

Original Paper

Effects of AI Assistance Timing on Pharmacists' Trust in Automated Pill Recognition Technology: Within-Participants Experimental Study

Jin Yong Kim¹, PhD; Brigid Rowell², MA; Megan Whitaker², MHI; Qiyuan Chen¹, MSE; Raed Al Kontar¹, PhD; Corey Lester², PharmD, PhD; Xi Jessie Yang¹, PhD

¹Industrial and Operations Engineering, University of Michigan, Ann Arbor, MI, United States

²College of Pharmacy, University of Michigan, Ann Arbor, MI, United States

Corresponding Author:

Xi Jessie Yang, PhD
Industrial and Operations Engineering
University of Michigan
1640 IOE, 1205 Beal Ave
Ann Arbor, MI 48109-1382
United States
Phone: 1 7347630541
Email: xijyang@umich.edu

Abstract

Background: Image-based pill verification systems demonstrate high model accuracy. However, their effectiveness in pharmacy practice depends on how pharmacists interact with the AI output. The timing of AI advice is one design factor that influences these interactions, yet its impact on pharmacists' moment-to-moment trust dynamics during medication verification requires further investigation.

Objective: This study aims to investigate how the timing and conditionality of AI assistance shape pharmacists' trust dynamics during medication verification.

Methods: Between April and December 2024, 50 licensed pharmacists completed a browser-based simulated medication dispensing task with 2 AI types: ex-ante advice (AI advice given concurrently with clinical information) and ex-post advice (AI advice given after an initial diagnosis). Ex-post advice was further divided into the involved ex-post and the not-involved ex-post conditions. The experiment used a within-participants design with varying AI types and AI recognition patterns (right fill-correct recognition, right fill-incorrect recognition, wrong fill-correct recognition, and wrong fill-incorrect recognition). The primary outcomes were trust adjustment magnitude and trust adjustment, which were analyzed using mixed-effects linear regression models.

Results: Trust adjustment magnitude differed significantly across AI assistance conditions. The involved ex-post condition led to the highest magnitude of trust adjustment, followed by ex-ante advice (mean difference 9.65, 95% CI 8.55-10.75; $P<.001$), with the not-involved ex-post condition showing the lowest magnitude (mean difference 3.33, 95% CI 2.63-4.03; $P<.001$). Analysis by each recognition pattern revealed significant differences in trust adjustment when the right drugs were incorrectly rejected. In this pattern, the involved ex-post condition led to larger trust decrements than both ex-ante advice (mean difference -2.13, 95% CI -3.83 to -0.44; $P=.008$) and not-involved ex-post conditions (mean difference -7.32, 95% CI -13.69 to -.95; $P=.009$). A marginal difference was observed between ex-ante advice and not-involved ex-post conditions (mean difference -5.19, 95% CI -11.55 to 1.17; $P=.051$). No significant differences were observed for other recognition patterns.

Conclusions: Both the timing and conditionality of AI assistance influenced pharmacists' trust dynamics. Disagreement-based AI interventions (involved ex-post) that incorrectly challenged pharmacists led to a substantial trust decrement, whereas the not-involved AI intervention resulted in more stable trust fluctuations. These findings highlight the importance of designing AI systems that align intervention strategies with user expertise and task demands to foster appropriate trust in safety-critical workflows.

Trial Registration: ClinicalTrials.gov NCT06245044; <https://clinicaltrials.gov/study/NCT06245044>

Keywords: artificial intelligence; human-computer interaction; medication verification; medication errors; trust in automation

Introduction

Background

Pharmacists play a critical role in ensuring that patients receive the correct medication, particularly during medication verification, in which the medication filled inside the bottle must match the prescription label [1]. Errors at this stage, such as dispensing the wrong drug, strength, or formulation, are a major source of preventable adverse drug events, with error rates that can exceed 1% in both community and hospital settings [2-4]. These errors adversely impact patient safety and trigger emergency room visits and hospitalizations [5]. Contributing factors include high workload, fatigue, technology limitations, lack of trust in automation, and the high cognitive demands placed on pharmacists who must manage multiple tasks simultaneously [6-9]. Because medication verification is both labor-intensive and cognitively demanding, tools that reduce dispensing errors have the potential to substantially improve patient outcomes and decrease health care costs [10].

To mitigate these risks, image-based pill verification systems have been introduced [11-15]. Advances in deep learning have further improved the performance of these systems, with recent models achieving accuracies greater than 98% [16]. However, the successful integration of AI systems into pharmacy workflows depends not only on model accuracy but also on how humans perceive, trust, interpret, and act on AI output [8,17,18]. Prior human factors and human-computer interaction (HCI) research highlights two design dimensions that strongly influence human-AI interaction: (1) the amount and type of model information presented and (2) the timing of AI assistance [19,20]. These dimensions shape cognitive load, situational awareness, and the development of trust in autonomous and AI systems, defined as the belief that an autonomous or AI agent will support an individual's goals under conditions of uncertainty [21].

The amount and type of model information can play a key role in promoting clinician understanding of clinical decision support systems (CDSSs). To overcome the barrier of the “black-box” nature of many of these CDSSs, in which model outputs provide little insight into how predictions are generated, research has emphasized the importance of enhancing system transparency through explanatory feedback and uncertainty visualizations [9,22-26]. For instance, Kompa et al [26] investigated how different levels of explanation in a CDSS supporting balance disorder diagnosis influenced health care professionals' trust and reliance. Participants used either a selective CDSS that presented examination results alone or a comprehensive CDSS that additionally provided explanatory context. Their findings showed that explanations positively affected trust, whereas limited explanations led users to question the system's reliability [26]. Kim et al [9]

examined how presenting model uncertainty information in a medication verification AI affected pharmacists' trust. Their pill characteristic plots conveyed match-mismatch information based on visual attributes such as imprint, color, shape, and score, and the uncertainty plot illustrated the predicted probability of each set of predictions summarized in a histogram [11]. By manipulating the presence of uncertainty visualization as a between-participants factor, they found that communicating uncertainty improved pharmacists' end trust, suggesting that transparency fosters more calibrated trust and active engagement with AI systems [9].

Another critical design factor influencing human-AI interaction is the timing of AI advice [20]. In clinical decision-making, AI systems may present recommendations before the human makes the initial decision (ex-ante AI advice) or after the human forms an initial judgment (ex-post AI advice) [19,27]. Ex-ante systems provide continuous feedback throughout the task [28,29], supporting real-time integration of AI insights, but increasing the risk of cognitive overload or anchoring when AI outputs are incorrect [30-32]. In contrast, ex-post systems require the human decision maker to commit to an initial decision before receiving AI advice [33-36]. While this approach can preserve human decision autonomy and reduce the anchoring effect, it may exacerbate confirmation bias [37-39], contribute to algorithm aversion [35,40,41], and introduce workflow interruptions when AI intervenes unnecessarily [42].

Empirical findings on the optimal timing of AI assistance remain mixed. Yin et al [27] investigated how the timing of AI advice affects doctors' diagnostic decision-making. In their experiment, they recruited and asked physicians to diagnose clinical cases. The authors compared ex-post advice with ex-ante advice. They found that physicians engaged more deeply with ex-post advice, resulting in higher diagnostic accuracy (task performance) compared with ex-ante advice. Conversely, Fogliato et al [43] explored a human-AI workflow sequence with veterinary radiologists, who reviewed X-ray images and made diagnostic judgments under two conditions: (1) a 1-step workflow (ex-ante advice), where AI inferences were shown before radiologists' diagnosis and (2) a 2-step workflow (ex-post advice), where radiologists made an initial diagnosis first, then saw the AI's recommendations. The authors reported that radiologists' diagnostic performance was better in the ex-ante advice system compared to the ex-post advice system. With the ex-post advice help, radiologists rarely revised their initial judgments and rated the AI as less useful [43].

Properties of Trust Dynamics

Trust in autonomous and AI systems has been extensively studied across domains, including health care [9,17,21]. Trust is a critical determinant of how people rely on autonomous systems and is strongly influenced by system performance [44-46].

Most existing literature has adopted a snapshot approach to trust measurement, typically assessing trust through posttask or end-of-experiment questionnaires. While informative, this approach overlooks that trust is a dynamic variable and evolves over the course of repeated, moment-to-moment interactions with autonomous systems [45,47,48].

Yang et al [45] identified and summarized 3 key trust dynamics properties. First, continuity suggests that trust at a given moment (i) is strongly associated with trust at the immediately preceding moment ($i - 1$). Trust tends to increase following automation successes and decrease following failures [9,46]. Second, negativity bias highlights the asymmetric impact of autonomy performance, wherein trust is difficult to build but can be lost quickly. Prior studies have shown that negative experiences with automation failures exert a stronger influence on trust than positive experiences [8,49-51]. Third, stabilization refers to the tendency for trust to stabilize over repeated interactions with the same system, as users gain familiarity and form more stable expectations [52].

