

Journal of Applied Social Psychology

Original Article

Underestimating Others’ Positive Attitude Toward AI Reduces Intentions to Adopt AI in Various Business Contexts

Submission ID 6375341e-cb72-43d7-a28d-0b716cff88ab

Submission Version Initial Submission

PDF Generation 09 Sep 2025 01:14:38 EST by Atypon ReX

Files for peer review

All files submitted by the author for peer review are listed below. Files that could not be converted to PDF are indicated; reviewers are able to access them online.

Name	Type of File	Size	Page
manuscript.docx	Anonymized Main Document - MS Word	1.4 MB	Page 3
supplementary materials.docx	Supplementary Material for Review	70.7 KB	Page 49

Abstract

This research, employing an intersubjective approach, explores how individuals' perceptions of others' AI attitudes influence their adoption willingness. Findings from an online survey ($N = 309$) and three lab studies (Study 2: $N = 261$; Study 3a: $N = 247$; Study 3b: $N = 253$), and a single-paper meta-analysis reveal that individuals consistently underestimate others' positive attitudes toward AI, leading to reduced willingness to adopt AI. Perceived efficacy and risks partially mediate this effect across diverse scenarios (e.g., home purchase, job interviews, online medical consultation), with a single-paper meta-analysis indicating medium effect sizes. Addressing biased perceptions of AI attitudes could promote adoption via the social norms approach.

Keywords: artificial intelligence, AI acceptance, intersubjective perception, outcome efficacy, perceived risk

Underestimating Others' Positive Attitude Toward AI Reduces Intentions to Adopt AI in Various Business Contexts

1. Introduction

Historically, artificial intelligence (AI) primarily served as tools for human tasks. However, technological advancements have propelled AI beyond its traditional role, allowing it to take on tasks previously reserved for humans. This shift represents a significant change, with AI evolving from a mere tool to a decision-maker in various fields like advertising (Lee & Cho, 2020), hiring (Kelan, 2023), healthcare (Wang et al., 2022), investments (Zhang et al., 2021), sales forecasting (Kawaguchi, 2021), and law enforcement (Berk, 2021). The rapid advancement of AI has led to a nuanced societal landscape characterized by contrasting attitudes. On one hand, there is “algorithm aversion,” where individuals hesitate to embrace AI despite its demonstrated superiority (Langer & Landers, 2021; Parent-Rocheleau & Parker, 2021). Conversely, some individuals show “algorithm appreciation,” as seen in studies like Logg et al.’s (2019). To address these contrasting views, scholars have conducted systematic reviews, focusing on AI features, user traits, and task specifics (Gerlich, 2023; Kelly et al., 2023; Ismatullaev & Kim, 2024).

However, an overlooked aspect in understanding individuals’ willingness to adopt AI is the influence of others. Given our social nature, human decision-making often succumbs to the sway of others’ attitudes or actions, known as social influence (Sammut & Bauer, 2021). Social psychologists have extensively studied social influence in various contexts, from simple laboratory tasks (Shank et al., 2019) to complex real-world decisions like pro-

environmental behaviors (Chen, Wan, & Yang, 2022), COVID-19 preventive measures (Chen et al., 2021), healthy eating (Ragelienė & Grønhoj, 2020), and online products rating (Sridhar & Srinivasan, 2012). Social influence emerges as a critical factor deserving further examination in the context of AI adoption. This study employs an intersubjective approach (Chiu et al., 2010; Wan et al., 2007) to systematically explore the effects of social influence on individuals' acceptance of AI across four empirical studies. Through this investigation, we seek to offer a novel contribution to the existing body of literature.

This paper advances AI acceptance research through an intersubjective approach in three key aspects. Firstly, this study shifts away from the predominant focus on user or AI characteristics and instead emphasizes individuals' perceptions of others' attitudes towards AI. A key aspect of this investigation involves assessing whether individuals accurately perceive others' attitudes toward AI within three important business domains: home purchasing, job interviews, and medical consultations. Secondly, it explores whether these perceptions, even if inaccurate, impact individuals' willingness to use AI. Lastly, the study integrates well-established psychological concepts—perceived risk and outcome efficacy—to elucidate the psychological mechanisms underlying how intersubjective perceptions of others' AI attitudes influence one's adoption inclination.

1.1 AI Acceptance

Understanding the factors shaping people's perception, acceptance, and utilization of AI is crucial amid the rapid development of AI. Kelly et al. (2023)'s research emphasizes transparency, compatibility, reliability, and simplification as key factors in bolstering users' attitudes and trust toward AI. Ismatullaev and Kim (2024) echo these findings, highlighting

the importance of addressing technological issues to alleviate concerns related to human factors like distrust and skepticism. Their reviews underscore the significance of factors such as perceived usefulness and performance expectancy in predicting user behavior across various industries, with trust and attitudes playing central roles in understanding AI acceptance.

While research typically focuses on user characteristics and AI attributes, recent studies by Gursoy et al. (2019) and Lin et al. (2020) introduce social influence into AI acceptance models. However, there is a notable gap in understanding how individuals perceive social influence in AI acceptance and how this shapes their own adoption and use of AI. This gap is crucial given that people's perceptions of others' behaviors or attitudes are often biased (Blanton et al., 2008), emphasizing the need for an approach that considers these biases in AI adoption.

1.2 Intersubjective Approach

The intersubjective approach rests on foundational principles, particularly relevant here:

- (1) Individuals conduct assessments and form perceptions of intersubjective reality, encompassing the prevalent behaviors and attitudes within their sociocultural contexts, with some of these perceptions diverging from personal values;
- (2) Individuals often act based on these perceptions, prioritizing them over personal values (Chiu et al., 2010). Notably, individuals' perceptions of others' attitudes or behaviors are prone to biases (Blanton et al., 2008). Empirical evidence supports these premises, revealing challenges in accurately assessing others' behaviors or attitudes in various contexts (Zhao & Epley, 2022; Chen, Wan, & Yang, 2022; Sparkman et al., 2022; Graupensperger et al., 2021).

Building on these insights, we hypothesize that individuals form intersubjective perceptions of others' attitudes toward AI—specifically, perceptions of prevailing positive or negative sentiments. We seek to assess the accuracy of these perceptions in reflecting actual attitudes toward AI. Although a systematic investigation into this area is lacking, circumstantial evidence suggests that individuals may perceive others as more negative toward AI than they are in reality, potentially influenced by stereotypes of AI users as less capable or intelligent (Diab et al., 2011; Eastwood et al., 2012). Consequently, we propose that individuals tend to underestimate the prevalence of positive attitudes toward AI (Hypothesis 1).

Furthermore, individuals not only hold biased perceptions but also use them to guide their behavior. For example, Chen, Wan, and Yang (2022) found a link between underestimating pro-environmental behavior and reduced engagement. Similarly, overestimating societal acceptance of negative behaviors increases personal engagement in those behaviors (Berry-Cabán et al., 2020). Building on this evidence, our hypothesis suggests that underestimating others' positive attitudes toward AI leads to reduced willingness to embrace AI use (Hypothesis 2).

1.3 Outcome Efficacy and Perceived Risk

Another focal point of this paper is unraveling how underestimated intersubjective perceptions of positive attitudes toward AI in others hinder individuals' intentions to utilize AI. We introduce two well-established concepts in AI adoption literature—outcome efficacy and perceived risk—to elucidate the adverse impact of underestimated intersubjective perception on AI usage intentions. Regarding outcome efficacy, studies suggest that

individuals' perceptions of others' behaviors or attitudes shape their beliefs about the effectiveness of their own actions. Kerr (1996) demonstrated, in the context of a public goods game, that the perception of widespread non-cooperation reduces individuals' beliefs in the effectiveness of their own cooperative behaviors. Similarly, in the pro-environmental domain, Fornara et al. (2011) showed that observing others' non-engagement in pro-environmental behaviors reduces individuals' efficacy for their own participation. We propose a similar scenario for AI usage, where underestimating others' positive attitudes toward AI diminishes individuals' perceptions of the effectiveness of using AI for problem-solving, thus reducing their willingness to use AI (Hypothesis 3a).

Additionally, we posit that perceived risk is another psychological factor explaining the impact of misperceptions of others' attitudes toward AI on adoption behavior. Scholars have observed that individuals' perceptions of others' behaviors or attitudes influence their judgments about the severity of risks. For example, Chen, Wan, and Yang (2022) found that individuals rely on perceptions of the prevalence of pro-environmental attitudes to inform their judgments about environmental risk. On the other hand, empirical evidence indicates that perceived risk significantly reduces individuals' willingness to adopt AI (Hasan et al., 2021; Pizam et al., 2024; Schwesig et al., 2023). Synthesizing these findings, we hypothesize that underestimating positive attitudes toward AI in others heightens individuals' perceived risk associated with AI adoption, thereby decreasing their willingness to adopt AI (Hypothesis 3b).

1.4 The Present Research

In this article, we present four empirical studies, comprising an online survey (Study 1)

and three in-lab experiments (Studies 2, 3a, and 3b). Adequate sample sizes have been determined using G*Power 3.1 (Faul et al., 2007). Studies 2, 3a, and 3b have been pre-registered on AsPredicted: Study 2 (https://aspredicted.org/R7N_P8W), Study 3a (https://aspredicted.org/HKW_F35), and Study 3b (https://aspredicted.org/W83_K38). All procedures, including measures, manipulations, and participant exclusions, are reported in these studies. In addition, prior to each survey or experiment, informed consent was obtained from all participants. Approval for this article has been granted by the ethics committee of the corresponding author's institution.

2. Study 1

2.1 Participants

Study 1 provides preliminary evidence for the hypotheses outlined in this paper, employing an online survey as the data collection method. A total of 320 questionnaires were distributed in China via the online platform “Credamo,” resulting in 309 participants after excluding 11 individuals who did not pass the attention tests (1. Which of the following cities is the capital of China? 2. What day of the week is it today?). The mean age of the participants was 30.67 ± 8.58 years, with 61.49% identifying as female. The post-hoc power analysis indicated that, with the given sample size, a paired-samples *t*-test ($d = -0.44$) achieved nearly 100% statistical power. Moreover, considering that the effect size in Study 1 was close to a medium effect size, subsequent studies adopted a medium effect size to calculate the required sample size.

2.2 Variables

The predictor variable in Study 1 was the perceived attitudes of others toward AI.

Initially, we assessed participants' personal attitudes toward AI (Cronbach's $\alpha = 0.82$) using four items adapted from Kim et al. (2023): (1) I am [1 = very negative, 7 = very positive] about AI in general; (2) I think AI is [1 = very bad, 7 = very good]; (3) I am [1 = very unfavorable; 7 = very favorable] about AI; (4) I [1 = strongly dislike, 7 = strongly like] AI. Subsequently, participants were prompted to estimate the attitudes that the majority of individuals participating in the survey held toward AI, thereby obtaining the predictor variable for Study 1, denoted as perceived others' attitudes (Cronbach's $\alpha = 0.85$). Higher mean scores indicate more positive attitudes toward AI. The subsequent analyses will investigate whether participants tend to underestimate others' positive attitudes toward AI. This will be achieved primarily by comparing participants' perceptions of others' attitudes toward AI with the actual attitudes that participants hold.

The outcome variable under investigation in Study 1 was participants' willingness to engage with AI. In Study 1, we selected three scenarios: home purchase, job interview, and online medical consultation (refer to Section 1 of the Supplementary Materials [SM] for comprehensive descriptions of these scenarios). Our choice was driven by practical considerations regarding the increasing integration of AI in critical business domains, including consumer decision-making (Ahn et al., 2022; Longoni & Cian, 2022; Xie et al., 2022), employee recruitment (Figuerola-Armijos et al., 2023; Lacroux & Martin-Lacroux, 2022), and healthcare services (Formosa et al., 2022; Reich et al., 2023). These areas are not only vital to business operations but also represent a growing body of scholarly research that highlights the importance of understanding AI's role in enhancing business practices. Thus, our focus on these application scenarios aligns with the broader business research agenda and

underscores their relevance in contemporary business contexts. Participants were tasked with expressing their willingness to engage with AI, represented by HomeFinder, JobFit, and HealthChat, respectively, in each scenario. This assessment employed a 7-point scale, ranging from 1 (indicative of strong reluctance) to 7 (indicative of strong willingness).

In Study 1, the mediating variables under consideration were perceived risk (Cronbach's $\alpha = 0.93$) and outcome efficacy (Cronbach's $\alpha = 0.74$), each assessed through four items. The questionnaire items gauging perceived risk were adapted from Hasan et al.'s (2021) research, focusing on individuals' perceptions regarding the risks and potential losses associated with AI utilization. Examples include statements like "Using AI is risky" and "Using AI will increase a lot of uncertainty." The items pertaining to outcome efficacy were adapted from Kim et al. (2021), exploring individuals' beliefs in the capacity of AI use to yield positive outcomes. Illustrative statements include "Using AI can help people get a lot of convenience" and "Using AI can improve people's quality of life." A comprehensive list of the questionnaire items is provided in Section 2 of the SM. Responses for all items were recorded on a 7-point scale, where 1 signifies complete disagreement, and 7 signifies complete agreement. Higher mean scores reflect elevated perceived risk and outcome efficacy among participants.

