

and radiology; and (♠) search and knowledge discovery—driven by large language models, natural language processing, and retrieval-augmented generation. The thematic map positioned the operational applications cluster (cluster ②) in the first quadrant (motor themes), indicating its high centrality and density—reflecting a mature, well-established research domain. The post-publication impact assessment cluster (cluster ③) was situated at the boundary between the first and third quadrants, suggesting that while the field is gaining centrality and relevance, it has not yet achieved sufficient internal cohesion and structural consolidation. The medical and healthcare cluster (cluster ④) appeared in the second quadrant (highly developed but isolated), indicating its specialized nature and limited integration with other thematic clusters. The search and knowledge discovery cluster (cluster ♠) was positioned in the fourth quadrant (basic and emerging themes), reflecting its rapid growth following the emergence of ChatGPT and other generative AI tools in late 2022. Notably, the ethical management cluster (cluster ①) was located in the lower-left quadrant (emerging or declining themes), revealing a critical paradox: while ethical discourse is highly frequent in the literature, it remains structurally underdeveloped—implying that principles are broadly accepted, yet practical implementation mechanisms are still immature. Temporal analysis of keyword dynamics further indicated a significant shift in research focus: during 2023–2025, scholarly attention moved from purely technical exploration toward normative regulation, policy formulation, and governance-oriented inquiry.

Conclusion: The findings reveal that the field of responsible AI in scholarly publishing has undergone a fundamental transformation from a technical-innovative concern to a governance-ethical paradigm. Five distinct conceptual clusters were identified, each representing a different layer of the scholarly publishing ecosystem: meta-governance (ethical principles and research integrity), operational-functional (data analysis, writing, peer review), post-publication evaluation (bibliometric and scientometric assessment), domain-specific applications (medicine and healthcare), and knowledge discovery (large language models and retrieval systems). A critical paradox emerged from the analysis: while ethical principles such as transparency, accountability, and trustworthiness are widely discussed and broadly accepted, their practical implementation mechanisms remain structurally underdeveloped. The thematic map positioned the ethical management cluster in the emerging or declining quadrant, indicating that despite high frequency in the literature, operational frameworks for enforcing these principles have not yet matured. This gap between accepted principles and enforceable mechanisms represents the central

challenge facing the scholarly publishing community. Furthermore, the analysis revealed a significant operational-normative disconnect: the operational applications cluster, while mature and well-established, shows limited structural linkage with the ethical management cluster. This suggests that AI tools are being extensively integrated into research workflows without adequate transparency, disclosure, or accountability mechanisms—a finding with important implications for research integrity. Similarly, the emergence of large language models in knowledge discovery highlights those ethical risks begin at the earliest stages of the research cycle, necessitating transparency requirements that extend beyond writing and publication. These findings suggest that future governance frameworks should adopt a layered, ecosystem-based approach that addresses normative, operational, and contextual dimensions simultaneously. Such frameworks must emphasize mandatory transparency and disclosure, discipline-sensitive policy differentiation, and systematic enhancement of AI literacy among all stakeholders.

Keywords: Artificial intelligence; Scientific publishing; Research integrity; Bibliometric analysis; AI governance; Responsible AI; Science mapping.

زود آید و دایره اش نشانه

*

,

مقدمه و بیان مسئله

طی دهه اخیر، تحول دیجیتال در علم با ظهور و توسعه سریع فناوری‌های مبتنی بر هوش مصنوعی، به‌ویژه مدل‌های زبانی بزرگ^۱، وارد مرحله‌ای کیفی جدید شده است. این فناوری‌ها با قابلیت‌هایی همچون تولید متن علمی، خلاصه‌سازی ادبیات، تحلیل داده‌ها و پشتیبانی از تصمیم‌گیری‌های علمی، به سرعت در حال ادغام با زیرساخت‌های زیست‌بوم پژوهش و نشر علمی هستند (Al- (Qudimat et al., ۲۰۲۵; Celik, ۲۰۲۵). ظهور ابزارهایی مانند ChatGPT، نه تنها فرآیند نگارش علمی را تسهیل کرده، بلکه دامنه تأثیر هوش مصنوعی را به مراحل مختلف چرخه حیات نشر، از جستجوی ادبیات و تولید دانش تا داوری، انتشار و ارزیابی تأثیر علمی، گسترش داده است. این تحول، نویدبخش افزایش کارایی، سرعت و دسترس‌پذیری در تولید و انتشار دانش علمی است (Daskaliuk et al., ۲۰۲۵; Upshall, ۲۰۱۹).

با این حال، گسترش سریع استفاده از هوش مصنوعی در نشر علمی، نگرانی‌های قابل توجهی را نیز به همراه داشته است. پژوهش‌ها نشان داده‌اند که استفاده از این فناوری‌ها می‌تواند با چالش‌هایی همچون تولید اطلاعات نادرست یا ساختگی، سوگیری الگوریتمی، ابهام در مسئولیت‌پذیری علمی و تهدید یکپارچگی علمی همراه باشد (Ateriya et al., ۲۰۲۵; Carobene et al., ۲۰۲۳; Javanbakht, ۲۰۲۴). به‌ویژه، قابلیت مدل‌های زبانی در تولید متون علمی شبه‌انسانی، پرسش‌های بنیادینی درباره اصالت علمی، نقش نویسندگان و اعتبار نظام داوری هم‌اکنون مطرح ساخته است. برآوردها نشان می‌دهد که ده‌ها هزار مقاله علمی در سال‌های اخیر با کمک مدل‌های زبانی تولید شده‌اند، که بیانگر نفوذ گسترده این فناوری در نظام تولید دانش است (Gray, ۲۰۲۴; Liang et al., ۲۰۲۴).

در پاسخ به این تحولات، مجلات علمی و ناشران بین‌المللی تلاش کرده‌اند با تدوین سیاست‌ها و دستورالعمل‌هایی، چارچوب‌هایی برای استفاده مسئولانه از هوش مصنوعی ارائه دهند. با این حال، شواهد نشان می‌دهد که این سیاست‌ها با ناهمگونی قابل توجهی همراه هستند و فاقد استانداردسازی و انسجام کافی هستند (Ganjavi et al., ۲۰۲۴; Leung et al., ۲۰۲۳). تفاوت در الزامات افشا، دامنه کاربرد مجاز و نحوه تنظیم مسئولیت‌ها، نه تنها موجب ابهام برای پژوهشگران شده، بلکه اثربخشی این سیاست‌ها را نیز محدود ساخته است. این وضعیت، نشان‌دهنده نبود یک چارچوب یکپارچه و مبتنی بر شواهد برای حکمرانی استفاده از هوش مصنوعی در نشر علمی است. نهادهای سیاست‌گذار باید به‌طور نظام‌مند، تأثیرات هوش مصنوعی را بر فعالیت‌های علمی روزمره، از جمله در تیم‌سازی انسان و هوش مصنوعی، کار، مسیرهای شغلی و آموزش ارزیابی کنند؛ نقطه‌ای که در آن، تغییرات مهمی ممکن است رخ دهد.

مرور مطالعات موجود نشان می‌دهد که تمرکز غالب پژوهش‌ها بر دو حوزه اصلی بوده است: نخست، نقش هوش مصنوعی در نگارش علمی، از جمله تأثیر آن بر سبک نوشتار، ساختار مقالات و فرآیند تولید متن (Cheng et al., ۲۰۲۴; Masukume, ۲۰۲۴) و دوم، تحلیل سیاست‌ها و دستورالعمل‌های ناشران در مواجهه با این فناوری (Ganjavi et al., ۲۰۲۳). درواقع، پژوهش‌ها اغلب به‌صورت جزیره‌ای به ابعاد مختلف پدیده نگریند (Ganjavi et al., ۲۰۲۴; Leung et al., ۲۰۲۳) و عمدتاً بر جنبه‌های محدودی همچون نگارش یا سیاست‌های ناشران متمرکز بوده‌اند، بدون آنکه پیوندهای مفهومی میان آن‌ها را آشکار سازند. تاکنون مطالعه‌ای که با رویکردی سیستمی و با استفاده از ابزارهای علم‌سنجی، ساختار مفهومی، شبکه موضوعی و پیوستگی میان ابعاد مختلف کاربرد هوش مصنوعی در زیست‌بوم نشر را ترسیم و واکاوی کند و بر این اساس، چارچوبی یکپارچه برای سیاست‌گذاری ارائه دهد، انجام نشده است. پژوهش حاضر با بهره‌گیری از ابزارهای تحلیل علم‌سنجی، به دنبال پر کردن این شکاف ساختاری است تا

۱. Large Language Models

نشان دهد چگونه خوشه‌های موضوعی پراکنده، در عمل به یکدیگر متصل شده‌اند. مدل سه‌لایه-پنج‌خوشه‌ای مستخرج از این پژوهش، سازوکارهای اتصال این خوشه‌ها را در قالب لایه‌های فراحاکمیتی، کارکردی و ارزیابی ترسیم می‌کند. بنابراین، شکاف اصلی در ادبیات موجود، فقدان یک تحلیل جامع، نظام‌مند و مبتنی بر شواهد کتاب‌سنجی است که بتواند ابعاد مختلف کاربرد هوش مصنوعی در نشر علمی (شامل جنبه‌های هنجاری، کارکردی و زمینه‌ای) را به‌صورت یکپارچه و با رویکردی سیستمی پوشش دهد. چنین تحلیلی می‌تواند زمینه‌ساز ارائه چارچوبی منسجم برای سیاست‌گذاری مسئولانه در این حوزه باشد. بر این اساس، پرسش اصلی پژوهش حاضر به این شرح مطرح می‌شود: ساختار مفهومی و موضوعی مطالعات به‌کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی چگونه است؟

پرسش‌های پژوهش

براساس پرسش کلی پژوهش و با هدف پر کردن شکاف‌های پژوهشی شناسایی‌شده، پرسش‌های ویژه این پژوهش به شرح زیر تعریف می‌شوند:

۱. شبکه هم‌واژگانی مطالعات به‌کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی شامل چه خوشه‌هایی است؟
۲. روند تحولات و پویایی موضوعات پژوهش‌های به‌کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی در طول سال‌های ۲۰۲۰ تا ۲۰۲۵ چگونه است؟
۳. روند راهبردی تکامل موضوعی مطالعات به‌کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی براساس سنجه‌های مرکزیت و چگالی به چه صورت است؟
۴. ساختار مفهومی مطالعات به‌کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی براساس تحلیل عاملی (تحلیل تناظر چندگانه) چگونه است؟

چارچوب نظری

نظریه تحول دیجیتال علم و پارادایم پژوهش تقویت‌شده با هوش مصنوعی^۱

تحول دیجیتال علم، یکی از مهم‌ترین دگرگونی‌های نظام تولید دانش در قرن بیست‌ویکم به شمار می‌آید. در این چارچوب، هوش مصنوعی نه تنها ابزار فناوری، بلکه یک عامل ساختاری در بازتعریف فرآیندهای پژوهش، تولید دانش و ارتباطات علمی تلقی می‌شود (Wang et al., ۲۰۲۳). ظهور نظام‌های مبتنی بر یادگیری ماشین و به‌ویژه مدل‌های زبانی بزرگ مانند ChatGPT، امکان خودکارسازی یا تقویت بسیاری از فعالیت‌های شناختی پژوهشگران، از جمله جستجوی ادبیات، تولید فرضیه، تحلیل داده‌ها و نگارش علمی را فراهم کرده است. (Al-Qudimat et al., ۲۰۲۵; Khalifa & Albadawy, ۲۰۲۴)

در چارچوب نظری «پژوهش تقویت‌شده با هوش مصنوعی»، هوش مصنوعی به‌عنوان یک «عامل شناختی مکمل» عمل می‌کند که ظرفیت‌های شناختی انسان را تقویت می‌کند، نه اینکه جایگزین آن شود. این رویکرد براساس نظریه شناخت توزیع‌شده^۲ است که بیان می‌کند فرآیندهای شناختی در تعامل میان انسان، ابزارها و محیط شکل می‌گیرند. (Checco et al., ۲۰۲۱) در این چارچوب، هوش مصنوعی مولد می‌تواند به‌عنوان یک «دستیار شناختی»^۳ عمل کند که وظایفی مانند ترکیب اطلاعات، استخراج الگوها و تولید پیش‌نویس‌های اولیه را تسهیل می‌کند.

مطالعات نشان داده‌اند که هوش مصنوعی مولد می‌تواند در مراحل مختلف چرخه پژوهش، از تولید ایده‌های تحقیقاتی تا تحلیل داده‌ها و نگارش مقاله، نقش مؤثری ایفا کند (Celik, ۲۰۲۵). برای مثال، این فناوری می‌تواند با تحلیل حجم عظیمی از متون علمی،

۱. AI-Augmented Research Paradigm

۲. Distributed Cognition

۳. Cognitive Assistant

شکاف‌های پژوهشی را شناسایی کرده یا فرضیه‌های جدیدی پیشنهاد دهد. (Dos Santos et al., ۲۰۲۳) همچنین، در مرحله نگارش، مدل‌های زبانی می‌توانند ساختار منطقی متن را بهبود بخشند، خطاهای زبانی را اصلاح و فرآیند تولید متن علمی را تسریع کنند (Onuoha, ۲۰۲۵). در مرحله پس از انتشار نیز، هوش مصنوعی می‌تواند در نمایه‌سازی، تحلیل استنادات و ارزیابی تأثیر علمی نقش ایفا کند.

با این حال، نقش هوش مصنوعی در پژوهش، به شدت وابسته به سطح پیچیدگی وظیفه است. مطالعات نشان داده‌اند که هوش مصنوعی در وظایف ساختاریافته مانند خلاصه‌سازی، دسته‌بندی و استخراج اطلاعات عملکرد بالایی دارد، اما در وظایفی که نیازمند قضاوت علمی پیچیده، ارزیابی نوآوری، یا استدلال مفهومی هستند، همچنان به نظارت انسانی نیاز دارد (Khlaif et al., ۲۰۲۳). (Kousha & Thelwall, ۲۰۲۴) این موضوع نشان می‌دهد که مدل بهینه تعامل، مدل «همکاری انسان-هوش مصنوعی» است، نه جایگزینی کامل انسان. بنابراین، براساس این چارچوب نظری، هوش مصنوعی مولد را می‌توان به عنوان یک فناوری توانمندساز^۱ در نظر گرفت که ساختار، سرعت و کارایی تولید دانش علمی را تغییر می‌دهد، اما همچنان در چارچوب نظارت انسانی عمل می‌کند.

رویکردهای لایه‌ای و بوم‌شناختی به حکمرانی هوش مصنوعی

مطالعات همچنین نشان داده‌اند که استفاده از هوش مصنوعی در چرخه نشر با چالش‌هایی همراه است، از جمله:

- تولید اطلاعات نادرست یا ساختگی
- سوگیری الگوریتمی
- کاهش شفافیت
- ابهام در مسئولیت علمی (Carobene et al., ۲۰۲۳; Gurnal & Rana, ۲۰۲۵).

این چالش‌ها نشان می‌دهد که ادغام هوش مصنوعی در چرخه نشر، نیازمند چارچوب‌های حکمرانی و سیاست‌گذاری مناسب است. رویکردهای نوین در ادبیات حکمرانی هوش مصنوعی نیز بر این نکته تأکید دارند که استفاده از هوش مصنوعی در نشر علمی، مستلزم تحلیل در چندین لایه و با توجه به زمینه است.

الف) همگرایی اصول اخلاقی و ضرورت لایه‌بندی هنجاری

پژوهش جوبین و همکاران (Jobin et al., ۲۰۱۹) با تحلیل ۸۴ سند راهنمای اخلاقی هوش مصنوعی از سراسر جهان نشان داد که علی‌رغم تنوع فرهنگی و نهادی، همگرایی جهانی حول پنج اصل اخلاقی شکل گرفته است: شفافیت، عدالت و انصاف، عدم ضرر، مسئولیت‌پذیری و حریم خصوصی. با این حال، آن‌ها همچنین دریافتند که واگرایی قابل‌توجهی در تفسیر، توجیه و شیوه اجرای این اصول وجود دارد. این یافته نشان می‌دهد که وجود اصول مشترک برای حکمرانی کافی نیست، بلکه نیاز به چارچوب‌های لایه‌ای است که بتواند سطوح مختلف تصمیم‌گیری (از ارزش‌های هنجاری تا رویه‌های اجرایی) را پوشش دهد. در همین راستا، فلورییدی و همکاران (Floridi et al., ۲۰۱۸) در چارچوب AI4People، پنج اصل اخلاقی را مبنا قرار داده و ۲۰ توصیه عملی برای سیاست‌گذاران ارائه دادند. این توصیه‌ها طیفی از ارزیابی و توسعه فناوری تا ایجاد مشوق‌ها و حمایت از کاربردهای خوب هوش مصنوعی را شامل می‌شود. اخیراً، لی (Lee, ۲۰۲۶) نظریه «پیمان اختیار^۲» را در چهار لایه معرفی کرده است: اختیار هنجاری (ارزش‌ها)، اختیار معرفتی (معیارهای شواهد)، اختیار منبع (اعتبار منابع) و اختیار داده (گزینش داده). این رویکرد نشان می‌دهد که حکمرانی مسئولانه هوش مصنوعی مستلزم شفافیت و قابلیت حسابرسی نه فقط در سطح خروجی‌ها بلکه در همه لایه‌ها است. این

۱. Enabling Technology

۲. Authority Stack

لایه‌بندی می‌تواند مبنایی نظری برای تحلیل جایگاه مفاهیمی همچون «صداقت آکادمیک»، «شفافیت» و «پاسخگویی» در فرایند نشر علمی فراهم آورد.

