

Research Acceleration vs. Authenticity Concern: Two Sides of the AI Coin in Academic Writing

Ali Akbar Saboury^{1,*}

Article Info

Article type:

Popularization of
Science

Article history:

Receive Date:

01 February 2026

Revise Date:

09 February 2026

Accept Date:

11 February 2026

Available online:

21 June 2026

Keywords:

Research Ethics

AI Regulation

AI-Assisted Writing

Large Language

Models (LLMs)

Academic Integrity

COPE

Artificial intelligence, as the most powerful technological tool of the present century, has transcended the traditional boundaries of knowledge production and scientific research. This technology offers capabilities such as extraordinary data processing, automation of repetitive tasks, and assistance in complex analyses, providing an unprecedented opportunity to accelerate and deepen scientific discoveries. However, its application in writing scientific articles has brought significant achievements along with serious ethical, methodological, and legal challenges. Issues such as reduced transparency and research integrity, reinforcement of existing biases, violation of intellectual property rights, and reduction of researchers' analytical skills are examples of these difficulties. Faced with this complexity, the global community is striving to establish a balance between encouraging innovation and managing the risks associated with this technology by developing legal frameworks and regulations such as the European Union Artificial Intelligence Act, the Bletchley Declaration, the Committee on Publication Ethics, and the recommendations of the Organization for Economic Cooperation and Development. This article examines the pros and cons of artificial intelligence in scientific writing and analyzes relevant international regulations to present a path for the responsible, transparent, and effective use of this tool in advancing science. Iran's national guidelines for AI use in research emphasize transparency and full disclosure of its application.

Cite this article: Saboury, Ali Akbar. (2026). Research Acceleration vs. Authenticity Concern: Two Sides of the AI Coin in Academic Writing. *Science Cultivation*, 16 (1).36-46.



© Author(s) retain the copyright and full publishing rights.

Publisher: Foundation for the Advancement of Science and Technology in Iran (FAST-IRAN) and Iran Society of Biophysical Chemistry (ISOBC)

* Author for Correspondence, Professor, Tel: 98 2166956984, Fax: 98 2166404680, E-mail: Saboury@ut.ac.ir

¹ Institute of Biochemistry and Biophysics, University of Tehran, Tehran, Iran

شتاب‌بخشی پژوهش و دغدغه اصالت: دو روی سکه هوش مصنوعی در مقاله‌نویسی

علی اکبر صبوری^{۱*}

چکیده

هوش مصنوعی به‌عنوان قدرتمندترین ابزار فناورانه قرن حاضر، مرزهای سنتی تولید دانش و پژوهش علمی را درنور دیده است. این فناوری با ارائه قابلیت‌هایی چون پردازش فوق‌العاده داده‌ها، خودکارسازی وظایف تکراری و کمک به تحلیل‌های پیچیده، فرصتی بی‌سابقه برای تسریع و تعمیق اکتشافات علمی فراهم کرده است. بااین‌حال، به‌کارگیری آن در نوشتن مقالات علمی، دستاوردهای قابل‌توجهی را با چالش‌های اخلاقی، روش‌شناختی و حقوقی جدی همراه ساخته است. مسائلی مانند کاهش شفافیت و اصالت پژوهش، تقویت سوگیری‌های موجود، نقض مالکیت فکری و کاهش مهارت‌های تحلیلی پژوهشگران، نمونه‌ای از این دشواری‌هاست. در مواجهه با این پیچیدگی، جامعه جهانی با تدوین چارچوب‌های قانونی و مقرراتی نظیر قانون هوش مصنوعی اتحادیه اروپا، بیانیه بلچلی، کمیته اخلاق نشر و توصیه‌های سازمان همکاری اقتصادی و توسعه، در تلاش است تا موازنه‌ای بین تشویق نوآوری و مدیریت ریسک‌های ناشی از این فناوری برقرار کند. این مقاله با بررسی دو سویه مزایا و معایب هوش مصنوعی در نگارش علمی و تحلیل مقررات بین‌المللی مرتبط، مسیر بهره‌گیری مسئولانه، شفاف و مؤثر از این ابزار را در خدمت پیشرفت علم ارائه می‌کند. راهنماهای ملی ایران بر استفاده از هوش مصنوعی در پژوهش بر شفافیت و افشای کامل چگونگی استفاده از آن نیز تأکید دارند.

اطلاعات مقاله

نوع مقاله:

مقاله ترویجی

تاریخچه مقاله:

تاریخ دریافت:

۱۲ بهمن ۱۴۰۴

تاریخ بازنگری:

۲۰ بهمن ۱۴۰۴

تاریخ پذیرش:

۲۲ بهمن ۱۴۰۴

تاریخ انتشار:

۳۱ خرداد ۱۴۰۵

کلیدواژه‌ها:

اخلاق پژوهشی

قانون‌گذاری هوش

مصنوعی

نوشتن با کمک

هوش مصنوعی

مدل‌های زبانی بزرگ

صداقت علمی

کمیته اخلاق نشر

استناد: صبوری، علی اکبر. (۱۴۰۵). شتاب‌بخشی پژوهش و دغدغه اصالت: دو روی سکه هوش مصنوعی در مقاله‌نویسی. *نشاء علم*، ۱۶ (۱)، ۳۶-۴۶.



ناشر: بنیاد پیشبرد علم و فناوری در ایران و انجمن بیوشیمی فیزیک ایران
© نویسندگان حق نشر و کلیه حقوق انتشار را برای خود حفظ می‌کنند.

* عهده‌دار مکاتبات، استاد ممتاز، تلفن: ۶۶۹۵۶۹۸۴ (۹۸۲۱)، دورنگار: ۶۶۴۰۴۶۸۰ (۹۸۲۱)، آدرس الکترونیکی saboury@ut.ac.ir

^۱ مرکز تحقیقات بیوشیمی و بیوفیزیک، دانشگاه تهران، ایران

۱. مقدمه: انقلابی دوگانه در آستانه علم

ظهور هوش مصنوعی به‌ویژه در شکل مدل‌های زبانی بزرگ، نقطه عطفی در تاریخ روش‌شناسی علمی محسوب می‌شود. سرعت، مقیاس و پیچیدگی پژوهش‌های علمی امروز، گاه از توانایی‌های تحلیلی محض انسانی فراتر رفته است. در این میان، هوش مصنوعی با وعده خودکارسازی وظایف وقت‌گیر، پردازش داده‌های کلان^۱ و کشف الگوهای پنهان، به یاری پژوهشگران شتافته و دریچه‌های جدیدی به‌سوی شناخت گشوده است. ازسوی دیگر، این ابزار قدرتمند، سؤالات بنیادینی درباره اصالت فکری، مسئولیت اخلاقی، شفافیت روش‌شناختی و آینده نقش پژوهشگر برانگیخته است. جامعه علمی در آستانه یک تحول پارادایمی قرار دارد؛ تحولی که موفقیت آن در گرو شناخت عمیق قابلیت‌ها و محدودیت‌های هوش مصنوعی، و مهم‌تر از آن، طراحی چارچوب‌های حکمرانی هوشمند و بین‌المللی است که بتواند نوآوری را بدون قربانی کردن صحت، امانت‌داری و اعتبار علم هدایت کند. این مقاله در پی واکاوی این جنبه‌های دوگانه و بررسی پاسخ‌های نظارتی در حال ظهور در سطح جهان است.

