

Volume 2	Issue 1	Apr – Jun 2026
Received	Accepted	Published
06 April 2026	12 June 2026	15 June 2026

When AI feels unfair: The role of emotional depletion and AI literacy in shaping workplace deviant behavior

DOI: [10.66578/btis.v2i2.29](https://doi.org/10.66578/btis.v2i2.29)

Shiva Roosta Rafiee

Department of Business Management, Kharazmi University, Tehran, Iran, 1599964515

Corresponding Author: Shiva Roosta Rafiee **Email:** shiva.roostarafiee@khu.ac.ir

ABSTRACT

This study examines how perceived algorithmic injustice influences employees' workplace deviant behavior. Drawing on Affective Events Theory, the study further investigates the mediating role of emotional depletion and the moderating role of AI literacy in shaping this relationship. Data were collected using a time-lagged survey design from employees working in knowledge-intensive organizations in major Canadian cities. A total of 398 questionnaires were distributed, resulting in 362 initial responses. After matching responses across two waves and removing unmatched cases and outliers, the final sample consisted of 335 employees. Structural equation modeling was employed to test the hypothesized relationships, including mediation and moderation effects. The results indicate that perceived algorithmic injustice significantly increases workplace deviant behavior both directly and indirectly through emotional depletion. Emotional depletion partially mediates this relationship, highlighting emotional strain as a key mechanism linking unfair AI-driven decisions to negative employee behaviors. Furthermore, AI literacy moderates the relationship between algorithmic injustice and emotional depletion, such that the relationship is weaker among employees with higher levels of AI literacy. The findings underscore the importance of fair and transparent AI systems in organizations. Managers should invest in enhancing employees' AI literacy through targeted training programs to reduce emotional strain and mitigate deviant behaviors associated with perceived algorithmic injustice. This study extends the literature on workplace deviance by introducing algorithmic injustice as a technology-driven source of perceived unfairness in AI-enabled work environments. It contributes by identifying emotional depletion as a key mediating mechanism and demonstrating the moderating role of AI literacy as a critical individual capability. By integrating both affective and capability-based perspectives, the study provides a more comprehensive explanation of how employees respond to AI-driven decisions.

Keywords: Algorithmic injustice; Emotional depletion; Workplace deviant behavior; AI literacy; Affective Events Theory; Artificial intelligence



INTRODUCTION

Artificial intelligence (AI) is increasingly embedded in organizational decision-making processes, shaping outcomes related to recruitment (United Nations Educational and Organization, 2021), performance evaluation (Wamba, 2022) and resource allocation (Hemphill, 2024). While these systems promise efficiency and objectivity, growing evidence suggests that employees do not always perceive AI-driven decisions as fair (Jácome de Moura Jr et al., 2024). When algorithmic systems are viewed as opaque, biased or difficult to understand, employees may develop perceptions of algorithmic injustice, which can have important implications for their attitudes and behaviors at work (Weber et al., 2025).

Perceptions of fairness have long been recognized as a critical determinant of employee behavior. Prior research shows that when employees perceive unfair treatment, they are more likely to experience negative emotions and engage in behaviors that violate organizational norms (Anderies, 2015, Yadav et al., 2019). However, much of this literature has focused on human decision-makers, with limited attention to AI-driven contexts. As organizations increasingly rely on algorithmic systems, it becomes essential to understand how employees respond to perceived unfairness arising from these technologies.

Drawing on Affective Events Theory (AET), which posits that workplace events trigger emotional reactions that subsequently influence behavior (Podsakoff et al., 2003), this study argues that algorithmic injustice represents a negative workplace event that can deplete employees' emotional resources. Specifically, when employees perceive AI-based decisions as unfair, they may experience emotional depletion, a state characterized by reduced emotional energy and psychological strain (Wang and Zhang, 2025a). This emotional state, in turn, can increase the likelihood of workplace deviant behavior, as employees may seek to restore balance or cope with perceived mistreatment.

Despite the growing relevance of AI in organizations, there remains a lack of understanding regarding the emotional mechanisms through which algorithmic injustice influences employee behavior (Taha Kandil, 2025). While prior studies have identified various emotional responses to workplace unfairness, the role of emotional depletion as a mediating mechanism in AI-driven environments remains underexplored. Addressing this gap is important, as it provides insight into how and why employees react negatively to algorithmic decisions, beyond purely cognitive evaluations (Wamba-Taguimdje et al., 2020, Buhmann and Fieseler, 2021).

In addition to emotional processes, individual capabilities may play a crucial role in shaping employee responses to AI systems (Raisch and Krakowski, 2021). In this regard, AI literacy defined as an individual's ability to understand, interpret and effectively use AI-based systems may serve as an important boundary condition (United Nations Educational and Organization, 2021). Employees with higher levels of AI literacy are more likely to comprehend how algorithmic decisions are made, which may reduce uncertainty and perceived unfairness (Jácome de Moura Jr et al., 2024). Consequently, they may be better equipped to regulate their emotional responses and experience lower levels of emotional depletion when confronted with algorithmic injustice. Prior research suggests that individual capabilities and coping mechanisms can buffer the negative effects of workplace stressors (Ravenscroft et al., 2012, Wynen et al., 2025), yet their role in AI-related contexts remains insufficiently examined.

This research broadens existing knowledge regarding to the literature in several important ways. First, it extends research on workplace fairness by introducing algorithmic injustice as a contemporary and technology-driven source of perceived unfairness. Second, it advances theoretical understanding by identifying emotional depletion as a key mechanism linking algorithmic injustice to workplace deviant behavior, thereby responding to calls for greater attention to emotional processes in organizational research. Third, it highlights the moderating role of AI literacy, demonstrating how employee capabilities can mitigate the negative consequences of perceived unfairness in AI-driven environments. In doing so, the study integrates both affective and capability-based perspectives to provide a more comprehensive understanding of employee behavior.

