

Leader AI Symbolization, AI Threat Perception, and Employee Well-Being: The Moderating Role of Employee Openness

Huijuan Deng

Taylor's University, Selangor 47500, Malaysia

ABSTRACT

Grounded in the framework of the signaling theory, our research examined the relationship between leader AI symbolization and employee well-being. We used a questionnaire survey ($N = 299$) to test all hypotheses. The results show that leader AI symbolization is positively related to employee well-being via employee AI threat perception. Moreover, employee openness moderates this relationship such that the effect is stronger when employee openness is high (vs. low). Our findings offer novel insights into the communicative and emotional labor required of leaders in the AI era, advancing both theoretical understanding and practical strategies for leading technological change.

KEYWORDS

AI symbolization; AI threat perception; Employee well-being; Employee openness

1 Introduction and Research Background

As artificial intelligence (AI) technologies continue to permeate organizational processes and decision-making, questions concerning their impact on employees have drawn increasing scholarly and managerial attention. While AI holds the promise of enhancing efficiency, innovation, and strategic foresight, it also evokes uncertainty and perceived threats among employees, particularly regarding job security, autonomy, and role identity (Yogeeswaran et al., 2016; Zhang et al., 2023). Amid this growing concern, organizational leaders are uniquely positioned to shape how employees interpret and react to AI-related change. However, existing research has largely focused on leaders' instrumental behaviors—such as championing AI adoption or allocating resources—while overlooking the symbolic and communicative dimensions of leadership that may shape employees' psychological responses (Rai et al., 2019; Bareis & Katzenbach, 2022).

To address this gap, the present research draws on signaling theory (Spence, 1973) to examine how leader AI symbolization—defined as the intentional use of symbolic behaviors and discourse to frame AI as beneficial and aligned with organizational values—influences employees' perceptions and well-being. In contexts of uncertainty, such as technological transformation, employees face information asymmetries and rely on leadership signals to interpret organizational intentions and forecast future outcomes (Taj, 2016; Morandini et al., 2023). We propose that leader AI symbolization acts as a salient signal that can reduce employees' perception of AI as a threat, and thereby enhance their psychological well-being.

Furthermore, this research explores how individual differences moderate the reception and effectiveness of these symbolic signals. Specifically, we investigate employee openness to experience as a boundary condition that shapes how leader cues are internalized and interpreted. Drawing on signaling theory's emphasis on both signal senders and receivers (Connelly et al., 2011), we argue that employees high in openness are more receptive to change-oriented signals and more likely to derive positive interpretations from leader symbolization. This leads us to propose a moderated mediation model in which the indirect effect of leader AI symbolization on employee well-being—via AI threat perception—is contingent upon the level of employee openness. The illustrative research model is depicted in Figure 1.

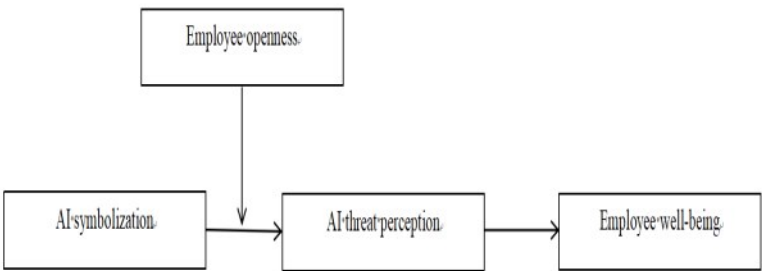


Fig.1 Theoretical Model of the Current Research

By integrating leadership communication, individual differences, and employee psychological outcomes within a unified theoretical framework, this research seeks to advance our understanding of how organizations can foster healthier, more adaptive employee responses amid AI-driven transformations.

2 Theory and Hypothesis Development

2.1 Leader AI Symbolization and AI Threat Perception

This research is grounded in signaling theory (Spence, 1978), which posits that under conditions of information asymmetry or uncertainty, individuals rely on observable behaviors—signals—to infer intentions and make evaluative judgments (Taj, 2016). Within organizational settings, particularly during technological transitions such as AI integration, employees often lack full visibility into the implications of such changes (Morandini et al., 2023). As a result, they depend on leaders' behaviors and communications as interpretive cues to assess risk, derive meaning, and shape their emotional and cognitive responses (He et al., 2023).

