



OPEN ACCESS

EDITED BY

Nicola Mucci,
University of Florence, Italy

REVIEWED BY

Federica Picasso,
University of Genoa, Italy
Sascha Schweitzer,
Shenzhen Technology University, China

*CORRESPONDENCE

Zsolt Péter Szabó
✉ szabo.zsolt@ppk.elte.hu

RECEIVED 30 June 2026

REVISED 22 July 2026

ACCEPTED 30 July 2026

PUBLISHED 17 August 2026

CITATION

Marót JL, Palcsó T and Szabó ZP (2026)
Source matters, AI anxiety less so:
comparing employee reactions
to human, AI, and hybrid performance
feedback and the limited role of AI
anxiety.
Front. Psychol. 17:1924487.
doi: 10.3389/fpsyg.2026.1924487

COPYRIGHT

© 2026 Marót, Palcsó and Szabó. This is
an open-access article distributed under
the terms of the [Creative Commons
Attribution License \(CC BY\)](#). The use,
distribution or reproduction in other
forums is permitted, provided the
original author(s) and the copyright
owner(s) are credited and that the
original publication in this journal is
cited, in accordance with accepted
academic practice. No use, distribution
or reproduction is permitted which does
not comply with these terms.

Source matters, AI anxiety less so: comparing employee reactions to human, AI, and hybrid performance feedback and the limited role of AI anxiety

Janka Laura Marót^{1,2,3}, Tamara Palcsó³ and
Zsolt Péter Szabó^{2,3*}

¹Doctoral School of Psychology, ELTE Eötvös Loránd University, Budapest, Hungary, ²Institute of Psychology, ELTE Eötvös Loránd University, Budapest, Hungary, ³Institute of Strategy and Management, Corvinus University of Budapest, Budapest, Hungary

Introduction: Large language models (LLMs) are increasingly integrated into human resource management practices and are becoming capable of supporting managerial tasks such as performance feedback. Although these technologies may improve the efficiency and quality of feedback delivery, it remains unclear whether knowledge of AI involvement influences recipients' reactions to it.

Methods: This work presents results from a vignette-based experiment ($N = 192$) examining employees' reactions to a hypothetical performance feedback scenario. Participants were randomly assigned to one of four disclosed feedback-production workflows (fully human-written, fully AI-generated, AI-generated and human-refined, or human-written and AI-refined), while the feedback text itself was identical across conditions. Participants evaluated the feedback on a range of emotional, motivational, relational, and performance-related outcomes while AI anxiety was measured as a covariate in the analysis.

Results: Results demonstrated that fully AI-generated feedback elicited less favorable affective, motivational, relational, and performance-related responses than fully human-written feedback, whereas hybrid feedback generally occupied an intermediate position, with human-written feedback refined by AI producing outcomes comparable to purely human feedback. AI anxiety showed only limited effects.

Discussion: These findings suggest that the perceived source of performance feedback influences how feedback is evaluated, whereas AI anxiety showed only limited associations with employees' reactions. The discussion addresses the theoretical implications of these findings for performance feedback, as well as their practical implications for organizations.

KEYWORDS

AI anxiety, AI-generated feedback, disclosure effect, feedback source, generative artificial intelligence, human-AI collaboration, performance feedback

1 Introduction

Online discussions about generative artificial intelligence (genAI) increasingly reveal ambivalence toward managers' use of large language models (LLMs) in employee feedback. Some employees worry that AI-generated feedback may undermine the interpersonal value of managerial communication. As one Reddit user put it, "I want to know what my manager thinks of my performance, not AI," capturing the concern that feedback should reflect a manager's own judgment rather than the output of an AI system. Others adopt a more instrumental view, arguing that AI is simply "a tool" that can help managers communicate more effectively. These contrasting reactions point to an emerging tension in organizational life: although employees may value feedback as a relational and evaluative signal from their manager, LLMs are increasingly capable of producing and revising human-like feedback messages (García Peñalvo and Vázquez Ingelmo, 2023; Meyer et al., 2024), and AI-based tools are becoming more common in managerial and human resource practices (Budhwar et al., 2023; Kellogg et al., 2020; Tambe et al., 2019).

This tension is further complicated by the fact that recipients usually see only the final feedback message, not the process through which it was produced. Managers may use LLMs to generate feedback, revise their own feedback, or refine AI-generated drafts, but this involvement may not be visible in the final text.

Against this backdrop, the present study examines whether explicit source disclosure changes employees' responses to otherwise identical feedback in a hypothetical workplace scenario. Specifically, through a vignette experiment, we examine a wide range of reactions including recipients' emotional (affective response and rumination), motivational (willingness to correct mistakes and work motivation), relational responses (relationship with-, seeking help from-, and psychological distance from the supervisor) and perceived performance- and organization-related (expected performance and organizational commitment) reactions to feedback across different source conditions. In addition, we examine the potential role of AI anxiety as a covariate.

1.1 Literature review

Effective workplace feedback is important but difficult because it requires managers to combine accurate performance information with relationally sensitive communication. Although feedback can guide future behavior, clarify performance expectations, and support work motivation (Ilgen et al., 1979), it does not necessarily lead to performance improvement and may even undermine performance (Kluger and DeNisi, 1996). One reason for these mixed effects is the cognitive and relational complexity of feedback delivery (London, 2003). Feedback effectiveness depends not only on the content of the message, but also on the supervisor-subordinate relationship, recipients' perceptions and reactions, goal-setting processes, and the way feedback is delivered (Heine et al., 2026; Lechermeier and Fassnacht, 2018; Sleiman et al., 2020).

When these factors are poorly managed, feedback may elicit negative affective responses, which can even facilitate learning under certain conditions (Zhao, 2011), but they also lead to increased turnover intention, increased counterproductive

work behavior, defensiveness, anger, interpersonal conflict, and disengagement rather than development (Bakker and Leiter, 2010; Belschak and Den Hartog, 2009; Besieux, 2017; Marót et al., 2024). Negative feedback experiences may weaken perceptions of supervisory support and relationship quality (Heine et al., 2026). Consequently, employees may become less willing to seek guidance or assistance from their supervisor (Berkovich and Eyal, 2018). Negative feedback experiences may also trigger rumination, causing employees to repeatedly think about the feedback after the interaction (Mendiratta and Singh, 2026). Together, these affective, cognitive, and relational reactions may ultimately influence broader work-related attitudes and outcomes, including organizational commitment (Belschak and Den Hartog, 2009). These potential consequences illustrate both the importance and the complexity of effective feedback delivery, creating a potential role for AI-supported feedback tools, particularly LLMs, while also raising questions about how such tools may affect the meaning and interpretation of managerial feedback.

Artificial intelligence-supported feedback tools may help managers address some of the cognitive and linguistic demands of feedback delivery. Research on algorithmic judgment suggests that people often perceive algorithmic recommendations as useful and reliable (Logg et al., 2019), and AI systems may improve assessment quality in well-structured tasks by making evaluations more accurate, consistent, fair, or potentially less biased than those of human managers (Tong et al., 2021; Qin et al., 2023). Given the cognitive constraints and potential biases that often limit supervisors, AI may also support the generation of more personalized and comprehensive feedback (Tong et al., 2021). Thus, from a practical perspective, LLMs may assist managers by helping them draft, structure, or refine feedback messages.

