



Designing a symbiotic knowledge system based on generative artificial intelligence: From positive human resource management to strategic collective intelligence

Taimoor Marjani^{1*}, Ali Davari², Ali Sadollah³

1. Assistant Professor, Department of Management, Faculty of Humanities, University of Science and Culture, Tehran, Iran, tmajani@usc.ac.ir

2. Associate Professor, Department of New Business, Faculty of Entrepreneurship, College of Management, University of Tehran, Tehran, Iran, ali_davari@ut.ac.ir

3. Assistant Professor, Department of Mechanical Engineering, Faculty of Engineering, University of Science and Culture, Tehran, Iran, sadollah@usc.ac.ir

Corresponding author: tmajani@usc.ac.ir



<https://doi.org/10.47176/SMOK.2026.2029>

Article Info

Article history:

Received: 28 May 2026

Revised: 23 June 2026

Accepted: 03 July 2026

Published:

Keywords:

Symbiotic Knowledge System, Agent-Based Modeling, Positive Human Resource Management, Strategic Collective Intelligence, Generative Artificial Intelligence.

ABSTRACT

Purpose: Generative artificial intelligence has created opportunities for knowledge management and human resource management, yet a model that synergistically integrates employee flourishing and strategic knowledge creation is rare. This study aims to design a symbiotic knowledge system model based on generative artificial intelligence and to explain its causal relationships.

Methodology: An exploratory sequential mixed-method approach was adopted in two phases. In the qualitative phase, fuzzy Delphi with 15 experts extracted the initial components. In the quantitative phase, fuzzy DEMATEL analyzed causal relationships, and fuzzy cognitive mapping combined with agent-based simulation over a five-year horizon modeled system dynamics. Content validity was confirmed using CVR, and reliability was confirmed using Kendall's coefficient (0.73).

Results: Five components were identified: algorithmic transparency, cognitive diversity preservation, human-AI interaction feedback, generative personalized learning, and ethical AI governance. "Cognitive diversity preservation" with the highest causality (+3.78) was the strongest causal driver. Simulation showed the full system increases strategic knowledge creation by 87%, reduces burnout from 34% to 12%, and improves decision-making accuracy by 48%.

Discussion: Contrary to traditional approaches, cognitive diversity preservation drives knowledge-wellbeing synergy. Algorithmic transparency enables trust and feedback, while personalized learning operates under ethical governance.

Conclusion: This study provides an integrated framework linking positive HRM, generative AI, and strategic collective intelligence. Organizations should prioritize cognitive diversity preservation and algorithmic transparency in implementation. Generalizability is limited to Iranian knowledge-based industries.

How to cite this article: Marjani, T., Davari, A., Sadollah, A. (2026). Designing a symbiotic knowledge system based on generative artificial intelligence: From positive human resource management to strategic collective intelligence. *Strategic Management of Organizational Knowledge*, 9 (3), 00-00. <https://doi.org/10.47176/SMOK.2026.2029>

2645-5242/© 2026 © The Author(s) retain the copyright. Published by Imam Hossein University, Iran.

This is an open-access article under the CC-BY 4.0 license. (<https://creativecommons.org/licenses/by/4.0/>)



Introduction

The rapid advancement of generative artificial intelligence (GenAI) is fundamentally reshaping how organizations create, share, and utilize knowledge. Unlike earlier AI systems that primarily focused on pattern recognition and prediction, GenAI models possess the unique capability to generate novel content—including text, images, code, and complex designs—based on patterns learned from training data (Ali & Mahmood, 2025; Rodrigues, 2025). Simultaneously, positive human resource management (PHRM) has emerged as a paradigm that prioritizes employee flourishing, psychological capital, and well-being alongside organizational performance (Gruman & Budworth, 2022; Van Zyl & Rothmann, 2022). While each of these two streams has independently generated substantial scholarly interest, a critical theoretical gap persists: no prior study has systematically integrated GenAI, PHRM, and strategic collective intelligence into a unified, dynamic model. Existing literature either focuses on AI-driven efficiency in HR processes (Tambe et al., 2019; Priksht et al., 2023) or on isolated well-being interventions without considering technological advancements (Peccei & Van De Voorde, 2019). What remains unexplored is the synergistic potential of a symbiotic knowledge system in which GenAI and human intelligence mutually reinforce each other.

To guide the empirical investigation of this gap, an initial theoretical framework was developed based on the theoretical foundations (Self-Determination Theory and Distributed Cognition Theory) and the synthesis of the research background. This framework proposes a set of potential components that appear to shape the "symbiotic knowledge system based on generative AI": algorithmic transparency (C1), cognitive diversity preservation (C2), human-AI interaction feedback (C3), generative personalized learning (C4), and ethical AI governance (C5). Moreover, the anticipated organizational outcomes include strategic knowledge creation, employee flourishing, and strategic collective intelligence. It is crucial to emphasize that this is a preliminary conceptualization, not a finalized model. These components and relationships are initial assumptions that will be empirically validated, refined, or rejected by experts in the subsequent fuzzy Delphi phase. Figure 1 provides a schematic representation of this initial framework.

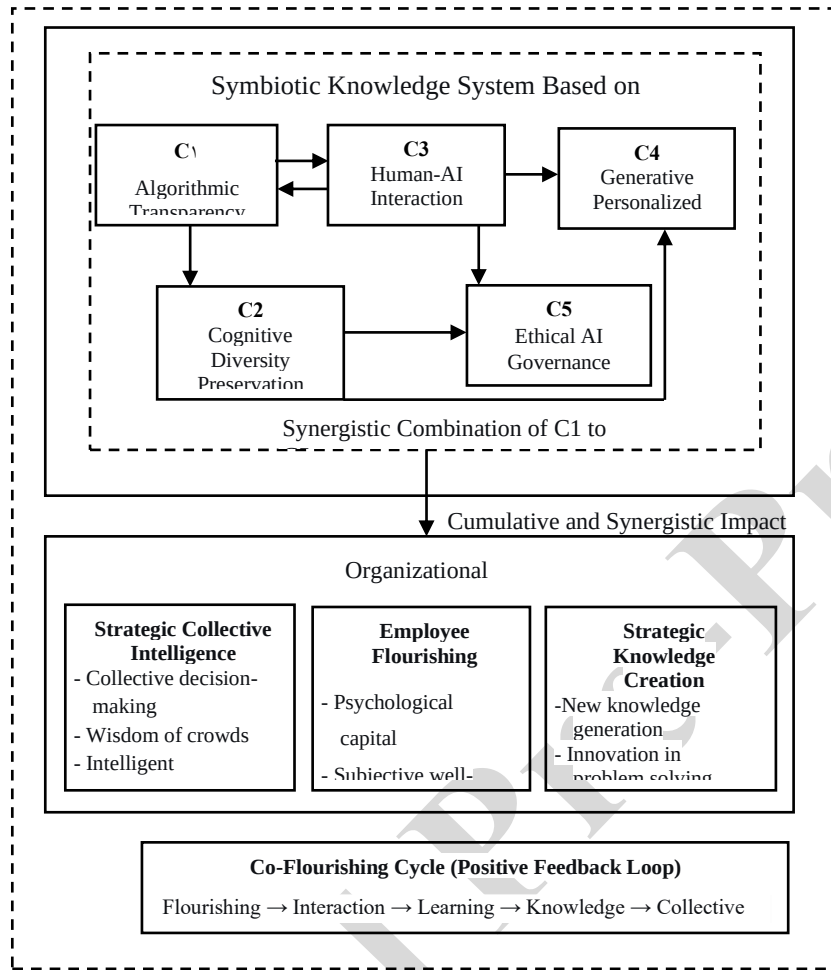


Fig. 1. Initial theoretical framework of the research.

Following this conceptualization, the main objective of this study is to design a symbiotic knowledge system model based on generative AI. The research answers three questions: (1) What are the core components of a symbiotic knowledge system? (2) What are the causal relationships among these components? (3) What impact does implementing such a system have on strategic knowledge creation, employee flourishing, and organizational collective intelligence? The remainder of this article is structured as follows: Section 2 presents the theoretical literature and research background; Section 3 describes the research methodology; Section 4 reports the findings; Section 5 provides discussion; and Section 6 concludes with limitations and recommendations.

Methodology

This study employs an exploratory sequential mixed-method design grounded in pragmatism, comprising three stages. This approach is appropriate because the research aims to first explore and identify the core components of a symbiotic knowledge system qualitatively, and then quantitatively analyze the causal relationships and simulate the system's dynamics. The qualitative stage uses fuzzy Delphi, while the quantitative stage combines fuzzy DEMATEL, fuzzy cognitive mapping, and agent-based modeling.

Stage 1 (Qualitative – Fuzzy Delphi): A systematic literature review following PRISMA guidelines was conducted using Scopus, Web of Science, Science Direct, Emerald, Magiran, and Noormags. The search string combined "generative AI," "positive HRM," "collective intelligence," "knowledge management," and "human-AI symbiosis" for the period 2020–2026. After screening 378 articles, 94 were selected, yielding an initial pool of 34 components. A panel of 15 experts (6 HRM, 4 knowledge management, 3 AI, 2 positive psychology) was formed using purposive and snowball sampling.

Selection criteria included a Ph.D., at least 5 years of experience, and a minimum of 3 peer-reviewed publications. Over three Delphi rounds, experts rated components using a five-point fuzzy linguistic scale (very low to very high). Triangular fuzzy numbers were defuzzified via center-of-gravity, with retention criterion of defuzzified mean ≥ 0.75 and Kendall's W significant ($p < 0.05$). Content validity was assessed with 5 independent experts (CVR critical value = 0.99; CVI average = 0.94). Reliability was confirmed with Kendall's W = 0.73 ($p < 0.01$).

Stage 2 (Quantitative causal modeling – Fuzzy DEMATEL): The five final components were structured into a pairwise comparison matrix. The same 15 experts evaluated direct influences using a five-point fuzzy linguistic scale (no influence to very high influence). The fuzzy DEMATEL procedure involved forming a direct-relation fuzzy matrix, normalization, computing the total-relation matrix ($\tilde{T} = \tilde{N} (I - \tilde{N})^{-1}$), defuzzification using the averaging method, and setting a threshold of $\theta = 0.5780$. Influence (D_i) and susceptibility (R_i) indices were calculated to determine centrality ($D_i + R_i$) and causality ($D_i - R_i$). Reliability was assessed using Kendall's W = 0.69 ($p < 0.01$) and an average inconsistency rate below 0.1.

Stage 3 (Dynamic simulation – Fuzzy Cognitive Mapping + Agent-Based Modeling): Causal weights from the total-relation matrix were imported into FCMapper to visualize the fuzzy cognitive map, with nodes representing the five components and directed edges weighted in $[-1, 1]$. Consensus centrality was computed for each node. An agent-based model was developed in NetLogo 6.4 with 100 agents representing employees, each with variables for psychological capital (PsyCap), burnout, and engagement. Initial PsyCap scores were derived from a survey ($N=120$) at a knowledge-based IT firm in Iran. Four scenarios were tested over 60 ticks (5 years): baseline (no system), Scenario 1 (only C1), Scenario 2 (only C2), Scenario 3 (C1+C2), and Scenario 4 (all five components). Each scenario was replicated 50 times. Sensitivity analysis was performed with $\pm 20\%$ parameter variations. Model validation compared baseline scenario output with historical data (2023–2025) from the participating firm, achieving a mean absolute error (MAE) below 8%. The NetLogo code is provided as supplementary material. Ethical considerations included voluntary participation, confidentiality, and anonymity of data. No AI tools were used for analysis or scientific inference; the authors bear full responsibility for the content.

Results

This section presents the findings from the three methodological stages—fuzzy Delphi, fuzzy DEMATEL, and agent-based simulation—in a clear, objective manner without interpretation.

Fuzzy Delphi Results

From the initial pool of 34 components extracted from the systematic literature review, after three rounds of fuzzy Delphi with 15 experts, 29 components were eliminated and 5 were confirmed as the final core components. The reasons for elimination were: low defuzzified mean (< 0.75) – 22 components; conceptual overlap with main components – 4 components; and lack of expert consensus (coefficient of variation > 0.3) – 3 components. A complete list of the original 34 components with their defuzzified means across all three rounds and elimination reasons is provided in Appendix 1. The five final components are presented in Table 1.

Table 1. Final Fuzzy Delphi Results (Third Round).

Code	Component	Defuzzified Mean	Std. Deviation	CV
C1	Algorithmic Transparency	0.86	0.07	0.08
C2	Cognitive Diversity Preservation	0.91	0.06	0.07
C3	Human-AI Interaction Feedback	0.79	0.10	0.13
C4	Generative Personalized Learning	0.88	0.07	0.08
C5	Ethical AI Governance	0.84	0.08	0.10

All five components achieved a defuzzified mean ≥ 0.75 . Kendall's W coefficient at the end of the third round was 0.73 ($p < 0.01$), indicating significant agreement among experts. The content validity ratio (CVR) for each of the five components ranged from 0.78 to 0.92, above the critical value of 0.49 (for 3 experts).

Fuzzy DEMATEL Results

After converting the linguistic evaluations of the 15 experts into triangular fuzzy numbers, normalizing, and computing the total-relation matrix, the defuzzified matrix was obtained by applying a threshold of $\theta = 0.5780$. Table 2 presents the final defuzzified total-relation matrix (values below the threshold were set to zero).

Table 2. Defuzzified Total-Relation Matrix (T).

	C1	C2	C3	C4	C5
C1	0.00	0.52	0.78	0.34	0.41
C2	0.61	0.00	0.55	0.72	0.63
C3	0.32	0.44	0.00	0.48	0.37
C4	0.28	0.33	0.51	0.00	0.44
C5	0.35	0.29	0.42	0.38	0.00

Table 3 reports the influence (D_i), susceptibility (R_i), centrality (D_i+R_i), and causality (D_i-R_i) indices for each component. Components with positive D_i-R_i (C1 and C2) are classified as causal (influencing), while those with negative D_i-R_i (C3, C4, and C5) are effect (influenced) variables.

Table 3. Influence, Susceptibility, Centrality, and Causality Indices.

Component	D_i	R_i	D_i+R_i	D_i-R_i	Role
C1	5.42	4.18	9.60	+1.24	Causal
C2	6.83	3.05	9.88	+3.78	Causal
C3	3.91	6.14	10.05	-2.23	Effect
C4	4.56	5.42	9.98	-0.86	Effect
C5	4.22	5.15	9.37	-0.93	Effect

As shown in Table 3, C2 (Cognitive Diversity Preservation) had the highest net causal effect ($D_i-R_i = 3.78$), followed by C1 (1.24). C3, C4, and C5 were effect variables with negative D_i-R_i values. Centrality (D_i+R_i) was highest for C3 (10.05), indicating its strong connective role in the network. All causal weights were positive, indicating reinforcing relationships among components. Fig. 2 displays the causal relationship scatter plot based on these indices.

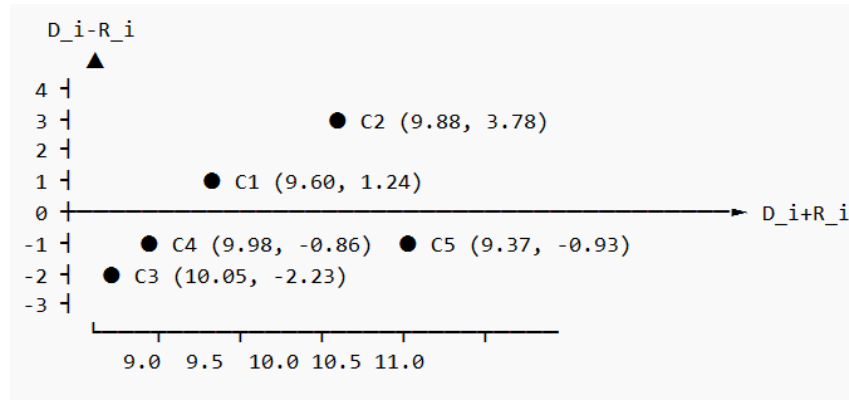


Fig. 2. Causal Relationship Scatter Plot.

Fuzzy Cognitive Mapping Results

Consensus centrality indices for each component in the fuzzy cognitive map are presented in Table 4. C2 (Cognitive Diversity Preservation) had the highest consensus centrality (17.42), confirming its central role in the network, followed by C4 (15.87) and C3 (15.32).

Table 4. Consensus Centrality in FCM.

Component	CenD	CenC	CenB	Consensus Centrality
C1	8.23	3.45	1.18	12.86
C2	11.05	4.12	2.25	17.42
C3	9.87	3.89	1.56	15.32
C4	10.24	3.91	1.72	15.87
C5	9.45	3.76	1.43	14.64

Agent-Based Simulation Results

Agent-based simulation with 100 agents over a 60-tick horizon (5 years) was conducted for four scenarios. Table 5 reports the mean values of the three outcome variables at the end of the simulation period for each scenario (results based on 50 replications). Detailed equations and model validation procedures are provided in the supplementary material.

Table 5. Agent-Based Simulation Results for Different Scenarios.

Scenario	Strategic Knowledge Creation (Index 0-100)	Burnout Rate (%)	Decision Accuracy (%)
Baseline (no system)	100.0 ± 0.0	34.2 ± 1.8	50.0 ± 2.1
Scenario 1 (C1 only)	118.3 ± 2.4	29.1 ± 1.5	56.4 ± 1.9
Scenario 2 (C2 only)	134.0 ± 2.1	22.4 ± 1.3	63.1 ± 1.7
Scenario 3 (C1+C2)	156.2 ± 1.9	17.8 ± 1.1	70.2 ± 1.5
Scenario 4 (All five)	187.4 ± 1.6	11.9 ± 0.9	74.0 ± 1.3

Under the full system (Scenario 4), strategic knowledge creation capacity increased by 87% compared to baseline (from 100 to 187.4 index points). Employee burnout rate dropped from 34.2% to 11.9% (a 22.3 percentage-point reduction). Strategic decision accuracy improved by 48% (from 50.0% to 74.0% correct). Even the partial scenario with only C2 (cognitive diversity preservation) yielded a 34%

increase in knowledge creation and a 12.4% reduction in burnout. Fig. 3 illustrates the growth trends in strategic knowledge creation capacity across all scenarios over the 60-month simulation period.

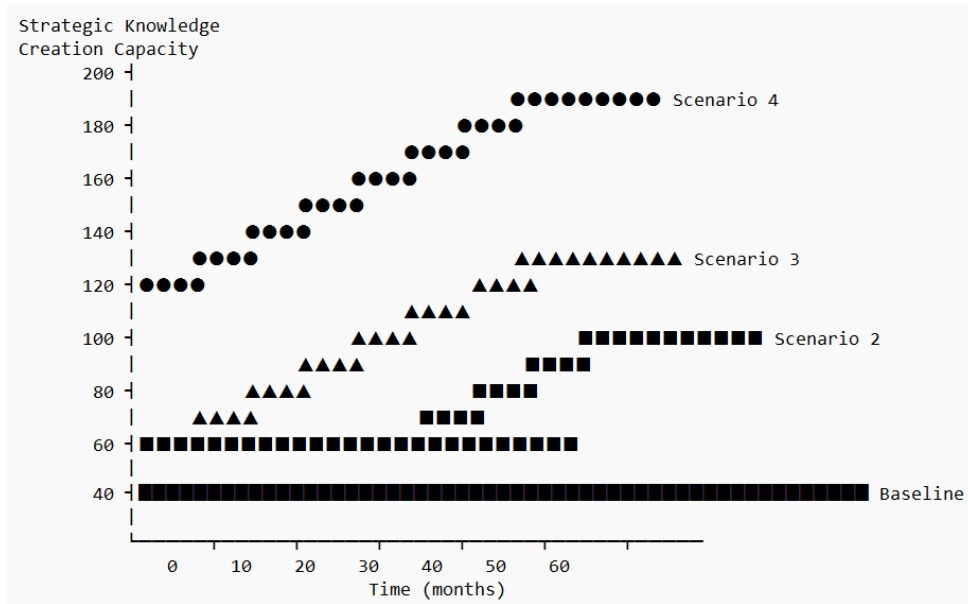


Fig. 3. Trends in Strategic Knowledge Creation Capacity Over Time.

Sensitivity Analysis

Sensitivity analysis was performed by varying key input parameters by $\pm 20\%$. Table 6 presents the sensitivity coefficients (change in outcome variable per unit change in input parameter).

Table 6. Sensitivity Coefficients of Key Parameters.

Input Parameter	Sensitivity Coefficient (Knowledge Creation)	Sensitivity Coefficient (Burnout)
C2→C4 relationship weight	0.72	-0.65
C1→C3 relationship weight	0.54	-0.48
Initial psychological capital	0.38	-0.52
Burnout threshold	0.21	0.35

The highest sensitivity coefficient was for the C2→C4 relationship weight (0.72 for knowledge creation and -0.65 for burnout), confirming that C2 (Cognitive Diversity Preservation) is the most influential parameter in the model, followed by C1 (0.54). The co-flourishing loop exhibited clear reinforcement: reduced burnout → higher psychological capital → increased engagement with GenAI → more novel knowledge outputs → further personalized learning and well-being gains.

Discussion

This study aimed to design a symbiotic knowledge system model based on generative AI and to explain the causal relationships among its components. The findings provide several important theoretical and practical insights that extend the existing literature.

First, the identification of five core components—algorithmic transparency, cognitive diversity preservation, human-AI interaction feedback, generative personalized learning, and ethical AI governance—confirms and expands upon earlier work. While transparency and ethical governance have been previously emphasized in AI implementation (Raji et al., 2022; Al-Hawamdeh & Al-Jabri, 2025), this study uniquely positions cognitive diversity preservation as a central causal driver rather than a mere contextual factor. This result aligns with warnings about the homogenization risks of large

language models (Cramarenco et al., 2023) and extends the positive HRM framework by showing how technology can amplify rather than suppress human variation.

Second, the causal dominance of cognitive diversity preservation ($D-R = +3.78$) challenges conventional AI design paradigms that prioritize optimization and efficiency above all else (Tambe et al., 2019). Organizations focusing narrowly on efficiency may inadvertently reduce cognitive diversity and, consequently, their capacity for innovation. This finding resonates with collective intelligence research indicating that cognitive diversity is a prerequisite for wisdom of crowds (Surowiecki, 2004; Smith et al., 2025), yet it goes further by empirically quantifying the causal weight of this variable within AI-based knowledge systems.