Present Study

Building on these findings, the present study investigates how the timing of AI assistance shapes pharmacists' trust dynamics during medication verification. We examine 2 modes of AI engagement: ex-post advice (AI advice given after an initial diagnosis) and ex-ante advice (AI advice given concurrently with clinical information). The ex-ante advice continuously displays predictive information throughout the task, and the ex-post advice intervenes only when its prediction conflicts with the pharmacist's initial decision. Ex-ante advice may support continuous decision-making but risks cognitive overload and anchoring, whereas ex-post advice may reduce unnecessary attentional demands but risks disruptive false alarms [30,32-34,36].

By comparing these timing strategies, we aim to identify how different modes of AI engagement influence trust adjustment in a safety-critical and expert-driven medication verification task.

In addition, prior empirical validations of trust dynamics properties (continuity, negativity bias, and stabilization) have relied primarily on simplified tasks and nonexpert participant populations [8,45,48]. By examining licensed pharmacists engaged in realistic medication verification tasks, the present study extends prior work and strengthens the ecological validity of trust dynamics properties in a safety-critical health care context.

Methods

Ethical Considerations

This study was exempt from institutional review board oversight by the University of Michigan (HUM#00241223). All participants provided electronic informed consent prior to the experiment, and data were collected anonymously. Participants received US \$100 upon completing the study.

Participants

Pharmacists were recruited via email through the Minnesota Pharmacy Practice-Based Research Network, Pharmacy Practice Enhancement and Active Research Link in Wisconsin, and the University of Michigan College of Pharmacy Pharmacist Preceptor Network. Participants were eligible for the study if they were licensed in the United States, were at least 18 years old, and had access to a laptop or desktop computer with a webcam. Pharmacists were excluded if they needed assistive technology to operate a computer, wore eyeglasses with more than 1 power, had uncorrected cataracts, intraocular implants, glaucoma, or permanently dilated pupils, and/or had eye movement or alignment abnormalities (eg, lazy eye, strabismus, nystagmus). Fifty licensed US pharmacists participated in the study between April and December 2024.

AI Model

The AI model used in this study is a Bayesian neural network (BNN) designed to predict the National Drug Code (NDC) for dispensed pills while estimating the prediction uncertainty [11,16]. This was achieved by applying random dropout to a ResNet-34 convolutional neural network architecture [53, 54]. Instead of outputting a single probability distribution for each class, the model generated 50 probability vectors for each image. The final NDC prediction was determined by averaging those vectors and selecting the class with the highest mean probability.

The model was trained on a dataset of 432,974 images representing 352 distinct NDC classes from a mail-order pharmacy in the United States, which uses a robotic system to count pills, fill and label vials, capture images of the contents, and seal the vials [55]. The dataset contains 1 year of robot-captured images of solid oral medication (ie, tablets and capsules) inside prescription vials, with each image associated with its corresponding NDC and annotated with attributes such as color, shape, size, manufacturer, tablet scoring, and imprint. The number of available images for each NDC varied from 3 to 12105, with a median of 540 (IQR 257-1291).

The medications in the dataset were categorized into 12 distinct colors: white (42.1%), yellow (12.3%), pink (9.1%), orange (7.1%), multicolor (5.9%), green (5.2%), red (5.2%), blue (4.8%), brown (3.8%), purple (3.1%), turquoise (0.7%), and gray (0.7%). Additionally, seven unique shapes were categorized: round (49.6%), oval (33.4%), capsule (16.2%), hexagon-6-sided (0.4%), triangle (0.3%), trapezoid (0.1%), and pentagon-5-sided (0.04%).

Experimental Apparatus

Overview

The study utilized Labvanced, a browser-based experimental platform, to conduct simulated medication dispensing tasks. Participants accessed the testbed through a web browser on their personal computers, which enabled the software to collect eye-tracking data through the webcam while also recording response decisions and reaction times for each task.

Experimental Task

In this experiment, participants performed a simulated medication verification task supported by varying types of imperfect automated pill recognition systems [9,12,13]. The user interface displayed a reference image of prescription medication, the prescription details, and an image of the filled medication. In the 2 AI-assistance conditions, the AI-generated decision support was also presented. The objective was to determine whether the filled medication matched the prescription (ie, reference image). If the 2 images matched, the correct response was to click “accept”; if they differed, the appropriate response was to click “reject.”

The experimental stimuli, including the reference NDC and corresponding filled medication images, were carefully selected from the dataset of 432,974 images. The selection

prioritized a broad representation of medication colors and shapes while excluding blurry images to ensure clarity. To reduce potential learning effects, each reference NDC was presented no more than twice during the experiment. To ensure sufficient exposure to both correct and incorrect AI predictions, we constructed a 300-image experimental dataset (representing 154 distinct NDCs) by selectively sampling from the larger dataset to increase the proportion of naturally occurring incorrect predictions, resulting in an overall AI accuracy of approximately 79%. NDCs could occur no more than twice to prevent bias. AI outputs were not artificially modified, and all predictions reflect the model’s original performance. Because this sampling approach preserved the model’s inherent error patterns, prediction errors were not uniformly distributed across medication characteristics (see Table 1).

Table 1. Distribution of medication characteristics for the experimental dataset (n=300).

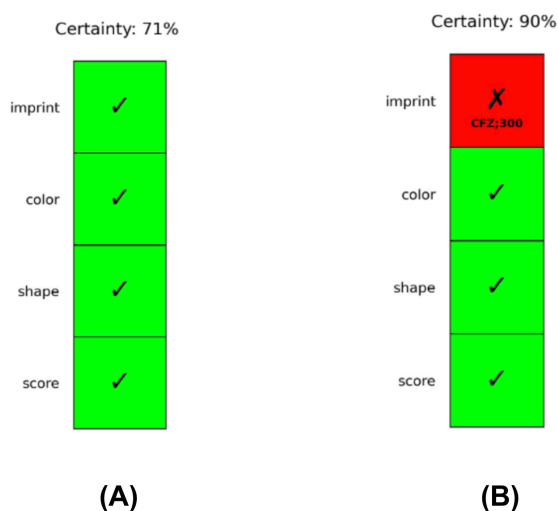
Characteristics	Filled, n (%)	Reference, n (%)
Color		
White	154 (51)	160 (53)
Yellow	42 (14)	34 (11)
Pink	24 (8)	22 (7)
Orange	21 (7)	25 (8)
Multicolor	4 (1)	4 (1)
Green	14 (5)	14 (5)
Red	7 (2)	7 (2)
Blue	12 (4)	12 (4)
Brown	14 (5)	14 (5)
Purple	4 (1)	4 (1)
Turquoise	4 (1)	4 (1)
Shape		
Round	181 (60)	187 (62)
Oval	80 (27)	72 (24)
Capsule	35 (12)	37 (12)
Hexagon (6-sided)	2 (1)	2 (1)
Triangle	2 (1)	2 (1)

To characterize the nature of AI errors, the visual similarity between the reference medication and the AI-predicted medication was examined. Most incorrect predictions (93%) occurred between medications sharing the same color and shape (ie, visually similar medications with different NDCs). A small proportion involved a shape mismatch (5%) or a color mismatch (2%), and none involved mismatches in both color and shape.

Two types of AI assistance, both powered by the same underlying BNN model, were integrated into the interface. The AI model communicated its confidence (or uncertainty) and feature matches through a checkbox plot indicating 4 key pill characteristics: imprint, color, shape, and score (Figure 1A). Matching characteristics were represented using green checkboxes with checkmarks. Mismatches were represented using red checkboxes marked with an “X,” along with a textual description of the AI-predicted characteristic (eg, the predicted imprint shown in Figure 1B). The AI’s confidence

was represented by the predicted probability of the highest-ranked NDC class.

Figure 1. Checkbox plots. (A) Checkbox plot indicating a match and (B) checkbox plot indicating a mismatch. Matching characteristics were represented using green checkboxes with checkmarks. Mismatches were represented using red checkboxes marked with an “X,” along with a textual description of the AI-predicted characteristic.