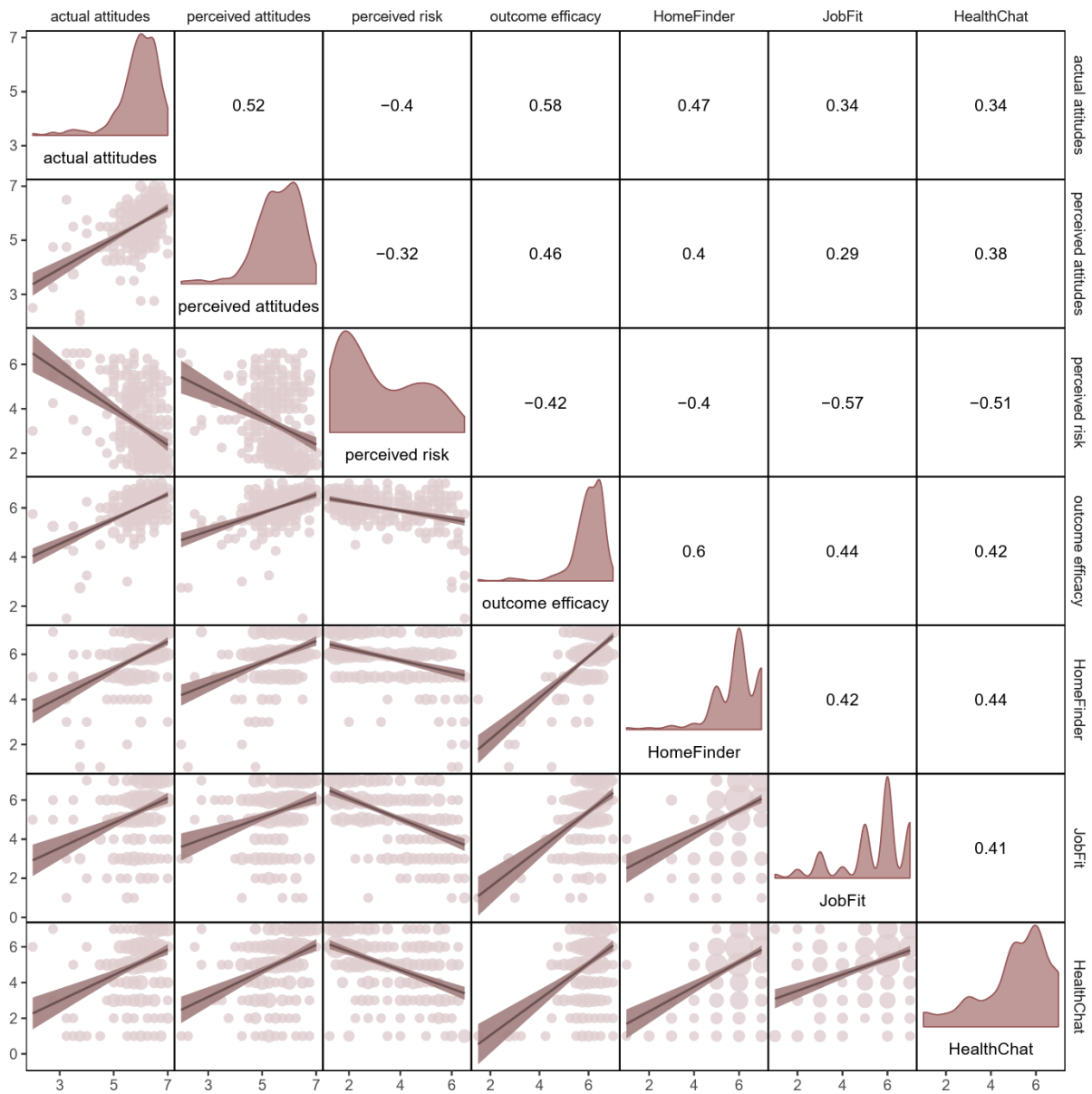
Furthermore, control variables for Study 1 encompassed gender (1 = male; 2 = female), age, education (1 = junior high school and below; 6 = doctoral), and annual income (1 = ¥50,000 and below; 7 = more than ¥500,000).

2.3 Results and Discussion

A common method bias test was initially conducted. We examined the presence of a

significant common method bias issue by comparing the fit indices of a one-factor model against a five-factor model (participants' actual attitudes, perceived others' attitudes, perceived risk, outcome efficacy, and willingness to use). The findings indicated that the fit indices of the five-factor model ($\chi^2 = 310.92$, $df = 142$, $\chi^2/df = 2.19$, CFI = 0.945, TLI = 0.934, RMSEA = 0.062) were significantly superior to those of the one-factor model ($\chi^2 = 1316.69$, $df = 152$, $\chi^2/df = 8.66$, CFI = 0.623, TLI = 0.576, RMSEA = 0.157) ($\Delta\chi^2 = 1005.78$, $\Delta df = 10$, $p < 0.001$). This suggests the absence of a notable common method bias problem. Figure 1 visually represents the correlation coefficients among the key variables in Study 1.

Figure 1
Correlation coefficients for key variables of Study 1.



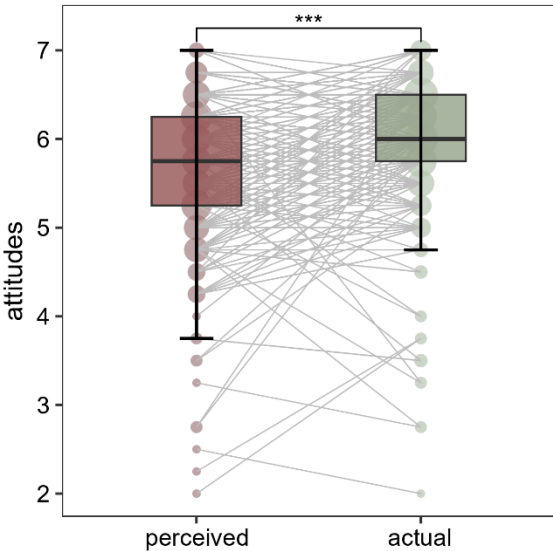
Note. $N = 309$, *** $p < 0.001$; in the lower-left part, the size of the circle denotes the number of participants; numbers in the upper-right part represent correlation coefficients, while the diagonal features a kernel density plot; the means and standard deviations for the variables are as follows: actual attitudes = 5.94 ± 0.75 , perceived attitudes = 5.60 ± 0.82 , perceived risk = 3.24 ± 1.54 , outcome efficacy = 6.02 ± 0.66 , HomeFinder = 5.93 ± 1.00 , JobFit = 5.43 ± 1.42 , HealthChat = 5.09 ± 1.58 .

2.3.1 Underestimation of Others’ Positive Attitudes toward AI Use

To assess whether participants underestimated the positive attitudes held by others toward AI, we conducted a paired-samples t -test comparing perceived others’ attitudes and participants’ actual attitudes. The results, as illustrated in Figure 2, clearly indicate that

participants’ estimates of others’ attitudes ($M = 5.60, SD = 0.82$) are significantly lower than the actual positive attitudes held by participants ($M = 5.94, SD = 0.75$) ($t(308) = -7.65, p < 0.001, d = -0.44, 95\%CI [-0.42, -0.25]$). This underestimation persisted across all four questionnaire items ($ts(308) = -6.64 \sim -4.90, ps < 0.001, ds = -0.38 \sim -0.28$). Therefore, Hypothesis 1 was substantiated.

Figure 2
Boxplot comparing participants’ perceptions of others’ attitudes toward AI and the actual attitudes held by participants.



Note. $N = 309, ***p < 0.001$. The figure includes circles representing the number of participants. The boxplot is detailed as follows: the horizontal line within the box represents the median, while the lower and upper edges of the box indicate the first quartile (Q1) and the third quartile (Q3), respectively, encompassing 50% of the data. The length of the box ($Q3 - Q1$) signifies the interquartile range (IQR), with lines extending beyond the box representing the upper limit ($Q3 + 1.5IQR$) and lower limit ($Q1 - 1.5IQR$). Points outside these limits are considered outliers. This description applies to subsequent boxplots without additional clarification.

We investigated the relationship between perceived others’ attitudes and individuals’ willingness to use AI across various scenarios. Utilizing perceived others’ attitudes as the

predictor variable and willingness to use AI in different contexts as the outcome variable, the regression results presented in Table 1 were derived while controlling for participants’ gender, age, education level, and annual income.

Firstly, participants’ perceived others’ attitudes exhibit a significant positive association with their willingness to use AI in home purchase ($B = 0.48, \beta = 0.40, t = 7.50, p < 0.001, 95\% \text{ CI } [0.35, 0.60]$), job interviews ($B = 0.47, \beta = 0.27, t = 5.00, p < 0.001, 95\% \text{ CI } [0.28, 0.65]$), and online medical consultation ($B = 0.66, \beta = 0.34, t = 6.62, p < 0.001, 95\% \text{ CI } [0.46, 0.85]$). This suggests that underestimating others’ positive attitudes toward AI is linked to participants’ diminished willingness to use AI, providing initial support for Hypothesis 2.

Secondly, the regression coefficients for control variables indicate a significant positive correlation between participants’ annual income and their willingness to use AI (home purchase: $B = 0.15, \beta = 0.17, t = 2.95, p = 0.003, 95\% \text{ CI } [0.05, 0.25]$; job interview: $B = 0.27, \beta = 0.22, t = 3.68, p < 0.001, 95\% \text{ CI } [0.13, 0.42]$; online medical consultation: $B = 0.24, \beta = 0.17, t = 3.06, p = 0.002, 95\% \text{ CI } [0.09, 0.39]$). Furthermore, more educated individuals were less receptive to using an AI interviewer in their job search ($B = -0.23, \beta = -0.11, t = -1.99, p = 0.048, 95\% \text{ CI } [-0.45, -0.002]$). Conversely, age significantly and positively predicted an individual’s willingness to use an AI doctor for an online consultation when sick ($B = 0.02, \beta = 0.12, t = 2.24, p = 0.026, 95\% \text{ CI } [0.003, 0.04]$). Nevertheless, gender-based differences were not discerned in any of the three scenarios.

Table 1

Regression analysis of perceived others’ attitudes on willingness to use AI across three

scenarios.

Variables	Home purchase		Job interview		Online medical consultation	
	β	95% CI	β	95% CI	β	95% CI
Predictor						
POA	0.40***	[0.35, 0.60]	0.27***	[0.28, 0.65]	0.34***	[0.46, 0.85]
Controls						
Gender	−0.04	[−0.30, 0.12]	−0.04	[−0.43, 0.18]	−0.05	[−0.49, 0.16]
Age	−0.07	[−0.02, 0.005]	0.05	[−0.01, 0.03]	0.12*	[0.003, 0.04]
Education	−0.04	[−0.22, 0.09]	−0.11*	[−0.45, −0.002]	0.01	[−0.22, 0.26]
Income	0.17**	[0.05, 0.25]	0.22***	[0.13, 0.42]	0.17**	[0.09, 0.39]
Adj- R^2		0.171		0.130		0.196

Note. $N = 309$; * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$; POA = perceived others’ attitudes.

While counterintuitive, the observation that older individuals exhibited a higher trust in AI aligns with existing research, which highlights that older (vs. younger) individuals tend to rely more on AI in mediation contexts due to reduced confidence in their own abilities (Ho et al., 2005). Conversely, the finding that better-educated individuals expressed greater resistance to AI in job searches appears incongruent with extant literature (Mahmud et al., 2022). We posit that this discrepancy may stem from the nature of jobs envisioned by highly educated individuals, which potentially involve more subjective assessments, thereby fostering a preference for human interviewers (Castelo et al., 2019).

2.3.2 Mediation Analyses

The investigation into how underestimation of others’ attitudes diminishes participants’ willingness to use AI delves into the mediating roles of perceived risk and outcome efficacy. Table 2 presents regression results utilizing perceived others’ attitudes as the predictor variable and perceived risk and outcome efficacy as the outcome variables. The findings indicate that a more positive estimation of others’ attitudes correlates with lower perceived

risk of AI ($B = -0.51, \beta = -0.27, t = -5.53, p < 0.001, 95\% \text{ CI } [-0.69, -0.33]$) and higher outcome efficacy ($B = 0.37, \beta = 0.47, t = 9.09, p < 0.001, 95\% \text{ CI } [0.29, 0.45]$).

Table 2

Regression analysis of perceived others’ attitudes on perceived risk and outcome efficacy.

Variables	Perceived risk		Outcome efficacy	
	β	95% CI	β	95% CI
Predictor				
POA	-0.27***	[-0.69, -0.33]	0.47***	[0.29, 0.45]
Controls				
Gender	0.09	[-0.02, 0.59]	0.01	[-0.12, 0.15]
Age	-0.14**	[-0.04, -0.01]	-0.06	[-0.01, 0.003]
Education	0.03	[-0.16, 0.28]	-0.002	[-0.10, 0.10]
Income	-0.36***	[-0.63, -0.34]	0.10	[-0.01, 0.12]
Adj- R^2	0.282		0.211	

Note. $N = 309$; ** $p < 0.01$, *** $p < 0.001$; POA = perceived others’ attitudes.

Subsequently, perceived risk and outcome efficacy were introduced into the regression model in Table 1, yielding the results in Table 3. A decrease in participants’ perceived risk across the three scenarios corresponds to heightened willingness to use AI (home purchase: $B = -0.10, \beta = -0.15, t = -2.69, p = 0.008, 95\% \text{ CI } [-0.17, -0.03]$; job interview: $B = -0.42, \beta = -0.45, t = -7.96, p < 0.001, 95\% \text{ CI } [-0.52, -0.31]$; online medical consultation: $B = -0.33, \beta = -0.32, t = -5.45, p < 0.001, 95\% \text{ CI } [-0.45, -0.21]$). Conversely, increased participant outcome efficacy aligns with strengthened willingness to use AI in all three scenarios (home purchase: $B = 0.72, \beta = 0.48, t = 8.93, p < 0.001, 95\% \text{ CI } [0.56, 0.88]$; job interview: $B = 0.52, \beta = 0.24, t = 4.40, p < 0.001, 95\% \text{ CI } [0.29, 0.75]$; online medical consultation: $B = 0.48, \beta = 0.20, t = 3.56, p < 0.001, 95\% \text{ CI } [0.21, 0.75]$).