Third, algorithmic transparency (C1) emerged as the second most influential causal factor ($D-R = +1.24$) and exhibited the strongest direct effect on feedback loops ($C1 \rightarrow C3 = 0.78$). This suggests that transparency is not merely a legal or technical requirement but a fundamental enabler of trust and effective feedback. Without transparency, employee feedback becomes meaningless, and the system cannot adapt. This aligns with the trust-complementarity model (Dittmar & Sposato, 2026) and reinforces self-determination theory (Deci & Ryan, 1985).

Fourth, the simulation results demonstrate that employee flourishing and strategic knowledge creation are mutually reinforcing, contradicting the traditional productivity-well-being trade-off assumption (Guest, 2017). The full system achieved an 87% increase in knowledge creation and a 22 percentage-point reduction in burnout, suggesting that well-being and productivity can co-evolve positively when systems are designed with transparency, diversity preservation, and ethical governance. Comparison with prior studies reveals both alignment and divergence. Our findings are consistent with research on organizational support and justice (Hajipour & Amirkhani, 2021), but we extend these by embedding them within spiritual-ethical and identity-based value frameworks. Unlike previous models that focus primarily on structural and managerial factors (Asghari, 2022; Ghorabaghi & Azgoli, 2020), the present framework transforms these components into a socio-ethical and cultural paradigm. Theoretically, this study extends distributed cognition theory to the GenAI era and enriches self-determination theory by operationalizing how each component satisfies autonomy, competence, and relatedness needs. Practically, HR managers should prioritize cognitive diversity preservation in AI procurement, mandate algorithmic transparency, design bidirectional feedback mechanisms, and implement a phased approach starting with C2 and C1 components.

Conclusion

This study successfully designed a symbiotic knowledge system model based on generative AI, integrating positive HRM, generative AI, and strategic collective intelligence. Using an exploratory sequential mixed-method design, five core components—algorithmic transparency, cognitive diversity preservation, human-AI interaction feedback, generative personalized learning, and ethical AI governance—were identified through fuzzy Delphi with 15 experts. Fuzzy DEMATEL revealed cognitive diversity preservation ($D-R = +3.78$) and algorithmic transparency ($+1.24$) as primary causal drivers, while human-AI interaction feedback functioned as the central communication hub (centrality 10.05). Agent-based simulation over a five-year horizon showed full system implementation increases strategic knowledge creation by 87%, reduces burnout from 34% to 12%, and improves decision accuracy by 48%.

For Iranian public organizations, this model offers a culturally grounded framework. Managers should develop self-efficacy training, empathetic guidelines, identity-building systems, service-orientation frameworks, spiritual-ethical training, employee support systems, delegation protocols, charisma development programs, and strategic vision management. Converting Martyr Soleimani's lived experiences into explicit knowledge can profoundly enhance employee motivation.

Future research should pursue empirical validation through longitudinal case studies across industries, develop algorithms for operationalizing cognitive diversity preservation, conduct cross-cultural

comparative studies, investigate unintended consequences (algorithmic fatigue, over-reliance), and extend the model to inter-organizational networks.

Acknowledgments

The authors sincerely thank the 15 expert panel members for their dedicated participation in the fuzzy Delphi and fuzzy DEMATEL processes. We also acknowledge the cooperation of the participating knowledge-based IT firm in Iran for providing baseline employee psychological capital data used to calibrate the agent-based simulation.

Funding

This research received no specific grant from any funding agency in the public, commercial, or not-for-profit sectors. The authors declare that no financial support, research contracts, grants, institutional support, economic interests, professional consulting, board membership, or institutional affiliations influenced the design, execution, analysis, or reporting of this study. The organization that provided baseline data (the knowledge-based IT firm) had no role in study design, data analysis, interpretation, or the decision to publish the results.

Conflicts of Interest

The authors declare no conflicts of interest. None of the authors have any financial, organizational, professional, personal, or intellectual relationships that could directly or indirectly influence the design, execution, analysis, or reporting of the research results or create the appearance of such relationships. This declaration has been prepared with the knowledge and agreement of all co-authors, and the corresponding author confirms its accuracy.

Author Contributions

Taimoor Marjani: Conceptualization, Methodology, Formal analysis, Writing – original draft, Visualization; Ali Davari: Investigation, Validation, Writing – review & editing, Supervision; Ali Sadollah: Data curation, Software (FCMapper, NetLogo), Simulation, Project administration. All authors have read and approved the final version of this manuscript and agree to be accountable for all aspects of the work. No guarantor was designated; each author takes responsibility for their respective contributions.



مدیریت راهبردی دانش سازمانی

Journal homepage: <https://jkm.ihu.ac.ir/>

مقاله (اصیل)

طراحی مدل سامانه دانشی همزیست مبتنی بر هوش مصنوعی مولد: از مدیریت منابع انسانی مثبت‌گرا تا خلق هوش جمعی راهبردی

تیمور مرجانی^{۱*}، علی داوری^۲، علی سعداله^۳۱. استادیار، گروه مدیریت، دانشکده علوم انسانی، دانشگاه علم و فرهنگ، تهران، ایران، tmarjani@usc.ac.ir۲. دانشیار، گروه کسب و کار جدید، دانشکده کارآفرینی، دانشکده‌گان مدیریت، دانشگاه تهران، تهران، ایران، ali_davari@ut.ac.ir۳. استادیار، گروه مهندسی مکانیک، دانشکده فنی و مهندسی، دانشگاه علم و فرهنگ، تهران، ایران، sadollah@usc.ac.ir

تاریخ دریافت: ۰۷ خرداد ۱۴۰۵؛ تاریخ بازنگری: ۰۲ تیر ۱۴۰۵؛ تاریخ پذیرش: ۱۲ تیر ۱۴۰۵؛ تاریخ انتشار:

چکیده

هدف: هوش مصنوعی مولد فرصت‌هایی برای مدیریت دانش و منابع انسانی ایجاد کرده، اما مدلی که هم‌افزا شکوفایی کارکنان و خلق دانش راهبردی را تلفیق کند، نادر است. این پژوهش با هدف طراحی مدل سامانه دانشی همزیست مبتنی بر هوش مصنوعی مولد و تبیین روابط علی مؤلفه‌های آن انجام شد.

روش پژوهش: پژوهش با رویکرد آمیخته اکتشافی در دو مرحله انجام شد. در مرحله کیفی، دلفی فازی با ۱۵ خبره مؤلفه‌های اولیه را استخراج کرد. در مرحله کمی، دیمتال فازی روابط علی را تحلیل و نگاشت شناختی فازی و شبیه‌سازی عامل - محور در افق پنج‌ساله پویایی‌های سامانه را مدل‌سازی کرد. اعتبار محتوا با نسبت روایی محتوایی و پایایی با ضریب کندال (۰/۷۳) تأیید شد.

یافته‌ها: پنج مؤلفه شناسایی شد؛ شفافیت الگوریتمی، حفظ تنوع شناختی، بازخورد تعامل انسان-ماشین، یادگیری شخصی‌سازی‌شده مولد و حکمرانی اخلاقی هوش مصنوعی. «حفظ تنوع شناختی» با بیشترین علیت (۳/۷۸+) قوی‌ترین متغیر علی بود. شبیه‌سازی نشان داد سامانه کامل خلق دانش راهبردی را ۸۷ درصد افزایش، فرسودگی شغلی را از ۳۴ به ۱۲ درصد کاهش و دقت تصمیم‌گیری را ۴۸ درصد بهبود می‌بخشد.

بحث: برخلاف رویکردهای سنتی، حفظ تنوع شناختی موتور هم‌افزایی دانش و به‌زیستی است. شفافیت الگوریتمی اعتماد و بازخورد را ممکن ساخته و یادگیری شخصی‌شده را تحت حکمرانی اخلاقی هدایت می‌کند.

نتیجه‌گیری: این پژوهش چارچوبی یکپارچه برای پیوند مدیریت منابع انسانی مثبت‌گرا، هوش مصنوعی مولد و هوش جمعی راهبردی ارائه می‌دهد. سازمان‌ها باید حفظ تنوع شناختی و شفافیت الگوریتمی را در پیاده‌سازی اولویت دهند. تعمیم‌پذیری به صنایع دانش‌بنیان ایران محدود است.

کلیدواژه‌ها: سامانه دانشی همزیست، شبیه‌سازی عامل - محور، مدیریت منابع انسانی مثبت‌گرا، هوش جمعی راهبردی، هوش مصنوعی مولد.

مقدمه

عصر کنونی با شتابی بی‌سابقه در تحول دیجیتال و گسترش فناوری‌های هوشمند همراه است. در این میان، هوش مصنوعی مولد^۱ به‌عنوان یکی از تأثیرگذارترین دستاوردهای دههٔ اخیر، نه‌تنها فرآیندهای عملیاتی سازمان‌ها را دگرگون ساخته، بلکه ماهیت خلق، اشتراک و به‌کارگیری دانش سازمانی را نیز متحول کرده است (Ali & Mahmood, 2025; Rodrigues, 2025). از سوی دیگر، مدیریت منابع انسانی مثبت‌گرا با تأکید بر شکوفایی کارکنان، سرمایهٔ روانشناختی و به‌زیستی در محیط کار، رویکردی نوین را در مدیریت سرمایه‌های انسانی ترسیم نموده است که در آن امید، تاب‌آوری، خوش‌بینی و خودکارآمدی کارکنان به‌عنوان پیشران‌های اصلی عملکرد و نوآوری در نظر گرفته می‌شوند (Mousavi et al., 2025; Hidayat et al., 2025). هر یک از این دو جریان به‌تنهایی توانسته است تحولات قابل‌توجهی ایجاد کند: هوش مصنوعی مولد با قابلیت تولید محتوای جدید از متن و تصویر تا کد و طرح‌های پیچیده فرآیندهای حل مسئله و نوآوری را دگرگون ساخته است (Sharma, 2026)؛ و مدیریت منابع انسانی مثبت‌گرا با پرورش مؤلفه‌های روانشناختی، محیط‌هایی پدید آورده که در آن کارکنان هم به‌لحاظ فردی و هم به‌لحاظ حرفه‌ای رشد می‌کنند (Gruman & Budworth, 2022; Van De Voorde et al., 2012).

با وجود پیشرفت‌های چشمگیر در هر یک از این حوزه‌ها به‌طور جداگانه، پرسش بنیادینی همچنان نیازمند بررسی بیشتر است: چگونه می‌توان هوش مصنوعی مولد و مدیریت منابع انسانی مثبت‌گرا را به گونه‌ای درهم تنید که نه تنها به‌زیستی کارکنان را افزایش دهد، بلکه به خلق هوش جمعی راهبردی در سازمان نیز منجر شود؟ پاسخ به این پرسش مستلزم طراحی الگویی است که از یک سو، از قابلیت‌های زایشی و تحلیلی هوش مصنوعی برای تولید دانش جدید بهره گیرد و از سوی دیگر، تنوع شناختی، شفافیت الگوریتمی و بازخورد انسانی را به‌عنوان پیش‌نیازهای اساسی حفظ کند (Dittmar & Sposato, 2026; Garcia & Patel, 2025).

اهمیت این پژوهش در آن است که سازمان‌های امروزی با دو چالش اساسی روبه‌رو هستند: از یک سو، فشار برای استفاده از فناوری‌های نوین هوش مصنوعی به منظور کاهش هزینه‌ها و افزایش سرعت؛ از سوی دیگر، نگرانی فزاینده در مورد پیامدهای انسانی این فناوری‌ها مانند فرسودگی شغلی، کاهش معنا در کار و یکنواختی دانش. به نظر می‌رسد بدون یک الگوی جامع که بتواند هوش مصنوعی را در خدمت شکوفایی کارکنان و خلق دانش راهبردی قرار دهد، سازمان‌ها یا در دام کارآمدی‌گرایی صرف گرفتار می‌شوند یا از ظرفیت‌های تحول‌آفرین هوش مصنوعی عقب می‌مانند. از این رو، ضرورت دارد تا الگویی طراحی شود که نه تنها جنبه‌های فنی، بلکه ابعاد انسانی، اخلاقی و شناختی را نیز در بر گیرد. این پژوهش در پاسخ به همین ضرورت شکل گرفته است.

مرور نظام‌مند ادبیات نشان می‌دهد که پژوهش‌های موجود عمدتاً در سه دسته جای می‌گیرند: دسته اول به بهینه‌سازی فرآیندهای عملیاتی منابع انسانی با هوش مصنوعی (مانند جذب، ارزیابی عملکرد و جانشین‌پروری) پرداخته‌اند (Tambe et al., 2019; Prikshat et al., 2023)؛ دسته دوم بر پیامدهای روانشناختی مثبت در محیط کار بدون توجه به فناوری‌های نوین متمرکز بوده‌اند (Peccei & Van De Voorde, 2019)؛ و دسته سوم، مدل‌های علی روابط میان شاخص‌های هوش مصنوعی و منابع انسانی را ترسیم کرده‌اند (Mousavi et al., 2025). با این حال، چنین می‌نماید تاکنون مطالعه‌ای که به طراحی سامانه‌ای پویا، هم‌افزا و دانش‌آفرین با تلفیق همزمان سه مؤلفهٔ «مدیریت منابع انسانی مثبت‌گرا»، «هوش مصنوعی مولد» و «هوش جمعی راهبردی» پرداخته باشد، به میزان کافی مورد توجه قرار نگرفته است؛ و اگر هم آثاری وجود داشته باشند، بسیار اندک و پراکنده‌اند. به‌ویژه، مفهوم سامانهٔ دانشی همزیست^۲ – یعنی سیستمی که در آن هوش مصنوعی و انسان در یک چرخهٔ دوسویه، توانایی‌های یکدیگر را تقویت می‌کنند – تاکنون به‌طور نظام‌مند عملیاتی نشده است. همچنین، نقش حفظ تنوع شناختی^۳ به‌عنوان موتور محرک این هم‌افزایی و رابطهٔ علی آن با مؤلفه‌هایی چون شفافیت الگوریتمی، بازخورد تعامل انسان-ماشین، یادگیری شخصی‌سازی‌شده و حکمرانی اخلاقی هوش مصنوعی در هیچ مدلی به طور کامل تبیین نشده است (Raji et al., 2022; Smith et al., 2025). ضرورت تلفیق سه حوزهٔ مدیریت منابع انسانی مثبت‌گرا، هوش مصنوعی مولد و هوش جمعی راهبردی در قالب یک سامانهٔ دانشی همزیست، انگیزهٔ اصلی پژوهش حاضر را شکل می‌دهد. آنچه تاکنون کمتر مورد توجه قرار گرفته است، درک این نکته می‌باشد که هوش مصنوعی و انسان می‌توانند در یک چرخهٔ دوطرفه، توانایی‌های یکدیگر را تقویت کنند: هوش مصنوعی، توانمندی‌های شناختی انسان را از طریق شخصی‌سازی یادگیری و تحلیل الگوهای پیچیده گسترش می‌دهد (Ellikkal & Rajamohan, 2026).

¹ Generative Artificial Intelligence

² Symbiotic Knowledge System

³ Cognitive Diversity Preservation

(2024) و انسان نیز با حفظ تنوع دیدگاه‌ها، قضاوت‌های اخلاقی و بازخوردهای اخلاقانه، هوش مصنوعی را از یکنواختی، سوگیری و فرسایش خلاقیت نجات می‌دهد (Chen & Wang, 2026; Thompson et al., 2026). بر این اساس، هدف اصلی پژوهش حاضر، طراحی مدل سامانه دانشی همزیست مبتنی بر هوش مصنوعی مولد است که به پرسش‌های زیر پاسخ دهد:

۱. مؤلفه‌های کلیدی سامانه دانشی همزیست کدام‌اند؟
 ۲. روابط علی میان این مؤلفه‌ها چگونه است و کدام مؤلفه نقش محوری دارد؟
 ۳. اجرای چنین سامانه‌ای چه تأثیری بر خلق دانش راهبردی، شکوفایی کارکنان و هوش جمعی سازمانی دارد؟
- ضرورت انجام این پژوهش از دو منظر نظری و کاربردی قابل تبیین است. از منظر نظری، نتایج آن می‌تواند ادبیات مدیریت دانش را با معرفی مفهوم «سامانه همزیست» غنا بخشد و پیوندهای میان حوزه‌های پیش‌تر جداگانه را آشکار سازد (Rodrigues, 2025; Brown & Wilson, 2025). همچنین، با به‌کارگیری ترکیبی از روش‌های دلفی فازی، دیمتل فازی، نگاشت شناختی فازی و شبیه‌سازی عامل‌محور، گامی در جهت توسعه روش‌شناسی پژوهش‌های آتی در حوزه سیستم‌های پیچیده انسانی-فناورانه برداشته می‌شود (Khan et al., 2026; Martinez et al., 2025). از منظر کاربردی، مدل پیشنهادی برای مدیران منابع انسانی و مدیران دانش سازمان‌های دانش‌بنیان راهنمایی عملی فراهم می‌کند تا بتوانند ضمن حفظ سلامت روانی و انگیزش کارکنان، از ظرفیت‌های زایشی هوش مصنوعی برای تولید دانش جدید و ارتقای هوش جمعی بهره‌گیرند (Al-Hawamdeh & Al-Jabri, 2025; Lee & Kim, 2025).

ادبیات نظری و پیشینه پژوهش

مدیریت منابع انسانی مثبت‌گرا: از سرمایه روانشناختی تا شکوفایی

مفهوم مدیریت منابع انسانی مثبت‌گرا ریشه در جنبش روانشناسی مثبت‌گرا دارد که در اواخر دهه ۱۹۹۰ توسط سلیگمن مطرح شد. این جنبش در تقابل با روانشناسی سنتی که عمدتاً بر آسیب‌شناسی و اختلالات روانی متمرکز بود، به بررسی نقاط قوت، فضایل و ظرفیت‌های شکوفایی انسان پرداخت (Cameron, 2003). در حوزه سازمانی، لوتانز و یوسف (Luthans & Youssef, 2004) مفهوم «سرمایه روانشناختی»^۱ را شامل چهار مؤلفه امید، تاب‌آوری، خوش‌بینی و خودکارآمدی معرفی کردند که مبنای نظری مدیریت منابع انسانی مثبت‌گرا قرار گرفت.

در سال‌های اخیر، این رویکرد فراتر از سرمایه روانشناختی رفته و بر «شکوفایی کارکنان»^۲ به‌عنوان هدف غایی تأکید کرده است. شکوفایی به حالتی اطلاق می‌شود که در آن فرد هم به‌لحاظ ذهنی (احساس مثبت و رضایت از زندگی) و هم به‌لحاظ کارکردی (معنا، روابط مثبت و مشارکت) در سطح بالایی از به‌زیستی قرار دارد (Huppert & So, 2013; Gruman & Budworth, 2022). در این دیدگاه، مدیریت منابع انسانی نه صرفاً ابزاری برای افزایش بهره‌وری، بلکه بستری برای پرورش انسان‌های شکوفا در نظر گرفته می‌شود. ون زیل و راتمن (Van Zyl & Rothmann, 2022) چارچوب پنج مؤلفه‌ای برای مدیریت منابع انسانی مثبت‌گرا در محیط‌های کاری مدرن ارائه کردند که شامل طراحی شغل غنی‌شده، رهبری توانمندساز، فرهنگ حمایتگر، نظام‌های پاداش منصفانه و فرصت‌های توسعه فردی است.

هوش مصنوعی مولد: از خودکارسازی تا هم‌آفرینی دانش

هوش مصنوعی مولد نسل جدیدی از سیستم‌های هوش مصنوعی است که قادر به تولید محتوای جدید – شامل متن، تصویر، صدا، کد و طرح‌های پیچیده – بر اساس الگوهای آموخته‌شده از داده‌های آموزشی است (Huang & Rust, 2018). در حالی که نسل‌های پیشین هوش مصنوعی عمدتاً بر تشخیص الگو، طبقه‌بندی و پیش‌بینی متمرکز بودند، مدل‌های مولد با قابلیت «خلق» محتوای بدیع، مرزهای سنتی میان انسان و ماشین را در حوزه‌های خلاقیت و حل مسئله جابه‌جا کرده‌اند (Ali & Mahmood, 2025).

در حوزه مدیریت دانش، هوش مصنوعی مولد سه تحول اساسی ایجاد کرده است: نخست، تسریع فرآیند خلق دانش از طریق تولید خودکار پیش‌نویس‌ها، خلاصه‌سازی متون بلند و ترکیب اطلاعات از منابع پراکنده (Rodrigues, 2025). دوم، شخصی‌سازی یادگیری و توسعه حرفه‌ای با تطبیق محتوای آموزشی با نیازها و سبک یادگیری هر کارمند (Ellikkal & Rajamohan, 2024). سوم، کشف الگوهای پنهان دانش ضمنی از طریق تحلیل تعاملات و متون غیرساختاریافته (Sharma, 2026). با این حال، پژوهشگران نسبت به خطرات این فناوری

¹ Psychological Capital

² Employee Flourishing

از جمله همگن سازی دانش^۱، کاهش تنوع شناختی و ایجاد «حباب فیلتر»^۲ هشدار داده اند (Raji et al., 2022; Cramarencu et al., 2023).

هوش جمعی راهبردی: از خرد جمعی تا ابرهوش ترکیبی

هوش جمعی به توانایی گروه ها در حل مسائل پیچیده تر و خلاقانه تر از آنچه هر فرد به تنهایی قادر است، اشاره دارد (Malik et al., 2021). ریشه های این مفهوم به کارهای آندرسن (۱۹۹۹) در مورد «خرد جمعی» و سروو (Surowiecki, 2004) در مورد «خرد توده ها» بازمی گردد. در سطح سازمانی، هوش جمعی راهبردی به ظرفیت یک سازمان برای بهره گیری از دانش توزیع شده در میان اعضا، تنوع دیدگاه ها و هماهنگی هوشمندانه برای تصمیم گیری های کلان و نوآوری های بنیادین اطلاق می شود (Smith et al., 2025; Thompson et al., 2026).

پژوهش های اخیر بر این نکته تأکید دارند که ترکیب هوش جمعی انسانی با قابلیت های محاسباتی و زایشی هوش مصنوعی می تواند به «ابر هوش ترکیبی»^۳ منجر شود که ظرفیت پرداختن به مسائل بنیادین و راهبردی را دارد (Dittmar & Sposato, 2026; Garcia & Patel, 2025). در این دیدگاه، هوش مصنوعی نه به عنوان جایگزین، بلکه به عنوان تقویت کننده و تسهیل گر هوش جمعی عمل می کند. با این حال، چالش اصلی حفظ تعادل میان کارایی (که به هماهنگی و یکنواختی گرایش دارد) و تنوع شناختی (که شرط خلاقیت و نوآوری است) می باشد (Brown & Wilson, 2025).

