

Psychometric Properties of the General Attitude toward Artificial Intelligence Scale

Esmail Sadri Damirchi 

Department of Counseling, University of Mohaghegh Ardabili, Ardabil, Iran. E-mail: E.sadri@uma.ac.ir

Nasser Abbasi 

Department of Counseling, University of Mohaghegh Ardabili, Ardabil, Iran. E-mail: n.abbasi@uma.ac.ir

Maryam Ghahremanloo 

Department of Counseling, University of Mohaghegh Ardabili, Ardabil, Iran. E-mail: m.ghahremanloo@uma.ac.ir

Ali Ghorbaninejad 

Department of Counseling, University of Mohaghegh Ardabili, Ardabil, Iran. E-mail: Alighorbaninejad@atu.ac.ir

Mohammadreza Noroozi Homayoon  *

Corresponding Author, Department of Counseling, University of Mohaghegh Ardabili, Ardabil, Iran. E-mail: mohammadreza.noroozi@uma.ac.ir

Abstract

This study examined the psychometric properties of the General Attitude toward Artificial Intelligence Scale (GAAIS) in an Iranian sample. The research followed a descriptive survey design, and 414 participants were selected through convenience sampling. Data were collected using the General Attitude toward Artificial Intelligence Scale (Shepman & Radway, 2020) and the Technology Readiness Scale (Chen et al., 2014). To analyze the data, this study employed descriptive statistics, Pearson correlation tests, Cronbach's alpha coefficients, and confirmatory factor analysis (CFA). The CFA results supported a two-factor structure (positive attitude and negative attitude toward AI). Pearson correlation analyses revealed significant positive relationships between both subscales of the General Attitude toward AI Scale and the four subscales of the Technology Readiness Scale, supporting concurrent validity. Test-retest reliability was assessed through two administrations, yielding correlation coefficients of 0.69 and 0.74, confirming good temporal stability. These findings demonstrate that the General Attitude toward AI Scale is a valid and reliable measure for assessing this construct in Iranian populations.

Keywords: Attitude toward Artificial Intelligence; Psychometric Properties; Artificial Intelligence

Cite this Article: Sadri Damirchi, S., Abbasi, N., Ghahremanloo, M., Ghorbaninejad, A., & Noroozi Homayoon, M. (2025). Psychometric Properties of the General Attitude toward Artificial Intelligence Scale. *Educational Measurement*, 16(61), 159-181. <https://doi.org/10.22054/jem.2025.82349.3576>



© 2016 by Allameh Tabataba'i University Press
Publisher: Allameh Tabataba'i University Press

Introduction

The growing integration of Artificial Intelligence (AI) across diverse domains, particularly in scientific research, has established it as a pivotal tool for accelerating and enhancing research processes. AI supports researchers in hypothesis generation, experimental design, collection and interpretation of complex datasets, and deriving insights that may remain inaccessible through conventional scientific approaches (Wang et al., 2023). Recent rapid advancements in AI technology have substantially influenced various social systems, including economic, political, educational, and psychological domains (Kaya et al., 2024).

AI technology plays a pivotal role in enhancing diagnostic and therapeutic procedures for both mental and physical health conditions, improving efficiency while minimizing human error (Gansser & Reich, 2021; Hartwig, 2021). However, these benefits are accompanied by significant risks and challenges, including diminished accountability and pressing ethical concerns that demand careful consideration (Bonnefon, Rahwan, & Shariff, 2023; Said et al., 2023). This paper systematically examines the dual aspects of AI implementation in psychology and mental health, highlighting the crucial need to evaluate public attitudes toward AI while developing psychometrically sound measurement tools (Grassini, 2023).

Methodology

This study utilized a survey-based research design. The target population comprised all Ardabil residents aged 25-60 during the 2023-2024 period. Through convenience sampling of technology-engaged individuals, we recruited 450 participants, with 414 questionnaires meeting inclusion criteria (98 men, 316 women). Participants received full disclosure of study objectives, confidentiality assurances, and provided informed consent prior to questionnaire completion.

Instruments

1. The General Attitude toward Artificial Intelligence Scale (GAAIS), developed by Schepman and Rodway (2020), measures individuals' attitudes toward AI. This 20-item instrument comprises two subscales: positive attitudes (12 items) and negative attitudes (8 items). Participants respond using a 5-point Likert scale ranging from "strongly disagree" to "strongly agree," with negative items reverse-scored. The total score is calculated by summing both subscale scores. In the

original validation study (Schepman & Rodway, 2022), the scale demonstrated strong internal consistency, with Cronbach's $\alpha = 0.88$ for the positive attitude subscale and $\alpha = 0.82$ for the negative attitude subscale.

2. The Technology Readiness Index (TRI), developed by Chen et al. (2014), assesses individuals' readiness to adopt new technologies through 18 items measuring four dimensions: optimism, innovativeness, discomfort, and insecurity. The scale employs a Likert-type response format ranging from "very low" to "very high." Following Nouri's (2015) validation, discomfort and insecurity subscales are reverse-scored, with total scores representing the sum of all subscale scores. Reliability analyses demonstrated strong internal consistency, with Cronbach's $\alpha = 0.82$ in Nouri's (2015) study and $\alpha = 0.79$ in the current study.

Results

This study evaluated overall attitudes toward artificial intelligence (AI) in an Iranian sample using the Positive Attitude and Negative Attitude subscales. As shown in Table 1, descriptive statistics revealed the following: the Positive Attitude subscale showed a mean of 18.84 (SD = 4.50), the Negative Attitude subscale a mean of 16.05 (SD = 3.90), and the total score a mean of 34.77 (SD = 5.59). The observed standard deviations suggest considerable variability in attitudes toward AI among participants.

The scale's validity was evaluated through three approaches: face validity, concurrent validity, and construct validity. Face validity was established via expert review by two counseling professors and five counseling Ph.D. students. Concurrent validity was demonstrated through significant correlations between the General Attitude toward AI Scale and the Technology Readiness Index (TRI), with $r = 0.26$ ($p < 0.01$) for positive attitudes and $r = 0.28$ ($p < 0.01$) for negative attitudes (Table 2). Construct validity was supported by strong correlations between the total score and subscales: $r = 0.72$ ($p < 0.01$) with the positive attitude subscale and $r = 0.60$ ($p < 0.01$) with the negative attitude subscale (Table 2). These results collectively indicate robust psychometric properties for the scale.

The scale's reliability was evaluated using Cronbach's alpha and test-retest methods. As shown in Table 3, Cronbach's alpha for positive attitude was 0.73 in the first administration and 0.79 in the second

administration. For the entire scale, Cronbach's alpha was 0.77, indicating satisfactory internal consistency. Test-retest reliability coefficients showed good stability over time, with 0.69 for positive attitude, 0.74 for negative attitude, and 0.72 for the total score.

Exploratory and confirmatory factor analyses were conducted to examine construct validity. The exploratory factor analysis yielded two factors accounting for 65.69% of the total variance (Table 5). Factor loadings indicated 12 items loaded on the first factor (positive attitude) and 8 items on the second factor (negative attitude), using varimax rotation with a 0.50 loading threshold. Confirmatory factor analysis showed good model fit (Table 6), with fit indices including CFI = 0.88, GFI = 0.91, and RMSEA = 0.063, all indicating acceptable model fit.

Discussion

"This study aimed to assess the psychometric properties of the General Attitude toward Artificial Intelligence Scale in an Iranian sample. The results of statistical analyses, including confirmatory factor analysis (CFA), confirmed the presence of two subscales: positive and negative attitudes toward artificial intelligence (AI). These findings are consistent with those of Schepman and Rodway (2020), suggesting that the two-factor structure of the scale can be effectively applied in the Iranian context. In this research, face validity was confirmed by experts in the field of counseling, indicating the clarity and cultural relevance of the scale in Iran.

Concurrent validity was also confirmed through comparison with the Technology Readiness Index, as the correlations between positive and negative attitudes toward AI were significant, supporting the scale's adequate validity. Construct validity was assessed through structural equation modeling (SEM) and CFA, with all fit indices demonstrating satisfactory results, confirming that the scale has appropriate construct validity. Additionally, reliability was verified using Cronbach's alpha and test-retest reliability, indicating the scale's internal consistency and temporal stability. Overall, the findings of this study confirm that the General Attitude toward Artificial Intelligence Scale, which includes both positive and negative attitude components, exhibits satisfactory validity and reliability in the Iranian population.

