

The Impact of Anthropomorphism of Collaborative Robots on Their Acceptance in the Automotive Industry: The Mediating Role of Perceived Safety

HamidReza Talaie*¹

1. Assistant Professor, Department of Industrial Management, Faculty of Administrative Sciences and Economics, Arak University, Arak, Iran

*. Corresponding Author: h-Talaie@araku.ac.ir

Received: 31 July 2025

Revised: 15 March 2026

Accepted: 23 March 2026

Abstract

This study examines the anthropomorphism of collaborative robots and its impact on perceived competence, perceived safety, perceived threat, and, ultimately, employee acceptance of these robots in the automotive industry. The research is applied, descriptive-correlational, and cross-sectional. The statistical population comprises employees in the automotive industry, from whom 384 individuals were randomly selected. The measurement model's reliability and validity were confirmed using Cronbach's alpha, composite reliability, AVE, cross-loadings, and the HTMT criterion. To test the mediating effects, mediation analysis was conducted, and the VAF index was calculated. Statistical analysis was conducted using structural equation modeling (SEM). The findings revealed that the anthropomorphic design of collaborative robots positively influences perceived safety and perceived competence, while it may also increase perceived threat. Nonetheless, perceived safety plays a significant role in the acceptance of collaborative robots and can amplify the impact of anthropomorphic design. The acceptance of collaborative robots remains a relatively new area of research and has not yet been thoroughly explored, particularly in the Iranian context. Results suggest that improving acceptance of collaborative robots requires reducing perceived threat and enhancing safety in their design.

Keywords: Anthropomorphism of collaborative robots, Acceptance of collaborative robots, Perceived competence, Perceived threat, Perceived safety.

Citation: Talaie, H. R. (2026). The Impact of Anthropomorphism of Collaborative Robots on Their Acceptance in the Automotive Industry: The Mediating Role of Perceived Safety. *Journal of Technology Development Management*, 13(1), 179–208.

<https://doi.org/10.22104/jtdm.2026.7788.3440>

تأثیر انسان‌نمایی ربات‌های همکار بر پذیرش آن‌ها در صنعت خودروسازی: نقش میانجی ایمنی

درک‌شده

حمیدرضا طلایی*

۱. استادیار گروه مدیریت صنعتی، دانشکده علوم اداری و اقتصاد، دانشگاه اراک، اراک، ایران

*. نویسنده مسئول: h-Talaie@araku.ac.ir

پذیرش: ۳ فروردین ۱۴۰۵

بازنگری: ۲۴ اسفند ۱۴۰۴

دریافت: ۹ مرداد ۱۴۰۴

چکیده

پژوهش حاضر بر انسان‌نما بودن ربات‌های همکار و تأثیر آن بر متغیرهای شایستگی درک‌شده، ایمنی درک‌شده، تهدید درک‌شده و در نهایت پذیرش این ربات‌ها توسط کارکنان در صنعت خودروسازی متمرکز است. پژوهش حاضر از نوع کاربردی، توصیفی-همبستگی و مقطعی است. جامعه آماری پژوهش کارمندان صنعت خودروسازی کشور هستند و ۳۸۴ نفر به عنوان نمونه و با نمونه‌گیری تصادفی در نظر شدند. روایی و پایایی با استفاده از شاخص‌های آلفای کرونباخ، پایایی ترکیبی، AVE، بارهای عرضی و آزمون HTMT تأیید شد. برای آزمون نقش میانجی، تحلیل میانجی‌گری براساس شاخص VAF محاسبه گردید. در تجزیه و تحلیل داده‌های آماری، از رویکرد مدلسازی معادلات ساختاری استفاده شده است. نتایج پژوهش نشان داد که طراحی انسان‌نمای ربات‌های همکار تأثیر مثبتی بر افزایش ایمنی درک‌شده و شایستگی درک‌شده دارد، اما در عین حال می‌تواند تهدید درک‌شده را نیز افزایش دهد. با این حال، ایمنی درک‌شده نقش مهمی در پذیرش ربات‌های همکار ایفا می‌کند و می‌تواند تأثیر طراحی انسان‌نما را تقویت کند. پذیرش ربات‌های همکار یکی از حوزه‌های جدید در ادبیات است که تاکنون ابعاد آن به طور عمیق به خصوص در داخل کشور بررسی نشده است. یافته‌ها نشان می‌دهند که برای ارتقای پذیرش ربات‌های همکار، تمرکز بر کاهش احساس تهدید و افزایش ایمنی در طراحی ضروری است.

کلمات کلیدی: انسان‌نما بودن ربات همکار، پذیرش ربات‌های همکار، شایستگی درک‌شده، تهدید درک‌شده، ایمنی درک‌شده.

مقدمه

صنعت ۵.۰ به‌عنوان یک مدل صنعتی جدید، تمرکز را از دیجیتالی‌سازی به توجه به سلامت جسمی و روانی انسان‌ها تغییر داده است (لیائو و همکاران^۱، ۲۰۲۳) و استفاده از ربات‌های همکار (کو-ربات‌ها) که قابلیت تعامل مستقیم با انسان‌ها را دارند، به یکی از ارکان این صنعت تبدیل شده است (کلبه و فیس^۲، ۲۰۲۵). کو-ربات‌ها به دلیل ویژگی‌هایی نظیر ایمنی، هوشمندی، هزینه پایین، و انعطاف‌پذیری در صنایع مختلف و مراقبت‌های بهداشتی کاربرد گسترده‌ای پیدا کرده‌اند و پیش‌بینی می‌شود که فروش آن‌ها در سال‌های آینده به شدت افزایش یابد (لیائو و همکاران^۳، ۲۰۲۴؛ کراس و همکاران^۴، ۲۰۲۴). با این حال، با وجود نگرش مثبت مدیران، برخی کارکنان نگرانی‌هایی نظیر ترس و طرد کردن این فناوری‌ها را دارند (لکمن و همکاران^۵، ۲۰۲۳). درک نگرش کارکنان نسبت به ربات‌ها برای پذیرش و کارایی بهتر این فناوری‌ها در سیستم‌های تولید ضروری است (پراسیدا و افساری^۶، ۲۰۲۲). مطالعات قبلی از مدل پذیرش فناوری برای تحلیل نگرش‌ها به ربات‌ها استفاده کرده‌اند، اما این مدل‌ها برای ارزیابی پذیرش کو-ربات‌ها مناسب نیستند، زیرا این ربات‌ها بیشتر به‌عنوان تجهیزات هوشمند برای جایگزینی نیروی انسانی در کارهای فیزیکی طراحی شده‌اند. بنابراین، نیاز به مدل‌های جدید برای ارزیابی پذیرش کو-ربات‌ها وجود دارد (بامگارتن و همکاران^۷، ۲۰۲۲). در این زمینه، انسان‌نما بودن ربات‌ها به‌عنوان یک عامل کلیدی در تعامل انسان-ربات شناخته می‌شود (کراس و همکاران، ۲۰۲۴). طراحی انسان‌نما به ربات‌ها اجازه می‌دهد تا ویژگی‌های مشابه انسان مانند اندام‌های انعطاف‌پذیر و حالات صورت را داشته باشند (لوتین و همکاران^۸، ۲۰۲۴)، که این ویژگی‌ها به بهبود تعاملات اجتماعی و راحت‌تر شدن پذیرش این ربات‌ها در محیط‌های کاری کمک می‌کند (لیائو و همکاران، ۲۰۲۳؛ ۲۰۲۴).

هدف تعامل انسانی-ربات صنعتی^۹ (IHRI) ادغام مزایای ربات‌ها و کارکنان است و شامل بهبود شرایط کاری اپراتورها و بهبود عملکرد تولید برای جبران ضعف‌های هر دو طرف است (گوالتیری و همکاران^{۱۰}، ۲۰۲۱). پژوهش حاضر بر پایه نظریه‌های کلیدی در حوزه تعامل انسان-ربات بنا شده است، که شامل پدیده دره غیر طبیعی^{۱۱} و

¹ Liao et al.

² Klebbe and Frieese

³ Liao et al.

⁴ Kraus et al.

⁵ Leichtmann et al.

⁶ Prassida and Asfari

⁷ Baumgartner et al.

⁸ Lutin et al.

⁹ Industrial human-robot interaction

¹⁰ Gualtieri et al.

¹¹ Uncanny valley

نظریه تهدید بین گروهی می‌شود. وقتی درجه انسان‌نمایی از یک نقطه بحرانی عبور کند، احساسات مردم نسبت به ربات‌ها از مثبت به منفی تغییر می‌کند و به قعر دره غیر طبیعی می‌رسد. سپس، با بهبود بیشتر درجه انسان‌نمایی، نگرش‌ها بهتر می‌شوند (لکمن و همکاران، ۲۰۲۳). علاوه بر اثر غیر طبیعی ناشی از شباهت ظاهری به انسان، محققان دریافتند که اثر غیر طبیعی ناشی از ربات‌های انسان‌نما ممکن است از درک احساسات و محیط آن‌ها ناشی شود. اپل و همکاران^۱ (۲۰۲۰) اشاره کردند که شایستگی درک شده ربات‌ها بیشتر از توانایی برنامه‌ریزی و خودکنترلی می‌تواند حس عجیبی ایجاد کند. ترس مردم از ربات‌ها نه تنها به ظاهر انسان‌نمای آن‌ها بستگی دارد، بلکه به شدت با درک احساسی، ادراک فکری، توانایی استدلال و سایر توانایی‌های مشابه انسان مرتبط است. به همین دلیل، کو-ربات‌ها که ترکیبی از هوش مصنوعی، حسگرهای با دقت بالا و طراحی حرکت روان را دارند، تمامی قابلیت‌های فوق را برای ایجاد احساسات منفی در کارکنان نسبت به ربات‌های همکار فراهم می‌آورند. با توجه به کاربرد گسترده کو-ربات‌ها در کارخانه‌ها، می‌توان آن‌ها را به‌عنوان یک نیروی کار دیگر در نظر گرفت که به طور اجتناب‌ناپذیری منجر به نگرانی کارکنان کارخانه در مورد فرصت‌های شغلی و درک تهدید از طرف کو-ربات‌ها خواهد شد (لیو و همکاران^۲، ۲۰۲۴). این پدیده ترس به خوبی توسط نظریه تهدید بین گروهی^۳ توضیح داده می‌شود.

تهدید بین گروهی^۴ به این معنا است که رفتارها، باورها و ویژگی‌های مختلف یک گروه توسط اهداف، توسعه و بقا گروه دیگر تهدید می‌شود. تهدید بین گروهی می‌تواند به تهدید واقعی و تهدید هویتی تقسیم شود. تهدید واقعی به تهدید منابع مادی، ایمنی و سلامت فیزیکی گروه و تهدید هویتی به تهدید یگانگی، ارزش و تمایز گروه اشاره دارد (لیائو و همکاران، ۲۰۲۳). مردم ممکن است کو-ربات‌های انسان‌نما را به‌عنوان یک گروه بیرونی درک کنند و به طور فعال در پی تمایز میان دو گروه باشند. از یک طرف، ربات‌هایی که درجه بالایی از انسان‌نمایی دارند، یگانگی گروه انسانی درون گروهی را کاهش می‌دهند و در نتیجه تهدید هویتی برای اعضای گروه درون گروهی ایجاد می‌کنند (پراسیدا و افساری، ۲۰۲۲). از طرف دیگر، مزایای ناشی از کاربرد کو-ربات‌ها به طور اجتناب‌ناپذیری باعث می‌شود که کارگران احساس تهدید واقعی مانند تهدید شغل‌های فعلی و فرصت‌های شغلی کنند، که منجر به تهدید روانی جایگزینی می‌شود (لوتین و همکاران، ۲۰۲۴). بنابراین، ورود کو-ربات‌ها به کارخانه‌ها ممکن است نه تنها تهدید واقعی برای مشاغل انسانی، فرصت‌های شغلی و منابع مادی ایجاد کند، بلکه با تار کردن مرزهای

¹ Appel et al.

² Liu et al.

³ Intergroup threat

⁴ Intergroup Threat

میان انسان‌ها و ربات‌ها، تهدید هویتی نیز به همراه داشته باشد که این نگرانی‌ها می‌توانند مانع اصلی پذیرش کو-ربات‌ها باشند. همچنین، زمانی که کارکنان حس کنند ربات‌های همکار از ایمنی کافی برخوردار هستند و احتمال بروز خطر یا حادثه در استفاده از آن‌ها کم است، احتمال پذیرش و استفاده از این فناوری افزایش می‌یابد (بامگارتن و همکاران، ۲۰۲۲).

