

WHEN EMPATHY IS NOT ENOUGH: EMOTIONAL AUTHENTICITY AND ROLE LEGITIMACY IN THE PROFESSIONAL ACCEPTANCE OF GENERATIVE AI CAREER SUPPORT

Zhiqian Zhang, Chen Chen, and Jongchang Ahn*

Department of Information Systems, Hanyang University, Republic of Korea

Received: 04/04/2026

Accepted: 20/05/2026

*Corresponding author: ajchang@hanyang.ac.kr

ABSTRACT

Generative AI is increasingly used where users seek information, reassurance, and emotional understanding. Yet empathic responses do not necessarily lead users to accept AI as a legitimate support source. Drawing on career support, a setting that combines practical uncertainty and emotional strain, this study examines how perceived AI empathy quality is converted into professional acceptance intention. We conceptualize acceptance as a bounded status-granting process shaped by two judgments: whether AI's expression of care feels emotionally authentic and whether AI is legitimate for a first-line support role. Using 350 valid responses from an anonymous vignette-based online survey, we tested a parallel mediation model with confirmatory factor analysis and structural equation modeling. Results show that perceived AI empathy quality positively predicts emotional authenticity and role legitimacy, which in turn mediate its relationship with professional acceptance intention. These effects remain robust after accounting for perceived competence and scenario realism. The findings suggest that empathic AI is accepted not simply because it sounds supportive, but because users regard its care as credible and its role as appropriately bounded. This study contributes to AI empathy, human AI interaction, and responsible digital service design by identifying relational and role-based mechanisms of acceptance.

KEYWORDS: *Empathic AI; career support; emotional authenticity; role legitimacy; professional acceptance; responsible AI; structural equation modeling*

1. INTRODUCTION

Generative AI is rapidly moving from backstage analytical functions to frontline interaction roles. Earlier research anticipated that AI would extend beyond routine tasks to activities involving judgment, explanation, and user support [1]. With the rise of generative systems, users now interact with technologies that not only retrieve information but also provide advice, reassurance, and emotionally supportive responses [2]. This shift raises a central question: can such capabilities lead users to accept AI as a credible source of professional support, rather than merely as a tool for information delivery?

Recent evidence suggests this question is increasingly relevant. In healthcare communication, AI-generated responses to patient inquiries have been rated higher in both quality and empathy than those written by physicians [3]. Other studies show that AI messages can enhance users' sense of being heard, although this effect diminishes when the content is explicitly identified as AI-generated [4]. These findings indicate that AI can produce language that users perceive as warm, attentive, and responsive.

However, a gap emerges between response quality and source acceptance. Empathic performance at the message level does not necessarily translate into acceptance of AI as a legitimate provider of support. Identical content is more likely to be perceived as genuinely supportive when attributed to a human rather than to AI [5], and users may still prefer human interaction even when they rate AI responses positively [6]. Disclosure of AI involvement can further reduce trust and raise concerns about legitimacy [7]. The key issue, therefore, is not whether AI can generate empathic responses, but why such responses often fail to establish AI as an acceptable source of professional support.

Existing research offers only partial explanations. Prior studies have examined how users evaluate responses in terms of warmth, supportiveness, and understanding [8], while another stream has focused on algorithm aversion and competence judgments [9]. Yet these perspectives do not fully explain why positive evaluations of AI-generated empathy often do not lead to broader acceptance. This paper addresses this gap by framing acceptance as a dual judgment process, in which users assess both the credibility of the care provided and the appropriateness of AI occupying a professional support role.

Career support provides a particularly useful context for examining this issue. Compared with routine service interactions, career-related decisions involve higher emotional stakes, including uncertainty, anxiety, and personal aspiration. At the same time, career support lacks the formal safeguards found in clinical settings, as it often centers on preliminary guidance and interpretation rather than regulated intervention. Prior research links career guidance to well-being [10], and the quality of the counsellor client relationship is consistently associated with outcomes [11]. As AI becomes more integrated into career development services, it introduces both opportunities for scalability and concerns regarding transparency, bias, limited relational depth, and unclear responsibility boundaries [12]. This combination makes career support an effective setting for understanding why empathic AI responses may be valued yet still require justification as a professional source.

From the perspective of sustainable digital service systems, AI-based career support can expand access to timely guidance for individuals facing career uncertainty. However, such scalability is viable only if users perceive both the care provided as credible and the role of AI as appropriate. This study, therefore, contributes to responsible and sustainable AI deployment by identifying the conditions under which empathic AI can be accepted as a bounded source of professional support.

Building on this foundation, the paper develops a parallel mediation model that conceptualizes AI acceptance as a process of granting limited professional status. While perceived empathy is an important starting point, it is not sufficient on its own. Users will also evaluate whether the response is sincere, specific, and original (referred to as emotional authenticity) and whether the AI can appropriately provide frontline career support (referred to as role legitimacy). Rather than simply adding mediating variables, this study explains why empathic responses do not automatically lead to acceptance. It shows that acceptance depends on a joint evaluation of relational credibility and appropriate role boundaries in contexts where emotional interaction and professional judgment intersect.

2. LITERATURE REVIEW AND THEORETICAL FOUNDATIONS

2.1. From Empathic Responses to Professional Acceptance

Generative AI has shifted artificial intelligence from backstage analytical tasks to visible interaction roles. The key issue is no longer limited to prediction or recommendation, but whether AI can engage in activities that involve explanation, response, and user support [13]. Recent studies show that AI systems can produce replies that are perceived as detailed, attentive, and empathetic [14]. While these findings demonstrate response-level capability, they do not explain when such responses are accepted as coming from a legitimate professional source.

This gap between response quality and source acceptance is increasingly evident. Signals that a message is AI-generated can shape how users interpret trustworthiness and social meaning [15]. In contexts with strong relational or symbolic value, users often continue to prefer human over AI providers [16]. This suggests that evaluating whether a response feels empathic is not the same as accepting AI as a provider of support. A more difficult question remains: under what conditions does empathic performance translate into acceptance of AI in a bounded professional role?

2.2. Emotional authenticity: social response and mind perception

One explanation for this gap lies in how users interpret the emotional meaning of AI responses. Research in human computer interaction shows that people often respond to computers in social ways, treating them as if they were social actors [17]. However, such reactions do not imply genuine relational endorsement. Users may respond with politeness or perceive empathy-like qualities without viewing the system as a truly caring agent.

Research on mind perception clarifies this distinction. Studies differentiate between perceived agency (the ability to act) and perceived experience (the ability to feel) [18]. In supportive contexts, perceived experience is particularly important. Even well-crafted responses may fail to convey credible care if users doubt that the source can genuinely understand or “feel” their situation. This helps explain why replacing human providers with AI can reduce perceived warmth, even when message quality remains high [19].

Related research on AI-generated content supports this view. Disclosure of AI authorship can reduce perceived authenticity and influence downstream reactions [20], while perceived authenticity in other contexts strengthens trust and engagement [21]. Studies of empathic chatbots further show that empathy improves user experience only when it is interpreted as

meaningful and credible [22]. Emotional authenticity, therefore, represents a key judgment through which supportive language becomes believable care.

2.3. Role Legitimacy: Legitimacy and Professional Role Boundaries

A second explanation lies in how users evaluate whether AI should occupy a given role. Role legitimacy concerns not the quality of a response, but whether AI is appropriate as a provider of support. Legitimacy theory defines this as a judgment that an actor’s behavior aligns with shared norms and expectations [23]. As AI becomes more visible in support contexts, acceptance depends not only on usefulness but also on whether its role is considered appropriate.