The results indicate that this scale can be a valid tool for assessing attitudes toward AI in various social and research contexts. However, this study has some limitations. The sample was limited to individuals

residing in Ardabil province, which may challenge the generalizability of the findings to the entire Iranian population. Additionally, future research is recommended to assess the scale in different populations and age groups to further ensure comprehensive validation and reliability testing.

Conclusion

This study evaluated the psychometric properties of the General Attitude toward Artificial Intelligence Scale in an Iranian sample. Confirmatory factor analysis supported the scale's two-factor structure (positive and negative attitudes toward AI), demonstrating adequate validity and reliability. Face validity was established through expert evaluation, while concurrent validity was confirmed via correlation analyses. Limitations include the regional restriction of the sample to Ardabil province, which may affect generalizability. Future studies should validate the scale across diverse demographic groups and geographic regions to strengthen its applicability.

Conflict of Interest

The authors declare no personal or organizational conflicts of interest regarding this study or its findings.

Acknowledgments

The authors gratefully acknowledge all participants for their valuable contributions to this study.

ویژگی‌های روان‌سنجی مقیاس نگرش کلی به هوش مصنوعی

اسماعیل صدری دمیرچی

گروه مشاوره، دانشگاه محقق اردبیلی، اردبیل، ایران. رایانامه:
E.sadri@uma.ac.ir

ناصر عباسی

گروه مشاوره، دانشگاه محقق اردبیلی، اردبیل، ایران. رایانامه:
n.abbasi@uma.ac.ir

مریم قهرمانلو

گروه مشاوره، دانشگاه محقق اردبیلی، اردبیل، ایران. رایانامه:
m.ghahremanloo@uma.ac.i

علی قربانی نژاد

گروه مشاوره، دانشگاه محقق اردبیلی، اردبیل، ایران. رایانامه:
Alighorbaninejad@atu.ac.ir

محمد رضا نوروزی همایون*

نویسنده مسئول، گروه مشاوره، دانشگاه محقق اردبیلی، اردبیل، ایران.
رایانامه: Mohammadreza.noroozi@uma.ac.ir

چکیده

هدف پژوهش حاضر محاسبه ویژگی‌های روان‌سنجی مقیاس نگرش کلی به هوش مصنوعی در نمونه‌ی ایرانی بود. روش پژوهش حاضر توصیفی از نوع پیمایشی بود. برای جمع‌آوری داده‌ها تعداد ۴۱۴ نفر به روش نمونه‌گیری در دسترس انتخاب شدند؛ و برای جمع‌آوری داده‌ها از پرسشنامه نگرش کلی به هوش مصنوعی شپمن و رادوی (۲۰۲۰) و مقیاس آمادگی برای پذیرش فناوری چن و همکاران (۲۰۱۴) استفاده گردید. برای تجزیه و تحلیل داده‌ها علاوه بر آمار توصیفی، از آزمون همبستگی پیرسون، ضریب آلفای کرونباخ و تحلیل عاملی تأییدی استفاده شد. نتایج تحلیل عاملی تأییدی، دو عامل نگرش مثبت و نگرش منفی به هوش مصنوعی را تأیید کرد. نتایج تحلیل ضریب همبستگی پیرسون برای بررسی روایی هم‌زمان دو خرده مقیاس نگرش کلی نسبت به هوش مصنوعی با چهار خرده مقیاس آمادگی برای پذیرش فناوری همبستگی مثبت و معناداری را نشان داد. پایایی بازآزمایی مقیاس نگرش نسبت به هوش مصنوعی بر اساس نتایج دوبار اجرای آزمون محاسبه و با ضرایب همبستگی ۰/۶۹ و ۰/۷۴ مورد تأیید قرار گرفت؛ بنابراین، مقیاس نگرش کلی به هوش مصنوعی برای سنجش این سازه در نمونه‌ی ایرانی از روایی و پایایی لازم برخوردار است.

کلیدواژه‌ها: نگرش به هوش مصنوعی، ویژگی‌های روان‌سنجی، هوش مصنوعی

استناد به این مقاله: صدری دمیرچی، اسماعیل، عباسی، ناصر، قهرمانلو، مریم، قربانی نژاد، علی، و نوروزی همایون، محمد رضا. (۱۴۰۴). ویژگی‌های روان‌سنجی مقیاس نگرش کلی به هوش مصنوعی. فصلنامه اندازه‌گیری تربیتی، ۱۶(۶۱)، ۱۵۹-۱۸۱. <https://doi.org/10.22054/jem.2025.82349.3576>

مقدمه

امروزه با توجه به کاربرد روزافزون هوش مصنوعی^۱ (AI) در اکتشافات علمی برای تقویت و تسریع تحقیقات ادغام می‌شود و به دانشمندان در ایجاد فرضیه‌ها، طراحی آزمایش‌ها، جمع‌آوری و تفسیر مجموعه داده‌های بزرگ و به دست آوردن بینشی که ممکن است تنها با استفاده از روش‌های علمی سنتی امکان‌پذیر نباشد، کمک می‌کند (H. Wang et al., 2023). در سال‌های اخیر پیشرفت‌های سریع در تکنولوژی‌های هوش مصنوعی، همه سیستم‌های اجتماعی از جمله اقتصاد، سیاست، علم و آموزش را تحت تأثیر قرار داده است (Kaya et al., 2024). پس ناگزیر جامعه در چند دهه آینده عمیقاً تحت تأثیر هوش مصنوعی قرار خواهد گرفت و در حوزه‌های متعددی از جمله مسائل روان‌شناسی نیز رخنه خواهد کرد (Schepman & Rodway, 2023; B. Wang et al., 2023). Gansser and Reich (2021) هوش مصنوعی را به‌عنوان تکنولوژی توصیف می‌کنند که برای تسهیل زندگی انسان و کمک به افراد در موقعیت‌های خاص توسعه یافته است. از نظر Darko و همکاران (2021) هوش مصنوعی تکنولوژی کلیدی انقلاب صنعتی چهارم است. هوش مصنوعی به‌طور کلی به‌عنوان ماشینی با توانایی انجام وظایفی مانند توانایی محاسبه، تجزیه و تحلیل، استدلال، یادگیری و کشف معنا تعریف می‌شود که در تسهیل امور نقش مهمی را ایفا می‌کند (Federspiel et al., 2023). برخی از پژوهشگران این حوزه تأکید می‌کنند که هوش مصنوعی باعث افزایش بهره‌وری و کارایی می‌شود، فرصت‌های جدیدی را ایجاد می‌کند، خطاهای انسانی را کاهش می‌دهد، مسائل و مشکلات پیچیده را حل می‌کند و کارهای سخت و طاقت‌فرسایی را انجام می‌دهد؛ بنابراین، این مزایای هوش مصنوعی ممکن است برای افراد وقت بیشتری در جهت یادگیری، آزمایش و کشف ایجاد کند که این امر به نوبه خود می‌تواند خلاقیت و کیفیت زندگی بشر را افزایش دهد (Kaya et al., 2024). موضوعاتی همچون کارکردهای اجرایی، شبکه‌های دل‌بستگی و تنظیم شناختی هیجان (Hatami Nejad, Sadeghi, et al., 2025; Hatami Nejad et al., 2024; Hatami Nejad, Sadri Damirchi, et al., 2025; Noroozi Homayoon, Akhavi Samarein, et al., 2025; Noroozi Homayoon, Hatami Nejad, & Sadri Damirchi, 2024; Noroozi Homayoon, Hatami Nejad, Sadri Damirchi, et al., 2024; Noroozi Homayoon, Sadeghi, et al., 2024; Noroozi Homayoon, Sadri Damirchi, et al., 2025) از جمله موضوعات مهم و اساسی در حوزه روان‌شناسی است که هوش مصنوعی