From a practical standpoint, the findings offer important implications for organizations seeking to implement AI systems effectively. As AI becomes more prevalent, organizations must not only ensure the technical accuracy of these systems but also address employees' perceptions of fairness and transparency. Moreover, investing in AI literacy through training and development programs can help employees better understand algorithmic processes, thereby reducing emotional strain and minimizing deviant behaviors. These insights are particularly relevant for organizations aiming to balance technological advancement with employee well-being and ethical considerations. Overall, this study develops and tests a model in which algorithmic injustice influences workplace deviant behavior through emotional depletion, with AI literacy acting as a moderating boundary condition. By examining both emotional and capability-based pathways, the study provides a timely and relevant perspective on employee behavior in the era of AI-driven decision-making (see Figure 1).

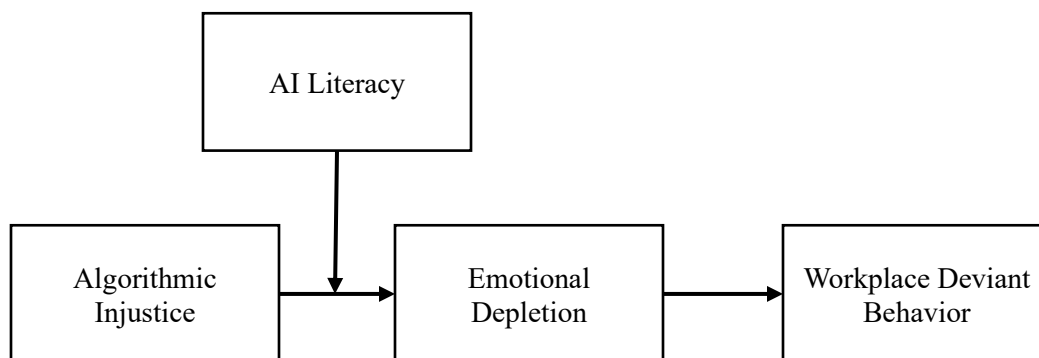


Figure 1 Conceptual model

LITERATURE REVIEW AND HYPOTHESES DEVELOPMENT

Algorithmic injustice and workplace deviant behavior

Organizational fairness has long been recognized as a fundamental driver of employee attitudes and behaviors. Employees evaluate fairness by comparing their inputs and outcomes, and when this balance is disrupted, negative reactions are likely to occur (Michel and Hargis, 2017). While traditional research has focused on human decision-makers, the increasing use of artificial

intelligence (AI) systems has introduced a new dimension of fairness perceptions algorithmic injustice. Algorithmic injustice refers to employees' perceptions that AI-driven decisions are biased, opaque or unfairly distributed.

Unlike traditional organizational injustice arising from human supervisors or managers, algorithmic injustice emerges from automated, data-driven, and often opaque decision-making systems. Employees may perceive algorithmic decisions differently because AI systems typically lack interpersonal explanation, emotional sensitivity, and opportunities for direct appeal or clarification. While some scholars argue that algorithmic systems can improve objectivity and reduce human bias through standardized decision processes, others contend that algorithmic opacity and embedded data biases may intensify perceptions of unfairness and discrimination. These contrasting perspectives suggest that employee reactions toward AI-driven decisions may differ from traditional human-based injustice contexts, thereby requiring separate theoretical examination within AI-enabled workplaces.

Recent studies suggest that employees often struggle to understand how algorithmic systems function, leading to uncertainty and skepticism regarding their fairness (Han et al., 2016b, Yadav et al., 2019). When AI systems are perceived as lacking transparency or accountability, employees may attribute unfair outcomes to the organization itself (Buhmann and Fieseler, 2021, Gupta and Khan, 2024). This perception is particularly critical because, according to the agent-system model of justice, individuals tend to direct their reactions toward the perceived source of unfair treatment (Han et al., 2016a, Cucino et al., 2021, Hanafy et al., 2025).

However, prior research presents conflicting perspectives regarding the fairness of AI-driven decision-making systems. Some scholars argue that algorithmic systems enhance consistency and reduce subjective human bias by relying on standardized data-driven processes. In contrast, other studies suggest that algorithmic systems may intensify perceptions of unfairness because employees often perceive them as opaque, difficult to challenge, and lacking interpersonal accountability. These contrasting perspectives indicate that employee responses toward AI-driven decisions remain psychologically complex and context-dependent, highlighting the need to better understand the emotional and behavioral consequences of perceived algorithmic injustice within organizations.

From the perspective of Affective Events Theory (AET), algorithmic injustice can therefore be conceptualized as a distinct technology-enabled workplace event that triggers emotional and behavioral responses (Podsakoff et al., 2003). Employees who perceive AI-driven decisions as unfair may experience frustration and dissatisfaction (Mikalef et al., 2023), which can manifest in behaviors that violate organizational norms (Agarwal et al., 2024). These behaviors, referred to as workplace deviant behavior, include actions such as reduced effort, withdrawal and intentional rule-breaking (Wang et al., 2024).

Prior research generally suggests that perceptions of injustice are positively associated with deviant behaviors, as employees attempt to restore equity or retaliate against perceived mistreatment (Anderies, 2015, Yadav et al., 2019). Extending this logic to AI-driven contexts, it is reasonable to expect that algorithmic injustice will similarly lead to deviant workplace behaviors. Therefore, the following hypothesis is proposed:

H1. Algorithmic injustice positively influences workplace deviant behavior.

**The mediating role of emotional depletion**

Although the direct relationship between perceived injustice and deviant behavior has been well established, less attention has been given to the underlying emotional mechanisms that explain this relationship, particularly in AI-driven environments. According to AET, workplace events first influence employees' emotional states, which then shape their behavioral responses (Podsakoff et al., 2003). In this context, emotional depletion represents a key psychological mechanism through which algorithmic injustice may influence workplace deviant behavior.