مدل‌های پیشین پیچیدگی انسان‌نمایی ربات‌های همکار را به طور یکپارچه بررسی نکرده‌اند. نظریه دره غیرطبیعی صرفاً بر شباهت ظاهری تمرکز دارد و نظریه تهدید بین گروهی عمدتاً بر جنبه‌های هویتی و واقعی تأکید می‌کند. مدل پژوهش حاضر با ترکیب این نظریه‌ها و اضافه کردن سازه ایمنی درک‌شده از پژوهش بامگارتن و همکاران (۲۰۲۲) به مدل پژوهش لیائو و همکاران (۲۰۲۳)، نشان می‌دهد که پذیرش در یک میدان نیروی روان‌شناختی رخ می‌دهد که در آن چندین ارزیابی به طور هم‌زمان در جریان هستند. منطق این مدل آن است که کارکنان در مواجهه با یک ربات انسان‌نما، یک ارزیابی جامع از هزینه-فایده و خطر-فرصت انجام می‌دهند؛ آیا این ربات توانا است؟ (شایستگی)، آیا موجودیت من را تهدید می‌کند؟ (تهدید)، و آیا تعامل با آن برای من بی‌خطر است؟ (ایمنی).

صنعت خودروسازی کشور با چالش‌های پیچیده‌ای در زمینه افزایش بهره‌وری، کاهش هزینه‌ها و بهبود کیفیت تولید روبه‌رو است (هادی رشمئی و همکاران، ۱۴۰۲؛ اسماعیلی و همکاران، ۱۴۰۰) و استفاده از ربات‌های همکار در خط تولید یکی از راه‌حل‌های نوین برای پاسخ به این چالش‌ها است (طلایی، ۱۴۰۴). با این حال، پذیرش تکنولوژی‌های پیشرفته در میان کارکنان، نیازمند توجه به جنبه‌های روانشناختی و اجتماعی آن است (خاتمی فیروز آبادی و همکاران، ۱۳۹۷). یکی از ویژگی‌های مهم ربات‌های همکار، انسان‌نما بودن آن‌ها است که می‌تواند تأثیر زیادی بر نحوه تعامل کارکنان با این فناوری‌ها داشته باشد. بنابراین، هدف این پژوهش بررسی تأثیر انسان‌نما بودن ربات‌های همکار بر پذیرش ربات‌های همکار در صنعت خودروسازی از طریق نقش میانجی شایستگی، ایمنی و تهدید درک شده است.

مبانی نظری

انسان‌نمایی به نسبت دادن ویژگی‌های انسانی، انگیزه‌ها، نیت‌ها، احساسات، تخیل یا رفتار واقعی به موجودات غیر انسانی اشاره دارد (جاوید و همکاران ۱، ۲۰۲۲). انسان‌نمایی می‌تواند درک افراد از رفتار موجودات غیر انسانی را بهبود بخشد و عدم قطعیت در تعامل انسان و ربات را کاهش دهد (لیائو و همکاران، ۲۰۲۴). افراد ترجیح می‌دهند

با ربات‌هایی که نشانه‌های اجتماعی دارند کار کنند، زیرا این نشانه‌ها حس دوستی را القا می‌کنند (کراس و همکاران، ۲۰۲۴). اما برخی بیان کرده‌اند که از دید افراد، ربات‌هایی که وظایف عملکردی انجام می‌دهند، بهتر است بیشتر به جای انسان‌نما بودن، مکانیزه باشند (لکمن و همکاران، ۲۰۲۳). مساله اصلی این است که آیا طراحی انسان‌نمای ربات‌های همکار می‌تواند سیستمی مؤثر ایجاد کند که به‌آسانی با تعاملات اجتماعی انسانی سازگار شود و برای پذیرش ربات‌ها توسط افراد مفید باشد. به نظر می‌رسد، در مقایسه با طراحی انسان‌نما، طراحی مکانیزه ربات‌های همکار بیشتر با انتظارات کارکنان مطابقت دارد و کارکنان تمایل بیشتری به پذیرش ربات‌های همکار مکانیزه برای کار با آن‌ها داشته باشند. اما، همچنان این رابطه مبهم است و بنابراین، فرضیه زیر پیشنهاد می‌شود:

فرضیه یک: انسان‌نما بودن ربات‌های همکار بر پذیرش آن‌ها توسط افراد تاثیر مثبت و معناداری دارد.

شایستگی درک‌شده به مهارت یا کارایی اشاره دارد و در این مطالعه به ارزیابی افراد از این موضوع اشاره دارد که آیا ربات توانایی انجام یک وظیفه خاص را دارد یا خیر. در فرآیند تعامل انسان و ربات، ظاهر ربات که به‌عنوان نقطه تماس اولیه مطرح است، اساس شکل‌گیری اولین برداشت‌ها و ارزیابی‌ها را تشکیل می‌دهد (کوپ و همکاران، ۲۰۲۲). مسئله ظاهر ربات نمی‌تواند جدا از بحث انسان‌نمایی ربات‌ها مورد بررسی قرار گیرد (تائسی و همکاران، ۲۰۲۳). انسان‌نمایی ربات‌ها می‌تواند بر درک افراد از شایستگی ربات‌ها تأثیر بگذارد. انسان‌نمایی در ربات‌های همکار نه تنها در ظاهر آن‌ها منعکس می‌شود، بلکه در جنبه‌های احساسی و رفتاری نیز نمود دارد (لیو و همکاران، ۲۰۲۴). با توجه به اینکه ربات‌های همکار از فناوری‌های هوشمندی مانند تشخیص تصویر و حسگرهای چندگانه استفاده می‌کنند، آن‌ها می‌توانند قابلیت‌هایی همچون حرکت‌های یکپارچه، کنترل نیرو، حفاظت در برابر برخورد و سایر توانایی‌ها را به دست آورند (کوهن و همکاران، ۲۰۲۲). بر خلاف ربات‌های صنعتی سنتی که تنها قادر به انجام وظایف ساده و تکراری هستند، کارکنان می‌توانند شایستگی شبه انسانی ربات‌های همکار را احساس کنند (پالیکا، ۲۰۲۲). بنابراین، فرضیه زیر پیشنهاد می‌شود:

فرضیه دو: انسان‌نما بودن ربات همکار بر شایستگی درک شده تاثیر مثبت و معناداری دارد.

تهدید درک‌شده به برداشت افراد از تهدیدهای واقعی یا خیالی اشاره دارد و شامل تهدیدات مرتبط با واقعیت و هویت می‌شود. تهدیدات واقعی به خطرهایی اشاره دارند که منابع مادی، ایمنی و سلامت جسمی گروه داخلی

1 Kopp et al.

2 Taesi et al.

3 Cohen et al.

4 Paliga

را تهدید می‌کنند، در حالی که تهدیدات هویتی به خطرهایی مرتبط می‌شوند که منحصربه‌فرد بودن، ارزش‌ها و تمایز گروه داخلی را تهدید می‌کنند (لیائو و همکاران، ۲۰۲۳). برخی از مطالعات نشان می‌دهند که انسان‌نمایی نقش مثبتی در بهبود تعامل انسان و ربات ایفا می‌کند و می‌تواند پیش‌بینی و درک رفتار ربات را برای انسان‌ها آسان‌تر کند و تعامل افراد با ربات‌ها را افزایش دهد (کراس و همکاران، ۲۰۲۴). با این حال، برخی تحقیقات معتقدند که مزایای انسان‌نمایی ممکن است تنها در محیط‌ها و جنبه‌های خاصی مشاهده شود (پراسیدا و افساری، ۲۰۲۲). ظاهر ربات‌های همکار نه تنها به صورت انسان‌نما طراحی شده است، بلکه در حرکات و توانایی‌هایشان نیز انسان‌گونه هستند و می‌توانند وظایف مختلف تولیدی را پشتیبانی کنند (کوپ و همکاران، ۲۰۲۲). این موضوع منحصربه‌فرد بودن و جایگاه ویژه کارکنان در کارخانه‌ها را به چالش می‌کشد و تهدیدی برای هویت آن‌ها محسوب می‌شود (کاسانو، ۲۰۲۳). علاوه بر این، با ورود به عصر صنعت ۵.۰ و افزایش انتظارات بازیگران بازار، حضور گسترده ربات‌های همکار در محیط کار بدون شک فرصت‌های شغلی کارکنان را تهدید می‌کند (گورتلر و همکاران، ۲۰۲۳). استفاده از ربات‌ها باعث جایگزینی کارکنان و محروم کردن آن‌ها از دسترسی به منابع مادی اجتماعی می‌شود (کوهن و همکاران، ۲۰۲۲). بر اساس این یافته‌ها، فرضیه زیر در این مطالعه پیشنهاد می‌شود:

فرضیه سه: انسان نما بودن ربات همکار بر تهدید درک شده تأثیر مثبت و معناداری دارد.

تأثیر انسان‌نمایی ربات‌های همکار بر ایمنی درک‌شده به این موضوع اشاره دارد که چگونه طراحی و رفتار انسان‌مانند ربات‌ها بر احساس امنیت و آرامش کارکنان در محیط کاری تأثیر می‌گذارد (تائسی و همکاران، ۲۰۲۳). ایمنی درک‌شده به این معناست که افراد چقدر ربات را به عنوان موجودی بی‌خطر و ایمن در تعاملات کاری تلقی می‌کنند. زمانی که یک ربات دارای ویژگی‌هایی مشابه انسان، مانند حرکات روان، رفتار دوستانه یا واکنش‌های احساسی است، کارکنان ممکن است احساس کنند که تعامل با این ربات طبیعی‌تر و کم‌خطر است (کاسانو، ۲۰۲۳). پیش‌بینی‌پذیری حرکات و رفتارهای ربات‌های انسان‌نما می‌تواند اضطراب ناشی از حرکات غیرمنتظره یا پیچیدگی‌های فنی را کاهش داده و احساس امنیت بیشتری ایجاد کند (پالیگا، ۲۰۲۲). به عنوان مثال، اگر حرکات ربات به آرامی و با الگوهای قابل درک انجام شود، کارکنان کمتر احساس خطر خواهند کرد. بنابراین، فرضیه زیر پیشنهاد می‌شود:

فرضیه چهار: انسان نما بودن ربات همکار بر ایمنی درک شده تأثیر مثبت و معناداری دارد.

1 Cusano

2 Guertler et al.

ربات‌های همکار دارای شایستگی‌های هوشمندی مانند توانایی ابراز احساسات، توانایی استدلال و توانایی ادراک هستند که قابل جایگزینی با توانایی‌های مورد نیاز کارکنان را دارد (کلبه و فیس، ۲۰۲۵). جایگزینی مهارت‌ها به معنای کاهش فرصت‌های شغلی برای کارکنان است و ممکن است منجر به کاهش پذیرش ربات‌های همکار توسط آن‌ها شود. علاوه بر این، ربات‌های همکار دارای قابلیت‌هایی فراتر از توانایی‌های انسانی مانند نگهداری طولانی مدت اجسام، محاسبات سریع و حافظه فوق‌العاده هستند که می‌توانند باعث ایجاد احساس از دست دادن کنترل در کارکنان شوند (لیائو و همکاران، ۲۰۲۴). هرچه درجه انسان‌نمایی ربات‌های همکار بیشتر باشد، احتمالاً شایستگی آن‌ها بیشتر توسط کارکنان درک می‌شود، اما این افزایش شایستگی احتمالاً منجر به کاهش پذیرش آن‌ها توسط کارکنان می‌شود. بر این اساس، فرضیه‌های زیر برای آزمون تجربی رابطه شایستگی درک‌شده و پذیرش ربات‌های همکار پیشنهاد می‌شود.

فرضیه پنج: شایستگی درک شده بر پذیرش ربات‌های همکار تاثیر مثبت و معناداری دارد.

فرضیه میانجی یک: انسان نما بودن ربات همکار بر پذیرش ربات‌های همکار از طریق نقش میانجی شایستگی درک شده تاثیر مثبت و معناداری دارد.