(ب) رویکرد بوم‌شناختی جامعه‌شناختی-فنی

در مقابل رویکردهای صرفاً هنجاری، هرناندز و همکاران (Domínguez Hernández et al., ۲۰۲۵) مفهوم «بوم‌شناسی جامعه‌شناختی-فنی هوش مصنوعی»^۱ را پیش نهاده‌اند. این رویکرد با الهام از «وارونه‌سازی زیرساختی» در مطالعات اجتماعی فناوری، هوش مصنوعی را شبکه‌ای در هم تنیده از کنشگران، رویه‌ها، جریان‌های منابع (داده، سرمایه، نیروی انسانی) و روابط قدرت می‌داند. از این منظر، سیاست‌گذاری برای کاربرد مسئولانه هوش مصنوعی در نشر علمی نمی‌تواند به تدوین اصول کلی اکتفا کند، بلکه باید زمینه‌های خاص (مانند تفاوت میان رشته پزشکی و علوم انسانی) و لایه‌های زیرساختی (مانند استخراج داده و سوگیری الگوریتمی) را نیز در نظر گیرد.

حکمرانی هوش مصنوعی و مدل تعامل انسان-هوش مصنوعی^۲

یکی از مهم‌ترین چارچوب‌های نظری برای تحلیل نقش هوش مصنوعی در پژوهش، نظریه «هوش ترکیبی»^۳ است. این نظریه بیان می‌کند که بهترین عملکرد، زمانی حاصل می‌شود که انسان و هوش مصنوعی به‌صورت مکمل عمل کنند، به‌طوری‌که هوش مصنوعی وظایف تکراری و پردازشی را انجام داده و انسان، مسئول قضاوت علمی و تصمیم‌گیری باشد (Salih et al., ۲۰۲۵). در این چارچوب، هوش مصنوعی به‌عنوان یک نظام پشتیبان تصمیم‌گیری^۴ عمل می‌کند، نه یک عامل تصمیم‌گیر مستقل. مطالعات نشان داده‌اند که این مدل می‌تواند کارایی فرآیندهای نشر را تا ۳۰-۴۰ درصد افزایش دهد، درحالی‌که کیفیت علمی حفظ می‌شود (Dorskaliuk et al., ۲۰۲۵). با این حال، به‌کارگیری مسئولانه از هوشواره‌ها نیازمند چارچوب‌های حکمرانی است که شامل موارد زیر باشد:

- شفافیت در استفاده از هوش مصنوعی
- مسئولیت‌پذیری پژوهشگران
- نظارت انسانی
- کنترل کیفیت خروجی‌ها (Leung et al., ۲۰۲۳).

این اصول، مبنای چارچوب‌های اخلاقی استفاده از هوش مصنوعی در نشر علمی را تشکیل می‌دهند. براساس چارچوب‌های نظری فوق، می‌توان نتیجه گرفت که هوش مصنوعی مولد، یک فناوری تحول‌آفرین است که تمامی مراحل چرخه پژوهش و نشر را تحت تأثیر قرار می‌دهد. تأثیر هوش مصنوعی در مراحل مختلف چرخه پژوهش، متفاوت است و به پیچیدگی وظیفه بستگی دارد. مدل بهینه استفاده از هوش مصنوعی، مدل همکاری انسان-هوش مصنوعی بوده و استفاده مؤثر از هوش مصنوعی، نیازمند چارچوب‌های حکمرانی، شفافیت و نظارت انسانی است.

پیشینه پژوهش

هوش مصنوعی، به‌ویژه مدل‌های زبانی بزرگ در مدت زمان کوتاهی به یکی از موضوعات محوری و بحث‌برانگیز در زیست‌بوم انتشارات علمی تبدیل شده‌اند. سرعت پذیرش این فناوری‌ها به حدی بوده است که موجی از مطالعات کتاب‌سنجی و علم‌سنجی را برای شناسایی، کمی‌سازی و تحلیل ابعاد مختلف آن به راه انداخته است. مرور نظام‌مند این مطالعات می‌تواند تصویری روشن از

۱. sociotechnical ecology of AI

۲. Human-AI Governance and Hybrid Intelligence Theory

۳. Hybrid Intelligence

۴. Decision Support System

وضعیت موجود، روندها و مهم‌تر از آن، خلأهای پژوهشی پیش رو ترسیم کند. تلاش می‌شود تا پژوهش‌های انجام‌شده در چارچوب مراحل پنج‌گانه چرخه حیات نشر (تحقیق، نگارش، داوری، انتشار و پس از انتشار) سازماندهی و ارائه شوند.

۱. نفوذ و ردپای هوش مصنوعی در مرحله نگارش علمی

بیشترین تمرکز مطالعات کتاب‌سنجی بر روی مرحله نگارش، متمرکز بوده است. این مطالعات با به کارگیری روش‌های گوناگون، از تحلیل کلمات کلیدی گرفته تا الگوریتم‌های پیشرفته تشخیص متن، به دنبال برآورد میزان نفوذ متون تولیدشده یا ویرایش‌شده توسط هوش مصنوعی بوده‌اند.

لیانگ و همکاران (Liang et al., ۲۰۲۴) با انجام اولین تحلیل کلان و نظام‌مند بر روی نزدیک به یک میلیون مقاله از پلتفرم‌های bioRxiv, arXiv و مجلات Nature، نشان دادند که استفاده از مدل‌های زبانی بزرگ در نگارش علمی از اواخر سال ۲۰۲۲ روندی صعودی و شتابان یافته است. به طور مشخص، بیشترین میزان استفاده در مقالات علوم رایانه (تا ۱۷/۵ درصد) و کمترین آن در مقالات ریاضی (تا ۶/۳ درصد) مشاهده شد. چنگ و همکاران (Cheng et al., ۲۰۲۴) نیز با بهره‌گیری از یک ابزار تشخیص پیشرفته، ضمن تأیید این نفوذ گسترده در سکوها پیش‌نشر، به یافته‌های ظریف‌تری دست یافتند. ایشان نشان دادند که میزان استفاده از هوش مصنوعی با روند جستجوهای ChatGPT همبستگی داشته، در میان نویسندگان با پیشینه زبانی و جغرافیایی مختلف متفاوت است (بیش از ۵ درصد در مقالات نویسندگان ایتالیایی و چینی) و بیشترین کاربرد آن در بخش‌های «تنظیم فرضیه» و «نتیجه‌گیری» مقالات است. در همین راستا، مطالعه کتاب‌سنجی مانگوبات و سعیدی (Mangubat & Saeedi, ۲۰۲۵) با تحلیل داده‌های پایگاه اسکوپوس در بازه ۲۰۱۵ تا ۲۰۲۵، تصویری کلان از روندهای نوظهور در این حوزه ارائه می‌دهد. یافته‌های آنان نشان می‌دهد که نرخ رشد سالانه انتشارات مرتبط با کاربرد هوش مصنوعی در نگارش پژوهشی دانشجویان ۳۲/۷۵ درصد بوده و کشورهایی مانند چین و ایالات متحده با سرمایه‌گذاری قابل توجه، پیشتاز تولید علم در این عرصه هستند. تحلیل هم‌رخدادی واژگان در این مطالعه، خوشه‌های موضوعی متمرکز بر ابزارهای هوش مصنوعی (از جمله ChatGPT، Grammarly و Quillbot)، «نگارش آکادمیک»، «دانشجویان» و «آموزش عالی» را آشکار ساخته که مؤید پذیرش فزاینده این فناوری‌ها در زیست‌بوم‌های دیجیتال آموزشی است. این یافته‌ها، همسو با مطالعات پیشین (Cheng et al., ۲۰۲۴; Liang et al., ۲۰۲۴)، بر ضرورت توجه سیاست‌گذاران آموزشی به تدوین چارچوب‌های شفاف برای به کارگیری مسئولانه هوشواره‌ها در فرایند نگارش پژوهشی تأکید می‌ورزد.

رویکرد دیگری که در این حوزه مورد استفاده قرار گرفته، تحلیل فراوانی واژگان شاخص است. گری (Gray, ۲۰۲۴) با استفاده از این روش، برآورد کرد که حداقل ۶۰،۰۰۰ مقاله (بیش از یک درصد از کل مقالات) در سال ۲۰۲۳ با کمک مدل‌های زبانی بزرگ نوشته شده‌اند. در مطالعه‌ای مشابه اما با ظرافت بیشتر، ماسوکومه (Masukume, ۲۰۲۴) نشان داد که فراوانی همزمان کلماتی مانند "delve"، "realm" و "underscore" که به عنوان نشانگان استفاده از چت‌بات‌های هوش مصنوعی شناخته می‌شوند، در سال‌های ۲۰۲۳ و ۲۰۲۴ نسبت به تمام دوره ۱۷۵ ساله پیش از آن، افزایش چشمگیر (تا ۸۵ برابر) داشته است. این یافته‌ها همگی بر تغییر سبک و محتوای نگارش علمی در نتیجه استفاده از این ابزارها صحنه می‌گذارند. میلر و همکاران (Miller et al., ۲۰۲۳) حتی پیش از همه‌گیر شدن ChatGPT نیز روندی افزایشی را در متون با احتمال بالای تولید توسط هوش مصنوعی در چکیده‌های نمایه‌شده در مدلاین از سال ۲۰۲۰ تا ۲۰۲۳ شناسایی کردند.

در کنار این تحلیل‌های کمی، پژوهش‌هایی نیز به بررسی پیامدهای کیفی این پدیده پرداخته‌اند. مدینا و همکاران (Medina et al., ۲۰۲۴) در مرور کتاب‌سنجی خود، پیامدهای استفاده از ChatGPT را در سه حوزه «کیفیت و توسعه مهارت‌های نوشتاری»، «یکپارچگی و اخلاق علمی» و «ارزشیابی آموزشی» دسته‌بندی کردند. مولینه و همکاران (Mulyanah et al., ۲۰۲۴) نیز در مرور نظام‌مند خود، شش محدودیت حیاتی برای ابزارهای دستیار نویسنده هوش مصنوعی احصا کردند که از آن جمله می‌توان به

عدم دقت علمی، ناتوانی در سنتز ایده‌های پیچیده و اصالت مشکوک اشاره نمود. این یافته‌ها با دیدگاه‌های کیفی دانشجویان دکتری در مطالعه عاد (۲۰۲۴, Ead) همخوانی دارد که بر مفید بودن این ابزارها در کنار خطر وابستگی بیش از حد و تضعیف تفکر انتقادی تأکید کرده‌اند.

۲. کاربردهای نوظهور هوش مصنوعی در مرحله تحقیق

اگرچه حجم پژوهش‌ها در این بخش کمتر است، اما شواهدی از کاربرد فزاینده هوش مصنوعی در مرحله تحقیق، به‌ویژه در زمینه جستجوی ادبیات، تحلیل داده و ارزیابی خودکار نگارش، مشاهده می‌شود. مطالعات کتاب‌سنجی ناکرووی و همکاران (Nakrowi, ۲۰۲۵) و دیودی و الوری (Dwivedi & Elluri, ۲۰۲۴) به موضوعات نوظهوری مانند ChatGPT، مدل‌های زبانی بزرگ و کاربردهای هوش مصنوعی مولد در حوزه‌های مختلف علمی اشاره دارند که بر ادغام این فناوری در لایه‌های زیرین فرآیند تحقیق دلالت دارد. در سطحی عملی‌تر، مطالعات لیتوینوا و همکاران (Litvinova et al., ۲۰۲۴) و شوئه (Xue, ۲۰۲۴) به حوزه ارزیابی خودکار نوشتار پرداخته‌اند. شوئه با رویکرد فراتحلیل نشان داد که ابزارهای ارزیابی خودکار نوشتار می‌توانند به‌طور معناداری به بهبود مهارت نوشتاری کمک کنند که بر نقش پررنگ آن‌ها در بازخورددهی و بهبود فرآیند تحقیق تأکید دارد. با این حال، به نظر می‌رسد که کاربرد هوش مصنوعی در مراحل اکتشافی و تحلیلی تحقیق، نسبت به کاربرد آن در نگارش، کمتر مورد کاوش قرار گرفته است.

۳. سیاست‌گذاری انتشاراتی در مواجهه با هوش مصنوعی: مرحله انتشار

یکی از حیاتی‌ترین و در عین حال چالش‌برانگیزترین حوزه‌ها، نحوه مواجهه ناشران و مجلات علمی با این فناوری است که ذیل مرحله انتشار، تعریف می‌شود. مطالعه شاخص و دوجانبه گنجوی و همکاران (Ganjavi et al., ۲۰۲۴) تصویری نگران‌کننده از وضعیت سیاست‌گذاری در این حوزه ترسیم کرده است. یافته‌های این پژوهش نشان می‌دهد که علیرغم افزایش آگاهی، هنوز بسیاری از ناشران برتر (حدود ۷۶ درصد) فاقد راهنمایی مدونی برای نویسندگان در مورد استفاده از هوش مصنوعی مولد هستند. در مقابل، مجلات برتر وضعیت بهتری دارند (۸۷ درصد دارای راهنمود). نکته حائز اهمیت، ناهمگونی شدید در محتوای این راهنمودهاست. اگرچه تقریباً همه آن‌ها (بیش از ۹۵ درصد) قید هوش مصنوعی به عنوان نویسنده را ممنوع کرده‌اند، اما در مورد حوزه‌های مجاز استفاده، الزام به افشا، محل دقیق افشا (روش‌شناسی، بخش تشکر و قدردانی، یا بخشی جدید) و نحوه ارجاع به این ابزارها، تنوع و گاه حتی تضاد (میان مجله و ناشر مادر) مشاهده می‌شود. گنجوی و همکاران (Ganjavi et al., ۲۰۲۳) این وضعیت را عدم استانداردسازی توصیف کرده و بر ضرورت تدوین راهنمودهای یکپارچه برای حفظ یکپارچگی خروجی‌های علمی تأکید می‌ورزند. مرور نظام‌مند مرتضوی شاهرودی و زارعی (۱۴۰۴) نیز نگرانی‌های حقوقی مرتبط با مالکیت معنوی و حق چاپ را به‌عنوان یکی از چالش‌های فراروی این حوزه یادآور شده است. احمدی و جمشیدی (۱۴۰۴) نیز با رویکرد علم‌سنجی و روش تحلیل هم‌اژگانی مقالات مرتبط با این حوزه، به این نتیجه دست یافتند که استفاده از هوش مصنوعی در نگارش علمی با رعایت ملاحظات اخلاقی قابل پذیرش است. اگرچه هوش مصنوعی نباید به عنوان نویسنده شناخته شود، اما نقش آن در توسعه اثر نیازمند شفافیت و ارزیابی مستند است. نویسندگان موظفند محتوای تولید شده را بازنویسی اساسی کرده و منابع را به درستی استناد دهند. نظارت تخصصی انسانی برای پیشگیری از خطا و سوگیری ضروری است.

مرور پیشینه نشان می‌دهد که تمرکز پژوهش‌ها عمدتاً بر مراحل نگارش و انتشار معطوف بوده است. در مقابل، دو مرحله حیاتی دیگر از چرخه حیات نشر، یعنی داوری و پس از انتشار، کمتر مورد توجه قرار گرفته‌اند. در مورد کاربرد هوش مصنوعی در فرآیند داوری (مرحله ۳)، تنها اشاره‌ای گذرا در مطالعه مرتضوی شاهرودی و زارعی (۱۴۰۴) به «تسریع در فرایندهای داوری علمی» به‌عنوان یک فرصت دیده می‌شود. این در حالی است که پتانسیل‌های هوش مصنوعی برای کمک به ویراستاران در یافتن داوران مناسب، غربالگری اولیه مقالات، یا حتی شناسایی موارد سوءرفتار پژوهشی، حوزه‌ای بکر برای پژوهش است. به‌طور مشابه، تأثیرات بلندمدت هوش مصنوعی بر مرحله پس از انتشار، یعنی بر دیده‌شدن، تأثیرگذاری و شاخص‌های علم‌سنجی مقالات، مغفول مانده است. اشاره مرتضوی

شاهرودی و زارعی (۱۴۰۴) به خطر «افزایش سرعت تولید مقالات بی کیفیت» و روندهای کلان ترسیم شده توسط مارال (Maral, ۲۰۲۴) که آینده نگارش علمی را متأثر از هوش مصنوعی می‌داند، می‌تواند سرآغازی برای پژوهش در این حوزه باشد. این خلأ پژوهشی، به‌ویژه زمانی پررنگ‌تر می‌شود که بدانیم کیفیت و اعتبار ادبیات علمی در بلندمدت، مستقیماً بر شاخص‌های پس از انتشار، تأثیر می‌گذارد.

مرور نظام‌مند پژوهش‌های کتاب‌سنجی در حوزه هوش مصنوعی و انتشارات علمی، نشان‌دهنده رشد تصاعدی توجه به این حوزه است. مطالعات موجود با موفقیت به کمی‌سازی نفوذ هوش مصنوعی در مرحله نگارش (Cheng et al., ۲۰۲۴; Masukume, ۲۰۲۴) و مستندسازی ناهمگونی سیاست‌ها در مرحله انتشار (Ganjavi et al., ۲۰۲۴; Leung et al., ۲۰۲۳) پرداخته‌اند. با این حال، مقایسه این مطالعات نشان می‌دهد که توجه پژوهشی عمدتاً در یک بُعد متمرکز بوده و برخی ابعاد کمتر مورد توجه قرار گرفته‌اند. این عدم توازن، دو محور احتمالی نیازمند بررسی بیشتر را پیش روی پژوهش حاضر قرار می‌دهد:

۱. محور لایه‌ای-مرحله‌ای: مرور ادبیات حاکی از آن است که تمرکز پژوهش‌ها عمدتاً بر لایه عملیاتی (نگارش و تحلیل داده) بوده است، درحالی‌که لایه هنجاری (اصول اخلاقی، پاسخگویی) و لایه زمینه‌ای (تفاوت‌های حوزه‌ای) به‌نظر می‌رسد به‌صورت مجزا و بدون پیوند سازوکاری با یکدیگر مطالعه شده‌اند. در صورتی که این گسست لایه‌ای در پژوهش حاضر تأیید شود، این امر می‌تواند مانعی برای طراحی سیاست‌هایی باشد که همزمان شفافیت هنجاری، کارایی عملیاتی و حساسیت زمینه‌ای را تأمین کنند.

