

AI Advisor or Human Expert? Regret, Disappointment, and Future Reliance After Following or Rejecting Advice

Adam Stachowicz¹, Katarzyna Idzikowska²,

Abstract

Decision making is increasingly supported by artificial intelligence (AI) systems as well as human experts. This study examined whether advice source (AI algorithm vs. human expert) and the decision to follow or reject advice influenced reported regret, disappointment, and willingness to use the same source again after a negative outcome. Participants ($N = 212$) completed three hypothetical decision scenarios concerning airline tickets, a winter jacket, and a stock market investment in a 3 (scenario; within subjects) $\times$ 2 (advice source; between subjects) $\times$ 2 (advice followed vs. rejected; between subjects) mixed design. In the advice-followed condition, the recommendation subsequently proved incorrect, whereas in the advice-rejected condition, the recommendation retrospectively proved correct despite the negative outcome. Regret and disappointment differed significantly across scenarios, but neither emotion showed a main effect of advice source. Significant advice source $\times$ advice-taking interactions were found for both emotions; however, the direct comparisons between algorithmic and expert advice within the advice-followed condition were not significant and did not support the predicted differences. Participants reported greater overall willingness to reuse expert than algorithmic advice. Nevertheless, the predicted lower willingness to reuse algorithmic rather than expert advice

¹ Adam Stachowicz – Kozminski University, Warsaw, Poland, e-mail: 51728@kozminski.edu.pl

² Katarzyna Idzikowska – Kozminski University, Center for Economic Psychology and Decision Making, Warsaw, Poland, e-mail: kidzikowska@kozminski.edu.pl, <https://orcid.org/0000-0002-0689-5285>

specifically after following an incorrect recommendation was not supported. The largest effect concerned the advice-taking condition: willingness to reuse the source was substantially higher after rejecting a recommendation that subsequently proved correct than after following a recommendation that subsequently proved incorrect. Because advice-taking was confounded with retrospective advice accuracy, their independent effects could not be determined. The hypothetical scenarios and low reliability of the Regret and Disappointment Scale indices limit the strength and generalizability of the conclusions.

Keywords: regret, disappointment, decisional advice, artificial intelligence algorithms, algorithm aversion

Introduction

Contemporary decision-making processes, both in the dynamically changing professional sphere and in the increasingly complex consumer world, are continuously supported by advanced technologies. Alongside traditional sources of expert knowledge such as financial advisors, marketing specialists, or physicians, algorithmic systems play an ever-growing role. These decision support systems (DSS), encompassing both advice generated by artificial intelligence (AI) algorithms and advice based on human expertise, find extremely broad application in diverse fields. In the financial sector, algorithms support credit-risk assessment, automate investment advisory (often referred to as automated advisory, so-called *robo-advisory*) and optimize asset portfolios (Markus et al., 2008; Straka, 2000). In marketing, recommender systems personalize offers for customers, algorithms optimize advertising campaigns in real time, and forecast consumer trends (Lambrecht & Tucker, 2019; Wang et al., 2025). Medicine benefits from AI support in diagnostic processes, in the analysis of medical images, and in personalizing therapy (Esmaeilzadeh et al., 2015; Lee et al., 2013; Tschandl et al., 2019). Human resources management also increasingly relies on algorithms for candidate selection, competence assessment, and predicting employee turnover (Dastin, 2018; Upadhyay & Khandelwal, 2018). In addition, algorithms used in selection processes, despite their efficiency, may be underappreciated by users, which already signals the complexity of the human–algorithm relationship (Rabinovitch et al., 2024).

These two main sources of decision support – algorithms and humans – are characterized by diametrically different properties, which inevitably affects their perception, acceptance, and ultimate use by decision-makers. Algorithmic support, on the one hand, is attractive due to its high

speed in processing huge data sets, its ability to identify complex patterns that the human mind might not detect, and the consistency of its recommendations, which eliminates the influence of subjective factors or momentary emotions. The scalability of algorithmic solutions allows for their wide application at relatively low marginal cost. However, as Madhavan and Wiegmann (2007) note in their concept of the “perfect automation schema,” people often hold implicit expectations regarding the infallibility of algorithms. This schema suggests that machine errors are perceived as more surprising and less forgivable than human mistakes (Daschner & Obermaier, 2022; Dzindolet et al., 2003). Consequently, even a single algorithmic error can lead to strong negative emotional reactions and a sharp drop in trust. On the other hand, human experts, despite being susceptible to cognitive biases and fatigue, offer greater flexibility, the ability to achieve a deeper, holistic understanding of the nuances of a specific decision problem, and the capacity to take into account subtle contextual, emotional, and interpersonal factors such as empathy or the understanding of non-verbal signals. It is precisely these “soft” aspects that often constitute the human advantage, especially in situations requiring negotiation, persuasion, or building long-term relationships.

As artificial intelligence algorithms take over increasingly complex and responsible advisory tasks, a thorough understanding of the psychological aspects of human–algorithm and human–expert interaction becomes crucial. An individual’s propensity to accept and effectively use advice, regardless of its source, is influenced by a whole constellation of factors. The most important include: the perceived accuracy of the recommendation itself, the characteristics of the source (its credibility, reputation, perceived impartiality), as well as the characteristics of the decision-maker, such as their level of self-confidence, experience in the field, or general attitude toward new technologies (Parasuraman & Colby, 2015; Saxena, 2024).

However, research on human–AI interaction indicates that trust placed in algorithms is often more fragile and is governed by somewhat different rules than interpersonal trust. Even in situations where algorithms prove to be objectively more effective, presenting higher predictive accuracy (Meehl, 1954), people may show resistance to using them, preferring advice from human experts (Dietvorst et al., 2015; Rabinovitch et al., 2024). This intriguing phenomenon, termed “algorithm aversion,” has been described and investigated in detail by Dietvorst et al. (2015). Their experimental research showed that people often lose trust in an algorithm and abandon its support much more quickly and more decisively after observing even a single error, compared to an analogous error committed by a human. The tendency to more harshly “punish” algorithms for

mistakes appears to be stronger, whereas human errors tend to be perceived as more natural, inevitable, and consequently more easily forgivable.

Renier et al. (2021) suggest that this disproportion stems in part from the different expectations we hold of machines and of humans. From algorithms, perceived as tools based on logic and data, we expect considerably greater precision, objectivity, and reliability. Their errors disrupt this image and may elicit stronger negative reactions, including frustration and a sense of being let down (Burton et al., 2020). In contrast, human fallibility is in a way inscribed in our nature (*errare humanum est*), which makes us more tolerant of mistakes made by other people. Moreover, Sunstein and Gaffe (2025) identify various mechanisms underlying this aversion, including the desire for agency, negative moral or emotional reactions to being judged by an algorithm, and the conviction of the unique knowledge of human experts.

An important role in shaping the relationship with an advisor (human or algorithmic) is also played by the perceived possibility of sharing responsibility for the final decision outcome. Aschauer et al. (2023) point out that managers may be more inclined to share responsibility with a human advisor than with an algorithmic system, which may influence the decision-making process and reactions to its consequences. Blame avoidance theory, invoked by these authors, suggests that decision-makers may strategically use advisors (especially human ones) to diffuse responsibility for the potential negative consequences of their choices. The question of whether, and to what extent, responsibility can be attributed to algorithms, especially those of a “black box” nature, remains the subject of intense debate both in ethics and in legal scholarship (Coeckelbergh, 2020; Matthias, 2004).

Emotions play a crucial, though often underappreciated, role in the complex process of decision-making. They influence both the choices themselves (through the mechanism of anticipating future emotional states, so-called *affective forecasting*), and the evaluation of obtained outcomes and reactions to those outcomes. In the context of negative consequences of decisions, behavioral decision theories, developed since the 1980s, pay particular attention to two distinct emotions – although in everyday understanding often confused – regret and disappointment. Fundamental work in this area includes, above all, Regret Theory by Loomes and Sugden (1982) as well as the independently developed concepts of Bell (1982, 1985). These theories posit that

people, when making decisions under uncertainty, not only maximize expected utility but also try to minimize potential future regret (Bell, 1982).