At the same time, AI-supported feedback may be problematic because managerial feedback is not merely informational, but also relational. As discussed above, feedback communicates not only what an employee did well or poorly, but also whether the manager paid attention, exercised judgment, took responsibility for the employee's development, was available, and showed support. For this reason, feedback generated or substantially shaped by AI may be interpreted differently from feedback written by a human manager, even when the content is similar. Consistent with this view, although employees accept AI-generated feedback when it is perceived as clear, useful, and accurate, it also creates greater psychological distance because employees do not perceive AI systems as social actors, which may weaken its relational impact (Hein et al., 2024). Prior research suggests that disclosing AI or algorithmic involvement can reduce trust, perceived credibility, fairness, and warmth, even when the quality of the output is comparable to human judgment (Nazaretsky et al., 2026; Schilke and Reimann, 2025; Tong et al., 2021). Thus, source disclosure may change how recipients evaluate otherwise identical feedback. This concern is especially relevant in contexts that require intuition, contextual understanding, or subjective judgment, where AI-based evaluations may be perceived as less fair and less trustworthy than human evaluations (Lee, 2018). Performance feedback is such a context, because recipients may expect the message to reflect the manager's own knowledge, judgment, and involvement.

Disclosure effects may also depend on individual differences in how people perceive and respond to AI. People vary in their trust

in AI systems and in their comfort with AI involvement in human decision-making (Park et al., 2024). This may be especially relevant in managerial feedback, where AI involvement can raise concerns about human judgment, responsibility, or care. One relevant individual difference is AI anxiety, which refers to apprehension or discomfort associated with AI technologies. Because higher AI anxiety is associated with more negative attitudes toward AI (Kaya et al., 2024), individuals higher in AI anxiety may react more negatively when feedback is explicitly presented as AI-generated or AI-supported.

These concerns do not mean that AI-supported feedback must be fully automated. In practice, managers may use AI in more limited ways, for example to draft or refine feedback while retaining responsibility for the final message. Such hybrid forms of feedback may preserve some advantages of AI, including efficiency and linguistic support, while maintaining human involvement in the evaluation process (Biswas et al., 2024; Jarrahi, 2018; Trunk et al., 2020). Similar approaches have shown promise in educational settings, where human–AI collaboration may reduce some of the negative reactions associated with fully automated feedback (Henderson et al., 2025; Tossell et al., 2024; Zawiah et al., 2023; Zhang et al., 2025). However, less is known about how hybrid feedback is perceived in workplace feedback contexts.

In sum, AI-supported feedback creates a dilemma for managerial feedback. AI may help managers produce feedback that is more consistent, comprehensive, and useful, but its involvement may also change how feedback is interpreted, especially when recipients expect performance feedback to reflect the manager's own judgment and involvement. Therefore, understanding AI-supported feedback requires exploring how the disclosed source of feedback influences employees' evaluations of it.

1.2 The present research

The preceding arguments suggest that AI-supported feedback should be examined in terms of how AI is involved in feedback production. We therefore distinguish between four feedback-production workflows: fully human-written feedback (H), fully AI-generated feedback (AI), AI-generated feedback refined by a human (AIxH), and human-written feedback refined by AI (HxAI).

Although AI-generated feedback may offer some advantages over human-written feedback, feedback is also a relational act, therefore, feedback disclosed as fully human-written should elicit more positive reactions than feedback disclosed as fully AI-generated.

H1: Fully human-written feedback will elicit more positive emotional, motivational, and perceived performance-related reactions than fully AI-generated feedback.

Hybrid feedback may occupy an intermediate position. Because such feedback involves AI, it may be evaluated less positively than fully human-written feedback. However, because it also involves human input, it may be evaluated more positively than fully AI-generated feedback. Because the relative evaluation of hybrid feedback in managerial feedback contexts remains unclear, we examined this question in an exploratory manner.

Exploratory RQ1: How do recipients respond to hybrid human–AI feedback relative to fully human-written and fully AI-generated feedback?

At the same time, the two hybrid workflows may not be psychologically equivalent. AI-generated feedback refined by a human may signal that the human manager had final oversight, whereas human-written feedback refined by AI may signal that the manager originated the evaluation but used AI to improve the wording. Because these workflows imply different forms of human and AI involvement, we examined their difference as a research question.

Exploratory RQ2: Do recipients respond differently to AI-generated feedback refined by a human compared with human-written feedback refined by AI?

Finally, reactions to AI involvement may depend on individual differences. In particular, individuals higher in AI anxiety may respond more negatively when feedback is presented as AI-generated or AI-supported. Because individual differences in attitudes toward AI may influence evaluations of AI-supported feedback, AI anxiety was measured and included as a covariate in the analyses. Although the role of AI anxiety in managerial feedback contexts remains insufficiently understood, its potential influence was considered exploratory.

Exploratory RQ3: Do differences in recipients' reactions to disclosed feedback sources remain when accounting for AI anxiety?

To address these questions, we conducted an experimental vignette study in which participants were presented with an identical feedback text embedded in a hypothetical workplace scenario, varying only in the disclosed source (human, AI, and hybrid workflows). By holding the feedback message constant, this experimental vignette design allowed us to isolate the effect of source disclosure on participants' immediate reactions. We examined a wide range of outcomes, including recipients' emotional (affective response and rumination), motivational (willingness to correct mistakes and work motivation), relational (relationship with-, seeking help from-, and psychological distance from the supervisor), and performance- and organization-related reactions (expected performance and organizational commitment). In addition, we explored the potential role of AI anxiety as a covariate. Although this design supports causal inference about immediate scenario-based reactions, it captures anticipated responses to a standardized hypothetical situation rather than employees' actual behavior or longer-term reactions to consequential feedback within an ongoing workplace relationship.

2 Materials and methods

Data collection was conducted in 2025 through an online Qualtrics survey. Participants provided informed consent at the beginning in accordance with the Declaration of Helsinki. Those

who did not provide consent were directed to the end of the survey and did not participate further. The study received ethical approval from the Research Ethics Committee of Corvinus University of Budapest (file number: KRH/256/2025).

Statistical analysis and data visualization were conducted in R version 4.6.1 (R Core Team., 2026). The tidyverse (Wickham et al., 2019), ggplot2 (Wickham, 2016), and sjPlot (Lüdtke, 2025) packages were used for data transformation, analysis, and visualization. Power analysis was conducted using G*Power 3.1.9.7 (Faul et al., 2009).

2.1 Participants

An *a priori* power analysis indicated that, assuming an expected effect size of $f = 0.25$ (power = 0.80, $\alpha = 0.05$), the required sample size was 179 participants. A total of 230 Hungarian participants were recruited through convenience sampling via social media platforms and the researchers' professional and personal networks. Participants were excluded if they did not provide informed consent ($n = 0$), or were currently unemployed ($n = 18$). To minimize low-quality responses, we included one attention check and excluded 20 participants who failed it. The final sample consisted of 192 employees, which met the sample size requirement indicated by the *a priori* power analysis. A step-by-step breakdown of participant counts and exclusions across conditions is provided in [Supplementary Material 1](#).

Of the final sample, 105 participants were female (54.69%) and 87 were male (45.31%). Participants ranged in age from 19 to 65 years ($M_{age} = 32.34$ years, $SD_{age} = 11.84$ years). Regarding education, 130 participants had a higher education degree (67.71%), 58 had completed high school or lower education (30.21%), and four selected "other" (2.08%). In terms of organizational position, 127 participants worked at the subordinate level (66.15%), 32 were middle managers (16.67%), 12 were executives (6.25%), and 21 selected "other" (10.94%).