۲. محور یکپارچگی بوم‌شناختی: مطالعات بررسی‌شده عمدتاً به‌صورت جزیره‌ای به پدیده نگاشته‌اند. به‌نظر می‌رسد مطالعه‌ای که با نگاهی سیستمی-بوم‌شناختی، تعامل میان کنشگران (نویسندگان، داوران، ناشران)، جریان‌های منابع (داده، اعتبار، نیروی کار) و لایه‌های حکمرانی را در کل چرخه نشر تحلیل کند و بر این اساس، چارچوبی سیاستی ارائه دهد، در ادبیات موجود یافت نشد. پژوهش حاضر این احتمال را مطرح می‌کند که فقدان نگاه بوم‌شناختی، یک خلأ قابل توجه در حوزه حکمرانی هوش مصنوعی در نشر علمی باشد.

بر این اساس، پژوهش حاضر در پی آن است که ضمن تحلیل کتاب‌سنجی مطالعات مرتبط، ابتدا وجود این دو محور خلأ را به‌صورت نظام‌مند بررسی کند و در صورت تأیید، به تدوین چارچوبی بپردازد که نه تنها مراحل نشر، بلکه لایه‌های حکمرانی و بوم‌شناسی جامعه‌شناختی-فنی آن را نیز پوشش دهد.

روش‌شناسی پژوهش

پژوهش حاضر از نوع مطالعات کاربردی بوده و با رویکرد علم‌سنجی و تحلیل شبکه اجتماعی و هم‌رخدادی واژگان انجام شده است. منبع و ابزار اصلی استخراج داده‌های پژوهش، پایگاه استنادی وب آو ساینس است. جستجوی منابع مرتبط، در ۱۲ آوریل ۲۰۲۶ با اعمال راهبرد جستجوی مندرج در جدول ۱ در بخش جستجوی پیشرفته این پایگاه انجام شد.

در طراحی استراتژی جستجوی این پژوهش، برای اطمینان از جامعیت، استاندارد بودن و دقت واژگان، از مجموعه‌ای از اصطلاحنامه‌های معتبر و شناخته‌شده در حوزه‌های مرتبط استفاده شد. بدین منظور، از اصطلاحنامه MeSH (Medical Subject Headings) متعلق به کتابخانه ملی پزشکی آمریکا برای شناسایی واژگان استاندارد در حوزه‌های اخلاق پژوهش و ارزیابی تأثیرگذاری علمی؛ اصطلاحنامه Emtree (Elsevier Thesaurus) پایگاه Scopus برای واژگان کنترل‌شده در حوزه‌های هوش مصنوعی، نشر علمی و علم‌سنجی؛ اصطلاحنامه‌های IEEE Thesaurus و ACM Computing Classification System (CCS) برای شناسایی واژگان مرتبط با هوش مصنوعی و نشر علمی؛ اصطلاحنامه‌های علمی و فنی ایرانداک برای واژگان استاندارد انگلیسی در حوزه‌های علم‌سنجی و هوش مصنوعی؛ و همچنین واژگان کنترل‌شده پایگاه‌های Web of Science و Scopus در حوزه علم‌سنجی و کتاب‌سنجی استفاده گردید. این اصطلاحنامه‌ها به‌عنوان منابع استاندارد، در انتخاب واژگان جستجو، شناسایی مترادف‌ها و یکدست‌سازی کلیدواژه‌ها مورد استفاده قرار گرفتند. علاوه بر این، از نظرات

متخصصان و بررسی کلیدواژه‌های اختصاص یافته به مقالات مرتبط نیز برای تکمیل و به‌روزرسانی استراتژی جستجو استفاده شد. به منظور طراحی یک استراتژی جستجوی جامع تلاش بر این بود واژگانی انتخاب شوند که به‌صورت کامل فرایند نشر علمی را پوشش دهند؛ به‌گونه‌ای که نه‌تنها مراحل پیش از انتشار (نگارش علمی، داوری قبل از چاپ و فرایند انتشار)، بلکه ارزیابی‌ها و بازخوردهای پس از انتشار نیز در محدوده تحلیل قرار گیرند.

به منظور یکسان‌سازی واژگان جستجو، کنترل خطاهای ناشی از پوشش اشکال مختلف یک کلمه و حصول اطمینان از جامعیت جستجو از الگوهای ریشه‌یابی واژگان (Stemming) استفاده شد. بدین صورت که علامت (*) در پایان ریشه واژگانی که دارای اشکال مختلفی بودند (اسم، فعل، صفت و غیره)، قرار گرفت. این امر سبب می‌شود که پایگاه مورد جستجو به‌صورت خودکار تمام شکل‌های مشتق شده از آن واژه را بازیابی کند. به منظور غلبه بر خطاهای معنایی واژگان، از عملگر OR برای ترکیب مترادفات و معادل‌ها استفاده شد. پس از طراحی اولیه راهبرد جستجو، به منظور اطمینان از جامعیت آن و عدم بازیابی موارد نامرتبط، یک نمونه از نتایج (۱۰۰ رکورد اول) به‌صورت دستی بررسی و پس از ارزیابی، راهبرد نهایی تنظیم شد.

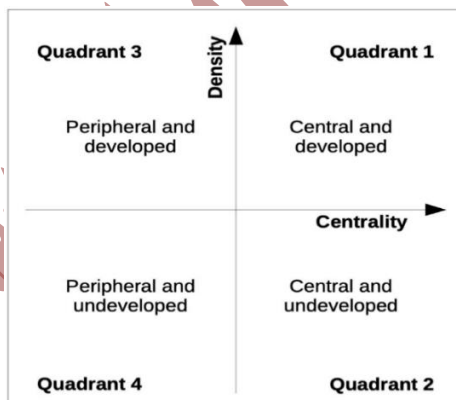
جدول ۱. راهبرد جستجو

منابع نهایی	نتایج بازیابی	کلیدواژه‌ها و محدودیت‌های اعمال شده	فیلد جستجو
۲۲۹۱۸	۲۴۳۱۱	("Artificial intelligence" OR "AI" OR "generative AI" OR "large language model*" OR "LLM") AND ("scientific publish*" OR "scholarly publish*" OR "research publish*" OR "academic publish*" OR "academic authorship" OR "scientific journal*" OR "scholarly journal*" OR "automated writing" OR "AI-assisted writing" OR "research design" OR "data analysis" OR "manuscript preparation" OR "scientific text production" OR "scientific writing" OR "academic writing" OR "scholarly writing" OR "research article*" OR "research writing" OR "writing process" OR "publish* process" OR "journal article*" OR "publication workflow*" OR "publication lifecycle*" OR "submission system*" OR "literature search*" OR "systematic review*" OR "literature review*" OR "evidence synthesis" OR "research workflow*" OR "research process" OR "research lifecycle*" OR "publication process" OR "scientific workflow*" OR "peer review*" OR "peer-review*" OR "review report" OR "review process" OR "reviewer recommendation" OR "journal polic*" OR "publisher polic*" OR "editorial workflow*" OR "editorial decision*" OR "journal management system*" OR "research evaluation" OR "research assessment" OR "citation analysis" OR "bibliometric*" OR "scientometric*" OR "altmetric*" OR "impact factor" OR "research metric*" OR "retract*" OR "research integrity" OR "correc* feedback" OR "post-publication*" OR "research ethic*" OR "publication ethic*" OR "academic integrity" OR "responsible use" OR "ethical implication*" OR "COPE" OR "guideline*"))	موضوع (TS)

منابع نهایی نتایج بازیابی	کلیدواژه‌ها و محدودیت‌های اعمال شده	فیلد جستجو
Refined By: Document Types: Article. Languages: English.		
Timespan: ۲۰۲۰-۰۱-۰۱ to ۲۰۲۵-۱۲-۳۱ (Index Date)		

با اعمال راهبرد جستجو تعداد ۲۴۳۱۱ مقاله بازیابی شد. معیارهای انتخاب منابع نهایی برای تحلیل عبارتند از: نوع منابع: مقالات پژوهشی، زبان: انگلیسی و بازه زمانی: ۲۰۲۵-۲۰۲۰. با حذف مقالات مروری، کنفرانسی، گزارش‌ها و کتاب‌ها، رکوردهایی که تکراری، مربوط به سال ۲۰۲۶ و یا دارای اطلاعات کتابشناختی ناقص بودند، در نهایت تعداد ۲۲۹۱۸ رکورد مبنای تحلیل نهایی قرار گرفت. به منظور غربالگری و حذف رکوردهای تکراری و دارای اطلاعات ناقص، فایل‌های استخراج شده از پایگاه‌ها را به محیط ویرایشگر نرم‌افزار Notepad++ منتقل و پس از انجام اصلاحات، مجدداً رکوردها به فرمت قابل قبول برای نرم‌افزارهای VOSviewer و پکیج Bibliometrix تبدیل و ذخیره شدند.

به منظور نگاشت ساختار موضوعی، سنجش بلوغ و تکامل و پویایی حوزه‌های پژوهشی به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی، به ترسیم نقشه موضوعی^۱ (نمودار راهبردی) پرداخته شد (شکل ۲). در این نمودار، از شاخص‌های تحلیل شبکه اجتماعی (مرکزیت^۲ و چگالی^۳) برای خوشه‌بندی کلمات کلیدی پرکاربرد مقالات و شناسایی کانون‌های پژوهشی استفاده شده که محور X مرکزیت و محور Y تراکم را نشان می‌دهد. بدین ترتیب، این نمودار شامل چهار بخش متفاوت با درجات مختلفی از تراکم و مرکزیت است. مرکزیت بالا نشان‌دهنده اهمیت بیش‌تر خوشه موضوعی در حوزه پژوهش است و تراکم و چگالی بیش‌تر نشان می‌دهد که خوشه موضوعی به خوبی توسعه یافته و بالغ‌تر است.



شکل ۲. چهار بخش از نقشه موضوعی یا نمودار راهبردی (Callon et al., ۱۹۹۱)

بخش اول نمودار، موضوعات پیشران (محرك)^۴ در ربع بالای سمت راست قرار دارند و نشان‌دهنده خوشه‌های بالغ با مرکزیت و چگالی بالا هستند. این مضامین کاملاً به پختگی و توسعه‌یافتگی رسیده‌اند و برای ساختار مفهومی و توسعه‌ی یک حوزه پژوهشی مهم هستند. این خوشه‌ها عمدتاً شامل موضوعاتی هستند که طی یک بازه طولانی توسط پژوهشگران در حال بررسی بوده‌اند. بخش دوم نمودار، موضوعات توسعه‌یافته و منفرد^۵ در ربع بالای سمت چپ قرار دارند و دارای خوشه‌هایی با چگالی بالا و مرکزیت کم هستند. بنابراین، مضامین مربوط به این بخش، علی‌رغم پختگی و توسعه‌یافتگی، منزوی و مجزا هستند و از اهمیت اندک یا حاشیه‌ای

^۱ Thematic map

^۲ Centrality

^۳ Density

^۴ Motor themes

^۵ Highly developed and isolated themes (Niche themes)

در حوزه پژوهشی مورد نظر برخوردارند. این موضوعات با وجود اینکه از واژگان تخصصی در حوزه‌ی مورد بررسی به شمار می‌روند، اما در آن رواج کم‌تری دارند. البته این خوشه‌ها قابلیت آن را دارند که در آینده وارد ناحیه پیشران شوند. بخش سوم نمودار شامل موضوعات پایه و بنیادین^۱ بوده که در ربع پایین سمت راست نقشه قرار می‌گیرند و حاوی کانون‌های پژوهشی با مرکزیت بالا، اما چگالی اندک هستند. این بخش موضوعات اولیه، پایه و عمومی را نشان می‌دهد که دارای اهمیت بالایی هستند، درحالی‌که توسعه چندان نیافته‌اند. درواقع، اگرچه این بخش شامل خوشه‌های مرکزی است که برای حوزه‌ی مورد مطالعه مهم هستند، اما این خوشه‌ها نابالغ‌اند و در عین حال دارای پتانسیل رشد در آینده هستند. بخش چهارم نمودار، موضوعات نوظهور یا نادیده گرفته‌شده^۲ در ربع پایین سمت چپ نمودار و متشکل از خوشه‌های موضوعی با چگالی و مرکزیت کم هستند. کانون‌های پژوهشی این حوزه ممکن است موضوعات علمی نوظهور و یا رو به زوال باشند. این موضوعات توسعه اندکی را تجربه کرده‌اند و از اهمیت کم‌تری برخوردارند و احتمالاً در آینده این شانس را دارند که از موضوعات مورد توجه پژوهشگران باشند یا کلاً از مسیر پژوهشی خارج شوند (Cobo et al., ۲۰۱۱-a; Cobo et al., ۲۰۱۱-b).

به منظور مصورسازی نقشه ساختار مفهومی^۳ از تحلیل هم‌واژگانی کلمات کلیدی در مطالعات این حوزه و پکیج Binlioshiny مربوط به به نرم‌افزار R استفاده شد. ابزار biblioshiny امکان تحلیل ساختار مفهومی با استفاده از سه رویکرد تحلیل عاملی و کاهش ابعاد شامل مقیاس‌گذاری چندبعدی^۴ (MDS)، تحلیل تناظر چندگانه^۵ (MCA) و تحلیل تناظری^۶ (CA) را فراهم می‌کند. در پژوهش حاضر به منظور ترسیم ساختار مفهومی، روش تحلیل تناظر چندگانه به کار گرفته شد که روشی برای تجزیه و تحلیل داده‌های چندمتغیره بوده که به منظور نمایش کلان داده‌ها و روابط میان آن‌ها در یک فضای با ابعاد کم‌تر به کار می‌رود. در این روش، درجه شباهت یا عدم تشابه میان هر جفت از داده‌ها محاسبه و نقشه‌ای از داده‌های حاصل در یک فضای دو یا سه‌بعدی ترسیم می‌شود که در آن داده‌های با ویژگی‌های مشابه نزدیک به هم قرار می‌گیرند (Borner et al., ۲۰۰۳). در تحلیل‌های علم‌سنجی از این روش برای مصورسازی روابط پیچیده میان اسناد یا اصطلاحات استفاده می‌شود. به عبارتی، تحلیل تناظر چندگانه با شناسایی خوشه‌هایی از اسناد یا اصطلاحات مرتبط، می‌تواند نقش مهمی در درک ساختار کلی داده‌ها داشته باشد.

برای ترسیم شبکه هم‌رخدادی واژگان از نرم‌افزار VOSviewer استفاده شد و Biblioshiny که یک افزونه تحت وب برای پکیج Bibliometrix است و توسط نرم‌افزار R فراخوانی می‌شود، امکان ترسیم نقشه‌ی موضوعی (نمودار راهبردی تکامل موضوعی) و ساختار مفهومی و نمودار پویایی حوزه‌های پژوهشی به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی را فراهم نمود. اطلاعات کتابشناختی مقالات بازیابی شده یکبار در قالب Tab delimited به منظور تحلیل‌های مربوط به نرم‌افزار VOSviewer، و یکبار دیگر در قالب Bibtex به منظور انجام تحلیل‌های مربوط به افزونه Biblioshiny انجام شد.

لازم به ذکر است که نرم‌افزار VOSviewer و پکیج Biblioshiny(R) برای پژوهشگران امکانی را فراهم می‌کنند که به منظور تحلیل بهتر و تهیه نقشه‌های دقیق‌تر، اصطلاحنامه‌ای از واژگان یکدست شده را به نرم‌افزار وارد کنند. این اصطلاحنامه پس از استخراج داده‌ها از پایگاه‌های اطلاعاتی و غربالگری آن‌ها تهیه شد و برای ترسیم نقشه‌های هم‌رخدادی واژگان، نمودار روند پویایی، نمودار راهبردی و نقشه ساختار مفهومی برای هر دو نرم‌افزار VOSviewer و پکیج Biblioshiny(R) به کار گرفته شد. در این راستا، در وهله نخست، سیاهه‌ای از تمام واژگان کلیدی منابع مورد تحلیل، از طریق نرم‌افزار VOSviewer استخراج و از فرمت .txt به فرمت .xlsx تبدیل شدند. سپس به تهیه اصطلاحنامه‌ای از واژگان یکدست شده اقدام گردید. بدین ترتیب که در اصطلاحنامه‌ی مذکور، تمام اختصارات به صورت شکل کامل کلمات نوشته شدند. به علاوه، کلیدواژه‌های مترادف، کلیدواژه‌های جمع

^۱ Basic and transversal themes^۲ Emerging or declining themes^۳ Conceptual Structure^۴ Multidimensional scaling analysis^۵ Multiple Correspondence Analysis^۶ correspondence analysis (CA)

تحلیل موضوعی مطالعات به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی با رویکرد علم‌سنجی **زودآیند ویرایش نشده**

و مفرد، و نیز کلیدواژه‌هایی که چند شکل نوشتاری داشتند، به شکل کلمات رایج‌تر تبدیل شدند. اصلاحنامه واژگان یکدست شده به فرمت .txt. در مرحله شبکه هم‌رخدادی واژگان به نرم‌افزار VOSviewer و برای مصورسازی نقشه موضوعی و ساختار مفهومی و نمودار تحولات موضوعی با پکیج Biblioshiny(R) به نرم‌افزار وارد شدند تا فرایند مصورسازی داده‌ها و ترسیم نقشه‌ها با دقت بیش‌تری صورت گیرد.

