



The Effect of Anthropomorphic Design on Information Trust and Adoption Intention in AI Chatbots

Yiyang Bai¹, Yichang Guo^{2,*} and Jialei Yu³

¹ Ealing International School, Dalian, 116000, China

² Sino-Korean New Media College of Zhongnan University of Economics and Law, Wuhan, 430073, China

³ School of Humanities and Social Sciences, Nanjing Forestry University, Nanjing, 210000, China

yichangguo2@gmail.com

Abstract. This study examines how human-like AI chatbots affect users' trust in their information and willingness to follow their advice. The study puts together the Stimulus-Organism-Response (SOR) model and the Technology Acceptance Model (TAM). These models help show two things: how helpful users find the AI, and how easy they find it to use, and how a person's own tendency to see things as human works as a moderator. The authors ran an experiment with two factors. One factor was how the AI talked—like a person or in a neutral way. The other factor was what the task was about—finding facts or dealing with emotions. A total of 189 participants with prior AI chatbot experience took part. The results show that the immediate effect of anthropomorphic focus on information trust was insignificant. However, trust strongly predicted adoption intention. Perceived usefulness served as a significant mediator, but the effect was negative—making the AI sound more human actually reduced users' perception of its functional value. Whether users found the AI easy to use did not bridge the link. The effect of their human-like view was only small. These findings suggest that adding human-like features does not automatically build user trust. Designers should apply anthropomorphic design with care, considering the task type. For information-seeking tasks, a more neutral style may be effective. For emotional support scenarios, a warmer tone may be helpful. A thoughtful approach that considers both the task and the user is likely to work better.

Keywords: Anthropomorphic Design, Information Trust, Adoption Intention, SOR Model, TAM Model

1 Introduction

People now use AI chatbots like ChatGPT and DeepSeek all the time. According to recent surveys, over 60% of internet users interact with AI chatbots at least once a week for information seeking, decision support, or even emotional engagement. This shows how deeply these tools have become part of daily life. So researchers really need to

understand what makes users trust AI. And the present study aims to know what makes users actually use the information AI gives.

To make AI seem less robotic, designers can add human-like touches. This might include adding emojis, using everyday language, or calling the user “you” rather than just stating information. Past work has shown that when AI talks this way, people react differently to it [1]. Other studies suggest that users put more trust in AI when it speaks in a more human-like tone [2]. But individuals still do not really know how this works. Does it happen because users start to think the AI is more helpful? Or because it feels easier to talk to? These questions are still unanswered.

To better understand this, the authors bring together two well-known models. One is the SOR. It explains how things outside a person affect what goes on inside their head, and then how that leads to behavior. The other model is the TAM. It focuses on two ideas: how useful and how easy the AI feels to users [3]. The thinking is that the way AI talks is the stimulus. That stimulus affects how useful and how easy users think the AI is. Those thoughts are the organism. Then those thoughts shape whether users trust the AI. Trust is the response. Finally, trust affects whether users take the advice given by the AI.

This study tested this with a 2×2 experiment. Two things were changed. One was whether the AI used human-like language or not. The other was the kind of task. Some tasks were about finding information. Others were about getting emotional support. A total of 189 people took part. They had all used AI chatbots before.

With this study, the authors hope to see the steps between AI design and user trust. The authors also want to show how the TAM can be used against the backdrop of information adoption.

2 Literature Review

The SOR gives the study its main structure. It says that things outside a person stimulus(S) affect what they think and feel inside organism(O), and that then leads to what they do response(R). Previous studies found that users’ internal states are key variables that affect their online behavioral intention. For example, trust has an important influence on information adoption intention [4]. Therefore, based on these studies, this research conducts an exploration using the SOR model.

In this study, external environmental stimuli correspond to the anthropomorphic design of artificial intelligence platforms. This kind of design changes two things: how useful users think the information is, and how easy they find it to use. These factors finally lead to users’ information adoption behavior.

