

AI-Based Decision Making: Not the Decision-Maker but the Outcome's Favorability Determines the Perception of University Topic Allocations

Christopher Esch, Farid Fares, Elisabeth Kals, and Christina U. Pfeuffer

Department of Psychology, Catholic University of Eichstätt-Ingolstadt

Artificial intelligence (AI) is increasingly used for decision making, but its perception compared to human decision-makers remains underexplored, especially in educational contexts. In this study, we investigated the influence of the decision-maker (AI vs. human) on fairness perception, trust, and emotional responses in participants ($N = 329$) who are allocated university course topics. While the allocation process was identical, participants were told that either a human lecturer or an AI made the decision based on the same pieces of information. Furthermore, we manipulated whether the corresponding decision-making process was explicitly communicated as fair or whether no comment was made regarding the process' fairness. Finally, we assessed how favorable students rated the outcome of the allocation process. Against our hypotheses, Bayesian evidence indicated that neither the decision-maker nor whether the decision-making process was communicated as fair had an impact on students' fairness perception, trust, or emotional responses. Students' evaluations of the university course topic allocation process were strongly associated with the favorability of the outcome. Given that an AI can better optimize allocations according to students' preferences than human decision-makers, these findings support a broader implementation of AI-based decision making in the context of allocation decisions at university.

Keywords: artificial intelligence, decision support, fairness, trust, outcome favorability

Artificial intelligence (AI) has become a pervasive force in decision making across various sectors like finance, health care, and education (Starke et al., 2022). The rapid adoption of AI technologies, exemplified by the exponential growth of ChatGPT reaching 100 million users in just 2 months (K. Hu, 2023), highlights their profound relevance. While AI systems promise efficiency and scalability, users' concerns about their fairness and trustworthiness persist, particularly when algorithms are used to support sensitive decisions (Chouldechova, 2016; Marcinkowski et al., 2020).

These concerns are not unfounded, as AI-based decision support systems, despite their potential to reduce human biases, can inadvertently perpetuate or exacerbate inequalities. For example, flawed training data or suboptimal design has led to cases such as the COMPAS algorithm disproportionately labeling Black criminal defendants as high-risk reoffenders (Chouldechova, 2016). Nonetheless, algorithms and AI systems are often not perceived as fair and trustworthy as would be appropriate based on their characteristics (e.g., Dietvorst et al., 2015; Saxena et al., 2019). This appears to be the case because algorithmic decision making

(i.e., based on predetermined rules and statistical models, M. K. Lee, 2018) and especially AI-based decision making (i.e., most often based on black box computations, Ali et al., 2023) are most often perceived as abstract and opaque, adding uncertainty and complexity to human–algorithm/human–AI interactions. This is especially the case when AI replaces human decision-makers in contexts like employment or education (Bankins et al., 2022). Although work on explainable AI and initiatives to promote AI literacy seek to address this opacity, their practical adoption remains limited (Doshi-Velez & Kim, 2017; Ng et al., 2021).

This further highlights the importance of understanding how fairness perceptions in AI-based decision making are shaped to support users in developing appropriate levels of trust to accept and adopt fair and trustworthy AI decision support (e.g., Kelly et al., 2023; Shin, 2021). Here, we assessed the corresponding impact of the decision-maker (human vs. AI) and the communicated fairness of an allocation decision in the educational context. Specifically, we assessed the perception of university course topic allocations by human lecturers versus a (supposed) AI.

This article was published Online First March 9, 2026.

Action Editor: Richard Landers was the action editor for this article.

ORCID iD: Christopher Esch  <https://orcid.org/0009-0001-6891-690X>.

Funding: The authors received no third-party funding for this research project.

Disclosures: The authors have no conflicts of interest to disclose.

Author Contributions: Christopher Esch and Farid Fares share joint first authorship.

Data Availability: The study's design, research question, and hypotheses were preregistered on the Open Science Framework (Fares et al., 2023; <https://osf.io/7xnr/overview>). The data, study materials, and analysis scripts are available on the Open Science Framework (Esch, Fares, Pfeuffer, & Kals, 2025; <https://osf.io/qe52u/overview>). Data were newly collected for this

project. Results of this project were previously presented in the form of a preprint on the Open Science Framework (Esch, Fares, Kals, & Pfeuffer, 2025; https://osf.io/preprints/osf/cwxgt_v1).

Open Access License: This work is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License (CC BY-NC-ND). This license permits copying and redistributing the work in any medium or format for noncommercial use provided the original authors and source are credited and a link to the license is included in attribution. No derivative works are permitted under this license.

Contact Information: Correspondence concerning this article should be addressed to Christopher Esch, Department of Psychology, Catholic University of Eichstätt-Ingolstadt, Ostenstraße 25, 85072 Eichstätt, Germany. Email: christopher.esch@ku.de

Fairness Perceptions in AI-Based Decision Making

Fairness perception pertains to how individuals assess the fairness of decision-making processes and outcomes (Greenberg, 1987; Grunenberg et al., 2024). In organizational contexts, four key dimensions of justice are recognized (Colquitt, 2001): distributive justice (fairness of the outcomes), procedural justice (fairness of the processes used to arrive at the outcomes), interpersonal justice (fairness of the treatment individuals receive), and informational justice (fairness in communication and explanations). All four dimensions of fairness perception also apply to AI-based decisions (Narayanan et al., 2024). In organizational contexts, these dimensions are often referred to as justice dimensions, reflecting normative standards and principles. In the following, we exclusively focus on the subjective perception of these aspects of justice, specifically how individuals feel about an allocation in the educational context. We therefore use the term “perceived fairness” in the following to refer to the subjective perception of justice (Colquitt et al., 2005).

Positive fairness perceptions can foster cooperation, satisfaction, and acceptance, while negative perceptions can harm well-being and lead to disruptive behaviors (human decision-makers: e.g., Colquitt et al., 2013; AI decision-makers: e.g., Bankins et al., 2022; Kelly et al., 2023). Procedural fairness is a particularly influential sub-dimension of fairness perception. Perceiving procedures as fair can, in turn, lead to more positive perceptions of outcomes, even if the outcomes are unfavorable (Brockner & Wiesenfeld, 1996; Folger et al., 1979; Grgić-Hlača et al., 2018). This *fair process effect* highlights the importance of communicating the fairness of procedures in order to enhance perceptions of fairness and acceptance of decisions. We assumed that perceived fairness plays a crucial role in shaping the perception of human versus AI-based allocation decisions in educational settings. Building on findings regarding the fair process effect, we manipulated whether participants received explicit communication about the fairness of the allocation process or not to further investigate its impact on perceived fairness.

Trust and Emotional Responses in AI-Based Decision Making

Trust in decision-makers, whether human or AI, similarly plays an important role in how individuals react to decision-making systems (Bianchi et al., 2014; M. K. Lee, 2018), as humans tend to apply social rules to computers (see Madhavan & Wiegmann, 2007). Based on the conceptual model of Mayer et al. (1995), which has been adapted to the context of AI and automation (see Kohn et al., 2021, for an overview), trust goes beyond a purely technical assessment of reliability. Instead, trust can be considered as a relational assessment based on factors such as ability (competence), integrity (adherence to principles), and benevolence (intent to do good; Solberg et al., 2022).

Moreover, fairness perceptions significantly affect emotions (Krehbiel & Cropanzano, 2000) and trust, as AI systems perceived as fair and transparent can build trust (Shin, 2021; Wang et al., 2020). However, perceived unfairness or errors can lead to *algorithm aversion* and decreased trust (Dietvorst et al., 2015). Unlike humans, algorithms are often viewed as incapable of learning from mistakes, which can further exacerbate distrust (Dietvorst et al., 2015). Especially in the educational context, this is critical: If students do not perceive the AI-based system to act in their best interest (see

benevolence), trust may erode, in turn, reducing acceptance and actual use (e.g., see Ayanwale & Ndlovu, 2024; Low et al., 2025) regardless of how mathematically optimal the system is. Thus, students’ perceptions of any implemented AI system are crucial for the system’s evaluation and actual use. At the same time, students’ evaluation of AI systems used by universities can affect the reputation of these educational institutions (Marcinkowski et al., 2020).

In line with the link between trust and fairness, emotional responses to AI-based decisions are likewise closely tied to how fair a process is perceived to be. Lack of transparency and interpersonal engagement can evoke negative emotions such as frustration, distrust, or dissatisfaction (Holmvall & Bobocel, 2008). The impersonal nature of algorithms may make individuals feel that their unique circumstances are overlooked, especially in social contexts (M. K. Lee & Baykal, 2017). Conversely, when AI systems are perceived as fair, efficient, and transparent, they can foster positive emotions like satisfaction (M. K. Lee, 2018).