The distinction between regret and disappointment, crucial for understanding the psychology of decision-making, is based on different processes of counterfactual comparisons – that is, thinking about alternative “what if...” scenarios (Kahneman & Tversky, 1982; Marcatto & Ferrante, 2008). As Marcatto and Ferrante (2008) explain, regret is an emotion resulting from comparing the actually obtained outcome with what could have been achieved had the decision-maker made a different choice. It is strongly linked to a sense of personal responsibility for the choice made (internal attribution) and to counterfactual thinking focused on one’s own action (“I regret choosing X instead of Y”). Disappointment, in turn, as the same authors and Zeelenberg et al.’s (2000) argue, arises from comparing the obtained outcome with what could have happened had external circumstances or the state of the world independent of the decision-maker turned out differently (e.g., had the market behaved as expected). This emotion is associated with attributing causes to external factors (external attribution) and counterfactual thinking focused on the situation (“I am disappointed that Z happened, even though I made the right decision”). The theoretical foundations of this distinction can also be found in Weiner’s (1985) attribution theory, which explains how people interpret the causes of successes and failures and how these interpretations affect their emotions and motivation.

Experiencing regret or disappointment has consequences for future decisions and behavior. As research by Zeelenberg et al. (2001) indicates, these emotions can lead to changes in decision strategies, increased risk aversion, lower overall satisfaction with the choices made, and modifications of subsequent consumer and investment behavior. Giorgetta (2012) emphasizes that understanding the factors determining the experience of these emotions is crucial for a fuller picture of decision-making processes.

Despite growing interest in trust in algorithms, algorithm aversion, and reactions to algorithmic errors (e.g., Burton et al., 2020; Renier et al., 2021; Daschner & Obermaier, 2022), relatively little attention has been paid to specific emotional reactions arising after negative outcomes involving algorithmic and human advice. Such outcomes may occur either because a decision-maker follows a recommendation that subsequently proves incorrect or because the decision-maker rejects a recommendation that would have produced a better result. These two situations may activate different counterfactual and attributional processes. Following incorrect

advice may focus attention on the decision to trust the advisor, whereas rejecting correct advice may make the forgone alternative particularly salient.

It is therefore important to examine whether regret and disappointment depend not only on whether the advisor is an AI algorithm or a human expert, but also on whether the recommendation was followed. Regret, which is associated with personal responsibility and choice-focused counterfactual thinking, may be particularly sensitive to the relationship between advice source and the decision to accept or reject the recommendation. Disappointment, which is more closely associated with external circumstances and violated expectations, may instead depend more strongly on the situational context and on expectations concerning the advisor's reliability.

The aim of this study was to investigate how the source of advice (an AI algorithm vs. a human expert) and whether the advice was followed influenced regret, disappointment, and future willingness to use the same source after a negative decision outcome. Based on the extended literature review presented, concerning both the specifics of interaction with algorithmic systems, the phenomenon of algorithm aversion, and the psychological mechanisms of regret and disappointment, the following research hypotheses were formulated:

H1: Following incorrect advice from a human expert elicits greater regret than following incorrect advice from an AI algorithm.

The rationale is the assumption that decision-makers may feel greater personal responsibility (internal attribution) for the choice of trusting a fallible human, which, according to theory (Marcatto & Ferrante, 2008), intensifies the feeling of regret. The perception of a lower possibility of 'sharing' responsibility with a human (Aschauer et al., 2023) may reinforce this effect.

H2: Following incorrect advice from an AI algorithm elicits greater disappointment than following incorrect advice from a human expert.

This hypothesis is based on the assumption that algorithms are often perceived as precise and data-driven systems (Renier et al., 2021), so their failure may be attributed to the external

source and experienced as a stronger violation of expectations. Such external attribution constitutes the basis of the emotion of disappointment (Marcatto & Ferrante, 2008).

H3: *Following incorrect algorithmic advice results in lower willingness to use the same source again than following incorrect advice from a human expert.*

This prediction follows directly from the phenomenon of algorithm aversion, which is strongly triggered after observing an error (Daschner & Obermaier, 2022; Dietvorst et al., 2015). It is posited that the subsequent willingness to rely on the source will be lower for algorithmic than for human advice. The advice-not-taken conditions were included to explore emotional and behavioral reactions to rejecting a recommendation that subsequently proved correct. No directional hypotheses were formulated for these conditions.

Method

Participants

366 individuals recruited online through social media took part in the study. Based on previously defined exclusion criteria, i.e., non-completion of the study, too short (< 225 seconds) or too long (> 1500 seconds) time to complete the survey, and failure to meet the control condition in the Regret and Disappointment Scale (RDS), 154 records were excluded. The final sample size for the analyses was 212 individuals (including 71.7% women and 28.3% men).

Participants' ages ranged from 19 to 53 years ($M = 27.11$, $SD = 6.8$). The majority of the sample consisted of young adults, with the most numerous groups aged 25 (17.0%), 24 (13.7%), and 23 (9.4%). In terms of education, individuals with a master's or engineering degree (52.4%) and with a bachelor's or engineering degree (31.6%) dominated. More than half of the respondents (53.3%) lived in large cities of over 500,000 inhabitants.

Study design and materials

The experiment employed a 3 (decision scenario: airline ticket purchase, winter jacket purchase, stock market investment; within-subjects factor) $\times$ 2 (advice source: AI algorithm vs. human

expert; between-subjects factor) $\times$ 2 (protagonist's behavior: advice taken vs. advice not taken; between-subjects factor) design.

The first between-subjects independent variable was the advice source, manipulated at two levels: advice generated by an AI algorithm vs. advice given by a human expert. The second between-subjects independent variable was the scenario protagonist's behavior toward the advice, also manipulated at two levels: advice taken vs. advice not taken. The within-subjects factor consisted of three decision scenarios, all of which were presented to each participant. Participants were randomly assigned to one of the four experimental conditions:

- (1) Algorithm-Taken (A-T),
- (2) Algorithm-Not taken (A-NT),
- (3) Expert-Taken (E-T),
- (4) Expert-Not taken (E-NT).

The stimulus material consisted of three hypothetical decision scenarios describing common life situations: airline ticket purchase, winter jacket purchase, and stock market investment (see Appendix B). Each scenario ended with a negative outcome for the protagonist. The content of the scenarios was adapted to four experimental conditions, differing in the indicated advice source (AI algorithm or human expert) and in the description of the protagonist's behavior (advice taken or advice not taken). For example, in the airline ticket scenario, the advice source manipulation was introduced as follows:

- In conditions in which the source of advice was an algorithm, participants read: "To avoid overpaying, before buying the tickets you use the free advice of an artificial intelligence algorithm that analyzes flight prices. The algorithm suggested that you hold off on the purchase (...)."
- In conditions in which the source of advice was a human expert, participants read: "To avoid overpaying, before buying the tickets you use the free advice of an expert who analyzes flight prices. The expert suggested that you hold off on the purchase (...)."

The manipulation of the protagonist's behavior was introduced as follows:

- In the advice-taken condition, read: "You followed the advice and postponed the ticket purchase." The consequence (in this specific scenario for this situation) was an increase in ticket prices. Thus, the recommendation that had been followed turned out to be incorrect.

- In the advice-not-taken condition, participants read: “You did not follow the advice and did not postpone the ticket purchase.” The consequence (in this specific scenario for this situation) was a drop in ticket prices, which meant that the protagonist had overpaid by buying earlier. Thus, the recommendation that had been rejected would have produced a better outcome and was therefore retrospectively correct.

Accordingly, the final outcome was negative in all four conditions, but the accuracy of the advice differed depending on whether it had been followed. In the advice-taken condition, participants evaluated a situation in which the protagonist followed an incorrect recommendation. In the advice-not-taken condition, they evaluated a situation in which the protagonist rejected a recommendation that subsequently proved correct. Analogous manipulations were used in the winter jacket and stock market investment scenarios.

The three scenarios were selected for several reasons. First, they represented relatively common and familiar situations involving uncertainty about future states of affairs, such as price changes, promotions, or market fluctuations. Second, they made it possible to present an unambiguously negative outcome while systematically manipulating the source of advice and the protagonist’s response to it. Third, the scenarios represented different decision domains – services, consumer goods, and personal finance – allowing for a preliminary examination of whether the observed effects generalized across contexts.

