

Kim, Byung-Jik; Kim, Min-Jik

Article

How artificial intelligence-induced job insecurity shapes knowledge dynamics: The mitigating role of artificial intelligence self-efficacy

Journal of Innovation & Knowledge (JIK)

Provided in Cooperation with:

Elsevier

Suggested Citation: Kim, Byung-Jik; Kim, Min-Jik (2024) : How artificial intelligence-induced job insecurity shapes knowledge dynamics: The mitigating role of artificial intelligence self-efficacy, Journal of Innovation & Knowledge (JIK), ISSN 2444-569X, Elsevier, Amsterdam, Vol. 9, Iss. 4, pp. 1-16,
<https://doi.org/10.1016/j.jik.2024.100590>

This Version is available at:

<https://hdl.handle.net/10419/327492>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.



<https://creativecommons.org/licenses/by-nc-nd/4.0/>



How artificial intelligence-induced job insecurity shapes knowledge dynamics: the mitigating role of artificial intelligence self-efficacy

Byung-Jik Kim^{a,b}, Min-Jik Kim^{c,*}

^a College of Business Administration, University of Ulsan, 93 Daehak-ro, Nam-gu, Ulsan 44610, South Korea

^b Department of Psychology, Yonsei University, 50, Yonsei-ro, Seodaemun-gu 03722, South Korea

^c School of Industrial Management, Korea University of Technology and Education, 1600, Chungjeol-ro, Dongnam-gu, Cheonan-si, Chungcheongnam-do 31253, South Korea

ARTICLE INFO

Article History:

Received 5 August 2024

Accepted 28 September 2024

Available online 14 November 2024

Keywords:

artificial intelligence-induced job insecurity
self-efficacy in artificial intelligence learning
psychological safety
knowledge-hiding behavior
moderated mediation model

M12

M15

O33

JEL codes:

M12 (Personnel Management

Executives

Executive Compensation)

M15 (IT Management)

O33 (Technological Change: Choices and

Consequences

Diffusion Processes)

ABSTRACT

This research examines the intricate relationships between artificial intelligence (AI)-induced job insecurity, psychological safety, knowledge-hiding behavior, and self-efficacy in AI learning within organizational contexts. As AI technologies increasingly permeate the workplace, comprehending their impact on employee behavior and organizational dynamics becomes crucial. Based on several theories, we use a time-lagged research design to propose and test a moderated mediation model. We collected data from 402 employees across various industries in South Korea at three different time points. Our findings reveal that AI-induced job insecurity positively relates to knowledge-hiding behavior, directly and indirectly, via reduced psychological safety. Moreover, we discover that self-efficacy in AI learning moderates the relationship between AI-induced job insecurity and psychological safety, such that high self-efficacy buffers the harmful influence of job insecurity on psychological safety. These results enhance the existing literature on organizational technological change by clarifying the psychological processes through which AI implementation influences employee behavior. Our study highlights the critical role of psychological safety as a mediator and self-efficacy as a moderator in this process. These insights present significant implications for managers and organizations navigating the challenges of AI integration. They emphasize the need for strategies that foster psychological safety and enhance members' confidence in their ability to adapt to AI technologies. Our research underscores the significance of considering both the technical and human aspects of AI implementation within organizational contexts.

© 2024 The Author(s). Published by Elsevier España, S.L.U. on behalf of Journal of Innovation & Knowledge.

This is an open access article under the CC BY-NC-ND license

(<http://creativecommons.org/licenses/by-nc-nd/4.0/>)

Introduction

Artificial intelligence (AI) has emerged as a powerful and influential factor in the contemporary corporate world, instigating substantial transformations across industries and fundamentally altering the concept of employment. Although AI offers unprecedented opportunities for innovation and efficiency, it simultaneously poses significant challenges, particularly regarding job security (Bankins et al., 2024; Clement et al., 2024; Wu et al., 2022). The concept of AI-induced job insecurity has become increasingly prominent as organizations increasingly adopt AI technologies, potentially displacing human workers or fundamentally changing job roles (Budhwar et al., 2022;

Brougham & Haar, 2018; Ključnikov et al., 2023; Kraus et al., 2023; Pereira et al., 2023). This phenomenon has significant implications for individuals, organizations, and society, necessitating a comprehensive knowledge of its impact on the workforce. The pervasive nature of AI-induced job insecurity is evident across various sectors, including manufacturing and service industries, where AI systems can perform tasks traditionally executed by humans (Budhwar et al., 2022; Frey & Osborne, 2017; Li et al., 2023; Wu et al., 2022; Xue et al., 2024). The rapid advancement of AI capabilities exacerbates the uncertainty regarding future employment prospects, rendering AI-induced job insecurity a critical concern for both scholars and practitioners.

The impact of AI-induced job insecurity on employees' responses within organizations is a significant area of study, considering its profound implications on workplace dynamics and organizational performance. In AI-driven work environments, employees' perceptions, attitudes, and behaviors may be substantially affected by their sense

Note: All authors declare that they have no conflict of interest in this study.

* Corresponding author.

E-mail addresses: kimbj8212@ulsan.ac.kr (B.-J. Kim), mkim@koreatech.ac.kr (M.-J. Kim).

<https://doi.org/10.1016/j.jik.2024.100590>

2444-569X/© 2024 The Author(s). Published by Elsevier España, S.L.U. on behalf of Journal of Innovation & Knowledge. This is an open access article under the CC BY-NC-ND license (<http://creativecommons.org/licenses/by-nc-nd/4.0/>)

of job security or its absence (Bankins et al., 2024; Li et al., 2023; Shoss, 2017; Zitar et al., 2023). Comprehending these responses is essential for multiple reasons. Firstly, employee responses to AI-induced job insecurity may impact individual performance, thereby potentially impacting overall productivity and innovation (Sverke et al., 2019; Wu et al., 2022). Secondly, these responses may influence organizational culture and climate, thereby influencing the collective attitude toward AI adoption and technological change (Obreja et al., 2024; Goran et al., 2017). Lastly, comprehending employee reactions is crucial for formulating effective strategies to mitigate adverse outcomes and foster a positive, adaptive workforce amid AI-driven transformations (Felicetti et al., 2024; Ghislieri et al., 2018; Pereira et al., 2023). Therefore, thoroughly examining how AI-induced job insecurity affects employee perceptions, attitudes, and behaviors is imperative for organizations seeking to successfully navigate AI integration's challenges and opportunities (Wu et al., 2022).

Although the significance of examining the impact of AI-induced job insecurity on employees is acknowledged, numerous research gaps persist in the literature. The first significant gap pertains to the insufficient focus on the effect of AI-induced job insecurity on employee knowledge-related behaviors, particularly knowledge-hiding behavior (Bankins et al., 2024; Pereira et al., 2023; Wu et al., 2022). In the contemporary knowledge-driven economy, the effective management and sharing of knowledge are crucial for organizational success and innovation (Wang & Noe, 2010). However, the potential correlation between AI-induced job insecurity and employees' propensity to hide knowledge remains largely unexplored. This gap is particularly concerning because knowledge-hiding behavior (KHB) can significantly hinder organizational learning, creativity, and overall performance (Connelly et al., 2012, 2015). As AI technologies increasingly permeate the workplace, comprehending the impact of job insecurity on employees' willingness to share or conceal knowledge is essential for maintaining a competitive edge and fostering a collaborative work environment.

The second research gap addresses the lack of analysis of the mediating processes and moderators involved in the dynamics between AI-induced job insecurity and KHB (Bankins et al., 2024; Pereira et al., 2023). Although direct relationships yield valuable insights, exploring the mechanisms through which AI-induced job insecurity influences KHB and the conditions that amplify or mitigate this relationship provides a comprehensive understanding of the issue. Notably, the potential mediating effect of employee psychological safety in this context has been inadequately examined. Psychological safety, defined as the collective perception that a team environment supports interpersonal risk-taking (Edmondson, 1999), may be pivotal in elucidating why AI-induced job insecurity encourages KHB. Investigating psychological safety as a mediator enables researchers to explore the psychological mechanisms linking job insecurity to knowledge hiding, providing essential insights for both theoretical and practical applications.

The third research gap pertains to the limited examination of the protective (moderating) role of self-efficacy-related variables, particularly in the context of AI learning. In workplaces that have adopted AI, an employee's confidence regarding their capability to master and adapt to AI technologies is likely to significantly influence their reactions to AI-induced job insecurity (Kim & Kim, 2024; Wu et al., 2022). Self-efficacy in AI learning is an essential personal resource that can alleviate the adverse influences of job insecurity on psychological safety and subsequent KHB (Bandura, 1997). Understanding this moderating effect is crucial for various reasons. Firstly, it can assist in identifying employees who may exhibit greater resilience in response to AI-induced job insecurity. Secondly, it can guide the development of targeted interventions and training programs to enhance employees' AI-related self-efficacy. Lastly, it enhances a more comprehensive model regarding the interplay between individual differences interact and organizational factors in influencing responses to AI implementation.

Based on the identified research gaps, this study addresses the following research questions:

1. How does AI-induced job insecurity influence KHB within organizations?
2. What is the role of psychological safety in mediating the relationship between AI-induced job insecurity and knowledge-hiding behavior?
3. How significantly does self-efficacy in AI learning moderate the relationship between AI-induced job insecurity and psychological safety?

By addressing these questions, our study aims to enhance the comprehension of the complex dynamics between AI implementation, employee perceptions, and knowledge management practices in contemporary organizations.

This study adopts a multifaceted approach based on solid theoretical foundations to address these research gaps. Drawing on the conservation of resources theory, social exchange theory, uncertainty reduction theory, the job demands-resources (JD-R) model, affective events theory (AET), organizational support theory, and social cognitive theory (Bandura, 1986), this paper proposes a moderated mediation model that explicates the AI-induced job insecurity-KHB link. Specifically, we examine the mediating role of psychological safety in the association, arguing that AI-induced job insecurity diminishes psychological safety, subsequently resulting in augmented KHB. Moreover, this research navigates the moderating effect of self-efficacy in AI learning, positing that elevated self-efficacy levels can mitigate the negative influences of AI-induced job insecurity on psychological safety. This comprehensive model fills the identified research voids and enriches our understanding of employee responses to AI deployment at work. By amalgamating these theoretical views and examining both mediating and moderating dynamics, this study delivers crucial insights to academics and practitioners addressing the complexities of AI integration in organizational contexts.

This study significantly contributes to the existing knowledge on AI-induced job insecurity and its organizational repercussions. First, it enhances our comprehension of the outcomes of AI-induced job insecurity by analyzing its influence on KHB. Consequently, it connects the dots between technological evolution and knowledge management research. Second, by exploring the mediating role of psychological safety, this research sheds light on the psychological mechanisms underpinning the impact of AI-induced job insecurity on employee actions. This exploration provides an enhanced understanding of the cognitive and emotional processes linking job insecurity to knowledge-related behaviors in AI-imbuéd settings. Third, this study investigates the moderating influence of self-efficacy in AI learning, underscoring the significance of individual differences in shaping reactions to AI-induced job insecurity and offering valuable insights for tailored intervention strategies. Lastly, based on the conservation of resources, social exchange, and social cognitive theories, the study's moderated mediation model establishes a comprehensive theoretical framework for interpreting the intricate interactions among organizational context, individual perceptions, and behavioral outcomes following AI-induced changes. Collectively, these contributions propel scholarly understanding in organizational behavior and human resource management. Moreover, they offer practical implications for organizations aiming to address the challenges of AI integration and foster a culture of knowledge sharing and psychological safety.

Theory and hypotheses

AI-induced job insecurity and employee KHB

We propose that AI-induced job insecurity may increase employee KHB. In recent years, "AI-induced job insecurity" has

emerged as a significant area of research, intersecting the fields of organizational psychology, labor economics, and technology management. This phenomenon indicates the perceived threat to job stability and continuity due to the rapid advancement and integration of AI technologies in the workplace (Brougham & Haar, 2018; Climent et al., 2024). AI-induced job insecurity embodies a distinct form of job insecurity, characterized by employees' concerns regarding the potential displacement or significant alteration of their roles due to AI and automation (AlQershi et al., 2023; Ključnikov et al., 2023; Kraus et al., 2023; Nam, 2019). Unlike conventional sources of job insecurity, such as economic downturns or organizational restructuring, AI-induced insecurity stems from technological advancements that may replicate or surpass human cognitive abilities across diverse tasks (Borges et al., 2021; Frey & Osborne, 2017). Recent research has highlighted the multifaceted nature of AI-induced job insecurity. It includes concerns regarding job loss and anxieties about skill obsolescence, reduced job autonomy, and changes in job content (Brougham & Haar, 2020; Li et al., 2023; Xue et al., 2024). This broader conceptualization aligns with the job preservation insecurity framework, which considers threats to valued job features beyond simple employment continuity (Shoss, 2017). The ramifications of AI-induced job insecurity surpass the conventional outcomes linked to general job insecurity. Studies have discovered that it can result in increased resistance to AI technologies, reduced job satisfaction, and decreased organizational commitment (Brougham & Haar, 2020; Ključnikov et al., 2023; Kraus et al., 2023). Moreover, AI-induced job insecurity has been associated with heightened levels of technostress and decreased psychological well-being (Califf et al., 2020).