Comparative Effects of AI-Based and Human Decision Making

Research comparing AI-based and human decision making reveals mixed findings, indicating that fairness perceptions are context-dependent (Starke et al., 2022). In some settings, AI decision making is perceived more positively due to its efficiency and objectivity, such as in warehouse assignments and university admissions (Bai et al., 2020; Marcinkowski et al., 2020). In the educational sector, Marcinkowski et al. (2020) found that students perceived algorithmic selection committees as procedurally and distributively fairer and less biased than human admission committees. Additionally, AI systems often score higher on consistency, a key component of procedural fairness (Kaibel et al., 2019).

This acceptance, however, is not universal in education. Some studies show that students may react negatively to algorithmic grading when the results conflict with their performance expectations. Kizilcec (2016) found that when such discrepancies occur, detailed transparency about the algorithm does not mitigate student distrust, a finding also shown in research on imperfect automated grading systems (Hsu et al., 2021). Furthermore, a recent study suggests a general reluctance toward AI in university grading, showing that students tend to rate human professors as fairer than AI graders (Jones-Jang et al., 2025). In other domains that are presumed to require human contributions like health care, work-related decisions, and criminal justice, human decision making is often perceived as fairer (Acikgoz et al., 2020; Harrison et al., 2020; M. K. Lee & Rich, 2021). In such contexts, human decision-makers tend to receive higher ratings in interpersonal fairness due to their ability to convey empathy and understanding (Schlicker et al., 2021).

Task complexity further influences fairness perceptions. Humans are preferred over AI decision-makers in high-complexity tasks requiring judgment and nuance, where distributive and interpersonal fairness are paramount (Nagtegaal, 2021). In contrast, AI is preferred in low-complexity tasks that prioritize consistency and data-driven decisions, aligning with procedural fairness. Other conditional factors like social group representation can affect perceptions. For instance, individuals may perceive AI decisions as fairer when their social group lacks representation in human decision making (Miller & Keiser, 2021).

Purpose of the Present Study

Prior research has established that fairness, trust, and emotional responses are critical factors shaping individuals' acceptance of decision-making processes (Colquitt et al., 2013; M. K. Lee, 2018). However, many prior studies showed conflicting and context-dependent findings (Starke et al., 2022). One area that remains especially underexplored in the field of AI decision making is how communicating the fairness of an allocation influences these perceptions (see M. K. Lee et al., 2019). Transparent communication about the fairness of a decision-making process may significantly impact how people react to both AI-based and human decisions. For this study, we examine this within the educational context of university courses.

Here, we manipulated the decision-maker (AI-based vs. human) and the fairness communicated (fairness communicated vs. no information) between student groups that were allocated topics of university courses based on their indicated preferences. We chose this allocation setting to systematically compare how AI-based allocations are perceived in terms of fairness, trust, and emotional reactions compared to allocations performed by humans. We focused on the educational setting of university topic allocations because, despite the relatively low task complexity, we assume that human nuance is perceived as important due to both the need to incorporate qualitative information beyond quantitative topic selection and the naturally social classroom setting. Based on these considerations, we formulate the following research question: Do fairness, trust, and emotional responses of students differ between AI-based versus human decision making, and does the communicated decisional fairness (information on decisional fairness vs. no information) affect this?

We posit that the decision-maker (AI-based vs. human), the fairness communication, and the favorability of outcomes are central determinants that affect fairness perceptions, trust, and emotional responses such as satisfaction and outrage in the context of university topic allocation decisions. Regarding the main effects of the decision-maker and fairness communicated, we hypothesize the following:

Hypothesis 1 (a–d): A human decision-maker is associated with increased (a) perceived fairness, (b) trust, and (c) satisfaction and decreased (d) outrage compared to an AI-based decision-maker.

Hypothesis 2 (a–d): Explicitly communicating the fairness of the decision-making process is associated with increased perceived fairness, trust, and satisfaction and decreased outrage compared to a condition in which no information is provided.

As for outcome favorability, we draw on the concept of the outcome (favorability) bias. This concept posits that positive (individual) results can override procedural (and/or technical) evaluations (see Lipshitz, 1989). In line with prior findings in the context of AI-based decision making (Bankins et al., 2022; Wang et al., 2020), we expect that a more favorable outcome leads to more positive perceptions (fairness, trust, emotional responses) of the allocation process compared to a less favorable outcome.

Hypothesis 3 (a–d): A favorable outcome in topic allocation is associated with increased fairness, trust, and satisfaction and decreased outrage compared to an unfavorable outcome.

Finally, we expect interactions between outcome favorability and decision-maker as well as between outcome favorability and communicated fairness. Specifically, we anticipate that fairness communication amplifies differences between decision-makers, whereas high outcome favorability may act to overshadow differences between conditions.

Hypothesis 4 (Decision-Maker $\times$ Communicated Fairness; a–d): When the topic allocation process is communicated as fair, the difference in perceptions (fairness, trust, satisfaction, outrage) between human and AI-based decision making is increased compared to when no information regarding fairness is provided.

Hypothesis 5 (Decision-Maker $\times$ Outcome Favorability; a–d): When the outcome is favorable, the difference in perceptions between human and AI-based decision making is reduced compared to when the outcome is unfavorable.

Hypothesis 6 (Communicated Fairness $\times$ Outcome Favorability; a–d): When the outcome is favorable, the difference in perceptions between communicated versus not communicated fairness is reduced compared to when the outcome is unfavorable.

To foreshadow the results, contrary to our hypotheses derived based on the literature, we found Bayesian evidence against an impact of decision-maker and communicated fairness. Only outcome favorability was the central factor associated with fairness perception, trust, and emotional responses in students allocated to university topics.

Method

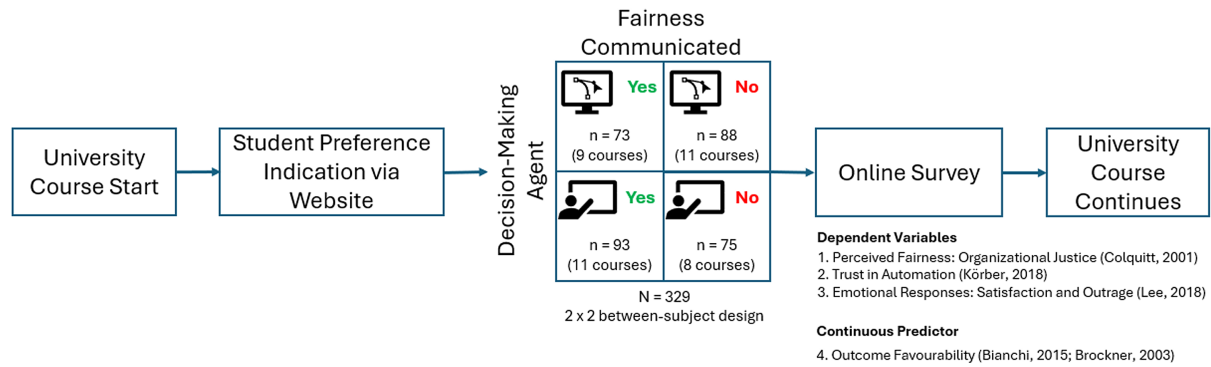
We examined the influence of the decision-maker (human vs. AI) and of the communicated fairness of a topic allocation process (fairness communicated: communicated as fair vs. no information) on student's perceived fairness, trust, and emotional responses (satisfaction, outrage) to topic allocation in their university courses. To do so, we conducted a 2 (decision-maker) $\times$ 2 (fairness communicated) between-subjects field experiment within university courses of the Catholic University of Eichstätt-Ingolstadt (see Figure 1 for the experiment structure).

Participants

As there was no prior study in the context of university course topic allocation to base our sample size estimation on, we opted for a Bayesian sequential testing approach (Schönbrodt et al., 2017). Our Bayesian stopping criterion was a Bayes factor (BF) of at least 3 in favor of or against the null hypothesis for the effect of decision-maker for the main dependent variables (perceived fairness, trust, satisfaction, outrage). Data collection at the Catholic University of Eichstätt-Ingolstadt started in the summer term of 2023, and the Bayesian stopping criterion was reached after four semesters.

