



Structural Modeling of Responsible Artificial Intelligence-Based Human Resource Management and Its Implications for Employee Flourishing

Abbasali Rastgar *

*Corresponding Author, Prof., Department of Management, Faculty of Economics, Management and Administrative Sciences, Semnan University, Semnan, Iran. E-mail: a_rastgar@semnan.ac.ir

Minasadat Mousavi 

Ph.D., Department of Management, Faculty of Economics, Management and Administrative Sciences, Semnan University, Semnan, Iran. E-mail: minamousavi@semnan.ac.ir

Abstract

Objective

This study aimed to identify, explain, and model the structural and hierarchical relationships between the components of responsible Artificial Intelligence (AI) in Human Resource Management (HRM) and their subsequent outcomes on employee flourishing within digital organizations. The core research problem addresses a critical gap in the extant literature: with the rapid integration of algorithms and intelligent decision-support systems into HRM, most organizations have disproportionately focused on technical efficiency, processing speed, and cost reduction, thereby neglecting the ethical dimensions and human consequences of these technologies. This study seeks to demonstrate how principles such as algorithmic transparency, accountability, fairness, privacy, and psychological safety can lead to positive work experiences and optimal employee outcomes within the context of human-machine interaction. Consequently, this research aims to explain how abstract AI ethical principles translate into psychological and behavioral outcomes for employees, thereby bridging the theoretical gap between normative AI principles and sustainable HRM practices.

Methods

Utilizing an exploratory-sequential mixed-methods design, this study was conducted within a pragmatic research paradigm. In the qualitative phase, semi-structured interviews were conducted with 17 dual-domain experts, comprising HRM specialists and AI or digital transformation experts, who were selected through purposive criterion and snowball

Citation: Rastgar, Abbasali & Mousavi, Minasadat (2026). Structural Modeling of Responsible Artificial Intelligence-Based Human Resource Management and Its Implications for Employee Flourishing. *Journal of Public Administration*, 18(3), 775-804 (in Persian)



sampling techniques. The qualitative data were analyzed using thematic analysis to extract relevant themes and dimensions. In the second (quantitative) phase, Interpretive Structural Modeling (ISM) was employed to map the hierarchical structure and interactions among the identified components. Prior to constructing the Structural Self-Interaction Matrix (SSIM), the dimensions and themes extracted from the qualitative phase were re-evaluated and validated by an expert panel in order to prevent the imposition of researcher bias and ensure the credibility of the findings. Furthermore, the analysis was grounded in Socio-Technical Systems (STS) theory to elucidate the alignment between the technical subsystem, comprising algorithms and infrastructure, and the social subsystem, encompassing humans and organizational interactions, in the implementation of responsible AI.

Results

The qualitative phase identified 15 sub-themes, which were subsequently categorized under the broader dimensions of responsible AI and employee flourishing. The ISM analysis revealed that implementing responsible AI in HR processes is not a one-dimensional technological intervention, but rather a multi-layered, systemic process structured across eight hierarchical levels. At the foundational levels, namely, the governance and technical layers, variables such as ethical policy-making, robust data infrastructure, and fair algorithmic design function as primary drivers. These foundational variables subsequently influence the middle layers, corresponding to the socio-psychological subsystem, by fostering organizational trust, enhancing perceived procedural justice, and mitigating ethical risks. Ultimately, at the highest level of the model, designated as the outcome subsystem, employee flourishing is realized through constructs such as meaning at work, active engagement, personal growth, and optimal performance. The findings collectively confirm that social and ethical components play a crucial mediating role in translating technical capabilities into meaningful human outcomes.

Conclusion

By proposing a three-layered structural-interpretive model, comprising governing-technical, socio-psychological, and outcome layers, this study addresses a significant gap in the literature at the intersection of responsible AI and HRM. The results demonstrate that the success and sustainability of intelligent technologies in organizations rely fundamentally on the joint optimization of both technical and social subsystems. In the absence of adequate attention to the psychological and ethical needs of employees, AI implementation may trigger work alienation and decrease organizational commitment, thereby undermining the very objectives it seeks to achieve. The proposed model serves as a strategic roadmap for organizational policymakers, HR managers, and system developers seeking to design human-centered technological solutions and promote human resource flourishing in the era of digital transformation. These insights offer both theoretical contributions and practical guidance for organizations striving to balance technological advancement with ethical responsibility and employee well-being.

Keywords: Responsible artificial intelligence, Employee thriving, Transparency and accountability, Organizational trust, Socio-technical systems theory.



مدل سازی ساختاری مدیریت منابع انسانی مبتنی بر هوش مصنوعی مسئولانه و پیامدهای آن بر شکوفایی کارکنان

عباسعلی رستگار*

* نویسنده مسئول، استاد، گروه مدیریت، دانشکده اقتصاد، مدیریت و علوم اداری، دانشگاه سمنان، سمنان، ایران. رایانامه: a_rastgar@semnan.ac.ir

میناسادات موسوی

دکتری، گروه مدیریت، دانشکده اقتصاد، مدیریت و علوم اداری، دانشگاه سمنان، سمنان، ایران. رایانامه: minamousavi@semnan.ac.ir

چکیده

هدف: پژوهش حاضر با هدف شناسایی، تبیین و مدل سازی روابط ساختاری و سلسله مراتبی میان مؤلفه های مدیریت منابع انسانی مبتنی بر هوش مصنوعی مسئولانه و پیامدهای آن بر شکوفایی کارکنان در بستر سازمان های دیجیتال انجام شد. مسئله اصلی پژوهش این است که با ورود پُرشتاب الگوریتم ها و سامانه های تصمیم یار هوشمند به قلمرو مدیریت منابع انسانی، بخش عمده ای از سازمان ها، فقط بر کارایی فنی، سرعت و کاهش هزینه های عملیاتی متمرکز شده و از ابعاد اخلاقی و پیامدهای انسانی این فناوری ها غفلت ورزیده اند. این مطالعه به دنبال آن است تا نشان دهد که چگونه اصولی مانند شفافیت الگوریتمی، پاسخ گوئی، عدالت، حریم خصوصی و امنیت روان شناختی، می توانند در تعامل انسان - ماشین، به ارتقای تجربه مثبت کاری و عملکرد بهینه کارکنان منجر شوند. در واقع، این پژوهش تلاش می کند تا نحوه انتقال اصول اخلاقی و فنی هوش مصنوعی به پیامدهای روان شناختی و رفتاری کارکنان را تبیین کند و با ارائه چارچوبی منسجم، به پر کردن شکاف نظری موجود میان اصول انتزاعی هوش مصنوعی مسئولانه و مدیریت منابع انسانی پایدار یاری دهد.

روش: این پژوهش از نظر جهت گیری کاربردی و از نظر استراتژی اجرا، با رویکرد آمیخته اکتشافی - تبیینی (متوالی) در قالب پارادایم پراگماتیسم انجام شد. در مرحله کیفی، به منظور استخراج مؤلفه ها و ابعاد اولیه، مصاحبه های نیمه ساختاریافته با ۱۷ نفر از خبرگان دوحوزه ای (متخصصان مدیریت منابع انسانی و متخصصان هوش مصنوعی و تحول دیجیتال) که به روش نمونه گیری هدفمند معیار محور و گلوله برفی انتخاب شده بودند، انجام گرفت. تحلیل داده های کیفی با استفاده از روش تحلیل مضمون تأملی انجام شد. در مرحله دوم (کمی)، برای کشف روابط ساختاریافته و سطح بندی مؤلفه ها، تکنیک مدل سازی ساختاری تفسیری به کار گرفته شد. ابعاد و مضامین استخراج شده از مرحله کیفی، پیش از ورود به ماتریس خودتأملی ساختاری، در پیل خبرگان بار دیگر ارزیابی و تأیید شدند تا از عدم تحمیل پیش فرض های پژوهشگر اطمینان حاصل شود. همچنین تحلیل ها با تکیه بر نظریه سیستم های اجتماعی - فنی صورت گرفت تا پیوند میان زیرسیستم فنی (الگوریتم ها و زیرساخت ها) و زیرسیستم اجتماعی (انسان ها و تعاملات سازمانی) در استقرار هوش مصنوعی مسئولانه مشخص شود.

استناد: رستگار، عباسعلی و موسوی، میناسادات (۱۴۰۵). مدل سازی ساختاری مدیریت منابع انسانی مبتنی بر هوش مصنوعی مسئولانه و پیامدهای آن بر شکوفایی کارکنان. مدیریت دولتی، ۱۸ (۳)، ۷۷۵-۸۰۴.

یافته‌ها: یافته‌های بخش کیفی نشان‌دهنده شناسایی ۱۵ مضمون فرعی در قالب ابعاد هوش مصنوعی مسئولانه و شکوفایی کارکنان بود. نتایج حاصل از تحلیل ساختاری تفسیری نشان داد که استقرار هوش مصنوعی مسئولانه در فرایندهای منابع انسانی، نه یک مداخله تک‌بعدی فناورانه، بلکه فرایندی چندلایه و سیستماتیک است که در ۸ سطح سلسله‌مراتبی سازمان‌دهی می‌شود. در لایه‌های زیرین و ریشه‌ای مدل (سطوح حاکمیتی و فنی)، متغیرهایی مانند خط‌مشی‌گذاری اخلاقی، زیرساخت‌های داده‌ای و طراحی الگوریتمی منصفانه قرار دارند که به‌عنوان محرک‌های اصلی عمل می‌کنند. این متغیرها در لایه‌های میانی (زیرسیستم اجتماعی - روان‌شناختی) از طریق تقویت اعتماد سازمانی، ادراک عدالت رویه‌ای و کاهش ریسک‌های اخلاقی، بر لایه‌های بالایی اثر می‌گذارند. در نهایت، در بالاترین سطح مدل (زیرسیستم پیامدی)، شکوفایی کارکنان در قالب سازه‌هایی چون معنا در کار، درگیری فعال، رشد فردی و عملکرد بهینه محقق می‌شود. یافته‌ها مؤید آن است که مؤلفه‌های اجتماعی و اخلاقی، در انتقال اثرهای قابلیت‌های فنی به پیامدهای انسانی نقش کلیدی و واسطه‌ای ایفا می‌کنند.

نتیجه‌گیری: این پژوهش با ارائه یک مدل ساختاری - تفسیری سه‌لایه (حاکمیتی - فنی، اجتماعی - روان‌شناختی و پیامدی)، خلأ موجود در ادبیات هوش مصنوعی مسئولانه و مدیریت منابع انسانی را پوشش می‌دهد. نتایج پژوهش نشان می‌دهد که موفقیت و پایداری فناوری‌های هوشمند در سازمان‌ها، مستلزم هم‌راستایی و بهینه‌سازی مشترک زیرسیستم‌های فنی و اجتماعی است. بدون توجه به نیازهای روان‌شناختی و اخلاقی کارکنان، پیاده‌سازی هوش مصنوعی می‌تواند به کاهش تعهد و احساس بیگانگی شغلی منجر شود. مدل پیشنهادی این پژوهش، به‌عنوان یک نقشه راهبردی، می‌تواند توسط سیاست‌گذاران سازمانی، مدیران منابع انسانی و توسعه‌دهندگان سیستم‌های هوشمند برای طراحی راه‌کارهای فناورانه انسان‌محور و ارتقای شکوفایی منابع انسانی در عصر تحول دیجیتال استفاده شود.

کلیدواژه‌ها: هوش مصنوعی مسئولانه، شکوفایی کارکنان، شفافیت و پاسخ‌گویی، اعتماد سازمانی، نظریه سیستم‌های اجتماعی - فنی.

مقدمه

در سال‌های اخیر، یکپارچه‌سازی هوش مصنوعی در فرایندهای مدیریت منابع انسانی، تحولات چشمگیری را در حوزه‌هایی مانند جذب و انتخاب، ارزیابی عملکرد، توسعه کارکنان و تصمیم‌گیری‌های مدیریتی ایجاد کرده است. هوش مصنوعی با قابلیت تحلیل کلان‌داده‌ها، شناسایی الگوهای رفتاری و پشتیبانی از تصمیم‌گیری، کارایی، سرعت و دقت بسیاری از فرایندهای منابع انسانی را ارتقا داده (مؤمنی، یعقوبی، روشن و پورعزت، ۱۴۰۵) و به تدریج، از یک ابزار عملیاتی، به یکی از مؤلفه‌های راهبردی مدیریت منابع انسانی تبدیل شده است (بودوار، مالک، دی‌سیلوا و تتویسوتان^۱، ۲۰۲۳). با این حال، گسترش استفاده از سامانه‌های الگوریتمی، چالش‌های اخلاقی، مدیریتی و اجتماعی جدیدی را نیز به همراه داشته است. مطالعات اخیر نشان می‌دهند که کاهش شفافیت در تصمیم‌گیری، بازتولید سوگیری‌های الگوریتمی، نقض حریم خصوصی، محدود شدن قابلیت توضیح‌پذیری تصمیم‌ها و کاهش نقش قضاوت انسانی، از مهم‌ترین چالش‌های استقرار هوش مصنوعی در مدیریت منابع انسانی محسوب می‌شوند. در ادبیات پژوهش، این چالش‌ها با ادراک کارکنان از سامانه‌های هوشمند به عنوان فرایندهایی غیرشفاف، فهم‌ناپذیر و تهدیدکننده استقلال فردی همراه گزارش شده‌اند (بانکینز، فرموسا و اوکین^۲، ۲۰۲۳؛ لیچ دوبالد و همکاران^۳، ۲۰۲۲).