KHB is a significant construct in organizational behavior and knowledge management research, defined as a deliberate effort by an individual to conceal or withhold information requested by another party (Connelly et al., 2012, 2015). Recently, this concept has garnered significant attention because of its potential negative impact on organizational learning, innovation, and performance (Fauzi, 2023). The consequences of KHB are predominantly negative and far-reaching. Studies have consistently associated knowledge hiding with reduced creativity and innovation, decreased team performance, and diminished organizational citizenship behaviors (Bogilović et al., 2017; Černe et al., 2014). Moreover, KHB can engender a cycle of reciprocal distrust and increased knowledge hiding, potentially resulting in a toxic organizational culture (Škerlavaj et al., 2018).

The relationship between AI-induced job insecurity and KHB is elucidated via the lens of conservation of resources (COR) (Hobfoll, 1989), social exchange (Blau, 1964), and uncertainty reduction theories (Berger & Calabrese, 1975). These theoretical frameworks establish a robust basis for explaining why employees experiencing AI-induced job insecurity may engage in KHB.

First, the COR theory (Hobfoll, 1989) provides a robust framework for comprehending the relationship between AI-induced job insecurity and KHB. It posits that individuals strive to obtain, retain, and protect valuable resources, and they experience stress when these resources are threatened or depleted (Hobfoll et al., 2018). In the context of AI implementation in the workplace, employees' job security and specialized knowledge can be regarded as critical resources. As AI technologies advance, potentially replicating or surpassing human cognitive abilities in numerous tasks (Frey & Osborne, 2017), employees may perceive their job security as jeopardized. According to COR theory, this perception of resource threat can trigger protective behaviors intended to conserve remaining resources (Halbesleben et al., 2014). In this context, employees' knowledge emerges as an increasingly valuable resource that distinguishes them from AI capabilities. KHB can be interpreted as a resource protection strategy within the COR framework. By engaging in KHB, employees attempt to maintain their perceived value and irreplaceability within the

organization, thus safeguarding their job security against AI-induced threats. In short, the COR theory predicts that as employees experience increased AI-induced job insecurity (a perceived threat to a valued resource), they will be more inclined to engage in KHB to protect their remaining valuable resources (i.e., their unique knowledge and skills).

Second, as proposed by Blau (1964), the social exchange theory posits that social behavior results from an exchange process to maximize benefits and minimize costs. Within organizational contexts, this theory implies that employees establish reciprocal relationships with their employers, anticipating a balance of mutual exchange over time. In AI implementation, social exchange theory provides a robust framework for understanding the relationship between AI-induced job insecurity and KHB. Employees may interpret the introduction of AI technologies that potentially threaten job security as violating the implicit psychological contract between them and their employer (Rousseau & McLean Parks, 1993). This contract typically encompasses expectations of job stability in exchange for loyalty and performance. AI-induced job insecurity can create a perceived imbalance in the exchange relationship, wherein employees may feel that their contributions and loyalty are not reciprocated with job security, resulting in feelings of unfairness or inequity. To address this perceived imbalance, employees may exhibit negative reciprocity behaviors. KHB can be regarded as one such behavior, where employees withhold valuable information to restore balance or safeguard their position within the organization (Černe et al., 2014). Employees hide knowledge to maintain their unique value proposition within the organization, potentially safeguarding their positions against AI-related job threats. This theoretical rationale directly supports our hypothesis that AI-induced job insecurity will increase KHB. When employees perceive a breach in their psychological contract due to AI implementation, they are more inclined to engage in KHB as a form of negative reciprocity and strategic value retention.

Third, the uncertainty reduction theory, introduced by Berger and Calabrese (1975), proposes that individuals have a fundamental need to minimize uncertainty in their environment to enhance predictability and understanding. This theory is particularly relevant in explaining the relationship between AI-induced job insecurity and KHB. The introduction of AI technologies generates considerable uncertainty in the workplace, particularly regarding future job roles, required skills, and individual value within the organization. The theory posits that elevated uncertainty motivates individuals to pursue information and engage in behaviors that help them regain control and predictability. In AI-induced job insecurity, employees' specialized knowledge becomes crucial for reducing uncertainty. By retaining exclusive knowledge, employees seek to uphold their unique value proposition within the organization, secure their ongoing relevance in an AI-augmented workplace, and exert influence over their future job prospects. Although uncertainty reduction theory typically focuses on information-seeking behaviors, information retention (through knowledge hiding) functions similarly in this context. By controlling the flow of their specialized knowledge, employees attempt to mitigate uncertainty regarding their future within the organization (Fong et al., 2018). This theoretical perspective supports our hypothesis that AI-induced job insecurity will increase KHB. Employees facing heightened uncertainty due to AI implementation are more likely to engage in KHB to reduce this uncertainty and maintain control over their professional future.

This integrated theoretical approach provides a robust and logical foundation for comprehending the intricate relationship between AI-induced job insecurity and KHB in modern organizations. By utilizing COR, social exchange, and uncertainty reduction theories, we can better comprehend the psychological mechanisms that influence employee behavior in response to the challenges posed by AI implementation in the workplace.

Hypothesis 1. AI-induced job insecurity positively correlates with KHB.

AI-induced job insecurity and psychological safety

This paper proposes that AI-induced job insecurity may decrease employee psychological safety. Psychological safety, which is defined as the collective perception that a team or organization provides a secure setting for interpersonal risk-taking (Edmondson, 1999), has gained prominence due to its significant influence on team dynamics, innovation, and organizational performance (Edmondson & Bransby, 2023). Fundamentally, psychological safety entails an individual's belief that they can freely express themselves without apprehension of negative repercussions on their self-image, status, or career trajectory (Kahn, 1990). In environments characterized by psychological safety, team members are motivated to voice concerns, share ideas, acknowledge errors, and solicit feedback (Newman et al., 2017). The implications of psychological safety are extensive and overwhelmingly positive, with research consistently demonstrating its role in enhancing learning behaviors, creativity, and innovation within teams (Edmondson & Bransby, 2023; Newman et al., 2017). Additionally, psychological safety correlates with enhanced job performance, increased employee engagement, and more effective knowledge sharing (Edmondson & Bransby, 2023; Frazier et al., 2017).

The interplay between AI-induced job insecurity and employee psychological safety can be explored through various theoretical perspectives, including the job demands-resources (JD-R) model, AET, and organizational support theory. Collectively, these frameworks elucidate why employees experiencing AI-induced job insecurity may witness a decline in psychological safety.

First, the JD-R model offers a detailed framework for examining the impact of workplace attributes on employee well-being and performance (Bakker & Demerouti, 2017). This model categorizes job characteristics into two primary categories: job demands and job resources. Job demands are elements of work that require sustained physical, cognitive, or emotional effort, potentially resulting in physiological and psychological strain. Conversely, job resources are elements that aid in achieving work goals, alleviating job demands, or promoting personal growth and development. In scenarios of AI-induced job insecurity, such insecurity is perceived as a significant job demand. The looming threat of job loss or role obsolescence due to AI advancements compels employees to utilize substantial cognitive and emotional resources to navigate this uncertainty. This escalation in demand can erode an individual's capacity to engage in behaviors associated with psychological safety, such as initiating dialogue, acknowledging mistakes, or soliciting feedback. Furthermore, the JD-R model proposes that elevated job demands coupled with inadequate resources can result in strain and burnout (Demerouti et al., 2001). Thus, AI-induced job insecurity not only elevates job demands but may also diminish perceived job resources such as job control and future career prospects, thus fostering a work environment where psychological safety is compromised.

Second, the AET proposes that workplace situations elicit emotional responses, subsequently impacting attitudes and behaviors. (Weiss & Cropanzano, 1996). The theory emphasizes the significance of time in the progression of emotional experiences and their outcomes, implying that the accumulation of affective experiences over time influences work attitudes and behaviors. In the context of applying AET to AI-induced job insecurity and psychological safety, the implementation or advancement of AI in the workplace can be regarded as a series of affective events. Every instance of AI implementation or news regarding AI advancements may trigger emotional responses, such as anxiety, fear, or uncertainty. According to AET, these emotional responses do not manifest in isolation but accumulate over time. The cumulative effect of these negative affective experiences can result in a persistent state of emotional strain, which

may manifest as a diminished sense of psychological safety. Employees who continually experience negative emotions associated with AI-induced job insecurity may exhibit a diminished propensity to participate in vulnerable behaviors, including acknowledging mistakes or seeking help, which are hallmarks of psychological safety (Edmondson & Lei, 2014).

Third, the organizational support theory implies that employees form generalized beliefs regarding the degree to which their organization values their contributions and cares for their welfare (Eisenberger et al., 1986). These perceptions of organizational support (POS) are pivotal in molding employees' relationships with their employers and influencing various work-related outcomes. Within the context of AI integration, employees' POS can be significantly influenced by how they perceive the organization's handling of this technological shift. If deploying AI technologies is perceived as jeopardizing job security without sufficient communication, training, or backing from the organization, employees may view this as an indication of the organization's disregard for their welfare. The perceived lack of support may undermine the foundational trust that is essential for psychological safety. Consequently, when employees perceive that their organization undervalues their contributions or disregards their future in light of AI advancements, their likelihood of engaging in interpersonal risk-taking or behaviors that expose their vulnerability diminishes (Kurtessis et al., 2017).

Integrating these theoretical perspectives establishes a robust foundation for our hypothesis that AI-induced job insecurity will negatively affect employee psychological safety. The JD-R model explains how AI-induced job insecurity functions as a job demand that can deplete resources essential for sustaining psychological safety. AET elucidates how ongoing concerns regarding AI and job security can accumulate, resulting in a persistent state of reduced psychological safety. Organizational support theory elucidates how perceptions of an organization's AI implementation strategy can undermine the trust and support essential for psychological safety to flourish.

Hypothesis 2. AI-induced job insecurity negatively correlates with psychological safety.

Psychological safety and KHB

This research proposes that decreased psychological safety among employees can significantly elevate their risk of KHB. Various theoretical frameworks, such as social exchange theory, the theory of planned behavior, and organizational learning theory, elucidate the psychological safety-KHB link. These perspectives offer a comprehensive foundation for understanding why employees with elevated psychological safety may be less inclined to engage in KHB.

Firstly, the social exchange theory, proposed by Blau (1964), posits that social behavior emerges from an exchange process. Individuals engage in interactions anticipating rewards or benefits in return. In organizational contexts, these exchanges encompass not only material benefits but also socio-emotional resources, such as trust, support, and recognition. In the context of psychological safety and KHB, this theory offers a robust structure for understanding the reciprocal nature of behaviors at work. When members perceive a significant level of psychological safety, they regard it as a valuable resource offered by their organization and colleagues. This engenders a sense of obligation to reciprocate, not via direct repayment, but through behaviors that benefit the organization and team members.

Psychological safety, characterized by an environment conducive to taking interpersonal risks, is considered an organization's investment in its employees. The norm of reciprocity (Gouldner, 1960) posits that individuals are motivated to reciprocate favorable treatment with positive responses. Therefore, employees experiencing psychological safety are more inclined to reciprocate by openly sharing their knowledge rather than hiding it. Moreover, social exchange theory

posits that relationships evolve over time into trusted, loyal, and mutually committed bonds (Cropanzano & Mitchell, 2005). In a psychologically safe environment, continuous interactions between employees and their organization/colleagues are likely to foster relationships that diminish the inclination toward knowledge hiding.