2A). Feedback was displayed after the participant made the decision, indicating the accuracy of the participant's selection and the AI's prediction (Figure 2B). On the other hand, the ex-post advice condition served as a “safeguard,” appearing only when the participant's response contradicted the AI's prediction. For example, if the participant initially selected “accept” (Figure 3A) when the AI predicted a mismatch, the AI intervened and prompted a reverification by displaying a message, such as: “You selected Accept. The AI disagrees. Confirm or change your response” (Figure 3B). The AI assistance was intentionally imperfect, occasionally providing incorrect recommendations. Although the AI model's accuracy was 98.5%, it was deliberately lowered to 79% to ensure sufficient exposure to both AI's errors and AI's correct guidance outputs. This manipulation enabled the systematic observation of pharmacists' responses to AI successes and failures. The selected reliability level was informed by prior human factors research suggesting that when automation reliability falls below approximately 70%, its costs may outweigh its benefits, highlighting the importance of studying user responses across a range of imperfect, yet still functional, automation performance levels [56].

The ex-ante advice condition continuously displayed the checkbox plot before the participant made a decision (Figure

Figure 2. The ex-ante advice condition displayed the checkbox plot before the participant made a decision. (A) Ex-ante advice selection and (B) ex-ante advice feedback. Feedback was displayed after the initial decision.

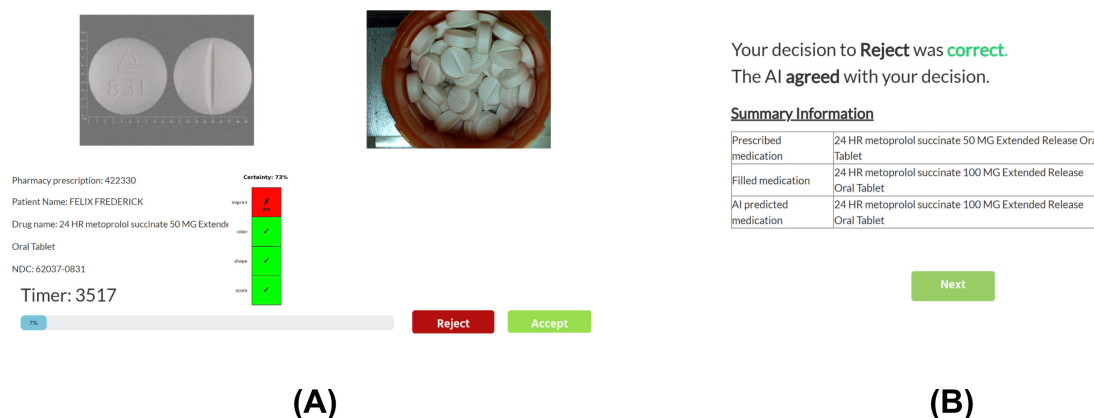
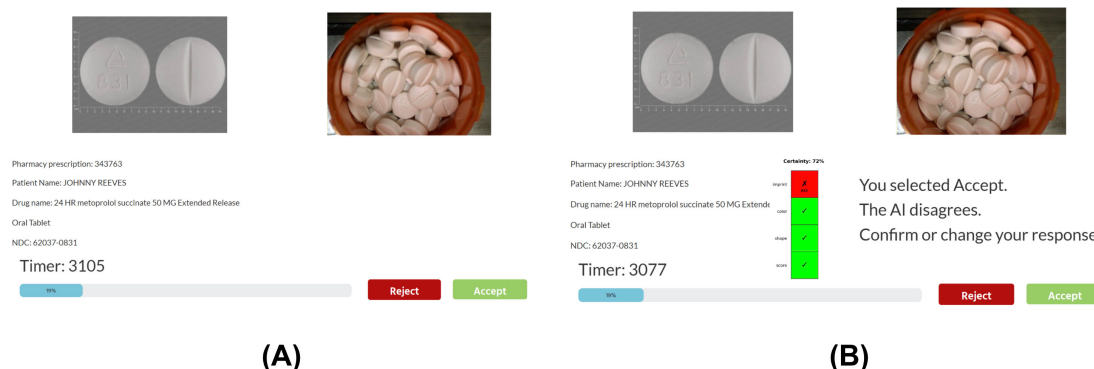


Figure 3. The ex-post advice condition served as a “safeguard,” appearing only when the participant's response contradicted the AI's prediction. (A) Ex-post advice initial selection and (B) ex-post advice intervention.



Participants completed 3 blocks of 100 medication verification trials, with each block corresponding to each

AI assistance condition, including a no-assistance control condition. In the AI-assisted blocks, participants received

trial-by-trial feedback on indicating whether their decision to accept or reject was correct and whether the AI's recommendation was accurate (ie, "Your decision to accept was correct. The AI agreed with your decision"). Following each AI-assisted trial, participants rated their trust in the AI using a visual analog scale ranging from 0 to 100, anchored by "Not at all trust" (0) on the far left to "Completely trust" (100) on the far right. Because the no-assistance control condition did not involve AI support, no trust ratings were obtained.

Experimental Design

The experiment used a within-participants design with varying AI types (ex-ante advice and ex-post advice) and AI recognition patterns (right fill-correct recognition, right fill-incorrect recognition, wrong fill-correct recognition, and wrong fill-incorrect recognition). Within the ex-post advice condition, AI assistance was further classified as the involved ex-post condition (ex-post-AI disagreed with the human decision) and the not-involved ex-post condition (ex-post-AI agreed with the human decision). The agreement conditions depended on participants' responses and therefore could not be predetermined. Block order was counterbalanced across participants using a Latin square design. Each AI type block consisted of 100 medication verification trials randomly selected by the experimental platform from the larger trial pool, including 60 right fill-correct recognition, 16 right fill-incorrect recognition, 19 wrong fill-correct recognition, and 5 wrong fill-incorrect recognition trials.

Measures

After each trial i , participants report their $trust(i)$ in the AI. We calculate a trust adjustment as:

$Trust\ adjustment(i) = Trust(i) - Trust(i - 1)$, where $i = 2, 3, \dots, 100$

Since the moment-to-moment trust is reported after each trial, 99 trust adjustments are obtained from each participant.

Participant demographics, trust propensity, and eye-tracking data were also collected. Trust propensity was evaluated as a potential covariate during model development, but was not a significant predictor and did not alter the results; therefore, it was excluded from the final models. Because the present study focuses on within-participant trust adjustment, demographic and eye-tracking analyses are beyond the scope of this paper and are reported in a companion manuscript.

Experimental Procedure

Before the experiment, participants met with a member of the research team via Zoom (Zoom Communications Inc.) to verify that technical and environmental requirements were met, including a functioning webcam, adequate lighting, and a quiet workspace. Once confirmed, participants received a secure link and password to access the Labvanced software. An instructional video then introduced the mock verification process using the study interface, explaining that the objective was to determine whether an image of a filled medication vial matched the prescribed reference image.

The video also outlined the 2 AI assistance features.

Prior to starting the verification tasks, participants completed a demographic questionnaire, the Occupational Fatigue Exhaustion/Recovery Scale (OFER), and the trust propensity survey, followed by calibration procedures for the eye-tracking software [57,58].

During each trial, a "reference image" (the prescribed medication) appeared on the left, and a "filled vial image" (the dispensed medication) appeared on the right, along with relevant prescription details (Figure 2A). Participants decided whether to accept or reject the filled medication based on this comparison. Each participant completed 100 mock verification trials under the ex-ante advice condition (Figure 2A), 100 trials under the ex-post advice condition (Figure 3B), and 100 trials without AI assistance, in a counterbalanced order. The 6 block orders were assigned approximately evenly across 50 participants (4 sequences: $n=8$ each; 2 sequences: $n=9$ each). After each block of 100 trials, participants completed postcondition surveys, including the trust survey on a 5-point Likert scale, the system usability scale, the NASA Task Load Index (NASA-TLX) survey, and nonmandatory free-response feedback questions [59-61].

Analysis

We analyzed trust adjustment magnitude and trust adjustment using mixed-effects linear regression models. The models included categorical predictors and random intercepts to account for repeated measures within participants. Predictors included AI type and AI recognition pattern, forming a 3×4 design. We first examined differences between AI types (ex-ante advice and ex-post advice) and then further decomposed the ex-post advice into the involved ex-post condition and the not-involved ex-post condition for a 3-level comparison between interaction conditions. This decomposition was prespecified to examine the trust-dynamics associated with disagreement-triggered versus agreement-based AI interactions. Because assignment to the involved and not-involved conditions depends on pharmacists' initial decisions, these comparisons were response-contingent and are treated as exploratory analyses rather than confirmatory tests.

AI recognition patterns consisted of 4 levels: right fill-correct recognition, right fill-incorrect recognition, wrong fill-correct recognition, and wrong fill-incorrect recognition. Because the study focused on pharmacists' trust in AI-assisted decision-making, we restricted the analysis to AI-assisted conditions and excluded data from the no-AI-assistance control block.