A total of 614 bachelor's and master's students in university courses of different subjects (39 courses total; subjects: Psychology, Journalism, Business Studies, and Education) at the Catholic University of Eichstätt-Ingolstadt participated in this study. The data of 285 student participants were excluded from the final analysis due to participation in a prior university course in which the study was

Figure 1
Structure and Time Course of the Experiment



conducted ($N = 186$), not consenting to data collection ($N = 7$), or incomplete responses to the online survey ($N = 67$). Last, one course ($N = 25$) was not considered for analysis, as the type of allocation procedure used in the course significantly deviated from the allocations used in the other courses (timeslot allocation instead of topic allocation). Thus, our final sample size consisted of 329 German bachelor's and master's students (250/76.0% female, 69/21.0% male, 5/1.5% nonbinary/third gender; 5/1.5% preferred not to answer the question; $M_{\text{age}} = 22.8$ years, $SD_{\text{age}} = 3.44$) who participated in a total of 39 different university courses. University courses were randomly allocated to the four between-subjects groups: Nine courses were allocated to an AI decision-maker and fairness was communicated ($N = 73$ participants), 11 courses to an AI decision-maker and no fairness communication ($N = 88$ participants), 11 courses to a human decision-maker and fairness was communicated ($N = 93$ participants), and finally, eight courses were allocated to a human decision-maker without communicating fairness ($N = 75$ participants; see also Table 1).

Participating lecturers (i.e., confederates) were recruited through personal contacts, word of mouth, and advertisements in a university group interested in using AI for teaching. Participants used the topic allocation tool within the scope of their regular university courses to indicate their topic priorities. After topics had been allocated according to the respective between-subjects condition of the university course, they were asked whether they would be willing to

participate in a short survey on the allocation tool (during their university course). Those willing to participate were asked to provide informed consent prior to taking part in the experiment. Participating students were debriefed both in writing at the end of the survey and orally by the confederate lecturers after completing the survey.

Design and Procedure

To assess the impact of the decision-maker (human vs. AI) and the communicated fairness (fairness communicated vs. no information) on the perception of a topic allocation process at university, we compared four between-subjects groups conducted in different university courses. Data collection took place at the beginning of each semester, that is, during one of the first sessions of the respective university courses. The experiment was integrated into the regular university course session in which topics (e.g., presentation or project topics) were allocated. The participating lecturers administered the field experiments by acting as confederates. The course of the experiment was such that confederate lecturers first conducted their course as usual until topics were to be allocated. Then, they mentioned that they would use a new online tool to inquire about students' topic preferences. Subsequently, students reached the online allocation tool via a link and indicated their topic preferences there. The confederate lecturers communicated the allocation decision in

Table 1
Means and Standard Deviations per Dependent Variable and Condition

Variable	Experimental groups			
	AI + fairness communicated: $N_{\text{courses}} = 9, N_{\text{participants}} = 73$	Human + fairness communicated: $N_{\text{courses}} = 11, N_{\text{participants}} = 93$	AI + no information: $N_{\text{courses}} = 11, N_{\text{participants}} = 88$	Human + no information: $N_{\text{courses}} = 8, N_{\text{participants}} = 75$
	$M (SD)$	$M (SD)$	$M (SD)$	$M (SD)$
Overall fairness	3.78 (0.99)	3.79 (1.06)	3.66 (2.04)	3.77 (1.01)
Procedural	2.99 (1.43)	3.26 (1.50)	3.08 (1.36)	3.17 (1.51)
Distributive	3.79 (1.38)	3.78 (1.51)	3.66 (1.49)	3.73 (1.52)
Informational	3.99 (1.26)	3.83 (1.12)	3.75 (1.23)	3.80 (1.21)
Interpersonal	4.38 (1.00)	4.28 (1.15)	4.14 (1.12)	4.40 (0.96)
Trust	3.98 (1.15)	4.02 (1.13)	3.72 (1.18)	4.13 (1.08)
Satisfaction	3.59 (1.30)	3.40 (1.50)	3.50 (1.41)	3.56 (1.54)
Outrage	2.10 (1.23)	2.01 (1.18)	2.13 (1.38)	1.77 (1.10)

Note. All scores range between 1 and 5. Values in parentheses denote the corresponding standard deviations. AI = artificial intelligence.

accordance with the course's decision-maker and communicated fairness condition. Then, students were asked whether they would be willing to participate in a short survey on the allocation tool and were provided with the link to the online survey. Participating students were debriefed at the end of the survey as well as by the confederate lecturers once all students had completed the survey.

For the human decision-maker condition, confederate lecturers informed students that they themselves would allocate the students to their topics on the basis of the priority information collected via the online topic allocation tool. In the AI-based condition, lecturers informed students that the topic allocation tool would automatically allocate students to their topics using an AI based on their indicated priorities, and the lecturer would have no say in the allocation process. Please note that topics were instantly allocated using the same algorithm (based on the stable marriage problem) in all conditions, ensuring equal underlying allocation principles throughout the conditions. Thus, topic allocation did not actually use and rely on an AI-based system. Moreover, participants only entered their preferences and did not have any opportunity to actively interact with the supposed AI system in the AI decision-maker condition. In the human decision-maker condition, lecturers implemented a short break (~2–3 min) after students used the online allocation tool, presumably to give them time for topic allocation. This was done to ensure that students believed that the lecturer had allocated the topics in the human decision-maker condition. To manipulate fairness communicated, students were either informed that the process used by the respective decision-maker was fair (fairness communicated condition) or they were not given any information regarding the fairness of the process (no information condition). The decision-maker used within each course and (if applicable) the fairness of the allocation process were mentioned to the students twice by the confederate lecturers: once at the beginning of the course and once just prior to the actual allocation procedure and the students' use of the online topic allocation tool. Additionally, students in the fairness communicated condition were shown a brief statement at the beginning of the survey, reminding them that the allocation was conducted in a fair manner.

Measures

Fairness Perception

We used Colquitt's (2001) Organizational Justice Scale to assess perceived fairness. We refined the questionnaire by reducing it to 10 items and adapting the wording to fit both AI and human decision-making contexts. The Organizational Justice Scale consists of four subscales on procedural, distributive, interpersonal, and informational justice, which are rated on a scale from *to a very small extent* (1) to *to a very large extent* (5). Following confirmatory factor analysis and the removal of one problematic item (Item 3—procedural justice) to improve factorial clarity, the final four-factor measurement model demonstrated a good (comparative fit index, Tucker–Lewis index, standardized root-mean-square residual) to poor (root-mean-square error of approximation) fit to the data, $\chi^2(21) = 89.44$, $p < .001$, comparative fit index = 0.966, Tucker–Lewis index = 0.942, root-mean-square error of approximation = 0.100, 90% CI [0.079, 0.121], standardized root-mean-square residual = 0.061 (see the Appendix for more details). Result patterns did not differ when including the excluded item. The internal

consistency for the scale was assessed using Cronbach's α , which resulted in a coefficient of $\alpha = .89$.

Trust

We used Körber's (2018) two-item trust subscale in the Trust in Automation questionnaire. We adapted the wording to address both AI and human decision-making contexts. The scale ranged from *strongly disagree* (1) to *strongly agree* (5). The economical two-item scale has a Pearson correlation coefficient of $r = 0.85$, indicating a good intercorrelation between the items.

Outcome Favorability

To determine participants' perception of the favorability of the topic allocation outcome, we asked participants, "How favorable was the topic allocation for you personally?" A similar form of this item was previously used in other studies (e.g., Bianchi et al., 2014; Brockner et al., 2003). The scale ranged from *not at all favorable* (1) to *very favorable* (5).

Emotions

To assess participants' affective responses regarding the topic allocation process, we asked, "To what extent did the topic allocation process make you feel [emotion]?" (see, e.g., Larsson, 2011; M. K. Lee, 2018; Weiss et al., 1999, for prior use of similar scales). We inquired about four emotions: satisfaction, outrage, happiness, and disappointment. The scales ranged from *to a very small extent* (1) to *to a very large extent* (5).

Data and Material Availability

We report how we determined our sample size following a Bayesian approach, and we report all data exclusions as well as all manipulations and dependent measures used in this study. This study's design, research question, and hypotheses were preregistered prior to the start of data collection on the Open Science Framework (Fares et al., 2023; <https://osf.io/t7xnr/overview>). The data, study materials, and analysis scripts are publicly available on the Open Science Framework (Esch, Fares, Pfeuffer, & Kals, 2025; <https://osf.io/qe52u/overview>). Data were analyzed using R (Version 4.4.0; R Core Team, 2024) and the package *BayesFactor* (Version 0.9.12-4.7, Morey & Rouder, 2023).