Emotional depletion refers to a state of reduced emotional energy and psychological fatigue resulting from prolonged exposure to stressors (Gong, 2025). When employees perceive algorithmic decisions as unfair, they may experience cognitive dissonance, frustration and a sense of helplessness. Coping with these negative experiences requires emotional resources, which can gradually become depleted over time.

Prior studies have demonstrated that perceived injustice significantly contributes to emotional strain, including burnout and emotional exhaustion (Yadav et al., 2019, Michel and Hargis, 2017). In turn, emotionally depleted individuals are more likely to engage in deviant behaviors as a coping mechanism or to restore perceived balance (Wang et al., 2015, Hosen et al., 2023). Emotional depletion reduces self-regulation capacity, making it more difficult for individuals to adhere to organizational norms and increasing the likelihood of counterproductive behaviors.

Furthermore, emotional depletion can be understood as a form of resource loss within the framework of conservation of resources theory, where individuals strive to protect their remaining resources when faced with stress (Carter et al., 2007). In such situations, engaging in deviant behavior may serve as a means of conserving or replenishing depleted emotional resources. Based on these arguments, emotional depletion is expected to mediate the relationship between algorithmic injustice and workplace deviant behavior. Accordingly, the following hypothesis is proposed:

H2. Emotional depletion mediates the relationship between algorithmic injustice and workplace deviant behavior.

The moderating role of AI literacy

While algorithmic injustice may lead to emotional depletion, not all employees respond to such perceptions in the same way (Stone et al., 2020, De Marcellis-Warin et al., 2022). Individual differences play an important role in shaping how employees interpret and react to workplace events (Dey et al., 2024, Gupta and Khan, 2024). In this regard, AI literacy is proposed as a critical moderating factor that can influence the strength of the relationship between algorithmic injustice and emotional depletion.

AI literacy refers to an individual's ability to understand, interpret and effectively interact with AI systems (Wang and Zhang, 2025b). Employees with higher levels of AI literacy are more likely to comprehend how algorithmic decisions are generated, including the limitations and potential biases of these systems. This understanding can reduce uncertainty and mitigate perceptions of unfairness, thereby weakening the emotional impact of algorithmic injustice.

According to AET, individual characteristics influence how employees appraise workplace events and regulate their emotional responses (Podsakoff et al., 2003). Employees with higher AI literacy may engage in more adaptive cognitive appraisals, interpreting algorithmic outcomes as system

limitations rather than intentional unfairness. This perspective can help reduce emotional strain and prevent the escalation of negative reactions.

Empirical research supports the idea that individual capabilities can buffer the effects of workplace stressors. For example, emotional intelligence and coping skills have been shown to reduce the impact of stress on emotional exhaustion (Ravenscroft et al., 2012, Chatterjee et al., 2023). Similarly, AI literacy can function as a cognitive resource that enables employees to better manage the challenges associated with AI-driven decision-making (Jiang et al., 2025). Therefore, it is expected that AI literacy will weaken the positive relationship between algorithmic injustice and emotional depletion. Specifically, employees with higher AI literacy will experience lower levels of emotional depletion when confronted with algorithmic injustice compared to those with lower AI literacy. Based on these arguments, the following hypothesis is proposed:

H3. AI literacy moderates the relationship between algorithmic injustice and emotional depletion such that the relationship is weaker for employees with higher levels of AI literacy.

METHOD

Participants and procedure

Data for this study were collected from employees working in knowledge-intensive organizations located in major Canadian cities (e.g., Toronto, Vancouver and Montreal), where AI-enabled systems are increasingly integrated into workplace decision-making processes (e.g., performance evaluation, task allocation and monitoring). This context is particularly appropriate for the present study, as such environments rely heavily on AI-driven systems, increasing the likelihood that employees develop perceptions of algorithmic injustice.

A time-lagged survey design was employed to reduce common method bias and strengthen causal inference (Podsakoff et al., 2003). Data were collected in two waves separated by a time lag of approximately three weeks. In the first wave, respondents provided information on demographic characteristics, algorithmic injustice and emotional depletion. In the second wave, the same participants were invited to report on AI literacy and workplace deviant behavior.

A total of 398 questionnaires were distributed using a convenience sampling approach, as a comprehensive sampling frame of employees working with AI-based systems was not accessible. Participants were recruited through professional networks and online platforms, including email invitations and professional networking sites such as LinkedIn, enabling access to individuals actively engaged in AI-supported work settings. Of the distributed questionnaires, 362 responses were received in the first wave, yielding an initial response rate of approximately 90.9%.

Responses from the two waves were matched using unique identifiers to ensure consistency. A total of 350 responses were successfully matched across both waves, while 12 responses could not be matched and were excluded from further analysis. Subsequently, multivariate outliers were identified using the Mahalanobis distance criterion (Kline, 2015), resulting in the removal of 15 additional responses. This process yielded a final sample of 335 usable responses, corresponding to an effective response rate of approximately 84.2%. No significant missing data were detected.

The sample size is considered adequate for structural equation modeling, as it exceeds the recommended minimum threshold of 200 observations for models of moderate complexity.

Furthermore, a sample size of 335 provides sufficient statistical power to detect medium effect sizes with a power level above 0.80, ensuring the robustness of the findings.

Regarding the demographic profile, 204 respondents (61%) were male and 131 (39%) were female. In terms of age distribution, 107 respondents (32%) were under 30 years, 137 (41%) were between 31 and 40 years and the remaining respondents were above 40 years. Most participants held at least a bachelor's degree, with 147 respondents (44%) possessing a master's degree or higher. In terms of work experience, 121 respondents (36%) had less than five years of experience, while 114 (34%) had between five and ten years of experience. Overall, the sample reflects a diverse and balanced representation of employees across different demographic groups.

Measures

All constructs were measured using previously validated scales adapted to fit the context of AI-driven organizational decision-making. Responses were recorded on a five-point Likert scale ranging from 1 = strongly disagree to 5 = strongly agree. All scales were adapted from previously validated measures and contextually modified to reflect AI-enabled organizational decision-making environments.