به منظور تحلیل اولیه اسناد منتخب، نسخه BibTeX داده‌های مربوط به اطلاعات کتابشناختی مقالات، دانلود و به نرم‌افزار biblioshiny وارد شدند. اطلاعات کلی و مهم در مورد مقالات مورد تحلیل در جدول ۲ ارائه شده است. این جدول نشان می‌دهد که نرخ رشد سالانه برای مقالات تحلیل شده ۷۱،۴۶ درصد بود که نشان‌دهنده‌ی افزایش قابل توجه و مداوم در انتشارات علمی در بازه زمانی مورد بررسی است. در زمینه‌ی به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی، همکاری میان پژوهشگران نسبتاً بالا است که می‌تواند نشان‌دهنده‌ی رویکرد چندرشته‌ای به این زمینه باشد. به‌طور کلی، شاخص همکاری ۳۰،۸۲ درصد نشان می‌دهد که مطالعات این حوزه حاصل تلاش‌های مشترک میان متخصصان و پژوهشگران است و تعداد کم‌تری از مقالات توسط یک نویسنده منتشر شده‌اند که این امر به‌طور بالقوه می‌تواند منجر به نتایج جامع‌تر و متنوع‌تری شود.

جدول ۲. اطلاعات اصلی مربوط به مقالات مورد تحلیل

توصیف	نتایج
پایگاه	وب آو ساینس
نوع منابع	مقالات نشریه
بازه زمانی	۲۰۲۰-۲۰۲۵
تعداد نشریات	۵۵۵۳
تعداد مقالات	۲۲۹۱۸
نرخ رشد سالانه مقالات	۷۱،۴۶ درصد
میانگین سنی مقالات	۲،۱۷
میانگین استناد در هر سند	۱۴،۹
کلمات کلیدی پایگاه (ID)	۱۹۱۱۷
کلمات کلیدی نویسنده (DE)	۴۵۶۵۰
تعداد نویسندگان	۹۰۴۰۷
مقالات تک نویسنده	۲۱۲۹
نویسندگان مشترک در هر سند	۵،۴۵
تألیفات مشترک بین‌المللی	۳۰،۸۲ درصد

در جدول ۳ روند انتشارات و استنادات دریافتی مقالات حوزه به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی در بازه زمانی ۲۰۲۰-۲۰۲۵ ارائه شده است.

جدول ۳. روند رشد و استنادات دریافتی مقالات حوزه به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی در بازه زمانی ۲۰۲۰-۲۰۲۵

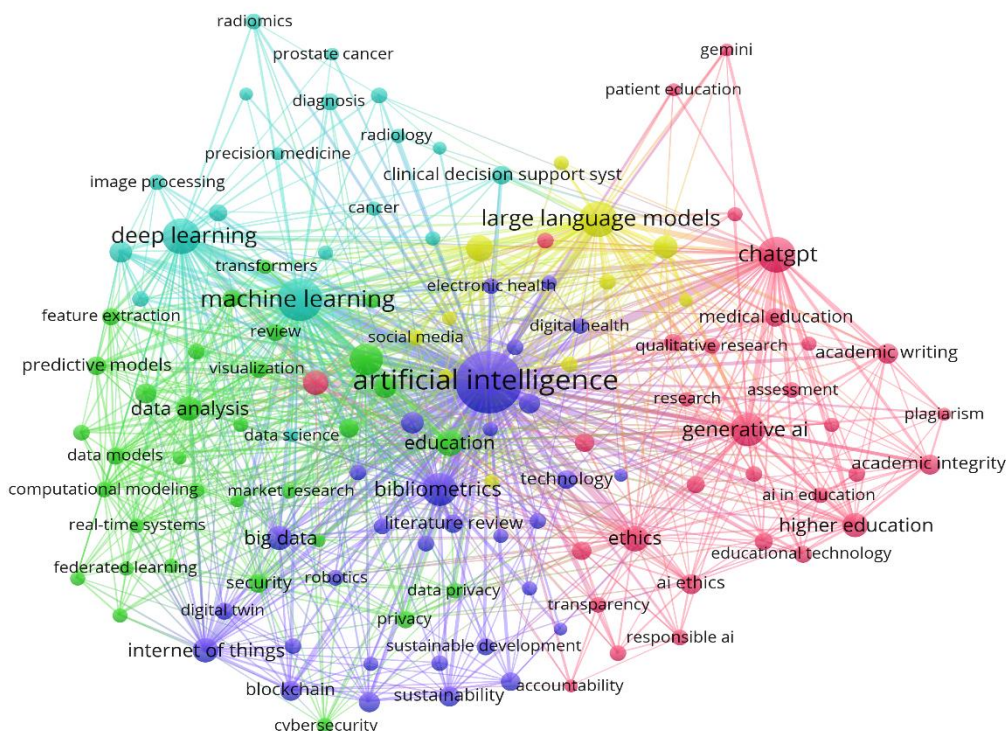
سال انتشار	تعداد مقالات	تعداد استنادات
۳۸۱۸۹	۶۹۰	۲۰۲۰
۵۴۱۰۷	۱۳۶۳	۲۰۲۱
۴۹۴۴۱	۲۳۳۷	۲۰۲۲
۷۱۷۶۳	۲۵۶۸	۲۰۲۳
۸۵۵۸۴	۵۷۳۴	۲۰۲۴
۴۲۴۵۷	۱۰۲۲۶	۲۰۲۵

جدول ۳ نشان می‌دهد که انتشاراتی که با هدف تحلیل به‌کارگیری هوشواره‌ها در انتشارات علمی در خلال سال‌های ۲۰۲۰ تا ۲۰۲۵ منتشر شده‌اند، به‌طور کلی روند رشد بالایی داشته‌اند. بر این اساس، در سطح جهانی، در اولین سال مورد بررسی (۲۰۲۰)، ۶۹۰ مقاله منتشر شده است، این در حالی است که در سال ۲۰۲۵، این تعداد به ۱۰۲۲۶ مقاله رسیده که رشد ۱۴۸۲ برابری را نشان می‌دهد. افزایش انتشارات این حوزه از سال ۲۰۲۰ تا ۲۰۲۵ را می‌توان تحت تأثیر ترکیبی از عوامل فناورانه، نهادی و سیاستی دانست. توسعه ابزارهای هوش مصنوعی مولد از اواخر سال ۲۰۲۲، به ویژه Chat GPT و سایر ابزارهای مشابه، موانع نگارش، جستجو، خلاصه‌سازی و ویرایش متن علمی را کاهش داد که این امر، حجم انتشارات را در حوزه‌های مرتبط با هوش مصنوعی مسئولانه و اخلاق نشر افزایش داد. به علاوه، گسترش انتشارات با دسترسی آزاد و پلتفرم‌های علمی دیجیتال، انتشارات را در سطح جهانی سریع‌تر و دسترس‌پذیرتر کرد. بنابراین، نگرانی‌ها در مورد صداقت علمی، حاکمیت هوش مصنوعی و سیاست ویراستاری افزایش یافته است. همین امر سبب شده که پژوهشگران مطالعات خود را در خصوص استفاده درست و مسئولانه از هوش مصنوعی افزایش دهند.

یافته‌های پژوهش

پاسخ به پرسش اول پژوهش: شبکه هم‌واژگانی مطالعات به‌کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی شامل چه خوشه‌هایی است؟

واژگان کلیدی پژوهش‌ها نقش مهمی در درک مضامین اصلی، روندها، موضوعات و حوزه‌های فعال در یک زمینه‌ی مطالعاتی خاص ایفا می‌کنند و به پژوهشگران و متخصصان یک حوزه این امکان را می‌دهند تا مسائل رایج و مفاهیم کلیدی را در متون آن حوزه شناسایی کنند. در شکل ۳ شبکه هم‌رخدادی واژگان کلیدی نویسندگان مقالات به‌کارگیری هوشواره‌ها در نظام نشر علمی، قابل نمایش است که در آن اندازه‌ی هر گره، بیانگر تعداد دفعات تکرار واژه‌ی مورد نظر در مقالات و فاصله بین گره‌ها بیانگر میزان استفاده از این کلیدواژه‌ها به‌صورت هم‌زمان است. بدین معنی که اگر دو کلیدواژه با یک پیوند به هم مرتبط شده‌اند، حداقل در یکی از مقالات تحلیل شده به‌صورت هم‌زمان به‌کار رفته‌اند و می‌توان اذعان داشت که با هم، هم‌رخدادی دارند. شکل ۳، نقشه هم‌واژگانی را برای ۱۲۵ کلیدواژه با حداقل ۶۰ رخداد به منظور ارائه‌ی روند توسعه مفهومی نشان می‌دهد.



شکل ۳. شبکه هم‌رخدادی واژگان در پژوهش‌های به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی

مطابق شکل ۳، شبکه هم‌رخدادی واژگان مطالعات به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی شامل ۵ خوشه است که هر کدام با رنگی خاص به نمایش گذاشته شده‌اند. مهم‌ترین خوشه‌ها که حاوی بیش‌ترین تعداد کلمات کلیدی هستند با رنگ‌های قرمز و سبز قابل نمایش‌اند. در جدول ۴ جزئیات اطلاعات مربوط به هر کدام از خوشه‌های موضوعی در شبکه علمی این حوزه حوزه شامل تعداد رخداد واژگان و قدرت ارتباطات هر واژه با سایر اصطلاحات به همراه بسامد آن‌ها به تفکیک هر خوشه ارائه شده است. در پژوهش حاضر مانند سایر مطالعات علم‌سنجی و تحلیل‌های هم‌واژگانی، نام‌گذاری خوشه‌ها براساس توجه به واژگان پربسامد و شاخص هر خوشه (یعنی واژگانی که در آن خوشه، بیشترین فراوانی و اهمیت را دارند و به‌عنوان «نماینده کانونی» خوشه عمل می‌کنند) و توجه به محتوای کلی خوشه (یعنی مفهوم اصلی و حاکم بر کل خوشه، که از ترکیب واژگان مختلف و روابط بین آن‌ها استخراج می‌شود) انجام شده است. معادل انگلیسی کلیدواژه‌های جدول ۴ به مقاله پیوست شده است (پیوست ۱).

جدول ۴. اطلاعات مربوط به شبکه هم‌واژگانی به تفکیک هر خوشه

خوشه	تعداد گره‌ها	عنوان خوشه	کلیدواژه‌های مربوط به هر خوشه به ترتیب بالاترین قدرت پیوند به همراه بسامد آن‌ها
۱	۳۳	مدیریت اخلاقی	چت‌جی‌بی‌تی (۱۴۹۸)، هوش مصنوعی مولد (۱۱۴۴)، اخلاق (۵۱۴)، آموزش عالی (۴۲۶)، هوش مصنوعی توضیح‌پذیر (۴۲۵)، صداقت هوش مصنوعی و آکادمیک (۲۴۵)، نگارش علمی (۲۳۵)، اخلاق هوش مصنوعی (۲۳۲)، آموزش پزشکی (۱۹۵)، هوش مصنوعی قابل‌اعتماد (۱۷۳)، یکپارچگی پژوهش فناوری آموزشی (۱۵۰)، همکاری انسان و هوش مصنوعی (۱۸۴)، ملاحظات اخلاقی (۱۳۵)، شفافیت (۱۰۷)، سنجش (۱۱۰)، پذیرش فناوری (۱۴۷)، حاکمیت هوش مصنوعی (۱۰۱)، ارزیابی (۱۱۰)، هوش مصنوعی مسئول (۹۸)، پرستاری (۱۳۰)، هوش مصنوعی در آموزش (۱۱۰)، سواد هوش مصنوعی (۱۰۹)، جمنیای (۷۰)، سرقت ادبی (۷۰)، مهندسی پرامیت (۸۰)، پژوهش (۸۰)، پاسخگویی (۶۹)، پژوهش کیفی (۱۱۶)، درگیری تحصیلی (۸۱)، اخلاق پژوهش (۱۱۰)، آموزش بیمار (۷۰)، ادراکات (۷۵)
۲	۳۲	کاربرد هوش	مرور نظام‌مند (۹۶۷)، آموزش (۵۷۰)، تحلیل داده‌ها (۴۳۳)، مدل‌های داده (۱۵۲)، دستورالعمل (۳۰۵)، امنیت (۱۹۷)، مدل‌های

خوشه	تعداد گره‌ها	عنوان خوشه	کلیدواژه‌های مربوط به هر خوشه به ترتیب بالاترین قدرت پیوند به همراه بسامد آن‌ها
سبز		مصنوعی در	پیش‌بینی‌کننده (۱۸۹)، دقت (۱۴۸)، مدل‌سازی محاسباتی (۱۲۲)، تحلیل پیش‌بینانه (۲۱۶)، حریم خصوصی (۱۵۷)، تصمیم‌گیری تحلیل داده، فرایند (۱۸۶)، استخراج ویژگی (۱۵۲)، سیستم عامل‌های بلادرنگ (۹۲)، بررسی یا دوری (۱۳۳)، مصورسازی (۱۳۳)، حریم خصوصی داده‌ها نگارش، انتشار و (۱۲۲)، تحقیقات بازار یا مطالعات بازاریابی (۷۷)، ناشر (۷۰)، بهینه‌سازی (۱۲۷)، رایانش ابری (۹۰)، داده‌کاوی (۱۰۰)، مانیتورینگ (۹۱)، دآوری علمی نشر دانشگاهی (۶۸)، امنیت سایبری (۱۰۷)، رایانش لبه‌ای (۷۷)، دآوری هم‌تا (۶۸)، مصورسازی داده‌ها (۸۴)، ویراستار (۷۳)، ثبات یا پایایی (۸۰)، یادگیری فدرال (۸۰)
۳		کاربرد هوش مصنوعی در	هوش مصنوعی (۸۷۸۶)، کتاب‌سنجی (۱۲۰۴)، اینترنت اشیاء (۴۲۷)، داده‌های بزرگ (۴۴۴)، بهداشت و درمان (۲۲۶)، بلاکچین (۲۳۴)، کووید ۱۹ (۲۷۳)، پایداری (۲۵۹)، صنعت ۴.۰ (۲۳۰)، فناوری (۲۰۴)، تحول دیجیتال (۱۸۴)، مرور ادبیات (۱۷۹)، دوقلوی دیجیتال (۱۴۵)، علم‌سنجی (۱۱۴)، اتوماسیون (۱۲۵)، سلامت دیجیتال (۱۲۶)، توسعه پایدار (۱۴۶)، سلامت الکترونیک (۱۱۹)، الگوریتم (۱۰۳)، روباتیک (۱۰۰)، تأثیرگذاری (۹۸)، واقعیت مجازی (۱۰۷)، فناوری دیجیتال (۱۲۲)، دیجیتال‌سازی (۱۰۳)، نوآوری (۹۳)، تحلیل استنادی (۸۴)، سلامت عمومی (۹۰)، فناوری اطلاعات (۷۴)، تله‌مدیسن یا پزشکی از راه دور (۸۷)، چالش‌ها (۶۶)، روندهای پژوهش (۶۹)
۴		کاربرد هوش مصنوعی در	یادگیری ماشین (۲۵۱۱)، یادگیری عمیق (۱۴۸۱)، شبکه‌های عصبی مصنوعی (۳۲۵)، سیستم‌های پشتیبانی تصمیم بالینی (۱۹۳)، الگوریتم‌های طبقه‌بندی (۱۳۹)، تشخیص (۱۵۲)، بنیایی کامپیوتری (۱۱۳)، پردازش تصویر (۱۰۱)، سیستم‌های پشتیبانی تصمیم (۹۱)، انتشارات پزشکی و علم داده (۶۹)، سرطان سینه (۱۲۸)، رادیولوژی (۸۰)، سرطان (۸۹)، کنترل کیفیت (۷۸)، تصویربرداری تشدید مغناطیسی (۶۶)، پزشکی سلامت شخصی سازی شده یا هدفمند (۶۷)، سرطان پروستات (۶۶)
۵		کاربرد هوش مصنوعی در فرایند جستجو و کشف دانش	مدل‌های زبانی بزرگ (۱۴۶۴)، پردازش زبان طبیعی (۴۸۴)، چت‌بات‌های هوش مصنوعی (۳۸۹)، متن‌کاوی (۱۲۴)، رسانه‌های اجتماعی (۱۱۷)، غربالگری (۱۰۴)، مدل‌سازی موضوعی (۷۶)، هوش مصنوعی محاوره‌ای (۶۷)، تحلیل احساسات (۶۸)، تولید تقویت‌شده بازایی یا RAG (۷۲)، گراف دانش (۷۵)، جستجو (۶۶)

جدول ۴ نشان می‌دهد که خوشه‌ی اول و دوم به واسطه‌ی جای دادن ۳۲ گره در خود، بزرگ‌ترین خوشه‌ی هم‌واژگانی مطالعات به‌کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی را تشکیل می‌دهند. تحلیل محتوای گره‌های موجود در خوشه اول (مانند اخلاق، اخلاق هوش مصنوعی، هوش مصنوعی قابل اعتماد، صداقت آکادمیک، ملاحظات اخلاقی، شفافیت، حکمرانی هوش مصنوعی، هوش مصنوعی مسئول، پاسخگویی و اخلاق پژوهش) نشان می‌دهد که خوشه‌ی اول با محوریت موضوعات مرتبط با مدیریت اخلاقی هوش مصنوعی و یکپارچگی پژوهش شکل گرفته است. در واقع، پرتکرارترین کلمات کلیدی این خوشه، اصطلاحات هنجاری و حاکمیتی هستند؛ بنابراین، خوشه مذکور، یکی از موضوعات مهم در حوزه به‌کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی است. زیرا بر استفاده اخلاقی از هوش مصنوعی تأکید کرده و به حفظ اعتماد در ارتباطات علمی کمک می‌کند. کاربرد و قدرت پیوند بالای کلیدواژه‌ی چت‌جی‌بی‌تی و پس از آن جمینای در این خوشه بیانگر تأکید بیش‌تر پژوهشگران بر این ابزارها برای پیشبرد فعالیت‌های علمی و پژوهشی بوده است.