۲. مزایای هوش مصنوعی در نوشتن و تولید مقالات علمی

هوش مصنوعی با تبدیل شدن به یک همکار پژوهشی، مزایای ملموسی را برای جامعه علمی به ارمغان آورده است که مهمترین آنها عبارت است از:

۱- افزایش کارایی و سرعت پژوهش: هوش مصنوعی می‌تواند کارهای زمان‌بری مانند مرور ادبیات پژوهش، خلاصه‌سازی مقالات، ترجمه متون تخصصی و حتی پیش‌نویس اولیه بخش‌هایی از مقاله را با سرعتی باورنکردنی انجام دهد. این امر زمان ارزشمند پژوهشگر را برای تفکر انتقادی، طراحی آزمایش و تفسیر عمیق‌تر نتایج آزاد می‌کند. ابزارهای هوش مصنوعی برخلاف انسان، هرگز خسته نمی‌شوند، نیاز به استراحت ندارند و می‌توانند به طور مداوم و بدون وقفه کار کنند. این ویژگی باعث می‌شود فرآیندهای پژوهشی وقت‌گیری که نیاز به پردازش پیوسته دارند (مانند تحلیل مجموعه‌داده‌های بسیار بزرگ، اجرای شبیه‌سازی‌های طولانی یا فرآیندهای غربالگری خودکار) با سرعت بسیار بیشتری نسبت به زمانی که فقط توسط نیروی انسانی انجام می‌شد، به‌پیش روند. در نتیجه، زمان کلی موردنیاز برای تکمیل یک پژوهش کاهش یافته و بازدهی افزایش می‌یابد.

۲- تحلیل داده‌های پیچیده و کلان: در عصر زیست‌شناسی ژنومی، ستاره‌شناسی رادیویی یا علوم اجتماعی کل‌نگر، حجم داده‌ها غالباً برای تحلیل انسانی غیرقابل‌مدیریت است. الگوریتم‌های هوش مصنوعی می‌توانند این داده‌های کلان را پردازش، الگوهای نامرئی را کشف و روابط پیچیده را مدل‌سازی کنند. این توانایی، به‌ویژه در رشته‌هایی مانند کشف دارو، پیش‌بینی تغییرات آب‌وهوایی و تحلیل شبکه‌های اجتماعی، منجر به بینش‌هایی کاملاً جدید شده است.

۳- کاهش خطاهای انسانی و سوگیری (در شرایط ایدئال): یکی از مزایای بالقوه هوش مصنوعی، انجام محاسبات و عملیات تکراری با دقت فوق‌العاده بالا است که می‌تواند خطاهای ناشی از خستگی یا سهل‌انگاری انسانی را کاهش دهد. همچنین، در تئوری، یک سیستم هوش مصنوعی که بر اساس داده‌های آموزشی بی‌طرفانه و جامع آموزش دیده باشد، می‌تواند در تصمیم‌گیری‌هایی مانند ارزیابی اولیه مقالات یا غربالگری داده‌ها، از تعصبات ناخودآگاه انسانی فاصله بگیرد. این ویژگی می‌تواند به عادلانه‌تر شدن فرآیند داوری و ارزیابی علمی کمک کند.

۴- دستیار شخصی‌سازی شده پژوهش: هوش مصنوعی می‌تواند به‌عنوان یک دستیار پژوهشی هوشمند عمل کند که با یادگیری حوزه تخصصی و سبک نویسنده، در ایده‌پردازی، بهبود ساختار نگارش، کنترل گرامر و استانداردسازی فرمت‌بندی ارجاعات کمک کند. این حمایت می‌تواند کیفیت صوری و حتی محتوایی مقالات را ارتقا دهد.

جدول ۱: مزایای کلیدی هوش مصنوعی در فرآیند پژوهش و تولید علمی

مزیت	شرح	مثال کاربردی
کارایی و سرعت	خودکارسازی وظایف اداری و مقدماتی پژوهش	مرور خودکار ادبیات، ترجمه، فرمت‌بندی
تحلیل داده‌های کلان	پردازش مجموعه‌داده‌های بسیار بزرگ و پیچیده	تحلیل تصاویر پزشکی، داده‌های ژنومی، شبیه‌سازی‌های آب‌وهوایی
دقت و عینیت (بالقوه)	کاهش خطاهای محاسباتی و عملیاتی تکراری	پردازش داده‌های آزمایشگاهی، تحلیل‌های آماری پیچیده
دستیاری هوشمند	کمک به سازماندهی، نگارش و بهبود مقاله	پیشنهاد ساختار، کنترل انسجام متنی، کمک به بازنویسی و ساختارسازی اصولی (پارافریز) متن

^۱ Big Data

کرد، بر عهده پژوهشگر انسانی باقی می‌ماند. همچنین، حقوق مالکیت فکری بر آثار تولیدشده یا کمک گرفته شده از هوش مصنوعی، و نیز مسئولیت خطاها یا عواقب منفی پژوهش، در هاله‌ای از ابهام قرار دارد.

۵- مشکلات شفافیت و تکثیرپذیری: بسیاری از مدل‌های پیشرفته هوش مصنوعی جعبه سیاه^۳ هستند؛ یعنی فرآیند استنتاج و تصمیم‌گیری آن‌ها برای انسان قابل‌درک نیست. این عدم شفافیت، ارزیابی منتقدانه روش‌شناسی و تکرارپذیری^۴ نتایج پژوهش را که از ارکان علم است، با مشکل مواجه می‌سازد. از نظر اصول اخلاقی، فقط نتایج قابل تکرار قابلیت انتشار دارند.

۴. چارچوب‌های قانونی و مقررات بین‌المللی هوش

مصنوعی در پژوهش

در پاسخ به چالش‌های پیچیده فناوری هوش مصنوعی، نهادها و کشورهای مختلف در حال تدوین مقرراتی برای حکمرانی بر توسعه و استفاده از آن هستند. این تلاش‌ها در سطح بین‌المللی در حال شکل‌گیری است و آگاهی از آنها برای هر پژوهشگری که قصد استفاده از این ابزار را دارد، ضروری است.

