



kharazmi University



Assessing Artificial Intelligence Literacy Among Students at Kharazmi University Based on Meta AI Literacy Scale

Atena Latifi¹ , Mohammad Zerehsaz² , Ali Azimi^{3✉}

1. MSc in Knowledge and Information Science, Kharazmi University, Tehran, Iran. Email: latifiatena743@gmail.com
 2. Department of Knowledge and Information Science, Kharazmi University, Tehran, Iran. Email: zerehsaz@khu.ac.ir
 3. Corresponding author, Department of Knowledge and Information Science, Kharazmi University, Tehran, Iran. Email: azimia@khu.ac.ir

Article Info	ABSTRACT
Article type: Research Article Article history: Received 20 September 2025 Revised in 30 November 2025 Accepted 7 December 2025 Published online 26 December 2025 Keywords: Artificial intelligence literacy, Data literacy, Meta AI Literacy Scale, Students, Privacy, Kharazmi University	Purpose: This study aimed to assess artificial intelligence (AI) literacy among students at Kharazmi University and to examine differences across components, dimensions, educational levels, and fields of study. The study conceptualized AI literacy as a multidimensional competency encompassing not only knowledge and tool use but also self-efficacy, self-management, ethical awareness, and the ability to engage critically with AI systems. Method: The study employed an applied descriptive survey design. The final sample consisted of 335 undergraduates, master's, and doctoral students selected through proportional stratified sampling with accessibility considerations. Data were collected using the Meta AI Literacy Scale (MAILS). Because the Persian version had been validated in previous Iranian studies, only qualitative face validity was examined in the present study through feedback on item clarity and comprehensibility. Reliability was assessed using Cronbach's alpha. Data were analyzed in SPSS 27 using descriptive statistics, one-sample t-test, Friedman test, one-way ANOVA, Tukey post hoc test, and Kruskal–Wallis test. Findings: The mean overall AI literacy score was 5.83 out of 10 and was significantly above the theoretical midpoint ($t = 11.02$, $p < .001$). Seventy-three percent of students were classified at a desirable level. Use and application had the highest mean, whereas creation had the lowest. Significant differences were found among components and dimensions. No significant differences were observed across educational levels; however, AI literacy differed significantly across fields of study, and computer science and engineering students reported the highest scores. Conclusion: Although students' perceived AI literacy was generally desirable, the gap between using ready-made tools and creating, evaluating, and solving problems with AI indicates a need for systematic, discipline-sensitive, and ethically grounded education. Future research should combine self-report measures with performance-based, longitudinal, and experimental assessments.

Cite this article: Latifi, A., Zerehsaz, M., & Azimi, A. (2026). Assessing Artificial Intelligence Literacy Among Students at Kharazmi University Based on the Meta AI Literacy Scale. *Human-Information Interaction*, 12(4), 1-21.





Kharazmi University



Introduction

Artificial intelligence is no longer a peripheral technological issue in higher education. It increasingly shapes how students search for information, write, learn, evaluate evidence, interact with digital systems, and imagine their future professional roles. The historical analogy with the introduction of Student's t-test in 1908 is useful here. When Gosset, under the pen name Student, introduced a method for drawing reliable conclusions from small samples, statistical reasoning moved from a purely specialist activity toward a more widely usable analytical practice. The spread of AI in universities has produced a comparable epistemic pressure: the tool is powerful, but its educational value depends on users' literacy. Students who use AI without understanding its limits, biases, probabilistic nature, and ethical implications may become faster learners but not necessarily more competent. Therefore, assessing AI literacy is not simply a measurement exercise; it is a way of examining whether universities are preparing students for a sociotechnical environment in which AI mediates knowledge, judgment, and action.

Purpose

The present study assessed the level of AI literacy among students at Kharazmi University. It examined whether students' AI literacy was at a desirable level, whether meaningful differences existed among the main components and dimensions of AI literacy, and whether AI literacy varied according to educational level and field of study. The study used the Meta AI Literacy Scale (MAILS), because this instrument conceptualizes AI literacy as a multidimensional construct rather than as mere tool familiarity. In this study, AI literacy included the ability to use and apply AI, understand AI concepts, detect AI-generated outputs, attend to AI ethics, and create AI-related outputs. It also included AI self-efficacy in learning and problem-solving, as well as AI self-management, including persuasion literacy and emotion regulation in interactions with AI systems.

Theoretical framing

Feenberg's critical theory of technology served as the main interpretive frame, treating AI as a socio-technical arrangement whose design and use embody values, institutional priorities, and relations of power. Stiegler's pharmakon perspective clarified AI's dual potential to support learning while also displacing memory, judgment, and responsibility when reliance becomes disproportionate. Postphenomenology was used to interpret technological mediation of problem framing and judgment; actor-network theory to interpret disciplinary differences as products of curricula, infrastructures, actors, and tools; Haraway's cyborg approach to understand the hybrid constitution of the learner; and Han's critique to illuminate cognitive pressure, transparency, and self-monitoring. The literary example of Orwell's *Nineteen Eighty-Four* was used as a heuristic for data surveillance and privacy. These perspectives were linked to MAILS by distinguishing operational literacy from critical, ethical, and self-regulatory competence.

Measurement model

The study's empirical framework was the Meta AI Literacy Scale (MAILS). MAILS was particularly appropriate because it operationalizes AI literacy beyond simple awareness or tool use. It includes an AI literacy component covering the use and application, understanding, detection, ethics, and creation of AI; an AI self-efficacy component covering problem-solving



kharazmi University



and learning; and an AI self-management component covering persuasion literacy and emotion regulation. Thus, MAILS connects knowledge and skill with psychological, ethical, and self-regulatory capacities. In the present study, this structure allowed the analysis to distinguish between students who can use AI tools and students who can critically understand, evaluate, create with, and manage themselves in AI-mediated environments.

Methods and Materials

The study adopted an applied descriptive survey design. The statistical population included students at Kharazmi University across undergraduate, master's, and doctoral levels. A proportional stratified sampling strategy was used, with accessibility considerations, and 335 valid responses were analyzed. The standardized Meta AI Literacy Scale (MAILS), administered on a 0-to-10 response scale, measured three components and nine dimensions. Because the Persian version had been validated in previous Iranian studies, content validity was not reassessed in the present study; only qualitative face validity was examined through feedback from 20 professors and students and the research supervisors on clarity, readability, and comprehensibility. Reliability was assessed using Cronbach's alpha, with a total-scale coefficient of 0.963. Data were analyzed in SPSS 27 using descriptive statistics, one-sample t-test, Friedman test, one-way ANOVA, Tukey post hoc test, and Kruskal–Wallis test. As a cross-sectional survey, the study identifies statistical differences and cannot establish causal effects.

Results: The findings showed that students' overall AI literacy mean was 5.83 out of 10, with a standard deviation of 1.88. The one-sample t-test indicated that this mean was significantly above the theoretical midpoint, $t(334) = 11.02$, $p < .001$. The effect size was moderate to relatively strong, with Cohen's $d = .602$ and Hedges' $g = .601$. Based on the adequacy classification, 243 students, or approximately 73 percent of the sample, were placed in the desirable level of AI literacy, while 92 students, or 27 percent, remained below the desirable threshold. This indicates a generally favorable but uneven AI literacy profile among Kharazmi University students.

Results also revealed important internal differences among AI literacy dimensions. The highest mean score was observed for use and application of AI, suggesting that students are relatively comfortable with using existing AI tools for learning, searching, writing, and everyday tasks. In contrast, the lowest mean score was observed for creating AI tools or AI-based outputs. The Friedman test confirmed significant differences among the main components and dimensions. This pattern is theoretically important: students appear stronger as AI users than as critical designers, creators, or problem solvers in AI-mediated environments. AI self-management also showed relatively high scores, particularly in persuasion literacy and emotion regulation, but AI self-efficacy, especially problem solving, was weaker than basic use.

Comparisons across educational levels showed no statistically significant difference in overall AI literacy among undergraduate, master's, and doctoral students. This suggests that progression to a higher formal academic level does not automatically lead to higher AI literacy. In contrast, overall AI literacy scores differed significantly across fields of study. Students in computer science and engineering reported the highest mean scores, followed by students in basic sciences and technical and engineering fields. In contrast, students in humanities and social and behavioral sciences reported lower means. This pattern may be



Kharazmi University

Journal of Human-Information Interaction

Online ISSN: 2423-7418

<https://hiikhu.ac.ir/>



associated with differences in curricular exposure to computational thinking, programming, data practices, and technical tools rather than educational level alone.

Discussion and Conclusion

The survey findings indicate that students' perceived AI literacy was relatively desirable but uneven: use and application were strongest, creation and problem-solving were weaker, educational levels did not differ significantly, and scores differed across fields of study. Feenberg's framework helps interpret the use-creation gap as a distinction between operating within a ready-made technical system and critically participating in its design; actor-network theory helps interpret field differences as associations with disciplinary curricula, infrastructures, and opportunities rather than as causal effects of field of study. Stiegler, Haraway, and Han further clarify why competent use should include preserving judgment, authorship, self-regulation, and responsibility. The present survey did not measure learning loss, agency decay, or the effects of AI use.

Implications

The findings have direct implications for curriculum design. AI literacy education should be embedded across disciplines rather than restricted to computer science programs. A staged model is recommended: conceptual and ethical foundations for all students; discipline-specific applications; advanced practice in prompt design, data reasoning, and output evaluation; and agency-preserving routines such as draft-first/AI-second work, comparison with independent sources, explicit uncertainty checks, and periodic AI-free transfer tasks. Students should learn not only how to use AI but also when not to use it, how to exit an AI-mediated workflow, and how to preserve plural sources, authorship, and responsibility. Universities should also teach privacy-aware data practices, regulatory awareness, statistical reasoning, and data citizenship so that students can protect confidential information and participate in institutional AI governance.

Ethical Considerations

Compliance with Research Ethics

Participation was voluntary and informed. Data confidentiality was maintained, and no identifying information about participants was disclosed in the research report.

Authors' Contributions

First Author: Conducting the research, data collection, quantitative analysis, and preparing the initial thesis-based draft. Second Author: Contributing to methodological consultation, scientific advice, and manuscript review. Third Author: Scientific supervision, theoretical and methodological guidance, final analysis and review, and correspondence as the corresponding author.

Conflict of Interest

The authors declare no conflict of interest.

Acknowledgements

This article is derived from the master's thesis of Atena Latifi in Knowledge and Information Science at Kharazmi University, conducted under the supervision of Dr. Ali Azimi and the advisory support of Dr. Mohammad Zerehsaz.

ارزیابی سواد هوش مصنوعی دانشجویان دانشگاه خوارزمی

بر اساس مقیاس سواد هوش مصنوعی متا

آتنا لطیفی^۱، محمد زره‌ساز^۲، علی عظیمی^۳

۱. کارشناسی ارشد علم اطلاعات و دانش‌شناسی، دانشگاه خوارزمی، تهران، ایران. رایانامه: latifiatena743@gmail.com

۲. عضو هیأت علمی گروه علم اطلاعات و دانش‌شناسی، دانشگاه خوارزمی، تهران، ایران. رایانامه: zerehsaz@khu.ac.ir

۳. نویسنده مسئول، گروه علم اطلاعات و دانش‌شناسی، دانشگاه خوارزمی، تهران، ایران. رایانامه: azimia@khu.ac.ir