To measure the key dependent variables – regret and disappointment – the Regret and Disappointment Scale (RDS) by Marcatto and Ferrante (2008) was chosen. The choice of this instrument was motivated by its strong theoretical grounding: the scale was designed specifically to operationalize and empirically distinguish regret and disappointment on the basis of counterfactual comparison theory and attribution processes, which is fully consistent with the theoretical framework of the present work. The RDS scale allows for capturing the key components of both emotions, including thinking focused on one’s own choice (regret) versus on external factors (disappointment) as well as the sense of responsibility (internal vs. external attribution). The scale was translated from English into Polish using a back-translation procedure. It was also necessary to adapt the wording to the narrative specificity of the hypothetical study scenarios, in

order to facilitate participants' identification with the protagonist and the evaluation of his or her possible emotional reactions.

The measure consisted of seven items. Six of them were rated on a 5-point Likert scale (1 – “Strongly agree” to 5 – “Strongly disagree”), measuring: general negative affective reaction, counterfactual thinking focused on the choice, counterfactual thinking focused on external factors, internal attribution, external attribution, and satisfaction. The seventh question forced a choice of the dominant counterfactual explanation (see Appendix A). In line with the methodology of the scale's authors, the Regret Index (the mean of statements 2 and 4) and the Disappointment Index (the mean of statements 3 and 5) were calculated. Both indices were reverse-scored for the present analyses, so that higher values of the indicator denote a higher level of experienced regret or disappointment respectively, whereas lower values signal a lower level of these emotions. It should, however, be noted that the reliability analysis of these indices in the study showed low internal consistency (Cronbach's alpha coefficients for both indices ranged from 0.259 to 0.507 depending on the scenario), which suggests caution in interpreting results based on these indicators. The low reliability of the indices increases measurement error and requires caution when interpreting the findings.

After each scenario, the declared future behavioral intentions to reuse a given advice source were measured. Future behavioral intention was measured using the statement: “In a similar situation, next time I would take the advice of the [algorithm/expert].” Responses were given on a 5-point Likert scale (1 – “Strongly agree” to 5 – “Strongly disagree”). This item was reverse-scored for analysis, so that higher values indicated a greater willingness to reuse the advice.

The questionnaire also contained items on the frequency of using algorithmic tools and the advice of human experts, preferences regarding the source of support in various situations, and standard demographic questions.

Procedure

The study was conducted in the form of an online questionnaire using the LimeSurvey platform. The procedure began with a presentation of the aim of the study and obtaining the participant's informed consent. Participants were then randomly assigned to one of the four experimental groups. The final sample comprised $n = 54$ participants in the Algorithm–Taken condition, $n = 48$ in the Algorithm–Not-Taken condition, $n = 56$ in the Expert–Taken condition, and $n = 54$ in the Expert–Not-Taken condition. Within the assigned group, each participant was presented in turn

with the three decision scenarios. After reading each scenario, respondents were asked to step into the role of the protagonist and answer the RDS scale questions and the measure of future behavioral intention. After completing this part, all participants answered items on the use of algorithms and expert advice and on preferences regarding the source of support, and finally demographic questions. Before the main study, a pilot study ($N = 8$) was conducted to verify the clarity of the instructions, the realism of the scenarios, and the readability of the questions, which allowed for minor adjustments to the final version of the instrument.

Results

To examine the effects of advice source, whether the advice was followed, and decision scenario on regret, disappointment, and future intentions, a series of mixed analyses of variance was conducted. Advice source and advice-taking were between-subjects factors, whereas scenario was a three-level within-subjects factor.

Levene's tests indicated that the homogeneity-of-variance assumption was met at the individual measurement points for regret and future intention. For disappointment, the assumption was met in Scenarios 2 and 3 but was violated in Scenario 1, $p = .006$. Box's tests of equality of covariance matrices were significant for all analyses, all $ps \leq .032$, and the results should therefore be interpreted with some caution.

Mauchly's test indicated that the assumption of sphericity was violated for all three dependent variables: regret, $W = .968$, $p = .035$, $\epsilon GG = .969$; disappointment, $W = .852$, $p < .001$, $\epsilon GG = .871$; and future intention, $W = .887$, $p < .001$, $\epsilon GG = .898$. Consequently, Greenhouse–Geisser corrections were applied to all within-subjects effects involving the scenario factor.

Effect of advice source and whether the advice was taken on the level of regret

Analysis of the level of regret (measured with the regret scale, where higher values denote higher regret) revealed a statistically significant main effect of the within-subjects factor scenario, $F(1.94, 403.19) = 21.87$, $p < .001$, $\eta^2 = .095$. *Post-hoc* comparisons (with Bonferroni correction) indicated that the mean level of regret was significantly lower in the winter jacket scenario ($M = 3.82$, $SD = 0.05$) than in the airline ticket scenario ($M = 4.05$, $SD = 0.05$, $p < .001$) and in the stock investment

scenario ($M = 4.18$, $SD = 0.05$, $p < .001$). The difference between the airline ticket and stock investment scenarios was not significant ($p = .082$).

A significant scenario $\times$ advice-taken interaction was also found, $F(1.94, 403.19) = 14.98$, $p < .001$, $\eta^2 = .067$. Analysis of the estimated marginal means suggests a complex pattern of this interaction. When the advice had been followed, regret was relatively similar in the airline ticket ($M = 4.13$, $SE = 0.06$) and stock investment scenarios ($M = 4.02$, $SE = 0.07$) and was slightly lower in the winter jacket scenario ($M = 3.94$, $SE = 0.07$). When the advice had not been followed, a different pattern emerged: regret was highest in the stock investment scenario ($M = 4.33$, $SE = 0.07$), lower in the airline ticket scenario ($M = 3.98$, $SE = 0.07$), and lowest in the winter jacket scenario ($M = 3.69$, $SE = 0.08$). This shows that the influence of the situational context on the level of regret was strongly dependent on whether action was taken in accordance with the advice. This also indicates that the pattern of changes in the level of regret across particular scenarios was different for the group that took the advice compared to the group that did not.

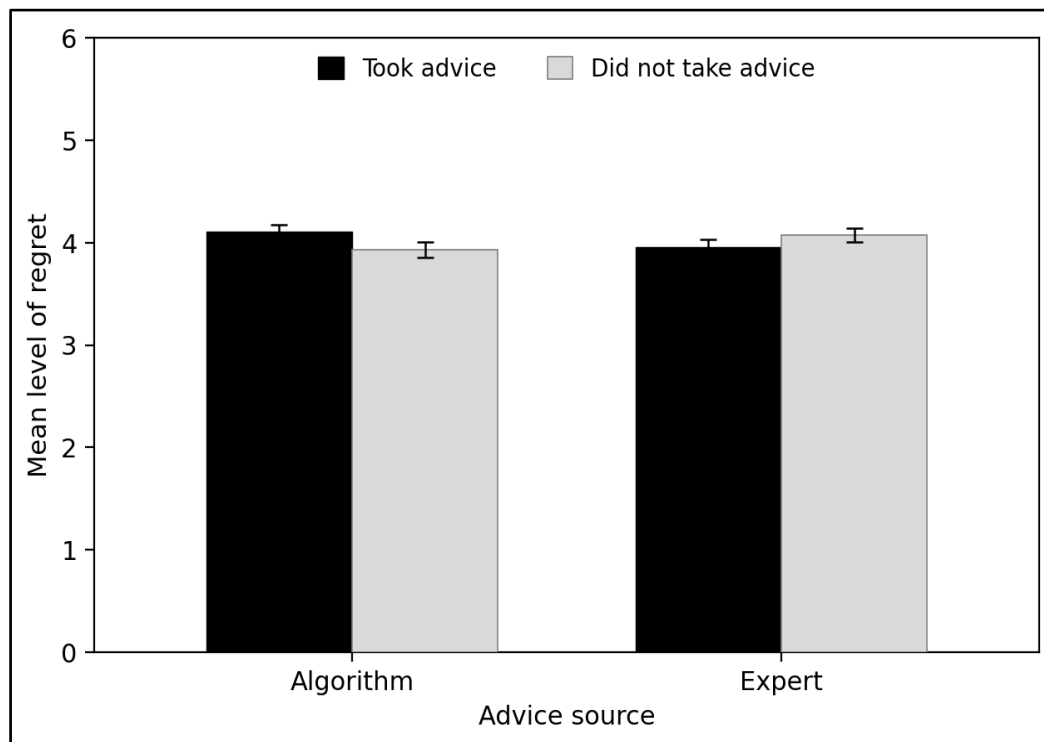
The scenario $\times$ advice source interaction was not significant, $F(1.94, 403.19) = 1.87$, $p = .156$, $\eta^2 = .009$. The three-way scenario $\times$ advice source $\times$ advice-taken interaction was also not significant, $F(1.94, 403.19) = 1.03$, $p = .356$, $\eta^2 = .005$.