Results

Given diverging findings regarding AI-based versus human allocation in different contexts (see Starke et al., 2022), we opted for a Bayesian approach to be able to attain both evidence in favor of as well as against the null hypothesis (see Table 1 for descriptive statistics of the dependent variables). Bayesian linear mixed model analyses were conducted per dependent variable. To calculate BFs and estimate the strength of evidence in favor of or against the null and alternative hypotheses, we compared models that did versus did not include an effect of interest (main effect or interaction). Our maximum model per dependent variable included fixed effects of decision-maker and fairness communicated, the continuous predictor outcome favorability (z -standardized), and all corresponding interactions of fixed effects and the continuous predictor. We

preregistered a random effects structure that included both by-course and by-participant random intercepts. But as a Bayesian linear mixed model including both random intercepts had to return a singular fit (overfitting), we only included by-course random intercepts in our final models. Thus, the maximum model per dependent variable was as follows: dependent variable $\sim$ decision-maker $\times$ fairness communicated $\times$ outcome favorability + (11 course). $BF_{10} > 3$ (in favor of) or $BF_{10} < 0.33$ (against) the alternative hypothesis were considered as conclusive (see M. D. Lee & Wagenmakers, 2014).

Fairness Perception

For the dependent variable overall fairness (Organizational Justice Questionnaire), we found conclusive evidence against effects of decision-maker (Hypothesis 1a), $BF_{10} = 0.19$ ($\pm 10.01\%$), and fairness communicated (Hypothesis 2a), $BF_{10} = 0.19$ ($\pm 0.91\%$). We further observed conclusive evidence for an effect of outcome favorability (Hypothesis 3a), $BF_{10} = 4.79e+42$ ($\pm 1.24\%$; see Figure 1 for a depiction of the influences of outcome favorability, decision-maker, and fairness communicated for all dependent variables). Perceived fairness increased with higher outcome favorability ratings. We found conclusive evidence against all interaction effects—Decision-Maker $\times$ Fairness Communicated (Hypothesis 4a), $BF_{10} = 0.25$ ($\pm 2.49\%$); Decision-Maker $\times$ Outcome Favorability (Hypothesis 5a), $BF_{10} = 0.09$ ($\pm 3.48\%$); Fairness Communicated $\times$ Outcome Favorability (Hypothesis 6a), $BF_{10} = 0.08$ ($\pm 3.16\%$); and the three-way interaction (Decision-Maker $\times$ Fairness Communicated $\times$ Outcome Favorability), $BF_{10} = 0.0006$ ($\pm 2.56\%$). To ascertain our results, we additionally assessed the four fairness dimensions separately (see Table 2). Our findings confirmed the same pattern found for overall fairness for all fairness dimensions (Figure 2).

Trust

For the dependent variable trust, we found conclusive evidence against effects of decision-maker (Hypothesis 1b), $BF_{10} = 0.27$ ($\pm 3.43\%$), and fairness communicated (Hypothesis 2b), $BF_{10} = 0.19$ ($\pm 0.68\%$). We observed conclusive evidence for an effect of outcome favorability (Hypothesis 3b), $BF_{10} = 9.42e+18$ ($\pm 2.28\%$). Perceived trust increased with higher outcome favorability ratings. We found conclusive evidence against the interaction effects—Decision-Maker $\times$ Outcome Favorability (Hypothesis 5b), $BF_{10} = 0.12$ ($\pm 1.62\%$); Fairness Communicated $\times$ Outcome Favorability (Hypothesis 6b), $BF_{10} = 0.20$ ($\pm 2.45\%$); and the three-way interaction (Decision-Maker $\times$ Fairness Communicated $\times$ Outcome Favorability), $BF_{10} = 0.006$ ($\pm 1.97\%$). The interaction effect between decision-maker and fairness communicated was inconclusive (Hypothesis 4b), $BF_{10} = 0.49$ ($\pm 2.08\%$).

Satisfaction

For the dependent variable satisfaction, we found conclusive evidence against effects of decision-maker (Hypothesis 1c), $BF_{10} = 0.18$ ($\pm 0.50\%$), and fairness communicated (Hypothesis 2c), $BF_{10} = 0.21$ ($\pm 0.63\%$). We observed conclusive evidence for an effect of outcome favorability (Hypothesis 3c), $BF_{10} = 2.04e+39$ ($\pm 0.62\%$). Perceived satisfaction increased with higher outcome favorability ratings. We found conclusive evidence against the interaction effects—Decision-

Maker $\times$ Outcome Favorability (Hypothesis 5c), $BF_{10} = 0.16$ ($\pm 1.35\%$); Fairness Communicated $\times$ Outcome Favorability (Hypothesis 6c), $BF_{10} = 0.25$ ($\pm 3.23\%$); and the three-way interaction (Decision-Maker $\times$ Fairness Communicated $\times$ Outcome Favorability), $BF_{10} = 0.003$ ($\pm 3.88\%$). The interaction effect between decision-maker and fairness communicated was inconclusive (Hypothesis 4c), $BF_{10} = 0.34$ ($\pm 1.72\%$).

Outrage

For the dependent variable outrage, we found conclusive evidence against effects of decision-maker (Hypothesis 1d), $BF_{10} = 0.25$ ($\pm 2.63\%$), and fairness communicated (Hypothesis 2d), $BF_{10} = 0.23$ ($\pm 0.70\%$). We observed conclusive evidence for an effect of outcome favorability (Hypothesis 3d), $BF_{10} = 6.25e+11$ ($\pm 0.47\%$). Perceived outrage decreased with higher outcome favorability ratings. We found conclusive evidence against the interaction effects—Decision-Maker $\times$ Outcome Favorability (Hypothesis 5d), $BF_{10} = 0.13$ ($\pm 4.38\%$); Fairness Communicated $\times$ Outcome Favorability (Hypothesis 6d), $BF_{10} = 0.13$ ($\pm 1.24\%$); and the three-way interaction (Decision-Maker $\times$ Fairness Communicated $\times$ Outcome Favorability), $BF_{10} = 0.002$ ($\pm 4.91\%$). The interaction effect between decision-maker and fairness communicated was inconclusive (Hypothesis 4d), $BF_{10} = 0.41$ ($\pm 4.73\%$).

Discussion

Our aim was to assess the perception of an AI-based versus human allocation process in an educational setting, specifically, course topic allocation at university. In a 2×2 between-subjects design, we investigated how the decision-maker (AI vs. human), the communicated fairness of the allocation process, and the outcome favorability affected students' perceived fairness, trust, and emotional responses.

Our results are clear-cut: Contrary to our expectations and the findings of some prior studies in different contexts (e.g., Bai et al., 2020; Marcinkowski et al., 2020; but see Plane et al., 2017; Suen et al., 2019), Bayesian evidence indicated that neither the decision-maker nor the communicated fairness (each considering individual fixed effects as well as interactions) had an impact on how fair and trustworthy students perceived a topic allocation process or on their corresponding emotional responses. That is, there was evidence against the notion that students differentially evaluated AI and human decision-makers for university course topic allocations. However, consistent across all assessed variables, the factor most strongly associated with students' evaluations of the topic allocation process was the favorability of the outcome (see also Wang et al., 2020). Importantly, outcome favorability should not be interpreted as a strictly objective measure of outcome quality but rather as individuals' postallocation evaluation. This suggests that, at least in the context of university topic allocations, the subjective experience of the result (i.e., outcome favorability) overshadows procedural aspects of the topic allocation process.