We conducted regression 2-tailed t tests to compare mean trust outcomes across AI types within each AI recognition pattern. To account for multiple pairwise comparisons, P values for post hoc comparisons were adjusted using the Benjamini-Hochberg false discovery rate procedure. We conducted all analyses using mixed-effects linear regression models implemented in the lme4 package, applying the Kenward-Roger method to estimate degrees of freedom [62]. All analyses were conducted in the R statistical software (version 4.2.2) [63].

Results

Overview

Observations with trust-adjustment values exceeding ± 3 SDs from the mean were excluded prior to analysis to reduce the influence of extreme outliers. To evaluate order effects, block order was included as a fixed-effect covariate. The likelihood ratio test showed that adding block order did not improve model fit for either outcome (trust adjustment magnitude: $\chi^2_9=8.88$; $P=.11$; trust adjustment: $\chi^2_9=6.25$; $P=.28$). Therefore, the results from the model excluding block order are presented.

Table 2. Mean and SD of trust adjustment magnitude and trust change for different AI recognition patterns. AI recognition pattern to confusion matrix representation^a.

AI type and AI recognition pattern (fill accuracy - (human) AI)	Trust adjustment magnitude	Trust adjustment	Observations, n
Ex-ante advice	5.94 (15.23)		
Right-correct		2.14 (9.03)	2868
Wrong-correct		1.60 (9.90)	1088
Right-incorrect		-15.98 (26.98)	76
Wrong-incorrect		-12.78 (19.85)	800
Ex-post advice	5.37 (15.23)		
Right-correct		1.83 (8.47)	2841
Wrong-correct		1.68 (7.64)	1045
Right-incorrect		-17.84 (29.24)	90
Wrong-incorrect		-9.57 (16.71)	798
Involved Ex-Post	15.44 (26.36)		
Right-(incorrect)-correct		0.07 (13.74)	99
Wrong-(incorrect)-correct		2.75 (7.57)	95
Right-(correct)-incorrect		-18.09 (29.33)	68
Wrong-(correct)-incorrect		-8.96 (16.70)	769
Not-involved Ex-Post	2.60 (8.31)		
Right-(correct)-correct		1.90 (8.21)	2742
Wrong-(correct)-correct		1.57 (7.64)	950
Right-(incorrect)-incorrect		-11.14 (26.06)	22
Wrong-(incorrect)-incorrect		-11.45 (16.97)	29

^aAI recognition pattern to confusion matrix representation. 472 right fill and correct recognition: correct rejection; wrong fill and correct recognition: correct detection; 473 right fill and incorrect recognition: false positive; wrong fill and incorrect recognition: false negative.

Trust Adjustment Magnitude Between Ex-Ante, Involved Ex-Post, and Not-Involved Ex-Post

When ex-post advice was further separated into involved and not-involved ex-post, the interaction condition ($F_{2,9565.1}=380$, $\eta^2=0.07$; $P<.001$) significantly affected the magnitude of trust adjustment.

A post hoc comparison between interaction conditions revealed that involved ex-post condition (EMM 15.78, 95% CI 13.62-17.94) led to significantly higher trust adjustment magnitude compared to ex-ante advice (EMM 6.13, 95% CI 4.15-8.11; $t_{9555.51}=21.00$, mean difference 9.65, 95% CI 8.55-10.75, Cohen $d=0.72$; $P<.001$) and not-involved ex-post condition (EMM 2.80, 95% CI 0.80-4.79; $t_{9556.27}=27.55$, mean difference 12.98, 95% CI 11.85-14.11, Cohen $d=0.97$; $P<.001$). Ex-ante advice led to a higher trust adjustment

Comparison Between AI Types

Trust Adjustment Magnitude Between Ex-Ante and Ex-Post Advice

AI type ($F_{1,9562.1}=3.45$, $\eta^2=0.0004$; $P=.06$) marginally affected the magnitude of trust adjustment. Ex-ante advice (estimated marginal mean [EMM] 6.13, 95% CI 4.28-7.97) led to marginally higher trust adjustment magnitude compared to ex-post advice (EMM 5.60, 95% CI 3.76-7.45; Table 2).

magnitude compared to not-involved ex-post condition ($t_{9554.52}=11.46$, mean difference 3.33, 95% CI 2.63-4.03, Cohen $d=.25$; $P<.001$; Table 2).

Comparison Between AI Types Within Each Recognition Pattern

Trust Adjustment Between Ex-Ante and Ex-Post Advice for Each Recognition Pattern

Recognition pattern ($F_{3,9555.9}=811.37$, $\eta^2=0.2$; $P<.001$) significantly affected trust adjustment; however, AI type did not influence trust adjustment ($F_{1,9548.6}=2.47$, $\eta^2=0.0003$; $P=.12$). However, a marginal interaction effect was observed ($F_{3,9550.7}=2.52$, $\eta^2=0.0008$; $P=.06$).

Because of the marginal interaction effect, we compare Ex-Ante and ex-post advice within each recognition pattern.

Significant differences were observed when the right drugs were incorrectly rejected. Ex-post advice (EMM -17.84 , 95% CI -19.06 to -16.62) showed larger trust decrement compared to ex-ante advice (EMM -15.97 , 95% CI -17.19 to -14.76 ; $t_{9556.76}=-2.67$, mean difference -1.87 , 95% CI -3.24 to -0.49 , Cohen $d=-0.13$; $P=.008$). For other recognition patterns, no significant differences were observed.

Trust Adjustment Between Ex-Ante, Involved Ex-Post, and Not-Involved Ex-Post for Each Recognition Pattern

Both interaction condition ($F_{2,9558.7}=3.89$, $\eta^2=0.0008$; $P=.03$) and recognition pattern ($F_{3,9553.7}=484.11$, $\eta^2=0.13$; $P<.001$) significantly affected trust adjustment. In addition, a significant interaction effect was observed ($F_{6,9570.8}=2.12$, $\eta^2=.0013$; $P=.048$).

Because of the significant interaction effect, we compared interaction conditions within each recognition pattern. Significant differences were observed when the right drugs were incorrectly rejected. Involved ex-post condition (EMM -18.11 , 95% CI -19.34 to -16.87) showed larger trust decrement compared to ex-ante advice (EMM -15.97 , 95% CI -17.19 to -14.76 ; $t_{9553.27}=-3.02$, mean difference -2.13 , 95% CI -3.83 to -0.44 , Cohen $d=-0.15$; $P=.008$) and not-involved ex-post condition (EMM -10.79 , 95% CI -15.96 to -5.61 ; $t_{9577.26}=-2.75$, mean difference -7.32 , 95% CI -13.69 to -0.95 , Cohen $d=-0.52$; $P=.009$). A marginal

difference was observed between ex-ante advice and not-involved ex-post condition ($t_{9569.61}=-1.94$, mean difference -5.19 , 95% CI -11.55 to 1.17 , Cohen $d=-0.37$; $P=.052$). For other recognition patterns, no significant differences were observed.

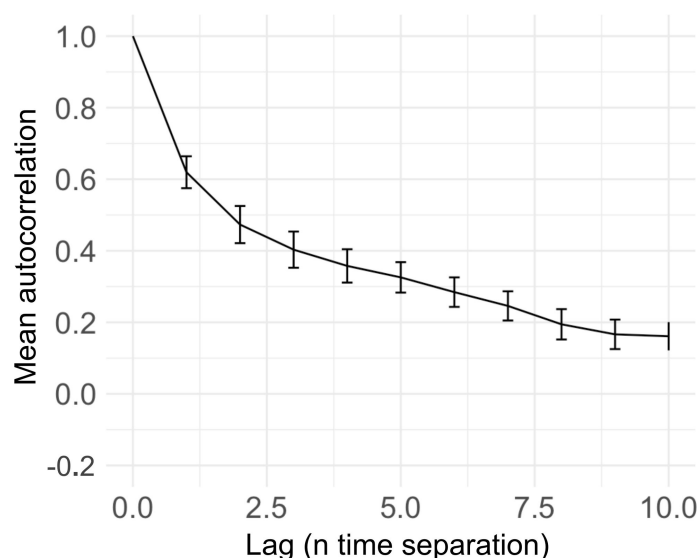
Evaluating Trust Properties: Continuity, Negativity Bias, and Stabilization

To explore the key characteristics of trust dynamics, the recognition patterns were grouped into 2 categories: correct AI predictions (right drug correctly approved and *wrong* drug correctly rejected) and incorrect AI predictions (right drug *incorrectly* rejected and *wrong* drug *incorrectly* accepted).

Continuity

Comparing instances of correct versus incorrect AI predictions revealed that successful AI predictions generally increased user trust, while failures led to decreases. In addition, trust levels at a given moment (t) were strongly correlated with trust from the preceding moment ($t - 1$). To assess how current trust ratings relate to past ones, we calculated the autocorrelation of trust ratings. As shown in Figure 4, trust at time t closely aligns with trust at time $t - 1$. However, as the time interval between ratings grew longer, the average autocorrelation declined, indicating that trust correlation weakens over time (Figure 4).

