

**How Supportive and Precise AI Communication Styles Influence Adherence to
AI Advice in High-Risk versus Low-Risk Advice-Seeking Context**

En-Hsien Su

Department of Communication, University of Washington

Thesis Advisor: Prof. Wang Liao

Abstract

As AI systems increasingly provide everyday advice, an important question emerges: what makes people trust and follow AI recommendations, especially when the stakes feel risky? Drawing on research on algorithm aversion and trust, this study examines how AI communication style and perceived risk jointly shape advice adherence. A $2 \times 2 \times 2$ between-subjects experiment manipulated AI communication style (supportive vs. precise), perceived risk (high vs. low), and advisory context (cybersecurity vs. public health). Results showed that cognitive-based trust grounded in competence, reliability, and professionalism was the strongest predictor of adherence across contexts. In contrast, affective-based trust predicted adherence only under low-risk conditions. While supportive language increased affective trust, precise communication showed only limited effects on cognitive trust. These findings suggest that in risky advisory situations, users prioritize competence over emotional reassurance, offering new insight into how AI systems can reduce algorithm aversion and encourage reliance on AI advice.

How Supportive and Precise AI Communication Styles Influence Adherence to AI Advice in High-Risk versus Low-Risk Advice-Seeking Context

AI systems are increasingly popular for everyday information- and advice-seeking. A well-documented puzzle in this space is algorithm aversion: people turn away from algorithmic recommendations even when algorithms outperform humans, especially after witnessing a single error (Dietvorst, Simmons, & Massey, 2015). A key mechanism underlying algorithm aversion is trust: people decide whether to follow a recommendation based not only on perceived accuracy, but also on whether they view the system as reliable, competent, and acceptable as an advising agent (Burton et al., 2020). Researchers have examined related influences such as error visibility and the perceived opacity or lack of transparency in algorithmic reasoning (Dzindolet et al., 2003; Hoff & Bashir, 2015). However, prior accounts largely treat trust as context-free, assuming that the same trust dynamics apply across decisions of varying consequence.

This study foregrounds perceived risk—the felt stakes and uncertainty of the decision (Griffin, Dunwoody, & Yang, 2008)—as a neglected factor in trust-related mechanisms of algorithm aversion. When perceived stakes are high, small ambiguities can loom large, and tolerance for algorithmic error may erode more quickly than tolerance for human error. When perceived stakes are low, this asymmetry may soften. At the same time, this study emphasizes communication styles as another key factor that shapes trust, because these styles can signal the core components of an advising agent’s trustworthiness—ability (whether the agent is competent), benevolence (whether the agent aligns with the user), and integrity (whether the agent is morally reliable; Mayer, Davis, & Schoorman, 1995).

Although prior work shows that communication style can shape perceptions of trustworthiness (Renn & Levine, 1991), existing research has not clarified how communication

style and perceived risk jointly influence trust in AI advising. Individuals tend to prioritize reliability and error avoidance in high-risk situations but place greater value on feeling understood and cared for in low-risk contexts (Hoff & Bashir, 2015). This pattern suggests that communication style may serve as a consequential cue for calibrating trust in algorithmic advice, particularly as users shift what they need from an adviser depending on the perceived risk.

However, it remains unclear whether sensitivity to communication style changes across different levels of perceived risk, and whether different styles differentially affect trust and algorithm aversion. By integrating communication style with perceived risk, this study aims to provide a more nuanced account of how users form trust in AI advisers through AI communication style as a designable cue, and how perceived risk alters user interpretations of these cues when deciding whether to trust and ultimately adhere to AI-generated advice.

To achieve the goal, this paper reports a $2 \times 2 \times 2$ between-subjects experiment manipulating the communication style of an AI advisor (supportive vs. precise), the perceived risk of the advice-giving context (high vs. low), and scenario context (cybersecurity vs. mask wearing) to assess their independent and joint effects on trust-related mechanisms (cognitive trust and affective trust) and subsequent advice adherence. The experiment is situated in everyday advisory contexts that are broadly relevant to U.S. adults, specifically cybersecurity account protection and public health–related mask-wearing decisions for which individuals realistically and frequently seek guidance.

The following literature review first examines algorithm aversion and the central role of trust in shaping whether users accept or reject algorithmic advice. Next, it reviews research on communication style, outlining how precise and supportive styles differentially signal ability

and benevolence, thereby shaping distinct forms of trust. Finally, it discusses perceived risk, highlighting how risk alters users' sensitivity to stylistic cues and the conditions under which trust becomes essential for advice adherence.

Trust, Algorithm Aversion, and Communication Style

People often respond differently to advice depending on whether it comes from a human or an algorithm (Dietvorst, Simmons, & Massey, 2015). Even when a machine performs well, users may begin to avoid its recommendations after observing a single mistake. This pattern, known as algorithm aversion, reflects a reduced willingness to rely on algorithmic advice. In other words, algorithm aversion describes a decline in users' reliance on algorithmic recommendations relative to human ones. In experimental research, this reduction is typically assessed through advice adherence, which measures whether a person follows or ignores a recommendation provided by an adviser (Burton et al., 2020). Because adherence is a direct behavioral outcome, it serves as a clear indicator of algorithm aversion: high adherence reflects low algorithm aversion, as users are willing to act on the algorithm's recommendation. In contrast, low adherence reflects stronger algorithm aversion, indicating that users are rejecting the algorithm's advice even when its accuracy remains unchanged.

Trust plays a foundational role in determining whether users rely on algorithmic advice, making it a central mechanism for understanding why algorithm aversion occurs. Burton et al. (2020) emphasize that trust is pivotal in determining whether users accept algorithmic recommendations. Building on this idea, prior work shows that users expect advising systems to appear predictable, competent, and aligned with their interests before accepting their suggestions (Burton et al., 2020). Other studies find that people reduce their reliance on an algorithm when they lack confidence in its reliability or intentions, even when its performance is strong (Hoff &

Bashir, 2015). In addition, users rely more on an algorithm when they see cues that others trust it, suggesting that social trust signals can increase willingness to follow algorithmic advice (Alexander, Blinder, & Zak, 2018). Taken together, these findings indicate that algorithm aversion is not merely a reaction to error visibility or opacity but is fundamentally a trust problem: adherence occurs only when users feel assured that the system is dependable and a legitimate partner in decision-making.

