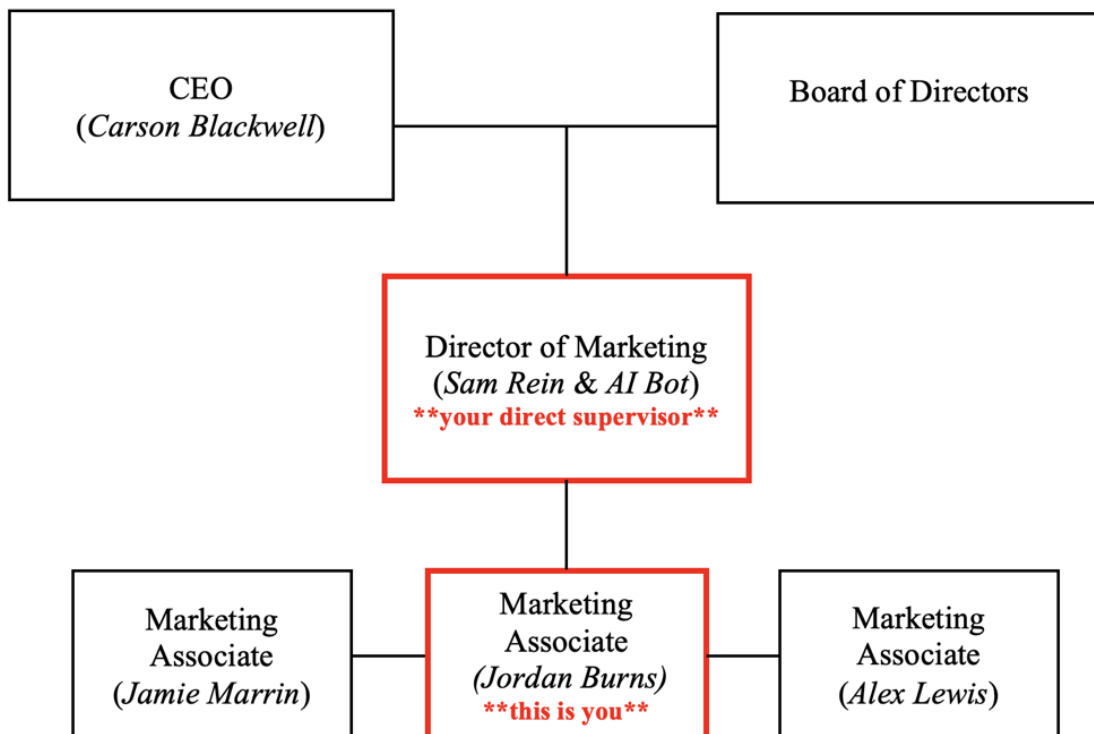
Human Leader



AI Leader



AI-augmented Leader



Appendix B

Study 1 Full Leader Descriptions

Human Leader

Sam is the Marketing Department Director and serves as your immediate supervisor. Sam is a **highly engaged and people-focused** leader who **values collaboration and team input**. With experience managing marketing and sales teams, Sam is known for **building strong connections with employees and fostering a positive team dynamic**.

In team meetings, Sam **encourages employees to share concerns** and ensures work is executed efficiently. Your colleagues describe Sam as **attentive and supportive, appreciating the opportunity to engage with you and your team**.

AI Leader

An AI software, referred to as Bot, is the Marketing Department Director and serves as your immediate supervisor. Bot is a **structured and data-driven leader** programmed to **optimize marketing and sales performance**. It is known for **using analytics to guide strategy, providing direction on tasks and strategic priorities**.

In team meetings, **Bot outlines what needs to be done and ensures that work is executed efficiently**. Your colleagues describe Bot as **focused and efficient, providing clear expectations to you and your team**.

AI-augmented Leader

Sam is the Marketing Department Director and serves as your immediate supervisor. They **incorporate AI-driven insights into their decision-making** using Bot, a structured and data-driven AI machine. Sam is known for **providing direction and fostering employee relationships through building strong connections with employees and using Bot's analytics to guide strategy**.

In team meetings, **Sam presents strategic priorities, often utilizing insights from Bot to support decision-making**. Your colleagues describe Sam as efficient, **providing expectations and engagement opportunities**.