Figure 4. Autocorrelation of trust as a function of time separation. The error bars represent ± 2 SEs.



Negativity Bias

To test for negativity bias, we compared the magnitude of trust increases following correct predictions with the magnitude of trust decreases following incorrect ones. The results indicated a significant difference between the two: correct AI predictions (mean 2.98, SD 7.33) produced smaller changes than incorrect ones (mean 15.80, SD 23.60) ($t_{4726}=30.59$, mean difference -12.7 , 95% CI -13.6 to -11.9 , standardized effect size $=-1.15$; $P<.001$).

Stabilization

To examine the stabilization property, linear mixed models were used to analyze trust adjustments over time (sequential success [or failure] occurrence).

The findings revealed that the magnitude of trust adjustments decreased significantly across occurrences for both correct predictions ($\beta=-.037$, 95% CI -0.055 to -0.019 , $t_{4021.10}=-4.10$; $P<.001$) and incorrect predictions ($\beta=-.402$, 95% CI -0.659 to -0.155 , $t_{827.46}=-3.20$; $P=.001$).

Discussion

Principal Findings

This study investigated how the timing of AI assistance shaped pharmacists' trust during AI-assisted medication dispensing. We implemented two AI advice timing conditions. In the ex-ante advice condition, the AI presented recognition predictions before pharmacists' initial decision. In the ex-post advice condition, the AI provided selective, postdecision intervention. The ex-post advice was further separated into the involved ex-post condition, which intervened when disagreed with the pharmacist, and the not-involved ex-post condition, which withheld intervention when in agreement.

Overall, differences in trust adjustment magnitude across AI types were modest, with ex-ante advice eliciting slightly higher trust adjustment magnitude than ex-post advice. However, when AI generated incorrect recommendations on initially right fills, ex-ante advice resulted in significantly smaller trust decrement than ex-post advice. Further differentiation of ex-post advice revealed an ordered pattern, where involved ex-post that actively disagreed in the decision process prompted the greatest trust adjustment magnitude, followed by ex-ante advice, with not-involved ex-post yielding the smallest magnitude. This ordering was consistent across the right-incorrect recognition pattern.

In addition, pharmacists' trust dynamics demonstrated ecological validity of trust dynamics properties within a realistic verification context. The findings indicate that pharmacists' trust in AI is dynamic and influenced by the timing, level of involvement, and situational performance of AI assistance during medication verification.

Interpretation of Findings

Comparison Between AI Types

Comparison between the ex-ante and ex-post advice revealed a marginally higher magnitude for ex-ante advice compared to the ex-post advice, which aligns with the findings of Yin et al [27] that the ex-post advice resulted in superior diagnostic quality. When we further decomposed the ex-post advice paradigm into involved ex-post and not-involved ex-post to examine more nuanced trust dynamics associated with AI agreement and disagreement, the involved ex-post condition yielded the highest magnitude of trust adjustment, followed by the ex-ante advice, and then the not-involved ex-post condition. This pattern suggests that interruptive, disagreement-based interventions may exert a strong influence on trust. The ex-ante advice also required pharmacists to evaluate the AI's accuracy during verification, leading to greater trust fluctuations than in the not-involved ex-post condition, which preserved workflow continuity and imposed no extra attentional demand.

Within Recognition Pattern Comparisons

We further analyzed trust adjustment behavior across AI recognition patterns, which considered both the correctness of

the filled medication (*right fill* or *wrong fill*) and the accuracy of the AI's prediction (*correct* or *incorrect*). Significant differences in trust adjustment across ex-ante and ex-post advice occurred when *the right fill was incorrectly rejected*. In this pattern, the ex-post advice led to a larger trust decrement compared to the ex-ante advice. For *the right fill was incorrectly rejected* pattern, significant differences in trust adjustment also occurred across interaction conditions. The involved ex-post condition led to the largest trust decrement compared to the other interaction conditions. The ex-ante advice showed a marginally greater trust decrement compared to the not-involved ex-post condition.

Understanding this pattern requires examining how each interaction condition responded in scenarios where AI incorrectly rejected the right drug. In the ex-ante advice condition, the AI's incorrect prediction was visible from the start. Mismatch indicators (ie, red checkboxes) may have falsely anchored pharmacists' attention toward an assumed mismatch and led to a moderate trust decrement [30,32]. In the involved ex-post condition, after pharmacists made the correct initial decision to accept the right drug, AI erroneously classified the mismatch and intervened with a false alert. This prompted the pharmacist to reconsider a correct initial decision, and the largest trust decrement was observed. This mirrors findings that incorrect ex-post alerts elicit strong trust decrements [42]. In contrast, the not-involved ex-post condition functioned as a passive agreement system. Under this condition, the AI silently agreed with the pharmacists' incorrect initial rejection of the right drug. Here, the error was shared between pharmacists and the AI, and the smallest trust decrement was observed.

Operationally, an incorrect rejection of the right drug represents an unnecessary interruption in the dispensing process. Following the AI's incorrect suggestion would require the pharmacist to reverify or refill a medication that was ready to be dispensed, introducing inefficiency. The larger trust decrement observed in this scenario likely does not stem from task consequence. When pharmacists' accurate initial decisions were falsely contradicted by the involved ex-post condition, the error was salient, yielding strong trust decrements. Conversely, when both pharmacists and AI were wrong (not-involved ex-post), error attribution was shared, producing the smallest trust adjustment.

Validating Trust Dynamics Properties

Beyond evaluating AI interaction types, this study validated 3 trust dynamics properties—continuity, negativity bias, and stabilization [8,45,64]—with an ecologically realistic pharmacy verification environment. Prior validations have often relied on simplified tasks and nonexpert populations. By involving licensed pharmacists performing authentic pill verification, this study strengthens the ecological validity of trust-dynamics in human-AI interaction.

With licensed pharmacists, we demonstrated that trust evolves in a temporally dependent manner, such that trust at a given moment is shaped by prior interactions, strengthening following correct AI performance and declining after errors. We also observed negativity bias, reflecting the asymmetric

influence of system performance, whereby trust is more sensitive to AI failures than to successes, making it difficult to build but easy to lose. Finally, pharmacists' trust exhibited stabilization, showing a tendency to converge over repeated interactions as users developed more consistent expectations of the system's behavior.

Implications

These findings underscore the need for careful design of conditional intervention systems in safety-critical workflows. Ex-post advice assistance, implemented as a disagreement-triggered intervention, carries risks of confirmation bias, low revision rates, and negative reactions to incorrect alerts [40, 43]. For pharmacists, who are expert decision makers and typically accurate, false AI interventions impose unnecessary cognitive costs. Frequent exposure to erroneous interventions may lead to over-skepticism toward the AI, undermining the intended safety benefit. Conversely, ex-ante advice avoids disruptive alerts but introduces anchoring and information-overload risks of Ex-Ante paradigms [30,31]. Both modalities present distinct trade-offs between autonomy support, cognitive load, and error-correction potential.

Designers of medication verification AI must consider these dynamics when determining when and how AI assistance should be delivered. For example, hybrid approaches, such as confidence-gated alerts, adaptive interventions, or selective suppression of low-confidence interruptions, may mitigate the negative impacts observed with the involved ex-post condition.

Limitations and Future Directions

We acknowledge several limitations of this study and propose directions for future research. First, the AI system's accuracy was experimentally reduced from its operational performance (98.5%) to 79% to ensure sufficient exposure to both correct and incorrect recommendations within a feasible number of trials. Although this approach enabled systematic examination of trust adjustments, it may limit generalizability to real-world settings where AI errors may be less frequent. In high-accuracy environments, errors may be more salient and unexpected, potentially eliciting different trust dynamics. Future research should investigate how trust evolves under near-realistic levels of AI reliability.

Second, in the involved ex-post condition, AI disagreement and workflow interruption were inherently coupled. Thus, the observed aggregated trust adjustment magnitude and trust decrement may reflect not only disagreement with the AI recommendation itself but also the cognitive and attentional disruption introduced by mid-task interventions. Future work should investigate the independent effects of disagreement content and workflow interruption on trust adjustment.

Third, the present analyses focused on self-reported trust adjustment, captured through trial-by-trial trust ratings, rather than subsequent behavioral adaptation to the AI support. As a result, the study does not explicitly highlight whether the trust pattern predicts behavioral responses, such as overriding AI recommendations, changes in decision strategies, or adjustments in reaction time. The asymmetric trust pattern observed in the study (negativity bias) may also relate to the broader literature on algorithm aversion, which suggests humans might respond differently to algorithmic mistakes, to an extent that they may disregard using systems even if they outperform humans [40,65]. Future work should investigate how moment-to-moment trust dynamics influence the subsequent decision-making behavior during repeated interactions with AI systems.