1. artificial intelligence

می‌تواند به مداخلات درمانگران و متخصصان این حوزه کمک شایانی نماید. تکنولوژی‌های مرتبط با هوش مصنوعی و استراتژی‌های آن می‌تواند تقریباً در همه حوزه‌هایی که به رفتار و هوش انسانی در زمینه تصمیم‌گیری، برنامه‌های مراقبت‌های بهداشتی، درمان بیماران مبتلا به بیماری‌های پزشکی، مدیریت کسب‌وکار مربوط می‌شوند، بسیار مفید باشد (Parvathy & Mishra, 2023). همچنین این تکنولوژی در تشخیص بیماری‌های روانی و جسمی نیز می‌تواند به پزشکان، روان‌پزشکان، روان‌شناسان و مشاوران در فرایند تشخیص بیماری‌های جسمانی و روانی کمک کند. به‌عنوان مثال، بررسی سلامت روان در حالت سنتی عمدتاً بر اساس گزارش ذهنی مراجع از حالات شناختی و عاطفی خود و سیر علائمی که تجربه می‌کنند صورت می‌گیرد که این به نوبه خود باید توسط یک درمانگر ارزیابی و تشخیص داده شود که البته ممکن است تشخیص درمانگر در طول زمان درست یا غلط باشد. این تشخیص‌ها را به غلط و یا در مدت زمان طولانی‌تر متوجه شود؛ اما با گسترش هوش مصنوعی این تکنولوژی می‌تواند در جهت شناسایی و تشخیص اختلال مراجعان، به درمانگر در جهت تشخیص درست و سریع علائم بالینی کمک کند. پس این تکنولوژی می‌تواند به‌عنوان ابزار انقلابی برای تحقیقات روان‌پزشکی، روان‌شناسی و پزشکی عمل کند (Alimour et al., 2023; Espejo et al., 2023; Olawade et al., 2024; Yaacob et al., 2023). افزایش خودافشایی‌های مراجع در جلسات روان‌درمانی خواهد شد (Lee et al., 2024)؛ اما فارغ از مزایایی که هوش مصنوعی دارد، این تکنولوژی می‌تواند خطرات و تهدیدات جدی را ایجاد کند (Said et al., 2023). چهره‌های برجسته‌ای مانند استیون هاو‌کینگ^۱ هشدار داده‌اند که پیشرفت هوش مصنوعی در نهایت می‌تواند موجودیت انسان را به خطر بیندازد (Kumar & Choudhury, 2022). در مارس سال ۲۰۲۳، ماسک^۲ و چندین چهره کلیدی در زمینه هوش مصنوعی طی نامه‌ای سرگشاده خواستار توقف موقت توسعه هوش مصنوعی شدند (Letter, 2023). در کنار موارد مطرح‌شده، برخی از خطرات و تهدیداتی که در مورد گسترش هوش مصنوعی مطرح شده است شامل: کاهش مسئولیت‌پذیری، مانع‌تراشی در روند قانونی و تصمیم‌گیری‌های اشتباه است (Weidinger et al., 2022). به‌عنوان مثال، ماشین‌های هوش مصنوعی می‌توانند به‌عنوان یک عنصر اخلاقی عمل کرده و تصمیماتی بگیرند که بر نتایج بیماران انسانی تأثیر گذار باشد و یا به نحوی عمل کنند که برخی از مسائل

1. Stephen Hawking

2. Musk

اخلاقی را نادیده بگیرند (Bonnefon et al., 2023). اگرچه افراد دیگری بر استفاده از هوش مصنوعی به عنوان یک فرصت جدید برای آینده بشری تأکید دارند که هوش مصنوعی می تواند فرصتی برای آینده نوع بشر باشد (O'Shaughnessy et al., 2023; Samuel, 2023)؛ بنابراین با توجه به مزایا و معایبی که هوش مصنوعی می تواند داشته باشد و همچنین وجود مقیاس ها و ابزارهای کمی که به طور خاص با هدف ارزیابی نگرش نسبت به هوش مصنوعی گسترش پیدا کرده است در حوزه های مختلف به ویژه در مسائل روان شناسی بررسی نگرش و دیدگاه های افراد در مورد هوش مصنوعی بسیار مهم و ضروری است (Grassini, 2023). به عنوان مثال در پژوهشی Stein و همکاران (2024) نشان دادند که متغیرهایی همچون سن پایین فرد و نگرش موافق در این حوزه می تواند پیش بینی کننده خوبی در جهت داشتن دیدگاه های مثبت تری نسبت به کاربرد هوش مصنوعی و تکنولوژی های مرتبط در این حوزه باشد، درحالی که حساسیت به باورهای توطئه می تواند به نگرش منفی تری در این حوزه مرتبط باشد. Hadlington و همکاران (2023) در پژوهش خود در زمینه نگرش به هوش مصنوعی، ابزار اندازه گیری جدیدی را ارائه کرده اند که توانایی ارزیابی نگرش فعلی فرد نسبت به هوش مصنوعی را نشان می دهد. در پژوهش Kieslich و همکاران (2024) تکنولوژی هوش مصنوعی برای ارزیابی تخصصی در حوزه های مختلف بررسی شده است که دارای اعتبار و روایی مناسبی می باشند. در پژوهشی Eyüp and Kayhan (2023) نیز نشان دادند که اضطراب و نگرش معلمان قبل از اشتغال به کار نسبت به هوش مصنوعی تفاوتی نداشتند. همچنین در پژوهش Robertson و همکاران (2023) مشخص شد که جهت اثرگذاری هرچه بیشتر این تکنولوژی در بیماران، نگرش های مختلف بیماران نسبت به هوش مصنوعی از اهمیت بسیار زیادی برخوردار است. پژوهش Schepman and Rodway (2020) نیز مقیاس نگرش های عمومی نسبت به هوش مصنوعی را معرفی کرد که یک مقیاس ۲۰ ماده ای با ساختار دو عاملی (نگرش های مثبت و منفی نسبت به هوش مصنوعی) است. با این حال، تمامی این مقیاس ها با یکسری از چالش هایی در زمینه عملیاتی شدن و کاربرد احتمالی آن ها در زمینه ارزیابی نگرش کلی مواجه هستند؛ بنابراین درحالی که راه حل های مبتنی بر هوش مصنوعی در حوزه سلامت روان امیدوارکننده هستند، تحقیقات بیشتری برای پرداختن به جنبه های اخلاقی و اجتماعی این تکنولوژی ها و همچنین ایجاد تحقیقات و اقدامات پزشکی کارآمد در این بخش نوآورانه مورد نیاز است (Omarov et

al., 2023). در پژوهش حاضر، با توجه به گسترش و رشد هوش مصنوعی در جوامع و وجود دیدگاه و نگرش‌های متفاوت نسبت به این تکنولوژی، لازم است تا به هنجاریابی پرسشنامه نگرش کلی نسبت به هوش مصنوعی پرداخته شود؛ به عبارت دیگر در هدف مهم مطالعه ما هنجاریابی ابزاری است که با آن می‌توان نگرش‌های کلی نسبت به هوش مصنوعی را در زمینه‌های عملی و پژوهشی اندازه‌گیری کرده و جنبه‌های مفهومی چنین ابزاری را کشف کرد.

روش

این پژوهش از نظر هدف کاربردی، از لحاظ گردآوری داده توصیفی-پیمایشی و از جهت تحلیل داده، همبستگی (تحلیل عاملی اکتشافی و تأییدی) بود. جامعه آماری این پژوهش شامل دانشجویان دانشگاه محقق اردبیلی و دانشگاه آزاد اسلامی واحد اردبیل بود که در سال تحصیلی ۱۴۰۳-۱۴۰۲ مشغول به تحصیل بودند و در بازه سنی ۲۵ تا ۶۰ سال قرار داشتند. این دو دانشگاه به عنوان مراکز علمی اصلی شهر اردبیل، نمایندگان مناسبی از جامعه شهری بودند. همچنین، این افراد به دلیل آشنایی بیشتر با مفاهیم فناوری و هوش مصنوعی و دسترسی فعال به تکنولوژی، برای اهداف این پژوهش انتخاب شدند؛ بنابراین، روش نمونه‌گیری به صورت هدفمند بود. در تحلیل عاملی تأییدی و مدل‌یابی معادلات ساختاری، حداقل حجم نمونه بر اساس متغیرهای پنهان تعیین می‌شود نه متغیرهای مشاهده‌پذیر. ۲۰ نمونه برای هر متغیر پنهان لازم است؛ ولی با توجه به اینکه در تحلیل معادلات ساختاری همیشه تأکید بر این است که کف نمونه نباید از ۲۰۰ نفر کمتر باشد، برای بال بردن اعتبار بیرونی پژوهش، تعداد ۴۵۰ پرسشنامه توزیع شد. در نهایت ۴۱۴ پرسشنامه قابل تحلیل تشخیص داده شد که شامل ۹۸ آقا و ۳۱۶ خانم بود. این تفاوت در تعداد آقایان و خانم‌ها به دلیل جامعه آماری بالای دانشجویان دختر و همکاری بیشتر این قشر بود.