در خوشه دوم، غلبه‌ی کلمات کلیدی روش‌محور مانند تحلیل داده‌ها، مدل‌های پیش‌بینی‌کننده، تحلیل پیش‌بینانه، استخراج ویژگی، داده‌کاوی، دقت، مدل‌سازی محاسباتی، مصورسازی داده‌ها، بهینه‌سازی و اصطلاحات مرتبط با طبقه‌بندی، بیانگر اهمیت کاربرد هوش مصنوعی در فرایند تحلیل داده است. کاربرد واژگانی مانند مرور نظام‌مند، دستورالعمل، بررسی یا دآوری، دآوری هم‌تا، ناشر، ویراستار، نشر دانشگاهی، مانیتورینگ، تصمیم‌گیری و پایایی نشان می‌دهد که به‌کارگیری مسئولانه هوشواره‌ها نه تنها برای تحلیل داده‌ها، بلکه برای فرایند نگارش علمی، کیفیت انتشار و دآوری علمی ضروری است. محتوای این خوشه، نمایانگر به‌کارگیری عملیاتی هوشواره‌ها به‌صورت مسئولانه در نظام نشر علمی است. برخلاف خوشه اول که عمدتاً اخلاقی و هنجاری است، این خوشه بر نحوه استفاده واقعی از هوش مصنوعی در فرایند پژوهش و انتشار تمرکز دارد و شامل مراحل می‌شود که در آن دانش علمی تولید، پردازش، ارزیابی و منتشر می‌شود.

بررسی کلیدواژه‌های موجود در خوشه سوم (شامل کتاب‌سنجی، علم‌سنجی، تأثیرگذاری، تحلیل استنادی، چالش‌ها و روندهای پژوهش) حکایت از آن دارد که سومین خوشه‌ی بزرگ هم‌واژگانی با محوریت مطالعات مرتبط با کاربرد هوش مصنوعی در ارزیابی‌های

تحلیل موضوعی مطالعات به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی با رویکرد علم‌سنجی **زودآیند ویرایش نشده**

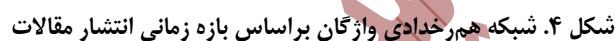
پس از انتشار و سنجش تأثیرگذاری علمی شکل گرفته که از اهمیت بالایی در نظام نشر علمی برخوردار است. زیرا اعتبار منابع منتشر شده به این مسئله بستگی دارد که تا چه حد ابزارهای هوش مصنوعی در فرایند مختلف نگارش و انتشار پژوهش‌ها به صورت مسئولانه استفاده شده است. این خوشه از آن جهت اهمیت دارد که هوش مصنوعی می‌تواند داده‌های کتاب‌سنجی و آلتمتریک (مانند تحلیل الگوهای استنادی، روندهای موضوعی، حوزه‌های نوظهور، شبکه‌های همکاری، نفوذ نویسندگان، کیفیت نشریات و عملکرد علمی کشورها در رشته‌های مختلف) را به صورت دقیق‌تر، کارتر و در مقیاس بزرگ‌تر نسبت به روش‌های دستی سنتی پردازش کند. بنابراین، تحلیل تأثیرگذاری مبتنی بر هوش مصنوعی می‌تواند مطالعات کتاب‌سنجی را با ادغام رویکردهای پیش‌بینانه و آینده‌نگرانه تقویت کرده و تصمیمات استراتژیک و مبتنی بر شواهد را در مدیریت و سیاست‌گذاری پژوهش‌ها بهبود بخشد.

نگاهی به برخی از پرکاربردترین کلیدواژه‌های موجود در خوشه چهارم، حکایت از آن دارد که این خوشه با محوریت مطالعات مرتبط با کاربرد هوش مصنوعی و ابزارهای یادگیری ماشین و شبکه‌های عصبی عمیق در انتشارات پزشکی و سلامت شکل گرفته است. با توجه به اینکه مطالعات پزشکی و سلامت مستقیماً بر تشخیص، درمان، ایمنی بیمار، سیاست‌گذاری بهداشت عمومی و تصمیم‌گیری‌های بالینی تأثیر می‌گذارند، استفاده از ابزارهای هوش مصنوعی در این حوزه پیامدهایی دارد که بسیار فراگیر است. خطاها در کاربرد هوش مصنوعی در مطالعات پزشکی می‌تواند عواقب جدی در دنیای واقعی داشته باشد. اگر ابزارهای هوش مصنوعی به طور نامناسب استفاده شوند، ممکن است سوگیری را افزایش و شفافیت را کاهش دهند و تفسیرهای نادرست ایجاد کنند. در حوزه پزشکی و سلامت، چنین مشکلاتی نه تنها بر کیفیت انتشار، بلکه بر قابلیت اطمینان استفاده بالینی بعدی نیز تأثیرگذار است. بنابراین، استفاده مسئولانه از هوش مصنوعی در نظام انتشارات سلامت ضروری است.

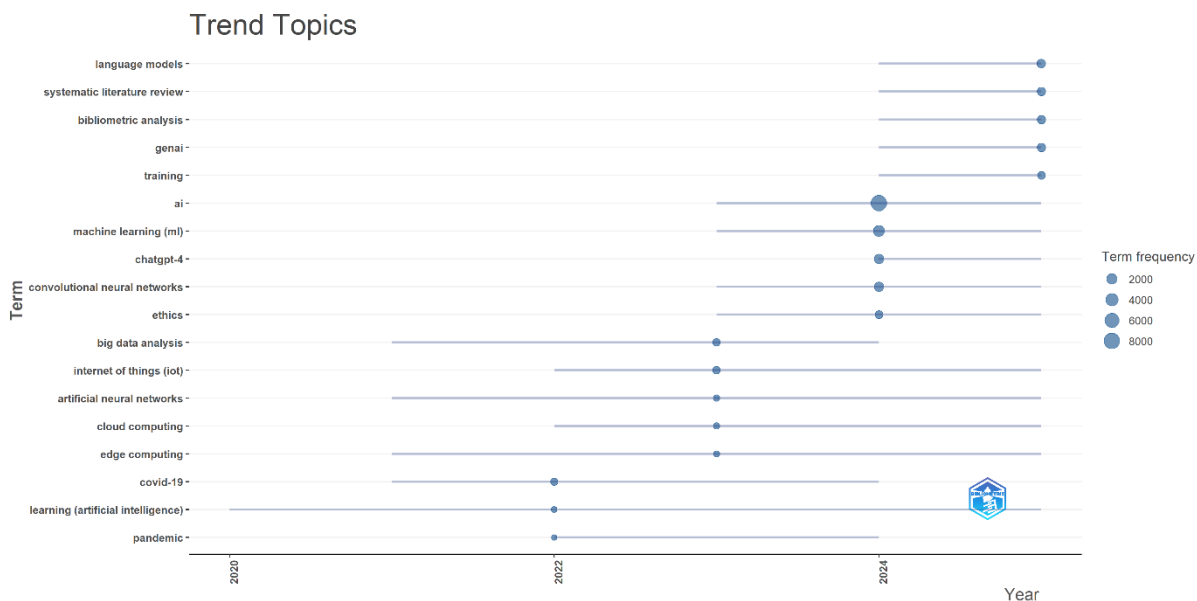
خوشه پنجم با محوریت مباحث و موضوعات مرتبط با کاربرد ابزارهای هوش مصنوعی و مدل‌های زبانی در فرایند جستجوی خودکار و کشف دانش، شکل گرفته است. حضور کلیدواژه‌هایی مانند مدل‌های زبانی بزرگ، پردازش زبان طبیعی، چت‌بات‌های هوش مصنوعی، متن‌کاوی، غربالگری، مدل‌سازی موضوعی، تحلیل احساسات، گراف دانش و جستجو در این خوشه شاهدی بر این مدعاست. یافتن، شناسایی و سازماندهی منابع مرتبط، یکی از تأثیرگذارترین مراحل نظام نشر علمی است و کیفیت آن می‌تواند علاوه بر بهبود سرعت و دقت در فرایند پژوهش، اعتبار یافته‌های مطالعات را افزایش دهد. قابلیت جستجو، غربالگری و کشف خودکار دانش مطلوب به پژوهشگران کمک می‌کند تا فراتر از تطبیق ساده کلمات کلیدی حرکت کنند و ارتباطات پنهان بین مفاهیم و اسناد را کشف کنند. در نتیجه، هوش مصنوعی می‌تواند زمان مورد نیاز برای جستجوی متون را کاهش دهد و در عین حال جامعیت بازایی منابع را افزایش دهد. بنابراین، کاربرد هوش مصنوعی در فرایند جستجو و کشف دانش، نیاز به شفافیت و توجه دارد.

پاسخ به پرسش دوم پژوهش: روند تحولات و پویایی موضوعات پژوهش‌های به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی در طول سال‌های ۲۰۲۰ تا ۲۰۲۵ به چه صورت است؟

شبکه هم‌رخدادی واژگان حوزه‌ی به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی براساس میانگین زمان انتشار پژوهش‌ها در شکل ۴ نشان داده شده است. رنگ هر گره نشان‌دهنده میانگین سال انتشار مطالعات مرتبط با آن موضوع است، به گونه‌ای که طیف رنگی از سرمه‌ای برای موضوعات قدیمی‌تر (۲۰۲۰) به آبی (۲۰۲۳) و زرد برای موضوعات جدیدتر (۲۰۲۵) تغییر می‌کند.



علاوه بر تحلیل شبکه هم‌رخدادی واژگان، به منظور بررسی روند تغییرات موضوعات پژوهش‌های حوزه به‌کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی، به تحلیل کلمات کلیدی پرکاربرد در پنج سال گذشته پرداخته و پویایی کلمات کلیدی در این زمینه براساس فیلد کلیدواژه‌های نویسندگان در شکل ۵، نمایش داده شده است. در این راستا، پنج کلیدواژه برتر در هر سال که دارای حداقل فراوانی بیست بودند، انتخاب شدند. این نمودار، بیانگر آن است که چه موضوعاتی برای مدت طولانی‌تری مورد پژوهش قرار گرفته‌اند و چه موضوعاتی اخیراً به لیست مطالعات اضافه شده‌اند و چگونه موضوعات پژوهش در طول زمان تغییر کرده‌اند.



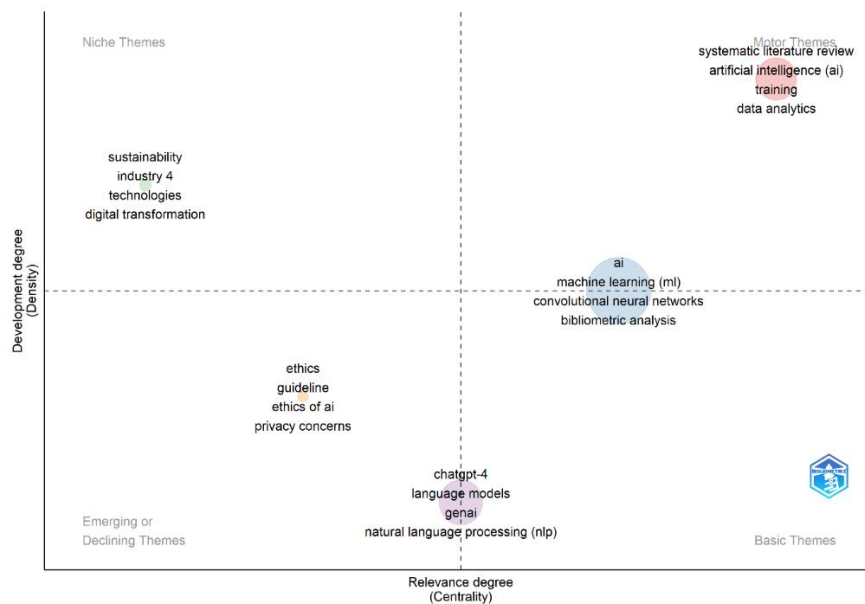
شکل ۵. پویایی موضوعات پژوهش‌های به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی طی سال‌های ۲۰۲۰-۲۰۲۵

شکل ۵ نشان می‌دهد که واژگان کلیدی برتر هر سال با هم متفاوت هستند که این تغییرات، نشانگر ماهیت در حال تحول این حوزه است. نگاهی به محبوب‌ترین کلیدواژه‌ها طی سال‌های اخیر نشان می‌دهد که در مطالعات این حوزه بر استفاده اخلاقی از ابزارهای هوش مصنوعی مولد و مدل‌های زبانی بر نگارش نظام‌مند پژوهش‌ها و ارزیابی‌های علمی پس از انتشار منابع تأکید بیشتری شده است. بنابراین، حوزه‌های مذکور می‌توانند در آینده نیز از موضوعات محبوب و مورد توجه پژوهشگران باشند.

پاسخ به پرسش سوم پژوهش: روند راهبردی تکامل موضوعی مطالعات به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی براساس سنجه‌های مرکزیت و چگالی به چه صورت است؟

به منظور خوشه‌بندی مضامین اصلی و مطالعه تکامل حوزه به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی، نقشه موضوعی^۱ (نمودار راهبردی) پژوهش‌های این حوزه بر مبنای شاخص‌های مرکزیت و چگالی ترسیم شد. شکل ۶ نقشه موضوعی پژوهش‌های حوزه به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی طی سال‌های ۲۰۲۰ تا ۲۰۲۵ را براساس کلمات کلیدی نویسندگان مقالات نشان می‌دهد. هر حباب در نقشه، یک خوشه را نشان می‌دهد و اندازه حباب‌ها در آن، بیانگر اهمیت خوشه‌های موضوعی در این حوزه (میزان رخداد کلمات مربوط به هر خوشه) است. درحالی‌که رنگ حباب‌ها نشان‌دهنده تمایز خوشه‌ها بوده که هر کدام براساس شباهت موضوعی میان کلیدواژگان مربوط به هر خوشه، تشکیل شده است. کلماتی که در هر خوشه نشان داده شده دارای رخداد بالاتری نسبت به سایر واژگان مربوط به آن خوشه است و در نتیجه می‌تواند معرف محتوای آن خوشه باشد.

^۱ Thematic map



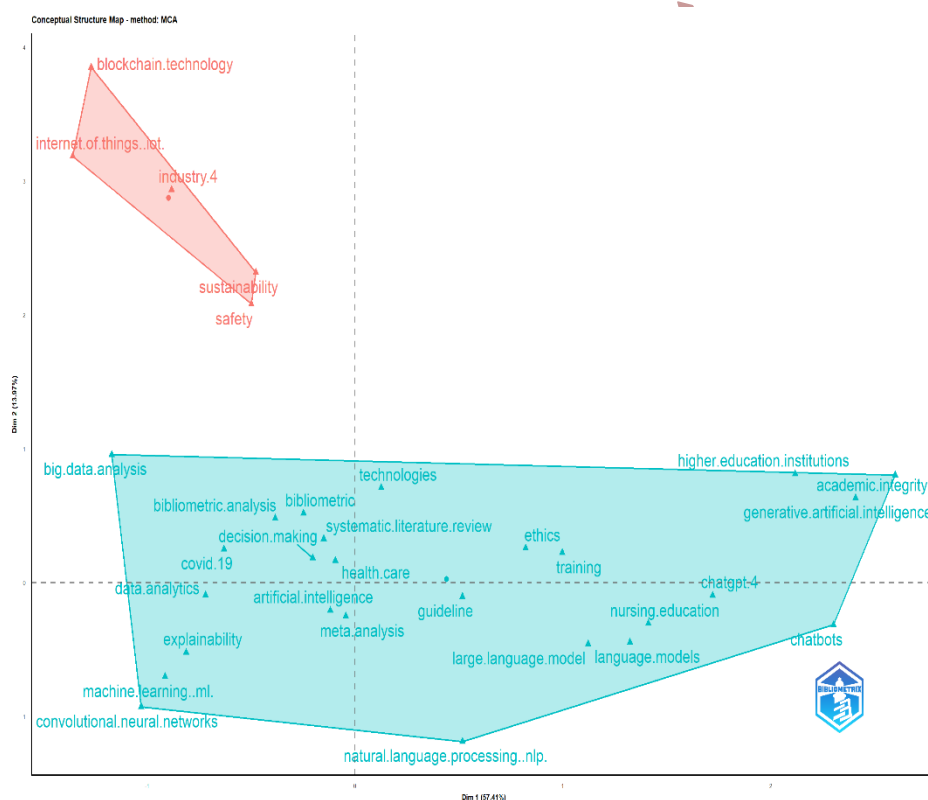
شکل ۶. روند راهبردی تکامل موضوعی پژوهش‌های به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی طی سال‌های ۲۰۲۰ تا ۲۰۲۵

براساس معیارهای مرکزیت و چگالی، پنج خوشه اصلی در حوزه به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی شناسایی شده است. در نقشه موضوعی پژوهش‌های این حوزه، موضوعات پیشران (محرك) که نشان‌دهنده مضامین تثبیت‌شده، تأثیرگذار و زمینه‌ساز پیشرفت در زمینه به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی است، حاوی یک خوشه‌ی واحد است که دربرگیرنده‌ی کلمات کلیدی مانند «مرور نظام‌مند متون، هوش مصنوعی، آموزش و تحلیل داده» است (مرکزیت ۰,۰۲۳ و چگالی ۰,۷۸۴). این کلیدواژه‌ها نشان می‌دهد که موضوعات پیشران شامل کاربرد ابزارهای هوش مصنوعی در فرایند تحلیل داده و نگارش نظام‌مند پژوهش‌ها هستند. خوشه‌ی موضوعات توسعه‌یافته و منفرد که به معنی حوزه‌های پژوهشی تخصصی و در عین حال کم‌تر رایج در این حوزه است، شامل کلمات کلیدی پایداری، صنعت ۴، فناوری‌ها و تحول دیجیتال است (مرکزیت ۰,۰۰۷ و چگالی ۰,۶۶۲). این خوشه بر چگونگی ادغام هوش مصنوعی در سیستم‌های انتشارات دیجیتالی و تحول فناورانه و سازمانی هوش مصنوعی مسئولانه در انتشار علمی تمرکز دارد. موضوعات مرتبط با کاربرد هوش مصنوعی، ابزارهای یادگیری ماشین و شبکه‌های عمیق در بهبود فرایند ارزیابی پژوهش‌ها و تأثیرگذاری علمی، با خصوصیات مضامین پیشران (محرك) و پایه (با توجه به قرارگیری خوشه در مرز بین این دو بخش نمودار) ظاهر شده‌اند (مرکزیت ۰,۰۲۲ و چگالی ۰,۵۹۶). خوشه‌ی مرتبط با استفاده از ابزارهای هوش مصنوعی مولد، مدل‌های زبانی، پردازش زبان طبیعی و پلتفرم چت‌جی‌پی‌تی در فرایند نگارش و نشر علمی، در مرز موضوعات پایه و نوظهور در نقشه قرار گرفته و در نتیجه می‌توان گفت که خصوصیات هر دو را داراست (مرکزیت ۰,۰۱۶ و چگالی ۰,۵۴۷). خوشه موضوعات نوظهور یا رو به زوال که بیانگر حوزه‌هایی هستند که ممکن است با گذشت زمان اهمیت خود را از دست بدهند یا به موضوعات مهم در آینده تبدیل شوند، بر نگرانی‌های اخلاقی، امنیتی و هنجاری در کاربرد ابزارهای هوش مصنوعی در نظام نشر علمی تأکید دارد (مرکزیت ۰,۰۰۸ و چگالی ۰,۵۵۸).

