

Original Paper

Trust in Generative AI for Health Information Consumption and the Effect of Learned Dependency: Randomized Controlled Experimental Study

Arif Ahmed, MS; Gondy Leroy, PhD; Agrim Sachdeva, PhD; Philip Harber, PhD; Stephen Rains, PhD; Seokjun Youn, PhD; Prosanta Barai, MS

University of Arizona, Tucson, AZ, United States

Corresponding Author:

Arif Ahmed, MS
University of Arizona
1200 E University Blvd
Tucson, AZ 85721
United States
Phone: 1 9408086282
Email: Arifahmed@arizona.edu

Abstract

Background: Generative AI (GenAI) is increasingly used by health information consumers to interpret medical content and support decision-making. Although these systems provide accessible and timely information, they may also produce inaccurate or misleading outputs. Effective use of GenAI, therefore, depends on users' ability to calibrate trust based on information accuracy. However, little is known about how learned dependency on GenAI (the habitual reliance on AI systems for solving problems) influences trust calibration in health information contexts.

Objective: This study examines how learned dependency on GenAI affects health information consumers' calibration of trust in AI-generated information, and whether text-based visual attention cues (TVCs), such as highlighting critical information in text, mitigate overreliance on incorrect outputs.

Methods: We conducted 2 randomized controlled experiments: the first involved 338 college students, and the second replicated the study with 563 Amazon Mechanical Turk participants. Both studies used a 2 × 2 between-participants design, manipulating (1) information accuracy (correct vs incorrect) and (2) TVCs (text highlight vs no text highlight). Participants evaluated AI-generated health information presented alongside source text. Trust was measured using a multi-item scale, and learned dependency on GenAI was assessed using a validated self-reported measure. Linear regression models were used to examine main and interaction effects.

Results: Across both experiments, information accuracy had a significant positive effect on trust, with participants expressing greater trust in correct than in incorrect AI-generated information (experiment 1: $B=2.107$, 95% CI 1.337-2.878; $P<.001$ and experiment 2: $B=0.203$, 95% CI 0.115-0.290; $P<.001$). Learned dependency was positively associated with trust in both experiments (experiment 1: $B=0.277$, 95% CI 0.033-0.521; $P=.03$ and experiment 2: $B=0.822$, 95% CI 0.715-0.929; $P<.001$), such that users with greater dependency trusted AI outputs more overall. Critically, the interaction between information accuracy and learned dependency was negative and significant in both experiments (experiment 1: $B=-0.399$, 95% CI -0.695 to -0.104 ; $P<.001$ and experiment 2: $B=-0.459$, 95% CI -0.577 to -0.340 ; $P<.001$), indicating that higher dependency reduces users' sensitivity to information inaccuracy. Text highlighting did not significantly affect trust in either experiment, nor did it moderate the relationship between learned dependency and trust.

Conclusions: This study demonstrates that while users generally trust accurate AI-generated health information more than inaccurate information, higher self-reported learned dependency is associated with weaker trust calibration, predicting greater susceptibility to incorrect outputs. TVCs, such as text highlighting, are insufficient to mitigate this effect. These findings highlight the need for more effective design interventions to support critical evaluation and reduce overreliance on GenAI in health information environments.

J Med Internet Res 2026;28:e98326; doi: [10.2196/98326](https://doi.org/10.2196/98326)

Keywords: learned dependency; generative AI; GenAI; trust calibration; text-based visual attention; automation bias; health information

Introduction

Background

AI, particularly generative AI (GenAI), has rapidly transformed the health care landscape, offering new tools for patient education, clinical decision support, and personalized medical recommendations. Since the release of ChatGPT in late 2022, public engagement with GenAI systems has surged, with a growing number of individuals turning to AI-powered chatbots for health-related queries [1,2]. These tools promise convenience and accessibility, but they also produce erroneous outputs, raising critical questions about how users evaluate and calibrate trust in AI-generated health information [3-14]. Studies demonstrate that AI chatbot responses to health queries frequently deviate from clinical guidelines, contain inaccuracies, and may mislead patients who are unaware of these limitations [15,16]. A key concern emerging from this pattern of use is that repeated and largely uncritical engagement with GenAI may foster overreliance. As users experience GenAI outputs as convenient and immediately accessible, they may progressively reduce their own evaluative effort, a process that, over time, can give rise to learned dependency. Learned dependency is a behavioral phenomenon defined as habitual reliance on external AI systems arising from repeated reinforcement, often at the expense of independent thinking and critical evaluation [17]. Research on overreliance on AI demonstrates that users frequently accept incorrect AI outputs, particularly when they perceive the system as competent [18-22]. This pattern, known as automation bias, is strengthened through repeated reliance, as habitual use of AI systems progressively reduces users' critical scrutiny of system outputs [18,22]. In health information contexts, where the consequences of misplaced trust may directly affect health decisions, understanding how learned dependency shapes trust calibration is a pressing but underexplored concern.

Despite the rapid integration of GenAI into health information ecosystems, existing research has not adequately examined how learned dependency shapes trust calibration in this context. Prior work has explored dependency-related phenomena among students [23] and professionals such as clinicians or software developers [24], yet health information consumers whose GenAI use directly bears on health comprehension and decision-making have remained largely overlooked. Prior research on online health information has extensively documented how source credibility, interface design, and information quality shape user trust [25], yet the role of habitual AI dependency in disrupting trust calibration has not been experimentally examined. Critically, no prior experimental study has tested whether learned dependency disrupts users' ability to distinguish accurate from inaccurate AI-generated health information, or whether interface-level interventions such as text-based visual attention cues (TVCs) can counteract such effects. Addressing this gap is essential

for designing AI-mediated health systems that enhance rather than erode users' critical evaluation [26].

Our study addresses this gap through 2 controlled experiments examining how learned dependency on GenAI affects trust calibration in AI-generated health information. We focus on trust as the key outcome because it is the primary cognitive mechanism through which users decide whether to accept or question AI-generated content, and because prior research consistently identifies trust as central to health information engagement and AI adoption [3,27-30]. We further test whether text highlighting, a practical interface-level attention cue, can mitigate dependency-driven overreliance on inaccurate AI outputs. By focusing specifically on trust as the outcome variable, this study makes a deliberate and bounded contribution: we do not examine downstream behavioral intentions or adoption behavior, which represent important but distinct constructs that warrant dedicated investigation in future work.

Theoretical Framework

Our study draws on 2 complementary theoretical frameworks: trust calibration theory and attention theory. Trust calibration theory holds that effective human-automation interaction requires users to adjust their trust in proportion to a system's actual reliability [27]. Well-calibrated trust enables users to rely on a system when it performs correctly and withhold trust when it errs. Critically, miscalibrated trust in either direction has consequences: overreliance leads users to accept erroneous outputs, while underreliance causes them to discard useful information [27,28]. This framework underpins our prediction about the connection between information accuracy and trust.

Building on these theoretical foundations, we introduce learned dependency as a behavioral phenomenon that may disrupt trust calibration. Learned dependency refers to the tendency for users to increasingly rely on AI-generated outputs after repeated exposure and perceived successful use. This construct is theoretically grounded in 2 related literatures. First, research on automation bias demonstrates that repeated reliance on automated systems leads users to defer to system outputs with reduced critical scrutiny, even when errors are present [18,20,21]. This effect is particularly pronounced in health contexts, where AI-generated advice has been shown to significantly reduce diagnostic accuracy when it is incorrect [21]. Second, cognitive offloading theory suggests that users who habitually delegate information evaluation to external systems progressively reduce their own cognitive investment in that task [22,31], a pattern that becomes self-reinforcing over repeated interactions. Together, these mechanisms suggest that learned dependency may be positively associated with overall trust in AI outputs while simultaneously weakening users' sensitivity to information inaccuracy, as dependent users process AI outputs with reduced analytical engagement [32].

Against this backdrop, interface-level attention cues such as text highlighting represent a practical and scalable intervention to support trust calibration. Highlighting key information increases its perceptual salience, which, according to the limited capacity model of mediated message processing [33], should free attentional resources for evaluative processing rather than information search. Eye-tracking research has shown that users consistently direct their visual attention toward highlighted text areas, demonstrating the von Restorff isolation effect for digital highlighting interfaces [34,35]. The von Restorff effect, also known as the isolation effect, suggests that when people are presented with several similar items, the item that stands out as different is more likely to capture attention and be remembered. In digital highlighting interfaces, this principle works by making highlighted information visually distinct, prompting users to pause and process it more carefully. Importantly, among available highlighting methods, bold text has been found to produce the shortest visual search times in graphical user interface tasks, outperforming color-only highlighting on search efficiency [36]. Whether such attentional cueing is sufficient to support trust calibration among highly dependent users, however, remains an open empirical question. We pose our research questions as follows:

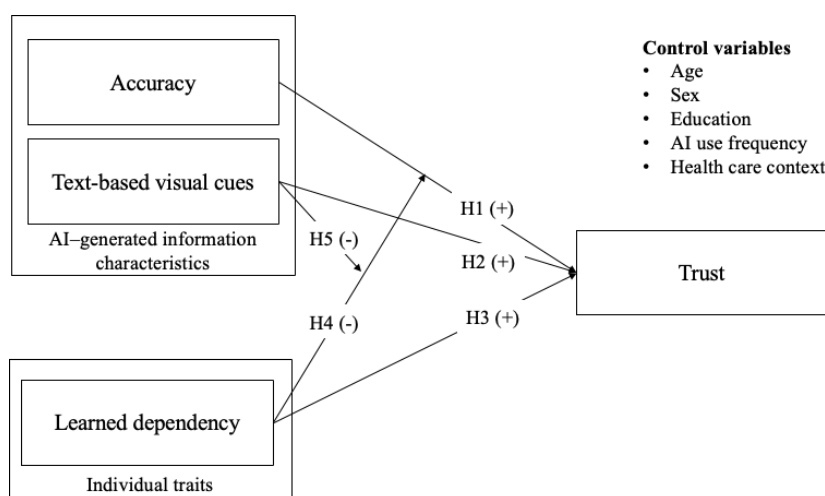
Research question 1: How do the accuracy of AI-generated health information and TVC (eg, highlighting text information) influence health information consumers' trust in AI-generated information?