Appendix C

Study 1 Colleague Emails - Transparency Manipulation

Human Leader & High Transparency

From:	Alex L. <alewis@gembmarketing.org>
To:	Jordan B. <jburns@gembmarketing.org>
Sent:	3/10/2025 10:25am
Subject:	PureFusion Product Launch – Check In

Hey Jordan,

I just wanted to check in—how are things going on your end? It was great to see you in the project meeting today. I really appreciated how Sam **took the time to explain the reasoning behind recent decisions** and made sure we understand how our work fits into the bigger picture. It makes things feel a lot clearer and more understandable.

Whenever priorities shift, Sam is **transparent about the factors driving those changes and encourage us to ask questions**, which makes it easier to stay aligned and feel confident in what we're doing. Plus, it's been great to have **open discussions about strategy and expectations** rather than just executing tasks without knowing the full context.

Let's catch up later and compare notes on the project—I'd love to hear your thoughts!

Best,
Alex

Alex Lewis
Associate, Marketing Team
GemB Marketing Solutions

Human Leader & Low Transparency

From:	Alex L. <alewis@gembmarketing.org>
To:	Jordan B. <jburns@gembmarketing.org>
Sent:	3/10/2025 10:25am
Subject:	PureFusion Product Launch – Check In

Hey Jordan,

I just wanted to check in—how are things going on your end? It was great to see you in the project meeting today. I noticed that Sam **didn't provide much explanation behind recent decisions**, which makes it harder to understand how our work connects to the bigger picture. It's been tough to keep track of shifting priorities when there's not a lot of insight into why things are changing.

Most of the time, we just roll with new directions, but **there's little detail on what's influencing decisions, and asking questions doesn't seem encouraged**. It would definitely help if we had **more clarity on expectations and strategy**, so we're not just executing tasks without the full picture.

Let's catch up later and compare notes on the project—I'd love to hear your thoughts!

Best,
Alex

Alex Lewis
Associate, Marketing Team
GemB Marketing Solutions

AI Leader & High Transparency

From: Alex L. <alewis@gembmarketing.org>
To: Jordan B. <jburns@gembmarketing.org>
Sent: 3/10/2025 10:25am
Subject: PureFusion Product Launch – Check In

Hey Jordan,

I just wanted to check in—how are things going on your end? It was great to see you in the project meeting today. I really appreciated how Bot **took the time to explain the reasoning behind recent decisions** and made sure we understand how our work fits into the bigger picture. It makes things feel a lot clearer and more understandable.

Whenever priorities shift, Bot is **transparent about the factors driving those changes and encourage us to ask questions**, which makes it easier to stay aligned and feel confident in what we're doing. Plus, it's been great to have **open discussions about strategy and expectations** rather than just executing tasks without knowing the full context.

Let's catch up later and compare notes on the project—I'd love to hear your thoughts!

Best,
Alex

Alex Lewis
Associate, Marketing Team
GemB Marketing Solutions

AI Leader & Low Transparency

From: Alex L. <alewis@gembmarketing.org>
To: Jordan B. <jburns@gembmarketing.org>
Sent: 3/10/2025 10:25am
Subject: PureFusion Product Launch – Check In

Hey Jordan,

I just wanted to check in—how are things going on your end? It was great to see you in the project meeting today. I noticed that Bot **didn't provide much explanation behind recent decisions**, which makes it harder to understand how our work connects to the bigger picture. It's been tough to keep track of shifting priorities when there's not a lot of insight into why things are changing.

Most of the time, we just roll with new directions, but **there's little detail on what's influencing decisions, and asking questions doesn't seem encouraged**. It would definitely help if we had **more clarity on expectations and strategy**, so we're not just executing tasks without the full picture.

Let's catch up later and compare notes on the project—I'd love to hear your thoughts!

Best,
Alex

Alex Lewis
Associate, Marketing Team
GemB Marketing Solutions

AI-augmented Leader & High Transparency

From:	Alex L. <alewis@gembmarketing.org>
To:	Jordan B. <jburns@gembmarketing.org>
Sent:	3/10/2025 10:25am
Subject:	PureFusion Product Launch – Check In

Hey Jordan,

I just wanted to check in—how are things going on your end? It was great to see you in the project meeting today. I really appreciated how Sam **took the time to explain the reasoning behind recent decisions** and made sure we understand Bot's insights and how our work fits into the bigger picture. It makes things feel a lot clearer and more understandable.

Whenever priorities shift, Sam is **transparent about the factors Bot indicates are driving those changes and encourage us to ask questions**, which makes it easier to stay aligned and feel confident in what we're doing. Plus, it's been great to have **open discussions about Bot's strategy and expectations** rather than just executing tasks without knowing the full context.