گسترش استفاده از ربات‌های همکار تهدیدی واقعی برای مشاغل انسانی، فرصت‌های شغلی و منابع مادی ایجاد می‌کند (کلبه و فیس، ۲۰۲۵). پیشرفت در سطح هوش و توانایی حرکت ربات‌ها مرزهای میان انسان و ربات را مبهم می‌کند و تهدیدی برای هویت انسانی به وجود می‌آورد. زمانی که افراد با تهدیدهایی مواجه می‌شوند، سیستم بازداری رفتاری آن‌ها وارد عمل شده و منجر به هوشیاری و اضطراب می‌شود. به عبارت دیگر، زمانی که افراد تهدیدهایی را از سوی ربات‌ها در زمینه‌های واقعیت و هویت درک می‌کنند، ممکن است اضطراب و هوشیاری شدیدی را تجربه کنند (گورتلر و همکاران، ۲۰۲۳). اضطراب و ترس از فناوری می‌تواند تمایل افراد به پذیرش را کاهش دهد و در نتیجه، مانعی برای پذیرش فناوری رباتیک توسط افراد ایجاد کند و احتمالاً هرچه درجه انسان‌نمایی ربات‌های همکار بیشتر باشد، کارکنان تهدید بیشتری از سوی این ربات‌ها احساس می‌کنند، که در نهایت منجر به کاهش پذیرش آن‌ها توسط کارکنان می‌شود. بر این اساس، فرضیه‌های زیر در این مطالعه پیشنهاد می‌شود:

فرضیه شش: تهدید درک شده بر پذیرش ربات‌های همکار تاثیر مثبت و معناداری دارد.

فرضیه میانجی دو: انسان نما بودن ربات همکار بر پذیرش ربات‌های همکار از طریق نقش میانجی تهدید درک شده تاثیر مثبت و معناداری دارد.

ایمنی درک‌شده به ارزیابی ذهنی کاربران از میزان امنیت و بی‌خطری ربات‌ها در محیط کاری اشاره دارد (کاسانو، ۲۰۲۳). زمانی که کارکنان احساس کنند ربات‌های همکار خطری برای آن‌ها یا محیط کاری ایجاد نمی‌کنند و حرکات و رفتار آن‌ها قابل پیش‌بینی و ایمن است، احتمال پذیرش و استفاده از این فناوری افزایش می‌یابد (بامگارتن و همکاران، ۲۰۲۲). ایمنی درک‌شده نه تنها نگرانی‌های فیزیکی، مانند برخورد یا آسیب را کاهش می‌دهد، بلکه اضطراب روانی ناشی از عدم اطمینان درباره تعامل با فناوری را نیز از بین می‌برد (کوهن و همکاران، ۲۰۲۲). مطالعات نشان داده‌اند که ایمنی درک‌شده نقش مثبتی در پذیرش فناوری‌های جدید ایفا می‌کند، زیرا افراد تمایل بیشتری به تعامل با فناوری‌هایی دارند که احساس امنیت را در آن‌ها تقویت می‌کنند (کوپ و همکاران، ۲۰۲۲). ویژگی‌هایی مانند حرکات روان، رفتار دوستانه، و واکنش‌های احساسی که از طراحی انسان‌نما ناشی می‌شوند، می‌توانند تعامل انسان و ربات را طبیعی‌تر کنند و به کاهش احساس خطر کمک کنند (تائسی و همکاران، ۲۰۲۳). افزایش ایمنی درک‌شده به واسطه طراحی انسان‌نما می‌تواند نقش مهمی در پذیرش ربات‌های همکار ایفا کند. به بیان دیگر، زمانی که طراحی انسان‌نما ایمنی درک‌شده را بهبود می‌بخشد، پذیرش فناوری نیز به صورت غیر مستقیم تحت تأثیر قرار می‌گیرد. بر این اساس، فرضیه‌های زیر در این مطالعه پیشنهاد می‌شود:

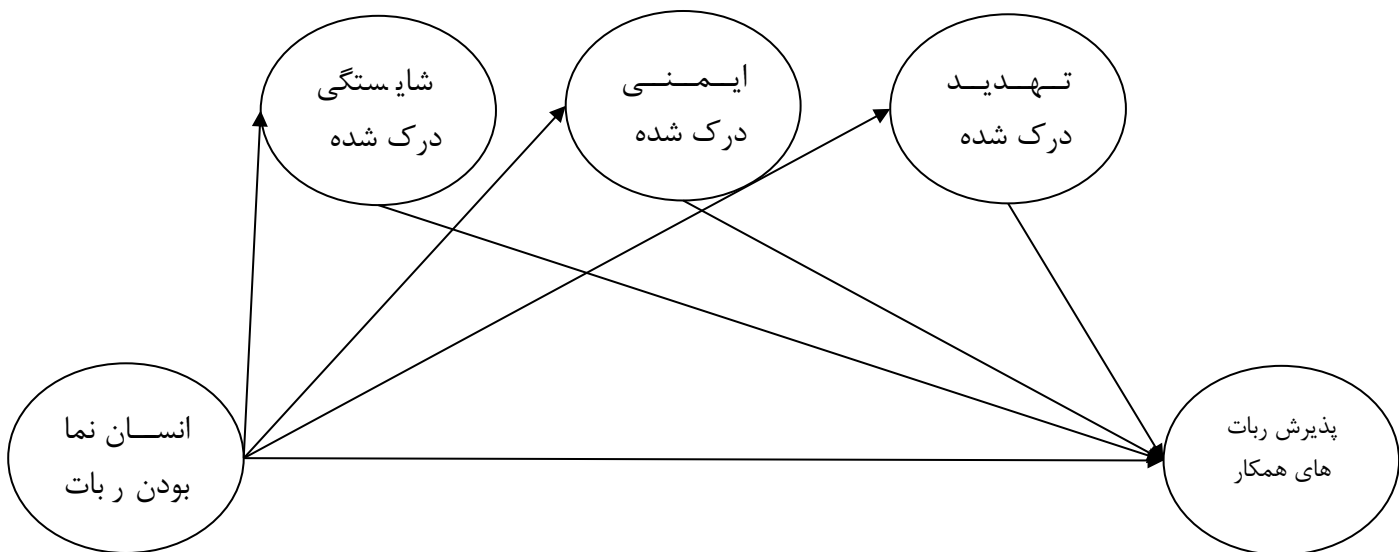
فرضیه هفت: ایمنی درک‌شده بر پذیرش ربات‌های همکار تأثیر مثبت و معناداری دارد.

فرضیه میانجی سه: انسان‌نما بودن ربات همکار بر پذیرش ربات‌های همکار از طریق نقش میانجی ایمنی درک‌شده تأثیر مثبت و معناداری دارد.

پیشینه تحقیق

با مرور انتقادی ادبیات مرتبط می‌توان مشاهده کرد که پژوهش‌های پیشین درباره پذیرش ربات‌های همکار عمدتاً بر مدل‌های کلاسیک پذیرش فناوری یا بر بررسی اثرات منفرد انسان‌نمایی تمرکز داشته‌اند و کمتر به سازوکارهای روان‌شناختی میانجی توجه کرده‌اند. مطالعاتی مانند لیائو و همکاران (۲۰۲۳؛ ۲۰۲۴) و بامگارتن و همکاران (۲۰۲۲) نشان داده‌اند که انسان‌نمایی می‌تواند ادراکاتی همچون شایستگی، ایمنی و تهدید را تحت تأثیر قرار دهد، اما نتایج این پژوهش‌ها درباره نقش نهایی این سازه‌ها در پذیرش ربات‌ها یکدست و همسو نیست. به‌ویژه، برخی پژوهش‌ها اثر مثبت شایستگی درک‌شده بر پذیرش را گزارش کرده‌اند، در حالی که شواهد نظری و تجربی دیگری نشان می‌دهد شایستگی بالا می‌تواند به احساس تهدید، جایگزین‌پذیری شغلی و از دست دادن کنترل منجر شود. از سوی دیگر، نقش ایمنی درک‌شده اغلب به صورت یک پیامد جانبی طراحی فنی ربات بررسی شده و کمتر به عنوان یک مسیر میانجی کلیدی در پذیرش تحلیل شده است. همچنین، بخش قابل توجهی از ادبیات موجود

در زمینه‌های غیر صنعتی یا در بافت‌های فرهنگی متفاوت انجام شده و شواهد تجربی محدودی درباره صنعت خودرو و زمینه فرهنگی ایران وجود دارد. بر این اساس، خلأ اصلی پژوهش حاضر در نبود یک مدل یکپارچه است که به‌طور هم‌زمان اثر انسان‌نما بودن ربات‌های همکار را از طریق سه سازه شایستگی، تهدید و ایمنی درک‌شده بر پذیرش کارکنان بررسی کند و نشان دهد کدام مسیرها در بافت صنعتی و فرهنگی ایران تقویت یا تضعیف می‌شوند. پژوهش حاضر با ارائه چنین مدلی تلاش می‌کند این خلأ را پوشش دهد. براساس فرضیه‌های مطرح شده، مدل مفهومی پژوهش در شکل شماره ۱ نمایش داده شده است.



شکل ۱: مدل مفهومی پژوهش برگرفته از لیائو و همکاران (۲۰۲۳) و بامگارتن و همکاران (۲۰۲۲)

روش پژوهش

پژوهش حاضر از حیث هدف، کاربردی و از حیث ماهیت، توصیفی-همبستگی است. جامعه آماری این پژوهش کارمندان صنعت خودروسازی کشور بودند. باتوجه به اینکه در مطالعات مربوط به مدلسازی معادلات ساختاری برای محاسبه حجم نمونه از رابطه $5q < n < 15q$ استفاده می‌شود (هیر جی آر و همکاران^۱، ۲۰۲۱) که در فرمول فوق q تعداد سوالات پرسشنامه و n اندازه نمونه است. تعداد سوالات پرسشنامه پژوهش ۲۰ سوال بوده است و حداقل حجم نمونه کافی ۱۰۰ نفر می‌باشد. در پژوهش حاضر نیز از رابطه فوق برای محاسبه حجم نمونه استفاده

¹ Hair Jr et al.

می‌شود و ۳۸۴ نفر به عنوان نمونه آماری و بر اساس نمونه‌گیری تصادفی در نظر گرفته می‌شوند، که این تعداد با جدول مورگان نیز سازگار است. برای جمع‌آوری داده‌ها از پرسشنامه‌های استاندارد لیانو و همکاران (۲۰۲۳) و بامگارتن و همکاران (۲۰۲۲) استفاده شده است. همه متغیرهای پژوهش در قالب طیف پنج درجه‌ای لیکرت (۱- کاملاً مخالفم تا ۵- کاملاً موافقم) بررسی شدند. تحلیل داده‌ها در دو بخش آمار توصیفی و استنباطی با استفاده از نرم‌افزارهای SPSS و SmartPLS نسخه ۳ انجام شد. روایی صوری پرسشنامه با استفاده از نظرات خبرگان صنعت خودرو و اساتید متخصص دانشگاه تأیید شد.

در روش مدلسازی معادلات ساختاری، سه سطح مدل بیرونی (روایی و پایایی)، مدل درونی (مدل ساختاری) و برازش کلی مدل ارزیابی می‌گردد. برای ارزیابی مدل اندازه‌گیری، پایایی سازه‌ها با استفاده از آلفای کرونباخ، پایایی ترکیبی و پایایی اشتراکی ρ_A و روایی همگرا با شاخص AVE (Communality) و بار عاملی بررسی شد. روایی واگرا نیز با استفاده از معیار فورنل لارکر، بارهای عرضی و آزمون HTMT مورد ارزیابی قرار گرفت. کیفیت مدل اندازه‌گیری و ساختاری با استفاده از شاخص‌های R^2 ، Q^2 و F^2 سنجیده شد و برازش کلی مدل با شاخص‌های SRMR و NFI بررسی گردید. به منظور بررسی روایی، از روایی همگرا به معنی همبستگی سوالات مرتبط با یک متغیر با همان متغیر استفاده شده که با دو معیار ضرایب بار عاملی (حد قابل قبول ۰/۴)، میانگین واریانس به اشتراک گذاشته شده AVE (Communality) (حد قابل قبول ۰/۵) و آماره عامل تورم واریانس (VIF) سنجیده می‌گردد. آماره عامل تورم واریانس (VIF) برای ارزیابی هم‌خطی بین شاخص‌های اندازه‌گیری یک سازه استفاده می‌شود. این معیار نشان می‌دهد که آیا شاخص‌های یک سازه به شدت با هم همبسته هستند یا نه. اگر مقدار VIF بیرونی یک شاخص بیشتر از ۵ باشد، نشان‌دهنده هم‌خطی چندگانه قابل توجهی است و ممکن است نیاز به بازنگری یا حذف شاخص‌های مشکل‌دار باشد. هدف اصلی از ارزیابی VIF بیرونی، تضمین این است که شاخص‌های اندازه‌گیری یک سازه مستقل از هم بوده و به صورت معتبر و پایا سازه را اندازه‌گیری می‌کنند (لاو و فونگ^۱، ۲۰۲۰). برای ارزیابی هم‌خطی بین سازه‌های پنهان مدل، از Inner VIF استفاده شد که برای اندازه‌گیری همبستگی میان متغیرهای پنهان مدل به کار گرفته می‌شود که نباید بیشتر از مقدار بحرانی ۵ باشد. این ارزیابی کمک می‌کند تا از هم‌خطی شدید بین سازه‌های پنهان جلوگیری شود و اعتبار مدل افزایش یابد (هیر جی آر و همکاران، ۲۰۲۱). برای سنجش روایی واگرا در سطح متغیر پنهان، از فورنل لارکر، بارهای عرضی و آزمون HTMT استفاده می‌شود. فورنل و لارکر برای بررسی روایی واگرا ماتریسی را پیشنهاد می‌دهند که قطر اصلی این ماتریس

¹ Law and Fong

جذر مقادیر AVE مربوط به هر یک از متغیرها است. در این ماتریس، باید همبستگی هر سازه با سازه‌های خود در مقایسه با سایر سازه‌ها بیشتر است تا روایی واگرا تأیید گردد (هیر جی آر و همکاران، ۲۰۲۱). همچنین، مقادیر HTMT بین سازه‌ها باید کمتر از مقدار بحرانی ۰/۸۵ باشد (هیر جی آر و همکاران، ۲۰۲۱). برای سنجش پایایی از معیارهای آلفای کرونباخ، پایایی ترکیبی و پایایی اشتراکی ρ_A استفاده شده است که حد قابل قبول این دو معیار ۰/۷ است. (النعمی و همکاران^۱، ۲۰۲۱).