روال کار بدین صورت بود که ابتدا نسبت به ترجمه فارسی گویه‌های مربوط به مقیاس نگرش کلی به هوش مصنوعی اقدام شد. برای این منظور ابتدا فرم انگلیسی مقیاس توسط پژوهشگران مطالعه حاضر به فارسی ترجمه شد. سپس کیفیت ترجمه، مطابقت فرهنگی گویه‌ها و روایی صوری و محتوایی توسط دو نفر از اساتید رشته مشاوره و همچنین ۵ نفر دانشجوی دکتری رشته مشاوره مورد بازبینی و تأیید قرار گرفت. بعد از آن، برای کشف ایرادات احتمالی در هنگام پاسخگویی، فرم اصلاح‌شده در اختیار ۵ فرد قرار گرفت و از آن‌ها خواسته

شد گویه‌ها را به دقت بخوانند و در مورد قابلیت فهم آن‌ها نظر دهند. پس از ویرایش نهایی، نسخه فارسی توسط یک مترجم به انگلیسی برگردانده شد و به دنبال آن فرم بازترجمه مقیاس توسط یک هیئت علمی دانشگاه و متخصص زبان انگلیسی با فرم انگلیسی اصلی مطابقت داده شد و مورد تأیید قرار گرفت. پس از اخذ مجوزهای لازم، لینک پرسشنامه جهت مطالعه مقدماتی پایایی در اختیار ۴۵ نفر از دانشجویان قرار گرفت. پس از قابل قبول بودن نتایج مطالعه مقدماتی، لینک پرسشنامه‌ها در اختیار آنان و سایر نمونه‌ها قرار گرفت. از جمله ملاک‌های ورود به پژوهش دانشجو بودن و داشتن رضایت آگاهانه جهت مشارکت در پژوهش بود و ملاک‌های خروج نیز، شامل عدم پاسخگویی دقیق به همه سؤالات پرسش‌نامه و یا انصراف از ادامه پاسخگویی توسط آزمودنی‌ها بود. در راستای رعایت اخلاق پژوهشی، به شرکت‌کنندگان در مورد اهداف پژوهش و نیز محرمانه بودن اطلاعات و پاسخ‌ها آگاهی داده شد و پرسش‌نامه‌ها که به صورت خودگزارشی و بدون درج نام و مشخصات هویتی بود در اختیار آن‌ها قرار گرفت. برای تحلیل داده‌ها از شاخص‌های توصیفی، ضریب همبستگی پیرسون، تحلیل عاملی اکتشافی و تأییدی با کمک نرم‌افزارهای SPSS27، LISREL استفاده شد.

ابزارها

۱- نگرش کلی نسبت به هوش مصنوعی^۱

این مقیاس توسط Schepman and Rodway (2020) برای ارزیابی نگرش افراد نسبت به هوش مصنوعی طراحی و مورد استفاده قرار گرفت. این مقیاس از دو خرده‌مقیاس و ۲۰ سؤال تشکیل شده است که شامل نگرش مثبت و نگرش منفی است، نگرش مثبت شامل ۱۲ سؤال و نگرش منفی شامل ۸ سؤال است که آزمودنشوندگان در پنج گزینه به آن‌ها پاسخ می‌دهند که شامل کاملاً مخالفم تا حدودی مخالفم، نظری ندارم تا حدودی موافقم و کاملاً موافقم ندارم. سؤالات نگرش منفی به صورت معکوس نمره‌گذاری می‌شود. نمره جمع کل این مقیاس از طریق جمع نمرات نگرش مثبت و منفی به دست می‌آید. در پژوهش Schepman and Rodway (2020) میزان آلفای کرونباخ برای نگرش مثبت ۰/۸۸ و برای نگرش منفی ۰/۸۲ به دست آمده است.

1. general attitudes towards Artificial Intelligence Scale

۲- مقیاس آمادگی پذیرش تکنولوژی^۱

پرسشنامه آمادگی پذیرش تکنولوژی توسط Chen و همکاران (2014) ساخته شد. پرسشنامه دارای ۱۸ سؤال در ۴ مؤلفه (خوش‌بینی، خلاق‌بودن، عدم سهولت و عدم امنیت) را موردسنجش قرار می‌دهد. این پرسشنامه توسط نوری (۱۳۹۴) اعتباریابی شده است. با استفاده از طیف لیکرت (خیلی کم تا خیلی زیاد) نمره‌گذاری شده است و نمره خرده مقیاس‌های عدم سهولت و عدم امنیت به صورت معکوس نمره‌گذاری می‌شود و برای به دست آوردن نمره کل نمره کل خرده مقیاس‌ها با هم جمع می‌شوند. در پژوهش نوری آلفای کرونباخ کل پرسشنامه ۰/۸۲ و در این پژوهش میزان آلفای کرونباخ ۰/۷۹ به دست آمد.

یافته‌ها

جدول ۱. آماره‌های توصیفی مقیاس نگرش کلی نسبت به هوش مصنوعی

متغیر	حداقل	حداکثر	میانگین	انحراف معیار
نگرش مثبت	۱۲	۲۹	۱۸/۸۴	۴/۵۰
نگرش منفی	۸	۲۴	۱۶/۰۵	۳/۹۰
نمره کل	۲۰	۵۰	۳۴/۷۷	۵/۵۹

همان‌طور که در جدول ۱ مشاهده می‌شود، میانگین و انحراف معیار خرده‌مقیاس نگرش مثبت به هوش مصنوعی به ترتیب برابر با ۱۸/۸۴ و ۴/۵۰، میانگین و انحراف معیار خرده‌مقیاس نگرش منفی نسبت به هوش مصنوعی به ترتیب برابر با ۱۶/۰۵، ۳/۹۰ و میانگین و انحراف معیار نمره کل نگرش کلی نسبت به هوش مصنوعی ۳۴/۷۷ و ۵/۵۹ است.

به منظور کسب اطمینان از روایی مقیاس، روایی صوری، هم‌زمان و سازه برای این مقیاس موردبررسی قرار گرفت. برای تعیین روایی صوری از نسخه‌ی فارسی مقیاس، فرمی تهیه گردید و در اختیار ۲ نفر از استاد‌های گروه مشاوره و ۵ نفر از دانشجویان دکتری تخصصی مشاوره قرار گرفت تا سؤالات آن را از لحاظ شفافیت، روانی و قابلیت درک و فهم بودن و متناسب با شرایط فرهنگی جامعه ما موردبررسی قرار گیرد. از نظر این افراد مقیاس موردنظر به لحاظ صوری ابزاری روا تشخیص داده شد. برای تعیین روایی هم‌زمان مقیاس نگرش کلی

نسبت به هوش مصنوعی از پرسشنامه آمادگی برای پذیرش تکنولوژی استفاده شد. نتایج در جدول ۲ ارائه شده است.

جدول ۲. ضریب همبستگی بین مقیاس نگرش کلی نسبت به هوش مصنوعی و مقیاس آمادگی برای پذیرش تکنولوژی

مقیاس	۱	۲	۳	۴	۵	۶	۷	۸
نگرش مثبت	۱							
نگرش منفی	۰/۳۸**	۱						
نمره کل	۰/۷۲**	۰/۶۰**	۱					
خوش بینی	۰/۴۱**	۰/۳۹**	۰/۳۳**	۱				
خلاق بودن	۰/۴۰**	۰/۳۴**	۰/۲۹**	۰/۶۲**	۱			
عدم سهولت	۰/۲۹**	۰/۳۰۲**	۰/۲۳**	۰/۱۹*	۰/۳۲**	۱		
عدم امنیت	۰/۲۴**	۰/۲۶۵**	۰/۳۵**	۰/۱۹۶*	۰/۱۷۲**	۰/۳۳**	۱	
آمادگی	۰/۲۶**	۰/۲۸**	۰/۲۹**	۰/۶۲**	۰/۸۱**	۰/۶۱**	۰/۵۷**	۱

* $P < 0/05$ ** $p < 0/01$

همان‌طور که در جدول ۲ مشاهده می‌شود، رابطه‌ی نمره کل مقیاس نگرش کلی نسبت به هوش مصنوعی و خرده مقیاس‌های آن با مقیاس آمادگی پذیرش تکنولوژی در سطح آلفای ۰/۰۱ معنادار است. همبستگی مقیاس نگرش کلی نسبت به هوش مصنوعی و خرده مقیاس‌های نگرش مثبت و نگرش منفی با مقیاس آمادگی پذیرش تکنولوژی به ترتیب ۰/۲۶، ۰/۲۸ و ۰/۲۹ به دست آمده است. این همبستگی به معنای روایی هم‌زمان مناسب مقیاس تلقی می‌شود. روایی سازه مقیاس نگرش کلی نسبت به هوش مصنوعی به شیوه‌ی محاسبه ضریب همبستگی مقیاس با خرده مقیاس‌های آن مورد بررسی قرار گرفت. همان‌طور که در جدول ۲ مشاهده می‌شود، ضرایب همبستگی نگرش کلی نسبت به هوش مصنوعی با دو خرده مقیاس نگرش مثبت و نگرش منفی به ترتیب ۰/۷۲ و ۰/۶۰ به دست آمد که در سطح آلفای ۰/۰۱ معنادار هستند. از طرفی ضرایب همبستگی بین خرده مقیاس نگرش مثبت و نگرش منفی ۰/۳۸ است که در مقایسه با ضرایب همبستگی کل مقیاس با خرده مقیاس‌ها، مقدار کمتری را نشان می‌دهد. به لحاظ نظری مقیاس باید با خرده مقیاس همبستگی بالایی داشته باشد و خرده مقیاس‌ها با هم همبستگی کمتری داشته باشند.