Let's catch up later and compare notes on the project—I'd love to hear your thoughts!

Best,
Alex

Alex Lewis
Associate, Marketing Team
GemB Marketing Solutions

AI-augmented Leader & Low Transparency

From:	Alex L. <alewis@gembmarketing.org>
To:	Jordan B. <jburns@gembmarketing.org>
Sent:	3/10/2025 10:25am
Subject:	PureFusion Product Launch – Check In

Hey Jordan,

I just wanted to check in—how are things going on your end? It was great to see you in the project meeting today. I noticed that Sam **didn't provide much explanation behind recent decisions and Bot's output**, which makes it harder to understand how our work connects to the bigger picture. It's been tough to keep track of shifting priorities when there's not a lot of insight into why things are changing.

Most of the time, we just roll with new directions and whatever Bot says, but **there's little detail on what's influencing decisions, and asking questions doesn't seem encouraged**. It would definitely help if Sam provided **more clarity on Bot's expectations and strategy**, so we're not just executing tasks without the full picture.

Let's catch up later and compare notes on the project—I'd love to hear your thoughts!

Best,
Alex

Alex Lewis
Associate, Marketing Team
GemB Marketing Solutions

Appendix D

Study 1 Leader Emails - Leader Type and Transparency Manipulation

Human Leader & High Transparency

From:	Sam R. <srein@gembmarketing.org>
To:	Jordan B. <jburns@gembmarketing.org>
Sent:	3/11/2025 9:21am
Subject:	PureFusion Launch – Please Read

Hi Jordan,

As we move forward with the PureFusion launch, I'd like you to take the lead on finalizing the market research report. Your perspective is valuable, and pulling together the latest trends will help us shape our marketing strategy. Let me know if you need anything while working on it—I appreciate your efforts.

These next steps are based on a combination of past market trends, real-time consumer engagement data, and team discussions. By reviewing these insights, we're ensuring that our approach to the PureFusion launch is backed by both strategy and data.

Looking forward to seeing what you put together!

Regards,
Sam

Sam Rein
Director, Marketing Team
GemB Marketing Solutions

Human Leader & Low Transparency

From:	Sam R. <srein@gembmarketing.org>
To:	Jordan B. <jburns@gembmarketing.org>
Sent:	3/11/2025 9:21am
Subject:	PureFusion Launch – Please Read

Hi Jordan,

As we move forward with the PureFusion launch, I'd like you to take the lead on finalizing the market research report. Your perspective is valuable, and pulling together the latest trends will help us shape our marketing strategy. Let me know if you need anything while working on it—I appreciate your efforts.

These next steps are set, and we will move forward as planned with the launch of PureFusion. These directions have been determined based on various factors, and staying on track with them will ensure progress.

Looking forward to seeing what you put together!

Regards,
Sam

Sam Rein
Director, Marketing Team
GemB Marketing Solutions

AI Leader & High Transparency

From:	Bot <AIbot@gembmarketing.org>
To:	Jordan B. <jburns@gembmarketing.org>
Sent:	3/11/2025 9:21am
Subject:	PureFusion Launch – Please Read

Jordan,

For the next step in the PureFusion launch, you need to finalize the market research report. Use the latest data to summarize consumer engagement trends. This information will be used to adjust our marketing strategy. Follow the outlined format and ensure all key metrics are included.

These next steps are based on a combination of past market trends, real-time consumer engagement data, and team discussions. By reviewing these insights, we're ensuring that our approach to the PureFusion launch is backed by both strategy and data.

Proceed with the task as directed,

Bot

AI Bot

*Director, Marketing Team
GemB Marketing Solutions*

AI Leader & Low Transparency

From:	Bot <AIbot@gembmarketing.org>
To:	Jordan B. <jburns@gembmarketing.org>
Sent:	3/11/2025 9:21am
Subject:	PureFusion Launch – Please Read

Jordan,

For the next step in the PureFusion launch, you need to finalize the market research report. Use the latest data to summarize consumer engagement trends. This information will be used to adjust our marketing strategy. Follow the outlined format and ensure all key metrics are included.

These next steps are set, and we will move forward as planned with the launch of PureFusion. These directions have been determined based on various factors, and staying on track with them will ensure progress.