In trust research, evaluations of an adviser's trustworthiness are commonly conceptualized along several dimensions. Mayer, Davis, and Schoorman's (1995) model identifies ability, benevolence, and integrity as core antecedents of trust—judgments that determine whether individuals believe a system is competent, aligned with their interests, and committed to acting reliably. Building on this framework, McAllister (1995) distinguishes between two forms of trust shaped by these judgments: cognitive-based trust, which reflects confidence in an adviser's competence and reliability, and affective-based trust, which reflects a sense of interpersonal care, concern, and psychological safety. Although conceptually distinct, these two forms of trust jointly influence whether individuals feel secure enough to rely on advice.

These dimensions of trust have direct implications for how users respond to an algorithmic adviser. When users perceive the adviser as high in ability or benevolence, both cognitive and affective trust can increase, making them more willing to accept its guidance. Conversely, when users judge the adviser as untrustworthy, adherence declines and manifests as algorithm aversion. This perspective helps explain why even a single visible mistake can sharply reduce reliance: the error functions as a breach of trust rather than merely a performance cue. At the same time, trust also accounts for the opposite pattern—algorithm appreciation. When users

view the system as reliable and aligned with their interests or values, they may readily accept or even over-rely on its recommendations (Logg, Minson, & Moore, 2019). Thus, adherence is shaped less by objective accuracy and more by how users interpret who the adviser is, positioning trust as the central psychological mechanism underlying both decreased and increased reliance on algorithmic advice.

Therefore, cognitive- and affective-based trust represent two distinct pathways through which users decide whether to rely on an algorithmic adviser. Cognitive-based trust should facilitate adherence by reinforcing perceptions of technical competence and reliability, whereas affective-based trust should promote adherence by creating a sense of interpersonal care and psychological safety. Because both forms of trust reduce the likelihood of discounting or avoiding recommendations, higher levels of either form of trust are expected to decrease behavioral manifestations of algorithm aversion. These expectations lead to the following hypotheses:

H1a. In the AI advising contexts, higher levels of cognitive-based trust will increase advice adherence

H1b. In the AI advising contexts, higher levels of affective-based trust will increase advice adherence

Cognitive- and affective-based trust reflect distinct psychological routes to advice adherence, and AI systems—lacking human social cues—may evoke these routes differently than human advisers. Identifying which mechanism most strongly drives adherence is therefore theoretically meaningful. This motivates the following research question.

RQ1. In AI advising contexts, which form of trust—cognitive-based or affective-based—is more effective in increasing adherence?

Having established trust as the central mechanism underlying algorithm aversion, the next question concerns the antecedents of trust. Prior research shows that people infer trust-related qualities from linguistic cues (Burgoon, Birk, & Pfau, 1990). Therefore, communication style may serve as a mechanism that informs trust, thereby influencing whether algorithmic advice is followed. Communication style refers to the patterned linguistic and interpersonal tendencies individuals display when interacting, including choices in tone, structure, and orientation toward the interlocutor (de Vries, Bakker-Pieper, & Oostenveld, 2009). People evaluate not only what is said but also how it is expressed, using stylistic cues to infer a source's intentions and characteristics (Burgoon, Birk, & Pfau, 1990). Thus, even when informational content is held constant, differences in communication style can influence perceived trustworthiness.

Two styles are particularly relevant to the formation of trust. *Preciseness* is characterized by structured expression, clarity, and careful wording (de Vries et al., 2009). Such expression suggests that the adviser is organized, competent, and less prone to error, thereby strengthening cognitive-based trust—the form of trust rooted in confidence in the adviser's ability and reliability (McAllister, 1995). In contrast, *supportiveness* emphasizes responsiveness, acknowledgment, and empathy toward the recipient's concerns (Brown & Levinson, 1987; Bickmore & Picard, 2005). These features convey benevolence and interpersonal alignment, fostering affective-based trust by making the adviser feel understanding, caring, and psychologically safe to rely on (McAllister, 1995).

A similar mechanism operates in interactions with AI systems. Even when users are aware that an AI lacks genuine emotion, they still infer professionalism or care based on linguistic cues (Reeves & Nass, 1996; Nass & Moon, 2000). Supportive language increases feelings of being understood, whereas precise language strengthens perceptions of reliability and expertise

(Bickmore & Picard, 2005). This is especially consequential because AI systems do not possess pre-existing relational credibility—users must rely on linguistic style to judge trustworthiness. From a processing perspective, when individuals face uncertainty or lack domain knowledge, they rely more on accessible linguistic cues as heuristic signals of credibility rather than engaging in systematic evaluation (Chaiken, 1980). Thus, communication style is not merely a surface feature; it functions as a core cue for trust formation and advice adherence.

Although prior research shows that communication style can shape perceptions of ability, benevolence, and trustworthiness, existing work has not clarified whether different styles lead to different types of trust—cognitive versus affective—and whether these trust types play different roles in shaping advice adherence in AI advising contexts. To address this gap, it is necessary to distinguish how precise and supportive styles map onto each trust pathway. Therefore, I propose the following hypotheses:

H2a: Regarding AI communication styles, preciseness will lead to higher levels of cognitive-based trust than supportiveness.

H2b: Regarding AI communication styles, supportiveness will lead to higher levels of affective-based trust than preciseness.

Perceived Risk, Trust, and Communication Style

Perceived risk shapes how individuals evaluate advice and determine whether to rely on it. When situations are perceived as uncertain and consequential, individuals place greater emphasis on the credibility of the information source, and trust becomes a prerequisite for adherence (Griffin, Dunwoody, & Neuwirth, 1999). In contexts of risk and uncertainty, trust serves to fill gaps left by incomplete information—particularly when the decision environment lacks transparency or when individuals are unable to fully verify information themselves

(Slovic, 1993; Siegrist & Cvetkovich, 2000; Earle & Siegrist, 2006), as perceived risk increases, trust becomes a more central determinant of whether individuals accept or reject guidance. By contrast, when perceived risk is low, individuals rely more on heuristic processing (Petty & Cacioppo, 1986; Chaiken, 1980). Prior work shows that people often draw inferences from surface-level cues available in the message or its presentation (Sundar, 2008). Although these studies do not examine communication style specifically, this framework suggests that communication style could similarly function as a peripheral cue—shaping impressions such as friendliness or professionalism—without exerting a strong influence on adherence when perceived vulnerability is low. Thus, while communication styles may be noticed, they are not decisive for decision-making when perceived risk is minimal.