در پاسخ به این چالش‌ها، مفهوم «هوش مصنوعی مسئولانه» به عنوان پارادایمی نوین برای تضمین اخلاق‌مداری، شفافیت، پاسخ‌گویی، عدالت، حفظ حریم خصوصی و نظارت انسانی در طراحی، توسعه و به کارگیری سامانه‌های هوشمند مطرح شده است (استال و همکاران^۴، ۲۰۲۳). در سال‌های اخیر، چارچوب‌های معتبر بین‌المللی از جمله چارچوب مدیریت ریسک هوش مصنوعی مؤسسه ملی استاندارد و فناوری آمریکا، قانون هوش مصنوعی اتحادیه اروپا^۵، اصول هوش مصنوعی مسئولانه گوگل^۶ و چارچوب استاندارد هوش مصنوعی مسئولانه مایکروسافت^۷ نیز بر اصولی همچون عدالت، شفافیت، پاسخ‌گویی، مدیریت ریسک، قابلیت توضیح‌پذیری و نظارت انسانی تأکید کرده‌اند. با وجود اشتراک این چارچوب‌ها در معرفی اصول حاکم بر توسعه مسئولانه هوش مصنوعی، تمرکز آن‌ها عمدتاً بر الزامات سیاستی، حاکمیتی و طراحی سامانه‌های هوشمند است و درباره نحوه سازمان‌یافتگی این اصول در بستر مدیریت منابع انسانی و روابط ساختاری میان آن‌ها، تبیین محدودی ارائه شده است. از این رو، با وجود تکرار این چارچوب‌های هنجاری، فقدان یک مدل ساختاری که بتواند این اصول را به تجربه زیسته کارکنان در محیط کار پیوند بزند، به شکافی بنیادین در ادبیات مدیریت منابع انسانی بدل شده است. این پژوهش با تمرکز بر همین شکاف، به واکاوی روابط ساختاری حاکم بر این مؤلفه‌ها می‌پردازد.

1. Budhwar, Malik, De Silva & Thevisuthan

2. Bankins, Formosa & O'Kane

3. Leicht-Deobald et al.

4. Stahl et al.

5. EU AI Act

6. Google Responsible AI

7. Microsoft Responsible AI Standard

استقرار هوش مصنوعی مسئولانه در مدیریت منابع انسانی، نه تنها ریسک‌های اخلاقی را کاهش می‌دهد، بلکه با ایجاد بستر اعتماد و ادراک عدالت رویه‌ای، می‌تواند نقش کلیدی در بهبود پیامدهای مثبت روان‌شناختی و عملکردی کارکنان، از جمله «شکوفایی کارکنان» ایفا کند (حاجی اسماعیلی، طهماسبی و باباشاهی، ۱۴۰۵؛ مالیک، بودوار و کاظمی^۱، ۲۰۲۳). شکوفایی کارکنان که بر اساس دیدگاه اسپریتزر، ساتکلیف، داتون، سوننشاین و گرانت^۲ (۲۰۰۵)، به تجربه هم‌زمان بهزیستی روان‌شناختی و عملکرد اثربخش در محیط کار اشاره دارد، در پژوهش حاضر فقط به‌عنوان یک سازه نظری در نظر گرفته نشده است؛ بلکه ابعاد آن در مرحله کیفی پژوهش از طریق تحلیل مضمون مصاحبه‌های نیمه‌ساختاریافته با خبرگان شناسایی، پالایش و اعتبارسنجی شده و سپس به‌عنوان ورودی مرحله مدل‌سازی ساختاری - تفسیری مورد استفاده قرار گرفته است. از این منظر، شکوفایی کارکنان مفهومی چندبعدی است که بایستی در بستر سازمان‌های دیجیتال و در تعامل میان فناوری، ساختارهای سازمانی و تجربه کارکنان مطالعه شود.

مرور پیشینه پژوهش نشان می‌دهد که با وجود گسترش مطالعات مرتبط با کاربرد هوش مصنوعی در مدیریت منابع انسانی، بخش عمده‌ای از پژوهش‌های پیشین بر مزایای فنی و اقتصادی این فناوری، از جمله افزایش کارایی، بهبود تصمیم‌گیری و خودکارسازی فرایندها، یا بر شناسایی چالش‌های اخلاقی ناشی از به‌کارگیری آن، نظیر سوگیری الگوریتمی، نقض حریم خصوصی و شفافیت تصمیمات، متمرکز بوده‌اند (چودوری و همکاران^۳، ۲۰۲۳؛ موسوی، رستگار و شفیع نیک‌آبادی، ۱۴۰۴؛ رستگار و موسوی، ۱۴۰۲). اگرچه این مطالعات شناخت ارزشمندی از ظرفیت‌ها و چالش‌های هوش مصنوعی در مدیریت منابع انسانی ارائه کرده‌اند، همچنان خلأ شایان توجهی در تبیین نحوه سازمان‌یافتگی مؤلفه‌های هوش مصنوعی مسئولانه، جایگاه نسبی هر یک از این مؤلفه‌ها در یک ساختار منسجم و روابط سلسله‌مراتبی میان ابعاد فنی، مدیریتی، اخلاقی و انسانی مشاهده می‌شود. همچنین، چارچوب‌های بین‌المللی هوش مصنوعی مسئولانه بر مجموعه‌ای از اصول و الزامات مشترک تأکید دارند؛ اما نحوه استقرار این اصول در بستر مدیریت منابع انسانی و ارتباط ساختاری آن‌ها با ابعاد شکوفایی کارکنان کمتر در کانون توجه قرار گرفته است.

برای مطالعه این مسئله، پژوهش حاضر از نظریه سیستم‌های اجتماعی - فنی^۴، به‌عنوان چارچوب نظری استفاده می‌کند. این نظریه که توسط چرنز^۵ (۱۹۷۶) و تریت^۶ (۱۹۸۱) توسعه یافته است، بر این اصل استوار است که اثربخشی نظام‌های سازمانی در گرو طراحی و بهینه‌سازی هم‌زمان زیرسیستم فنی و زیرسیستم اجتماعی است. بر اساس اصول طراحی سیستم‌های اجتماعی - فنی، از جمله سازگاری^۷، حداقل مشخص‌سازی بحرانی^۸، کنترل انحراف^۹، جریان

1. Malik, Budhwar & Kazmi

2. Spreitzer, Sutcliffe, Dutton, Sonenshein & Grant

3. Chowdhury et al.

4. Socio-Technical Systems Theory

5. Cherns

6. Trist

7. Compatibility

8. Minimal Critical Specification

9. Variance Control

اطلاعات^۱ و مشارکت در طراحی^۲، فناوری و عوامل انسانی نباید به صورت مستقل طراحی و مدیریت شوند، بلکه لازم است در قالب یک نظام یکپارچه و هماهنگ مورد توجه قرار گیرند (چرنز، ۱۹۷۶؛ تریست، ۱۹۸۱). از این منظر، هوش مصنوعی مسئولانه، فقط مجموعه‌ای از الزامات فنی یا اخلاقی نیست، بلکه سازوکاری برای ایجاد هم‌راستایی میان قابلیت‌های فناوریانه، الزامات حاکمیتی و نیازهای انسانی در نظام مدیریت منابع انسانی محسوب می‌شود. بنابراین، نظریه سیستم‌های اجتماعی - فنی، بستر مناسبی برای مطالعه روابط ساختاری میان مؤلفه‌های هوش مصنوعی مسئولانه و ابعاد شکوفایی کارکنان فراهم می‌آورد (ورونتیس و همکاران^۳، ۲۰۲۲).

بر این اساس، پژوهش حاضر با هدف شناسایی، تبیین و مدل‌سازی روابط ساختاری میان مؤلفه‌های مدیریت منابع انسانی مبتنی بر هوش مصنوعی مسئولانه و ابعاد شکوفایی کارکنان انجام شده است. این پژوهش با بهره‌گیری از رویکرد آمیخته اکتشافی - تبیینی، ابتدا از طریق مصاحبه‌های نیمه‌ساختاریافته با خبرگان، مؤلفه‌های کلیدی را شناسایی و تحلیل می‌کند و سپس با استفاده از تکنیک مدل‌سازی ساختاری - تفسیری، جایگاه و روابط سلسله‌مراتبی این مؤلفه‌ها را توضیح می‌دهد. نوآوری نظری پژوهش حاضر در ارائه یک مدل ساختاری از آرایش مؤلفه‌های هوش مصنوعی مسئولانه در مدیریت منابع انسانی، توسعه کاربرد نظریه سیستم‌های اجتماعی - فنی در محیط‌های کاری مبتنی بر هوش مصنوعی و تبیین جایگاه نسبی مؤلفه‌های فنی، حاکمیتی، مدیریتی و انسانی در یک چارچوب یکپارچه است. همچنین، نتایج این پژوهش، ضمن غنای ادبیات مدیریت منابع انسانی و هوش مصنوعی مسئولانه، می‌تواند مبنایی برای طراحی سیاست‌ها، سازوکارهای حاکمیتی و برنامه‌های استقرار مسئولانه هوش مصنوعی در سازمان‌ها را فراهم آورد.

پیشینه نظری پژوهش

نگرگاه نظری: نظریه سیستم‌های اجتماعی - فنی

برای تبیین شبکه روابط ساختاری میان هوش مصنوعی و پیامدهای انسانی، پژوهش حاضر بر «نظریه سیستم‌های اجتماعی - فنی^۴» (تریست، ۱۹۸۱؛ چرنز، ۱۹۷۶) تمرکز کرده است. بر اساس اصل بنیادین این نظریه، یعنی «بهینه‌سازی مشترک^۵»، کارایی و اثربخشی سیستم‌های سازمانی، فقط زمانی محقق می‌شود که طراحی و استقرار زیرسیستم‌های «فنی» (مانند الگوریتم‌های هوش مصنوعی) همگام، هم‌تراز و در یک پیوند ساختاری با نیازهای زیرسیستم‌های «اجتماعی» (کارکنان، ساختارهای انسانی و ارزش‌های روان‌شناختی) صورت پذیرد. ادبیات پیشین نشان می‌دهد که تمرکز صرف بر کارایی فنی هوش مصنوعی (دترمینیسم تکنولوژیک)، غالباً به شکست در پیاده‌سازی، بیگانگی کاری و مقاومت کارکنان منجر می‌شود.

1. Information Flow

2. Participation in Design

3. Vrontis et al.

4. Socio-Technical Systems Theory - STS

5. Joint Optimization

از این جنبه، «هوش مصنوعی مسئولانه»، به عنوان یک سازوکار میانجی^۱ و پیونددهنده عمل می کند که با هم راستا کردن منطق الگوریتمیک با ارزش های انسانی، زیرسیستم های اجتماعی و فنی را به تعادل ساختاری می رساند (اسچولته و استریت^۲، ۲۰۲۵). در تبیین دقیق تر این مفهوم بر مبنای اصول نه گانه طراحی سیستم های اجتماعی - فنی (چرنز، ۱۹۷۶)، این «تعادل ساختاری» تجلی «اصل سازگاری^۳» و بهینه سازی مشترک است؛ به این معنا که منطق طراحی الگوریتم ها باید با ارزش ها و نیازهای روان شناختی منابع انسانی هم سو باشد. علاوه بر این، در این مدل، مؤلفه هایی نظیر «پیشگیری از سوگیری و رعایت انصاف» پیوند تنگاتنگی با «اصل کنترل انحراف^۴» دارند. بر اساس این اصل، خطاها و انحراف ها باید در نزدیک ترین نقطه به مبدأ وقوع، شناسایی و اصلاح شوند؛ بنابراین شفافیت الگوریتم ها به عنوان یک تکنولوژی میانجی، به کاربران اجازه می دهد تا سوگیری ها را در همان مرحله پردازش داده مهار کنند. در نهایت، «توانمندسازی مدیریتی و نظارت انسانی» بازتابی از «مشارکت در طراحی^۵» است؛ زیرا مدیران توانمند شده، فقط مصرف کننده فناوری نیستند، بلکه از طریق اعمال کنترل انسانی^۶ در تنظیم و بازطراحی پیوسته سیستم در بستر اجتماعی سازمان مشارکت فعال دارند تا به بازتولید «ساختارهای خودسازمان ده^۷» کمک کنند.

بر اساس منطق نظریه سیستم های اجتماعی - فنی، پیامدهای انسانی ناشی از فناوری، نه از ویژگی های فنی فناوری به تنهایی و نه از ویژگی های فردی کارکنان به صورت مستقل ناشی می شود، بلکه حاصل کیفیت تعامل و هم ترازای ساختاری میان این دو زیرسیستم است. در بستر مدیریت منابع انسانی مبتنی بر هوش مصنوعی، مؤلفه هایی نظیر شفافیت الگوریتمی، تبیین پذیری، عدالت، پاسخ گوئی و نظارت انسانی به عنوان عناصر زیرسیستم فنی عمل می کنند. این عناصر از طریق شکل دهی ادراک کارکنان نسبت به انصاف، اعتماد، کنترل ادراک شده و امنیت روان شناختی، بر زیرسیستم اجتماعی اثر می گذارند. از این جنبه، شکوفایی کارکنان یک متغیر وابسته صرف به فناوری نیست، بلکه نتیجه هم راستایی موفق میان الزامات فنی سامانه های هوش مصنوعی و نیازهای انسانی کارکنان در یک شبکه سیستمی است.

تکامل هوش مصنوعی در مدیریت منابع انسانی و ضرورت رویکرد مسئولانه

بررسی سیر تحول فناوری در سازمان ها نشان می دهد که ورود هوش مصنوعی به حوزه مدیریت منابع انسانی، موجب دگرگونی اساسی در منطق تصمیم گیری، طراحی فرایندها و شیوه های تعامل با نیروی کار شده است. پژوهشگران (مانند خایر، ماهاداسا، تولی و آنده^۸، ۲۰۲۰) تأیید کرده اند که کاربرد هوش مصنوعی در جذب، ارزیابی عملکرد و آموزش، سرعت و دقت تصمیم ها را به شدت افزایش می دهد. با این حال، نسل جدید مطالعات هشدار می دهند که پدیده هایی نظیر «جعبه سیاه الگوریتمی» باعث شده اند که کارکنان، تصمیم های ماشین را فرایندهایی فهم ناپذیر و چالش ناپذیر ادراک کنند که این امر به طور مستقیم عاملیت انسانی و اعتماد سازمانی را هدف قرار می دهد.