Table 3
Regression analysis of perceived others’ attitudes, perceived risk and outcome efficacy on willingness to use.

Variables	Home purchase		Job interview		Online medical consultation	
	β	95% CI	β	95% CI	β	95% CI
Predictor						
POA	0.13*	[−0.04, 0.28]	0.04	[−0.12, 0.24]	0.16**	[0.11, 0.52]
Mediators						
Perceived risk	−0.15**	[−0.17, −0.03]	−0.45***	[−0.52, −0.31]	−0.32***	[−0.45, −0.21]
Outcome efficacy	0.48***	[0.56, 0.88]	0.24***	[0.29, 0.75]	0.20***	[0.21, 0.75]
Controls						
Gender	−0.03	[−0.25, 0.11]	−0.005	[−0.28, 0.25]	−0.02	[−0.38, 0.22]
Age	−0.06	[−0.02, 0.004]	−0.001	[−0.02, 0.02]	0.09	[−0.001, 0.04]
Education	−0.04	[−0.19, 0.08]	−0.10*	[−0.39, −0.01]	0.02	[−0.18, 0.26]
Income	0.07	[−0.03, 0.15]	0.03	[−0.09, 0.17]	0.04	[−0.10, 0.21]
Adj- R^2		0.397		0.372		0.329

Note. $N = 309$; * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$; POA = perceived others’ attitudes.

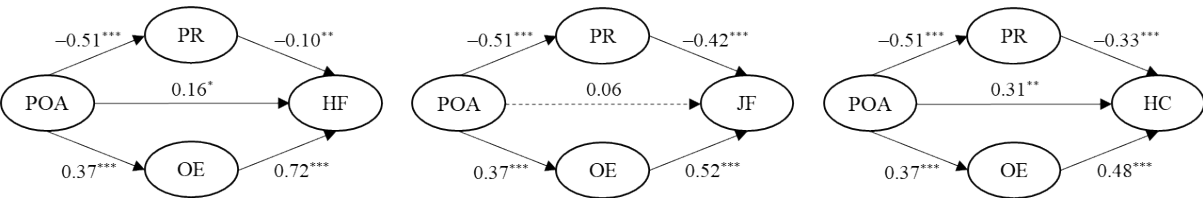
The obtained results suggest a mediation model where perceived risk and outcome efficacy play mediating roles in the influence of perceived others’ attitudes on one’s willingness to use AI, as illustrated in Figure 3. To rigorously analyze the effects, total, direct, and indirect effects were calculated in SPSS 25.0 using PROCESS 4.1 (Preacher & Hayes, 2004) with 5,000 bootstrap samples (Model = 4), and the results are presented in Table 4.

Regarding the indirect effects of perceived risk and outcome efficacy, although both were significant across all three scenarios, outcome efficacy demonstrated a considerably larger impact than perceived risk in home purchase. The difference between the two indirect effects was significantly greater than 0 (Effect = 0.22, $SE = 0.08$, 95% CI [0.09, 0.39]).

However, this difference was not significant in the job interview scenario (Effect = -0.02 , $SE = 0.08$, 95% CI $[-0.19, 0.14]$) and the online medical consultation scenario (Effect = 0.01 , $SE = 0.08$, 95% CI $[-0.17, 0.16]$). This indicates that, especially in the context of home purchase, it is predominantly outcome efficacy rather than perceived risk that shapes an individual's inclination to depend on AI. In summary, these results reveal nuanced variations across scenarios; nonetheless, in a broader sense, the aforementioned findings support Hypotheses 3a and 3b.

Figure 3

Mediation models in Study 1.



Note. $N = 309$, $^{*}p < 0.05$, $^{**}p < 0.01$, $^{***}p < 0.001$; POA = perceived others' attitudes, PR = perceived risk, OE = outcome efficacy, HF = HomeFinder, JF = JobFit, HC = HealthChat.

Table 4

Total, direct, and indirect effects across three scenarios.

Effects	Home purchase	Job interview	Online consultation
Total effect	0.48 [0.35, 0.60]	0.47 [0.28, 0.65]	0.66 [0.46, 0.85]
Direct effect	0.16 [0.04, 0.28]	0.06 [-0.12, 0.24]	0.31 [0.11, 0.52]
Indirect effect: perceived risk	0.05 [0.01, 0.10]	0.21 [0.11, 0.33]	0.17 [0.08, 0.27]
Indirect effect: outcome efficacy	0.27 [0.14, 0.44]	0.19 [0.07, 0.31]	0.18 [0.04, 0.30]
Outcome efficacy – perceived risk	0.22 [0.09, 0.39]	$-0.02 [-0.19, 0.14]$	0.01 [-0.17, 0.16]

Finally, we conducted a robustness test of the aforementioned results, specifically

examining whether they remain consistent after controlling for participants' actual attitudes toward AI, following the approach of Wan et al. (2007). The test revealed that the results remained largely unchanged, as detailed in Section 3 of the SM.

3. Study 2

The primary objective of Study 2 was to replicate the findings of Study 1 in a controlled laboratory setting, thereby establishing causal evidence for the identified relationship between intersubjective perception and willingness to use AI.

3.1 Participants

Methodologically, Study 2 employed various statistical analyses, such as multiple regression, independent samples *t*-tests, and paired samples *t*-tests. Sample sizes for Study 2 were determined based on the requirements of independent samples *t*-tests, known for their relatively higher sample size demands. Power analyses indicated that, assuming a medium effect size of $d = 0.50$ and a significance level of $\alpha = 0.05$, the independent samples *t*-test (two-tailed) necessitated a minimum of 210 participants to achieve a statistical power of 95%. The recruitment process yielded a total of 270 college students and 261 valid participants were included after excluding those who failed the attention tests. This comprised 130 participants in the underestimation cue group and 131 in the overestimation cue group. The mean age of the participants was 22.32 ± 2.13 years, with 62.84% identifying as female.

3.2 Variables

Study 2 employed a one-way, two-level between-participants design, wherein participants were subjected to random assignment in either the underestimation cue group

(coded as 1) or the overestimation cue group (coded as 0). The independent variable in Study 2 pertained to participants' categorical grouping. The dependent variable gauged participants' willingness to engage with AI across three distinct scenarios: home purchase, job interview, and online medical consultation. The study incorporated perceived risk and outcome efficacy as mediator variables, alongside gender and age as control variables. Notably, all measurement procedures aligned with the methodological framework established in Study 1.

3.3 Procedure

We commenced the experimental procedure in Study 2 by initially assessing participants' personal attitudes toward AI (Cronbach's $\alpha = 0.76$). Subsequently, participants were tasked with estimating the attitudes toward AI held by the majority of individuals partaking in the experiment (Cronbach's $\alpha = 0.81$). Following this initial phase, participants were randomly assigned to either the underestimation cue group or the overestimation cue group. Participants in the underestimation cue group were explicitly informed that, based on our prior research findings, they were likely to have underestimated the positive attitudes of the majority toward AI. Conversely, participants in the overestimation cue group were similarly informed that, according to our previous research, they were likely to have overestimated the positive attitudes of the majority toward AI. Subsequently, participants were given the opportunity to reassess the attitudes of others based on the provided cue (Cronbach's $\alpha = 0.97$).

The subsequent steps involved measuring participants' perceived risk (Cronbach's $\alpha = 0.82$) and outcome efficacy (Cronbach's $\alpha = 0.70$). After that, participants were presented with three distinct scenarios—home purchase, job interview, and online medical consultation—in a randomized order, consistent with the structure of Study 1. For each

scenario, participants expressed their willingness to engage with AI. Following the scenario assessments, participants answered an attention-testing question related to the cued message, and demographic variables (gender and age) were recorded.

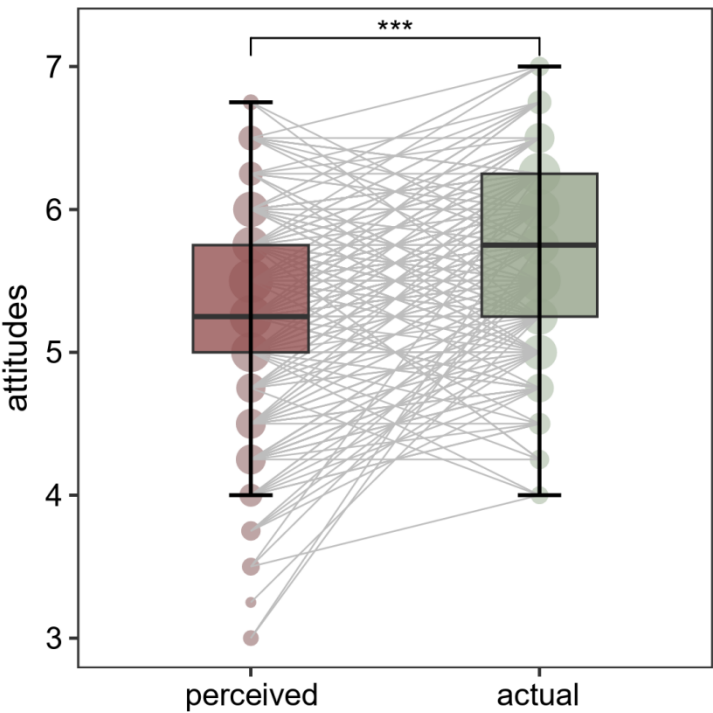
3.4 Results and Discussion

The initial examination of group randomization involved independent samples t -tests, indicating no significant differences between the underestimation cue group ($n = 130$, $M_{\text{perceived}} = 5.25$, $SD_{\text{perceived}} = 0.67$; $M_{\text{actual}} = 5.69$, $SD_{\text{actual}} = 0.62$) and the overestimation cue group ($n = 131$, $M_{\text{perceived}} = 5.22$, $SD_{\text{perceived}} = 0.72$; $M_{\text{actual}} = 5.64$, $SD_{\text{actual}} = 0.62$) in both their actual attitudes ($t(259) = 0.69$, $p = 0.489$, 95% CI $[-0.10, 0.20]$, $BF_{01} = 5.87$) and perceived others' attitudes ($t(259) = 0.35$, $p = 0.723$, 95% CI $[-0.14, 0.20]$, $BF_{01} = 6.94$) before the intervention. The Bayes Factors ($BF_{01} = 5.87$ and $BF_{01} = 6.94$, respectively) supported the absence of significant differences.

Subsequently, a paired-samples t -test assessed whether participants underestimated positive attitudes toward AI. The results, presented in Figure 4, revealed that participants' estimates of others' attitudes were significantly lower than the actual positive attitudes held by participants ($N = 261$, $M_{\text{perceived}} = 5.24$, $SD_{\text{perceived}} = 0.69$; $M_{\text{actual}} = 5.67$, $SD_{\text{actual}} = 0.62$) ($t(260) = -9.16$, $p < 0.001$, $d = -0.57$, 95% CI $[-0.52, -0.34]$). This reaffirmed the support for Hypothesis 1, aligning with the consistent pattern observed in Study 1.

Figure 4

Boxplot comparing participants' actual attitudes toward AI with their estimates of others' attitudes.



Note. $N = 261$, *** $p < 0.001$.

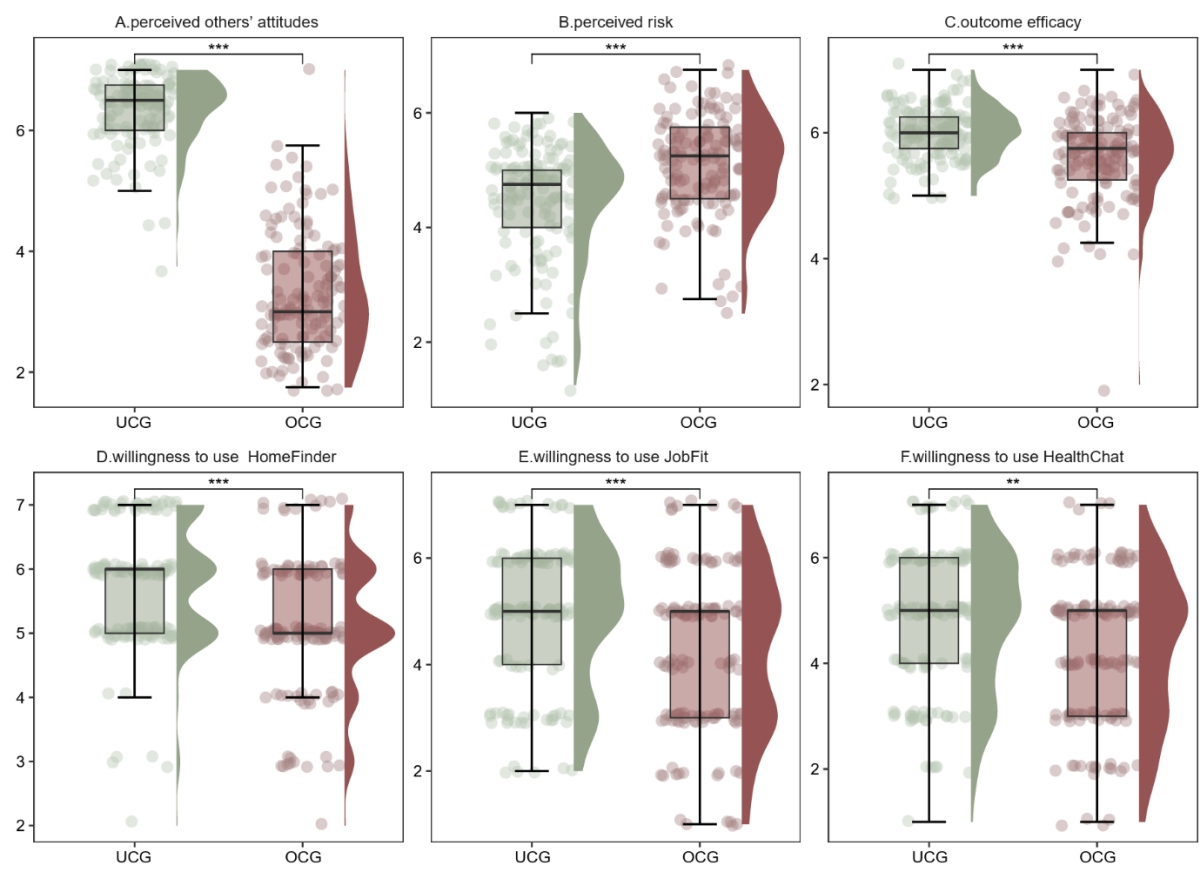
A manipulation validity test was conducted to assess the effectiveness of the experimental manipulation, specifically examining the significant difference between participants’ estimates of others’ attitudes in the two groups after the manipulation (refer to Figure 5). Results revealed that the underestimation cue group ($M = 6.37$, $SD = 0.58$) exhibited significantly higher estimates of others’ attitudes compared to the overestimation cue group ($M = 3.34$, $SD = 0.99$) ($t(209.13) = 30.17$, $p < 0.001$, $d = 3.73$, 95% CI [2.83, 3.23]), confirming the validity of the experimental manipulation in Study 2.