The analysis of between-subjects effects showed no main effect of advice source, $F(1, 208) = 0.002$, $p = .967$, $\eta^2 < .001$, and no main effect of whether the advice had been followed, $F(1, 208) = 0.14$, $p = .707$, $\eta^2 = .001$. However, the advice source $\times$ advice-taken interaction was significant, $F(1, 208) = 3.94$, $p = .049$, $\eta^2 = .019$.

As shown in Figure 1, when the protagonist had followed the advice, regret was somewhat higher after algorithmic advice ($M = 4.10$, $SE = 0.07$) than after expert advice ($M = 3.96$, $SE = 0.07$). When the protagonist had rejected the advice, the pattern was reversed: regret was somewhat higher after expert advice ($M = 4.07$, $SE = 0.07$) than after algorithmic advice ($M = 3.93$, $SE = 0.08$). Neither of the simple comparisons between the algorithm and expert conditions was

individually significant, however: $p = .146$ in the advice-taken condition and $p = .179$ in the advice-not-taken condition. The crossover interaction should therefore be interpreted cautiously.

Figure 1. Mean level of experienced regret depending on the advice source and whether the advice was taken



Source: Own elaboration.

The presented results concerning the level of regret do not directly confirm hypothesis H1, which assumed an overall higher level of regret experienced after receiving unsuccessful advice from a human expert compared to algorithmic advice. The analysis did not show a significant main effect for the advice source ($p = .967$). What was observed, however, was a statistically significant interaction of advice source with whether the advice was taken ($p = .049$). This interaction reveals a more complex picture: the level of regret was higher after advice from the algorithm (compared to the expert) in the situation where participants (in the scenario) took the advice. The opposite tendency, consistent with the direction predicted in H1 (higher regret after advice from the expert), occurred only in the conditions in which participants did not take the advice received. Thus,

hypothesis H1 was not confirmed in its general form, and the effect of the advice source on the level of regret turned out to depend on the decision to use it.

Effect of advice source and whether the advice was taken on the level of disappointment

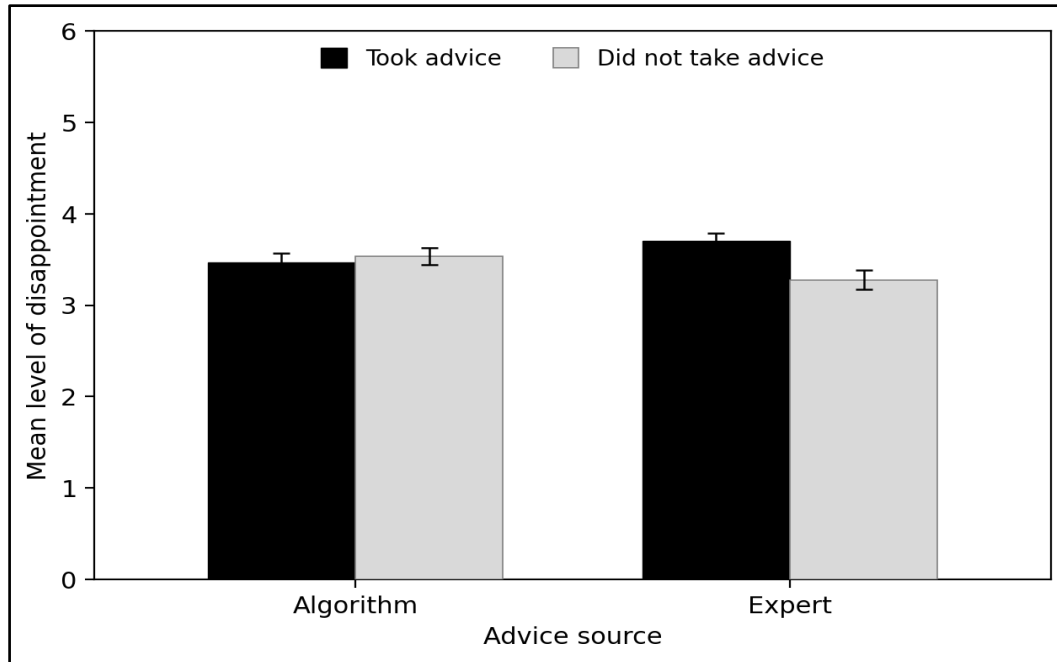
In the case of disappointment (measured with the Disappointment Index, where higher values denote a higher level of disappointment), ANOVA also revealed a significant main effect of the factor scenario, $F(1.94, 403.19) = 33.92, p < .001, \eta^2 = .140$. *Post-hoc* comparisons (with Bonferroni correction) revealed that the level of disappointment was significantly highest in the stock investment scenario ($M = 3.74, SE = 0.06$), lower in the airline ticket scenario ($M = 3.46, SE = 0.06, p < .001$ compared with scenario 3) and lowest in the winter jacket scenario ($M = 3.30, SE = 0.06$). All three pairwise differences were statistically significant, all $ps < .001$.

In contrast to the within-subjects pattern, none of the interactions involving the scenario factor reached significance: scenario $\times$ advice source ($p = .650$), scenario $\times$ advice-taken ($p = .304$), or the three-way interaction ($p = .981$). For the between-subjects effects, neither the main effect of advice source ($p = .912$) nor that of advice-taken ($p = .052$) was significant; the advice source $\times$ advice-taken interaction, however, was significant, $F(1, 208) = 6.88, p = .009, \eta^2 = .032$.

As illustrated in Figure 2, after expert advice, disappointment was significantly higher when the advice had been followed ($M = 3.71, SE = 0.09$) than when it had been rejected ($M = 3.28, SE = 0.09$), $p = .001$. In the algorithm condition, disappointment was similar when the advice had been followed ($M = 3.47, SE = 0.09$) and when it had been rejected ($M = 3.54, SE = 0.10$), $p = .645$. Comparisons between advice sources did not reach significance either when the advice had been followed, $p = .071$, or when it had been rejected, $p = .059$. Thus, the level of disappointment

depended not on the advice source alone but on its combination with whether the advice was followed.

Figure 2. Mean level of disappointment depending on the advice source and whether the advice was taken



Source: own elaboration.

The obtained results of the analysis of the level of disappointment do not provide support for hypothesis H2. According to H2, it was predicted that erroneous advice from an algorithm would result in a higher level of experienced disappointment than analogous advice given by a human expert. The analysis of variance carried out, however, did not show a statistically significant main effect of the factor advice source ($p = .912$) on the declared level of disappointment. Thus H2, framed as a main effect of advice source, was not supported. It should be noted, however, that a significant advice source $\times$ advice-taken interaction did emerge ($p = .009$): the source of the advice influenced disappointment only in combination with whether the advice had been followed, rather than as a main effect.