Leaders play a crucial signaling role in shaping organizational narratives (Denning, 2011). When leaders symbolically endorse AI—through discourse, behaviors, and strategic framing—they convey that AI is both benign and valuable to the organization (Bareis & Katzenbach, 2022). According to signaling theory, employees interpret these positive cues as evidence that AI represents an opportunity rather than a threat. Such signals reduce uncertainty and anxiety, thereby lowering employees' AI threat perception. Conversely, absent or ambiguous signals may foster threat-laden assumptions. Therefore, we hypothesize,

Hypothesis 1: Leader AI symbolization is negatively related to employee AI threat perception.

2.2 Leader AI Symbolization, AI Threat Perception and Employee Well-Being

From a signaling perspective, once a signal (leader symbolization) is internalized, it affects downstream psychological outcomes. If employees perceive AI as a threat to their job security, role identity, or competence, they are likely to experience stress, anxiety, and emotional exhaustion—all known detractors of well-being. Thus, threat perception acts as a cognitive stressor, mediating the impact of environmental signals on individual outcomes. Therefore, we hypothesize,

Hypothesis 2: Employee AI threat perception is negatively related to employee well-being.

Combining the theoretical rationale underlying Hypotheses 1 and 2, we propose that employee perceptions of AI-related threat serve as a key psychological mechanism through which leader AI symbolization influences employee well-being. Specifically, when leaders convey clear and positive signals about AI, such communications are likely to attenuate employees' threat perceptions, thereby fostering improved emotional and psychological functioning in the workplace. Therefore, we hypothesize,

Hypothesis 3: Leader AI symbolization have a positive indirect relationship with employee well-being via employee AI threat perception.

2.3 The Moderating Role of Employee Openness

Signaling theory also acknowledges that the characteristics of signal receivers influence how signals are interpreted (Connelly et al., 2011). Among these receiver attributes, openness to experience plays a particularly salient role in shaping individual responses to change-oriented signals. Employees who score high on openness to experience tend to be more curious, imaginative, and receptive to novel or unconventional ideas (Seppälä et al., 2013). Consequently, they are more inclined to perceive leader AI symbolization—defined as leaders' intentional use of AI as a symbolic representation of future-oriented change and innovation—as both credible and meaningful. This heightened receptivity enhances the salience of the signal and fosters a more constructive interpretation, ultimately contributing to a greater reduction in perceived AI-related threat. Such employees are more likely to interpret AI not as a threat to their professional roles, but as an opportunity for personal growth, skill development, and career evolution within the organizational context.

Conversely, employees low in openness to experience tend to favor familiar routines and established norms, and often exhibit resistance to change (Oreg, 2003). For these individuals, leader AI symbolization may be received with skepticism or even outright distrust. They are more likely to interpret such signals as ambiguous, inauthentic, or potentially threatening, thereby diminishing the signal's intended reassuring effect. As a result, these employees may be more inclined to perceive AI as a disruptive force rather than a collaborative tool, thereby maintaining—or even intensifying—their perception of AI as a threat. Therefore, we hypothesize,

Hypothesis 4: Employee openness moderates the negative relationship between leader AI symbolization and employee AI threat perception such that the relationship is stronger when employee openness is high (vs. low).

Building on the mediating mechanism outlined in Hypothesis 3 and the moderating effect specified in Hypothesis 4, we argue that employee openness to experience plays a pivotal role in shaping the indirect relationship between leader

AI symbolization and employee well-being via AI threat perception. Consistent with signaling theory, employees interpret leaders' symbolic communications as cues for understanding the implications of organizational changes such as AI integration. When leaders present AI in a positive and future-oriented manner, such signaling can mitigate employees' threat perceptions, thereby promoting greater psychological well-being. However, the effectiveness of this signaling process is not uniform across individuals. Employees high in openness to experience are more receptive to change and more inclined to process leader cues in a constructive manner, resulting in a more substantial reduction in perceived AI threat. Consequently, the positive downstream impact on well-being is amplified for these individuals. In contrast, employees low in openness may be less responsive to these symbolic cues, thereby weakening the strength of the indirect effect. Accordingly, we propose that employee openness to experience moderates the mediated relationship between leader AI symbolization and well-being through AI threat perception, such that the mediation is stronger for individuals high in openness. Therefore, we hypothesize,

Hypothesis 5: Employee openness moderates the positive indirect relationship between leader AI symbolization and employee well-being via employee AI threat perception such that the indirect relationship is stronger when employee openness is high (vs. low).