Subsequently, the impact of the experimental manipulation on perceived risk, outcome efficacy, and willingness to use AI was examined (see Figure 5). Independent samples t -tests revealed that (1) perceived risk was significantly lower in the underestimation cue group ($M = 4.43$, $SD = 1.04$) compared to the overestimation cue group ($M = 5.09$, $SD = 0.87$) ($t(259) =$

$-5.53, p < 0.001, d = 0.69, 95\% \text{ CI } [-0.89, -0.42]$), and outcome efficacy was significantly higher in the underestimation cue group ($M = 6.00, SD = 0.41$) than in the overestimation cue group ($M = 5.61, SD = 0.69$) ($t(212.75) = 5.59, p < 0.001, d = 0.69, 95\% \text{ CI } [0.25, 0.53]$); and (2) the underestimation cue group exhibited higher willingness to use AI in home purchases, job interviews, and online medical consultation (HomeFinder: $M = 5.72, SD = 1.00$; JobFit: $M = 4.95, SD = 1.37$; HealthChat: $M = 4.85, SD = 1.36$) compared to the overestimation cue group (HomeFinder: $M = 5.11, SD = 1.05$; JobFit: $M = 4.26, SD = 1.56$; HealthChat: $M = 4.27, SD = 1.46$) (HomeFinder: $t(259) = 4.80, p < 0.001, d = 0.60, 95\% \text{ CI } [0.36, 0.86]$; JobFit: $t(255.35) = 3.81, p < 0.001, d = 0.47, 95\% \text{ CI } [0.34, 1.05]$; HealthChat: $t(259) = 3.32, p = 0.001, d = 0.41, 95\% \text{ CI } [0.24, 0.92]$). These findings suggest that individuals' perceptions of others' attitudes influence perceived risk, outcome efficacy, and willingness to use AI across various scenarios. Specifically, more positive perceptions of others' attitudes correlate with lower perceived risk, higher outcome efficacy, and increased willingness to use AI.

Figure 5

Comparison of the differences between the two groups on key variables.



Note. $N = 261$, $**p < 0.01$, $***p < 0.001$; UCG = underestimation cue group, OCG = overestimation cue group.

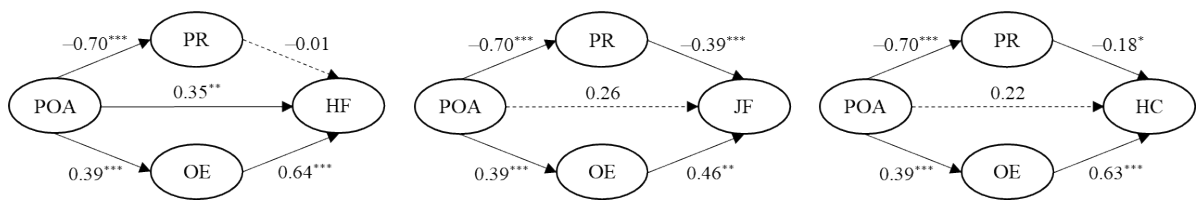
Furthermore, we conducted a comprehensive analysis using PROCESS 4.1 (Preacher & Hayes, 2004) with 5,000 bootstrap samples (Model = 4) to replicate the analytic process employed in Study 1. The results from Study 2 closely mirrored those of Study 1, reinforcing the robustness of our findings. Figure 6 illustrates the mediation models corresponding to Study 2, while Table 5 delineates the total, direct, and indirect effects across the three scenarios.

Specifically, we observed that (1) outcome efficacy continued to exert substantial indirect effects in Study 2 for home purchase (Effect = 0.25, $SE = 0.06$, 95% CI [0.14, 0.37]), job interviews (Effect = 0.18, $SE = 0.07$, 95% CI [0.05, 0.32]), and online medical consultation (Effect = 0.25, $SE = 0.07$, 95% CI [0.12, 0.38]). (2) In the Study 2 home

purchase scenario, the indirect effect of outcome efficacy remained significantly larger than perceived risk (Effect = 0.24, *SE* = 0.08, 95% CI [0.09, 0.40]); however, this difference was not significant in the job interview (Effect = −0.09, *SE* = 0.12, 95% CI [−0.33, 0.13]) and online counseling (Effect = 0.12, *SE* = 0.10, 95% CI [−0.07, 0.31]) scenarios. (3) Unlike in Study 1, the indirect effect of perceived risk in home purchase scenario was not significant (Effect = 0.01, *SE* = 0.04, 95% CI [−0.08, 0.09]).

Figure 6

Mediation models in Study 2.



Note. *N* = 261, **p* < 0.05, ***p* < 0.01, ****p* < 0.001; POA = perceived others' attitudes (overestimation cue group = 0; underestimation cue group = 1), PR = perceived risk, OE = outcome efficacy, HF = HomeFinder, JF = JobFit, HC = HealthChat.

Table 5

Total, direct, and indirect effects across three scenarios.

Effects	Home purchase	Job interview	Online consultation
Total effect	0.61 [0.35, 0.86]	0.71 [0.35, 1.07]	0.59 [0.24, 0.94]
Direct effect	0.35 [0.09, 0.61]	0.26 [−0.12, 0.64]	0.22 [−0.15, 0.58]
Indirect effect: perceived risk	0.01 [−0.08, 0.09]	0.27 [0.13, 0.44]	0.13 [0.002, 0.26]
Indirect effect: outcome efficacy	0.25 [0.14, 0.37]	0.18 [0.05, 0.32]	0.25 [0.12, 0.38]
Outcome efficacy – perceived risk	0.24 [0.09, 0.40]	−0.09 [−0.33, 0.13]	0.12 [−0.07, 0.31]

Similar to Study 1, the outcomes remained largely consistent subsequent to the inclusion of participants' actual attitudes toward AI as a covariate in the mediation model, as delineated

in Section 4 of the SM.

In summary, Study 2, given its experimental design, offers causal insights supporting Hypotheses 1 and 2 in this paper. However, in the case of Hypotheses 3a and 3b, the study exclusively furnishes evidence affirming the “perceived others’ attitudes → perceived risk/outcome efficacy” pathway. Notably, it does not provide substantiation for the pathway “perceived risk/outcome efficacy → willingness to use AI,” a gap we intend to address in Study 3.

4. Study 3

While Study 2 presented experimental evidence highlighting the influence of intersubjective perception on the willingness to use AI, its experimental manipulation exclusively focused on participants’ perception of others’ attitudes. Consequently, it confines its examination to the causal relationships of “perceived others” attitudes → perceived risk/outcome efficacy” and “perceived others’ attitudes → willingness to use AI” in isolation. To address this limitation, Study 3 employed the manipulation-of-mediation-as-a-moderator (MMM) design proposed by Ge (2023) to rigorously scrutinize the mediating effects of perceived risk and outcome efficacy. This approach aims to establish a more robust basis for causal inference pertaining to “perceived risk/outcome efficacy → intentions to use AI” and “perceived others’ attitudes → intentions to use AI.” Study 3 comprised two distinct experiments, Study 3a and Study 3b, each strategically manipulating participants’ perceived others’ attitudes alongside perceived risk or outcome efficacy, respectively, to fulfill the aforementioned objectives.

4.1 Participants

Studies 3a and 3b were conducted with a uniform experimental design, employing a 2×2 between-participants ANOVA. The applied statistical methods encompassed independent samples *t*-tests, paired samples *t*-tests, and ANOVA. Sample size determination was predicated on the ANOVA, which typically demands a substantial participant pool. Power analyses conducted with G*Power 3.1 indicated that, considering a medium effect size of $f = 0.25$ and a significance level of $\alpha = 0.05$, a 2×2 between-participants ANOVA ($df = 1$; number of groups = 4) necessitated a minimum of 210 participants to achieve a statistical test power of 95%. Following the exclusion of 13 participants who did not meet the attention test criteria, Study 3a enlisted 247 college students, distributed as follows: 62 in the underestimation cue-low risk group, 63 in the underestimation cue-control group, 61 in the overestimation cue-low risk group, and 61 in the overestimation cue-control group. The mean age of the participants was 22.20 ± 1.79 years, with 58.70% being female. In Study 3b, an initial recruitment involved 260 college students, and following the exclusion of 7 participants, a total of 253 valid participants constituted the final sample. The distribution encompassed 64 participants allocated to the underestimation cue-high efficacy group, 63 to the underestimation cue-control group, 62 to the overestimation cue-high efficacy group, and 64 to the overestimation cue-control group. The mean age of participants in Experiment 3b was 21.92 ± 1.79 years, with 62.85% being female.

4.2 Variables

Study 3a employed a 2 (underestimation cue vs. overestimation cue) $\times 2$ (low risk vs. control) between-participants design. The manipulation of participants' perceived attitudes

toward others paralleled that in Study 2. For the manipulation of perceived risk, participants were presented with risk information. Specifically, in Study 3a, perceived risk manipulation hinged on whether participants received information regarding the rate of misdiagnosis by AI physicians, designating a low-risk group if the corresponding information was presented and a control group if the information was omitted. In Study 3b, following a 2 (underestimation cue vs. overestimation cue) $\times$ 2 (high efficacy vs. control) between-participants design, outcome efficacy was manipulated based on whether participants were exposed to patient feedback regarding the effectiveness of AI doctor recommendations. Analogously, a high-performance group received the pertinent information, while a control group did not. The dependent variable for both studies was participants' willingness to use an AI doctor for online consultation, with perceived risk and outcome efficacy serving as mediating variables, respectively.

4.3 Procedure

In Studies 3a and 3b, our initial step involved assessing participants' attitudes toward AI (Cronbach's $\alpha = 0.80$ in both studies) and soliciting their estimations of the attitudes of the majority of fellow participants in the experiment toward AI (Cronbach's $\alpha = 0.81$ in both studies). Subsequently, participants were randomly assigned to either an underestimation cue group or an overestimation cue group, allowing them to reassess others' attitudes based on the provided cues (Study 3a: Cronbach's $\alpha = 0.98$; Study 3b: Cronbach's $\alpha = 0.97$).

In Study 3a, participants were further randomized into a low-risk group and a control group. Those in the low-risk group were informed that the official published report indicated a very low misdiagnosis rate for the AI doctor HealthChat's performance over the past six

months. In contrast, participants in the control group were presented with information unrelated to misdiagnosis rates. Subsequently, participants evaluated the perceived risk of using an AI doctor for online consultation (“How risky do you think it is to make an online consultation to the AI doctor HealthChat?”; 1 = no risk at all, 7 = very high risk) and indicated their willingness to use the AI doctor.

In Study 3b, participants were assigned to a high-efficacy group or a control group. Participants in the high-efficacy group were informed that patient feedback revealed the effectiveness of AI doctor HealthChat’s advice to the majority of patients over the past six months. Conversely, the control group received irrelevant information. Participants in this study assessed the perceived effectiveness of the advice provided by the AI doctor (“To what extent do you think the advice provided by the AI doctor HealthChat is effective?”; 1 = not effective at all, 7 = completely effective) and expressed their willingness to use the AI doctor.

At the conclusion of both experiments, participants were tasked with responding to questions pertaining to cued messages as an attention test. Additionally, we collected demographic information including the gender and age of the participants. It is noteworthy that the measures concerning participants’ actual and perceived others’ attitudes toward AI, as well as the measures of willingness to use AI, remained consistent with those employed in Study 1.