Technology keeps advancing, and computers now play a big part in everyday life. People often tend to give human qualities to non-human things. This happens not just in daily life but also in areas like how people interact with computers. Anthropomorphism is a common psychological phenomenon. It is often described as the tendency to attribute human traits, purposes and mental states to artificial things. It is not only common in daily life but also extends to practical applications such as human-computer

interaction. Nicolas Spatola et al. conducted experiments on the psychological mechanisms and functional values of personification on subjects from multiple countries through online questionnaires and psychological scales. The study shows that the stronger an individual's "mentalizing" ability (the ability to understand others' minds) is, the easier it is for people to endow robots with cognitive and emotional states in real-time interactions [5]. Therefore, in this study, the user's perception of the anthropomorphic AI platform is affected by many factors during human-computer interaction. This demonstrates the anthropomorphic design of an artificial intelligence platform may exert both beneficial and detrimental impacts: For one thing, it enhances user participation and trust; for another, it can also lead to unrealistic user expectations or ethical issues.

In a quantitative study concentrating on AI chatbot users in e-commerce contexts, Sharareh Shahidi Hamedani et al. came out showing that anthropomorphic design helps consumers' adoption behavior. Such design also plays a mediating role between adoption and perceived trust [6].

The measurement of information trust dates back to the pioneering study on source credibility by McCroskey [7]. This classic research is still applicable in information interaction on AI platforms. This study partly uses the scale developed by McCroskey to measure information trust [7]. This scale has good reliability and validity. Although the McCroskey's Credibility Scale was originally developed in the context of interpersonal communication, it is also applicable to information interaction mediated by artificial intelligence. Because when users evaluate content generated by artificial intelligence, the two core dimensions of credibility and professional ability are both very important [7].

According to research investigating user behavior on social commerce platforms, Lu, B., Fan, W. and his co-authors found that feeling like someone is there—called social presence—goes with higher trust and more buying [4]. That is to say, when users perceive that they are interacting with a socially existing entity rather than a mechanical system, their level of trust can be substantially enhanced. This will further increase their willingness to purchase [4]. Therefore, social presence is a key mediating variable that they are interacting with a socially existing entity rather than a mechanical system.

Sharareh Shahidi Hamedani et al. carried out a quantitative study. The study focused on consumers who had used AI chatbots before. They found one important factor that can predict users' information adoption behavior. This factor is users' belief that AI chatbots can help them complete e-commerce tasks efficiently. These tasks include consultation, ordering, and after-sales service [6]. Users' opinion that the chatbot is simple to utilize and does not need extra learning effort also directly shapes adoption behavior. It also has a positive relationship with perceived usefulness [6]. They are consistent with the core assumption of TAM. The assumption is that functional perception determines adoption behavior. In another study, researchers explored the important reasons why students in higher education choose to use ChatGPT. They found that the feeling of anthropomorphism helps improve students' willingness to adopt information, but it does so by increasing user trust. It also makes the TAM cognitive path better at predicting behavioral intention [8]. This study shows that in adopting AI information, anthropomorphism is not just a social stimulus (S). It is also an important factor that

influences information adoption behavior (R) through the TAM cognitive process (O). As mentioned before, the study by Sharareh Shahidi Hamedani and others found that the perception of anthropomorphic design connects information adoption behavior and perceived trust [6]. Also, many studies have shown that anthropomorphic features can influence user trust.

Eliasson & Engström used 30 pieces of content generated by ChatGPT (fifteen of them were restaurant reviews and fifteen were travel posts). They asked participants to rate these contents [3]. AI chatbots used an informal communication style. This style includes a conversational tone, first-person expressions, and colloquial language. When this style was used, participants gave much higher scores. The scores were about perceived anthropomorphism and trust. This was true for restaurant reviews and travel posts created by AI. This directly proves that communication style has a positive effect on user trust [3]. Communication style is a specific anthropomorphic design cue. Meanwhile, regression analysis shows that perceived anthropomorphism can predict user trust greatly and positively. This means anthropomorphic design boosts trust mainly by making AI more human-like. This also offers direct proof that trust is a key cognitive and emotional state [3].

It is worth our attention. The mediating effect of anthropomorphic design connects information adoption and perceived trust. This effect may be influenced by many different practical situations. The study by Eliasson and Engström suggests that the effectiveness of informal anthropomorphic style in enhancing trust may depend on specific contexts. This style has a more significant effect in hedonic and experiential fields. Typical fields include restaurant reviews and travel sharing. This context sensitivity shows the importance of proper design. The AI communication style must match the nature of the information task. This rule is directly related to the design of this study. The nature of information or tasks should be considered carefully. This rule applies to the design of AI information presentation [3].