Future intentions to take advice

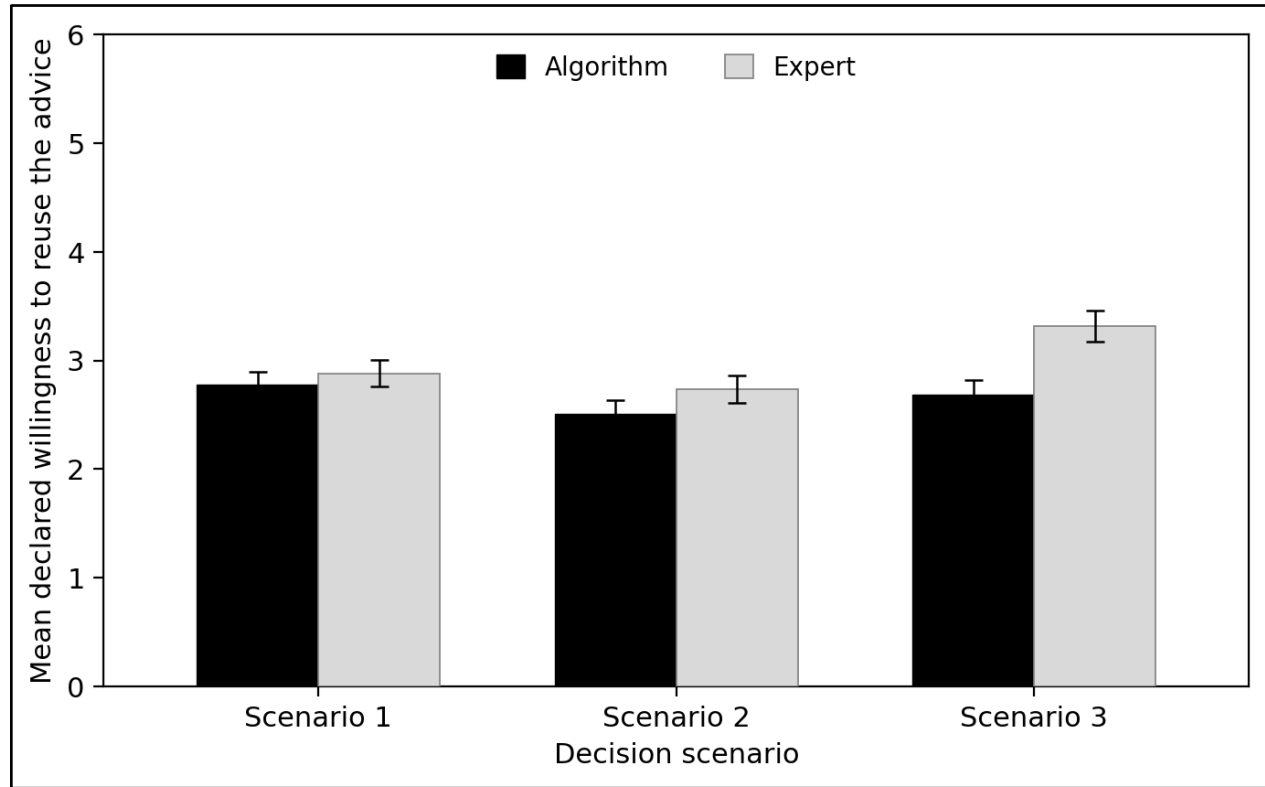
The analysis of the declared intention to reuse the advice in the future (where higher values denote a greater willingness to reuse it) revealed a significant main effect of the factor scenario, $F(1.94,$

403.19) = 13.56, $p < .001$, $\eta^2 = .061$. *Post-hoc* comparisons (with Bonferroni correction) revealed that the declared willingness to reuse the advice was significantly higher in the stock investment scenario ($M = 3.04$, $SE = 0.07$) than in the winter jacket ($M = 2.66$, $SE = 0.07$, $p < .001$), and also higher in the airline ticket scenario ($M = 2.86$, $SE = 0.06$) than in scenario 2 ($p = .004$). The difference between the airline ticket and stock investment scenarios was not significant, $p = .094$.

A significant scenario $\times$ advice source interaction effect was also found, $F(1.94, 403.19) = 6.75$, $p = .001$, $\eta^2 = .031$. For algorithmic advice, willingness to reuse the advice was relatively stable across the airline ticket ($M = 2.77$), winter jacket ($M = 2.51$), and stock investment scenarios ($M = 2.69$). For expert advice, willingness was distinctly higher in the stock investment scenario ($M = 3.32$) than in the airline ticket ($M = 2.88$) and winter jacket scenarios ($M = 2.74$). This suggests that the readiness to trust an expert again after an error was particularly high in the context of financial decisions, whereas for the algorithm the situational context was of lesser importance for future intentions. This means that the pattern of changes in declared intentions across the successive scenarios differed depending on whether the advice came from an algorithm or from an

expert. No significant scenario $\times$ advice-taken interaction ($p = .075$) or three-way interaction ($p = .378$) was found.

Figure 3. Mean declared willingness to reuse the advice depending on the advice source and the decision scenario



Source: own elaboration.

The analysis of between-subjects effects revealed a significant main effect of advice source, $F(1, 208) = 8.34, p = .004, \eta^2 = .039$. Participants reported greater willingness to use expert advice again ($M = 3.00, SE = 0.07$) than algorithmic advice ($M = 2.71, SE = 0.07$).

A considerably stronger main effect was found for whether the advice had been followed, $F(1, 208) = 323.86, p < .001, \eta^2 = .609$. Participants were substantially more willing to use the source again when the protagonist had rejected a recommendation that subsequently proved correct ($M = 3.75, SE = 0.07$) than when the protagonist had followed a recommendation that subsequently proved incorrect ($M = 1.95, SE = 0.07$). Because advice-taking was fully confounded with retrospective advice accuracy, this effect reflects a contrast between following an incorrect

recommendation and rejecting a recommendation that subsequently proved correct, rather than a pure effect of following versus rejecting advice.

The advice source $\times$ advice-taken interaction was not significant, $F(1, 208) = 1.11$, $p = .293$, $\eta^2 = .005$. Within the advice-taken condition, willingness to reuse the advice did not differ significantly between the algorithm ($M = 1.86$, $SE = 0.10$) and the expert ($M = 2.04$, $SE = 0.10$), $p = .188$. Within the advice-not-taken condition, willingness was higher for expert advice ($M = 3.95$, $SE = 0.10$) than for algorithmic advice ($M = 3.56$, $SE = 0.11$), $p = .007$. Because the overall interaction was not significant, the difference between the significance levels of these simple comparisons should not itself be interpreted as evidence that the source effect differed across advice-taking conditions.

The significant main effect of advice source is consistent with a general tendency toward lower willingness to rely on algorithmic than human advice. However, H3 specifically predicted that following incorrect algorithmic advice would reduce future willingness more strongly than following incorrect expert advice. This prediction was not supported: the source difference within the advice-taken condition was not significant, and the advice source $\times$ advice-taken interaction was also not significant. The results therefore provide evidence consistent with a general source preference but not with the specific post-error mechanism proposed in H3.

Discussion

The present study examined emotional and behavioral reactions to negative decision outcomes involving advice from an AI algorithm or a human expert. Specifically, it investigated whether reported regret, disappointment, and willingness to use the same source again depended on the source of the recommendation and on whether the protagonist had followed or rejected it. In the advice-taken conditions, the recommendation subsequently proved incorrect, whereas in the advice-not-taken conditions, the rejected recommendation subsequently proved correct. The findings indicate that reactions to advice cannot be explained by advice source alone. Instead, they depended on the combination of advice source, the decision to follow or reject the recommendation, and the situational context.

The central aim of the study was to compare regret and disappointment after negative outcomes involving algorithmic and human advice. Advice source did not have a significant main effect on either emotion. Thus, H1, which predicted greater regret after following incorrect advice

from a human expert, and H2, which predicted greater disappointment after following incorrect advice from an AI algorithm, were not supported. The significant interactions between advice source and whether the recommendation had been followed nevertheless indicate that emotional reactions were shaped by a more complex combination of these factors.

For regret, the advice source $\times$ advice-taking interaction reached statistical significance. When the recommendation had been followed and subsequently proved incorrect, regret was numerically higher after algorithmic than after expert advice. When the recommendation had been rejected and subsequently proved correct, the numerical pattern was reversed: regret was higher after expert than after algorithmic advice. However, neither of the simple comparisons between the algorithm and the expert conditions was statistically significant. The interaction should therefore be interpreted cautiously and cannot be treated as evidence that one source consistently elicited greater regret than the other.

The observed pattern is broadly consistent with the assumption that regret is linked to personal responsibility and choice-focused counterfactual thinking (Marcatto & Ferrante, 2008). Regret may depend not only on the identity of the advisor but also on how the decision-maker interprets their own role in accepting or rejecting the recommendation. Following advice that subsequently proves incorrect may direct attention toward the decision to trust the source, whereas rejecting advice that subsequently proves correct may make the forgone alternative particularly salient. The significant scenario $\times$ advice-taking interaction further indicates that these counterfactual processes varied across decision contexts. In particular, regret was especially high in the stock investment scenario when the recommendation had been rejected, possibly because the forgone financial gain and the failure to achieve the stated investment goal created a particularly salient counterfactual comparison. This interpretation remains tentative, however, because the scenarios differed on several dimensions, including domain, stakes, and the nature of the negative consequence.

A similar complexity emerged for disappointment. Although advice source did not have a significant main effect, the advice source $\times$ advice-taking interaction was significant. After expert advice, disappointment was greater when the recommendation had been followed and subsequently proved incorrect than when it had been rejected and subsequently proved correct. In the algorithm condition, disappointment was similar across the advice-taken and advice-not-taken conditions. Importantly, however, the direct comparisons between algorithmic and expert advice were not

statistically significant within either advice-taking condition. H2 was therefore not supported: in the condition in which the recommendation was followed and proved incorrect, disappointment was numerically higher after expert advice than after algorithmic advice, contrary to the predicted direction.