اطلاعات مقاله	چکیده
نوع مقاله: مقاله پژوهشی	هدف: پژوهش حاضر با هدف ارزیابی سواد هوش مصنوعی دانشجویان دانشگاه خوارزمی و تحلیل تفاوت آن بر حسب مؤلفه‌ها، ابعاد، مقطع تحصیلی و گروه رشته تحصیلی انجام شد. تمرکز پژوهش بر این بود که دانشجویان تا چه اندازه از دانش، مهارت، خودکارآمدی و خودمدیریتی لازم برای مواجهه آگاهانه، انتقادی و مسئولانه با ابزارهای هوش مصنوعی برخوردارند.
تاریخ دریافت: ۱۴۰۴/۰۶/۲۹	روش: پژوهش از نظر هدف کاربردی و از نظر شیوه گردآوری داده‌ها، توصیفی - پیمایشی بود. جامعه آماری شامل دانشجویان دانشگاه خوارزمی در مقاطع کارشناسی، کارشناسی ارشد و دکتری بود. نمونه‌هایی پژوهش ۳۳۵ نفر بود که با نمونه‌گیری طبقه‌ای نسبی و با ملاحظه اصل دسترسی انتخاب شدند. داده‌ها با استفاده از مقیاس استاندارد سواد هوش مصنوعی متا گردآوری و با نرم‌افزار SPSS نسخه ۲۷ تحلیل شد. با توجه به اعتبارسنجی نسخه فارسی مقیاس در پژوهش‌های پیشین، در پژوهش حاضر صرفاً روایی صوری کیفی آن از نظر وضوح و فهم‌پذیری گویه‌ها بررسی شد و پایایی با آلفای کرونباخ سنجیده شد.
تاریخ بازنگری: ۱۴۰۴/۰۹/۰۹	یافته‌ها: میانگین سواد کلی هوش مصنوعی دانشجویان ۵/۸۳ از ۱۰ بود و آزمون t تک‌نمونه‌ای نشان داد که این میانگین به‌طور معناداری بالاتر از حد متوسط نظری است ($t = 11.02, p < .001$). ۷۳ درصد دانشجویان در سطح مطلوب قرار گرفتند. بیشترین میانگین مربوط به بُعد استفاده و کاربرد هوش مصنوعی و کمترین میانگین مربوط به بُعد ایجاد ابزارهای جدید بود. آزمون فریدمن نشان داد میان مؤلفه‌ها و ابعاد مختلف تفاوت معنادار وجود دارد. تفاوت معناداری میان مقاطع تحصیلی مشاهده نشد، اما تفاوت بر حسب رشته تحصیلی معنادار بود و دانشجویان علوم و مهندسی کامپیوتر بالاترین میانگین را کسب کردند.
کلیدواژه‌ها: سواد هوش مصنوعی، مقیاس سواد هوش مصنوعی، متا، دانشجویان، عاملیت انسانی، حریم خصوصی، دانشگاه خوارزمی	نتیجه‌گیری: اگرچه کلیت سواد ادراک‌شده هوش مصنوعی دانشجویان مطلوب ارزیابی شد، با این حال شکاف میان استفاده از ابزارهای آماده و توانایی خلق، ارزیابی انتقادی و حل مسئله با هوش مصنوعی نشان می‌دهد که آموزش دانشگاهی نباید به مهارت مصرف ابزار محدود شود. طراحی برنامه‌های آموزشی میان‌رشته‌ای، تقویت اخلاق، راستی‌آزمایی و ارزیابی انتقادی، و توجه به تفاوت‌های رشته‌ای برای توسعه سواد هوش مصنوعی مسئولانه ضروری است. سنجش‌های آینده باید انتقال یادگیری پس از حذف حمایت هوش مصنوعی، توان راستی‌آزمایی، مالکیت شناختی بر محصول، توان توضیح مسیر استدلال و قابلیت امتناع از استفاده نامناسب را ارزیابی کنند. همچنین، سواد داده و حریم خصوصی باید بخشی صریح از آموزش دانشگاهی باشد تا دانشجویان بتوانند منشأ، گردآوری، اشتراک‌گذاری و بهره‌برداری از داده‌های شخصی و پژوهشی را ارزیابی و از محرمانگی اطلاعات حفاظت کنند.

استناد: لطیفی، آتنا؛ زره‌ساز، محمد؛ و عظیمی‌وقار، علی (۱۴۰۵). ارزیابی سواد هوش مصنوعی دانشجویان دانشگاه خوارزمی بر اساس مقیاس سواد هوش مصنوعی متا. *تعامل انسان و اطلاعات*، ۱۲ (۴)، ۲۱-۱.

مقدمه

در تاریخ علم، برخی ابزارها فقط یک فناوری یا روش تازه نیستند، بلکه شیوه اندیشیدن، تصمیم‌گیری و تولید دانش را تغییر می‌دهند. آزمون t استیودنت نمونه‌ای روشن از این نوع تحول است. ویلیام سیلی گاوس^۱ که به دلیل محدودیت‌های انتشار در کارخانه گینس با نام مستعار «Student» مقاله خود را منتشر کرد، در سال ۱۹۰۸ با مقاله «خطای محتمل میانگین» امکان استنباط آماری معتبر از نمونه‌های کوچک را فراهم ساخت (گاوس، ۱۹۰۸). اهمیت این رخداد در خود فرمول آماری خلاصه نمی‌شد؛ بلکه در این بود که پژوهشگر می‌توانست با ابزاری نو، میان داده محدود و تصمیم علمی پیوندی معتبر برقرار کند. همان‌گونه که آزمون t در آغاز قرن بیستم افق تازه‌ای برای تحلیل آماری گشود، هوش مصنوعی در آغاز قرن بیست و یکم نیز افق تازه‌ای برای یادگیری، پژوهش و تصمیم‌گیری دانشگاهی گشوده است؛ با این تفاوت که این بار مسئله فقط «داشتن ابزار» نیست، بلکه «سواد استفاده از ابزار» است.

هوش مصنوعی در آموزش عالی امروز همان واکنش دوگانه‌ای را برانگیخته که بسیاری از فناوری‌های دانشی در آغاز ظهور خود ایجاد کرده‌اند: از یک سو امید به افزایش سرعت، دقت، دسترسی و خلاقیت؛ و از سوی دیگر نگرانی درباره اتکای بیش از حد، سوگیری، تضعیف تفکر انتقادی، ابهام در مسئولیت اخلاقی و نابرابری در دسترسی. گزارش‌های جدید نشان می‌دهد استفاده دانشجویان از ابزارهای هوش مصنوعی مولد در آموزش عالی با سرعتی چشمگیر افزایش یافته است؛ برای نمونه، پیمایش دانشجویی HEPI در سال ۲۰۲۵ از افزایش گسترده استفاده دانشجویان از هوش مصنوعی در تکالیف، توضیح مفاهیم، خلاصه‌سازی و آماده‌سازی فعالیت‌های یادگیری خبر می‌دهد (فریمن، ۲۰۲۵). در چنین وضعیتی، مسئله اصلی دانشگاه‌ها دیگر فقط ممنوعیت یا مجاز بودن استفاده از هوش مصنوعی نیست؛ بلکه توانمند کردن دانشجویان برای استفاده آگاهانه، انتقادی و مسئولانه از آن است.

مفهوم سواد نیز در همین زمینه باید بازخوانی شود. سواد در معنای کلاسیک خود به توانایی خواندن و نوشتن اشاره داشت، اما در جهان معاصر به حوزه‌هایی مانند سواد دیجیتال، سواد داده، سواد الگوریتمی، سواد رسانه‌ای و اکنون سواد هوش مصنوعی گسترش یافته است. سواد هوش مصنوعی به توانایی فهم اصول، کاربردها، محدودیت‌ها و پیامدهای هوش مصنوعی و نیز توانایی استفاده، ارزیابی و تعامل مسئولانه با سامانه‌های هوشمند اشاره دارد (لانگ و ماگرکو، ۲۰۲۰؛ نگ و همکاران، ۲۰۲۱). این تعریف نشان می‌دهد که دانشجو فقط کاربر ابزار نیست؛ بلکه باید بتواند درباره خروجی هوش مصنوعی پرسش کند، خطاها و سوگیری‌ها را تشخیص دهد، پیامدهای اخلاقی را بسنجد و در فرایند یادگیری و پژوهش، قضاوت انسانی خود را حفظ کند.

تحول تاریخی سواد و پیدایش سواد داده

شمارزو^۲ (۲۰۲۳) سواد را فقط توان خواندن و نوشتن نمی‌داند، بلکه آن را توانایی، اعتماد و تمایل فرد برای درگیر شدن با زبان به‌منظور کسب، ساخت و ارتباط معنا تعریف می‌کند؛ تعریفی که تفکر انتقادی و فهم اطلاعات پیچیده در بافت‌های گوناگون را نیز دربر می‌گیرد. از این منظر، تاریخ سواد را می‌توان زنجیره‌ای از نقاط عطف دانست: اختراع چاپ و دسترسی انبوه به کتاب، گسترش کتابخانه‌های عمومی، آموزش اجباری، ظهور رادیو و تلویزیون و سپس رایانه و اینترنت. هر مرحله، هم دامنه دسترسی به دانش را افزایش داده و هم شکل مشارکت اقتصادی، سلامت، تحرک اجتماعی و شهروندی را دگرگون کرده است. در امتداد این تاریخ، سواد داده و سواد هوش مصنوعی مرحله‌ای تازه از تحول سوادند؛ زیرا موضوع دیگر صرفاً خواندن متن نیست، بلکه فهم چگونگی تولید داده، تبدیل آن به شاخص و مدل، و اثرگذاری ارزیابی‌های تحلیلی بر تصمیم و کنش انسان است. بنابراین، سواد داده حلقه واسط میان دسترسی به اطلاعات و توان داوری درباره الگوریتمی است که آن اطلاعات را انتخاب، رتبه‌بندی یا تفسیر می‌کند.

گسترش آموزش حوزه‌های علوم، فناوری، مهندسی و ریاضیات که اصطلاحاً STEM^۳ نامیده می‌شود پاسخی مهم به رشد هوش مصنوعی بوده است، اما به‌تنهایی برای شکل‌گیری سواد جامع کافی نیست. آموزش فنی می‌تواند فهم الگوریتم و

¹ William Sealy Gosset

² Schmarzo

³ Science, Technology, Engineering, and Mathematics

برنامه‌نویسی را افزایش دهد، ولی لزوماً به آگاهی از حریم خصوصی، پیامدهای اجتماعی، سوگیری، مسئولیت اخلاقی یا توان مخالفت با سامانه نمی‌انجامد. تأکید شمارزو (۲۰۲۳) بر آموزش همه گروه‌های سنی و رشته‌ای و تبدیل افراد به «شهروندان علم داده» نشان می‌دهد سواد هوش مصنوعی باید از انحصار متخصصان خارج و به شایستگی عمومی برای مشارکت در تصمیم‌های فناورانه تبدیل شود. با ظهور هوش مصنوعی مولد، خود مفهوم سواد نیز با تنشی تازه روبه‌رو شده است. در تلقی عملیاتی، فرد باسواد کسی است که بتواند ابزار مناسب را انتخاب کند، دستور مؤثر بنویسد و خروجی قابل استفاده بگیرد؛ اما در تلقی عاملیت‌محور، سواد زمانی تحقق می‌یابد که فرد همچنان مؤلف هدف، معیار، داوری و مسئولیت نهایی باقی بماند. از این‌رو، توان تولید یک متن یا پاسخ با حضور هوش مصنوعی لزوماً به معنای تملک دانش، یادگیری پایدار یا توان بازتولید مستقل آن نیست. چارچوب‌های جدید یونسکو و سازمان همکاری و توسعه اقتصادی (OECD) نیز بر نگرش انسان‌محور، قضاوت انتقادی، خلق مسئولانه و تصمیم آگاهانه درباره استفاده یا عدم استفاده از هوش مصنوعی تأکید دارند و تصریح می‌کنند که صرف تعامل با ابزار، به خودی خود شایستگی سواد هوش مصنوعی را ایجاد نمی‌کند (میائو و همکاران، ۲۰۲۴؛ سازمان همکاری و توسعه اقتصادی و کمیسیون اروپا، ۲۰۲۶).