1. Mediating Mechanism
2. Schulte & Streit
3. Compatibility Principle
4. Variance Control
5. Design Participation
6. Human-in-the-loop
7. Self-organizing Structures
8. Khair, Mahadasa, Tuli & Ande

در پاسخ به این بحران، مفهوم «هوش مصنوعی مسئولانه» ظهور کرده است. هوش مصنوعی مسئولانه در مدیریت منابع انسانی به مجموعه‌ای از سازوکارها، اصول و رویه‌های حاکمیتی اطلاق می‌شود که تضمین می‌کنند طراحی، استقرار و استفاده از سامانه‌های هوش مصنوعی با معیارهای شفافیت، پاسخ‌گویی، تبیین‌پذیری، عدالت، حفظ حریم خصوصی، نظارت انسانی و قابلیت اصلاح و کنترل هم‌سو باشد. این مفهوم بر مبنای چارچوب‌های معتبر بین‌المللی از جمله چارچوب مدیریت ریسک هوش مصنوعی مؤسسه ملی استاندارد و فناوری آمریکا^۱، قانون هوش مصنوعی اتحادیه اروپا^۲ و ادبیات هوش مصنوعی مسئولانه در مدیریت منابع انسانی توسعه یافته است. تریپاتی و کومار^۳ (۲۰۲۵) هوش مصنوعی مسئولانه را نه یک افزونه فنی، بلکه یک «الزام اخلاقی - استراتژیک» می‌دانند که چهار رکن اساسی دارد: عدالت الگوریتمی، شفافیت تصمیم‌ها، پاسخ‌گویی و حفظ کرامت انسانی. استقرار این الگو، تضمین‌کننده ایجاد تعادل میان مزایای فناورانه و مسئولیت‌های اجتماعی در اکوسیستم منابع انسانی است (چانگ و که^۴، ۲۰۲۴).

شکوفایی کارکنان در عصر الگوریتم‌ها

ادبیات نوین مدیریت منابع انسانی، از تمرکز سنتی بر متغیرهای تقلیلی نظیر «رضایت شغلی» عبور کرده است (واعظی و پاشایی، ۱۴۰۴) و بر مفهوم کلان‌تر «شکوفایی کارکنان» تأکید دارد. شکوفایی، سازه‌ای چندبعدی است که احساس معنا در کار، سلامت روان، درگیری فعال و رشد فردی را دربرمی‌گیرد (والتون، کیمپیمکی و ساولا^۵، ۲۰۲۶). در واقع، شکوفایی کارکنان، به‌عنوان یک وضعیت روان‌شناختی مثبت و پویا در نظر گرفته شده است که در آن، افراد علاوه‌بر عملکرد اثربخش، احساس رشد، یادگیری، معناداری و پیشرفت مستمر را در محیط کار تجربه می‌کنند. این تعریف با دیدگاه «شکوفایی در کار» ارائه‌شده توسط اسپریتزر و همکاران (۲۰۰۵) هم‌سو است که شکوفایی را حاصل ترکیب یادگیری و سرزندگی در محیط کار می‌دانند. بر این اساس، مؤلفه‌هایی نظیر رشد حرفه‌ای، معنا در کار، امنیت روان‌شناختی و مشارکت فعال به‌عنوان بازتاب‌های کلیدی شکوفایی کارکنان در بستر هوش مصنوعی مسئولانه شناسایی شدند. در پژوهش حاضر نیز، این ابعاد چهارگانه (رشد حرفه‌ای، معنا در کار، امنیت روان‌شناختی و مشارکت فعال) بر اساس یکپارچه‌سازی چارچوب‌های نظری اسپریتزر و همکاران (۲۰۰۵) و والتون و همکاران (۲۰۲۶) به‌عنوان شالوده مفهومی سازه شکوفایی کارکنان صورت‌بندی شده‌اند.

مطالعات نشان می‌دهد که صرف پیاده‌سازی فناوری‌های هوشمند، تضمین‌کننده شکوفایی و رفاه کارکنان نیست؛ بلکه نحوه ادراک آن‌ها از «عدالت» و «حریم خصوصی» در تعامل با الگوریتم‌هاست که تعیین می‌کند آیا فناوری به شکوفایی منجر می‌شود یا به بیگانگی کاری می‌انجامد (صادقی^۶، ۲۰۲۴؛ سومرو، فن، سوهو، سومرو و شیخ^۷، ۲۰۲۴). در

1. NIST AI RMF
2. EU AI Act
3. Tripathi & Kumar
4. Chang & Ke
5. Valtonen, Valtonen, Kimpimäki & Savela
6. Sadeghi
7. Soomro, Fan, Sohu, Soomro & Shaikh

چارچوب نظریه سیستم‌های اجتماعی - فنی، هوش مصنوعی مسئولانه از طریق ایجاد شفافیت در منطق تصمیم‌گیری الگوریتم‌ها، کاهش ادراک سوگیری، افزایش قابلیت پاسخ‌گویی و حفظ نظارت انسانی، شرایطی را فراهم می‌کند که کارکنان احساس اعتماد، انصاف و امنیت روان‌شناختی بیشتری داشته باشند. این سازوکارها به‌عنوان حلقه اتصال میان زیرسیستم فنی و زیرسیستم اجتماعی عمل کرده و زمینه لازم برای شکل‌گیری مشارکت فعال، یادگیری مستمر، معنا در کار و در نهایت شکوفایی کارکنان را فراهم می‌سازند.

مرور انتقادی مطالعات پیشین و تبیین شکاف پژوهشی

مرور نظام‌مند ادبیات در تقاطع هوش مصنوعی و مدیریت منابع انسانی نشان می‌دهد که با وجود رشد فزاینده مطالعات، پیکره دانشی موجود با سه نارسایی بنیادین نظری و متدولوژیک مواجه است:

غلبه پارادایم تکنومحور و سوگیری کارکردگرا^۱: جریان غالب پژوهش‌ها تاکنون، هوش مصنوعی را در چارچوب مکتب کارکردگرایی و فقط به‌عنوان ابزاری برای بهینه‌سازی فرایندها و ارتقای بهره‌وری سازمانی ارزیابی کرده است (بودوار و همکاران، ۲۰۲۳؛ ورونیتس و همکاران، ۲۰۲۲) و تأثیرهای عمیق این فناوری بر تجربه زیسته، امنیت روانی و شکوفایی روان‌شناختی کارکنان تا حد زیادی نادیده گرفته شده است (اسچولته و استریت، ۲۰۲۵).

خلاً در تبیین شبکه نومولوژیک و سازوکارهای ساختاری^۲: بخش عمده‌ای از مطالعاتی که به پیامدهای اخلاقی و انسانی هوش مصنوعی پرداخته‌اند، در سطح مرور مفهومی متوقف شده یا به آزمون‌های هم‌بستگی خطی ساده بسنده کرده‌اند (پریکشات، مالک و بودوار^۳، ۲۰۲۳). بر این اساس، خلاً تئوریک جدی در درک چگونگی اثرگذاری ابعاد «هوش مصنوعی مسئولانه» بر «شکوفایی کارکنان» وجود دارد. در واقع، هنوز مدلی که بتواند این ارتباط را به‌صورت نظام‌مند، در قالب یک «سلسله‌مراتب ساختاری» تبیین کند، توسعه نیافته است (تریپاتی و کومار، ۲۰۲۵).

مطالعات پیشین به‌ندرت سازوکارهای میانجی و فرایندهای انتقال را که از طریق آن‌ها ویژگی‌های فنی سامانه‌های هوش مصنوعی به پیامدهای انسانی و روان‌شناختی تبدیل می‌شوند، مورد توجه قرار داده‌اند. در نتیجه، نحوه تعامل زیرسیستم‌های فنی و اجتماعی که هسته مرکزی نظریه سیستم‌های اجتماعی - فنی و اصول نه‌گانه چرنز را تشکیل می‌دهد، همچنان به‌طور کافی در مدل‌های ساختاری تبیین نشده است.

نکته مهم دیگر آن است که اگرچه چارچوب‌های معتبر بین‌المللی، نظیر چارچوب مدیریت ریسک هوش مصنوعی مؤسسه ملی استاندارد و فناوری آمریکا، قانون هوش مصنوعی اتحادیه اروپا، اصول هوش مصنوعی گوگل و چارچوب هوش مصنوعی مسئولانه شرکت مایکروسافت، سهم مهمی در توسعه ادبیات هوش مصنوعی مسئولانه داشته‌اند؛ اما تمرکز اصلی آن‌ها بر طراحی، حاکمیت، مدیریت ریسک و کنترل پیامدهای نامطلوب سامانه‌های هوش مصنوعی است. این چارچوب‌ها مجموعه‌ای از اصول هنجاری و الزامات اجرایی نظیر عدالت، شفافیت، پاسخ‌گویی، امنیت، حفظ حریم

خصوصی و نظارت انسانی را ارائه می‌کنند؛ اما به‌طور عمده در سطح «آنچه باید انجام شود» (اصول هنجاری) باقی می‌مانند.

جدول ۱. تقابل رویکرد چارچوب‌های بین‌المللی با رویکرد پژوهش حاضر

شاخص مقایسه	چارچوب‌های بین‌المللی	مدل پیشنهادی پژوهش حاضر
تمرکز اصلی	مدیریت ریسک و حاکمیت	مدیریت منابع انسانی و شکوفایی
ماهیت رویکرد	هنجاری و الزامی	ساختاری و عملکردی
هدف نهایی	پیشگیری از آسیب و کنترل ریسک	ارتقای تجربه زیسته و شکوفایی کارکنان
واحد فنی	سیستم و زیرساخت فنی	تعامل سیستم و انسان
پرسش کلیدی	چه اصولی باید رعایت شوند؟	این اصول چگونه به پیامد انسانی تبدیل می‌شوند؟

همان‌طور که در جدول ۱ تبیین شد، شکاف اصلی نه در کمبود اصول، بلکه در فقدان «سازوکار انتقال» نهفته است. در واقع، چارچوب‌های مذکور بر «محتوای» اصول تمرکز دارند، در حالی که پژوهش حاضر بر «ساختار» اثرگذاری آن‌ها متمرکز است.

در مقابل، هنوز در ادبیات موجود توضیح روشنی دربارهٔ این موضوع وجود ندارد که این اصول از طریق چه سازوکارهای سازمانی، اجتماعی و روان‌شناختی به پیامدهای انسانی مطلوب در محیط کار منجر می‌شوند. به‌ویژه در حوزهٔ مدیریت منابع انسانی، بخش عمدهٔ پژوهش‌های موجود یا بر شناسایی اصول هوش مصنوعی مسئولانه متمرکز بوده‌اند یا اثرهای مستقیم فناوری را بر پیامدهای سازمانی بررسی کرده‌اند. در نتیجه، دانش موجود هنوز فاقد یک مدل ساختاری منسجم است که بتواند مسیرهای اثرگذاری مؤلفه‌های فنی هوش مصنوعی مسئولانه بر متغیرهای انسانی را در قالب یک نظام فنی - اجتماعی تبیین کند. این خلأ زمانی اهمیت بیشتری پیدا می‌کند که شکوفایی کارکنان به‌عنوان یک سازه چندبعدی و پیچیده، نه پیامد مستقیم فناوری، بلکه حاصل تعامل میان ویژگی‌های فنی سامانه‌های هوش مصنوعی و ادراکات، نگرش‌ها و تجربهٔ زیسته کارکنان تلقی شود. بر این اساس، نوآوری نظری پژوهش حاضر در معرفی اصول جدید هوش مصنوعی مسئولانه نیست، بلکه در تبیین روابط ساختاری میان اصول شناخته‌شده هوش مصنوعی مسئولانه و پیامدهای انسانی آن در بستر مدیریت منابع انسانی نهفته است. پژوهش حاضر با اتکا بر نظریه سیستم‌های اجتماعی - فنی و بهره‌گیری از مدل‌سازی ساختاری - تفسیری، می‌کوشد تا شکاف موجود میان چارچوب‌های حاکمیتی هوش مصنوعی مسئولانه و پیامدهای رفتاری و روان‌شناختی کارکنان را پر کند و با ترسیم «زنجیرهٔ تعاملات»، سازوکارهای انتقال اثرهای مؤلفه‌های فنی به شکوفایی کارکنان را آشکار سازد.

یکانه‌انگاری روش‌شناختی^۱ و فقر مطالعات اکتشافی – بومی

پدیده هوش مصنوعی مسئولانه ماهیتی به‌شدت نوظهور، چندبعدی و وابسته به بافتار^۲ دارد. با این حال، پیشینه تجربی غالباً بر رویکردهای کمی – قیاسی تکیه کرده است. این رویکرد از کشف مؤلفه‌های پنهان مرتبط با ارزش‌ها و نگرش‌های بومی نیروی کار عاجز است (بونجاک، سرنا و پوپوویچ^۳، ۲۰۲۳).

موضع پژوهش و معماری روش‌شناختی^۴

در پاسخ مستقیم به نارسایی‌های فوق، پژوهش حاضر با عبور از طرح‌های تک‌روشی، یک استراتژی «آمیخته اکتشافی – تبیینی متوالی»^۵ را اتخاذ کرده است. با استفاده از مطالب فوق، این معماری پژوهشی ابتدا در مرحله کیفی، از طریق تحلیل مضمون عمیق، ابعاد پنهان و بومی هوش مصنوعی مسئولانه را شناسایی می‌کند. سپس در مرحله کمی، با بهره‌گیری از تکنیک مدل‌سازی ساختاری تفسیری، از تحلیل‌های آماری تخت رایج فراتر رفته و امکان ترسیم شبکه روابط ساختاری و سطح‌بندی مؤلفه‌ها را فراهم می‌آورد. این رویکرد، گامی بدیع در جهت گذار از «توصیف پدیده» به «تبیین سازوکارهای ساختاری» هوش مصنوعی مسئولانه در ارتقای شکوفایی کارکنان محسوب می‌شود.

روش‌شناسی پژوهش

پژوهش حاضر در چارچوب پارادایم پراگماتیسم و با بهره‌گیری از طرح آمیخته اکتشافی – تبیینی متوالی انجام شده است (کرسول و پلانوکلازک^۶، ۲۰۱۸). با توجه به فقدان چارچوب‌های نظری و مدل‌های ساختاری جامع برای تبیین روابط میان هوش مصنوعی مسئولانه و شکوفایی کارکنان، این رویکرد امکان می‌دهد تا در مرحله نخست، ابعاد، مؤلفه‌ها و سازوکارهای مرتبط با پدیده مورد مطالعه به صورت اکتشافی شناسایی شوند و در مرحله دوم، ساختار روابط و الگوی تعامل میان این مؤلفه‌ها مورد بررسی قرار گیرد. در مطالعاتی با ماهیت اکتشافی و توسعه مدل، دانش تخصصی و تجربیات خبرگان، یکی از معتبرترین منابع برای شناسایی، پالایش و سازمان‌دهی مفاهیم پیچیده و چندبعدی محسوب می‌شود. از این رو، در مرحله مدل‌سازی، از قضاوت گروهی خبرگان آشنا با حوزه‌های مدیریت منابع انسانی، هوش مصنوعی و اخلاق فناوری استفاده شد تا ساختار روابط میان مؤلفه‌های شناسایی شده استخراج شود. این رویکرد با مبانی روش مدل‌سازی ساختاری – تفسیری سازگار بوده و امکان تبیین روابط سلسله‌مراتبی، الگوهای تعامل و وابستگی متقابل میان متغیرها را در حوزه‌های نوظهور و دارای پیچیدگی مفهومی فراهم می‌سازد.