In AI advising interactions, this dynamic is especially pronounced. Because AI systems lack a pre-existing reputation or relational history, users must rely on communication style as a key signal for evaluating trustworthiness. Communication style therefore functions as a trust cue that users interpret to infer whether the AI is competent or aligned with their interests (Terwel et al., 2009). The social amplification of risk framework suggests that in high-risk contexts, individuals become more sensitive to interpersonal and stylistic cues—such as tone, stance, and expressiveness—because these cues help them judge whether a source is dependable and worthy of reliance (Kasperson et al., 1988). At the same time, heightened perceived risk is often accompanied by increased anxiety, which amplifies the motivation to seek reassurance and clarity (Gadarian & Albertson, 2014). Under such uncertainty, individuals also rely more on affective evaluations to determine whether a source feels “safe” to trust (Slovic & Peters, 2006), suggesting that supportive communication may shape the emotional meaning of an interaction. Together, these perspectives suggest that high-risk contexts should amplify the influence of

communication style on trust. Specifically, a precise communication style should be more effective under high risk in fostering cognitive-based trust, as it conveys competence and reduces uncertainty, whereas a supportive communication style should be more effective under high risk in fostering affective-based trust, as it provides psychological safety and signals alignment with the recipient. Thus, perceived risk is expected to moderate the effects of communication style on trust.

H3a. The positive effect of preciseness on cognitive-based trust will be stronger when perceived risk is high versus low.

H3b. The positive effect of supportiveness on affective-based trust will be stronger when perceived risk is high versus low.

Because trust functions as the proximal psychological predictor of adherence, these moderated effects should extend to behavioral outcomes:

H4. The indirect effect of communication style on adherence through trust will be stronger when perceived risk is high versus low.

This raises a broader question about how perceived risk moderates the relationship between different types of trust and advice adherence.

RQ2. How do cognitive-based and affective-based trust differentially influence advice adherence in high-risk versus low-risk AI advising contexts?

Methods

Participants

A total of 227 participants were recruited through Prolific, an online platform commonly used for social scientific research. Eligibility criteria include being at least 18 years old, residing in the United States at the time of participation, and demonstrating English proficiency, because

the scenarios were written in English and required comprehension of health and cybersecurity contexts. Individuals currently affiliated with the University of Washington were excluded due to funding restrictions. Participants who failed an attention-check item were also excluded from the final sample. All responses were collected anonymously through Qualtrics, and participants were compensated \$1 for completing the study, with the median completion time approximately 5 minutes. The study was approved by the IRB at the author's institution.

Procedures

The experiment involves a 2 (precise versus supportive style) $\times$ 2 (high versus low risk) $\times$ 2 (cybersecurity versus mask-wearing scenario) between-subjects design. Participants first read one scenario randomly selected from the two contexts: In the cybersecurity scenario, the protagonist received a security alert from Google, whereas in the mask-wearing scenario, the protagonist observed change in mask-wearing behaviors in public transits during a pandemic. Each of the context further varies in the level of risk.

In the high-risk cybersecurity condition, participants read: "Anna is in her 60s. She just received a security alert from Google that there had been unauthorized attempts to access her accounts. Her login credentials had also been exposed in several confirmed data breaches. To prevent leaking sensitive data, including her credit card information, she must reset her password and complete security verification immediately. She wonders what she should do." In the low-risk condition, participants read: "Anna is a college senior. She recently received a notification from Google about increasing incidents of data breaching and unauthorized attempts to access people's online accounts. Google encourages her to complete security verification and enable multi-step authentication to safeguard account security. She wonders what she should do."

In the high-risk mask-wearing condition, participants read: "Anna is a recent retiree in

her seventies. The city where she lives experienced a respiratory disease outbreak several months ago, and some of her friends became seriously ill. Anna has been wearing facial mask carefully. Since last week, Anna has noticed that during her daily bus trips to the grocery store, more and more passengers have started wearing masks again. She has also heard that another outbreak is coming. Anna wonders what she should do.” In the low-risk condition, participants read: “Anna is a college senior in her twenties. The city where she lives experienced a respiratory disease outbreak several months ago, and some of her friends became ill from it. Anna has been wearing facial mask carefully. Since last week, Anna has noticed that during her daily bus trips to the grocery store, only a few passengers were wearing masks. She heard that the outbreak is fading. Anna wonders what she should do.”

After reading the scenario, participants proceeded to review a conversation transcript between the protagonist and an AI chatbot. In the cybersecurity scenario, the protagonist started the conversation with “I received a security notification about unusual activity on Google account. What are some actions I can take to protect my account?” In the mask wearing scenario, the protagonist stated “I’ve noticed some changes in how people around me are wearing face masks during my bus trips. How should I respond in this situation?”

The content of AI’s advice remained constant within each scenario, while communication style varied in terms of being precise or supportive. In the cybersecurity scenario, the precise AI response reads: “Based on the information provided, take the following steps: 1. Review the security alert, identify the required actions and deadline. 2. Change password and complete security verifications. 3. Check your account activity for unfamiliar login attempts. 4. Report any suspicious activity. 5. Keep a record of the actions taken for future reference.” In contrast, the supportive AI response reads: “I understand how receiving a security alert like this could

make you feel uneasy, especially when it involves your personal accounts. Before sharing my take on this, may I hear a bit more about your thoughts and hunch? Are you mainly worried that someone accessed your account, unsure whether the alert is real, or trying to figure out what to do first? This would help us dissect the problem more carefully. Here are few general steps that people in a similar situation can do: 1. Take a careful look at the security alert so we can understand what action is required and whether there is a deadline. 2. Once you have done that, change your password and complete the security verification in your account settings to protect your account. 3. You can also check your account activity to see if there are any unfamiliar access attempts. If anything seems suspicious, reporting it through the platform's official support channel can help resolve the issue. 4. Also, I suggest you keep a record of what you've done. This may help you feel more at ease because we can see what has already been handled. If you want, you may tell me more about what the alert said, what worries you most, and whether anything unusual has happened recently. I can help create a more personalized and targeted strategy for your situation."