در مقاله حاضر، نظریه انتقادی فناوری فینبرگ^۴ چارچوب اصلی تفسیر است. بر اساس این دیدگاه، فناوری پدیده‌ای خنثی و بیرون از جامعه نیست، بلکه در طراحی و کاربرد آن ارزش‌ها، منافع نهادی و روابط قدرت تثبیت می‌شود (فینبرگ، ۱۹۹۱). از این‌رو، هوش مصنوعی در دانشگاه فقط ابزار تسریع نوشتن، جست‌وجو یا تحلیل داده نیست؛ بلکه سامانه‌ای اجتماعی - فنی است که می‌تواند معیارهای شایستگی، شیوه داوری و رابطه دانشجو با دانش را بازآرایی کند. دیدگاه «فارماکون» استیگلر این دوگانگی را تکمیل می‌کند: فناوری هم می‌تواند داربست یادگیری و گسترش‌دهنده توانایی باشد و هم، در صورت اتکای نامتناسب، بخشی از حافظه، داوری و مسئولیت را از یادگیرنده به سامانه منتقل کند (استیگلر، ۲۰۱۳). بنابراین، سواد هوش مصنوعی باید هم توان بهره‌گیری و هم توان تشخیص حدود، خطرها و زمان امتناع از کاربرد فناوری را دربر گیرد. این چارچوب اصلی با چند رویکرد مکمل روشن‌تر می‌شود. پس‌پدیدارشناسی وربیک^۵ بر میانجی‌گری فناوری در دیدن مسئله، صورت‌بندی پرسش و ارزیابی پاسخ تأکید دارد (وربیک، ۲۰۱۱). نظریه کنشگر - شبکه لاتور نشان می‌دهد کاربست هوش مصنوعی حاصل شبکه‌ای از دانشجو، استاد، ابزار، داده، الگوریتم، برنامه درسی و زیرساخت است، نه صرفاً ویژگی فردی کاربر (لاتور، ۲۰۰۵). رویکرد سایبورگ هاراوی نیز مرز ساده انسان و ماشین را به چالش می‌کشد و هویت یادگیرنده را در پیوند با ابزارها و شبکه‌ها می‌فهمد (هاراوی، ۱۹۸۵). در مقابل، نقد هان^۶ از جامعه شفافیت و فرسودگی یادآور می‌شود که دسترسی دائمی، سرعت و کارآمدی ظاهری می‌تواند با فشار شناختی، خودنظارتی و خستگی همراه شود (هان، ۲۰۱۵). این رویکردها نظریه‌های رقیب نیستند، بلکه ابعاد مکمل یک پرسش واحد را توضیح می‌دهند: دانشجو چگونه می‌تواند در تعامل با هوش مصنوعی، فهم، خودمدیریتی و مسئولیت خود را حفظ کند؟

بر پایه این پیوند نظری، تمایز میان «بهبود عملکرد در حضور ابزار» و «رشد قابلیت درونی یادگیرنده» اهمیت می‌یابد. شکاف عملکرد - قابلیت زمانی مطرح است که کیفیت یا سرعت خروجی با کمک هوش مصنوعی افزایش یابد، اما فرد پس از حذف حمایت نتواند مسئله را مستقل صورت‌بندی، راه‌حل را بازسازی یا از آن دفاع کند. فرسایش عاملیت نیز به انتقال تدریجی تعیین هدف، انتخاب معیار، تولید پاسخ و راستی‌آزمایی از انسان به سامانه اشاره دارد.

ادبیات جدید نیز همین نگاه چندبعدی را تأیید می‌کند. لانگ و ماگرکو (۲۰۲۰) در طرح شایستگی‌های سواد هوش مصنوعی نشان دادند که این سواد شامل آشنایی مفهومی، توان استفاده، درک محدودیت‌ها، تشخیص کاربردها و فهم پیامدهای اجتماعی است. نگ و همکاران (۲۰۲۱) نیز در مرور اکتشافی خود چهار قلمرو دانستن و فهمیدن، استفاده و کاربرد، ارزیابی و خلق، و مسائل اخلاقی را برای سواد هوش مصنوعی برجسته کردند. مرور نظام‌مند المطرفی و همکاران (۲۰۲۴) نشان داد مطالعات سال‌های ۲۰۱۹ تا ۲۰۲۳ سواد هوش مصنوعی را به‌عنوان سازه‌ای چندبعدی و نیازمند آموزش هدفمند در سطوح

⁴ Feenberg

⁵ Verbeek

⁶ Latour

⁷ Han

مختلف آموزشی مفهوم‌سازی کرده‌اند. همچنین پینسکی و بنلیان (۲۰۲۴) در مرور جامع خود تأکید کردند که سواد هوش مصنوعی برای کاربران باید هم شامل اجزای دانشی و عملی و هم پیامدهای رفتاری و شناختی باشد.

شواهد تجربی جدید نشان می‌دهد رابطه هوش مصنوعی و یادگیری شرطی و وابسته به طراحی تعامل است. باستانی و همکاران (۲۰۲۵) در یک آزمایش میدانی در آموزش ریاضی نشان دادند دسترسی بدون حفظ آموزشی می‌تواند عملکرد حین تمرین را افزایش دهد، اما عملکرد مستقل بعدی را تضعیف کند؛ در مقابل، کستین و همکاران (۲۰۲۵) گزارش کردند یک معلم‌یار هوش مصنوعی که بر اصول یادگیری فعال، داربست‌بندی، بازخورد هدفمند و خودآهنگی بنا شده بود، یادگیری دانشجویان را نسبت به کلاس فعال افزایش داد. بنابراین، پرسش علمی مناسب «هوش مصنوعی مفید است یا مضر؟» نیست، بلکه این است که چه نوع طراحی، در چه مرحله‌ای و برای کدام هدف، یادگیری را تقویت یا جایگزین می‌کند.

این تمایز در پژوهش‌های مربوط به عاملیت و تفکر انتقادی نیز مشاهده شده است. لی و همکاران (۲۰۲۵) در مطالعه ۳۱۹ کارآموز دریافتند اعتماد بیشتر به هوش مصنوعی با تلاش کمتر برای تفکر انتقادی همراه بود و نقش تفکر از تولید راه‌حل به راستی‌آزمایی، ادغام پاسخ و نظارت بر کار منتقل می‌شد. فراتحلیل واکارو و همکاران (۲۰۲۴) نیز نشان داد ترکیب انسان و هوش مصنوعی به‌طور متوسط الزاماً از بهترین انسان یا هوش مصنوعی بهتر نیست و نوع تکلیف و تقسیم کار تعیین‌کننده است. در تازه‌ترین مطالعه آزمایشی، لی و همکاران (۲۰۲۶) نشان دادند استفاده منفعلانه و کپی‌محور از هوش مصنوعی، خودکارآمدی، مالکیت روان‌شناختی و معنای کار را کاهش می‌دهد، در حالی که الگوی همکاری فعال—نوشتن پیش‌نویس توسط فرد و سپس اصلاح با کمک هوش مصنوعی—این پیامدها را تا حد زیادی مهار می‌کند. این یافته‌ها ضرورت گذار از «سواد استفاده» به «سواد اتکای متناسب و عاملیت‌محور» را برجسته می‌سازند.

یکی از تلاش‌های مهم برای سنجش تجربی این سازه، مقیاس سواد هوش مصنوعی متا یا MAILS^۸ است. کارولوس و همکاران (۲۰۲۳) این مقیاس را با تکیه بر مدل‌های شایستگی سواد هوش مصنوعی و شایستگی‌های روان‌شناختی طراحی کردند تا سنجش سواد هوش مصنوعی فقط به دانستن چند مفهوم فنی یا گزارش میزان استفاده از ابزار محدود نشود. اهمیت MAILS در آن است که سواد هوش مصنوعی را به‌صورت سازه‌ای چندلایه می‌فهمد: لایه نخست، توان شناختی و عملی کار با هوش مصنوعی است؛ لایه دوم، باور فرد به توانایی یادگیری و حل مسئله در مواجهه با هوش مصنوعی است؛ و لایه سوم، توان خودمدیریتی در برابر پیامدهای اقناعی، هیجانی و شناختی فناوری است (کارولوس و همکاران، ۲۰۲۳؛ کوچ و همکاران، ۲۰۲۴الف).

در ساختار MAILS، مؤلفه «سواد هوش مصنوعی» پنج بُعد استفاده و کاربرد، درک مفاهیم، تشخیص هوش مصنوعی، اخلاق هوش مصنوعی و ایجاد هوش مصنوعی را دربر می‌گیرد. این پنج بُعد نشان می‌دهد فرد تا چه حد می‌تواند ابزارهای هوشمند را به‌کار گیرد، منطق پایه آن‌ها را بفهمد، خروجی یا محتوای مبتنی بر هوش مصنوعی را تشخیص دهد، پیامدهای اخلاقی و اجتماعی آن را در نظر بگیرد و در سطحی خلاقانه با ابزارها یا خروجی‌های هوش مصنوعی کار کند. مؤلفه «خودکارآمدی» نیز دو بُعد حل مسائل مرتبط با هوش مصنوعی و یادگیری را شامل می‌شود و نشان می‌دهد دانشجو تا چه اندازه خود را قادر به فهم تغییرات سریع، حل خطاها، اصلاح راهبردها و یادگیری مستمر در این حوزه می‌داند. مؤلفه «خودمدیریتی» نیز با دو بُعد سواد متقاعدسازی و تنظیم احساسات، به جنبه‌های انتقادی و روان‌شناختی تعامل با هوش مصنوعی می‌پردازد؛ یعنی اینکه دانشجو بتواند پیام‌های اقناعی، سوگیری‌ها، فشار شناختی و واکنش‌های هیجانی خود را هنگام مواجهه با سامانه‌های هوشمند مدیریت کند. از این رو، MAILS برای پژوهش حاضر مناسب بود، زیرا امکان می‌داد تفاوت میان «استفاده از هوش مصنوعی» و «سواد انتقادی و خودتنظیم‌گرانه در مواجهه با هوش مصنوعی» به‌صورت تجربی سنجیده شود.

رمان (۱۹۸۴) جرج اورول استعاره‌ای برای توضیح اهمیت سواد داده و حریم خصوصی فراهم می‌آورد. اورول فضای نظارت فراگیر، تبلیغات و بازنویسی حقیقت را به تصویر می‌کشد (اورول، ۱۹۴۹). در عصر کلان‌داده، این نگرانی از نظارت جمعی فراتر می‌رود؛ زیرا دانه‌بندی و پیوند داده‌ها امکان بازسازی الگوهای جست‌وجو، حرکت، ارتباط و ترجیح هر فرد را فراهم می‌کند. در نتیجه، سامانه هوشمند می‌تواند نه فقط ناظر، بلکه واسطه‌ای شخصی‌سازی‌شده برای انتخاب آنچه فرد می‌بیند، می‌پرسد و باورپذیر می‌یابد باشد.

⁸ Meta AI Literacy Scale (MAILS)

این استعاره با تحلیل‌های معاصر درباره سوگیری و قدرت الگوریتمی پیوند می‌خورد. اونیل (۲۰۱۶) و شمارزو (۲۰۲۳) نشان می‌دهند مدل‌های به‌ظاهر عینی در حوزه‌هایی مانند استخدام، اعتبار و پذیرش دانشگاهی می‌توانند سوگیری‌های تاریخی را بازتولید کنند. از این‌رو، سواد هوش مصنوعی در این پژوهش فقط شناخت دقت کلی مدل نیست؛ بلکه توان پرسش درباره منشأ داده، توزیع هزینه خطا، گروه‌های کمترنمایان، امکان اعتراض و قابلیت توضیح، اصلاح یا رد خروجی را نیز شامل می‌شود. بدین ترتیب، رمان «۱۹۸۴»، دیدگاه فارماکون استیگلر، نقد هان و رویکرد هاراوی همگی به ضرورت حفظ عاملیت، تنوع و مسئولیت انسانی در زیست‌بوم داده‌ای اشاره دارند.

در ایران نیز دو مطالعه پیشین از مقیاس MAILS استفاده کرده و ساختار یا ویژگی‌های روان‌سنجی آن را بررسی کرده‌اند: امیدی و جامه‌بزرگ (۱۴۰۴) در میان دانشجویان دانشگاه علامه طباطبائی و حاجی‌انوری و رمضانی (۱۴۰۳) در میان اعضای هیأت علمی. با وجود این مطالعات، شواهد درباره وضعیت چندبعدی سواد هوش مصنوعی در دانشگاه‌های جامع و مقایسه هم‌زمان مؤلفه‌ها، مقاطع و گروه‌های رشته‌ای همچنان محدود است. این خلأ به‌ویژه از آن رو اهمیت دارد که دانشجویان رشته‌های انسانی، اجتماعی، مهندسی، علوم پایه و علوم کامپیوتر با فرصت‌های متفاوت یادگیری و مواجهه با هوش مصنوعی وارد دانشگاه می‌شوند.

بر این اساس، پژوهش حاضر با تمرکز بر دانشجویان دانشگاه خوارزمی، سطح ادراک شده سواد هوش مصنوعی آنان را بر اساس مقیاس MAILS ارزیابی می‌کند و می‌کوشد نشان دهد وضعیت مؤلفه‌ها و ابعاد مختلف چگونه است، آیا میانگین کلی از سطح متوسط نظری فراتر می‌رود و آیا میان مقاطع و گروه‌های رشته‌ای تفاوت معناداری مشاهده می‌شود. این مطالعه در پی تعیین اثر علی مقطع یا رشته نیست؛ بلکه الگوی توصیفی تفاوت‌ها را برای طراحی برنامه‌های آموزشی، کارگاه‌های دانشگاهی و سیاست‌گذاری سواد هوش مصنوعی مسئولانه فراهم می‌کند.