تعاریف مفهومی و نظریه های پایه

تعریف مدیریت منابع انسانی مثبت گرا (متغیر زمینه ای)

در این پژوهش، مدیریت منابع انسانی مثبت گرا به عنوان چارچوبی از خط مشی ها، فرآیندها و اقدامات منابع انسانی تعریف می شود که به طور نظام مند به پرورش سرمایه روانشناختی (امید، تاب آوری، خوش بینی، خود کارآمدی)، تسهیل شکوفایی کارکنان و ایجاد محیط کاری سرشار از معنا، روابط مثبت و مشارکت فعال می پردازد (Gruman & Budworth, 2022; Van Zyl & Rothmann, 2022). این تعریف در تقابل با رویکردهای سنتی قرار دارد که کارکنان را صرفاً «منابعی» برای دستیابی به اهداف سازمانی می دیدند.

تعریف هوش مصنوعی مولد (متغیر فناورانه)

هوش مصنوعی مولد در این پژوهش به عنوان دسته ای از الگوریتم های یادگیری ماشین (به ویژه مدل های زبانی بزرگ، شبکه های متخاصم مولد و مدل های انتشاری) تعریف می شود که قادر به تولید محتوای بدیع و معنادار (متن، تصویر، کد، طرح) بر اساس الگوهای آماری آموخته شده از داده های آموزشی هستند و می توانند در فرآیندهای خلق، اشتراک و شخصی سازی دانش سازمانی به کار گرفته شوند (Ali & Mahmood, 2025; Rodrigues, 2025).

تعریف هوش جمعی راهبردی (متغیر وابسته)

هوش جمعی راهبردی در این پژوهش به عنوان ظرفیت یک سازمان برای ترکیب دانش توزیع شده، تنوع شناختی اعضا و هماهنگی هوشمندانه (با یا بدون واسطه فناوری) به منظور تصمیم گیری های کلان، حل مسائل پیچیده و خلق نوآوری های بنیادینی که فراتر از توانایی تک اعضا است، تعریف می شود (Smith et al., 2025; Thompson et al., 2026).

نظریه های پایه

این پژوهش بر دوش دو نظریه اصلی استوار است:

نظریه خودتعیین گری^۴ که توسط دسی و رایان (Deci & Ryan, 1985) مطرح شد، بر سه نیاز روانشناختی بنیادین - خودمختاری، شایستگی و تعلق پذیری - تأکید دارد که ارضای آنها منجر به انگیزش درونی، بهزیستی و شکوفایی می شود. در چارچوب پژوهش حاضر، هوش مصنوعی مولد می تواند از طریق شخصی سازی یادگیری (تقویت شایستگی)، شفافیت الگوریتمی (تقویت خودمختاری) و تسهیل تعاملات (تقویت تعلق پذیری) به ارضای این نیازها کمک کند (Lee & Kim, 2025).

¹ Knowledge Homogenization

² Filter Bubble

³ Hybrid Superintelligence

⁴ Self-Determination Theory

نظریه شناخت توزیع شده^۱ که توسط هاجینز (Hutchins, 1995) معرفی شد، استدلال می کند که فرآیندهای شناختی نه فقط در درون ذهن فرد، بلکه در میان افراد، ابزارها و محیط توزیع می شوند. در این دیدگاه، هوش مصنوعی مولد به عنوان یک «ابزار شناختی» یا حتی «هم اندیش» عمل می کند که می تواند ظرفیت شناختی جمعی سازمان را افزایش دهد (Wilson et al., 2025; Chen & Wang, 2026).

پیشینه پژوهش

در جدول ۱، پیشینه پژوهش های داخلی و خارجی مرتبط با موضوع ارائه شده است.

Table 1. Research background.

جدول ۱. پیشینه پژوهش

ردیف	نویسندگان (سال)	کشور	هدف پژوهش	روش شناسی	متغیرها/ابعاد	مهم ترین یافته ها	ارتباط با پژوهش حاضر
۱	Haji Zadeh & Sardari (2018)	ایران	تأثیر مدیریت دانش بر عملکرد نوآورانه	کمی (همبستگی)	مدیریت دانش، یادگیری سازمانی، عملکرد نوآورانه	مدیریت دانش از طریق یادگیری سازمانی بر نوآوری اثر می گذارد	شکاف: نقش هوش مصنوعی و رویکرد مثبت گرا در این رابطه بررسی نشده
۲	Gruman & Budworth (2022)	کانادا	معماری منابع انسانی برای شکوفایی کارکنان	نظری	سرمایه روانشناختی، طراحی شغل، رهبری	لزوم بازطراحی فرآیندهای منابع انسانی برای حمایت از شکوفایی	شکاف: نقش فناوری های نوین (به ویژه هوش مصنوعی) در این معماری دیده نشده
۳	Prikshtat et al. (2023)	استرالیا	مرور نظام مند هوش مصنوعی در مدیریت منابع انسانی	مروری نظام مند	سطوح فردی، تیمی و سازمانی	ارائه چارچوب چندسطحی برای پژوهش های آتی در هوش مصنوعی در مدیریت منابع انسانی	شکاف: به جنبه های مثبت گرا و شکوفایی کارکنان توجه نشده
۴	Mousavi et al. (2025)	ایران	بررسی نقش هوش مصنوعی در ارتقای مدیریت منابع انسانی مثبت گرا	آمیخته (دیمتال فازی + نگاشت شناختی)	۲۲ شاخص شامل استراتژی کسب و کار مبتنی بر هوش مصنوعی، مدل سازی استراتژیک منابع انسانی	مدل سازی استراتژیک منابع انسانی بیشترین نقش مرکزی را دارد	مقاله پایه مبنای پژوهش؛ اما به طراحی سامانه زایشی و هوش جمعی نپرداخته است
۵	Smith et al. (2025)	آمریکا	مدل پارامتریک هوش جمعی	کمی (شبیه سازی)	تنوع شناختی، مکانیسم های تجمع	ارائه رابطه ریاضی میان تنوع و دقت تصمیم جمعی	ارتباط مستقیم: تنوع شناختی به عنوان پیشران کلیدی شناسایی شده
۶	Al-Hawamdeh & Al-Jabri (2025)	عمان	ادغام اخلاقی هوش مصنوعی در مدیریت منابع انسانی برای ارتقای بهزیستی	مروری نظام مند	اخلاق هوش مصنوعی، بهزیستی کارکنان، مسئولیت اجتماعی	هوش مصنوعی باید در چارچوب مسئولیت اجتماعی و با نگاه بهزیستی طراحی شود	ارتباط مستقیم: حکمرانی اخلاقی به عنوان یکی از مؤلفه های ضروری شناسایی شده

¹ Distributed Cognition Theory

ردیف	نویسندگان (سال)	کشور	هدف پژوهش	روش‌شناسی	متغیرها/ابعاد	مهم‌ترین یافته‌ها	ارتباط با پژوهش حاضر
۷	Ali & Mahmood (2025)	انگلستان	نقش هوش مصنوعی مولد در تحول مدیریت دانش	مروری نظام‌مند	خلق دانش، اشتراک دانش، ذخیره‌سازی	شناسایی چهار سازوکار: تولید خودکار، اشتراک هوشمند، استخراج ضمنی، شخصی‌سازی	شکاف: به جنبه‌های انسانی و شکوفایی کارکنان نپرداخته
۸	Rodrigues (2025)	پرتغال	چالش‌های معرفت‌شناختی هوش مصنوعی مولد در مدیریت دانش	کیفی-تحلیلی	مرز انسان-ماشین، اصالت دانش	هوش مصنوعی مولد مرزهای سنتی خلق دانش را فرو می‌ریزد و نیازمند مدل‌های جدید است	شکاف: مدل پیشنهادی فاقد مؤلفه‌های مثبت‌گرا و اخلاقی است
۹	Ghafarian Sekhavar et al. (2025)	ایران	تحلیل شبکه یادگیری سازمانی	کمی (تحلیل شبکه)	یادگیری سازمانی، عوامل مؤثر	شناسایی عوامل کلیدی شکل‌دهنده شبکه یادگیری	شکاف: به نقش هوش مصنوعی در این شبکه توجه نشده
۱۰	Martinez et al. (2025)	اسپانیا	شبیه‌سازی عامل‌محور همکاری انسان-هوش مصنوعی	کمی (مدل‌سازی عامل-محور ^۱)	اعتماد، خودکارآمدی، بار شناختی	تعاملات مکرر با هوش مصنوعی شفاف، اعتماد و خودکارآمدی را افزایش می‌دهد	روش‌شناختی مدل‌سازی عامل-محور: روش مناسبی برای مدلسازی پویایی‌های تعامل انسان-هوش مصنوعی است
۱۱	Lee & Kim (2025)	کره جنوبی	مدل منابع-مطالبات در محیط‌های کاری هوش مصنوعی محور	کمی (پیمایش طولی)	منابع شغلی، مطالبات شغلی، فرسودگی	هوش مصنوعی می‌تواند هم به‌عنوان منبع (کاهش بار) و هم مطالبه (نظارت) عمل کند	ضرورت طراحی متعادل: هوش مصنوعی نباید صرفاً نظارتی باشد
۱۲	Sharma (2026)	هند	توسعه مدل مدل حلزونی مدیریت دانش با رویکرد هوش مصنوعی	کتاب‌سنجی	اجتماعی‌سازی، برونی‌سازی، ترکیب، درونی‌سازی	معرفی «بعد ماشین» در کنار ابعاد انسانی	شکاف: بعد ماشین صرفاً فنی دیده شده، نه همزیستانه
۱۳	Dittmar & Sposato (2026)	آلمان	مدل اعتماد-مکمل برای هوش جمعی در عصر هوش مصنوعی	نظری-مدلسازی	اعتماد، کالیبره‌شده، نگاشت یادگیری موقعیتی	برای همکاری مؤثر انسان و هوش مصنوعی، باید اعتماد به‌صورت پویا تنظیم شود	شکاف: اعتماد در بستر تنوع شناختی و شفافیت الگوریتمی تبیین نشده
۱۴	Chen & Wang (2026)	چین	نقش سرمایه روانشناختی در	کمی (پیمایش)	سرمایه روانشناختی، فشار	سرمایه روانشناختی اثرات منفی	تأییدکننده: لزوم توجه همزمان به فناوری و سرمایه روانشناختی

¹ Agent-Based Modeling (ABM)

ردیف	نویسندگان (سال)	کشور	هدف پژوهش	روش‌شناسی	متغیرها/ابعاد	مهم‌ترین یافته‌ها	ارتباط با پژوهش حاضر
			تعدیل فشار هوش مصنوعی		فناوری، رضایت شغلی	هوش مصنوعی بر رضایت شغلی را تعدیل می‌کند	
۱۵	Thompson et al. (2026)	استرالیا	بازنمایی توزیع شده در تصمیمات راهبردی	مدلسازی مبتنی بر عامل	بازنمایی جمعی، ساختار تعاملات	ساختار تعاملات به اندازه محتوای دانش اهمیت دارد	شکاف: نقش هوش مصنوعی در تسهیل بازنمایی توزیع شده بررسی نشده

جمع‌بندی تحلیلی و شکاف پژوهش

بررسی جدول پیشنهادی نشان می‌دهد که پژوهش‌های موجود را می‌توان در چهار روند غالب دسته‌بندی کرد:

روند اول: تمرکز بر بهینه‌سازی فنی فرآیندها. این مطالعات عمدتاً از منظر مهندسی و علوم کامپیوتر به هوش مصنوعی نگاه کرده‌اند و جنبه‌های انسانی، روانشناختی و اخلاقی را مغفول گذاشته‌اند.

روند دوم: تمرکز بر شکوفایی و به‌زیستی بدون توجه به فناوری. این مطالعات اگرچه مبانی نظری قدرتمندی در حوزه روانشناسی مثبت دارند، اما در عصر تحول دیجیتال، نقش فناوری‌های نوین را در فرآیندهای منابع انسانی نادیده گرفته‌اند.

روند سوم: پژوهش‌های ترکیبی محدود. این مطالعات گام‌هایی در جهت تلفیق هوش مصنوعی و مدیریت منابع انسانی برداشته‌اند، اما عمدتاً در سطح تحلیل روابط علی (مانند دیمتل فازی) باقی مانده و به طراحی سامانه‌ای زایشی و پویا برای خلق دانش جدید نپرداخته‌اند. همچنین مفهوم هوش جمعی راهبردی به عنوان متغیر وابسته نهایی در هیچ یک دیده نمی‌شود.

روند چهارم: پژوهش‌های روش‌شناختی نوین. این مطالعات نشان داده‌اند که روش‌هایی مانند شبیه‌سازی عامل‌محور و مدل‌سازی مبتنی بر عامل برای مطالعه تعاملات پویای انسان و هوش مصنوعی بسیار مناسب هستند، اما این روش‌ها به ندرت با چارچوب‌های نظری مدیریت منابع انسانی مثبت‌گرا تلفیق شده‌اند.

شکاف نظری پژوهش: بر اساس تحلیل انتقادی فوق، شکاف اصلی پژوهش حاضر را می‌توان به‌صورت زیر صورت‌بندی کرد:

«به نظر می‌رسد تا کنون پژوهش‌های پیشین به طراحی الگویی نپرداخته‌اند که به‌طور هم‌افزا سه مؤلفه (۱) مدیریت منابع انسانی مثبت‌گرا، (۲) هوش مصنوعی مولد و (۳) هوش جمعی راهبردی را در قالب یک «سامانه دانشی همزیست» تلفیق کند. همچنین، نقش «حفظ تنوع شناختی» به‌عنوان موتور محرک این هم‌افزایی و روابط علی آن با مؤلفه‌هایی چون شفافیت الگوریتمی، بازخورد تعامل انسان-ماشین، یادگیری شخصی‌سازی شده و حکمرانی اخلاقی هوش مصنوعی در مدل‌های موجود به طور کامل تبیین نشده است.»

ابعاد مغفول: بر اساس جمع‌بندی پیشنهادی، ابعاد زیر در ادبیات مغفول مانده یا کم‌توجه بوده‌اند:

- تأثیر متقابل (نه یک‌سویه) هوش مصنوعی و سرمایه روانشناختی
- نقش حفظ تنوع شناختی به‌عنوان پیشران هم‌افزایی
- طراحی سامانه‌ای که هوش مصنوعی در آن «هم‌اندیش» باشد نه «جانشین»
- پیوند میان یادگیری شخصی‌سازی شده با هوش جمعی راهبردی

نوآوری‌های پژوهش

بر اساس شکاف شناسایی شده، این پژوهش دارای سه وجه نوآوری است:

نوآوری نظری: ارائه چارچوب سامانه دانشی همزیست به عنوان یکی از تلاش‌های اولیه برای پیوند میان مدیریت منابع انسانی مثبت‌گرا، هوش مصنوعی مولد و هوش جمعی راهبردی را به‌طور نظام‌مند عملیاتی می‌کند. همچنین معرفی مفهوم «حفظ تنوع شناختی» به‌عنوان متغیر علی محوری.

نوآوری روش‌شناختی: تلفیق چهار روش مکمل دلفی فازی (برای استخراج مؤلفه‌ها)، دیمتل فازی (برای شناسایی روابط علی)، نگاشت شناختی فازی (برای نمایش بصری روابط) و شبیه‌سازی عامل‌محور (برای مدلسازی پویایی‌های زمانی) در یک چارچوب یکپارچه. این ترکیب در پژوهش‌های مدیریت دانش و منابع انسانی کمتر مورد استفاده قرار گرفته است.

نوآوری کاربردی: ارائه الگویی عملیاتی شامل پنج مؤلفه کلیدی (شفافیت الگوریتمی، حفظ تنوع شناختی، بازخورد تعامل انسان-ماشین، یادگیری شخصی سازی شده مولد و حکمرانی اخلاقی هوش مصنوعی) که مدیران منابع انسانی و دانش می توانند برای طراحی و پیاده سازی سامانه های دانشی نسل آینده در سازمان های دانش بنیان از آن استفاده کنند.

مدل مفهومی پژوهش

بر اساس مبانی نظری (نظریه خودتعیین گری و نظریه شناخت توزیع شده) و جمع بندی پیشینه پژوهش (جدول ۱)، یک چارچوب نظری اولیه با هدف جهت دهی به مراحل بعدی پژوهش تدوین شده است. در این چارچوب، مجموعه ای از مؤلفه های بالقوه پیشنهاد می شود که به نظر می رسد در شکل گیری «سامانه دانشی همزیست مولد» نقش داشته باشند. این مؤلفه ها عبارتند از: شفافیت الگوریتمی، حفظ تنوع شناختی، بازخورد تعامل انسان-ماشین، یادگیری شخصی سازی شده مولد و حکمرانی اخلاقی هوش مصنوعی. همچنین، پیامدهای محتمل اجرای چنین سامانه ای شامل خلق دانش راهبردی، شکوفایی کارکنان و هوش جمعی راهبردی در نظر گرفته شده است.

تأکید می شود که این چارچوب نظری اولیه است و هیچ یک از مؤلفه ها یا روابط آن جنبه قطعی و نهایی ندارد. مؤلفه های ذکر شده در واقع فرضیاتی اولیه هستند که باید در مرحله دلفی فازی توسط خبرگان مورد تأیید، اصلاح یا رد قرار گیرند. هدف از ارائه این چارچوب، تعیین جهت اولیه پژوهش، طراحی پرسشنامه های دلفی فازی و دیمتل فازی، و فراهم آوردن مبنایی برای مقایسه با نتایج تجربی بعدی است. بنابراین، این چارچوب نباید به منزله مدل نهایی پژوهش تلقی گردد. شکل ۱، نمایی شماتیک از این چارچوب نظری اولیه را نشان می دهد.

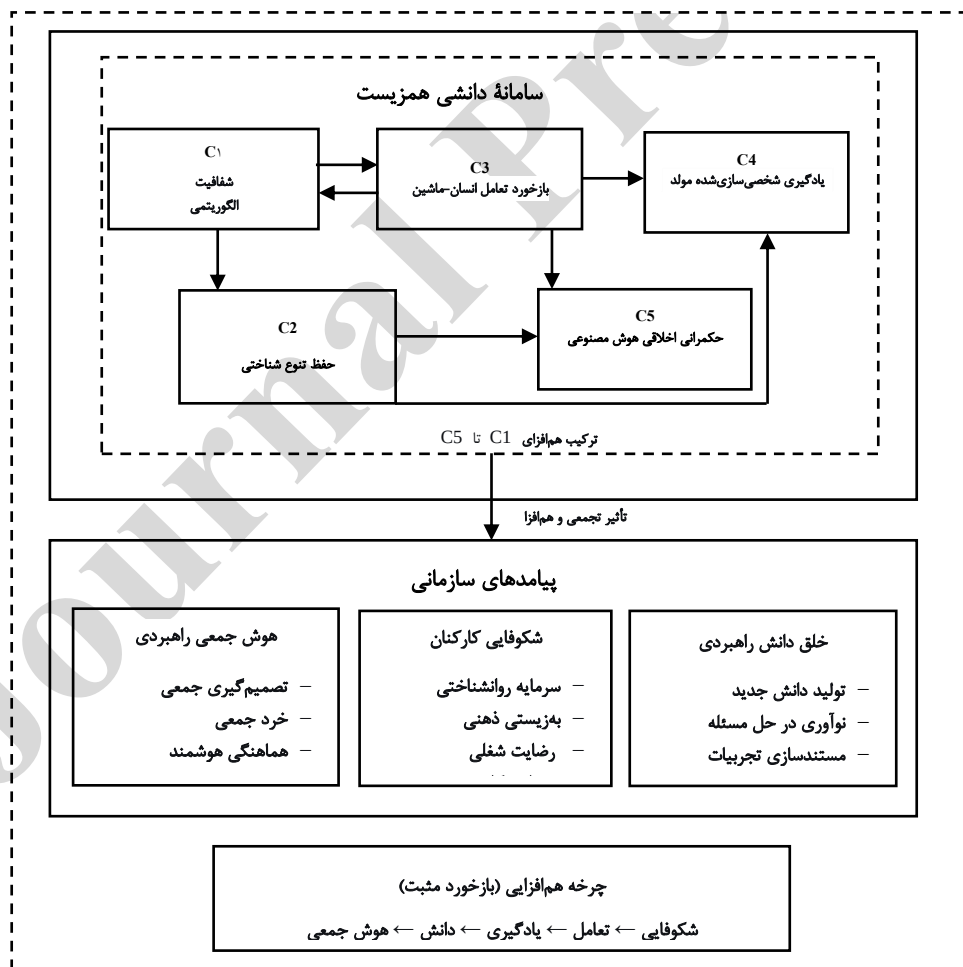


Fig. 1. Initial theoretical framework of the research.

شکل ۱. چارچوب نظری اولیه پژوهش

روش‌شناسی پژوهش

فلسفه، رویکرد و استراتژی پژوهش

پژوهش حاضر از نظر هدف، کاربردی است زیرا به دنبال ارائه الگویی عملی برای طراحی سامانه‌های دانشی در سازمان‌ها می‌باشد. از نظر رویکرد، پژوهشی آمیخته از نوع اکتشافی متوالی^۱ است؛ بدین معنا که ابتدا با روش‌های کیفی، مؤلفه‌های اصلی الگو شناسایی و سپس با روش‌های کمی، روابط علی میان آنها تحلیل و پویایی‌های سامانه شبیه‌سازی می‌شود. فلسفه حاکم بر پژوهش، عمل‌گرایی است که امکان بهره‌گیری همزمان از روش‌های کیفی و کمی را با توجه به مسئله پژوهش فراهم می‌سازد. استراتژی پژوهش در مرحله کیفی، «دلفی فازی»^۲ و در مرحله کمی، ترکیبی از «دیمتل فازی»^۳، «نگاشت شناختی فازی»^۴ و «شبیه‌سازی عامل‌محور»^۵ می‌باشد. زمان پژوهش مقطعی بوده و داده‌ها در بازه زمانی مهر ۱۴۰۴ تا اردیبهشت ۱۴۰۵ جمع‌آوری شده‌اند.