مرحله اول (اکتشاف کیفی و استخراج مؤلفه‌ها)

جامعه آماری این مرحله، خبرگان دوحوزه‌ای (بخش منابع انسانی و توسعه‌دهندگان/سیاست‌گذاران هوش مصنوعی) بود.

1. Methodological Monism
2. Context-dependent
3. Bunjak, Černe & Popovič
4. Methodological Positioning
5. Sequential Exploratory Mixed-Methods Design
6. Creswell & Plano Clark

انتخاب مشارکت‌کنندگان بر اساس رویکرد نمونه‌گیری «هدفمند معیارمحور»^۱ سپس «گلوله برفی» انجام شد. معیارهای ورود شامل داشتن حداقل ۱۵ سال سابقه کار تخصصی و درگیری عملی/پژوهشی با سیستم‌های هوشمند بود. فرایند جمع‌آوری داده‌ها از طریق مصاحبه‌های نیمه‌ساختاریافته صورت گرفت و پس از انجام ۱۷ مصاحبه، پژوهش به «اشباع نظری» رسید؛ نقطه‌ای که داده‌های جدید، کد مفهومی تازه‌ای تولید نکردند. داده‌های متنی با استفاده از رویکرد بازنگری شده و تأملی تحلیل مضمون استخراج شدند (براون و کلارک^۲، ۲۰۲۱). برای تضمین «اعتبارپذیری»^۳، از تکنیک «بررسی توسط اعضا» استفاده شد. همچنین، به‌منظور کنترل سوگیری پژوهشگر و تأیید «پایایی بین‌کدگذاران» از ضرب کاپای کوهن استفاده شد. مقدار محاسباتی این شاخص برای دو کدگذار مستقل برابر با ۰/۸۳ به‌دست آمد که نشان‌دهنده توافق بالا و مطلوب است.

مرحله دوم (مدل‌سازی ساختاری تفسیری)

به‌منظور گذار یکپارچه از فاز کیفی به کمی و عملیاتی‌سازی دقیق متغیرهای پژوهش، ابعاد اولیه استخراج‌شده از ادبیات نظری (به‌ویژه مؤلفه‌های چندگانه شکوفایی کارکنان)، فقط به‌عنوان فهرست اولیه متغیرها مطرح شدند و پیش از ورود به مدل ساختاری - تفسیری، در پیل خبرگان و در فرایند تحلیل مضمون کیفی، بررسی، تعدیل و تأیید شدند. بر این اساس، این متغیرها به‌صورت پیش‌فرض در نظر گرفته نشدند؛ بلکه ابتدا در فاز کیفی و از طریق تحلیل مضمون مصاحبه‌ها، توسط پیل خبرگان مورد ارزیابی، بومی‌سازی و غربالگری قرار گرفتند. تنها پس از تأیید نهایی این مؤلفه‌ها، به‌عنوان ورودی‌های قطعی تکنیک «مدل‌سازی ساختاری تفسیری» مورد استفاده قرار گرفتند این تکنیک، یک متدولوژی سیستماتیک برای شناسایی روابط ساختاریافته مستقیم و غیرمستقیم میان متغیرهای یک سیستم پیچیده است (سوشیل^۴، ۲۰۱۲). شایان ذکر است که روابط حاصل از این تکنیک بیانگر ساختار ارتباطات و سطح‌بندی میان متغیرها بوده و نباید به‌عنوان روابط علی میان متغیرها تفسیر شوند. در این مرحله، یک پیل خبره متشکل از ۱۷ نفر از متخصصان ارشد، ماتریس خودتأملی ساختاری را تکمیل کرد. به‌منظور حفظ انسجام مفهومی میان مرحله شناسایی مؤلفه‌ها و مرحله مدل‌سازی ساختاری، بخشی از خبرگان در هر دو مرحله پژوهش مشارکت داشتند. این رویکرد در مطالعات اکتشافی و مدل‌سازی مبتنی بر خبرگان با هدف اطمینان از تداوم دانش تخصصی و تفسیر منسجم مفاهیم به‌کار گرفته می‌شود. با این حال، برای کاهش احتمال سوگیری ناشی از توافق گروهی، تلاش شد از خبرگان دارای سوابق علمی، تخصصی و تجربیات حرفه‌ای متنوع در حوزه‌های مدیریت منابع انسانی، هوش مصنوعی و اخلاق فناوری استفاده شود. این ماتریس بر اساس مقایسه‌های زوجی مفهومی برای تعیین جهت روابط (V, A, X, O) میان ابعاد هوش مصنوعی مسئولانه و شکوفایی کارکنان تنظیم شد. برای اطمینان از هم‌گرایی و اجماع نظرهای خبرگان در تکمیل ماتریس‌ها، از «ضرب

1. Criterion-based Purposive Sampling
2. Braun & Clarke
3. Credibility
4. Sushil

هماهنگی کندال^۱ استفاده شد و مقدار آن در پژوهش حاضر ۰/۸۴ به دست آمد. از آنجا که مقادیر بالای ۰/۷ نشان دهنده توافق قوی میان ارزیابان است (فیلد^۲، ۲۰۱۸)، پایایی ابزار در مرحله مدل سازی به طور قطعی تأیید شد. در نهایت، با تشکیل ماتریس های دسترسی اولیه و اعمال «قانون تراگذاری» در ریاضیات، ماتریس دسترسی نهایی شکل گرفت و مبنای ترسیم گراف جهت دار و سطح بندی سلسله مراتبی متغیرها شد.

تجزیه و تحلیل داده ها

در گام نخست، داده های میدانی از طریق مصاحبه های عمیق و نیمه ساختاریافته با ۱۷ نفر از خبرگان گردآوری شد. این پنل متشکل از متخصصانی با میانگین ۲۱ سال سابقه حرفه ای متمرکز و تحصیلات عالی (۲۹ درصد دارای مدرک دکتری) بود. همان طور که در جدول ۲ مشاهده می شود، توزیع متوازن تخصص ها در حوزه های مدیریت منابع انسانی، هوش مصنوعی، و اخلاق فناوری، امکان تثلیث^۳ منابع داده و بررسی چندوجهی پدیده را فراهم ساخت. فرایند نمونه گیری تا رسیدن به «اشباع نظری» ادامه یافت؛ به طوری که از مصاحبه پانزدهم به بعد، هیچ کد مفهومی جدیدی که تغییر معناداری در ساختار مدل ایجاد کند، شناسایی نشد.

جدول ۲. اطلاعات جمعیت شناختی خبرگان پژوهش

متغیر	طبقه بندی	تعداد	درصد
جنسیت	مرد	۱۱	۶۴/۷
	زن	۶	۳۵/۳
سطح تحصیلات	کارشناسی ارشد	۱۲	۷۰/۶
	دکتری	۵	۲۹/۴
حوزه تخصصی اصلی	مدیریت منابع انسانی	۷	۴۱/۲
	مدیریت رفتار سازمانی	۳	۱۷/۶
	اخلاق فناوری و سیاست گذاری	۴	۲۳/۵
	هوش مصنوعی و تحلیل داده	۳	۱۷/۶
نوع تجربه حرفه ای	دانشگاهی	۶	۳۵/۳
	اجرایی / مدیریتی	۵	۲۹/۴
	دانشگاهی - اجرایی	۶	۳۵/۳

تحلیل متون پیاده سازی شده با اتکا بر رویکرد شش مرحله ای و تأملی تحلیل مضمون (براون و کلارک، ۲۰۲۱) انجام گرفت. خروجی این فرایند رفت و برگشتی، استخراج ۱۹۶ کد اولیه بود که پس از پالایش مستمر، تقطیر معنایی و حذف هم پوشانی ها، در قالب ۱۵ مضمون فرعی و ۵ مضمون اصلی (تم) ساختاریافته شدند (جدول ۳).

1. Kendall's W
2. Field
3. Triangulation

جدول ۳. مضامین اصلی، فرعی و پایه مدل پژوهش

مضامین اصلی	مضامین فرعی	مضامین پایه
حاکمیت و تنظیم‌گری اخلاقی	پیشگیری از سوگیری داده‌ای	وجود سوگیری تاریخی در داده‌های استخدام، بازتولید تبعیض جنسیتی در الگوریتم‌ها، بازتولید تبعیض سنی در ارزیابی عملکرد، وابستگی مدل به داده‌های ناقص، نبود پایش عدالت الگوریتمی، ضرورت تست برابری فرصت، بررسی پیامدهای ناهمگون تصمیمات، کنترل متغیرهای تبعیض‌آمیز، اصلاح دیتاست‌های قدیمی، تحلیل خطای سیستماتیک، مقایسه نتایج بین گروه‌ها، حساسیت مدل به ویژگی‌های جمعیت‌شناختی، ضرورت ممیزی الگوریتم، شفاف‌سازی منبع داده‌ها، استانداردسازی داده‌های ورودی، حذف داده‌های غیرمرتبط، اعتبارسنجی داده‌های آموزشی، بازنگری دوره‌ای مدل
	پاسخ‌گویی و مسئولیت‌پذیری	مشخص بودن مسئول انسانی تصمیم، امکان پیگیری تصمیم الگوریتمی، ثبت مستندات تصمیمات، سازوکار اعتراض کارکنان، کمیته نظارت اخلاقی، گزارش‌دهی شفاف مدیریتی، تعیین مرجع پاسخ‌گو، جلوگیری از فرافکنی مسئولیت به ماشین، تعریف حدود اختیارات سیستم، وجود نظام رسیدگی به شکایات، مستندسازی فرایند تصمیم، ارزیابی اثرهای اجتماعی تصمیمات
	چارچوب سیاست‌گذاری	تدوین آیین‌نامه استفاده از هوش مصنوعی، تعریف اصول اخلاقی سازمان، انطباق با قوانین کار، رعایت مقررات حفاظت داده، وجود سیاست محرمانگی، آموزش اخلاق هوش مصنوعی به مدیران، تدوین کد رفتاری، پایش انطباق با استانداردها، ارزیابی ریسک فناوری، تحلیل اثرهای بلندمدت، طراحی چارچوب کنترل داخلی، تعریف خطوط قرمز تصمیم‌گیری، تعیین سطح مداخله انسانی، هم‌راستاسازی با ارزش‌های سازمان، ارزیابی تبعات شهرت سازمان، توجه به عدالت بین‌نسلی، در نظر گرفتن پیامدهای فرهنگی، پایش اثرهای نابرابر فناوری، بررسی اکثریت‌ها بر اقلیت‌ها، مستندسازی تغییرات الگوریتم، تعیین مالکیت داده‌ها، تعیین سطح دسترسی کاربران، ارزیابی امنیت سایبری، کنترل نشت اطلاعات، پیشگیری از سوءاستفاده داده، تعیین چرخه عمر داده، حذف داده‌های مازاد، به‌روزرسانی دوره‌ای سیاست‌ها، بازبینی سیاست‌های داخلی، ممیزی بیرونی سیستم، گزارش اثرات اخلاقی سالانه، ارزیابی ریسک تبعیض، تحلیل اثرات تصمیم بر رفاه کارکنان، تعریف شاخص‌های اخلاقی عملکرد
طراحی و معماری فناوریانه مسئولانه	تبیین‌پذیری الگوریتم	قابلیت توضیح منطق تصمیم‌گیری، امکان مشاهده معیارهای ارزیابی، وضوح شاخص‌های رتبه‌بندی داوطلبان، قابلیت ارائه گزارش تحلیلی به کارکنان، نمایش وزن متغیرها در مدل، تبیین علت رد یا پذیرش فرد، توضیح دلایل کاهش امتیاز عملکرد، امکان مشاهده روند تصمیم، استفاده از مدل‌های قابل توضیح، اجتناب از الگوریتم‌های جعبه سیاه، ترجمه خروجی فنی به زبان فهمیدنی، طراحی داشبورد شفاف، مستندسازی منطق مدل، آموزش کارکنان درباره نحوه تحلیل خروجی، امکان مقایسه نتایج انسانی و ماشینی
	قابلیت کنترل و اصلاح	امکان اصلاح دستی خروجی، توقف سیستم در شرایط خطا، تنظیم سطح خودکارسازی، تعریف آستانه‌های هشدار، شناسایی نتایج غیرعادی، امکان بازتنظیم پارامترها، بررسی خطای پیش‌بینی، بازبینی تصمیم‌های حساس، آزمایش سناریوهای جایگزین، ارزیابی دقت الگوریتم، اعتبارسنجی متقابل مدل، آزمون دوره‌ای عملکرد، امکان لغو تصمیم نهایی