جدول ۳. ضرایب آلفای کرونباخ و همبستگی بین نمره‌های مقیاس نگرش نسبت به هوش مصنوعی در دو نوبت اجرا ($n=45$) و کل نمونه ($n=414$)

متغیر	نوبت اول			نوبت دوم			بازآزمایی	آلفای کل
	M	SD	ضریب آلفا	M	SD	ضریب آلفا		
نگرش مثبت	۱۸/۲۳	۳/۹۷	۰/۷۳	۱۸/۸۱	۴/۰۳	۰/۷۹	۰/۶۹	۰/۷۷۲
نگرش منفی	۱۵/۹۸	۳/۲۳	۰/۷۰	۱۶/۱۷	۳/۶۴	۰/۷۴	۰/۷۴	۰/۷۲۳
نگرش کلی	۳۵/۵۶	۵/۶۳	۰/۷۹	۳۴/۹۸	۵/۵۴	۰/۷۶	۰/۷۲	۰/۷۸

همان‌طور که در جدول ۳ مشاهده می‌شود، ضرایب آلفای کرونباخ خرده‌مقیاس نگرش مثبت، نگرش منفی و برای کل مقیاس نگرش کلی نسبت به هوش مصنوعی برای نمونه ۴۵ نفری از آزمودنی‌ها در نوبت اول به ترتیب ۰/۷۳، ۰/۷۰ و ۰/۷۹ و در نوبت دوم به ترتیب ۰/۷۹، ۰/۷۴ و ۰/۷۶ به دست آمد. ضرایب آلفای کرونباخ خرده‌مقیاس نگرش مثبت، نگرش منفی و همچنین کل مقیاس نگرش کلی نسبت به هوش مصنوعی برای کل نمونه ۴۱۴ نفری به ترتیب تقریباً ۰/۷۷ و ۰/۷۲ و ۰/۷۸ به دست آمد. این ضرایب نشانه‌ی همسانی درونی رضایت‌بخش مقیاس نگرش کلی نسبت به هوش مصنوعی است. همچنین، ضرایب همبستگی بین نمره‌های ۴۵ نفر از پاسخ‌دهنده‌ها در دو مرحله اجرای آزمون با فاصله‌ی ۱۶ روز برای سنجش پایایی بازآزمایی مقیاس نگرش کلی نسبت به هوش مصنوعی محاسبه شده است. این ضرایب برای نگرش مثبت ۰/۶۹، برای نگرش منفی ۰/۷۴ و برای کل مقیاس ۰/۷۲ محاسبه شد که در سطح آلفای ۰/۰۱ درصد معنادار است. این نشانه پایایی بازآزمایی رضایت‌بخش مقیاس است.

به‌منظور بررسی ساختار عاملی مقیاس نگرش کلی نسبت به هوش مصنوعی سؤالات پرسشنامه به شیوه تحلیل عاملی اکتشافی به شیوه مؤلفه‌های اصلی مورد تحلیل قرار گرفت. برای اطمینان از مناسب بودن داده‌ها از ضریب KMO و آزمون بارتلت استفاده شد. مقدار KMO همواره بین ۱ و ۰ در نوسان است. در صورتی که مقدار آن کمتر از ۰/۵ باشد، داده‌ها برای تحلیل عاملی مناسب نخواهد بود و اگر مقدار آن بین ۰/۵ تا ۰/۶۹ باشد باید با احتیاط به اجرای تحلیل عاملی پرداخت، اما در صورتی که مقدار آن بزرگ‌تر از ۰/۷ باشد، همبستگی موجود بین داده‌ها برای تحلیل عاملی مناسب است. آزمون بارتلت نیز این فرضیه که ماتریس همبستگی‌های مشاهده‌شده متعلق به جامعه‌ای با متغیرهای ناهمبسته است، را می‌آزماید.

جدول ۴. ضریب KMO و آزمون بارتلت برای ماتریس همبستگی متغیرها

۰/۷۴۵	مقدار KMO
۲۳۲۸/۷۰۶	مقدار بارتلت
۰/۰۰۰	سطح معناداری

همان‌طور که در جدول ۴ مشاهده می‌شود، مقدار KMO برابر است با ۰/۷۴۵ و آزمون بارتلت نیز مورد تأیید قرار گرفت ($P < ۰/۰۱$)؛ بنابراین متغیرهای پژوهش برای تحلیل عاملی مناسب بودند.

پس از اطمینان از کیفیت ماتریس همبستگی سؤالات پرسشنامه و همچنین کفایت نمونه‌برداری محتوایی، سؤالات مقیاس نگرش کلی نسبت به هوش مصنوعی به شیوه واریماکس چرخش داده شد. نتایج حاصل از تحلیل عاملی اکتشافی نشان داد، دو عامل، در مجموع ۶۵/۶۹ درصد از واریانس کل را تبیین می‌کنند. در این تحلیل مقدار بار عاملی مورد قبول ۰/۵ تعیین گردید. به جز سؤال ۵ بقیه سؤالات مقیاس از بار عاملی ۰/۵۰ به بالا برخوردار است. در جدول ۵ محتوا و بار عاملی عوامل مقیاس نگرش کلی نسبت به هوش مصنوعی آمده است.

جدول ۵. بار عاملی سؤالات مقیاس نگرش کلی نسبت به هوش مصنوعی

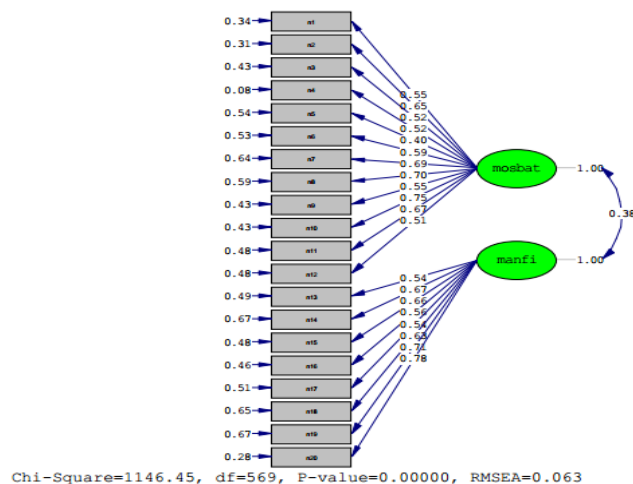
عوامل		گویه‌ها
۲	۱	
۰/۵۵		هوش مصنوعی کاربردهای مفید زیادی دارد.
۰/۶۵		آنچه هوش مصنوعی می‌تواند انجام دهد مرا تحت تأثیر قرار داده است.
۰/۵۲		هوش مصنوعی می‌تواند تأثیر مثبتی بر روی رفاه و آسایش مردم داشته باشد.
۰/۵۲		هوش مصنوعی جالب و هیجان‌انگیز است.
۰/۴۰		هوش مصنوعی می‌تواند فرصت‌های اقتصادی جدیدی برای کشور فراهم نماید.
۰/۵۹		سیستم‌های مبتنی بر هوش مصنوعی بهتر از انسان‌ها می‌توانند عمل کنند.
۰/۶۹		بخش اعظم جامعه از آینده‌ای سرشار از هوش مصنوعی بهره‌مند خواهد بود.
۰/۷۰		من مایل به استفاده از سیستم‌های مبتنی بر هوش مصنوعی در زندگی روزمره خود هستم.
۰/۵۵		برای معاملات روزمره، تعامل با یک سیستم مبتنی بر هوش مصنوعی را ترجیح می‌دهم.
۰/۷۵		برای انجام بسیاری از کارهای روزمره، یک دستیار هوش مصنوعی بهتر از یک کارمند است.
۰/۶۷		من مایلم هوش مصنوعی را در شغلم بکار بگیرم.
۰/۵۱		سیستم‌های هوش مصنوعی می‌توانند کمک کنند تا مردم حس خوشحالی بیشتری داشته باشند.