Research question 2: How does learned dependency on GenAI influence health information consumers' trust in AI-generated information, and how does it moderate the effects of information accuracy and TVC on trust?

Hypotheses Development and Research Framework

Building on the integrative framework presented above, we examine how information accuracy, TVC, and learned dependency on GenAI jointly shape health information consumers' trust in AI-generated content. Rooted in trust calibration theory and attention theory, our framework proposes that habitual reliance on GenAI reduces scrutiny in information evaluation, with important consequences for how users respond to accurate and inaccurate AI-generated outputs. We have developed the research framework (Figure 1), and the corresponding hypotheses are presented below.

Figure 1. Research framework.



Trust calibration theory predicts that users should, under normal conditions, trust accurate information more than inaccurate information. When AI-generated health information is correct, users who evaluate it critically are expected to recognize its quality and report higher trust. We therefore hypothesize the following:

H1: Information accuracy is positively associated with health information consumers' trust in AI-generated health information.

Attention theory suggests that directing users' attention to specific portions of AI-generated health information through TVC can support more careful evaluation of that content. In this study, text highlighting is applied to key portions of the AI-generated health text regardless of whether those portions contain accurate or inaccurate information, depending on the experimental condition. By directing attentional resources

toward the highlighted content, users are better positioned to compare it against the source clinical text and detect potential inaccuracies. Text highlighting is, therefore, proposed as a design intervention that supports trust calibration rather than simply increasing or decreasing trust; that is, users who receive highlighting should be more sensitive to information accuracy than those who do not. Therefore, we propose the following hypothesis:

H2: The presence of TVC (text highlighting) is positively associated with health information consumers' trust in AI-generated health information.

Research on automation bias and cognitive offloading suggests that users who have developed habitual reliance on AI systems tend to engage with AI outputs with reduced analytical effort. As a result, highly dependent users are expected to extend greater overall trust in AI-generated

information, regardless of its accuracy. This leads to the following hypothesis:

H3: Learned dependency on GenAI is positively associated with health information consumers' trust in AI-generated health information.

Although H3 predicts a positive main effect of learned dependency on trust, trust calibration theory further predicts that this dependency may distort the relationship between information accuracy and trust. Specifically, users with high learned dependency are expected to process AI outputs with reduced scrutiny, weakening their ability to distinguish between accurate and inaccurate content. Based on this reasoning, we hypothesize the following:

H4: Higher learned dependency is negatively associated with the positive relationship between information accuracy and health information consumers' trust in AI-generated health information, such that the effect of information accuracy on trust is weaker when learned dependency is higher.

Finally, although TVCs may generally support trust calibration (H2), their effectiveness may be attenuated among users with high learned dependency. Highly dependent users may continue to accept AI outputs with limited scrutiny even when key information is visually highlighted, as habitual reliance reduces their motivation to engage in effortful evaluation. We therefore hypothesize:

H5: TVCs weaken the trust-distorting effect of learned dependency on trust calibration such that the negative moderating effect of learned dependency on the relationship between information accuracy and trust (H4) is attenuated when TVCs are present.

Figure 1 illustrates the research framework guiding this study. Information accuracy and TVC are positioned as AI-generated information characteristics that directly influence trust. Learned dependency is positioned as an individual trait that both directly influences trust (H3) and moderates the effects of information accuracy (H4) and TVC (H5) on trust. The framework reflects a 2×2 experimental design in which information accuracy and TVC are manipulated between participants, while learned dependency is measured as a continuous individual difference variable.

Our study contributes to the literature in several ways. First, it advances research on human-AI interaction by introducing learned dependency as a behavioral phenomenon that can influence trust calibration in AI-generated health information. While prior work has examined trust in AI systems [19,37,38], automation bias [18,20], cognitive forcing functions [19], and GenAI dependency [26] separately, no prior study has experimentally tested their joint influence on trust calibration in health information contexts.

Second, the study tests whether text highlighting is a practical interface-level attention cue that can attenuate the trust-distorting effect of learned dependency. Building on prior work on cognitive forcing functions [19] and typographical cues, specifically TVC [35,36,39-41], the study extends

this literature to GenAI health information contexts and provides evidence on the limits of surface-level attentional interventions.

Third, the study offers practical implications for the design of GenAI systems used in health information environments. As consumers increasingly rely on GenAI tools for medical guidance and interpretation of health information, understanding how trust and dependency shape user behavior becomes critical to the safety of health information consumers [3-6]. The findings highlight the need for interface designs and decision-support mechanisms that promote critical evaluation and reduce the risk of overreliance on AI-generated health information.

Collectively, these contributions deepen our understanding of how GenAI shapes trust calibration in health information contexts, positioning learned dependency as a key mechanism through which habitual AI use influences users' evaluative judgment.

Methods

Study Design

We conducted 2 randomized controlled experiments to examine the impact of GenAI on human trust calibration. The experiments evaluated participants' trust in AI-generated health information using a between-participants design with randomized assignment to conditions.

Ethical Considerations

Our study was approved by the institutional review board (IRB) prior to data collection (IRB protocol: STUDY00002235). All participants provided informed consent before beginning the survey, covering the study's purpose, estimated completion time, voluntary participation, and data storage procedures. No personally identifiable information was collected. Data were stored securely on the Qualtrics platform and were accessible only to the research team. The health information stimuli were obtained from the MIMIC-III (Medical Information Mart for Intensive Care) clinical database, a credentialed-access dataset, under a data use agreement held by the research team. The CHERRIES (Checklist for Reporting Results of Internet E-Surveys) checklist (Checklist 1) covers the online survey or experiment administered via Qualtrics for both experiments. Where responses differ between experiments, both are reported separately.

Experiment 1

Recruitment of Participants

Experiment 1 used a convenience sampling approach. Participants were junior-level students enrolled in a business course at a university. The use of junior-level business students is appropriate because trust calibration and learned dependency are not domain-specific constructs [17, 26]. The course instructor informed students about the study in advance, and participants received 10 extra credit points

upon completion. The study was conducted online, beginning with a pretask demographic questionnaire that collected data on participants' age, sex, education level, race, ethnicity, and major field of study. In addition, participants completed a 6-item GenAI dependency scale, adapted from prior research on media and technology reliance [42] (see [Multimedia Appendix 1](#) for the full item list). Responses were recorded on a 5-point Likert scale ranging from 1 ("strongly disagree") to 5 ("strongly agree"), with the items designed to assess participants' habitual reliance on GenAI tools for decision-making and information gathering. Following the initial survey, participants were assigned to the experimental condition, which assessed trust in AI-curated health information. All instructions and materials were delivered electronically through the study platform (Qualtrics).

Experimental Design

This experiment used a 2×2 between-participants factorial experimental design to examine how information accuracy and TVC influence health information consumers' trust calibration in AI-generated health information and how learned dependency on GenAI moderates these relationships. The 2 independent variables were (1) information accuracy (correct vs incorrect AI-generated health information) and (2) TVC (text highlighting via bold text vs no highlighting). The dependent variable was trust in AI-generated health information. Learned dependency was measured as a continuous individual difference variable. Participants were randomly assigned to one of 20 experimental groups (2 accuracy conditions $\times$ 2 TVC conditions $\times$ 5 distinct health texts), with randomization handled automatically by Qualtrics. Using 5 texts reduces the risk that observed effects are an artifact of one particular stimulus, increasing confidence that the accuracy and highlighting effects generalize across content.

Stimuli

Discharge instructions were sourced from the MIMIC-III clinical database, a credentialed-access, deidentified dataset widely used in clinical research. Five summaries were randomly selected to represent common inpatient medical conditions. Each original summary was presented alongside a version generated by GPT-4o (default version and default settings). For the accuracy manipulation, GPT-4o-generated summaries either preserved the original clinical content (correct condition) or contained deliberate alterations in critical medical details (incorrect condition) at 3 locations in the text. In the TVC condition, key portions of the AI-generated text were highlighted using boldface [35,36]. No formal manipulation check was included, as doing so would contaminate the trust measurement. This is acknowledged as a limitation.