In the mask-wearing scenario, the precise AI response reads: "Based on the information provided, take the following steps: 1. Keep monitoring mask-wearing behaviors in public spaces, as well as people showing symptoms. 2. Maintain safe distance from others. 3. Wash your hands regularly with soap for at least 20 seconds. 4. Avoid touching your face, especially eyes, nose, and mouth. 5. Keep wearing a face mask in crowded environments. 6. Consistently applying these preventive practices over the next 6 months." In contrast, the supportive AI response reads: "I can understand why this situation might feel a little concerning, especially when you notice changes in how people wear masks during your bus trips. Before sharing my take on this, may I hear a bit more about your situation and concerns? Are you mostly worried about getting sick,

feeling unsure about what is normal, or just trying to decide what precautions make sense? Here are few general steps that people in a similar situation can do: 1. Pay attention to conditions in shared public spaces, especially on public transportation or in crowded indoor environments. 2. Keep some distance from others in crowded places to lower the chance of close contact with someone who may be sick. 3. Wash your hands regularly with soap and water for at least 20 seconds, especially after returning home or touching shared surfaces. 4. Avoid touching your eyes, nose, or mouth after contact with public surfaces. 5. Wear a face mask in crowded indoor spaces to help reduce exposure. Keeping these precautions over the next six months may help protect your health and give you more peace of mind. If you have chronic health conditions, that would also be important to consider when choosing the best precautions. If you want, you may tell me more about your specific concerns, how often you take the bus, and what situations make you feel most uncomfortable. I can provide more specific suggestions to help you navigate through this situation.”

After reading the transcript, participants completed a survey with measures of the key variables. Attention check items are embedded within the trust section, and responses from participants who fail the attention check or complete the survey at implausibly fast speeds (≤ 1 minute) were excluded prior to analysis.

Measurements

Unless otherwise specified, all variables were measured a 7-point Likert scale. Items within each scale are presented in randomized order to reduce response bias.

For *advice adherence*, participants responded to the question, “If you were Anna, how likely would you follow the AI’s advice?” in a 7-point scale ranging from extremely unlikely (1) to extremely likely (7) ($M = 5.62$, $SD = 1.258$).

Trust in the AI advisor was measured based on McAllister's (1995). Specifically *cognitive-based trust* was measured as participants' 7-level agreement to statements such as "the AI appears to approach Anna's issue with professionalism," "based on the information provided, I have confidence in this AI's competence in handling this issue," "I could rely on this AI to provide careful and consistent guidance," and "most people in a similar situation would consider this AI trustworthy." Items were averaged as a composite score ($M = 5.70$, $SD = 1.10$, Cronbach's $\alpha = .90$). *Affective-based trust* was similarly measured using statements including "I could openly share my concerns with this AI," "if I shared my problems with this AI, I believe it would respond in a constructive and caring way," "interacting with this AI would make me feel emotionally supported in this situation," and "I believe this AI would want to listen if I talked about difficulties related to my situation." Items were averaged as a composite score ($M = 5.13$, $SD = 1.34$, Cronbach's $\alpha = .89$).

In addition to the variables involved in the hypotheses and RQs additional perceptions were measured for manipulation checks. For the manipulation of communication style, *perceived preciseness* ($M = 5.73$, $SD = 1.09$, Cronbach's $\alpha = .87$) was measured using items conceptually adapted from de Vries et al. (2009), with the following items: "the AI's response presented the information in a clear and well-organized structure," "the AI's response showed careful consideration in how the information was explained," "the AI's response focused on important information rather than trivial or superficial details," and "the AI's response communicated its message using only the amount of detail that was necessary." *Perceived supportiveness* ($M = 5.52$, $SD = 1.27$, Cronbach's $\alpha = .90$). was assessed using items adapted from politeness theory and relational communication frameworks (Brown & Levinson, 1987), including "the AI's response showed that it noticed and acknowledged Anna's situation and concerns," "the AI

responded in a way that reflected an understanding of what mattered to Anna,” and “the AI communicated as if it were working together with Anna to address the situation.”

For the manipulation of risk, *perceived risk* was measured using two questions on a zero to ten scale ($M = 5.83$, $SD = 2.24$, $r = .714$, $p < .001$), including “how risky does Anna’s situation feel to you?” with 0 = “not risky at all” to 10 = “extremely risky,” and “if you were Anna, how likely do you think something negative would happen?” with 0 = “very unlikely” to 10 = “very likely.”

Results

Manipulation Checks

The risk manipulation successfully influenced participants’ perceived risk, $t(225) = 4.36$, $p < .001$, Cohen’s $d = .58$. Participants in the high-risk conditions reported higher perceived risk on average ($M = 6.44$, $SD = 2.28$) than participants in the low-risk conditions ($M = 5.19$, $SD = 2.01$). The communication style manipulation successfully influenced participants’ perceived preciseness, $t(225) = 2.44$, $p = .015$, Cohen’s $d = .32$. Participants in the precise condition reported higher perceived preciseness on average ($M = 5.90$, $SD = 0.85$) than participants in the supportive condition ($M = 5.55$, $SD = 1.27$). At the same time, there is marginal evidence that the manipulation also influenced perceived supportiveness, $t(225) = 1.94$, $p = .054$, Cohen’s $d = .26$. Participants in the supportive condition reported higher perceived supportiveness on average ($M = 5.69$, $SD = 1.22$) than participants in the precise condition ($M = 5.36$, $SD = 1.31$).