پرسش‌های پژوهش

۱. سطح بسندگی سواد هوش مصنوعی دانشجویان دانشگاه خوارزمی چگونه است؟
۲. آیا میان مؤلفه‌ها و ابعاد مختلف سواد هوش مصنوعی دانشجویان تفاوت معناداری وجود دارد؟
۳. آیا بین مقاطع مختلف تحصیلی دانشجویان از نظر سطح سواد هوش مصنوعی تفاوت معناداری وجود دارد؟
۴. آیا بین گروه‌های رشته تحصیلی مختلف دانشجویان از نظر سطح سواد هوش مصنوعی تفاوت معناداری وجود دارد؟

روش پژوهش

پژوهش حاضر از نظر هدف، کاربردی و از نظر نحوه گردآوری داده‌ها، توصیفی و از نوع پیمایشی بود. جامعه آماری پژوهش شامل دانشجویان دانشگاه خوارزمی در سه مقطع کارشناسی، کارشناسی ارشد و دکتری بود. جامعه آماری شامل ۶۶۱۲ دانشجوی کارشناسی، ۴۲۷۶ دانشجوی کارشناسی ارشد و ۹۶۴ دانشجوی دکتری بود. نمونه پژوهش با روش طبقه‌ای نسبی بر اساس مقطع تحصیلی و با ملاحظه سهولت دسترسی انتخاب شد. حجم نمونه پیش‌بینی شده پیش از ریزش ۳۷۲ نفر بود و پس از حذف پاسخ‌های ناقص یا نامعتبر، داده‌های ۳۳۵ نفر در تحلیل نهایی باقی ماند. توزیع نمونه پس از ریزش شامل ۱۸۸ دانشجوی کارشناسی، ۱۱۷ دانشجوی کارشناسی ارشد و ۳۰ دانشجوی دکتری بود (جدول ۱).

جدول ۱. حجم جامعه و نمونه بر حسب مقطع تحصیلی

مقطع تحصیلی	حجم جامعه آماری	نمونه قبل از ریزش	نمونه بعد از ریزش
کارشناسی	۶۶۱۲	۲۰۸	۱۸۸
کارشناسی ارشد	۴۲۷۶	۱۳۴	۱۱۷
دکتری	۹۶۴	۳۰	۳۰

ابزار گردآوری داده‌ها، پرسشنامه استاندارد سواد هوش مصنوعی متا (MAILS) بود. این ابزار با مقیاس پاسخ‌دهی ۰ تا ۱۰ اجرا شد؛ به این معنا که نمره بالاتر نشان‌دهنده سطح بالاتر ادراک شده از توانایی، آگاهی یا مهارت در همان بُعد است. MAILS سه مؤلفه اصلی و نه بُعد را پوشش می‌دهد. مؤلفه سواد هوش مصنوعی شامل استفاده و کاربرد، درک مفاهیم، تشخیص هوش مصنوعی، اخلاق هوش مصنوعی و ایجاد هوش مصنوعی است؛ مؤلفه خودکارآمدی شامل حل مسائل مرتبط با

هوش مصنوعی و یادگیری است؛ و مؤلفه خودمدیریتی شامل سواد متقاعدسازی و تنظیم احساسات است. در این پژوهش، نمره هر بُعد از میانگین گویه‌های مربوط به همان بُعد، نمره هر مؤلفه از میانگین ابعاد زیرمجموعه، و نمره کلی سواد هوش مصنوعی از میانگین کل ابعاد محاسبه شد.

انتخاب MAILS از آن جهت اهمیت داشت که این ابزار، سواد هوش مصنوعی را صرفاً به سطح استفاده ابزاری از سامانه‌های هوشمند تقلیل نمی‌دهد. در این مقیاس، دانش مفهومی، توان تشخیص، اخلاق، آفرینش، خودکارآمدی، یادگیری، سواد متقاعدسازی و تنظیم احساسات در کنار هم سنجیده می‌شوند. بنابراین، نتایج پژوهش می‌تواند نشان دهد دانشجویان در استفاده روزمره از ابزارهای هوش مصنوعی تا چه حد توانمندند یا اینکه در ابعاد عمیق‌تر، مانند فهم سازوکارها، ارزیابی انتقادی، حل مسئله، خلق با هوش مصنوعی و مدیریت هیجان و اقناع الگوریتمی نیز از سطح کافی برخوردارند. با این همه، باید تأکید کرد که MAILS یک ابزار خودگزارشی است و به‌تنهایی نمی‌تواند نشان دهد دانشجو در غیاب ابزار نیز قادر به انجام تکلیف است، آموخته‌ها را به موقعیت تازه انتقال می‌دهد، خروجی را به‌درستی راستی‌آزمایی می‌کند یا مالکیت شناختی خود را حفظ کرده است. بنابراین، یافته‌های حاضر سطح ادراک‌شده از شایستگی را می‌سنجند و نباید به‌عنوان شاهد مستقیم افزایش یادگیری یا نبود فرسایش عاملیت تفسیر شوند.

با توجه به اینکه ساختار و ویژگی‌های روان‌سنجی مقیاس MAILS در پژوهش‌های پیشین ایرانی (امیدی و جامه‌بزرگ، ۱۴۰۴؛ حاجی‌انوری و رضانی، ۱۴۰۳) بررسی و تأیید شده بود، روایی محتوایی آن در پژوهش حاضر مجدداً ارزیابی نشد. در این مطالعه صرفاً روایی صوری کیفی بررسی شد؛ بدین منظور، نسخه فارسی پرسشنامه در اختیار ۲۰ نفر از دانشجویان جامعه هدف و استادان راهنما و مشاور قرار گرفت و بازخورد آنان درباره وضوح، روانی، فهم‌پذیری و تناسب واژگان گویه‌ها در اصلاح نهایی لحاظ شد.

پایایی ابزار با آلفای کرونباخ محاسبه شد. ضریب آلفای کرونباخ برای کل پرسشنامه ۰/۹۶۳ بود. ضرایب پایایی مؤلفه‌ها و ابعاد نیز در دامنه قابل قبول تا بسیار مطلوب قرار داشت؛ به‌گونه‌ای که آلفای مؤلفه سواد هوش مصنوعی ۰/۹۵۴، خودکارآمدی ۰/۹۱۷ و خودمدیریتی ۰/۸۹۳ گزارش شد (جدول ۲). این ضرایب نشان می‌دهد ابزار از همسانی درونی مناسبی برای سنجش سازه‌های مورد نظر برخوردار بوده است.

داده‌ها با نرم‌افزار SPSS نسخه ۲۷ تحلیل شدند. ابتدا شاخص‌های توصیفی شامل فراوانی، درصد، میانگین، انحراف معیار، حداقل و حداکثر محاسبه شد. سپس پیش‌فرض‌های آماری با آزمون کولموگوروف - اسمیرنوف و آزمون برابری واریانس‌ها بررسی شد. برای بررسی مطلوبیت سطح کلی سواد هوش مصنوعی از آزمون t تک‌نمونه‌ای استفاده شد. برای مقایسه مؤلفه‌ها و ابعاد از آزمون فریدمن، برای مقایسه مقاطع تحصیلی از تحلیل واریانس یک‌طرفه و آزمون تعقیبی توکی، و برای برخی مقایسه‌های رتبه‌ای بر حسب مقطع و رشته تحصیلی از آزمون کروسکال - والیس استفاده شد.

جدول ۲. پایایی مؤلفه‌ها و ابعاد پرسشنامه سواد هوش مصنوعی متا

مؤلفه/بُعد	نام مؤلفه/بُعد	تعداد ابعاد/گویه‌ها	آلفای کرونباخ
بُعد	سواد هوش مصنوعی	۵ بُعد	۰/۹۵۴
	استفاده و کاربرد	۶ گویه	۰/۹۳۸
	درک مفاهیم	۵ گویه	۰/۹۰۲
	تشخیص	۳ گویه	۰/۷۸۴
	اخلاق	۳ گویه	۰/۸۲۷
	ایجاد ابزارهای جدید	۲ گویه	۰/۸۶۳
مؤلفه	خودکارآمدی	۲ بُعد	۰/۹۱۷
بُعد	یادگیری/به‌روز بودن	۲ گویه	۰/۹۶۲
	حل مسئله	۲ گویه	۰/۸۴۸
مؤلفه	خودمدیریتی	۲ بُعد	۰/۸۹۳
بُعد	سواد متقاعدسازی	۳ گویه	۰/۸۵۰
	تنظیم هیجان و احساسات	۳ گویه	۰/۸۴۶
کل	کل پرسشنامه	۲۹ گویه	۰/۹۶۳

یافته‌های پژوهش

ویژگی‌های جمعیت‌شناختی پاسخ‌دهندگان

از مجموع ۳۳۵ پاسخ‌دهنده، ۲۳۲ نفر زن، ۱۰۰ نفر مرد و ۳ نفر مایل به پاسخ‌گویی به پرسش جنسیت نبودند (جدول ۳). از نظر مقطع تحصیلی، بیش از نیمی از پاسخ‌دهندگان در مقطع کارشناسی قرار داشتند. از نظر سنی، بیشترین سهم مربوط به گروه ۱۸ تا ۲۲ سال بود. همچنین بیشترین فراوانی رشته‌ای به گروه علوم اجتماعی و رفتاری اختصاص داشت.

جدول ۳. ویژگی‌های جمعیت‌شناختی نمونه پژوهش

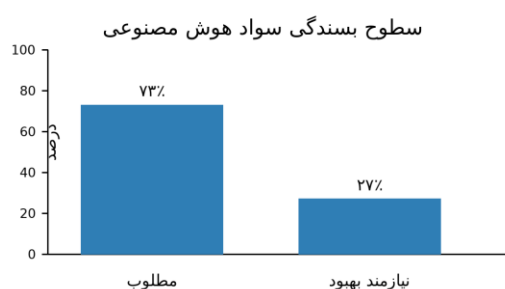
متغیر	طبقه	فراوانی	درصد
جنسیت	زن	۲۳۲	۶۹/۲
	مرد	۱۰۰	۲۹/۸
	مایل به پاسخ نیستم	۳	۰/۸
مقطع	کارشناسی	۱۸۸	۵۶/۱
	کارشناسی ارشد	۱۱۷	۳۴/۹
	دکتری	۳۰	۸/۹
سن	۱۸ تا ۲۲ سال	۱۵۹	۴۷/۵
	۲۳ تا ۲۷ سال	۱۱۳	۳۳/۷
	۲۸ تا ۳۲ سال	۳۹	۱۱/۶
	۳۳ سال و بالاتر	۲۴	۷/۲
رشته	علوم اجتماعی و رفتاری	۱۴۴	۴۳/۰
	علوم پایه	۷۲	۲۱/۵
	علوم انسانی	۴۸	۱۴/۳
	فنی و مهندسی	۴۸	۱۴/۳
	علوم و مهندسی کامپیوتر	۲۳	۶/۹

سطح کلی و بسندگی سواد هوش مصنوعی

نتایج (جدول ۴) نشان داد میانگین سواد کلی هوش مصنوعی دانشجویان برابر با ۵/۸۳ از ۱۰ است. آزمون t تک‌نمونه‌ای نشان داد این میانگین به‌طور معناداری بالاتر از حد متوسط نظری است ($t = ۱۱,۰۲$, $df = ۳۳۴$, $p < ۰,۰۱$). مقدار اندازه اثر کوهن ۰/۶۰۲ و مقدار هجز ۰/۶۰۱ بود. بر اساس نمودار بسندگی (شکل ۲)، ۷۳ درصد دانشجویان در سطح مطلوب و ۲۷ درصد در سطح نامطلوب یا نیازمند تقویت قرار گرفتند. این نتیجه بیانگر ارزیابی نسبتاً مثبت دانشجویان از شایستگی‌های خود است و با یافته امیدی و جامعه‌بزرگ (۱۴۰۴) در نمونه‌ای دیگر از دانشجویان ایرانی همخوانی دارد؛ با این حال، به دلیل خودگزارشی بودن ابزار، نباید آن را معادل عملکرد واقعی یا یادگیری پایدار تلقی کرد.

جدول ۴. نتایج آزمون t تک‌نمونه‌ای و شاخص‌های بسندگی

سواد هوش مصنوعی



شکل ۲. توزیع سطح بسندگی سواد هوش مصنوعی دانشجویان.