عوامل		گویه‌ها
۲	۱	
۰/۵۴		هوش مصنوعی برای جاسوسی از مردم استفاده می‌شود.
۰/۶۷		هوش مصنوعی ممکن است کنترل مردم را به دست بگیرد.
۰/۶۶		من وقتی به استفاده‌هایی که در آینده از هوش مصنوعی می‌شود فکر می‌کنم، از ناراحتی و ترس به خودم می‌لرزم
۰/۵۶		من فکر می‌کنم هوش مصنوعی خطرناک است.
۰/۵۴		سازمان‌ها از هوش مصنوعی استفاده‌ی غیراخلاقی می‌کنند
۰/۶۸		سیستم‌های هوش مصنوعی به نظر من خطاهای زیادی را مرتکب می‌شوند.
۰/۷۱		اگر کاربرد هوش مصنوعی بیشتر و بیشتر شود، افرادی مثل من صدمه خواهند دید.
۰/۷۸		به نظر من هوش مصنوعی گمراه‌کننده و اشتباه/فاسد است

طبق جدول ۵ بار عاملی سؤالات ۱ تا ۱۲ بیشترین بار عاملی را بر روی عامل اول داشتند. بار عاملی سؤالات ۱۳ تا ۲۰ نیز بیشترین بار عاملی را بر روی عامل دوم داشتند.

در ادامه به منظور بررسی روایی سازه، از روش تحلیل عاملی تأییدی مرتبه اول استفاده شد. در مدل‌های تحلیل عاملی تأییدی مرتبه اول، نمرات هر مورد مطالعه در یک متغیر، درواقع منعکس‌کننده وضعیت آن مورد در یک عامل زیربنایی‌تر است که به دلیل پنهان بودنش امکان اندازه‌گیری مستقیم آن وجود ندارد؛ اما این عامل زیربنایی و پنهان خود از ابعاد عامل پنهان دیگری محسوب نمی‌شود و درواقع تنها یک لایه از متغیر یا متغیرهای پنهان وجود دارد. در یک مدل عاملی مرتبه اول، سه نوع متغیر وجود دارد. این متغیرها شامل متغیر پنهان بیرونی، متغیر مشاهده‌شده بیرونی و متغیر خطا هستند. در مدل‌های عاملی مرتبه اول سه نوع پارامتر آزاد وجود دارد که محقق اغلب مایل به برآورد آن‌هاست. این پارامترها شامل پارامتر فی (یعنی واریانس‌ها و کوواریانس‌های عامل‌ها یا متغیرهای پنهان بیرونی)، پارامتر لامدا (کوواریانس یا ضریب همبستگی بین هر متغیر آشکار با متغیر پنهان) و پارامتر تتا دلتا (واریانس‌ها و کوواریانس‌های متغیرهای خطای مرتبط با متغیرهای آشکار و اندازه‌گیری) هستند (قاسمی، ۱۳۹۲). برای برآورد مدل از شاخص‌های مجذور خی دو (X^2)، شاخص نسبت مجذور خی دو به درجه آزادی (X^2/df)، شاخص نیکویی برازش (GFI)، شاخص نیکویی برازش انطباقی (AGFI)، شاخص برازش مقایسه‌ای (CFI)، شاخص برازش هنجار شده (NFI)، شاخص نیکویی برازش هنجار نشده (NNFI)، خطای ریشه مجذور میانگین تقریب (RMSEA) و باقیمانده مجذور میانگین (RMR) استفاده شد.

شکل ۱. مدل استاندارد تحلیل عاملی تأییدی مقیاس نگرش کلی نسبت به هوش مصنوعی



جدول ۶. مقادیر شاخص‌های برازش الگوی تحلیل عاملی تأییدی مقیاس نگرش کلی نسبت به هوش مصنوعی

RMR	RMSEA	NNFI	NFI	AGFI	GFI	CFI	X^2/df	X^2
۰/۰۴۹	۰/۰۶۳	۰/۸۶	۰/۸۶	۰/۸۹	۰/۹۱	۰/۸۸	۵۶۹	۱۱۴۶/۴۵

بحث و نتیجه‌گیری

پژوهش حاضر با هدف تعیین ویژگی‌های روان‌سنجی مقیاس نگرش کلی نسبت به هوش مصنوعی در جامعه ایران انجام گرفت. یافته‌های پژوهش ویژگی‌های روان‌سنجی مقیاس نگرش کلی نسبت به هوش مصنوعی در نمونه‌های ایرانی را تأیید کرد. نتایج مربوط به تحلیل عاملی تأییدی با روش تحلیل عاملی وجود دو خرده مقیاس یعنی نگرش مثبت و نگرش منفی نسبت به هوش مصنوعی را تأیید کرد. این یافته با نتایج پژوهش‌های انجام گرفته Schepman and Rodway (2020)، در زمینه‌ی ساختار عاملی مقیاس نگرش کلی نسبت به هوش مصنوعی همخوانی دارد.

جهت تعیین روایی صوری از نسخه‌ی فارسی مقیاس، فرمی تهیه گردید و در اختیار ۲ نفر از استادان رشته مشاوره دانشگاه و ۵ نفر از دانشجویان دکتری تخصصی مشاوره قرار گرفت تا سؤالات آن را از نظر شفافیت، روانی و قابل درک و فهم بودن و متناسب با شرایط

فرهنگی جامعه‌ی ما مورد بررسی قرار گیرد. از نظر این افراد مقیاس مذکور به لحاظ صوری ابزاری روا تشخیص داده شد، روایی هم‌زمان مقیاس از طریق اجرای هم‌زمان با مقیاس آمادگی برای پذیرش تکنولوژی محاسبه شد. ضرایب همبستگی میانگین نمره‌های آزمودنی‌ها در خرده مقیاس‌های نگرش مثبت و منفی نسبت به هوش مصنوعی معنادار بود. بر اساس این نتایج، مقیاس نگرش کلی نسبت به هوش مصنوعی از روایی کافی برخوردار است. روایی سازه مقیاس نیز از طریق تحلیل عاملی تأییدی مورد بررسی قرار گرفت؛ و نشان داد تمامی شاخص‌های برازندگی مدل اندازه‌گیری مقیاس نگرش کلی نسبت به هوش مصنوعی از برازش مناسب و در نتیجه، روایی سازه‌ی قابل قبولی برخوردار است. این یافته با نتایج پژوهش انجام شده توسط Schepman and Rodway (2020) در زمینه‌ی روایی مقیاس نگرش کلی نسبت به هوش مصنوعی همخوانی دارد.

همسانی درونی سؤالات مقیاس نگرش کلی نسبت به هوش مصنوعی برحسب ضرایب آلفای کرونباخ محاسبه شد و مورد تأیید قرار گرفت. پایایی بازآزمایی مقیاس برحسب محاسبه‌ی ضرایب همبستگی بین نمره‌های تعدادی از شرکت‌کنندگان در دو مرحله به فاصله‌ی ۱۶ روز برای خرده مقیاس نگرش مثبت و منفی معنادار به دست آمد. این ضرایب نشانه‌ی پایایی بازآزمایی رضایت‌بخش مقیاس است. این یافته‌ها با نتایج پژوهش‌های انجام شده توسط Schepman and Rodway (2020)، در زمینه‌ی پایایی مقیاس نگرش کلی نسبت به هوش مصنوعی همخوانی دارد.

به منظور بررسی روایی سازه، از روش مدل معادلات ساختاری از نوع تحلیل عاملی تأییدی استفاده شد. برای برآورد مدل از شاخص‌های مجذور خی دو (X^2)، شاخص نسبت مجذور خی دو به درجه آزادی (X^2/df)، شاخص نیکویی برازش (GFI)، شاخص نیکویی برازش انطباقی (AGFI)، شاخص برازش مقایسه‌ای (CFI)، شاخص برازش هنجار شده (NFI)، شاخص نیکویی برازش هنجار نشده (NNFI)، خطای ریشه مجذور میانگین تقریب (RMSEA) و باقیمانده مجذور میانگین (RMR) استفاده شد. نتایج نشان داد مقیاس نگرش کلی نسبت به هوش مصنوعی از روایی سازه‌ای مناسبی برخوردار است.