Procedure

Participants completed the study in approximately 10 minutes, with a maximum time limit of 20 minutes. In the trust evaluation task, participants were presented with 2 texts side by side: the original MIMIC-III discharge instructions and the corresponding GPT-4o-generated version ([Figure 2](#)). The AI-generated version varied in 2 key ways depending on the assigned condition: (1) accurate or inaccurate information and (2) presence or absence of bold-text highlighting. Qualtrics automatically randomized participants to one of the 20 experimental groups, with back-navigation disabled. After reviewing both texts, participants completed the trust scales.

Figure 2. Examples of the experimental conditions.

Below are the discharge instructions provided to a right flank pain patient upon discharge from the emergency department. The text on the left-hand side is the original discharge instruction, while the text on the right-hand side is generated by GenAI based on the original information. Discharge Instructions: It was a pleasure taking care of you while you were in the hospital. You were admitted to the hospital for shortness of breath and not making urine. Your heart was not pumping well which caused you to have extra fluid. We gave you medication to help you to urinate out this fluid however your kidneys were not working properly and you needed dialysis to do the job of your kidneys. You got a special IV to be used for dialysis. You will need to continue going to dialysis three times a week. Your blood pressure was also so low so we did not give you your home blood pressure medications while you were in the hospital. Discharge Instruction organized by GenAI It was a pleasure taking care of you while you were in the hospital. You were admitted to the hospital for shortness of breath and not making urine. Your heart was not pumping well which caused you to have extra fluid. We gave you medication to help you urinate out this fluid, however, your kidneys were not working properly and you needed dialysis to do the job of your kidneys. You got a special IV to be used for dialysis. You will need to continue going to dialysis three times a week. Your blood pressure was also so low that we did not give you your home blood pressure medications while you were in the hospital.	Below are the discharge instructions provided to a right flank pain patient upon discharge from the emergency department. The text on the left-hand side is the original discharge instruction, while the text on the right-hand side is generated by GenAI based on the original information. Discharge Instructions: It was a pleasure taking care of you while you were in the hospital. You were admitted to the hospital for shortness of breath and not making urine. Your heart was not pumping well which caused you to have extra fluid. We gave you medication to help you to urinate out this fluid however your kidneys were not working properly and you needed dialysis to do the job of your kidneys. You got a special IV to be used for dialysis. You will need to continue going to dialysis three times a week. Your blood pressure was also so low so we did not give you your home blood pressure medications while you were in the hospital. Discharge Instruction organized by GenAI It was a pleasure taking care of you while you were in the hospital. You were admitted to the hospital for shortness of breath and not making urine. Your heart was not pumping well which caused you to have extra fluid. We gave you medication to help you urinate out this fluid, however, your kidneys were not working properly and you needed dialysis to do the job of your kidneys. You got a special IV to be used for dialysis. You will need to continue going to dialysis three times a week. Your blood pressure was also so low that we did not give you your home blood pressure medications while you were in the hospital.
Condition 1: correct information and no-highlight	Condition 2: correct information and highlight
Below are the discharge instructions provided to a leg skin infection patient upon discharge from the emergency department. The text on the left-hand side is the original discharge instruction, while the text on the right-hand side is generated by GenAI based on the original information. Discharge Instructions: Dear Ms. [**Known lastname 2251**], It was a pleasure taking care of you here at [**Hospital1 18**]. You were admitted to the hospital because you were found to have an infection of the skin on your left leg. While you were in the emergency room your blood pressure was on the lower side so you were in the ICU for a night to make sure it didn't drop further and it was stable (your blood pressure at baseline runs very low and you were asymptomatic throughout your ICU stay). You were then transferred to the regular medical floor where you were stable. Discharge Instructions organized by GenAI Dear Ms. [Known Lastname 2251], It was a pleasure taking care of you here at [Hospital1 18]. You were admitted to the hospital because you were found to have an infection of the lungs on your right leg. While you were in the emergency room, your blood pressure was on the higher side, so you were in the ICU for three nights to make sure it didn't spike further and it was monitored closely (your blood pressure at baseline runs very high, and you were experiencing symptoms throughout your ICU stay). You were then transferred to the surgical ward where you were still under observation.	Below are the discharge instructions provided to a leg skin infection patient upon discharge from the emergency department. The text on the left-hand side is the original discharge instruction, while the text on the right-hand side is generated by GenAI based on the original information. Discharge Instructions: Dear Ms. [**Known lastname 2251**], It was a pleasure taking care of you here at [**Hospital1 18**]. You were admitted to the hospital because you were found to have an infection of the skin on your left leg. While you were in the emergency room your blood pressure was on the lower side so you were in the ICU for a night to make sure it didn't drop further and it was stable (your blood pressure at baseline runs very low and you were asymptomatic throughout your ICU stay). You were then transferred to the regular Discharge Instructions organized by GenAI Dear Ms. [Known Lastname 2251], It was a pleasure taking care of you here at [Hospital1 18]. You were admitted to the hospital because you were found to have an infection of the lungs on your right leg. While you were in the emergency room, your blood pressure was on the higher side, so you were in the ICU for three nights to make sure it didn't spike further and it was monitored closely (your blood pressure at baseline runs very high, and you were
Condition 3: incorrect information and no-highlight	Condition 4: incorrect information and highlight

Measures

Trust in AI-generated health information was measured using a 3-item scale adapted from Lee and See [27] and modified for the GenAI health context: (1) “The GenAI-generated information is trustworthy,” (2) “The GenAI’s output is accurate,” and (3) “The GenAI is well prepared to handle public health data.” Responses were on a 5-point Likert scale (1=strongly disagree to 5=strongly agree). The scale demonstrated acceptable reliability ($\alpha=0.724$, composite reliability [CR]=0.729). Learned dependency was mean-centered prior to constructing interaction terms.

Experiment 2

Overview

To enhance the external validity of experiment 1 and extend the findings to a more diverse population, we conducted a replication study (experiment 2) using Amazon Mechanical Turk (MTurk). Experiment 2 retained the same 2×2 between-participants factorial design (information accuracy $\times$ TVC) and tested an identical set of hypotheses (H1-H5). Three key differences were introduced to address limitations of experiment 1. First, the learned dependency scale was expanded from 6 to 17 items to improve psychometric coverage of the construct. Second, the accuracy manipulation was strengthened from 3 to 6 altered locations within the AI-generated text. Third, participants were recruited from MTurk rather than a university setting, enabling a broader and more occupationally diverse sample.

Recruitment of Participants

Participants were recruited via MTurk, a widely used crowdsourcing platform for behavioral research. Eligible participants were required to be US-based MTurk workers with an approval rating above 95% and were compensated US \$0.50 upon completion. All instructions and materials were delivered electronically through Qualtrics.

Experimental Design

Experiment 2 retained the same 2×2 between-participants factorial design as experiment 1, manipulating information accuracy (correct vs incorrect) and TVC (text highlighting vs no highlighting), with participants randomly assigned to 1 of 4 conditions across 5 health texts (20 experimental groups). The accuracy manipulation was strengthened relative to experiment 1: incorrect information was introduced at 6 locations within the AI-generated health text (compared with 3 in experiment 1), with all 6 locations highlighted in the TVC condition.

Stimuli

The same MIMIC-III-derived discharge instructions and GPT-4o-generated counterparts used in experiment 1 served as stimuli, with the accuracy manipulation extended to 6 altered locations per text (compared to 3 in experiment 1) to create a stronger experimental treatment.

Procedure

The procedure mirrored experiment 1: participants viewed the original discharge instructions alongside the GPT-4o-generated version, which varied in accuracy and in the presence or absence of bold-text highlighting, with condition assignment automatically randomized by Qualtrics.

Measures

Trust in AI-generated health information was measured using a 4-item scale ($\alpha=0.796$, $CR=0.800$, $AVE=0.501$), following confirmatory factor analysis (CFA), which identified and removed 2 items with weak standardized loadings from the original 6-item pool, which combined 3 items adapted from Lee and See [27] with 3 items adapted from Turel and Cui [28]. The retained items were “The GenAI’s output is accurate,” “The GenAI is well prepared to handle public health data,” “I trust AI tools,” and “AI tools are trustworthy.” The learned dependency on GenAI construct was assessed using an expanded 17-item scale combining the 6 items from experiment 1 with 11 items adapted from Goh et al [26], capturing a broader range of dependency-related behaviors and attitudes ($\alpha=0.932$, $CR=0.930$). Responses for both scales were recorded on a 5-point Likert scale (1=strongly disagree to 5=strongly agree). Learned dependency was mean-centered prior to constructing interaction terms.

Reliability and Validity

Trust Measure Items

In experiment 1, the 3-item trust scale adapted from Lee and See [27] demonstrated acceptable internal consistency (Cronbach $\alpha=0.724$, $CR=0.729$). Because a 3-item single-factor model is just-identified ($df=0$), global fit indices cannot be meaningfully assessed; standardized factor loadings ranged from 0.594 to 0.808 ($AVE=0.478$). The AVE falls marginally below the 0.50 threshold, which is acceptable when CR exceeds 0.70 [43]. In experiment 2, the trust scale was expanded to 6 items by adding 3 items adapted from Turel and Cui [28]. Following CFA, items 1 and 4 were removed due to weak standardized loadings (0.511 and 0.420, respectively). The retained 4-item scale demonstrated excellent model fit ($\chi^2_2=1.612$, $\chi^2/df=0.806$, comparative fit index [CFI]=1.000, root mean square error of approximation [RMSEA]=0.000) and strong psychometric properties ($\alpha=0.796$, $CR=0.800$, $AVE=0.501$). Full CFA results for both experiments are reported in [Multimedia Appendix 1](#).