برای سنجش کیفیت مدل اندازه‌گیری از معیارهای Q^2 و R^2 استفاده شده است. معیار Q^2 قدرت پیش‌بینی مدل را مشخص می‌سازد. یک مقدار Q^2 بزرگتر از صفر برای یک متغیر پنهان درون‌زا نشان می‌دهد که مدل مسیر دارای ارتباط پیش‌بینی کننده برای این سازه است. معیار R^2 برای متصل کردن بخش اندازه‌گیری و بخش ساختاری مدل‌سازی معادلات ساختاری به کار می‌رود و نشان از تأثیری دارد که یک متغیر مستقل بر یک متغیر وابسته می‌گذارد. معیار R^2 تنها برای سازه‌های وابسته مدل محاسبه می‌گردد و در مورد سازه‌های برون‌زا (مستقل) مقدار این معیار صفر است. مقدار این شاخص بین صفر تا یک می‌باشد و اگر از ۰/۶ بیشتر باشد نشان می‌دهد، متغیرهای مستقل تا حد زیادی توانسته‌اند تغییرات متغیر وابسته را تبیین کنند. هر چه مقدار R^2 مربوط به سازه‌های درون‌زای یک مدل بیشتر باشد، نشان از برازش بهتر مدل است (هیر جی آر و همکاران، ۲۰۲۱). در این پژوهش، برازش کلی مدل با شاخص ریشه میانگین مربعات باقیمانده استاندارد (SRMR) و شاخص تناسب به‌هنگار (NFI) مورد سنجش قرار گرفته است. شاخص SRMR بین صفر تا یک است و هر قدر که کوچکتر باشد بیانگر برازش بیشتر کل مدل است. خط برش این شاخص ۰/۸ است. به عبارت دیگر چنانچه SRMR یک مدل ۰/۸ یا کمتر باشد بیانگر برازش کلی بالای مدل است و هر قدر که بیشتر از ۰/۸ باشد بیانگر برازش کمتر مدل است. شاخص تناسب به‌هنگار (NFI) نیز بالای ۰/۷ قابل قبول و بالاتر از ۰/۹ مطلوب است (النعمی و همکاران، ۲۰۲۱). به‌منظور آزمون نقش میانجی، ابتدا اثر مستقیم انسان‌نما بودن ربات همکار بر پذیرش بدون حضور متغیرهای میانجی آزمون گردید. سپس مدل با حضور هر یک از متغیرهای میانجی اجرا و اثر غیر مستقیم با استفاده از روش Bootstrapping بررسی شد. در نهایت، شاخص VAF برای تعیین نوع میانجی‌گری محاسبه و تفسیر می‌گردد. این رویکرد امکان تمایز میان میانجی‌گری کامل، جزئی یا عدم میانجی‌گری را فراهم می‌سازد. با توجه به ماهیت اکتشافی مدل، وجود سازه‌های میانجی متعدد، و تمرکز بر پیش‌بینی و تبیین واریانس، استفاده از رویکرد PLS-SEM در این پژوهش استفاده شد.

¹ AlNuaimi et al.

یافته‌ها

آمار توصیفی

در بخش آمار توصیفی پژوهش، اطلاعات جمعیت‌شناختی ۳۸۴ پاسخ‌دهنده مورد بررسی قرار گرفت. از نظر سنی، بیشترین تعداد پاسخ‌دهندگان در بازه سنی ۳۱ تا ۴۰ سال قرار داشتند (۴۳٪)، در حالی که ۲۷٪ در بازه ۲۰ تا ۳۰ سال، ۱۹٪ در بازه ۴۱ تا ۵۰ سال، و تنها ۱۱٪ بالای ۵۱ سال بودند. از نظر سطح تحصیلات، اکثریت پاسخ‌دهندگان دارای مدرک کارشناسی بودند (۵۵٪)، در حالی که ۲۵٪ دارای مدرک کارشناسی ارشد و بالاتر و ۲۰٪ دارای مدرک زیر دیپلم یا دیپلم بودند. همچنین، بررسی جنسیت نشان داد که ۷۶٪ از پاسخ‌دهندگان مرد و ۲۴٪ زن بودند. نتایج نشان‌دهنده غلبه جمعیت مرد، با تحصیلات دانشگاهی، و عمدتاً در بازه سنی میانی در جامعه آماری پژوهش است.

ارزیابی مدل اندازه‌گیری

ارزیابی مدل اندازه‌گیری در سه بخش پایایی، روایی همگرا، روایی واگرا بیان می‌شود.

پایایی

پایایی تمام سازه‌ها در جدول شماره ۱ ارائه شده است. تمامی متغیرهای پژوهش از مقادیر پایایی بالاتر از ۰/۷ برخوردار هستند که نشان از پایایی قابل قبول پرسشنامه دارد.

جدول ۱: پایایی سازه‌ها

سازه‌ها	پایایی	پایایی ترکیبی	پایایی اشتراکی rho_A
انسان نما بودن ربات همکار	۰/۸۰۹	۰/۸۷۵	۰/۸۰۹
ایمنی درک شده	۰/۸۳۱	۰/۸۸۸	۰/۸۳۲
تهدید درک شده	۰/۸۲۱	۰/۸۸۲	۰/۸۲۶
شایستگی درک شده	۰/۷۷۰	۰/۸۵۳	۰/۷۷۶
پذیرش ربات همکار	۰/۷۷۱	۰/۸۵۲	۰/۷۹۳

روایی همگرا

جدول ۲، مقادیر بار عاملی، مقدار VIF و AVE (Communality) سازه‌ها و سوالات پرسشنامه را نشان می‌دهد. بار عاملی تمامی سوالات (گویه) ضرایب بالاتر از ۰/۴ دارند، که نشان از روایی مناسب پرسشنامه این پژوهش دارد. همچنین، مقدار VIF برای تمام گویه‌ها کمتر از ۵ است، که نشان می‌دهد شاخص‌های اندازه‌گیری هر سازه مستقل

از هم بوده و به صورت معتبر اندازه گیری می کنند. در مورد (AVE (Communality، مقدار بحرانی عدد ۰/۵ است و روایی همگرا قابل قبول است.

جدول ۲: مقادیر بار عاملی، مقدار VIF و AVE (Communality)

ردیف	سازه	گویه	بار عاملی	VIF	AVE (Communality)
۱	انسان بودن ربات همکار	گاهها ممکن است این رباتها را با یک انسان واقعی اشتباه بگیرم.	۰/۸۱۵	۱/۸۱	۰/۶۳۶
۲		حرکات این رباتها به گونه ای است که می تواند منجر به تعامل طبیعی و مشابه انسان می شود.	۰/۷۹۷	۱/۷۲	
۳		رفتارهای این رباتها من را گاهها به یاد رفتارهای انسانی می اندازد.	۰/۷۸۰	۱/۵۷	
۴		من حس می کنم این رباتها در حال تبدیل به شبه انسان شدن هستند.	۰/۷۹۷	۱/۶۶	
۵	شیستگی درک شده	این رباتها توانایی انجام وظایف محوله به شکلی مؤثر را دارند.	۰/۷۳۵	۱/۴۸	۰/۵۹۲
۶		احساس می کنم که می توان به عملکرد این رباتها اعتماد کرد.	۰/۸۰۸	۱/۶۰	
۷		این رباتها به نظر می رسد توانمندی لازم برای انجام کارها را دارند.	۰/۷۶۴	۱/۵۱	
۸		این رباتها به گونه ای عمل می کنند که نشان دهنده تجربه و تخصص در وظایف محوله است.	۰/۷۶۹	۱/۵۳	
۹	رهبر درک شده	استفاده فزاینده از رباتهای همکار ممکن است فرصت های شغلی انسانی را کاهش دهد.	۰/۸۵۴	۲/۱۴	۰/۶۵۱
۱۰		پیشرفت در فناوری رباتهای همکار ممکن است موجب تضعیف جایگاه انسان در محیط کار شود.	۰/۸۰۱	۱/۹۱	
۱۱		حرکات و توانایی های رباتهای همکار می تواند موجب نگرانی از کاهش امنیت کاری شود.	۰/۷۴۸	۱/۵۵	
۱۲		پیشرفت های اخیر در رباتیک همکار ممکن است بر درک ما از معنای انسان بودن تأثیر بگذارد.	۰/۸۲۱	۱/۷۳	
۱۳	ایمنی درک شده	این رباتها به گونه ای طراحی شده اند که در هنگام کار احتمال خطر برای کاربران را به حداقل می رسانند.	۰/۷۵۲	۱/۴۸	۰/۶۶۵
۱۴		احساس می کنم حرکات این رباتها به شکلی قابل پیش بینی و ایمن انجام می شود.	۰/۸۴۲	۲/۲۷	
۱۵		عملکرد این رباتها باعث کاهش نگرانی من از خطرات احتمالی در محیط کار می شود.	۰/۸۶۰	۲/۴۳	
۱۶		به نظر می رسد که سیستمها و قطعات این رباتها، استانداردهای ایمنی را رعایت می کنند.	۰/۸۰۵	۱/۶۶	

۰/۵۹۱	۱/۳۷	۰/۷۴۱	استفاده از این ربات‌ها می‌تواند بهره‌وری من را در انجام وظایف کاری افزایش دهد.	۳۰ ربات‌های همکار	۱۷
	۱/۸۴	۰/۸۵۶	احساس می‌کنم کار کردن با این ربات‌ها تجربه‌ای مفید و رضایت‌بخش خواهد بود.		۱۸
	۱/۸۴	۰/۷۹۱	این ربات‌ها می‌توانند به‌طور مستقل وظایف مورد نیاز را انجام داده و به کاهش بار کاری کمک کنند.		۱۹
	۱/۵۱	۰/۶۷۵	حضور این ربات‌ها در محیط کاری به نظر می‌رسد که همکاری مؤثری را ایجاد کند.		۲۰

روایی واگرا

به‌منظور بررسی روایی واگرا، آزمون بارهای عرضی انجام شد. نتایج نشان داد که بار عاملی هر گویه بر سازه مربوط به خود به‌طور معناداری بالاتر از بار آن بر سایر سازه‌ها است. این یافته نشان می‌دهد که هر گویه به‌درستی سازه مورد نظر را اندازه‌گیری کرده و تداخل مفهومی معناداری میان سازه‌ها وجود ندارد (جدول ۳).