۴-۱. اتحادیه اروپا و قانون پیشگامانه هوش مصنوعی (AI Act)^۵: قانون هوش مصنوعی اتحادیه اروپا در سال ۲۰۲۱ میلادی پیشنهاد شد و پس از طی روند تصویب، سرانجام در اوت ۲۰۲۴ به طور رسمی لازم‌الاجرا شد. مقررات اصلی آن به تدریج تا سال ۲۰۲۷ اجرایی خواهند شد. این قانون، نخستین چارچوب نظارتی جامع را ایجاد کرده است و همان‌طور که جدول ۲ نشان می‌دهد، رویکردی مبتنی بر ریسک دارد [۱]. براین اساس، سیستم‌های هوش مصنوعی باتوجه به میزان خطری که برای سلامت، ایمنی و حقوق افراد ایجاد می‌کنند، دسته‌بندی و تنظیم می‌شوند. هدف این قانون تضمین ایمنی، شفافیت و محوریت انسان در توسعه و استفاده از سیستم‌های هوش مصنوعی در بازار اتحادیه اروپا است. هسته اصلی این قانون، یک رویکرد مبتنی بر ریسک است که سیستم‌های هوش مصنوعی را بر اساس میزان خطری که برای ایمنی، سلامت و حقوق اساسی افراد ایجاد می‌کنند، دسته‌بندی و

۳. معایب استفاده از هوش مصنوعی در پژوهش و

تولید علمی

علی‌رغم مزایای چشمگیر، وابستگی بی‌قیدوشرط به هوش مصنوعی در پژوهش، خطرات و چالش‌های نگران‌کننده‌ای را به همراه دارد، از جمله:

۱- تهدید اصالت پژوهشی: بزرگ‌ترین خطر، تولید محتوای غیراصیل، تقلیدی یا کاملاً ساختگی و تخیلی^۱ توسط هوش مصنوعی است. این سیستم‌ها بر اساس الگوهای آماری عمل می‌کنند و قادر به درک حقیقت یا ایجاد دانش نو به معنای انسانی آن نیستند. مقالات تولیدشده ممکن است ظاهری فریبنده اما فاقد بینش اصیل داشته باشند. همچنین، استفاده بدون ذکر منبع از محتوای تولیدشده توسط هوش مصنوعی، مصداق سرقت علمی^۲ و تقلب در پژوهش محسوب می‌شود.

۲- تشدید سوگیری و نابرابری: برخلاف ادعای عینیت، سیستم‌های هوش مصنوعی اغلب سوگیری‌های موجود در داده‌های آموزشی خود را بازتولید و حتی تشدید می‌کنند. اگر داده‌های آموزشی عمدتاً از مقالات پژوهشگران غربی یا متعلق به مؤسسات خاصی گردآوری شده باشد، خروجی سیستم نیز همان دیدگاه‌ها را ترجیح خواهد داد. این امر می‌تواند به حاشیه‌رانی بیشتر صداها و علمی از کشورهای درحال توسعه یا حوزه‌های پژوهشی کمتر متداول بینجامد.

۳- کاهش مهارت‌های تحلیلی و انتقادی پژوهشگران: اتکای بیش‌ازحد به هوش مصنوعی برای تحلیل داده، نوشتن و حتی استدلال، می‌تواند به فرسایش مهارت‌های بنیادین پژوهشگران مانند تفکر انتقادی، نگارش عمیق و طراحی روش‌شناسی منجر شود. این مسئله در درازمدت می‌تواند خلاقیت و استقلال فکری جامعه علمی را تضعیف کند.

۴- چالش‌های اخلاقی و حقوقی: هوش مصنوعی فاقد قضاوت اخلاقی و حس مسئولیت‌پذیری است. تصمیم‌گیری درباره اینکه چه پژوهشی اخلاقی است یا چگونه باید نتایج حساس را گزارش

^۱- Hallucination

^۲- Plagiarism

^۳- Black Box

^۴- Reproducibility

^۵- Artificial Intelligence Act (AI Act)

تنظیم می‌کند. این قانون چهار سطح ریسک اصلی دارد که برای هر کدام قوانین خاصی وضع شده است (جدول ۱). این قانون تأکید ویژه‌ای بر شفافیت دارد و برای «سیستم‌های عمومی هوش مصنوعی» (مانند مدل‌های زبانی بزرگ)، الزام به افشای اطلاعات جامع درباره داده‌های آموزشی و مطابقت با حقوق مالکیت فکری را مقرر می‌کند [۱].

۲-۴. ایالات متحده آمریکا و رویکرد بخش‌محور: برخلاف رویکرد یکپارچه اروپا، ایالات متحده فاقد یک قانون جامع فدرال در زمینه هوش مصنوعی است و عمدتاً از رویکرد بخش‌محور^۱ پیروی می‌کند. بااین‌حال، «فرمان اجرایی رئیس‌جمهور در زمینه امنیت، ایمنی و اعتماد هوش مصنوعی» (اکتبر ۲۰۲۳) گام مهمی برای هماهنگی تلاش‌هاست [۲]. این فرمان به آژانس‌های فدرال دستور می‌دهد تا استانداردها و خط‌مشی‌های ایمنی را توسعه دهند، از حریم خصوصی حمایت کنند و برابری و حقوق مدنی را ارتقا بخشند. این دستورالعمل‌ها اگرچه به طور مستقیم پژوهش دانشگاهی را تنظیم نمی‌کنند، اما جهت‌گیری کلی حکومت فدرال را نشان می‌دهند و بر محیط تحقیقاتی تأثیر می‌گذارند.

۳-۴. تلاش‌های چندجانبه و چارچوب‌های بین‌المللی: جامعه جهانی برای ایجاد هماهنگی در حال تلاش است. بیانیه بلجلی^۲ که در نوامبر ۲۰۲۳ در پایان نخستین اجلاس جهانی/ایمنی هوش مصنوعی در بریتانیا منتشر شد، سندی مهم است. این اعلامیه توسط ۲۸ کشور از جمله چین، ایالات متحده و کشورهای عضو اتحادیه اروپا امضا شد. موضوع اصلی این بیانیه، شناسایی فرصت‌ها و ریسک‌های مشترک جهانی ناشی از پیشرفته‌ترین و قدرتمندترین مدل‌های هوش مصنوعی مانند ChatGPT است [۳-۴]. امضاکنندگان توافق کردند که برای درک و کاهش این ریسک‌ها از جمله در حوزه‌هایی مانند سایبری و زیست‌فناوری، نیاز به همکاری بین‌المللی فوری و توسعه تحقیقات علمی مشترک وجود دارد. هدف نهایی تضمین توسعه و استفاده ایمن، شفاف و مسئولانه از هوش مصنوعی به نفع همه انسانها است.