Learned Dependency on GenAI Items

In experiment 1, the 6-item learned dependency scale adapted from prior research on technology reliance [42] demonstrated good model fit ($\chi^2_9=24.995$, $\chi^2/df=2.777$, CFI=0.977, RMSEA=0.068) and acceptable reliability ($\alpha=0.782$, $CR=0.786$, $AVE=0.416$), though items 1 and 2 showed weak standardized loadings (0.109 and 0.489, respectively), noted as a limitation. In experiment 2, the scale was expanded to 17 items by combining the 6 experiment 1 items with 11 items adapted verbatim from Goh et al [26]. This expanded scale demonstrated excellent internal consistency and CR ($\alpha=0.932$, $CR=0.930$) with acceptable model fit (CFI=0.915, RMSEA=0.082), consistent with the original validation by Goh et al [26]. The AVE (0.462) falls marginally below the 0.50 threshold, which is common for large item pools where individual items capture distinct behavioral facets of the construct. Response distributions showed minimal straight-lining (0.8% for the dependency scale; 1.0% for the trust scale), and all variance inflation factors were below 3.2, indicating no multicollinearity concerns. Discriminant validity was assessed using the Fornell-Larcker criterion by comparing the interconstruct correlation between trust and learned dependency against the square root of AVE for each construct (Table S10 in [Multimedia Appendix 1](#)). In experiment 1, discriminant validity was not established (interconstruct $r=0.697$ exceeded $\sqrt{AVE}$ for both constructs), indicating overlap between the trust and dependency measures in this sample. In experiment 2, discriminant validity was supported ($r=0.643$, below $\sqrt{AVE}$ for both constructs), following CFA-based refinement of both scales. Full CFA results are reported in [Multimedia Appendix 1](#).

Results

Demographics

A total of 338 participants completed the first experimental study (Table 1 and Figure 3). The sample was nearly evenly split by sex (n=171, 50.59% male participants vs

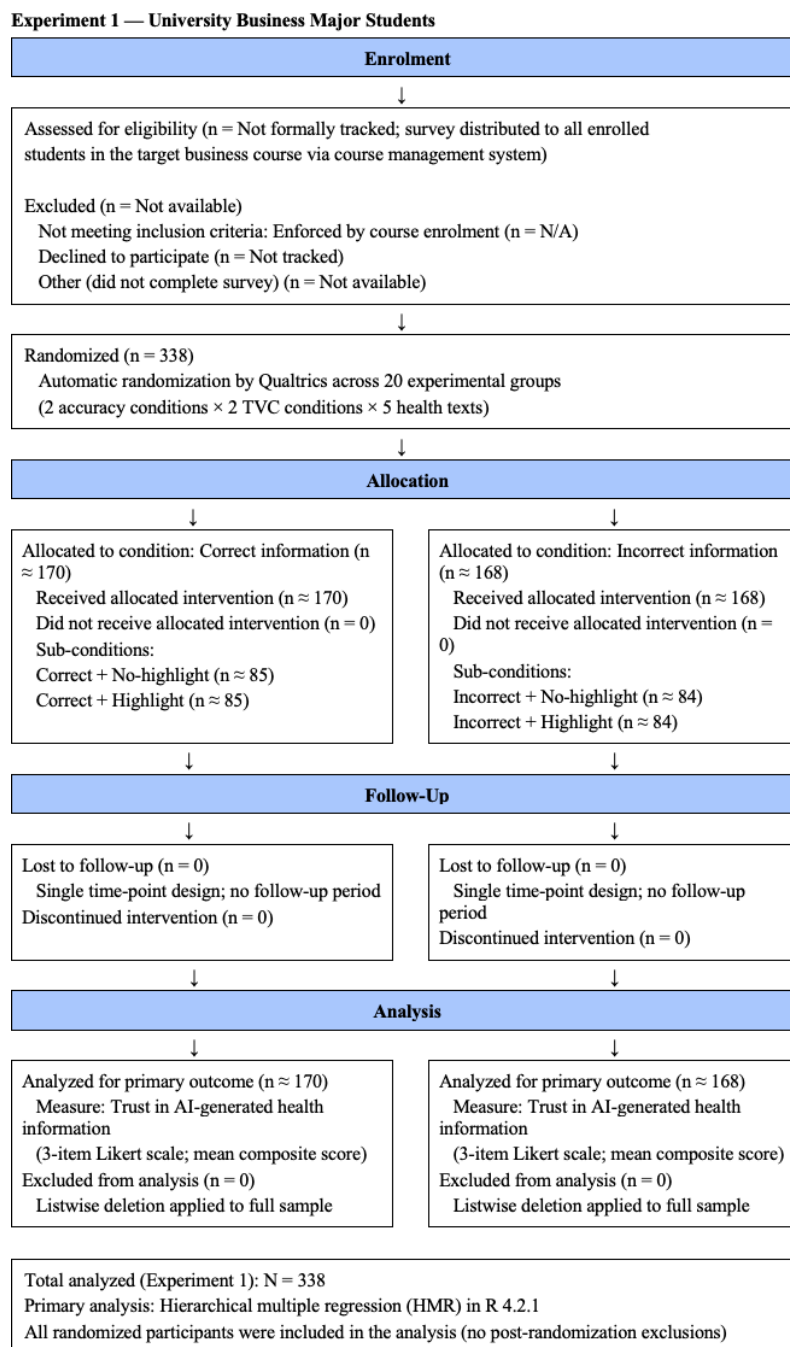
n=167, 49.41% female participants) and predominantly aged 18 to 30 (n=336, 99.41%). Most participants were identified as White (n=266, 78.70%), and 69 (20.41%) participants were identified as Hispanic or Latino. Educational attainment was relatively low, with 80.47% (n=272) having completed only high school. Reflecting the recruitment context, 97.93% (n=331) of participants were business or economics majors.

Table 1. Demographic information.

Characteristics	Experiment 1 (n=338), n (%)	Experiment 2 (n=563), n (%)
Sex		
Male	171 (50.59)	439 (77.98)
Female	167 (49.41)	122 (21.67)
Other	0 (0)	2 (0.36)
Age (y)		
18-20	102 (30.18)	1 (0.18)
21-30	234 (69.23)	348 (61.81)
31-40	2 (0.59)	182 (32.33)
41-50	0 (0)	13 (2.31)
51 years and older	0 (0)	19 (3.37)
Race		
White	266 (78.70)	538 (95.56)
Asian	29 (8.58)	7 (1.24)
Black	7 (2.07)	7 (1.24)
American Indian	3 (0.89)	3 (0.53)
Hawaiian	0 (0)	1 (0.18)
More than one race	23 (6.80)	0 (0)
American Indian, White	0 (0)	4 (0.71)
White, more than one race	0 (0)	2 (0.36)
Asian, White	0 (0)	1 (0.18)
Black, White	0 (0)	1 (0.18)
Unknown	3 (0.89)	0 (0)
Ethnicity		
Hispanic or Latino	69 (20.41)	133 (23.62)
Not Hispanic or Latino	263 (77.81)	426 (75.67)
Unknown	6 (1.78)	4 (0.71)
Education		
Less than high school	0 (0)	2 (0.36)
High school	272 (80.47)	3 (0.53)
Associate's degree	43 (12.72)	13 (2.31)
Bachelor's degree	23 (6.80)	410 (72.82)
Master's degree	0 (0)	125 (22.20)
Other professional degree	0 (0)	1 (0.18)
Doctorate	0 (0)	9 (1.60)
Major field of study		
Science, technology, engineering, and mathematics (STEM)	4 (1.18)	180 (31.97)
Health or medicine	0 (0)	110 (19.54)
Business or economics	331 (97.93)	109 (19.36)
Social sciences or humanities	0 (0)	66 (11.72)
Arts or design	0 (0)	39 (6.93)
Education or teaching	0 (0)	48 (8.53)
Law or criminology	0 (0)	5 (0.89)