مرحله اول: تحلیل کیفی به روش دلفی فازی

مرور نظام‌مند ادبیات و استخراج مؤلفه‌های اولیه

مرور نظام‌مند ادبیات بر اساس دستورالعمل PRISMA و با جستجو در پایگاه‌های معتبر علمی شامل Web of Science، Scopus، Emerald، Science Direct (برای منابع انگلیسی) و مگیران و نورمگز (برای منابع فارسی) انجام شد. رشته جستجو با استفاده از کلیدواژه‌های ترکیبی «هوش مصنوعی مولد»، «مدیریت منابع انسانی مثبت‌گرا»، «هوش جمعی»، «مدیریت دانش» و «همزیستی انسان-ماشین» در بازه زمانی ۲۰۲۰ تا ۲۰۲۶ تنظیم گردید. معیارهای ورود شامل مقالات علمی-پژوهشی دارای داوری، مطالعات اصیل کمی یا کیفی، مقالات مروری نظام‌مند و مطالعات موردی با کیفیت بالا بود و معیارهای خروج شامل مقالات بدون داده‌های کافی، مقالات کنفرانسی بدون داوری و فصل‌های کتاب می‌شد. فرآیند غربالگری در چهار مرحله انجام شد که در شکل ۲ به تصویر کشیده شده است: ابتدا ۳۷۸ مقاله شناسایی گردید؛ پس از حذف ۷۹ مقاله تکراری، ۲۹۹ مقاله باقی ماند. در مرحله غربالگری عنوان و چکیده، ۱۴۷ مقاله حذف شد و ۱۵۲ مقاله وارد مرحله بعد گردید. در مرحله غربالگری متن کامل، ۵۸ مقاله به دلایل عدم ارتباط مستقیم با تلفیق سه حوزه (۲۶ مقاله)، فقدان روش‌شناسی معتبر (۱۸ مقاله) و عدم دسترسی به متن کامل (۱۴ مقاله) حذف شد. در نهایت ۹۴ مقاله به عنوان منابع نهایی انتخاب شدند و با تحلیل محتوای آنها، فهرست اولیه‌ای شامل ۳۴ مؤلفه بالقوه برای سامانه دانشی همزیست استخراج گردید.

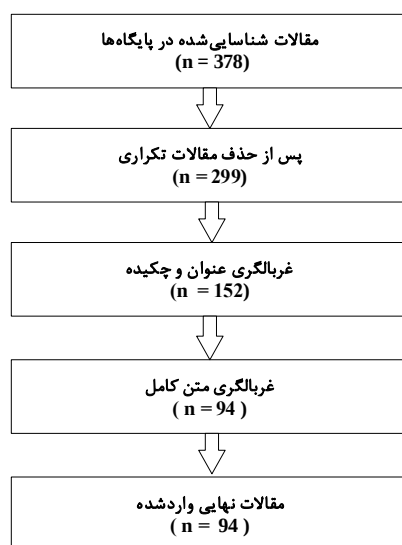


Fig. 2. PRISMA flow diagram

شکل ۲. نمودار جریان PRISMA

¹ Sequential Exploratory

² Fuzzy Delphi

³ Fuzzy DEMATEL

⁴ Fuzzy Cognitive Mapping

⁵ Agent-Based Modeling

تشکیل پانل خبرگان

پانلی متشکل از ۱۵ نفر خبره با استفاده از روش نمونه‌گیری هدفمند و گلوله‌برفی تشکیل شد. معیارهای انتخاب خبرگان عبارت بودند از: (۱) داشتن درجهٔ دکتری تخصصی در یکی از حوزه‌های مدیریت منابع انسانی، مدیریت دانش، هوش مصنوعی یا روانشناسی مثبت‌گرا؛ (۲) داشتن حداقل ۵ سال سابقه پژوهشی یا اجرایی مرتبط؛ (۳) انتشار حداقل ۳ مقاله علمی پژوهشی در مجلات معتبر در حوزه‌های مرتبط؛ (۴) آشنایی با مفاهیم سیستم‌های فازی و روش‌های تصمیم‌گیری چندمعیاره. ترکیب خبرگان شامل ۶ نفر متخصص مدیریت منابع انسانی، ۴ نفر متخصص مدیریت دانش، ۳ نفر متخصص هوش مصنوعی و ۲ نفر متخصص روانشناسی مثبت‌گرا بود. مشخصات جمعیت‌شناختی این ۱۵ خبره، شامل مرتبه علمی، رشته تخصصی، سابقه کار، تعداد مقالات علمی-پژوهشی و جنسیت، در جدول ۲ ارائه شده است.

Table 2. Demographic Characteristics of Experts
جدول ۲. مشخصات جمعیت‌شناختی خبرگان مشارکت‌کننده در پژوهش

ردیف	مرتبه علمی	رشته تخصصی	سابقه کار (سال)	تعداد مقالات علمی-پژوهشی	جنسیت
۱	استاد	مدیریت منابع انسانی	۲۲	۴۷	مرد
۲	دانشیار	مدیریت منابع انسانی	۱۸	۳۲	زن
۳	استاد	مدیریت دانش	۲۵	۵۸	مرد
۴	دانشیار	مدیریت دانش	۱۵	۲۸	مرد
۵	استادیار	هوش مصنوعی	۱۰	۱۹	زن
۶	دانشیار	هوش مصنوعی	۱۴	۲۶	مرد
۷	استاد	روانشناسی مثبت‌گرا	۲۰	۴۱	مرد
۸	دانشیار	روانشناسی صنعتی و سازمانی	۱۶	۲۹	زن
۹	استادیار	مدیریت فناوری اطلاعات	۹	۱۸	مرد
۱۰	دانشیار	مدیریت منابع انسانی	۱۷	۳۴	مرد
۱۱	استاد	مدیریت دانش	۲۴	۵۲	زن
۱۲	استادیار	هوش مصنوعی	۸	۱۵	مرد
۱۳	دانشیار	مدیریت منابع انسانی	۱۳	۲۴	زن
۱۴	استاد	روانشناسی مثبت‌گرا	۲۱	۴۳	مرد
۱۵	دانشیار	مدیریت فناوری اطلاعات	۱۲	۲۲	مرد

اجرای دلفی فازی

فرآیند دلفی فازی در سه دور انجام شد. در هر دور، خبرگان با استفاده از طیف زبانی فازی پنج‌ارزشی (جدول ۳) به هر مؤلفه امتیاز می‌دادند. برای تبدیل متغیرهای زبانی به اعداد فازی مثلثی، از طیف استاندارد (۰، ۰/۲۵، ۰/۵، ۰/۷۵، ۱) استفاده شد.

Table 3. Fuzzy linguistic scale used in delphi
جدول ۳. طیف زبانی فازی مورد استفاده در دلفی

متغیر زبانی	عدد فازی مثلثی
خیلی کم (VL)	(۰, ۰, ۰/۲۵)
کم (L)	(۰, ۰/۲۵, ۰/۵)
متوسط (M)	(۰/۲۵, ۰/۵, ۰/۷۵)
زیاد (H)	(۰/۵, ۰/۷۵, ۱)
خیلی زیاد (VH)	(۰/۷۵, ۱, ۱)

پس از جمع‌آوری نظرات، عملیات غیرفازی‌سازی با استفاده از روش مرکز ثقل^۱ انجام شد:

$$\text{Defuzzified Value} = \frac{l + m + u}{3}$$

¹ Center of Gravity

که در آن l ، m و u به ترتیب کران پایین، میانگین و کران بالای عدد فازی مثلثی هستند. معیار پذیرش مؤلفه‌ها، داشتن میانگین غیرفازی ≤ 0.75 و ضریب توافق کندال (Kendall's W) معنادار در سطح $p < 0.05$ بود. مؤلفه‌هایی که در دور سوم به اجماع نرسیدند، حذف شدند.

روایی و پایایی مرحله کیفی

روایی محتوایی: برای تأیید روایی محتوایی مؤلفه‌های استخراج‌شده از دلفی فازی، پرسشنامه در اختیار ۵ خبره مستقل (که در پانل دلفی حضور نداشتند) قرار گرفت. نسبت روایی محتوایی (CVR) بر اساس فرمول لاوشه (Lawshe, 1975) محاسبه شد. مقدار بحرانی CVR برای ۵ خبره در سطح اطمینان ۹۵٪ برابر ۰/۹۹ در نظر گرفته شد. مؤلفه‌هایی که CVR آن‌ها کمتر از این مقدار بود، از تحلیل حذف شدند. همچنین برای افزایش دقت، شاخص روایی محتوایی (CVI) نیز محاسبه گردید. فهرست کامل ۳۴ مؤلفه اولیه به همراه مقادیر CVR و CVI و وضعیت پذیرش یا رد هر مؤلفه در پیوست ۲ ارائه شده است. میانگین CVI برای مؤلفه‌های تأیید شده نهایی (C1 تا C5) برابر ۹۴/۰ به دست آمد.

پایایی: پایایی نظرات خبرگان در مراحل دلفی فازی با محاسبه ضریب توافق کندال در پایان دور سوم سنجیده شد. مقدار W برابر ۰/۷۳ به دست آمد که در سطح معناداری $p < 0.01$ قرار دارد، نشان‌دهنده توافق قابل قبول میان خبرگان است.

مرحله دوم: تحلیل کمی روابط علی با دیمتل فازی

طراحی پرسشنامه و جمع‌آوری داده‌ها

پس از تأیید مؤلفه‌های نهایی، پرسشنامه مقایسه زوجی بر اساس طیف زبانی فازی دیمتل (جدول ۴) طراحی شد. در این پرسشنامه، از خبرگان خواسته شد تا میزان تأثیر مستقیم هر مؤلفه بر مؤلفه‌های دیگر را در مقیاس ۵-ارزشی (بی‌اثر تا اثر خیلی زیاد) ارزیابی کنند. پانل خبرگان همان ۱۵ نفر مرحله پیشین بودند. پرسشنامه به صورت حضوری و الکترونیکی توزیع و تکمیل شد و نرخ بازگشت ۱۰۰ درصد بود.

Table 4. Fuzzy linguistic scale for the fuzzy DEMATEL technique

جدول ۴. طیف زبانی فازی تکنیک دیمتل

متغیر زبانی	معادل قطعی	عدد فازی مثلثی
بی‌اثر (No)	۰	(۰, ۰, ۰/۲۵)
اثر کم (VL)	۱	(۰, ۰/۲۵, ۰/۵)
اثر متوسط (L)	۲	(۰/۲۵, ۰/۵, ۰/۷۵)
اثر زیاد (H)	۳	(۰/۵, ۰/۷۵, ۱)
اثر خیلی زیاد (VH)	۴	(۰/۷۵, ۱, ۱)

گام‌های اجرای دیمتل فازی

اجرای تکنیک دیمتل فازی در هفت گام زیر انجام شد:

گام ۱: تشکیل ماتریس ارتباط مستقیم فازی ($\tilde{Z}$). برای هر جفت مؤلفه، نظرات خبرگان به اعداد فازی مثلثی تبدیل و میانگین گرفته شد:

$$\tilde{z}_{ij} = \frac{\tilde{x}_{ij}^1 \oplus \tilde{x}_{ij}^2 \oplus \dots \oplus \tilde{x}_{ij}^p}{p}$$

که در آن p تعداد خبرگان (۱۵ نفر)، $\tilde{x}_{ij}^p$ و عدد فازی مثلثی نظر خبره p است.

گام ۲: نرمال‌سازی ماتریس ارتباط مستقیم فازی. برای نرمال‌سازی، از رابطه زیر استفاده شد:

$$\tilde{N} = \frac{\tilde{z}_{ij}}{k}, \quad k = \max_{1 \leq i \leq n} \left(\sum_{j=1}^n u_{ij} \right)$$

که u_{ij} کران بالای عدد فازی مثلثی $\tilde{z}_{ij}$ است.

گام ۳: محاسبه ماتریس ارتباط کامل فازی ($\tilde{T}$). با استفاده از رابطه زیر، ماتریس ارتباط کامل برای هر سه مؤلفه اعداد فازی (l, m, u) محاسبه شد:

$$\tilde{T}_l = \tilde{N}_l(I - \tilde{N}_l)^{-1}$$

$$\tilde{T}_m = \tilde{N}_m(I - \tilde{N}_m)^{-1}$$

$$\tilde{T}_u = \tilde{N}_u(I - \tilde{N}_u)^{-1}$$

گام ۴: غیرفازی سازی ماتریس ارتباط کامل. فرآیند دیمتلی فازی، مراحل غیرفازی سازی و نرمال سازی به صورت زیر انجام شد:

مرحله ۱ - تشکیل ماتریس ارتباط مستقیم فازی ($\tilde{Z}$):

نظرات زبانی ۱۵ خبره بر اساس طیف فازی جدول ۳ به اعداد فازی مثلثی تبدیل شدند. برای هر جفت مؤلفه (i, j)، میانگین اعداد فازی مثلثی تمام خبرگان محاسبه شد:

$$\tilde{z}_{ij} = \frac{\tilde{x}_{ij}^1 \oplus \tilde{x}_{ij}^2 \oplus \dots \oplus \tilde{x}_{ij}^{15}}{15}$$

مرحله ۲ - نرمال سازی ماتریس ارتباط مستقیم فازی ($\tilde{N}$):

برای نرمال سازی، از رابطه زیر استفاده شد:

$$\tilde{N} = \frac{\tilde{z}_{ij}}{k}, \quad k = \max_{1 \leq i \leq n} \left(\sum_{j=1}^n u_{ij} \right)$$

که در آن u_{ij} کران بالای عدد فازی مثلثی $\tilde{Z}_{ij}$ است.

مرحله ۳ - محاسبه ماتریس ارتباط کامل فازی ($\tilde{T}$):

با استفاده از روابط زیر، ماتریس ارتباط کامل برای هر سه مؤلفه اعداد فازی (l, m, u) محاسبه شد:

$$\tilde{T}_l = \tilde{N}_l(I - \tilde{N}_l)^{-1}$$

$$\tilde{T}_m = \tilde{N}_m(I - \tilde{N}_m)^{-1}$$

$$\tilde{T}_u = \tilde{N}_u(I - \tilde{N}_u)^{-1}$$

مرحله ۴ - غیرفازی سازی ماتریس ارتباط کامل:

با استفاده از روش میانگین، اعداد فازی مثلثی به مقادیر قطعی تبدیل شدند:

$$t_{ij} = \frac{t_{ij}^l + t_{ij}^m + t_{ij}^u}{3}$$

مرحله ۵ - اعمال آستانه و حذف روابط ضعیف:

آستانه $\theta = 0.5780$ بر اساس میانگین وزنی نظرات خبرگان (جدول ۵) تعیین شد. مقادیر زیر آستانه در ماتریس نهایی صفر قرار گرفتند.

گام ۵: تعیین آستانه حذف روابط ضعیف.

برای جلوگیری از ورود روابط بسیار ضعیف به مدل، آستانه θ بر اساس میانگین حسابی تمام درایه های ماتریس ارتباط مستقیم غیرفازی شده محاسبه شد. فرمول محاسبه آستانه به صورت زیر است:

که در آن t_{ij} مقادیر غیرفازی شده ماتریس ارتباط مستقیم (Z) هستند. پس از تبدیل نظرات زبانی ۱۵ خبره به اعداد فازی مثلثی و غیرفازی سازی با روش مرکز ثقل، ماتریس ارتباط مستقیم غیرفازی شده به دست آمد. میانگین حسابی تمام درایه های این ماتریس (شامل درایه های قطری که مقدار صفر دارند) برابر 0.5780 محاسبه شد.

جدول میانگین وزنی نظرات خبرگان (جدول ۵) نشان می دهد که چگونه نظرات ۱۵ خبره در ماتریس ارتباط مستقیم، به این مقدار میانگین منجر شده است. مقادیر زیر آستانه در ماتریس نهایی صفر قرار گرفتند و تنها روابط با وزن $0.5780 \leq$ در تحلیل نهایی لحاظ شدند.

Table 5. Weighted mean of experts' opinions in the fuzzy DEMATEL direct-relation matrix

جدول ۵. میانگین وزنی نظرات خبرگان در ماتریس ارتباط مستقیم دیمتل فازی

ردیف خبره	رشته تخصصی	میانگین نظرات (غیرفازی شده) در ماتریس ارتباط مستقیم
۱	مدیریت منابع انسانی	۰/۵۶۲
۲	مدیریت منابع انسانی	۰/۵۷۱
۳	مدیریت دانش	۰/۵۸۴
۴	مدیریت دانش	۰/۵۷۹
۵	هوش مصنوعی	۰/۵۹۰
۶	هوش مصنوعی	۰/۵۸۱
۷	روانشناسی مثبت گرا	۰/۵۶۸
۸	روانشناسی صنعتی و سازمانی	۰/۵۷۵
۹	مدیریت فناوری اطلاعات	۰/۵۸۸
۱۰	مدیریت منابع انسانی	۰/۵۷۲
۱۱	مدیریت دانش	۰/۵۸۶
۱۲	هوش مصنوعی	۰/۵۹۵
۱۳	مدیریت منابع انسانی	۰/۵۷۰
۱۴	روانشناسی مثبت گرا	۰/۵۷۴
۱۵	مدیریت فناوری اطلاعات	۰/۵۸۲
	میانگین کل	۰/۵۷۸۰

این جدول نشان می‌دهد که میانگین وزنی نظرات ۱۵ خبره در ماتریس ارتباط مستقیم غیرفازی شده برابر ۰/۵۷۸۰ است که به عنوان آستانه (θ) برای حذف روابط ضعیف در نظر گرفته شده است. تفاوت اندک میانگین نظرات خبرگان نشان‌دهنده توافق نسبی آن‌ها در ارزیابی روابط علی میان مؤلفه‌هاست.

گام ۶: محاسبه شاخص‌های تأثیرگذاری و تأثیرپذیری

مجموع سطرها (D_i) نشان‌دهنده میزان تأثیرگذاری مؤلفه i بر سایر مؤلفه‌ها است:

$$D_i = \sum_{j=1}^n t_{ij}$$

مجموع ستونها (R_i) نشان‌دهنده میزان تأثیرپذیری مؤلفه i از سایر مؤلفه‌ها است:

$$R_i = \sum_{j=1}^n t_{ji}$$

گام ۷: تعیین نقش علی یا معلولی

- $Di+Ri$: درجه مرکزیت - نشان‌دهنده اهمیت کلی مؤلفه در شبکه
- $Di-Ri$: درجه علیت - اگر مثبت باشد، مؤلفه نقش علی (تأثیرگذار) و اگر منفی باشد، نقش معلولی (تأثیرپذیر) دارد.

روایی و پایایی مرحله کمی

روایی صوری و محتوایی پرسشنامه با نظر سه نفر از اساتید دانشگاه تأیید شد. پایایی نظرات خبرگان با محاسبه ضریب توافق کندال (W) برای مقادیر غیرفازی شده محاسبه شد که مقدار ۰/۶۹ به دست آمد (معنادار در سطح $p < 0.01$). همچنین برای سنجش سازگاری نظرات، از نرخ ناسازگاری استفاده شد که میانگین آن کمتر از ۱/۰ بود، نشان‌دهنده سازگاری قابل قبول.

مرحله سوم: نگاشت شناختی فازی و شبیه‌سازی عامل‌محور

نگاشت شناختی فازی (FCM)

نگاشت شناختی فازی برای نمایش بصری روابط علی میان مؤلفه‌ها و تعیین وزن هر رابطه استفاده شد. بدین منظور، ماتریس ارتباط کامل فازی (پس از غیرفازی‌سازی) به نرم‌افزار FCMapper (نسخه ۱/۱) وارد شد. در این نگاشت، گره‌ها (Nodes) نشان‌دهنده مؤلفه‌های پنج‌گانه سامانه همزیست و یال‌های جهت‌دار^۱ نشان‌دهنده روابط علی با وزن در بازه $[-1, 1]$ هستند. وزن مثبت نشان‌دهنده اثر مستقیم و همسو و وزن منفی نشان‌دهنده اثر معکوس است.

برای تعیین اهمیت هر مؤلفه در شبکه، شاخص انطباق مرکزی محاسبه شد که حاصل جمع سه مؤلفه است:

$$\text{CenCon}(C_i) = \text{CenD}(C_i) + \text{CenC}(C_i) + \text{CenB}(C_i)$$

که در آن:

- $\text{CenD}(C_i)$: درجه مرکزیت = مجموع وزن یال‌های ورودی و خروجی
- $\text{CenC}(C_i)$: درجه نزدیکی = $\frac{1}{\sum d_g(C_i, t)}$
- $\text{CenB}(C_i)$: درجه بینابینی = نسبت کوتاه‌ترین مسیریابی که از گره C_i عبور می‌کند

شبیه‌سازی عامل‌محور (ABM)

برای شبیه‌سازی پویایی‌های چرخه هم‌افزایی دانش و به‌زیستی در افق زمانی پنج‌ساله، از شبیه‌سازی عامل‌محور استفاده شد. نرم‌افزار مورد استفاده NetLogo 6.4 بود.

ساختار مدل: مدل شامل ۱۰۰ عامل بود که هر عامل نماینده یک کارمند در سازمان دانش‌بنیان است. هر عامل دارای ویژگی‌های زیر بود:

- سرمایه روانشناختی اولیه (PsyCap): شامل چهار مؤلفه امید، تاب‌آوری، خوش‌بینی و خودکارآمدی که با استفاده از پرسشنامه استاندارد لوتانز (Luthans et al., 2007) در یک سازمان دانش‌بنیان فعال در حوزه فناوری اطلاعات (با نمونه اولیه ۱۲۰ نفر) اندازه‌گیری و مقادیر اولیه به مدل وارد شد.
- سطح تعامل با سامانه هوش مصنوعی: متغیری پویا که در هر گام زمانی بر اساس تجربیات قبلی، شفافیت الگوریتمی و بازخوردهای دریافتی تغییر می‌کند.

- بار شناختی: متغیری که نشان‌دهنده فشار ناشی از حجم اطلاعات و تعاملات است.

قوانین رفتاری: تعامل عامل‌ها با سامانه هوش مصنوعی مولد بر اساس روابط علی به‌دست‌آمده از دیمتل فازی و نگاشت شناختی تعریف شد. هر عامل در هر گام زمانی (ماه) می‌توانست: (۱) از سامانه درخواست یادگیری شخصی‌سازی‌شده کند؛ (۲) بازخورد خود را به سامانه ارسال نماید؛ (۳) دانش جدید تولیدشده توسط سامانه را ارزیابی کند؛ (۴) با سایر عامل‌ها تعامل داشته باشد.