The TAM model rests on two main ideas: how helpful users find the AI and how easy it is to use [9]. This study gives a definition to perceived usefulness. It stands for the degree to which users think subjectively. Users believe that using an AI platform can improve their efficiency in processing information. It can also help users make better decisions. Perceived ease of use means the level of effort that people feel to use the AI platform. That is, whether the platform is easy to operate and understand [9].

People's attitudes toward technology often shape what they decide to do. For example, Al-Abdullatif looked at university students in Saudi Arabia [10]. The results showed that when students had positive feelings about using chatbots in education, they were more willing to use them [10]. This also supports the theoretical hypothesis of this study. The hypothesis takes information trust as a direct factor. Information trust is a special positive attitude. It focuses on the reliability of information. This factor affects information adoption intention directly.

Past research has tied AI anthropomorphism to how users take in information. This earlier work gives later studies a good base to build on. One clear finding stands out: adding human-like touches to AI systems tends to work. These elements include conversational tone, personalized language, emojis, and virtual-image visual design. This design can greatly raise users' trust in AI. It can further improve their willingness to

use the platform [3]. This proves the effectiveness and rationality of anthropomorphic design on AI platforms. Researchers have started to study the operating mechanism of this path. Some researchers found a related result. Anthropomorphism can improve users' social presence from an emotional angle. Social presence means feeling like someone else is there with you. It helps build emotional connection and trust with the other party [4]. Other scholars adopted the classic TAM from a cognitive angle. They proved that anthropomorphism can improve two user perceptions. They showed that when AI seems more human, users see it as more helpful and easier to use. These factors further influence users' attitudes and behaviors [8]. It more clearly explains the core question about the working mechanism of anthropomorphism. On the dependent variable level, research is shifting from "usage intention" to the more specific "information adoption intention and information adoption behavior. AI platforms have increasingly served as primary information sources. Studies point out that widespread reliance on AI-generated content for decision-making across domains [11]. Understanding the mechanisms underlying users' trust in and adoption of AI-provided information has emerged as a pressing research priority. Although there are many research results, existing studies are still fragmented [12]. Scholars Lopez-Lopez, D., & Iniesta, M. B. pointed out that AI is reshaping the consumption journey. However, the academic field lacks a holistic knowledge framework [12]. To fill this gap, the authors set up a model that looks at both mediation and moderation. This helps show how AI design shapes whether users take its advice, and when it works or not. This study innovatively combines the SOR model and the TAM model. Under the SOR framework, the mediating variables are perceived usefulness and perceived ease of use. It more clearly explains the core question about the working mechanism of anthropomorphism. This study also takes anthropomorphism tendency as a moderating variable [9]. It tests the moderating effect of users' anthropomorphism tendency on the research path. This can clarify the boundary conditions of anthropomorphic design. It also explains the differences in AI anthropomorphism perception among different users. This study takes information adoption intention as the dependent variable. It proves the effective path of the TAM model in this research [13]. The measurement scale of "anthropomorphic style" in this study is partly based on the perceived humanization scale developed by Go and Sundar [14]. According to the research context, we have adjusted some items slightly. This way can measure the anthropomorphic differences of AI more accurately. The differences mainly lie in language affinity and emoji use.

This study proves the effectiveness of anthropomorphism again. It also further explores its specific internal mechanism. The research results may offer theoretical and practical implications for AI platforms. These implications target the information interaction design of AI platforms. They help improve user trust and information adoption to the fullest. This works in different specific task scenarios.

3 Results

3.1 Reliability Analysis

Five core scales were checked for internal consistency. To check if the scales gave steady results, the authors ran Cronbach's alpha. Table 1 holds the numbers.

Table 1. Cronbach's α Coefficients for Each Scale

Scale	Items	Number of Items	Cronbach's α
Perceived Usefulness (PU)	PU1–PU4	4	0.954
Perceived Ease of Use (PEOU)	PEOU1–PEOU4	4	0.952
Trust in Information (TR)	TR1–TR4	4	0.953
Willingness to Adopt Information (AD)	AD1–AD4	4	0.939
Attribution Tendencies (AT)	AT1–AT6	6	0.818

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

Table 1 shows that the five scales had Cronbach's α between 0.818 and 0.954, all exceeding the standard of 0.7. This means the scales gave steady results, the measurement results are reliable, and they are suitable for subsequent statistical analysis.