This issue becomes more pronounced in professional settings. Professions are not defined solely by skills; they are structured systems of expertise, jurisdiction, and role boundaries [24].

Career support operates within such a system, involving advice, interpretation, and emotional guidance, along with expectations about who should provide these services.

Empirical research reflects these concerns. Individuals assess whether actions and actors are appropriate within social and cognitive norms [25], and studies on generative AI adoption show that legitimacy influences usage beyond perceived usefulness [26]. In healthcare, users may resist AI in core roles despite recognizing its efficiency [27]. Similarly, research on high-empathy service contexts finds that users prefer human providers, although hybrid human–AI arrangements can reduce resistance [28]. These findings suggest that hesitation toward AI often stems not from its inability to generate useful responses, but from uncertainty about its rightful place in a professional role. Role legitimacy captures this dimension of role authorization.

2.4. Why Career Support is a Revealing Context

Career support provides a particularly suitable context for examining these issues because it combines practical guidance with emotional strain. Individuals typically seek such support during periods of uncertainty, failure, or transition. These situations involve not only information needs but also emotional concerns. Research shows that career interventions can influence decision outcomes and subsequent actions [29], highlighting the practical importance of such support.

Career counseling research further underscores its relational dimension. Meta-analyses demonstrate that the quality of the working relationship between counselor and

WHEN EMPATHY IS NOT ENOUGH: EMOTIONAL AUTHENTICITY AND ROLE LEGITIMACY IN THE PROFESSIONAL ACCEPTANCE OF GENERATIVE AI CAREER SUPPORT

client is associated with outcomes [30], and that career counseling can affect both career-related and mental health outcomes [31]. At the same time, digital career tools expand access but raise concerns about guidance quality, transparency, and role boundaries [32].

This combination of emotional relevance and limited institutional structure makes career support an effective setting for studying AI acceptance. In this context, users may value empathic responses while still questioning whether AI should serve as a professional support provider. Both relational credibility and role legitimacy are therefore directly visible rather than embedded in general attitudes toward technology.

2.5. Research Gap and Model

Existing research highlights three key points. First, AI can produce responses that appear empathic in service interactions [33]. Second, source attribution and disclosure influence how users evaluate AI-generated communication [34]. Third, legitimacy becomes critical when AI enters socially visible roles with defined responsibilities [35]. However, current work does not fully explain why empathic response quality does not translate into professional acceptance in career support settings.

This study addresses this gap by framing the issue as a conversion problem from empathic performance to bounded professional acceptance. Perceived empathy quality is necessary but not sufficient. Users evaluate AI at two levels: relationally, by assessing whether the response feels credible and specific; and structurally, by assessing whether AI is appropriate for a first-line support role. Professional acceptance is therefore conceptualized as the joint outcome of relational credibility and role authorization, rather than as a simple extension of trust or competence.

2.6. Boundary Distinctions among Focal Constructs

To avoid overlap with related concepts such as trust or competence, the model distinguishes constructs based on the type of judgment they represent. Professional acceptance is treated as a dual process involving relational credibility and role authorization. Emotional authenticity captures the first judgment, while role legitimacy captures the second. Table 1 summarizes these distinctions and clarifies how the focal constructs differ from related concepts.

The contribution of this framework lies not in introducing new variables but in explaining how these two judgments jointly translate empathic response quality into bounded professional acceptance within AI-based career support.

Table 1. Boundary distinctions among focal constructs

Construct	Core question	What it captures	What it does not capture	Closest neighboring construct
Perceived AI empathy quality (PEQ)	Does the response understand and fit the user's situation?	Emotional recognition, situational fit, supportive responsiveness	Whether the care feels credible or whether AI belongs in the role	Emotional authenticity; perceived competence
Emotional authenticity (EA)	Does the care feel credible rather than formulaic?	Credible, sincere, specific care expression	Whether AI is normatively appropriate for the support role	Trust; perceived empathy
Role legitimacy (RL)	Does AI belong in the first-line support position?	Role fit, boundary appropriateness, normative acceptability	Whether the response feels sincere or emotionally real	Legitimacy; perceived competence
Professional acceptance intention (PAI)	Will the user include AI in the future support process?	Ongoing acceptance, first-line source acceptance, support-process inclusion	Preference for AI over human counselors	Trust; relative source preference
Perceived competence (PC)	Can AI handle the support task effectively?	Effectiveness, problem identification, actionable advice	Role authorization or credible care	Role legitimacy; perceived usefulness

Note. PEQ = perceived AI empathy quality; EA = emotional authenticity; RL = role legitimacy; PAI = professional acceptance intention; PC = perceived competence.

3. Theory Development and Hypotheses

Building on the preceding review, this study frames professional acceptance as a conversion process rather than a simple attitude outcome. While perceived AI empathy quality is important, it does not directly translate into acceptance. Before users accept AI as a source of support, they make two additional judgments: whether the care expressed in the response is believable, and whether AI is appropriate for the support role. These judgments correspond to emotional authenticity and role legitimacy, respectively.

Although both are expected to contribute to professional acceptance, they operate through distinct mechanisms.

Figure 1 presents the proposed conceptual model. It specifies two parallel pathways through which perceived AI empathy quality influences professional acceptance intention: a relational pathway via emotional authenticity and a role-based pathway via role legitimacy. The following sections develop the hypotheses corresponding to each path.

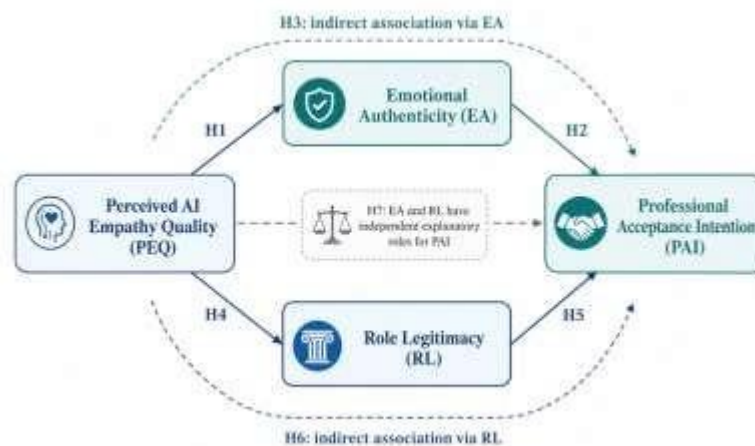


Figure 1. Proposed conceptual model of professional acceptance of empathic AI in career support.

Note. PEQ = perceived AI empathy quality; EA = emotional authenticity; RL = role legitimacy; PAI = professional acceptance intention. H3 and H6 denote indirect associations; H7 denotes the independent explanatory roles of EA and RL.

3.1. Perceived AI Empathy Quality and Emotional Authenticity

In career support contexts, users encounter the response before forming a stable view of the source. As a result, the quality and tone of the response play a central role in shaping interpretation. When a message reflects an understanding of the user's emotional state, addresses the specific situation, and offers contextually appropriate support, users are more likely to perceive the care as credible rather than formulaic.

Prior research supports this relationship. Perceived empathy influences evaluations of service providers, including chatbots, even in brief interactions [36]. Studies on authenticity further show that authenticity emerges during the interpretation of a message rather than being assigned afterward [37]. In career support settings, higher perceived AI empathy quality should therefore increase emotional authenticity.

H1: Perceived AI empathy quality is positively associated with emotional authenticity.