Second, the theory of planned behavior (Ajzen, 1991) offers a cognitive structure for understanding intentional behavior. It indicates that behavioral intentions are influenced by three crucial factors: attitudes toward the conduct, subjective norms, and perceived behavioral control. Utilizing this theory to examine the psychological safety-KHB link offers valuable insights. Psychological safety influences an individual's attitude toward knowledge sharing. In a psychologically safe environment, employees are likely to perceive knowledge sharing more positively because the perceived risks of sharing information are minimized. Furthermore, psychological safety influences the perceived social pressure to either engage in or abstain from knowledge hiding. Open communication and collaboration are often valued and expected in psychologically safe environments, establishing a norm that discourages knowledge hiding. Additionally, psychological safety enhances individuals' perception of their ability to share knowledge without negative consequences. This increased sense of control over the knowledge-sharing outcomes can reduce the perceived need for knowledge hiding. This theory proposes that the elements collectively affect a member's intention to participate in KHB. Psychological safety creates a cognitive framework that discourages KHBs by positively impacting attitudes, subjective norms, and perceived behavioral control.

Third, the organizational learning theory (Argyris & Schön, 1978) highlights the significance of ongoing learning and adaptation for organizational success. This theory distinguishes between single-loop learning, which entails correcting errors within existing frameworks, and double-loop learning, which involves questioning and modifying underlying assumptions and norms. Psychological safety is essential for facilitating both forms of learning. It encourages employees to report errors and share information regarding day-to-day operations, thereby enabling the organization to efficiently detect and correct mistakes. In psychologically safe environments, employees feel more comfortable challenging existing norms and proposing innovative ideas, and this facilitates deeper organizational learning and adaptation. Thus, psychological safety can decrease members' KHB.

Integrating these theoretical perspectives provides a robust foundation for our hypothesis that employee psychological safety is negatively associated with KHB. Social exchange theory suggests that psychological safety engenders a sense of obligation to reciprocate with positive actions, such as knowledge sharing, thus discouraging KHB. The theory of planned behavior indicates that psychological safety influences attitudes, norms, and perceived control, rendering knowledge hiding less attractive and less essential. Organizational learning theory posits that psychological safety is crucial for building an environment where open knowledge sharing is essential for collective learning and adaptation. Consequently, such an environment reduces the inclination toward knowledge hiding.

Hypothesis 3. Psychological safety is negatively related to KHB.

Mediating role of psychological safety in the ai-induced job insecurity-KHB link

This paper suggests that psychological safety mediates the AI-induced job insecurity-KHB link. The context-attitude-behavior (CAB) structure offers a robust theoretical ground for understanding the relationship between AI-induced job insecurity, psychological safety, and KHB. The CAB framework posits that context is an antecedent to both attitudes and behaviors, whereas attitudes mediate the relationship between context and behavior. This sequential

process emphasizes the significance of considering the broader context in which attitudes are formed and behaviors occur (Johns, 2006).

In the CAB model, context encompasses the situational factors that influence the incidence and interpretation of workplace behaviors. This encompasses physical, social, and task environments (Mowday & Sutton, 1993). The framework suggests that these contextual factors shape individual attitudes by influencing perceptions, interpretations, and evaluations of various environmental stimuli. In this framework, attitudes are characterized as persistent structures of beliefs, emotions, and behavioral inclinations toward socially important entities, collectives, occurrences, or symbols (Hogg & Vaughan, 2018). They serve as a mediating mechanism by which contextual factors influence behavior. The CAB framework proposes that attitudes are more proximally associated with behavior than context, serving as a bridge between environmental stimuli and individual actions. Behavior, the final component of the framework, indicates members' observable actions in response to their environment and internal states. The CAB model suggests that although context can directly influence behavior, a considerable portion of this influence is mediated by attitudes.

In this study, AI-induced job insecurity is the contextual factor, psychological safety is the attitudinal component, and KHB is the behavioral outcome. The CAB framework implies that the influence of AI-induced job insecurity on KHB is indirect; rather, it is mediated by the formation of attitudes, specifically the perception of psychological safety.

AI-induced job insecurity engenders a work context marked by uncertainty and potential threats to one's professional future (Brougham & Haar, 2018). This environmental condition can profoundly influence employees' attitudes toward their workplace, particularly their sense of psychological safety. This concept refers to how individuals perceive the potential outcomes of taking social risks within their work environment (Edmondson & Lei, 2014). In AI-induced job insecurity, employees may perceive greater risks associated with open communication and vulnerability, potentially reducing psychological safety.

The attitudinal component of psychological safety is essential in fostering behavioral outcomes. According to the CAB framework, attitudes act as a proximal predictor of behavior, mediating the influence of contextual factors (Shin & Hur, 2019). In this scenario, reduced psychological safety would increase KHB as a protective mechanism. Members engage in KHB when they perceive potential threats or a lack of safety in sharing information (Connelly et al., 2012).

This mediation hypothesis aligns with the sequential nature of the CAB framework, indicating that the influence of AI-induced job insecurity on KHB operates through its impact on psychological safety. AI-induced job insecurity reduces psychological safety, fostering an attitudinal climate that is more conducive to KHB.

Hypothesis 4. Psychological safety mediates the AI-induced job insecurity-KHB link.

Moderating influence of self-efficacy in ai learning in the ai-induced job insecurity-psychological safety link

We posit that self-efficacy in AI learning functions as a moderating factor, mitigating the harmful influence of AI-induced job insecurity on psychological safety. Based on our prior discussions, it is evident that AI-induced job insecurity may decrease psychological safety levels. However, it is crucial to recognize that the influence of AI-induced job insecurity on psychological safety varies across various organizational contexts.

The concept of "self-efficacy in AI learning" is an emerging area of study that integrates the well-established construct of self-efficacy with the rapidly evolving domain of AI education and training. This specific form of self-efficacy indicates individual conviction in their

capability to learn, understand, and apply AI concepts and technologies effectively (Kim & Kim, 2024; Kim et al., 2024). As AI increasingly infiltrates diverse sectors of society and the economy, the importance of AI literacy has grown significantly. In this context, self-efficacy in AI learning is crucial in determining individuals' engagement with AI education, their unwavering determination in the presence of obstacles, and ultimately their mastery of AI-related skills (Hatlevik et al., 2018; Kim & Kim, 2024; Kim et al., 2024). The consequences of high self-efficacy in AI learning are predominantly favorable. Prior studies have demonstrated that members with higher AI learning self-efficacy are more likely to engage in AI-related courses and activities, demonstrate greater persistence in learning complex AI concepts, and achieve enhanced learning outcomes. Moreover, high self-efficacy in AI learning is associated with heightened interest in AI-related careers and a more optimistic outlook on the role of AI in future workplaces (Scherer et al., 2019).

The moderating effect of self-efficacy in AI learning on AI-induced job insecurity-employee psychological safety link is understood via the lenses of social cognitive theory (Bandura, 1986) and the JD-R model. These theoretical frameworks provide a strong foundation for understanding how individuals' beliefs in their ability to learn and adapt to AI technologies can buffer the harmful influences of AI-induced job insecurity on psychological safety.

The social cognitive theory asserts that individuals' beliefs in their ability to perform specific tasks significantly shapes their behavior, motivation, and emotional states (Bandura, 1997). Within AI implementation, employees with high self-efficacy in AI learning tend to perceive AI-related changes as challenges to be conquered rather than threats to be evaded. Such positive cognitive assessments can mitigate the adverse effects of AI-induced job insecurity on psychological safety.

The JD-R model reinforces this mitigating effect by proposing that personal resources, such as self-efficacy, can aid individuals in coping job demands (Bakker & Demerouti, 2017). Thus, AI-induced job insecurity is categorized as a job demand, whereas self-efficacy in AI learning acts as a personal resource that counterbalances the harmful impact of this demand on psychological safety.

When employees exhibit high self-efficacy in AI learning, the detrimental influence of AI-induced job insecurity on psychological safety is likely to be considerably diminished. Elevated self-efficacy in AI learning cultivates a growth mindset (Dweck, 2006). Employees with this mindset perceive AI-related challenges as opportunities for personal and professional growth rather than threats to their job security. This perspective enables them to sustain a higher degree of psychological safety, even amidst AI-related uncertainties. Furthermore, high self-efficacy fosters proactive behavior (Parker et al., 2010). Employees who possess confidence in their capability to master AI-related skills are more inclined to take the initiative to learn and integrate new AI technologies. This proactive stance can foster a sense of control over their work environment, a crucial element of psychological safety (Edmondson & Lei, 2014). Additionally, high self-efficacy in AI learning can enhance performance and adaptability (Compeau & Higgins, 1995). As employees effectively acquire new AI skills, they reinforce their belief in their ability to adapt to technological changes. This positive reinforcement cycle can help sustain psychological safety by diminishing the perceived threat of AI to job security.

For example, in a financial services organization implementing AI-driven customer service chatbots, a member with high self-efficacy in AI learning may perceive this change as an opportunity to enhance their skills. They may actively endeavor to understand the AI system, propose enhancements, and identify ways to complement the AI's capabilities with their human expertise. This proactive approach, derived from their high self-efficacy, enables them to maintain psychological safety despite the potential threat to traditional customer service roles.

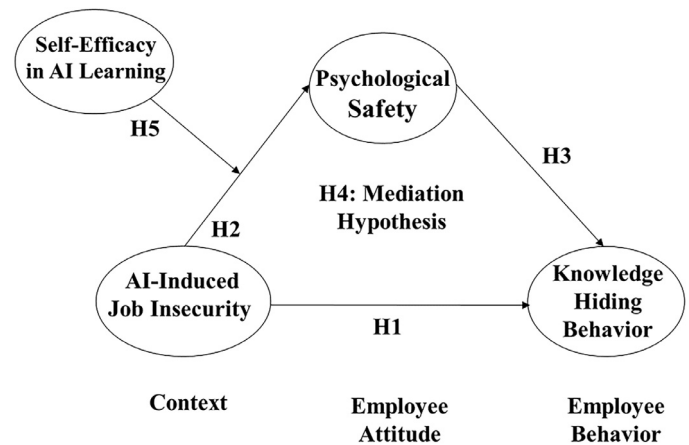


Fig. 1. Theoretical Model.

Conversely, when employees possess low self-efficacy in AI learning, AI-induced job insecurity's harmful influences on psychological safety are likely more pronounced. Low self-efficacy in AI learning frequently relates to a fixed mindset (Dweck, 2006). Employees with this mindset may perceive AI advancements as insurmountable challenges, exacerbating their job insecurity and significantly reducing their psychological safety. Moreover, low self-efficacy can result in avoidance behaviors (Bandura, 1997). Employees who question their ability to learn AI-related skills may refrain from interacting with new technologies, resulting in a self-fulfilling prophecy that renders their skills obsolete. This avoidance may heighten their perceived vulnerability, further diminishing their psychological safety. Furthermore, low AI learning self-efficacy can increase stress and anxiety (Shu et al., 2011). The perceived inability to cope with AI-related changes can cause chronic stress, which directly undermines psychological safety by fostering a persistent state of threat and uncertainty.

For instance, in the same financial services organization, a member with low self-efficacy in AI learning may view the implementation of AI chatbots as a direct threat to their employment. They may exhibit reluctance to engage with the new technology, refrain from expressing their concerns or ideas for improvement, and encounter increased anxiety regarding their future roles. This defensive stance, stemming from their low self-efficacy, can result in a marked decrease in psychological safety due to their difficulty exploring AI-driven changes.

Based on the above arguments, we suggest the following hypothesis (Please See Fig. 1).

Hypothesis 5. Self-efficacy in AI learning moderates the AI-induced job insecurity-psychological safety link, such that high self-efficacy decreases the negative impact of AI-induced job insecurity on psychological safety.

Method

Participants and procedure

The study encompassed a heterogeneous cohort of workers, aged 20 and older, from various corporations across South Korea. The data collection was structured into three distinct stages using a 3-wave time-lagged research design to comprehensively explore changes over time. Data were gathered using Macromil Embrain, a leading online research platform with a proven track record in facilitating academic studies. This platform hosts a diverse panel of over 5.49 million potential respondents spanning multiple demographics and industries. Macromil Embrain employs rigorous quality control measures, including:

- 1. Verification of participant identities using multi-step authentication processes
- 2. Regular updates and validation of participant information
- 3. Sophisticated algorithms for the detection and prevention of fraudulent responses
- 4. Adherence to international data protection regulations

The platform's reliability and validity have been established through multiple peer-reviewed studies published in reputable journals across diverse disciplines (e.g., Kim & Kim, 2024; Kim et al., 2024). Utilizing this platform guarantees a diverse and representative sample while maintaining high data quality standards. During the online sign-up phase, participants confirmed their current employment status and completed a verification step, which necessitated the provision of either a mobile number or an email address to enhance the security of the data collection process. The effectiveness of utilizing digital surveys to obtain a wide-ranging and varied sample has been well-supported by its demonstrated validity and reliability in prior research (Landers & Behrend, 2015).