Disappointment also differed substantially between the decision scenarios. It was highest in the stock investment scenario and lowest in the winter jacket scenario. This pattern is compatible with the conceptualization of disappointment as an emotion associated with violated expectations and unfavorable states of the world (Marcatto & Ferrante, 2008). The stock investment scenario involved a 35% loss, a forgone 25% gain, and failure to accumulate funds for a mortgage down payment, which may have made its negative outcome especially salient. However, the scenarios also differed in how consequences were presented: the airline ticket scenario used absolute monetary amounts, the jacket scenario used changes in percentage discounts, and the investment scenario used percentage returns together with a salient long-term goal. The observed differences may therefore reflect perceived severity, numerical format, outcome salience, or other scenario-specific features rather than the decision domain itself. Because each domain was represented by only one scenario, these explanations cannot be distinguished in the present study.

The findings concerning future willingness to use the advice source again require a distinction between a general source preference and the specific post-error mechanism proposed in H3. Across conditions, participants reported greater willingness to reuse expert advice than algorithmic advice. This pattern is broadly consistent with research showing that people may prefer human judgment to algorithmic recommendations and may be more reluctant to rely on algorithmic systems (Dietvorst et al., 2015; Renier et al., 2021). However, the advice source $\times$ advice-taking interaction was not significant. Moreover, within the advice-taken condition – the condition in which the recommendation was incorrect – willingness to reuse the source did not differ significantly between the algorithm and the human expert. H3, which predicted lower subsequent willingness specifically after following incorrect algorithmic rather than incorrect expert advice, was therefore not supported.

The main effect of advice source should consequently not be interpreted as direct evidence that an algorithmic error produced a stronger decline in trust than a comparable human error. The study did not include a baseline measure of trust or willingness, and thus did not assess change over time. Furthermore, the general difference between sources was observed across both advice-

taking conditions, including the condition in which the rejected recommendation subsequently proved correct. The result is more accurately described as a general preference for future reliance on human rather than algorithmic advice within the scenarios used in the study.

The largest difference in future willingness emerged between the two advice-taking conditions. Participants were substantially more willing to use the source again after rejecting a recommendation that subsequently proved correct than after following a recommendation that subsequently proved incorrect. This finding should be interpreted as a contrast between two compound outcome structures rather than as evidence that following advice per se reduces subsequent reliance. In the advice-followed condition, accepting the recommendation was combined with the recommendation proving inaccurate and contributing to the negative outcome. In the advice-rejected condition, rejecting the recommendation was combined with the recommendation proving accurate and with the realization that following it would have produced a better result.

Several mechanisms may therefore account for the observed difference. Participants may have updated their evaluation of the source on the basis of its retrospective accuracy, attributed the loss more directly to a source whose recommendation had been followed, or learned counterfactually that the rejected recommendation should have been accepted. Differences in personal responsibility associated with accepting versus rejecting advice may also have contributed. Because advice-taking and retrospective accuracy were not manipulated independently, the present study cannot determine the relative contribution of these mechanisms.

Taken together, the results suggest that emotional responses and future behavioral intentions were related to advice in different ways. Advice source had no main effect on regret or disappointment but was associated with future willingness to use the source again. This divergence indicates that a preference for human advice does not necessarily result from stronger immediate emotional reactions to algorithmic advice. Future reliance may instead depend on broader beliefs about the credibility, flexibility, accountability, or capacity for improvement of human and algorithmic advisors. These mechanisms were not directly measured in the present study and should be examined in future research.

The different patterns are compatible with the theoretical distinction between regret and disappointment, although the low reliability of the RDS indices requires substantial caution. Regret varied with the relationship between scenario and advice-taking and showed a crossover interaction

between advice source and advice-taking. Disappointment was strongly affected by scenario and was particularly elevated when incorrect expert advice had been followed. These results suggest that regret and disappointment may reflect partly different counterfactual and attributional processes. However, given the weak internal consistency of the two-item indices, the observed differences should not be treated as strong evidence regarding the precise psychological mechanisms underlying the emotions.

Relevance to contemporary AI advice

The algorithmic advisor used in the present study was described as a relatively narrow predictive system that generated a discrete recommendation. This operationalization differs from contemporary generative AI and large language models such as ChatGPT, Claude, or Gemini, which provide conversational, elaborate, and often human-like responses. These systems may be perceived not merely as computational tools but as interactive advisors capable of explaining and justifying their recommendations (Osborne & Bailey, 2025; Porsdam Mann et al., 2023). Generalizing the present findings to such systems therefore requires caution.

Recent theoretical approaches suggest that preferences for algorithms or humans depend on both the perceived capability of the advisor and the extent to which a task requires personalization. Jussupow et al. (2024) argued that algorithm aversion and algorithm appreciation may reflect context-dependent manifestations of the same underlying preference structure. Similarly, the capability–personalization framework proposed by Qin et al. (2025) suggests that algorithms are more likely to be preferred when they are perceived as highly capable and when the task does not require individualized understanding. The present scenarios differed in their likely demands for numerical analysis, subjective judgment, and personalization. However, because perceived capability and personalization requirements were not measured, the observed scenario differences cannot be directly attributed to these mechanisms. Future studies should assess or experimentally manipulate them.

The conversational style of contemporary LLMs may also alter reactions to their advice. Anthropomorphic linguistic cues can increase perceived credibility and encourage social responses to an artificial system, but they may also raise expectations and intensify negative reactions when the system fails (Cohn et al., 2024; Crolic et al., 2022; Peter et al., 2025). The neutral description of an “artificial intelligence algorithm” used in the present study did not include conversational

interaction, explanations, or human-like language. Future research should therefore compare neutral predictive algorithms with anthropomorphically presented AI advisors and examine whether the presentation style changes responsibility attribution, emotional reactions, and subsequent reliance.

The type of algorithmic error may also be important. The scenarios in the present study described incorrect predictions about future prices or market outcomes. Generative AI systems may instead produce hallucinations: plausible but factually fabricated information (Christensen et al., 2025; Rejón-Guardia et al., 2026). These error types may produce different attributional and emotional responses. An incorrect prediction can be attributed partly to an inherently uncertain external environment, whereas a hallucination originates in information generated by the system itself. Comparing predictive errors with hallucinated information may therefore help clarify when algorithmic failures elicit regret, disappointment, blame, or withdrawal of trust.

Practical implications

The practical implications of the findings should be formulated cautiously. The study does not demonstrate that users are necessarily less tolerant of AI errors than of human errors, because willingness to reuse algorithmic and expert advice did not differ significantly within the condition in which the recommendation was followed and proved incorrect. It does, however, indicate a general tendency toward lower future reliance on algorithmic than human advice across the scenarios.

Organizations implementing AI-based decision support should therefore consider not only the objective accuracy of their systems but also the broader beliefs users hold about algorithmic and human advisors. Transparent communication about the system's capabilities, limitations, and uncertainty may help users develop more calibrated expectations. When an AI system provides an incorrect recommendation, explanations of why the error occurred and what corrective action has been taken may support the rebuilding of appropriately calibrated trust. Such strategies should not aim to maximize reliance indiscriminately, but rather to help users decide when algorithmic support is and is not warranted.

The scenario effects indicate that reactions to advisory systems may be sensitive to the broader decision context. However, they should not be interpreted as evidence of differences between decision domains, because domain was confounded with outcome magnitude, numerical

format, salience, and the nature of the negative consequence. Establishing domain-specific effects would require multiple scenarios within each domain and outcomes matched or pretested for perceived severity.

Limitations and future research

The study has several limitations. First, participants responded to hypothetical scenarios rather than making consequential decisions. The measures therefore captured imagined or scenario-based emotional reactions and declared behavioral intentions, which may be less intense and less predictive of actual behavior than reactions to real outcomes. Future research should use incentivized experimental tasks, repeated interactions, or field studies involving genuine consequences.

Second, the sample consisted predominantly of young and well-educated adults. Age, education, digital competence, and prior experience with AI may influence expectations of algorithmic systems and willingness to rely on them. More demographically diverse samples are needed to determine the generalizability of the findings.