Proceed with the task as directed,

Bot

AI Bot

*Director, Marketing Team
GemB Marketing Solutions*

AI-augmented Leader & High Transparency

From:	Sam R. <srein@gembmarketing.org>
To:	Jordan B. <jburns@gembmarketing.org>
Sent:	3/11/2025 9:21am
Subject:	PureFusion Launch – Please Read

Hi Jordan,

For the next step in the PureFusion launch, I'd like you to finalize the market research report. Bot's latest data shows key consumer engagement trends, and your analysis will help us refine our strategy. Using both the data and your insights will make sure we're on the right track.

These next steps are based on a combination of past market trends, real-time consumer engagement data, and team discussions. By reviewing these insights, we're ensuring that our approach to the PureFusion launch is backed by both strategy and data.

Looking forward to seeing what you put together.

Regards,
Sam

Sam Rein
Director, Marketing Team
GemB Marketing Solutions

AI-augmented Leader & Low Transparency

From:	Sam R. <srein@gembmarketing.org>
To:	Jordan B. <jburns@gembmarketing.org>
Sent:	3/11/2025 9:21am
Subject:	PureFusion Launch – Please Read

Hi Jordan,

For the next step in the PureFusion launch, I'd like you to finalize the market research report. Bot's latest data shows key consumer engagement trends, and your analysis will help us refine our strategy. Using both the data and your insights will make sure we're on the right track.

These next steps are set, and we will move forward as planned with the launch of PureFusion. These directions have been determined based on various factors, and staying on track with them will ensure progress.

Looking forward to seeing what you put together.

Regards,
Sam

Sam Rein
Director, Marketing Team
GemB Marketing Solutions

Appendix E

Situational Context Manipulations Across All Studies

Consideration

Consider the following context:

Imagine that the GemB wants to launch PureFusion **in a way that supports and values all employees**. The focus is on **making sure that team members feel included, appreciated, and encouraged to contribute their ideas**. Supervisors should help employees develop their skills, recognize their hard work, and create a supportive work environment.

However, an issue has arisen within your team: *Jamie, one of your colleagues, has taken credit for an idea that Alex originally proposed for the product launch.*

This situation has created tension among team members, making it difficult for them to collaborate effectively. Your supervisor is needed to step in, address the conflict fairly, ensure all voices are valued, and rebuild trust within the team.

Initiating Structure

Consider the following context:

Imagine that the company wants to launch PureFusion by **making sure everything is well-organized, efficient, and clearly structured**. The focus is on **setting clear goals, defining roles, planning ahead, and making sure operations run smoothly**. Supervisors should provide clear instructions, check on progress, and solve problems quickly.

However, an issue has arisen: *a data breach has exposed confidential product details, and customers may now have access to information about the launch before it was intended to be released.*

Your supervisor is needed to take control of the situation by immediately coordinating a response plan, assigning clear roles to team members, and closely monitoring all communication with customers and stakeholders to prevent similar security issues in the future.

Appendix F

Manipulation Check Measures

Leader Type (Study 1 only)

Instructions: *Based on your Supervisor, please rate each item on 7-point Likert scale (1 = Strongly Disagree, 7 = Strongly Agree).*

1. My supervisor is considerate in their interactions with my colleagues and I.
2. My supervisor is friendly.
3. My supervisor seems to care about the employees in our department.

Transparency

Instructions: *Based on your Supervisor, please rate each item on 7-point Likert scale (1 = Strongly Disagree, 7 = Strongly Agree).*

1. My supervisor is open to giving and receiving feedback.
2. My supervisor typically provides their reasoning for decisions.
3. My supervisor's words and actions don't always align. (Reverse-Scored)
4. My supervisor keeps employees informed about things happening in the organization.

Situational Context (Study 2 only)

Instructions: Considering the above situation, discuss the level to which the supervisor is required to express each of the following behaviors. **7-point Likert scale** (1 = Strongly Disagree, 7 = Strongly Agree).

1. This situation required my leader to support employees.
2. This situation required my leader to develop employees.
3. This situation required my leader to recognize employees' work.
4. This situation required my leader to empower employees.
5. This situation required my leader to clarify roles.
6. This situation required my leader to plan ahead.
7. This situation required my leader to monitor operations.
8. This situation required my leader to solve problems.

Appendix G

Outcome Measures

Satisfaction with the Leader

Instructions: *Based on your Supervisor, please rate each item on 7-point Likert scale (1 = Strongly Disagree, 7 = Strongly Agree).*

1. I would be satisfied working under this supervisor.
2. I would feel comfortable having them as my supervisor.
3. This supervisor would look out for my best interest.