3 Method

3.1 Sample and Procedures

In this study, we recruited participants through Credamo (<https://credamo.cn>), a professional online survey platform in China. All participants were working employees from various companies. We collected employee surveys at multiple time points and used the last four digits of participants' phone numbers (which they provided at the end of each survey) to match surveys completed by the same employee across different time periods. To ensure anonymity, we only recorded participants' name initials and the last four digits of their phone numbers. Through this approach, we initially recruited 340 employees to participate.

We collected data through three waves of online surveys. In the first page of each survey, we noted that their responses would only be used for third-party academic research and explained the importance of careful responses to our research. We also assured them of anonymity and confidentiality. To further encourage participation, we also compensated each employee with 2 yuan RMB (approximately 0.27 USD) for finishing each survey, respectively. To improve quality of survey data, we added one attention check item in each wave of the survey (e.g., "Please choose 'Strongly agree' for this question", Meade & Craig, 2012).

At Time 1 (T1), employees rated leader AI symbolization, openness, and reported demographic information. At Time 2 (T2), approximately two weeks later, we sent another survey to the employees who has returned the T1 survey, they rated AI threat perception. At Time 3 (T3), approximately two weeks later, we sent another survey to the employees who has returned the T2 survey, they rated employee well-being.

After data matching, we excluded careless responses which failed our attention check items or selected the same option consecutively or recklessly filled out demographic information, we also excluded groups which only had one sample left. Finally, we collected data from 299 employees. Among the 299 employees, 31.10% were male. On average, they were 32.47 years old ($SD = 7.62$), had 16.29 years of education ($SD = 1.64$), and had worked in their current organization for 7.01 years ($SD = 5.97$).

3.2 Measures

The surveys were conducted in Mandarin Chinese, though the measures we used were all scales adapted from English literatures. Therefore, we translated all the scales following Brislin's (1980) translation – back translation procedure. Specifically, a bilingual researcher translated the scales from English to Chinese, which were subsequently back-translated into English by another bilingual researcher. The two researchers conducted their translations independently. Upon completion, minor discrepancies were observed between the original English version and the back-translated items. These inconsistencies were resolved through subsequent discussions.

AI symbolization (T1). We adapted the six-item AI symbolization scale developed by He et al. (2023). Sample item: "My supervisor initiates discussions related to artificial intelligence (AI) with me.", "My supervisor expresses his or her interest in AI to me.", and "My supervisor displays many AI-related items in his or her workplace (e.g. AI-related posters, humanoid robot models, etc.)." (ranging from 1 = "strongly disagree" to 5 = "strongly agree") ($\alpha = .88$).

Employee openness (T1). We used the four-item individual openness scale developed by Meng et al. (2021) to measure employee openness. Sample item: "At work, I like to hear a variety of different ideas.", "At work, I have a rich imagination.", "At work, I like to think about problems." (ranging from 1 = "strongly disagree" to 5 = "strongly agree") ($\alpha = .67$).

AI threat perception (T2). We used the five-item AI threat perception scale developed by Yogeewaran et al. (2016). Sample item: “The increased use of AI in our working lives will lead to human unemployment.”, “AI will replace human jobs.”, “In the long run, AI poses a direct threat to human security and well-being.” (ranging from 1 = “strongly disagree” to 5 = “strongly agree”) ($\alpha = .89$).

Employee well-being (T3). We used the six-item employee well-being scale developed by Zheng et al. (2015). Sample item: “I feel mostly satisfied with what I specifically did.”, “Overall, I am satisfied with the work I perform.”, and “My job is very interesting.” (ranging from 1 = “strongly disagree” to 5 = “strongly agree”) ($\alpha = .82$).

Control variables. Following previous research which suggested that demographic variables may influence employee well-being (Tay et al., 2014). We controlled for employee demographic variables, including employee’s gender, age, education year, job tenure. Gender (coded as 0 = male, 1 = female) was included due to documented differences in workplace stress and well-being across genders (e.g., Nelson & Burke, 2002). Age (in years) was accounted for, as older employees may experience different levels of job satisfaction and stress compared to younger workers (e.g., Kooij et al., 2011). Education years (total years of formal education) was controlled because higher education may correlate with better coping strategies or job expectations (Ross & Wu, 1995). Finally, job tenure (years in the current position) was adjusted for, as longer tenure may relate to increased organizational commitment or, conversely, burnout. By statistically controlling for these variables, we aimed to isolate the unique effects of our focal predictors on employee well-being.