مضامین پایه	مضامین فرعی	مضامین اصلی
رمزگذاری داده‌های کارکنان، محدودسازی دسترسی به اطلاعات، تفکیک داده‌های شخصی و شغلی، ذخیره‌سازی ایمن اطلاعات، اطلاع‌رسانی درباره نحوه استفاده از داده، اخذ رضایت آگاهانه کارکنان، حذف داده‌های غیرضروری، پایش دسترسی غیرمجاز، تعیین مدت نگهداری داده، شفافیت در انتقال داده، رعایت مقررات حفاظت داده، جلوگیری از اشتراک‌گذاری غیرمجاز	حفاظت داده و حریم خصوصی	
اطلاع‌رسانی قبل از اجرای سیستم، توضیح اهداف استفاده از هوش مصنوعی، جلسات پرسش و پاسخ کارکنان، شفاف‌سازی پیامدهای شغلی، اطلاع‌رسانی درباره تغییرات سیستم، ارائه گزارش‌های دوره‌ای، پاسخ‌گویی مدیران به نگرانی‌ها، ایجاد کانال ارتباطی مستقیم، انتشار راهنمای کاربری، اطلاع از معیارهای ارتقا، شفافیت در ارزیابی عملکرد	ارتباطات شفاف	
آموزش سواد داده‌ای مدیران، آموزش تحلیل گزارش‌های هوش مصنوعی، توسعه مهارت تصمیم‌گیری داده‌محور، آشنایی مدیران با محدودیت‌های الگوریتم، درک ریسک‌های فناوری، تقویت مهارت نظارت انسانی، آموزش اخلاق هوش مصنوعی، ارتقای شایستگی دیجیتال، تقویت تفکر انتقادی نسبت به خروجی سیستم	توانمندسازی مدیریتی	اجرای سازمانی و فرایندی
استفاده از هوش مصنوعی به‌عنوان ابزار کمکی، تأیید نهایی انسانی تصمیمات حساس، ترکیب قضاوت انسانی و داده، بررسی شرایط خاص فردی، امکان نادیده گرفتن خروجی سیستم، انعطاف در تصمیم‌گیری، ارزیابی کیفی در کنار ارزیابی کمی، مشورت گروهی قبل از تصمیم نهایی، استفاده مرحله‌ای از هوش مصنوعی، مداخله انسانی در موارد بحرانی، حفظ نقش سرپرست مستقیم، توازن بین اتوماسیون و قضاوت انسانی	تصمیم‌ترکیبی انسان - ماشین	
احساس عدالت رویه‌ای، احساس عدالت توزیعی، عدالت تعاملی، پذیرش تصمیمات سازمان، ادراک بی‌طرفی سیستم، احساس برابری فرصت، کاهش احساس تبعیض، مقایسه منصفانه عملکرد	ادراک انصاف	سازوکارهای روان‌شناختی
اعتماد به فناوری، اعتماد به مدیریت، اعتماد به صحت داده‌ها، باور به خیرخواهی سازمان، کاهش بدبینی نسبت به سیستم، اطمینان از صحت تصمیم، پذیرش نتایج الگوریتم، تمایل به استفاده از سیستم	اعتماد سازمانی	
کاهش اضطراب فناورانه، کاهش ترس از جایگزینی، احساس امنیت شغلی، اطمینان از حفظ کرامت انسانی، کاهش نگرانی درباره نظارت افراطی، احساس کنترل بر مسیر شغلی، آرامش در مواجهه با فناوری، کاهش استرس ناشی از ارزیابی دیجیتال	امنیت روان‌شناختی	
یادگیری شخصی‌سازی‌شده، پیشنهاد دوره‌های متناسب با مهارت، شناسایی شکاف مهارتی، برنامه‌ریزی مسیر شغلی، دریافت بازخورد تحلیلی، افزایش مهارت‌های تخصصی، توسعه شایستگی‌های نرم، تسهیل پیشرفت شغلی	رشد حرفه‌ای	
احساس پیشرفت حرفه‌ای، درک ارزش کار، احساس تأثیرگذاری، هم‌راستایی با اهداف سازمان، تعلق سازمانی، انگیزش درونی، رضایت از مسیر رشد	معنا در کار	پیامدهای توسعه‌ای و شکوفایی
افزایش تعامل شغلی، مشارکت در تصمیم‌گیری، تمایل به ارائه ایده، افزایش انرژی کاری، تعهد سازمانی، رفتار شهروندی سازمانی، همکاری بین‌واحدی، تمایل به ماندگاری در سازمان، احساس مالکیت نسبت به کار، انگیزه برای بهبود عملکرد، اشتیاق به یادگیری مستمر، تعامل مثبت با مدیر، احساس دیده شدن، احساس ارزشمندی، ابتکار عمل فردی، افزایش مسئولیت‌پذیری، تمایل به نوآوری، مشارکت در پروژه‌های توسعه‌ای، بازخوردگیری فعال، تلاش برای خودبهبودگری، هم‌سویی با فرهنگ دیجیتال سازمان	درگیری و مشارکت فعال	

تحلیل عمیق یافته‌های کیفی نشان می‌دهد که استقرار هوش مصنوعی مسئولانه، از یک مداخله صرفاً فناورانه فراتر رفته و به‌مثابه یک دگردیسی عمیق در اکوسیستم منابع انسانی سازمان عمل می‌کند. نتایج حاکی از آن است که غلبه بر پدیده «جعبه سیاه» از طریق شفافیت و تبیین‌پذیری الگوریتم‌ها، شرط لازم برای اعتمادسازی است؛ اما شرط کافی نیست. این شفافیت تنها زمانی به «امنیت روان‌شناختی» منجر می‌شود که با پاسخ‌گویی اخلاقی و صیانت از حریم خصوصی درهم‌آمیزد. در چنین بستر امنی است که سازوکارهای روان‌شناختی کارکنان فعال شده و با احیای احساس عاملیت و ادراک انصاف، مسیر برای شکوفایی، رشد حرفه‌ای و مشارکت فعال آنان هموار می‌شود.

یافته‌های پژوهش

به‌منظور کشف روابط ساختاریافته و سطح‌بندی ۱۵ مضمون فرعی استخراج‌شده، از تکنیک مدل‌سازی ساختاری تفسیری استفاده شد. در گام اول، ماتریس خودتعاملی ساختاری بر اساس خرد جمعی پنل خبرگان تشکیل شد (جدول ۴). در این ماتریس، جهت روابط میان جفت‌متغیرهای i و j با استفاده از نمادهای استاندارد مشخص شد: V (متغیر i بر j تأثیر دارد)، A (متغیر j بر i تأثیر دارد)، X (رابطه دوسویه) و O (عدم وجود رابطه).

جدول ۴. ماتریس ساختار روابط درونی متغیرها

	۱۵	۱۴	۱۳	۱۲	۱۱	۱۰	۹	۸	۷	۶	۵	۴	۳	۲	۱	
۱	O	O	O	O	V	V	V	A	O	O	A	X	A	A	-	
۲	V	O	O	V	V	V	X	A	X	V	A	A	A	-	V	
۳	V	V	V	V	V	V	V	V	V	V	V	V	-	V	V	
۴	O	O	O	O	V	V	V	A	V	O	A	-	A	V	X	
۵	O	O	O	O	V	V	V	A	O	O	-	V	A	V	V	
۶	O	O	O	V	V	O	O	A	O	-	O	O	A	A	O	
۷	V	O	O	V	V	V	X	A	-	O	O	A	A	X	O	
۸	V	V	V	V	V	V	V	-	V	V	V	V	A	V	V	
۹	V	O	O	O	V	V	-	A	X	O	A	A	A	X	A	
۱۰	V	V	O	V	V	-	A	A	A	O	A	A	A	A	A	
۱۱	V	V	O	X	-	A	A	A	A	A	A	A	A	A	A	
۱۲	V	O	O	-	X	A	O	A	A	A	O	O	A	A	O	
۱۳	V	V	-	O	O	O	O	A	O	O	O	O	A	O	O	
۱۴	V	-	A	O	A	A	O	A	O	O	O	O	A	O	O	
۱۵	-	A	A	A	A	A	A	A	A	O	O	O	A	A	O	

در گام بعد، بر اساس قواعد تبدیل باینری، ماتریس خودتعاملی ساختاری به «ماتریس دسترسی اولیه» با درایه‌های صفر و یک تبدیل شد؛ بدین صورت که قطر اصلی ماتریس برابر با ۱ در نظر گرفته شد.

جدول ۵. ماتریس دسترسی اولیه

۱۵	۱۴	۱۳	۱۲	۱۱	۱۰	۹	۸	۷	۶	۵	۴	۳	۲	۱	
۰	۰	۰	۰	۱	۱	۱	۰	۰	۰	۰	۱	۰	۰	۱	۱
۱	۰	۰	۱	۱	۱	۱	۰	۱	۱	۰	۰	۰	۱	۱	۲
۱	۱	۱	۱	۱	۱	۱	۱	۱	۱	۱	۱	۱	۱	۱	۳
۰	۰	۰	۰	۱	۱	۱	۰	۱	۰	۰	۱	۰	۱	۱	۴
۰	۰	۰	۰	۱	۱	۱	۰	۰	۰	۱	۱	۰	۱	۱	۵
۰	۰	۰	۱	۱	۰	۰	۰	۰	۱	۰	۰	۰	۰	۰	۶
۱	۰	۰	۱	۱	۱	۱	۰	۱	۰	۰	۰	۰	۱	۰	۷
۱	۱	۱	۱	۱	۱	۱	۱	۱	۱	۱	۱	۰	۱	۱	۸
۱	۰	۰	۰	۱	۱	۱	۰	۱	۰	۰	۰	۰	۱	۰	۹
۱	۱	۰	۱	۱	۱	۰	۰	۰	۰	۰	۰	۰	۰	۰	۱۰
۱	۱	۰	۱	۱	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۱۱
۱	۰	۰	۱	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۱۲
۱	۱	۱	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۱۳
۱	۱	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۱۴
۱	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۱۵

مهم ترین گام در مدل سازی ساختاری - تفسیری، کشف روابط پنهان سیستمی است. بر همین اساس، ماتریس اولیه با استفاده از الگوریتم محاسباتی «فلوید - وارشال»^۱ در نرم افزار اکسل مورد تحلیل قرار گرفت تا «اصل تراگذاری» بر آن اعمال شود. منطق ریاضی این اصل بیان می کند که اگر عامل F_i بر عامل F_j اثرگذار باشد و عامل F_j نیز بر عامل F_k تأثیر بگذارد، آنگاه عامل F_i نیز به صورت غیرمستقیم بر عامل F_k اثر خواهد داشت. با تکرار این الگوریتم تا رسیدن به ثبات ساختاری، ماتریس دسترسی نهایی شکل گرفت که دربردارنده کلیه روابط مستقیم و غیرمستقیم متغیرهاست (جدول ۵) (آلاوامله، الطوال، لعلو و جامع^۲، ۲۰۲۳؛ داس، کیرون و پرامود^۳، ۲۰۲۴). به منظور سطح بندی متغیرها و ترسیم ساختار سلسله مراتبی مدل، مجموعه های «دسترسی»^۴ و «پیشایندی»^۵ برای هر متغیر استخراج شد. مجموعه دسترسی شامل تمامی عواملی است که یک عامل می تواند به آن ها دسترسی مستقیم یا غیرمستقیم داشته باشد و مجموعه پیشایندی شامل تمامی عواملی است که بر آن عامل اثرگذار هستند. بر اساس الگوریتم کانونی مدل سازی ساختاری - تفسیری، در هر تکرار، متغیرهایی که «مجموعه دسترسی» آن ها با «مجموعه اشتراک» کاملاً برابر بود، به عنوان عناصر بالاترین سطح سیستم شناسایی شدند (وبر^۶، ۲۰۲۱). دلیل ریاضی این امر آن است که این متغیرها، تأثیرپذیرترین عناصر سیستم هستند و توانایی اثرگذاری بر سایر سطوح را ندارند. پس از شناسایی این متغیرها، آن ها از فهرست حذف شدند و فرایند تا سطح بندی تمامی ۱۵ متغیر ادامه یافت (جدول ۶).

1. Floyd-Warshall Algorithm
2. Alawamleh, Al-Twal, Lahlouh & Jame
3. Das, Kiron & Pram
4. Reachability Set
5. Antecedent Set
6. Weber

جدول ۶. مجموعه‌های دستیابی، پیش‌بینی و اشتراک شاخص‌ها برای تعیین سطوح

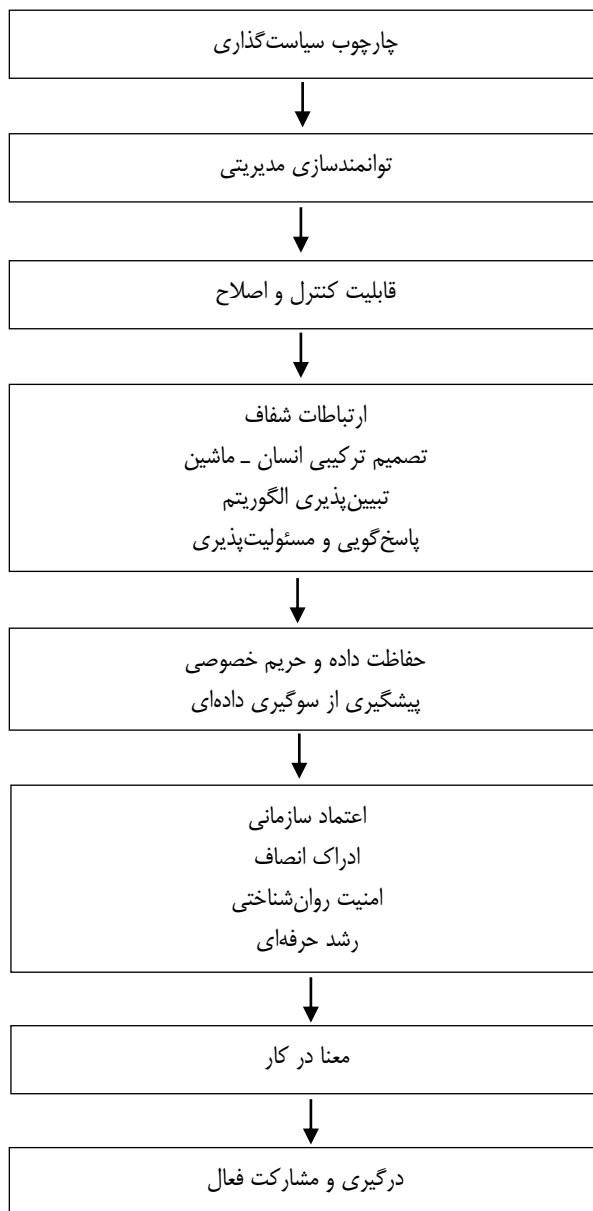
شاخص	مجموعه دستیابی	مجموعه پیش‌بینی	اشتراک
۱	F۱, F۲, F۴, F۷, F۹, F۱۰, F۱۱, F۱۲, F۱۴, F۱۵	F۱, F۲, F۳, F۴, F۵, F۷, F۸, F۹	F۱, F۲, F۴, F۷, F۹
۲	F۱, F۲, F۴, F۶, F۷, F۹, F۱۰, F۱۱, F۱۲, F۱۴, F۱۵	F۱, F۲, F۳, F۴, F۵, F۷, F۸, F۹	F۱, F۲, F۴, F۷, F۹
۳	F۱-F۱۵	F۳	F۳
۴	F۱, F۲, F۴, F۶, F۷, F۹, F۱۰, F۱۱, F۱۲, F۱۴, F۱۵	F۱, F۲, F۳, F۴, F۵, F۷, F۸, F۹	F۱, F۲, F۴, F۷, F۹
۵	F۱, F۲, F۴, F۵, F۶, F۷, F۹, F۱۰, F۱۱, F۱۲, F۱۴, F۱۵	F۲, F۵, F۸	F۵
۶	F۶, F۱۰, F۱۱, F۱۲, F۱۴, F۱۵	F۲, F۳, F۴, F۵, F۶, F۷, F۸, F۹	F۶
۷	F۱, F۲, F۴, F۶, F۷, F۹, F۱۰, F۱۱, F۱۲, F۱۴, F۱۵	F۱, F۲, F۳, F۴, F۵, F۷, F۸, F۹	F۱, F۲, F۴, F۷, F۹
۸	F۱, F۲, F۴, F۶, F۷, F۹, F۱۰, F۱۱, F۱۲, F۱۴, F۱۵	F۳, F۸	F۸
۹	F۱, F۲, F۴, F۶, F۷, F۹, F۱۰, F۱۱, F۱۲, F۱۴, F۱۵	F۱, F۲, F۳, F۴, F۵, F۶, F۷, F۸, F۹	F۹
۱۰	F۱۰, F۱۱, F۱۲, F۱۳, F۱۴, F۱۵	F۱-F۱۲	F۱۰, F۱۱, F۱۲
۱۱	F۱۰, F۱۱, F۱۲, F۱۳, F۱۴, F۱۵	F۱-F۱۲	F۱۰, F۱۱, F۱۲
۱۲	F۱۰, F۱۱, F۱۲, F۱۳, F۱۴, F۱۵	F۱-F۱۲	F۱۰, F۱۱, F۱۲
۱۳	F۱۳, F۱۴, F۱۵	F۳, F۸, F۱۳	F۱۳
۱۴	F۱-F۱۴	F۹, F۱۰, F۱۱, F۱۲, F۱۳, F۱۴, F۱۵	F۱۴
۱۵	F۱-F۱۵	F۸, F۹, F۱۰, F۱۱, F۱۲, F۱۳, F۱۴, F۱۵	F۱۵