شاخص	مقدار
میانگین سواد کلی	۵/۸۳
انحراف معیار	۱/۸۸
t	۱۱/۰۲
درجه آزادی	۳۳۴
سطح معناداری	$> ۰/۰۰۱$
تفاوت میانگین با حد متوسط نظری	۱/۱۲
اندازه اثر کوهن	۰/۶۰۲
اندازه اثر هجز	۰/۶۰۱
سطح مطلوب	۲۴۳ نفر؛ ۷۳ درصد
سطح نامطلوب/نیازمند تقویت	۹۲ نفر؛ ۲۷ درصد

توصیف مؤلفه‌ها و ابعاد سواد هوش مصنوعی

یافته‌های توصیفی نشان داد در سطح مؤلفه‌ها، خودمدیریتی با میانگین $۶/۸۳$ و سواد هوش مصنوعی با میانگین $۶/۴۸$ بالاتر از خودکارآمدی با میانگین $۵/۱۷$ قرار گرفته‌اند. در سطح ابعاد، استفاده و کاربرد هوش مصنوعی با میانگین $۷/۰۱$ بالاترین مقدار و ایجاد ابزارهای جدید هوش مصنوعی با میانگین $۲/۵۴$ پایین‌ترین مقدار را داشت. این الگو نشان می‌دهد پاسخ‌دهندگان در استفاده از ابزارهای موجود، توانمندی ادراک‌شده بیشتری از خلق، توسعه و حل مسئله گزارش کرده‌اند. از داده‌های پیمایشی حاضر نمی‌توان نتیجه گرفت که استفاده از ابزار علت ضعف در ایجاد یا حل مسئله بوده است.

جدول ۵. میانگین و انحراف معیار مؤلفه‌ها و ابعاد سواد هوش مصنوعی

متغیر	میانگین	انحراف معیار	حداقل	حداکثر
سواد کلی هوش مصنوعی	۵/۸۳	۱/۸۸	۰/۶۶	۱۰
مؤلفه سواد هوش مصنوعی	۶/۴۸	۱/۹۶	۰/۵۴	۱۰
استفاده و کاربرد	۷/۰۱	۲/۲۶	۰	۱۰
درک مفاهیم	۵/۹۹	۲/۲۷	۰	۱۰
تشخیص	۶/۴۴	۲/۲۶	۰	۱۰
اخلاق	۶/۴۵	۲/۱۶	۰/۳۳	۱۰
ایجاد	۲/۵۴	۲/۸۱	۰	۱۰
مؤلفه خودکارآمدی	۵/۱۷	۲/۵۹	۰	۱۰
یادگیری/به‌روز بودن	۵/۴۳	۲/۷۲	۰	۱۰
حل مسئله	۴/۹۰	۲/۷۸	۰	۱۰
مؤلفه خودمدیریتی	۶/۸۳	۲/۰۲	۰/۳۳	۱۰
سواد متقاعدسازی	۶/۸۳	۲/۱۷	۰	۱۰
تنظیم هیجان و احساسات	۶/۸۳	۲/۲۲	۰	۱۰

مقایسه مؤلفه‌ها و ابعاد سواد هوش مصنوعی

برای بررسی تفاوت میان مؤلفه‌ها و ابعاد، از آزمون فریدمن استفاده شد (جدول ۶). نتایج نشان داد میان سه مؤلفه اصلی سواد هوش مصنوعی، خودکارآمدی و خودمدیریتی تفاوت معنادار وجود دارد ($\chi^2 = ۱۸۸,۹۷۶, p < ,۰۰۱$). همچنین در سطح ابعاد نه‌گانه تفاوت معنادار مشاهده شد ($\chi^2 = ۹۳۰,۴۰۶, p < ,۰۰۱$). بر اساس میانگین رتبه‌ها، استفاده و کاربرد در بالاترین رتبه و ایجاد هوش مصنوعی در پایین‌ترین رتبه قرار گرفت. این نتیجه نامتوازن بودن نمرات ابعاد را نشان می‌دهد و مبنایی برای اولویت‌بندی آموزشی فراهم می‌کند؛ اما به‌تنهایی علت این تفاوت‌ها یا اثر یک مداخله آموزشی را مشخص نمی‌سازد. (جدول ۷)

جدول ۶. نتایج آزمون فریدمن برای مؤلفه‌ها و ابعاد

مقایسه	درجه آزادی	آماره χ^2 دو	سطح معناداری
مقایسه سه مؤلفه اصلی	۲	۱۸۸/۹۷۶	$۰/۰۰۱ >$
مقایسه ابعاد نه‌گانه	۸	۹۳۰/۴۰۶	$۰/۰۰۱ >$

جدول ۷. میانگین رتبه ابعاد در آزمون فریدمن

بعد	میانگین رتبه	بعد	میانگین رتبه
استفاده و کاربرد	۶/۷۷	درک مفاهیم	۴/۹۴
تنظیم هیجان و احساسات	۶/۱۹	یادگیری/به‌روز بودن	۴/۲۹
سواد متقاعدسازی	۶/۰۱	حل مسئله	۳/۴۸
تشخیص	۵/۸۴	ایجاد	۱/۷۷
اخلاق	۵/۷۱		

مقایسه سواد هوش مصنوعی بر حسب مقطع تحصیلی

نتایج تحلیل واریانس یک طرفه نشان داد تفاوت میان دانشجویان کارشناسی، کارشناسی ارشد و دکتری از نظر سواد کلی هوش مصنوعی معنادار نیست ($F = 0.067, p = 0.935$) (جدول ۸). آزمون توکی نیز هر سه مقطع را در یک گروه همگن قرار داد. نتایج آزمون کروسکال - والیس برای سه مؤلفه اصلی نیز تفاوت معناداری میان مقاطع نشان نداد. بنابراین، در نمونه حاضر، بالاتر بودن مقطع تحصیلی با نمره بالاتر سواد هوش مصنوعی همراه نبود. این نتیجه رابطه علی میان مقطع و سواد هوش مصنوعی را رد یا اثبات نمی‌کند و ممکن است با عواملی مانند آموزش مستقیم، تجربه قبلی و فرصت کاربرد مرتبط باشد.

جدول ۸. مقایسه سواد کلی هوش مصنوعی بر حسب مقطع تحصیلی

منبع واریانس	مجموع مجزورات	df	میانگین مجزورات	F	p
بین گروه‌ها	۳۹۲/۶۳۰	۲	۱۹۶/۳۱۵	۰/۰۶۷	۰/۹۳۵
درون گروه‌ها	۹۷۴۵۷۰/۸۵۰	۳۳۲	۲۹۳۵/۴۵۴	-	-
کل	۹۷۴۹۶۳/۴۸۱	۳۳۴	-	-	-
مقطع تحصیلی	تعداد		میانگین نمرات	گروه همگن توکی	
کارشناسی ارشد	۱۲۰		۱۷۶/۵۰	۱	
کارشناسی	۱۸۵		۱۷۷/۷۲	۱	
دکتری	۳۰		۱۸۰/۴۷	۱	

مقایسه سواد هوش مصنوعی بر حسب رشته تحصیلی

بر حسب رشته تحصیلی، تفاوت معناداری در نمره سواد کلی هوش مصنوعی مشاهده شد ($F = 7.900, p < 0.01$) (جدول ۹-۱۱). دانشجویان علوم و مهندسی کامپیوتر با میانگین ۷/۵۳ بالاترین نمره را گزارش کردند و پس از آنان دانشجویان علوم پایه، فنی و مهندسی، علوم انسانی و علوم اجتماعی و رفتاری قرار گرفتند. آزمون کروسکال - والیس نیز تفاوت رشته‌ای در مؤلفه‌ها و ابعاد را نشان داد. این الگو می‌تواند با تفاوت فرصت‌های آموزشی، برنامه درسی و میزان مواجهه با داده و ابزارهای محاسباتی مرتبط باشد؛ با این حال، طرح پیمایشی حاضر اجازه نمی‌دهد رشته تحصیلی به عنوان علت این تفاوت معرفی شود.

جدول ۹. میانگین سواد کلی هوش مصنوعی بر حسب گروه رشته تحصیلی

گروه رشته تحصیلی	میانگین	انحراف معیار
علوم و مهندسی کامپیوتر	۷/۵۳	۱/۳۲
علوم پایه	۶/۶۰	۱/۷۳
فنی و مهندسی	۶/۴۲	۱/۹۴
علوم انسانی	۵/۷۳	۱/۸۴
علوم اجتماعی و رفتاری	۵/۶۹	۱/۸۱
کل نمونه	۶/۱۲	۱/۸۶

جدول ۱۰. نتایج آزمون آنوا برای مقایسه سواد کلی هوش مصنوعی بر اساس رشته تحصیلی

منبع تغییر	مجموع مجزورات	df	میانگین مجزورات	F	p
بین گروه‌ها	۱۰۱/۳۰۹	۴	۲۵/۳۲۷	۷/۹۰۰	۰/۰۰۱>
درون گروه‌ها	۱۰۵۷/۹۸۱	۳۳۰	۳/۲۰۶	-	-
کل	۱۱۵۹/۲۹۱	۳۳۴	-	-	-

جدول ۱۱. میانگین مؤلفه‌ها و ابعاد به تفکیک گروه رشته تحصیلی

مؤلفه/بعد	علوم اجتماعی و رفتاری	علوم انسانی	فنی و مهندسی	علوم پایه	علوم و مهندسی کامپیوتر
سواد هوش مصنوعی	۶/۰۴	۶/۰۴	۶/۷۷	۷/۰۳	۷/۸۴
استفاده	۶/۵۳	۶/۵۷	۷/۴۸	۷/۴۷	۸/۶۰
درک مفاهیم	۵/۶۳	۵/۴۵	۶/۳۲	۶/۴۱	۷/۴۵
تشخیص	۶/۰۰	۶/۰۳	۶/۸۲	۶/۹۲	۷/۸۶
اخلاق	۵/۹۹	۶/۱۰	۶/۴۴	۷/۳۲	۷/۴۵
ایجاد	۲/۱۰	۲/۱۵	۳/۵۰	۲/۳۷	۴/۷۴
خودکارآمدی	۴/۶۹	۴/۵۹	۵/۵۴	۵/۶۰	۷/۲۷
حل مسائل	۴/۳۸	۴/۲۷	۵/۳۱	۵/۴۰	۷/۰۹
یادگیری	۴/۹۹	۴/۹۱	۵/۷۶	۵/۸۱	۷/۴۶
خودمدیریتی	۶/۴۶	۶/۷۳	۶/۷۸	۷/۴۶	۷/۵۳
سواد متقاعدسازی	۶/۵۵	۶/۶۷	۶/۵۱	۷/۴۶	۷/۶۲
تنظیم احساسات	۶/۳۷	۶/۷۸	۷/۰۵	۷/۴۶	۷/۴۳

بحث

یافته نخست پژوهش نشان داد میانگین سواد کلی هوش مصنوعی دانشجویان دانشگاه خوارزمی ۵/۸۳ از ۱۰ بود و ۷۳ درصد پاسخ‌دهندگان در سطح مطلوب قرار گرفتند. این نتیجه با گزارش امیدی و جامه‌بزرگ (۱۴۰۴) درباره مطلوب بودن نسبی سواد هوش مصنوعی دانشجویان همخوان است. با این حال، نمره حاصل از MAILS بیانگر شایستگی ادراک شده است؛ بنابراین، مطلوب بودن میانگین کلی به معنای کفایت کامل، عملکرد مستقل یا یادگیری پایدار نیست. یافته حاضر باید همراه با تفاوت میان مؤلفه‌ها و ماهیت خودگزارشی و مقطعی داده‌ها تفسیر شود.

یافته دوم نشان داد استفاده و کاربرد با میانگین ۷/۰۱ بالاترین و ایجاد هوش مصنوعی با میانگین ۲/۵۴ پایین‌ترین نمره را داشت؛ خودکارآمدی نیز با میانگین ۵/۱۷ از دو مؤلفه دیگر پایین‌تر بود. بنابراین، پاسخ‌دهندگان در کار با ابزارهای آماده بیش از خلق، حل مسئله و سازگاری فعال با مسائل جدید احساس توانمندی کرده‌اند. این الگو با منطق چندبعدی MAILS سازگار است؛ زیرا استفاده، فهم، تشخیص، اخلاق، ایجاد، خودکارآمدی و خودمدیریتی را شایستگی‌های متمایز می‌داند و اجازه نمی‌دهد نمره بالای کاربرد به همه ابعاد تعمیم داده شود (کارولوس و همکاران، ۲۰۲۳؛ کوچ و همکاران، ۲۰۲۴الف).