Minor wording modifications were made to existing measurement scales to align the items with AI-driven workplace decision-making contexts. The adaptation process focused on preserving the original conceptual meaning of each construct while contextualizing references toward AI-enabled organizational systems. To ensure clarity and contextual appropriateness, the adapted items were reviewed by academic researchers familiar with organizational behavior and AI-related workplace studies prior to data collection.

Algorithmic injustice: Algorithmic injustice was measured using a five-item scale adapted from organizational justice literature (Michel and Hargis, 2017). The wording of the items was modified to capture perceptions of fairness in AI-driven decision-making. This scale has been widely used and validated in organizational research (Michel and Hargis, 2017). Negatively worded items were reverse-coded to ensure consistency in interpretation. A sample item includes: "Decisions made by AI systems in my organization are fair." All items were reverse-coded so that higher scores indicate higher levels of perceived algorithmic injustice. The Cronbach's alpha for this scale was 0.88, indicating good internal consistency.

Emotional depletion: Emotional depletion was assessed using a five-item scale derived from the Maslach Burnout Inventory (Gong, 2025), which has been extensively validated in prior studies (Zeng et al., 2025). A sample item includes: "*I feel emotionally drained from my work.*" The reliability coefficient (Cronbach's alpha) for this scale was 0.84.

AI literacy: AI literacy was measured using a three-item scale adapted from research on AI competencies (Wang and Zhang, 2025b). This construct reflects employees' ability to understand, interpret and interact with AI systems. A sample item includes: "*I understand how AI systems make decisions in my workplace.*" The Cronbach's alpha for this scale was 0.81.

Workplace deviant behavior: Workplace deviant behavior was measured using an eight-item scale developed by Wang et al. (2024), which has been widely validated in organizational behavior research. This scale captures behaviors that violate organizational norms and harm the organization. A sample item includes: *“I intentionally come to work late without permission.”* The reliability coefficient for this scale was **0.90**. Self-reported measures were considered appropriate because workplace deviant behavior is often covert in nature and may not be accurately captured through supervisor or peer ratings (Wang et al., 2024, Zhao et al., 2025).

Control variables

Consistent with prior research, age, gender, education level and work experience were included as control variables. These variables were controlled because demographic characteristics have been shown to influence perceptions of fairness, emotional responses and deviant behaviors (Zhao et al., 2025, Anderies, 2015).

Data analysis strategy

The data were analyzed using structural equation modeling (SEM) with AMOS version 24, following a two-step approach involving measurement and structural model assessment (Shmueli et al., 2019). In the first stage, confirmatory factor analysis (CFA) was conducted to evaluate the reliability and validity of the measurement model, including assessments of factor loadings, composite reliability and average variance extracted. In the second stage, the structural model was tested to examine the hypothesized relationships among the constructs. To test the mediating effect, a bootstrapping procedure with 5,000 resamples and a 95% confidence interval was employed, as recommended for robust estimation of indirect effects (Preacher and Hayes, 2008). This approach allows for more accurate estimation of mediation effects without relying on normal distribution assumptions.

SEM was employed to examine the measurement model, direct relationships, and mediation effects because it enables simultaneous estimation of multiple relationships while accounting for measurement error. However, hierarchical regression analysis was used to test the moderating effect because interaction-based moderation analysis through mean-centered variables and interaction terms is widely recommended and provides clearer interpretation of conditional effects. The combined use of SEM and hierarchical regression is consistent with prior organizational behavior research examining mediation and moderation within the same study design.

The moderating effect of AI literacy was examined using hierarchical regression analysis. All predictor variables were mean-centered prior to analysis to reduce multicollinearity, after which an interaction term between algorithmic injustice and AI literacy was computed and entered into the regression model. The significance of the interaction term and the change in explained variance (ΔR^2) were used to assess the moderating effect. This approach is consistent with prior organizational behavior research examining interaction-based moderation effects.

Common method bias

Given the use of self-reported data, both procedural and statistical remedies were applied to address common method bias. Procedurally, the time-lagged design and separation of measurement waves (independent and mediating variables in Wave 1; moderating and dependent

variables in Wave 2) helped reduce respondents' ability to infer relationships among constructs (Podsakoff et al., 2003). In addition, questionnaire items were randomized to further minimize response bias.

Statistically, Harman's single-factor test was conducted, and the first factor accounted for 38.42% of the total variance, which is below the recommended threshold of 50%, indicating that common method bias is not a serious concern. In addition, a common latent factor approach was applied in the measurement model, and the results showed no substantial change in factor loadings. Additionally, variance inflation factors (VIF) were examined, and all values were below the threshold of 3.3, indicating that common method bias is unlikely to bias the results.

Ethical considerations

Participation in the study was voluntary, and respondents were assured of the confidentiality and anonymity of their responses. No personally identifiable information was collected, and participants were informed that their responses would be used solely for academic research purposes. These procedures were followed to ensure compliance with standard ethical research practices. This study was conducted in accordance with institutional research ethics guidelines, and ethical approval was obtained prior to data collection.

RESULT

Preliminary analysis

A total of 398 questionnaires were distributed among employees working in knowledge-intensive organizations across major Canadian cities (e.g., Toronto, Vancouver and Montreal). Of these, 362 responses were received, yielding an initial response rate of 90.9%. After matching responses across the two survey waves, 350 responses were retained, while 12 unmatched responses were excluded. Subsequently, 15 multivariate outliers were removed using the Mahalanobis distance criterion (Kline, 2015), resulting in a final sample of 335 usable responses and an effective response rate of 84.2%.

Prior to hypothesis testing, the dataset was screened for missing values, outliers and normality. No significant missing data were detected. Skewness and kurtosis values ranged between -1.21 and 1.34 , indicating that the data met the assumptions required for multivariate analysis (Shmueli et al., 2019).