The research aimed to collect time-lagged data from actively employed individuals in South Korean companies, addressing the limitations frequently associated with cross-sectional studies. The advanced digital tools employed in this study facilitated precise tracking of participant engagement throughout the survey period, ensuring consistent participation across all data collection phases. Data collection intervals were meticulously scheduled to occur every five to six weeks, with each survey session remaining open for two to three days to maximize response rates. To maintain data integrity, we enforced strict protocols, including using geo-IP restriction traps to prevent any overly rapid responses, thus safeguarding the quality and credibility of the research outcomes.

The survey administrators proactively contacted potential participants to invite them to participate in the study. The team explicitly informed potential respondents that their participation was entirely voluntary and reassured them of the confidentiality and exclusive use of their data for research purposes. Strict ethical standards were upheld throughout the process upon obtaining informed consent from willing participants. Participants were compensated for their time with a monetary reward of about \$8 or \$9.

To minimize the potential for sample bias, the team employed a stratified random sampling technique. The method entailed randomly selecting participants from specific strata, thereby mitigating biases associated with demographic and occupational variables, including gender, age, position, educational level, and industry sector. The survey team meticulously tracked participant engagement across multiple digital platforms to guarantee consistent participation for the same individuals throughout all three survey waves (Refer to Table 1).

Regarding participation, the initial survey phase garnered responses from 833 employees. The number diminished to 591 the second time, with the final phase recording responses from 405 employees. Following the data collection phase, a rigorous cleaning was conducted to eliminate any incomplete entries. Ultimately, the study's final sample size comprised 402 respondents who completed all phases of the survey, resulting in a final response rate of 48.26%. The sample size calculation was guided by scholarly advice, including using G*Power for statistical power analysis.

Several methodological and theoretical considerations drove the decision to employ a time-lagged design. First, this approach mitigates common method bias, a significant concern in behavioral research (Podsakoff et al., 2003). By temporally separating the measurement of predictor and outcome variables, we minimize the potential for spurious correlations due to simultaneous measurement. Second, the time-lagged design aligns with our model's theoretical temporal sequence. AI-induced job insecurity (measured at Time 1) is anticipated to influence psychological safety (Time 2),

Table 1
Descriptive Characteristics of the Sample.

Characteristic	Percent
Gender	
Men	51.7%
Women	48.3%
Age (years)	
20–29	21.4%
30–39	23.6%
40–49	27.1%
50–59	27.9%
Education	
High school or below	11.9%
Community college	19.2%
Bachelor's degree	56.0%
Master's degree or higher	12.9%
Position	
Staff	42.3%
Assistant manager	16.2%
Manager or deputy general manager	23.7%
Department/general manager or director and above	17.9%
Firm Size	
1–9 employees	16.2%
10–29 employees	16.9%
30–49 employees	10.2%
50–99 employees	16.7%
100–149 employees	7.7%
150–299 employees	7.0%
300–449 employees	4.0%
500–999 employees	6.5%
1,000–4,999 employees	8.7%
Above 5,000 employees	6.2%
Industry Type	
Manufacturing	21.6%
Services	14.4%
Construction	12.2%
Health and welfare	16.4%
Information services and telecommunications	8.9%
Education	16.2%
Financial/insurance	2.2%
Consulting and advertising	0.7%
Others	7.2%

thereby affecting KHB (Time 3). This design enables us to capture the unfolding of these processes over time, yielding a more accurate representation of the causal relationships hypothesized in our model (Mitchell & James, 2001). Lastly, due to the dynamic nature of technological change and its effects on employee perceptions and behaviors, the time-lagged approach enables us to capture the evolving nature of AI-induced job insecurity and its consequences. This design is particularly appropriate for examining the long-term effects of technological changes on organizational behavior (Zaheer et al., 1999).

Measures

During the initial survey phase, we queried participants on their experiences of AI-induced job insecurity and the level of self-efficacy in AI learning they received. The second survey focused on assessing their perceptions of psychological safety. In the third survey, we collected data concerning the participants' levels of knowledge-hiding behavior. All variables were measured using multi-item scales rated on a 5-point Likert scale.

AI-Induced job insecurity (Point in Time 1, as reported by workers)

To measure AI-induced job insecurity, we adapted a scale originally developed by Kraimer et al. (2005): job insecurity scale. The original scale was modified to address the context of AI implementation in the workplace. The adaptation ensures that the measure captures the unique aspects of job insecurity arising from AI advancements. This adaptation preserves the essence of Kraimer et al.'s (2005) original scale while tailoring it to the specific context of

AI in the workplace. This allows for a more precise measurement of job insecurity related to AI advancements. The adapted scale comprises five items: “The introduction of artificial intelligence technologies has made my job at work insecure”, “If my current organization faces economic problems, my job will be replaced by artificial intelligence”, “Even if I want to, I will not be able to keep my present job due to the adoption of artificial intelligence”, “If the organization’s economic situation worsens, the organization will introduce artificial intelligence technology, and my current job may disappear”, and “Even if I want to, I am not confident that I can continue to work at my organization due to the introduction of artificial intelligence technology”. These items were designed to reflect the perceived threat of job loss or instability attributable to AI implementation, aligning with the original scale’s focus on job insecurity within the context of technological change. Cronbach’s alpha was reported to be 0.954.

Self-efficacy in AI learning (Point in Time 1, as reported by workers)

This study employed a four-item scale from prior research that adapted Bandura’s self-efficacy scale for application in AI learning (Bandura, 1986). The scale was developed based on prior research (Kim & Kim, 2024). Specifically, this study analyzed the following items: “I am confident in my ability to learn artificial intelligence technology appropriately in my work”, “I am able to learn artificial intelligence technology to perform my job well, even when the situation is challenging”, “I can develop my competencies needed for my job through AI technology learning”, and “I will be able to learn important information and skills from my AI training”. Cronbach’s alpha was reported to be 0.938.

Psychological safety (Point in Time 2, as reported by workers)

Using seven questions drawn from Edmondson’s (1999) psychological safety measure, we assessed employees’ sense of psychological safety. Employees’ views on psychological safety are assessed using this scale. The sample includes the following items: “It is safe to take a risk in this organization”, “I am able to bring up problems and tough issues in this organization”, “It is easy for me to ask other members of this organization for help”, and “No one in this organization would deliberately act in a way that undermines my efforts”. The items were selected from previous studies conducted in South Korea (Kim et al., 2021; Kim et al., 2019). Cronbach’s alpha was reported to be 0.799.

KHB (Point in Time 3, as reported by workers)

The extent of knowledge-hiding behavior was evaluated using five items of the KHB scale, which comprises eleven items (Connelly et al., 2012). The current paper shortened the full items because the five items were empirically validated by previous studies conducted in the South Korean context (Jeong et al., 2022, 2023). The items included: “I offer him/her some other information instead of what he/she really wants”, “I agree to help him/her but instead give him/her information different from what s/he wants”, “I pretend that I do not know the information”, “I explain that the information is confidential and only available to people on a particular project”, and “I tell him/her that my boss would not let anyone share this knowledge”. Cronbach’s alpha was reported to be 0.931.

Control variables

Based on prior studies (Connelly et al., 2012, 2015), this study utilized several control factors, including employee tenure, gender, job position, and educational level, to adjust for their possible influences on the outcomes. The data was gathered during the initial survey phase. The rationale for including these specific control variables is based on their established correlation with KHB in existing literature, with the objective of minimizing omitted variable bias and enhancing the clarity of the results of our primary variables of interest.

Statistical analysis

To examine the connections among the research variables, we utilized SPSS version 28 to conduct a Pearson correlation analysis after data collection. Anderson and Gerbing (1988) proposed a two-stage methodology, which we meticulously adhered to: a measurement model and a structural model. To examine the structural model using moderated mediation analysis, we employed the AMOS 26 program and utilized the maximum likelihood estimate technique. After confirming the measurement model with Confirmatory Factor Analysis (CFA), we adhered to established Structural Equation Modeling procedures.

To evaluate the appropriateness of the model in accurately representing the empirical data, we employed several model fit indices, including the Comparative Fit Index (CFI), the Tucker-Lewis Index (TLI), and the Root Mean Square Error of Approximation (RMSEA). Previously published academic criteria considered model fits acceptable if CFI and TLI values exceeded 0.90 and RMSEA was below 0.06.

Results

Descriptive statistics

Our research revealed significant correlations among the key variables of interest: AI-induced job insecurity, self-efficacy in AI learning, psychological safety, and KHB. Table 2 shows these correlations.

Measurement model

To examine the sufficiency of our measuring model, we performed a CFA on all questions to assess the discriminant validity of our four primary variables: AI-induced job insecurity, self-efficacy in AI learning, psychological safety, and KHB. We conducted a series of chi-square difference tests to compare the 4-factor model (AI-induced job insecurity, self-efficacy in AI learning, psychological safety, and KHB) with alternative models. The 3-factor model had a chi-square value of 1633.060 with 131 degrees of freedom, a CFI of 0.737, a TLI of 0.692, and an RMSEA of 0.169. The 2-factor model had a chi-square value of 1958.358 with 133 degrees of freedom, a CFI of 0.680, a TLI of 0.632, and an RMSEA of 0.185. The one-factor model had a chi-square value of 3346.416 with 134 degrees of freedom, a CFI of 0.437, a TLI of 0.357, and an RMSEA of 0.245. The fit indices of each factor model implied that the 4-factor model outperformed the other models regarding fit. The chi-square value was 217.966 with 128 degrees of freedom. Additionally, the CFI was 0.984, the TLI was 0.981, and the RMSEA was 0.042. Further chi-square tests confirmed that the four research variables exhibited sufficient discriminant validity.

Structural model

We employed a moderated mediation model to investigate our hypotheses, integrating mediation and moderation analyses. Within this framework, we explored the extent to which AI-induced job insecurity affects KHB, with a particular focus on the mediating role of psychological safety. Furthermore, self-efficacy in artificial intelligence learning was considered a moderating variable that could weaken the decreasing impact of AI-induced job insecurity on psychological safety.

To operationalize our moderation analysis, this study constructed an interaction term by multiplying the variables representing AI-induced job insecurity and self-efficacy in AI learning. To minimize the effects of multicollinearity and maintain the integrity of our correlations, we initially centered these variables around their mean values. This technique not only mitigated multicollinearity but also

Table 2
Correlation among Research Variables.

	Mean	S.D.	1	2	3	4	5	6	7
1. Gender_T1	1.48	0.50	-						
2. Education_T1	2.70	0.84	-0.02	-					
3. Tenure_T1	66.22	72.34	-0.07	0.05	-				
4. Position_T1	2.43	1.52	-0.32**	0.28**	0.32**	-			
5. AIJL_T1	2.37	1.02	-0.11*	0.07	-0.05	0.01	-		
6. SE_T1	3.36	0.85	-0.06	0.10*	0.02	0.01	0.00	-	
7. PS_T2	3.27	0.58	-0.04	0.02	0.08	0.16**	-0.16**	0.11*	-
8. KHB_T3	2.17	0.88	-0.16**	-0.03	-0.12*	0.09	0.37**	-0.04	-0.21**

Notes: * $p < 0.05$. ** $p < 0.01$. S.D. means standard deviation, AIJL means AI-induced Job Insecurity, SE means self-efficacy in AI learning, PS means psychological safety, and KHB means knowledge-hiding behavior. As for gender, males are coded as 1 and females as 2. As for positions, general manager or higher are coded as 5, deputy general manager and department manager 4, assistant manager 3, clerk 2, and others below clerk as 1. As for education, the “below high school diploma” level is coded as 1, the “community college” level as 2, the “bachelor’s” level as 3, and the “master’s degree or more” level is coded as 5.

preserved correlation strength, thus enhancing the reliability of our moderation analysis (Brace et al., 2003).

To evaluate the degree of potential multicollinearity, we computed variance inflation factors (VIF) and appropriate tolerance levels using Brace et al.’s (2003) method. The results indicated that both AI-induced job insecurity and self-efficacy in AI learning exhibited VIF scores of 1.001, with tolerance indices closely aligned at 0.999. These metrics demonstrated that our variables exhibited no significant multicollinearity concerns, with VIF values substantially below the commonly accepted threshold of 10 and tolerance levels significantly above the minimum criterion of 0.2 (Brace et al., 2003).