این یافته را می‌توان با چارچوب نظری مقاله تفسیر کرد. در نظریه انتقادی فینبرگ، فاصله استفاده و ایجاد، تفاوت میان عمل در چارچوب فنی از پیش طراحی شده و مشارکت آگاهانه در نقد یا بازطراحی آن است. دیدگاه فارماکون استیگر توضیح می‌دهد که همان ابزاری که دسترسی و کارایی را افزایش می‌دهد، در صورت اتکای نامتناسب می‌تواند بخشی از داوری و مسئولیت را به سامانه واگذار کند. هاروی یادآور می‌شود که یادگیرنده در تعامل با ابزار شکل می‌گیرد، نه در جدایی کامل از آن؛ و نقد هان نشان می‌دهد سرعت و دسترسی دائمی ممکن است نیاز به خودمدیریتی، تنظیم هیجان و مقاومت در برابر فشار شناختی را افزایش دهد.

سواد عملیاتی، یادگیری و شکاف عاملیت

در پرتو تمایز میان عملکرد و قابلیت، الگوی نمرات می‌تواند بیانگر غلبه سواد عملیاتی بر ابعاد عاملیت‌محور باشد؛ یعنی دانشجو چگونگی گرفتن خروجی را بهتر از ایجاد، حل مسئله یا بازسازی مستقل راه حل ارزیابی کرده است. مطالعات آزمایشی باستانی و همکاران (۲۰۲۵)، کستین و همکاران (۲۰۲۵) و لی و همکاران (۲۰۲۶) در این مقاله شواهد بیرونی برای توضیح اهمیت طراحی تعامل انسان و هوش مصنوعی‌اند، نه بخشی از داده‌های پژوهش حاضر. این مطالعات نشان داده‌اند پیامدهای یادگیری می‌تواند بر حسب ارائه مستقیم پاسخ، داربست‌بندی آموزشی یا همکاری فعال متفاوت باشد.

یافته سوم نشان داد میان دانشجویان کارشناسی، کارشناسی ارشد و دکتری تفاوت معناداری در سواد کلی هوش مصنوعی مشاهده نشد. این نتیجه به معنای یکسان بودن کامل دانشجویان یا بی‌اثر بودن آموزش دانشگاهی نیست؛ بلکه فقط نشان می‌دهد در نمونه حاضر، بالاتر بودن مقطع با نمره بالاتر MAILS همراه نبوده است. ممکن است سواد هوش مصنوعی بیش از مقطع رسمی، با آموزش مستقیم، تجربه شخصی، نوع تکالیف و فرصت دسترسی مرتبط باشد. از این‌رو، نتایج از گنجاندن آموزش سواد هوش مصنوعی در همه مقاطع حمایت می‌کند، اما رابطه علی عوامل یادشده باید در طرح‌های طولی یا آزمایشی بررسی شود.

یافته چهارم نشان داد نمرات سواد هوش مصنوعی بر حسب رشته تحصیلی تفاوت معناداری داشت و دانشجویان علوم و مهندسی کامپیوتر بالاترین میانگین را گزارش کردند. نظریه کنشگر - شبکه این تفاوت را به‌صورت محصول شبکه‌ای از برنامه درسی، استادان، همکلاسی‌ها، پروژه‌ها، ابزارها، داده‌ها و زیرساخت دسترسی تفسیر می‌کند (لاتور، ۲۰۰۵). در نتیجه، تفاوت رشته‌ای را می‌توان با ناهمسانی فرصت‌های مواجهه و یادگیری مرتبط دانست، نه اثر قطعی خود رشته. این تفسیر همچنین نشان می‌دهد کاهش شکاف رشته‌ای فقط با توصیه فردی به دانشجویان ممکن نیست و به مداخله در برنامه درسی و زیرساخت نیاز دارد.

توانمندسازی فرهنگی و شهروندی داده

چارچوب «شهروند علم داده» شمارزو (۲۰۲۳) به یافته رشته‌ای و ضعف نسبی ایجاد و حل مسئله جهت کاربردی می‌دهد. دانشجو لازم نیست دانشمند داده باشد، اما باید بتواند در تعریف مسئله، ارزیابی پیامد، تشخیص شاخص‌های ارزش و نظارت اخلاقی بر کاربرد داده مشارکت کند. از این‌رو، نمره بالاتر گروه‌های فنی نباید به انحصار گفت‌وگوی اخلاقی و سیاستی در این رشته‌ها منجر شود؛ بلکه آموزش عمومی باید امکان مشارکت همه گروه‌ها را در تصمیم‌های فناورانه فراهم کند. این برداشت با نتیجه پژوهش حاضر درباره تفاوت رشته‌ای نیز سازگار است: اگر فرصت‌های برنامه درسی، پروژه‌های عملی، دسترسی به ابزار و فرهنگ رشته‌ای بخشی از شبکه شکل‌دهنده سواد باشند، کاهش شکاف فقط با توصیه فردی به دانشجو ممکن نیست و به بازطراحی فرصت‌های یادگیری، همکاری میان‌رشته‌ای و مشارکت دانشجویان غیر فنی در بحث‌های داده، اخلاق و حکمرانی هوش مصنوعی نیاز دارد. در چنین الگویی، دانشجوی علوم انسانی یا اجتماعی نیز باید بتواند درباره معیار عدالت، هزینه خطا، معنا و پیامد اجتماعی خروجی الگوریتم داوری کند و دانشجوی فنی نیز باید نسبت به محدودیت داده، سوگیری و مسئولیت نهادی حساس باشد.

نمره نسبتاً بالای خودمدیریتی در کنار پایین‌تر بودن خودکارآمدی نیز باید با احتیاط تفسیر شود. این الگو نشان می‌دهد دانشجویان توانایی ادراک شده خود را در مدیریت اقناع و هیجان بالاتر از حل مسئله و یادگیری ارزیابی کرده‌اند، اما پرسشنامه رفتار واقعی آنان را در موقعیت فشار، سرعت یا اتکای بالا اندازه‌گیری نکرده است. پیوند نقد هان با ابعاد تنظیم هیجان و سواد متقاعدسازی نشان می‌دهد دسترسی دائمی و انتظار پاسخ فوری می‌تواند فرد را به پذیرش سریع خروجی، خودنظارتی و خستگی شناختی سوق دهد. از سوی دیگر، رویکرد سایبورگ هارای یادآور می‌شود که هدف، بازگشت به جدایی کامل انسان و ماشین نیست؛ هویت یادگیرنده همواره در تعامل با ابزار، متن، استاد و شبکه شکل می‌گیرد. مسئله اصلی حفظ امکان داوری، مخالفت و بازتعریف هدف در این پیوند است. دیدگاه فارماکون استیگلر نیز این دوگانگی را روشن می‌کند: همان فناوری می‌تواند حافظه و خلاقیت را پشتیبانی کند یا در صورت واگذاری نامتناسب وظایف، تمرین یادآوری، استدلال و مسئولیت را کاهش دهد. بنابراین، آموزش خودمدیریتی باید شامل مکث پیش از پذیرش پاسخ، ثبت دلیل اعتماد یا تردید، تفکیک پیشنهاد از تصمیم نهایی و تمرین خروج از فرایند هوش مصنوعی باشد.

در مجموع، چهار یافته اصلی پژوهش - مطلوب بودن نسبی نمره کلی، نامتوازن بودن ابعاد، نبود تفاوت معنادار میان مقاطع و وجود تفاوت میان رشته‌ها - با ترکیب مدل اندازه‌گیری MAILS و چارچوب‌های فلسفه فناوری قابل تفسیر است. MAILS محل تفاوت‌ها را آشکار می‌کند؛ فینبرگ و استیگلر فاصله استفاده و عاملیت را توضیح می‌دهند؛ لاتور تفاوت رشته‌ای را در سطح شبکه تحلیل می‌کند؛ و هارای و هان ضرورت توجه به هویت، خودمدیریتی و فشار شناختی را روشن می‌سازند.

محفظه معرفتی و محدودشدن افق دانایی

در محیطی که پاسخ‌های روان، شخصی‌سازی‌شده و ازپیش‌سنتز شده به سرعت در دسترس‌اند، فرد ممکن است حجم زیادی از اطلاعات دریافت کند، اما افق دانایی خود را به همان منابع و صورت‌بندی‌های پیشنهادی سامانه محدود سازد. مفهوم «محفظه معرفتی الگوریتمی» برتری استفاده و کاربرد بر ایجاد و حل مسئله، ضرورت توجه به این موضوع را در آموزش و پژوهش آینده مطرح می‌کند. رمان «۱۹۸۴» استعاره‌ای برای فهم وضعیتی است که در آن واسطه اطلاعاتی فقط محتوا را منتقل نمی‌کند، بلکه مرز پرسش‌های قابل طرح و روایت‌های قابل مشاهده را نیز شکل می‌دهد. از این رو، سواد عاملیت‌محور باید توان خروج از محیط مدل، مقایسه با منابع مستقل، جست‌وجوی دیدگاه مخالف، حفظ زبان و تجربه محلی و تشخیص غیبت یا کم‌نمایی گروه‌ها در داده را دربر گیرد.

حریم خصوصی و محرمانگی اطلاعات

بعد اخلاق در پژوهش حاضر میانگین ۶/۴۵ داشت. داده محصول جانبی بسیاری از کنش‌های روزمره است و پیوند تراکنش‌ها، جست‌وجوها، مکالمات، موقعیت مکانی و پرونده‌های سازمانی می‌تواند برای استنباط گرایش‌های فردی به کار رود. در این سطح، استعاره «۱۹۸۴» از نظارت آشکار به نظارت داده‌محور و شخصی‌سازی‌شده منتقل می‌شود: سامانه می‌تواند بر پایه ردپای دیجیتال نه فقط رفتار گذشته، بلکه ترجیح، آسیب‌پذیری یا احتمال تصمیم آینده را برآورد کند. دانشجو باید بداند متن ورودی به چت‌بات ممکن است حاوی داده شخصی، اطلاعات پژوهشی منتشر نشده، داده هویتی مشارکت‌کنندگان، محتوای محرمانه سازمانی یا ترکیبی از داده‌های ظاهراً بی‌خطر باشد که در کنار هم هویت فرد را آشکار می‌کنند.

پیامد آموزشی این بحث آن است که حفظ محرمانگی نباید به یک جمله تشریفاتی محدود شود. سواد هوش مصنوعی باید تشخیص داده حساس، ارزیابی ضرورت اشتراک‌گذاری، حذف یا ناشناس‌سازی شناسه‌ها، بررسی شرایط استفاده و سیاست حریم خصوصی و پرهیز از بارگذاری داده محرمانه در سامانه‌های عمومی را دربر گیرد. افزون بر این، دانشجو باید تفاوت میان ناشناس‌سازی و مستعارسازی، امکان بازشناسایی از طریق ترکیب داده‌ها، محل ذخیره‌سازی و مدت نگهداری اطلاعات و حق حذف یا اعتراض را بشناسد. این توانایی‌ها با مفهوم عاملیت پیوند دارند، زیرا کاربر را از پذیرش منفعلانه به تصمیم آگاهانه درباره نوع داده، هدف پردازش، نهاد دریافت‌کننده، امکان استفاده ثانویه و مدت نگهداری منتقل می‌کنند. دانشگاه نیز باید مسئولیت را فقط بر دوش فرد نگذارد و ابزارهای مجاز، طبقه‌بندی داده و مسیر گزارش خطا یا نقض محرمانگی را روشن سازد. چارچوب «منشور حقوق هوش مصنوعی» در روایت شمارزو (۲۰۲۳) این دلالت را تکمیل می‌کند: حفاظت در برابر سامانه نایمن، مصونیت از تبعیض الگوریتمی، اختیار نسبت به استفاده از داده و آگاهی از نقش سامانه خودکار در نتایج اثرگذار بر فرد. این اصول با نظریه انتقادی فینبرگ پیوند دارند، زیرا نشان می‌دهند طراحی فنی همواره انتخابی اجتماعی و قابل مناقشه است؛ با فارماکون استیگر نیز سازگارند، زیرا بهره‌مندی از ظرفیت فناوری مستلزم مراقبت از پیامدهای آن است. بنابراین، پیشنهاد‌های مربوط به حریم خصوصی، توضیح‌پذیری، رضایت و امکان اعتراض از یافته مستقیم پرسشنامه فراتر می‌روند، اما با بعد اخلاق، هدف آموزش سواد مسئولانه و ضرورت حفظ عاملیت انسانی سازگارند. در برنامه آموزشی، این اصول باید به تمرین‌های عینی مانند ارزیابی سیاست یک ابزار، بازنویسی یک ورودی برای حذف داده حساس و تحلیل اینکه چه کسی از خطای الگوریتم زیان می‌بیند تبدیل شوند.