همچنین، «توصیه‌های شورای سازمان همکاری اقتصادی و توسعه (OECD)^۳ در مورد هوش مصنوعی» که در سال ۲۰۱۹ به‌روزرسانی شد، یک چارچوب سیاستی بین‌المللی برای «هوش مصنوعی قابل‌اعتماد» ارائه می‌دهد. این اصول شامل شفافیت و افشا، انصاف، پاسخگویی و طراحی ایمن است و بیش از پنجاه کشور آن را تصویب کرده‌اند [۵] که معیارهای اخلاقی و حرفه‌ای مهمی را برای جامعه علمی جهانی تعیین می‌کنند.

جدول ۲: سطح ریسک و الزامات قانونی هوش مصنوعی اتحادیه اروپا

سطح ریسک	مثال کاربردی	الزامات قانونی
ریسک غیرقابل‌تبریر	سیستم‌های دستکاری ناپیدا (مانند سوءاستفاده از آسیب‌پذیری افراد)، سیستم‌های نمره‌دهی اجتماعی (Social Scoring)، شناسایی بیومتریک زنده در فضاهای عمومی برای مجریان قانون (با استثنائات محدود).	ممنوعیت کامل (این ممنوعیت‌ها از فوریه ۲۰۲۵ اعمال شد)
ریسک بالا	سیستم‌های ایمنی در زیرساخت‌های حیاتی (مانند حمل‌ونقل، دستگاه‌های پزشکی، استخدام، مهاجرت، اجرای قانون و دادرسی قضایی)	سخت‌گیرانه‌ترین مقررات؛ از جمله ارزیابی ریسک، کیفیت داده‌ها، مستندسازی دقیق، نظارت انسانی و کسب نشان CE (از اوت ۲۰۲۶ قابل اجرا).
ریسک محدود	چت‌بات‌ها و سیستم‌های تولید محتوا (مانند دیپ‌فیک).	الزامات شفافیت؛ کاربران باید آگاه باشند که با یک سیستم هوش مصنوعی در تعامل هستند (از اوت ۲۰۲۶ قابل اجرا)
ریسک کم یا صفر	فیلترهای هرزنامه یا بازی‌های ویدیویی مبتنی بر هوش مصنوعی	بدون مقررات خاص

^۱- Sectoral

^۲- Bletchley Declaration

^۳- Organization for Economic Co-operation and Development (OECD)

سوگیری در الگوریتم‌ها، برای حفظ یکپارچگی علمی ضروری است.

۳- گفت‌وگوی مستمر بین‌المللی و هماهنگی مقرراتی: اختلاف در رویکردهای نظارتی (مانند اتحادیه اروپا و آمریکا) می‌تواند همکاری علمی جهانی را مختل کند. تقویت گفت‌وگو در پلتفرم‌هایی مانند OECD و سازمان ملل برای همگرایی در اصول کلیدی، امری حیاتی است.

۴- تقویت نقش پژوهشگر به‌عنوان رهبر اخلاقی: در نهایت، هوش مصنوعی یک ابزار است و قضاوت نهایی درباره ارزش، اخلاق‌مندی و جهت‌گیری علم با پژوهشگر انسانی است. تقویت آموزش اخلاق پژوهش و تفکر انتقادی در کنار آموزش فناوری، کلید هدایت این تحول در مسیری درست است.

علم در تعامل پویا با فناوری‌های جدید بالیده است. هوش مصنوعی نیز آزمونی بزرگ برای تاب‌آوری، انعطاف‌پذیری جامعه علمی است. با پذیرش فرصت‌ها، شناخت خطرات و ایجاد چارچوب‌های مسئولانه حکمرانی، می‌توان از این فناوری نه به‌عنوان جایگزینی برای نبوغ انسانی، که به‌عنوان قدرتمندترین هم‌پیمان برای گشودن رازهای ناشناخته جهان بهره گرفت. آینده علم، آینده همکاری هم‌افزا بین هوش انسانی و هوش ماشینی خواهد بود.

۶. بد اخلاقی‌هایی در استفاده از هوش مصنوعی

استفاده غیرمسئولانه از هوش مصنوعی در نوشتن مقالات علمی، طیفی از بد اخلاقی‌های پژوهشی جدی را به وجود آورده که بنیان اعتماد و اصالت علم را تهدید می‌کند. بارزترین این تخلفات، سرقت علمی پنهان^۳ و تولید محتوای غیراصیل است. زمانی که پژوهشگر از خروجی یک مدل زبانی بزرگ^۴ (LLM)، بدون ذکر منبع یا با ادعای نویسنده‌گی خود استفاده می‌کند، مرتکب سرقت

۴-۴. پیامدهای مقرراتی برای پژوهش علمی: این قوانین و چارچوب‌ها تأثیر مستقیمی بر پژوهش علمی دارند. اصل کلیدی که از دل همه آنها برمی‌آید، «شفافیت» است. این بدان معناست که پژوهشگران موظف‌اند به‌وضوح نمایان^۱ سازند که چگونه، کجا و با کدام ابزار هوش مصنوعی در فرآیند تحقیق، تحلیل داده یا نگارش مقاله کمک گرفته‌اند. بسیاری از نشریات علمی معتبر اکنون این افشا را اجباری کرده‌اند. همچنین، آگاهی از قوانین مربوط به حفاظت از داده‌ها (مانند GDPR در اروپا)^۲ هنگام استفاده از هوش مصنوعی برای پردازش داده‌های شخصی در پژوهش، حیاتی است.