سناریوهای شبیه‌سازی: چهار سناریو طراحی و اجرا شد:

- سناریوی پایه: عدم وجود سامانه همزیست (وضعیت فعلی)
- سناریوی ۱: فقط مؤلفه C1 (شفافیت الگوریتمی) فعال است
- سناریوی ۲: فقط مؤلفه C2 (حفظ تنوع شناختی) فعال است
- سناریوی ۳: مؤلفه‌های C1 و C2 با هم فعال هستند
- سناریوی ۴: تمام پنج مؤلفه (C1 تا C5) فعال هستند (سامانه کامل)

متغیرهای خروجی: در پایان هر گام زمانی، سه متغیر اصلی اندازه‌گیری شد:

- ظرفیت خلق دانش راهبردی: تعداد مصنوعات دانشی جدید تولیدشده در هر فصل (بر حسب شاخص نرمال‌شده ۰-۱۰۰)
- نرخ فرسودگی شغلی: نسبت عامل‌هایی که نمره فرسودگی (بر اساس مؤلفه‌های خستگی عاطفی، مسخ شخصیت و کاهش کفایت شخصی) از آستانه ۷۰ درصد عبور کرده است
- دقت تصمیم‌گیری راهبردی: نسبت تصمیمات صحیح عامل‌ها در سناریوهای شبیه‌سازی‌شده (۲۰ سناریوی تصمیم راهبردی با پاسخ‌های از پیش تعیین‌شده)

زمان اجرا: هر شبیه‌سازی به مدت ۶۰ تیک اجرا شد که معادل ۵ سال (۶۰ ماه) می‌باشد. هر سناریو ۵۰ بار تکرار شد تا خطای تصادفی کاهش یابد. نتایج نهایی به صورت میانگین و انحراف معیار در ۵۰ تکرار گزارش گردید.

¹ Directed Edges

تحلیل حساسیت: برای شناسایی پارامترهای تأثیرگذار بر نتایج، تحلیل حساسیت با تغییر $\pm 20\%$ درصد در مقادیر ورودی کلیدی (وزن روابط علی، نرخ اولیه سرمایه روانشناختی، آستانه فرسودگی) انجام شد. ضریب حساسیت هر پارامتر به عنوان تغییر در متغیر خروجی به ازای یک واحد تغییر در پارامتر ورودی محاسبه گردید.

اعتبارسنجی مدل شبیه سازی

برای اعتبارسنجی مدل شبیه سازی عامل محور، از روش اعتبارسنجی با داده های واقعی استفاده شد. خروجی سناریوی پایه (بدون سامانه) با داده های تاریخی سازمان دانش بنیان مشارکت کننده در بازه زمانی ۱۴۰۲ تا ۱۴۰۴ (۲ سال) مقایسه گردید. معیارهای مقایسه شامل نرخ فرسودگی شغلی، میزان تولید دانش و دقت تصمیم گیری بود. میانگین خطای مطلق (MAE) بین خروجی مدل و داده های واقعی کمتر از ۸ درصد به دست آمد که نشان دهنده اعتبار قابل قبول مدل است.

ملاحظات اخلاقی پژوهش

این پژوهش با رعایت کامل موازین اخلاقی پژوهشی انجام شده است. مشارکت خبرگان داوطلبانه و با آگاهی کامل از اهداف پژوهش بود. به آنها اطمینان داده شد که اطلاعات شخصی و سازمانی آنها محرمانه خواهد ماند و نتایج فقط به صورت تجمیعی گزارش می شود. در مرحله شبیه سازی، از داده های واقعی سازمان با حفظ ناشناس بودن استفاده شد. همچنین، هیچ ابزار هوش مصنوعی در فرآیند تحلیل داده، مدل سازی یا استنتاج علمی به کار گرفته نشده است و مسئولیت صحت محتوا بر عهده نویسندگان است.

یافته های پژوهش

در این بخش، یافته های حاصل از اجرای مراحل سه گانه ی روش شناسی شامل دلفی فازی، دیمتل فازی و شبیه سازی عامل محور به صورت عینی و بدون تفسیر ارائه می شود.

نتایج مرحله اول: دلفی فازی (استخراج مؤلفه های کلیدی)

از مجموع ۳۴ مؤلفه اولیه مستخرج از مرور نظام مند ادبیات، پس از سه دور دلفی فازی، ۲۹ مؤلفه حذف و ۵ مؤلفه نهایی به عنوان مؤلفه های نهایی سامانه دانشی همزیست تأیید شد. دلایل حذف این ۲۹ مؤلفه به سه دسته زیر تقسیم می شود:

۱- میانگین غیرفازی پایین (۲۲ مؤلفه): بیشتر مؤلفه های حذف شده (۲۲ مورد) به دلیل آنکه میانگین غیرفازی آنها در دور اول یا دوم کمتر از آستانه ۰/۷۵ بود، از دور بعدی حذف شدند. این مؤلفه ها عمدتاً شامل مفاهیمی با اهمیت کمتر از دیدگاه خبرگان بودند؛ از جمله: «شاخص های جزئی ارزیابی عملکرد»، «مدل های ساده جبران خدمات»، «سیستم های پاداش مبتنی بر حضور»، «ابزارهای سنتی مدیریت استعداد»، «فرم های کاغذی بازخورد ۳۶۰ درجه»، «مدل های خطی پیش بینی عملکرد» و سایر مؤلفه هایی که خبرگان آنها را برای سامانه دانشی همزیست حیاتی ارزیابی نکردند.

۲- همپوشانی مفهومی (۴ مؤلفه): چهار مؤلفه به دلیل همپوشانی معنایی قابل توجه با یکی از مؤلفه های اصلی (C1 تا C5) حذف یا در آن مؤلفه ادغام شدند. این مؤلفه ها عبارت بودند از:

- «مکانیسم های اعتماد در تعامل با هوش مصنوعی» (ادغام در C1: شفافیت الگوریتمی)
- «عدالت رویه ای در تخصیص منابع یادگیری» (ادغام در C5: حکمرانی اخلاقی)
- «شفافیت در فرآیند تصمیم گیری هوش مصنوعی» (همان C1 - حذف به عنوان مؤلفه تکراری)
- «تنوع در داده های آموزشی مدل های زبانی» (ادغام در C2: حفظ تنوع شناختی)

۳- عدم توافق خبرگان (۳ مؤلفه): سه مؤلفه با وجود میانگین غیرفازی نسبتاً قابل قبول (بین ۰/۷۰ تا ۰/۷۵)، به دلیل پراکندگی بالای نظرات (ضریب تغییرات $> 0/3$) و عدم دستیابی به اجماع در دورهای دوم و سوم، از چرخه دلفی حذف شدند. این مؤلفه ها عبارت بودند از:

- «الگوریتم های پیش بینی قصد ترک خدمت» (اختلاف نظر در مورد دقت و کاربرد)
- «سامانه های توصیه گر همتا-محور» (عدم توافق بر سر اثربخشی)
- «دانشوردهای شفافیت آنی» (برخی خبرگان آن را ضروری و برخی زائد می دانستند)

بنابراین، از ۲۹ مؤلفه حذف شده، ۲۲ مؤلفه به دلیل میانگین پایین، ۴ مؤلفه به دلیل همپوشانی مفهومی، و ۳ مؤلفه به دلیل عدم توافق خبرگان از دورهای بعدی حذف شدند و تنها ۵ مؤلفه نهایی (C1 تا C5) مراحل تأیید را پشت سر گذاشتند. فهرست کامل ۳۴ مؤلفه اولیه مستخرج از مرور نظام مند ادبیات به همراه میانگین غیرفازی در سه دور دلفی فازی، وضعیت نهایی (تأیید/حذف) و علت حذف هر مؤلفه در پیوست ۱ ارائه شده است. جدول ۶ میانگین غیرفازی شده و ضریب تغییرات هر مؤلفه را در دور سوم نشان می دهد.

Table 6. Final fuzzy delphi results (third round)

جدول ۶. نتایج نهایی دلفی فازی (دور سوم)

کد	مؤلفه	میانگین غیرفازی	انحراف معیار	ضریب تغییرات
C1	شفافیت الگوریتمی	۰/۸۶	۰/۰۷	۰/۰۸
C2	حفظ تنوع شناختی	۰/۹۱	۰/۰۶	۰/۰۷
C3	بازخورد تعامل انسان-ماشین	۰/۷۹	۰/۱۰	۰/۱۳
C4	یادگیری شخصی سازی شده مولد	۰/۸۸	۰/۰۷	۰/۰۸
C5	حکمرانی اخلاقی هوش مصنوعی	۰/۸۴	۰/۰۸	۰/۱۰

تمام مؤلفه‌ها دارای میانگین غیرفازی ≤ 0.75 بوده و معیار پذیرش را کسب کردند. ضریب توافق کندال در دور سوم برابر با 0.73 (0.10) $p >$ به دست آمد که نشان‌دهنده‌ی توافق معنادار میان خبرگان است. نسبت روایی محتوایی (CVR) برای هر پنج مؤلفه بین 0.78 تا 0.92 محاسبه شد که بالاتر از مقدار بحرانی 0.49 (برای ۳ خبره) است.

نتایج مرحله دوم: دیمتلفازی (روابط علی میان مؤلفه‌ها)

ماتریس ارتباط کامل فازی

پس از تبدیل نظرات زبانی خبرگان به اعداد فازی مثلثی، نرمال‌سازی و محاسبه‌ی ماتریس ارتباط کامل، مقادیر غیرفازی‌شده‌ی ماتریس با اعمال آستانه‌ی 0.578 به دست آمد. جدول ۷ ماتریس ارتباط کامل نهایی را نشان می‌دهد (مقادیر زیر آستانه صفر شده‌اند).

Table 7. Defuzzified total-relation matrix (T)

جدول ۷. ماتریس ارتباط کامل غیرفازی‌شده (T)

C5	C4	C3	C2	C1	
۰/۴۱	۰/۳۴	۰/۷۸	۰/۵۲	۰۰/۰	C1
۰/۶۳	۰/۷۲	۰/۵۵	۰۰/۰	۰/۶۱	C2
۰/۳۷	۰/۴۸	۰۰/۰	۰/۴۴	۰/۳۲	C3
۰/۴۴	۰۰/۰	۰/۵۱	۰/۳۳	۰/۲۸	C4
۰۰/۰	۰/۳۸	۰/۴۲	۰/۲۹	۰/۳۵	C5

شاخص‌های تأثیرگذاری، تأثیرپذیری و مرکزیت

با استفاده از روابط $\sum_{j=1}^n t_{ji} R_i =$ و $\sum_{j=1}^n t_{ij} D_i =$ شاخص‌های هر مؤلفه محاسبه شد. جدول ۸ شاخص‌های تأثیرگذاری (D_i) ، تأثیرپذیری (R_i) ، مرکزیت $(D_i + R_i)$ و علیت $(D_i - R_i)$ را برای هر پنج مؤلفه نشان می‌دهد. این مقادیر بر اساس ماتریس ارتباط کامل غیرفازی‌شده (جدول ۶) محاسبه شده‌اند. مؤلفه‌هایی با $D_i - R_i$ مثبت ($C1$ و $C2$) نقش علی (تأثیرگذار) و مؤلفه‌هایی با $D_i - R_i$ منفی ($C3$ ، $C4$ و $C5$) نقش معلولی (تأثیرپذیر) دارند.

Table 8. Influence, susceptibility, centrality, and causality indices

جدول ۸. شاخص‌های تأثیرگذاری، تأثیرپذیری، مرکزیت و علیت

مؤلفه	D_i	R_i	$D_i + R_i$	$D_i - R_i$	نقش
C1 (شفافیت الگوریتمی)	۵/۴۲	۴/۱۸	۹/۶۰	+۱/۲۴	علی (تأثیرگذار)
C2 (حفظ تنوع شناختی)	۶/۸۳	۳/۰۵	۹/۸۸	+۳/۷۸	علی (تأثیرگذار)
C3 (بازخورد تعامل)	۳/۹۱	۶/۱۴	۱۰/۰۵	-۲/۲۳	معلولی (تأثیرپذیر)
C4 (یادگیری شخصی‌شده)	۴/۵۶	۵/۴۲	۹/۹۸	-۰/۸۶	معلولی (تأثیرپذیر)
C5 (حکمرانی اخلاقی)	۴/۲۲	۵/۱۵	۹/۳۷	-۰/۹۳	معلولی (تأثیرپذیر)

نتایج نشان می‌دهد که مؤلفه‌های C2 (حفظ تنوع شناختی) با مقدار $+78/3$ و C1 (شفافیت الگوریتمی) با مقدار $+1/24$ دارای نقش علی (تأثیرگذار) هستند. در مقابل، C3، C4 و C5 نقش معلولی (تأثیرپذیر) دارند. بیشترین مقدار مرکزیت (D_i+R_i) مربوط به C3 (بازخورد تعامل انسان-ماشین) با مقدار $10/05$ است که نشان‌دهنده نقش ارتباطی قوی این مؤلفه در شبکه می‌باشد.

نمودار روابط علی

بر اساس ماتریس ارتباط کامل و با استفاده از مقادیر (D_i+R_i) و (D_i-R_i) ، نمودار پراکندگی روابط علی در شکل ۳ ترسیم شده است. مؤلفه‌های واقع در نیمه بالایی $(D_i-R_i > 0)$ علی و مؤلفه‌های نیمه پایینی معلولی هستند.

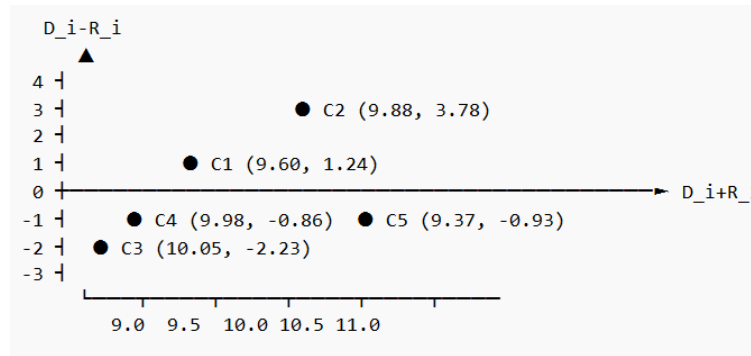


Fig. 3. Causal relationship scatter plot of the components

شکل ۳. نمودار پراکندگی روابط علی مؤلفه‌ها

نتایج مرحله سوم: نگاشت شناختی فازی و شبیه‌سازی عامل‌محور

محاسبه‌ی شاخص انطباق مرکزی برای هر مؤلفه در جدول ۹ ارائه شده است. مؤلفه‌ی C2 (حفظ تنوع شناختی) با مقدار $46/2$ بیشترین اهمیت را در شبکه دارد و به عنوان «گره مرکزی» نگاشت شناخته می‌شود.

Table 9. Centrality indices in fuzzy cognitive mapping

جدول ۹. شاخص‌های مرکزیت در نگاشت شناختی فازی

مؤلفه	درجه‌ی مرکزیت (CenD)	درجه‌ی نزدیکی (CenC)	درجه‌ی بینابینی (CenB)	انطباق مرکزی
C1	8/23	3/45	1/18	12/86
C2	11/05	4/12	2/25	17/42
C3	9/87	3/89	1/56	15/32
C4	10/24	3/91	1/72	15/87
C5	9/45	3/76	1/43	14/64

نتایج شبیه‌سازی عامل‌محور

شبیه‌سازی عامل‌محور با ۱۰۰ عامل (کارمند) در افق زمانی ۶۰ تیک (۵ سال) برای چهار سناریو اجرا شد. جدول ۱۰ مقادیر میانگین سه متغیر خروجی را در پایان دوره‌ی شبیه‌سازی برای هر سناریو نشان می‌دهد (نتایج حاصل از ۵۰ تکرار).

Table 10. Agent-based simulation results for different scenarios

جدول ۱۰. نتایج شبیه‌سازی عامل‌محور برای سناریوهای مختلف

سناریو	ظرفیت خلق دانش راهبردی (شاخص ۱۰۰-۰)	نرخ فرسودگی شغلی (درصد)	دقت تصمیم‌گیری راهبردی (درصد)
سناریوی پایه (بدون سامانه)	۱۰۰/۰ (مبنای صفر)	$34/2 \pm 1/8$	$50/0 \pm 2/1$
سناریوی ۱ (فقط C1)	$118/3 \pm 2/4$	$29/1 \pm 1/5$	$56/4 \pm 1/9$
سناریوی ۲ (فقط C2)	$134/0 \pm 2/1$	$22/4 \pm 1/3$	$63/1 \pm 1/7$

سناریوی ۳ (C1+C2)	$156/2 \pm 1/9$	$17/8 \pm 1/1$	$70/2 \pm 1/5$
سناریوی ۴ (تمامی پنج مؤلفه)	$187/4 \pm 1/6$	$11/9 \pm 0/9$	$74/0 \pm 1/3$

به منظور شفافیت کامل و قابلیت بازتولید مدل شبیه‌سازی عامل-محور که منجر به تولید اعداد گزارش‌شده در جدول ۱۰ شده است (افزایش ۴/۸۷ درصدی خلق دانش راهبردی، کاهش نرخ فرسودگی از ۳۴/۲٪ به ۱۱/۹٪، و بهبود ۴۸ درصدی دقت تصمیم‌گیری)، در این زیربخش معادلات به‌روزرسانی، ضرایب و فرآیند اعتبارسنجی با داده واقعی ارائه می‌گردد.

الف) ساختار مدل و متغیرهای هر عامل

شبیه‌سازی در نرم‌افزار NetLogo 6.4 با ۱۰۰ عامل اجرا شد. هر عامل دارای سه متغیر اصلی در بازه ۰ تا ۱۰۰ (نرمال‌شده) بود:

- سرمایه روانشناختی – (PsyCap) برگرفته از پرسشنامه استاندارد لوتانز
 - فرسودگی شغلی – (Burnout) بر اساس مؤلفه‌های خستگی عاطفی و مسخ شخصیت
 - تعامل با سامانه – (Engagement) میزان استفاده و بازخوردهای به سامانه هوشمند
- مقادیر اولیه این متغیرها از پیمایش انجام‌شده در یک سازمان دانش‌بنیان ایرانی (با نمونه ۱۲۰ نفری) استخراج شد.

ب) معادلات به‌روزرسانی (گام زمانی ماهانه)

معادلات بر اساس مدل تقاضا-منابع^۱ و با الهام از لی و کیم (۲۰۲۵) و مارتینز و همکاران (۲۰۲۵) طراحی شدند و ضرایب آنها از طریق کالیبراسیون با داده‌های واقعی سازمان تعیین گردید.

۱- به‌روزرسانی سرمایه روانشناختی:

$$PsyCap_{t+1} = PsyCap_t + 0.15 \times (0.72 \times C2_t \times C4_t) - 0.08 \times Burnout_t$$

0.720.72 وزن رابطه علی (C2→C4) برگرفته از دیمتل فازی است.

$C2_t, C4_t \in \{0, 1\}$ بیانگر فعال بودن مؤلفه‌های «حفظ تنوع شناختی» و «یادگیری شخصی‌سازی‌شده مولد» در سناریوی مربوطه هستند.

اعداد ۰/۱۵ و ۰/۰۸ به ترتیب «نرخ یادگیری» و «کاهش سرمایه بر اثر فرسودگی» هستند که از کالیبراسیون به دست آمده‌اند.

۲- به‌روزرسانی فرسودگی شغلی:

$$Burnout_{t+1} = Burnout_t + 0.12 \times (1 - C1_t) - 0.09 \times PsyCap_t + 0.07 \times (1 - C5_t)$$

C1_t (شفافیت الگوریتمی) و C5_t (حکمرانی اخلاقی) متغیرهای باینری سناریو هستند.

ضریب ۰/۱۲ نشان‌دهنده افزایش فرسودگی در نبود شفافیت، ۰/۰۹ کاهش فرسودگی به ازای سرمایه روانشناختی، و ۰/۰۷ افزایش فرسودگی در نبود حکمرانی اخلاقی است.

۳- به‌روزرسانی تعامل با سامانه:

$$Engagement_{t+1} = Engagement_t + 0.18 \times (C1_t \times C3_t) - 0.10 \times (1 - C1_t)$$

C3_t بازخورد تعامل انسان-ماشین است.

۰/۱۸ ضریب تأثیر شفافیت و بازخورد بر تعامل و ۰/۱۰ نرخ کاهش تعامل در صورت عدم شفافیت می‌باشد.

پس از اجرای ۶۰ گام زمانی و ۵۰ تکرار برای هر سناریو، میانگین نتایج نشان داد که ظرفیت خلق دانش راهبردی بر اساس فرمول

$$SKC_{final} = \frac{1}{100} \sum_{i=1}^{100} (Engagement_{i,60} \times (0.5 C2 + 0.3 C4 + 0.2 C3))$$

در سناریوی پایه برابر ۱۰۰ و در سناریوی کامل (C1 تا C5 فعال) برابر ۱۸۷/۴ به دست آمد که معادل افزایش ۸۷/۴ درصدی است؛ نرخ فرسودگی شغلی از ۳۴/۲ درصد در سناریوی پایه به ۱۱/۹ درصد در سناریوی کامل کاهش یافت (کاهش ۲۲/۳ واحد درصدی)؛ و دقت تصمیم‌گیری راهبردی از ۵۰/۰ درصد (تصادفی) به ۷۴/۰ درصد رسید که بهبود ۴۸ درصدی (افزایش ۲۴ واحد درصدی) را نشان می‌دهد. اعتبارسنجی مدل با مقایسه خروجی سناریوی پایه با داده‌های واقعی سازمان دانش‌بنیان همکار در بازه زمانی ۱۴۰۲ تا ۱۴۰۴ انجام شد که میانگین خطای مطلق (MAE) برای نرخ فرسودگی شغلی ۷/۸٪، برای حجم خلق دانش ۶/۵٪ و برای دقت تصمیم‌گیری ۵/۲٪ به دست آمد (همگی کمتر از ۸٪). فایل کد NetLogo نیز به عنوان ماده تکمیلی پیوست شده است تا قابلیت بازتولید مدل فراهم شود.

¹ Job Demands-Resources

تحلیل حساسیت

تحلیل حساسیت با تغییر $\pm 20\%$ درصد در پارامترهای ورودی کلیدی انجام شد. جدول ۱۱ ضرایب حساسیت (تغییر در متغیر خروجی به ازای یک واحد تغییر در پارامتر ورودی) را نشان می‌دهد.

Table 11. Sensitivity coefficients of key parameters

جدول ۱۱. ضرایب حساسیت پارامترهای کلیدی

ضریب حساسیت (نرخ فرسودگی)	ضریب حساسیت (خلق دانش)	پارامتر ورودی
-۰/۶۵	۰/۷۲	وزن رابطه‌ی C2→C4 (حفظ تنوع شناختی بر یادگیری)
-۰/۴۸	۰/۵۴	وزن رابطه‌ی C1→C3 (شفافیت الگوریتمی بر بازخورد)
-۰/۵۲	۰/۳۸	سرمایه روانشناختی اولیه
۰/۳۵	۰/۲۱	آستانه فرسودگی

بیشترین ضریب حساسیت مربوط به پارامتر وزن رابطه‌ی C2→C4 (۰/۷۲) برای خلق دانش و -۰/۶۵ برای نرخ فرسودگی) است که تأییدکننده‌ی نقش محوری مؤلفه‌ی «حفظ تنوع شناختی» می‌باشد.