به طور کلی این پژوهش، ویژگی‌های روان‌سنجی مقیاس نگرش کلی نسبت به هوش مصنوعی را در نمونه‌ای از ایرانی‌ها تأیید کرد. با توجه به یافته‌های این پژوهش می‌توان نتیجه گرفت که همسو با پژوهش Schepman and Rodway (2020) گویه‌های این مقیاس

مناسب انتخاب شده‌اند و این مقیاس بدون تغییر و حذف احتمالی برخی گویه‌ها ساختار خود را حفظ کرد؛ بنابراین، این مقیاس که نگرش مثبت و منفی افراد را نسبت به هوش مصنوعی به‌خوبی مورد ارزیابی قرار می‌دهد، در جامعه‌ی ایران روایی و پایایی مناسبی دارد و می‌تواند در موقعیت‌های اجتماعی، پژوهشی و ... مورد استفاده قرار گیرد و زمینه‌ها برای پژوهش‌های بیشتر در مورد نگرش افراد در مورد هوش مصنوعی را فراهم کند. هر پژوهشی برخی محدودیت‌های دارد و این پژوهش نیز عاری از این محدودیت‌ها نیست. ازجمله آن‌که این پژوهش در نمونه دانشجویان دانشگاه محقق اردبیلی و دانشگاه آزاد اسلامی واحد اردبیل اجرا شده است. با توجه به این‌که این افراد ممکن است آشنایی بیشتری با مفاهیم هوش مصنوعی و فناوری داشته باشند، تعمیم نتایج به کل کشور با محدودیت همراه است. علاوه بر این، تفاوت در ترکیب جنسیتی نمونه ممکن است نتایج را تحت تأثیر قرار دهد؛ بنابراین، پیشنهاد می‌شود این ابزار در نمونه‌های دیگر، گروه‌های سنی مختلف، و با روش‌های ارزیابی متنوع‌تری مورد بررسی قرار گیرد تا قابلیت تعمیم و اعتبار آن افزایش یابد.

مشارکت نویسندگان

در این پژوهش، تمامی نویسندگان در فرآیند طراحی مطالعه، جمع‌آوری داده‌ها، تحلیل آماری، نگارش پیش‌نویس اولیه و بازنگری نهایی مقاله به‌طور مؤثر مشارکت داشته‌اند.

تعارض منافع

یافته‌های این مطالعه هیچ‌گونه تضاد با منافع شخصی یا سازمانی ندارد.

سپاسگزاری

از کلیه‌ی شرکت‌کنندگان در پژوهش که با شرکت خود در پژوهش، به روند اجرای طرح کمک کردند، سپاسگزاری می‌شود.

منابع

- حاتمی نژاد، محمد، صدری دمیرچی، اسماعیل، اخوی ثمرین، زهرا، جعفری مرادلو، مهدی و نوروزی همایون، محمدرضا. (۱۴۰۳). مقایسه اثربخشی روایت‌درمانی و درمان هیجان‌مدار بر طرح‌واره‌های ناسازگار اولیه، تکانش‌گری و افسردگی در دانش‌آموزان دارای افکار خودکشی. *مطالعات روان‌شناختی*، ۲۰(۳)، ۷-۲۳. doi: 10.22051/psy.2024.47061.2955
- نوروزی همایون، محمدرضا، حاتمی نژاد، محمد و صدری دمیرچی، اسماعیل. (۱۴۰۲). اثربخشی درمان گروهی سایکودراما و بازی‌درمانی شناختی رفتاری بر روی کارکردهای اجرایی (حافظه فعال، بازداری پاسخ، انعطاف‌پذیری شناختی و خودتنظیمی هیجانی) در دانش‌آموزان پسر دارای اختلال اضطراب اجتماعی. *عصب روان‌شناسی*، ۹(۳۵)، ۱-۱۷. doi: 10.30473/clpsy.2024.68482.1710
- نوروزی همایون، محمدرضا، حاتمی نژاد، محمد، صدری دمیرچی، اسماعیل و صادقی، مسعود. (۱۴۰۳). روابط ساختاری ناگویی هیجانی و کمال‌گرایی با نقش میانجی حساسیت اضطرابی در پیش‌بینی اختلال وسواسی جبری دانش‌آموزان. *روان‌شناسی بالینی*، ۱۶(۳)، ۶۵-۷۴. doi: 10.22075/jcp.2024.34232.2897
- نوروزی همایون، محمدرضا، صادقی، مسعود، صدری دمیرچی، اسماعیل و حاتمی نژاد، محمد. (۱۴۰۳). رابطه حس انسجام و حمایت اجتماعی ادراک‌شده با تاب‌آوری: نقش میانجی راهبردهای انطباقی تنظیم شناختی هیجانی سالمندان. *روان‌شناسی پیری*، ۱۰(۴)، ۴۲۷-۴۵۵. doi: 10.22126/jap.2025.11026.1801

References

- Alimour, S. A., Alnono, E., Aljasmi, S., El Farran, H., Alqawasmi, A. A., Alrabeei, M. M., & Aburayya, A. (2024). The quality traits of artificial intelligence operations in predicting mental healthcare professionals' perceptions: A case study in the psychotherapy division. *Journal of Autonomous Intelligence*, 7(4). <https://doi.org/https://jai.front-sci.com/index.php/jai/article/view/1438>
- Bonnefon, J. F., Rahwan, I., & Shariff, A. (2023). The moral psychology of Artificial Intelligence. *Annual Review of Psychology*, 75. <https://doi.org/https://doi.org/10.1146/annurev-psych-030123-113559>
- Darko, A., Chan, A. P., Adabre, M. A., Edwards, D. J., Hosseini, M. R., & Ameyaw, E. E. (2020). Artificial intelligence in the AEC industry: Scientometric analysis and visualization of research activities. *Automation in Construction*, 112, 103081. <https://doi.org/http://dx.doi.org/10.1016%2Fj.autcon.2020.103081>
- Espejo, G., Reiner, W., & Wenzinger, M. (2023). Exploring the Role of Artificial Intelligence in Mental Healthcare: Progress, Pitfalls, and Promises. *Cureus*, 15(9). <https://doi.org/https://doi.org/10.7759/cureus.44748>
- Eyüp, B., & Kayhan, S. (2023). Pre-Service Turkish Language Teachers' Anxiety and Attitudes Toward Artificial Intelligence. *International Journal of Education*