Hypothesis Testing and Research Questions

H1a and 1b predicted that higher cognitive-based trust and affective-based trust would increase advice adherence. To test these hypotheses and address RQ1, both trust dimensions

were entered a regression as predictors of advice adherence. The model was significant, $F(2, 224) = 65.32, p < .001, R^2 = .368$. Cognitive-based trust significantly predicted advice adherence, $B = .628, SE = .080, t = 7.81, p < .001, 95\% CI [.470, .787]$, supporting H1a. However, affective-based trust was not a significant predictor, $B = .080, SE = .066, t = 1.20, p = .230, 95\% CI [-.051, .210]$. Thus, H1b was not supported.

RQ1 also asked which of the two types of trust has a stronger effect on advice adherence. Based on the above-mentioned regression model, a linear hypothesis test compared the affective-based trust coefficient with the cognitive-based trust coefficient. The coefficient difference ($-.549$) was significant, $SE = .133, p < .001, 95\% CI [-.811, -.286]$. Cognitive-based trust was a stronger predictor of advice adherence than affective-based trust.

H2a predicted that the precise communication style would lead to higher cognitive-based trust than the supportive communication style. A linear regression using the communication style manipulation to predict cognitive-based trust showed that the effect was in the expected direction but only marginally significant, $B = -.256, SE = .145, \beta = -.117, t(225) = -1.76, p = .079, R^2 = .014$. Thus, H2a was only marginally supported. H2b predicted that the supportive communication style would lead to higher affective-based trust than the precise communication style. This hypothesis was supported, as the communication style manipulation significantly altered affective-based trust, $B = .356, SE = .176, \beta = .134, t(225) = 2.02, p = .044, R^2 = .018$. Participants in the supportive condition reported higher affective-based trust than those in the precise condition.

H3a predicted that the effect of communication style on cognitive-based trust would be stronger under high-risk conditions than under low-risk conditions. A factorial ANOVA was conducted with communication style, risk level, scenario context, and their interactions

predicting cognitive-based trust. As shown in Figure 1, the profile plot suggested that cognitive-based trust was somewhat higher for the precise message than the supportive message under low-risk conditions, while this difference was smaller under high-risk conditions. However, the risk $\times$ style interaction was not significant, $F(1, 219) = 2.22, p = .138$. Therefore, H3a was not supported. It is noteworthy that the main effect of communication style was marginally significant, $F(1, 219) = 3.60, p = .059$, suggesting that precise messages tended to produce slightly higher cognitive-based trust overall.

H3b predicted that the effect of supportive communication style on affective-based trust would be stronger under high-risk conditions than under low-risk conditions. A factorial ANOVA was conducted with communication style, risk level, scenario context, and their interactions predicting affective-based trust. As shown in Figure 2, supportive messages appeared to produce higher affective-based trust than precise messages, especially under high-risk conditions. However, the risk $\times$ style interaction was not significant, $F(1, 219) = 1.89, p = .170$. Therefore, H3b was not supported. The main effect of communication style was marginally significant, $F(1, 219) = 3.61, p = .059$, suggesting that supportive messages tended to elicit higher affective-based trust overall.

H4 predicted that the indirect effect of communication style on advice adherence through trust would be stronger under high perceived risk than under low perceived risk. This hypothesis was tested using PROCESS Model 59, with communication style as the independent variable, advice adherence as the dependent variable, affective-based trust and cognitive-based trust as parallel mediators, and perceived risk as the moderator. The conditional indirect effect through affective-based trust was not significant under either low risk, effect = .035, 95% bootstrap CI $[-.134, .177]$, or high risk, effect = $-.036$, 95% bootstrap CI $[-.210, .062]$. The index of

moderated mediation for the affective-based trust pathway was also not significant, index = $-.071$, 95% bootstrap CI $[-.286, .118]$. For the cognitive-based trust pathway, the indirect effect was significant under low risk, effect = $-.215$, 95% bootstrap CI $[-.458, -.025]$, but not under high risk, effect = $-.042$, 95% bootstrap CI $[-.375, .248]$. However, the index of moderated mediation was not significant, index = $.174$, 95% bootstrap CI $[-.204, .564]$. Therefore, H4 was not supported, although there is some signal that perceived risk may moderated the indirect effect of communication style on advice adherence through cognitive-based trust.

RQ2 asked how cognitive-based and affective-based trust differentially influenced advice adherence in high-risk versus low-risk AI advising contexts. Using the final adherence model from the same PROCESS analysis, RQ2 examined whether the effects of cognitive-based and affective-based trust differed across risk levels. Affective-based trust significantly predicted advice adherence under low-risk conditions, $B = .287$, $SE = .093$, $t = 3.09$, $p = .002$, 95% CI $[.104, .470]$. However, this effect was not significant under high-risk conditions, $B = -.062$, $SE = .099$, $t = -.63$, $p = .531$, 95% CI $[-.257, .133]$. The interaction between affective-based trust and risk was significant, $B = -.349$, $SE = .136$, $t = -2.57$, $p = .011$, indicating that the relationship between affective-based trust and adherence differed by risk level. In contrast, cognitive-based trust significantly predicted advice adherence under both low-risk conditions, $B = .455$, $SE = .115$, $t = 3.97$, $p < .001$, 95% CI $[.229, .681]$, and high-risk conditions, $B = .734$, $SE = .119$, $t = 6.19$, $p < .001$, 95% CI $[.500, .967]$. The interaction between cognitive-based trust and risk did not reach conventional significance, $B = .279$, $SE = .165$, $t = 1.69$, $p = .092$. Overall, these results suggest that cognitive-based trust was a more stable predictor of advice adherence across risk contexts, whereas affective-based trust predicted adherence only in low-risk contexts.

Discussion

This study examined how AI communication style and perceived risk shape trust and advice adherence in everyday AI advising contexts. Drawing on algorithm aversion research, the study conceptualized advice adherence as more than a response to the content of an AI recommendation; it is also a reflection of whether users trust the AI advisor enough to rely on it. The findings partially support the proposed hypotheses. Most importantly, the results show that trust in AI advising is not a single mechanism. Cognitive-based trust, which reflects competence, reliability, and professionalism, was the strongest and most stable predictor of advice adherence across risk contexts. In contrast, affective-based trust played a more limited role, predicting adherence only in low-risk contexts and becoming nonsignificant once cognitive-based trust was included in the same model. These findings suggest that algorithm aversion is fundamentally tied to competence-based trust: users are more likely to follow AI advice when the system appears dependable and capable, rather than merely caring or emotionally supportive.