2.2 Procedure

Participants first completed self-reported questions on work status and basic demographic characteristics. They were then presented with a vignette describing a hypothetical workplace situation in which they were asked to imagine receiving performance feedback from a fictional supervisor. Participants were randomly assigned to one of four conditions that differed only in the disclosed feedback-production process: the supervisor was described as either personally writing all feedback (fully human-written, $n = 42$), using AI-generated feedback without modification (fully AI-generated, $n = 56$), reviewing and approving AI-generated feedback before sending it (AI-generated refined by human, $n = 44$), or personally writing feedback and subsequently refining it with the assistance of AI (human-written refined by AI, $n = 50$). After this, all participants were presented with the same "excerpt [...] from their most recent performance review." The feedback contained both positive ("strong in meeting deadlines and completing administrative tasks") and developmental elements ("pay more conscious attention to joint discussions and communication"). Participants subsequently completed measures assessing their

reactions to the feedback. The full vignettes are presented in [Supplementary Material 2](#).

2.3 Measures

Immediately after the vignette, participants answered three open-ended questions designed to capture their initial thoughts, feelings, and intended actions in response to the feedback. They were asked to describe in one or two sentences what thoughts came to mind when they received the feedback, how they felt about the feedback, and how they would act after receiving it.

Participants then completed a manipulation check to assess whether they correctly understood the disclosed source of the feedback. They answered the question, "According to the situation, How does your manager provide feedback?" Response options were: "personally themselves," "using artificial intelligence, without modification," "using artificial intelligence, with personal review and approval," and "personally, then refined and corrected using artificial intelligence."

The next block consisted of the main reaction measures. Self-reported single-item outcomes were rated on 7-point Likert-type scales. The assessed outcomes were general affective response ("How would that feedback make you feel?"; 1 = Very bad, 7 = Very good), rumination ("How much would you ruminate on this feedback throughout the day?"; 1 = Not at all, 7 = All the time), willingness to correct the mistake ("After this feedback, To what extent do you intend to correct the mistake?"; 1 = I would not intend to at all, 7 = I would fully intend to), openness to ask for help from the supervisor ("After this feedback, How likely do you think it is that you would go to your supervisor with professional questions?"; 1 = Not likely at all, 7 = Entirely likely), organizational commitment ("How would this feedback affect your commitment to the organization?"; 1 = It would greatly reduce it, 7 = It would greatly increase it), relationship with the supervisor ("How would this feedback affect the relationship with your supervisor?"; 1 = It would greatly worsen it, 7 = It would greatly improve it), expected performance ("How would this feedback affect your performance?"; 1 = It would greatly worsen it, 7 = It would greatly improve it), and work motivation ("How would this feedback affect your work motivation?"; 1 = It would greatly reduce it, 7 = It would greatly increase it). Perceived psychological distance from the feedback source was assessed with a four-item scale based on the theoretical framework of Hein et al. (2024). An example item is: "I feel that the person giving feedback is emotionally close to me." Items were rated on a 5-point Likert-type scale ranging from 1 = not at all to 5 = completely. The scale showed high internal consistency ($\alpha = 0.89$).

Finally, AI anxiety was assessed as an exploratory measure using the AI Anxiety Scale (AIAS; Wang and Wang, 2022). Because no validated Hungarian version was available, the items were translated using the back-translation method (Brislin, 1970). The scale includes 21 items rated on a 7-point Likert-type scale ranging from 1 = not at all to 7 = completely and covers four dimensions: learning anxiety (8 items; e.g., "I feel anxious when I have to learn how to use AI technologies/products"), job replacement anxiety (6 items; e.g., "I am afraid that AI technologies/products will replace humans"), sociotechnical blindness (4 items; e.g., "I am afraid that AI technology/products could be misused")

and AI configuration anxiety [3 items; e.g., “I find humanoid AI technologies/products (e.g., humanoid robots) frightening.”]. Although the AIAS comprises multiple dimensions, it was developed to capture a broader AI anxiety construct. Accordingly, subsequent analyses were conducted using the overall AIAS score rather than the individual subscale scores (Wang and Wang, 2022). Cronbach’s alpha coefficients indicated excellent reliability for the total scale ($\alpha = 0.94$).

2.4 Data analysis

To address the study objectives, a manipulation check was conducted using a chi-square test to examine whether participants correctly recognized the disclosed source of the feedback. Differences in reactions to feedback across the four source conditions were examined using multivariate analysis of variance (MANOVA). MANOVA was used instead of a series of separate ANOVAs because it accounts for potential correlations among the outcome variables (Grice and Iwasaki, 2007). AI anxiety was included in the model as a covariate to test whether reactions to different feedback sources varied with individual differences in AI anxiety. *Post hoc* pairwise comparisons of estimated marginal means (EMMs) were conducted using Bonferroni-adjusted *p*-values.

3 Results

3.1 Preliminary analysis

The manipulation check confirmed effectiveness: a chi-square test indicated a significant association between the disclosed and perceived source of feedback, $\chi^2(9) = 162.28$, $p < 0.001$, suggesting that participants were able to correctly identify ($n = 123$, 64.06%) the feedback source above chance level (25%).

Pre-manipulation differences in sociodemographic variables were examined to assess the random allocation of participants to conditions. The analysis revealed no statistically significant baseline differences between the four conditions for gender ($p = 0.808$) and position ($p = 0.421$), except for level of education ($p = 0.029$). Therefore, all analyses were verified by including education level as a control variable, which did not yield any meaningful changes in the outcomes.

3.2 Main analyses

A multivariate analysis of variance (MANOVA) was conducted to test whether the disclosed source of the feedback (human, AI, AI-generated refined by a human, and human-written refined by AI) influenced participants’ reactions across nine outcome variables: willingness to correct the mistake, openness to ask for help from the supervisor, organizational commitment, relationship with the supervisor, expected performance, general affective response, rumination, motivation, and perceived psychological distance. Using Pillai’s trace, the overall multivariate effect of feedback source was significant, $V = 0.37$, $F(27, 546) = 2.86$, $p < 0.001$.

Follow-up univariate ANOVAs revealed significant, mostly small-to-moderate effects of feedback source on affective response, $F(3, 188) = 5.16$, $p = 0.002$, $\eta^2 = 0.08$; willingness to correct the mistake, $F(3, 188) = 10.45$, $p < 0.001$, $\eta^2 = 0.14$; openness to ask for help from the supervisor, $F(3, 188) = 13.17$, $p < 0.001$, $\eta^2 = 0.17$; organizational commitment, $F(3, 188) = 11.67$, $p < 0.001$, $\eta^2 = 0.16$; relationship with the supervisor, $F(3, 188) = 16.81$, $p < 0.001$, $\eta^2 = 0.21$; expected performance, $F(3, 188) = 6.70$, $p < 0.001$, $\eta^2 = 0.10$; motivation, $F(3, 188) = 9.43$, $p < 0.001$, $\eta^2 = 0.13$; and perceived psychological distance, $F(3, 188) = 20.12$, $p < 0.001$, $\eta^2 = 0.24$. No significant effect was found for rumination, $F(3, 188) = 0.57$, $p = 0.64$, $\eta^2 = 0.01$.