Measurement model

A confirmatory factor analysis (CFA) was conducted using AMOS version 24 to evaluate the measurement model. The hypothesized four-factor model demonstrated an acceptable fit to the data ($\chi^2/df = 2.08$, CFI = 0.948, GFI = 0.912, RMSEA = 0.058 and RMR = 0.066), exceeding recommended thresholds (Shmueli et al., 2019, Kline, 2015). All standardized factor loadings were significant and ranged from 0.61 to 0.86, exceeding the recommended threshold of 0.50. Composite reliability (CR) values ranged from 0.78 to 0.89, and average variance extracted (AVE) values ranged from 0.50 to 0.66, supporting convergent validity. Discriminant validity was assessed using the Fornell–Larcker criterion. As shown in Table 1, the square root of AVE for each construct exceeds its correlations with other constructs

Table 1. Discriminant validity (Fornell–Larcker criterion)

Variables	1	2	3	4
1. Algorithmic injustice	0.801			
2. Emotional depletion	0.27	0.729		
3. Workplace deviant behavior	0.44	0.29	0.752	
4. AI literacy	−0.16	−0.23	−0.19	0.768

Note: Diagonal values represent the square root of AVE.

In addition to the Fornell–Larcker criterion, discriminant validity was further assessed using the heterotrait–monotrait ratio (HTMT). As shown in Table 1A, all HTMT values were below the recommended threshold of 0.85, providing additional evidence of discriminant validity among the study constructs.

Table 1A. HTMT values

Variables	1	2	3	4
1. Algorithmic injustice	—			
2. Emotional depletion	0.48			
3. Workplace deviant behavior	0.54	0.43		
4. AI literacy	0.29	0.34	0.31	

Note: All HTMT values are below the recommended threshold of 0.85.

Descriptive statistics and correlations

Descriptive statistics and Pearson correlations for all study variables are presented in Table 2. Algorithmic injustice is positively correlated with emotional depletion ($r = 0.27$, $p < 0.01$) and workplace deviant behavior ($r = 0.44$, $p < 0.01$). Emotional depletion is also positively associated with workplace deviant behavior ($r = 0.29$, $p < 0.01$).

AI literacy is negatively correlated with algorithmic injustice ($r = -0.16$, $p < 0.05$), emotional depletion ($r = -0.23$, $p < 0.01$) and workplace deviant behavior ($r = -0.19$, $p < 0.01$), indicating that higher AI literacy is associated with lower perceived injustice, reduced emotional strain and fewer deviant behaviors.

Table 2. Descriptive statistics and correlations

Variables	1	2	3	4	Mean	SD
1. Algorithmic injustice	1				3.08	0.94
2. Emotional depletion	0.27**	1			3.29	0.88
3. Workplace deviant behavior	0.44**	0.29**	1		2.88	0.97
4. AI literacy	−0.16*	−0.23**	−0.19**	1	3.15	0.83

Note: * $p < 0.05$, ** $p < 0.01$.

Structural model results

The structural model was evaluated using structural equation modeling with bootstrapping (5,000 resamples and a 95% confidence interval). The results are presented in Table 3. Algorithmic injustice has a significant positive effect on workplace deviant behavior ($\beta = 0.46$, $p < 0.001$), supporting H1. In addition, algorithmic injustice significantly predicts emotional depletion ($\beta = 0.39$, $p = 0.003$). Emotional depletion also has a positive effect on workplace deviant behavior ($\beta = 0.24$, $p < 0.001$). Control variables (age, gender, education level and work experience) were included in the model. The results indicate that none of these variables had a statistically significant effect on workplace deviant behavior, as all coefficients were small and non-significant (β values ranged between -0.07 and 0.05 , all $p > 0.05$). The findings demonstrate that the hypothesized relationships are robust beyond demographic influences.

Table 3. Structural model results

Hypothesis	Path	β	SE	CR	p	LLCI	ULCI
H1	Algorithmic injustice → Workplace deviant behavior	0.46	0.043	10.72	<0.001	0.372	0.541
—	Algorithmic injustice → Emotional depletion	0.39	0.131	2.98	0.003	0.142	0.641
—	Emotional depletion → Workplace deviant behavior	0.24	0.058	4.12	<0.001	0.131	0.347

Mediation analysis

The mediating role of emotional depletion was examined using a bootstrapping procedure with 5,000 resamples and a 95% confidence interval. The results are presented in Table 4. The direct effect of algorithmic injustice on workplace deviant behavior remained significant ($\beta = 0.46$, $p < 0.001$). In addition, the indirect effect via emotional depletion was positive and significant ($\beta = 0.09$, $p = 0.010$), with the confidence interval not including zero (LLCI = 0.028, ULCI = 0.167). These findings indicate partial mediation, supporting H2.

Table 4. Mediation results

Path	β	SE	p	LLCI	ULCI
Direct: Algorithmic injustice → Workplace deviant behavior	0.46	0.043	<0.001	0.372	0.541
Indirect: Algorithmic injustice → Emotional depletion → Workplace deviant behavior	0.09	0.036	0.010	0.028	0.167

Moderation analysis

The moderating role of AI literacy was examined using hierarchical regression analysis. The results are presented in Table 5, and the interaction effect is further visualized in Figure 2 to facilitate interpretation. The interaction term between algorithmic injustice and AI literacy was found to be negative and significant ($\beta = -0.29$, $p = 0.020$), indicating that AI literacy weakens the positive relationship between algorithmic injustice and emotional depletion. These results support H3.

Table 5. Moderation results

Variables	β	SE	p	LLCI	ULCI
Algorithmic injustice	0.36	0.118	0.006	0.128	0.592
AI literacy	-0.26	0.109	0.018	-0.474	-0.046
Algorithmic injustice $\times$ AI literacy	-0.29	0.124	0.020	-0.532	-0.041
R ² (Step 1)	0.17				
R ² (Step 2)	0.20				
ΔR^2	0.03				

As shown in Table 5, the interaction term between algorithmic injustice and AI literacy is negative and significant, supporting the moderating hypothesis. To further interpret the interaction effect, a simple slope analysis was conducted and is illustrated in Figure 2. As illustrated in Figure 2, the simple slope analysis shows that the positive relationship between algorithmic injustice and emotional depletion is significantly stronger at low levels of AI literacy and weaker at high levels of AI literacy. Specifically, the slope is steeper for employees with low AI literacy, indicating greater emotional depletion when perceiving algorithmic injustice. In contrast, the slope is flatter for employees with high AI literacy, suggesting a buffering effect. These findings further support the moderating role of AI literacy.