جدول ۷. تعیین سطح شاخص‌های مدل‌سازی ساختاری - تفسیری

شماره سطح	شاخص‌های هر سطح	کارکرد اصلی
سطح ۱	درگیری و مشارکت فعال (۱۵)	وابسته: تأثیرپذیری بسیار بالا و تأثیرگذاری بسیار کم. این شاخص خروجی نهایی مدل است.
سطح ۲	معنا در کار (۱۴)	وابسته: متغیرهای هدف که بیشتر تحت تأثیر سطوح پایین‌تر قرار دارند.
سطح ۳	ادراک انصاف (۱۰)، اعتماد سازمانی (۱۱)، امنیت روان‌شناختی (۱۲)، رشد حرفه‌ای (۱۳)	رابط / پیوندی: متغیرهای میانی که هم تأثیرپذیری و هم تأثیرگذاری نسبتاً بالایی دارند و نقش انتقال‌دهنده را ایفا می‌کنند.
سطح ۴	حفاظت داده و حریم خصوصی (۶)، پیشگیری از سوگیری داده‌ای (۱)	رابط / پیوندی: متغیرهای واسطه‌ای در زنجیره تأثیرات.
سطح ۵	ارتباطات شفاف (۷)، تصمیم‌ترکیبی انسان-ماشین (۹)، تبیین‌پذیری الگوریتم (۴)، پاسخ‌گویی و مسئولیت‌پذیری (۲)	رابط / پیوندی: متغیرهای واسطه‌ای که به سمت استقلال میل می‌کنند.
سطح ۶	قابلیت کنترل و اصلاح (۵)	مستقل / نفوذی: تأثیرگذاری بالا و تأثیرپذیری کم.
سطح ۷	توانمندسازی مدیریتی (۸)	مستقل / نفوذی: متغیرهای پایه‌ای با قدرت هدایت بالا.
سطح ۸	چارچوب سیاست‌گذاری (۳)	مستقل: سنگ بنای مدل. این شاخص کلیدی‌ترین محرک سیستم است که بر تمام شاخص‌های دیگر تأثیر می‌گذارد اما از هیچ‌کدام تأثیر نمی‌پذیرد.

خروجی نهایی این فرایند، استخراج یک مدل ساختاری هشت‌سطحی است (جدول ۷).

مدل نهایی پژوهش که با استفاده از روش مدل‌سازی ساختاری - تفسیری توسعه یافته است، ساختار سلسله‌مراتبی مؤلفه‌های مرتبط با اخلاق هوش مصنوعی را در ۸ سطح ترسیم می‌کند (شکل ۱). این مدل، الگوی ارتباطات و تعاملات ساختاری میان متغیرها را در قالب یک نظام سطح‌بندی شده نشان می‌دهد؛ به گونه‌ای که مؤلفه‌های زیربنایی در سطوح پایین‌تر و مؤلفه‌های وابسته‌تر در سطوح بالاتر قرار گرفته‌اند.



شکل ۱. مدل نهایی پژوهش

در قاعده مدل (سطح ۸)، «چارچوب سیاست‌گذاری» به عنوان پیشران ساختاری و بنیادی‌ترین مؤلفه قرار دارد. این مؤلفه در نظام روابط، دارای بیشترین قدرت نفوذ بوده و بستر اولیه برای تعاملات سایر مؤلفه‌ها در سیستم را فراهم

می‌آورد. در سطوح ۷ و ۶، مؤلفه‌های «توانمندسازی مدیریتی» و «قابلیت کنترل و اصلاح» قرار گرفته‌اند که در شبکه روابط، نقش پیونددهنده سیاست‌های کلان به الزامات عملیاتی و فنی را ایفا می‌کنند. سطوح میانی (۵، ۴ و ۳)، به‌عنوان کانون تعاملات ساختاری در مدل شناخته می‌شوند؛ به‌گونه‌ای که مؤلفه‌هایی همچون «پاسخ‌گویی» و «تبیین‌پذیری» در پیوند با سازوکارهای فنی نظیر «پیشگیری از سوگیری» قرار دارند. این سازوکارها در ساختار مدل، با مؤلفه‌های ادراکی - روان‌شناختی مانند «ادراک انصاف» و «اعتماد سازمانی» در یک شبکه ارتباطی متقابل قرار می‌گیرند. در نهایت، در سطوح ۲ و ۱، مؤلفه‌های وابسته به سایر متغیرها جای گرفته‌اند. «معنا در کار» به‌عنوان یک پیامد روان‌شناختی کلیدی، در بالاترین سطح مدل، جایگاه وابسته‌تری نسبت به سایر مؤلفه‌ها دارد و با مؤلفه «درگیری و مشارکت فعال» در یک سطح از تحلیل ساختاری قرار می‌گیرد. این مدل نشان می‌دهد که استقرار هوش مصنوعی در سازمان‌ها مستلزم یک رویکرد جامع و چندلایه است که از تدوین سیاست‌های ساختاری در سطح پایه آغاز شده و تا برقراری ارتباط با ادراک انصاف و اعتماد در سطح فردی امتداد می‌یابد.

نتیجه‌گیری و پیشنهادها

پژوهش حاضر با هدف ارائه یک مدل یکپارچه برای پیاده‌سازی هوش مصنوعی مسئولانه در مدیریت منابع انسانی و تبیین جایگاه آن در شکوفایی کارکنان انجام شد. یافته کلیدی پژوهش، تدوین یک مدل ساختاری - تفسیری هشت‌سطحی است که نشان می‌دهد دستیابی به شکوفایی کارکنان، از طریق هوش مصنوعی نه یک پیامد مستقیم فناوری، بلکه نتیجه زنجیره تعاملات ساختاری پیچیده است که از سطح حاکمیتی آغاز شده و از طریق سازوکارهای میانجی به پیامدهای روان‌شناختی و رفتاری در سطوح بالاتر منتهی می‌شود. این یافته نشان می‌دهد که رابطه میان هوش مصنوعی و شکوفایی کارکنان، رابطه‌ای خطی و فناورانه نیست، بلکه رابطه‌ای چندلایه و وابسته به هم‌ترازی ساختاری میان زیرسیستم فنی و زیرسیستم اجتماعی سازمان است. این یافته، دقیقاً همان شکاف تبیینی مطرح شده در بخش پیشینه را پر می‌کند؛ چرا که نشان می‌دهد چگونه اصول هنجاری هوش مصنوعی مسئولانه، از طریق این زنجیره تعاملات، از سطح «الزامات سیاستی» فراتر می‌رود و به «تجربیات زیسته و پیامدهای انسانی» در محیط کار تبدیل می‌شوند.

بر خلاف برخی پژوهش‌ها که نقطه شروع را مستقیماً فناوری می‌دانند (برای نمونه تامبه، کاپلی و یاکوبوویچ^۱، ۲۰۱۹)، مدل حاضر نشان می‌دهد که «چارچوب سیاست‌گذاری» (سطح ۸) به‌عنوان یک متغیر مستقل، بنیادی‌ترین عنصر در ساختار سیستم است. در واقع، برخلاف رویکردهای صرفاً هنجاری چارچوب‌های بین‌المللی (مانند میکروسافت و گوگل) که بر «محتوای» اصول تمرکز دارند، یافته‌های این پژوهش نشان می‌دهد که اهمیت اصلی در «ساختار عملیاتی‌سازی» این اصول نهفته است؛ به این معنا که سیاست‌گذاری تنها زمانی به تغییر در رفتار کارکنان منجر می‌شود که به‌عنوان ریشه یک زنجیره ساختاری عمل کند. این نتیجه، از جنبه نظریه سیستم‌های اجتماعی - فنی قابل تبیین

است؛ زیرا در این چارچوب، پیش نیاز بهینه سازی مشترک میان فناوری و انسان، وجود سازوکارهای حاکمیتی است که جهت گیری ارزشی و اخلاقی سیستم را تعیین می کنند. این یافته با ادبیات حاکمیت فناوری نیز هم سو است که استدلال می کند بدون وجود خط مشی های روشن، شفاف و اخلاقی، استقرار فناوری های پیشرفته با بی نظمی نهادی و تضعیف اعتماد سازمانی همراه خواهد بود (وون کروگ^۱، ۲۰۱۲) در این معنا، چارچوب سیاست گذاری فقط یک سند اداری نیست، بلکه یک پیشران ساختاری است که بستر نهادی لازم برای «توانمندسازی مدیریتی» (سطح ۷) و «قابلیت کنترل و اصلاح» (سطح ۶) را فراهم می کند. پژوهش حاضر بدین ترتیب استدلال می کند که مسئولیت پذیری هوش مصنوعی پیش از آنکه در سطح الگوریتم ها تعریف شود، باید در سطح استراتژی و سیاست های کلان سازمان نهادینه شود.

قلب مدل در سطوح میانی قرار دارد؛ جایی که متغیرهای رابط، انتقال الزامات حاکمیتی به تجربه زیسته کارکنان را ممکن می سازند. از جنبه نظری، این بخش از مدل با اصول طراحی سیستم های اجتماعی - فنی چرنز (۱۹۷۶) و ادبیات عدالت سازمانی پیوند می خورد. در چارچوب عدالت سازمانی، کارکنان فقط به نتایج تصمیمات سازمانی توجه نمی کنند، بلکه شیوه اتخاذ این تصمیمات و امکان مشارکت یا بازبینی آن ها را نیز در ارزیابی خود دخیل می دانند. در محیط های مبتنی بر هوش مصنوعی، زمانی که منطق تصمیمات الگوریتمی قابل درک و تبیین باشد، امکان بازبینی و اصلاح وجود داشته باشد و سازوکارهای کنترل سوگیری فعال باشند، عدالت رویه ای تقویت می شود و ادراک انصاف در میان کارکنان شکل می گیرد. در ساختار مدل، این سازوکارها پیش نیازهای سطح ۳، یعنی «اعتماد سازمانی»، «امنیت روان شناختی» و «ادراک انصاف» را تشکیل می دهند. این نتیجه با نظریه تبادل اجتماعی (بلاو^۲، ۱۹۵۴) نیز هم خوان است؛ نظریه ای که بر مبادله متقابل میان رفتار منصفانه سازمان و واکنش های مثبت کارکنان تأکید دارد. در این چارچوب، شفافیت، پاسخ گویی و پیشگیری از سوگیری، به منزله سیگنال هایی از احترام و ارزش گذاری سازمان به کارکنان عمل می کنند و به شکل گیری اعتماد، تعهد و مشارکت منجر می شوند. بنابراین، اعتماد در مدل حاضر نه یک پیش فرض ایستا، بلکه یک سازه برساخته در شبکه روابط است که از طریق عدالت، شفافیت و پاسخ گویی تولید و بازتولید می شود.