3.3 Results

Means, standard deviations, and correlations among variables are presented in Table 1. We first conducted a series of confirmatory factor analyses (CFAs) to ensure the four key constructs in our model (i.e., AI symbolization, employee openness, AI threat perception, and employee well-being) had satisfactory discriminant validity. As shown in Table 2, the hypothesized four-factor structure demonstrated a good fit to the data ($\chi^2[183] = 390.31$, $p < .001$; SRMR = .05, RMSEA = .06, CFI = .93; TLI = .92; Hu & Bentler, 1999) and was also superior to alternative models.

Table 1 Descriptive Statistics and Correlation Matrix

Variables	Mean	SD	1	2	3	4	5	6	7	8
AI symbolization	3.73	0.82	(.88)							
Employee openness	3.96	0.60	.62***	(.67)						
AI threat perception	2.70	0.92	-.29***	-.30***	(.89)					
Employee well-being	4.08	0.58	.63***	.67***	-.32***	(.82)				
Employee gender	1.69	0.46	-.03	-.06	.002	-.05	-			
Employee age	32.47	7.62	.20**	.21***	-.13*	.25***	.04	-		
Employee education	16.29	1.64	.10	.12*	-.10	.07	.08	.11*	-	
Employee tenure	7.01	5.97	.13*	.18**	-.14*	.24***	-.01	.83***	-.03	-

Note. N = 299. For employee gender, 0 = male; 1 = female. Correlations for all variables represent group-mean centered relationships. *** $p < .001$, ** $p < .01$, * $p < .05$.

Table 2 Results of Model Fit Estimates

Model	Factors	χ^2	df	RMSEA	SRMR	CFI	TLI	$\Delta\chi^2$ (Δ df)
Baseline	Four factors	390.31	183	.06	.05	.93	.92	-
Alternatives								
Model 1	Three factors. AI symbolization and employee openness combined	473.65	186	.07	.06	.90	.89	37.65 (3)***
Model 2	Three factors. AI symbolization and AI threat perception combined	1125.37	186	.13	.11	.67	.65	735.06 (3)***
Model 3	Two factors. AI symbolization, AI threat perception and employee openness combined	1197.46	188	.13	.11	.66	.62	807.15 (5)***
Model 4	One factor. All variables combined	1304.77	189	.14	.12	.63	.59	914.46 (6)***

Note. N = 299; *** $p < .001$.

3.4 Tests of Hypotheses

As shown in Model 1 of Table 3, the negative relationship between AI symbolization and AI threat perception was significant ($b = -.30$, $p < .001$). Thus, Hypothesis 1 was supported. Moreover, as shown in Model 3 of Table 3, AI threat perception was significantly related to employee well-being ($b = -.09$, $p < .01$). We ran Bootstrap test to examine the indirect effect (Preacher & Hayes, 2008). The Bootstrap results showed that the indirect effect of AI symbolization on employee well-being via AI threat perception was significant (estimate = .03, 95% CI = [.008, .051]). Thus, Hypothesis 2 was supported.

As shown in Model 2 from Table 3, the interactional effect of AI symbolization and employee openness on AI threat perception was significant ($b = -.32, p < .001$). In addition, the relationship between AI symbolization and AI threat perception was significant and negative when employee openness was high (estimate = $-.48$, 95% CI = $[-.706, -.257]$), but was not significant when employee openness was low (estimate = $-.09$, 95% CI = $[-.254, .071]$; see Figure 2), consistent with Hypothesis 3.

The moderated mediation hypothesis results showed that the indirect effect was significant when employee openness was high (estimate = $.04$, 95% CI $[.013, .081]$), but was not significant when employee openness was low (estimate = $.01$, 95% CI $[-.001, .022]$), the difference was significant (estimate = 0.03 , 95% CI $[.001, .060]$). Thus, Hypothesis 4 was supported.