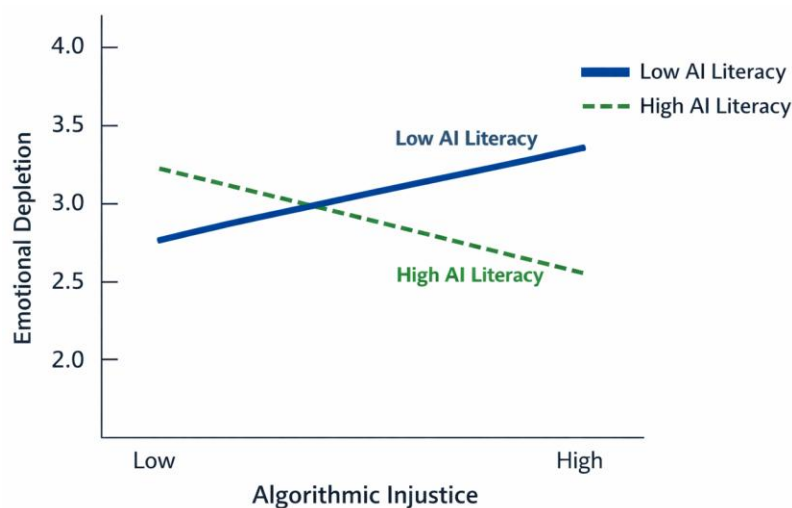


Figure 2. Moderation effect of AI literacy on the relationship between algorithmic injustice and emotional depletion

Source: Authors' own work



DISCUSSION

This study examined the impact of perceived algorithmic injustice on workplace deviant behavior, while uncovering the underlying emotional mechanism and boundary condition shaping this relationship. Drawing on Affective Events Theory (AET), a model incorporating both mediating and moderating mechanisms was developed and tested within the context of AI-driven organizational decision-making.

The findings provide strong support for the proposed framework. The findings demonstrate that perceived algorithmic injustice represents a significant organizational stressor capable of increasing workplace deviant behavior. This finding suggests that employees may interpret opaque or biased AI-driven decisions as organizational mistreatment, thereby increasing the likelihood of retaliatory or counterproductive workplace responses. The mediating role of emotional depletion further highlights the importance of emotional resource loss in explaining how employees react to AI-driven unfairness. This finding reinforces the argument that employees react not only to perceived injustice itself, but also to the psychological strain generated by repeated exposure to algorithmic decision-making.

Third, AI literacy moderates the relationship between algorithmic injustice and emotional depletion, suggesting that employees with stronger AI-related competencies are better equipped to cognitively interpret algorithmic decisions and manage associated emotional strain. This finding highlights AI literacy not merely as a technical skill, but as an important psychological resource that shapes employee adaptation within AI-enabled workplaces. This finding indicates that individuals with greater understanding of AI systems are better able to cognitively process algorithmic decisions, thereby reducing emotional strain. Rather than viewing AI-driven decisions as purely technical processes, the findings demonstrate that employees interpret algorithmic outcomes as emotionally meaningful workplace events capable of shaping subsequent behavioral reactions.

Theoretical implications

This study makes several important contributions to the literature. First, it extends research on organizational justice by introducing algorithmic injustice as a technology-driven form of perceived unfairness. While prior studies have primarily focused on human decision-makers, the present findings demonstrate that fairness perceptions in AI-driven contexts can similarly influence employee behavior. By conceptualizing algorithmic injustice as a distinct organizational stressor, this study responds to recent calls for research examining the human implications of AI adoption. The findings also reinforce the central assumptions of Affective Events Theory by demonstrating that algorithmic injustice functions as a negative affective workplace event capable of triggering emotional resource depletion. Consistent with AET, employees do not merely respond cognitively to AI-driven unfairness; rather, their emotional reactions become central mechanisms shaping subsequent workplace behavior. The findings therefore suggest that AI-enabled organizational systems influence employee conduct indirectly through affective processes, highlighting the importance of emotions in understanding employee responses to algorithmic decision-making. This extends AET into technology-enabled work environments where

organizational events increasingly originate from automated and algorithmic systems rather than traditional interpersonal interactions.

Second, this study advances Affective Events Theory by empirically identifying emotional depletion as a key mediating mechanism linking algorithmic injustice to workplace deviant behavior. Although AET emphasizes the role of emotions, empirical research has often overlooked specific emotional processes in technology-enabled environments. By demonstrating the mediating role of emotional depletion, this study provides a more nuanced understanding of how negative workplace events translate into behavioral outcomes.

Third, the study contributes by integrating a capability-based perspective through the inclusion of AI literacy as a moderating variable. Unlike prior research that emphasizes affective or leadership-based drivers of workplace behavior, this study highlights the role of cognitive capabilities in shaping employee responses to AI-driven decisions. The findings demonstrate that AI literacy functions as a critical personal resource that mitigates emotional strain, thereby extending research on individual differences and coping mechanisms. Unlike prior studies that primarily focus on affective or leadership-based drivers of workplace deviance, this study introduces algorithmic injustice as a technology-driven antecedent and integrates both emotional and capability-based mechanisms within a single framework. By jointly examining emotional and capability-based mechanisms, this study offers a more comprehensive explanation of employee behavior in AI-enabled work environments, bridging two streams of research that are often examined in isolation.