Findings from the mediation evaluation

To find the most appropriate mediation model, we employed a chi-square difference test to determine if a full mediation model was more suitable than a partial mediation one. The direct association between AI-induced job insecurity and KHB was not examined in the full mediation model, in contrast to the partial mediation one. Both models had a strong fit, with the full mediation model displaying fit indices of ($\chi^2 = 338.914$ (df = 157), CFI = 0.960, TLI = 0.951, RMSEA = 0.054), and the partial mediation model showing ($\chi^2 = 300.161$ (df = 156), CFI = 0.968, TLI = 0.961, RMSEA = 0.048). Nevertheless, our analysis revealed that the partial mediation model demonstrated a better fit than that of the full mediation model ($\Delta\chi^2 [1] = 38.753$, $p < 0.01$). This preference for the partial mediation model has both theoretical and practical implications. Theoretically, it posits that although psychological safety plays a crucial mediating role, AI-induced job insecurity also directly affects KHB. This finding aligns with the COR theory (Hobfoll, 1989), indicating that employees may resort to knowledge hiding as a direct response to perceived threats to their job security, alongside an indirect path through reduced psychological safety.

Practically, the partial mediation model implies that interventions to reduce KHB should address both AI-induced job insecurity and psychological safety. Although enhancing psychological safety is crucial, organizations must also directly address employees’ concerns about job security amid AI implementation to effectively mitigate

KHBs. Additionally, we utilized control variables, including tenure, gender, educational level, and occupational position. However, these variables did not achieve statistical significance in influencing KHB, with the exception of position ($\beta = 0.148$, $p < 0.01$) and tenure ($\beta = -0.138$, $p < 0.01$).

Our analysis, which included these control variables, revealed a significant direct link between AI-induced job insecurity and KHB in the partial mediation model ($\beta = 0.313$, $p < 0.001$), thus supporting Hypothesis 1. This finding underscores the adequacy of the partial mediation model, which ultimately became the preferred model due to its significant direct pathway from AI-induced job insecurity on KHB and its superior fit indices. This comparison validates the relevance of the partial mediation model, demonstrating that the influence of AI-induced job insecurity on KHB is both direct and mediated by psychological safety.

Moreover, our findings corroborate Hypothesis 2, indicating a strong negative impact of AI-induced job insecurity on psychological safety ($\beta = -0.253$, $p < 0.001$), and Hypothesis 3, which reveals that psychological safety substantially decreases KHB ($\beta = -0.226$, $p < 0.001$). Table 3 depicts the results.

Bootstrapping for testing mediation effect

To examine the mediating effect of psychological safety on the relationship between AI-induced job insecurity and KHB (Hypothesis 4), we employed a robust bootstrapping technique, as recommended by Shrout and Bolger (2002). This approach presents several advantages compared to traditional mediation methods.

First, we utilized a substantial bootstrap sample of 10,000 resamples, which enhances the stability and reliability of our estimates, regardless of the original sample’s distribution. Second, we calculated 95% bias-corrected confidence intervals for the indirect effect. This method yields a more accurate estimate of the true population parameter and is more robust against Type I errors than normal theory approaches. Third, to determine the statistical significance of the indirect effect of AI-induced job insecurity on KHB through psychological safety, we examined whether the 95% CI encompassed zero. A CI that does not contain zero indicates a significant indirect effect.

Table 3
Results of Structural Model.

Hypothesis	Path (Relationship)	Unstandardized Estimate	S.E.	Standardized Estimate	Supported
1	AI-induced Job Insecurity → KHB	0.251	0.041	0.313***	Yes
2	AI-induced Job Insecurity → Psychological Safety	-0.094	0.021	-0.253***	Yes
4	Psychological Safety → KHB	-0.489	0.124	-0.226***	Yes
5	AI-induced Job Insecurity × Self-Efficacy in AI Learning	0.108	0.022	0.274***	Yes

Notes: ** $p < 0.01$. *** $p < 0.05$. Estimate indicates standardized coefficients. S.E. means standard error. KHB means knowledge-hiding behavior.

Table 4
Direct, Indirect, and Total Effects of the Final Research Model.

Model (Hypothesis 4)	Direct Effect	Indirect Effect	Total Effect
AI-induced Job Insecurity → Psy- chological Safety → KHB	0.313	0.057	0.370

Notes: All values are standardized. KHB means knowledge-hiding behavior.

Our analysis indicated that the bootstrapped 95% CI for the indirect effect ranged from 0.027 to 0.101, which does not include zero. This result strongly supports Hypothesis 4, confirming the mediating role of psychological safety in the relationship between AI-induced job insecurity and KHB. Table 4 presents these insights.

Findings from the moderation evaluation

To examine the moderating effect of self-efficacy in AI learning on the relationship between AI-induced job insecurity and psychological safety (Hypothesis 5), we conducted a hierarchical multiple regression analysis, following the procedures outlined by Aiken and West (1991). Our approach comprised several key steps:

First, we centered the predictor variables (AI-induced job insecurity and self-efficacy in AI learning) by deducting the mean value of each variable from all observations. This procedure mitigates multicollinearity and aids in interpreting the main effects in the presence of a significant interaction. Second, we generated an interaction term by multiplying the centered variables for AI-induced job insecurity and self-efficacy in AI learning. Third, we conducted a three-step hierarchical regression: a. Step 1: We inputted the control variables. b. Step 2: We added the centered main effect variables (AI-induced job insecurity and self-efficacy in AI learning). c. Step 3: We incorporated the interaction term. Fourth, to determine the presence of a moderation effect, we examined the statistical significance of the interaction term. Lastly, we created an interaction plot to visually represent the moderation effect, which aids in interpreting and communicating our findings (Please see Figs. 2 and 3).

The moderation analysis revealed a significant interaction effect between AI-induced job insecurity and self-efficacy in AI learning on psychological safety ($\beta = 0.274$, $p < 0.001$). This finding aligns with social cognitive theory (Bandura, 1986), which indicates that individuals with higher self-efficacy in AI learning are better equipped to maintain psychological safety even amid AI-induced job insecurity. Theoretically, this interaction effect indicates that self-efficacy in AI learning acts as a personal resource that mitigates the negative impact of job insecurity on psychological safety. This supports the job demands-resources model (Bakker & Demerouti, 2017), wherein self-efficacy functions as a resource that aids employees in coping with the demands posed by AI-induced job insecurity.

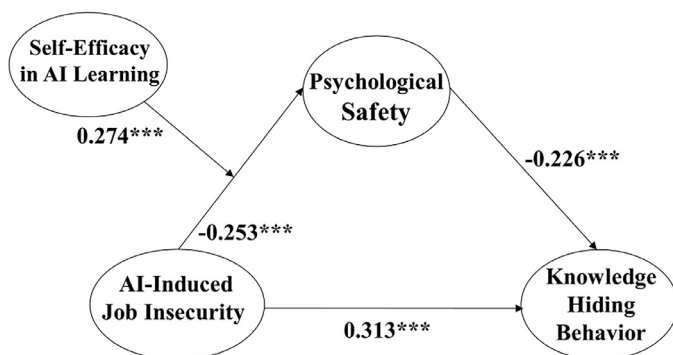


Fig. 2. Coefficient Values of our Research Model (** $p < 0.01$, *** $p < 0.001$. All values are standardized).

These results highlight the potential of enhancing employees' self-efficacy in AI learning as a strategic approach to mitigate the negative effects of AI-induced job insecurity on psychological safety. Organizations implementing AI technologies should consider investing in AI learning programs to enhance technical skills and boost employees' confidence in their ability to adapt to AI-driven changes.

Discussion

This research offers profound insights into the intricate interplay between AI-induced job insecurity, psychological safety, KHB, and self-efficacy in AI learning within organizational contexts. Our results enhance the existing body of knowledge regarding technological disruption in the workplace and expand our insight into employee behavior during AI-driven transformations.

The validation of our first hypothesis, which asserts a positive correlation between AI-induced job insecurity and KHB, extends and enhances prior research on job insecurity and its impact on employee behavior (Shoss, 2017; Connelly et al., 2012). This outcome indicates that the threat posed by AI may intensify knowledge hiding among employees; thus, it aligns with the COR theory (Hobfoll, 1989). Employees who perceive AI as a threat to their job security may resort to KHB as a defensive mechanism to safeguard their perceived value within the company. This finding broadens the application of COR theory to the dynamics of AI-driven workplace changes, demonstrating how technological progress can instigate resource-protective behaviors in unprecedented ways.

Our second hypothesis was also supported; it suggested a negative link between AI-induced job insecurity and psychological safety. This outcome contributes to the expanding literature regarding the psychological effects of technological evolution in the workplace (Edmondson & Lei, 2014; Frazier et al., 2017). The perceived threats from AI impact not only job security but also erode the psychological conditions that are essential for open communication and risk-taking within teams. This relationship underscores the extensive consequences of AI implementation that extend beyond direct job concerns, indicating a potential degradation of team dynamics and organizational culture. The result highlights the necessity for a comprehensive strategy in managing AI integration that addresses organizational change's technical and psychological facets.

The affirmation of our third hypothesis, which predicts a negative correlation between psychological safety and KHB, emphasizes the significance of psychological safety in promoting knowledge-sharing and collaborative behaviors within organizations. This finding builds upon prior research (Edmondson, 1999; Newman et al., 2017) and extends it to contexts involving AI-induced changes. It indicates that when employees experience psychological security, they are less inclined to engage in knowledge hiding behavior, even in the face of perceived threats from AI. This outcome accentuates the vital role of organizational climate in counteracting negative behaviors arising from technological uncertainties and highlights psychological safety as a safeguard against the adverse effects of AI-induced job insecurity.

The mediation analysis supporting the fourth hypothesis provides a nuanced understanding of the mechanisms by which AI-induced job insecurity affects employee behavior. By demonstrating that psychological safety partially mediates the relationship between AI-induced job insecurity and KHB, this analysis builds upon earlier research regarding the mediating role of psychological states in workplace behavior (Cheng & Chan, 2008; De Clercq et al., 2019). It reveals that although AI-induced job insecurity directly influences knowledge hiding, a substantial portion of this effect is mediated through diminished psychological safety. This insight offers a potential intervention for organizations seeking to curtail knowledge hiding in the context of AI deployment, underscoring the significance of

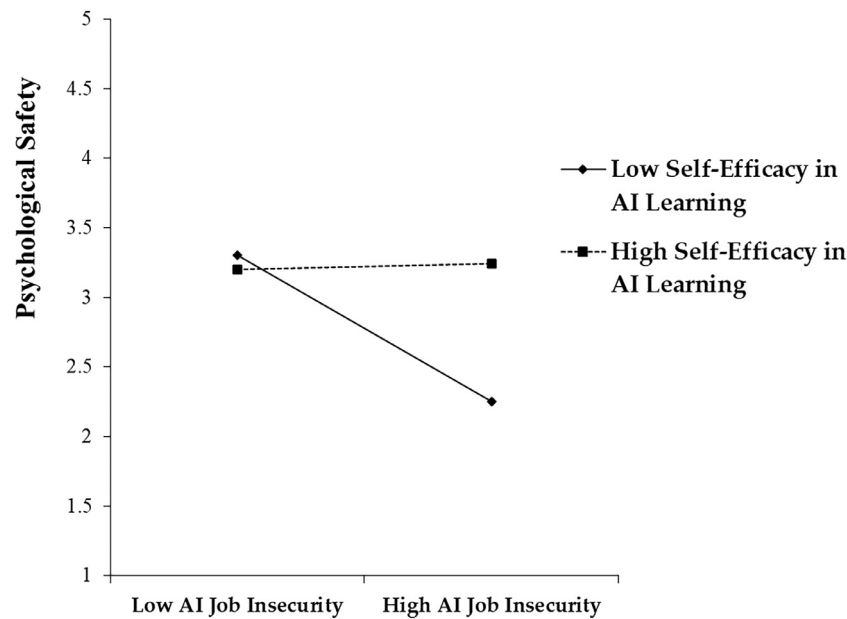


Fig. 3. Moderating Effect of Self-Efficacy in AI Learning in the AI-induced job insecurity-Psychological Safety link.

maintaining a psychologically safe environment during technological transitions.

The corroboration of our fifth hypothesis, which posited that self-efficacy in AI learning moderates the relationship between AI-induced job insecurity and psychological safety, expands social cognitive theory (Bandura, 1986) into the context of AI adoption within organizations. This finding aligns with prior research on the buffering effects of self-efficacy in high-stress work environments (Schaubroeck & Merritt, 1997) and underscores its relevance amid technological shifts. It indicates that employees with greater confidence in their ability to learn and adapt to AI technologies are more adept at preserving their psychological safety, even amid job insecurity. This result stresses the significance of nurturing AI learning self-efficacy as a strategic approach for maintaining a psychologically safe work environment during technological transitions.