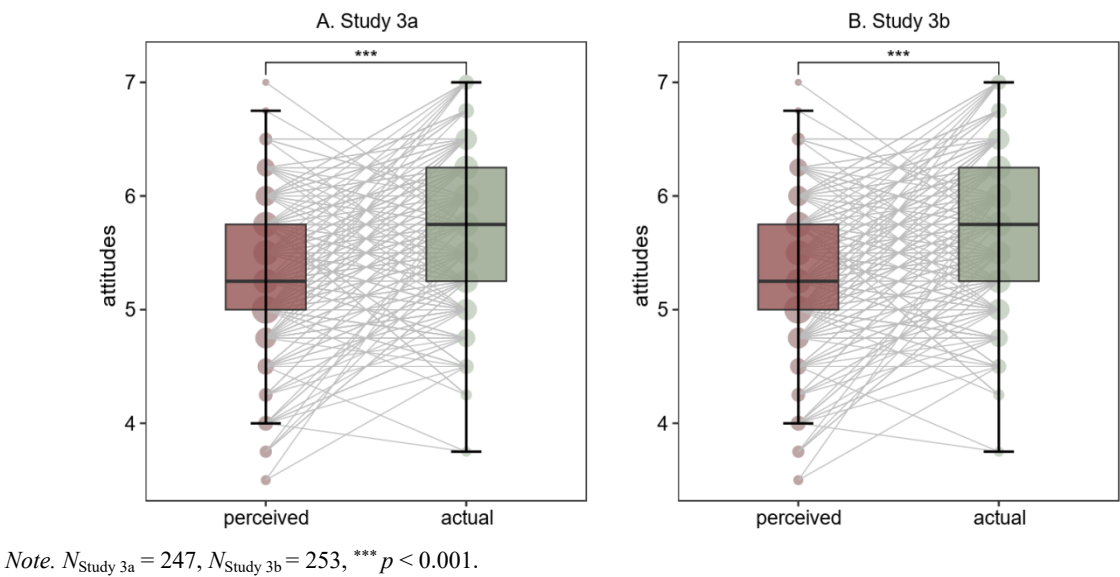
4.4 Results and Discussion

The randomization of the groupings underwent initial scrutiny, with ANOVAs conducted on the pre-intervention participants’ actual attitudes and perceived others’ attitudes in Study 3a. Results indicated non-significant main effects for manipulation of message cues (actual

attitudes: $F(1, 243) = 0.51, p = 0.477$; perceived others' attitudes: $F(1, 243) = 0.35, p = 0.554$) and risk manipulation (actual attitudes: $F(1, 243) = 0.11, p = 0.741$; perceived others' attitudes: $F(1, 243) = 0.20, p = 0.654$). Additionally, the interaction effect between the two factors was also found to be non-significant (actual attitudes: $F(1, 243) = 1.13, p = 0.290$; perceived others' attitudes: $F(1, 243) = 0.05, p = 0.832$). Similar results were obtained in Study 3b for manipulation of message cues (actual attitudes: $F(1, 249) = 0.03, p = 0.866$; perceived others' attitudes: $F(1, 249) = 0.15, p = 0.700$) and efficacy manipulation (actual attitudes: $F(1, 249) = 0.11, p = 0.737$; perceived others' attitudes: $F(1, 249) = 0.09, p = 0.767$), with the interaction effect also not reaching significance (actual own attitudes: $F(1, 249) = 1.01, p = 0.317$; perceived others' attitudes: $F(1, 249) = 2.22, p = 0.138$). Secondly, paired-samples t -tests revealed that participants in both experiments significantly underestimated the positive attitudes that others actually held toward AI (Study 3a: $N = 247, M_{\text{perceived}} = 5.28, SD_{\text{perceived}} = 0.63; M_{\text{actual}} = 5.70, SD_{\text{actual}} = 0.66, t(246) = -8.97, p < 0.001, d = -0.57, 95\% \text{ CI } [-0.51, -0.33]$; Study 3b: $N = 253, M_{\text{perceived}} = 5.19, SD_{\text{perceived}} = 0.70; M_{\text{actual}} = 5.65, SD_{\text{actual}} = 0.67, t(252) = -9.36, p < 0.001, d = -0.59, 95\% \text{ CI } [-0.56, -0.37]$), as illustrated in Figure 7.

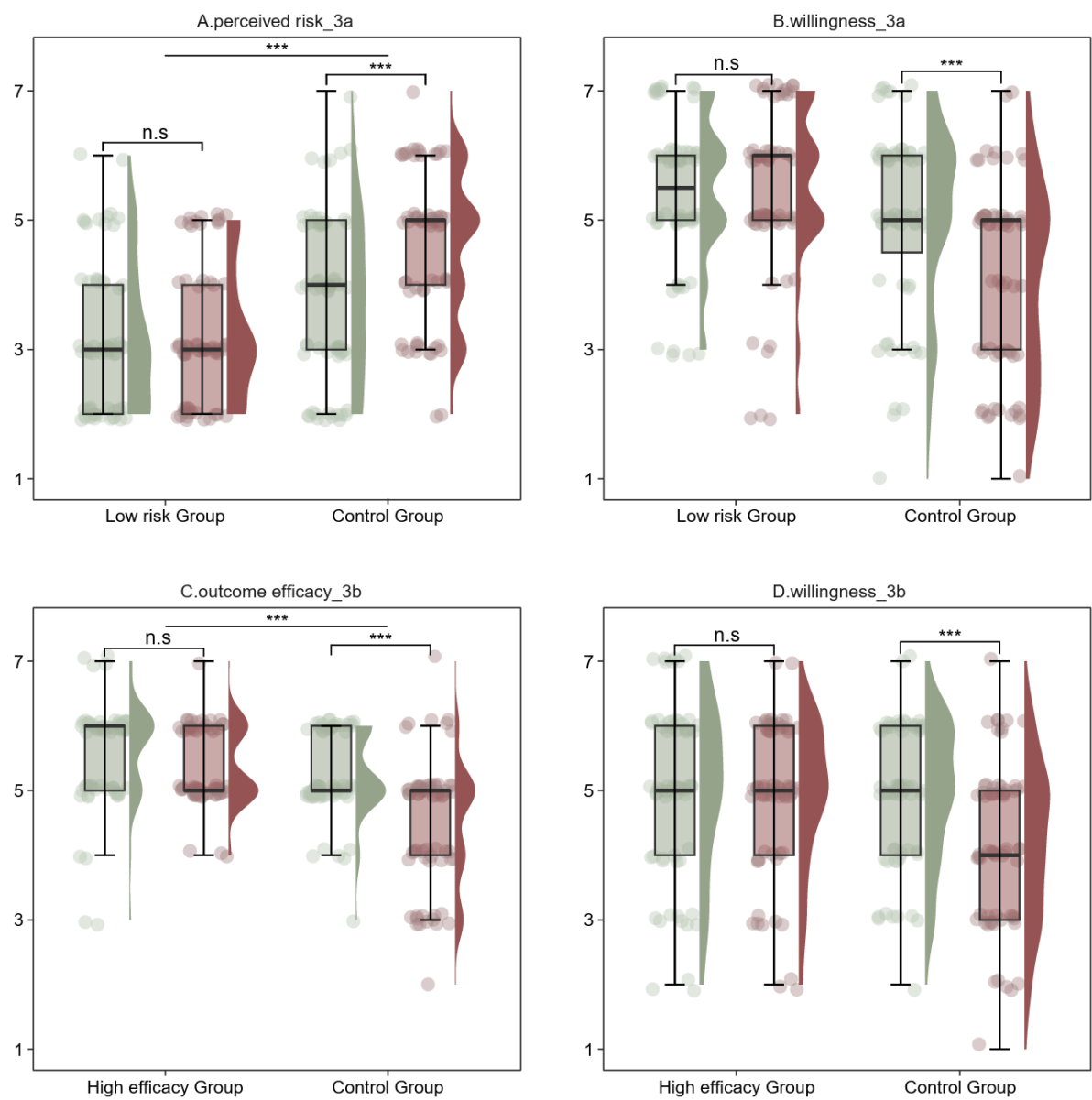
Figure 7

Boxplots comparing pre-intervention perceived others' attitudes and actual attitudes for Study 3a and Study 3b.



Subsequently, independent samples t -tests were employed to assess whether the experimental manipulation induced changes in participants' perceptions of others' attitudes. The results demonstrated a significant difference, with the underestimation cue group ($n = 125$, $M = 6.29$, $SD = 0.56$) exhibiting markedly higher estimates of others' attitudes than the overestimation cue group ($n = 122$, $M = 3.22$, $SD = 0.90$) ($t(201.78) = 32.15$, $p < 0.001$, $d = 4.10$, 95% CI [2.88, 3.26]). This implies the effectiveness of Study 3a's manipulation of others' attitudes. Similar outcomes were observed in Study 3b: the underestimation cue group ($n = 127$, $M = 6.27$, $SD = 0.66$) estimated others' attitudes significantly higher than the overestimation cue group ($n = 126$, $M = 3.26$, $SD = 0.93$) ($t(225.71) = 29.60$, $p < 0.001$, $d = 3.73$, 95% CI [2.81, 3.21]).

Figure 8
Manipulation and mediation tests for Studies 3a and 3b.



Note. $N_{\text{Study 3a}} = 247$, $N_{\text{Study 3b}} = 253$, *** $p < 0.001$, n.s. = non-significant; the underestimation cue group is represented by the color green, while the overestimation cue group is denoted by the color red.

In Study 3a, following Ge’s (2023) methodological guidelines for the MMM design, we conducted a 2×2 between-participants ANOVA to assess variances in perceived risk among the four groups (refer to Figure 8A), with a focus on two primary outcomes: Firstly, we examined the significance of the main effect of the risk manipulation to gauge its effectiveness; secondly, we evaluated whether a notable difference existed in perceived risk

between the cued underestimation and cued overestimation groups under the control condition, aiming to ascertain a potential causal link between the independent variable and the unmanipulated mediator. The findings revealed (1) a noteworthy main effect of risk manipulation (low-risk group: $n = 123$, $M = 3.22$, $SD = 1.11$; control group: $n = 124$, $M = 4.17$, $SD = 1.30$; $F(1, 243) = 40.48$, $p < 0.001$, $\eta_p^2 = 0.14$), indicating the efficacy of the risk manipulation in Study 3a; additionally, both the directional cueing main effect ($F(1, 243) = 7.97$, $p = 0.005$, $\eta_p^2 = 0.03$) and the interaction effect between the two ($F(1, 243) = 6.12$, $p = 0.014$, $\eta_p^2 = 0.02$) were statistically significant. (2) Further analysis revealed that in the control condition, perceived risk in the cued underestimation group ($n = 63$, $M = 3.78$, $SD = 1.31$) was significantly lower than that in the cued overestimation group ($n = 61$, $M = 4.57$, $SD = 1.18$) ($F(1, 243) = 14.08$, $p < 0.001$, $\eta_p^2 = 0.05$), suggesting that cueing direction does influence individuals' perceived risk of an AI doctor; however, this difference was not significant in the low-risk condition ($F(1, 243) = 0.06$, $p = 0.806$).

Subsequently, another 2×2 between-participants ANOVA was deployed to scrutinize differences in willingness to use AI across the four participant groups (see Figure 8B). Significant main effects were discerned for both manipulation of message cues ($F(1, 243) = 6.71$, $p = 0.010$, $\eta_p^2 = 0.03$) and risk manipulation ($F(1, 243) = 23.83$, $p < 0.001$, $\eta_p^2 = 0.09$), accompanied by a consequential interaction effect ($F(1, 243) = 7.43$, $p = 0.007$, $\eta_p^2 = 0.03$). Further dissection through simple effects analyses unveiled that (1) under low-risk conditions, the distinction between the underestimation cue group ($n = 62$, $M = 5.42$, $SD = 1.14$) and the overestimation cue group ($n = 61$, $M = 5.44$, $SD = 1.31$) was statistically nonsignificant ($F(1, 243) = 0.01$, $p = 0.924$); (2) under control conditions, the

underestimation cue group ($n = 63$, $M = 5.05$, $SD = 1.44$) exhibited markedly greater willingness to use AI compared to the overestimation cue group ($n = 61$, $M = 4.13$, $SD = 1.50$) ($F(1, 243) = 14.18$, $p < 0.001$, $\eta_p^2 = 0.06$). This signifies that the manipulation of perceived risk nullified the impact of message cues regarding others' attitudes on willingness to use AI.

Parallel analyses were conducted on the data from Study 3b. Initially, a 2×2 between-participants ANOVA was employed to assess variations in outcome efficacy across the four groups (refer to Figure 8C). The findings indicated that (1) the main effect of the efficacy manipulation was significant (high-efficacy group: $n = 126$, $M = 5.48$, $SD = 0.71$; control group: $n = 127$, $M = 4.83$, $SD = 0.92$; $F(1, 249) = 43.13$, $p < 0.001$, $\eta_p^2 = 0.15$), affirming the effectiveness of Study 3b's manipulation of outcome efficacy; additionally, both the main effect of cue direction ($F(1, 249) = 25.55$, $p < 0.001$, $\eta_p^2 = 0.09$) and the interaction effect between the two were statistically significant ($F(1, 249) = 6.06$, $p = 0.015$, $\eta_p^2 = 0.02$). (2) Further analysis revealed that in the control condition, the outcome efficacy of the cued underestimation group ($n = 63$, $M = 5.21$, $SD = 0.65$) was significantly higher than that of the cued overestimation group ($n = 64$, $M = 4.47$, $SD = 1.01$) ($F(1, 249) = 28.35$, $p < 0.001$, $\eta_p^2 = 0.10$), indicating that cueing direction influences individuals' perceptions of outcome efficacy; however, no significant differences were observed in the high-performance condition ($F(1, 249) = 3.35$, $p = 0.068$).