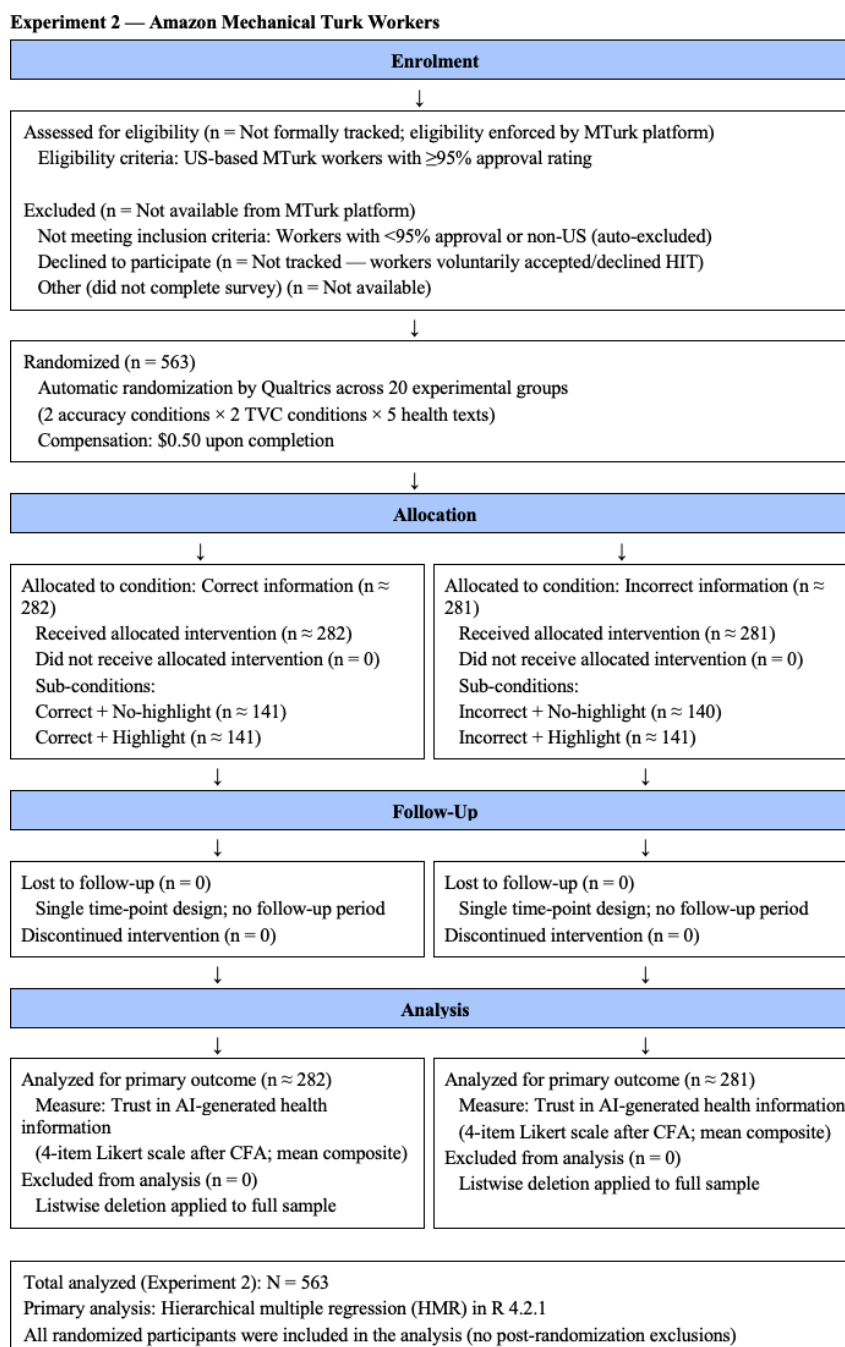
Characteristics	Experiment 1 (n=338), n (%)	Experiment 2 (n=563), n (%)
Agriculture or environmental studies	0 (0)	2 (0.36)
Other	3 (0.89)	4 (0.71)

Figure 3. CONSORT flow diagram - experiment 1.



A total of 563 participants completed the second experimental study (Table 1 and Figure 4). The sample was predominantly male (n=439, 77.98%) and primarily aged 21 to 40 (n=530, 94.14%), with a mean age of 31.7 (SD 6.5) years. Most participants identified as White (n=538, 95.56%), and 23.62% (n=133) identified as Hispanic or Latino. Educational attainment was considerably higher than in experiment 1,

with 72.82% (n=410) holding a bachelor's degree and 22.20% (n=125) holding a master's degree. Major fields of study were more varied, led by STEM (science, technology, engineering, and mathematics; n=180, 31.97%), health or medicine (n=110, 19.54%), and business or economics (n=109, 19.36%).

Figure 4. CONSORT flow diagram - experiment 2.

Data Analysis

All statistical analyses were conducted using R version 4.2.1. Hierarchical multiple regression was the primary analytical approach, with variables entered in 3 blocks: block 1 included control variables (health care context, age, education, and sex), block 2 added the main effects (information accuracy, TVC, and learned dependency), and block 3 added the 2 interaction terms (information accuracy × learned dependency; TVC × learned dependency). Learned dependency was mean-centered prior to constructing interaction terms to reduce multicollinearity. Incremental model fit was assessed using ΔR^2 and the associated *F*-test at each block. Common method bias was assessed using the Harman single-factor test. To assess the psychometric properties of the multi-item

constructs, a CFA was conducted for the trust scale and the learned dependency scale using the lavaan package in R. Model fit was evaluated using standard indices, including standardized root mean square residual, CFI, Tucker-Lewis Index, and RMSEA, and convergent and discriminant validity were assessed using standardized factor loadings, AVE, CR, and Cronbach α . Full CFA results, including factor loadings and fit indices, are reported in [Multimedia Appendix 1](#).

Regression

To examine the effects of information accuracy, TVC, and learned dependency on trust in AI-generated health information, we conducted linear regression analyses. [Table 2](#) presents the regression coefficients, and [Table 3](#) summarizes hypothesis testing results. The results of experiment 1

(Table 2) were consistent with H1; information accuracy had a significant positive effect on trust ($B=2.107$, 95% CI 1.337-2.878; $P<.001$). The results show that participants reported substantially higher trust in accurate AI-generated information than in inaccurate AI-generated information, indicating that, on average, results are consistent with users being more sensitive to information accuracy when comparing correct and incorrect outputs. In contrast, TVC (highlighting) did not have a significant effect on trust ($B=0.149$, 95% CI -0.622 to 0.920 ; $P=.70$), providing no support for H2. This suggests that directing users' attention through

highlighting does not meaningfully influence their trust in AI-generated information. Supporting H3, learned dependency on GenAI was positively associated with trust ($B=0.277$, 95% CI 0.033-0.521; $P=.03$). Individuals with higher levels of dependency exhibited greater overall trust in AI-generated outputs. Among the control variables, most effects were not statistically significant. However, exposure to certain text conditions (health care context 4: $B=-0.449$; $P=.01$; health care context 5: $B=-0.610$; $P<.001$) was associated with lower trust, suggesting potential context-related effects.

Table 2. Regression results.

Variable	Experiment 1		Experiment 2	
	<i>B</i> (95% CI)	<i>P</i> value	<i>B</i> (95% CI)	<i>P</i> value
Intercept	2.213 (0.615 to 3.812)	.007	3.624 (3.337 to 3.911)	<.001
Control				
Health care context 2	−0.230 (−0.576 to 0.116)	.19	0.003 (−0.136 to 0.142)	.97
Health care context 3	−0.165 (−0.508 to 0.179)	.35	0.011 (−0.127 to 0.149)	.88
Health care context 4	−0.449 (−0.797 to −0.101)	.01	0.007 (−0.145 to 0.130)	.92
Health care context 5	−0.610 (−0.954 to −0.265)	<.001	−0.068 (−0.208 to 0.071)	.34
GenAI ^a dependence measure	0.277 (0.033 to 0.521)	.03	0.822 (0.715 to 0.929)	<.001
Education				
High school	0.087 (−0.247 to 0.422)	.61	0.889 (0.281 to 1.497)	.004
Less than high school	0.342 (−1.673 to 2.357)	.74	0.678 (−0.127 to 1.483)	.10
Associate's degree	N/A ^b	NA	−0.228 (−0.520 to 0.065)	.13
Master's degree	N/A	NA	0.018 (−0.093 to 0.129)	.75
Other professional degree	N/A	NA	0.088 (−0.948 to 1.125)	.87
Doctorate	N/A	NA	0.099 (−0.449 to 0.251)	.58
Age	0.002 (−0.058 to 0.062)	.95	0.014 (0.006 to 0.022)	<.001
Sex	0.005 (−0.035 to 0.082)	.85	0.017 (−0.127 to 0.093)	.76
Predictors				
Information: correct	2.107 (1.337 to 2.878)	<.001	0.203 (0.115 to 0.290)	<.001
TVC ^c : No highlight	0.149 (−0.622 to 0.920)	.70	0.013 (−0.074 to 0.101)	.76
Interactions				
Information: correct × GenAI dependence	−0.399 (−0.695 to −0.104)	.008	−0.459 (−0.577 to −0.340)	<.001
TVC: no highlight × GenAI dependence	−0.009 (−0.305 to 0.287)	.95	−0.063 (−0.179 to 0.054)	.29
Information: correct × TVC: no highlight × GenAI dependence	−0.023 (−0.618 to 0.573)	.94	0.142 (−0.095 to 0.379)	.24

^aGenAI: generative AI.

^bNot applicable.

^cTVC: text-based visual attention cue.

Table 3. Results of hypotheses.