3.2. Emotional Authenticity and Professional Acceptance Intention

Professional acceptance requires more than a favorable reaction to a single response. It involves considering whether the source could play a role in future support. Such acceptance is unlikely if the response feels emotionally hollow or insincere. A message may be well-structured and receive a positive evaluation, yet still fail to establish the basis for ongoing engagement if it lacks credibility.

Evidence from related research supports this reasoning. Perceived authenticity has been shown to strengthen downstream responses and engagement [38]. Similarly, human-AI collaboration can improve acceptance by making the source appear more credible and less purely artificial [39]. In career support contexts, emotional authenticity should therefore both directly increase professional acceptance intention and mediate the effect of perceived empathy quality.

H2: Emotional authenticity is positively associated with professional acceptance intention.

H3: Emotional authenticity mediates the relationship between perceived AI empathy quality and professional acceptance intention.

3.3. *Perceived AI Empathy Quality and Role Legitimacy*

Role legitimacy addresses a different dimension of evaluation. It concerns whether AI is appropriate for the role of a first-line support provider, rather than whether a response feels sincere. Although this judgment is normative, it is informed by observed performance. Users infer what role AI can occupy based on how it behaves in interaction.

When AI responses are limited or mechanical, users are likely to view it as a tool. However, when AI demonstrates sensitivity to emotional context and provides support that fits the situation, it becomes more plausible to see it as capable of a broader support role. Research on social presence shows that systems displaying socially responsive behavior can shift how users interpret service interactions [40]. Similarly, studies of health chatbots indicate that expressions of empathy can influence evaluations of both the message and the system [41]. These findings suggest that higher perceived empathy quality can enhance perceptions of role legitimacy.

H4: Perceived AI empathy quality is positively associated with role legitimacy.

3.4. *Role Legitimacy and Professional Acceptance Intention*

Even when AI performs well, users may hesitate to accept it if they do not view it as appropriate for the role. Professional acceptance depends not only on performance but also on perceived fit within established norms and expectations. Legitimacy theory suggests that acceptance increases when an actor's behavior aligns with shared beliefs about appropriate roles [42]. From a social judgment perspective, legitimacy shapes how users evaluate an actor's standing and acceptability [43].

This issue is particularly relevant in career support, where advice, interpretation, and emotional reassurance are closely intertwined. Research shows that AI involvement can alter how users interpret communication and social relationships, making the source itself part of the evaluation [44]. In professional services, disclosure of AI use can reduce satisfaction by signaling lower role appropriateness [45]. Studies on the "word-of-machine" effect further show that acceptance varies depending on the task and context [46]. Together, these findings suggest that users are unlikely to accept AI as a first-line support source unless it is seen as appropriate for that role.

H5: Role legitimacy is positively associated with professional acceptance intention.

H6: Role legitimacy mediates the relationship between perceived AI empathy quality and professional acceptance intention.

3.5. *Dual Mechanisms and a Rival Explanation*

Emotional authenticity and role legitimacy represent two distinct mechanisms and should not be reduced to a general notion of trust. Emotional authenticity addresses whether the care expressed in a response is believable, whereas role legitimacy concerns whether AI is appropriate for the support role. These judgments can diverge: a response may feel sincere but still come from a source seen as inappropriate, or AI may be accepted as a tool while its care remains unconvincing.

At the same time, perceived competence provides an important alternative explanation. Users may revise their evaluation of AI competence when they become aware of its non-human origin [47], and research on algorithm aversion shows that acceptance depends on task characteristics and judgment type [48]. Because the responses examined in this study combine emotional reassurance with actionable guidance, acceptance could be driven by perceived capability rather than relational or role-based judgments.

To address this possibility, perceived competence is included as a rival explanation in robustness analyses rather than as part of the main model. This approach allows us to assess whether emotional authenticity and role legitimacy retain independent explanatory power when competence is taken into account.

H7: Emotional authenticity and role legitimacy each have an independent positive association with professional acceptance intention after controlling for perceived competence and the other mediator.

4. MATERIALS AND METHODS

4.1. *Methodological Positioning and Overall Design*

This study adopts a vignette-based anonymous online survey within a structured three-stage workflow: scale development and pretesting, single-session survey administration, and statistical analysis. Rather than examining general attitudes toward AI, the study focuses on how individuals interpret a specific AI-generated response in a career-support context and whether such interpretation is associated with professional acceptance.

The vignette approach serves both practical and theoretical purposes. It enables a controlled and context-specific interaction without requiring

live intervention, which is appropriate when the research interest lies in user judgment rather than observed behavior in high-stakes settings. The survey instrument was organized into ordered blocks, such as screening, background, vignette exposure, proximal perceptions, and outcome variables, to ensure procedural clarity. However, this structure does not provide the causal leverage of longitudinal or experimental designs.

The study uses a single AI source and a fixed response to isolate the internal formation of professional acceptance rather than comparative evaluation across sources. Because the response combines emotional support with actionable guidance, perceived competence is retained as a rival explanation to distinguish between being understood and being seen as capable of handling the task. Figure 2 summarizes the overall research design and analytical workflow.

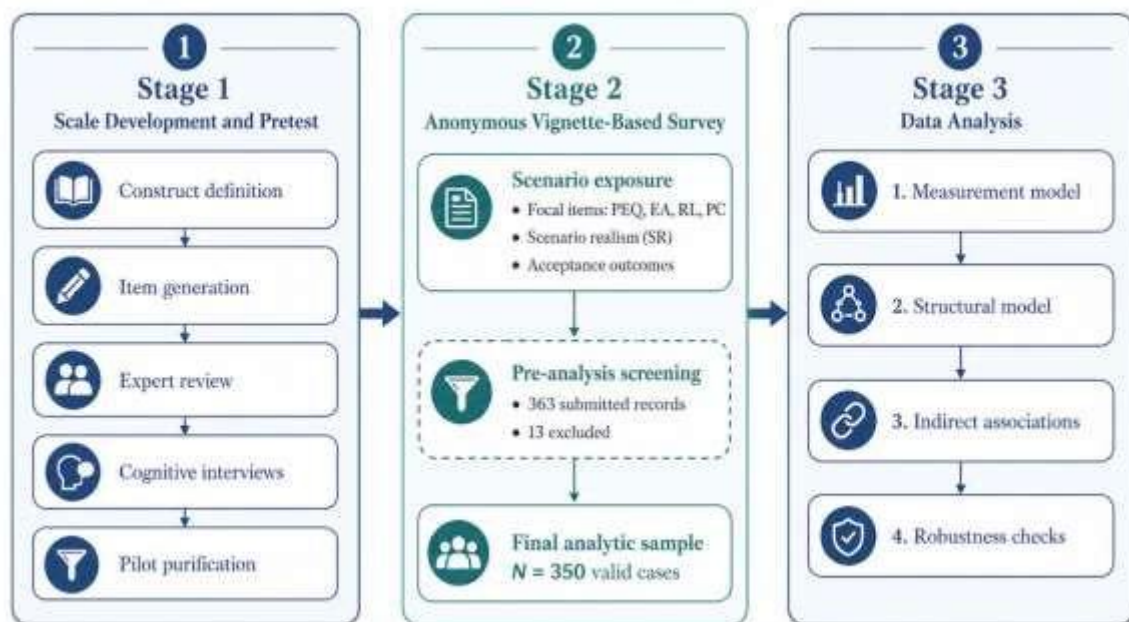


Figure 2. Research design flowchart.

Note. The figure summarizes measure development, anonymous vignette-based survey administration, pre-analysis data screening, and SEM analyses. SR = scenario realism; PEQ = perceived AI empathy quality; EA = emotional authenticity; RL = role legitimacy; PC = perceived competence.