Moreover, a 2×2 between-participants ANOVA on willingness to use AI across the four groups revealed significant main effects for both manipulation of message cues ($F(1, 249) = 7.47$, $p = 0.007$, $\eta_p^2 = 0.03$) and efficacy manipulation ($F(1, 249) = 5.84$, $p = 0.016$, $\eta_p^2 =$

0.02), accompanied by a consequential interaction effect ($F(1, 249) = 4.88, p = 0.028, \eta_p^2 = 0.02$). Further simple effects analyses disclosed that (1) under high-efficacy conditions, the distinction between the underestimation cue group ($n = 64, M = 4.98, SD = 1.34$) and the overestimation cue group ($n = 62, M = 4.90, SD = 1.18$) was statistically nonsignificant ($F(1, 249) = 0.14, p = 0.712$); (2) under control conditions, the underestimation cue group ($n = 63, M = 4.95, SD = 1.11$) demonstrated markedly greater willingness to use compared to the overestimation cue group ($n = 64, M = 4.19, SD = 1.27$) ($F(1, 249) = 12.27, p < 0.001, \eta_p^2 = 0.05$). These findings are depicted in Panel F of Figure 8. This denotes that the manipulation of outcome efficacy nullified the influence of cue information on willingness to use.

In summary, the manipulation of participants' estimations regarding others' attitudes exerts an influence on the intention to use AI by modifying perceived risk/outcome efficacy. Nevertheless, when participants are furnished with information pertaining to risk/outcome efficacy, the pathway "perceived others' attitudes $\rightarrow$ perceived risk/outcome efficacy" exhibits partial impedance. Consequently, the willingness to use AI in both groups becomes impervious to the effects of the independent variable, yielding no statistically significant difference. This observation underscores the mediating role of perceived risk/outcome efficacy between perceived others' attitudes and the inclination to use AI.

6. General Discussion

This study adopts an approach that examines how individuals perceive others' attitudes toward AI and how these perceptions influence their acceptance of AI across various business contexts. A key idea here is that people's perceptions of others' attitudes may not always

match reality, leading to decisions made without complete information. The empirical findings strongly support this idea. Specifically, individuals consistently tend to underestimate their peers' positive attitudes toward AI, which in turn reduces their willingness to engage with AI. Additionally, the research shows that when individuals underestimate others' positive attitudes, they also perceive AI as less effective while seeing more risks associated with its use. These combined factors negatively impact people's willingness to embrace AI.

6.1 Theoretical Implications

This study significantly advances our understanding of AI acceptance, departing from the typical focus on AI features and user traits. Prior research has explored various aspects such as AI transparency (Bodimani, 2024), decision complexity (Stein et al., 2020), and psychological effects like loss of control (Dietvorst et al., 2018) and impacts on self-esteem (Lee, 2018). Moreover, tasks demanding high cognitive and emotional engagement have been linked to reduced AI usage (Bigman & Gray, 2018). This paper introduces a novel perspective, elucidating that a biased perception—specifically, the underestimation of others' positive attitudes toward AI—constitutes a crucial factor influencing individuals' inclination to adopt AI. Biased perceptions of others' attitudes and their consequential impact on individual behavior align with prevalent patterns observed across various social phenomena (see Blanton et al., 2008 for a review), including pro-environmental behavior (Sparkman et al., 2022), alcoholism (Amialchuk & Sapci, 2021), safe driving (Geber et al., 2021), and eating out after COVID-19 (Palacios et al., 2022).

The phenomenon wherein individuals tend to underestimate the positive attitudes of

others toward AI may find explanation through the theory of pluralistic ignorance (Sargent & Newman, 2021). In the context of AI acceptance, pluralistic ignorance posits that individuals may harbor favorable sentiments toward AI privately, yet erroneously perceive others as holding negative attitudes. This misconception often stems from the stereotype that AI users are perceived as less competent or intelligent (Diab et al., 2011; Eastwood et al., 2012). Despite having positive experiences or views regarding AI, individuals may refrain from vocalizing these sentiments if they believe that the prevailing group opinion is negative. Consequently, this reticence to express favorable attitudes can perpetuate a mistaken belief that the majority harbors skepticism or negativity toward AI, thereby fostering a cycle of pluralistic ignorance.

Notably, we conducted a mini meta-analysis of the findings presented in this article (see Section 5 of the SM for details), which revealed medium effect sizes for both the underestimation of others' positive attitudes toward AI and its subsequent impact on individual acceptance. In contrast, previous studies on similar topics often report small effect sizes. This observation regarding effect size may be linked to the relatively recent emergence of AI, despite significant advancements in the field. Unlike well-established phenomena such as pro-environmental behaviors, alcohol consumption, or driving habits, AI is still a nascent concept. Its limited integration into individuals' daily lives complicates the accurate assessment of others' true attitudes toward AI. As a result, individuals' estimations of others' positive attitudes toward AI often diverge significantly from reality.

Furthermore, the relatively limited commonality of AI applications in daily life contributes to the absence of well-defined norms governing its usage. The dearth of

established rules on AI utilization, coupled with the high level of unfamiliarity surrounding its integration into everyday activities, aligns with research on social norms. In such contexts marked by heightened unfamiliarity, the perceived attitudes or behaviors of the majority hold greater sway over an individual's own behavior (McDonald & Crandall, 2015). This phenomenon is evident in our study, where the perception of others' positive attitudes toward AI exerts a medium effect size on an individual's acceptance of AI.

6.2 Practical Implications

Researchers have recognized the myriad advantages of promoting the adoption of AI in business contexts and have proposed various strategies to achieve this objective. Some approaches include enlightening potential algorithm users about the learning capabilities of algorithms, emphasizing their ability to improve based on past errors (Berger et al., 2021; Reich et al., 2023). Additionally, strategies involve users understanding the limitations of their own capabilities (Filiz et al., 2021), comprehending the decision-making process or rationale behind AI (Litterscheidt & Streich, 2020; Sharan & Romano, 2020), or providing users with a sense of control or choice in AI utilization (Dietvorst et al., 2018). A systematic review by Burton et al. (2020) has distilled five methods to promote AI adoption, including the development of algorithmic literacy, human-in-the-loop decision-making, context-specific behavioral design, engaging intuition, and aiding ecological rationality. Significantly, the context-specific behavioral design aligns with the outcomes of this study. For instance, Alexander et al. (2018) suggest leveraging perceived social consensus to improve AI acceptance. This paper, extending beyond, underscores the relevance of its findings, revealing biased perceptions of social consensus regarding AI use, predominantly leaning

towards the lower end. This sets the stage for intervention through the well-established social psychology strategy: the social norms approach.

The social norms approach aims to induce positive behavioral changes by rectifying individuals' biased perceptions. This paper's findings affirm the existence of misperceptions about others' positive attitudes towards AI, suggesting the potential applicability of the social norms approach to promote the AI adoption. While this paper does not include a longitudinal intervention study, its findings suggest the alignment of AI adoption with the prerequisites of the social norms approach. The existing body of research on this approach offers a degree of optimism regarding its effectiveness in promoting AI use.

The foundational mechanism of the social norms approach lies in providing the target group with accurate information about others' behaviors and attitudes. Researchers have delineated two methods to achieve this objective. The first method involves delivering normative messages top-down by authoritative figures. The second method entails enabling the target group to acquire authentic information through norm dialogues with similar peers (Shank et al., 2019). Both methods have proven effective in inducing behavioral change within the target group. In the context of AI acceptance, policymakers can employ judicious methods to extensively survey people's genuine attitudes toward AI and disseminate the survey results through widely accepted channels. The findings of this paper indicate that this strategy can alter people's perceptions of others' positive attitudes toward AI, consequently enhancing their own intentions to use AI. Another viable approach involves organizing periodic communication sessions between existing and potential users of AI. As validated by Shank et al.'s (2019) investigation on cooperation, the straightforward approach of norm talks

often yields favorable outcomes.

6.3 Limitations and Future Directions

Before drawing conclusions, it is imperative to acknowledge several limitations and suggest avenues for future research. Primarily, the studies conducted for this paper were exclusively carried out in China, and the monocultural nature of the sample raises concerns about the generalizability of the conclusions to other cultural contexts. Notably, scholars have emphasized the influential role of culture in AI adoption (Duan et al., 2019). One pertinent dimension, cultural tightness, as conceptualized by Gelfand et al. (2011), is particularly relevant to this study. Cultural tightness refers to the degree of tolerance within a culture for deviations from normative behaviors. In a tight culture like China, individuals may be more motivated to accurately infer the behavior and attitudes of the majority to conform to societal norms and avoid potential consequences. Conversely, in a relatively loose culture, there might be less incentive for individuals to speculate about the thoughts or attitudes of others. While the hypothesis posits a potentially deeper cognitive bias toward positive attitudes toward AI in a loose culture, the limitation of a single-culture sample in this paper precluded empirical testing. Future cross-cultural research is encouraged to validate this perspective.

Secondly, all dependent variables in this paper pertained to verbally reported intentions to use AI. The well-documented intention-behavior gap (Conner & Norman, 2022) in the literature raises concerns about the applicability of the findings to actual AI use. Therefore, future research should aim to validate the results using objective behavioral records as the dependent variable. Furthermore, the use of university students as participants in the experiments, with the exception of non-student participants in Study 1, introduces a potential

limitation. Scholars in China have highlighted that individuals, whether they are college students (Studies 2, 3a, and 3b) or participants recruited through online platforms (Study 1), often exhibit elevated socio-economic statuses characterized by superior educational backgrounds and higher income levels than the broader general public (Chen, Yang, et al., 2022). Consequently, it is crucial to investigate whether the conclusions of this paper hold true for individuals with lower socio-economic statuses, necessitating further verification. Finally, while the three scenarios selected in this study hold significance in individuals' daily experiences, it is acknowledged that this selection may be somewhat arbitrary. Hence, the generalizability of the findings to other contexts remains uncertain. Future research endeavors are encouraged to validate our conclusions across a broader spectrum of real-world scenarios.

References

- Ahn, J., Kim, J., & Sung, Y. (2022). The effect of gender stereotypes on artificial intelligence recommendations. *Journal of Business Research*, 141, 50–59. <https://doi.org/10.1016/j.jbusres.2021.12.007>
- 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.2018.07.026>
- Amialchuk, A., & Sapci, O. (2021). The influence of normative misperceptions on alcohol-related problems among school-age adolescents in the US. *Review of Economics of the Household*, 19(2), 453–472. <https://doi.org/10.1007/s11150-020-09481-3>
- Berger, B., Adam, M., Rühr, A., & Benlian, A. (2021). Watch me improve—Algorithm aversion and demonstrating the ability to learn. *Business & Information Systems Engineering*, 63(1), 55–68. <https://doi.org/10.1007/s12599-020-00678-5>
- Berk, R. A. (2021). Artificial intelligence, predictive policing, and risk assessment for law enforcement. *Annual Review of Criminology*, 4, 209–237. <https://doi.org/10.1146/annurev-criminol-051520-012342>
- Berry-Cabán, C. S., Orchowski, L. M., Wimsatt, M., Winstead, T. L., Klaric, J., Prisock, K., Metzger, E., & Kazemi, D. (2020). Perceived and collective norms associated with sexual violence among male soldiers. *Journal of Family Violence*, 35, 339–347. <https://doi.org/10.1007/s10896-019-00096-6>
- Bigman, Y. E., & Gray, K. (2018). People are averse to machines making moral decisions. *Cognition*, 181, 21–34. <https://doi.org/10.1016/j.cognition.2018.08.003>
- Blanton, H., Köblitz, A., & McCaul, K. D. (2008). Misperceptions about norm misperceptions: Descriptive, injunctive, and affective ‘social norming’ efforts to change health behaviors. *Social and Personality Psychology Compass*, 2(3), 1379–1399. <https://doi.org/10.1111/j.1751-9004.2008.00107.x>
- Bodimani, M. (2024). Assessing the impact of transparent AI systems in enhancing user trust and privacy. *Journal of Science & Technology*, 5(1), 50–67. <https://www.thesciencebrigade.com/jst/article/view/68>
- 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>
- Castelo, N., Bos, M. W., & Lehmann, D. R. (2019). Task-dependent algorithm aversion. *Journal of Marketing Research*, 56(5), 809–825. <https://doi.org/10.1177/0022243719851788>
- Chen, S., Liu, J., & Hu, H. (2021). A norm-based conditional process model of the negative impact of optimistic bias on self-protection behaviors during the COVID-19 pandemic in three Chinese cities. *Frontiers in Psychology*, 12, 659218. <https://doi.org/10.3389/fpsyg.2021.659218>
- Chen, S., Wan, F., & Yang, S. (2022). Normative misperceptions regarding pro-environmental behavior: Mediating roles of outcome efficacy and problem awareness. *Journal of Environmental Psychology*, 84, 101917. <https://doi.org/10.1016/j.jenvp.2022.101917>