۵. آینده‌نگری

هوش مصنوعی در نگارش و تولید مقالات علمی، همچون یک تیغ دو لبه عمل می‌کند. از یک‌سو، با افزایش کارایی، تحلیل داده‌های پیچیده و کمک به فرآیند خلاقیت، پتانسیل آن را دارد که دورانی طلایی از اکتشافات علمی را رقم بزند. از سوی دیگر، خطرات جدی مربوط به اصالت، انصاف، شفافیت و مسئولیت‌پذیری را به همراه دارد که اگر مدیریت نشوند، می‌توانند بنیان‌های اعتماد عمومی به علم را تخریب کنند. مسیر به جلو نه در طرد این فناوری بسیار مهم که در ادغام خردمندانه و مسئولانه آن در اکوسیستم پژوهش نهفته است. این امر مستلزم اقدامات چندجانبه است:

- ۱- تدوین و اجرای راهبردهای نهادی: دانشگاه‌ها و ناشران علمی باید راهنمایی‌هایی روشن و اجباری برای بیان شفاف استفاده از هوش مصنوعی، تعریف سرقت علمی در این زمینه و آموزش مهارت‌های سواد دیجیتال-انتقادی به پژوهشگران وضع کنند.
- ۲- توسعه ابزارهای نظارتی فنی: ایجاد ابزارهای قابل اعتماد برای شناسایی محتوای تولیدشده توسط هوش مصنوعی و آزمون

^۱- Disclose

^۲- General Data Protection Regulation (GDPR)

مقررات عمومی حفاظت از داده‌ها (GDPR)، یک قانون جامع اتحادیه اروپا برای حفظ حریم خصوصی است که در سال ۲۰۱۸ میلادی اجرایی شد. این قانون، استانداردهایی را برای جمع‌آوری، پردازش و ذخیره‌سازی داده‌های شخصی (مانند نام، آدرس ایمیل یا اطلاعات مکانی) افراد مقیم منطقه اقتصادی اروپا تعیین می‌کند. مهم‌ترین هدف از ارائه این مقررات، اعطای کنترل بیشتر به افراد بر داده‌های شخصی خود است. این قانون بر پایه اصولی مانند رضایت آگاهانه و صریح، شفافیت، محدودیت هدف و امنیت داده‌ها استوار است و به افراد حق‌هایی مانند دسترسی به داده‌ها، درخواست اصلاح، انتقال و حتی حذف آنها (حق فراموشی) می‌دهد. حوزه شمول این قانون بسیار گسترده است و صرفاً محدود به شرکت‌های مستقر در اتحادیه اروپا نیست، بلکه هر سازمانی در سراسر جهان که کالا یا خدماتی به ساکنان این منطقه ارائه می‌دهد یا داده‌های آنان را پردازش می‌کند را نیز شامل می‌شود. تخلف از این قانون می‌تواند جریمه‌های سنگینی تا بیست میلیون یورو یا چهار درصد از درآمد سالانه سازمان متخلف در پی داشته باشد. برای مطالعه دقیق‌تر می‌توانید به متن کامل و رسمی این مقررات [۶] که به زبان انگلیسی و به‌صورت آنلاین منتشر شده است مراجعه کنید.

^۳- Stealth Plagiarism

^۴- Large Language Model (LLM)

علمی شده است. این مسئله با پیچیدگی بیشتری همراه می‌شود، زیرا ابزارهای کنونی کشف سرقت علمی اغلب در شناسایی متون تولیدشده توسط هوش مصنوعی ناتوان هستند.

مدل‌های زبانی بزرگ نوعی از هوش مصنوعی هستند که روی حجم عظیمی از داده متنی آموزش دیده‌اند. این مدل‌ها (مانند Claude, Gemini, ChatGPT) توانایی فهم، تولید و پردازش زبان طبیعی را دارند. آنها می‌توانند متنی روان و مرتبط با موضوع تولید کنند، به سؤالات پاسخ دهند و یا متن ورودی را خلاصه یا ترجمه کنند. این مدل‌ها دقیقاً همان ابزار فنی اصلی هستند که امکان نوشتن مقاله با کمک هوش مصنوعی را فراهم می‌کنند. این ارتباط از طریق قابلیت‌های کلیدی زیر شکل می‌گیرد: تولید متن، درک و خلاصه‌سازی، ترجمه، پاسخ به سؤال، سازماندهی و ساختاردهی. LLM یک «همکار» است، نه «نویسنده». با وجود توانایی‌های بالا، LLM فاقد قضاوت علمی، درک واقعی و مسئولیت اخلاقی است. بنابراین نقطه قوت آن این است که یک دستیار توانمند برای افزایش بهره‌وری در وظایف مکانیکی و ایده‌پردازی اولیه است و نقطه‌ضعف آن غیرقابل‌اعتماد آن برای تولید دانش اصیل، تحلیل داده، استدلال علمی پیچیده یا ارائه نتیجه‌گیری‌های معتبر می‌باشد. این مدل‌ها مستعد تولید اطلاعات نادرست و تخیلی و تکرار سوگیری‌های موجود در داده‌های آموزشی خود، ساختن منابع، داده‌ها، اسنادها و حتی کل پژوهش‌های علمی به‌ظاهر معتبر اما کاملاً ساختگی هستند. در جمع‌بندی، LLM موتور اصلی هوش مصنوعی‌هایی است که در نوشتن مقاله کمک می‌کنند. آنها می‌توانند روند نوشتن را با خودکارسازی برخی مراحل تسریع کنند، اما کیفیت، صحت، اصالت و اخلاق‌مندی محتوای نهایی، کاملاً در گرو نظارت، بازبینی و تخصص پژوهشگر انسانی است. این مدل‌ها با تکیه بر این اطلاعات کذب، نه تنها اعتبار مقاله را از بین می‌برد، بلکه زنجیره دانش را آلوده می‌سازد و امکان تکثیرپذیری پژوهش را از بین می‌برد به بیان ساده، LLM قلم سریع و باسوادی است که شما آن را هدایت می‌کنید، نه یک پژوهشگر جایگزین.

استفاده اخلاقی و مؤثر از مدل‌های زبانی بزرگ در پژوهش، مستلزم رعایت اصل شفافیت مطلق و حفظ حاکمیت فکری پژوهشگر است. این ابزار را می‌توان به‌عنوان یک دستیار پژوهشی

هوشمند در مراحل مکانیکی و مقدماتی به کار گرفت، به شرطی که نقش آن صریحاً در بخش «روش‌شناسی» یا «تقدیر و تشکر» مقاله افشا شود. به‌عنوان مثال، از LLM می‌توان برای خلاصه‌سازی و اولویت‌بندی متون علمی در مرحله مرور ادبیات، ایجاد پیش‌نویس اولیه برای بخش‌های توصیفی، پاراگراف‌بندی اصولی جملات و پیشنهاد ساختار منطقی برای مقاله، ترجمه یا ویرایش زبان استفاده کرد. با این حال، هیچ‌یک از خروجی‌های LLM نباید بدون بازبینی عمیق و مسئولانه پژوهشگر پذیرفته شود. وظایفی مانند طراحی روش تحقیق، تحلیل داده‌ها، تفسیر نتایج، استدلال‌های پیچیده و نتیجه‌گیری نهایی که مستلزم قضاوت تخصصی، تفکر انتقادی و مسئولیت‌پذیری اخلاقی هستند، باید به طور کامل توسط خود پژوهشگر انجام شوند. استفاده نادرست و غیرشفاف از این ابزارها نه تنها مصداق سرقت علمی است، بلکه به اعتبار پژوهش و پژوهشگر خدشه وارد می‌سازد.