4.2. Research Context and Sample

The empirical setting is front-end career support, which combines practical guidance with emotional strain. Individuals typically seek such support during periods of uncertainty, failed applications, or career transitions. In these situations, users require not only information but also interpretation, reassurance, and a sense of being understood. Prior research shows that career decision-making often involves both content-related challenges and emotional difficulties [49].

The target sample includes final-year undergraduates, recent graduates, graduate students engaged in job search, and early-career workers within approximately three years of labor market entry. These groups are closely aligned with the focal scenario and are therefore well positioned to provide informed evaluations. Inclusion criteria required respondents to be at least 18 years old, to have recent experience with career-related decisions or job search, and to complete the anonymous questionnaire.

4.3. Construct operationalization and scale development

The model includes four focal latent constructs: perceived AI empathy quality (PEQ), emotional authenticity (EA), role legitimacy (RL), and professional acceptance intention (PAI). Perceived competence (PC) is included as a rival explanation in robustness analyses. All constructs are modeled as reflective latent variables and

measured using seven-point Likert type scales, consistent with established practices in behavioral and information-systems research [50].

Item development followed a structured process: construct definition, item generation, content validation, cognitive interviewing, pilot testing, and final refinement. Cognitive interviews were particularly useful for identifying ambiguities in item interpretation [51].

WHEN EMPATHY IS NOT ENOUGH: EMOTIONAL AUTHENTICITY AND ROLE LEGITIMACY IN THE PROFESSIONAL ACCEPTANCE OF GENERATIVE AI CAREER SUPPORT

Maintaining clear construct boundaries was critical, given the theoretical emphasis on distinguishing relational credibility (EA) from response quality (PEQ) and role-based judgment (RL).

The questionnaire includes two categories of items. The first consists of focal measures used in the structural model (PEQ, EA, RL, PAI, and PC).

The second includes non-model items such as screening questions, background variables, scenario realism checks, and supplementary outcomes (e.g., trust and relative source preference). These items support data quality and contextual validity but are not part of the core hypothesis tests. Supplementary Table S2 provides the full questionnaire.

Table 2. Presents construct definitions and operationalization details.

Construct	Abbrev.	Working definition	Theoretical role	No. of items	Operational note
Perceived AI empathy quality	PEQ	Overall judgment that the AI response understood the user's emotional state, grasped the situation, and responded with appropriate support.	Antecedent	5	Measures understanding, responsiveness, and emotional fit; excludes utility and competence.
Emotional authenticity	EA	Judgment that the care expressed in the AI response feels credible, sincere, and not merely formulaic.	Parallel mediator	5	Measures whether the care feels believable; excludes role appropriateness.
Role legitimacy	RL	Judgment that AI is appropriate for occupying a first-line career-support role in the focal situation.	Parallel mediator	5	Measures role fit and boundary legitimacy; excludes sincerity.
Professional acceptance intention	PAI	Willingness to include the AI source in one's support process and treat it as an acceptable first-line support source.	Outcome	4	Measures ongoing acceptance rather than relative source choice.
Perceived competence	PC	Judgment that the AI is capable of handling the focal support task effectively.	Rival explanation	4	Included only in robustness checks.

4.4. Data Collection Procedure

Data were collected through anonymous online questionnaire. The survey sequence included informed consent, eligibility screening, background information, vignette exposure, and subsequent evaluation of the AI response. Respondents assessed PEQ, EA, RL, PC, and scenario realism, followed by professional acceptance intention and supplementary outcomes.

The single-vignette design ensured a constant interaction context, allowing the analysis to focus on variation in interpretation rather than variation in stimulus content. No personally identifiable information was collected. Platform-generated metadata were used only for data-quality checks and duplicate detection.

Because the study relies on a single-session survey, it is best characterized as a vignette based SEM design with ordered measurement blocks. Procedural remedies, such as anonymity, construct separation, and structured flow, were applied to reduce common method bias, although such bias cannot be fully eliminated [52].

4.5. Sample Size Planning and Data Quality Control

Sample size planning was based on model complexity rather than heuristic rules. Required sample size in SEM depends on factor loadings, path coefficients, and overall model structure [53]. Considering the expected medium effect sizes, a minimum of approximately 300 valid responses was deemed sufficient [54]. The final sample of 350 valid cases exceeded this threshold.

Data quality was ensured through both survey design and pre-analysis screening. Of the 363 initial responses, 13 were removed due to indicators of low-quality data, including extremely short completion times, invariant responses, or patterned answering. The final dataset contained no missing values, duplicate submissions, or systematic response biases.

4.6. Statistical Analysis Strategy

The analysis proceeded in two stages. The first stage focused on measurement validation. Item quality was assessed using reliability analysis, standardized factor loadings, cross-loading

diagnostics, and theoretical consistency. Parallel analysis was used for factor retention, as it provides a more robust criterion than traditional eigenvalue-based rules [55].

The second stage involved confirmatory factor analysis (CFA) and structural equation modeling (SEM) using maximum likelihood estimation. Measurement validity was evaluated through composite reliability, average variance extracted, convergent validity, and discriminant validity. Given the conceptual proximity between constructs, discriminant validity (particularly between EA and RL) was carefully assessed using

the HTMT criterion [56]. Model fit was evaluated using multiple indices (CFI, TLI, RMSEA, SRMR) rather than relying on a single cutoff [57-59].

After establishing the measurement model, the structural model was estimated. The analysis first tested the parallel mediation model without the rival explanation and then incorporated perceived competence to examine the robustness of the results. Indirect effects were assessed using bootstrap confidence intervals, which are appropriate for mediation analysis [60]. Figure 3 presents the proposed SEM model corresponding to Hypotheses H1-H7.

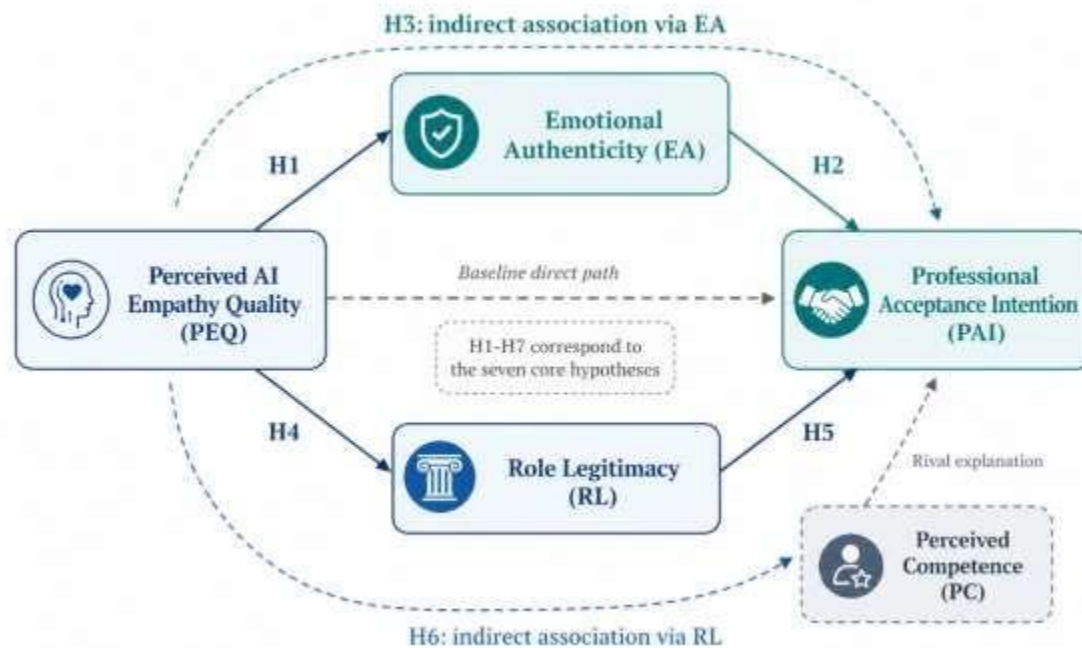


Figure 3. Proposed SEM model.