Bonferroni-adjusted pairwise comparisons showed that fully AI-generated feedback was evaluated less favorably than fully human feedback on affective response (EMMs = 3.86 vs. 4.76, $p = 0.002$), willingness to correct the mistake (EMMs = 4.43 vs. 6.02, $p < 0.001$), openness to ask for help (EMMs = 3.43 vs. 5.31, $p < 0.001$), organizational commitment (EMMs = 3.46 vs. 4.57, $p < 0.001$), relationship with the supervisor (EMMs = 3.23 vs. 4.69, $p < 0.001$), expected performance (EMMs = 4.29 vs. 5.19, $p < 0.001$), motivation (EMMs = 3.66 vs. 4.79, $p < 0.001$), and perceived psychological distance (EMMs = 1.67 vs. 2.89, $p < 0.001$). Fully AI-generated feedback was also evaluated less favorably than human-written feedback refined by AI on affective response (EMMs = 3.86 vs. 4.48, $p = 0.046$), willingness to correct the mistake (EMMs = 4.43 vs. 5.56, $p = 0.001$), openness to ask for help (EMMs = 3.43 vs. 4.82, $p < 0.001$), organizational commitment (EMMs = 3.46 vs. 4.38, $p < 0.001$), relationship with the supervisor (EMMs = 3.23 vs. 4.22, $p < 0.001$), expected performance (EMMs = 4.29 vs. 4.82, $p = 0.042$), motivation (EMMs = 3.66 vs. 4.68, $p < 0.001$), and perceived psychological distance (EMMs = 1.67 vs. 2.54, $p < 0.001$). When compared to AI-generated feedback refined by a human, fully AI-generated feedback was again associated less favorably on openness to ask for help (EMMs = 3.43 vs. 4.57, $p = 0.002$), relationship with the supervisor (EMMs = 3.23 vs. 3.98, $p = 0.003$), and perceived psychological distance (EMMs = 1.67 vs. 2.15, $p = 0.023$). No other pairwise differences were significant. Figure 1 depicts these differences, and Table 1 summarizes them (see more details in Supplementary Materials 3, 4).

3.3 Robustness analysis among participants who passed the manipulation check

To examine the robustness of the findings, the main analyses were repeated among participants who correctly identified the disclosed feedback source in the manipulation check ($n = 123$). The multivariate effect of feedback source remained significant, $V = 0.70$, $F(27, 339) = 3.83$, $p < 0.001$. The univariate results followed the same overall pattern as those obtained in the full sample. Significant effects emerged for all outcomes except rumination, and the effect sizes were generally larger, ranging from $\eta^2 = 0.19$ to 0.42. The main pairwise differences were also replicated, although several additional significant comparisons emerged. Overall, the robustness analysis supported the main pattern of findings while indicating more pronounced differences

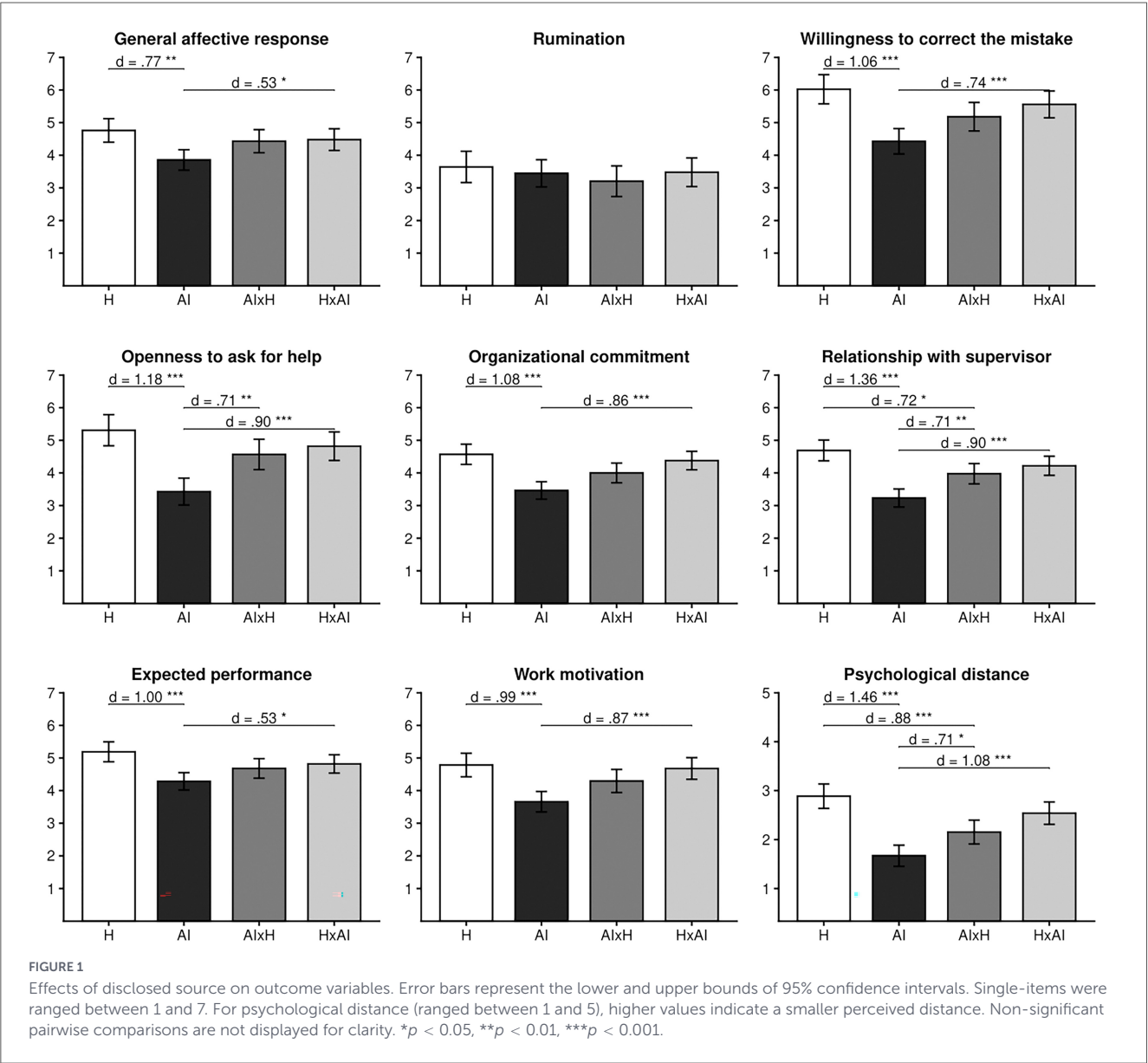


TABLE 1 Summary of the main analysis results.