3.2 Manipulation Checks

Personification Manipulation Test. To verify the validity of the anthropomorphism manipulation, independent samples t-tests were conducted on the two manipulation items—M1 (degree of human-likeness) and M2 (degree of friendliness)—between the anthropomorphism group (Group A+B, $n = 92$) and the non-anthropomorphism group (Group C+D, $n = 97$). The results are shown in Table 2.

Table 2. Personification Manipulation Test: Differences between the Personification Group and the Non-Personification Group on M1 and M2

Variable	Group with Personification M (SD)	Non-personification Group M (SD)	t	df	p
M1 (Human-like Quality)	4.19	3.12	6.276	187	$< .001^{***}$
M2 (Affectivity)	4.55	3.39	5.909	187	$< .001^{***}$

Note. M = mean. SD not shown due to space constraints. *** $p < .001$.

On M1 and M2, the human-like AI group came out higher than the neutral group, according to the t-test. So the way the AI talked worked as planned. Users in the human-like group saw it as more human.

Testing the Narrative Framework Manipulation. To find out if users saw the tasks as different, a chi-square test was run. A chi-square test (χ^2) was used to compare the differences in responses to Question M3 (judgment of narrative type) between the factual group (Groups A+C, $n = 95$) and the emotional group (Groups B+D, $n = 94$) (See Table 3).

Table 3. Test of Narrative Framework Manipulation: Differences in M3 Choices Between the Factual Group and the Emotional Group

Group	Selected "Factual" (n)	Selected "Emotional Type" (n)	χ^2	df	p
Factual Group (A+C)	46	49	0.047	1	.828
Emotional Group (B+D)	48	46	—	—	—

Note. χ^2 = chi-square statistic.

The chi-square test gave a p-value of .828, which means the two groups did not differ much in how they answered Question M3: $\chi^2(1) = 0.047$, $p = .828$. The participants' subjective judgments regarding narrative types did not show the expected differentiation; the narrative framework manipulation failed to successfully distinguish participants' subjective perceptions. This may be related to the presentation of the experimental materials, the wording of the questions, or differences in participants' understanding of the narrative framework. In subsequent analyses, the effects of the narrative framework as a between-group variable should be interpreted with caution.

3.3 Descriptive Statistics

For the five main measures, the study got the average and standard deviation for each of the four groups: Group A (with personification + factual narrative), Group B (with personification + emotional narrative), Group C (without personification + factual narrative), and Group D (without personification + emotional narrative).

As shown in Table 4, most average scores fell above 4, which is the middle point of the 7-point scale. The anthropomorphic groups (A and B) and the non-anthropomorphic groups (Groups C and D) had similar mean scores for information trust (TR). Preliminary descriptive statistics did not reveal significant between-group differences, which require further verification through hypothesis testing.

Table 4. Descriptive Statistics of Key Variables Across Experimental Groups (M ± SD)

Group	n	Per- ceived Use- fulness (PU)	Per- ceived Ease of Use (PEOU)	Trust in Information TR	Willing- ness to Adopt Information AD	At- tribu- tion Tenden- cies AT
Group A (Personification + Facts)	46	5. 11 (1. 36)	5. 10 (1. 30)	4. 79 (1. 16)	4. 86 (1. 15)	4. 16 (0. 87)
Group B (Includes Per- sonification + Emotion)	46	4. 76 (1. 33)	4. 68 (1. 28)	4. 41 (1. 07)	4. 57 (1. 09)	4. 12 (0. 91)
Group C (No Personifi- cation + Facts)	49	4. 66 (1. 33)	4. 54 (1. 24)	4. 22 (1. 11)	4. 30 (0. 98)	4. 01 (0. 93)
Group D (No Personifi- cation + Emo- tions)	48	6. 06 (0. 65)	5. 61 (0. 98)	5. 18 (0. 94)	5. 24 (0. 81)	4. 58 (1. 14)

Note: Standard deviation (SD) is shown in parentheses. PU = perceived usefulness; PEOU = perceived ease of use; TR = information trust; AD = information adoption intention; AT = anthropomorphic tendency. All scales used a 7-point Likert scale (1 = strongly disagree, 7 = strongly agree).

3.4 Hypothesis Testing

H1: The Effect of Personification on Information Trust. The first hypothesis (H1) predicted that making the AI more human-like would lead to higher trust in the information. The authors ran a t-test to see how the average scores for trust in information compared between the two groups. The human-like AI group had an average of 4. 60 (SD = 1. 33), while the neutral group had 4. 70 (SD = 1. 27).