Note. H1-H7 correspond to the seven core hypotheses. Dashed lines indicate the baseline direct path and the rival-explanation path assessed in robustness checks. The structural paths should be interpreted as theory-guided associations.

5. RESULTS

5.1. Data screening and sample profile

The anonymous online survey initially yielded 363 responses. Thirteen records were excluded during pre-analysis screening due to data-quality concerns, including extremely short completion times, invariant responses across item blocks, and patterned answering indicative of low engagement. These exclusions were conducted prior to any model estimation or hypothesis testing.

The final analytic sample consisted of 350 valid responses, all meeting the eligibility criteria of being at least 18 years old and having recent experience with career-related decisions or transitions. No item-level missing values were observed. Completion times ranged from 127 to 221 seconds ($M = 168.51$, $SD = 24.12$). No duplicate submissions or full-block invariant response patterns were retained. Table 3 summarizes the sample characteristics and data-quality checks.

Table 3. Summarizes the sample characteristics and data-quality checks

Characteristic	n / M	% / SD
Initial returned questionnaires	363	100
Invalid records excluded before analysis	13	3.6

WHEN EMPATHY IS NOT ENOUGH: EMOTIONAL AUTHENTICITY AND ROLE LEGITIMACY IN THE PROFESSIONAL ACCEPTANCE OF GENERATIVE AI CAREER SUPPORT

Final analytic sample	350	96.4
Completed screening criteria	350	100
Duplicate submissions retained in analytic file	0	0
Full-block invariant responding retained	0	0
Missing item values	0	0
Completion time (seconds)	168.51	24.12
Career stage: senior undergraduate	65	18.6
Career stage: recent graduate	76	21.7
Career stage: graduate student	97	27.7
Career stage: early-career worker	92	26.3
Career stage: other	20	5.7
Used AI for career-related issues in past six months	275	78.6
Currently in job-search/transition/uncertainty/setback stage	239	68.3
AI tool familiarity (1–7)	4.84	1.12
Comfort using AI for important personal decisions (1–7)	4.34	1.24

Note. Percentages for the initial-return and exclusion rows are based on 363 submitted records; all other percentages are based on the final analytic sample (N = 350). The 13 excluded records were removed before model estimation and hypothesis testing because they showed one or more data-quality problems, including very short completion times, invariant item-block responses, or clearly patterned responding. All scale variables used seven-point response formats.

5.2. Measurement Model, Reliability, and Discriminant Validity

The five-factor measurement model (PEQ, EA, RL, PAI, and PC) demonstrated excellent fit to the data: $\chi^2(220) = 238.55$, $\chi^2/df = 1.08$, CFI = 0.996, TLI = 0.995, RMSEA = .016, and SRMR = 0.039. These indices were interpreted jointly, following recommended SEM practice [59].

All standardized factor loadings were strong, ranging from 0.718 to .830. As re-reported in Table 4, Cronbach's alpha and composite reliability values exceeded 0.85 across all constructs, and average variance extracted (AVE) values were above 0.57, supporting both internal consistency and convergent validity.

Table 4. Descriptive statistics, reliability, and convergent validity.

Construct	Items	Mean	SD	α	CR	AVE	Loading range
PEQ	5	5.27	0.8	0.87	0.871	0.574	0.718–0.785
EA	5	5.15	0.86	0.878	0.878	0.591	0.736–0.793
RL	5	5.15	0.9	0.884	0.884	0.605	0.732–0.818
PAI	4	5.04	0.93	0.872	0.873	0.631	0.766–0.830
PC	4	5.12	0.86	0.855	0.857	0.6	0.720–0.818

Note. α = Cronbach's alpha; CR = composite reliability; AVE = average variance extracted.

Discriminant validity was also well supported (Table 5). The square root of AVE for each construct exceeded its correlations with other constructs, and the highest HTMT value (0.653) remained below the 0.85 threshold. This result is

theoretically important because the model depends on distinguishing perceived empathy quality, emotional authenticity, and role legitimacy as separate constructs rather than treating them as a single positive evaluation of AI.

Table 5. Correlations and discriminant validity.

Construct	1	2	3	4	5
1. PEQ	0.758	0.653	0.573	0.56	0.274
2. EA	0.570***	0.769	0.251	0.56	0.197

3. RL	0.503***	0.222***	0.778	0.48	0.126
4. PAI	0.488***	0.490***	0.422***	0.795	0.365
5. PC	0.235***	0.171**	0.107*	0.316***	0.774

Note. Diagonal values are the square roots of AVE. Lower-triangle values are Pearson correlations. Upper-triangle values are HTMT ratios. * $p < 0.05$; ** $p < 0.01$; *** $p < 0.001$.

Further support comes from alternative model comparisons (Table 6). Models that merged PEQ with EA, EA with RL, or RL with PC showed substantially poorer fit, and the one-factor model

performed poorly. These findings confirm that the focal constructs are empirically distinct and align with the study's theoretical boundary definitions.

Table 6. Alternative measurement-model comparisons.

Model	χ^2	df	CFI	TLI	RMSEA	SRMR
Five-factor	238.55	220	0.996	0.995	0.016	0.039
PEQ+EA merged	606.47	224	0.909	0.897	0.07	0.071
EA+RL merged	1089.09	224	0.793	0.766	0.105	0.111
RL+PC merged	857.21	224	0.849	0.829	0.09	0.11
One-factor	2120.84	230	0.548	0.502	0.153	0.142

5.3. Structural Model and Hypothesis Tests

The item-level structural model also demonstrated excellent fit: $\chi^2(222) = 239.05$, $\chi^2/df = 1.08$, CFI = 0.996, TLI = 0.995, RMSEA = 0.015, and SRMR = 0.039.

As specified in Table 7, perceived AI empathy quality (PEQ) significantly predicted both emotional authenticity (EA) ($\beta = 0.570$, $p < 0.001$) and role legitimacy (RL) ($\beta = 0.503$, $p < 0.001$), supporting H1 and H4.

Consistent with the proposed dual-path model, both EA ($\beta = 0.339$, $p < 0.001$) and RL ($\beta = 0.266$, $p < 0.001$) were positively associated with professional acceptance intention (PAI), supporting H2 and H5. The direct path from PEQ to PAI remained significant but reduced ($\beta =$

0.160, $p = 0.007$), indicating partial rather than full mediation.

When perceived competence (PC) was introduced as a rival explanation, EA ($\beta = 0.328$, $p < 0.001$) and RL ($\beta = 0.269$, $p < .001$) remained significant, while PC also showed a positive association with PAI ($\beta = 0.204$, $p < 0.001$). This pattern supports H7 and indicates that professional acceptance is not reducible to perceived competence alone.

Overall, the structural results support the central theoretical claim: empathic response quality does not directly translate into acceptance, but operates through relational (EA) and role-based (RL) judgments.

Table 7. Structural path estimates.