در این فرایند، تقویت سطوح میانی به ایجاد «ساختارهای خودسازمان ده»^۳ کمک می کند؛ ساختارهایی که امکان حفظ عاملیت انسانی و تنظیم تعادل میان سیستم فنی و سیستم انسانی را فراهم می سازند. اهمیت این سازوکارها در تجربه های واقعی سازمانی نیز قابل مشاهده است. برای مثال، تجربه شرکت آمازون در توسعه سامانه غربالگری رزومه مبتنی بر هوش مصنوعی نشان داد که آموزش الگوریتم بر داده های تاریخی سوگیرانه می تواند به بازتولید تبعیض جنسیتی منجر شود. این تجربه آشکار ساخت که حتی سامانه هایی با عملکرد فنی مناسب، در صورت فقدان سازوکارهای شفافیت، نظارت انسانی و کنترل سوگیری، با کاهش ادراک انصاف و تضعیف اعتماد ذی نفعان همراه خواهند بود. از جنبه عدالت سازمانی، چنین وضعیتی به تضعیف عدالت رویه ای می انجامد؛ از جنبه نظریه تبادل اجتماعی، پیام بی توجهی سازمان به دغدغه های کارکنان را منتقل می کند؛ و از جنبه سیستم های اجتماعی - فنی، نشانه شکست در بهینه سازی

1. von Krogh

2. Blau

3. Self-organizing Structures

مشترک و غلبه زیرسیستم فنی بر زیرسیستم اجتماعی است. یافته‌های پژوهش حاضر نیز نشان می‌دهد که پیشگیری از سوگیری، تبیین‌پذیری و پاسخ‌گویی، پیش‌نیازهای اساسی شکل‌گیری اعتماد سازمانی و امنیت روان‌شناختی در شبکه روابط هستند. بنابراین، تجربه‌آمازون به‌عنوان یک شاهد تجربی، دقیقاً تأییدکننده مدل حاضر است: فقدان سازوکارهای انتقال در سطوح میانی (عدالت و شفافیت)، باعث می‌شود که حتی پیشرفته‌ترین زیرسیستم فنی، نتواند به شکوفایی زیرسیستم اجتماعی منجر شود و در نهایت در شکست بهینه‌سازی مشترک متوقف شود.

در نهایت، متغیرهای وابسته مدل یعنی «معنا در کار» (سطح ۲) و «درگیری و مشارکت فعال» (سطح ۱) به‌عنوان پیامدهای نهایی این زنجیره ساختاری ظاهر می‌شوند. این نتیجه دیدگاه‌هایی را که هوش مصنوعی را ذاتاً تهدیدی برای عاملیت انسانی می‌دانند، به چالش می‌کشد. پژوهش حاضر نشان می‌دهد که اگر هوش مصنوعی از طریق سطوح ۸ تا ۳ و در چارچوب یک معماری مسئولانه پیاده‌سازی شود، می‌تواند با خودکارسازی وظایف تکراری، ارائه بازخوردهای توسعه‌ای و افزایش قابلیت تصمیم‌گیری آگاهانه، نه تنها عاملیت انسانی را تضعیف نکند، بلکه به شکوفایی آن کمک کند (اسپریتزر و همکاران، ۲۰۰۵).

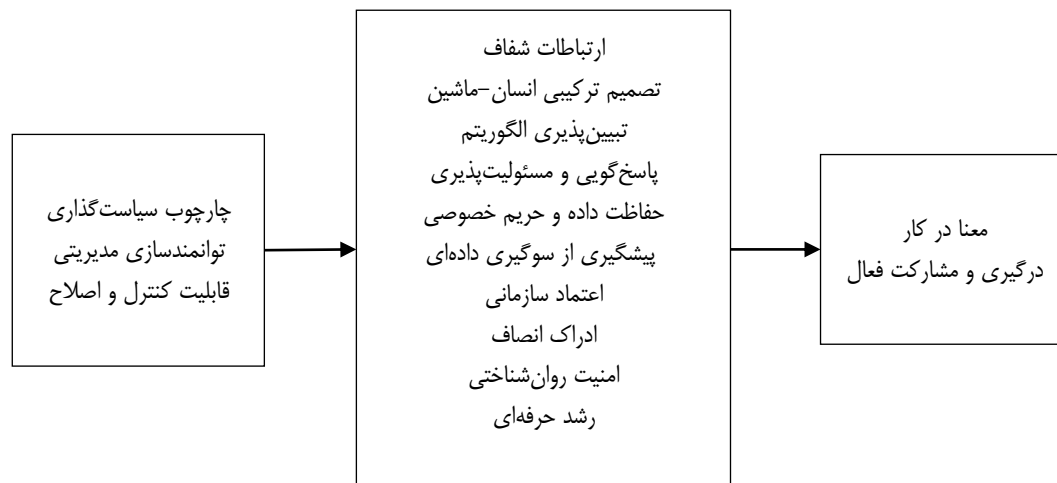
اگرچه مدل نهایی در قالب هشت سطح سلسله‌مراتبی ارائه شده است، به‌منظور ارتقای قابلیت تفسیر نظری و تسهیل کاربرد مدیریتی، این سطوح در قالب سه لایه مفهومی کلان‌بازنمایی شده‌اند. مبنای این لایه‌بندی، منطق نظری «نظریه سیستم‌های اجتماعی - فنی» است که بر ضرورت هم‌ترازی و بهینه‌سازی مشترک میان زیرسیستم فنی و زیرسیستم اجتماعی در سازمان تأکید دارد. بر این اساس، سطوح پایین‌تر مدل که نقش پیشران‌های ساختاری را ایفا می‌کنند در قالب لایه «حاکمیتی - فنی» (سطوح ۸ تا ۶) صورت‌بندی شده‌اند. این لایه به‌عنوان زیرسیستم فنی شامل مؤلفه‌هایی نظیر چارچوب‌های سیاست‌گذاری، سازوکارهای حاکمیت داده و الگوریتم، توانمندسازی مدیریتی و قابلیت کنترل، نظارت و اصلاح سامانه‌های هوش مصنوعی است که بستر نهادی و فناورانه استقرار هوش مصنوعی مسئولانه را فراهم می‌سازند.

در سطح بعدی، لایه «اجتماعی - روان‌شناختی» (سطوح ۵ تا ۳) قرار دارد که نمایانگر زیرسیستم اجتماعی سازمان است. این لایه بیانگر سازوکارهایی است که از طریق آن‌ها ویژگی‌های زیرسیستم فنی به ادراکات، نگرش‌ها و تجربه‌های کارکنان منتقل می‌شود. مؤلفه‌هایی نظیر عدالت ادراک‌شده، شفافیت تصمیم‌های الگوریتمی، اعتماد به سامانه‌های هوش مصنوعی و امنیت روان‌شناختی در این سطح قرار می‌گیرند و نقش میانجی را در تبدیل الزامات فنی و حاکمیتی به پیامدهای انسانی ایفا می‌کنند.

در نهایت، لایه «پیامدی» (سطوح ۲ تا ۱) قرار دارد که منعکس‌کننده برون‌دادهای انسانی و سازمانی حاصل از هم‌ترازی موفق زیرسیستم‌های فنی و اجتماعی است. در این سطح، پیامدهایی نظیر معنا در کار، مشارکت فعال کارکنان و در نهایت شکوفایی کارکنان به‌عنوان نتایج نهایی پیاده‌سازی مسئولانه هوش مصنوعی در مدیریت منابع انسانی ظهور می‌یابد.

بنابراین، بازنمایی سه‌لایه‌ای ارائه‌شده در این پژوهش تفسیری نظری از ساختار روابط نفوذ استخراج‌شده از مدل‌سازی ساختاری - تفسیری است که با اتکا به منطق نظریه سیستم‌های اجتماعی - فنی صورت‌بندی شده است. این

تفسیر مفهومی، بدون ایجاد تغییر در ساختار سلسله‌مراتبی مدل، امکان درک بهتر منطق علی میان سطوح مختلف را فراهم ساخته و در مرحله نهایی تحلیل نیز با در نظر گرفتن دیدگاه‌های خبرگان حوزه مدیریت منابع انسانی و هوش مصنوعی مورد تأیید قرار گرفته است. از این رو، این لایه‌بندی علاوه بر تقویت انسجام نظری مدل، کاربرپذیری آن را برای سیاست‌گذاران و مدیران در مسیر پیاده‌سازی مسئولانه هوش مصنوعی در سازمان‌ها افزایش می‌دهد.



شکل ۲. مدل تسهیل‌یافته نهایی پژوهش

این پژوهش چندین مشارکت نظری مهم دارد:

- یکپارچه‌سازی ادبیات: این مطالعه با ارائه یک مدل سلسله‌مراتبی، سه حوزه مجزای ادبیات را به هم پیوند می‌زند: حاکمیت فناوری اطلاعات (سطوح بالا)، عدالت سازمانی (سطوح میانی)، و رفتار سازمانی مثبت‌گرا (سطوح پایین). این مدل تبیین می‌کند که چگونه مفاهیم کلان حاکمیتی در یک پیوند ساختاری با پیامدهای خرد روان‌شناختی قرار می‌گیرند.
- عمق‌بخشی به نقش اعتماد: در حالی که بسیاری از مطالعات بر اهمیت «اعتماد به هوش مصنوعی» تأکید می‌کنند، این پژوهش با تعیین جایگاه «اعتماد سازمانی» در سطوح میانی مدل، نشان می‌دهد که اعتماد، بخشی از یک شبکه تعاملی گسترده‌تر است که سازمان برای مدیریت منصفانه فناوری ایجاد می‌کند.
- ارائه مدل ساختاری برای هوش مصنوعی مسئولانه: این پژوهش فراتر از فهرست کردن اصول اخلاقی، یک مدل ساختاری ارائه می‌دهد که جایگاه و اولویت ساختاری هر یک از این اصول را در مسیر تحقق نتایج مرتبط با کارکنان مشخص می‌کند.
- بسط چارچوب‌های هوش مصنوعی مسئولانه به حوزه پیامدهای انسانی: یکی از مشارکت‌های نظری مهم این پژوهش، انتقال کانون توجه از اصول حاکمیتی هوش مصنوعی مسئولانه به جایگاه آن‌ها در تجربیات انسانی

در محیط کار است. مدل ارائه شده آشکار می‌سازد که مؤلفه‌های فنی هوش مصنوعی مسئولانه زمانی با شکوفایی کارکنان هم‌راستا می‌شوند که از طریق اعتماد سازمانی، ادراک انصاف و امنیت روان‌شناختی در تجربه کارکنان بازتولید شوند. از این منظر، پژوهش حاضر فهم موجود از هوش مصنوعی مسئولانه را از سطح طراحی و حاکمیت فناوری به سطح تعاملات انسانی و پیامدهای رفتاری در سازمان گسترش می‌دهد.

یافته‌های این پژوهش راهنمایی‌های روشنی برای مدیران و رهبران سازمان‌ها فراهم می‌کند:

- از سطوح بنیادین شروع کنید: مدیران ارشد باید پیش از استقرار سیستم‌های هوش مصنوعی، ابتدا یک «چارچوب سیاست‌گذاری» جامع، شفاف و اخلاقی تدوین کنند که به‌عنوان پیشران اصلی سیستم عمل کند.
- روی سازوکارهای ساختاری سرمایه‌گذاری کنید: سازمان‌ها باید به دنبال ابزارهایی باشند که در ساختار عملیاتی خود دارای قابلیت «تبیین‌پذیری» و «کنترل انسانی» هستند.
- اعتماد را در شبکه روابط مدیریت کنید: هدف اصلی باید ایجاد بستری برای پذیرش و اعتماد کارکنان باشد. مدیران باید از هوش مصنوعی به‌عنوان ابزاری برای توانمندسازی و رشد حرفه‌ای استفاده کنند.

علاوه‌براین، یافته‌های پژوهش نشان می‌دهد که مسئولیت استقرار هوش مصنوعی مسئولانه صرفاً بر عهده واحد فناوری اطلاعات نیست و نیازمند همکاری میان مدیران منابع انسانی، مدیران فناوری و مدیران ارشد سازمان است. از جنبه اجرایی، سازمان‌ها می‌توانند کمیته‌های حاکمیت هوش مصنوعی تشکیل دهند، ممیزی‌های دوره‌ای سوگیری الگوریتمی را انجام دهند، سازوکارهای اعتراض و بازبینی تصمیمات مبتنی بر هوش مصنوعی را برای کارکنان فراهم سازند و شاخص‌های اعتماد، ادراک انصاف و امنیت روان‌شناختی را به‌عنوان بخشی از نظام پایش عملکرد استقرار هوش مصنوعی مورد ارزیابی قرار دهند. چنین اقداماتی می‌تواند به کاهش مقاومت کارکنان در برابر فناوری و افزایش پذیرش و اثربخشی سامانه‌های هوش مصنوعی کمک کند.

هر پژوهشی با محدودیت‌هایی همراه است که باید در تفسیر یافته‌ها مورد توجه قرار گیرند.

نخست، مدل پیشنهادی این پژوهش بر پایه قضاوت و دانش تخصصی خبرگان توسعه یافته است. اگرچه بهره‌گیری از خبرگان در مطالعات اکتشافی و روش مدل‌سازی ساختاری - تفسیری رویکردی پذیرفته‌شده برای شناسایی و ساختاردهی پدیده‌های پیچیده و نوظهور محسوب می‌شود، اما ممکن است همه پیچیدگی‌ها و تفاوت‌های موجود در محیط‌های واقعی سازمانی را منعکس نکند. از این رو، یافته‌های پژوهش باید در چارچوب یک مدل ساختاری مبتنی بر ادراکات و قضاوت‌های تخصصی تفسیر شوند.

محدودیت دیگر، به ویژگی‌های پنل خبرگان پژوهش برمی‌گردد. نخست اینکه حجم نمونه خبرگان نسبتاً محدود بوده و مشارکت بخشی از آن‌ها در هر دو مرحله پژوهش، اگرچه به حفظ انسجام مفهومی کمک کرده، اما ممکن است احتمال شکل‌گیری «سوگیری توافقی»^۱ را افزایش داده باشد. همچنین، تمرکز نمونه‌گیری در این پنل بیشتر بر

تخصص‌های کارکردی (مدیریت منابع انسانی، توسعه هوش مصنوعی و اخلاق فناوری) بوده و تنوع صنایع مورد بررسی در آن محدود است. از آنجا که استقرار هوش مصنوعی در صنایع مختلف (نظیر بانکداری، نظام سلامت، تولید و...) می‌تواند چالش‌ها و پیامدهای متفاوتی برای شکوفایی کارکنان به همراه داشته باشد، این مسئله تعمیم‌پذیری دقیق نتایج به بافتارهای خاص صنعتی را با محدودیت مواجه می‌سازد. از این‌رو، پیشنهاد می‌شود در پژوهش‌های آتی علاوه بر استفاده از گروه‌های مستقل خبرگان در مراحل مختلف، مدل پیشنهادی با مشارکت طیف گسترده‌تری از ذی‌نفعان در صنایع مختلف مورد بررسی و آزمون قرار گیرد.