- Chen, S., Yang, S., Wang, H., & Wan, F. (2022). Subjective social class positively predicts altruistic punishment. *Acta Psychologica Sinica*, 54(12), 1548–1561. <https://doi.org/10.3724/SP.J.1041.2022.01548>
- Chiu, C. Y., Gelfand, M. J., Yamagishi, T., Shteynberg, G., & Wan, C. (2010). Intersubjective culture: The role of intersubjective perceptions in cross-cultural research. *Perspectives on Psychological Science*, 5(4), 482–493. <https://doi.org/10.1177/1745691610375562>
- Conner, M., & Norman, P. (2022). Understanding the intention-behavior gap: The role of intention strength. *Frontiers in Psychology*, 13, 923464. <https://doi.org/10.3389/fpsyg.2022.923464>
- Diab, D. L., Pui, S. Y., Yankelevich, M., & Highhouse, S. (2011). Lay perceptions of selection decision aids in US and non-US samples. *International Journal of Selection and Assessment*, 19(2), 209–216. <https://doi.org/10.1111/j.1468-2389.2011.00548.x>
- Dietvorst, B.J., Simmons, J.P., Massey, C. (2018). Overcoming algorithm aversion: People will use imperfect algorithms if they can (even slightly) modify them. *Management Science*, 64(3), 1155–1170. <https://doi.org/10.1287/mnsc.2016.2643>
- Duan, Y., Edwards, J. S., & Dwivedi, Y. K. (2019). Artificial intelligence for decision making in the era of Big Data—Evolution, challenges and research agenda. *International Journal of Information Management*, 48, 63–71. <https://doi.org/10.1016/j.ijinfomgt.2019.01.021>
- Eastwood, J., Snook, B., & Luther, K. (2012). What people want from their professionals: Attitudes toward decision-making strategies. *Journal of Behavioral Decision Making*, 25(5), 458–468. <https://doi.org/10.1002/bdm.741>
- Faul, F., Erdfelder, E., Lang, A. G., & Buchner, A. (2007). G* power 3: A flexible statistical power analysis program for the social, behavioral, and biomedical sciences. *Behavior Research Methods*, 39(2), 175–191. <https://doi.org/10.3758/BF03193146>
- Figueroa-Armijos, M., Clark, B. B., & da Motta Veiga, S. P. (2023). Ethical perceptions of AI in hiring and organizational trust: The role of performance expectancy and social influence. *Journal of Business Ethics*, 186(1), 179–197. <https://doi.org/10.1007/s10551-022-05166-2>
- Filiz, I., Judek, J. R., Lorenz, M., & Spiwoks, M. (2021). Reducing algorithm aversion through experience. *Journal of Behavioral and Experimental Finance*, 31, 100524. <https://doi.org/10.1016/j.jbef.2021.100524>
- Formosa, P., Rogers, W., Griep, Y., Bankins, S., & Richards, D. (2022). Medical AI and human dignity: Contrasting perceptions of human and artificially intelligent (AI) decision making in diagnostic and medical resource allocation contexts. *Computers in Human Behavior*, 133, 107296. <https://doi.org/10.1016/j.chb.2022.107296>
- Fornara, F., Carrus, G., Passafaro, P., & Bonnes, M. (2011). Distinguishing the sources of normative influence on proenvironmental behaviors: The role of local norms in household waste recycling. *Group Processes & Intergroup Relations*, 14(5), 623–635. <https://doi.org/10.1177/1368430211408149>
- Geber, S., Baumann, E., Czerwinski, F., & Klimmt, C. (2021). The effects of social norms among peer groups on risk behavior: A multilevel approach to differentiate perceived and collective norms. *Communication Research*, 48(3), 319–345.

<https://doi.org/10.1177/0093650218824213>

- Gelfand, M. J., Raver, J. L., Nishii, L., Leslie, L. M., Lun, J., Lim, B. C., Duan, L., Almaliach, A., Ang, S., Arnadottir, J., Aycan, Z., Boehnke, K., Boski, P., Cabecinhas, R., Chan, D., Chhokar, J., D'Amato, A., Ferrer, M., Fischlmayr, I. C., ... Yamaguchi, S. (2011). Differences between tight and loose cultures: A 33-nation study. *Science*, 332(6033), 1100–1104. <https://doi.org/10.1126/science.1197754>
- Gerlich, M. (2023). Perceptions and acceptance of artificial intelligence: A multi-dimensional study. *Social Sciences*, 12(9), 502. <https://doi.org/10.3390/socsci12090502>
- Ge, X. (2023). Experimentally manipulating mediating processes: Why and how to examine mediation using statistical moderation analyses. *Journal of Experimental Social Psychology*, 109, 104507. <https://doi.org/10.1016/j.jesp.2023.104507>
- Graupensperger, S., Lee, C. M., & Larimer, M. E. (2021). Young adults underestimate how well peers adhere to COVID-19 preventive behavioral guidelines. *The Journal of Primary Prevention*, 42(3), 309–318. <https://doi.org/10.1007/s10935-021-00633-4>
- Gursoy, D., Chi, O. H., Lu, L., & Nunkoo, R. (2019). Consumers acceptance of artificially intelligent (AI) device use in service delivery. *International Journal of Information Management*, 49, 157–169. <https://doi.org/10.1016/j.ijinfomgt.2019.03.008>
- Hasan, R., Shams, R., & Rahman, M. (2021). Consumer trust and perceived risk for voice-controlled artificial intelligence: The case of Siri. *Journal of Business Research*, 131, 591–597. <https://doi.org/10.1016/j.jbusres.2020.12.012>
- Ho, G., Wheatley, D., & Scialfa, C. T. (2005). Age differences in trust and reliance of a medication management system. *Interacting with Computers*, 17(6), 690–710. <https://doi.org/10.1016/j.intcom.2005.09.007>
- Ismatullaev, U. V. U., & Kim, S. H. (2024). Review of the factors affecting acceptance of AI-infused systems. *Human Factors*, 66(1), 126–144. <https://doi.org/10.1177/00187208211064707>
- Kawaguchi, K. (2021). When will workers follow an algorithm? A field experiment with a retail business. *Management Science*, 67(3), 1670–1695. <https://doi.org/10.1287/mnsc.2020.3599>
- Kelan, E. K. (2023). Algorithmic inclusion: Shaping the predictive algorithms of artificial intelligence in hiring. *Human Resource Management Journal*. <https://doi.org/10.1111/1748-8583.12511>
- Kelly, S., Kaye, S. A., & Oviedo-Trespalacios, O. (2023). What factors contribute to the acceptance of artificial intelligence? A systematic review. *Telematics and Informatics*, 77, 101925. <https://doi.org/10.1016/j.tele.2022.101925>
- Kerr, N. L. (1996). “Does my contribution really matter?”: Efficacy in social dilemmas. *European Review of Social Psychology*, 7(1), 209–240. <https://doi.org/10.1080/14792779643000029>
- Kim, D., Hwang, K., Lee, E., & Yoon, Y. (2023). The effect of artificial intelligence (AI) robot characteristics and dialectical thinking on AI robot adoption intention. *Journal of Consumer Behaviour*. <https://doi.org/10.1002/cb.2201>
- Kim, J., Merrill Jr., K., & Collins, C. (2021). AI as a friend or assistant: The mediating role of perceived usefulness in social AI vs. functional AI. *Telematics and Informatics*, 64, 101694. <https://doi.org/10.1016/j.tele.2021.101694>

- Lacroux, A., & Martin-Lacroux, C. (2022). Should I trust the artificial intelligence to recruit? Recruiters' perceptions and behavior when faced with algorithm-based recommendation systems during resume screening. *Frontiers in Psychology*, 13. <https://doi.org/10.3389/fpsyg.2022.895997>
- Langer, M., & Landers, R. N. (2021). The future of artificial intelligence at work: A review on effects of decision automation and augmentation on workers targeted by algorithms and third-party observers. *Computers in Human Behavior*, 123, 106878. <https://doi.org/10.1016/j.chb.2021.106878>
- Lee, H., & Cho, C. H. (2020). Digital advertising: Present and future prospects. *International Journal of Advertising*, 39(3), 332–341. <https://doi.org/10.1080/02650487.2019.1642015>
- Lee, M. K. (2018). Understanding perception of algorithmic decisions: Fairness, trust, and emotion in response to algorithmic management. *Big Data & Society*, 5(1), 2053951718756684. <https://doi.org/10.1177/2053951718756684>
- Lin, H., Chi, O. H., & Gursoy, D. (2020). Antecedents of customers' acceptance of artificially intelligent robotic device use in hospitality services. *Journal of Hospitality Marketing & Management*, 29(5), 530–549. <https://doi.org/10.1080/19368623.2020.1685053>
- Litterscheidt, R., & Streich, D. J. (2020). Financial education and digital asset management: What's in the black box? *Journal of Behavioral and Experimental Economics*, 87, 101573. <https://doi.org/10.1016/j.socrec.2020.101573>
- 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>
- Longoni, C., & Cian, L. (2022). Artificial intelligence in utilitarian vs. hedonic contexts: The “word-of-machine” effect. *Journal of Marketing*, 86(1), 91–108. <https://doi.org/10.1177/0022242920957347>
- Mahmud, H., Islam, A. N., Ahmed, S. I., & Smolander, K. (2022). What influences algorithmic decision-making? A systematic literature review on algorithm aversion. *Technological Forecasting and Social Change*, 175, 121390. <https://doi.org/10.1016/j.techfore.2021.121390>
- McDonald, R. I., & Crandall, C. S. (2015). Social norms and social influence. *Current Opinion in Behavioral Sciences*, 3, 147–151. <https://doi.org/10.1016/j.cobeha.2015.04.006>
- Palacios, J., Fan, Y., Yoeli, E., Wang, J., Chai, Y., Sun, W., Rand, D. G., & Zheng, S. (2022). Encouraging the resumption of economic activity after COVID-19: Evidence from a large scale-field experiment in China. *Proceedings of the National Academy of Sciences of the United States of America*, 119(5), e2100719119. <https://doi.org/10.1073/pnas.2100719119>
- Parent-Rochelleau, X., & Parker, S. K. (2021). Algorithms as work designers: How algorithmic management influences the design of jobs. *Human Resource Management Review*, 32(3), 100838. <https://doi.org/10.1016/j.hrmr.2021.100838>
- Pizam, A., Ozturk, A. B., Hacikara, A., Zhang, T., Balderas-Cejudo, A., Buhalis, D., Fuchs, G., Hara, T., Meira, J., Revilla, R. G. M., Sethi, D., Shen, Y., & State, O. (2024). The

- role of perceived risk and information security on customers' acceptance of service robots in the hotel industry. *International Journal of Hospitality Management*, 117, 103641. <https://doi.org/10.1016/j.ijhm.2023.103641>
- Preacher, K. J., & Hayes, A. F. (2004). SPSS and SAS procedures for estimating indirect effects in simple mediation models. *Behavior Research Methods, Instruments, & Computers*, 36(4), 717–731. <https://doi.org/10.3758/BF03206553>
- Ragelienė, T., & Grønhøj, A. (2020). The influence of peers' and siblings' on children's and adolescents' healthy eating behavior. A systematic literature review. *Appetite*, 148, 104592. <https://doi.org/10.1016/j.appet.2020.104592>
- Reich, T., Kaju, A., & Maglio, S. J. (2023). How to overcome algorithm aversion: Learning from mistakes. *Journal of Consumer Psychology*, 33(2), 285–302. <https://doi.org/10.1002/jcpy.1313>
- Sammur, G., & Bauer, M. W. (2021). *The psychology of social influence: Modes and modalities of shifting common sense*. Cambridge, UK: Cambridge University Press. <https://doi.org/10.1017/9781108236423>
- Sargent, R. H., & Newman, L. S. (2021). Pluralistic ignorance research in psychology: A scoping review of topic and method variation and directions for future research. *Review of General Psychology*, 25(2), 163–184. <https://doi.org/10.1177/1089268021995168>
- Schwesig, R., Brich, I., Buder, J., Huff, M., & Said, N. (2023). Using artificial intelligence (AI)? Risk and opportunity perception of AI predict people's willingness to use AI. *Journal of Risk Research*, 26(10), 1053–1084. <https://doi.org/10.1080/13669877.2023.2249927>
- Shank, D. B., Kashima, Y., Peters, K., Li, Y., Robins, G., & Kirley, M. (2019). Norm talk and human cooperation: Can we talk ourselves into cooperation? *Journal of Personality and Social Psychology*, 117(1), 99–123. <https://doi.org/10.1037/pspi0000163>
- Sharan, N. N., & Romano, D. M. (2020). The effects of personality and locus of control on trust in humans versus artificial intelligence. *Heliyon*, 6(8), e04572. <https://doi.org/10.1016/J.HELIYON.2020.E04572>
- Sparkman, G., Geiger, N., & Weber, E. U. (2022). Americans experience a false social reality by underestimating popular climate policy support by nearly half. *Nature Communications*, 13, 4779. <https://doi.org/10.1038/s41467-022-32412-y>
- Sridhar, S., & Srinivasan, R. (2012). Social influence effects in online product ratings. *Journal of Marketing*, 76(5), 70–88. <https://doi.org/10.1509/jm.10.0377>
- Stein, J. P., Appel, M., Jost, A., & Ohler, P. (2020). Matter over mind? How the acceptance of digital entities depends on their appearance, mental prowess, and the interaction between both. *International Journal of Human-Computer Studies*, 142, 102463. <https://doi.org/10.1016/j.ijhcs.2020.102463>
- Wan, C., Chiu, C. Y., Tam, K. P., Lee, S. L., Lau, I. Y. M., & Peng, S. (2007). Perceived cultural importance and actual self-importance of values in cultural identification. *Journal of Personality and Social Psychology*, 92(2), 337–354. <https://doi.org/10.1037/0022-3514.92.2.337>
- Wang, G., Badal, A., Jia, X., Maltz, J. S., Mueller, K., Myers, K. J., Niu, C., Vannier, M., Yan, P., Yu, Z., & Zeng, R. (2022). Development of metaverse for intelligent healthcare. *Nature Machine Intelligence*, 4(11), 922–929.