Practical implications

The findings of this study offer several important implications for managers and organizations implementing AI systems. First, organizations should recognize that the adoption of AI is not solely a technical decision but also a behavioral and psychological challenge. Perceptions of algorithmic injustice can trigger emotional strain and deviant behaviors, potentially undermining both employee well-being and organizational performance.

Second, organizations should prioritize transparency and fairness in AI-driven decision-making processes. Providing clear explanations of how AI systems operate and ensuring accountability in decision outcomes can reduce perceptions of unfairness and minimize negative employee reactions. Such governance-oriented practices can help organizations balance technological efficiency with employee well-being and ethical responsibility. The findings also raise important implications regarding AI governance and organizational accountability. As organizations increasingly rely on AI-enabled systems, ethical governance mechanisms become critical for maintaining employee trust and minimizing negative behavioral consequences. Organizations should ensure that AI-driven decisions remain transparent, explainable, and subject to appropriate human oversight. The implementation of explainable AI systems may help employees better understand how algorithmic outcomes are generated, thereby reducing perceptions of opacity and unfairness. In addition, organizations should establish ethical review procedures, fairness auditing mechanisms, and employee feedback channels to ensure that AI systems operate in a responsible and accountable manner.

Third, the results highlight AI literacy as a strategic organizational capability. Investing in training and development programs that enhance employees' understanding of AI systems can reduce emotional depletion and mitigate deviant behaviors. By improving AI literacy, organizations

enable employees to better interpret algorithmic decisions, thereby reducing uncertainty and emotional strain.

Finally, managers should actively monitor employees' emotional well-being in AI-enabled work environments. Providing support mechanisms, such as stress management programs or employee assistance initiatives, can help prevent emotional depletion from escalating into counterproductive workplace behaviors.

Limitations and future research

Despite its contributions, this study has several limitations that should be acknowledged. First, the use of a convenience sampling approach and the focus on knowledge-intensive organizations may limit the generalizability of the findings. Future research could examine the proposed relationships across different industries and cultural contexts to enhance external validity.

Second, although the time-lagged design reduces concerns related to common method bias, the study relies on self-reported data. Future studies could incorporate multi-source data, such as supervisor or peer evaluations, to provide a more robust assessment of workplace deviant behavior. Third, this study focused on emotional depletion as the primary mediating mechanism. However, other psychological processes, such as trust in AI systems, perceived control or cognitive appraisal, may also play important roles. Future research could explore additional mediators to provide a more comprehensive understanding of the underlying mechanisms.

Fourth, while AI literacy was examined as a moderating factor, other contextual variables may also influence employee responses to algorithmic decision-making. Factors such as ethical climate, leadership style or organizational transparency may shape how employees interpret and react to AI-driven decisions. Future research could extend the model by incorporating these variables.

REFERENCES

- AGARWAL, P., SWAMI, S. & MALHOTRA, S. K. 2024. Artificial intelligence adoption in the post COVID-19 new-normal and role of smart technologies in transforming business: a review. *Journal of Science and Technology Policy Management*, 15, 506-529.
- ANDERIES, J. M. 2015. Understanding the dynamics of sustainable social-ecological systems: human behavior, institutions, and regulatory feedback networks. *Bulletin of mathematical biology*, 77, 259-280.
- BUHMANN, A. & FIESELER, C. 2021. Towards a deliberative framework for responsible innovation in artificial intelligence. *Technology in Society*, 64, 101475.
- CARTER, C. R., KAUFMANN, L. & MICHEL, A. 2007. Behavioral supply management: a taxonomy of judgment and decision-making biases. *International Journal of Physical Distribution & Logistics Management*, 37, 631-669.
- CHATTERJEE, S., CHAUDHURI, R., VRONTIS, D. & GIOVANDO, G. 2023. Digital workplace and organization performance: Moderating role of digital leadership capability. *Journal of Innovation & Knowledge*, 8, 100334.
- CUCINO, V., DEL SARTO, N., DI MININ, A. & PICCALUGA, A. 2021. Empowered or engaged employees? A fuzzy set analysis on knowledge transfer professionals. *Journal of Knowledge Management*, 25, 1081-1104.