جدول ۳: جدول بارهای عرضی گویه‌ها

پذیرش ربات همکار	شایستگی درک شده	تهدید درک شده	ایمنی درک شده	انسان نما بودن ربات همکار	
۰/۴۰۰	۰/۶۸۵	۰/۶۳۰	۰/۵۱۸	۰/۸۱۵	q1
۰/۳۹۲	۰/۶۶۵	۰/۵۶۵	۰/۵۴۵	۰/۷۹۷	q2
۰/۴۹۰	۰/۶۲۰	۰/۵۹۸	۰/۵۵۶	۰/۷۸۰	q3
۰/۴۹۶	۰/۶۵۵	۰/۶۱۷	۰/۵۱۷	۰/۷۹۷	q4
۰/۳۷۶	۰/۷۳۵	۰/۵۳۲	۰/۴۳۵	۰/۵۴۲	q5
۰/۴۳۲	۰/۸۰۸	۰/۶۶۵	۰/۶۵۸	۰/۷۱۹	q6
۰/۳۹۹	۰/۷۶۴	۰/۵۷۴	۰/۵۳۸	۰/۶۲۷	q7
۰/۳۸۹	۰/۷۶۹	۰/۶۴۷	۰/۵۱۶	۰/۶۲۹	q8
۰/۴۶۹	۰/۶۸۶	۰/۸۵۴	۰/۵۷۴	۰/۶۳۸	q9
۰/۴۳۳	۰/۶۱۸	۰/۸۰۱	۰/۶۲۷	۰/۵۵۷	q10
۰/۳۸۳	۰/۶۱۰	۰/۷۴۸	۰/۵۰۲	۰/۶۱۳	q11
۰/۵۶۳	۰/۶۳۱	۰/۸۲۱	۰/۷۲۷	۰/۶۳۲	q12
۰/۵۳۲	۰/۶۳۷	۰/۶۹۹	۰/۷۵۲	۰/۵۷۶	q13
۰/۶۳۷	۰/۴۹۷	۰/۶۳۳	۰/۸۴۲	۰/۵۱۵	q14
۰/۶۳۹	۰/۵۲۷	۰/۵۰۷	۰/۸۶۰	۰/۵۰۱	q15
۰/۶۱۶	۰/۶۳۷	۰/۶۳۰	۰/۸۰۵	۰/۵۹۰	q16

پذیرش ربات همکار	شایستگی درک شده	تهدید درک شده	ایمنی درک شده	انسان نما بودن ربات همکار	
۰/۷۴۱	۰/۴۲۲	۰/۵۱۴	۰/۶۱۹	۰/۵۳۸	q17
۰/۸۵۶	۰/۵۷۵	۰/۵۸۵	۰/۶۶۹	۰/۵۶۵	q18
۰/۷۹۱	۰/۳۰۱	۰/۳۱۵	۰/۵۴۵	۰/۲۸۴	q19
۰/۶۷۵	۰/۲۵۰	۰/۲۸۴	۰/۴۰۴	۰/۲۳۸	q20

در ادامه، برای بررسی دقیق‌تر روایی واگرا، آزمون نسبت همبستگی ناهمگون HTMT مورد استفاده قرار گرفت (جدول ۴). نتایج نشان داد که تمامی مقادیر HTMT بین سازه‌های پژوهش کمتر از مقدار بحرانی ۰/۸۵ هستند که حاکی از تأیید روایی واگرا بین سازه‌ها است.

جدول ۴: آزمون نسبت همبستگی ناهمگون HTMT

پذیرش ربات همکار	شایستگی درک شده	تهدید درک شده	ایمنی درک شده	انسان نما بودن ربات همکار	
					انسان نما بودن ربات همکار
				۰/۸۱۸	ایمنی درک شده
			۰/۷۱۲	۰/۷۱۹	تهدید درک شده
		۰/۶۵۴	۰/۸۰۱	۰/۶۵۴	شایستگی درک شده
۰/۶۴۲	۰/۶۸۸	۰/۵۴۳	۰/۶۶۷		پذیرش ربات همکار

روایی واگرا در سطح متغیر پنهان توسط ماتریس فورنل و لارکر قابل تشخیص است. مطابق جدول ۵، همبستگی یک متغیر با سازه‌های خود در مقایسه با سایر متغیرها بیشتر است که نشان از روایی واگرا قابل قبول پرسشنامه دارد.

جدول ۵: جدول فورنل - لارکر

پذیرش ربات همکار	شایستگی درک شده	تهدید درک شده	ایمنی درک شده	انسان نما بودن ربات همکار	
					انسان نما بودن ربات همکار
			۰/۸۱۶	۰/۶۷۰	ایمنی درک شده
		۰/۸۰۷	۰/۷۵۷	۰/۷۵۷	تهدید درک شده
	۰/۷۶۹	۰/۷۹۸	۰/۷۰۵	۰/۷۲۳	شایستگی درک شده
۰/۷۶۹	۰/۵۲۰	۰/۵۷۷	۰/۷۴۴	۰/۵۵۸	پذیرش ربات همکار

ارزیابی کیفیت مدل

به‌منظور ارزیابی کیفیت پیش‌بینی مدل اندازه‌گیری و مدل ساختاری، از روش Blindfolding و شاخص Q^2 استفاده شد که در ادبیات PLS-SEM به‌ترتیب به‌عنوان CV-Communality و CV-Redundancy شناخته می‌شوند. شاخص CV-Com یا Q^2 مربوط به Communality بیانگر توان مدل اندازه‌گیری در بازتولید مقادیر مشاهده‌شده متغیرهای درون‌زا است و مقادیر بالاتر از صفر نشان‌دهنده کیفیت پیش‌بینی قابل قبول مدل اندازه‌گیری می‌باشد. همچنین، شاخص CV-Red یا Q^2 Redundancy برای ارزیابی کیفیت پیش‌بینی مدل ساختاری محاسبه شد (جدول ۶). نتایج نشان داد که مقادیر Q^2 برای تمامی سازه‌های درون‌زا مثبت است که بیانگر قدرت پیش‌بینی مناسب مدل ساختاری و تأیید کفایت روابط علی ترسیم‌شده در مدل مفهومی پژوهش است.

جدول ۶: مقادیر Q^2

ردیف	سازه‌های درون‌زا	Q^2
۱	ایمنی درک شده	۰/۲۷۹
۲	تهدید درک شده	۰/۳۵۱
۳	شایستگی درک شده	۰/۳۷۸
۴	پذیرش ربات همکار	۰/۳۰۶

ارزیابی مدل ساختاری

برای ارزیابی هم‌خطی بین متغیرهای پنهان مدل، Inner VIF محاسبه شد. تمامی مقادیر Inner VIF برای متغیرهای پنهان مدل در محدوده مجاز قرار داشتند و کمتر از ۵ بودند، که نشان‌دهنده عدم وجود هم‌خطی شدید میان متغیرهای پنهان است (جدول ۷). نتایج نشان می‌دهد که مدل به‌طور مناسب و معتبر می‌تواند رابطه‌های میان متغیرهای پنهان را ارزیابی کند و هیچ تأثیری بر دقت نتایج حاصل از تخمین‌های مدل نداشته است.

جدول ۷: مقادیر Inner VIF

پذیرش ربات همکار	شایستگی درک شده	تهدید درک شده	ایمنی درک شده	انسان نما بودن ربات همکار	
۳/۴۶۰	۱/۰۰	۱/۰۰	۱/۰۰	انسان نما بودن ربات همکار	
۲/۵۴۶				ایمنی درک شده	
۳/۵۶۷				تهدید درک شده	
۳/۹۹۴				شایستگی درک شده	
				پذیرش ربات همکار	

به منظور ارزیابی توان تبیین مدل ساختاری، ضریب تعیین R^2 برای سازه‌های درون‌زا محاسبه شد (جدول ۸). ضریب R^2 نشان می‌دهد چه میزان از واریانس سازه‌های وابسته توسط متغیرهای مستقل مدل تبیین می‌شود. نتایج نشان داد که مقادیر R^2 به دست آمده برای سازه‌های درون‌زا در سطوح قابل قبول قرار دارند که بیانگر توان تبیین مناسب مدل و کفایت روابط ساختاری ترسیم شده در تبیین متغیرهای وابسته پژوهش است.

جدول ۸: مقادیر R^2

ردیف	متغیرهای درون‌زا	R^2
۱	ایمنی درک شده	۰/۴۴۹
۲	تهدید درک شده	۰/۵۷۴
۳	شایستگی درک شده	۰/۶۷۷
۴	پذیرش ربات همکار	۰/۵۷۶

در ادامه، برای بررسی شدت تأثیر هر یک از مسیرهای ساختاری، اندازه اثر F^2 محاسبه شد (جدول ۹). شاخص F^2 میزان تغییر در مقدار R^2 یک سازه درون‌زا را در اثر حذف یک متغیر پیش‌بین نشان می‌دهد و معیاری برای ارزیابی اهمیت نسبی مسیرها محسوب می‌شود. نتایج نشان داد که اندازه اثر مسیرهای کلیدی مدل در محدوده قابل قبول قرار دارد که حاکی از نقش معنادار این مسیرها در تبیین سازه‌های درون‌زا و تأیید ساختار علی مدل پژوهش است.

جدول ۹: اندازه اثر F^2

پذیرش ربات همکار	شایستگی درک شده	تهدید درک شده	ایمنی درک شده	انسان نما بودن ربات همکار	
۰/۰۲۸	۲/۱۰۰	۱/۳۴۶	۰/۸۱۵	انسان نما بودن ربات همکار	
۰/۴۶۱				ایمنی درک شده	
۰/۰۰۰				تهدید درک شده	
۰/۰۱۴				شایستگی درک شده	
				پذیرش ربات همکار	

آزمون فرضیه‌های مستقیم

برای آزمون فرضیه از معیار آماره t استفاده می‌شود. در نرم‌افزار Smart PLS این ضرایب با استفاده از دستور Bootstrapping محاسبه می‌شوند. آماره t باید از ۱٫۹۶ بیشتر باشند تا در سطح معناداری ۹۵٪ معنادار بودن فرضیه را تأیید کرد. نتایج حاصل از آزمون فرضیه‌های پژوهش در جدول ۱۰ خلاصه شده است. همان‌گونه که در

جدول ۱۰ و شکل ۲ مشاهده می‌شود، ۵ فرضیه مستقیم تأیید و ۲ فرضیه مستقیم تأثیر شایستگی درک‌شده و تهدید درک‌شده بر پذیرش ربات‌های همکار تأیید نشده‌اند.

جدول ۱۰: نتایج حاصل از آزمون فرضیه‌ها

ردیف	مسیر	آماره t	ضریب مسیر	نتیجه
۱	انسان نما بودن ربات همکار بر پذیرش ربات‌های همکار	۳/۵۴۷	۰/۲۰۴	تأیید
۲	انسان نما بودن ربات همکار بر شایستگی درک‌شده	۴۱/۷۲۹	۰/۸۲۳	تأیید
۳	انسان نما بودن ربات همکار بر تهدید درک‌شده	۲۵/۳۵۰	۰/۷۵۷	تأیید
۴	انسان نما بودن ربات همکار بر ایمنی درک‌شده	۱۷/۶۱۳	۰/۶۷۰	تأیید
۵	شایستگی درک‌شده بر پذیرش ربات‌های همکار	۱/۹۳۹	-۰/۱۵۴	عدم تأیید
۶	تهدید درک‌شده بر پذیرش ربات‌های همکار	۰/۰۵۹	۰/۰۰۵	عدم تأیید
۷	ایمنی درک‌شده بر پذیرش ربات‌های همکار	۱۵/۶۸۵	۰/۷۱۳	تأیید

با وجود اینکه مقادیر عامل تورم واریانس درونی (Inner VIF) برای متغیرهای پنهان مدل در محدوده مجاز و کمتر از ۵ قرار دارد و نشان‌دهنده عدم وجود هم‌خطی شدید بین متغیرهای پنهان است، مشاهده ضریب منفی در مسیر شایستگی به پذیرش می‌تواند به دلیل اثر سرکوب باشد. اثر سرکوب زمانی رخ می‌دهد که متغیرهای دیگر به طور غیر مستقیم بر مسیر مورد نظر تأثیر می‌گذارند. به عبارت دیگر، متغیرهای دیگر مدل ممکن است بر این مسیر تأثیر گذاشته و باعث معکوس شدن ضریب شوند، حتی زمانی که همبستگی مثبت میان دو سازه وجود دارد. این پدیده در مدل‌های پیچیده با تعاملات زیاد ممکن است مشاهده شود.

آزمون فرضیه‌های میانجی

تحلیل میانجی‌گری به صورت گام‌به‌گام و مطابق رویه استاندارد انجام شد. در گام نخست، اثر مستقیم انسان‌نما بودن ربات همکار بر پذیرش ربات‌های همکار بدون حضور متغیرهای میانجی آزمون شد که این مسیر معنادار بود. در گام دوم، مدل با حضور هر یک از متغیرهای میانجی به صورت جداگانه اجرا گردید. نتایج نشان داد که اثر غیر مستقیم انسان‌نما بودن بر پذیرش از طریق ایمنی درک‌شده معنادار است، در حالی که اثر غیر مستقیم از طریق شایستگی و تهدید درک‌شده معنادار نبود. شاخص VAF برای مسیر ایمنی درک‌شده برابر با ۰/۸۵ به دست آمد که نشان‌دهنده نقش میانجی کامل ایمنی درک‌شده در رابطه بین انسان‌نما بودن و پذیرش ربات‌های همکار است. در مقابل، مقادیر VAF برای شایستگی و تهدید درک‌شده کمتر از ۰/۲ بود که بیانگر عدم ایفای نقش میانجی توسط این متغیرها است. این یافته‌ها نشان می‌دهد که ایمنی درک‌شده سازوکار اصلی انتقال اثر انسان‌نما بودن بر

پذیرش ربات‌های همکار در صنعت خودرو ایران است. نتایج تحلیل میانجی‌گری متغیرهای شایستگی، تهدید و ایمنی درک‌شده در جدول ۱۱ ارائه شده است.