The second major finding concerns how communication style shapes different forms of trust. Consistent with H2b, supportiveness increased affective-based trust, suggesting that language that acknowledges users' concerns and signals care can make an AI advisor feel more emotionally responsive. However, H2a was not supported: preciseness showed only a weak tendency to increase cognitive-based trust. This pattern suggests that supportive and precise styles do not operate as simple opposites. A supportive AI response can still contain useful and organized advice, and a precise response may not automatically create stronger perceptions of competence unless its clarity, structure, and expertise are especially salient. In this sense, communication style matters, but its effects are uneven. Supportive style appears more effective at shaping emotional impressions of the AI, while precise style may require stronger cues of

expertise or evidence quality to reliably increase cognitive-based trust.

The third major finding concerns the role of perceived risk. Contrary to H3a, H3b, and H4, perceived risk did not significantly strengthen the effects of communication style on trust, nor did it significantly moderate the indirect effect of communication style on adherence through either form of trust. In other words, high-risk contexts did not simply make users more responsive to precise or supportive communication styles. However, RQ2 showed that risk still mattered in a more specific way: it changed which type of trust translated into advice adherence. Cognitive-based trust predicted adherence in both low-risk and high-risk contexts, whereas affective-based trust predicted adherence only in low-risk contexts. This suggests that when the stakes feel higher, users may become less persuaded by emotional support and more dependent on whether the AI appears knowledgeable, reliable, and competent. Thus, risk does not seem to primarily alter how communication style induces trust; rather, it alters which form of trust becomes useful for deciding whether to follow AI advice.

Limitations and Future Directions

Several limitations should be noted. First, the supportiveness manipulation was only marginally significant. Although participants in the supportive condition reported higher perceived supportiveness than those in the precise condition, the difference did not reach the conventional significant level. This suggests that the supportive and precise AI response may not have been perceived as fully distinct. One possible reason is that the supportive responses still contained concrete advice, structure, and action steps, while the precise responses were still polite and reasonable. Future studies could create stronger contrast between communication styles by making supportive responses more explicitly relational, emotionally validating, and collaborative, while keeping precise responses more concise, structured, and competence

focused. At the same time, future designs should carefully hold the core response recommendation constant so that differences in trust can be attributed to style rather than informational content. Future studies could also avoid treating preciseness and supportiveness as opposite ends of a single dimension. In the present study, the design compared supportive and precise responses as separate styles, but the manipulation suggests that an AI response can be both supportive and precise at the same time. A more refined future design could use a 2×2 model that independently manipulates supportiveness and preciseness, such as high versus low supportiveness crossed with high versus low preciseness. This would allow researchers to examine whether the most effective AI advice is not simply supportive or precise, but a combination of relational care and clear, competence-focused guidance.

Second, the study measured advice adherence intention rather than actual behavioral adherence. Although intention is useful for examining algorithm aversion in controlled experimental settings, it may not fully capture how people behave when they face real consequences. Participants were asked to imagine what Anna should do, rather than making decisions about their own accounts, health, or personal safety. Future research could examine behavioral outcomes more directly, such as whether users click a security-check link, choose to enable two-factor authentication, save an AI-generated plan, or select whether to follow a recommended precaution. This behavioral measure would provide stronger evidence about when communication style and trust reduce algorithm aversion.

Third, the study used hypothetical everyday advisory scenarios in only two domains: cybersecurity and mask wearing. These contexts are useful because they are familiar and broadly relevant, but they may not represent all forms of high-risk AI advising. In more consequential domains, such as medical triage, financial decisions, immigration advising, or legal guidance,

users may respond differently to precision, supportiveness, and risk. Future research should test whether the same pattern holds across domains where the stakes are more personally consequential. This would help clarify whether cognitive-based trust remains the dominant predictor of adherence when advice involves more serious or irreversible outcomes.

Fourth, the role of perceived risk should be interpreted more narrowly than originally expected. Rather than directly strengthening the indirect pathway from communication style to adherence through trust, risk appeared to shape which type of trust mattered for adherence. Future research could examine perceived risk more precisely by measuring related factors such as vulnerability, anxiety, decision urgency, and perceived consequence. These variables may help explain why affective-based trust becomes less predictive when users perceive the situation as high risk.

Finally, this study relied on U.S.-based adult participants recruited through Prolific, which may limit the generalizability of the findings. Future research should examine whether the same trust patterns appear across different cultural groups, age groups, and levels of AI familiarity. For example, users in Western contexts may place stronger emphasis on clarity, autonomy, and competence-based trust, whereas East Asian users may be more sensitive to relational cues, politeness, emotional reassurance, and harmony-oriented communication. These cultural differences may shape whether cognitive-based or affective-based trust plays a stronger role in AI advice adherence. Cross-cultural studies could therefore clarify whether cognitive-based trust is universally central to reducing algorithm aversion, or whether affective-based trust becomes more important in cultural contexts where relational communication is more valued.

Conclusion

This study contributes to research on algorithm aversion by showing that users'

willingness to follow AI advice depends not only on the advice itself, but also on the type of trust the AI advisor produces. Across everyday advisory contexts, cognitive-based trust emerged as the most consistent predictor of advice adherence, suggesting that users are more likely to rely on AI when they perceive it as competent, reliable, and professional. Affective-based trust played a more limited role, predicting adherence only in low-risk contexts. Although perceived risk did not significantly strengthen the indirect effect of communication style on adherence through trust, it helped reveal that different forms of trust matter under different risk conditions. These findings suggest that AI advising systems should be designed not only to sound supportive, but also to clearly communicate competence, reliability, and actionable guidance. By distinguishing cognitive-based and affective-based trust, this study provides a more nuanced account of how communication style and perceived risk shape algorithmic advice adherence.