Third, the reliability of the Regret and Disappointment Scale indices was low. Because each index consisted of only two items, the resulting scores were particularly vulnerable to measurement error. The weak internal consistency substantially limits conclusions about differences in regret and disappointment. Future research should use more reliable multi-item measures, validate the Polish version of the instrument, and consider combining self-report measures with behavioral or physiological indicators. The possibility that translation, scenario adaptation, or cultural interpretation contributed to the low reliability should be examined empirically rather than assumed.

Fourth, the manipulation of advice-taking was fully confounded with retrospective advice accuracy. In the advice-taken condition, accepting the recommendation was always combined with the recommendation proving incorrect, whereas in the advice-not-taken condition, rejecting the recommendation was always combined with the recommendation proving correct. Consequently, effects involving advice-taking cannot be attributed independently to the act of following or rejecting advice, the accuracy of the recommendation, causal attribution of the outcome to the source, or the salience of the forgone alternative. This limitation applies not only to future willingness but also to the interactions involving advice-taking observed for regret and

disappointment. Future studies should manipulate advice-taking and recommendation accuracy orthogonally, creating conditions in which followed and rejected recommendations can prove either correct or incorrect.

Fifth, the study included only three scenarios and treated the algorithm as a homogeneous category. Each decision domain was represented by a single scenario, and the scenarios differed not only in content but also in outcome magnitude, numerical format, salience, and stakes. The airline ticket scenario expressed the loss in absolute monetary amounts, the jacket scenario in percentage discounts, and the investment scenario in percentage returns and linked the outcome to a salient mortgage-related goal. Because perceived severity was neither measured nor equated, scenario effects cannot be attributed specifically to decision domain. Future studies should include multiple scenarios within each domain and pretest or match the perceived severity and salience of their outcomes. They should also manipulate specific features of the AI system, such as transparency, explainability, autonomy, and interaction style.

Sixth, individual differences such as general attitudes toward AI, technology readiness, risk propensity, attributional style, and prior use of advisory systems were not incorporated as moderators. Including these variables could help explain why some people show algorithm aversion while others display algorithm appreciation.

Finally, the rapid development of generative AI limits the temporal and ecological generalizability of a manipulation based on a neutrally described predictive algorithm. Contemporary users may associate AI advice with conversational LLMs rather than with a discrete recommendation generated by an impersonal system. Future research should employ actual AI outputs, manipulate disclosure of the source, and vary the degree of anthropomorphic presentation.

Conclusion

The present study suggests that reactions to advice depend on more than whether the source is an AI algorithm or a human expert. Advice source did not have a direct main effect on reported regret or disappointment, although both emotions were shaped by interactions between source and whether the recommendation had been followed. Neither H1 nor H2 was supported, and the specific post-error difference predicted in H3 was also not confirmed.

Participants nevertheless expressed a general preference for relying on expert rather than algorithmic advice in the future. The strongest result concerned the distinction between following

an incorrect recommendation and rejecting a recommendation that subsequently proved correct. This effect is best interpreted as counterfactual learning from the outcome, although the design does not permit the independent effects of advice-taking and advice accuracy to be separated. Overall, the findings highlight the importance of jointly considering advice source, the decision to accept or reject a recommendation, its subsequent accuracy, and the situational context. They also show that emotional reactions and future reliance need not follow the same pattern. Further research using more reliable measures, orthogonal manipulations of advice-taking and accuracy, consequential decisions, and contemporary conversational AI systems is needed to clarify how people respond when algorithmic and human advice succeeds or fails.

References

- Aschauer, F., Sohn, M., & Hirsch, B. (2023). Managerial advice-taking: Sharing responsibility with (non)human advisors trumps decision accuracy. *European Management Review*, 21(1), 186–203. <https://doi.org/10.1111/emre.12575>
- Bell, D. E. (1982). Regret in decision making under uncertainty. *Operations Research*, 30(5), 961–981. <https://doi.org/10.1287/opre.30.5.961>
- Bell, D. E. (1985). Disappointment in decision making under uncertainty. *Operations Research*, 33(1), 1–27. <https://doi.org/10.1287/opre.33.1.1>
- Burton, J. W., Stein, M. K., & Jensen, T. B. (2020). A systematic review of algorithm aversion in augmented decision making. *Journal of Behavioral Decision Making*, 33(2), 220–239. <https://doi.org/10.1002/bdm.2155>
- Christensen, J., Hansen, J. M., & Wilson, P. (2025). Understanding the role and impact of generative artificial intelligence (AI) hallucination within consumers' tourism decision-making processes. *Current Issues in Tourism*, 28(4), 545–560. <https://doi.org/10.1080/13683500.2023.2300032>
- Coeckelbergh, M. (2020). Artificial intelligence, responsibility attribution, and a relational justification of explainability. *Science and Engineering Ethics*, 26(4), 2051–2068. <https://doi.org/10.1007/s11948-019-00146-8>
- Cohn, M., Pushkarna, M., Olanubi, G. O., Moran, J. M., Padgett, D., Mengesha, Z., & Heldreth, C. (2024). Believing anthropomorphism: Examining the role of anthropomorphic cues on trust in large language models. In: *Extended Abstracts of the CHI Conference on Human Factors in Computing Systems*, 1–15. <https://doi.org/10.48550/arXiv.2405.06079>
- Crolic, C., Thomaz, F., Hadi, R., & Stephen, A. T. (2022). Blame the bot: Anthropomorphism and anger in customer–chatbot interactions. *Journal of Marketing*, 86(1), 132–148. <https://doi.org/10.1177/00222429211045687>
- Daschner, S., & Obermaier, R. (2022). Algorithm aversion? On the influence of advice accuracy on trust in algorithmic advice. *Journal of Decision Systems*, 31(sup1), 77–97. <https://doi.org/10.1080/12460125.2022.2070951>
- Dastin, J. (2018, October 10). *Amazon scraps secret AI recruiting tool that showed bias against women*. Reuters. <https://www.reuters.com/article/us-amazon-com-jobs-automation->

insight/amazon-scrapes-secret-ai-recruiting-tool-that-showed-bias-against-women-
idUSKCN1MK08G

Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, *144*(1), 114–126. <https://doi.org/10.1037/xge0000033>

Dzindolet, M. T., Peterson, S. A., Pomranky, R. A., Pierce, L. G., & Beck, H. P. (2003). The role of trust in automation reliance. *International Journal of Human-Computer Studies*, *58*(6), 697–718. [https://doi.org/10.1016/s1071-5819\(03\)00038-7](https://doi.org/10.1016/s1071-5819(03)00038-7)

Esmailzadeh, P., Sambasivan, M., Kumar, N., & Nezakati, H. (2015). Adoption of clinical decision support systems in a developing country: Antecedents and outcomes of physician's threat

to perceived professional autonomy. *International Journal of Medical Informatics*, 84(8), 548–560. <https://doi.org/10.1016/j.ijmedinf.2015.03.007>

Giorgetta, C. (2012). *The emotional side of decision-making: Regret and disappointment*. Saarbrücken: Lambert Academic Publishing.

Jussupow, E., Benbasat, I., & Heinzl, A. (2024). An integrative perspective on algorithm aversion and appreciation in decision-making. *MIS Quarterly*, 48(4), 1575–1590. <https://doi.org/10.25300/MISQ/2024/18512>

Kahneman, D., & Tversky, A. (1982). The simulation heuristic. In: D. Kahneman, P. Slovic, & A. Tversky (Eds.), *Judgment under uncertainty: Heuristics and biases* (pp. 201–208). Cambridge: Cambridge University Press. <https://doi.org/10.1017/CBO9780511809477.015>

Lambrecht, A., & Tucker, C. (2019). Algorithmic bias? An empirical study of apparent gender-based discrimination in the display of STEM career ads. *Management Science*, 65(7), 2966–2981. <https://doi.org/10.1287/mnsc.2018.3093>

Lee, C. S., Nagy, P. G., Weaver, S. J., & Newman-Toker, D. E. (2013). Cognitive and system factors contributing to diagnostic errors in radiology. *American Journal of Roentgenology*, 201(3), 611–617. <https://doi.org/10.2214/ajr.12.10375>

Loomes, G., & Sugden, R. (1982). Regret theory: An alternative theory of rational choice under uncertainty. *The Economic Journal*, 92(368), 805–824. <https://doi.org/10.2307/2232669>

Madhavan, P., & Wiegmann, D. A. (2007). Similarities and differences between human-human and human-automation trust: An integrative review. *Theoretical Issues in Ergonomics Science*, 8(4), 277–301. <https://doi.org/10.1080/14639220500337708>

Marcatto, F., & Ferrante, D. (2008). The Regret and Disappointment Scale: An instrument for assessing regret and disappointment in decision making. *Judgment and Decision Making*, 3(1), 87–99. <https://doi.org/10.1017/s193029750000019x>

Markus, M. L., Dutta, A., Steinfield, C. W., & Wigand, R. T. (2008). The computerization movement in the US home mortgage Industry: Automated underwriting from 1980 to 2004. In: K.