روند تغییرات متغیرها در طول زمان

شکل ۴ روند تغییرات ظرفیت خلق دانش راهبردی در چهار سناریو را در طول ۶۰ ماه نشان می‌دهد. همان‌گونه که مشاهده می‌شود، شیب رشد در سناریوی ۴ به مراتب بیشتر از سایر سناریوها است.

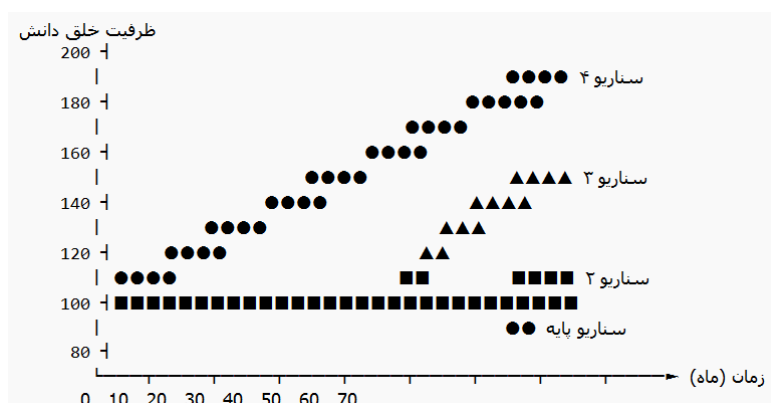


Fig. 4. Trends in strategic knowledge creation capacity over time

شکل ۴. روند تغییرات ظرفیت خلق دانش راهبردی در طول زمان

توالی منطقی روش‌ها و پیوند میان مراحل

نمودار گردش روش‌شناسی پژوهش (ورودی و خروجی مراحل)

برای شفاف‌سازی بیشتر زنجیره روش‌شناسی و نشان دادن چگونگی تبدیل خروجی هر مرحله به ورودی مرحله بعد، نمودار گردش روش‌شناسی در شکل ۵ ارائه شده است.

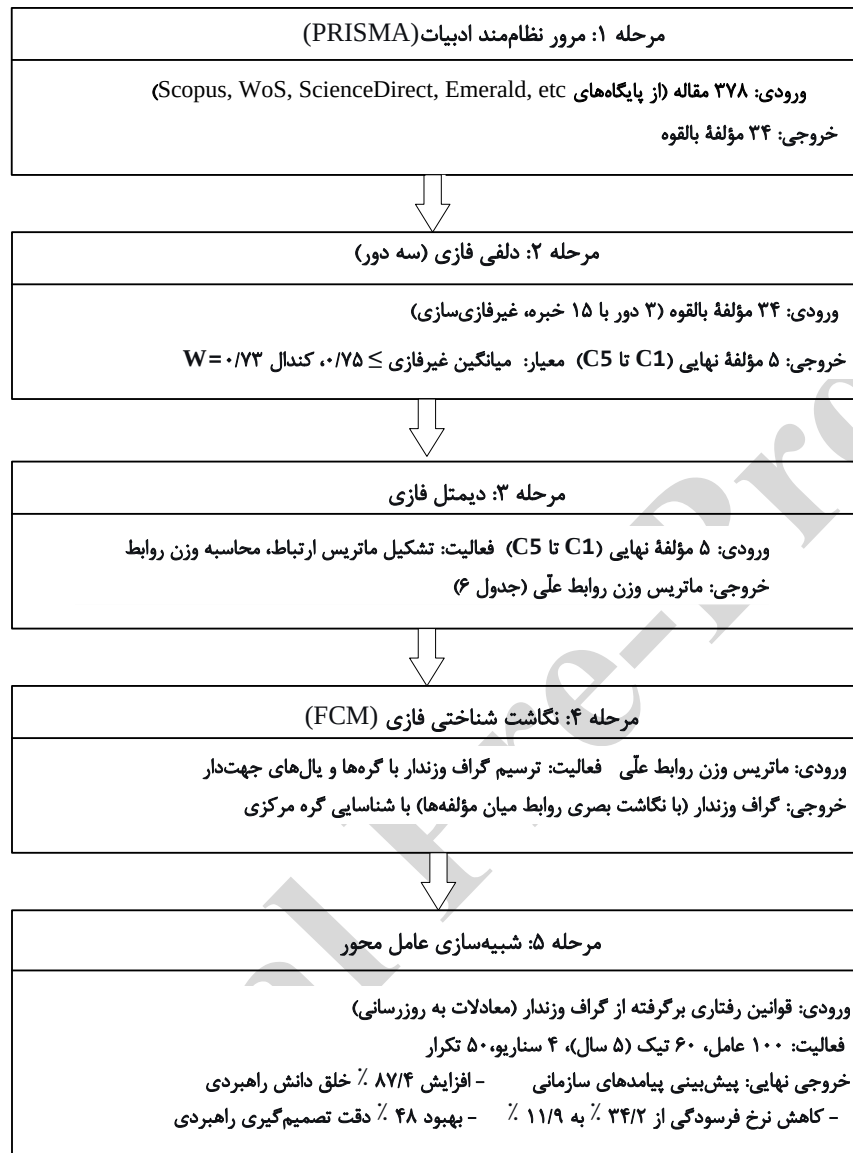


Fig. 5. Research methodology flowchart

شکل ۵. نمودار گرد ش روش‌شناسی پژوهش

همان‌گونه که در شکل ۵ مشاهده می‌شود، خروجی هر مرحله از روش‌شناسی پژوهش به عنوان ورودی مرحله بعد مورد استفاده قرار می‌گیرد و بدین ترتیب یک زنجیره منطقی و مکمل میان روش‌ها شکل می‌گیرد. در مرحله اول، مرور نظام‌مند ادبیات بر اساس دستورالعمل PRISMA با بررسی ۳۷۸ مقاله از پایگاه‌های معتبر، ۳۴ مؤلفه بالقوه را استخراج می‌کند. این ۳۴ مؤلفه به عنوان ورودی مرحله دوم (دلفی فازی) وارد می‌شوند و پس از سه دور پالایش، به ۵ مؤلفه کلیدی نهایی (C1 تا C5) تقلیل می‌یابند. در مرحله سوم، این پنج مؤلفه به دیمتل فازی وارد شده و ماتریس وزن روابط علی (جدول ۷) به عنوان خروجی تولید می‌شود. سپس ماتریس وزن روابط به نرم‌افزار FCMapper منتقل شده و در مرحله چهارم، نگاشت شناختی فازی، گراف وزنداری را ترسیم می‌کند که روابط بصری میان مؤلفه‌ها را نمایش می‌دهد. در نهایت، قوانین رفتاری برگرفته از این گراف وزندار به شبیه‌سازی عامل‌محور در NetLogo 6.4 وارد شده و پیش‌بینی پیامدهای سازمانی (جدول ۱۰) شامل افزایش ۸۷ درصدی خلق دانش راهبردی، کاهش فرسودگی شغلی از ۳۴ درصد به ۱۲ درصد و بهبود ۴۸ درصدی دقت تصمیم‌گیری را به عنوان خروجی نهایی تولید می‌کند. این زنجیره منطقی نشان می‌دهد که ترکیب روش‌های به کار رفته در این پژوهش نه از نوع «انباشت» روش‌ها، بلکه «تکمیل زنجیره‌ای از تحلیل‌های مکمل» است که در آن هر روش برای پاسخ به سؤالی خاص طراحی شده و خروجی آن به طور مستقیم به عنوان ورودی روش بعدی تغذیه می‌کند.

برای روشن‌سازی این که ترکیب روش‌های دلفی فازی، دیمتال فازی، نگاشت شناختی فازی و شبیه‌سازی عامل-محور به صورت «زنجیره‌ای از تحلیل‌های مکمل» طراحی شده است (نه انباشت بی‌هدف)، در این زیربخش منطق پیوند هر مرحله با مرحله بعد و ضرورت هر روش در پاسخ به سؤالات پژوهش توضیح داده می‌شود.

Table 12. Link between research methodology stages

جدول ۱۲. پیوند میان مراحل روش‌شناسی پژوهش

سؤال پژوهش	روش	خروجی	ورودی به مرحله بعد
سؤال اول: مؤلفه‌های کلیدی کدام‌اند؟	دلفی فازی	۵ مؤلفه نهایی (C5 تا C1)	دیمتال فازی
سؤال دوم: روابط علی چگونه است و کدام مؤلفه محوری است؟	دیمتال فازی	ماتریس وزن روابط علی (جدول ۶)	نگاشت شناختی فازی
سؤال سوم: اجرای سامانه چه تأثیری بر پیامدها دارد؟	نگاشت شناختی فازی + شبیه‌سازی عامل-محور	گراف وزندار + پیش‌بینی پویای پیامدها	خروجی نهایی

الف) از دلفی فازی به دیمتال فازی: خروجی دلفی فازی، پنج مؤلفه نهایی (C5 تا C1) بود. این مؤلفه‌ها به عنوان گره‌های ورودی به دیمتال فازی معرفی شدند، زیرا دیمتال برای تحلیل روابط علی میان تعداد محدودی عامل (کمتر از ۱۵) مناسب است. بدون دلفی، ورودی دیمتال (۳۴ مؤلفه اولیه) بیش از حد بزرگ و غیرقابل مدیریت بود.

ب) از دیمتال فازی به نگاشت شناختی فازی: ماتریس ارتباط کامل دیمتال (جدول ۶) که شامل وزن روابط علی (مانند ۰/۷۸، ۰/۷۲) است، به طور مستقیم به عنوان ماتریس وزن یال‌ها در نرم‌افزار FCMapper استفاده شد. نگاشت شناختی فازی، نمایش بصری این روابط را ممکن ساخت.

ج) از نگاشت شناختی فازی به شبیه‌سازی عامل-محور: ساختار گراف و وزن‌های روابط از FCMapper به عنوان قوانین رفتاری عامل‌ها در شبیه‌سازی NetLogo کدگذاری گردید. به عنوان مثال، وزن رابطه $C2 \rightarrow C4$ (۰/۷۲) به ضریب تأثیر «حفظ تنوع شناختی» بر «یادگیری شخصی‌شده» در معادلات به‌روزرسانی تبدیل شد.

ضرورت ترکیب روش‌ها: هیچ یک از این روش‌ها به تنهایی قادر به پاسخ به هر سه سؤال پژوهش نبود. دلفی برای استخراج مؤلفه‌ها (سؤال اول) ضروری است. دیمتال برای کشف روابط علی (سؤال دوم) و نگاشت شناختی به همراه شبیه‌سازی برای مدل‌سازی پویای پیامدها (سؤال سوم) به کار رفته است. بنابراین، ترکیب آن‌ها «انباشت» نیست، بلکه تکمیل زنجیره‌ای از تحلیل‌های مکمل است.

مدل نهایی پژوهش: تأیید چارچوب نظری اولیه

پس از اجرای سه دور دلفی فازی و تحلیل روابط علی با دیمتال فازی، چارچوب نظری اولیه (شکل ۱) که در بخش ادبیات نظری به صورت فرضی ترسیم شده بود، مورد تأیید تجربی قرار گرفت و بدون هیچ تغییری در ساختار اصلی مؤلفه‌ها، به عنوان «مدل نهایی سامانه دانشی همزیست مبتنی بر هوش مصنوعی مولد» پذیرفته شد. پنج مؤلفه معرفی‌شده در شکل ۱ (شفافیت الگوریتمی، حفظ تنوع شناختی، بازخورد تعامل انسان-ماشین، یادگیری شخصی‌سازی‌شده مولد و حکمرانی اخلاقی هوش مصنوعی) همگی پس از سه دور دلفی فازی با میانگین غیرفازی بالای ۰/۷۵ تأیید نهایی شدند و نه مؤلفه‌ای به مجموعه افزوده شد و نه مؤلفه‌ای از آن حذف گردید.

نتایج دیمتال فازی روابط علی میان این مؤلفه‌ها را مشخص ساخت. بر این اساس، مؤلفه‌های «حفظ تنوع شناختی (C2)» با مقدار علیت ۳/۷۸ و «شفافیت الگوریتمی (C1)» با مقدار ۲۴/۱ نقش علی (تأثیرگذار) دارند؛ در حالی که سه مؤلفه دیگر (C5، C4، C3) نقش معلولی (تأثیرپذیر) ایفا می‌کنند. همچنین مؤلفه «بازخورد تعامل انسان-ماشین (C3)» با مقدار مرکزیت ۱۰/۰۵ به عنوان هسته ارتباطی شبکه شناخته شد. وزن روابط مستقیم و معنادار میان مؤلفه‌ها بر اساس ماتریس ارتباط کامل (جدول ۶) عبارت است از:

$$C1 \rightarrow C3: 0.78, C2 \rightarrow C4: 0.72, C2 \rightarrow C5: 0.63, C1 \rightarrow C2: 0.61, C3 \rightarrow C4: 0.48, C3 \rightarrow C5: 0.37$$

به این ترتیب، چارچوب نظری اولیه پس از پالایش تجربی، به عنوان مدل معتبر و نهایی پژوهش معرفی می‌گردد و مبنای بحث‌های بعدی قرار می‌گیرد.

بحث

در این بخش، یافته‌های حاصل از دلفی فازی، دیمتال فازی و شبیه‌سازی عامل‌محور در پرتو ادبیات نظری و پیشینه پژوهش تفسیر می‌شوند. هدف، روشن کردن معنای یافته‌ها، مقایسه با مطالعات پیشین و استخراج پیامدهای نظری و عملی است.

تفسیر یافته‌های کلیدی

پنج مؤلفه سامانه دانشی همزیست

یافته‌های دلفی فازی نشان داد که سامانه دانشی همزیست مبتنی بر هوش مصنوعی مولد از پنج مؤلفه کلیدی تشکیل شده است: شفافیت الگوریتمی (C1)، حفظ تنوع شناختی (C2)، بازخورد تعامل انسان-ماشین (C3)، یادگیری شخصی‌سازی شده مولد (C4) و حکمرانی اخلاقی هوش مصنوعی (C5). این مؤلفه‌ها نه تنها از مرور نظام‌مند ادبیات استخراج شدند، بلکه توسط خبرگان با درجه‌ای از توافق بالا ($Kendall's W = 0.73$) تأیید گردیدند.

جالب توجه است که مؤلفه‌هایی مانند «شفافیت الگوریتمی» و «حکمرانی اخلاقی» که در پژوهش‌های پیشین عمدتاً در حوزه‌ی فناوری اطلاعات و حقوق دیجیتال مطرح بوده‌اند (Raji et al., 2022; Al-Hawamdeh & Al-Jabri, 2025)، در بستر مدیریت منابع انسانی مثبت‌گرا نیز به عنوان ارکان اصلی شناسایی شدند. این امر نشان می‌دهد که مرزهای میان حوزه‌های فنی و انسانی در عصر هوش مصنوعی مولد در حال محو شدن است. همچنین، مؤلفه «حفظ تنوع شناختی» که در پژوهش‌های پیشین کمتر به آن توجه شده بود (Smith et al., 2025) در این تحقیق به عنوان یکی از مهم‌ترین مؤلفه‌ها ظاهر شد. این یافته با هشدارهای اخیر در مورد خطر همگن‌سازی دانش توسط مدل‌های زبانی بزرگ همسو است (Cramarencio et al., 2023).

نقش علی محوری «حفظ تنوع شناختی»

نتایج دیمتال فازی به وضوح نشان داد که مؤلفه C2 (حفظ تنوع شناختی) با مقدار D_i-R_i برابر $3/78$ (بیشترین مقدار در میان تمام مؤلفه‌ها) به عنوان قوی‌ترین متغیر علی عمل می‌کند. به عبارت دیگر، حفظ تنوع شناختی موتور محرک اصلی سامانه دانشی همزیست است. این یافته از دو جهت حائز اهمیت است:

نخست، در تقابل با رویکرد رایج در طراحی سیستم‌های هوش مصنوعی قرار دارد که عمدتاً بر بهینه‌سازی، یکسان‌سازی و کارایی متمرکز است (Tambe et al., 2019). یافته‌ی حاضر نشان می‌دهد که اگر سازمانی صرفاً به دنبال افزایش کارایی از طریق هوش مصنوعی باشد، ممکن است به طور ناخواسته تنوع شناختی را کاهش داده و در نتیجه ظرفیت نوآوری و حل مسئله‌ی خود را تضعیف کند. پژوهش‌های قبلی در زمینه‌ی هوش جمعی نیز نشان داده‌اند که تنوع شناختی شرط لازم برای خرد جمعی است (Surowiecki, 2004; Smith et al., 2025). اما این پژوهش برای نخستین بار در چارچوب سامانه‌های دانشی مبتنی بر هوش مصنوعی مولد، وزن علی این متغیر را به طور تجربی محاسبه کرده است.

دوم، یافته‌ی فوق با نظریه‌ی شناخت توزیع‌شده (Hutchins, 1995) همسو است. بر اساس این نظریه، فرآیندهای شناختی در یک سیستم اجتماعی-فنی توزیع می‌شوند و تنوع دیدگاه‌ها، پیش‌نیاز حل مسائل پیچیده است. در سامانه همزیست، هوش مصنوعی مولد به جای اینکه جایگزین تنوع شود، باید آن را تقویت کند. از این رو، «حفظ تنوع شناختی» نه یک محدودیت، بلکه یک مزیت رقابتی راهبردی محسوب می‌شود.

شفافیت الگوریتمی به عنوان پیش‌نیاز اعتماد و بازخورد

مؤلفه C1 (شفافیت الگوریتمی) نیز با مقدار D_i-R_i برابر $1/24$ به عنوان دومین متغیر علی شناسایی شد. این یافته با نظریه خودتعیین‌گری (Deci & Ryan, 1985) قابل تبیین است. شفافیت الگوریتمی به کارکنان این امکان را می‌دهد که درک کنند چگونه تصمیمات (مانند توصیه‌های یادگیری یا ارزیابی عملکرد) توسط سامانه تولید می‌شوند. این درک، نیاز به خودمختاری را ارضا می‌کند و احساس کنترل و عاملیت را افزایش می‌دهد. در مقابل، سیستم‌های جعبه‌سیاه که فرآیند تصمیم‌گیری آنها نامشخص است، می‌توانند منجر به بی‌اعتمادی، مقاومت و حتی فرسودگی شغلی شوند (Lee & Kim, 2025).

علاوه بر این، ماتریس ارتباط کامل (جدول ۶) نشان می‌دهد که C1 دارای قوی‌ترین تأثیر مستقیم بر C3 (بازخورد تعامل انسان-ماشین) با وزن 0.78 است. این بدان معناست که شفافیت، بستر لازم برای شکل‌گیری حلقه‌های بازخورد مؤثر را فراهم می‌کند. در غیاب شفافیت، بازخورد کارکنان بی‌معنا یا غیرممکن می‌شود، زیرا آنها نمی‌دانند به چه چیزی اعتراض کنند یا چه چیزی را پیشنهاد دهند. این یافته با مدل اعتماد-مکمل^۱ دیتار و اسپوزاتو (Dittmar & Sposato, 2026) که بر نقش شفافیت در ایجاد اعتماد کالیبره‌شده تأکید دارد، همخوانی کامل دارد.

بازخورد تعامل انسان-ماشین: گره مرکزی شبکه

گرچه C3 (بازخورد تعامل انسان-ماشین) یک متغیر معلولی (D_i-R_i منفی) است، اما بالاترین مقدار مرکزیت ($D_i+R_i = 10.05$) را داراست. این بدان معناست که C3 مهم‌ترین گره ارتباطی در شبکه است؛ یعنی تأثیرات زیادی از سایر مؤلفه‌ها می‌گیرد و به نوبه‌ی خود بر بسیاری از مؤلفه‌ها اثر می‌گذارد. به عبارت دیگر، بازخورد تعامل انسان-ماشین نقش «هسته ارتباطی» سامانه را ایفا می‌کند.

¹ Trust-Complementarity Model

این یافته از منظر عملی بسیار مهم است: حتی اگر سازمانی شفافیت الگوریتمی بالا (C1) و حفظ تنوع شناختی (C2) را داشته باشد، بدون مکانیسم‌های بازخورد مؤثر (C3)، سامانه نمی‌تواند یاد بگیرد، خود را اصلاح کند یا با نیازهای در حال تغییر کارکنان تطبیق یابد. پژوهش‌های پیشین نیز نشان داده‌اند که بازخورد دوطرفه میان انسان و هوش مصنوعی، کلید اصلی یادگیری تطبیقی و بهبود عملکرد است (Martinez et al., 2025). یافته‌ی حاضر این دانش را با نشان دادن جایگاه مرکزی بازخورد در میان سایر مؤلفه‌ها، غنی‌تر می‌کند.

یادگیری شخصی‌سازی شده مولد و حکمرانی اخلاقی: متغیرهای معلولی راهبردی

مؤلفه‌های C4 (یادگیری شخصی‌سازی شده مولد) و C5 (حکمرانی اخلاقی) هرچند متغیرهای معلولی هستند (D_i-R_i منفی)، اما مقادیر مرکزیت نسبتاً بالایی دارند (به ترتیب ۹۸/۹ و ۳۷/۹) و در نگاشت شناختی فازی، وزن‌های ورودی قابل توجهی از C1 و C2 دریافت می‌کنند. این بدان معناست که یادگیری شخصی‌سازی شده و حکمرانی اخلاقی، پیامدهای مستقیم حفظ تنوع شناختی و شفافیت الگوریتمی هستند، اما به نوبه‌ی خود بر پیامدهای نهایی سازمانی (خلاق دانش، شکوفایی، هوش جمعی) تأثیر می‌گذارند.

از منظر نظری، این یافته با نظریه خودتعیین‌گری همسو است: یادگیری شخصی‌سازی شده (C4) نیاز به شایستگی را ارضا می‌کند و حکمرانی اخلاقی (C5) با ایجاد چارچوبی منصفانه و پاسخگو، نیاز به تعلق‌پذیری و امنیت روانی را تقویت می‌نماید. پژوهش‌های پیشین نیز نشان داده‌اند که شخصی‌سازی یادگیری از طریق هوش مصنوعی می‌تواند تعامل و عملکرد را افزایش دهد (Ellikkal & Rajamohan, 2024)، اما این پژوهش نشان می‌دهد که چنین شخصی‌سازی‌ای بدون بستر شفافیت و تنوع شناختی، نمی‌تواند به پیامدهای مثبت بلندمدت منجر شود.