Table 5 H1 Hypothesis Test: Independent-Samples t-Test of the Effect of Personification on Information Trust

Dependent Variable	Group with Personification M (SD)	Non-personification Group M (SD)	t	df	p
Trust in Information (TR_mean)	4. 60 (1. 13)	4. 69 (1. 13)	-0. 565	187	. 573

Note. *** $p < .001$. Group with personification $n = 92$, group without personification $n = 97$.

The t-test showed no clear gap in trust between the two groups, $t(187) = -0.565$, $p = .573$; H1 is not supported. This result indicates that, in the experimental context of this study, the anthropomorphic design of the AI interface did not significantly enhance users' trust in the information provided by the AI. One reason might be that while the human-like style changes how easy and how useful users find the AI, trust also depends on other things, the formation of trust is also moderated by various factors such as the quality of the information content, users' prior knowledge, and individual differences, resulting in the direct effect of anthropomorphism on trust failing to reach statistical significance (See Table 5).

H2a/H2b: Mediating Effects of Perceived Usefulness and Perceived Usability. H2a looked at whether perceived usefulness acted as a link between anthropomorphism and trust. H2b did the same, but for perceived ease of use. The study used Bootstrap with 5,000 resamples and set the confidence level at 95% was used to test these two parallel mediation paths, with the results shown in Table 6.

Table 6 Findings from the Mediation Effects of H2a/H2b (Bootstrap 5,000 times, 95% CI)

Path	Path Coefficient a	p	Path Coefficient b	p	Indirect effect	Direct effect c'	95% Bootstrap CI
X→PU→TR (H2a)	-0. 418	. 029*	0. 625	<. 001***	-0. 261	0. 084	[-0. 510, -0. 029]
X→PEOU→TR (H2b)	-0. 178	. 336	0. 673	<. 001***	-0. 120	0. 084	[-0. 367, 0. 126]

Note. X = Personification (1 = Yes, 0 = No); TR = Trust in Information; PU means how helpful users think the AI is. PEOU means how easy it feels to use. a = the link from X to the mediator; b = the link from the mediator to TR. (after controlling for X). Direct effect $c' = 0.084$, $p = .442$ (direct effect of personification on TR, controlling for both PU and PEOU). Total effect $c = -0.093$, $p = .573$. * $p < .05$, *** $p < .001$.

Bootstrap mediation test results indicate that the mediation path of Perceived Usefulness (PU) (H2a) is significant: path a1 (Personification $\rightarrow$ PU) = -0.418 , $p = .029$; path b1 (PU $\rightarrow$ TR, controlling for X) = 0.625 , $p < .001$; Indirect effect = -0.261 , 95% Bootstrap CI = $[-0.510, -0.029]$; since the confidence interval stayed away from 0, the mediating effect is significant. This result indicates that perceived usefulness acts as a key link in the process by which personification influences information trust, but the direction is negative—that is, under conditions of personification in this study, perceived usefulness is actually slightly lower, which in turn affects trust scores.

The mediating path of PEOU (H2b) was not clear: Path a2 (Personification $\rightarrow$ PEOU) = -0.178 , $p = .336$; Indirect effect = -0.120 , 95% Bootstrap CI = $[-0.367, 0.126]$, the interval covers 0, and the mediating effect is not significant. H2b is not supported.

H3: Moderating Effect of Anthropomorphism Tendency. H3 asked whether a person's natural tendency to see things as human would change how anthropomorphism affects trust. Hierarchical regression analysis was employed. After centering anthropomorphism (X) and the mean of the tendency toward anthropomorphism (AT), an interaction term ($X \times AT$) was introduced to check how it shaped things (See Table 7).