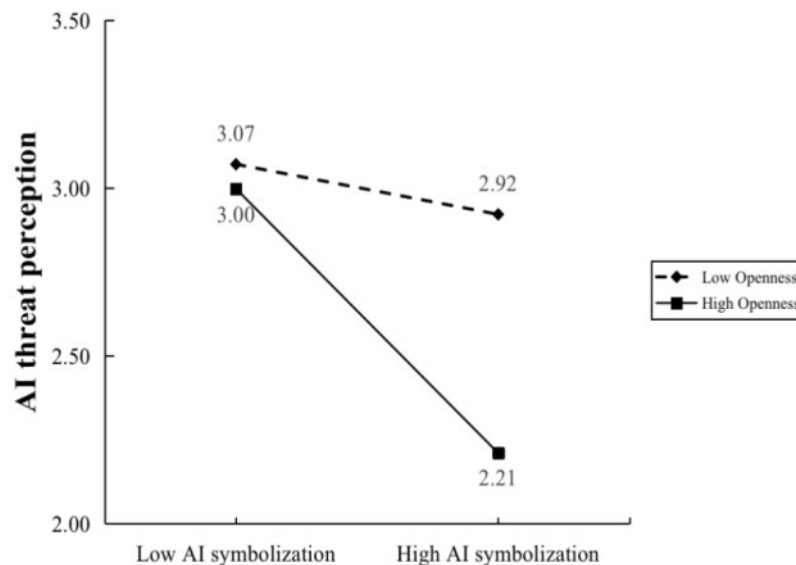


Fig.2 The Moderating Role of Employee Openness in the Relationship between AI Symbolization and AI Threat Perception

Table 3 Hypothesis Tests Results

Variable	AI threat perception		Employee well-being
	Model 1	Model 2	Model 3
Employee gender	-.01	.03	-.04
Employee age	.01	.01	.01
Employee education	-.05	-.04	.01
Employee tenure	-.02	-.02	.01
AI symbolization	-.30***	-.27***	.40***
Employee openness		-.32**	
AI symbolization × Employee openness		-.32***	
AI threat perception			-.09**
R2	.10	.15	.44
F	6.49***	7.62***	37.58***

Note. N= 299. For employee gender, 0 = male; 1 = female. *** $p < .001$, ** $p < .01$, * $p < .05$.

4 General Discussion

Based on the signaling theory, we propose that leader AI symbolization is positively related to employee well-being via employee AI threat perception. Moreover, employee openness moderates this relationship such that the effect is stronger when employee openness is high (vs. low). To test our hypotheses, we conducted a multi-wave questionnaire research, the results supported our hypotheses.

5 Theoretical Contributions

This research makes several theoretical contributions to the literature on artificial intelligence and organizational behavior. First, it extends prior work by illuminating the downstream effects of leader AI symbolization—an underexamined yet conceptually significant construct. Whereas existing studies have largely focused on leader behaviors in facilitating AI adoption (e.g., Rai et al., 2019), this study shifts the focus toward symbolic expressions, such as framing AI as an opportunity or strategic asset, which function as salient signals shaping employee interpretations and affective responses. By uncovering the indirect effect of leader AI symbolization on employee well-being through AI threat

perception, our findings contribute a more nuanced perspective on the symbolic management of technological change (Ashforth & Humphrey, 1993; Dutton et al., 1994).

Second, this study advances the literature on AI threat perception by identifying leader communication and individual differences—specifically, openness to experience—as meaningful antecedents. While prior research has predominantly conceptualized AI threat perception as a static response to technological characteristics or the risk of job displacement (e.g., Schepman & Rodway, 2020; Zhang et al., 2023), our findings suggest that such perceptions are socially constructed and dynamic. Specifically, threat perceptions are shaped by symbolic leadership signals and moderated by employee-level traits such as openness. This perspective expands the theoretical understanding of threat perception by emphasizing its relational and contextual malleability.

6 Practical Implications

This research yields two key implications for organizational practice amid ongoing AI-driven transformation. First, leaders should engage in deliberate AI symbolization to shape employee perceptions effectively. Our findings highlight the critical role of not only what leaders implement, but how they frame and communicate AI-related initiatives. Symbolic expressions that portray AI as a collaborative partner in innovation—rather than a disruptive threat—function as positive signals that alleviate employee anxiety and enhance psychological well-being. These results are consistent with prior research demonstrating that leader framing influences employee attitudes toward organizational change (Fiss & Zajac, 2006). Accordingly, leadership development programs should incorporate training in symbolic communication tailored to the demands of digital transformation.

Second, the effectiveness of AI-related communication can be enhanced by tailoring strategies to individual differences among employees. Not all employees interpret leadership signals uniformly. This study demonstrates that openness to experience moderates the effect of leader AI symbolization on threat perception. Employees high in openness are more receptive to positive, future-oriented messages about AI, whereas those low in openness may require additional support and more nuanced, targeted messaging. Recognizing and addressing such individual-level variability—through tools such as personality assessments or personalized communication approaches—can promote smoother AI adoption and foster greater psychological safety across the workforce (Judge et al., 2003).