Hypothesis/path	Path	Model 1 β	p	Model 2 β	p	Conclusion
H1	PEQ $\rightarrow$ EA	0.57	< 0.001	0.57	< 0.001	Supported
H4	PEQ $\rightarrow$ RL	0.503	< 0.001	0.503	< 0.001	Supported
H2	EA $\rightarrow$ PAI	0.339	< 0.001	0.328	< 0.001	Supported
H5	RL $\rightarrow$ PAI	0.266	< 0.001	0.269	< 0.001	Supported
Baseline	PEQ $\rightarrow$ PAI	0.16	0.007	0.117	0.046	
Rival	PC $\rightarrow$ PAI	—	—	0.204	< 0.001	Significant
H7	EA and RL unique effects with PC	—	—	Both significant	—	

Note. Model 1 estimates PAI from PEQ, EA, and RL. Model 2 adds PC as a rival explanation. Coefficients are standardized.

5.4. Indirect Effects and Robustness Checks

Bootstrap analyses (5,000 resamples) confirmed both mediation pathways (Table 8). The indirect effect through emotional authenticity was significant in both the main model ($\beta = 0.194$,

95% CI [0.134, 0.259]) and the model including perceived competence ($\beta = 0.187$, 95% CI [0.130, 0.251]), supporting H3.

WHEN EMPATHY IS NOT ENOUGH: EMOTIONAL AUTHENTICITY AND ROLE LEGITIMACY IN THE PROFESSIONAL ACCEPTANCE OF GENERATIVE AI CAREER SUPPORT

Similarly, the indirect effect through role legitimacy was significant in both specifications (main model: $\beta = 0.134$, 95% CI [0.077, 0.197]; PC model: $\beta = 0.135$, 95% CI [0.078, 0.201]), supporting H6. The total indirect effect remained stable and significant across models.

These findings reinforce the interpretation that professional acceptance is a conversion process: users do not simply carry forward their evaluation of empathy, but reinterpret it through judgments of credible care and appropriate role positioning.

Table 8. Structural path estimates.

Effect	Model 1 estimate	95% CI	Model 2 estimate	95% CI	Conclusion
H3: PEQ → EA → PAI	0.194	[0.134, 0.259]	0.187	[0.130, 0.251]	Supported
H6: PEQ → RL → PAI	0.134	[0.077, 0.197]	0.135	[0.078, 0.201]	Supported
Total indirect effect	0.328	[0.242, 0.422]	0.323	[0.241, 0.414]	Significant
Direct effect	0.16	[0.047, 0.264]	0.117	[0.009, 0.223]	Reduced

Note. Confidence intervals are percentile bootstrap intervals based on 5,000 resamples. Estimates are standardized indirect associations. Model 2 includes PC as a rival explanation.

Robustness checks further support this conclusion. Regression analyses using composite scores and additional control variables (AI familiarity, comfort with AI, prior AI use, career uncertainty, and scenario realism) produced

consistent results (Table 9). Emotional authenticity, role legitimacy, and perceived competence remained significant predictors of professional acceptance, while the direct effect of PEQ became marginal.

Table 9. Structural path estimates.

Specification	N	PEQ → PAI	EA → PAI	RL → PAI	PC → PAI	R ²
Controls + scenario realism	350	0.113 (p = 0.056)	0.320 (p < 0.001)	0.263 (p < 0.001)	0.202 (p < .001)	0.403
Excluding low-SR cutoff cases	295	0.079 (p = 0.227)	0.366 (p < 0.001)	0.250 (p < 0.001)	0.182 (p < .001)	0.36

Note. Coefficients are standardized estimates based on scale composites. The full-sample model includes AI familiarity, comfort using AI for important personal decisions, prior career-related AI use, current career setback status, and scenario realism. The low-SR sensitivity analysis excludes cases at or below the lowest-decile cutoff for scenario realism; ties at the cutoff are excluded, yielding N = 295. SR = scenario realism.

A sensitivity analysis excluding respondents with low perceived scenario realism yielded similar results. The stability of coefficients across specifications suggests that the findings are not driven by respondents who found the vignette unrealistic or disengaging.

6. DISCUSSION: CONTRIBUTIONS, IMPLICATIONS, AND BOUNDARY CONDITIONS

6.1. Theoretical contributions

This study advances research on AI-based support by showing that professional acceptance of empathic AI is best understood as a bounded status-granting process, rather than a direct reaction to supportive language. The findings demonstrate that perceived AI empathy quality is associated with professional acceptance intention through two distinct pathways—emotional authenticity and role legitimacy—and that both remain significant even after accounting for perceived competence and scenario realism. In

other words, even when AI produces responses that appear empathic, users do not automatically accept it as a support provider. They must first interpret the care as credible and the role as appropriate before granting AI a limited position in the support process.

This reframing contributes to AI empathy research by shifting the focus from re-sponse-level evaluation to interpretive conditions of acceptance. Prior work has largely examined whether AI can generate language that appears empathic. The present study shows that such performance is only the starting point. Empathic responses must be translated into relational credibility and role authorization before they can support professional acceptance.

Second, the study conceptualizes professional acceptance as a dual-judgment process. Emotional authenticity and role legitimacy are empirically distinct and retain independent explanatory power even when perceived competence is considered. This distinction is theoretically important because it demonstrates

that users do not evaluate AI solely in terms of usefulness or capability. Instead, they make two separate judgments: whether the care feels believable and whether AI belongs in the support role. By identifying these mechanisms, the study extends AI acceptance research beyond general constructs such as trust, usefulness, or competence.

The relative strength of the two pathways further refines this contribution. The in-direct effect through emotional authenticity is consistently larger than that through role legitimacy, suggesting that users first need to interpret AI responses as credible care. However, authenticity alone is not sufficient: role legitimacy remains necessary to authorize AI as a first-line support provider. Together, these findings clarify the sequential logic of acceptance without reducing it to a single evaluative dimension.

Third, the study positions career support as a theoretically revealing intermediate context. Unlike purely informational services, career support involves emotional strain, self-interpretation, and future-oriented decision-making. At the same time, it lacks the institutional protections of clinical settings. This combination makes both relational credibility and role boundaries especially salient. By situating the model in this context, the study highlights how AI acceptance depends on the interaction between emotional meaning and professional role expectations.

6.2. Practical Implications

The findings suggest that effective AI career-support systems require more than well-phrased empathic responses. They require design strategies that address how users interpret both the care provided and the role of the system.

First, systems should prioritize authenticity-oriented design rather than optimizing generic empathic wording. Because emotional authenticity mediates the relationship between perceived empathy and acceptance, responses should move beyond formulaic reassurance. This includes referencing the user's specific situation, acknowledging uncertainty, and linking emotional validation to context-sensitive next steps. The goal is not simply to state that the system "understands," but to demonstrate understanding through situationally grounded support.

Second, systems should incorporate role-boundary design. Role legitimacy significantly influences acceptance, indicating that users care about whether AI is appropriately positioned. A more defensible approach is to present AI as a first-line support assistant that provides initial clarification, emotional containment, and

preparatory guidance, rather than as a decision-maker. Clear role definitions, transparent limits, and guidance on when to seek human support can strengthen perceived legitimacy more effectively than anthropomorphic design.

Third, organizations should implement escalation design by embedding AI within a broader support ecosystem. While perceived competence contributes to acceptance, it does not replace the need for authenticity and legitimacy. For complex or high-stakes situations, systems should offer structured pathways for referral or human handoff. This is particularly important when users experience sustained uncertainty, repeated setbacks, or signs of distress.

From a sustainability perspective, these design principles are critical. AI-based career support can expand access to timely guidance, but scalability is responsible only when systems clearly communicate their limits, maintain credible interaction, and integrate with human support structures. In this way, AI can enhance service capacity without undermining professional standards or user trust.