Table 7. Regression Analysis of the Moderating Effect of AT on the Path from Personification to Trust

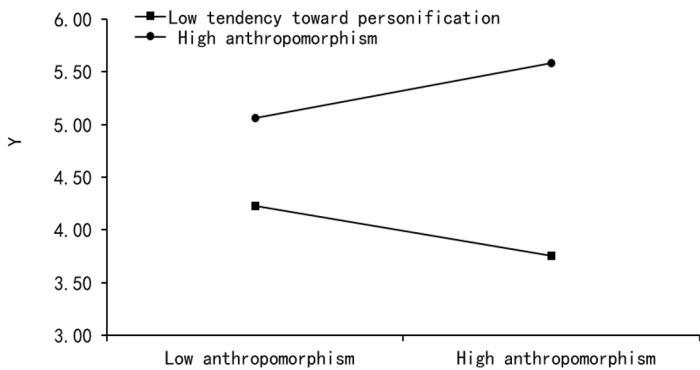
Predictor Variables	B	SE	β	t	p
Constant term	4.658	0.091	—	51.10	< .001
Personification (X, Centralization)	0.012	0.131	.005	0.090	.930
Anthropomorphism (AT, Centralization)	0.665	0.067	.511	9.958	< .001***
$X \times AT$ (Interaction Term)	0.251	0.143	.095	1.760	.080

Note. $R^2 = .324$, $F(3, 185) = 29.58$, $p < .001$. B = the raw coefficient; SE = standard error; β = standardized version. Both independent and moderator variables were mean-centered. * $p < .05$, *** $p < .001$.

The results of the moderation effect of a regression test show that the regression coefficient for the interaction term ($X \times AT$) is $B = 0.251$, $t(185) = 1.760$, $p = .080$. While marginally significant, it did not meet the conventional significance threshold ($p < .05$), and H3 was not statistically supported. However, the tendency to anthropomorphize (AT) itself clearly boosted information trust ($B = 0.665$, $p < .001$), meaning that the stronger an individual's tendency to anthropomorphize, the higher their level of trust in information provided by AI.

To further describe the trends in the interaction pattern, simple slope plots were plotted for the high anthropomorphism (+1 SD) and low anthropomorphism (−1 SD) groups, respectively (Figure 1). The results show that under high anthropomorphism conditions, the trust scores of the anthropomorphism group ($M = 5.57$) edged out the

neutral group by a little ($M = 5.31$), with a positive slope (slope = 0.258); whereas under low anthropomorphism conditions, the trust scores for the anthropomorphism group ($M = 3.77$) fell a bit short of the neutral group ($M = 4.01$), with a negative slope (slope = -0.234). The directions of the two slopes were opposite, consistent with the expected direction of the moderation hypothesis, but the statistical power of the interaction effect was limited.



Note. High/low anthropomorphism conditions represent the mean ± 1 SD. The vertical axis shows the mean information trust score (TR_mean) (1–7 points), and the horizontal axis represents the anthropomorphism condition.

Fig. 1. Simple slope plot of the moderating effect of anthropomorphism tendency (H3) (Picture credit: Original).

H4: The Effect of Information Trust on Willingness to Adopt Information. The last hypothesis (H4) tested whether trusting the AI would make users more willing to follow its advice. A linear regression analysis was conducted with information trust (TR_mean) as the independent variable and willingness to adopt information (AD_mean) as the outcome; Table 8 holds the results.

Table 8 Test of Hypothesis H4: Linear Regression Analysis of Information Trust on Willingness to Adopt Information

Predictor	B	SE	β	t	p
Constant term	1.144	0.148	—	5.960	< .001
Information Trust (TR_mean)	0.774	0.031	.816	19.296	< .001***

Note. $R^2 = .666$, $F(1, 187) = 617.52$, $p < .001$. β = standardized coefficient. *** $p < .001$.

The regression results showed that information trust strongly predicts the intention to adopt information, with $B = 0.774$, $\beta = .816$, $t(187) = 24.850$, $p < .001$, and the model explaining variance $R^2 = .666$. H4 is supported, indicating that the higher the

users' trust in the information provided by AI, the stronger their intention to adopt that information. Trust is the core mechanism driving information adoption, which matches the Information Adoption Model.

3.5 Summary of Hypothesis Testing Results

To facilitate an overall understanding of the testing of each hypothesis in this study, the results of all hypothesis tests as Table 9 shows.

Table 9 Hypothesis Results Summary

Hypothesis	Hypothesis	Test Results	Conclusion
H1	Personification lifts information trust	$t(187) = -0.565, p = .573$	Not supported
H2a	Perceived Usefulness mediates the effect of personification on trust	Indirect effect = -0.261 , 95% CI $[-0.510, -0.029]$	Supported (negative mediation)
H2b	The Effect of Perceived Usability Mediated by Personification on Trust	Mediated effect = -0.120 , 95% CI $[-0.367, 0.126]$	Not supported
H3	Personification Tendencies Moderate the Effect of Personification on Trust	$B = 0.251, t = 1.760, p = .080$	Not significant (marginal)
H4	Information trust has a positive effect on willingness to adopt information	$\beta = .816, t(187) = 19.296, p < .001$	Supported

Note. $p < .05, p < .01, p < .001$.