These findings contribute to the growing body of research on the comparative effects of AI-based versus human decision making (see Starke et al., 2022). While previous studies have shown mixed results, with some favoring human decision making, particularly in complex or sensitive domains (e.g., Acikgoz et al., 2020; M. K. Lee & Rich, 2021), the present study suggests that AI-based decision making can be equally acceptable in certain contexts, such as course topic allocation, where the task type may be viewed as less complex

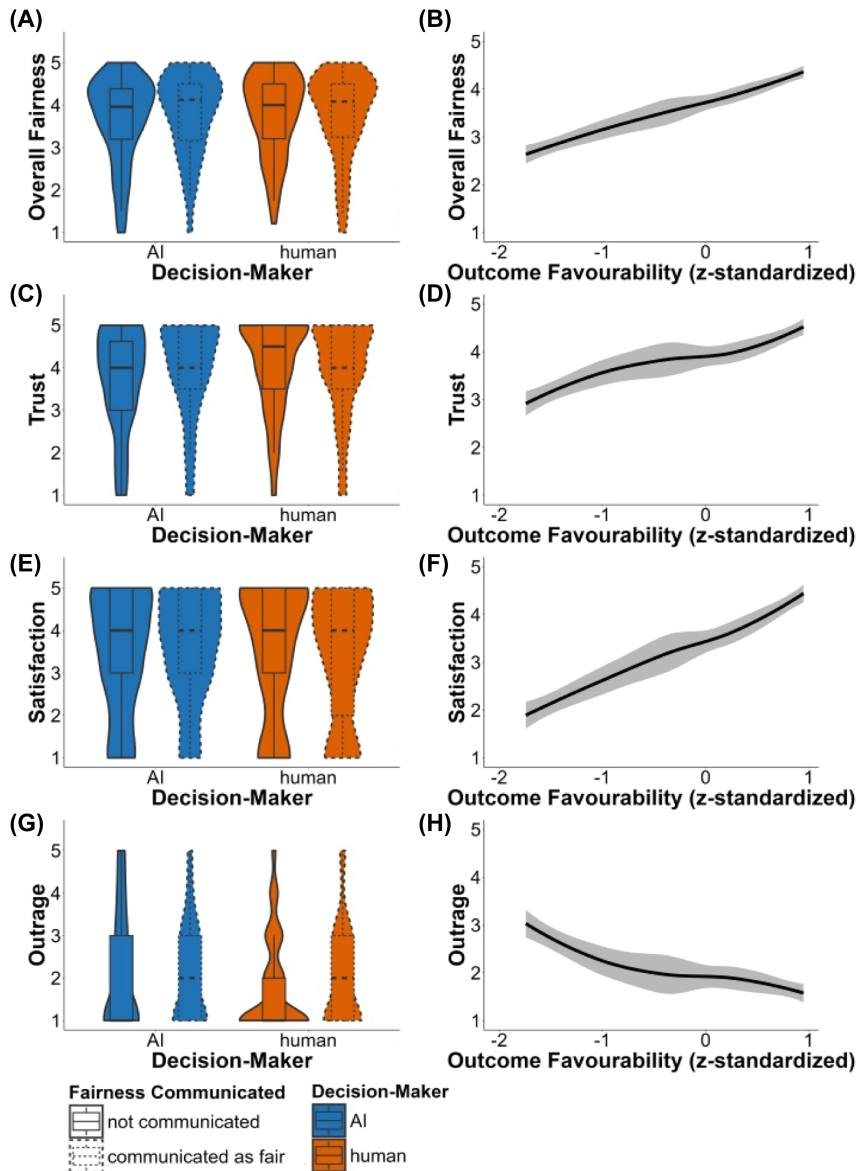
Table 2*Bayes Factors in Favor of the Alternative Hypothesis per Dependent Variable*

Effect	Dependent variable	BF ₁₀ (error rate)	Interpretation (++ to --)
H1: Decision-maker	Overall fairness	0.19 (±10.01%)	--
	Procedural	0.19 (±0.77%)	--
	Distributional	0.17 (±1.64%)	--
	Informational	0.18 (±1.43%)	--
	Interpersonal	0.20 (±0.42%)	--
	Trust	0.27 (±3.43%)	--
	Satisfaction	0.18 (±0.50%)	--
	Outrage	0.25 (±2.63%)	--
H2: Fairness communicated	Overall fairness	0.19 (±0.91%)	--
	Procedural	0.18 (±2.71%)	--
	Distributional	0.18 (±0.46%)	--
	Informational	0.19 (±0.75%)	--
	Interpersonal	0.20 (±0.82%)	--
	Trust	0.19 (±0.68%)	--
	Satisfaction	0.21 (±0.63%)	--
	Outrage	0.23 (±0.70%)	--
H3: Outcome favorability	Overall fairness	4.79e+42 (±1.24%)	++
	Procedural	5.49e+23 (±1.37%)	++
	Distributional	7.65e+51 (±0.57%)	++
	Informational	1.32e+09 (±1.02%)	++
	Interpersonal	2.64e+14 (±1.18%)	++
	Trust	9.42e+18 (±2.28%)	++
	Satisfaction	2.04e+39 (±0.62%)	++
	Outrage	6.25e+11 (±0.47%)	++
H4: Decision-Maker × Fairness Communicated	Overall fairness	0.25 (±2.49%)	--
	Procedural	0.24 (±2.91%)	--
	Distributional	0.25 (±2.91%)	--
	Informational	0.30 (±1.81%)	--
	Interpersonal	0.49 (±8.00%)	-
	Trust	0.49 (±2.08%)	-
	Satisfaction	0.34 (±1.72%)	-
	Outrage	0.41 (±4.73%)	-
H5: Decision-Maker × Outcome Favorability	Overall fairness	0.09 (±3.48%)	--
	Procedural	0.23 (±2.00%)	--
	Distributional	0.09 (±1.28%)	--
	Informational	0.15 (±2.29%)	--
	Interpersonal	0.12 (±1.42%)	--
	Trust	0.12 (±1.62%)	--
	Satisfaction	0.16 (±1.35%)	--
	Outrage	0.13 (±4.38%)	--
H6: Fairness Communicated × Outcome Favorability	Overall fairness	0.08 (±3.16%)	--
	Procedural	0.10 (±2.08%)	--
	Distributional	0.11 (±28.77%)	--
	Informational	0.14 (±1.49%)	--
	Interpersonal	0.12 (±1.03%)	--
	Trust	0.20 (±2.45%)	--
	Satisfaction	0.25 (±3.23%)	--
	Outrage	0.13 (±1.24%)	--
Decision-Maker × Fairness Communicated × Outcome Favorability	Overall fairness	0.0006 (±2.56%)	--
	Procedural	0.003 (±7.66%)	--
	Distributional	0.0002 (±3.37%)	--
	Informational	0.002 (±12.97%)	--
	Interpersonal	0.006 (±11.92%)	--
	Trust	0.006 (±1.97%)	--
	Satisfaction	0.003 (±3.88%)	--
	Outrage	0.002 (±4.91%)	--

Note. Results in bold indicate conclusive results; “++” and “--” refer to conclusive evidence respectively for or against an effect; “+” and “-” refer to inconclusive tendencies toward or against an effect. H = hypothesis.

Figure 2

Interactions Between Decision-Maker and Fairness Communicated (Left Column) and Influence of Outcome Favorability (Right Column)



Note. Panels A, C, E, and G illustrate the effects of decision-maker and fairness communicated per dependent variable. The horizontal lines in the violin plots denote the respective mean. Panels B, D, F, and H illustrate the influence of outcome favorability on the respective dependent variable. The gray bands show the 95% confidence interval. AI = artificial intelligence.

and more mechanical (i.e., the decision is perceived as being more strongly based on data and not emotional judgment; see M. K. Lee, 2018). Consequently, future research should systematically investigate contexts that vary in perceived task type (e.g., more mechanical vs. human; see M. K. Lee, 2018) and perceived relevance to clarify how these factors shape acceptance, trust, and fairness perceptions of AI-based versus human decision making.

The present decision-maker manipulation was strictly label-based and did not involve direct interaction with or observable behavior of

a functioning AI-based system beyond observing one allocation outcome, which was deliberately equated between conditions. Therefore, our findings reflect the impact of perceived decision-maker identity unconfounded by actual AI functionality. This is advantageous to clearly and unambiguously answer this study's research question regarding the role of the decision-maker. However, it simultaneously limits the generalizability of the present findings to contexts that involve more complex human-AI interaction. Furthermore, it has to be noted that, for practical constraints

given by the number of students and the distribution of study programs at the Catholic University of Eichstätt-Ingolstadt where this field study took place in the scope of actual university courses, participants mainly came from social sciences and humanities study programs. Future studies should address this limitation by further broadening the scope of academic disciplines covered and increasing the diversity of participating students. Moreover, please also note that the present study did not account for individual differences, for instance, in students' preferences for human versus AI decision making, their prior experiences with AI-based decision making (see algorithm aversion; Dietvorst et al., 2015), or students' AI literacy (Araujo et al., 2020). These appear to be highly relevant additional aspects to consider in future studies. Similarly, future studies could use vignette settings to manipulate outcome favorability also in the low range and assess whether the decision-maker might play a role when outcomes are very unfavorable (a contrasting interaction as found by Yalcin et al., 2022).