6.3. Boundary Conditions and Future Research

Several limitations define the boundary conditions of this study and suggest directions for future research.

First, the study relies on a single vignette and a fixed AI response. While this design improves internal consistency, it limits generalizability across contexts and response styles. Future research should examine whether the proposed model holds across different types of career challenges (e.g., job-search failure, transition uncertainty, interview anxiety) and across variations in AI response style, including differences in emotional intensity, specificity, and role framing. Experimental designs that manipulate these features could provide stronger causal evidence.

Second, the study focuses on acceptance intention rather than observed behavior. Although intention is a meaningful indicator, future work could examine actual choices between AI and human support, repeated use, or escalation behavior. Designs that compare AI-only, human-only, and hybrid conditions would be particularly useful for testing whether role legitimacy influences real decision-making.

Third, the analysis centers on individual-level perceptions. Future research could extend the model by incorporating broader institutional factors, such as platform reputation, regulatory context, cultural expectations, and professional norms. Role legitimacy, in particular, may depend not only on the AI response itself but also on the

organization deploying the system and the institutional environment in which it operates.

7. Conclusions

Generative AI can now produce language that appears supportive, but the professional meaning of that support is not automatically established. This study shows that perceived AI empathy quality is associated with professional acceptance intention through two distinct mechanisms: emotional authenticity and role legitimacy. These mechanisms remain significant even when perceived competence and scenario realism are considered, indicating that users do not accept AI as a support provider simply because it performs well at the response level.

Instead, professional acceptance emerges when users interpret the care as credible, view the AI role as appropriate, and consider the system sufficiently capable for the task. In the context of career support, this results in a bounded form of professional acceptance, in which AI may be accepted as an initial support node but not as a full substitute for human professionals.

The implication is not that AI should replace human career counselors, but that it can play a meaningful role within a structured support system. When emotional authenticity, role boundaries, competence, and escalation pathways are clearly defined, AI can contribute to expanding access to career support while maintaining responsible and sustainable service design.

Supplementary Materials: Text S1: Vignette stimulus and AI response used in the survey; Figure S1: Full measurement and structural model; Table S1: Hypothesis–construct–item mapping; Table S2: Full questionnaire aligned with the final field version.

REFERENCES

- Abbott, A., *The System of Professions: An Essay on the Division of Expert Labor*. 2014: University of Chicago Press.
- Ayers, J.W., et al., "Comparing physician and artificial intelligence chatbot responses to patient questions posted to a public social media forum." *JAMA Internal Medicine*, 2023. 183(6): p. 589-596.
- Bitektine, A., "Toward a theory of social judgments of organizations: The case of legitimacy, reputation, and status." *Academy of Management Review*, 2011. 36(1): p. 151-179.
- Böhm, R., et al., "People devalue generative AI's competence but not its advice in addressing societal and personal challenges." *Communications Psychology*, 2023. 1(1): p. 32.
- Brown, S.D., et al., "Critical ingredients of career choice interventions: More analyses and new hypotheses." *Journal of Vocational Behavior*, 2003. 62(3): p. 411-428.
- Bui, H.T., V. Filimonau, and H. Sezerel, "AI-thenticity: Exploring the effect of perceived authenticity of AI-generated visual content on tourist patronage intentions." *Journal of Destination Marketing & Management*, 2024. 34: p. 100956.
- Bunduchi, R., D.-A. Sitar-Tăut, and D. Mican, "A legitimacy-based explanation for user acceptance of controversial technologies: The case of Generative AI." *Technological Forecasting and Social Change*, 2025. 215: p. 124095.
- Castelo, N., M.W. Bos, and D.R. Lehmann, "Task-dependent algorithm aversion." *Journal of Marketing Research*, 2019. 56(5): p. 809-825.

Author Contributions: Conceptualization, Z.Z.; methodology, Z.Z., J.A.; software, Z.Z., C.C.; validation, Z.Z., C.C. and J.A.; formal analysis, Z.Z.; investigation, Z.Z., C.C.; resources, Z.Z., J.A.; data curation, Z.Z.; writing—original draft preparation, Z.Z.; writing—review and editing, Z.Z., C.C. and J.A.; visualization, Z.Z., C.C.; supervision, J.A.; project administration, C.C, J.A. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Institutional Review Board Statement: The study was conducted in accordance with the Declaration of Helsinki. Ethical review and approval were waived for this study due to the anonymous and non-interventional nature of the online survey, which collected no direct personal identifiers and involved no clinical procedure or intervention.

Informed Consent Statement: All participants were informed about the purpose of the study, the anonymous handling of responses, and their right to withdraw before submitting the questionnaire. Informed consent was obtained from all participants prior to participation.

Data Availability Statement: The data presented in this study are available on request from the corresponding author. The data are not publicly available due to privacy and ethical restrictions related to anonymous survey responses.

Acknowledgments: The authors would like to thank all survey participants for their time and responses.

Conflicts of Interest: The authors declare no conflicts of interest.