7 Limitations and Future Directions

This research has several limitations that suggest promising directions for future research. First, the use of a cross-sectional survey design restricts the ability to draw causal inferences. Future research could incorporate experimental methods, such as scenario-based experiments or field interventions, to more rigorously test the causal impact of leader AI symbolization on employee outcomes. Second, while this research focused on AI threat perception as a mediator and employee openness as a moderator, the model could be expanded by exploring alternative psychological mechanisms—such as trust in leadership or perceived job insecurity—as well as additional boundary conditions like technological self-efficacy or organizational culture. These extensions would deepen our understanding of how leader signaling influences employee responses to AI.

References

- [1] ASHFORTH B E, HUMPHREY R H. Emotional labor in service roles: The influence of identity[J]. *Academy of Management Review*, 1993, 18(1): 88-115.
- [2] BAREIS J, KATZENBACH C. Talking AI into being: The narratives and imaginaries of national AI strategies and their performative politics[J]. *Science, Technology & Human Values*, 2022, 47(1): 35-60.
- [3] BRISLIN R W. Cross-cultural research methods[M]. Boston: Springer, 1980.
- [4] CONNELLY B L, CERTO S T, REUTZEL C R. Signaling theory: A review and assessment[J]. *Journal of Management*, 2011, 37(1): 39-67.
- [5] DENNING S. The secret language of leadership. How leaders inspire action through narrative[J]. *Strategic Direction*, 2009, 25(1).
- [6] DUTTON J E, DUKERICH J M, HARQUAIL C V. Organizational images and member identification[J]. *Administrative Science Quarterly*, 1994, 39(2): 239-263.
- [7] FISS P C, ZAJAC E J. The symbolic management of strategic change: Sensegiving via framing and decoupling[J]. *Academy of Management Journal*, 2006, 49(6): 1173-1193.
- [8] HE G, LIU P, ZHENG X, et al. Being proactive in the age of AI: Exploring the effectiveness of leaders' AI symbolization in stimulating employee job crafting[J]. *Management Decision*, 2023, 61(10): 2896-2919.
- [9] JUDGE T A, EREZ A, BONO J E, et al. The core self-evaluations scale: Development of a measure[J]. *Personnel Psychology*, 2003, 52(3): 717-742.
- [10] KOOLIJ D T, DE LANGE A H, JANSEN P G, et al. Age and work-related motives: Results of a meta-analysis[J]. *Journal of Organizational Behavior*, 2011, 32(2): 197-225.
- [11] MENG Y, YU B, LI C, et al. Psychometric properties of the Chinese version of the organization big five scale[J]. *Frontiers in Psychology*, 2021,

12: 781369.

- [12] MORANDINI S, FRABONI F, DE ANGELIS M, et al. The impact of artificial intelligence on workers' skills: Upskilling and reskilling in organisations [J]. *Informing Science*, 2023, 26: 39-68.
- [13] NELSON D L, BURKE R J. Gender, work, stress, and health?(pp. xii-260)[M]. Washington: American Psychological Association, 2002.
- [14] OREG S. Resistance to change: Developing an individual differences measure[J]. *Journal of Applied Psychology*, 2003, 88(4): 680-693.
- [15] RAI A, CONSTANTINIDES P, SARKER S. Next-generation digital platforms: Toward human–AI hybrids[J]. *MIS Quarterly*, 2019, 43(1).
- [16] ROSS C E, WU C L. The links between education and health[J]. *American Sociological Review*, 1995: 719-745.
- [17] SCHEPMAN A, RODWAY P. Initial validation of the general attitudes towards Artificial Intelligence Scale[J]. *Computers in Human Behavior Reports*, 2020, 1: 100014.
- [18] SEPPÄLÄ E M, ROSSOMANDO T, DOTY J R. Social connection and compassion: Important predictors of health and well-being[J]. *Social Research: An International Quarterly*, 2013, 80(2): 411–430.
- [19] SPENCE M. Job market signaling[J]. *Quarterly Journal of Economics*, 1978, 87(3): 355–374.
- [20] TAY L, LI M, MYERS D, et al. Religiosity and subjective well-being: An international perspective[J]. *Religion and Spirituality Across Cultures*, 2014: 163-175.
- [21] YOGESWARAN K, ZŁOTOWSKI J, LIVINGSTONE M, et al. The interactive effects of robot anthropomorphism and robot ability on perceived threat and support for robotics research[J]. *Journal of Human-Robot Interaction*, 2016, 5(2): 29–47.
- [22] ZHENG X, ZHU W, ZHAO H, et al. Employee well-being in organizations: Theoretical model, scale development, and cross-cultural validation [J]. *Journal of Organizational Behavior*, 2015, 36(5): 621-644.