بد اخلاقی دیگر، تخریب فرآیند داوری و ارزیابی علمی است. پژوهشگران ممکن است از هوش مصنوعی برای تولید انبوه و سریع مقالات با حداقل تغییرات (مقاله‌سازی افزونی^۱ یا تکه‌تکه‌سازی) استفاده کنند تا شاخص‌های کمی خود را به شکلی مصنوعی افزایش دهند، امری که بار داوران و کیفیت کلی ادبیات علمی را تحت تأثیر منفی قرار می‌دهد. مقاله‌سازی افزونی، یک سوءرفتار پژوهشی است که در آن داده‌های حاصل از یک مطالعه واحد به شکلی غیرضروری به چندین مقاله کوچک‌تر یا «تکه» تقسیم می‌شود تا شمار انتشارات پژوهشگر به طور مصنوعی افزایش یابد. هدف اصلی از این کار، دستیابی به «حداقل واحد قابل‌انتشار» و انباشتن رزومه است. این عمل به دلایل زیر غیراخلاقی و نادرست محسوب می‌شود [۷]:

الف) تحریف ادبیات علمی: خوانندگان و تحلیلگران (به‌ویژه در فراتحلیل‌ها) ممکن است تصور کنند این مقالات حاصل از مطالعات مستقل با نمونه‌های جداگانه هستند، درحالی‌که داده‌ها یکسان است.

ب) اتلاف وقت و منابع: وقت داوران، ویراستاران و خوانندگان به هدر می‌رود و جای مقالات اصیل دیگر گرفته می‌شود.

ج) محروم کردن جامعه علمی از درک جامع: ارائه نتایج به‌صورت تکه‌تکه شده، مانع از درک پیام اصلی و یکپارچه مطالعه می‌شود.

^۱- Salami Slicing

۲- کمک به فرآیند ایده‌پردازی و سازماندهی: از هوش مصنوعی می‌توان برای خلاصه‌سازی ادبیات موجود، پیشنهاد کلیدواژه‌های مرتبط، یا تولید ایده برای ساختار منطقی مقاله (مانند پیشنهاد عناوین بخش‌ها) بهره برد. این کار می‌تواند نقطه شروعی برای کار عمیق‌تر نویسنده باشد.

۳- کمک‌های فنی خاص: در حوزه‌هایی مانند کدنویسی، تحلیل‌های آماری اولیه، یا شبیه‌سازی ساده، هوش مصنوعی می‌تواند در تولید کد یا شناسایی الگوها کمک کند. در علوم انسانی نیز می‌تواند در ترجمه متون یا دسته‌بندی اولیه داده‌های کیفی نقش داشته باشد.

۸. شرایط ضروری برای مجاز شمردن هر گونه استفاده:

۱- بیان صریح^۱: نویسنده موظف است در بخشی از مقاله (اغلب در قسمت «روش‌شناسی» یا «تقدیر و تشکر») به‌وضوح ذکر کند که از کدام ابزار هوش مصنوعی، برای کدام قسمت‌های فرآیند پژوهش (مثلاً ویرایش زبان، تولید ایده اولیه، تحلیل داده) استفاده شده است. این شفافیت، شرط اخلاقی اولیه است.

۲- بازبینی، تأیید و مسئولیت‌پذیری نهایی: نویسنده باید مسئولیت کامل محتوای نهایی مقاله را برعهده گیرد. این بدان معناست که هر خروجی هوش مصنوعی باید از نظر دقت واقعی، اصالت استدلال، صحت داده‌ها و ارجاعات به‌دقت بازبینی، تصحیح و تأیید شود. نویسنده در قبال هر خطا یا اظهارنظر در مقاله مسئول است.

۳- ممنوعیت تولید محتوای اصیل: استفاده از هوش مصنوعی برای نوشتن بخش‌های اصیل مقاله مانند چکیده، مقدمه، تحلیل نتایج، بحث و نتیجه‌گیری غیرمجاز است. این بخش‌ها باید حاصل تفکر انتقادی، درک عمیق و قلم منحصربه‌فرد نویسنده باشد.

۴- رعایت دستورالعمل‌های ناشر و مؤسسه: پیش از هر اقدامی، باید دستورالعمل‌های خاص ناشر مجله هدف و مؤسسه آموزشی یا پژوهشی مربوطه درباره استفاده از هوش مصنوعی مطالعه و رعایت شود.

در مجموع، استفاده مجاز از هوش مصنوعی، استفاده‌ای است که کارایی فنی را افزایش دهد؛ اما نبوغ، قضاوت و صداقت فکری پژوهشگر را در کانون فرآیند تولید علم نگه دارد. هوش مصنوعی

باین‌حال، شرایط خاصی هم وجود دارد که انتشار چند مقاله از یک مطالعه بزرگ قابل قبول است، مانند زمانی که هر مقاله فرضیه، روش‌شناسی یا جامعه آماری کاملاً متفاوت و مستقل را بررسی می‌کند (مثلاً در مطالعات همه‌گیر شناختی بسیار بزرگ). در این موارد، شفافیت کامل و ارجاع متقابل بین مقالات الزامی است. همچنین، برخی ممکن است از هوش مصنوعی برای تولید یا دست‌کاری تصاویر، نمودارها و داده‌های آزمایشگاهی (سوءرفتار داده‌ای) به‌منظور اثبات ادعاهای پژوهشی خود استفاده کنند. این اقدام، که با ظهور ابزارهای تولید تصویر پیچیده‌تر شده، نقض فاحش اصول بنیادین صداقت و صحت در علم است. این بد اخلاقی‌ها زمانی تشدید می‌شوند که پژوهشگران، هوش مصنوعی را به‌عنوان جانشینی برای تفکر انتقادی و درک عمیق موضوع پژوهش قرار می‌دهند. این امر منجر به تولید مقالاتی می‌شود که اگرچه از نظر صوری روان و بی‌عیب به نظر می‌رسند، اما فاقد بینش، نوآوری اصیل و ارزش علمی واقعی هستند و در بلندمدت به فرسایش مهارت‌های تحلیلی جامعه پژوهشی می‌انجامد.