Outcome	EMM (H)	EMM (AI)	EMM (AIxH)	EMM (HxAI)	F (3, 188)	p	η^2
Affective response	4.76	3.86	4.43	4.48	5.16	0.002	0.08
Rumination	3.64	3.45	3.20	3.48	0.57	0.640	0.01
Willingness to correct the mistake	6.02	4.43	5.18	5.56	10.45	<0.001	0.14
Openness to ask for help	5.31	3.43	4.57	4.82	13.17	<0.001	0.17
Organizational commitment	4.57	3.46	4.00	4.38	11.67	<0.001	0.16
Relationship with the supervisor	4.69	3.23	3.98	4.22	16.81	<0.001	0.21
Expected performance	5.19	4.29	4.68	4.82	6.70	<0.001	0.10
Work motivation	4.79	3.66	4.30	4.68	9.43	<0.001	0.13
Psychological distance	2.89	1.67	2.15	2.54	20.12	<0.001	0.24

H = fully human-written, AI = fully AI-generated, AIxH = AI-generated and refined by a human, and HxAI = human-written and refined by AI. Higher scores indicate more positive reactions, except for rumination. For psychological distance, higher scores indicate lower perceived distance. η^2 = eta squared.

among participants who correctly identified the disclosed source. Detailed results are presented in [Supplementary Material 5](#).

3.4 AI anxiety as a covariate

Artificial intelligence anxiety was assessed after the outcome measures to avoid drawing participants' attention to AI-related concerns before they evaluated the feedback. Given this measurement order, condition assignment may have influenced participants' reported AI anxiety. Before conducting the covariate analysis, we therefore examined whether AIAS scores differed across conditions. No significant difference was found, $F(3, 188) = 0.55, p = 0.651$.

Artificial intelligence anxiety was subsequently included as a covariate in the model examining its association with participants' reactions to feedback. Using Pillai's trace, the multivariate effect of feedback source remained significant, $V = 0.37, F(27, 543) = 2.85, p < 0.001$ (see more details in [Supplementary Material 6](#)). In contrast, the multivariate effect of AI anxiety was not significant, $V = 0.08, F(9, 179) = 1.73, p = 0.085$.

At the univariate level, AI anxiety was significantly associated only with rumination, $F(1, 187) = 14.76, p < 0.001, \eta^2 = 0.07$, indicating that participants reporting higher AI anxiety also reported greater rumination following feedback. No other significant effects of AI anxiety emerged for outcomes (all $ps > 0.301$, see more details in [Supplementary Material 7](#)).

Taken together, these results suggest that although AI anxiety was associated with greater rumination overall, it did not systematically alter the pattern of reactions to the different feedback sources.

4 Discussion

Building on the growing body of research examining the integration of AI into performance appraisal processes ([Biswas et al., 2024; Jarrahi, 2018; Pei et al., 2024; Tong et al., 2021](#)), the present study aimed to investigate how the knowledge of the feedback source influences reactions to feedback in an organizational context. Specifically, we compared how employees react to feedback when the source (fully human-written, fully AI-generated, AI-generated refined by a human, and human-written refined by AI) is explicitly revealed. In addition, we explored AI anxiety as a covariate in the relationship between disclosed feedback sources and reactions.

Consistent with our hypotheses and prior research ([Nazaretsky et al., 2026; Tong et al., 2021](#)), fully human-written feedback elicited more favorable reactions than fully AI-generated feedback (H1). Specifically, human feedback elicits more positive affective responses, greater willingness to correct mistakes, higher openness to ask for help, stronger organizational commitment, better relationship with supervisors, higher motivation, and lower perceived psychological distance. These findings support the disclosure effect: once the source of feedback is known, it alters its effect ([Tong et al., 2021](#)).

These findings also extend the broader workplace feedback literature by demonstrating that recipients' reactions depend

not only on the content of the message but also on how the feedback is delivered and who is perceived to be responsible for it ([Ilgen et al., 1979; London, 2003; Lechermeier and Fassnacht, 2018](#)). Although participants received identical feedback, source disclosure influenced emotional, motivational, organizational, and relational outcomes. Differences in openness to ask for help, relationship with the supervisor, and psychological distance align with research emphasizing that feedback is embedded in the supervisor-employee relationship and may shape employees' willingness to seek subsequent guidance and support ([Berkovich and Eyal, 2018; Heine et al., 2026](#)).

One possible explanation for this pattern may lie in perceptions of source credibility. Prior research suggests that feedback can be effective only in influencing behavior if it is accepted and internalized by the recipient ([Ashford, 1986](#)). Feedback acceptance is closely related to the extent to which the receiver perceives the source as credible and professionally competent, as feedback from credible sources has been shown to increase acceptance and reduce defensiveness ([Ilgen et al., 1979](#)). Regarding supervisors, credibility has been shown to not only enhance feedback acceptance, but motivation, and ultimately performance as well, partly because it signals trust and reflects the quality of the supervisor-subordinate relationship ([Son and Kim, 2016; Steelman and Rutkowski, 2004; van de Ridder et al., 2015](#)).

From this perspective, the negative reactions toward disclosed AI involvement may reflect not the content of the feedback itself, but changes in how the source is perceived. In human interactions, a credible source is perceived as someone who has knowledge about the recipient (expertise), is willing to share their knowledge with the recipient (trustworthiness), and cares about the recipient (goodwill) ([Perloff, 2003](#)). While AI systems may fulfill the expertise dimension due to their technical advantages, their ability to convey trustworthiness and goodwill remains questionable. On the one hand, AI systems are designed to generate responses based on their "best available knowledge," even if this means hallucinations can occur ([Kaate et al., 2025](#)). On the other hand, it is not clear whether AI should be designed to simulate relational qualities such as trustworthiness and goodwill, as doing so may make users vulnerable to its influence ([Carroll, 2025; Henderson et al., 2025](#)). When supervisors incorporate AI into the feedback process, the boundaries of authorship and responsibility may become blurred, potentially making the credibility of the feedback source more ambiguous. As a result, it becomes a key question how to design AI-supported feedback that preserves source credibility and delivers effective feedback.

Our study also reinforces prior research on perceptions of human-AI collaboration ([Biswas et al., 2024; Jarrahi, 2018; Trunk et al., 2020](#)), suggesting that hybrid feedback may offer a promising solution. Although the mean values indicate a clear pattern in which purely human feedback is evaluated most positively, followed by human-written feedback refined by AI, then AI-generated feedback refined by human, and finally fully AI-generated feedback, statistical analyses reveal a more nuanced picture. Addressing RQ1, human-written feedback refined by AI did not differ significantly from purely human feedback and was evaluated more positively than fully AI-generated feedback on all outcomes except rumination. AI-generated feedback refined by a human also outperformed fully AI-generated feedback, although these differences were limited to relational dimensions,

including openness to ask for help, relationship quality with the supervisor, and psychological distance. At the same time, addressing RQ2, no significant differences emerged between the two hybrid workflows, suggesting that recipients do not meaningfully distinguish between different forms of human–AI collaboration despite their theoretically distinct patterns of human and AI involvement. Interpreted through the source credibility perspective outlined above, these findings may indicate that hybrid forms of feedback evoke source-related perceptions similar to those associated with purely human feedback. Importantly, the few remaining differences emerged between purely human feedback and AI-generated feedback refined by a human, and were concentrated in relational dimensions, suggesting that AI involvement may primarily affect the interpersonal component of the feedback process even when overall evaluations remain largely similar. This pattern is consistent with the possibility that preserving human authorship while using AI primarily as a refinement tool may be less disruptive to relationship-related perceptions than workflows in which AI is perceived as the originator of the feedback, while still allowing managers to draw on the broader potential of AI to improve consistency, comprehensiveness, and linguistic quality (Qin et al., 2023; Tong et al., 2021). Regarding our research question about how individual differences influence reactions to disclosed AI involvement in feedback processes, AI anxiety did not appear to shape reactions to feedback in the present study (RQ3). However, it is possible that other AI-related variables, such as broader attitudes toward AI or technology, may play a more important role. In line with the source credibility perspective discussed above, it is also possible that reactions to AI-supported feedback are linked less to AI itself and more to how AI involvement influences perceptions of the feedback source and the relational context surrounding feedback delivery.