Collectively, these findings enhance the understanding of employee behavior amid AI-driven workplace changes. They illuminate the interconnected dynamics of job insecurity, psychological safety, and knowledge management practices amid technological disruption. Further, they highlight the critical role of individual differences, particularly self-efficacy in AI learning, in shaping responses to these challenges.

Theoretical implications

This study offers significant theoretical implications that advance our knowledge of AI-induced job insecurity and its organizational consequences. Firstly, our research broadens the application of COR theory (Hobfoll, 1989) to the domain of workplace transformations induced by AI. We reveal a complex chain of resource-protective behaviors by conceptualizing AI-induced job insecurity as threatening employees' valued resources. This extension transcends mere resource preservation to illustrate a nuanced, multi-stage process wherein employees reassess their knowledge as a critical, differentiating resource in response to AI implementation threats. The subsequent engagement in KHB as a resource protection strategy demonstrates a more dynamic understanding of resource valuation and protection in rapidly evolving technological environments. This expanded application of COR theory the ongoing discourse regarding resource dynamics in contemporary workplaces (Halbesleben et al.,

2014), providing a more fluid and context-sensitive model of resource conservation.

Second, our study significantly advances the literature on psychological safety by establishing it as a crucial mediating mechanism between AI-induced job insecurity and KHB. This finding builds upon Edmondson's (1999) research by demonstrating the relevance of psychological safety amid technological disruption, thereby broadening its applicability beyond traditional team dynamics. By delineating a clear pathway from organizational-level technological changes to individual-level psychological states and behaviors, our research bridges the gap between macro-level organizational changes and micro-level psychological processes. It reveals how reduced psychological safety can trigger a negative spiral, where insecurity results in decreased safety, subsequently promoting knowledge hiding, which may further erode team trust and safety. This mediation model provides a deeper understanding of how organizational climate factors are affected by and respond to external technological threats.

Third, our research contributes significantly to social cognitive theory (Bandura, 1986) by showing self-efficacy's moderating role in AI learning. The finding highlights the significance of domain-specific self-efficacy in technologically dynamic environments, transcending general self-efficacy concepts. We demonstrate how high self-efficacy in AI learning can function as a psychological shield against the adverse consequences of AI-induced employment uncertainty, revealing its role as a personal resource amid technological threats. This insight indicates that self-efficacy in AI learning may be a key factor in employee adaptability, potentially influencing an individual's capacity to succeed in AI-driven workplaces. This contribution enhances our comprehension of how individual differences in capability beliefs can shape responses to technological threats, enriching the literature on employee adaptability (Pulakos et al., 2000) by identifying a specific form of self-efficacy that is particularly relevant in AI adoption.

Lastly, our moderated mediation model offers a comprehensive, multi-level framework that integrates COR, social exchange, and social cognitive theories. This integration advances organizational behavior research by providing a model that encompasses organizational, team, and individual levels, addressing the demands for increasingly complex, multi-level approaches in organizational research (Kozlowski & Klein, 2000). The model illustrates the

dynamic interactions between contextual factors (AI implementation), psychological states (job insecurity, psychological safety), individual differences (self-efficacy), and behaviors (knowledge hiding). It implicitly incorporates a temporal dimension, indicating how initial AI implementation can result in a sequence of psychological and behavioral responses over time. This comprehensive paradigm provides a more comprehensive perspective for analyzing the intricate interplay between organizational context, individual perceptions, and behavioral outcomes in rapidly evolving work environments. Consequently, it provides a theoretical foundation for understanding the complex, frequently non-linear effects of technological change on employee behavior.

Practical implications

This paper provides substantial practical implications for senior leadership teams, practitioners, and practitioners tackling the obstacles of AI adoption in their organizations.

Firstly, our findings underscore the essential need to proactively manage employee perceptions of job security during AI implementation. Top management teams should develop comprehensive communication strategies that explicitly articulate the organization's vision for AI integration, emphasizing how AI will augment rather than replace human capabilities. This approach aligns with [Rahmani et al.'s \(2023\)](#) research, which indicate that transparent communication regarding technological changes can significantly reduce employee anxiety and resistance. To address employee concerns and provide clarity on how AI will impact job roles, leaders should consider conducting regular town hall meetings, workshops, and one-on-one sessions. Additionally, organizations may consider adopting a "human-AI collaboration" framework, as proposed by [Wilson and Daugherty \(2018\)](#). This framework emphasizes the complementary strengths of human workers and AI systems. By framing AI implementation as an opportunity for employee upskilling and role enhancement instead of a threat, organizations can mitigate the negative effects of AI-induced job insecurity on knowledge-sharing behaviors and overall productivity.

Secondly, our research highlights the vital role of psychological safety in mediating the AI-induced job insecurity-KHB link. Practitioners should prioritize the creation and maintenance of a psychologically safe work environment, particularly during periods of technological transition. This can be achieved via several targeted interventions. For instance, leaders should model vulnerability and a willingness to learn, fostering an environment where employees can express concerns and share ideas without fear of repercussion. Implementing structured feedback mechanisms, such as "blameless post-mortems" after project completions or AI integrations, can foster a culture of continuous learning and improvement. Furthermore, organizations may consider establishing cross-functional AI task forces, including employees from various levels and departments. This inclusive approach can enhance psychological safety by demonstrating that all perspectives are valued in the AI integration process. Consequently, this may reduce KHBs and promote collaborative problem-solving.

Thirdly, our findings regarding the moderating role of self-efficacy in AI learning offer valuable insights for talent development and training strategies. Organizations should invest in comprehensive AI literacy programs customized for various job roles and skill levels. These programs should not only focus on technical skills but also the enhancement of members' confidence in their ability to learn and adapt to AI technologies. Based on [Bandura's \(1997\)](#) sources of self-efficacy, these programs can incorporate mastery experiences (hands-on AI projects), vicarious experiences (showcasing success stories of peers adapting to AI), verbal persuasion (positive reinforcement from leaders), and management of physiological states (stress management techniques). Additionally, organizations may consider

implementing AI mentorship programs, which entails pairing AI-proficient employees with those less confident in their AI skills. This peer-to-peer learning approach can enhance self-efficacy and facilitate knowledge sharing within the organization ([Noe et al., 2014](#)).

Lastly, our research highlights the need for a holistic approach to change management during AI implementation. Top management teams should regard AI integration not only as a technological upgrade but also as a complex organizational change process that impacts employee psychology, team dynamics, and organizational culture. To this end, organizations may consider adopting a socio-technical systems approach to AI implementation, as proposed by [Pasmore et al. \(2019\)](#). This approach emphasizes the interdependence of technical and social systems within organizations and advocates for the concurrent design of both systems during technological changes. Practically, this may entail establishing cross-functional teams tasked with both the technical implementation of AI and the management of its social implications. These teams should formulate integrated strategies that address technological aspects (e.g., AI system design, data management) alongside human factors (e.g., job redesign, skill development, cultural adaptation). By concurrently addressing technical and human elements, organizations can create a more resilient and adaptive workforce that can effectively leverage AI technologies while minimizing negative psychological and behavioral outcomes.

Limitations and suggestions for future research

Although this study offers valuable insights into the connections between AI-induced job instability, psychological safety, KHB, and self-efficacy in AI learning, recognizing certain limitations and proposing recommendations for future research is crucial.

Firstly, despite utilizing a 3-wave time-lagged design, the data we collected can be considered cross-sectional. This limitation hinders our ability to definitely demonstrate causal correlations. This snapshot approach may not fully capture the dynamic nature of technological change and its impact on members' attitudes and behaviors. [Podsakoff et al. \(2019\)](#) noted that cross-sectional designs can result in inflated correlations and ambiguous causal ordering. To address this limitation, future research should consider longitudinal designs that track changes in AI-induced job insecurity, psychological safety, and KHB over extended periods. Such studies can offer deeper perspectives on the temporal dynamics of these relationships and potentially uncover reciprocal effects or feedback loops. For instance, researchers should employ latent growth modeling or cross-lagged panel designs to examine the influence of initial levels of KHB on subsequent perceptions of job insecurity or psychological safety ([Little, 2013](#)).

Second, the scope of our study was limited to South Korea, potentially restricting the applicability of our findings to diverse cultural settings. Cultural elements such as power distance, uncertainty avoidance, and collectivism vs. individualism may significantly influence our observed relationships. [Taras et al. \(2010\)](#) demonstrated in their meta-analysis that cultural values can significantly moderate relationships in organizational behavior. To assess the robustness of these relationships across various cultural contexts, future studies should replicate this research in diverse cultural settings. Cross-cultural comparisons may provide valuable insights into how cultural factors moderate the relationships between AI-induced job insecurity, psychological safety, and KHB. Additionally, multi-country studies may assist in identifying universal aspects of employee responses to AI integration versus culture-specific reactions, following the approach proposed by [Tsui et al. \(2007\)](#) for cross-national organizational research.

Third, the study exclusively employed self-reported measures, which were susceptible to common method bias and social desirability effects. This is particularly pertinent for sensitive issues such as

job insecurity and KHB, where respondents may be inclined to provide socially desirable responses. Podsakoff et al. (2012) highlight that common method variance can significantly influence observed relationships in organizational research. Future research may enhance its rigor by incorporating multiple data sources, such as supervisor ratings of KHB, objective measures of AI implementation, or peer assessments of psychological safety. Additionally, experimental designs or scenario-based studies, as proposed by Aguinis and Bradley (2014), may aid in mitigating some of the limitations associated with self-reported measures and provide more robust evidence for the causal relationships posited in our model.

Fourth, despite focusing on AI-induced job insecurity and self-efficacy in AI learning, this study did not explore other potentially relevant AI-related variables, such as attitudes towards AI, perceived usefulness of AI, or the extent of AI integration in specific job roles. Recent research by Gursoy et al. (2019) has emphasized the significance of considering a wider array of AI-related perceptions within organizational contexts. Future research should expand the scope of AI-related variables to provide an enhanced comprehension of the influence of AI on employee behavior. For instance, studies should examine how different types or levels of AI integration in various job roles uniquely influence job insecurity and subsequent behaviors, building upon the findings of Makarius et al. (2020) regarding AI-augmented human resource management. Additionally, exploring employees' past experiences with AI or their general technological readiness may yield valuable insights into individual differences in responses to AI-induced job insecurity.

Fifth, this study focused on KHB as an outcome variable. However, AI-induced job insecurity may influence a broader range of work behaviors and attitudes. As Shoss (2017) suggested in her integrative review of job insecurity, the consequences of an unstable job may be far-reaching and diverse. Future studies should explore the possibility of broadening the scope of outcome variables to encompass other relevant behaviors such as innovative work behavior, organizational citizenship behavior, or counterproductive work behavior. Additionally, examining the impact on job satisfaction, organizational commitment, or turnover intentions may provide a deeper understanding of the results of AI-induced job insecurity. Such an approach would align with recent calls for more comprehensive examinations of employee responses to technological change (Brougham & Haar, 2018).

Sixth, although our study considered individual-level variables and team-level psychological safety, it did not extensively explore organizational-level factors that could influence the observed relationships. Kozlowski and Klein (2000) highlighted that organizational phenomena frequently involve processes that span multiple levels of analysis. Future research should incorporate organizational-level variables such as organizational culture, leadership styles, or formal AI policies and practices. Multi-level studies that concurrently examine individual, team, and organizational factors may yield a more nuanced comprehension of the interactions between different levels of analysis in influencing employee responses to AI integration. For instance, researchers should explore how organizational transparency regarding AI implementation plans moderates the AI-induced job insecurity-psychological safety link, drawing on recent studies of organizational communication during technological change (Bordia et al., 2021).

Seventh, although using an online research platform offered numerous advantages regarding reach and efficiency, we acknowledge potential limitations regarding sample representativeness. Online sampling may introduce biases, including self-selection bias and coverage bias (Bethlehem, 2010). To address these concerns and enhance the representativeness of our sample, we implemented the following measures:

1. Stratified Random Sampling: We employed a stratified random sampling technique to guarantee representation across key

demographic variables such as age, gender, education level, and industry sector. This approach alleviated potential biases associated with the overrepresentation of certain groups in online panels (Terhanian et al., 2016).