۷. استفاده‌های مجاز از هوش مصنوعی در نوشتن

مقالات علمی

استفاده مسئولانه و شفاف از هوش مصنوعی در نوشتن مقالات علمی، در صورتی که به‌مثابه یک ابزار کمکی و نه یک نویسنده جایگزین در نظر گرفته شود، می‌تواند مجاز و حتی ارزشمند باشد. کلید این مسئولیت‌پذیری، رعایت اصل شفافیت مطلق و حفظ حاکمیت فکری پژوهشگر است.

استفاده‌های مجاز و سازنده از هوش مصنوعی شامل موارد زیر است:

۱- بهبود صوری و زبانی: استفاده از هوش مصنوعی برای ویرایش دستور زبان، اصلاح ساختار جملات، بهبود روانی نگارش و یکدست‌سازی سبک نوشتار (به‌ویژه برای نویسندگان غیرانگلیسی‌زبان) کاملاً قابل قبول است. این ابزار می‌تواند در قالب‌بندی اولیه ارجاعات بر اساس استانداردهای مختلف (مانند APA یا Vancouver) نیز کمک کند، مشروط بر اینکه صحت نهایی همه استنادات توسط نویسنده بررسی شود.

^۱- Disclosure

می‌تواند یک دستیار مفید باشد، اما هرگز نمی‌تواند جایگاه مؤلف و مسئول اثر علمی را پر کند.

۹. استفاده‌ها غیرمجاز از هوش مصنوعی در نوشتن

مقالات علمی

استفاده از هوش مصنوعی در نوشتن مقاله علمی، در موارد زیر کاملاً غیرمجاز و نقض آشکار اخلاق پژوهشی محسوب می‌شود:

۱- تولید محتوای اصیل و استدلال‌های علمی

غیرمجازترین استفاده، واگذاری تفکر، تحلیل و نگارش اصیل علمی به هوش مصنوعی است. این شامل:

- نوشتن بخش‌های محتوایی مانند چکیده، مقدمه، تحلیل نتایج، بحث یا نتیجه‌گیری توسط هوش مصنوعی.
- تولید ایده‌ها، فرضیه‌ها یا تفسیرهای علمی نو بدون تفکر انتقادی عمیق پژوهشگر.
- خلق داده‌ها، نتایج آزمایشگاهی یا یافته‌های شبیه‌سازی‌شده.
- هوش مصنوعی هرگز نباید برای ساخت یا دست‌کاری داده‌های تجربی استفاده شود.

۲- انجام وظایفی که نیاز به قضاوت تخصصی انسانی دارند:

- ارزیابی و انتخاب منابع علمی: اتکا به هوش مصنوعی برای قضاوت درباره کیفیت، ارتباط یا اعتبار منابع پژوهشی.
- استنتاج روابط علی یا تفسیر معناداری آماری بدون نظارت و تأیید متخصص حوزه.
- تصمیم‌گیری‌های اخلاقی در طراحی پژوهش یا تفسیر پیامدهای اجتماعی یافته‌ها.

۳- تقلب آشکار و تخریب یکپارچگی علمی:

- تولید متن‌های خلاقانه برای فرار از شناسایی سرقت علمی (پارافریز^۱ غیراصیل و اتوماتیک). منظور از پارافریز غیراصیل و اتوماتیک، بازنویسی یک متن توسط نرم‌افزار یا هوش

مصنوعی است که صرفاً کلمات و ساختار جملات را به شکل مکانیکی تغییر می‌دهد، بدون اینکه نویسنده درک مفهومی از متن داشته یا محتوای جدیدی به آن بیفزاید.

- ایجاد مقاله‌های ساختگی کامل یا تولید انبوه مقالات کم محتوا (مقاله‌سازی سَلَمی) فقط برای افزایش رزومه.

- جعل هویت علمی: استفاده از هوش مصنوعی برای نوشتن مقاله‌ای که ادعا می‌کند نویسنده‌ای با تخصص خاص (مثلاً پزشک یا مهندس) آن را نوشته، درحالی‌که نویسنده فاقد آن صلاحیت است.

خط قرمز مطلق: نقض مسئولیت‌پذیری و شفافیت:

هرگونه استفاده مخفیانه و غیرشفاف از هوش مصنوعی که در آن نقش هوش مصنوعی در هر مرحله از پژوهش افشا نشود. پژوهشگر مسئولیت کامل محتوا را نپذیرد و آن را به هوش مصنوعی نسبت دهد.

دلیل اصلی غیرمجاز بودن این موارد آن است که هوش مصنوعی: فاقد قضاوت علمی، درک مفهومی و مسئولیت‌پذیری اخلاقی است. مستعد تولید اطلاعات نادرست یا ساختگی (هالوسینیشن^۲) است. سوگیری‌های موجود در داده‌های آموزشی را تقویت می‌کند.

مهارت‌های تحلیلی و نوآوری اصیل پژوهشگر را تضعیف می‌کند. استفاده غیرمجاز، تنها یک «تقلب فنی» نیست، بلکه نادرستی به ماهیت جستجوی علمی است که بر پایه صداقت، اصالت فکری و مسئولیت‌پذیری انسان‌ها بنا شده است.

۱۰. راهنمای کمیته اخلاق نشر بین‌المللی در استفاده

از هوش مصنوعی

بر اساس آخرین راهنمای کمیته اخلاق نشر^۳ (COPE) سال ۲۰۲۵، در مورد هوش مصنوعی، محورهای اصلی زیر مورد تأکید قرار گرفته‌اند [۸]:

۱- پارافریز (Paraphrase) به معنای بازنویسی یک متن یا ایده با کلمات و ساختار جمله‌ای جدید است، طوری که معنای اصلی آن به طور کامل و دقیق حفظ شود. در فارسی به آن «بازنویسی» یا «بیان مجدد» نیز می‌گویند.

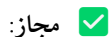
² Hallucination

³- Committee On Publication Ethics (COPE)

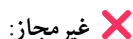
کمیته اخلاق نشر (COPE)، یک نهاد بین‌المللی غیرانتفاعی است که در سال ۱۹۹۷ میلادی برای ارتقای استانداردهای اخلاقی در انتشار پژوهش‌های علمی تأسیس شد. هدف اصلی این کمیته حمایت از سردبیران، ناشران و پژوهشگران در مواجهه با چالش‌های اخلاقی نشر و تدوین راهنماهای مناسب است. این کمیته برای تخلفات رایج مانند سرقت علمی، جعل داده، ارسال همزمان مقاله به چند مجله و نویسندگی جعلی، چارچوبها و نمودارهای تصمیم‌گیری ارائه می‌کند تا رسیدگی به این موارد به صورت منصفانه و شفاف انجام شود. کمیته با ارائه آموزش، مشاوره و برگزاری سمینار، به حفظ سلامت و اعتماد در فرآیند انتشار علمی کمک می‌کند و امروزه هزاران نشریه معتبر از سراسر جهان عضو آن هستند. برای مطالعه دقیق‌تر و دسترسی به منابع رسمی آن، به تارنمای کمیته مراجعه کنید: (<https://publicationethics.org>)

- ۱- شفافیت و بیان شفاف^۱
- الزام به بیان کامل استفاده از هرگونه ابزار هوش مصنوعی در فرآیند پژوهش

نمونه‌های عملی از راهنمای COPE:



- استفاده از ChatGPT برای بهبود گرامر و روان‌سازی متن.
- به کارگیری AI برای ترجمه متون (با ذکر منبع).
- استفاده از ابزارهای AI برای شناسایی الگو در داده‌ها.