4.1 Contributions

Beyond comparing AI and human feedback, our study contributes to the evaluation of hybrid feedback categories, extending prior evidence from educational settings (Henderson et al., 2025; Tossell et al., 2024; Zawiah et al., 2023; Zhang et al., 2025) to organizational contexts, showing that - except for relational dimensions - they do not perform worse than human feedback. Furthermore, our findings extend the literature on factors (such as attitudes towards AI - Park et al., 2024 or anthropocentrism - Wang et al., 2026) shaping perceptions of AI use in performance appraisal by examining the role of AI anxiety.

From a practical perspective, our findings reinforce Wang et al. (2026) argument that fully replacing managers with AI, specifically in the context of feedback delivery, is not advisable. However, the hybrid approaches, particularly using AI as a reviewing or support tool, appear to be a viable option for improving both the quality and efficiency of performance appraisal processes. AI can help managers keep track of the many criteria outlined in prior studies (Heine et al., 2026; Lechermeier and Fassnacht, 2018; Sleiman et al., 2020), process large volumes of performance-related data, and provide data-driven insights to support more informed evaluations, while also assisting in the refinement of written feedback and improving its constructiveness.

Implementing AI-supported feedback systems may also contribute to the long-term development of managers' feedback-giving skills, as they receive continuous (constructive) input on how to improve their feedback. However, these potential effects need further investigation, as they were not examined in the present study.

4.2 Limitations and future research directions

Although the results are clear, this study has several limitations. First, our hypothetical vignette design cannot fully capture the complexity of real workplace dynamics. Although holding the text constant strengthens causal identification of the labeling effect, participants reported anticipated reactions to a salient disclosure rather than behavior following feedback from an actual supervisor. The large differences observed for some comparisons (up to $d = 1.46$) may therefore partly reflect label salience or demand characteristics and may overstate effects likely in field settings. Therefore, future research could benefit from examining feedback processes in more naturalistic, real-world settings involving actual managers and employees, while carefully considering ethical constraints.

Second, the use of single-item measures limits the precision of construct measurement. Therefore, our findings primarily reflect immediate reactions to feedback rather than deeper psychological constructs or long-term outcomes such as future performance. Future studies could address this limitation by employing validated multi-item scales to capture psychological constructs more reliably.

Third, given that the strongest effects were observed in relational dimensions (e.g., openness to ask for help, relationship with the supervisor, psychological distance), future research should further explore these aspects, as prior work highlights the central role of relational context in determining feedback effectiveness (Heine et al., 2026). In the present study, we focused primarily on employees' reactions to feedback rather than on how they directly perceived the feedback source itself. Nevertheless, perceptions related to source credibility may represent an important underlying mechanism shaping responses to AI-supported feedback. Although prior research has already examined the credibility of AI-supported or AI-generated communication (e.g., Khan and Mishra, 2024; Ou et al., 2026; Shin, 2022), these studies have primarily been conducted outside workplace feedback contexts. Future studies could therefore benefit from directly examining how employees evaluate the credibility and perceived responsibility of AI-involved feedback sources in organizational contexts. Longitudinal research would be particularly valuable to examine how these perceptions evolve over time and influence longer-term outcomes.

Fourth, all participants were recruited in Hungary through convenience sampling. This non-probability approach creates self-selection concerns: people reached through the researchers' networks or motivated to complete a feedback-related survey may differ systematically from the broader employee population. Also, the cultural differences in attitudes toward artificial intelligence may influence responses to AI-generated feedback. Prior research in the Hungarian context has shown that the majority of individuals hold ambivalent attitudes toward AI, simultaneously recognizing its potential benefits and risks, while only a minority can be

characterized as clearly positive or clearly negative toward the technology (Horváth and Molnár, 2026). Consequently, reactions to AI-generated feedback observed in the present studies may partly reflect the specific attitudinal profile of Hungarian employees and may differ in populations with higher levels of AI acceptance or skepticism. In addition, due to the relatively small sample size, smaller effect sizes may not have been detected. Replication in larger and more demographically and culturally diverse samples is therefore needed.

Finally, two measurement limitations should be considered. First, the manipulation check assessed whether participants could correctly report the disclosed feedback-production workflow, but not whether they believed that the feedback had actually been produced in the stated way. Participants' perceptions of the credibility of the disclosure may therefore have influenced their responses. Future studies should include an additional question assessing whether participants believed the disclosed source information. Second, AI anxiety was measured after the manipulation. Although AIAS scores did not differ significantly across conditions in the present sample, it remains unclear whether the observed pattern would have been the same if AI anxiety had been measured earlier. Future studies should therefore assess AI anxiety in a separate measurement wave so that participants are less likely to associate the scale with the purpose of the study. Future research should also explore additional moderating and mediating factors, such as attitudes toward AI or personality traits, that may shape responses to AI-supported feedback. It would also be valuable to examine how results change when the source of feedback is not "widely known" but only "assumed" by employees. In the present study, the disclosed feedback source was clearly defined, however, in real organizational contexts, transparency around feedback processes may be limited (Reddy Yanamala, 2023).

5 Conclusion

This study reinforces the disclosure effect: when the source of feedback is known, human feedback is preferred over AI-generated feedback. Importantly, no significant differences were found between human-written and human-written-AI-refined hybrid feedback, suggesting that appropriate human-AI collaboration represents an effective alternative in performance appraisal contexts, enabling more efficient feedback processes while maintaining quality. Altogether, these results suggest that organizations and managers should not view AI as a replacement for human judgment, but rather as a complementary tool that can enhance feedback processes when carefully integrated and appropriately managed.

Data availability statement

The datasets presented in this study can be found in online repositories. The names of the repository/repositories and accession number(s) can be found below: Open Science

Framework: https://osf.io/whu35/overview?view_only=aec95000601a4c5b8c3b81bd326683fd.

Ethics statement

The studies involving humans were approved by Research Ethics Committee of Corvinus University of Budapest (file number: KRH/256/2025). The studies were conducted in accordance with the local legislation and institutional requirements. The participants provided their written informed consent to participate in this study.

Author contributions

JM: Conceptualization, Data curation, Formal analysis, Methodology, Project administration, Validation, Visualization, Writing – original draft, Writing – review & editing, Investigation, Resources. TP: Conceptualization, Formal analysis, Investigation, Methodology, Validation, Writing – original draft, Data curation. ZS: Conceptualization, Formal analysis, Methodology, Validation, Writing – original draft, Supervision, Writing – review & editing.

Funding

The author(s) declared that financial support was received for this work and/or its publication. Financial support for the publication of this article is provided by the Faculty of Education and Psychology, Eötvös Loránd University.

Acknowledgments

We thank Lilla Török for support in collecting data.

Conflict of interest

The author(s) declared that this work was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.

Generative AI statement

The author(s) declared that Generative AI was used in the creation of this manuscript. To enhance the clarity and conciseness

of the manuscript, assist in shortening the text, and support the development of R codes. The authors did not use AI tools for generating scientific content or drawing scientific conclusions, and firmly reject any unethical use of AI for research purposes. However, the authors believe that it helps fostering the equality of native and non-native English speakers in order for them to have the same opportunities. The use of AI was limited to linguistic refinement and technical support. All content was critically reviewed and edited by the authors, who take full responsibility for the final version of the manuscript.