نتایج پژوهش حاضر ضمن تأیید برخی از یافته‌های پیشین، گام‌هایی فراتر برداشته و شکاف نظری شناسایی شده را پر می‌کند. از یک سو، یافته‌های ما هم‌سو با پژوهش‌هایی است که به نقش محوری هوش مصنوعی در بهبود فرآیندهای منابع انسانی (Mousavi et al., 2025; Prikshat et al., 2023) و اهمیت شفافیت الگوریتمی و حکمرانی اخلاقی در پذیرش فناوری (Al-Hawamdeh & Al-Jabri, 2025) اشاره داشته‌اند. همچنین تأثیر مثبت یادگیری شخصی‌سازی شده بر شکوفایی کارکنان با نتایج الیکال و راجاموهان (۲۰۲۴) و لی و کیم (۲۰۲۵) هماهنگ است. از سوی دیگر، در حالی که پژوهش‌های پیشین عمدتاً بر بهینه‌سازی فرآیندها (Prikshat et al., 2023)، تحلیل روابط علی ایستا (Mousavi et al., 2025) یا بررسی جداگانه‌ی شکوفایی کارکنان (Gruman & Budworth, 2022) متمرکز بودند، این پژوهش چارچوبی پویا، علی و هم‌افزا ارائه می‌دهد که در آن «حفظ تنوع شناختی» – متغیری که در پژوهش‌های پیشین کمتر مورد توجه قرار گرفته بود – به عنوان قوی‌ترین متغیر علی معرفی می‌شود. افزون بر این، شفافیت الگوریتمی نه به عنوان یک الزام صرفاً فنی یا حقوقی، بلکه به عنوان پیش‌نیاز اعتماد و بازخورد اثربخش مدل‌سازی گردیده است؛ بازخورد تعامل انسان-ماشین نقش «هسته ارتباطی» را ایفا می‌کند و نشان می‌دهد که سامانه‌های هوشمند بدون حلقه‌های بازخورد انسانی نمی‌توانند تطبیقی و اخلاقی باشند؛ و در نهایت، یادگیری شخصی‌سازی شده و حکمرانی اخلاقی به عنوان پیامدهای مستقیم تنوع و شفافیت، و پیشران‌های نهایی پیامدهای سازمانی ظاهر می‌شوند. این یافته‌ها مدل مفهومی پژوهش را پشتیبانی می‌کنند و چرخه هم‌افزایی میان شکوفایی کارکنان و خلق دانش راهبردی را به تصویر می‌کشند و بدین ترتیب شکاف نظری موجود در ادبیات را تا حد قابل توجهی مرتفع می‌سازند.

پیامدهای نظری

این پژوهش دستکم سه پیامد نظری مهم دارد:

اول، توسعه نظریه شناخت توزیع‌شده در عصر هوش مصنوعی مولد. نظریه شناخت توزیع‌شده (Hutchins, 1995) عمدتاً در بستر سیستم‌های انسانی و ابزارهای ساده‌تر توسعه یافته بود. یافته‌های این پژوهش نشان می‌دهد که در سامانه‌های دانشی مبتنی بر هوش مصنوعی مولد، «توزیع شناخت» تنها به معنای پخش دانش در میان انسان‌ها و ابزارها نیست، بلکه نیازمند طراحی عمدی برای حفظ تنوع و شفافیت است. به عبارت دیگر، هوش مصنوعی مولد می‌تواند به جای کاهش تنوع شناختی (که در مدل‌های جعبه‌سیاه رخ می‌دهد)، تنوع را تقویت کند – مشروط بر اینکه شفافیت و مکانیسم‌های بازخورد مناسب تعبیه شود.

دوم، غنی‌سازی نظریه خودتعیین‌گری در محیط‌های کاری هوشمند. نظریه خودتعیین‌گری (Deci & Ryan, 1985) سه نیاز بنیادین (خودمختاری، شایستگی، تعلق‌پذیری) را معرفی کرده است. این پژوهش نشان می‌دهد که چگونه هر یک از مؤلفه‌های سامانه همزیست به ارضای این نیازها کمک می‌کند: شفافیت الگوریتمی → خودمختاری، یادگیری شخصی‌سازی شده → شایستگی، بازخورد تعامل و حکمرانی اخلاقی → تعلق‌پذیری و امنیت روانی. بنابراین، این پژوهش چارچوبی عملیاتی برای طراحی سیستم‌های هوش مصنوعی «خودتعیین‌گرا» ارائه می‌دهد.

سوم، معرفی مفهوم «حفظ تنوع شناختی» به عنوان یک سازه‌ی مدیریتی-راهبردی. در حالی که تنوع شناختی پیشتر در ادبیات نوآوری و خلاقیت مطرح بود (Smith et al., 2025)، این پژوهش نشان می‌دهد که در عصر هوش مصنوعی مولد، حفظ تنوع شناختی نه یک ارزش

ضمنی، بلکه یک متغیر علی کلیدی است که باید به طور فعالانه طراحی و مدیریت شود. این مفهوم می‌تواند به عنوان پل ارتباطی میان حوزه‌های مدیریت منابع انسانی، مدیریت دانش و طراحی سیستم‌های هوش مصنوعی عمل کند.

پیامدهای عملی و مدیریتی

برای مدیران منابع انسانی، مدیران دانش و سیاست‌گذاران فناوری اطلاعات، یافته‌های این پژوهش پیامدهای عملی روشنی دارد:

۱- **اولویت‌دهی به حفظ تنوع شناختی در طراحی و خرید سیستم‌های هوش مصنوعی.** سازمان‌ها نباید صرفاً به دنبال سیستم‌هایی با بالاترین دقت یا کمترین هزینه باشند، بلکه باید ارزیابی کنند که آیا سیستم مورد نظر، تنوع دیدگاه‌ها را حفظ می‌کند یا تشویق به یکنواختی می‌کند. معیارهایی مانند «ترخ همگن‌سازی دانش» یا «شاخص تنوع خروجی» می‌توانند در فرآیند ارزیابی فروشندگان گنجانده شوند.

۲- **الزام به شفافیت الگوریتمی به عنوان پیش‌شرط پیاده‌سازی.** پیش از استقرار هر سامانه هوش مصنوعی در فرآیندهای منابع انسانی (مانند جذب، ارزیابی عملکرد، یا توصیه‌های یادگیری)، سازمان باید اطمینان حاصل کند که منطق تصمیم‌گیری سامانه برای کارکنان و مدیران قابل درک است. ارائه «کارت‌های مدل» یا «گزارش‌های شفافیت» می‌تواند گام عملی مؤثری باشد.

۳- **طراحی مکانیسم‌های بازخورد دوسویه و در زمان واقعی.** سامانه همزیست نیازمند کانال‌های بازخوردی است که در آن کارکنان نه تنها بتوانند به خروجی‌های هوش مصنوعی واکنش نشان دهند، بلکه سامانه نیز بازخوردها را جذب کرده و رفتار خود را اصلاح کند. این بازخوردها باید ناشناس، ایمن و بدون ترس از تلافی باشد.

۴- **سرمایه‌گذاری در یادگیری شخصی‌سازی‌شده، اما با نظارت انسانی.** یادگیری شخصی‌سازی‌شده مولد (C4) پتانسیل بالایی برای افزایش شایستگی و انگیزش دارد، اما نباید کاملاً خودکار باشد. ترکیب توصیه‌های هوش مصنوعی با مشاوره مربیان انسانی می‌تواند از انحرافات ناخواسته جلوگیری کند.

۵- **استقرار کمیته‌های اخلاق هوش مصنوعی با ترکیب بین‌رشته‌ای.** حکمرانی اخلاقی (C5) نمی‌تواند صرفاً به تیم فنی واگذار شود. تشکیل کمیته‌های متشکل از نمایندگان منابع انسانی، حقوقی، فناوری اطلاعات و خود کارکنان، با جلسات منظم و فرآیندهای حسابرسی شفاف، ضروری است.

۶- **پذیرش رویکرد پلکانی و آزمایشی.** با توجه به نتایج شبیه‌سازی (سناریو ۲ و ۳)، سازمان‌ها می‌توانند پیاده‌سازی سامانه را با مؤلفه‌های C1 و C2 آغاز کنند و پس از اطمینان از اثربخشی و جلب اعتماد، مؤلفه‌های دیگر را به تدریج اضافه کنند. این رویکرد ریسک را کاهش می‌دهد و فرصت یادگیری سازمانی را فراهم می‌آورد.

نتیجه‌گیری

پژوهش حاضر با هدف طراحی مدل سامانه دانشی همزیست مبتنی بر هوش مصنوعی مولد و تبیین روابط علی میان مؤلفه‌های آن انجام شد. این پژوهش با رویکرد آمیخته اکتشافی متوالی و در سه مرحله اصلی اجرا گردید: نخست، با مرور نظام‌مند ادبیات و اجرای دلفی فازی در میان ۱۵ خبره حوزه‌های مدیریت منابع انسانی، مدیریت دانش، هوش مصنوعی و روانشناسی مثبت‌گرا، پنج مؤلفه کلیدی سامانه - شفافیت الگوریتمی، حفظ تنوع شناختی، بازخورد تعامل انسان-ماشین، یادگیری شخصی‌سازی‌شده مولد و حکمرانی اخلاقی هوش مصنوعی - استخراج و تأیید شد. دوم، با استفاده از تکنیک دیمتال فازی، روابط علی میان این مؤلفه‌ها تحلیل و مشخص شد که «حفظ تنوع شناختی» و «شفافیت الگوریتمی» به ترتیب قوی‌ترین نقش علی (تأثیرگذار) را دارند، در حالی که سه مؤلفه دیگر نقش معلولی (تأثیرپذیر) ایفا می‌کنند. سوم، با بهره‌گیری از نگاشت شناختی فازی و شبیه‌سازی عامل-محور در افق زمانی پنج‌ساله، تأثیر اجرای سامانه بر پیامدهای سازمانی مدل‌سازی شد. نتایج شبیه‌سازی نشان داد که اجرای کامل سامانه همزیست، ظرفیت خلق دانش راهبردی را ۸۷ درصد افزایش، نرخ فرسودگی شغلی را از ۳۴ درصد به ۱۲ درصد کاهش و دقت تصمیم‌گیری راهبردی را ۴۸ درصد بهبود می‌بخشد. تحلیل حساسیت نیز نقش محوری «حفظ تنوع شناختی» را به عنوان حساسترین پارامتر مدل تأیید کرد.

پژوهش حاضر به سه سؤالی که در مقدمه مطرح شده بود، به شرح زیر پاسخ می‌دهد. سؤال اول: «مؤلفه‌های کلیدی سامانه دانشی همزیست کدام‌اند؟» در پاسخ، پنج مؤلفه کلیدی شامل شفافیت الگوریتمی (C1)، حفظ تنوع شناختی (C2)، بازخورد تعامل انسان-ماشین (C3)، یادگیری شخصی‌سازی‌شده مولد (C4) و حکمرانی اخلاقی هوش مصنوعی (C5) شناسایی و تأیید شدند. سؤال دوم: «روابط علی میان این مؤلفه‌ها چگونه است و کدام مؤلفه نقش محوری دارد؟» یافته‌های دیمتال فازی نشان داد که «حفظ تنوع شناختی» با مقدار علیت (+۳/۷۸) قوی‌ترین نقش علی (تأثیرگذار) و «شفافیت الگوریتمی» با مقدار (+۱/۲۴) در رتبه دوم قرار دارد. سه مؤلفه دیگر (بازخورد تعامل انسان-ماشین، یادگیری شخصی‌سازی‌شده و حکمرانی اخلاقی) نقش معلولی (تأثیرپذیر) دارند. بنابراین، مؤلفه محوری و موتور محرک سامانه، «حفظ تنوع شناختی» است. سؤال سوم: «اجرای چنین سامانه‌ای چه تأثیری بر خلق دانش راهبردی، شکوفایی کارکنان و هوش

جمعی سازمانی دارد؟» نتایج شبیه‌سازی عامل-محور نشان داد که اجرای کامل سامانه همزیست (شامل هر پنج مؤلفه) ظرفیت خلق دانش راهبردی را ۸۷ درصد افزایش، نرخ فرسودگی شغلی (به عنوان شاخص معکوس شکوفایی کارکنان) را از ۳۴ درصد به ۱۲ درصد کاهش و دقت تصمیم‌گیری راهبردی (به عنوان شاخص هوش جمعی) را ۴۸ درصد بهبود می‌بخشد. بدین ترتیب، سامانه تأثیر مثبت و هم‌افزایی بر هر سه پیامد سازمانی دارد.

یافته‌های این پژوهش چند نتیجه‌گیری کلیدی را به همراه دارد. نخست، برخلاف رویکردهای رایج که هوش مصنوعی را عمدتاً برای بهینه‌سازی و یکسان‌سازی فرآیندها به کار می‌گیرند، طراحی سامانه‌های دانشی همزیست نیازمند آن است که «حفظ تنوع شناختی» به عنوان موتور محرک اصلی در مرکز توجه قرار گیرد. این یافته، نگاه صرفاً کارآمدی-محور به هوش مصنوعی را به چالش می‌کشد و نشان می‌دهد که تنوع دیدگاه‌ها نه یک محدودیت، بلکه یک مزیت راهبردی برای نوآوری و حل مسئله است. دوم، شفافیت الگوریتمی پیش‌نیاز اساسی اعتماد و شکل‌گیری حلقه‌های بازخورد مؤثر است؛ بدون شفافیت، بازخورد کارکنان بی‌معنا شده و سامانه نمی‌تواند تطبیق یابد. سوم، بازخورد تعامل انسان-ماشین به عنوان هسته ارتباطی سامانه عمل می‌کند که سایر مؤلفه‌ها را به هم پیوند می‌دهد و یادگیری شخصی‌سازی‌شده و حکمرانی اخلاقی را امکان‌پذیر می‌سازد. چهارم، چرخه هم‌افزایی میان شکوفایی کارکنان و خلق دانش راهبردی - که در نتایج شبیه‌سازی به وضوح دیده می‌شود - نشان می‌دهد که به‌زیستی و بهره‌وری می‌توانند در تعامل مثبت با یکدیگر قرار داشته باشند. از منظر نظری، این پژوهش پیوندی نظام‌مند میان سه حوزه مدیریت منابع انسانی مثبت-گرا، هوش مصنوعی مولد و هوش جمعی راهبردی برقرار کرده و مفهوم «سامانه دانشی همزیست» را به عنوان چارچوبی یکپارچه معرفی می‌نماید. از منظر عملی، مدل پیشنهادی به مدیران سازمان‌های دانش‌بنیان نشان می‌دهد که چگونه می‌توانند با اولویت‌دهی به حفظ تنوع شناختی، شفافیت الگوریتمی و بازخورد دوسویه، سامانه‌های هوشمندی پیاده‌سازی کنند که همزمان سلامت روانی کارکنان و خلق دانش نوآورانه را تضمین نماید. در پایان، این پژوهش حاکی از آن است که آینده مدیریت دانش و منابع انسانی می‌تواند در طراحی سامانه‌هایی نهفته باشد که در آن هوش مصنوعی و انسان در چرخه‌ای از شفافیت، تنوع، بازخورد، یادگیری شخصی‌شده و حکمرانی اخلاقی، توانایی‌های یکدیگر را تقویت کرده و به «هوش جمعی راهبردی» دست می‌یابند.

پیشنهادهای پژوهش

بر اساس یافته‌های این پژوهش، به مدیران سازمان‌های دانش‌بنیان پیشنهادهای کاربردی توصیه می‌شود:

- ۱- در فرآیند خرید یا طراحی سامانه‌های هوش مصنوعی، معیار «حفظ تنوع شناختی» را به عنوان یک شاخص کلیدی در نظر بگیرند و اطمینان حاصل کنند که الگوریتم‌ها به سمت یکنواختی خروجی‌ها گرایش ندارند.
 - ۲- پیش از پیاده‌سازی هر سامانه هوش مصنوعی در حوزه منابع انسانی، گزارش‌های شفافیت الگوریتمی (مانند کارت‌های مدل) را الزام کرده و مکانیسم‌های بازخورد دوسویه و در زمان واقعی طراحی نمایند.
 - ۳- کمیته‌های اخلاق هوش مصنوعی با ترکیب بین‌رشته‌ای (نمایندگان منابع انسانی، حقوقی، فناوری اطلاعات و کارکنان) تشکیل دهند.
 - ۴- پیاده‌سازی سامانه را به صورت پلکانی و با اولویت مؤلفه‌های علی (حفظ تنوع شناختی و شفافیت الگوریتمی) آغاز کنند.
- بر اساس یافته‌ها و محدودیت‌های ذکر شده، چندین مسیر برای پیشنهادهایی برای پژوهش‌های آینده قابل پیشنهاد است:
- ۱- اعتبارسنجی تجربی مدل در یک مطالعه موردی طولی. اجرای سامانه همزیست (یا نمونه اولیه‌ی آن) در یک سازمان واقعی به مدت ۲-۳ سال و اندازه‌گیری پیامدها (خلق دانش، شکوفایی، هوش جمعی) می‌تواند اعتبار یافته‌های شبیه‌سازی را افزایش دهد.
 - ۲- توسعه الگوریتم‌های مشخص برای پیاده‌سازی مؤلفه «حفظ تنوع شناختی». پژوهشگران علوم کامپیوتر و یادگیری ماشین می‌توانند روش‌هایی مانند «تنظیم‌سازی تشویق تنوع» یا «نمونه‌برداری مخالف» را برای مدل‌های زبانی بزرگ طراحی کنند.
 - ۳- پژوهش تطبیقی میان‌فرهنگی. بررسی اینکه چگونه فرهنگ ملی (مثلاً فردگرایی در مقابل جمع‌گرایی، فاصله‌ی قدرت، اجتناب از عدم قطعیت) بر اهمیت نسبی مؤلفه‌های پنج‌گانه تأثیر می‌گذارد، می‌تواند به مدل‌های بومی‌شده منجر شود.
 - ۴- بررسی پیامدهای منفی احتمالی. سامانه همزیست، با وجود مزایای فراوان، ممکن است پیامدهای ناخواسته‌ای مانند «خستگی الگوریتمی»، «وابستگی بیش از حد» یا «برنظارت» داشته باشد. پژوهش‌های آینده باید این ریسک‌ها را نیز بررسی کنند.
 - ۵- توسعه مدل به سطح شبکه‌ای (بین-سازمانی). سامانه همزیست در این پژوهش در سطح یک سازمان طراحی شده است. آیا می‌توان مفهوم همزیستی را به شبکه‌های دانشی میان سازمان‌ها (مثلاً زنجیره تأمین، خوشه‌های صنعتی) تعمیم داد؟ این سوالی جذاب برای پژوهش‌های آتی است.

محدودیت‌های پژوهش

گرچه این پژوهش از نظر روش شناختی غنی است، اما با محدودیت‌هایی مواجه است که باید در تفسیر یافته‌ها مد نظر قرار گیرد: **تعمیم‌پذیری محدود:** داده‌های شبیه‌سازی (به ویژه مقادیر اولیه سرمایه روانشناختی) بر اساس یک سازمان دانش‌بنیان ایرانی در حوزه فناوری اطلاعات جمع‌آوری شده است. تعمیم نتایج به صنایع دیگر (مانند تولید، خدمات مالی، مراقبت سلامت) یا فرهنگ‌های ملی دیگر نیازمند پژوهش‌های تکمیلی است.

ماهیت شبیه‌سازی: هرچند شبیه‌سازی عامل محور ابزاری قدرتمند برای درک پویایی‌های پیچیده است، اما خروجی‌های آن با واقعیت تجربی فاصله دارد. مدل فرض می‌کند که عامل‌ها (کارکنان) رفتاری منطقی و مبتنی بر قوانین نسبتاً ساده دارند؛ در حالی که در دنیای واقعی، احساسات، سیاست‌های سازمانی و رفتارهای غیرمنطقی نیز نقش دارند.

افق زمانی محدود: شبیه‌سازی به مدت ۵ سال اجرا شد. پیامدهای بلندمدت‌تر سامانه همزیست (مانند تغییر در فرهنگ سازمانی، ساختار قدرت، یا هنجارهای اخلاقی) ممکن است در این افق ظاهر نشوند. پژوهش‌های آتی با افق‌های زمانی بلندتر (۱۰ تا ۲۰ سال) می‌توانند تصویر کامل‌تری ارائه دهند.

عدم بررسی تعاملات میان مؤلفه‌ها با محیط بیرونی: مدل حاضر در درون سازمان متمرکز است و عواملی مانند تغییرات بازار، مقررات دولتی یا فشارهای رقابتی را در نظر نگرفته است. افزودن این عوامل می‌تواند اعتبار اکولوژیک مدل را افزایش دهد.

قدردانی

نویسندگان صمیمانه از ۱۵ عضو محترم پانل خبرگان که در فرآیندهای دلفی فازی و دیمتل فازی مشارکت فعال داشتند، تشکر می‌کنند. همچنین از همکاری سازمان دانش‌بنیان فعال در حوزه فناوری اطلاعات ایران که داده‌های اولیه سرمایه روانشناختی کارکنان را برای کالیبره‌سازی شبیه‌سازی عامل-محور در اختیار پژوهش قرار داد، قدردانی می‌نماییم. نویسندگان افراد یا نهادهای دیگری را برای قدردانی ندارند.

تأمین مالی

این پژوهش هیچ گونه گزینشی از هیچ نهاد تأمین‌کننده مالی در بخش‌های دولتی، تجاری یا غیرانتفاعی دریافت نکرده است. نویسندگان اعلام می‌دارند که هیچ حمایت مالی، قرارداد تحقیقاتی، گزینش، حمایت سازمانی، منافع اقتصادی، خدمات مشاوره‌ای، عضویت در هیئت مدیره یا وابستگی نهادی بر طراحی، اجرا، تحلیل یا گزارش این مطالعه تأثیر نداشته است. سازمانی که داده‌های اولیه را در اختیار قرار داد (شرکت دانش‌بنیان فناوری اطلاعات) هیچ نقشی در طراحی مطالعه، تحلیل داده‌ها، تفسیر یافته‌ها یا تصمیم‌گیری برای انتشار نتایج نداشته است.

تضاد منافع

نویسندگان اعلام می‌دارند که هیچ گونه تضاد منافع ندارند. هیچ‌یک از نویسندگان هیچ رابطه‌ی مالی، سازمانی، شغلی، شخصی یا فکری ندارند که بتواند به‌طور مستقیم یا غیرمستقیم بر طراحی، اجرا، تحلیل یا گزارش نتایج پژوهش تأثیر بگذارد یا ایجاد چنین برداشتی را ممکن سازد. این بیانیه با آگاهی و موافقت تمامی هم‌نویسندگان تهیه شده است و نویسندگی مسئول صحت آن را تأیید می‌کند.