برآیند نظری و تجربی این بحث را می‌توان در قالب الگویی سه‌سطحی از سواد هوش مصنوعی صورت‌بندی کرد. سطح نخست، «سواد عملیاتی» است و توان انتخاب و استفاده از ابزار، نوشتن دستور، دریافت خروجی و به‌کارگیری آن در تکالیف روزمره را دربر می‌گیرد. بالاترین میانگین بعد استفاده و کاربرد در پژوهش حاضر نشان می‌دهد دانشجویان در این سطح، دست‌کم بر اساس ارزیابی خود، وضعیت مساعدتری دارند. سطح دوم، «سواد انتقادی و خودتنظیم‌گرانه» است که فهم مفاهیم، تشخیص خروجی‌های هوش مصنوعی، ملاحظه اخلاق، سواد متقاعدسازی و تنظیم هیجان را شامل می‌شود. این سطح به دانشجو امکان می‌دهد درباره منبع، محدودیت، سوگیری، عدم قطعیت و پیامدهای یک پاسخ پرسش کند و واکنش شناختی و هیجانی خود را در برابر اقناع الگوریتمی مدیریت نماید. سطح سوم، «سواد خلاق و عاملیت‌محور» است که ایجاد، حل مسئله، بازطراحی راهبرد، حفظ مالکیت شناختی و توان ادامه کار پس از حذف حمایت هوش مصنوعی را در بر می‌گیرد. پایین بودن بعد ایجاد و پایین‌تر بودن خودکارآمدی نسبت به دو مؤلفه دیگر، اهمیت تقویت این سطح را برجسته می‌کند؛ اما چون داده‌ها

خودگزارشی و مقطعی‌اند، این الگو فقط تفسیری برای سازمان‌دهی یافته‌هاست و آزمون مستقیم سه سطح یا مسیر علی میان آنها به شمار نمی‌رود.

هر یک از چارچوب‌های نظری مقاله بخشی از این الگو را روشن می‌کند. نظریه انتقادی فناوری فینبرگ توضیح می‌دهد که چرا تسلط بر کاربری ابزار، بدون فهم «کد فنی» و امکان نقد و بازطراحی آن، برای سواد کامل کافی نیست؛ از این منظر، فاصله میان استفاده و ایجاد، فاصله میان سازگاری با طراحی موجود و مشارکت آگاهانه در شکل‌دهی فناوری است. مفهوم فارماکون نزد استیگلر نیز ماهیت دوگانه همین فرایند را نشان می‌دهد: هوش مصنوعی می‌تواند هم حافظه، خلاقیت و دسترسی به دانش را پشتیبانی کند و هم، در صورت واگذاری مداوم صورت‌بندی مسئله و داوری، ظرفیت‌های مستقل را کم‌ترین سازد. پس‌اپدیدارشناسی یادآور می‌شود که ابزار تنها پاسخ را منتقل نمی‌کند، بلکه نحوه دیدن مسئله و معیار پذیرفتن پاسخ را میانجی‌گری می‌کند؛ بنابراین، ابعاد تشخیص، اخلاق و خودمدیریتی باید در کنار کاربرد آموزش داده شوند. نظریه کنشگر - شبکه نیز یافته تفاوت رشته‌ای را در شبکه‌ای از برنامه درسی، استاد، همکلاسی، زیرساخت، داده و فرهنگ رشته‌ای قرار می‌دهد و مانع از آن می‌شود که تفاوت مشاهده‌شده صرفاً به ویژگی ذاتی افراد یا «اثر» رشته نسبت داده شود.

رویکرد سایبورگ هاروی به این جمع‌بندی بعد هویتی می‌افزاید. دانشجوی امروز در بسیاری از فعالیت‌های علمی با جست‌وجوگرها، سامانه‌های توصیه‌گر و ابزارهای مولد پیوند خورده است و محصول یادگیری او در تعامل انسان و ماشین شکل می‌گیرد. مسئله، بازگشت به جدایی کامل انسان و فناوری نیست، بلکه حفظ امکان نسبت‌دادن هدف، انتخاب، داوری و مسئولیت است. نقد هان نیز نشان می‌دهد که سرعت، دسترس‌پذیری دائمی و الزام ضمنی به بهره‌وری می‌تواند دانشجو را به پذیرش سریع خروجی، خودنظارتی و فشار شناختی سوق دهد. از این‌رو، تنظیم هیجان و سواد متقاعدسازی در MAILS فقط مؤلفه‌هایی روان‌شناختی نیستند، بلکه بخشی از توان مقاومت در برابر شتاب، تأییدطلبی و پذیرش بی‌درنگ پاسخ به شمار می‌آیند.

نتیجه‌گیری

نتایج این پیمایش نشان داد سواد ادراک‌شده هوش مصنوعی دانشجویان دانشگاه خوارزمی در مجموع نسبتاً مطلوب، اما در مؤلفه‌ها و ابعاد ناموازن است. شکاف اصلی میان نمره بالای استفاده و کاربرد و نمره پایین‌تر ایجاد، حل مسئله و خودکارآمدی مشاهده شد. همچنین، میان مقاطع تحصیلی تفاوت معناداری دیده نشد، در حالی که نمرات بر حسب رشته تحصیلی متفاوت بود.

پیام سیاستی این الگو آن است که آموزش عالی نباید به معرفی ابزارها و آموزش کار با چت‌بات‌ها محدود شود. شواهد آزمایشی پیشین نشان می‌دهد پیامدهای یادگیری به طراحی تعامل وابسته است و ارائه مستقیم پاسخ، داربست‌بندی و همکاری فعال می‌توانند نتایج متفاوتی داشته باشند (باستانی و همکاران، ۲۰۲۵؛ کستین و همکاران، ۲۰۲۵؛ لی و همکاران، ۲۰۲۶).

در نمونه حاضر، بالاتر بودن مقطع تحصیلی با نمره بالاتر سواد هوش مصنوعی همراه نبود. از این‌رو، دانشگاه باید آموزش این سواد را از کارگاه‌های پراکنده به برنامه‌ریزی درسی همه مقاطع ارتقا دهد. برای رشته‌های غیر فنی، آموزش می‌تواند از کاربرد ابزار آغاز شود، اما باید به فهم داده و الگوریتم، اخلاق، سوگیری، حریم خصوصی، مالکیت فکری و اعتبارسنجی خروجی گسترش یابد. تفاوت رشته‌ای نیز ضرورت برنامه‌های متناسب با بافت هر رشته و دسترسی عادلانه به فرصت‌های یادگیری را نشان می‌دهد.

بر اساس یافته‌های این پژوهش، پیشنهاد می‌شود دانشگاه خوارزمی و سایر دانشگاه‌ها برنامه‌ای چندسطحی برای ارتقای سواد هوش مصنوعی طراحی کنند: سطح نخست، آشنایی مفهومی، داده‌ای و اخلاقی برای همه دانشجویان؛ سطح دوم، کاربردهای رشته‌ای؛ سطح سوم، حل مسئله، ارزیابی خروجی و اتکای متناسب؛ و سطح چهارم، عاملیت و تنوع معرفتی. در سطح چهارم، دانشجو باید بیاموزد پیش از مراجعه به ابزار مسئله و معیار خود را بنویسد، ابتدا پیش‌نویس یا راه‌حل اولیه تولید کند، خروجی مدل را با منابع مستقل مقایسه کند، عدم قطعیت و دیدگاه مخالف را جست‌وجو کند و در موقعیت‌های نامناسب از استفاده امتناع ورزد. این «سواد امتناع و خروج» مکمل سواد استفاده است و امکان بازگشت به عمل مستقل را حفظ می‌کند.

در سطح سیاست دانشگاهی، برنامه سواد هوش مصنوعی باید با نظام روشن حفاظت از داده همراه شود: طبقه‌بندی و محرمانگی داده‌های آموزشی و پژوهشی، تعیین ابزارهای مجاز برای داده حساس، آموزش خواندن سیاست‌های حریم خصوصی، الزام به ناشناس‌سازی داده مشارکت‌کنندگان، اطلاع‌رسانی درباره پردازش خودکار و پیش‌بینی مسیر گزارش خطا یا تبعیض. این اقدامات، ترجمان دانشگاهی آگاهی از حریم خصوصی و اصول مطرح‌شده در منشور حقوق هوش مصنوعی‌اند و مانع می‌شوند مسئولیت حفاظت از داده صرفاً به انتخاب فردی دانشجو واگذار شود (شمارزو، ۲۰۲۳).

از نظر نظری، نتایج نشان می‌دهد ساختار MAILS برای سازمان‌دهی آموزش نیز قابل استفاده است: مسیر شناختی - عملی برای فهم، تشخیص، اخلاق، کاربرد و ایجاد؛ مسیر خودکارآمدی برای حل مسئله و یادگیری مستمر؛ و مسیر خودمدیریتی برای سواد متقاعدسازی، تنظیم هیجان و مقاومت در برابر پذیرش غیرانتقادی. نظریه انتقادی فینبرگ و دیدگاه فارماکون استیگر این ساختار را از آموزش مصرف مؤثر ابزار به آموزش مداخله آگاهانه و مسئولانه گسترش می‌دهند؛ لاتور، هاروی و هان نیز ابعاد شبکه‌ای، هویتی و خودتنظیم‌گرانه آن را روشن می‌کنند.

برای پژوهش و ارزیابی آینده، پیشنهاد می‌شود خودسنجی با چهار دسته شاخص عملکردی تکمیل شود: نخست، کیفیت اجرای مستقل پیش از استفاده از هوش مصنوعی؛ دوم، کیفیت و فرایند عملکرد در حضور هوش مصنوعی؛ سوم، انتقال با تأخیر و اجرای تکلیف هم‌ارز پس از حذف ابزار؛ و چهارم، شاخص‌های عاملیت شامل دقت راستی‌آزمایی، کالیبراسیون اعتماد، توان توضیح مسیر استدلال، تنوع منابع، یادآوری محتوا و احساس مالکیت بر محصول. مقایسه شرایط بدون هوش مصنوعی، هوش مصنوعی واکنشی و هوش مصنوعی پیش‌دستانه می‌تواند روشن کند که سامانه در کجا نقش داربست یادگیری و در کجا نقش جانشین شناخت را ایفا می‌کند.

در کاربرد روزمره، هوش مصنوعی مولد باید مانند «دستیار پژوهشی سخت‌کوش اما خطاپذیر» مدیریت شود، نه مرجع نهایی حقیقت. بر مبنای پیشنهاد شمارزو (۲۰۲۳)، کاربر باسواد باید منبع و مبنای پاسخ را مطالبه کند، پرسش‌های مرتبه دوم و سوم بپرسد، پاسخ را با منابع معتبر و سامانه‌های دیگر مقایسه کند، احتمال و عدم قطعیت را در نظر گیرد، هزینه خطاهای مثبت و منفی را بسنجد و پیش از ادغام خروجی در تصمیم، سوگیری تأییدی و پیامدهای ناخواسته را بررسی کند. این الگو، استفاده از هوش مصنوعی را از دریافت پاسخ به فرایند فعال آزمون، نقد و مسئولیت‌پذیری تبدیل می‌کند.

پژوهش حاضر در یک دانشگاه، با طرح مقطعی - پیمایشی، نمونه‌گیری همراه با ملاحظه دسترسی و ابزار خودگزارشی انجام شد؛ بنابراین، تعمیم نتایج محدود است و استنباط علی، سنجش عملکرد واقعی، یادگیری با تأخیر یا فرسایش عاملیت امکان‌پذیر نیست. مطالعات آینده باید از نمونه‌های چنددانشگاهی، طرح‌های طولی، سنجش‌های عملکردمحور و در صورت بررسی اثر مداخله، طرح‌های آزمایشی یا شبه‌آزمایشی استفاده کنند.

ملاحظات اخلاقی

مشارکت دانشجویان آگاهانه و داوطلبانه بود. محرمانگی داده‌ها رعایت شد و اطلاعات هویتی مشارکت‌کنندگان در گزارش پژوهش افشا نشده است.

مشارکت نویسندگان

نویسنده اول: اجرای پژوهش، گردآوری داده‌ها، تحلیل آماری و تهیه پیش‌نویس اولیه. نویسنده دوم: مشارکت در طراحی روش، مشاوره علمی و بازبینی متن. نویسنده سوم: نظارت علمی، راهبری نظری و روش‌شناختی، تحلیل نهایی و تهیه پیش‌نویس اولیه، بازبینی نهایی و مکاتبات مقاله به عنوان نویسنده مسئول.