Kraemer & M. Elliott (Eds.), *Computerization movements and technology diffusion: From mainframes to ubiquitous computing* (pp. 115–144). Medford, New Jersey: Information Today.

Matthias, A. (2004). The responsibility gap: Ascribing responsibility for the actions of learning automata. *Ethics and Information Technology*, 6(3), 175–183. <https://doi.org/10.1007/s10676-004-3422-1>

Meehl, P. E. (1954). *Clinical versus statistical prediction: A theoretical analysis and a review of the evidence*. Minneapolis: University of Minnesota Press. <https://doi.org/10.1037/11281-000>

Osborne, M. R., & Bailey, E. R. (2025). Me vs. the machine? Subjective evaluations of human- and AI-generated advice. *Scientific Reports*, 15, 3980. <https://doi.org/10.1038/s41598-025-86623-6>

Parasuraman, A., & Colby, C. L. (2015). An updated and streamlined technology readiness index: TRI 2.0. *Journal of Service Research*, 18(1), 59–74. <https://doi.org/10.1177/1094670514539730>

Peter, S., Riemer, K., & West, J. D. (2025). The benefits and dangers of anthropomorphic conversational agents. *Proceedings of the National Academy of Sciences*, 122(22), e2415898122. <https://doi.org/10.1073/pnas.2415898122>

Porsdam Mann, S., Earp, B. D., Nyholm, S., Danaher, J., Möller, N., Bowman-Smart, H., Hatherley, J., Koplin, J., Plozza, M., Rodger, D., Treit, P. V., Renard, G., McMillan, J., &

Savulescu, J. (2023). Generative AI entails a credit–blame asymmetry. *Nature Machine Intelligence*, 5(5), 472–475. <https://doi.org/10.1038/s42256-023-00653-1>

Qin, X., Zhou, X., Chen, C., Wu, D., Zhou, H., Dong, X., Cao, L., & Lu, J. G. (2025). AI aversion or appreciation? A capability-personalization framework and a meta-analytic review. *Psychological Bulletin*, 151(5), 580–599. <https://doi.org/10.1037/bul0000477>

Rabinovitch, H., Budescu, D. V., & Bereby-Meyer, Y. (2024). Algorithms in selection decisions: Effective, but unappreciated. *Journal of Behavioral Decision Making*, 37(2), e2368. <https://doi.org/10.1002/bdm.2368>

Rejón-Guardia, F., Molinillo, S., & Anaya-Sánchez, R. (2026). AI hallucinations in tourism: How errors impact consumer trust and recommendation acceptance. *Journal of Consumer Behaviour*, 25(2), 923–938. <https://doi.org/10.1002/cb.70105>

Renier, L. A., Schmid Mast, M., & Bekbergenova, A. (2021). To err is human, not algorithmic – Robust reactions to erring algorithms. *Computers in Human Behavior*, 124, 106879. <https://doi.org/10.1016/j.chb.2021.106879>

Saxena, S. (2024). Human-AI interaction: psychological perspectives. *International Journal of Advanced Multidisciplinary Scientific Research*, 1(10), 205–209. <https://doi.org/10.31426/ijamsr.2018.1.10.1028>

Straka, J. W. (2000). A shift in the mortgage landscape: The 1990s move to automated credit evaluations. *Journal of Housing Research*, 11(2), 207–232.

Sunstein, C. R., & Gaffe, J. H. (2025). An anatomy of algorithm aversion. *Science and Technology Law Review*, 26(1). <https://doi.org/10.52214/stlr.v26i1.13339>

Tschandl, P., Codella, N., Akay, B. N., Argenziano, G., Braun, R. P., Cabo, H., & Rosendahl, C. (2019). Comparison of the accuracy of human readers versus machine-learning algorithms for pigmented skin lesion classification: an open, web-based, international, diagnostic study. *The Lancet Oncology*, 20(7), 938–947. [https://doi.org/10.1016/s1470-2045\(19\)30333-x](https://doi.org/10.1016/s1470-2045(19)30333-x)

Upadhyay, A. K., & Khandelwal, K. (2018). Applying artificial intelligence: Implications for recruitment. *Strategic HR Review*, 17(5), 255–258. <https://doi.org/10.1108/shr-07-2018-0051>

Wang, L., Li, X., Zhu, H., & Zhao, Y. (2025). A dimensional exploration and scale development study of algorithm aversion. *Journal of the Operational Research Society*, 76(8), 1531–1552. <https://doi.org/10.1080/01605682.2024.2419544>

Weiner, B. (1985). An attributional theory of achievement motivation and emotion. *Psychological Review*, 92(4), 548–573. <https://doi.org/10.1037/0033-295x.92.4.548>

Zeelenberg, M., Inman, J. J., & Pieters, R. G. M. (2001). What we do when decisions go awry: Behavioral consequences of experienced regret. In: E. U. Weber, J. Baron, & G. Loomes (Eds.), *Conflict and tradeoffs in decision making* (pp. 136–155). Cambridge: Cambridge University Press. <https://doi.org/10.1002/bdm.424>

Zeelenberg, M., van Dijk, W. W., Manstead, A. S. R., & van der Pligt, J. (2000). On bad decisions and disconfirmed expectancies: The psychology of regret and disappointment. *Cognition and Emotion*, 14(4), 521–541. <https://doi.org/10.1080/026999300402781>

Appendix A: Measurement instruments used after scenarios

The following scales and questions, preceded by an instruction, were used to assess participants' reactions after reading each scenario.

Instruction for the participant: "Let us imagine that the described situation actually happened. Read each question carefully and rate the extent to which you agree with the following statements (for questions 1–6) or choose one answer (for question 7). Your answer should reflect your feelings."

Rating Likert scale for questions 1–6:

- Strongly agree
- Agree
- Neither agree nor disagree
- Disagree
- Strongly disagree

Statements (questions 1–6):

1. I would be sorry about what happened to me.
2. I would wish that my choice had been different.
3. I would wish that events I had no influence over had turned out differently.
4. I feel responsible for what happened to me.
5. Events I had no influence over are the cause of what happened to me.
6. I am satisfied with what happened to me.

Forced-choice question (question 7): 7. In your opinion, what could have led to a more favorable outcome in the described situation involving the ticket purchase? (a) If I had decided otherwise and not followed the algorithm's advice. (b) If the circumstances had been different.

Appendix B: Decision scenarios used in the study

Each participant was presented with three decision scenarios concerning: (1) the purchase of airline tickets, (2) the purchase of a winter jacket, and (3) a stock market investment. The scenarios were adapted to four experimental conditions resulting from the combination of two between-subjects factors: advice source (AI algorithm vs. human expert) and the protagonist's response to the advice (advice followed vs. advice rejected).

The advice source was manipulated by replacing references to an “artificial intelligence algorithm” with references to a “human expert.” The recommendation itself remained the same within each scenario. The advice-taking manipulation determined both the protagonist’s decision and the subsequent outcome. When the recommendation was followed, it subsequently proved incorrect; when it was rejected, it subsequently proved correct, although the protagonist experienced a negative outcome.

Table B1. Summary of Differences Between Conditions

Scenario	Recommendation	Advice followed	Advice rejected
Airline tickets	Postpone the purchase because prices may fall.	The protagonist postpones the purchase. The price increases from PLN 542 to PLN 778.	The protagonist buys immediately. The price subsequently falls from PLN 778 to PLN 542.
Winter jacket	Postpone the purchase because the discount may increase.	The protagonist postpones a purchase at a 20% discount. The discount later decreases to 10%.	The protagonist buys at a 10% discount. The discount later increases to 20%.
Stock investment	Invest in shares of company X.	The protagonist invests in X. X loses 35%, whereas Y gains 25%.	The protagonist rejects the advice and invests in Y. Y loses 35%, whereas X gains 25%.