References

- Alexander, V., Blinder, C., & Zak, P. J. (2018). Why trust an algorithm? Performance, cognition, and neurophysiology. *Computers in Human Behavior*, 89, 279–288.
<https://doi.org/10.1016/j.chb.2018.08.010>
- Bickmore, T. W., & Picard, R. W. (2005). Establishing and maintaining long-term human– computer relationships. *ACM Transactions on Computer-Human Interaction*, 12(2), 293–327.
<https://doi.org/10.1145/1067860.1067867>
- Burgoon, J. K., Birk, T., & Pfau, M. (1990). Nonverbal behaviors, persuasion, and credibility. *Human Communication Research*, 17(1), 140–169. <https://doi.org/10.1111/j.1468-2958.1990.tb00065.x>
- Brown, P., & Levinson, S. C. (1987). *Politeness: Some universals in language usage*. Cambridge University Press. <https://doi.org/10.1017/CBO9780511813085>
- Burton, J. W., Stein, M.-K., & Jensen, T. B. (2020). A systematic review of algorithm aversion in augmented decision making. *Journal of Behavioral Decision Making*, 33(2), 220–239.
<https://doi.org/10.1002/bdm.2155>
- Cacioppo, J. T., Petty, R. E., Kao, C. F., & Rodriguez, R. (1986). Central and peripheral routes to persuasion: An individual difference perspective. *Journal of Personality and Social Psychology*, 51(5), 1032–1043. <https://doi.org/10.1037/0022-3514.51.5.1032>
- Chaiken, S. (1980). Heuristic versus systematic information processing and the use of source versus message cues in persuasion. *Journal of Personality and Social Psychology*, 39(5), 752–766.
<https://doi.org/10.1037/0022-3514.39.5.752>
- De Vries, R. E., Bakker-Pieper, A., Konings, F. E., & Schouten, B. (2011). The Communication Styles Inventory (CSI): A six-dimensional behavioral model of communication styles and its relation with personality. *Communication Research*, 40(4), 506–532.
<https://doi.org/10.1177/0093650211413571>

- De Vries, R. E., Bakker-Pieper, A., & Oostenveld, W. (2010). Leadership = communication? The relations of leaders' communication styles with leadership styles, knowledge sharing and leadership outcomes. *Journal of Business and Psychology*, 25(3), 367–380.
<https://doi.org/10.1007/s10869-009-9140-5>
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114–126. <https://doi.org/10.1037/xge0000033>
- Dzindolet, M. T., Peterson, S. A., Pomranky, R. A., Pierce, L. G., & Beck, H. P. (2003). The role of trust in automation reliance. *International Journal of Human-Computer Studies*, 58(6), 697–718.
[https://doi.org/10.1016/S1071-5819\(03\)00038-7](https://doi.org/10.1016/S1071-5819(03)00038-7)
- Earle, T. C., & Siegrist, M. (2006). Morality information, performance information, and the distinction between trust and confidence. *Journal of Applied Social Psychology*, 36(2), 383–416.
<https://doi.org/10.1111/j.0021-9029.2006.00012.x>
- Gadarian, S. K., & Albertson, B. (2014). Anxiety, immigration, and the search for information. *Political Psychology*, 35(2), 133–164. <https://doi.org/10.1111/pops.12034>
- Griffin, R. J., Dunwoody, S., & Neuwirth, K. (1999). Proposed model of the relationship of risk information seeking and processing to the development of preventive behaviors. *Environmental Research*, 80(2), S230– S245. <https://doi.org/10.1006/enrs.1998.3940>
- Griffin, R. J., Dunwoody, S., & Yang, Z. J. (2012). Linking risk messages to information seeking and processing (College of Communication Faculty Research and Publications, No. 93). Marquette University. https://epublications.marquette.edu/comm_fac/93
- Hoff, K. A., & Bashir, M. (2015). Trust in automation. *Human Factors*, 57(3), 407–434.
<https://doi.org/10.1177/0018720814547570>
- Kasperson, R. E., Renn, O., Slovic, P., Brown, H., Emel, J., Goble, R., Kasperson, J. X., & Ratick, S.

(1988). The social amplification of risk: A conceptual framework. *Risk Analysis*, 8(2), 177–187.

<https://doi.org/10.1111/j.1539-6924.1988.tb01168.x>

Kasperson, R. E., Stallen, P. J. M., Renn, O., & Levine, D. (1992). Credibility and trust in risk communication. In R. E. Kasperson & P. J. M. Stallen (Eds.), *Communicating risks to the public: International perspectives* (pp. 175–217). Springer. https://doi.org/10.1007/978-94-009-1952-5_10

Logg, J. M., Minson, J. A., & Moore, D. A. (2019). Algorithm appreciation: People prefer algorithmic to human judgment. *Organizational Behavior and Human Decision Processes*, 151, 90–103. <https://doi.org/10.1016/j.obhdp.2018.12.005>

Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of Management Review*, 20(3), 709–734. <https://doi.org/10.5465/amr.1995.9508080335>

McAllister, D. J. (1995). Affect- and cognition-based trust as foundations for interpersonal cooperation in organizations. *Academy of Management Journal*, 38(1), 24–59. <https://doi.org/10.5465/256727>

Nass, C., & Moon, Y. (2000). Machines and mindlessness: Social responses to computers. *Journal of Social Issues*, 56(1), 81–103. <https://doi.org/10.1111/0022-4537.00153>

Reeves, B., & Nass, C. (1996). *The media equation: How people treat computers, television, and new media like real people and places*. Cambridge University Press.