4 Discussion

This study focused on AI design affects whether users trust and use information. It put special focus on information trust. It also checked whether users thought the AI was helpful and easy to use, and whether those views linked to trust, and whether a person's own tendency to see things as human changed the results. The results did not support the idea that anthropomorphic design boosts trust. Simply making AI sound like a person did not seem to make users trust it more. Firstly, the research findings and existing theories present a certain degree of contradiction and complexity. H1 was not supported. Simply making the AI sound like a person did not make users trust it more. This view is different from that of some early studies. Those early studies thought anthropomorphism can boost trust by raising social presence. There is a possible explanation for this difference. This explanation is based on the specific context of AI information services. Users judge information credibility in a certain way. They may focus more on

the information itself. The key points include content quality, accuracy and logic of the information. They care less about social cues from the AI interface. People may see anthropomorphism as a shallow aesthetic design. They do not take it as a sign of real AI ability. In this case, anthropomorphism can hardly help build deep trust. So future research should look more closely at when and where anthropomorphism works best. However, this study found a significant mediating effect of perceived usefulness (PU) (H2a supported), although in a negative direction. This indicates that under the experimental conditions of this study, anthropomorphic design may instead make some users perceive a decrease in its usefulness, thereby weakening trust. This might be related to the design of the experimental materials and users' stereotype that anthropomorphic AI is more "entertaining than practical". The specific mechanism remains to be studied. In contrast, the strong predictive effect of information trust on adoption intention (H4 supported) is highly consistent with the classic information adoption model, once again verifying that trust is the cornerstone of promoting behavioral intention. Additionally, the main effect of anthropomorphism tendency (AT) on trust is significant, but its moderating effect (H3) only reaches a marginal significance level, indicating that the influence of individual differences may indeed exist, but the statistics of this study are not sufficient to stably capture the interaction effect with the experimental manipulation. The abnormal situation where group D (no anthropomorphism + emotional narrative) has generally higher scores on all variables in the descriptive statistics also suggests that the narrative framework or other uncontrolled factors may have caused unexpected confusion, which is worth further exploration.

These findings point to a few things designers can think about. First, take a function-first method in anthropomorphic design. Practitioners should know the truth. Simply chasing a personified and friendly interface look cannot build deep user trust directly. This is especially true in high-stakes decision-making scenarios. Designers should prioritize ensuring the core value and accuracy of the information content. Second, enhance perceived usefulness. Since perceived usefulness is a key mediator of trust, the design should more clearly convey the efficiency and ability of AI tools to solve practical problems, highlighting their functional value, and avoid anthropomorphic elements overshadowing the main purpose, causing users to perceive it as. Distract users from the core task, which could otherwise lead to perceptions of superficiality. " Third, pay attention to individual differences. Although the moderating effect is not ideal, users with a high anthropomorphism tendency do show a more positive response trend. Where possible, provide customizable interaction styles (such as allowing users to choose "professional mode" or "companion mode"), which may better match the cognitive tendencies and psychological needs of different users.

Finally, this study also falls short in some ways, which points to what future studies can look into. First, the effectiveness of the experimental manipulation is questionable. The manipulation check of the narrative framework was not successful, which may have led to the "narrative type" between-group variable not achieving the expected effect, affecting the internal validity of the experiment. Future research needs to optimize the manipulation of materials and verification methods. Second, the limitations of the sample and context. This study utilized a student sample and was conducted in a single context. The external validity of the conclusions needs to be verified in different user

groups (such as older users, professionals) and diverse social scenarios (such as financial consultation, educational tutoring). Third, the deepening of measurement and mechanism. This study mainly relied on self-report scales. Future research could incorporate objective indicators such as eye-tracking metrics and skin conductance responses to more precisely depict the implicit cognitive and emotional processes of users, thereby mitigating the social desirability bias inherent in self-report scales. Fourth, exploring boundary conditions. Future studies could further investigate when anthropomorphism is beneficial, when it is not, and even when it has adverse side effects, for instance, by examining the moderating effects of variables such as task type (heuristic tasks analytical tasks) and information complexity, thereby constructing a more precise theoretical model.