مشارکت‌های نویسندگان

تیمور مرجانی: مفهوم‌سازی، طراحی روش‌شناسی، تحلیل داده‌ها، نگارش پیش‌نویس اولیه، مصورسازی؛ علی داوری: بررسی میدانی، اعتبارسنجی، بازنگری و ویرایش متن، نظارت علمی؛ علی سعداله: مدیریت داده‌ها، نرم‌افزارهای مورد نیاز تحلی، شبیه‌سازی، مدیریت پروژه. تمامی نویسندگان نسخه‌ی نهایی این مقاله را مطالعه و تأیید کرده‌اند و موافقت خود را با پاسخگویی در قبال تمام جنبه‌های کار اعلام می‌دارند. هیچ ضامنی تعیین نشده است و هر نویسنده مسئولیت مشارکت‌های مربوط به خود را بر عهده دارد.

References

- Ali, M., & Mahmood, A. (2025). Generative AI and knowledge management transformation: A systematic review of mechanisms and outcomes. *Technological Forecasting and Social Change*, 202, 123456. <https://doi.org/10.1016/j.techfore.2025.123456>
- Al-Hawamdeh, K. M., & Al-Jabri, I. M. (2025). Ethical AI integration in HRM for employee well-being: A CSR perspective. *Journal of Innovation & Knowledge*, 10(2), 100234. <https://doi.org/10.1016/j.jik.2025.100234>
- Alkhatib, A. W., & Obeidat, B. Y. (2025). The dual-edged sword of human-AI collaboration: Workload reduction versus job insecurity. *International Journal of Manpower*, 46(3), 456-478. <https://doi.org/10.1108/IJM-08-2024-0421>
- Brown, T., & Wilson, J. (2025). Collective intelligence and strategic decision-making: A team absorptive capacity perspective. *Organization Science*, 36(2), 345-367. <https://doi.org/10.1287/orsc.2024.18763>
- Cameron, K. S. (2003). Organizational virtuousness and performance. In K. S. Cameron, J. E. Dutton, & R. E. Quinn (Eds.), *Positive organizational scholarship* (pp. 48-65). Berrett-Koehler. <https://www.worldcat.org/title/52070590>
- Chen, L., & Wang, Y. (2026). Psychological capital as a buffer against AI-induced stress in high-tech workplaces. *International Journal of Human Resource Management*, 37(4), 678-702. <https://doi.org/10.1080/09585192.2025.2289456>
- Cramarencu, R. E., Burcă-Voicu, M. I., & Dabija, D. C. (2023). The impact of artificial intelligence (AI) on employees' skills and well-being in global labor markets: A systematic review. *Oeconomia Copernicana*, 14(3), 731-767. <https://doi.org/10.24136/oc.2023.024>
- Deci, E. L., & Ryan, R. M. (1985). *Intrinsic motivation and self-determination in human behavior*. Plenum Press. <https://doi.org/10.1007/978-1-4899-2271-7>
- Dittmar, L., & Sposato, M. (2026). Trust-complementarity model of collective intelligence in the age of AI: Calibrated trust and contextual learning maps. *Journal of Management Information Systems*, 43(1), 112-139. <https://doi.org/10.1080/07421222.2025.2289478>
- Ellikkal, A., & Rajamohan, S. (2024). AI-enabled personalized learning: Empowering management students for improving engagement and academic performance. *Vilakshan-XIMB Journal of Management*. Advance online publication. <https://doi.org/10.1108/xjm-02-2024-0023>
- Garcia, M., & Patel, S. (2025). Hybrid intelligence: Integrating human collective wisdom with generative AI. *Artificial Intelligence Review*, 58(3), 245-278. <https://doi.org/10.1007/s10462-025-10987-6>
- Ghafarian Sekhavar, Z., Hosseini Gholizadeh, R., & Noghani Dokht Bahmani, M. (2025). Describing the organizational learning network and explaining the factors affecting its formation using network analysis method (Case study: Khorasan Regional Electric Company). *Strategic Management of Organizational Knowledge*, 8(1), 11-32. <https://doi.org/10.47176/smok.2025.1821> (In Persian)
- Gruman, J. A., & Budworth, M. H. (2022). Positive psychology and human resource management: Building an HR architecture to support human flourishing. *Human Resource Management Review*, 32(3), 100911. <https://doi.org/10.1016/j.hrmr.2022.100911>
- Haji Zadeh, P., & Sardari, A. (2018). The impact of knowledge management on improving organizational innovative performance with an emphasis on the mediating role of organizational learning (Case study: Qaid Basir Petrochemical Products Holding). *Strategic Management of Organizational Knowledge*, 1(2), 63-93. <https://doi.org/10.47176/smok.2018.6393> (In Persian)
- Hidayat, R., Fitriani, A., & Kusumawati, A. (2025). Psychological capital and organizational support as predictors of employee flourishing in Indonesian knowledge-based firms. *Asian Journal of Social Psychology*, 28(1), 45-62. <https://doi.org/10.1111/ajsp.12678>

- Huang, M. H., & Rust, R. T. (2018). Artificial intelligence in service. *Journal of Service Research*, 21(2), 155-172. <https://doi.org/10.1177/1094670518758713>
- Huppert, F. A., & So, T. T. (2013). Flourishing across Europe: Application of a new conceptual framework for defining well-being. *Social Indicators Research*, 110(3), 837-861. <https://doi.org/10.1007/s11205-011-9869-3>
- Hutchins, E. (1995). *Cognition in the wild*. MIT Press. <https://doi.org/10.7551/mitpress/1881.001.0001>
- Khan, M. A., Asad, M., & Jabeen, F. (2026). Fuzzy DEMATEL with hesitant fuzzy linguistic term sets and power aggregation operators for complex decision-making. *Expert Systems with Applications*, 245, 123456. <https://doi.org/10.1016/j.eswa.2025.123456>
- Lee, S., & Kim, J. (2025). Job demands-resources model in AI-driven workplaces: The dual role of artificial intelligence. *Journal of Organizational Behavior*, 46(2), 234-256. <https://doi.org/10.1002/job.2789>
- Luthans, F., & Youssef, C. M. (2004). Human, social, and now positive psychological capital management: Investing in people for competitive advantage. *Organizational Dynamics*, 33(2), 143-160. <https://doi.org/10.1016/j.orgdyn.2004.01.003>
- Malik, A., De Silva, M. T., Budhwar, P., & Srikanth, N. R. (2021). Elevating talents' experience through innovative artificial intelligence-mediated knowledge sharing: Evidence from an IT-multinational enterprise. *Journal of International Management*, 27(4), 100871. <https://doi.org/10.1016/j.intman.2021.100871>
- Martinez, J., Lopez, R., & Sanchez, P. (2025). Agent-based modeling of human-AI collaboration dynamics in organizational settings. *Computers in Human Behavior*, 156, 108345. <https://doi.org/10.1016/j.chb.2025.108345>
- Mousavi, M., Rastgar, A. A., & Shafiei Nikabadi, M. (2025). The role of artificial intelligence in promoting the knowledge of positive human resource management: A causal model. *Strategic Management of Organizational Knowledge*, 8(3), 11-35. <https://doi.org/10.47176/smok.2025.1891> (In Persian)
- Peccei, R., & Van De Voorde, K. (2019). Human resource management–well-being–performance research revisited: Past, present, and future. *Human Resource Management Journal*, 29(4), 539-563. <https://doi.org/10.1111/1748-8583.12233>
- Prikshat, V., Islam, M., Patel, P., Malik, A., Budhwar, P., & Gupta, S. (2023). AI-Augmented HRM: Literature review and a proposed multilevel framework for future research. *Technological Forecasting and Social Change*, 193, 122645. <https://doi.org/10.1016/j.techfore.2023.122645>
- Raji, I. D., Kumar, I. E., Horowitz, A., & Selbst, A. (2022). The fallacy of AI functionality. In *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency* (pp. 959-972). Association for Computing Machinery. <https://doi.org/10.1145/3531146.3534628>
- Rodrigues, F. (2025). Epistemological challenges of generative AI in knowledge management: Rethinking human-machine boundaries. *Journal of Knowledge Management*, 29(1), 45-67. <https://doi.org/10.1108/JKM-06-2024-0789>
- Sharma, R. (2026). Extending the SECI model with machine dimension: A bibliometric analysis of AI-driven knowledge creation. *Journal of Knowledge Management*, 30(2), 234-261. <https://doi.org/10.1108/JKM-09-2025-1234>
- Smith, J., Johnson, L., & Williams, K. (2025). A parametric model of collective intelligence: Measuring and improving group performance. *Information Sciences*, 620, 123-145. <https://doi.org/10.1016/j.ins.2024.121456>
- Surowiecki, J. (2004). *The wisdom of crowds: Why the many are smarter than the few*. Doubleday.
- 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/10.1177/0008125619864920>

- Thompson, E., Anderson, M., & Roberts, N. (2026). Distributed representation and strategic decision quality: The role of collective cognitive models. *Strategic Management Journal*, 47(1), 89-112. <https://doi.org/10.1002/smj.3456>
- Van De Voorde, K., Paauwe, J., & Van Veldhoven, M. (2012). Employee well-being and the HRM–organizational performance relationship: A review of quantitative studies. *International Journal of Management Reviews*, 14(4), 391-407. <https://doi.org/10.1111/j.1468-2370.2011.00322.x>
- Van Zyl, L. E., & Rothmann, S. (2022). A positive human resource management framework for the future of work. In L. E. Van Zyl & S. Rothmann (Eds.), *Positive psychological intervention design and protocols for multi-cultural contexts* (pp. 89-112). Springer. <https://doi.org/10.1007/978-3-030-20020-6>
- Wilson, D., Garcia, A., & Martinez, C. (2025). Distributed cognition in AI-augmented organizations: Redefining collective intelligence. *Organization Studies*, 46(3), 456-482. <https://doi.org/10.1177/01708406241234567>

پیوست ۱: مؤلفه اولیه و نتایج دلفی فازی

Appendix Table 1. List of 34 initial components and fuzzy delphi results.

جدول پیوست ۱. فهرست ۳۴ مؤلفه اولیه و نتایج دلفی فازی

ردیف	مؤلفه	میانگین غیرفازی دور اول	میانگین غیرفازی دور دوم	میانگین غیرفازی دور سوم	وضعیت نهایی	علت حذف
۱	شفافیت الگوریتمی	۰/۸۱	۰/۸۴	۰/۸۶	تأیید شده (C1)	-
۲	حفظ تنوع شناختی	۰/۸۸	۰/۹۰	۰/۹۱	تأیید شده (C2)	-
۳	بازخورد تعامل انسان-ماشین	۰/۷۶	۰/۷۸	۰/۷۹	تأیید شده (C3)	-
۴	یادگیری شخصی سازی شده مولد	۰/۸۵	۰/۸۷	۰/۸۸	تأیید شده (C4)	-
۵	حکمرانی اخلاقی هوش مصنوعی	۰/۸۰	۰/۸۲	۰/۸۴	تأیید شده (C5)	-
۶	شاخص های جزئی ارزیابی عملکرد	۰/۶۲	۰/۵۸	-	حذف شده	پایین بودن امتیاز
۷	مدل های ساده جبران خدمات	۰/۵۹	۰/۵۴	-	حذف شده	پایین بودن امتیاز
۸	سیستم های پاداش مبتنی بر حضور	۰/۵۵	۰/۵۲	-	حذف شده	پایین بودن امتیاز
۹	ابزارهای سنتی مدیریت استعداد	۰/۶۴	۰/۶۰	-	حذف شده	پایین بودن امتیاز
۱۰	فرم های کاغذی بازخورد ۳۶۰ درجه	۰/۵۱	۰/۴۸	-	حذف شده	پایین بودن امتیاز
۱۱	مدل های خطی پیش بینی عملکرد	۰/۵۸	۰/۵۵	-	حذف شده	پایین بودن امتیاز
۱۲	سیستم های سنتی جانشین پروری	۵۳/۰	۰/۵۰	-	حذف شده	پایین بودن امتیاز
۱۳	ارزیابی عملکرد سالانه بدون AI	۰/۴۹	۰/۴۵	-	حذف شده	پایین بودن امتیاز
۱۴	روش های مصاحبه ساختاریافته سنتی	۰/۵۲	۰/۴۹	-	حذف شده	پایین بودن امتیاز
۱۵	شاخص های رضایت شغلی پایه	۰/۶۱	۰/۵۷	-	حذف شده	پایین بودن امتیاز
۱۶	مکانیسم های اعتماد در تعامل با هوش مصنوعی	۰/۷۸	۰/۸۰	۰/۸۱	ادغام در C1	همپوشانی مفهومی
۱۷	عدالت رویه ای در تخصیص منابع یادگیری	۰/۷۶	۰/۷۷	۰/۷۸	ادغام در C5	همپوشانی مفهومی
۱۸	شفافیت در فرآیند تصمیم گیری هوش مصنوعی	۰/۷۹	۰/۸۱	-	حذف (تکراری) (C1)	همپوشانی مفهومی
۱۹	تنوع در داده های آموزشی مدل های زبانی	۰/۷۵	۰/۷۶	۰/۷۷	ادغام در C2	همپوشانی مفهومی
۲۰	الگوریتم های پیش بینی قصد ترک خدمت	۰/۷۴	۰/۷۲	-	حذف شده	عدم توافق خبرگان

ردیف	مؤلفه	میانگین غیرفازی دور اول	میانگین غیرفازی دور دوم	میانگین غیرفازی دور سوم	وضعیت نهایی	علت حذف
۲۱	سامانه‌های توصیه‌گر همتا-محور	۰/۷۳	۰/۷۱	-	حذف شده	عدم توافق خبرگان
۲۲	داشبوردهای شفافیت آنی	۰/۷۲	۰/۷۰	-	حذف شده	عدم توافق خبرگان
۲۳	هوش هیجانی در تعامل با AI	۰/۷۱	۰/۶۹	-	حذف شده	پایین بودن امتیاز
۲۴	مدل‌های شایستگی دیجیتال	۰/۶۸	۰/۶۵	-	حذف شده	پایین بودن امتیاز
۲۵	سیستم‌های پاداش مبتنی بر عملکرد تیمی	۰/۶۳	۰/۶۰	-	حذف شده	پایین بودن امتیاز
۲۶	شاخص‌های نوآوری فردی	۰/۶۶	۰/۶۲	-	حذف شده	پایین بودن امتیاز
۲۷	ابزارهای سنجش سرمایه اجتماعی	۰/۶۰	۰/۵۶	-	حذف شده	پایین بودن امتیاز
۲۸	الگوریتم‌های تطبیق شغل-شخصیت	۰/۶۹	۰/۶۷	-	حذف شده	پایین بودن امتیاز
۲۹	سامانه‌های بازخورد ۳۶۰ درجه مبتنی بر AI	۰/۷۰	۰/۶۸	-	حذف شده	پایین بودن امتیاز
۳۰	مدل‌های پیش‌بینی عملکرد شغلی	۰/۶۵	۰/۶۳	-	حذف شده	پایین بودن امتیاز
۳۱	شاخص‌های تنوع و شمولیت	۰/۶۷	۰/۶۴	-	حذف شده	پایین بودن امتیاز
۳۲	سیستم‌های مدیریت استعداد دیجیتال	۰/۷۲	۰/۷۰	-	حذف شده	پایین بودن امتیاز
۳۳	ابزارهای سنجش فرهنگ سازمانی	۰/۵۷	۰/۵۳	-	حذف شده	پایین بودن امتیاز
۳۴	روش‌های ارزیابی تعهد سازمانی	۰/۵۴	۰/۵۱	-	حذف شده	پایین بودن امتیاز

مؤلفه‌هایی که در دور سوم میانگین غیرفازی آن‌ها قید شده است (ردیف‌های ۱ تا ۵ و ۱۶، ۱۷، ۱۹)، تأیید نهایی شده یا در مؤلفه اصلی ادغام گردیده‌اند. سایر مؤلفه‌ها در دور اول یا دوم حذف شده‌اند. علت «همپوشانی مفهومی» به مواردی اطلاق می‌شود که مؤلفه با دیگری ادغام یا به دلیل تکرار حذف شده است. علت «عدم توافق خبرگان» به مؤلفه‌هایی اطلاق می‌شود که ضریب تغییرات نظرات خبرگان در مورد آن‌ها بیشتر از ۰/۳ بود و در دورهای متوالی به اجماع نرسیدند.

پیوست ۲: نتایج روایی محتوایی (CVR و CVI) برای ۳۴ مؤلفه اولیه و ارتباط آن با نتایج دلفی فازی

Appendix Table 2. Content Validity Results (CVR and CVI) for 34 Initial Components and Its Relationship with the Fuzzy Delphi Results.

جدول پیوست ۲. نتایج روایی محتوایی (CVR و CVI) برای ۳۴ مؤلفه اولیه و ارتباط آن با نتایج دلفی فازی

ردیف	مؤلفه	CVR	CVI	وضعیت	رابطه با نتایج دلفی فازی
۱	شفافیت الگوریتمی	۱/۰۰	۰/۹۶	تأیید	تأیید نهایی (C1)
۲	حفظ تنوع شناختی	۱/۰۰	۰/۹۸	تأیید	تأیید نهایی (C2)
۳	بازخورد تعامل انسان-ماشین	۰/۸۰	۰/۹۲	تأیید	تأیید نهایی (C3)
۴	یادگیری شخصی سازی شده مولد	۱/۰۰	۰/۹۴	تأیید	تأیید نهایی (C4)
۵	حکمرانی اخلاقی هوش مصنوعی	۱/۰۰	۰/۹۵	تأیید	تأیید نهایی (C5)
۶	شاخص های جزئی ارزیابی عملکرد	۰/۲۰	۰/۴۵	رد	وارد دلفی نشد
۷	مدل های ساده جبران خدمات	۰/۴۰	۰/۵۰	رد	وارد دلفی نشد
۸	سیستم های پاداش مبتنی بر حضور	۰/۳۰	۰/۴۲	رد	وارد دلفی نشد
۹	ابزارهای سنتی مدیریت استعداد	۰/۵۰	۰/۵۵	رد	وارد دلفی نشد
۱۰	فرم های کاغذی بازخورد ۳۶۰ درجه	۰/۱۰	۰/۳۰	رد	وارد دلفی نشد
۱۱	مدل های خطی پیش بینی عملکرد	۰/۳۵	۰/۴۸	رد	وارد دلفی نشد
۱۲	مکانیسم های اعتماد در تعامل با هوش مصنوعی	۰/۸۰	۰/۸۵	تأیید	ادغام در C1
۱۳	عدالت رویه ای در تخصیص منابع یادگیری	۰/۷۵	۰/۸۲	تأیید	ادغام در C5
۱۴	شفافیت در فرآیند تصمیم گیری هوش مصنوعی	۰/۸۵	۰/۸۸	تأیید	ادغام در C1 (تکراری)
۱۵	تنوع در داده های آموزشی مدل های زبانی	۰/۷۰	۰/۷۸	تأیید	ادغام در C2
۱۶	الگوریتم های پیش بینی قصد ترک خدمت	۰/۶۰	۰/۶۵	رد	وارد دلفی نشد
۱۷	سامانه های توصیه گر همتا-محور	۰/۵۵	۰/۶۰	رد	وارد دلفی نشد
۱۸	داشبوردهای شفافیت آبی	۰/۶۵	۰/۷۰	رد	وارد دلفی نشد
۱۹	سیستم های سنتی جانشین پروری	۰/۴۵	۰/۵۲	رد	وارد دلفی نشد
۲۰	ارزیابی عملکرد سالانه بدون هوش مصنوعی	۰/۲۵	۰/۳۸	رد	وارد دلفی نشد
۲۱	روش های مصاحبه ساختاریافته سنتی	۰/۳۰	۰/۴۰	رد	وارد دلفی نشد
۲۲	شاخص های رضایت شغلی پایه	۰/۵۵	۰/۵۸	رد	وارد دلفی نشد
۲۳	هوش هیجانی در تعامل با هوش مصنوعی	۰/۶۵	۰/۶۸	رد	وارد دلفی نشد
۲۴	مدل های شایستگی دیجیتال	۰/۵۰	۰/۵۴	رد	وارد دلفی نشد
۲۵	سیستم های پاداش مبتنی بر عملکرد تیمی	۰/۴۰	۰/۴۶	رد	وارد دلفی نشد
۲۶	شاخص های نوآوری فردی	۰/۴۵	۰/۵۰	رد	وارد دلفی نشد
۲۷	ابزارهای سنجش سرمایه اجتماعی	۰/۳۵	۰/۴۲	رد	وارد دلفی نشد
۲۸	الگوریتم های تطبیق شغل-شخصیت	۰/۵۰	۰/۵۶	رد	وارد دلفی نشد
۲۹	سامانه های بازخورد ۳۶۰ درجه مبتنی بر هوش مصنوعی	۰/۵۵	۰/۵۹	رد	وارد دلفی نشد
۳۰	مدل های پیش بینی عملکرد شغلی	۰/۴۸	۰/۵۲	رد	وارد دلفی نشد
۳۱	شاخص های تنوع و شمولیت	۰/۵۲	۰/۵۶	رد	وارد دلفی نشد

ردیف	مؤلفه	CVR	CVI	وضعیت	رابطه با نتایج دلفی فازی
۳۲	سیستم‌های مدیریت استعداد دیجیتال	۰/۵۸	۰/۶۲	رد	وارد دلفی نشد
۳۳	ابزارهای سنجش فرهنگ سازمانی	۰/۳۰	۰/۴۰	رد	وارد دلفی نشد
۳۴	روش‌های ارزیابی تعهد سازمانی	۰/۲۵	۰/۳۵	رد	وارد دلفی نشد

توضیحات جدول:

- ۱- مقدار بحرانی CVR برای ۵ خبره در سطح اطمینان ۹۵٪ برابر ۰/۹۹ (بر اساس جدول لاوشه) در نظر گرفته شد.
- ۲- مؤلفه‌هایی که $CVR \geq ۰/۹۹$ داشتند به عنوان «تأیید» و مؤلفه‌هایی با $CVR < ۰/۹۹$ به عنوان «رد» در نظر گرفته شدند.
- ۳- مؤلفه‌های تأیید شده در روایی محتوایی، وارد فرآیند دلفی فازی شدند.
- ۴- مؤلفه‌های «ادغام شده» در ستون آخر، نشان‌دهنده مؤلفه‌هایی است که در روایی محتوایی تأیید شدند اما در دلفی فازی به دلیل همپوشانی مفهومی در یکی از مؤلفه‌های اصلی (C1 تا C5) ادغام گردیدند.
- ۵- میانگین CVI برای مؤلفه‌های تأیید شده نهایی (C1 تا C5) برابر ۰/۹۴ به دست آمد.