پاسخ به پرسش چهارم پژوهش: ساختار مفهومی مطالعات به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی براساس تحلیل عاملی به چه صورت است؟

یکی دیگر از راه‌های تحلیل موضوعات پژوهش در یک حوزه‌ی علمی، مصورسازی نقشه ساختار مفهومی با استفاده از تحلیل

هم‌واژگانی کلمات کلیدی در مطالعات آن حوزه است. در پژوهش حاضر به منظور نمایش ساختار مفهومی پژوهش‌های به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی، روش تحلیل تناظر چندگانه به کار گرفته شد. شکل ۷ نقشه ساختار مفهومی حاصل از تحلیل تناظر چندگانه و تحلیل کلمات کلیدی نویسندگان را نشان می‌دهد. در این نقشه، کلیدواژه‌هایی که نزدیک به نقطه‌ی مرکزی قرار گرفته‌اند، واژگانی هستند که در سال‌های اخیر بیش‌تر مورد توجه قرار گرفته‌اند و کلمات نزدیک به لبه، موضوعاتی هستند که کم‌تر در پژوهش‌های اخیر مورد استفاده قرار گرفته و یا در موضوعات دیگر گنجانده شده‌اند (Xie et al., ۲۰۲۰). در این نقشه، نزدیک‌تر بودن کلمات کلیدی به هم، نشان‌دهنده‌ی آن است که مقالات بیش‌تری آن‌ها را با هم ارائه می‌دهند و دور بودن واژگان از هم به معنی آن است که پژوهش‌های کم‌تری از این منظر به آن‌ها می‌پردازند. همان‌طور که شکل ۷ نشان می‌دهد، حوزه‌هایی مانند مراقبت سلامت، هوش مصنوعی، فراتحلیل، تصمیم‌گیری، کوید ۱۹، مرور نظام‌مند متون، دستورالعمل، تحلیل داده، کتابسنجی و پاسخگویی، با توجه به فاصله کوتاه‌تری که با نقطه مرکزی نمودار دارند، طی سال‌های اخیر بیش‌تر مورد توجه پژوهشگران بوده‌اند. نزدیک بودن دو کلیدواژه‌ی صداقت آکادمیک و هوش مصنوعی مولد، نشان‌دهنده رخداد بیش‌تر آن‌ها با هم در پژوهش‌های این حوزه بوده است.



شکل ۷. نقشه ساختار مفهومی مطالعات به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی با استفاده از تحلیل تناظر چندگانه

شکل ۷ بیانگر آن است که با استفاده از الگوریتم تحلیل تناظر چندگانه، دو خوشه‌ی موضوعی در حوزه به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی، شناسایی شده است. خوشه سبز رنگ به شفافیت و استفاده اخلاقی و مسئولانه از ابزارهای هوش مصنوعی در تحلیل داده‌های بزرگ، فرایند نگارش، انتشار و ارزیابی‌های پس از انتشار پژوهش‌ها با تأکید بر حوزه‌های پزشکی و سلامت اشاره دارد. خوشه قرمز رنگ که در فاصله‌ی دورتری از نقطه مرکز نمودار قرار گرفته و نسبت به خوشه قبلی از اهمیت کم‌تری برخوردار است، بر موضوعات مرتبط با فناوری‌های مؤثر بر پیشرفت سیستم انتشارات دیجیتالی و مسائل اخلاقی و امنیتی در استفاده از آن‌ها مانند فناوری بلاکچین، اینترنت اشیاء و سایر فناوری‌های صنعت ۴ تأکید می‌کند.

یافته‌های فوق، نقشه‌ای توصیفی از ساختار و توزیع موضوعی پژوهش‌های حوزه هوشواره‌ها در نشر علمی را ارائه می‌دهد. در بخش بعدی، با عبور از سطح توصیفی داده‌ها، الگوهای پنهان و تناقض‌های ساختاری مستتر در این شبکه موضوعی مورد استنتاج تحلیلی قرار گرفته و در قالب سه چالش مفهومی بنیادین، تبیین می‌شود.

بحث و نتیجه‌گیری

تحلیل علم‌سنجی ۲۲،۹۱۸ مقاله در بازه‌ی ۲۰۲۰ تا ۲۰۲۵ نشان می‌دهد که حوزه‌ی «به‌کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی» از یک دغدغه‌ی فنی-نوآورانه به یک پارادایم حکمرانی-اخلاقی تبدیل شده است. یافته‌های این پژوهش که از طریق شبکه‌ی هم‌واژگانی، نقشه‌ی موضوعی، تحلیل تناظر چندگانه و پویایی زمانی استخراج شده‌اند، نه تنها ساختار مفهومی این حوزه را ترسیم می‌کنند، بلکه شکاف‌های سیاستی و عملیاتی نظام نشر جهانی را نیز آشکار می‌سازند. در ادامه، تفسیر تحلیلی یافته‌ها براساس پرسش‌های پژوهش، همسویی با چارچوب نظری لایه‌ای-بوم‌شناختی، دلالت‌های سیاستی، محدودیت‌ها و مسیرهای پژوهشی آتی ارائه می‌شود.

تفسیر یافته‌ها در چارچوب نظری لایه‌ای-بوم‌شناختی

۱. خوشه اول: مدیریت اخلاقی هوش مصنوعی و یکپارچگی پژوهش - تجسم «لایه فراحاکمیتی»
پاسخ به پرسش اول نشان داد که شبکه‌ی هم‌واژگانی از پنج خوشه‌ی معنایی تشکیل شده است. خوشه‌ی اول با محوریت مفاهیمی چون «صدافت آکادمیک»، «شفافیت»، «پاسخگویی»، «هوش مصنوعی قابل اعتماد» و «حکمرانی هوش مصنوعی»، نمایانگر لایه هنجاری-حکمرانی است. این یافته با گزارش جوبین و همکاران (Jobin et al., ۲۰۱۹) مبنی بر همگرایی جهانی حول پنج اصل اخلاقی (شفافیت، عدالت، عدم ضرر، مسئولیت‌پذیری و حریم خصوصی) همسو است و مواضع کمیته اخلاق نشر علمی (COPE) که «مسئولیت‌پذیری» را ستون اصلی نویسندگی علمی می‌داند را تأیید می‌کند.
با این حال، پاسخ به پرسش دوم (نقشه‌ی موضوعی) آشکار ساخت که این خوشه در ربع پایین-چپ (نوظهور/در حال زوال) قرار دارد. این جایگاه نشان می‌دهد که علی‌رغم پذیرش گسترده‌ی اصول، شیوه‌های اجرایی و عملیاتی‌سازی آن‌ها هنوز به بلوغ نرسیده است. این تناقض ظریف - فراوانی بالا در مقابل بلوغ ساختاری پایین - زنگ خطری سیاستی است: اخلاق نباید به «ضمیمه‌ای تشریفاتی» در سیاست‌های ناشران تقلیل یابد. از منظر نظریه‌ی «پیمان اختیار»، لی (Lee, ۲۰۲۶) این وضعیت را می‌توان چنین تفسیر کرد: درحالی که «لایه اختیار هنجاری» (ارزش‌ها و اصول) تا حد زیادی تثبیت شده است، «لایه اختیار معرفتی» (معیارهای شواهد برای تشخیص استفاده مسئولانه)، «لایه اختیار منبع» (اعتبار ناشران و ابزارهای هوش مصنوعی) و «لایه اختیار داده» (گزینش داده‌های آموزشی) هنوز با ابهامات جدی مواجه هستند. به عبارتی، جامعه علمی می‌داند چه اصولی را باید رعایت کند، اما هنوز سازوکار روشنی برای اینکه چگونه رعایت می‌شود و چگونه باید حسابرسی شود، ندارد. تحلیل عمیق‌تر این خوشه نشان می‌دهد که ظهور واژه‌هایی مانند «سواد هوش مصنوعی» و «همکاری انسان و هوش مصنوعی» حاکی از آن است که جامعه‌ی جهانی در حال حرکت از ممنوعیت مطلق به مدیریت مسئولانه است. این تحول، نشان‌دهنده‌ی پذیرش هوش مصنوعی به عنوان بُعدی جدید از شایستگی حرفه‌ای پژوهشگران است - همان‌طور که آرمیتاژ (Armitage, ۲۰۲۵) استدلال می‌کند: «توانایی کار با ChatGPT اکنون بخشی از مهارت‌های ضروری نویسندگی علمی است، همان‌طور که تسلط بر Word یا EndNote بوده است.»

۲. خوشه دوم: کاربردهای عملیاتی (تحلیل داده، نگارش، انتشار و داوری) - تجسم «لایه کارکردی»
خوشه‌ی دوم - با تمرکز بر «تحلیل داده»، «مرور نظام‌مند»، «داوری همتا»، «تصمیم‌گیری»، «دستورالعمل» و «بهینه‌سازی» - بزرگ‌ترین خوشه از نظر تعداد گره‌هاست و نمایانگر لایه کارکردی-عملیاتی است. این یافته، سیاست‌های الزویر (Elsevier, ۲۰۲۴) و تیلور و فرانسیس (Taylor & Francis, ۲۰۲۴) را تقویت می‌کند که استفاده از هوش مصنوعی در «فرآیند تحقیق» را مشروع می‌دانند، مشروط بر افشای کامل.

در پاسخ به پرسش دوم، نقشه‌ی موضوعی نشان داد که این خوشه در ربع اول (موضوعات پیش‌ران) قرار دارد. این جایگاه نشان‌دهنده‌ی بلوغ و نهادینه‌شدگی هوش مصنوعی در فرایندهای میانی پژوهش است. حضور همزمان واژگانی چون «مرور نظام‌مند»، «داوری همتا» و «دستورالعمل» در کنار «تحلیل داده» و «مدل‌های پیش‌بینی‌کننده» حاکی از آن است که هوش مصنوعی از ابزار کمکی به بخشی جدایی‌ناپذیر از زیرساخت روزمره‌ی پژوهش تبدیل شده (Ganjavi et al., ۲۰۲۴; Liang et al., ۲۰۲۴) است. این یافته با شواهد موجود در ادبیات همخوانی دارد که نشان می‌دهد ابزارهای هوش مصنوعی اکنون در تمام مراحل چرخه علمی - از جستجوی ادبیات و تحلیل داده‌ها تا نگارش مقاله و داوری - به کار گرفته می‌شوند (Gartlehner et al., ۲۰۲۵) و حتی نهادهای معتبری مانند گروه کاکرین، استفاده از آن‌ها را به‌عنوان ابزاری برای تضمین کیفیت در کنار قضاوت انسانی تأیید کرده‌اند (Salman et al., ۲۰۲۵). بنابراین، هوش مصنوعی دیگر صرفاً یک «یار کمکی» نیست، بلکه به بخشی از زیرساخت روزمره و استاندارد پژوهش تبدیل شده است. با این حال، تحلیل تناظر چندگانه (پرسش سوم) و همچنین بررسی قدرت پیوند میان خوشه‌های در شبکه‌ی هم‌واژگانی نشان می‌دهد که پیوند معناداری میان این خوشه و خوشه‌ی اول (اخلاقی) وجود ندارد. این گسست، شکاف عملیاتی-هنجاری را آشکار می‌سازد: استفاده‌ی گسترده از هوش مصنوعی بدون الگوهای شفاف افشا و پاسخگویی در جریان است. از منظر بوم‌شناسی جامعه‌شناختی-فنی هرماندز و همکاران (Domínguez Hernández et al., ۲۰۲۵)، این وضعیت نشان می‌دهد که «رویه‌های استاندارد» و «جریان‌های منبع» در این لایه به بلوغ رسیده، اما همچنان نیازمند شفافیت و قابلیت حسابرسی در لایه فراحاکمیتی است. تحلیل مفهومی خروجی‌های علم‌سنجی این خوشه نشان می‌دهد که حضور واژه‌هایی مانند «یادگیری فدرال» و «رایانش لبه‌ای» حاکی از آن است که هوش مصنوعی دیگر تنها یک ابزار کمکی نیست، بلکه بخشی جدایی‌ناپذیر از زیرساخت‌های پژوهشی شده است.

۳. خوشه سوم: ارزیابی تأثیرگذاری علمی - تجسم «مرحله پس از انتشار»

خوشه‌ی سوم، مشتمل بر کلیدواژه‌هایی مانند «کتاب‌سنجی»، «علم‌سنجی»، «تحلیل استنادی» و «تأثیرگذاری»، براساس مراحل چرخه پژوهش، به‌عنوان نماینده‌ی مرحله‌ی پس از انتشار شناسایی می‌شود. این انتساب بر پایه‌ی ماهیت ابزاری این کلیدواژه‌ها (که همگی به ارزیابی کمی و کیفی مقالات منتشرشده مربوط می‌شوند) و همچنین جدایی ساختاری این خوشه از خوشه‌های مرتبط با نگارش و انتشار در نقشه‌ی موضوعی، صورت گرفته است. استخراج این خوشه به‌عنوان یک خوشه مستقل، یکی از مهم‌ترین یافته‌های این پژوهش محسوب می‌شود، چراکه مرور ادبیات نشان می‌دهد مطالعات پیشین (Ganjavi et al., ۲۰۲۴; Liang et al., ۲۰۲۴) عمدتاً بر مراحل نگارش و انتشار متمرکز بوده و توجه چندانی به ارزیابی تأثیر پس از انتشار، به‌ویژه در ارتباط با نقش هوش مصنوعی، نداشته‌اند.

در پاسخ به پرسش دوم، نقشه‌ی موضوعی نشان داد که این خوشه در ربع اول (موضوعات با مرکزیت بالا) و ربع سوم (موضوعات نوظهور با تراکم پایین) قرار دارد. در چارچوب علم‌سنجی، جایگاه در ربع این دو ربع نشان‌دهنده‌ی وضعیتی است که یک حوزه از یک سو در حال کسب اهمیت و ارتباط با سایر حوزه‌هاست، اما از سوی دیگر، هنوز به تراکم و پیوستگی درونی کافی برای تبدیل شدن به یک حوزه کاملاً توسعه‌یافته و تثبیت‌شده دست نیافته است. (Cobo et al., ۲۰۱۱-a) این یافته با شواهد موجود در ادبیات همخوانی دارد که نشان می‌دهد کاربرد هوش مصنوعی در کتاب‌سنجی و تحلیل استنادی، علیرغم رشد تصاعدی و تحولات روش‌شناختی سریع، هنوز با چالش‌های بنیادین در زمینه‌هایی مانند شفافیت، سوگیری و یکپارچگی مواجه است و به مرحله‌ای از «بلوغ کامل» که با ثبات ساختاری و اجماع نظری همراه باشد، نرسیده است (Haan, ۲۰۲۵).

مهم‌تر آنکه، تحلیل شبکه‌ی هم‌واژگانی نشان می‌دهد که پیوند ساختاری چندانی میان کلیدواژه‌های خوشه‌ی ارزیابی تأثیر و خوشه‌ی اخلاقی در ادبیات موجود وجود ندارد. هرچند این «نبود پیوند واژگانی» را نمی‌توان معادل «نبود توجه مفهومی» دانست و بررسی آن نیازمند تحلیل محتوایی عمیق‌تری است، اما می‌تواند به‌عنوان نشانه‌ای اولیه از احتمال وجود شکاف میان این دو حوزه تلقی شود.