Trust in the Leader

Instructions: *Based on your Supervisor, please rate each item on 7-point Likert scale (1 = Strongly Disagree, 7 = Strongly Agree).*

1. I can freely share my ideas, feelings, and hopes with this supervisor.
2. I can talk freely to this supervisor about difficulties I am having at work and know that they will want to listen.
3. If the supervisor left the organization, I would feel a sense of loss.
4. If I shared my problems with this supervisor, I know they would respond constructively and caringly.
5. I would have to say that we have both made considerable emotional investments in our working relationship.
6. This supervisor approaches their job with professionalism and dedication.
7. Given this supervisor's track record, I see no reason to doubt their competence and preparation for the job.
8. I can rely on this supervisor not to make my job more difficult by careless work.
9. My colleagues respect this supervisor.
10. My colleagues consider this supervisor to be trustworthy.
11. I would feel more confident in this supervisor if they were monitored more closely. (Reverse-Scored)

Leader Effectiveness

Instructions: *Based on your Supervisor and the situation above, please rate each item on 7-point Likert scale (1 = Strongly Disagree, 7 = Strongly Agree).*

1. My supervisor would be effective in managing this situation successfully.
2. My supervisor's leadership style aligns with the specific needs of this situation.
3. My supervisor has the necessary skills and expertise to handle the challenges in this situation.
4. My supervisor would provide the right amount and type of guidance to meet the demands of this situation.
5. My supervisor's leadership approach would create challenges in achieving success here. (Reverse-Scored)

Implicit Leader Theory Dimensions (Study 2 only)

Instructions: *Please read each descriptor below and rate how well it describes [a **human leader** / an **AI leader** / an **AI-supported leader (a human leader who uses AI in their day-to-day work)**]. Use a **scale from 1 to 10**, where 1 = 'Not at All Characteristic of [insert leader here]' and 10 = 'Extremely Characteristic of [insert leader here].' Consider each item carefully and provide your rating based on your perception of [insert leader here].*

1. Caring
2. Sympathetic
3. Compassionate
4. Kind
5. Empathetic
6. Selfless
7. Friendly
8. Sensitive
9. Motivated
10. Dedicated
11. Focused
12. Determined
13. Good decision maker
14. Goal oriented
15. Handles stress
16. Controlling
17. Pushy
18. Intimidating
19. Domineering
20. Coercive
21. Demanding
22. Risky
23. Power-hungry
24. Charismatic
25. Sociable
26. Dynamic
27. Bold
28. Commanding
29. Assertive
30. Authoritative
31. Tough
32. Strong
33. Firm
34. Creative
35. Innovative
36. Clever
37. Courageous
38. Well-groomed
39. Well-dressed
40. Masculine

- 41. Feminine
- 42. Tall
- 43. Male
- 44. Female
- 45. Attractive
- 46. Educated
- 47. Intellectual
- 48. Intelligent

**Items 1-8 are related to Sensitivity; 9-15 Dedication; 16-23 Tyranny; 24-27 Charisma; 28-33 Strength; 34-37 Creativity; 38-39 Well-groomed; 40-45 Masculinity; 46-48 Intelligence*

Appendix H

Covariate Measures

Propensity to Trust

Instructions: Please rate each item on **7-point Likert scale** (1 = Strongly Disagree, 7 = Strongly Agree).

Factor 1 – Trusting Others

1. Other people are out to get as much as they can for themselves. (Reverse-Scored)
2. Other people cannot be relied upon. (Reverse-Scored)
3. I have little faith in other people's promises. (Reverse-Scored)
4. Other people are primarily interested in their own welfare despite what they say. (Reverse-Scored)
5. In these competitive times, I have to be alert; otherwise, others will take advantage of me. (Reverse-Scored)
6. Other people who are friendly toward me are disloyal behind my back. (Reverse-Scored)
7. Other people lie to get ahead. (Reverse-Scored)
8. Other people are only concerned with their own well-being. (Reverse-Scored)
9. Other people let you down. (Reverse-Scored)

Factor 2 – Others' Reliability and Integrity

10. Other people can be relied upon to do what they say they will do.
11. Those in authority are likely to say what they really believe.
12. Other people answer public opinion polls honestly.
13. Experts can be relied upon to tell the truth about the limits of their knowledge.
14. Witnesses tell the truth in all circumstances.
15. Other people do what they say they will do.
16. Other people live by the idea that honesty is the best policy.