2. Quota Controls: We implemented quota controls to align our sample's demographic composition with that of the general working population in South Korea, utilizing recent labor force statistics.
3. Screening Questions: We utilized carefully designed screening questions to verify that participants satisfied our inclusion criteria, particularly concerning their exposure to AI technologies in their workplaces.
4. Data Quality Checks: To identify and omit low-quality responses, we implemented rigorous data quality checks, such as attention checks, speeding checks, and consistency checks (Meade & Craig, 2012).
5. Multiple Contact Methods: To mitigate coverage bias, we employed multiple contact methods (e.g., email, SMS) on the online platform to engage potential participants.

Despite these measures, we acknowledge that certain limitations persist. Our sample may underrepresent individuals with limited internet access or those less comfortable with online surveys. Additionally, using an online platform may have resulted in a sample with higher digital literacy than the general population.

Scholars should address these constraints and explore the suggested areas for future work to further enhance their understanding of employee behavior in the workplace when it comes to integrating AI. Such research will be essential for organizations to navigate the challenges and opportunities posed by AI technologies while ensuring a productive and psychologically healthy workforce.

Conclusion

This research offers significant insights into the intricate dynamics of AI-induced job insecurity and its effects on employee behavior within organizations. This study analyzes the interconnections between AI-induced job insecurity, psychological safety, knowledge-hiding behavior, and self-efficacy in AI learning, thereby enhancing the comprehension of the challenges organizations encounter in the era of AI integration.

Our findings reveal that AI-induced job insecurity contributes to increased knowledge-hiding behavior. This behavior occurs both directly and indirectly due to diminished psychological safety. This underscores the potential adverse effects of AI implementation on organizational knowledge sharing and collaborative efforts. Additionally, our results indicate that self-efficacy in AI learning serves as a protective factor, cushioning the detrimental effects of AI-induced job insecurity on psychological safety.

This study contributes to the literature by integrating self-efficacy in AI learning as a significant moderating variable in the relationship between AI-induced job insecurity and psychological safety. Unlike previous research that predominantly concentrated on general self-efficacy or computer self-efficacy, our study presents a domain-specific construct that is particularly relevant in the context of AI implementation. This approach provides a more nuanced comprehension of how individual differences in AI-related capabilities can influence employees' responses to technological changes in the workplace.

Our findings reveal that self-efficacy in AI learning mitigates the negative effects of AI-induced job insecurity on psychological safety. This insight extends social cognitive theory (Bandura, 1986) and the job demands-resources model (Bakker & Demerouti, 2017) by demonstrating how a distinct type of self-efficacy can serve as a personal resource amid technological uncertainties. This novel perspective offers both theoretical advancements and practical implications for overseeing the human aspects of AI integration with organizations.

Future research should explore the role of self-efficacy in AI learning across diverse organizational contexts. Specifically, researchers may investigate:

1. The development trajectory of self-efficacy in AI learning over time, examining the impact of exposure to AI technologies and training interventions on its growth.
2. The potential reciprocal relationships between self-efficacy in AI learning and actual AI-related performance, exploring how initial beliefs influence behavior and how successful experiences reinforce self-efficacy.
3. The generalizability of our findings across various cultural contexts, industries, and levels of AI implementation to understand the boundary conditions of the moderating effect of self-efficacy in AI learning.
4. The impact of organizational interventions aimed at enhancing self-efficacy in AI learning, assessing their effectiveness in mitigating the negative effects of AI-induced job insecurity.
5. The influence of leadership styles and organizational culture on fostering self-efficacy in AI learning and its subsequent effects on employee attitudes and behaviors during AI implementation.

These prospective research avenues would expand upon our findings and enhance the understanding of the psychological processes involved in adapting to AI technologies in the workplace.

In conclusion, as AI transforms workplace environments, understanding and mitigating its psychological impacts on employees becomes imperative for organizational success. Organizations can create an environment that adapts to and thrives amid technological advancements by cultivating psychological safety, boosting self-efficacy in AI learning, and addressing concerns related to job insecurity. This research establishes a foundation for future studies and practical measures aimed at navigating the intricate terrain of AI integration in professional settings.

Declaration of generative AI and AI-assisted technologies in the writing process

During the preparation of this work the author(s) used Claude 3.5 Sonnet in order to proofread. After using this tool/service, the author (s) reviewed and edited the content as needed and take(s) full responsibility for the content of the publication.

Contributors

Dr. Byung-Jik Kim contributed by writing the original draft of the manuscript and in the conceptualization, data collection, formal analysis, and methodology. Dr. Min-Jik Kim contributed to the conceptualization, analysis, revision, and in editing the manuscript. All authors not only contributed to the writing of the final manuscript, but also have read and agreed to the published version of the manuscript. All members of this research team contributed to the management or administration of this paper.

Ethics declarations

Data availability statements

The authors affirm that all data produced or examined during this study are incorporated in this published publication.

Ethical approval

The survey conducted in this study adhered to the standards outlined in the Declaration of Helsinki. The Institutional Review Board of

Yonsei University granted ethical approval (No:202208-HR-2975-01).

Informed consent

Every participant was provided with information about the objective and extent of this study, as well as the intended utilization of the data. The respondents willingly and freely participated in the study, with their identities kept confidential and their involvement being optional. Prior to participating in the survey, all subjects included in the study provided informed consent.

Declaration of competing interest

The authors declare no conflict of interest.

CRediT authorship contribution statement

Byung-Jik Kim: Writing – original draft, Software, Methodology, Formal analysis, Data curation, Conceptualization. **Min-Jik Kim:** Writing – review & editing, Validation, Supervision, Project administration, Methodology, Investigation, Formal analysis, Conceptualization.

Funding/Acknowledgement

None.

REFERENCES

- Aguinis, H., & Bradley, K. J. (2014). Best practice recommendations for designing and implementing experimental vignette methodology studies. *Organizational Research Methods*, 17(4), 351–371.
- Aiken, L. S., & West, S. G. (1991). *Multiple regression: Testing and interpreting interactions*. Sage Publications.
- Ajzen, I. (1991). The theory of planned behavior. *Organizational Behavior and Human Decision Processes*, 50(2), 179–211.
- AlQershi, N., Saufi, R. B. A., Yaziz, M. F. B. A., Permarupan, P. Y., Muhammad, N. M. N., Yusoff, M. N. H. B., & Ramayah, T. (2023). The threat of robots to career sustainability, and the pivotal role of knowledge management and human capital. *Journal of Innovation & Knowledge*, 8(3), 100386.
- Anderson, J. C., & Gerbing, D. W. (1988). Structural equation modeling in practice: A review and recommended two-step approach. *Psychological Bulletin*, 103(3), 411.
- Argyris, C., & Schön, D. A. (1978). *Organizational learning: A theory of action perspective*. Addison-Wesley.
- Bakker, A. B., & Demerouti, E. (2017). Job demands–resources theory: Taking stock and looking forward. *Journal of Occupational Health Psychology*, 22(3), 273–285.
- Bandura, A. (1986). *Social foundations of thought and action: A social cognitive theory*. Prentice-Hall.
- Bandura, A. (1997). *Self-efficacy: The exercise of control*. W.H. Freeman and Company.
- Bankins, S., Ocampo, A. C., Marrone, M., Restubog, S. L. D., & Woo, S. E. (2024). A multi-level review of artificial intelligence in organizations: Implications for organizational behavior research and practice. *Journal of Organizational Behavior*, 45(2), 159–182.
- Berger, C. R., & Calabrese, R. J. (1975). Some explorations in initial interaction and beyond: Toward a developmental theory of interpersonal communication. *Human Communication Research*, 1(2), 99–112.
- Bethlehem, J. (2010). Selection bias in web surveys. *International Statistical Review*, 78(2), 161–188.
- Blau, P. M. (1964). *Exchange and power in social life*. John Wiley & Sons.
- Bogilović, S., Černe, M., & Škerlavaj, M. (2017). Hiding behind a mask? Cultural intelligence, knowledge hiding, and individual and team creativity. *European Journal of Work and Organizational Psychology*, 26(5), 710–723.
- Bordia, P., Restubog, S. L. D., Jimmieson, N. L., & Irmer, B. E. (2021). Haunted by the past: Effects of poor change management history on employee attitudes and turnover. *Group & Organization Management*, 36(2), 191–222.
- Borges, A. F., Laurindo, F. J., Spínola, M. M., Gonçalves, R. F., & Mattos, C. A. (2021). The strategic use of artificial intelligence in the digital era: Systematic literature review and future research directions. *International journal of information management*, 57, 102225.
- Brace, N., Kemp, R., & Snelgar, R. (2003). *A guide to data analysis using SPSS for windows*. New York: Palgrave Macmillan.
- Brougham, D., & Haar, J. (2018). Smart technology, artificial intelligence, robotics, and algorithms (STARA): Employees' perceptions of our future workplace. *Journal of Management & Organization*, 24(2), 239–257.