به عنوان یک فرضیه برای پژوهش‌های آینده، این پرسش مطرح است که آیا پژوهش‌های حوزه‌ی ارزیابی تأثیر، به درستی به نقش هوش مصنوعی در تولید محتوای مورد ارزیابی توجه کرده‌اند؟ پرسشی که پاسخ به آن مستلزم تحلیل محتوایی مقالات این حوزه است. تحلیل این خوشه نشان می‌دهد که درحالی که ناشران بزرگ مثل اشپرنگر نیچر از هوش مصنوعی برای شناسایی سرقت ادبی یا تشخیص دستکاری تصویر استفاده می‌کنند، این ابزارها همزمان می‌توانند شاخص‌های تأثیرگذاری را دستکاری کنند (Ibrahim et al., ۲۰۲۵). این دوگانگی - هوش مصنوعی به عنوان هم «نگهبان» و هم «تهدید» - نیازمند سیاست‌های دقیق‌تری است که در ادبیات فعلی، کمتر مورد توجه قرار گرفته است.

۴. خوشه چهارم: انتشارات پزشکی و سلامت - تجسم «لایه زمینه‌ای-حوزه‌ای»

تشکیل یک خوشه مجزا برای پزشکی و سلامت (با تمرکز بر «سیستم‌های پشتیبان تصمیم بالینی»، «تشخیص»، «رادیولوژی» و «سرطان»)، مؤید رویکرد بوم‌شناختی هراندز و همکاران (Domínguez Hernández et al., ۲۰۲۵) است که بر زمینه‌مندی حکمرانی هوش مصنوعی تأکید می‌کند. در پاسخ به پرسش دوم، نقشه‌ی موضوعی نشان داد که این خوشه در ربع دوم (توسعه‌یافته منفرد) قرار دارد. این جایگاه نشان می‌دهد که علی‌رغم تمایز و اهمیت این حوزه، هنوز به‌طور یکپارچه با سایر لایه‌ها پیوند نخورده است. در حوزه پزشکی، خطاهای ناشی از کاربرد نامناسب هوش مصنوعی می‌تواند پیامدهای مستقیم و گاه جبران‌ناپذیری بر سلامت و زندگی بیماران داشته باشد. این ویژگی، این حوزه را از سایر حوزه‌های علمی (مانند ریاضیات یا علوم انسانی) متمایز نموده و نیازمند سیاست‌های سختگیرانه‌تر، شفافیت بالاتر و نظارت انسانی مستمر است. سیاست‌های JAMA و Lancet که استفاده از هوش مصنوعی را فقط برای «بهبود زبان» مجاز می‌دانند، مستقیماً با این یافته همسو است (Bagenal et al., ۲۰۲۴).

با این حال، تحلیل مفهومی خروجی‌های علم‌سنجی نشان می‌دهد که چنین رویکرد محافظه‌کارانه‌ای، در مواردی که نسبت فایده به ریسک به وضوح مثبت است، ممکن است زمینه‌ی بهره‌گیری از ظرفیت‌های این فناوری را محدود کند. این نگرانی با نقدهای وارد بر رویکرد «کاهش ریسک تا حد ممکن» در پیش‌نویس قانون هوش مصنوعی اتحادیه اروپا همخوانی دارد و کارشناسان بر غیرعملی بودن و بار نظارتی نامتناسب آن تأکید کرده و خواستار بازنگری مبتنی بر شواهد روزآمد و رویکردی معقول‌تر شده‌اند (Fraser & Bello y Villarino, ۲۰۲۳). چنان‌که مرور نظام‌مند آپوستولو و همکاران (Apostolou et al., ۲۰۲۶) نشان می‌دهد الگوریتم‌های هوش مصنوعی می‌توانند با استفاده از داده‌های تصویربرداری و نشانگرهای زیستی، سیگنال‌های ضعیف پیش‌بالینی سرطان لوزالمعده را شناسایی کنند - با این حال، موانع جدی همچون اعتبارسنجی خارجی محدود، نبود تحلیل کافی از سودمندی بالینی و ناهمگونی در پروتکل‌های بالینی، موجب شده است تا سیاست‌های نظارتی فعلی رویکردی محافظه‌کارانه و مبتنی بر «غیرقابل اعتماد» پنداشتن کلی این فناوری را در پیش گیرند. حضور کلیدواژه‌هایی نظیر «تشخیص» و «سیستم‌های پشتیبانی تصمیم بالینی» در این خوشه نشان می‌دهد که هوش مصنوعی نه تنها در فرایند نشر مقالات پزشکی، بلکه در محتوای بالینی آن مقالات نیز نقش پیدا کرده است. این هم‌پوشانی مرزهای پژوهشی و بالینی، لزوم تدوین راهنماهای حوزه‌ویژه را دوچندان می‌کند.

۵. خوشه پنجم: جستجو و کشف دانش - تجسم «مرحله آغازین چرخه پژوهش»

خوشه‌ی پنجم با محوریت «مدل‌های زبانی بزرگ»، «پردازش زبان طبیعی» و «تولید تقویت‌شده بازیابی»، اگرچه کوچک‌ترین خوشه است، اما از نظر مفهومی راهبردی‌ترین لایه‌ی آغازین چرخه‌ی پژوهش را پوشش می‌دهد. قرارگیری این خوشه در ربع چهارم نقشه‌ی موضوعی (نوظهور) و همچنین تحلیل پویایی زمانی که نشان‌دهنده‌ی جدیدترین موضوعات است، با ظهور فناوری‌هایی مانند چت‌جی‌پی‌تی همبستگی کامل دارد (Brown et al., ۲۰۲۵). این جایگاه نشان می‌دهد که جستجوی هوشمند ادبیات به سرعت در حال تبدیل شدن به یک ابزار پایه برای پژوهشگران است؛ با این حال، چالش‌های جدی در مورد «سوگیری الگوریتمی» و «عدم شفافیت در معیارهای بازیابی» پابرجاست.

تحلیل این خوشه، ادبیات موجود را از منظری نو گسترش می‌دهد: درحالی که سیاست‌های ناشران عمدتاً بر «نگارش» متمرکز است، این خوشه نشان می‌دهد که خطرات اخلاقی هوش مصنوعی از همان ابتدای چرخه پژوهش - یعنی زمانی که پژوهشگر از ابزارهایی

تحلیل موضوعی مطالعات به کارگیری مسئولانه هوشواره‌ها در فرایند نشر علمی با رویکرد علم‌سنجی **زودآیند ویرایش نشده**

مانند Elicit یا Semantic Scholar برای شناسایی منابع استفاده می‌کند - آغاز می‌شود. پژوهش‌های حسابرسی نشان داده‌اند که این ابزارها با چالش‌هایی مانند «شکاف ارجاعی» (عدم استناد به منابع خوانده‌شده) و «حباب‌های موضوعی» (ارائه اطلاعات تأییدکننده باورهای کاربر) مواجه هستند (Jacob et al., ۲۰۲۵). بنابراین، نیاز به شفافیت در فرآیند جستجو، موضوعی حیاتی است که تاکنون در سیاست‌های بین‌المللی ناشران به آن پرداخته نشده است.

پیشنهادهای اجرایی پژوهش

براساس یافته‌های پژوهش، پیشنهادهای اجرایی در چهار سطح، تدوین شده‌اند:

پیشنهاد	مبنای اولیه	اولویت
سطح اول: برای ناشران و مجلات علمی		
ایجاد «بیانیه شفافیت هوش مصنوعی» در هر مقاله	خوشه ۱ (شفافیت، پاسخگویی) و ضعف پیوند با خوشه ۲	بالا
الزام به ثبت ابزار هوش مصنوعی استفاده‌شده (نام، نسخه، پرامپت‌های خوشه ۱ (هوش مصنوعی قابل اعتماد، مهندسی پرامپت) کلیدی) در بخش روش‌شناسی		بالا
تدوین راهنمای تخصصی برای داوران در مواجهه با متون تولید/ویرایش‌شده خوشه ۲ (داوری همتا، دستورالعمل) با هوش مصنوعی		متوسط
بازنگری در شاخص‌های ارزیابی با افزودن «نمره شفافیت هوش مصنوعی» خوشه ۳ (کتاب‌سنجی، تأثیرگذاری) و شکاف پیوند با خوشه ۱		متوسط
سطح دوم: برای سیاست‌گذاران علم و معاونت‌های پژوهشی		
طراحی و اجباری کردن «آموزش سواد هوش مصنوعی» برای همه خوشه ۱ (سواد هوش مصنوعی، آموزش عالی) و نقشه موضوعی پژوهشگران		بالا
تعریف شاخص‌های ملی ارزیابی پژوهش با لحاظ «نسبت مشارکت انسانی خوشه ۳ (ارزیابی تأثیرگذاری، علم‌سنجی) و خوشه ۱ (صداقت در تولید دانش) آکادمیک»		بالا
تشکیل کارگروه‌های تخصصی (مانند پزشکی و سلامت) به منظور تدوین خوشه ۴ (پزشکی و سلامت) - تأیید لزوم توجه به بافت و حوزه سیاست‌های افشای هوش مصنوعی		بالا
حمایت مالی از پروژه‌های توسعه ابزارهای متن‌باز برای تشخیص خودکار خوشه ۲ (داوری همتا، اتوماسیون) و خوشه ۵ (پردازش زبان طبیعی) محتوای تولیدشده با هوش مصنوعی در داوری		متوسط
سطح سوم: برای پژوهشگران و نویسندگان		
افشای صریح و دقیق استفاده از هوش مصنوعی در بخش روش‌شناسی (نه خوشه ۱ (شفافیت، صداقت آکادمیک) فقط در تشکر)		بالا
استفاده از هوش مصنوعی فقط برای وظایف ساختاریافته (خلاصه‌سازی، خوشه ۲ و چارچوب نظری پژوهش تقویت‌شده (همکاری انسان- ویرایش زبانی، استخراج) و اجتناب از سپردن قضاوت علمی به آن (هوش مصنوعی)		بالا
مستندسازی پرامپت‌های مورد استفاده در فرایند جستجو و مرور نظام‌مند (به خوشه ۵ (مهندسی پرامپت، تولید تقویت‌شده بازایی، جستجو) ویژه در خوشه ۵)		متوسط
دریافت تأیید اخلاقی از کمیته اخلاق پژوهش در صورت استفاده از هوش خوشه ۴ (اخلاق پژوهش، حریم خصوصی) مصنوعی در تحلیل داده‌های حساس (پزشکی)		متوسط
سطح چهارم: برای توسعه‌دهندگان ابزارهای هوش مصنوعی		
طراحی قابلیت «گواهی شفافیت» در ابزارهای نوشتاری هوش مصنوعی خوشه ۱ (هوش مصنوعی قابل اعتماد، شفافیت)		بالا
ایجاد خروجی استاندارد از تاریخچه تعامل و پرامپت‌ها برای الحاق به مقاله خوشه ۱ (مهندسی پرامپت، پاسخگویی)		متوسط
توسعه نسخه‌های حوزه‌ویژه (پزشکی، حقوق، علوم انسانی) با پایگاه‌های خوشه ۴ و لزوم زمینه‌مندی (بوم‌شناسی جامعه‌شناختی-فنی) دانش معتبر و محدود		بالا

پیشنهاد برای پژوهش‌های آتی

مسیرهای زیر برای پژوهش‌های آتی، پیشنهاد می‌شود:

- عملیاتی‌سازی اصول اخلاقی: طراحی و اعتبارسنجی چارچوب «افشای استاندارد هوش مصنوعی» از طریق پژوهش ترکیبی (مرور نظام‌مند + دلفی + مطالعه موردی) و شناسایی موانع اجرایی از طریق مصاحبه‌های عمیق با ذی‌نفعان ایرانی.
- یکپارچه‌سازی ارزیابی و شفافیت: طراحی «شاخص یکپارچگی هوش مصنوعی» و بررسی همبستگی آن با استنادات دریافتی از طریق تحلیل استنادی تطبیقی بین مقالات دارای افشا و فاقد افشا.
- حوزه‌ویژنه‌نگاری و زمینه‌مندی: مقایسه سیاست‌های ناشران پزشکی در برابر علوم انسانی و انجام مطالعه فراترکیب درباره واکنش جامعه پزشکی به استفاده از هوش مصنوعی در نگارش بالینی.
- روش‌شناختی و ابزاری: توسعه ابزار متن‌باز برای تشخیص خودکار افشاجاری هوش مصنوعی و طراحی پلتفرم «بازیابی هوشمند شفاف» مبتنی بر گراف دانش با قابلیت نمایش معیارهای رتبه‌بندی.
- پیگیری روندها و آینده‌پژوهی: تکرار این مطالعه در سال ۲۰۲۷ برای ردیابی تحول ساختار مفهومی و سناریوپردازی از آینده نشر علمی در افق ۴۰۳۵ با محوریت تعامل انسان-هوش مصنوعی.

تقدیر و تشکر

این مقاله برگرفته از طرح پژوهشی با عنوان: «طراحی الگوی کاربردی هوش مصنوعی در نظام نشر علمی ایران» است که با حمایت مالی مؤسسه تحقیقات سیاست علمی کشور انجام شده و بدین وسیله از این مؤسسه سپاسگزاری می‌شود.

تعارض منافع

نویسندگان اعلام می‌دارند که در خصوص انتشار این مقاله تضاد منافع وجود ندارد. علاوه بر این، موضوعات اخلاقی، ازجمله سرقت ادبی، رضایت آگاهانه، سوء رفتار، جعل داده‌ها، انتشار و ارسال مجدد و مکرر و همچنین، سیاست مجله در قبال استفاده از هوش مصنوعی از سوی نویسندگان رعایت شده است.

فهرست منابع

- احمدی، ح. و جمشیدی، م.ج. (۱۴۰۴). چالش‌ها و تعارضات اخلاقی به‌کارگیری هوش مصنوعی در نگارش علمی: یک تحلیل هم‌واژگانی. کتابداری و اطلاع‌رسانی، ۲۸ (۳)، ۳۴-۵. <https://doi.org/10.30481/lis.2025.556316.2307>
- مرتضوی شاهرودی، م.ع. و زارعی، ع. (۱۴۰۴). اصول و چالش‌های اخلاقی استفاده از هوش مصنوعی در تحقیقات علمی. فصلنامه اخلاقی‌پژوهی، ۷ (۴)، ۲۶-۵. <https://doi.org/10.22034/ethics.2025.51545.1701>
- Ahmadi, H. and Jamshidi, M.J. (۲۰۱۴). Challenges and ethical conflicts of using artificial intelligence in scientific writing: A synonym analysis. *Library and Information Science*, ۲۸(۳), ۵-۳۴. <https://doi.org/10.30481/lis.2025.556316.2307> [In Persian].
- Al-Qudimat, A. R., Fares, Z. E., Elaarag, M., Osman, M., Al-Zoubi, R. M., & Aboumarzouk, O. M. (۲۰۲۵). Advancing medical research through artificial intelligence: progressive and transformative strategies: a literature review. *Health Science Reports*, ۸(۲), e۷۰۲۰۰. <https://doi.org/10.1002/hsr2.70200>
- Ateriya, N., Sonwani, N. S., Thakur, K. S., Kumar, A., & Verma, S. K. (۲۰۲۵). Exploring the ethical landscape of AI in academic writing. *Egyptian Journal of Forensic Sciences*, ۱۹(۱), ۳۶. <https://doi.org/10.1186/s41935-025-00453-1>

- Apostolou, V., Papageorgiou, G., & Tjortjis, C. (۲۰۲۶). Artificial intelligence for early detection of pancreatic cancer: pre-diagnostic detection across imaging, biomarkers, and EHRs: a systematic review. *Evolutionary Intelligence*, ۱۹(۳), ۸۹. <https://doi.org/10.1007/s12065-026-01202-6>
- Armitage, R. (۲۰۲۵). Your brain on ChatGPT. *British Journal of General Practice*, ۷۵(۷۵۸), ۴۱۰. <https://doi.org/10.3399/bjgp25Xv43181>
- Bagenal, J., Biamis, C., Boillot, M., Brierley, R., Chew, M., Dehnel, T., Frankish, H., Grainger, E., Pope, J., Prowse, J., Samuel, D., Slogrove, A. L., Stacey, J., Thapaliya, G., Trethewey, F., Wang, H. H., Varley-Reeves, J., & Kleinert, S. (۲۰۲۴). Generative AI: ensuring transparency and emphasising human intelligence and accountability. *Lancet (London, England)*, ۴۰۴(۱۰۴۶۸), ۲۱۴۲-۲۱۴۳. [https://doi.org/10.1016/S.0140-6736\(24\)02615-1](https://doi.org/10.1016/S.0140-6736(24)02615-1)
- Borner, K., Chen, C., & Boyack, K. W. (۲۰۰۳). Visualizing knowledge domains. *Annual Review of Information Science and Technology*, ۳۷, ۱۷۹-۲۵۵. <https://doi.org/10.1002/aris.1440370106>
- Brown, A., Roman, M., & Devereux, B. (۲۰۲۵). A systematic literature review of retrieval-augmented generation: Techniques, metrics, and challenges. *arXiv preprint arXiv: 2508.07401*. <https://doi.org/10.48550/arXiv.2508.07401>
- Callon, M., Courtial, J.P. & Laville, F. (۱۹۹۱). Co-word analysis as a tool for describing the network of interactions between basic and technological research: The case of polymer chemistry. *Scientometrics*, ۲۲, ۱۵۵-۲۰۵. <https://doi.org/10.1007/BF02019280>
- Carobene, A., Padoan, A., Cabitza, F., Banfi, G., & Plebani, M. (۲۰۲۳). Rising adoption of artificial intelligence in scientific publishing: evaluating the role, risks, and ethical implications in paper drafting and review process. *Clinical Chemistry and Laboratory Medicine (CCLM)*, ۶۲(۵), ۸۳۵-۸۴۳. <https://doi.org/10.1080/033373024220252445890>
- Checchio, A., Bracciale, L., Loreti, P., Pinfield, S., & Bianchi, G. (۲۰۲۱). AI-assisted peer review. *Humanities and Social Sciences Communications*, ۸(۱), ۱-۱۱. <https://doi.org/10.1057/s41599-020-00703-8>
- Cheng, H., Sheng, B., Lee, A., Chaudhary, V., Atanasov, A. G., Liu, N., ... & Zheng, Y. F. (۲۰۲۴). Have ai-generated texts from llm infiltrated the realm of scientific writing? a large-scale analysis of preprint platforms. *bioRxiv*. <https://doi.org/10.1101/2024.03.25.086710>
- Celik, S. U. (۲۰۲۵). Integrating artificial intelligence into scientific writing: a narrative review for clinical and surgical researchers. *The American Journal of Surgery*, ۱۱۶۶۵۷. <https://doi.org/10.1016/j.amjsurg.2025.116657>
- Cobo, M. J., Lopez-Herrera, A. G., Herrera-Viedma, E., & Herrera, F. (۲۰۱۱a). An approach for detecting, quantifying, and visualizing the evolution of a research field: A practical application to the Fuzzy Sets Theory field. *Journal of Informetrics*, ۹(۱), ۱۴۶-۱۶۶. <https://doi.org/10.1016/j.joi.2010.10.002>
- Cobo, M. J., Lopez-Herrera, A. G., Herrera-Viedma, E., & Herrera, F. (۲۰۱۱b). Science mapping software tools: Review, analysis, and cooperative study among tools. *Journal of the American Society for Information Science and Technology*, ۶۲(۷), ۱۳۸۲-۱۴۰۲. <https://doi.org/10.1002/asi.21020>
- Domínguez Hernández, A., Perini, A. M., Hadjiloizou, S., Borda, A., Mahomed, S., & Leslie, D. (۲۰۲۵). Towards a sociotechnical ecology of artificial intelligence: power, accountability,