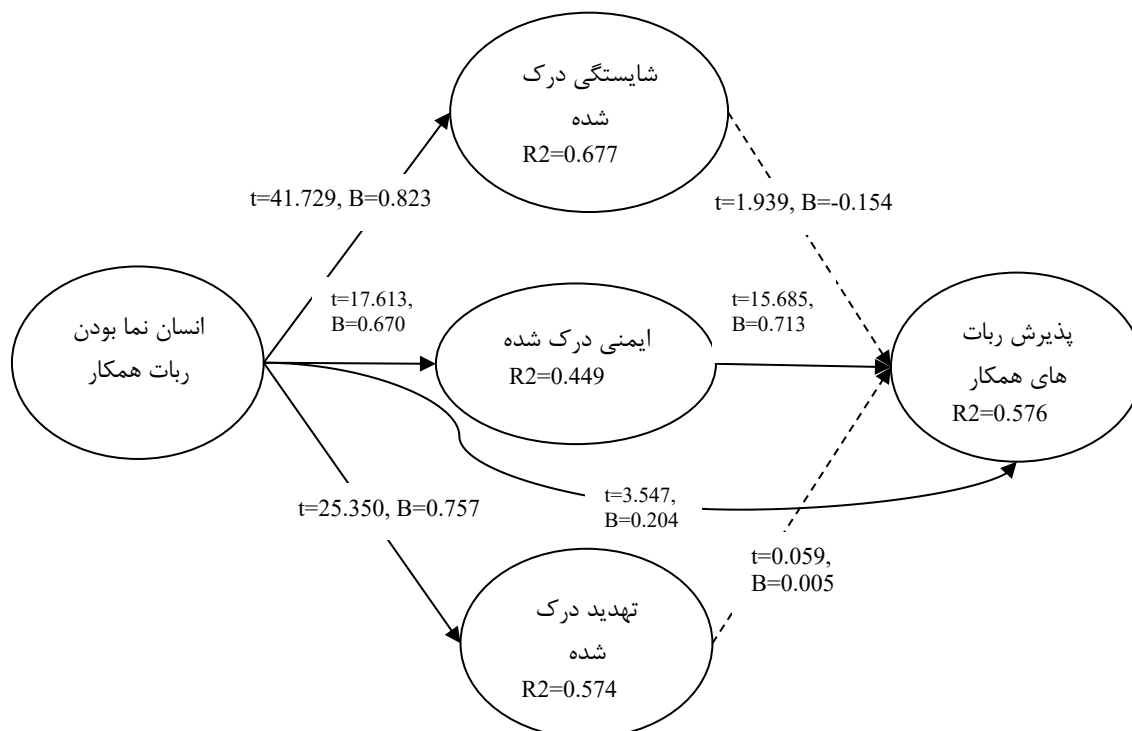
جدول ۱۱: تحلیل میانجی‌گری متغیرهای شایستگی، تهدید و ایمنی درک‌شده

متغیر میانجی	اثر مستقیم بدون میانجی	اثر غیر مستقیم	اثر کل	VAF	نوع میانجی
شایستگی درک‌شده	۰/۵۵۸ ***	۰/۱۲۶ - (غیر معنادار)	۰/۴۳۲	۰/۲۲	عدم میانجی
تهدید درک‌شده	۰/۵۵۸ ***	۰/۰۰۳ (غیر معنادار)	۰/۵۶۱	۰/۰۱	عدم میانجی
ایمنی درک‌شده	۰/۵۵۸ ***	۰/۴۷۷ ***	۰/۵۵۸	۰/۸۵	میانجی کامل

شاخص VAF از تقسیم اثر غیر مستقیم بر اثر کل محاسبه شد. بر اساس معیارهای پیشنهادی، مقدار VAF بالاتر از ۰.۸ نشان‌دهنده میانجی‌گری کامل، بین ۰/۲ تا ۰/۸ میانجی‌گری جزئی و کمتر از ۰/۲ بیانگر عدم نقش میانجی است.

برازش کلی مدل

در تحلیل برازش مدل، مقدار SRMR برابر با ۰/۰۷۲ به دست آمده است، که نشان از برازش مناسب کلی مدل پژوهش دارد. مقدار آماره NIF نیز ۰/۸۷ به دست آمده است که نشان از برازش قابل قبول مدل دارد.



شکل ۲: مدل ساختاری پژوهش

بحث و نتیجه‌گیری

مدل مفهومی پژوهش حاضر صرفاً در پی گردآوری سه مسیر میانجی نبود، بلکه به دنبال کشف معمای تعامل هم‌زمان آن‌ها در یک بافت صنعتی واقعی بود. انسان‌نمایی ربات همکار به طور هم‌زمان باعث افزایش شایستگی درک‌شده، تهدید درک‌شده، و ایمنی درک‌شده می‌شود. این یافته به تنهایی گویای ماهیت پیچیده انسان‌نمایی در صنعت است. اما نکته محوری و تمایز این مدل، در تحلیل مسیرها آشکار می‌شود. از میان سه مسیر ادراک، تنها ایمنی درک‌شده است که به پذیرش ربات همکار منجر می‌شود.

در بافت صنعت خودروی ایران، ارزیابی خطر (ایمنی) به طور کامل بر ارزیابی فرصت (شایستگی) و ارزیابی تهدید غلبه می‌کند. به عبارت دیگر، کارکنان ایرانی، با توجه به پیشینه صنعتی، نگرانی‌های شغلی مزمن، و شاید تجارب پیشین، یک اولویت‌بندی شناختی منحصربه‌فرد دارند. آن‌ها ابتدا و بیش از هر چیز می‌پرسند آیا با این موجود جدید در امان هستیم؟ تنها پس از دریافت پاسخ مثبت به این سؤال از طریق نشانه‌های ایمنی در طراحی انسان‌نما، احتمالاً زمینه برای پذیرش فراهم می‌شود. توانایی ربات (شایستگی) در این میان، نه یک فرصت، بلکه می‌تواند به عنوان تشدیدکننده نگرانی عمل کند. این یافته نشان می‌دهد که در غیاب ایمنی درک‌شده، سایر ویژگی‌های مثبت ربات نه تنها کمکی به پذیرش نمی‌کنند، بلکه می‌توانند به دلیل ایجاد احساس جایگزین‌پذیری، اثر معکوس نیز داشته باشند.

تمایز اصلی مدل پژوهش حاضر در دو سطح ساختاری و نتیجه‌محور قابل تبیین است. نوآوری مدل پژوهش در یکپارچه‌سازی هم‌زمان سه مسیر شایستگی، تهدید و ایمنی است که با پژوهش‌های چون لیائو و همکاران (۲۰۲۳) تفاوت دارد. همچنین، در حالی که مدل‌های قبلی اغلب شایستگی را موتور محرک پذیرش می‌دانستند، مدل پژوهش حاضر به‌ویژه در بافت فرهنگی و صنعتی ایران، ایمنی درک‌شده را به عنوان مسیر غالب پذیرش ربات‌های همکار معرفی می‌کند. این یافته، نظریه‌های موجود را گسترش می‌دهد و پیشنهاد می‌کند که طراحی برای پذیرش در صنایع با نگرانی‌های شغلی بالا، ابتدا باید طراحی برای ایمنی و نه صرفاً طراحی برای شباهت یا طراحی برای عملکرد باشد. پژوهش حاضر چارچوبی برای درک این نکته فراهم می‌کند که پذیرش فناوری در محیط کار، یک معادله صرفاً فایده‌محور نیست، بلکه یک معادله بقا و ایمنی است که ابتدا باید حل شود.

پژوهش‌های پیشین عمدتاً به بررسی تأثیر مستقیم انسان‌نمایی بر پذیرش ربات‌های همکار بسنده کرده‌اند و یا نقش یک متغیر میانجی را آزموده‌اند. آنچه در این میان مغفول مانده، بررسی هم‌زمان سه سازوکار روان‌شناختی شایستگی، تهدید و ایمنی درک‌شده به عنوان مسیرهای انتقال‌دهنده اثر انسان‌نمایی بر پذیرش ربات‌های همکار

است. مدل پژوهش، پس از آزمون تجربی با رویکرد مدلسازی معادلات ساختاری، نشان‌دهنده یک ساختار یکپارچه است که اثر انسان‌نما بودن ربات‌های همکار را بر پذیرش از طریق سه مسیر میانجی شایستگی درک‌شده، تهدید درک‌شده، و ایمنی درک‌شده بررسی می‌کند (شکل ۲). برای ارزیابی اهمیت نسبی مولفه‌ها، از اندازه اثر F^2 (جدول ۹)، ضرایب مسیر β (جدول ۱۰)، و شاخص VAF (جدول ۱۱) استفاده شد. ایمنی درک‌شده بیشترین اهمیت را دارد، با ضریب مسیر $0/713$ (قوی‌ترین مسیر مستقیم به پذیرش) و اندازه اثر $0/461$ (اثر متوسط تا بزرگ)، و نقش میانجی کامل که نشان می‌دهد 85% از اثر انسان‌نما بودن بر پذیرش از این مسیر منتقل می‌شود. در مقابل، شایستگی درک‌شده با اثر کوچک و منفی و تهدید درک‌شده با اثر ناچیز اهمیت کمتری دارند و نقش میانجی آن‌ها تأیید نشد. این یافته‌ها نشان می‌دهد که در بافت صنعت خودرو ایران، تمرکز بر کاهش تهدید و افزایش شایستگی کافی نیست؛ بلکه ایمنی درک‌شده کلید اصلی ادغام موفق ربات‌ها است.

پژوهش حاضر نشان داد که در صنعت خودروسازی ایران، صرفاً انسان‌نما کردن ربات‌های همکار (شبیه‌سازی ظاهر و حرکت انسان) برای پذیرش نهایی آن‌ها توسط کارکنان کافی نیست. یافته محوری آن است که احساس ایمنی کارکنان در تعامل با ربات، تعیین‌کننده‌ترین عامل در پذیرش است و سایر ادراکات مثبت مانند شایستگی ربات، به تنهایی به پذیرش منجر نمی‌شود. به عبارت دیگر، در خطوط تولید خودرو، کارکنان وقتی ربات همکار را می‌پذیرند که نخست احساس کنند این ربات برای آن‌ها بی‌خطر است. بنابراین، استقرار موفق ربات‌های همکار در صنعت خودرو بیش از آنکه نیازمند نمایش قابلیت‌های فنی پیشرفته باشد، نیازمند ایجاد اعتماد از طریق تضمین ایمنی فیزیکی و روانی کارکنان است. در حقیقت، خروجی محوری این مدل، شناسایی ایمنی درک‌شده به عنوان سازوکار اصلی انتقال اثر انسان‌نما بودن به پذیرش است. این خروجی پیشنهاد می‌کند که مدیران صنعت خودرو، در طراحی ربات‌ها، اولویت را به ویژگی‌هایی مانند حرکات پیش‌بینی‌پذیر و پروتکل‌های ایمنی بدهند تا پذیرش کارکنان را افزایش دهند. این مدل یکپارچه، نه تنها خلأ نظری را پر می‌کند، بلکه راهنمایی عملی برای سیاست‌گذاری در صنعت 50% ارائه می‌دهد، جایی که تعامل انسان-ربات مرکزی است.

یافته‌های پژوهش حاضر با نظریه‌های مبنایی دره غیر طبیعی و تهدید بین گروهی نیز قابل بررسی است. تأیید اثر مثبت انسان‌نمایی بر ایمنی درک‌شده و نقش میانجی کامل آن در پذیرش ربات‌های همکار، گسترش نظریه دره غیر طبیعی را پیشنهاد می‌دهد. در زمینه صنعت خودرو ایران، انسان‌نمایی نه تنها احساس ناخوشایندی و عجیب بودن ایجاد می‌کند، بلکه می‌تواند با تأکید بر پیش‌بینی‌پذیری و ایمنی، این اثر منفی را تعدیل کند. این یافته، دیدگاه نظریه دره غیر طبیعی را که عمدتاً بر جنبه‌های منفی تمرکز دارد، به چالش می‌کشد و پیشنهاد می‌کند که ایمنی درک‌شده می‌تواند به عنوان یک عامل محافظ شناختی عمل کند؛ موضوعی که در مطالعات

پیشین کمتر مورد بررسی قرار گرفته است. به عبارت دیگر، این پژوهش مدل نظری را از تمرکز صرف بر ظاهر فیزیکی به سوی طراحی شناختی-ایمنی گسترش می‌دهد. همچنین، عدم تأیید نقش میانجی شایستگی درک‌شده و تهدید درک‌شده، با وجود اثر مثبت انسان‌نمایی بر این دو متغیر، نظریه تهدید بین گروهی را در زمینه صنعتی به چالش می‌کشد. یافته‌ها پیشنهاد می‌کند که در محیط‌های کاری با نگرانی‌های شغلی بالا، تهدید و شایستگی ممکن است به عنوان محرک‌های مستقل عمل کنند، اما پذیرش را به طور مستقیم کاهش ندهند؛ زیرا ایمنی درک‌شده اثر آن‌ها را تحت‌الشعاع قرار می‌دهد. این امر ایمنی را به عنوان مسیر غالب اضافه می‌کند که می‌تواند تهدید هویتی و واقعی را کاهش دهد.