- Brougham, D., & Haar, J. (2020). Technological disruption and employment: The influence on job insecurity and turnover intentions: A multi-country study. *Technological Forecasting and Social Change*, 161, 120276.
- Budhwar, P., Malik, A., De Silva, M. T., & Thevisuthan, P. (2022). Artificial intelligence – challenges and opportunities for international HRM: a review and research agenda. *The International Journal of Human Resource Management*, 33(6), 1065–1097.
- Califf, C. B., Sarker, S., & Sarker, S. (2020). The bright and dark sides of technostress: A mixed-methods study involving healthcare IT. *MIS Quarterly*, 44(2), 809–856.
- Černe, M., Nerstad, C. G., Dysvik, A., & Škerlavaj, M. (2014). What goes around comes around: Knowledge hiding, perceived motivational climate, and creativity. *Academy of Management Journal*, 57(1), 172–192.
- Cheng, G. H. L., & Chan, D. K. S. (2008). Who suffers more from job insecurity? A meta-analytic review. *Applied Psychology*, 57(2), 272–303.
- Climent, R. C., Haftor, D. M., & Staniewski, M. W. (2024). AI-enabled business models for competitive advantage. *Journal of Innovation & Knowledge*, 9(3), 100532.
- Compeau, D. R., & Higgins, C. A. (1995). Computer self-efficacy: Development of a measure and initial test. *MIS Quarterly*, 19(2), 189–211.
- Connelly, C. E., & Zweig, D. (2015). How perpetrators and targets construe knowledge hiding in organizations. *European Journal of Work and Organizational Psychology*, 24(3), 479–489.
- Connelly, C. E., Zweig, D., Webster, J., & Trougakos, J. P. (2012). Knowledge hiding in organizations. *Journal of Organizational Behavior*, 33(1), 64–88.
- Cropanzano, R., & Mitchell, M. S. (2005). Social exchange theory: An interdisciplinary review. *Journal of Management*, 31(6), 874–900.
- De Clercq, D., Haq, I. U., & Azeem, M. U. (2019). Time-related work stress and counterproductive work behavior: Invigorating roles of deviant personality traits. *Personnel Review*, 48(7), 1756–1781.
- Demerouti, E., Bakker, A. B., Nachreiner, F., & Schaufeli, W. B. (2001). The job demands-resources model of burnout. *Journal of Applied Psychology*, 86(3), 499–512.
- Dweck, C. S. (2006). *Mindset: The new psychology of success*. Random House.
- Edmondson, A. (1999). Organizational safety and learning behavior in work teams. *Administrative Science Quarterly*, 44(2), 350–383.
- Edmondson, A. C., & Bransby, D. P. (2023). Psychological safety comes of age: Observed themes in an established literature. *Annual Review of Organizational Psychology and Organizational Behavior*, 10(1), 55–78.
- Edmondson, A. C., & Lei, Z. (2014). Psychological safety: The history, renaissance, and future of an interpersonal construct. *Annual Review of Organizational Psychology and Organizational Behavior*, 1(1), 23–43.
- Eisenberger, R., Huntington, R., Hutchison, S., & Sowa, D. (1986). Perceived organizational support. *Journal of Applied Psychology*, 71(3), 500–507.
- Fauzi, M. A. (2023). Knowledge hiding behavior in higher education institutions: A scientometric analysis and systematic literature review approach. *Journal of Knowledge Management*, 27(2), 302–327.
- Felicetti, A. M., Cimino, A., Mazzoleni, A., & Ammirato, S. (2024). Artificial intelligence and project management: An empirical investigation on the appropriation of generative Chatbots by project managers. *Journal of Innovation & Knowledge*, 9(3), 100545.
- Fong, P. S., Men, C., Luo, J., & Jia, R. (2018). Knowledge hiding and team creativity: The contingent role of task interdependence. *Management Decision*, 56(2), 329–343.
- Frazier, M. L., Fainshmidt, S., Klinger, R. L., Pezeshkan, A., & Vacheva, V. (2017). Psychological safety: A meta-analytic review and extension. *Personnel Psychology*, 70(1), 113–165.
- Frey, C. B., & Osborne, M. A. (2017). The future of employment: How susceptible are jobs to computerisation? *Technological Forecasting and Social Change*, 114, 254–280.
- Ghislieri, C., Molino, M., & Cortese, C. G. (2018). Work and organizational psychology looks at the fourth industrial revolution: How to support workers and organizations? *Frontiers in Psychology*, 9, 2365.
- Goran, J., LaBerge, L., & Srinivasan, R. (2017). Culture for a digital age. *McKinsey Quarterly*, 10, 1–10.
- Gouldner, A. W. (1960). The norm of reciprocity: A preliminary statement. *American Sociological Review*, 25(2), 161–178.
- Gursoy, D., Chi, O. H., Lu, L., & Nunkoo, R. (2019). Consumers acceptance of artificially intelligent (AI) device use in service delivery. *International Journal of Information Management*, 49, 157–169.
- Halbesleben, J. R., Neveu, J. P., Paustian-Underdahl, S. C., & Westman, M. (2014). Getting to the "COR": Understanding the role of resources in conservation of resources theory. *Journal of Management*, 40(5), 1334–1364.
- Hatlevik, O. E., Throndsen, I., Loi, M., & Gudmundsdottir, G. B. (2018). Students' ICT self-efficacy and computer and information literacy: Determinants and relationships. *Computers & Education*, 118, 107–119.
- Hobfoll, S. E. (1989). Conservation of resources: A new attempt at conceptualizing stress. *American Psychologist*, 44(3), 513–524.
- Hobfoll, S. E., Halbesleben, J., Neveu, J. P., & Westman, M. (2018). Conservation of resources in the organizational context: The reality of resources and their consequences. *Annual Review of Organizational Psychology and Organizational Behavior*, 5, 103–128.
- Hogg, M. A., & Vaughan, G. M. (2018). *Social Psychology* (8th ed.). Pearson.
- Jeong, J., Kim, B. J., & Kim, M. J. (2022). The impact of job insecurity on knowledge-hiding behavior: The mediating role of organizational identification and the buffering role of coaching leadership. *International Journal of Environmental Research and Public Health*, 19(23), 16017.
- Jeong, J., Kim, B. J., & Lee, J. (2023). The effect of job insecurity on knowledge hiding behavior: The mediation of psychological safety and the moderation of servant leadership. *Frontiers in Public Health*, 11, 1108881.
- Johns, G. (2006). The essential impact of context on organizational behavior. *Academy of Management Review*, 31(2), 386–408.
- Kahn, W. A. (1990). Psychological conditions of personal engagement and disengagement at work. *Academy of Management Journal*, 33(4), 692–724.
- Kim, B. J., & Kim, M. J. (2024). The influence of work overload on cybersecurity behavior: A moderated mediation model of psychological contract breach, burnout, and self-efficacy in AI learning such as ChatGPT. *Technology in Society*, 77, 102543.
- Kim, B. J., Jeon, N., Sohn, H., Lee, N., & Kim, M. J. (2024). *The impact of corporate social responsibility on employee burnout: The crucial role of work overload*. Corporate Social Responsibility and Environmental Management Advanced online publication.
- Kim, B. J., Park, S., & Kim, T. H. (2019). The effect of transformational leadership on team creativity: Sequential mediating effect of employee's psychological safety and creativity. *Asian Journal of Technology Innovation*, 27(1), 90–107.
- Kim, M. J., Chung, N., Lee, C. K., & Preis, M. W. (2021). The impact of innovation and gratification on authentic experience, subjective well-being, and behavioral intention in tourism virtual reality: The moderating role of technology readiness. *Telecommunications and Informatics*, 49, 101349.
- Ključnikov, A., Popkova, E. G., & Sergi, B. S. (2023). Global labour markets and workplaces in the age of intelligent machines. *Journal of Innovation & Knowledge*, 8(4), 100407.
- Kozlowski, S. W. J., & Klein, K. J. (2000). A multilevel approach to theory and research in organizations: Contextual, temporal, and emergent processes. In Klein, K. J., & Kozlowski, S. W. J. (Eds.), *Multilevel theory, research, and methods in organizations: Foundations, extensions, and new directions* (pp. 3–90). Jossey-Bass.
- Kraimer, M. L., Wayne, S. J., Liden, R. C., & Sparrowe, R. T. (2005). The role of job security in understanding the relationship between employees' perceptions of temporary workers and employees' performance. *Journal of Applied Psychology*, 90(2), 389–398.
- Kraus, S., Ferraris, A., & Bertello, A. (2023). The future of work: How innovation and digitalization re-shape the workplace. *Journal of Innovation & Knowledge*, 8(4), 100438.
- Kurtessiz, J. N., Eisenberger, R., Ford, M. T., Buffardi, L. C., Stewart, K. A., & Adis, C. S. (2017). Perceived organizational support: A meta-analytic evaluation of organizational support theory. *Journal of Management*, 43(6), 1854–1884.
- Landers, R. N., & Behrend, T. S. (2015). An inconvenient truth: Arbitrary distinctions between organizational, Mechanical Turk, and other convenience samples. *Industrial and Organizational Psychology*, 8(2), 142–164.
- Li, P., Bastone, A., Mohamad, T. A., & Schiavone, F. (2023). How does artificial intelligence impact human resources performance: evidence from a healthcare institution in the United Arab Emirates. *Journal of Innovation & Knowledge*, 8(2), 100340.
- Little, T. D. (2013). *Longitudinal structural equation modeling*. Guilford Press.
- Makarius, E. E., Larson, B. Z., & Vroman, S. R. (2020). The new world of work: Job search and career management in the age of AI and digitization. *Organizational Dynamics*, 50(1), 100794.
- Meade, A. W., & Craig, S. B. (2012). Identifying careless responses in survey data. *Psychological Methods*, 17(3), 437–455.
- Mitchell, T. R., & James, L. R. (2001). Building better theory: Time and the specification of when things happen. *Academy of Management Review*, 26(4), 530–547.
- Mowday, R. T., & Sutton, R. I. (1993). Organizational behavior: Linking individuals and groups to organizational contexts. *Annual Review of Psychology*, 44(1), 195–229.
- Nam, T. (2019). Technology usage, expected job sustainability, and perceived job insecurity. *Technological Forecasting and Social Change*, 138, 155–165.
- Newman, A., Donohue, R., & Eva, N. (2017). Psychological safety: A systematic review of the literature. *Human Resource Management Review*, 27(3), 521–535.
- Noe, R. A., Clarke, A. D., & Klein, H. J. (2014). Learning in the twenty-first-century workplace. *Annual Review of Organizational Psychology and Organizational Behavior*, 1(1), 245–275.
- Obreja, D. M., Rughinis, R., & Rosner, D. (2024). Mapping the conceptual structure of innovation in artificial intelligence research: A bibliometric analysis and systematic literature review. *Journal of Innovation & Knowledge*, 9(1), 100465.
- Parker, S. K., Bindl, U. K., & Strauss, K. (2010). Making things happen: A model of proactive motivation. *Journal of Management*, 36(4), 827–856.
- Pasmore, W., Winby, S., Mohrman, S. A., & Vanasse, R. (2019). Reflections: Sociotechnical systems design and organization change. *Journal of Change Management*, 19(2), 67–85.
- Pereira, V., Hadjilias, E., Christofi, M., & Vrontis, D. (2023). A systematic literature review on the impact of artificial intelligence on workplace outcomes: A multi-process perspective. *Human Resource Management Review*, 33(1), 100857.
- Podsakoff, P. M., MacKenzie, S. B., & Podsakoff, N. P. (2012). Sources of method bias in social science research and recommendations on how to control it. *Annual Review of Psychology*, 63, 539–569.
- Podsakoff, P. M., MacKenzie, S. B., & Podsakoff, N. P. (2019). Construct measurement and validation procedures in MIS and behavioral research: Integrating new and existing techniques. *MIS Quarterly*, 43(2), 609–642.
- 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(5), 879–903.
- Pulakos, E. D., Arad, S., Donovan, M. A., & Plamondon, K. E. (2000). Adaptability in the workplace: Development of a taxonomy of adaptive performance. *Journal of Applied Psychology*, 85(4), 612–624.
- Rahmani, D., Zeng, C., Chen, M. H., Fletcher, P., & Goke, R. (2023). Investigating the effects of online communication apprehension and digital technology anxiety on organizational dissent in virtual teams. *Computers in Human Behavior*, 144, 107719.
- Rousseau, D. M., & McLean Parks, J. (1993). The contracts of individuals and organizations. In L. Cummings & B. Staw (Eds.), *Research in Organizational Behavior*: (Vol. 15, pp. 1–43). Greenwich, CT: JAI Press.

- Schaubroeck, J., & Merritt, D. E. (1997). Divergent effects of job control on coping with work stressors: The key role of self-efficacy. *Academy of Management Journal*, 40(3), 738–754.
- Scherer, R., Siddiq, F., & Tondeur, J. (2019). The technology acceptance model (TAM): A meta-analytic structural equation modeling approach to explaining teachers' adoption of digital technology in education. *Computers & Education*, 128, 13–35.
- Shin, Y., & Hur, W. M. (2019). Linking flight attendants' job crafting and OCB from a JD-R perspective: A daily analysis of the mediation of job resources and demands. *Journal of Air Transport Management*, 79, 101681.
- Shoss, M. K. (2017). Job insecurity: An integrative review and agenda for future research. *Journal of Management*, 43(6), 1911–1939.
- Shrout, P. E., & Bolger, N. (2002). Mediation in experimental and nonexperimental studies: new procedures and recommendations. *Psychological Methods*, 7(4), 422–445.
- Shu, Q., Tu, Q., & Wang, K. (2011). The impact of computer self-efficacy and technology dependence on computer-related technostress: A social cognitive theory perspective. *International Journal of Human-Computer Interaction*, 27(10), 923–939.
- Škerlavaj, M., Connelly, C. E., Cerne, M., & Dysvik, A. (2018). Tell me if you can: Time pressure, prosocial motivation, perspective taking, and knowledge hiding. *Journal of Knowledge Management*, 22(7), 1489–1509.
- Sverke, M., Låstad, L., Hellgren, J., Richter, A., & Näswall, K. (2019). A meta-analysis of job insecurity and employee performance: Testing temporal aspects, rating source, welfare regime, and union density as moderators. *International Journal of Environmental Research and Public Health*, 16(14), 2536.
- Taras, V., Kirkman, B. L., & Steel, P. (2010). Examining the impact of culture's consequences: A three-decade, multilevel, meta-analytic review of Hofstede's cultural value dimensions. *Journal of Applied Psychology*, 95(3), 405–439.
- Terhanian, G., Bremer, J., Olmsted, J., & Guo, J. (2016). A process for developing an optimal model for reducing bias in nonprobability samples: The quest for accuracy continues in online survey research. *Journal of Advertising Research*, 56(1), 14–24.
- Tsui, A. S., Nifadkar, S. S., & Ou, A. Y. (2007). Cross-national, cross-cultural organizational behavior research: Advances, gaps, and recommendations. *Journal of Management*, 33(3), 426–478.
- Wang, S., & Noe, R. A. (2010). Knowledge sharing: A review and directions for future research. *Human Resource Management Review*, 20(2), 115–131.
- Weiss, H. M., & Cropanzano, R. (1996). Affective events theory: A theoretical discussion of the structure, causes and consequences of affective experiences at work. *Research in Organizational Behavior*, 18, 1–74.
- Wilson, H. J., & Daugherty, P. R. (2018). Collaborative intelligence: Humans and AI are joining forces. *Harvard Business Review*, 96(4), 114–123.
- Wu, T. J., Li, J. M., & Wu, Y. J. (2022). Employees' job insecurity perception and unsafe behaviours in human-machine collaboration. *Management Decision*, 60(9), 2409–2432.
- Xue, Z., Hou, Y., Cao, G., & Sun, G. (2024). How does digital transformation drive innovation in Chinese agribusiness: Mechanism and micro evidence. *Journal of Innovation & Knowledge*, 9(2), 100489.
- Zaheer, S., Albert, S., & Zaheer, A. (1999). Time scales and organizational theory. *Academy of Management Review*: (pp. 725–741)24.
- Zirar, A., Ali, S. I., & Islam, N. (2023). Worker and workplace Artificial Intelligence (AI) coexistence: Emerging themes and research agenda. *Technovation*, 124, 102747.