Conclusions

This study examined pharmacists' trust dynamics in response to variations in AI assistance timing and conditionality during the medication verification task. By comparing ex-ante advice and ex-post advice paradigms, and further decomposing ex-post advice into involved and not-involved conditions, the findings demonstrate that both timing and conditionality influence trust. In particular, ex-ante advice resulted in greater overall trust adjustment magnitude than in ex-post advice. However, when further decomposed, disagreement-based interventions that incorrectly challenge right expert decisions produced the largest trust decrements and greatest trust fluctuations.

These results extend prior work on AI advice timing by showing that the effectiveness of ex-post advice depends on the nature of AI involvement. While ex-post advice has been shown to enhance overall decision quality, the present findings indicate that the degree to which the AI agrees or disagrees with users' initial decision plays a critical role in trust adjustments beyond the AI advice timing alone. For expert users performing safety-critical tasks, such as pharmacists performing medication verification, erroneous AI interruptions may influence trust more negatively than continuous AI support. This highlights the need to align AI intervention strategies with user expertise and task demands.

Finally, this work also validates the view that trust in AI is a dynamic, interaction-driven construct that evolves through repeated encounters with system behavior. By validating trust dynamics properties in an ecologically realistic pharmacy verification task, this study contributes to human-AI interaction research that emphasizes the importance of capturing moment-to-moment trust adjustments.

Acknowledgments

During the preparation of this work, the authors used ChatGPT (OpenAI, 2026) to improve readability and language. After using this tool/service, the authors reviewed and edited the content as needed and take full responsibility for the content of the publication.

Funding

This study is supported by the National Institutes of Health's National Library of Medicine grant R01LM013624.

Data Availability

The datasets generated during this study are available on the ClinicalTrials.gov website under the identifier NCT06245044.

Authors' Contributions

JYK handled data curation, investigation, formal analysis, validation, visualization, writing – original draft, and writing – review and editing. BR handled investigation, project administration, and writing – review and editing. MW conducted investigation and writing – review and editing. QC supervised software, validation, and writing – review and editing. RAK contributed to conceptualization, funding acquisition, methodology, and supervision. CL contributed to conceptualization, funding acquisition, methodology, project administration, supervision, validation, and writing – review and editing. XJY managed conceptualization, funding acquisition, methodology, supervision, validation, writing – original draft, and writing – review and editing.

Conflicts of Interest

None declared.

References

1. Karia A, Norman R, Robinson S, et al. Pharmacist's time spent: Space for Pharmacy-based Interventions and Consultation Time (SPICE)—an observational time and motion study. *BMJ Open*. Mar 2, 2022;12(3):e055597. [doi: [10.1136/bmjopen-2021-055597](https://doi.org/10.1136/bmjopen-2021-055597)] [Medline: [35236731](https://pubmed.ncbi.nlm.nih.gov/35236731/)]
2. Szeinbach S, Seoane-Vazquez E, Parekh A, Herderick M. Dispensing errors in community pharmacy: perceived influence of sociotechnical factors. *Int J Qual Health Care*. Aug 2007;19(4):203-209. [doi: [10.1093/intqhc/mzm018](https://doi.org/10.1093/intqhc/mzm018)] [Medline: [17567597](https://pubmed.ncbi.nlm.nih.gov/17567597/)]
3. Um IS, Clough A, Tan ECK. Dispensing error rates in pharmacy: a systematic review and meta-analysis. *Res Social Adm Pharm*. Jan 2024;20(1):1-9. [doi: [10.1016/j.sapharm.2023.10.003](https://doi.org/10.1016/j.sapharm.2023.10.003)] [Medline: [37848350](https://pubmed.ncbi.nlm.nih.gov/37848350/)]
4. Campbell PJ, Patel M, Martin JR, et al. Systematic review and meta-analysis of community pharmacy error rates in the USA: 1993–2015. *BMJ Open Qual*. 2018;7(4):e000193. [doi: [10.1136/bmjopen-2017-000193](https://doi.org/10.1136/bmjopen-2017-000193)] [Medline: [30306141](https://pubmed.ncbi.nlm.nih.gov/30306141/)]
5. Reiner G, Pierce SL, Flynn J. Wrong drug and wrong dose dispensing errors identified in pharmacist professional liability claims. *J Am Pharm Assoc* (2003). 2020;60(5):e50–e56. [doi: [10.1016/j.japh.2020.02.027](https://doi.org/10.1016/j.japh.2020.02.027)] [Medline: [32217085](https://pubmed.ncbi.nlm.nih.gov/32217085/)]
6. Chui MA, Look KA, Mott DA. The association of subjective workload dimensions on quality of care and pharmacist quality of work life. *Res Social Adm Pharm*. 2014;10(2):328-340. [doi: [10.1016/j.sapharm.2013.05.007](https://doi.org/10.1016/j.sapharm.2013.05.007)] [Medline: [23791360](https://pubmed.ncbi.nlm.nih.gov/23791360/)]
7. Chui MA, Mott DA. Community pharmacists' subjective workload and perceived task performance: a human factors approach. *J Am Pharm Assoc* (2003). 2012;52(6):e153–e160. [doi: [10.1331/JAPhA.2012.11135](https://doi.org/10.1331/JAPhA.2012.11135)] [Medline: [23229977](https://pubmed.ncbi.nlm.nih.gov/23229977/)]
8. Kim JY, Lester C, Yang XJ. Beyond binary decisions: evaluating the effects of AI error type on trust and performance in AI-assisted tasks. *Hum Factors*. Oct 2025;67(10):1062-1083. [doi: [10.1177/00187208251326795](https://doi.org/10.1177/00187208251326795)] [Medline: [40104968](https://pubmed.ncbi.nlm.nih.gov/40104968/)]
9. Kim JY, Marshall VD, Rowell B, et al. The effects of presenting AI uncertainty information on pharmacists' trust in automated pill recognition technology: exploratory mixed subjects study. *JMIR Hum Factors*. Feb 11, 2025;12:e60273. [doi: [10.2196/60273](https://doi.org/10.2196/60273)] [Medline: [39932773](https://pubmed.ncbi.nlm.nih.gov/39932773/)]
10. Larios Delgado N, Usuyama N, Hall AK, et al. Fast and accurate medication identification. *NPJ Digit Med*. 2019;2(1):10. [doi: [10.1038/s41746-019-0086-0](https://doi.org/10.1038/s41746-019-0086-0)] [Medline: [31304359](https://pubmed.ncbi.nlm.nih.gov/31304359/)]
11. Zheng Y, Rowell B, Chen Q, et al. Designing human-centered AI to prevent medication dispensing errors: focus group study with pharmacists. *JMIR Form Res*. Dec 25, 2023;7(1):e51921. [doi: [10.2196/51921](https://doi.org/10.2196/51921)] [Medline: [38145475](https://pubmed.ncbi.nlm.nih.gov/38145475/)]
12. Tsai CC, Kim JY, Chen Q, et al. Effect of artificial intelligence helpfulness and uncertainty on cognitive interactions with pharmacists: randomized controlled trial. *J Med Internet Res*. Jan 31, 2025;27:e59946. [doi: [10.2196/59946](https://doi.org/10.2196/59946)] [Medline: [39888668](https://pubmed.ncbi.nlm.nih.gov/39888668/)]
13. Lester C, Rowell B, Zheng Y, et al. Effect of uncertainty-aware AI models on pharmacists' reaction time and decision-making in a web-based mock medication verification task: randomized controlled trial. *JMIR Med Inform*. Apr 18, 2025;13(1):e64902. [doi: [10.2196/64902](https://doi.org/10.2196/64902)] [Medline: [40249341](https://pubmed.ncbi.nlm.nih.gov/40249341/)]
14. Heo J, Kang Y, Lee S, Jeong DH, Kim KM. An accurate deep learning-based system for automatic pill identification: model development and validation. *J Med Internet Res*. Jan 13, 2023;25:e41043. [doi: [10.2196/41043](https://doi.org/10.2196/41043)] [Medline: [36637893](https://pubmed.ncbi.nlm.nih.gov/36637893/)]
15. Palenychka R, Lakhssassi A, Palenychka M. Verification of medication dispensing using the attentive computer vision approach. *Proc IEEE Int Symp Circuits Syst*. 2018:1-4. [doi: [10.1109/ISCAS.2018.8351202](https://doi.org/10.1109/ISCAS.2018.8351202)]