دوم، روش مدل‌سازی ساختاری - تفسیری با هدف شناسایی، ساختاردهی و سطح‌بندی روابط میان متغیرهای یک سیستم پیچیده به کار می‌رود و ماهیت آن مبتنی بر تحلیل قضاوت خبرگان است. بنابراین، روابط شناسایی‌شده در این پژوهش بیانگر الگوهای ساختاری، جهت‌دار و نفوذی میان مؤلفه‌ها بوده و نباید به منزله اثبات روابط علی در معنای آماری یا تجربی آن تلقی شوند. افزون بر این، مدل نهایی در مرحله کنونی مبتنی بر دانش و قضاوت خبرگان بوده و با داده‌های میدانی کارکنان یا سازمان‌ها مورد آزمون قرار نگرفته است؛ از این‌رو، فاقد اعتبارسنجی آماری و تحلیل برازش بوده و بیشتر ماهیتی اکتشافی و مفهومی دارد. همچنین، روش مدل‌سازی ساختاری-تفسیری کلاسیک امکان تحلیل حساسیت مبتنی بر تغییر وزن متغیرها را فراهم نمی‌کند، زیرا وزن عددی برای مؤلفه‌ها تعریف نمی‌شود. در نتیجه، پایداری ساختار مدل در برابر تغییرات احتمالی در شدت روابط میان مؤلفه‌ها مورد ارزیابی قرار نگرفت. پیشنهاد می‌شود پژوهش‌های آتی با استفاده از داده‌های تجربی و بهره‌گیری از روش‌هایی نظیر مدل‌سازی معادلات ساختاری، مدل‌سازی ساختاری - تفسیری فازی، میک‌مک فازی یا نگاشت شناختی فازی، اعتبار تجربی، برازش، پایداری و حساسیت مدل پیشنهادی را مورد بررسی قرار دهند.

سوم، پژوهش حاضر ماهیتی مقطعی دارد و روابط شناسایی‌شده را در یک مقطع زمانی مشخص بررسی کرده است. در حالی که فرایند استقرار و پذیرش هوش مصنوعی در سازمان‌ها ماهیتی پویا و تکاملی دارد، مطالعات طولی می‌توانند درک عمیق‌تری از چگونگی تحول این روابط در طول زمان ارائه دهند.

چهارم، مدل پیشنهادی در یک بستر فرهنگی و سازمانی خاص توسعه یافته است. از این‌رو، تعمیم نتایج به سایر صنایع، کشورها و زمینه‌های فرهنگی باید با احتیاط صورت گیرد. اجرای پژوهش‌های مشابه در محیط‌های سازمانی و فرهنگی متفاوت می‌تواند به ارزیابی پایداری و تعمیم‌پذیری ساختار پیشنهادی کمک کند.

بر این اساس، پیشنهاد می‌شود پژوهش‌های آینده با بهره‌گیری از داده‌های میدانی کارکنان و مدیران در سازمان‌های مختلف، مدل پیشنهادی را از طریق روش‌های کمی نظیر مدل‌سازی معادلات ساختاری مورد آزمون و تحلیل برازش قرار دهند تا جنبه‌های کاربردی و قابلیت تعمیم مدل ارتقا یابد. همچنین بررسی نقش متغیرهای زمینه‌ای و تعدیل‌گر، از جمله سواد دیجیتال کارکنان، آمادگی سازمانی برای پذیرش هوش مصنوعی، و شبکه‌های رهبری، می‌تواند به توسعه و غنای بیشتر مدل ارائه‌شده کمک کند.

در مجموع، یافته‌های این پژوهش نشان می‌دهد که استقرار مسئولانه هوش مصنوعی در مدیریت منابع انسانی صرفاً یک اقدام فناورانه نیست، بلکه فرایندی چندلایه و فنی - اجتماعی است که موفقیت آن به کیفیت تعامل میان

زیرسیستم فنی و زیرسیستم اجتماعی سازمان وابسته است. در چارچوب نظریه سیستم‌های اجتماعی - فنی، مؤلفه‌هایی نظیر شفافیت، تبیین‌پذیری، عدالت الگوریتمی، پاسخ‌گویی و نظارت انسانی از طریق شکل‌دهی ادراک کارکنان نسبت به اعتماد، انصاف و امنیت روان‌شناختی عمل می‌کنند. این سازوکارهای اجتماعی به نوبه خود زمینه لازم برای مشارکت فعال، معنا در کار و شکوفایی کارکنان را فراهم می‌سازند. از این منظر، شکوفایی کارکنان پیامد مستقیم فناوری نیست، بلکه حاصل هم‌ترازی موفق میان الزامات فنی سامانه‌های هوش مصنوعی و نیازهای انسانی کارکنان است. بنابراین، مدل سلسله‌مراتبی ارائه‌شده می‌تواند به‌عنوان چارچوبی راهنما برای طراحی و استقرار نظام‌های هوش مصنوعی انسان‌محور و مسئولانه در سازمان‌ها مورد استفاده قرار گیرد.

منابع

- حاجی اسماعیلی، حسین؛ طهماسبی، رضا و باباشاهی، جبار (۱۴۰۵). طراحی چارچوب توسعه استعداد فراگیر درسازمان (مورد مطالعه: سازمان تأمین اجتماعی). *مدیریت دولتی*، ۱۸(۱)، ۲۰۳-۲۳۳.
- رستگار، عباسعلی و موسوی، میناسادات (۱۴۰۲). بررسی اثر تنوع قومی اعضای تیم استارت‌آپ‌های ایرانی بر حجم سرمایه‌گذاری جمع‌آوری‌شده (مورد مطالعه: شرکت‌های حاضر در نمایشگاه GITEX دبی). *مدیریت دولتی*، ۱۵(۴)، ۷۵۱-۷۶۵.
- موسوی، میناسادات؛ رستگار، عباسعلی و شفیعی نیک‌آبادی، محسن (۱۴۰۴). هوش مصنوعی در مدیریت منابع انسانی با رویکرد مثبت‌گرا: چارچوبی برای شکوفایی کارکنان. *مطالعات مدیریت کسب‌وکار هوشمند*، ۱۴(۵۴)، ۱۷۵-۲۱۶.
- مؤمنی، امیررضا؛ یعقوبی، نورمحمد؛ روشن، سیدعلیقلی و پورعزت، علی اصغر (۱۴۰۵). طراحی الگوی پذیرش هوش مصنوعی در حکمرانی هوشمند با استفاده از رویکرد فراترکیب. *مدیریت دولتی*، ۱۸(۱)، ۲۶۳-۲۹۵.
- واعظی، سید کمال و پاشایی، فرانک (۱۴۰۴). کاربرد هوش مصنوعی برای تولید تلنگرهای رفتاری هوشمند. *مدیریت دولتی*، ۱۷(۴)، ۹۶۲-۹۳۸.

References

- Alawamleh, M., Al-Twal, A., Lahlouh, L. & Jame, R. O. (2023). Interpretive structural modelling of organizational innovation factors: An emerging market perspective. *Journal of Open Innovation: Technology, Market, and Complexity*, 9(2), 100067.
- Bankins, S., Formosa, P. & O’Kane, P. (2024). The ethical implications of artificial intelligence (AI) in human resource management: A review and framework. *Human Resource Management Review*, 34(1), 100940.
- Blau, P. M. (1964). *Exchange and power in social life*. Wiley.
- Braun, V. & Clarke, V. (2021). One size fits all? What counts as quality practice in (reflexive) thematic analysis? *Qualitative research in psychology*, 18(3), 328-352

- Budhwar, P., Malik, A., De Silva, M. T. & Thevisuthan, P. (2023). Artificial intelligence – challenges and opportunities for international HRM: a review and research agenda. *The International Journal of Human Resource Management*, 34(10), 1968-2007.
- Bunjak, A., Černe, M. & Popovič, A. (2023). Artificial intelligence and employees' well-being: A systematic literature review and future research agenda. *Journal of Business Research*, 162, 113886.
- Chang, Y. L. & Ke, J. (2024). Socially responsible artificial intelligence empowered people analytics: a novel framework towards sustainability. *Human Resource Development Review*, 23(1), 88-120.
- Cherns, A. (1976). The principles of sociotechnical design. *Human relations*, 29(8), 783-792.
- Chowdhury, S., Dey, P. K., Joel-Edgar, S., Bhattacharya, S., Rodriguez-Espindola, O., Abadie, A. & Truong, L. (2023). Unlocking the value of artificial intelligence in human resource management through AI capability framework. *Human Resource Management Review*, 33(1), 100899.
- Creswell, J. W. & Plano Clark, V. L. (2023). Revisiting mixed methods research designs twenty years later. *Handbook of mixed methods research designs*, 1(1), 21-36.
- Das, S. S., Pramod, V. R., & Kiron, K. (2024). Interpretive structural modeling (ISM): A literature review. *International Journal of Creative Research Thoughts*, 12(5), 863-868.
- European Commission. (2024). *Artificial Intelligence Act: Regulation of the European Parliament and of the Council laying down harmonised rules on artificial intelligence*. Official Journal of the European Union.
- Field, A. (2018). *Discovering statistics using IBM SPSS statistics*. Sage.
- Haji Esmaeili, H., Tahmasebi, R. & Babashahi, J. (2026). Designing an Inclusive Talent Development Framework in the Organization (Case Study: Social Security Organization). *Journal of Public Administration*, 18(1), 203-233. <https://doi.org/10.22059/jipa.2025.399679.3749> (in Persian)
- Khair, M. A., Mahadasa, R., Tuli, F. A. & Ande, J. R. P. K. (2020). Beyond human judgment: exploring the impact of artificial intelligence on HR decision-making efficiency and fairness. *Global Disclosure of Economics and Business*, 9(2), 163-176.
- Leicht-Deobald, U., Busch, T., Schank, C., Weibel, A., Scherer, S. & Schildhauer, T. (2022). The challenges of algorithm-based HR decision-making for personal integrity. *Journal of Business Ethics*, 179(3), 661-683.
- Malik, A., Budhwar, P. & Kazmi, B. A. (2023). Artificial intelligence (AI)-assisted HR decision making: A multi-level framework. *Human Resource Management Review*, 33(1), 100910.
- Momeni, A., Yaghoubi, N. M., Roshan, S. A. & Pourezzat, A. A. (2026). The Design of an Artificial Intelligence Adoption Model in Smart Governance Using a Meta-Synthesis Approach. *Journal of Public Administration*, 18(1), 263-295. doi: 10.22059/jipa.2025.398124.3737 (in Persian)

- Mousavi, M., Rastgar, A.A. & Shafiei Nikabadi, M. (2025). Artificial Intelligence in Human Resource Management with a Positive Approach: A Framework for Employee Flourishing. *Business Intelligence Management Studies*, 14(54), 175-216. doi: 10.22054/ims.2025.86244.2626 (in Persian)
- National Institute of Standards and Technology (NIST). (2023). Artificial Intelligence Risk Management Framework (AI RMF 1.0). U.S. Department of Commerce. <https://doi.org/10.6028/NIST.AI.100-1>
- Prikshat, V., Malik, A. & Budhwar, P. (2023). AI-augmented HRM: Antecedents, assimilation and multilevel consequences. *Human Resource Management Review*, 33(1), 100886.
- Rastgar, A. & Mousavi, M. (2023). Investigating the Effect of Ethnic Diversity of Iranian Startup Team Members on the Investment Capital Raised (Study Case: Startups at Dubai Gitex Exhibition). *Journal of Public Administration*, 15(4), 751-765. doi: 10.22059/jipa.2023.355429.3295 (in Persian)
- Sadeghi, S. (2024). *Employee Well-being in the Age of AI: Perceptions, Concerns, Behaviors, and Outcomes*. arXiv preprint arXiv:2412.04796.
- Schulte, P. A. & Streit, J. M. (2025). Psychosocial risks and ethical implications of technology: considerations for decent work. *Annals of Work Exposures and Health*, 69(4), 360-376.
- Soomro, S., Fan, M., Sohu, J. M., Soomro, S. & Shaikh, S. N. (2024). *AI adoption: a bridge or a barrier? The moderating role of organizational support in the path toward employee well-being*. Kybernetes.
- Spreitzer, G., Sutcliffe, K., Dutton, J., Sonenshein, S. & Grant, A. M. (2005). A socially embedded model of thriving at work. *Organization Science*, 16(5), 537-549. <https://doi.org/10.1287/orsc.1050.0153>
- Stahl, B. C., Antoniou, J., Bhalla, N., Brooks, L., Jansen, P., ... & Way, C. (2023). A systematic review of artificial intelligence impact assessments. *Artificial Intelligence Review*, 56(3), 2097-2127.
- Sushil. (2012). Interpreting the interpretive structural model. *Global Journal of Flexible Systems Management*, 13(2), 87-106.
- Tambe, P., Cappelli, P. & Yakubovich, V. (2019). Artificial intelligence in human resources management: Challenges and a path forward. *California Management Review*, 61(4), 15-42. <https://doi.org/10.1177/0008125619867910>
- Tripathi, A. & Kumar, V. (2025). Ethical practices of artificial intelligence: A management framework for responsible AI deployment in businesses. *AI and Ethics*, 1-12.
- Trist, E. L. (1981). The evolution of socio-technical systems: A conceptual framework and an action research program. *Occasional paper*, 2(1981), 1981.
- Vaezi, S. K. & Pashaei, F. (2025). Application of artificial intelligence for Generating Smart Behavioral Nudges. *Journal of Public Administration*, 17(4), 964-988. doi: 10.22059/jipa.2025.395752.3700 (in Persian)

- Valtonen, A., Kimpimäki, J. P. & Savela, N. (2026). Employee wellbeing: a computational review on the consequences of workplace automation. *Technovation*, 152, 103424.
- von Krogh, G. (2012). How does social software change knowledge management? Toward a strategic research agenda. *The Journal of Strategic Information Systems*, 21(2), 154–164. <https://doi.org/10.1016/j.jsis.2012.04.003>
- Vrontis, D., Christofi, M., Pereira, V., Tarba, S., Makrides, A. & Trichina, E. (2023). Artificial intelligence, robotics, advanced technologies and human resource management: a systematic review. *Artificial intelligence and international HRM*, 172-201. 55.
- Weber, R. (2021). Constructs and Indicators: An Ontological Analysis. *MIS Quarterly*, 45(4).