Factor 3 – Risk Aversion

17. It is important to 'save for a rainy day.'
18. I prefer a modest but safe return on my savings rather than a higher return that is uncertain.
19. "It is better to be safe than sorry."
20. If I had money to invest, I would look for security rather than spectacular returns.

AI Literacy

Instructions: Please rate each item on **7-point Likert scale** (1 = Strongly Disagree, 7 = Strongly Agree).

1. I can distinguish between smart devices and non-smart devices.
2. I do not know how AI technology can help me. (Reverse-Scored)

3. I can identify the AI technology employed in the applications and products I use.
4. I can skillfully use AI applications or products to help me with my daily work.
5. It is usually hard for me to learn to use a new AI application or product. (Reverse-Scored)
6. I can use AI applications or products to improve my work efficiency.
7. I can evaluate the capabilities and limitations of an AI application or product after using it for a while.
8. I can choose a proper solution from various solutions provided by a smart agent.
9. I can choose the most appropriate AI application or product from a variety for a particular task.
10. I always comply with ethical principles when using AI applications or products.
11. I am never alert to privacy and information security issues when using AI applications or products. (Reverse-Scored)
12. I am always alert to the abuse of AI technology.

Personality

Instructions: Please rate each item, based on how well they describe you, using a **7-point Likert scale** (*1 = Strongly Disagree, 7 = Strongly Agree*).

I am someone who...

1. Is talkative
2. Tends to find fault with others
3. Does a thorough job
4. Is depressed, blue
5. Is original, comes up with new ideas
6. Is reserved
7. Is helpful and unselfish with others
8. Can be somewhat careless
9. Is relaxed, handles stress well
10. Is curious about many different things
11. Is full of energy
12. Starts quarrels with others
13. Is a reliable worker
14. Can be tense
15. Is ingenious, a deep thinker
16. Generates a lot of enthusiasm
17. Has a forgiving nature
18. Tends to be disorganized
19. Worries a lot
20. Has an active imagination
21. Tends to be quiet
22. Is generally trusting
23. Tends to be lazy
24. Is emotionally stable, not easily upset
25. Is inventive
26. Has an assertive personality

27. Can be cold and aloof
28. Perseveres until the task is finished
29. Can be moody
30. Values artistic, aesthetic experiences
31. Is sometimes shy, inhibited
32. Is considerate and kind to almost everyone
33. Does things efficiently
34. Remains calm in tense situations
35. Prefers work that is routine
36. Is outgoing, sociable
37. Is sometimes rude to others
38. Makes plans and follows through with them
39. Gets nervous easily
40. Likes to reflect, play with ideas
41. Has few artistic interests
42. Likes to cooperate with others
43. Is easily distracted
44. Is sophisticated in art, music, or literature

Functionality and Trust in AI Technologies (Study 2 only)

Instructions: Please rate each item on a *5-point Likert scale* (*1 = Strongly Disagree, 5 = Strongly Agree*).

Factor 1 – Trust in Voice Assistance

1. I trust that AI can offer information and services that are in my best interest.
2. I trust that my personal data is protected from potential abuse when using AI.
3. I trust that my privacy is protected when using AI.
4. I trust that authorities exert effective control over organizations and companies providing AI services.

Factor 2 – Perceived Usefulness

5. Using AI would enable me to accomplish tasks more quickly.
6. Using AI would improve my performance when accomplishing tasks.
7. Using AI for accomplishing tasks would increase my productivity.
8. Using AI would enhance my effectiveness when accomplishing tasks.
9. I find AI useful for accomplishing tasks.

Factor 3 – Attitude Toward AI Technologies

10. I feel positive toward AI.
11. I feel that using AI is pleasant.
12. Using AI is a good idea.
13. Using AI is a smart way to get things done.

Factor 4 – Human-Like Trust in Smart Technologies

- 14. Smart technologies care about our well-being.
- 15. Smart technologies are sincerely concerned about addressing the problems of human users.
- 16. Smart technologies try to be helpful and do not operate out of selfish interest.
- 17. Smart technologies are truthful in their dealings.
- 18. Smart technologies keep their commitments and deliver on their promises.
- 19. Smart technologies are honest and do not abuse the information and advantages they have over their users.

Factor 5 – Functionality Trust in Smart Technologies

- 20. Smart technologies work well.
- 21. Smart technologies have the features necessary to complete key tasks.
- 22. Smart technologies are competent in their area of expertise.
- 23. Smart technologies are reliable.
- 24. Smart technologies are dependable.