باتوجه به نتایج به‌دست آمده، بین انسان نما بودن ربات همکار بر پذیرش ربات‌های همکار رابطه معناداری وجود داشت. نتیجه حاصل از این فرضیه با مطالعه‌ای که توسط لیائو و همکاران (۲۰۲۳) صورت گرفت، مطابقت ندارد. طراحی انسان‌نما در ربات‌های همکار می‌تواند تأثیر مثبتی بر پذیرش آن‌ها در محیط‌های صنعتی داشته باشد. در صنعت خودروسازی، جایی که تعامل انسان و فناوری به شکلی گسترده رخ می‌دهد، طراحی انسان‌نما به کارکنان کمک می‌کند که با ربات‌ها راحت‌تر ارتباط برقرار کنند. ویژگی‌های انسان‌نما، مانند رفتار مشابه انسان یا حرکات قابل پیش‌بینی، باعث می‌شود کارکنان احساس نزدیکی بیشتری به ربات‌ها داشته باشند و در نتیجه راحت‌تر آن‌ها را به‌عنوان بخشی از محیط کاری بپذیرند. این امر نشان‌دهنده نقش طراحی در کاهش مقاومت اولیه نسبت به فناوری‌های جدید است.

تفاوت نتایج بین پژوهش حاضر و مطالعه لیائو و همکاران (۲۰۲۳) می‌تواند به دلایل متعددی مربوط باشد که به عوامل زمینه‌ای، فرهنگی، صنعتی و طراحی پژوهش مرتبط هستند. تحقیق حاضر در صنعت خودروسازی انجام شده است، در حالی که مطالعه لیائو و همکاران در یک زمینه صنعتی دیگر انجام شده است. صنعت خودروسازی معمولاً با فناوری‌های پیشرفته و ربات‌های صنعتی آشنا است، و کارکنان ممکن است بیشتر به مزایا و کاربردهای این ربات‌ها واقف باشند. در چنین محیطی، ویژگی‌های انسان‌نما در ربات‌ها می‌تواند منجر به افزایش پذیرش شود، زیرا کارکنان این ویژگی‌ها را به‌عنوان عاملی مثبت برای بهبود همکاری و تعامل می‌بینند. پذیرش فناوری در جوامع مختلف می‌تواند به شدت تحت تأثیر عوامل فرهنگی باشد. در فرهنگ‌هایی که نوآوری و فناوری مورد استقبال قرار می‌گیرد، طراحی انسان‌نما ممکن است به‌عنوان عامل مثبتی درک شود. اما در فرهنگ‌هایی که نگرانی‌هایی مانند تهدید شغلی یا جایگزینی انسان توسط فناوری بیشتر است، ویژگی‌های انسان‌نما می‌توانند مقاومت و اضطراب ایجاد کنند. احتمالاً تفاوت فرهنگی بین جامعه مورد مطالعه و جامعه مورد بررسی لیائو و همکاران بر نتایج تأثیرگذار بوده است. در تحقیق حاضر ممکن است کارکنان ویژگی‌های انسان‌نما را به‌عنوان

عاملی برای بهبود تعامل و کاهش استرس کاری در نظر گرفته باشند. این در حالی است که در مطالعه لیائو و همکاران، احتمالاً ادراک کارکنان از انسان‌نما بودن بیشتر به سمت تهدید یا جایگزینی شغلی متمایل بوده است. این تفاوت در ادراک می‌تواند ناشی از نوع طراحی ربات‌ها، آموزش‌های ارائه‌شده، و سطح آشنایی کارکنان با این فناوری باشد. ربات‌های مورد استفاده در پژوهش حاضر ممکن است از نظر طراحی و قابلیت‌ها متفاوت از ربات‌های مورد بررسی در پژوهش لیائو و همکاران (۲۰۲۳) باشند.

یافته‌ها نشان می‌دهد که هرچه ربات‌ها از نظر طراحی انسان‌نما باشند، کارکنان بیشتر به توانمندی و قابلیت‌های آن‌ها اعتماد می‌کنند. در صنعت خودروسازی، ربات‌هایی که رفتار و ظاهرشان مشابه انسان است، ممکن است توسط کارکنان به‌عنوان ابزارهایی توانمند و قابل اعتماد دیده شوند. این شایستگی درک‌شده شامل قابلیت‌هایی همچون دقت در انجام وظایف، قابلیت پیش‌بینی رفتار، و هماهنگی با نیازهای کاری است. در نتیجه، طراحی انسان‌نما تأثیر مثبتی بر اعتبار و قابلیت درک‌شده ربات‌ها دارد، هرچند این تأثیر ممکن است در پذیرش کلی پیچیدگی‌هایی ایجاد کند. نتیجه فرضیه با مطالعه لیائو و همکاران (۲۰۲۳) مطابقت دارد.

تأثیر انسان‌نما بودن ربات همکار بر تهدید درک‌شده نیز تأیید شده است. یافته‌ها نشان می‌دهد که طراحی انسان‌نما، در کنار افزایش تعامل‌پذیری، می‌تواند باعث افزایش احساس تهدید در کارکنان شود. در صنعت خودروسازی، ربات‌هایی که بیش‌ازحد مشابه انسان به نظر می‌رسند، ممکن است به‌عنوان تهدیدی برای امنیت شغلی، هویت حرفه‌ای، و جایگاه انسانی کارکنان تلقی شوند. قابلیت‌هایی مانند دقت بالا، توانایی انجام وظایف پیچیده، و عدم نیاز به استراحت، ممکن است نگرانی‌هایی را در مورد جایگزینی انسان‌ها ایجاد کند. این یافته نشان می‌دهد که طراحی انسان‌نما علاوه بر مزایای تعامل‌پذیری، می‌تواند چالش‌هایی روان‌شناختی نیز ایجاد کند. نتیجه حاصل از این فرضیه با مطالعه لیائو و همکاران (۲۰۲۳) و کلبه و فیس (۲۰۲۵) مطابقت دارد. همچنین، فرضیه تأثیر انسان‌نما بودن ربات همکار بر ایمنی درک‌شده تأیید شده است. طراحی انسان‌نما می‌تواند حس ایمنی کارکنان را تقویت کند، به‌ویژه زمانی که ربات‌ها رفتارهایی مانند حرکات روان و پیش‌بینی‌پذیر دارند. نتیجه حاصل از این فرضیه با مطالعه بامگارتن و همکاران (۲۰۲۲) مطابقت دارد.

فرضیه تأثیر شایستگی درک‌شده بر پذیرش ربات‌های همکار رد شده است. با وجود اینکه طراحی انسان‌نما شایستگی درک‌شده ربات‌ها را افزایش می‌دهد، این شایستگی به‌تنهایی منجر به پذیرش آن‌ها نمی‌شود. در صنعت خودروسازی، ممکن است کارکنان ربات‌های بسیار توانمند را به‌عنوان تهدیدی برای موقعیت شغلی خود ببینند. این یافته نشان می‌دهد که حتی اگر ربات‌ها توانمند به نظر برسند، احساس تهدید یا از دست دادن کنترل می‌تواند پذیرش آن‌ها را کاهش دهد. نتیجه حاصل از این فرضیه با مطالعه لیائو و همکاران (۲۰۲۳) مطابقت دارد. اگرچه

شایستگی بالای ربات‌ها می‌تواند مزایای عملی مانند بهبود بهره‌وری و کارایی را نشان دهد، اما در عین حال می‌تواند باعث ایجاد احساس تهدید در میان کارکنان شود. در صنعت خودروسازی، که اغلب به نیروی انسانی ماهر متکی است، کارکنان ممکن است توانمندی بالای ربات‌ها را به‌عنوان تهدیدی برای امنیت شغلی خود تلقی کنند. در محیط‌های کارگاهی با نگرانی مزمن درباره امنیت شغلی، هر چه ربات شایسته‌تر و مستقل‌تر ادراک شود، معنای ضمنی برای کارکنان این است که سازمان کمتر به مهارت انسانی نیاز دارد. این برداشت می‌تواند حس جایگزین‌پذیری، از دست دادن کنترل بر فرایند کار، و تضعیف هویت حرفه‌ای را فعال کند. نتیجه این می‌شود که شایستگی بالا به جای جذابیت، تهدید نمادین تولید می‌کند و پذیرش را کاهش می‌دهد، مخصوصاً وقتی سیاست‌های رسمی درباره حفظ شغل، بازآموزی و مسیرهای مهارتی روشن نباشد. این تبیین دقیقاً با واقعیت‌های رایج در کارخانه‌های ایران هماهنگ است. کارکنان باهوش هستند و وقتی ربات خیلی خوب کار می‌کند، اولین سوال آن‌ها این نیست که چقدر بهره‌وری بالا می‌رود، سوال این است که نوبت حذف من چه زمانی است. احساس تهدید می‌تواند پذیرش فناوری را کاهش دهد و اثر مثبت شایستگی درک‌شده را خنثی کند. نزدیکی آماره آزمون به ۱/۹۶ نشان می‌دهد رابطه مورد بررسی احتمالاً به عوامل دیگری وابسته است که در مدل فعلی لحاظ نشده‌اند. به‌عنوان مثال، متغیرهایی مانند حمایت سازمانی، آموزش کارکنان، یا سیاست‌های تضمین شغلی ممکن است نقش تعدیل‌کننده‌ای داشته باشند که در صورت در نظر گرفتن آن‌ها، رابطه بین شایستگی درک‌شده و پذیرش ربات‌ها معنادار شود. همچنین، ممکن است کارکنان توانمندی بالای ربات‌ها را نه به‌عنوان یک فرصت برای کمک به انجام وظایف خود، بلکه به‌عنوان تهدیدی برای کاهش کنترل بر فرآیندهای کاری یا جایگزینی خود تلقی کنند. این نوع ادراک می‌تواند در برخی محیط‌های صنعتی، مانند صنعت خودروسازی، تأثیرگذاری شایستگی درک‌شده را محدود کند. اگر سیاست‌های سازمانی واضحی برای کاهش نگرانی‌های کارکنان از جایگزینی شغلی وجود داشته باشد، یا اگر آموزش‌هایی برای نشان دادن نحوه همکاری موثر با ربات‌ها ارائه شود، این متغیر می‌تواند به یک عامل مثبت برای پذیرش تبدیل شود.

فرضیه تأثیر تهدید درک‌شده بر پذیرش ربات‌های همکار رد شده است. برخلاف انتظار، احساس تهدید مستقیم از سوی ربات‌ها تأثیری معنادار بر پذیرش آن‌ها نداشته است. در صنعت خودروسازی، ممکن است کارکنان به دلایل مختلفی مانند حمایت‌های سازمانی یا رویکردهای مثبت مدیریتی، تهدید را کم‌اهمیت تلقی کنند یا توانایی کنار آمدن با آن را داشته باشند. تأثیر ایمنی درک‌شده بر پذیرش ربات‌های همکار تأیید شده است. حس ایمنی تأثیر قابل‌توجهی در پذیرش ربات‌های همکار دارد. کارکنانی که احساس می‌کنند تعامل با ربات‌ها ایمن است، راحت‌تر آن‌ها را می‌پذیرند.