Conclusion and Practical Implications

Our findings clearly illustrate that the subjective evaluation of the allocation outcome, rather than the decision-maker or fairness communications regarding the allocation process, affected students' perception and evaluation of the digitally mediated university topic allocation processes. Given the increasing use of AI systems in educational settings (e.g., Marcinkowski et al., 2020), these findings have important practical implications: Our study suggests that, at least for tasks comparable to topic allocation, universities may consider adopting efficient and optimized AI-based decision-making systems, without having to fear negative student perceptions. Here, we demonstrated that even when decision quality was equated between human and AI decision-makers (same underlying allocation principles) and therefore could not confound students' perceptions, there was evidence against differences in students' perceptions and evaluations of the decision process between human and AI decision-makers. Indeed, AI systems often excel at prioritizing and optimizing resources beyond a human decision-maker's capabilities (when aimed at mathematically optimal solutions unlike in this study). This may further boost satisfaction in realistic settings due to, on average, more optimal AI-based allocations when implemented transparently. Yet, the actual usage and thus success of AI decision making heavily relies on its human users' perceptions, in particular their acceptance, of the AI-based decision process (e.g., Wanner et al., 2022). This makes positive user perceptions of AI systems essential irrespective of the corresponding AI systems' mathematical optimality. Nevertheless, the present findings should not be misinterpreted as a normative endorsement of AI systems in general or in educational settings. When considering whether to delegate a task to an AI, the respective AI's fit for the task and potential concerns (e.g., data security and data privacy concerns; see, e.g., criteria for trustworthy AI of the European Union; see High-Level Expert Group on Artificial Intelligence, 2019) always need to be taken into consideration. In this regard, as AI becomes more prevalent in higher education, institutions also need to closely monitor and evaluate the sociocultural and ethical implications of this change.

Despite our data showing no effect of communicating procedural fairness, prior research consistently highlights the importance of transparent process communication for fostering fairness, trust,

positive emotional responses, and ultimately acceptance and adoption. Further research will be necessary to determine under which conditions each level of transparency is required and how it should best be communicated. Hence, we recommend that universities (and all organizations implementing AI decision making) nevertheless aim to transparently convey procedural aspects of allocations to maintain positive student perceptions, which remain vital for the successful implementation of AI in education and beyond.

References

- Acikgoz, Y., Davison, K. H., Compagnone, M., & Laske, M. (2020). Justice perceptions of artificial intelligence in selection. *International Journal of Selection and Assessment*, 28(4), 399–416. <https://doi.org/10.1111/ijsa.12306>
- Ali, S., Abuhmed, T., El-Sappagh, S., Muhammad, K., Alonso-Moral, J. M., Confalonieri, R., Guidotti, R., Del Ser, J., Díaz-Rodríguez, N., & Herrera, F. (2023). Explainable Artificial Intelligence (XAI): What we know and what is left to attain trustworthy artificial intelligence. *Information Fusion*, 99, Article 101805. <https://doi.org/10.1016/j.inffus.2023.101805>
- Araujo, T., Helberger, N., Kruijemeier, S., & de Vreese, C. H. (2020). In AI we trust? Perceptions about automated decision-making by artificial intelligence. *AI & Society*, 35(3), 611–623. <https://doi.org/10.1007/s00146-019-00931-w>
- Ayanwale, M. A., & Ndlovu, M. (2024). Investigating factors of students' behavioral intentions to adopt chatbot technologies in higher education: Perspective from expanded diffusion theory of innovation. *Computers in Human Behavior Reports*, 14, Article 100396. <https://doi.org/10.1016/j.chbr.2024.100396>
- Bai, B., Dai, H., Zhang, D., Zhang, F., & Hu, H. (2020). The impacts of algorithmic work assignment on fairness perceptions and productivity: Evidence from field experiments. *SSRN Electronic Journal*. Advance online publication. <https://doi.org/10.2139/ssrn.3550887>
- Bankins, S., Formosa, P., Griep, Y., & Richards, D. (2022). AI decision making with dignity? Contrasting workers' justice perceptions of human and AI decision making in a human resource management context. *Information Systems Frontiers*, 24(3), 857–875. <https://doi.org/10.1007/s10796-021-10223-8>
- Bianchi, E. C., Brockner, J., van den Bos, K., Seifert, M., Moon, H., van Dijke, M., & De Cremer, D. (2014). Trust in decision-making authorities dictates the form of the interactive relationship between outcome fairness and procedural fairness. *Personality and Social Psychology Bulletin*, 41(1), 19–34. <https://doi.org/10.1177/0146167214556237>
- Brockner, J., Heuer, L., Magner, N., Folger, R., Umphress, E., van den Bos, K., Vermunt, R., Magner, M., & Siegel, P. (2003). High procedural fairness heightens the effect of outcome favorability on self-evaluations: An attributional analysis. *Organizational Behavior and Human Decision Processes*, 91(1), 51–68. [https://doi.org/10.1016/S0749-5978\(02\)00531-9](https://doi.org/10.1016/S0749-5978(02)00531-9)
- Brockner, J., & Wiesenfeld, B. M. (1996). An integrative framework for explaining reactions to decisions: Interactive effects of outcomes and procedures. *Psychological Bulletin*, 120(2), 189–208. <https://doi.org/10.1037/0033-2909.120.2.189>
- Chouldechova, A. (2016). *Fair prediction with disparate impact: A study of bias in recidivism prediction instruments* (Version 1). arXiv. <https://doi.org/10.48550/ARXIV.1610.07524>
- Colquitt, J. A. (2001). On the dimensionality of organizational justice: A construct validation of a measure. *Journal of Applied Psychology*, 86(3), 386–400. <https://doi.org/10.1037/0021-9010.86.3.386>
- Colquitt, J. A., Greenberg, J., & Zapata-Phelan, C. P. (2005). What is organizational justice? A historical overview. In J. Greenberg & J. A. Colquitt (Eds.), *Handbook of organizational justice* (pp. 3–56). Lawrence Erlbaum Associates. <https://doi.org/10.4324/9780203774847>
- Colquitt, J. A., Scott, B. A., Rodell, J. B., Long, D. M., Zapata, C. P., Conlon, D. E., & Wesson, M. J. (2013). Justice at the millennium, a decade later: A