Hypothesis	Predicted relationship	Direction	Experiment 1: <i>B</i> ^a value	Result	Experiment 2: <i>B</i> value	Result
H1	Information accuracy → trust	Positive	2.107	Supported ^b	0.203	Supported ^c
H2	TVC ^d → trust	Positive	0.149	Not supported	0.013	Not supported
H3	Learned dependency → trust	Positive	0.277	Supported ^e	0.822	Supported ^c
H4	Information accuracy × learned dependency → trust	Negative	−0.399	Supported ^b	−0.459	Supported ^c

Hypothesis	Predicted relationship	Direction	Experiment 1: B^a value	Result	Experiment 2: B value	Result
H5	Information accuracy × learned dependency × TVC → trust	Positive	−0.023	Not supported	0.142	Not supported

^a B represents unstandardized regression coefficient.

^bSignificance level: $P < .01$.

^cSignificance level: $P < .001$.

^dTVC: text-based visual attention cue.

^eStatistically significant level: $P < .05$.

To assess whether learned dependency influences trust calibration, interaction terms were included in the model (Table 2). Consistent with H4, the interaction between information accuracy and learned dependency was negative and significant ($B = -0.399$, 95% CI -0.695 to -0.104 ; $P = .008$). This finding indicates that higher dependency exhibits reduced sensitivity to information inaccuracy. In other words, highly dependent users are more likely to trust incorrect AI-generated information than less dependent users, reflecting a pattern consistent with diminished trust calibration, as described in the automation bias literature [20]. In contrast, the interaction between TVC, learned dependency, and accuracy was not significant ($B = -0.023$, 95% CI -0.618 to 0.573 ; $P = .94$), providing no support for H5 (Table 3). This suggests that attention-based interventions, such as highlighting critical information, do not mitigate the effect of dependency on trust. Overall, the results provide strong support for the role of learned dependency in shaping trust in AI-generated health information.

Experiment 2 (Table 2) replicated this pattern using the refined 4-item trust scale. Consistent with H1, information accuracy again had a significant positive effect on trust ($B = 0.203$, 95% CI 0.115 – 0.290 ; $P < .001$). TVC again showed no significant effect on trust ($B = 0.013$, 95% CI -0.074 to 0.101 ; $P = .76$), providing no support for H2. Supporting H3, learned dependency was positively and more strongly associated with trust than in experiment 1 ($B = 0.822$, 95% CI 0.715 – 0.929 ; $P < .001$). The interaction between information accuracy and learned dependency was again negative and significant ($B = -0.459$, 95% CI -0.577 to -0.340 ; $P < .001$), providing further support for H4 and indicating that the trust-distorting effect of learned dependency replicated in a larger, more diverse sample with a strengthened accuracy manipulation. The interaction between TVC, learned dependency, and accuracy remained nonsignificant ($B = 0.142$, 95% CI -0.095 to 0.379 ; $P = .24$), again providing no support for H5. Taken together, the cross-sample replication of H1, H3, and H4, alongside the consistent null findings for H2 and H5, strengthens confidence in the robustness of these effects (Table 3).

Although users generally trust accurate information more than inaccurate information, increased dependency weakens this distinction. Furthermore, text highlighting does not significantly influence trust or reduce overreliance among highly dependent users.

Discussion

Principal Findings

This study examined how characteristics of AI-generated health information and users' behavioral reliance on GenAI influence trust in AI-generated health information. By integrating trust calibration theory and attention theory, our findings provide insights into how users evaluate AI-generated health information and how behavioral dependency may shape these evaluations.

First, the results suggest that information accuracy remains an important determinant of trust in AI-generated health information, which is consistent with trust calibration theory [27]. Across both studies, participants generally expressed greater trust in accurate AI-generated responses than in inaccurate ones. This finding indicates that users can calibrate their trust based on the quality of the information provided by AI systems.

However, the findings also reveal that learned dependency on GenAI plays a critical role in shaping users' trust judgments. Individuals who reported higher levels of dependency on AI systems exhibited greater trust in AI-generated health information overall. More importantly, dependency weakened users' ability to differentiate between accurate and inaccurate information. This pattern suggests that dependence on AI tools may reduce users' ability to critically evaluate AI outputs, increasing the likelihood of overreliance on GenAI-generated information. From the perspective of human-automation interaction, this pattern is consistent with automation bias, in which users defer to automated systems even when the outputs may be incorrect, a pattern observed across both studies. These findings are consistent with theoretical mechanisms described in the automation bias literature [18,20,21] and the cognitive offloading perspective [22,31].

The results also provide insights into the effectiveness of the TVC in supporting users' evaluation of AI-generated health information. Although attention theory suggests that directing users' attention to critical information may improve information processing, the TVC intervention did not significantly affect trust calibration across both studies. This finding suggests that interface cues, such as text highlighting directed at specific portions of AI-generated information, may not be sufficient to change how users evaluate its credibility, including in experiment 2, where the TVC was substantially

strengthened (doubled). Users who already depend heavily on AI systems may continue to rely on AI output regardless of such interface cues.

Taken together, these findings highlight an important behavioral tension in AI-assisted health information environments. Although users appear capable of recognizing differences in information accuracy, behavioral reliance on GenAI can weaken this evaluative process. As a result, dependency on AI tools may increase users' vulnerability to trusting inaccurate AI-generated health information. These dynamics raise important considerations for the design and governance of GenAI systems in highly sensitive contexts such as health care, where incorrect information may have adverse consequences for health-related decision-making.

Limitations

Several limitations should be considered when interpreting the findings of this study. First, while the inclusion of experiment 2 using a diverse MTurk population sample substantially improves upon the student-only design of experiment 1, MTurk samples themselves may differ from the broader general population in terms of digital literacy and online task experience. Future research should examine these relationships in clinical populations and among older adults with active health concerns. Second, the attention-based intervention across both experiments focused on text highlighting to direct users' attention to key information. Importantly, the manipulation was strengthened in experiment 2 by highlighting 6 locations rather than 3, yet the null finding for H5 persisted, providing strong evidence that highlighting alone is insufficient to counteract dependency-driven trust distortion regardless of intervention intensity. Future studies should test more substantive cognitive interventions such as uncertainty indicators, confidence cues, or AI reasoning explanations. Third, the scenarios used in both experiments represent controlled and simplified representations of health information. Real-world health decision-making contexts often involve higher stakes, emotional factors such as fear, urgency, stress, or hope, and richer multimodal information formats including colors, images, and video, all of which may influence how individuals interpret and trust AI-generated information. Fourth, this study focuses on trust as the proximal outcome and does not measure downstream behavioral intentions or actual health decisions, which represent important next steps for this research. Fifth, no formal manipulation check was included for the accuracy manipulation, which was a deliberate design choice to avoid contaminating the trust measurement, but this limits independent verification of the manipulation's effectiveness. Sixth, learned dependency was measured through self-report, which may be subject to social desirability bias. Future research should complement self-report measures with behavioral or longitudinal indicators. Seventh, both experiments used a controlled side-by-side presentation that differs from naturalistic health information seeking, limiting ecological validity. Eighth, the visual attention manipulation used bold text only; the null H5 finding should be interpreted as specific to this implementation rather than as evidence that all visual attention interventions would fail.

Although bold-text highlighting is an established form of TVC in the information processing and multimedia learning literature [39-41], it represents a relatively subtle form compared to richer visual modalities such as images, color overlays, animations, or graphics. Accordingly, the conclusions of this study apply specifically to text-based health information delivery contexts and may not generalize to multimodal AI-generated health information that incorporates images, graphics, or other rich visual elements. Ninth, in experiment 1, the trust and learned dependency scales did not demonstrate discriminant validity (interconstruct correlation exceeded $\sqrt{\text{AVE}}$ for both constructs), suggesting these constructs may not have been empirically distinguishable with the shorter scales used in that experiment. This concern was resolved in experiment 2 following CFA-based scale refinement, but it remains a limitation of the experiment 1 measurement model and should be considered when interpreting the H3 and H4 findings from that experiment.

Future Directions

Future research should pursue the following specific directions. First, future studies should test more substantive cognitive interventions beyond text highlighting, specifically uncertainty indicators that display the AI's confidence or uncertainty alongside its output, to determine whether these engage users at a deeper evaluative level. Second, longitudinal designs tracking the same users over extended GenAI use would allow the examination of how dependency develops over time. Third, future studies should examine these relationships in clinical populations with active health concerns. Fourth, behavioral outcome measures such as verification-seeking should complement self-reported trust. Fifth, multimodal stimulus designs incorporating color, images, graphics, and variable typography should be tested to determine whether richer visual interventions overcome the null H5 finding. Sixth, future studies should examine how health literacy, medical knowledge, and digital literacy moderate the dependency-trust relationship and include more diverse and clinically relevant populations.

Conclusions

Across 2 controlled experiments with distinct populations, our findings provide robust empirical evidence that both characteristics of AI-generated information and users' behavioral traits influence trust calibration in GenAI systems used for health information consumption. Information accuracy positively influences trust, and learned dependency consistently increases overall trust while reducing sensitivity to information inaccuracy. Critically, text highlighting failed to counteract dependency-driven trust distortion in either experiment, including experiment 2, where the visual intervention was substantially strengthened by highlighting incorrect information in 6 locations rather than 3. This pattern of null findings, replicated under conditions of stronger treatment, constitutes a substantive finding in its own right: the trust-distorting effect of learned dependency is robust and cannot be mitigated by surface-level attentional cueing alone. These findings highlight the importance of designing AI-enabled health information systems that promote

appropriate trust calibration through more substantive cognitive interventions. As GenAI tools become increasingly integrated into digital health ecosystems, understanding how dependency shapes evaluative judgment will be critical for ensuring safe and responsible use of AI-generated health information [16,44].