<https://doi.org/10.1038/s42256-022-00549-6>

- Xie, Z., Yu, Y., Zhang, J., & Chen, M. (2022). The searching artificial intelligence: Consumers show less aversion to algorithm-recommended search product. *Psychology & Marketing*, 39(10), 1902–1919. <https://doi.org/10.1002/mar.21706>
- Zhang, L., Pentina, I., & Fan, Y. (2021). Who do you choose? Comparing perceptions of human vs robo-advisor in the context of financial services. *Journal of Services Marketing*, 35(5), 634–646. <https://doi.org/10.1108/JSM-05-2020-0162>
- Zhao, X., & Epley, N. (2022). Surprisingly happy to have helped: Underestimating prosociality creates a misplaced barrier to asking for help. *Psychological Science*, 33(10), 1708–1731. <https://doi.org/10.1177/09567976221097615>

Supplementary Materials

Underestimating Others' Positive Attitude Toward AI Reduces Intentions to Adopt AI in Various Business Contexts

1. Study 1: Three Scenarios of Artificial Intelligence Application

[Home purchase] Imagine you are in the current situation of needing to purchase a home. A real estate website presents HomeFinder, an AI chatting application. Through this platform, you can engage in a conversation with the AI, and during the chat, it will furnish you with pertinent information about homes tailored to your requirements, offering advice on the home-buying process. On a scale of 1 to 7, where 1 denotes a strong unwillingness and 7 indicates a high willingness, to what extent would you be open to utilizing the AI HomeFinder for obtaining home information and advice on purchasing a home?

[Job interview] Envision yourself as a prospective graduate gearing up for an interview with Company X. Company X employs an AI application known as JobFit, designed to assess the competencies of each candidate and determine the alignment between the candidate's profile and the requirements of the position. During the interview, you submit your CV and engage in a 15-minute conversation with the AI JobFit. Upon conclusion of the interview, the AI JobFit will furnish you with the results and present several reasons for either passing or failing the interview. On a scale of 1 to 7, where 1 signifies a strong unwillingness and 7 indicates a high willingness, how open would you be to utilizing the AI JobFit in the interview process for Company X?

[Online medical consultation] Envision a scenario where you experience discomfort at night but are currently unable to visit a hospital. A medical application offers an online consultation service featuring an AI doctor. Through a chat box, you can input your symptoms, ask questions, and engage in a conversation with the AI doctor, HealthChat. HealthChat will provide you with diagnostic results and offer advice on medication and medical treatment. On a scale of 1 to 7, where 1 reflects a strong unwillingness and 7 signifies a high willingness, how inclined would you be to consult with the AI doctor

HealthChat online?

2. Scales Employed in Study 1 to Measure the Two Mediating Variables

To what extent do you agree with the following statements? Higher numbers indicate higher levels of agreement, 1 = strongly disagree, 7 = strongly agree.

(1) Perceived risk

- Using AI is risky.
- Using AI will increase a lot of uncertainty.
- There are some losses associated with the use of AI.
- Using AI can reveal personal information about the user.

(2) Outcome efficacy

- Using AI can help people get a lot of convenience.
- Using AI can improve people’s quality of life.
- Using AI helps people make more rational decisions.
- Using AI can make life more efficient.

3. Regression Results for Study 1 After Controlling for Participants’ Actual Attitudes Toward AI

The outcomes of the regression analyses following adjustment for participants’ actual attitudes toward AI are as follows: (1) The total effect of perceived others’ attitudes on willingness to use AI remains significantly positive across all scenarios (home purchase: Effect = 0.27, SE = 0.07, 95% CI [0.13, 0.41]; job interview: Effect = 0.27, SE = 0.11, 95% CI [0.06, 0.48]; online medical consultation: Effect = 0.51, SE = 0.11, 95% CI [0.28, 0.73]). (2) Additionally, the mediation effects of perceived risk (home purchase: Effect = 0.03, SE = 0.01, 95% CI [0.002, 0.06]; job interview: Effect = 0.12, SE = 0.05, 95% CI [0.03, 0.23]; online medical consultation: Effect = 0.09, SE = 0.04, 95% CI [0.02, 0.19]) and outcome efficacy (home purchase: Effect = 0.12, SE = 0.07, 95% CI [0.01, 0.28]; job interview: Effect = 0.10, SE = 0.05, 95% CI [0.01, 0.21]; online medical consultation: Effect = 0.09, SE = 0.05, 95% CI [0.004, 0.18]) remained significantly positive across the three scenarios. Detailed

results and graphs are presented in Table S1 and Figure S1, respectively.

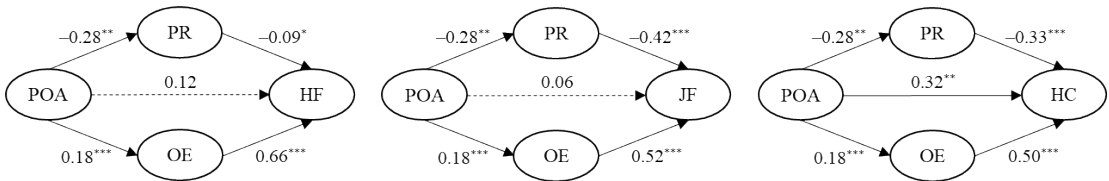
Table S1

Total, direct, and indirect effects across three scenarios.

Effects	Home purchase	Job interview	Online consultation
Total effect	0.27 [0.13, 0.41]	0.27 [0.06, 0.48]	0.51 [0.28, 0.73]
Direct effect	0.12 [-0.01, 0.25]	0.06 [-0.13, 0.25]	0.32 [0.11, 0.54]
Indirect effect: perceived risk	0.03 [0.002, 0.06]	0.12 [0.03, 0.23]	0.09 [0.02, 0.19]
Indirect effect: outcome efficacy	0.12 [0.01, 0.28]	0.10 [0.01, 0.21]	0.09 [0.004, 0.18]
Outcome efficacy – perceived risk	0.10 [-0.01, 0.26]	-0.02 [-0.14, 0.10]	-0.0003 [-0.13, 0.11]

Figure S1

Mediation models in Study 1.



Note. N = 309, * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$; POA = perceived others' attitudes, PR = perceived risk, OE = outcome efficacy, HF = HomeFinder, JF = JobFit, HC = HealthChat.

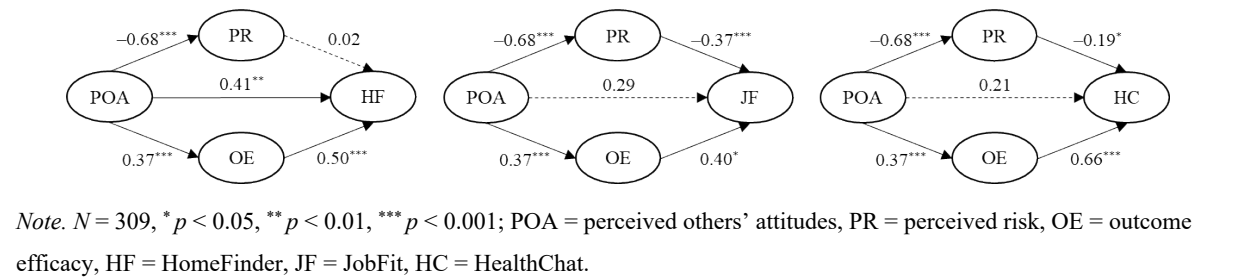
4. Regression Results for Study 2 After Controlling for Participants' Actual Attitudes Toward AI

The regression findings subsequent to controlling for participants' actual attitudes toward AI are as follows: The overall impact of perceived others' attitudes on the inclination to utilize AI persisted significantly positive across all scenarios (house purchase: Effect = 0.57, $SE = 0.12$, 95% CI [0.33, 0.81]; job interview: Effect = 0.68, $SE = 0.18$, 95% CI [0.33, 1.04]; online medical consultation: Effect = 0.58, $SE = 0.18$, 95% CI [0.23, 0.92]). The indirect influence of perceived risk remained significant solely in job interviews (Effect = 0.25, $SE = 0.08$, 95% CI [0.11, 0.42]) and online medical consultations (Effect = 0.13, $SE = 0.07$, 95% CI [0.01, 0.26]). Conversely, the indirect effect of outcome efficacy persisted significantly across all three scenarios (home purchase: Effect = 0.18, $SE = 0.07$, 95% CI [0.06, 0.32]; job interview: Effect = 0.15, $SE = 0.07$, 95% CI [0.02, 0.28]; online consultation: Effect = 0.24, $SE = 0.07$, 95% CI [0.11, 0.38]). A detailed presentation of these findings is provided in Table S2 and Figure S2.

Table S2
Total, direct, and indirect effects across three scenarios.

Effects	Home purchase	Job interview	Online consultation
Total effect	0.57 [0.33, 0.81]	0.68 [0.33, 1.04]	0.58 [0.23, 0.92]
Direct effect	0.41 [0.15, 0.67]	0.29 [-0.10, 0.67]	0.21 [-0.17, 0.58]
Indirect effect: perceived risk	-0.02 [-0.11, 0.06]	0.25 [0.11, 0.42]	0.13 [0.01, 0.26]
Indirect effect: outcome efficacy	0.18 [0.06, 0.32]	0.15 [0.02, 0.28]	0.24 [0.11, 0.38]
Outcome efficacy – perceived risk	0.20 [0.05, 0.36]	-0.11 [-0.31, 0.10]	0.12 [-0.07, 0.30]

Figure S2
Mediation models in Study 2.



5. Single-Article Meta-Analysis

The primary findings of this study are encapsulated in three principal insights: firstly, individuals tend to underestimate the affirmative attitudes of their peers towards AI; secondly, this underestimation significantly attenuates their individual proclivity to adopt AI; and thirdly, the mediating mechanisms of perceived risk and outcome efficacy play pivotal roles in the association between perceived others' attitudes and the willingness to employ AI. To systematically scrutinize the extent of this underestimation, its repercussions on individual AI adoption proclivities, and the intervening functions of perceived risk and outcome efficacy, we conducted a single-paper meta-analysis for each of these central findings in Study 4 (McShane & Böckenholt, 2017).

For the degree of underestimation, four effect sizes (d) were derived from paired-sample t -tests comparing participants' actual attitudes and perceived others' attitudes in Study 1, Study 2, Study 3a, and Study 3b. In investigating the nexus between perceived others' attitudes and the willingness to adopt AI, eight effect sizes were selected for uniformity: correlation coefficients (r) between perceived others' attitudes and willingness to adopt AI in

each scenario in Study 1, correlation coefficients (r) between perceived others' attitudes measured post-intervention and willingness to adopt AI across the three scenarios in Study 2, and correlation coefficients (r) between perceived others' attitudes measured post-intervention and willingness to adopt AI in the online consultation scenario in the control groups of Study 3a and Study 3b. To ascertain the mediating roles of perceived risk and outcome efficacy, we utilized perceived others' attitudes as the predictor variable and calculated standardized mediating effects using PROCESS 4.1 (Preacher & Hayes, 2004) (Bootstrap $N = 5,000$, Model = 4). This encompassed seven effect sizes for each mediator. Notably, in Study 3a and Study 3b, our focus was solely on control groups where perceived risk/outcome efficacy was unaltered.

The single-paper meta-analyses, executed using R 4.3.1 (package "meta"; Balduzzi et al., 2019), yielded the following outcomes: (1) The underestimation of others' positive attitudes towards AI was observed at a medium effect size level ($d = -0.53$, 95%CI $[-0.60, -0.47]$, $z = -16.37$, $p < 0.001$) (small effect size: $d = 0.20$; medium effect size: $d = 0.50$; large effect size: $d = 0.80$) (Cohen, 1988). (2) Analogously, the correlation coefficient between perceived attitudes of others and the individual willingness to adopt AI reached the medium effect size level ($r = 0.31$, 95%CI $[0.26, 0.35]$, $z = 13.36$, $p < 0.001$) (small effect size: $r = 0.10$; medium effect size: $r = 0.30$; large effect size: $r = 0.50$) (Cohen, 1988). (3) Regarding the mediating effects, perceived risk (effect = 0.07, 95%CI $[0.03, 0.10]$, $z = 3.78$, $p < 0.001$) accounted for approximately half of the indirect effect observed for outcome efficacy (effect = 0.13, 95%CI $[0.09, 0.17]$, $z = 6.33$, $p < 0.001$).

Richard et al. (2003) conducted an extensive meta-analysis encompassing 322 studies with substantial sample sizes ($N > 8,000,000$). Their findings indicated a relatively modest mean effect size in psychosocial studies ($r = 0.21$). In contrast, the pivotal findings in this study yielded medium effect sizes for the two primary insights. This implies that, despite being underexplored in prior research, the biased perception of others' attitudes toward AI emerges as a significant and influential factor in shaping AI acceptance.

References

- Balduzzi, S., Rücker, G., & Schwarzer, G. (2019). How to perform a meta-analysis with R: a practical tutorial. *Evidence Based Mental Health*, 22(4), 153–160. <https://doi.org/10.1136/ebmental-2019-300117>
- Cohen, J. (1988). *Statistical power analysis for the behavioral sciences*. Hillsdale, NJ: Erlbaum. <https://doi.org/10.4324/9780203771587>
- McShane, B. B., & Böckenholt, U. (2017). Single-paper meta-analysis: Benefits for study summary, theory testing, and replicability. *Journal of Consumer Research*, 43(6), 1048–1063. <https://doi.org/10.1093/jcr/ucw085>
- Richard, F. D., Bond, C. F., Jr., & Stokes-Zoota, J. J. (2003). One hundred years of social psychology quantitatively described. *Review of General Psychology*, 7(4), 331–363. <https://doi.org/10.1037/1089-2680.7.4.331>