Siegrist, M., & Cvetkovich, G. (2000). Perception of hazards: The role of social trust and knowledge. *Risk Analysis*, 20(5), 713–720. <https://doi.org/10.1111/0272-4332.205064>

Slovic, P. (1993). Perceived risk, trust, and democracy. *Risk Analysis*, 13(6), 675–682. <https://doi.org/10.1111/j.1539-6924.1993.tb01329.x>

Slovic, P., & Peters, E. (2006). Risk perception and affect. *Current Directions in Psychological Science*, 15(6), 322–325. <https://doi.org/10.1111/j.1467-8721.2006.00461.x>

Sundar, S. S. (2008). The MAIN model: A heuristic approach to understanding technology effects on

credibility. In M. J. Metzger & A. J. Flanagin (Eds.), *Digital media, youth, and credibility* (pp. 73–100). The MIT Press. <https://doi.org/10.1162/dmal.9780262562324.073>

Terwel, B. W., Harinck, F., Ellemers, N., & Daamen, D. D. L. (2009). Competence-based and integrity-based trust as predictors of acceptance of carbon dioxide capture and storage (CCS). *Risk Analysis*, 29(8), 1129–1140. <https://doi.org/10.1111/j.1539-6924.2009.01256.x>

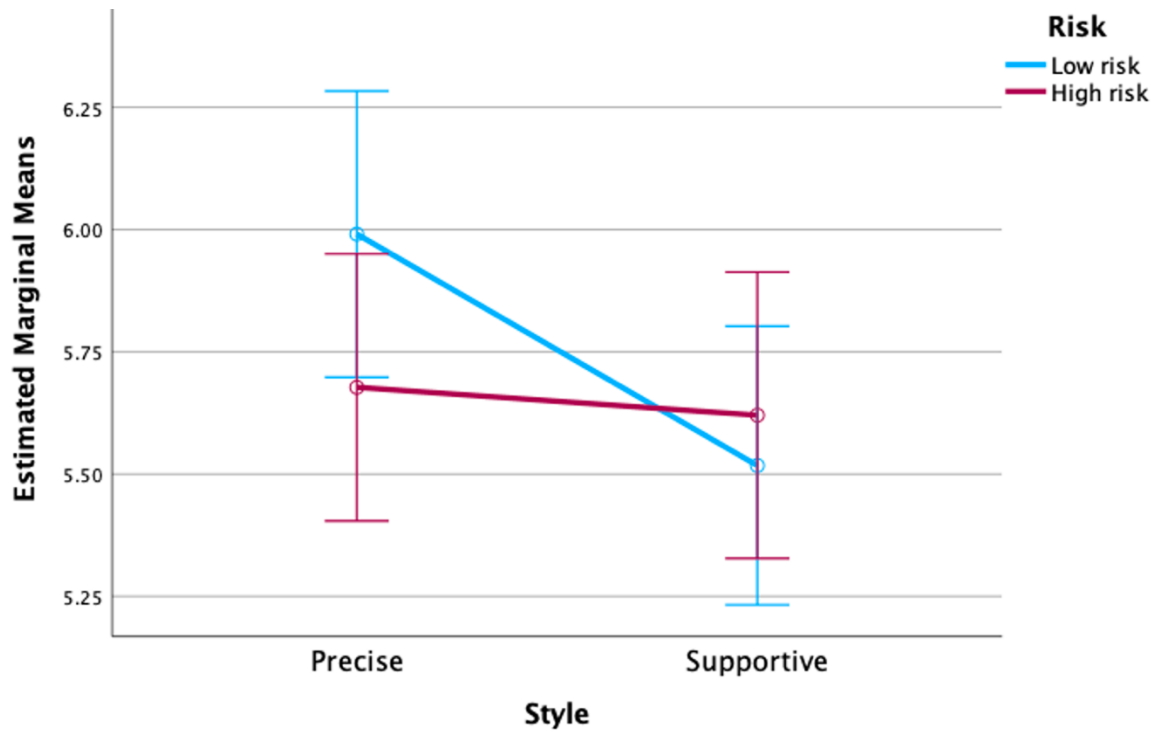


Figure 1

Estimated marginal means of cognitive-based trust by communication style and risk level. Error bars represent 95% confidence intervals.

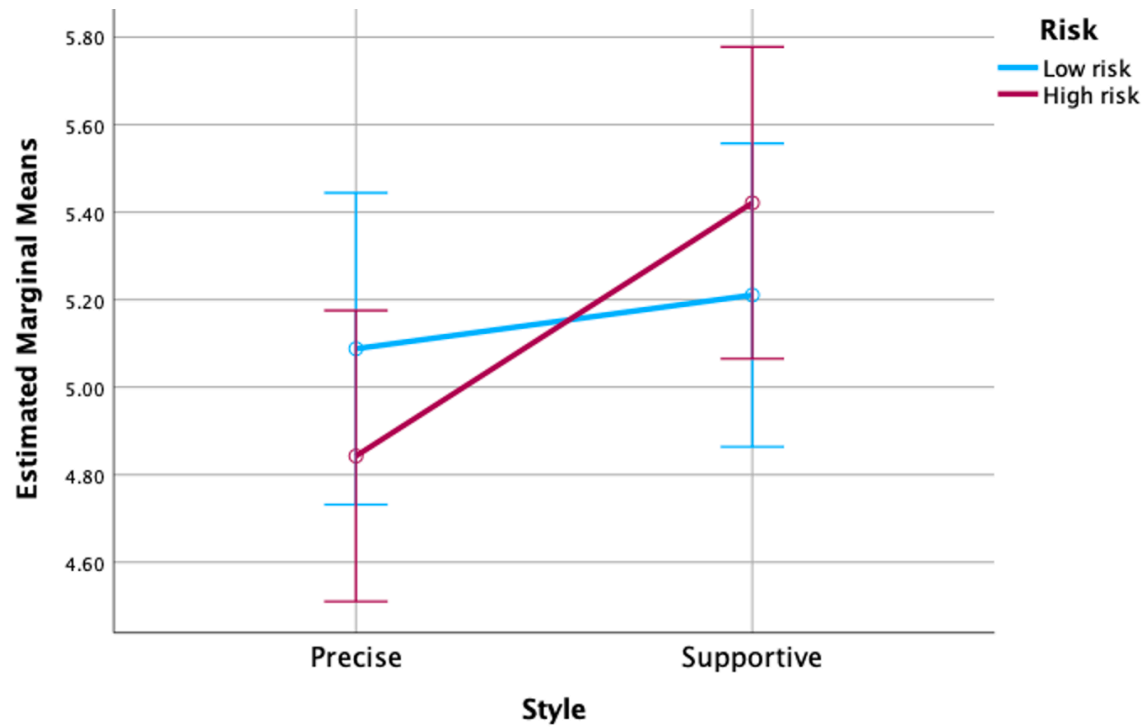


Figure 2

Estimated marginal means of affective-based trust by communication style and risk level. Error bars represent 95% confidence intervals.