Acknowledgments

We would like to acknowledge the contributions of the participants of this study.

During the preparation of this manuscript, Claude Sonnet 4.6 (Anthropic) and Grammarly (Superhuman Platform Inc) were used solely to assist with paraphrasing and grammar checking. The authors reviewed and verified all generated content and take full responsibility for the final manuscript.

Funding

The research reported in this paper was supported by the National Library of Medicine of the National Institutes of Health under award number R01LM011975. The content is solely the responsibility of the authors and does not necessarily represent the official views of the National Institutes of Health.

Data Availability

The study was conducted under the University of Arizona Institutional Review Board (IRB) protocol STUDY00002235, which was determined to be exempt. The dataset does not contain identifiable information collected from human participants. However, according to the data availability statement that we acquired from the IRB for this study, we are prohibited from making the data available to the general public. We have added the consent form as supplementary material. We can only make the study materials and survey materials available, and these are provided as supplementary material.

Authors' Contributions

Conceptualization: AA, PH, SR

Data curation: AA

Formal analysis: AA

Funding acquisition: GL

Methodology: AA, GL, PH, SR

Research framework: AA, AS

Supervision: GL

Validation: AA, GL, PH

Writing – original draft: AA

Writing – review & editing: AA, GL, AS, PH, SR, SY, PB

Conflicts of Interest

None declared.

Multimedia Appendix 1

Summary of measures and measurement model.

[\[DOCX File \(Microsoft Word File\), 25 KB-Multimedia Appendix 1\]](#)

Checklist 1

CHERRIES checklist.

[\[PDF File \(Adobe File\), 169 KB-Checklist 1\]](#)

Checklist 2

CONSORT checklist.

[\[DOCX File \(Microsoft Word File\), 26 KB-Checklist 2\]](#)

References

1. Esmaeilzadeh P, Maddah M, Mirzaei T. Using AI chatbots (e.g., ChatGPT) in seeking health-related information online: the case of a common ailment. *Comput Hum Behav Artif Hum*. Mar 2025;3:100127. [doi: [10.1016/j.chbah.2025.100127](https://doi.org/10.1016/j.chbah.2025.100127)]
2. Ayre J, Cvejic E, McCaffery KJ. Use of ChatGPT to obtain health information in Australia, 2024: insights from a nationally representative survey. *Med J Aust*. Mar 3, 2025;222(4):210-212. [doi: [10.5694/mja2.52598](https://doi.org/10.5694/mja2.52598)] [Medline: [39901778](https://pubmed.ncbi.nlm.nih.gov/39901778/)]
3. Choudhury A, Shamszare H. Investigating the impact of user trust on the adoption and use of ChatGPT: survey analysis. *J Med Internet Res*. Jun 14, 2023;25:e47184. [doi: [10.2196/47184](https://doi.org/10.2196/47184)] [Medline: [37314848](https://pubmed.ncbi.nlm.nih.gov/37314848/)]

4. Ayo-Ajibola O, Davis RJ, Lin ME, Riddell J, Kravitz RL. Characterizing the adoption and experiences of users of artificial intelligence-generated health information in the United States: cross-sectional questionnaire study. *J Med Internet Res*. Aug 14, 2024;26:e55138. [doi: [10.2196/55138](https://doi.org/10.2196/55138)] [Medline: [39141910](https://pubmed.ncbi.nlm.nih.gov/39141910/)]
5. Kadian A, Merhi MI, Gupta M, Ghafoori A, Dennehy D. I can't take this anymore! Understanding the relationship between personality traits and tolerance of generative AI hallucinations. *IEEE Trans Eng Manage*. 2025;72:3864-3876. [doi: [10.1109/TEM.2025.3606352](https://doi.org/10.1109/TEM.2025.3606352)]
6. Bates DW, Levine D, Syrowatka A, et al. The potential of artificial intelligence to improve patient safety: a scoping review. *NPJ Digit Med*. Mar 19, 2021;4(1):54. [doi: [10.1038/s41746-021-00423-6](https://doi.org/10.1038/s41746-021-00423-6)] [Medline: [33742085](https://pubmed.ncbi.nlm.nih.gov/33742085/)]
7. Davenport T, Kalakota R. The potential for artificial intelligence in healthcare. *Future Healthc J*. Jun 2019;6(2):94-98. [doi: [10.7861/futurehosp.6-2-94](https://doi.org/10.7861/futurehosp.6-2-94)] [Medline: [31363513](https://pubmed.ncbi.nlm.nih.gov/31363513/)]
8. Topol E. *Deep Medicine: How Artificial Intelligence Can Make Healthcare Human Again*. Basic Books; 2019. ISBN: 9781541644649
9. Carayon P, Hoonakker P. Human factors and usability for health information technology: old and new challenges. *Yearb Med Inform*. Aug 2019;28(1):71-77. [doi: [10.1055/s-0039-1677907](https://doi.org/10.1055/s-0039-1677907)] [Medline: [31419818](https://pubmed.ncbi.nlm.nih.gov/31419818/)]
10. Fraser H, Coiera E, Wong D. Safety of patient-facing digital symptom checkers. *Lancet*. Nov 24, 2018;392(10161):2263-2264. [doi: [10.1016/S0140-6736\(18\)32819-8](https://doi.org/10.1016/S0140-6736(18)32819-8)] [Medline: [30413281](https://pubmed.ncbi.nlm.nih.gov/30413281/)]
11. Ji Z, Lee N, Frieske R, et al. Survey of hallucination in natural language generation. *ACM Comput Surv*. 2023;55(12):1-38. [doi: [10.1145/3571730](https://doi.org/10.1145/3571730)]
12. Amann J, Blasimme A, Vayena E, Frey D, Madai VI, Precise4Q consortium. Explainability for artificial intelligence in healthcare: a multidisciplinary perspective. *BMC Med Inform Decis Mak*. Nov 30, 2020;20(1):310. [doi: [10.1186/s12911-020-01332-6](https://doi.org/10.1186/s12911-020-01332-6)] [Medline: [33256715](https://pubmed.ncbi.nlm.nih.gov/33256715/)]
13. Longoni C, Bonezzi A, Morewedge CK. Resistance to medical artificial intelligence. *J Consum Res*. Dec 2019;46(4):629-650. [doi: [10.1093/jcr/ucz013](https://doi.org/10.1093/jcr/ucz013)]
14. Bickmore TW, Trinh H, Olafsson S, et al. Patient and consumer safety risks when using conversational assistants for medical information: an observational study of Siri, Alexa, and Google Assistant. *J Med Internet Res*. Sep 4, 2018;20(9):e11510. [doi: [10.2196/11510](https://doi.org/10.2196/11510)] [Medline: [30181110](https://pubmed.ncbi.nlm.nih.gov/30181110/)]
15. Walker HL, Ghani S, Kuemmerli C, et al. Reliability of medical information provided by ChatGPT: assessment against clinical guidelines and Patient Information Quality Instrument. *J Med Internet Res*. Jun 30, 2023;25:e47479. [doi: [10.2196/47479](https://doi.org/10.2196/47479)] [Medline: [37389908](https://pubmed.ncbi.nlm.nih.gov/37389908/)]
16. Yau JYS, Saadat S, Hsu E, et al. Accuracy of prospective assessments of 4 large language model chatbot responses to patient questions about emergency care: experimental comparative study. *J Med Internet Res*. Nov 4, 2024;26:e60291. [doi: [10.2196/60291](https://doi.org/10.2196/60291)] [Medline: [39496149](https://pubmed.ncbi.nlm.nih.gov/39496149/)]
17. Bornstein RF. The dependent personality: developmental, social, and clinical perspectives. *Psychol Bull*. Jul 1992;112(1):3-23. [doi: [10.1037/0033-2909.112.1.3](https://doi.org/10.1037/0033-2909.112.1.3)] [Medline: [1529038](https://pubmed.ncbi.nlm.nih.gov/1529038/)]
18. Parasuraman R, Riley V. Humans and automation: use, misuse, disuse, abuse. *Hum Factors*. Jun 1997;39(2):230-253. [doi: [10.1518/001872097778543886](https://doi.org/10.1518/001872097778543886)]
19. Bućinca Z, Malaya MB, Gajos KZ. To trust or to think: cognitive forcing functions can reduce overreliance on AI in AI-assisted decision-making. *Proc ACM Hum Comput Interact*. 2021;5(CSCW1):1-21. [doi: [10.1145/3449287](https://doi.org/10.1145/3449287)]
20. Abdelwanis M, Alarafati HK, Tammam MMS, Simsekler MCE. Exploring the risks of automation bias in healthcare artificial intelligence applications: a Bowtie analysis. *J Saf Sci Resil*. Dec 2024;5(4):460-469. [doi: [10.1016/j.jnlssr.2024.06.001](https://doi.org/10.1016/j.jnlssr.2024.06.001)]
21. Gaube S, Suresh H, Raue M, et al. Do as AI say: susceptibility in deployment of clinical decision-aids. *NPJ Digit Med*. Feb 19, 2021;4(1):31. [doi: [10.1038/s41746-021-00385-9](https://doi.org/10.1038/s41746-021-00385-9)] [Medline: [33608629](https://pubmed.ncbi.nlm.nih.gov/33608629/)]
22. Passi S, Vorvoreanu M. Overreliance on AI: literature review. Microsoft Corporation; 2022. URL: <https://www.microsoft.com/en-US/research/wp-content/uploads/2022/06/Aether-Overreliance-on-AI-Review-Final-6.21.22.pdf> [Accessed 2026-08-26]
23. Giordano V, Spada I, Chiarello F, Fantoni G. The impact of ChatGPT on human skills: a quantitative study on Twitter data. *Technol Forecast Soc Change*. Jun 2024;203:123389. [doi: [10.1016/j.techfore.2024.123389](https://doi.org/10.1016/j.techfore.2024.123389)]
24. Chu Y, Liu P. Automation complacency on the road. *Ergonomics*. Nov 2023;66(11):1730-1749. [doi: [10.1080/00140139.2023.2210793](https://doi.org/10.1080/00140139.2023.2210793)] [Medline: [37139680](https://pubmed.ncbi.nlm.nih.gov/37139680/)]
25. Sillence E, Briggs P, Harris PR, Fishwick L. How do patients evaluate and make use of online health information? *Soc Sci Med*. May 2007;64(9):1853-1862. [doi: [10.1016/j.socscimed.2007.01.012](https://doi.org/10.1016/j.socscimed.2007.01.012)] [Medline: [17328998](https://pubmed.ncbi.nlm.nih.gov/17328998/)]
26. Goh AYH, Hartanto A, Majeed NM. Generative artificial intelligence dependency: scale development, validation, and its motivational, behavioral, and psychological correlates. *Comput Hum Behav Rep*. Dec 2025;20:100845. [doi: [10.1016/j.chbr.2025.100845](https://doi.org/10.1016/j.chbr.2025.100845)]