5 Conclusion

This study focused on how AI's way of talking shapes user trust and how willing users are to take its advice, and also examined whether they would adopt the suggestions. The authors stated that the two models were combined. One of them was the "stimulus - organism - response" model.

The direct effect of anthropomorphic design on trust was insignificant. So just making the AI sound like a person is not enough to build trust. But trust did strongly predict adoption intention. That means trust still really matters in this process. Perceived usefulness was a significant mediator. But the effect was negative. In the study, making the AI sound more human actually made users see it as less useful. That was surprising. Whether users thought the AI was easy to use did not act as a link. And the effect of anthropomorphic tendency was only marginal.

These results matter for both theory and practice. The study shows that the TAM works when it comes to users taking in information. This extends its use to AI systems and helps explain how users react to AI that sounds human. This understanding goes beyond traditional technology adoption. Moreover, it also indicates that anthropomorphic design does not always follow people's expectations.

For designers, this information is not merely suggesting that behaviors resembling those of humans are always correct. In fact, in certain situations, this approach may be the opposite. Research shows that when using anthropomorphic design, caution is necessary, depending on the type of task. For information search tasks, a more neutral style may be more effective. While providing emotional support, a warm tone can be helpful without causing users to have a distorted perception of the usefulness of artificial intelligence.

Authors Contribution

All the authors contributed equally and their names were listed in alphabetical order.

References

1. Schimmelpfennig R, Díaz M, Prabhakaran V, Davani A: Humanlike AI design increases anthropomorphism but yields divergent outcomes on engagement and trust globally. *arXiv* (2025)
2. Go E, Sundar S S: Humanizing chatbots: The effects of visual, identity and conversational cues on humanness perceptions. *Computers in Human Behavior*, **97**, 304–316 (2019)
3. Davis F D: Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, **13**(3), 319–340 (1989)
4. Lu B, Fan W, Zhou M: Social presence, trust, and social commerce purchase intention: An empirical research. *Computers in Human Behavior*, **56**, 225–237 (2016)
5. Spatola N, Marchesi S, Wykowska A: Different models of anthropomorphism across cultures and ontological limits in current frameworks: The integrative framework of anthropomorphism. *Frontiers in Robotics and AI*, **9**, 863319 (2022)
6. Unveiling the impact of AI-chatbot attributes and anthropomorphic cues in e-commerce: *Pakistan Journal of Commerce and Social Sciences*, **19**(3), 441–467 (2025)
7. McCroskey J C: Scales for the measurement of ethos. *Speech Monographs*, **33**(1), 65–72 (1966)
8. Eliasson O, Engström O: Perceived humanness and trust in AI chatbots given different communication styles. *Lund University, Lund* (2025)
9. Waytz A, Cacioppo J, Epley N: Who sees human? The stability and importance of individual differences in anthropomorphism. *Perspectives on Psychological Science*, **5**(3), 219–232 (2010)
10. Al-Abdullatif A M: Modeling students' perceptions of chatbots in learning: Integrating technology acceptance with the value-based adoption model. *Education Sciences*, **13**(11), 1151 (2023)
11. Vasconcelos H, Jörke M, Grunde-McLaughlin M, Gerstenberg T, Bernstein M S, Krishna R: Explanations can reduce overreliance on AI systems during decision-making. *Proceedings of the ACM on Human-Computer Interaction*, **7**(CSCW1), Article 129 (2023)
12. Lopez-Lopez D, Iniesta M B: The impact of conversational AI on consumer decision-making: A systematic review and cluster analysis. *International Journal of Engineering Business Management*, **17**, 18479790251351889 (2025)
13. Camilleri M A, Camilleri A C: The acceptance and usage of ChatGPT: An information adoption model perspective. *Proceedings of the International Conference on Communications and Future Internet*, 61–66 (2024)
14. Eun Go, S. Shyam Sundar, Humanizing chatbots: The effects of visual, identity and conversational cues on humanness perceptions, *Computers in Human Behavior*, Volume 97, 2019, Pages 304–316, ISSN 0747-5632 (2019)

Open Access This chapter is licensed under the terms of the Creative Commons Attribution-NonCommercial 4.0 International License (<http://creativecommons.org/licenses/by-nc/4.0/>), which permits any noncommercial use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if changes were made.

The images or other third party material in this chapter are included in the chapter's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the chapter's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.