- meta-analytic test of social exchange and affect-based perspectives. *Journal of Applied Psychology*, 98(2), 199–236. <https://doi.org/10.1037/a0031757>
- 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>
- Doshi-Velez, F., & Kim, B. (2017). *Towards a rigorous science of interpretable machine learning*. arXiv. <https://doi.org/10.48550/ARXIV.1702.08608>
- Esch, C., Fares, F., Kals, E., & Pfeuffer, C. U. (2025, July 21). *AI-based decision-making: Not the decision-maker but the outcome's favourability determines the perception of university topic allocations*. https://doi.org/10.31219/osf.io/cwxgt_v1
- Esch, C., Fares, F., Pfeuffer, C. U., & Kals, E. (2025, September 16). *Algorithmic vs. human decision-making: How does the method of decision-making and the level of communicated fairness influence justice perception?—A higher education perspective*. <https://osf.io/qe52u>
- Fares, F., Esch, C., Kals, E., & Pfeuffer, C. U. (2023, March 9). *Algorithmic vs. human decision-making: How does the method of decision-making and the level of communicated fairness influence justice perception?—A higher education perspective*. <https://doi.org/10.17605/OSF.IO/T7XNR>
- Folger, R., Rosenfield, D., Grove, J., & Corkran, L. (1979). Effects of “voice” and peer opinions on responses to inequity. *Journal of Personality and Social Psychology*, 37(12), 2253–2261. <https://doi.org/10.1037/0022-3514.37.12.2253>
- Greenberg, J. (1987). A taxonomy of organizational justice theories. *The Academy of Management Review*, 12(1), 9–22. <https://doi.org/10.2307/257990>
- Grgić-Hlača, N., Zafar, M. B., Gummad, K. P., & Weller, A. (2018). Beyond distributive fairness in algorithmic decision making: Feature selection for procedurally fair learning. *Proceedings of the AAAI Conference on Artificial Intelligence*, 32(1), 51–60. <https://doi.org/10.1609/aaai.v32i1.11296>
- Grunenberg, M., Prantl, J., Heidt, K., & Kals, E. (2024). Vertrauen in der Zusammenarbeit von Organisationen: Die Bedeutung organisationaler Gerechtigkeit und affektiven Erlebens [Trust in interorganizational collaboration: The role of organizational justice and affect]. *Gruppe. Interaktion. Organisation. Zeitschrift für Angewandte Organisationspsychologie*, 55(1), 33–46. <https://doi.org/10.1007/s11612-024-00728-6>
- Harrison, G., Hanson, J., Jacinto, C., Ramirez, J., & Ur, B. (2020). An empirical study on the perceived fairness of realistic, imperfect machine learning models. In M. Hildebrandt, C. Castillo, E. Celis, S. Ruggieri, L. Taylor, & G. Zanfir-Fortuna (Eds.), *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency* (pp. 392–402). Association for Computing Machinery. <https://doi.org/10.1145/3351095.3372831>
- High-Level Expert Group on Artificial Intelligence. (2019). *Ethics guidelines for trustworthy AI*. European Commission. <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>
- Holmvall, C. M., & Bobocel, D. R. (2008). What fair procedures say about me: Self-construals and reactions to procedural fairness. *Organizational Behavior and Human Decision Processes*, 105(2), 147–168. <https://doi.org/10.1016/j.obhdp.2007.09.001>
- Hsu, S., Li, T. W., Zhang, Z., Fowler, M., Zilles, C., & Karahalios, K. (2021). Attitudes surrounding an imperfect AI autograder. *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems* (pp. 1–15). Association for Computing Machinery. <https://doi.org/10.1145/3411764.3445424>
- Hu, K. (2023). *ChatGPT sets record for fastest-growing user base—Analyst note*. Reuters. <https://www.reuters.com/technology/chatgpt-sets-record-fastest-growing-user-base-analyst-note-2023-02-01/>
- Hu, L., & Bentler, P. M. (1999). Cutoff criteria for fit indexes in covariance structure analysis: Conventional criteria versus new alternatives. *Structural Equation Modeling: A Multidisciplinary Journal*, 6(1), 1–55. <https://doi.org/10.1080/10705519909540118>
- Jones-Jang, S. M., Chung, M., Choi, J., Kim, N., & Lee, S. (2025). Fairness perceptions of AI in grading systems: Examining how discontent with the status quo and outcome favorability reduce AI reluctance. *Computers and Education: Artificial Intelligence*, 8, Article 100419. <https://doi.org/10.1016/j.caeai.2025.100419>
- Kaibel, C., Koch-Bayram, I., Biemann, T., & Mühlenbock, M. (2019). Applicant perceptions of hiring algorithms—Uniqueness and discrimination experiences as moderators. *Academy of Management Proceedings*, 2019(1), Article 18172. <https://doi.org/10.5465/ambpp.2019.210>
- Kelly, S., Kaye, S.-A., & Oviedo-Trespalacios, O. (2023). What factors contribute to the acceptance of artificial intelligence? A systematic review. *Telematics and Informatics*, 77, Article 101925. <https://doi.org/10.1016/j.tele.2022.101925>
- Kizilcec, R. F. (2016). How much information? Effects of transparency on trust in an algorithmic interface. *Proceedings of the 2016 CHI Conference on Human Factors in Computing Systems* (pp. 2390–2395). Association for Computing Machinery. <https://doi.org/10.1145/2858036.2858402>
- Kohn, S. C., de Visser, E. J., Wiese, E., Lee, Y.-C., & Shaw, T. H. (2021). Measurement of trust in automation: A narrative review and reference guide. *Frontiers in Psychology*, 12, Article 604977. <https://doi.org/10.3389/fpsyg.2021.604977>
- Körber, M. (2018). Theoretical considerations and development of a questionnaire to measure trust in automation. In S. Bagnara, R. Tartaglia, S. Albolino, T. Alexander, & Y. Fujita (Eds.), *Proceedings of the 20th Congress of the International Ergonomics Association (IEA 2018)* (pp. 13–30). Advances in Intelligent Systems and Computing. https://doi.org/10.1007/978-3-319-96074-6_2
- Krehbiel, P. J., & Cropanzano, R. (2000). Procedural justice, outcome favorability and emotion. *Social Justice Research*, 13(4), 339–360. <https://doi.org/10.1023/A:1007670909889>
- Larsson, G. (2011, November 9–10). *The Emotional Stress Reaction Questionnaire (ESRQ): Measurement of stress reaction level in field conditions in 60 seconds* [Conference session]. 2nd Conference of the International Society of Military Science, Stockholm, Sweden. <https://urn.kb.se/resolve?urn=urn:nbn:se:fhs:diva-2595>
- Lee, M. D., & Wagenmakers, E.-J. (2014). *Bayesian cognitive modeling: A practical course*. Cambridge University Press. <https://doi.org/10.1017/CBO9781139087759>
- Lee, M. K. (2018). Understanding perception of algorithmic decisions: Fairness, trust, and emotion in response to algorithmic management. *Big Data & Society*, 5(1), 1–16. <https://doi.org/10.1177/2053951718756684>
- Lee, M. K., & Baykal, S. (2017). Algorithmic mediation in group decisions. In C. P. Lee (Ed.), *Proceedings of the 2017 ACM Conference on Computer Supported Cooperative Work and Social Computing* (pp. 1035–1048). Association for Computing Machinery. <https://doi.org/10.1145/2998181.2998230>
- Lee, M. K., Jain, A., Cha, H. J., Ojha, S., & Kusbit, D. (2019). Procedural justice in algorithmic fairness. *Proceedings of the ACM on Human-Computer Interaction*, 3, 1–26. <https://doi.org/10.1145/3359284>
- Lee, M. K., & Rich, K. (2021). Who is included in human perceptions of AI? Trust and perceived fairness around healthcare AI and cultural mistrust. In Y. Kitamura, A. Quigley, K. Isbister, T. Igarashi, P. Bjørn, & S. Drucker (Eds.), *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems* (pp. 1–14). Association for Computing Machinery. <https://doi.org/10.1145/3411764.3445570>
- Lipshitz, R. (1989). “Either a medal or a corporal”: The effects of success and failure on the evaluation of decision making and decision makers. *Organizational Behavior and Human Decision Processes*, 44(3), 380–395. [https://doi.org/10.1016/0749-5978\(89\)90015-0](https://doi.org/10.1016/0749-5978(89)90015-0)
- Low, M. P., Soliman, M., Al Balushi, M. K., & Leong, C.-M. (2025). Beyond conventional adoption: New insights into AI facilitators in higher education institutions. *Sustainable Futures*, 9, Article 100781. <https://doi.org/10.1016/j.sfr.2025.100781>