تعارض منافع

نویسندگان اعلام می‌کنند که تعارض منافی وجود ندارد.

حمایت مالی

این پژوهش حمایت مالی مستقلی دریافت نکرده است.

دسترسی به داده‌ها

داده‌های پژوهش با رعایت محرمانگی مشارکت‌کنندگان و در صورت درخواست موجه علمی، از طریق نویسندگان مسئول قابل پیگیری است.

سپاسگزاری

این مقاله مستخرج از پایان‌نامه کارشناسی‌ارشد آتنا لطیفی با راهنمایی دکتر علی عظیمی‌وقار و مشاوره دکتر محمد زره‌ساز در رشته علم اطلاعات و دانش‌شناسی دانشگاه خوارزمی است. نویسندگان از دانشجویانی که در فرایند گردآوری داده‌ها همکاری کردند سپاسگزاری می‌کنند.

منابع

امیدی، مصطفی؛ و جامه‌بزرگ، زهرا (۱۴۰۴). ارزیابی سواد هوش مصنوعی و تحلیل ساختار مفهومی آن در میان دانشجویان دانشگاه علامه طباطبائی. مدیریت، آموزش و توسعه در عصر دیجیتال، ۲(۲)، ۱-۱۸. <https://doi.org/10.22034/hel.2024.2036769.1985>
حاجی‌نوری، لادن؛ و رمضانی، عباس (۱۴۰۳). بررسی وضعیت سواد کاربست و عوامل مؤثر بر پذیرش هوش مصنوعی در بین اعضای هیأت علمی. نامه آموزش عالی، ۱۷(۶۸)، ۱۰۶-۱۳۱. <https://doi.org/10.61838/medda.192>

References

- Almatrafi, O., Johri, A., & Lee, H. (2024). A systematic review of AI literacy conceptualization, constructs, and implementation and assessment efforts (2019–2023). *Computers and Education Open*, 6, 100173. <https://doi.org/10.1016/j.caeo.2024.100173>
- Annapureddy, R., Fornaroli, A., & Gatica-Perez, D. (2025). Generative AI literacy: Twelve defining competencies. *Digital Government: Research and Practice*, 6(1), 1–21. <https://doi.org/10.1145/3685680>
- Bandura, A. (1997). Self-efficacy: The exercise of control. W. H. Freeman.
- Bastani, H., Bastani, O., Sungu, A., Ge, H., Kabakcı, Ö., & Mariman, R. (2025). Generative AI without guardrails can harm learning: Evidence from high school mathematics. *Proceedings of the National Academy of Sciences*, 122(26), e2422633122. <https://doi.org/10.1073/pnas.2422633122>
- Carolus, A., Koch, M. J., Straka, S., Latoschik, M. E., & Wienrich, C. (2023). MAILS – Meta AI literacy scale: Development and testing of an AI literacy questionnaire based on well-founded competency models and psychological change- and meta-competencies. *Computers in Human Behavior: Artificial Humans*, 1(2), 100014. <https://doi.org/10.1016/j.chbah.2023.100014>
- Cheng, M., Lee, C., Khadpe, P., Yu, S., Han, D., & Jurafsky, D. (2026). Sycophantic AI decreases prosocial intentions and promotes dependence. *Science*, 391(6792), eaec8352. <https://doi.org/10.1126/science.aec8352>
- Digital Education Council. (2025). Digital Education Council AI Literacy Framework. Retrieved from <https://www.digitaleducationcouncil.com/post/digital-education-council-ai-literacy-framework>
- Doshi, A. R., & Hauser, O. P. (2024). Generative artificial intelligence enhances creativity but reduces the diversity of novel content. *Science Advances*, 10(28), eadn5290. <https://doi.org/10.1126/sciadv.adn5290>
- Dubois, M., Ududec, C., Summerfield, C., & Luettgau, L. (2026). Ask don't tell: Reducing sycophancy in large language models. *arXiv*. <https://doi.org/10.48550/arXiv.2602.23971>
- Fanous, A., Goldberg, J., Agarwal, A. A., Lin, J., Zhou, A., Daneshjou, R., & Koyejo, S. (2025). SycEval: Evaluating LLM sycophancy. *arXiv*. <https://doi.org/10.48550/arXiv.2502.08177>
- Feenberg, A. (1991). *Critical theory of technology*. Oxford University Press. pp. 18- 42.
- Freeman, J. (2025). Student generative AI survey 2025. Higher Education Policy Institute.

- Hajianvari, L., & Ramezani, A. (2024). Investigating AI literacy, application, and factors influencing artificial intelligence acceptance by faculty members. *Higher Education Letter*, 17(68), 106–131. <https://doi.org/10.22034/hel.2024.2036769.1985> [In Persian].
- Han, B.-C. (2015). *The burnout society* (E. Butler, Trans.). Stanford University Press.
- Han, B.-C. (2015). *The transparency society* (E. Butler, Trans.). Stanford University Press.
- Haraway, D. (1985). A manifesto for cyborgs: Science, technology, and socialist feminism in the 1980s. *Socialist Review*, 80, 65–108.
- Hornberger, M., Bewersdorff, A., & Nerdel, C. (2025). A multinational assessment of AI literacy among university students. *Computers and Education: Artificial Intelligence*, 8, 100371.
- Jin, Y., Martínez-Maldonado, R., Gašević, D., & Yan, L. (2025). GLAT: The generative AI literacy assessment test. *Computers and Education: Artificial Intelligence*, 9, 100436. <https://doi.org/10.1016/j.caeai.2025.100436>
- Jin, Y., Yang, K., Martínez-Maldonado, R., Gašević, D., & Yan, L. (2025). The agency gap: How generative AI literacy shapes independent writing after AI support. *arXiv*. <https://doi.org/10.48550/arXiv.2507.04398>
- Kestin, G., Miller, K., Klales, A., Milbourne, T., & Ponti, G. (2025). AI tutoring outperforms in-class active learning: An RCT introducing a novel research-based design in an authentic educational setting. *Scientific Reports*, 15, 17458. <https://doi.org/10.1038/s41598-025-97652-6>
- Koch, M. J., Carolus, A., Wienrich, C., & Latoschik, M. E. (2024a). Meta AI literacy scale: Further validation and development of a short version. *Heliyon*, 10(21), e39686. <https://doi.org/10.1016/j.heliyon.2024.e39686>
- Koch, M. J., Wienrich, C., Straka, S., Latoschik, M. E., & Carolus, A. (2024b). Overview and confirmatory and exploratory factor analysis of AI literacy scale. *Computers and Education: Artificial Intelligence*, 7, 100310. <https://doi.org/10.1016/j.caeai.2024.100310>
- Kosmyna, N., Hauptmann, E., Yuan, Y. T., Situ, J., Liao, X.-H., Beresnitzky, A. V., Braunstein, I., & Maes, P. (2025). Your brain on ChatGPT: Accumulation of cognitive debt when using an AI assistant for essay writing task. *arXiv*. <https://doi.org/10.48550/arXiv.2506.08872>
- Latour, B. (2005). *Reassembling the social: An introduction to actor-network-theory*. Oxford University Press.
- Lee, E. H., Yin, Y., Jia, N., & Waksalak, C. J. (2026). Relying on AI at work reduces self-efficacy, ownership, and meaning while active collaboration mitigates the effects. *Scientific Reports*, 16, 13583. <https://doi.org/10.1038/s41598-026-42312-6>
- Lee, H.-P., Sarkar, A., Tankelevitch, L., Drosos, I., Rintel, S., Banks, R., & Wilson, N. (2025). The impact of generative AI on critical thinking: Self-reported reductions in cognitive effort and confidence effects from a survey of knowledge workers. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems* (Article 1121, pp. 1–22). Association for Computing Machinery. <https://doi.org/10.1145/3706598.3713778>
- Long, D., & Magerko, B. (2020). What is AI literacy? Competencies and design considerations. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (pp. 1–16). Association for Computing Machinery. <https://doi.org/10.1145/3313831.3376727>
- Miao, F., Shiohira, K., & Lao, N. (2024). AI competency framework for students. UNESCO. <https://unesdoc.unesco.org/ark:/48223/pf0000391105>
- Ng, D. T. K., Leung, J. K. L., Chu, S. K. W., & Qiao, M. S. (2021). Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence*, 2, 100041. <https://doi.org/10.1016/j.caeai.2021.100041>
- OECD & European Commission. (2026). *Empowering learners for the age of AI: An AI literacy framework for primary and secondary education*. OECD Publishing. <https://doi.org/10.1787/65cd27d4-en>
- OECD. (2026). *OECD digital education outlook 2026: Exploring effective uses of generative AI in education*. OECD Publishing. <https://doi.org/10.1787/062a7394-en>
- Omidi, M., & Jamebozorg, Z. (2025). Assessing AI literacy and analyzing its conceptual structure among students at Allameh Tabataba'i University. *Management, Education and Development in the Digital Age*, 2(2), 1–18. <https://doi.org/10.61838/medda.192> [In Persian].
- Orwell, G. (1949). *Nineteen eighty-four*. Secker & Warburg.

- O'Dea, X., Ng, D. T. K., O'Dea, M., & Shkuratsky, S. (2024). Factors affecting university students' generative AI literacy: Evidence and evaluation in the UK and Hong Kong contexts. *Policy Futures in Education*, 24(1), 13–34. <https://doi.org/10.1177/14782103241287401>
- O'Neil, C. (2016). *Weapons of math destruction: How big data increases inequality and threatens democracy*. Crown.
- Pinski, M., & Benlian, A. (2024). AI literacy for users – A comprehensive review and future research directions of learning methods, components, and effects. *Computers in Human Behavior: Artificial Humans*, 2(2), 100062. <https://doi.org/10.1016/j.chbah.2024.100062>
- Schmarzo, B. (2023). *AI & data literacy: Empowering citizens of data science*. Packt Publishing.
- Sharma, M., Tong, M., Korbak, T., Duvenaud, D., Asbell, A., Bowman, S. R., Cheng, N., Durmus, E., Hatfield-Dodds, Z., Johnston, S. R., Kravec, S., Maxwell, T., McCandlish, S., Ndousse, K., Rausch, O., Schiefer, N., Yan, D., Zhang, M., & Perez, E. (2023). Towards understanding sycophancy in language models. *arXiv*. <https://doi.org/10.48550/arXiv.2310.13548>
- Stiegler, B. (2013). *What makes life worth living: On pharmacology* (D. Ross, Trans.). Polity.
- WS, G. (1908). The probable error of a mean. *Biometrika*, 6(1), 1–25. <https://doi.org/10.2307/2331554>
- Vaccaro, M., Almaatouq, A., & Malone, T. (2024). When combinations of humans and AI are useful: A systematic review and meta-analysis. *Nature Human Behavior*, 8, 2293–2303. <https://doi.org/10.1038/s41562-024-02024-1>
- Verbeek, P.-P. (2011). *Moralizing technology: Understanding and designing the morality of things*. University of Chicago Press.
- Wei, J., Huang, D., Lu, Y., Zhou, D., & Le, Q. V. (2023). Simple synthetic data reduces sycophancy in large language models. *arXiv*. <https://doi.org/10.48550/arXiv.2308.03958>
- White House Office of Science and Technology Policy. (2022). *Blueprint for an AI Bill of Rights: Making automated systems work for the American people*. The White House. <https://www.whitehouse.gov/ostp/ai-bill-of-rights/>
- Wright, D., Masud, S., Moore, J., Yadav, S., Antoniak, M., Park, C. Y., & Augenstein, I. (2025). Epistemic diversity and knowledge collapse in large language models. *arXiv*. <https://doi.org/10.48550/arXiv.2510.04226>
- Yang, Y., Zhang, Y., Sun, D., He, W., & Wei, Y. (2025). Navigating the landscape of AI literacy education: Insights from a decade of research (2014–2024). *Humanities and Social Sciences Communications*, 12, 374. <https://doi.org/10.1057/s41599-025-04583-8>
- Zhai, C., Wibowo, S., & Li, L. D. (2024). The effects of over-reliance on AI dialogue systems on students' cognitive abilities: A systematic review. *Smart Learning Environments*, 11, 28. <https://doi.org/10.1186/s40561-024-00316-7>
- Zimmerman, B. J. (2000). Attaining self-regulation: A social cognitive perspective. In M. Boekaerts, P. Pintrich, & M. Zeidner (Eds.), *Handbook of self-regulation* (pp. 13–39). Academic Press.