برخی از روابط میانجی مانند نقش شایستگی و تهدید درک شده در رابطه میان انسان نما بودن و پذیرش، تأیید نشدند. تأثیر انسان نما بودن ربات همکار بر پذیرش ربات های همکار از طریق شایستگی درک شده رد شده است. اگرچه طراحی انسان نما شایستگی درک شده را افزایش می دهد، شایستگی به تنهایی نمی تواند نقش میانجی در پذیرش ایفا کند. کارکنان ممکن است توانمندی ربات ها را تهدیدی برای جایگاه خود بدانند، که این مسئله پذیرش آن ها را کاهش می دهد. همچنین، تأثیر انسان نما بودن ربات همکار بر پذیرش ربات های همکار از طریق تهدید درک شده رد شده است. طراحی انسان نما با افزایش احساس تهدید، تأثیر معناداری بر پذیرش نداشته است که نشان می دهد که احساس تهدید، اگرچه وجود دارد، نمی تواند به تنهایی بر پذیرش کارکنان تأثیر بگذارد. اما، تأثیر انسان نما بودن ربات همکار بر پذیرش ربات های همکار از طریق ایمنی درک شده تأیید شده است. از این رو، احساس ایمنی یک عامل کلیدی در پذیرش ربات های همکار است. طراحی انسان نما با ایجاد حس ایمنی، می تواند کارکنان را به پذیرش ربات ها ترغیب کند.

به عنوان کاربردهای مدیریتی پژوهش حاضر برای صنعت خودرو ایران می توان به موارد زیر اشاره نمود. یافته های پژوهش نشان می دهد که برای پذیرش موفق ربات های همکار در خطوط تولید خودرو، مدیران باید طراحی تجربه ایمن را بر طراحی انسان نما اولویت دهند. به این منظور، پیشنهاد می شود پیش از استقرار گسترده ربات ها در سالن های تولید، پروتکل های ایمنی به صورت شفاف و قابل مشاهده برای کارگران اجرا شود؛ از جمله خط کشی محدوده حرکتی ربات، نصب سیستم توقف اضطراری در دسترس، و آموزش عملی سناریوهای مواجهه با خطر. همچنین به جای تأکید صرف بر هوشمندی و شایستگی ربات (که می تواند احساس تهدید شغلی ایجاد کند)، در ارتباط با کارکنان بر نقش کمکی ربات در کاهش کارهای سخت و خطرناک و ایمنی بالای تعامل با آن تمرکز شود. استفاده از اپراتورهای باتجربه خط تولید به عنوان سفیران ایمنی که تجربه کار با ربات را به زبان ساده به همکاران منتقل می کنند، نسبت به بخشنامه های رسمی اثربخشی بیشتری در فرهنگ کاری صنعت خودرو ایران خواهد داشت. به عنوان گام های عملیاتی، ابتدا پایلوت کوچک در ایستگاه های کم ریسک با ۱-۲ کو-ربات آغاز شود. سرعت حرکات اولیه کاهش یابد و با نشانه های بصری ساده پیش بینی پذیری افزایش یابد. آموزش از طریق کارگران سفیر (۲-۳ نفر از هر شیفت) انجام شود تا اعتماد همکاران جلب گردد. همزمان، بیانیه رسمی تضمین عدم جایگزینی شغلی صادر شود. موفقیت با شاخص های ایمنی ادراک شده و تعداد توقف های اضطراری سنجیده شود. این رویکرد با محدودیت های بودجه ای و فرهنگی ایران سازگار است و می تواند پذیرش را بدون هزینه بالا افزایش دهد.

این پژوهش چند محدودیت دارد که باید در تفسیر نتایج لحاظ شود. نخست، داده‌ها به صورت مقطعی جمع‌آوری شده‌اند، بنابراین استنباط علی از مسیرهای مدل با احتیاط همراه است و ممکن است ادراک کارکنان پس از تجربه واقعی و طولانی‌تر کار با ربات‌های همکار تغییر کند. دوم، سنجش همه سازه‌ها با پرسشنامه خودگزارشی انجام شده است؛ در نتیجه احتمال تورش روش مشترک وجود دارد و پیشنهاد می‌شود در مطالعات آتی علاوه بر ادراکات، شاخص‌های رفتاری پذیرش مانند میزان استفاده واقعی، سطح تعامل انسان و ربات، و شاخص‌های ایمنی ثبت‌شده نیز وارد مدل شود. سوم، نمونه پژوهش با نمونه‌گیری در دسترس انتخاب شده است، لذا تعمیم یافته‌ها به سایر صنایع یا حتی شرکت‌های خودرویی با ساختار نیروی انسانی متفاوت باید با احتیاط انجام گیرد. همچنین، مدل پژوهش تنها سه میانجی را بررسی کرده و متغیرهای دیگر مانند اعتماد به فناوری، خودکارآمدی استفاده از ربات، حمایت سازمانی، یا تجربه قبلی با اتوماسیون در مدل گنجانده نشده‌اند که این امر ممکن است بخشی از واریانس پذیرش را توضیح ندهد. در ادامه، برای توسعه مدل، پیشنهاد می‌شود نقش تعدیل‌کننده اعتماد به فناوری آزمون شود تا مشخص گردد آیا رابطه شایستگی درک شده و پذیرش در میان کارکنانی با اعتماد بالاتر به فناوری تقویت و در میان کارکنان کم اعتماد تضعیف یا حتی معکوس می‌شود؛ زیرا اعتماد به فناوری می‌تواند نگرانی‌های مرتبط با جایگزینی شغلی و از دست دادن کنترل را کاهش دهد و شایستگی را از یک محرک تهدیدآمیز به یک عامل فرصت‌ساز تبدیل کند، ضمن آنکه می‌توان این تعدیل‌گر را در کنار خودکارآمدی یا تجربه قبلی کار با اتوماسیون نیز بررسی کرد. علاوه بر این، برای تحقیقات آتی در زمینه پذیرش ربات‌های همکار، پیشنهاد می‌شود مطالعات تطبیقی در صنایع مختلف انجام شود تا تفاوت‌های زمینه‌ای در پذیرش ربات‌ها بررسی گردد. همچنین، تحلیل نقش فرهنگ سازمانی، آموزش و آمادگی کارکنان می‌تواند به درک عمیق‌تری از عوامل تأثیرگذار بر پذیرش کمک کند. بررسی جنبه‌های روان‌شناختی و اجتماعی تعامل با ربات‌های انسان‌نما و ارزیابی تأثیر ویژگی‌های فنی ربات‌ها، مسیرهای پژوهشی مهمی را پیش روی محققان قرار می‌دهد.

منابع

- AlNuaimi, B. K., Khan, M., & Ajmal, M. M. (2021). The role of big data analytics capabilities in greening e-procurement: A higher order PLS-SEM analysis. *Technological Forecasting and Social Change*, 169, 120808. <https://doi.org/10.1016/j.techfore.2021.120808>
- Appel, M., Izydorczyk, D., Weber, S., Mara, M., & Lischetzke, T. (2020). The uncanny of mind in a machine: Humanoid robots as tools, agents, and experiencers. *Computers in Human Behavior*, 102, 274-286.
- Baumgartner, M., Kopp, T., & Kinkel, S. (2022). Analysing factory workers' acceptance of collaborative robots: a web-based tool for company representatives. *Electronics*, 11(1), 145.
- Cohen, Y., Shoval, S., Faccio, M., & Minto, R. (2022). Deploying cobots in collaborative systems: major considerations and productivity analysis. *International Journal of Production Research*, 60(6), 1815-1831.
- Cusano, N. (2023). Cobot and Sobot: for a new ontology of collaborative and social robots. *Foundations of Science*, 28(4), 1143-1155.
- Esmacili, A., Ghazinoori, S., Naghizadeh, M., Bamdad Soufi, J. and Manteghi, M. (2021). Operational Capabilities as One of the Preconditions for Presence in the Global Supply Network of the Autoparts Industry; A Multi-case Analysis. *Journal of Technology Development Management*, 9(1), 95-133. doi: 10.22104/jtdm.2021.4661.2750 [In Persian].
- Gualtieri, L., Rauch, E., & Vidoni, R. (2021). Emerging research fields in safety and ergonomics in industrial collaborative robotics: A systematic literature review. *Robotics and Computer-Integrated Manufacturing*, 67, 101998.
- Guertler, M., Tomidei, L., Sick, N., Carmichael, M., Paul, G., Wambsganss, A., ... & Hussain, S. (2023). When is a robot a cobot? Moving beyond manufacturing and arm-based cobot manipulators. *Proceedings of the Design Society*, 3, 3889-3898.
- Hair Jr, J. F., Hult, G. T. M., Ringle, C. M., Sarstedt, M., Danks, N. P., Ray, S., ... & Ray, S. (2021). An introduction to structural equation modeling. *Partial least squares structural equation modeling (PLS-SEM) using R: a workbook*, 1-29. https://doi.org/10.1007/978-3-030-80519-7_1
- Javaid, M., Haleem, A., Singh, R. P., Rab, S., & Suman, R. (2022). Significant applications of Cobots in the field of manufacturing. *Cognitive Robotics*, 2, 222-233.
- Khatami Firouzabadi, S. M. A., Tabatabaeian, S. H. and Dashti, M. A. (2019). Key success factors in the process of implementing advanced technologies in the industrial firms: Pieces of evidence

- from the automotive industry in Iran. *Journal of Technology Development Management*, 6(4), 89-126. doi: 10.22104/jtdm.2019.3000.2013 [In Persian].
- Klebbe, R., & Frieese, C. (2025). Exploring the Role of Robots in Inpatient Care: Caregivers' Perspectives on the Development and Evaluation of Collaborative Robot Applications—A Qualitative Research Approach. *International Journal of Social Robotics*, 1-18.
- Kopp, T., Baumgartner, M., & Kinkel, S. (2022). How linguistic framing affects factory workers' initial trust in collaborative robots: The interplay between anthropomorphism and technological replacement. *International Journal of Human-Computer Studies*, 158, 102730.
- Kraus, J., Miller, L., Klumpp, M., Babel, F., Scholz, D., Merger, J., & Baumann, M. (2024). On the role of beliefs and trust for the intention to use service robots: an integrated trustworthiness beliefs model for robot acceptance. *International Journal of Social Robotics*, 16(6), 1223-1246.
- Law, L., & Fong, N. (2020). Applying partial least squares structural equation modeling (PLS-SEM) in an investigation of undergraduate students' learning transfer of academic English. *Journal of English for Academic Purposes*, 46, 100884. <https://doi.org/10.1016/j.jeap.2020.100884>.
- Leichtmann, B., Hartung, J., Wilhelm, O., & Nitsch, V. (2023). New short scale to measure workers' attitudes toward the implementation of cooperative robots in industrial work settings: Instrument development and exploration of attitude structure. *International Journal of Social Robotics*, 15(6), 909-930.
- Liao, S., Lin, L., & Chen, Q. (2023). Research on the acceptance of collaborative robots for the industry 5.0 era--the mediating effect of perceived competence and the moderating effect of robot use self-efficacy. *International Journal of Industrial Ergonomics*, 95, 103455.
- Liao, S., Lin, L., Chen, Q., & Pei, H. (2024). Why not work with anthropomorphic collaborative robots? The mediation effect of perceived intelligence and the moderation effect of self-efficacy. *Human Factors and Ergonomics in Manufacturing & Service Industries*, 34(3), 241-260.
- Liao, S., Chen, C., Yao, Y., & Chen, Q. (2024). Would You be Willing to Share Knowledge with a Collaborative Robot? The Mediating Effect of Perceived Agency and the Moderating Effect of Identity Threat. *International Journal of Human-Computer Interaction*, 1-26.
- Liu, L., Zou, Z., & Greene, R. L. (2024). The effects of type and form of collaborative robots in manufacturing on trustworthiness, risk perceived, and acceptance. *International Journal of Human-Computer Interaction*, 40(10), 2697-2710.
- Lutin, E., Elprama, S. A., Cornelis, J., Leconte, P., Van Doninck, B., Witters, M., ... & Jacobs, A. (2024). Pilot Study on the Relationship Between Acceptance of Collaborative Robots and Stress. *International Journal of Social Robotics*, 16(6), 1475-1488.

- Paliga, M. (2022). Human–cobot interaction fluency and cobot operators’ job performance. The mediating role of work engagement: A survey. *Robotics and Autonomous Systems*, 155, 104191.
- Prassida, G. F., & Asfari, U. (2022). A conceptual model for the acceptance of collaborative robots in industry 5.0. *Procedia Computer Science*, 197, 61-67.
- Rashmei, M. H., Hosseini Shakib, M. and Khamseh, A. (2023). The pattern of Industry 4.0 technologies adoption in Iran's auto parts manufacturing companies. *Journal of Technology Development Management*, 11(3), 169-210. doi: 10.22104/jtdm.2024.6743.3274 [In Persian].
- Taesi, C., Aggogeri, F., & Pellegrini, N. (2023). COBOT applications—recent advances and challenges. *Robotics*, 12(3), 79.
- Talaie, H. (2025). The Role of Industry 5.0 in Mitigating Supply Chain Risks and Boosting Performance: A Pathway to Sustainable Competitive Advantage. *Strategic Value Chain Management*, 2(4), 55-80. doi: 10.22075/svcm.2025.39107.1052 [In Persian].