- and governance in a global context. *AI and Ethics*, 7(1). <https://doi.org/10.1007/s43681-020-00902-6>
- Dos Santos, Á. O., da Silva, E. S., Couto, L. M., Reis, G. V. L., & Belo, V. S. (2023). The use of artificial intelligence for automating or semi-automating biomedical literature analyses: A scoping review. *Journal of Biomedical Informatics*, 142, 104389. <https://doi.org/10.1016/j.jbi.2023.104389>
- Doskaliuk, B., Zimba, O., Yessirkepov, M., Klishch, I., & Yatsyshyn, R. (2020). Artificial Intelligence in Peer Review: Enhancing Efficiency While Preserving Integrity. *Journal of Korean Medical Science*, 40(7), e92. <https://doi.org/10.3346/jkms.2020.40.e92>
- Dwivedi, R., & Elluri, L. (2024). Exploring generative artificial intelligence research: A bibliometric analysis approach. *IEEE Access*, 12, 119884-119902. <https://doi.org/10.1109/ACCESS.2024.3450629>
- Ead, H. A. (2024). Exploring the impact of artificial intelligence on academic writing: perspectives of Ph. D. students in the faculty of science at cairo university. *SunText Review of Medical & Clinical Research*, 9(4). <https://doi.org/10.51737/2766-4813.2024.108>
- Elsevier. (2024). Artificial intelligence (AI) policy for authors and editors. <https://www.elsevier.com/about/policies/artificial-intelligence>
- Floridi et al. (2018). AI4People—an ethical framework for a good AI society. *Minds and Machines*, 28, 689-707. <https://doi.org/10.1007/s11023-018-9482-0>
- Fraser, H., & Bello y Villarino, J.-M. (2023). Acceptable risks in europe's proposed AI act: reasonableness and other principles for deciding how much risk management is enough. *European Journal of Risk Regulation*, 14(2), 431-446. <https://dx.doi.org/10.2139/ssrn.4516917>
- Ganjavi, C., Eppler, M. B., Pekcan, A., Biedermann, B., Abreu, A., Collins, G. S., Gill, I. S., & Cacciamani, G. E. (2024). Publishers' and journals' instructions to authors on use of generative artificial intelligence in academic and scientific publishing: bibliometric analysis. *BMJ*, 384. <https://dx.doi.org/10.1136/bmj-2023-077192>
- Ganjavi, C., Eppler, M. B., Pekcan, A., Biedermann, B., Abreu, A., Collins, G. S., Gill, I. S., & Cacciamani, G. E. (2023). Bibliometric analysis of publisher and journal instructions to authors on generative-AI in academic and scientific publishing. *arXiv preprint arXiv:2307.11918*. <https://doi.org/10.48550/arXiv.2307.11918>
- Gartlehner, G., Nussbaumer-Streit, B., Hamel, C., Garritty, C., Griebler, U., King, V. J., Devane, D., & Kamel, C. (2020). Responsible integration of artificial intelligence in rapid reviews: a position statement from the cochrane rapid reviews methods group. *Cochrane Evidence Synthesis and Methods*, 7(6), e70063. <https://doi.org/10.1002/cesm.70063>
- Gray, A. (2024). ChatGPT" contamination": estimating the prevalence of LLMs in the scholarly literature. *arXiv preprint arXiv:2403.16887*. <https://doi.org/10.48550/arXiv.2403.16887>
- Gurnal, P., & Rana, L. (2020). Artificial intelligence and publishing ethics: a narrative review and SWOT analysis. *Cureus*, 11(5), e84098. <https://doi.org/10.7759/cureus.84098>
- Haan, S. (2020). SemanticCite: citation verification with AI-powered full-text analysis and evidence-based reasoning. *arXiv preprint arXiv:2011.16198*. <https://doi.org/10.48550/arXiv.2011.16198>
- Ibrahim, H., Liu, F., Zaki, Y., & Rahwan, T. (2020). Citation manipulation through citation mills and pre-print servers. *Scientific Reports*, 10(1), 5480. <https://doi.org/10.1038/s41598-020-88709-7>

- Jacob, C., Kerrigan, P., & Bastos, M. (۲۰۲۵). The chat-chamber effect: trusting the AI hallucination. *Big Data & Society*, ۱۲(۱). <https://doi.org/10.1177/20539517241306345>
- Javanbakht, A. (۲۰۲۴). In context: AI will write your paper: the very different future of research and scientific writing in the age of artificial intelligence. *Journal of the American Academy of Child & Adolescent Psychiatry*, ۶۳(۷), ۶۸۱-۶۸۳. <https://doi.org/10.1016/j.jaac.2024.03.013>
- Jobin, A., Ienca, M., & Vayena, E. (۲۰۱۹). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, ۱(۹), ۳۸۹-۳۹۹. <https://doi.org/10.1038/s42256-019-0088-2>
- Khalifa, M., & Albadawy, M. (۲۰۲۴). Using artificial intelligence in academic writing and research: an essential productivity tool. *Computer Methods and Programs in Biomedicine Update*, ۵, ۱-۱۱. <https://doi.org/10.1016/j.cmpbup.2024.100145>
- Khlaif, Z. N., Mousa, A., Hattab, M. K., Itmazi, J., Hassan, A. A., Sanmugam, M., & Ayyoub, A. (۲۰۲۳). The potential and concerns of using AI in scientific research: ChatGPT performance evaluation. *JMIR Medical Education*, ۹(۱), e۴۷۰۴۹. <https://doi.org/10.2196/47049>
- Kousha, K., & Thelwall, M. (۲۰۲۴). Artificial intelligence to support publishing and peer review: A summary and review. *Learned Publishing*, ۳۷(۱), ۴-۱۲. <https://doi.org/10.1002/leap.۱۵۷۰>
- Lee, S. (۲۰۲۶). AI Integrity: a new paradigm for verifiable AI governance. *arXiv*: ۲۶۰۴.۱۱۰۶۵. <https://doi.org/10.48550/arXiv.2604.11065>
- Leung, T. I., de Azevedo Cardoso, T., Mavragani, A., & Eysenbach, G. (۲۰۲۳). Best practices for using AI tools as an author, peer reviewer, or editor. *Journal of Medical Internet Research*, ۲۵(۱), e۵۱۵۸۴. <https://doi.org/10.2196/51584>
- Liang, W., Zhang, Y., Wu, Z., Lepp, H., Ji, W., Zhao, X., ... & Zou, J. Y. (۲۰۲۴). Mapping the increasing use of LLMs in scientific papers. *arXiv preprint arXiv:2404.01268*. <https://doi.org/10.48550/arXiv.2404.01268>
- Litvinova, T. A., Mikros, G. K., & Dekhnich, O. V. (۲۰۲۴). Writing in the era of large language models: a bibliometric analysis of research field. *Научный результат. Вопросы теоретической и прикладной лингвистики*, ۱۰(۴), ۵-۱۶. <https://doi.org/10.184113/2313-8912-2024-10-4-01>
- Mangubat, F. M., & Saeedi, K. H. (۲۰۲۵). A bibliometric analysis of research trends in the application of artificial intelligence by college students in research writing. *Discover Education*, ۴(۱), ۶۷. <https://doi.org/10.1007/s44217-025-01067-4>
- Maral, M. (۲۰۲۴). Bibliometric analysis of global research on scientific writing. *DESIDOC Journal of Library & Information Technology*, ۴۴(۳). <https://doi.org/10.14429/djlit.44.03.19534>
- Masukume, G. (۲۰۲۴). The impact of AI on scientific literature: a surge in AI-associated words in academic and biomedical writing. *MedRxiv*, ۲۰۲۴-۰۵. <https://doi.org/10.1101/2024.05.31.24308296>
- Medina, W. A., Flores Velásquez, C. H., Olivares Zegarra, S. D. R., Morales Romero, G. P., Quispe Andía, A., Aybar Bellido, I. E., ... & Bravo, N. A. (۲۰۲۴). Using ChatGPT in university academic writing: A bibliometric review study on the implications for writing reports, papers, essays, and theses. *International Journal of Learning, Teaching and Educational Research*, ۲۳(۱۲), ۳۶۰-۳۸۱. <https://doi.org/10.26803/ijlter.23.12.19>
- Miller, L. E., Bhattacharyya, D., Miller, V. M., Bhattacharyya, M., & Miller, V. (۲۰۲۳). Recent trend in artificial intelligence-assisted biomedical publishing: a quantitative bibliometric analysis. *Cureus*, ۱۹(۵). <https://doi.org/10.7759/cureus.39224>

- Mortazavi Shahroudi, S. M. A., & Zarei, E. (۲۰۲۰). The principles and ethical challenges of using artificial intelligence in scientific research. *Journal of Moral Studies*, ۱(۴), ۵-۲۶. <https://doi.org/10.22034/ethics.2020.010405.1701> [In Persian].
- Mulyanah, E. Y., Hidayat, S., & Yuhana, Y. (۲۰۲۴). Enhancing the academic writing process using artificial intelligence: a bibliometric. *Humanities & Language: International Journal of Linguistics, Humanities, and Education*, ۱(۵), ۲۹۹-۳۰۷. <https://doi.org/10.32734/djd9kz89>
- Nakrowi, Z. S., & Lumettu, A. (۲۰۲۰). Two decades of academic writing assessment in higher education: a bibliometric and technological trend analysis of Scopus (۲۰۰۰-۲۰۲۰). *International Journal of Learning, Teaching and Educational Research*, ۲۴(۱۰), ۲۷۹-۳۰۶. <https://doi.org/10.26803/ijlter.24.10.13>
- Onuoha, C. E. (۲۰۲۰). Impact of artificial intelligence (AI) on the quality, efficiency, and transparency of the scholarly publishing process. *Trends in Scholarly Publishing*, ۴(۱), ۱۵-۲۱. <https://doi.org/10.21124/tsp.2020.10.21>
- Salih, M. H. B., Gul, T., Bukhari, S. S., Liaqat, I., & Majeed, A. (۲۰۲۰). AI-Enhanced scholarly communication: transforming peer review, knowledge dissemination, and academic publishing workflows in the digital era. *International Journal of Ethical AI Application*, ۱(۴), ۱۳-۱۹. <https://doi.org/10.64229/ijqp.7e84>
- Salman, H.A., Ahmad, M.A., Ibrahim, R., & Mahmood, J. (۲۰۲۰). Systematic analysis of generative AI tools integration in academic research and peer review. *Online Journal of Communication and Media Technologies*, ۱۹(۱), e۲۰۲۰۰۲. <https://doi.org/10.30930/ojcm/10832>
- Taylor & Francis. (۲۰۲۴). AI and machine learning: publishing policies. <https://authorserv-ices.taylorandfrancis.com/artificial-intelligence-policy>
- Upshall, M. (۲۰۱۹). Using AI to solve business problems in scholarly publishing. *Insights*, ۳۲(۱). <https://doi.org/10.1629/uksg.460>
- Wang, H., Fu, T., Du, Y., Gao, W., Huang, K., Liu, Z., ... & Zitnik, M. (۲۰۲۳). Scientific discovery in the age of artificial intelligence. *Nature*, ۶۲۰(۷۹۷۲), ۴۷-۶۰. <https://doi.org/10.1038/s41586-023-06221-2>
- Xie, H., Zhang, Y., Wu, Z., & Lv, T. (۲۰۲۰). A bibliometric analysis on land degradation: Current status, development, and future directions. *Land*, ۹(۱), ۲۸. <https://doi.org/10.3390/land9010028>
- Xue, Y. (۲۰۲۴). Towards automated writing evaluation: A comprehensive review with bibliometric, scientometric, and meta-analytic approaches. *Education and Information Technologies*, ۲۹(۱۵), ۱۹۵۵۳-۱۹۵۹۴. <https://doi.org/10.1007/s10639-024-12096-0>

پیوست ۱. معادل انگلیسی کلیدواژه‌های جدول ۴

Cluster	Total items	Most frequent keyword (occurrences)	Keywords of clusters in order of highest link strength
۱ (Red)	۳۲	Ethical management of artificial intelligence and research integrity	chatgpt (۱۴۹۸), generative ai (۱۱۴۴), ethics (۵۱۴), higher education (۴۲۶), explainable ai (۴۲۵), academic integrity (۲۴۵), academic writing (۲۳۵), ai ethics (۲۳۲), medical education (۱۹۵), trustworthy ai (۱۷۲), educational technology (۱۵۰), human-ai collaboration (۱۸۴), ethical considerations (۱۳۵), transparency (۱۰۷), measurement (۱۱۰), technology adoption (۱۴۷), ai governance (۱۰۱), evaluation (۱۱۰), responsible ai (۹۸), nursing (۱۳۰), ai in education (۱۱۰), ai literacy (۱۰۹), gemini (۷۰), plagiarism (۷۰), prompt engineering (۸۰), research (۸۰), accountability (۶۹), qualitative research (۱۱۶), academic engagement (۸۱), research ethics (۱۱۰), patient education (۷۰), perceptions (۷۵)
۲ (Green)	۳۲	Application of artificial intelligence in data analysis, writing, publishing	Systematic review (۹۶۷), education (۵۷۰), data analysis (۴۳۳), data models (۱۵۲), guidelines (۳۰۵), security (۱۹۷), predictive models (۱۸۹), accuracy (۱۴۸), computational modeling (۱۲۲), predictive analytics (۲۱۶), privacy

Cluster	Total items	Most frequent keyword (occurrences)	Keywords of clusters in order of highest link strength
		and scientific review process	(۱۵۷), decision making (۱۸۶), feature extraction (۱۵۲), real-time systems (۹۲), review (۱۳۳), visualization (۱۳۳), data privacy (۱۲۲), market research (۷۷), publisher (۷۰), optimization (۱۲۷), cloud computing (۹۰), data mining (۱۰۰), monitoring (۹۱), academic publishing (۶۸), cybersecurity (۱۰۷), edge computing (۷۷), peer review (۶۸), data visualization (۸۴), editor (۷۳), reliability (۸۰), federated learning (۸۰)
۳ (Blue)	۳۱	Application of artificial intelligence in scientific impact assessment	artificial intelligence (۸۷۸), bibliometrics (۱۲۰۴), internet of things (۴۲۷), big data (۴۴۴), healthcare (۲۲۶), blockchain (۲۳۴), covid-۱۹ (۲۷۳), sustainability (۲۵۹), industry ۴.۰ (۲۳۰), technology (۲۰۴), digital transformation (۱۸۴), literature review (۱۷۹), digital twin (۱۴۵), scientometrics (۱۱۴), automation (۱۲۵), digital health (۱۲۶), sustainable development (۱۴۶), electronic health (۱۱۹), algorithm (۱۰۳), robotics (۱۰۰), impact (۹۸), virtual reality (۱۰۷), digital technology (۱۲۲), digitalization (۱۰۳), innovation (۹۳), citation analysis (۸۴), public health (۹۰), information technology (۷۴), telemedicine (۸۷), challenges (۶۶), research trends (۶۹)
۴ (Turquoise)	۱۸	Application of artificial intelligence in medical and health publications	machine learning (۲۵۱۱), deep learning (۱۴۸۱), artificial neural networks (۳۲۵), clinical decision support systems (۱۹۳), classification algorithms (۱۳۹), diagnosis (۱۵۲), computer vision (۱۱۳), image processing (۱۰۱), decision support systems (۹۱), data science (۶۹), breast cancer (۱۲۸), radiology (۸۰), cancer (۸۹), quality control (۷۸), magnetic resonance imaging (۶۶), precision medicine (۶۷), prostate cancer (۶۶)
۵ (Yellow)	۱۲	Application of artificial intelligence in the process of search and knowledge discovery	Large Language Models (۱۴۶۴), Natural Language Processing (۴۸۴), AI Chatbots (۳۸۹), Text Mining (۱۲۴), Social Media (۱۱۷), Screening (۱۰۴), Topic Modeling (۷۶), Conversational AI (۶۷), Sentiment Analysis (۶۸), Retrieval Augmented Generation (RAG) (۷۲), Knowledge Graph (۷۵), Search (۶۶)

زودآیند ویرایش نشده