- and Literacy Studies, 11(4), 43-56.
<https://doi.org/https://doi.org/10.7575/aiac.ijels.v.11n.4p.43>
- Federspiel, F., Mitchell, R., Asokan, A., Umana, C., & McCoy, D. (2023). Threats by artificial intelligence to human health and human existence. *BMJ Global Health*, 8(5). <https://doi.org/https://doi.org/10.1136/bmjgh-2022-010435>
- Gansser, O. A., & Reich, C. S. (2021). A new acceptance model for artificial intelligence with extensions to UTAUT2: An empirical study in three segments of application. *Technology in Society*, 65, 101535. <https://doi.org/https://doi.org/10.1016/j.techsoc.2021.101535>
- Grassini, S. (2023). Development and validation of the AI attitude scale (AIAS-4): a brief measure of the general attitude toward artificial intelligence. *Frontiers in Psychology*, 14. <https://doi.org/https://doi.org/10.3389/fpsyg.2023.1191628>
- Hadlington, L., Binder, J., Gardner, S., Karanika-Murray, M., & Knight, S. (2023). The use of artificial intelligence in a military context: development of the attitudes toward AI in defense (AAID) scale. *Frontiers in Psychology*, 14, 1164810. <https://doi.org/https://doi.org/10.3389/fpsyg.2023.1164810>
- Hatami Nejad, M., Sadeghi, M., Sadri Damirchi, E., & Noroozi Homayoon, M. (2025). Examining the Influence of Alexithymia, Gender, and Age on Drug Use among Iranian Students: the Mediating Role of Emotion Regulation Difficulties. *Scientific Reports*, 15(1), 3650. <https://doi.org/10.1038/s41598-025-87394-w> [In Persian]
- Hatami Nejad, M., Sadri Damirchi, E., Akhavi Samarein, Z., Jafari Moradlo, M., & Noroozi Homayoon, M. (2024). Comparing the Effectiveness of Narrative Therapy and Emotion Focused Therapy on Primary Maladaptive Schemas, Impulsivity and Depression in Students with Suicidal Thoughts. *Journal of Psychological Studies*, 20(3), 7-23. <https://doi.org/10.22051/psy.2024.47061.2955> [In Persian]
- Hatami Nejad, M., Sadri Damirchi, E., Sadeghi, M., & Noroozi Homayoon, M. (2025). Developing a model of experiential avoidance based on childhood trauma and victimization, mediated by insecure attachment styles. *European Journal of Trauma & Dissociation*, 9(1), 100513. <https://doi.org/https://doi.org/10.1016/j.ejtd.2025.100513>
- Kaya, F., Aydin, F., Schepman, A., Rodway, P., Yetişensoy, O., & Demir Kaya, M. (2024). The roles of personality traits, AI anxiety, and demographic factors in attitudes toward artificial intelligence. *International Journal of Human-Computer Interaction*, 40(2), 497-514. <https://doi.org/https://doi.org/10.1080/10447318.2022.2151730>
- Kieslich, K., Marco, L., & Došenović, P. (2024). Ever Heard of Ethical AI? Investigating the Salience of Ethical AI Issues among the German Population. *International Journal of Human-Computer Interaction*, 40(11), 2986-2999. <https://doi.org/10.1080/10447318.2023.2178612>
- Kumar, S., & Choudhury, S. (2022). Humans, super humans, and super humanoids: debating Stephen Hawking's doomsday AI forecast. *AI Ethics*, 1-10. <https://doi.org/https://dx.doi.org/10.2139/ssrn.4203612>
- Lee, J., Lee, D., & Lee, J. G. (2024). Influence of rapport and social presence with an AI psychotherapy chatbot on users' self-disclosure. *International Journal of Human-Computer Interaction*, 40(7), 1620-1631. <https://doi.org/https://dx.doi.org/10.2139/ssrn.4063508>
- Letter, P. G. A. E. A. O. (2023). Pause Giant AI Experiments: An Open Letter. In. Noroozi Homayoon, M., Akhavi Samarein, Z., Sadeghi, M., Hatami Nejad, M., & Jafari Moradlo, m. (2025). Comparing the Efficacy of Emotion-Focused Therapy and Transcranial Direct Current Stimulation On Impulsivity, Emotion

- Regulation, And Suicidal Ideation In Young People With Borderline Personality Disorder. *Journal of Research in Psychopathology*, 6(1), 33-42. <https://doi.org/10.22098/jrp.2024.14488.1221>
- Noroozi Homayoon, M., Hatami Nejad, M., & Sadri Damichi, E. (2024). The Effectiveness of Psychodrama Group Therapy and Cognitive Behavioral Play Therapy on Executive Functions (Working Memory, Response Inhibition, Cognitive Flexibility and Emotional Self Regulation) in Male Students with Social Anxiety Disorder. *Neuropsychology*, 9(35), 1-17. <https://doi.org/10.30473/clpsy.2024.68482.1710> [In Persian]
- Noroozi Homayoon, M., Hatami Nejad, M., Sadri Damirchi, E., & Sadeghi, M. (2024). Structural Relationships between Alexithymia and Perfectionism with the Mediating Role of Anxiety Sensitivity in Predicting Obsessive-Compulsive Disorder in Students. *Journal of Clinical Psychology*, 16(3), 47-65. <https://doi.org/10.22075/jcp.2024.34232.2897> [In Persian]
- Noroozi Homayoon, M., Sadeghi, M., Sadri Damirchi, Esmail, & Hatami Nejad, M. (2024). The Relationship between Sense of Coherence and Perceived Social Support with Resilience: The Mediating Role of Adaptive Cognitive-Emotional Regulation Strategies in Older Adults. *Aging Psychology*, 10(4), 427-405. <https://doi.org/10.22126/jap.2025.11026.1801> [In Persian]
- Noroozi Homayoon, M., Sadri Damirchi, E., Sadeghi, M., & Hatami Nejad, M. (2025). Comparing the Effectiveness of Cognitive-Emotional Regulation Training and Transcranial Direct Current Stimulation on Difficulty in Regulating Emotions, Rumination and Quality of Life in Divorced Women. *Women's Health Bulletin*, 12(1), 21-29. <https://doi.org/10.30476/whb.2024.103315.1300>
- O'Shaughnessy, M. R., Schiff, D. S., Varshney, L. R., Rozell, C. J., & Davenport, M. A. (2023). What governs attitudes toward artificial intelligence adoption and governance? *Science and Public Policy*, 50(2), 161-176. <https://doi.org/https://doi.org/10.1093/scipol/scac056>
- Olawade, D. B., Wada, O. Z., Odetayo, A., David-Olawade, A. C., Asaolu, F., & Eberhardt, J. (2024). Enhancing mental health with Artificial Intelligence: Current trends and future prospects. *Journal of Medicine, Surgery, and Public Health*, 100099. <https://doi.org/http://dx.doi.org/10.1016/j.glmedi.2024.100099>
- Omarov, B., Narynov, S., & Zhumanov, Z. (2023). Artificial Intelligence-Enabled Chatbots in Mental Health: A Systematic Review. *Computers, Materials & Continua*, 74(3). <https://doi.org/https://doi.org/10.32604/cmc.2023.034655>
- Parvathy, V., & Mishra, D. (2023). Impact of the Role of Artificial Intelligence on Mental Health. International Conference on Smart Trends in Computing and Communications, Singapore.
- Robertson, C., Woods, A., Bergstrand, K., Findley, J., Balser, C., & Slepian, M. J. (2023). Diverse patients' attitudes towards Artificial Intelligence (AI) in diagnosis. *PLOS Digital Health*, 2(5), e0000237. <https://doi.org/10.1371/journal.pdig.0000237>
- Said, N., Potinteu, A. E., Brich, I., Buder, J., Schumm, H., & Huff, M. (2023). An artificial intelligence perspective: How knowledge and confidence shape risk and benefit perception. *Computers in Human Behavior*, 107855. <https://doi.org/http://dx.doi.org/10.1016/j.chb.2023.107855>
- Samuel, J. (2023). Two keys for surviving the inevitable AI invasion. In.
- Schepman, A., & Rodway, P. (2020). Initial validation of the general attitudes towards Artificial Intelligence Scale. *Computers in Human Behavior Reports*, 1, 100014. <https://doi.org/https://doi.org/10.1016/j.chbr.2020.100014>

- Schepman, A., & Rodway, P. (2023). The General Attitudes towards Artificial Intelligence Scale (GAAIS): Confirmatory validation and associations with personality, corporate distrust, and general trust. *International Journal of Human-Computer Interaction*, 39(13), 2724-2741. <https://doi.org/http://dx.doi.org/10.1080/10447318.2022.2085400>
- Stein, J. P., Messingschlager, T., Gnambs, T., Hutmacher, F., & Appel, M. (2024). Attitudes towards AI: measurement and associations with personality. *Scientific Reports*, 14(1), 2909. <https://doi.org/https://doi.org/10.1038/s41598-024-53335-2>
- Wang, B., Rau, P. L. P., & Yuan, T. (2023). Measuring user competence in using artificial intelligence: validity and reliability of artificial intelligence literacy scale. *Behaviour & Information Technology*, 42(9), 1324-1337. <https://doi.org/http://dx.doi.org/10.1080/0144929X.2022.2072768>
- Wang, H., Fu, T., Du, Y., Gao, W., Huang, K., Liu, Z., & Zitnik, M. (2023). Scientific discovery in the age of artificial intelligence. *Nature*, 620(7972), 47-60. <https://doi.org/https://doi.org/10.1038/s41586-023-06221-2>
- Weidinger, L., Uesato, J., Rauh, M., Griffin, C., Huang, P. S., Mellor, J., & Gabriel, I. (2022). Taxonomy of risks posed by language models. Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency.
- Yaacob, H., Hossain, F., Shari, S., Khare, S. K., Ooi, C. P., & Acharya, U. R. (2023). Application of artificial intelligence techniques for brain-computer interface in mental fatigue detection: a systematic review (2011-2022). *IEEE Access*. <https://doi.org/http://dx.doi.org/10.1109/ACCESS.2023.3296382>