- Davenport, T., et al., "How artificial intelligence will change the future of marketing." *Journal of the Academy of Marketing Science*, 2020. 48(1): p. 24-42.
- Deephhouse, D.L. and M. Suchman, "Legitimacy in organizational institutionalism." *The Sage Handbook of Organizational Institutionalism*, 2008. 49(77): p. 273-289.
- Dietvorst, B.J., J.P. Simmons, and C. Massey, "Algorithm aversion: people erroneously avoid algorithms after seeing them err." *Journal of Experimental Psychology: General*, 2015. 144(1): p. 114.
- Fornell, C. and D.F. Larcker, "Evaluating structural equation models with unobservable variables and measurement error." *Journal of Marketing Research*, 1981. 18(1): p. 39-50.
- Gati, I. and L. Asulin-Peretz, "Internet-based self-help career assessments and interventions: Challenges and implications for evidence-based career counseling." *Journal of Career Assessment*, 2011. 19(3): p. 259-273.
- Gati, I. and N. Levin, "Counseling for career decision- making difficulties: Measures and methods." *The Career Development Quarterly*, 2014. 62(2): p. 98-113.
- Granulo, A., C. Fuchs, and S. Puntoni, "Preference for human (vs. robotic) labor is stronger in symbolic consumption contexts." *Journal of Consumer Psychology*, 2021. 31(1): p. 72-80.
- Grayson, K. and R. Martinec, "Consumer perceptions of iconicity and indexicality and their influence on assessments of authentic market offerings." *Journal of Consumer Research*, 2004. 31(2): p. 296-312.
- Gray, H.M., K. Gray, and D.M. Wegner, "Dimensions of mind perception." *Science*, 2007. 315(5812): p. 619-619.
- Hayton, J.C., D.G. Allen, and V. Scarpello, "Factor retention decisions in exploratory factor analysis: A tutorial on parallel analysis." *Organizational Research Methods*, 2004. 7(2): p. 191-205.
- Henseler, J., C.M. Ringle, and M. Sarstedt, "A new criterion for assessing discriminant validity in variance-based structural equation modeling." *Journal of the Academy of Marketing Science*, 2015. 43(1): p. 115-135.
- Hermann, E. and S. Puntoni, "Artificial intelligence and consumer behavior: From predictive to generative AI." *Journal of Business Research*, 2024. 180: p. 114720.
- Hohenstein, J., et al., "Artificial intelligence in communication impacts language and social relationships." *Scientific Reports*, 2023. 13(1): p. 5487.
- Hu, L.T. and P.M. Bentler, "Cutoff criteria for fit indexes in covariance structure analysis: Conventional criteria versus new alternatives." *Structural Equation Modeling: A Multidisciplinary Journal*, 1999. 6(1): p. 1-55.
- Huang, M.-H. and R.T. Rust, "Artificial intelligence in service." *Journal of Service Research*, 2018. 21(2): p. 155-172.
- Jakesch, M., et al., "AI-mediated communication: How the perception that profile text was written by AI affects trustworthiness." In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems*. 2019.
- Juquelier, A., I. Poncin, and S. Hazée, "Empathic chatbots: A double-edged sword in customer experiences." *Journal of Business Research*, 2025. 188: p. 115074.
- Kirk, C.P. and J. Givi, "The AI-authorship effect: Understanding authenticity, moral disgust, and consumer responses to AI-generated marketing communications." *Journal of Business Research*, 2025. 186: p. 114984.
- Li, Y., et al., "Humans as teammates: The signal of human-AI teaming enhances consumer acceptance of chatbots." *International Journal of Information Management*, 2024. 76: p. 102771.
- Li, Y., Y. Chang, and Y. Li, "1+1<2? Unveiling the impact of AI-assisted disclosure on service satisfaction in professional services." *International Journal of Information Management*, 2025. 84: p. 102937.
- Liu, B. and S.S. Sundar, "Should machines express sympathy and empathy? Experiments with a health advice chatbot." *Cyberpsychology, Behavior, and Social Networking*, 2018. 21(10): p. 625-636.
- Longoni, C. and L. Cian, "Artificial intelligence in utilitarian vs. hedonic contexts: The 'word-of-machine' effect." *Journal of Marketing*, 2022. 86(1): p. 91-108.
- Longoni, C., A. Bonezzi, and C.K. Morewedge, "Resistance to medical artificial intelligence." *Journal of Consumer Research*, 2019. 46(4): p. 629-650.
- Luo, X., et al., "Frontiers: Machines vs. humans: The impact of artificial intelligence chatbot disclosure on customer purchases." *Marketing Science*, 2019. 38(6): p. 937-947.
- Lv, X., et al., "Artificial intelligence service recovery: The role of empathic response in hospitality customers' continuous usage intention." *Computers in Human Behavior*, 2022. 126: p. 106993.
- Lv, X., et al., "AI service may backfire: Reduced service warmth due to service provider transformation." *Journal of Retailing and Consumer Services*, 2025. 85: p. 104282.
- MacKenzie, S.B., P.M. Podsakoff, and N.P. Podsakoff, "Construct measurement and validation procedures in MIS and behavioral research: Integrating new and existing techniques." *MIS Quarterly*, 2011. 35(2): p. 293-A5.
- Markovitch, D.G., R.A. Stough, and D. Huang, "Consumer reactions to chatbot versus human service: an investigation in the role of outcome valence and perceived empathy." *Journal of Retailing and Consumer Services*, 2024. 79: p. 103847.
- Milot- Lapointe, F. and N. Arifouline, "A Meta- Analysis of the Effectiveness of Individual Career Counseling on Career and Mental Health Outcomes." *Journal of Employment Counseling*, 2025. 62(1): p. 49-57.
- Milot-Lapointe, F., Y. Le Corff, and N. Arifouline, "A meta-analytic investigation of the association between working alliance and outcomes of individual career counseling." *Journal of Career Assessment*, 2021. 29(3): p. 486-501.
- Morhart, F., et al., "Brand authenticity: An integrative framework and measurement scale." *Journal of Consumer Psychology*, 2015. 25(2): p. 200-218.
- Nass, C., J. Steuer, and E.R. Tauber, "Computers are social actors." In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*. 1994.
- Pandya, S.S. and J. Wang, "Artificial intelligence in career development: a scoping review." *Human Resource Development International*, 2024. 27(3): p. 324-344.

WHEN EMPATHY IS NOT ENOUGH: EMOTIONAL AUTHENTICITY AND ROLE LEGITIMACY IN THE PROFESSIONAL ACCEPTANCE OF GENERATIVE AI CAREER SUPPORT

- Peng, C., et al., "The effect of required warmth on consumer acceptance of artificial intelligence in service: The moderating role of AI-human collaboration." *International Journal of Information Management*, 2022. 66: p. 102533.
- Podsakoff, P.M., et al., "Common method biases in behavioral research: a critical review of the literature and recommended remedies." *Journal of Applied Psychology*, 2003. 88(5): p. 879.
- Preacher, K.J. and A.F. Hayes, "Asymptotic and resampling strategies for assessing and comparing indirect effects in multiple mediator models." *Behavior Research Methods*, 2008. 40(3): p. 879-891.
- Reis, H.T., E.P. Lemay Jr, and C. Finkenauer, "Toward understanding understanding: The importance of feeling understood in relationships." *Social and Personality Psychology Compass*, 2017. 11(3): p. e12308.
- Robertson, P.J., "The well-being outcomes of career guidance." *British Journal of Guidance & Counselling*, 2013. 41(3): p. 254-266.
- Rubin, M., et al., "Comparing the value of perceived human versus AI-generated empathy." *Nature Human Behaviour*, 2025: p. 1-15.
- Schilke, O. and M. Reimann, "The transparency dilemma: How AI disclosure erodes trust." *Organizational Behavior and Human Decision Processes*, 2025. 188: p. 104405.
- Sharma, A., et al., "Human-AI collaboration enables more empathic conversations in text-based peer-to-peer mental health support." *Nature Machine Intelligence*, 2023. 5(1): p. 46-57.
- Suchman, M.C., "Managing legitimacy: Strategic and institutional approaches." *Academy of Management Review*, 1995. 20(3): p. 571-610.
- Suddaby, R., A. Bitektine, and P. Haack, "Legitimacy." *Academy of Management Annals*, 2017. 11(1): p. 451-478.
- Tost, L.P., "An integrative model of legitimacy judgments." *Academy of Management Review*, 2011. 36(4): p. 686-710.
- Van Doorn, J., et al., "Domo arigato Mr. Roboto: Emergence of automated social presence in organizational frontlines and customers' service experiences." *Journal of Service Research*, 2017. 20(1): p. 43-58.
- Vandenberg, R.J. and C.E. Lance, "A review and synthesis of the measurement invariance literature: Suggestions, practices, and recommendations for organizational research." *Organizational Research Methods*, 2000. 3(1): p. 4-70.
- Wang, Y.A. and M. Rhemtulla, "Power analysis for parameter estimation in structural equation modeling: A discussion and tutorial." *Advances in Methods and Practices in Psychological Science*, 2021. 4(1): p. 2515245920918253.
- Wenger, J.D., C.D. Cameron, and M. Inzlicht, "People choose to receive human empathy despite rating AI empathy higher." *Communications Psychology*, 2026.
- Whiston, S.C., J. Rossier, and P.M.H. Barón, "The working alliance in career counseling: A systematic overview." *Journal of Career Assessment*, 2016. 24(4): p. 591-604.
- Willis, G.B., *Cognitive Interviewing: A Tool for Improving Questionnaire Design*. 2004: Sage Publications.
- Wolf, E.J., et al., "Sample size requirements for structural equation models: An evaluation of power, bias, and solution propriety." *Educational and Psychological Measurement*, 2013. 73(6): p. 913-934.
- Yin, Y., N. Jia, and C.J. Wakslak, "AI can help people feel heard, but an AI label diminishes this impact." *Proceedings of the National Academy of Sciences*, 2024. 121(14): p. e2319112121.