Any alternative text (alt text) provided alongside figures in this article has been generated by Frontiers with the support of artificial intelligence and reasonable efforts have been made to ensure accuracy, including review by the authors wherever possible. If you identify any issues, please contact us.

Publisher's note

All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated organizations, or those of the publisher, the editors and the reviewers. Any product that may be evaluated in this article, or claim that may be made by its manufacturer, is not guaranteed or endorsed by the publisher.

Supplementary material

The Supplementary Material for this article can be found online at: <https://www.frontiersin.org/articles/10.3389/fpsyg.2026.1924487/full#supplementary-material>

References

- Ashford, S. J. (1986). Feedback-seeking in individual adaptation: A resource perspective. *Acad. Manag. J.* 29, 465–487. doi: 10.2307/256219
- Bakker, A., and Leiter, M. P. (2010). "Where to go from here? Integration and future research on work engagement," in *Work engagement: A Handbook Of Essential Theory And Research*, eds A. B. Bakker and M. P. Leiter. (London: Psychology Press (Taylor and Francis Group)).
- Belschak, F. D., and Den Hartog, D. N. (2009). Consequences of positive and negative feedback: The impact on emotions and extra-role behaviors. *App. Psychol.* 58, 274–303. doi: 10.1111/j.1464-0597.2008.00336.x
- Berkovich, I., and Eyal, O. (2018). Help me if you can: Psychological distance and help-seeking intentions in employee-supervisor relations. *Stress Health.* 34, 425–434. doi: 10.1002/smi.2802
- Besieux, T. (2017). Why I hate feedback: Anchoring effective feedback within organizations. *Bus.Horiz.* 60, 435–439. doi: 10.1016/j.bushor.2017.03.001
- Biswas, M. I., Talukder, Md. S., and Khan, A. R. (2024). Who do you choose? Employees' perceptions of artificial intelligence versus humans in performance feedback. *China Acc. Fin. Rev.* 26, 512–532. doi: 10.1108/CAFR-08-2023-0095
- Brislin, R. W. (1970). Back-Translation for cross-cultural research. *J. Cross-Cul. Psychol.* 1, 185–216. doi: 10.1177/135910457000100301
- Budhwar, P., Chowdhury, S., Wood, G., Aguinis, H., Bamber, G. J., Beltran, J. R., et al. (2023). Human resource management in the age of generative artificial intelligence: Perspectives and research directions on ChatGPT. *J. Hum. Resour. Manag.* 33, 606–659. doi: 10.1111/1748-8583.12524
- Carroll, N. G. (2025). Artificial goodwill and human vulnerability: The case for building merely reliable, rather than trustworthy, artificial intelligence technologies. *Philo. Technol.* 38:65. doi: 10.1007/s13347-025-00881-w
- Faul, F., Erdfelder, E., Buchner, A., and Lang, A. G. (2009). Statistical power analyses using G*Power 3.1: Tests for correlation and regression analyses. *Behav Res Methods.* 41, 1149–1160. doi: 10.3758/BRM.41.4.1149
- García Peñalvo, F. J., and Vázquez Ingelmo, A. (2023). What do we mean by genai? A systematic mapping of the evolution, trends, and techniques involved in generative Ai. *international. J. Interact. Mult. Artif. Intell.* 8, 7–16. doi: 10.9781/ijimai.2023.07.006
- Grice, J. W., and Iwasaki, M. (2007). A truly multivariate approach to manova. *Appl. Mult. Res.* 12:199. doi: 10.22329/amr.v12i3.660
- Hein, I., Cecil, J., and Lerner, E. (2024). Acceptance and motivational effect of AI-driven feedback in the workplace: An experimental study with direct replication. *Front. Organ. Psychol.* 2:1468907. doi: 10.3389/forgp.2024.1468907
- Heine, E. C. E., Stouten, J., and Liden, R. C. (2026). Performance feedback: A critical systematic review. *J. Organ. Behav.* 47, 312–361. doi: 10.1002/job.70033
- Henderson, M., Bearman, M., Chung, J., Fawns, T., Buckingham Shum, S., Matthews, K. E., et al. (2026). Comparing generative AI and teacher feedback: Student perceptions of usefulness and trustworthiness. *Ass. Eval. High. Educ.* 51, 863–878. doi: 10.1080/02602938.2025.2502582
- Horváth, K., and Molnár, L. (2026). AI-related attitude segments in hungary: How personal values influence the way we perceive the technology. *Vezet. Budapest Manag. Rev.* 57, 2–18. doi: 10.14267/VEZTUD.2026.03.01
- Ilgén, D. R., Fisher, C. D., and Taylor, M. S. (1979). Consequences of individual feedback on behavior in organizations. *J. App. Psychol.* 64, 349–371. doi: 10.1037/0021-9010.64.4.349
- Jarrahi, M. H. (2018). Artificial intelligence and the future of work: Human-AI symbiosis in organizational decision making. *Busin. Horiz.* 61, 577–586. doi: 10.1016/j.bushor.2018.03.007
- Kaate, I., Salminen, J., Jung, S.-G., Xuan, T. T. T., Häyhänen, E., Azem, J. Y., et al. (2025). "You always get an answer": Analyzing users' interaction with AI-generated personas given unanswerable questions and risk of hallucination," in *Proceedings of the 30th International Conference on Intelligent User Interfaces* (New York, NY: Association for Computing Machinery), 1624–1638. doi: 10.1145/3708359.3712160
- Kaya, F., Aydin, F., Schepman, A., Rodway, P., Yetişensoy, O., and Demir Kaya, M. (2024). The roles of personality traits AI anxiety, and demographic factors in attitudes toward artificial intelligence.. *Int. J. Hum. Comp. Int.* 40, 497–514. doi: 10.1080/10447318.2022.2151730
- Kellogg, K. C., Valentine, M. A., and Christin, A. (2020). Algorithms at work: The new contested terrain of control. *Acad. Manag. Annals* 14, 366–410. doi: 10.5465/annals.2018.0174
- Khan, A. W., and Mishra, A. (2024). AI credibility and consumer-AI experiences: A conceptual framework. *J. Theor. Prac.* 34, 66–97. doi: 10.1108/JSTP-03-2023-0108
- Kluger, A. N., and DeNisi, A. (1996). The effects of feedback interventions on performance: A historical review, a meta-analysis, and a preliminary feedback intervention theory. *Psychol. Bull.* 119, 254–284. doi: 10.1037/0033-2909.119.2.254
- Lechermeier, J., and Fassnacht, M. (2018). How do performance feedback characteristics influence recipients' reactions? A state-of-the-art review on feedback source, timing, and valence effects. *Manag. Rev. Q.* 68, 145–193. doi: 10.1007/s11301-018-0136-8
- Lee, M. K. (2018). Understanding perception of algorithmic decisions: Fairness, trust, and emotion in response to algorithmic management. *Big Data Soc* 5:2053951718756684. doi: 10.1177/2053951718756684
- Logg, J. M., Minson, J. A., and Moore, D. A. (2019). Algorithm appreciation: People prefer algorithmic to human judgment. *Organ. Behav. Hum. Decis. Process.* 151, 90–103. doi: 10.1016/j.obhdp.2018.12.005
- London, M. (2003). *Job Feedback: Giving, Seeking, and Using Feedback for Performance Improvement*, 2nd Edn. New York, NY: Psychology Press. doi: 10.4324/9781410608871
- Lüdtke, D. (2025). *sjPlot: Data Visualization for Statistics in Social Science*. Vienna: Comprehensive R Archive Network. doi: 10.32614/CRAN.package.sjPlot
- Marót, J. L., Szabó, Z. P., Török, L., and Orosz, G. (2024). Hogyan mondjam el neked, hogy ne vegyem el a kedvedet? *Magyar Pszichológiai Szemle* 79, 531–551. doi: 10.1556/0016.2024.00091
- Mendiratta, H., and Singh, R. M. (2026). When feedback hurts: Influence of feedback on rumination and performance. *Soc. Psychol. Educ.* 29:22. doi: 10.1007/s11218-026-10173-6
- Meyer, J., Jansen, T., Schiller, R., Liebenow, L. W., Steinbach, M., Horbach, A., et al. (2024). Using LLMs to bring evidence-based feedback into the classroom: AI-generated feedback increases secondary students' text revision, motivation, and positive emotions. *Com. Edu. Artif. Int.* 6:100199. doi: 10.1016/j.caeai.2023.100199