27. Lee JD, See KA. Trust in automation: designing for appropriate reliance. *Hum Factors*. 2004;46(1):50-80. [doi: [10.1518/hfes.46.1.50.30392](https://doi.org/10.1518/hfes.46.1.50.30392)] [Medline: [15151155](https://pubmed.ncbi.nlm.nih.gov/15151155/)]
28. Turel O, Cui T. Apologizing artificial intelligence: designing and evaluating effective AI apologies after errors. *AI & Soc*. 2026;41(8):7347-7364. [doi: [10.1007/s00146-026-03067-w](https://doi.org/10.1007/s00146-026-03067-w)]
29. Kauttonen J, Rousi R, Alamäki A. Trust and acceptance challenges in the adoption of AI applications in health care: quantitative survey analysis. *J Med Internet Res*. Mar 21, 2025;27:e65567. [doi: [10.2196/65567](https://doi.org/10.2196/65567)] [Medline: [40116853](https://pubmed.ncbi.nlm.nih.gov/40116853/)]
30. Wichmann J, Gesk TS, Leyer M. Acceptance of AI in health care for short- and long-term treatments: pilot development study of an integrated theoretical model. *JMIR Form Res*. Jul 18, 2024;8:e48600. [doi: [10.2196/48600](https://doi.org/10.2196/48600)] [Medline: [39024565](https://pubmed.ncbi.nlm.nih.gov/39024565/)]
31. Kahneman D. *Attention and Effort*. Prentice-Hall; 1973. ISBN: 0130505188
32. Metzger MJ, Flanagin AJ. Credibility and trust of information in online environments: the use of cognitive heuristics. *J Pragmat*. Dec 2013;59:210-220. [doi: [10.1016/j.pragma.2013.07.012](https://doi.org/10.1016/j.pragma.2013.07.012)]
33. Lang A. The limited capacity model of mediated message processing. *J Commun*. Mar 2000;50(1):46-70. [doi: [10.1111/j.1460-2466.2000.tb02833.x](https://doi.org/10.1111/j.1460-2466.2000.tb02833.x)]
34. Steinfeld N. How do users examine online messages to determine if they are credible? An eye-tracking study of digital literacy, visual attention to metadata, and success in misinformation identification. *Soc Media Soc*. 2023;9(3). [doi: [10.1177/20563051231196871](https://doi.org/10.1177/20563051231196871)]
35. Chi EH, Gumbrecht M, Hong L. Visual foraging of highlighted text: an eye-tracking study. In: Jacko JA, editor. *Human-Computer Interaction HCI Intelligent Multimodal Interaction Environments HCI 2007*. Springer; 2007:589-598. [doi: [10.1007/978-3-540-73110-8_64](https://doi.org/10.1007/978-3-540-73110-8_64)]
36. Doi T. Effective highlighting modes of graphical user interfaces in visual information search. *Univ Access Inf Soc*. 2023;22(1):111-119. [doi: [10.1007/s10209-021-00827-x](https://doi.org/10.1007/s10209-021-00827-x)]
37. Turel O, Kalhan S. Prejudiced against the machine? Implicit associations and the transience of algorithm aversion. *MIS Q*. Dec 1, 2023;47(4):1369-1394. [doi: [10.25300/MISQ/2022/17961](https://doi.org/10.25300/MISQ/2022/17961)]
38. Alarcon GM, Capiola A. Explicating the trust process for effective human interaction with artificial intelligence and machine learning systems. *Front Comput Sci*. 2025;7:1662185. [doi: [10.3389/fcomp.2025.1662185](https://doi.org/10.3389/fcomp.2025.1662185)]
39. Fowler RL, Barker AS. Effectiveness of highlighting for retention of text material. *J Appl Psychol*. 1974;59(3):358-364. [doi: [10.1037/h0036750](https://doi.org/10.1037/h0036750)]
40. Lorch EP, Klusewitz MA. Effects of typographical cues on reading and recall of text. *Contemp Educ Psychol*. Jan 1995;20(1):51-64. [doi: [10.1006/ceps.1995.1003](https://doi.org/10.1006/ceps.1995.1003)]
41. Schneider S, Beege M, Nebel S, Rey GD. A meta-analysis of how signaling affects learning with media. *Educ Res Rev*. Feb 2018;23:1-24. [doi: [10.1016/j.edurev.2017.11.001](https://doi.org/10.1016/j.edurev.2017.11.001)]
42. Ye JH, Zhang M, Nong W, Wang L, Yang X. The relationship between inert thinking and ChatGPT dependence: an I-PACE model perspective. *Educ Inf Technol*. Feb 2025;30(3):3885-3909. [doi: [10.1007/s10639-024-12966-8](https://doi.org/10.1007/s10639-024-12966-8)]
43. Fornell C, Larcker DF. Evaluating structural equation models with unobservable variables and measurement error. *J Mark Res*. Feb 1981;18(1):39-50. [doi: [10.1177/002224378101800104](https://doi.org/10.1177/002224378101800104)]
44. Ayers JW, Poliak A, Dredze M, et al. Comparing physician and artificial intelligence chatbot responses to patient questions posted to a public social media forum. *JAMA Intern Med*. Jun 1, 2023;183(6):589-596. [doi: [10.1001/jamainternmed.2023.1838](https://doi.org/10.1001/jamainternmed.2023.1838)] [Medline: [37115527](https://pubmed.ncbi.nlm.nih.gov/37115527/)]

Abbreviations

CFA: confirmatory factor analysis
CFI: comparative fit index
CHERRIES: Checklist for Reporting Results of Internet E-Surveys
CR: composite reliability
GenAI: generative AI
IRB: institutional review board
MIMIC-III : Medical Information Mart for Intensive Care
MTurk: Amazon Mechanical Turk
RMSEA: root mean square error of approximation
STEM: science, technology, engineering, and mathematics
TVC: text-based visual attention cue

Edited by Ivan Steenstra; peer-reviewed by Goodness Nzeigwe, Mariana de Oliveira, Sajid Anwar, Suwen Ge, Yue Luo; submitted 15.Apr.2026; final revised version received 06.Aug.2026; accepted 07.Aug.2026; published 16.Sep.2026

Please cite as:

Ahmed A, Leroy G, Sachdeva A, Harber P, Rains S, Youn S, Barai P

Trust in Generative AI for Health Information Consumption and the Effect of Learned Dependency: Randomized Controlled Experimental Study

J Med Internet Res 2026;28:e98326

URL: <https://www.jmir.org/2026/1/e98326>

doi: [10.2196/98326](https://doi.org/10.2196/98326)

© Arif Ahmed, Gondy Leroy, Agrim Sachdeva, Philip Harber, Stephen Rains, Seokjun Youn, Prosanta Barai. Originally published in the Journal of Medical Internet Research (<https://www.jmir.org>), 16.Sep.2026. This is an open-access article distributed under the terms of the Creative Commons Attribution License (<https://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution, and reproduction in any medium, provided the original work, first published in the Journal of Medical Internet Research (ISSN 1438-8871), is properly cited. The complete bibliographic information, a link to the original publication on <https://www.jmir.org/>, as well as this copyright and license information must be included.