16. Chen Q, Al Kontar R, Nouiehed M, Yang XJ, Lester C. Rethinking cost-sensitive classification in deep learning via adversarial data augmentation. *INFORMS J Data Sci.* 2025;4(1):1-19. [doi: [10.1287/ijds.2022.0033](https://doi.org/10.1287/ijds.2022.0033)] [Medline: [42610103](https://pubmed.ncbi.nlm.nih.gov/42610103/)]
17. Gombolay M, Yang XJ, Hayes B, et al. Robotic assistance in the coordination of patient care. *Int J Rob Res.* 2018;37(10):1300-1316. [doi: [10.1177/0278364918778344](https://doi.org/10.1177/0278364918778344)]
18. Wang EH, Gross CP, Tilburt JC, et al. Shared decision making and use of decision AIDS for localized prostate cancer: perceptions from radiation oncologists and urologists. *JAMA Intern Med.* May 2015;175(5):792-799. [doi: [10.1001/jamainternmed.2015.63](https://doi.org/10.1001/jamainternmed.2015.63)] [Medline: [25751604](https://pubmed.ncbi.nlm.nih.gov/25751604/)]
19. Steyvers M, Kumar A. Three challenges for AI-assisted decision-making. *Perspect Psychol Sci.* 2024;19(5):722-734. [doi: [10.1177/17456916231181102](https://doi.org/10.1177/17456916231181102)] [Medline: [37439761](https://pubmed.ncbi.nlm.nih.gov/37439761/)]
20. Zhang Q, Yang XJ, Robert LP. What and when to explain? A survey of the impact of explanation on attitudes toward adopting automated vehicles. *IEEE Access.* 2021;9:159533-159540. [doi: [10.1109/ACCESS.2021.3130489](https://doi.org/10.1109/ACCESS.2021.3130489)]
21. 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/)]
22. von Eschenbach WJ. Transparency and the black box problem: why we do not trust AI. *Philos Technol.* 2021;34(4):1607-1622. [doi: [10.1007/s13347-021-00477-0](https://doi.org/10.1007/s13347-021-00477-0)]
23. Kelly CJ, Karthikesalingam A, Suleyman M, Corrado G, King D. Key challenges for delivering clinical impact with artificial intelligence. *BMC Med.* Oct 29, 2019;17(1):195. [doi: [10.1186/s12916-019-1426-2](https://doi.org/10.1186/s12916-019-1426-2)] [Medline: [31665002](https://pubmed.ncbi.nlm.nih.gov/31665002/)]
24. Cai CJ, Winter S, Steiner D, Wilcox L, Terry M. "Hello AI": uncovering the onboarding needs of medical practitioners for human-AI collaborative decision-making. *Proc ACM Hum-Comput Interact.* Nov 7, 2019;3(CSCW):1-24. [doi: [10.1145/3359206](https://doi.org/10.1145/3359206)]
25. Begoli E, Bhattacharya T, Kusnezov D. The need for uncertainty quantification in machine-assisted medical decision making. *Nat Mach Intell.* 2019;1:20-23. [doi: [10.1038/s42256-018-0004-1](https://doi.org/10.1038/s42256-018-0004-1)]
26. Kompa B, Snoek J, Beam AL. Second opinion needed: communicating uncertainty in medical machine learning. *NPJ Digit Med.* Jan 5, 2021;4(1):4. [doi: [10.1038/s41746-020-00367-3](https://doi.org/10.1038/s41746-020-00367-3)] [Medline: [33402680](https://pubmed.ncbi.nlm.nih.gov/33402680/)]
27. Yin J, Ngiam KY, Tan SSL, Teo HH. Designing AI-based work processes: how the timing of AI advice affects diagnostic decision making. *Manage Sci.* Nov 2025;71(11):9361-9383. [doi: [10.1287/mnsc.2022.01454](https://doi.org/10.1287/mnsc.2022.01454)]
28. Brennan M, Puri S, Ozrazgat-Baslanti T, et al. Comparing clinical judgment with the MySurgeryRisk algorithm for preoperative risk assessment: a pilot usability study. *Surgery.* May 2019;165(5):1035-1045. [doi: [10.1016/j.surg.2019.01.002](https://doi.org/10.1016/j.surg.2019.01.002)] [Medline: [30792011](https://pubmed.ncbi.nlm.nih.gov/30792011/)]
29. Tschandl P, Rinner C, Apalla Z, et al. Human-computer collaboration for skin cancer recognition. *Nat Med.* 2020;26(8):1229-1234. [doi: [10.1038/s41591-020-0942-0](https://doi.org/10.1038/s41591-020-0942-0)] [Medline: [32572267](https://pubmed.ncbi.nlm.nih.gov/32572267/)]
30. Tversky A, Kahneman D. Judgment under uncertainty: heuristics and biases. *Science.* Sep 27, 1974;185(4157):1124-1131. [doi: [10.1126/science.185.4157.1124](https://doi.org/10.1126/science.185.4157.1124)] [Medline: [17835457](https://pubmed.ncbi.nlm.nih.gov/17835457/)]
31. Bućinca Z, Malaya MB, Gajos KZ. To trust or to think. *Proc ACM Hum-Comput Interact.* Apr 13, 2021;5(CSCW1):1-21. [doi: [10.1145/3449287](https://doi.org/10.1145/3449287)]
32. Rastogi C, Zhang Y, Wei D, Varshney KR, Dhurandhar A, Tomsett R. Deciding fast and slow: the role of cognitive biases in AI-assisted decision-making. *Proc ACM Hum-Comput Interact.* Mar 30, 2022;6(CSCW1):1-22. [doi: [10.1145/3512930](https://doi.org/10.1145/3512930)]
33. Aoki T, Yamada A, Aoyama K, et al. Clinical usefulness of a deep learning-based system as the first screening on small-bowel capsule endoscopy reading. *Dig Endosc.* May 2020;32(4):585-591. [doi: [10.1111/den.13517](https://doi.org/10.1111/den.13517)] [Medline: [31441972](https://pubmed.ncbi.nlm.nih.gov/31441972/)]
34. Dratsch T, Chen X, Rezazade Mehrizi M, et al. Automation bias in mammography: the impact of artificial intelligence BI-RADS suggestions on reader performance. *Radiology.* May 2023;307(4):e222176. [doi: [10.1148/radiol.222176](https://doi.org/10.1148/radiol.222176)] [Medline: [37129490](https://pubmed.ncbi.nlm.nih.gov/37129490/)]
35. 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/)]
36. Green B, Chen Y. The principles and limits of algorithm-in-the-loop decision making. *Proc ACM Hum-Comput Interact.* Nov 7, 2019;3(CSCW):1-24. [doi: [10.1145/3359152](https://doi.org/10.1145/3359152)]
37. Yaniv I, Choshen-Hillel S. Exploiting the wisdom of others to make better decisions: suspending judgment reduces egocentrism and increases accuracy. *Behav Decis Mak.* Dec 2012;25(5):427-434. [doi: [10.1002/bdm.740](https://doi.org/10.1002/bdm.740)]
38. Klayman J. Varieties of confirmation bias. *Psychol Learn Motiv.* 1995;32:385-418. [doi: [10.1016/S0079-7421\(08\)60315-1](https://doi.org/10.1016/S0079-7421(08)60315-1)]
39. Nickerson RS. Confirmation bias: a ubiquitous phenomenon in many guises. *Rev Gen Psychol.* Jun 1998;2(2):175-220. [doi: [10.1037/1089-2680.2.2.175](https://doi.org/10.1037/1089-2680.2.2.175)]