- 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>
- Marcinkowski, F., Kieslich, K., Starke, C., & Lünich, M. (2020). Implications of AI (un-)fairness in higher education admissions. In M. Hildebrandt, C. Castillo, E. Celis, S. Ruggieri, L. Taylor, & G. Zanfir-Fortuna (Eds.), *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency* (pp. 122–130). Association for Computing Machinery. <https://doi.org/10.1145/3351095.3372867>
- Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of Management Review*, 20, 709–734. <https://doi.org/10.5465/amr.1995.9508080335>
- Miller, S. M., & Keiser, L. R. (2021). Representative bureaucracy and attitudes toward automated decision making. *Journal of Public Administration Research and Theory*, 31(1), 150–165. <https://doi.org/10.1093/jopart/muaa019>
- Morey, R. D., & Rouder, J. N. (2023). *Bayes factor: Computation of Bayes factors for common designs* (R package Version 0.9.12-4.6). <https://github.com/richarddmorey/bayesfactor>
- Nagtegaal, R. (2021). The impact of using algorithms for managerial decisions on public employees' procedural justice. *Government Information Quarterly*, 38(1), Article 101536. <https://doi.org/10.1016/j.giq.2020.101536>
- Narayanan, D., Nagpal, M., McGuire, J., Schweitzer, S., & De Cremer, D. (2024). Fairness perceptions of artificial intelligence: A review and path forward. *International Journal of Human-Computer Interaction*, 40(1), 4–23. <https://doi.org/10.1080/10447318.2023.2210890>
- Ng, D. T. K., Leung, J. K. L., Chu, S. K. W., & Qiao, M. S. (2021). Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence*, 2, Article 100041. <https://doi.org/10.1016/j.caeai.2021.100041>
- Plane, A. C., Redmiles, E. M., Mazurek, M. L., & Tschantz, M. C. (2017). Exploring user perceptions of discrimination in online targeted advertising. In E. Kirda & T. Ristenpart (Chairs), *Proceedings of the USENIX Security Symposium* [Symposium]. Vancouver, Canada.
- R Core Team. (2024). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing.
- Saxena, N. A., Huang, K., DeFilippis, E., Radanovic, G., Parkes, D. C., & Liu, Y. (2019). How do fairness definitions fare? *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society* (pp. 99–106). Association for Computing Machinery. <https://doi.org/10.1145/3306618.3314248>
- Schlicker, N., Langer, M., Ötting, S. K., Baum, K., König, C. J., & Wallach, D. (2021). What to expect from opening up 'black boxes'? Comparing perceptions of justice between human and automated agents. *Computers in Human Behavior*, 122, Article 106837. <https://doi.org/10.1016/j.chb.2021.106837>
- Schönbrodt, F. D., Wagenmakers, E.-J., Zehetleitner, M., & Perugini, M. (2017). Sequential hypothesis testing with Bayes factors: Efficiently testing mean differences. *Psychological Methods*, 22(2), 322–339. <https://doi.org/10.1037/met0000061>
- Shin, D. (2021). The effects of explainability and causability on perception, trust, and acceptance: Implications for explainable AI. *International Journal of Human-Computer Studies*, 146, Article 102551. <https://doi.org/10.1016/j.ijhcs.2020.102551>
- Solberg, E., Kaarstad, M., Eitheim, M. H. R., Bisio, R., Reegård, K., & Bloch, M. (2022). A Conceptual model of trust, perceived risk, and reliance on AI decision aids. *Group & Organization Management*, 47(2), 187–222. <https://doi.org/10.1177/10596011221081238>
- Starke, C., Baleis, J., Keller, B., & Marcinkowski, F. (2022). Fairness perceptions of algorithmic decision-making: A systematic review of the empirical literature. *Big Data & Society*, 9(2), 1–21. <https://doi.org/10.1177/20539517221115189>
- Suen, H.-Y., Chen, M. Y.-C., & Lu, S.-H. (2019). Does the use of synchrony and artificial intelligence in video interviews affect interview ratings and applicant attitudes? *Computers in Human Behavior*, 98, 93–101. <https://doi.org/10.1016/j.chb.2019.04.012>
- Wang, R., Harper, F. M., & Zhu, H. (2020). Factors influencing perceived fairness in algorithmic decision-making: Algorithm outcomes, development procedures, and individual differences. *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems (CHI '20)* (pp. 1–15). <https://doi.org/10.1145/3313831.3376813>
- Wanner, J., Herm, L.-V., Heinrich, K., Janiesch, C., & Zschech, P. (2022). The effect of transparency and trust on intelligent system acceptance: Evidence from a user-based study. *Electronic Markets*, 32(4), 2079–2102. <https://doi.org/10.1007/s12525-022-00593-5>
- Weiss, H. M., Suckow, K., & Cropanzano, R. (1999). Effects of justice conditions on discrete emotions. *Journal of Applied Psychology*, 84(5), 786–794. <https://doi.org/10.1037/0021-9010.84.5.786>
- Yalcin, G., Lim, S., Puntoni, S., & van Osselaer, S. M. J. (2022). Thumbs up or down: Consumer reactions to decisions by algorithms versus humans. *Journal of Marketing Research*, 59(4), 696–717. <https://doi.org/10.1177/00222437211070016>

Appendix

Factor Analysis and Further Results

Confirmatory Factor Analysis Results

This confirmatory factor analysis aimed to evaluate the hypothesized four-factor structure of the adapted Organizational Justice Scale (Colquitt, 2001), covering procedural, distributive, informational, and interactional fairness. The analysis was conducted using the *lavaan* package in R, based on the data from all 329 participants.

Overall, model fit was mixed (see Table A1). The comparative fit index (0.966) and Tucker–Lewis index (0.942) indicated good to very good fit, and the standardized root-mean-square residual (0.061) was well below the recommended cutoff. However, the root-mean-square error of approximation (0.100) was above common thresholds, suggesting a marginal-to-poor fit. Taken together, these indicators suggest the model provides a reasonable approximation of the data structure, though with some room for improvement.

To improve model clarity, one item from the procedural fairness subscale (Item 3) was removed due to poor fit. All remaining factor loadings were statistically significant and of acceptable to high magnitude, indicating that the items represent their intended constructs well (see Table A2).

In sum, the confirmatory factor analysis supports the use of the adapted four-factor structure for the assessment of fairness perceptions across AI and human decision-making contexts. While the root-mean-square error of approximation points to limitations in fit, the overall pattern of results is consistent with theoretical expectations.

Further Results for Emotions

The additional analyses of happiness and disappointment revealed patterns consistent with the emotion variables of satisfaction

Table A1

Summary of Fit Indices for the Revised Four-Factor Model of Organizational Justice

Index	Value	Recommended cutoff ^a
χ^2	89.441	
<i>df</i>	21	
$p(\chi^2)$	<0.001	>.05
CFI	0.966	≥.95 (or ≥.90)
TLI	0.942	≥.95 (or ≥.90)
RMSEA	0.100	≤.06 (or ≤.08)
90% CI [RMSEA]	[0.079, 0.121]	Lower ≤ .05, upper < .08
SRMR	0.061	≤.08

Note. $N = 329$. χ^2 = chi-square goodness-of-fit statistic; CFI = comparative fit index; TLI = Tucker–Lewis index; RMSEA = root-mean-square error of approximation; CI = confidence interval; SRMR = standardized root-mean-square residual.

^aCutoffs are based on common recommendations (e.g., see L. Hu & Bentler, 1999); the RMSEA indicates marginal-to-poor fit, while CFI and SRMR suggest good fit.

and outrage, respectively. More specifically, outcome favorability emerged as the central factor positively associated with happiness and disappointment. In contrast, there was evidence against any main effects of decision-maker and fairness communicated and no evidence for any two- or three-way interactions.

Happiness

For the dependent variable happiness, we found conclusive evidence against effects of decision-maker, $BF_{10} = 0.18 (\pm 0.84\%)$, and fairness communicated, $BF_{10} = 0.25 (\pm 7.09\%)$. We observed conclusive evidence for an effect of outcome favorability, $BF_{10} = 5.58e+37 (\pm 0.69\%)$. Perceived happiness increased with higher outcome favorability ratings. We found conclusive evidence against the interaction effects—Decision-Maker $\times$ Outcome Favorability, $BF_{10} = 0.13 (\pm 2.32\%)$; Decision-Maker $\times$ Fairness Communicated, $BF_{10} = 0.26 (\pm 2.34\%)$; and the three-way interaction (Decision-Maker $\times$ Fairness Communicated $\times$ Outcome Favorability), $BF_{10} = 0.013 (\pm 2.47\%)$. The interaction effect between fairness communicated and outcome favorability was inconclusive, $BF_{10} = 0.75 (\pm 1.33\%)$.

Table A2

Parameter Estimates for the Revised Four-Factor Model: Factor Loadings

Latent factor	Indicator	<i>B</i>	<i>SE</i>	<i>z</i>	<i>p</i>	β (Std.all)
Procedural fairness	Item 1	1.000 ^a				0.476
	Item 2	1.682	0.220	7.646	<.001	0.751
Distributive fairness	Item 1	1.000 ^a				0.938
	Item 2	1.003	0.037	27.329	<.001	0.946
Informational fairness	Item 1	1.000 ^a				0.881
	Item 2	0.940	0.046	20.260	<.001	0.853
Interactional fairness	Item 3	0.952	0.044	21.688	<.001	0.891
	Item 1	1.000 ^a				0.811
	Item 2	1.215	0.078	15.676	<.001	0.889

Note. All estimates are significant at $p < .001$. *B* = unstandardized loading; *SE* = standard error; *z* = *z*-statistic; β (Std.all) = completely standardized loading.

^aParameter fixed to 1.0 for model identification.

Disappointment

For the dependent variable disappointment, we found conclusive evidence against effects of decision-maker, $BF_{10} = 0.18 (\pm 0.59\%)$, and fairness communicated, $BF_{10} = 0.19 (\pm 0.45\%)$. We observed conclusive evidence for an effect of outcome favorability, $BF_{10} = 2.42e+27 (\pm 0.83\%)$. Perceived disappointment decreased with higher outcome favorability ratings. We found conclusive evidence against the interaction effects—Decision-Maker $\times$ Outcome Favorability, $BF_{10} = 0.09 (\pm 7.12\%)$; Fairness Communicated $\times$ Outcome Favorability, $BF_{10} = 0.09 (\pm 1.12\%)$; and the three-way interaction (Decision-Maker $\times$ Fairness Communicated $\times$ Outcome Favorability), $BF_{10} = 0.0008 (\pm 4.49\%)$. The interaction effect between decision-maker and fairness communicated was inconclusive, $BF_{10} = 0.43 (\pm 3.20\%)$.

Received September 26, 2025

Revision received January 14, 2026

Accepted January 14, 2026 ■