- Nazaretsky, T., Mejia-Domenzain, P., Swamy, V., Frej, J., and Käser, T. (2026). Who gives feedback matters: Student biases towards human and ai-generated formative feedback. *J. Com. Assisted Learn* 42:e70153. doi: 10.1111/jcal.70153
- Ou, M., Zheng, H., Zeng, Y., and Hansen, P. (2026). Trust it or not: Understanding users' motivations and strategies for assessing the credibility of AI-generated information. *New Media and Soc.* 28, 507–529. doi: 10.1177/14614448241293154
- Park, J., Woo, S. E., and Kim, J. (2024). Attitudes towards artificial intelligence at work: Scale development and validation. *J. Occup. Organ. Psychol.* 97, 920–951. doi: 10.1111/joop.12502
- Pei, J., Wang, H., Peng, Q., and Liu, S. (2024). Saving face: Leveraging artificial intelligence-based negative feedback to enhance employee job performance. *Hum. Manag.* 63, 775–790. doi: 10.1002/hrm.22226
- Perloff, R. M. (2003). *The Dynamics of Persuasion: Communication and Attitudes in the 21st Century*, 2nd Edn. Mahwah, NJ: Lawrence Erlbaum Associates.
- Qin, S., Jia, N., Luo, X., Liao, C., and Huang, Z. (2023). Perceived fairness of human managers compared with artificial intelligence in employee performance evaluation. *J. Manag. Inform. Syst.* 40, 1039–1070. doi: 10.1080/07421222.2023.2267316
- R Core Team (2026). *R: A Language and Environment for Statistical Computing*. Vienna: R Foundation for Statistical Computing. doi: 10.32614/R.manuals
- Reddy Yanamala, K. K. (2023). Transparency privacy, and accountability in ai-enhanced hr processes. *J. Adv. Comp. Syst.* 3, 10–18. doi: 10.69987/JACS.2023.30302
- Schilke, O., and Reimann, M. (2025). The transparency dilemma: How AI disclosure erodes trust. *Organ. Behav. Hum. Dec. Proc.* 188:104405. doi: 10.1016/j.obhdp.2025.104405
- Shin, D. (2022). How do people judge the credibility of algorithmic sources? *AI Soc.* 37, 81–96. doi: 10.1007/s00146-021-01158-4
- Sleiman, A. A., Sigurjonsdottir, S., Elnes, A., Gage, N. A., and Gravina, N. E. (2020). A quantitative review of performance feedback in organizational settings (1998–2018). *J. Organ. Behav. Manag.* 40, 303–332. doi: 10.1080/01608061.2020.1823300
- Son, S., and Kim, D.-Y. (2016). The role of perceived feedback sources' learning-goal orientation on feedback acceptance and employees' creativity. *J. Lead. Organ. Stud.* 23, 82–95. doi: 10.1177/1548051815613732
- Steelman, L. A., and Rutkowski, K. A. (2004). Moderators of employee reactions to negative feedback. *J. Manag. Psychol.* 19, 6–18. doi: 10.1108/02683940410520637
- Tambe, P., Cappelli, P., and Yakubovich, V. (2019). Artificial Intelligence in human resources management: Challenges and a path forward. *California Manag. Rev.* 61, 15–42. doi: 10.1177/0008125619867910
- Tong, S., Jia, N., Luo, X., and Fang, Z. (2021). The Janus face of artificial intelligence feedback: Deployment versus disclosure effects on employee performance. *Strat. Manag. J.* 42, 1600–1631. doi: 10.1002/smj.3322
- Tossell, C. C., Tenhundfeld, N. L., Momen, A., Cooley, K., and De Visser, E. J. (2024). Student perceptions of ChatGPT use in a college essay assignment: Implications for learning, grading, and trust in artificial intelligence. *IEEE Trans. Learn. Technol.* 17, 1069–1081. doi: 10.1109/TLT.2024.3355015
- Trunk, A., Birkel, H., and Hartmann, E. (2020). On the current state of combining human and artificial intelligence for strategic organizational decision making. *Bus. Res.* 13, 875–919. doi: 10.1007/s40685-020-00133-x
- van de Ridder, J. M. M., Berk, F. C. J., Stokking, K. M., and Ten Cate, O. T. J. (2015). Feedback providers' credibility impacts students' satisfaction with feedback and delayed performance. *Med Teach.* 37, 767–774. doi: 10.3109/0142159X.2014.970617
- Wang, Y., Jiao, Y., and Wang, Q. (2026). Who's the better mentor? How AI vs human supervisor developmental feedback influences feedback acceptance. *Bus. Proc. Manag. J.* 32, 861–886. doi: 10.1108/BPMJ-12-2024-1217
- Wang, Y.-Y., and Wang, Y.-S. (2022). Development and validation of an artificial intelligence anxiety scale: An initial application in predicting motivated learning behavior. *Int. Learn. Env.* 30, 619–634. doi: 10.1080/10494820.2019.1674887
- Wickham, H. (2016). *ggplot2: Elegant Graphics for Data Analysis*. New York, NY: Springer-Verlag.
- Wickham, H., Averick, M., Bryan, J., Chang, W., McGowan, L., François, R., et al. (2019). Welcome to the tidyverse. *J. Open Source Software* 4:1686. doi: 10.21105/joss.01686
- Zawiah, M., Al-Ashwal, F. Y., Gharaibeh, L., Abu Farha, R., Alzoubi, K. H., Abu Hammour, K., et al. (2023). ChatGPT and clinical training: perception. *J. Multidiscip. Healthc.* 16, 4099–4110. doi: 10.2147/JMDH.S439223
- Zhang, C., Shao, Y., Yuan, Y., and Shen, W. (2025). Artificial intelligence reshapes creativity: A multidimensional evaluation. *Psych J.* 14, 831–840. doi: 10.1002/pchj.70042
- Zhao, B. (2011). Learning from errors: The role of context, emotion, and personality. *J. Organiz. Behav.* 32, 435–463. doi: 10.1002/job.696