40. Dietvorst BJ, Simmons JP, Massey C. Algorithm aversion: people erroneously avoid algorithms after seeing them err. *J Exp Psychol Gen*. Feb 2015;144(1):114-126. [doi: [10.1037/xge0000033](https://doi.org/10.1037/xge0000033)] [Medline: [25401381](https://pubmed.ncbi.nlm.nih.gov/25401381/)]
41. Wang D, Yang Q, Abdul A, Lim BY. Designing theory-driven user-centric explainable AI. *Proc 2019 CHI Conf Hum Factors Comput Syst*. 2019:1-15. [doi: [10.1145/3290605.3300831](https://doi.org/10.1145/3290605.3300831)]
42. Lehman CD, Wellman RD, Buist DSM, et al. Diagnostic accuracy of digital screening mammography with and without computer-aided detection. *JAMA Intern Med*. Nov 2015;175(11):1828-1837. [doi: [10.1001/jamainternmed.2015.5231](https://doi.org/10.1001/jamainternmed.2015.5231)] [Medline: [26414882](https://pubmed.ncbi.nlm.nih.gov/26414882/)]
43. Fogliato R, Chappidi S, Lungren M, et al. Who goes first? Influences of human-AI workflow on decision making in clinical imaging. *Proc ACM Conf Fairness Accountab Transpar*. Jun 21, 2022:1362-1374. [doi: [10.1145/3531146.3533193](https://doi.org/10.1145/3531146.3533193)]
44. de Visser EJ, Peeters MMM, Jung MF, et al. Towards a theory of longitudinal trust calibration in human-robot teams. *Int J of Soc Robotics*. 2020;12(2):459-478. [doi: [10.1007/s12369-019-00596-x](https://doi.org/10.1007/s12369-019-00596-x)]
45. Yang XJ, Guo Y, Schemanske C. From trust to trust dynamics: combining empirical and computational approaches to model and predict trust dynamics. In: *Human-Automation Interaction: Transportation*. Springer; 2022:253-265. [doi: [10.1007/978-3-031-10784-9_15](https://doi.org/10.1007/978-3-031-10784-9_15)]
46. Lee J, Moray N. Trust, control strategies and allocation of function in human-machine systems. *Ergonomics*. Oct 1992;35(10):1243-1270. [doi: [10.1080/00140139208967392](https://doi.org/10.1080/00140139208967392)] [Medline: [1516577](https://pubmed.ncbi.nlm.nih.gov/1516577/)]
47. Wischniewski M, Krämer N, Müller E. Measuring and understanding trust calibrations for automated systems: a survey of the state-of-the-art and future directions. *Proc 2023 CHI Conf Hum Factors Comput Syst*. 2023:1-16. [doi: [10.1145/3544548.3581119](https://doi.org/10.1145/3544548.3581119)]
48. Yang XJ, Schemanske C, Searle C. Toward quantifying trust dynamics: how people adjust their trust after moment-to-moment interaction with automation. *Hum Factors*. Aug 2023;65(5):862-878. [doi: [10.1177/00187208211034716](https://doi.org/10.1177/00187208211034716)] [Medline: [34459266](https://pubmed.ncbi.nlm.nih.gov/34459266/)]
49. Manzey D, Reichenbach J, Onnasch L. Human performance consequences of automated decision aids: the impact of degree of automation and system experience. *J Cogn Eng Decis Mak*. 2012;6(1):57-87. [doi: [10.1177/1555343411433844](https://doi.org/10.1177/1555343411433844)]
50. Lee JD, Moray N. Trust, self-confidence, and operators' adaptation to automation. *Int J Hum Comput Stud*. Jan 1994;40(1):153-184. [doi: [10.1006/ijhc.1994.1007](https://doi.org/10.1006/ijhc.1994.1007)]
51. Yang XJ, Wickens CD, Hölttä-Otto K. How users adjust trust in automation: contrast effect and hindsight bias. *Proc Hum Factors Ergon Soc Annu Meet*. 2016;60:196-200. [doi: [10.1177/1541931213601044](https://doi.org/10.1177/1541931213601044)]
52. Yang XJ, Unhelkar VV, Li K, Shah JA. Evaluating effects of user experience and system transparency on trust in automation. *Proc ACM/IEEE Int Conf Hum Robot Interact*. Mar 6, 2017:408-416. [doi: [10.1145/2909824.3020230](https://doi.org/10.1145/2909824.3020230)]
53. Gal Y, Ghahramani Z. Dropout as a Bayesian approximation: representing model uncertainty in deep learning. Presented at: *Proceedings of the 33rd International Conference on Machine Learning*; Jun 19-24, 2016:1050-1059; New York City, NY. URL: <https://proceedings.mlr.press/v48/gal16.pdf> [Accessed 2026-08-27]
54. He K, Zhang X, Ren S, Sun J. Deep residual learning for image recognition. *IEEE Conf Comput Vis Pattern Recognit*. 2016:770-778. [doi: [10.1109/CVPR.2016.90](https://doi.org/10.1109/CVPR.2016.90)]
55. Lester CA, Li J, Ding Y, Rowell B, Kontar RA. Performance evaluation of a prescription medication image classification model: an observational cohort. *NPJ Digit Med*. Jul 27, 2021;4(1):118. [doi: [10.1038/s41746-021-00483-8](https://doi.org/10.1038/s41746-021-00483-8)] [Medline: [34315995](https://pubmed.ncbi.nlm.nih.gov/34315995/)]
56. Wickens CD, Dixon SR. The benefits of imperfect diagnostic automation: a synthesis of the literature. *Theor Issues Ergon Sci*. May 2007;8(3):201-212. [doi: [10.1080/14639220500370105](https://doi.org/10.1080/14639220500370105)]
57. Merritt SM, Heimbaugh H, LaChapell J, Lee D. I trust it, but I don't know why: effects of implicit attitudes toward automation on trust in an automated system. *Hum Factors*. Jun 2013;55(3):520-534. [doi: [10.1177/0018720812465081](https://doi.org/10.1177/0018720812465081)] [Medline: [23829027](https://pubmed.ncbi.nlm.nih.gov/23829027/)]
58. Winwood PC, Winefield AH, Dawson D, Lushington K. Development and validation of a scale to measure work-related fatigue and recovery: the Occupational Fatigue Exhaustion/Recovery Scale (OFER). *J Occup Environ Med*. Jun 2005;47(6):594-606. [doi: [10.1097/01.jom.0000161740.71049.c4](https://doi.org/10.1097/01.jom.0000161740.71049.c4)] [Medline: [15951720](https://pubmed.ncbi.nlm.nih.gov/15951720/)]
59. Muir BM, Moray N. Trust in automation. Part II. Experimental studies of trust and human intervention in a process control simulation. *Ergonomics*. Mar 1996;39(3):429-460. [doi: [10.1080/00140139608964474](https://doi.org/10.1080/00140139608964474)] [Medline: [8849495](https://pubmed.ncbi.nlm.nih.gov/8849495/)]
60. Hart SG. NASA-Task Load Index (NASA-TLX); 20 Years Later. *Proc Hum Factors Ergon Soc Annu Meet*. Oct 2006;50(9):904-908. [doi: [10.1177/154193120605000909](https://doi.org/10.1177/154193120605000909)]
61. Brooke J. SUS—a quick and dirty usability scale. In: Jordan PW, Thomas B, McClelland IL, Weerdmeester B, editors. *Usability Evaluation in Industry*. CRC Press; 1996:189-194. [doi: [10.1201/9781498710411](https://doi.org/10.1201/9781498710411)]

62. Bates D, Mächler M, Bolker B, Walker S. Fitting linear mixed-effects models using lme4. J Stat Softw. 2015;67:1-48. [doi: [10.18637/jss.v067.i01](https://doi.org/10.18637/jss.v067.i01)]
63. The R Project for Statistical Computing. 2023. URL: <https://www.R-project.org/> [Accessed 2026-08-27]
64. Guo Y, Yang XJ. Modeling and predicting trust dynamics in human–robot teaming: a Bayesian inference approach. Int J Soc Robotics. Dec 2021;13(8):1899-1909. [doi: [10.1007/s12369-020-00703-3](https://doi.org/10.1007/s12369-020-00703-3)]
65. Bogert E, Schecter A, Watson RT. Humans rely more on algorithms than social influence as a task becomes more difficult. Sci Rep. Apr 13, 2021;11(1):8028. [doi: [10.1038/s41598-021-87480-9](https://doi.org/10.1038/s41598-021-87480-9)] [Medline: [33850211](https://pubmed.ncbi.nlm.nih.gov/33850211/)]

Abbreviations

BNN: Bayesian neural network
CDSS: clinical decision support system
EMM: estimated marginal mean
HCI: human-computer interaction
NASA-TLX: NASA Task Load Index
NDC: National Drug Code
OFER: Occupational Fatigue Exhaustion/Recovery Scale

Edited by James Waterson, Stephanie Law; peer-reviewed by Amir Reza Ashraf, Nicolas Rudolf; submitted 03.Feb.2026; final revised version received 21.Jul.2026; accepted 06.Aug.2026; published 11.Sep.2026

Please cite as:

Kim JY, Rowell B, Whitaker M, Chen Q, Al Kontar R, Lester C, Yang XJ
Effects of AI Assistance Timing on Pharmacists' Trust in Automated Pill Recognition Technology: Within-Participants Experimental Study
JMIR Hum Factors 2026;13:e92822
URL: <https://humanfactors.jmir.org/2026/1/e92822>
doi: [10.2196/92822](https://doi.org/10.2196/92822)

© Jin Yong Kim, Brigid Rowell, Megan Whitaker, Qiyuan Chen, Raed Al Kontar, Corey Lester, Xi Jessie Yang. Originally published in JMIR Human Factors (<https://humanfactors.jmir.org>), 11.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 JMIR Human Factors, is properly cited. The complete bibliographic information, a link to the original publication on <https://humanfactors.jmir.org>, as well as this copyright and license information must be included.