- DE MARCELLIS-WARIN, N., MARTY, F., THELISSON, E. & WARIN, T. 2022. Artificial intelligence and consumer manipulations: from consumer's counter algorithms to firm's self-regulation tools. *AI and Ethics*, 2, 259-268.
- DEY, P. K., CHOWDHURY, S., ABADIE, A., VANN YAROSON, E. & SARKAR, S. 2024. Artificial intelligence-driven supply chain resilience in Vietnamese manufacturing small-and medium-sized enterprises. *International Journal of Production Research*, 62, 5417-5456.
- GONG, T. 2025. Algorithmic management and gig workers: engagement, exhaustion and citizenship behavior. *Management Decision*, 1-22.
- GUPTA, Y. & KHAN, F. M. 2024. Role of artificial intelligence in customer engagement: a systematic review and future research directions. *Journal of Modelling in Management*, 19, 1535-1565.
- HAN, S.-H., SEO, G., LI, J. & YOON, S. W. 2016a. The mediating effect of organizational commitment and employee empowerment: How transformational leadership impacts employee knowledge sharing intention. *Human Resource Development International*, 19, 98-115.
- HAN, S. H., SEO, G., YOON, S. W. & YOON, D.-Y. 2016b. Transformational leadership and knowledge sharing: Mediating roles of employee's empowerment, commitment, and citizenship behaviors. *Journal of Workplace Learning*, 28, 130-149.
- HANAFY, H. A., AL-HAJLA, A. H. & ELSHARNOUBY, M. H. 2025. Empowering leadership and employee innovation: unraveling the roles of psychological empowerment and knowledge sharing. *Journal of Humanities and Applied Social Sciences*, 7, 366-391.
- HEMPHILL, T. A. 2024. Copyright protection, artistic imagery, and the adoption of responsible artificial intelligence principles. *Journal of Ethics in Entrepreneurship and Technology*, 4, 2-6.
- HOSEN, M., OGBEIBU, S., LIM, W. M., FERRARIS, A., MUNIM, Z. H. & CHONG, Y.-L. 2023. Knowledge sharing behavior among academics: Insights from theory of planned behavior, perceived trust and organizational climate. *Journal of Knowledge Management*, 27, 1740-1764.
- JÁCOME DE MOURA JR, P., DOS SANTOS JUNIOR, C. D., PORTO-BELLINI, C. G. & DIAS JUNIOR, J. J. L. 2024. The over-concentration of innovation and firm-specific knowledge in the artificial intelligence industry. *Journal of the Knowledge Economy*, 15, 20547-20577.
- JIANG, L., XUAN, Y. & ZHANG, K. 2025. Unlocking innovation potential: the impact of artificial intelligence transformation on enterprise innovation capacity. *European Journal of Innovation Management*, 28, 4112-4131.
- KLINE, R. B. 2015. The mediation myth. *Basic and Applied Social Psychology*, 37, 202-213.
- MICHEL, J. S. & HARGIS, M. B. 2017. What motivates deviant behavior in the workplace? An examination of the mechanisms by which procedural injustice affects deviance. *Motivation and Emotion*, 41, 51-68.
- MIKALEF, P., ISLAM, N., PARIDA, V., SINGH, H. & ALTWAIJRY, N. 2023. Artificial intelligence (AI) competencies for organizational performance: A B2B marketing capabilities perspective. *Journal of Business Research*, 164, 113998.
- PODSAKOFF, P. M., MACKENZIE, S. B., LEE, J.-Y. & PODSAKOFF, N. P. 2003. Common method biases in behavioral research: a critical review of the literature and recommended remedies. *Journal of applied psychology*, 88, 879.
- PREACHER, K. J. & HAYES, A. F. 2008. Asymptotic and resampling strategies for assessing and comparing indirect effects in multiple mediator models. *Behavior research methods*, 40, 879-891.



- RAISCH, S. & KRAKOWSKI, S. 2021. Artificial intelligence and management: The automation–augmentation paradox. *Academy of management review*, 46, 192-210.
- RAVENSCROFT, A., SCHMIDT, A., COOK, J. & BRADLEY, C. 2012. Designing social media for informal learning and knowledge maturing in the digital workplace. *Journal of Computer Assisted Learning*, 28, 235-249.
- SHMUELI, G., SARSTEDT, M., HAIR, J. F., CHEAH, J.-H., TING, H., VAITHILINGAM, S. & RINGLE, C. M. 2019. Predictive model assessment in PLS-SEM: guidelines for using PLSpredict. *European journal of marketing*, 53, 2322-2347.
- STONE, M., ARAVOPOULOU, E., EKINCI, Y., EVANS, G., HOBBS, M., LABIB, A., LAUGHLIN, P., MACHTYNGER, J. & MACHTYNGER, L. 2020. Artificial intelligence (AI) in strategic marketing decision-making: a research agenda. *The Bottom Line*, 33, 183-200.
- TAHA KANDIL, T. 2025. Artificial intelligence (AI) and alleviating supply chain bullwhip effects: social network analysis-based review. *Journal of Global Operations and Strategic Sourcing*, 18, 5-35.
- UNITED NATIONS EDUCATIONAL, S. & ORGANIZATION, C. 2021. Recommendation on the ethics of artificial intelligence. United Nations Educational.
- WAMBA-TAGUIMDJIE, S.-L., FOSSO WAMBA, S., KALA KAMDJOU, J. R. & TCHATCHOUANG WANKO, C. E. 2020. Influence of artificial intelligence (AI) on firm performance: the business value of AI-based transformation projects. *Business process management journal*, 26, 1893-1924.
- WAMBA, S. F. 2022. Impact of artificial intelligence assimilation on firm performance: The mediating effects of organizational agility and customer agility. *International Journal of Information Management*, 67, 102544.
- WANG, H.-K., YEN, Y.-F. & TSENG, J.-F. 2015. Knowledge sharing in knowledge workers: The roles of social exchange theory and the theory of planned behavior. *Innovation*, 17, 450-465.
- WANG, S. & ZHANG, H. 2025a. Enhancing environmental, social, and governance performance through artificial intelligence supply chains in the energy industry: Roles of innovation, collaboration, and proactive sustainability strategy. *Renewable Energy*, 245, 122855.
- WANG, S. & ZHANG, H. 2025b. Generative artificial intelligence and internationalization green innovation: Roles of supply chain innovations and AI regulation for SMEs. *Technology in Society*, 82, 102898.
- WANG, Y., WANG, Z. & LI, J. 2024. Does algorithmic control facilitate platform workers' deviant behavior toward customers? The ego depletion perspective. *Computers in Human Behavior*, 156, 108242.
- WEBER, M., HEIN, A., WEKING, J. & KRCMAR, H. 2025. Orchestration logics for artificial intelligence platforms: From raw data to industry-specific applications. *Information Systems Journal*, 35, 1015-1043.
- WYNEN, J., BOON, J. & VAN DONINCK, D. 2025. Information-sharing in a turbulent Public sector: is more always better? Exploring information overload among civil servants during complex workplace changes. *Public Management Review*, 1-28.
- YADAV, M., CHOUDHARY, S. & JAIN, S. 2019. Transformational leadership and knowledge sharing behavior in freelancers: A moderated mediation model with employee engagement and social support. *Journal of Global Operations and Strategic Sourcing*, 12, 202-224.

- ZENG, B., CHOTIA, V., GHOSH, V. & CHENG, J. 2025. Digital antecedents and mechanisms towards sustainable digital innovation ecosystems: examining the role of circular supply chain resilience. *Technological Forecasting and Social Change*, 218, 124220.
- ZHAO, H., YUAN, B., WEN, M. & ZHANG, T. 2025. Algorithmic ethics in AI-enabled workplace systems for food delivery employees. *The Service Industries Journal*, 1-35.