- انتشار متون تولیدشده توسط AI بدون بازبینی و ویرایش.
 - استفاده از AI برای تولید داده‌های پژوهشی.
 - استناد به ChatGPT به عنوان مرجع علمی.
- این راهنما به صورت پویا به‌روزرسانی می‌شود و آخرین نسخه آن مستقیماً از وبسایت COPE در دسترس است [۸].

۱۱. قوانین موجود در ایران

- بر اساس «راهنمای استفاده از ابزارهای هوش مصنوعی در پژوهش» وزارت علوم، تحقیقات و فناوری، استفاده از هوش مصنوعی به عنوان یک ابزار کمکی در تحقیقات علمی و نوشتن مقاله علمی مجاز است، مشروط بر رعایت اصل **شفافیت مطلق** و حفظ **مسئولیت نهایی انسانی**. این راهنما استفاده اخلاقی را به دو دسته مجاز و غیرمجاز تقسیم می‌کند [۹]:

استفاده‌های مجاز شامل اموری است که در آن هوش مصنوعی نقش دستیار را دارد، مانند: ایده‌پردازی اولیه، خلاصه‌سازی مقالات، ترجمه و ویرایش متن، تحلیل داده، تولید کد ساده یا کمک در ساختاردهی نگارش. باین‌حال، محقق را موظف به‌وضوح در بخش «تقدیر و تشکر» یا «بیانیه استفاده از هوش مصنوعی» مقاله می‌نماید. نوع ابزار و نحوه استفاده از آن باید افشا شده و خروجی آن به‌دقت بازبینی و تأیید شود.

- ذکر دقیق ابزارها و نسخه‌ها (مانند ChatGPT-4, Mid Journey)
- تشریح نقش AI در هر مرحله از پژوهش (ایده‌پردازی، طراحی، تحلیل داده، نگارش)

۲- مسئولیت‌پذیری و پاسخگویی^۲

- پژوهشگران مسئول نهایی تمام محتوای تولیدشده توسط AI هستند.
- عدم پذیرش AI به عنوان نویسنده به دلیل ناتوانی در پذیرش مسئولیت.
- ضرورت بررسی دقیق خروجی‌های AI از نظر خطا، سوگیری و سرقت علمی.

۳- یکپارچگی محتوایی^۳

- ممنوعیت استفاده از AI برای تولید یا دستکاری داده‌های پژوهشی.
 - اجتناب از تولید تصاویر، نمودارها یا نتایج مصنوعی توسط AI.
 - حفظ اصالت اندیشه و جلوگیری از وابستگی بیش‌ازحد به AI.
- ۴- حریم خصوصی و امنیت داده^۴
- ممنوعیت بارگذاری داده‌های حساس در سکوهاي AI عمومی.
 - رعایت مقررات حفاظت داده (مانند GDPR) در استفاده از AI.
 - ارزیابی ریسک امنیتی سکوهاي AI قبل از استفاده.

۵- سوگیری و انصاف^۵

- شناسایی و گزارش سوگیری‌های احتمالی در الگوریتم‌های AI.
- اجتناب از تداوم تبعیض از طریق سیستم‌های AI.
- تنوع داده‌های آموزشی و توجه به نمایندگی گروه‌های مختلف.

۶- مالکیت فکری و استناد^۶

- عدم استناد به AI به عنوان منبع اولیه.
- رعایت حقوق مالکیت محتوای استفاده‌شده برای آموزش AI.
- شفافیت در تعارض منافع هنگام استفاده از ابزارهای AI تجاری.

۷- نظارت و ارزیابی مستمر^۶

¹- Transparency and Disclosure

²- Accountability

³- Content Integrity

⁴- Privacy and Data Security

⁴- Bias and Fairness

⁵- Intellectual Property and Citation

⁶- Oversight and Continuous Evaluation

انتقادی، طراحی خلاقانه آزمایش ها و تعمیق تحلیل ها فراهم آورد و در نتیجه چرخه تولید علم را شتاب بخشد.

با این حال، استفاده از هوش مصنوعی در نوشتن مقالات علمی، بدون چالش نیست. تهدید اصالت و یکتایی پژوهش، تشدید سوگیری های موجود، کاهش مهارت های تحلیلی پژوهشگران، مسائل پیچیده اخلاقی و حقوقی و مشکل شفافیت و تکثیر پذیری، از جمله دشواری های جدی پیش روی جامعه علمی هستند. هسته اصلی این چالش ها، حفظ اعتماد، اصالت و مسئولیت پذیری در فرآیند علم سازی است. هوش مصنوعی فاقد قضاوت اخلاقی، درک واقعی و مسئولیت پذیری ذاتی است و بنابراین نمی تواند و نباید جایگاه مؤلف و مسئول نهایی یک اثر علمی را پر کند.

در پاسخ به این چالش ها، تلاش های نظارتی و قانونی در سطح بین المللی در حال شکل گیری است. قانون هوش مصنوعی اتحادیه اروپا (AI Act) با رویکرد مبتنی بر ریسک، بیانیه بلجلی با تأکید بر همکاری بین المللی برای مدیریت ریسک های فناوری های پیشرفته و راهنمای سازمان همکاری اقتصادی و توسعه (OECD) برای هوش مصنوعی قابل اعتماد، نمونه هایی از این تلاش ها هستند. در سطح ملی ما نیز، دستورالعمل هایی مانند «راهنمای استفاده از ابزارهای هوش مصنوعی در پژوهش» وزارت علوم، تحقیقات و فناوری ایران، سعی در تعریف چارچوبی برای استفاده مسئولانه دارند. اصل کلیدی مشترک در همه این مقررات، **شفافیت مطلق** است. پژوهشگر ملزم است به صراحت افشا کند که چگونه، در کدام مرحله و با کدام ابزار هوش مصنوعی از این فناوری کمک گرفته است.

نتیجه آنکه، آینده مطلوب در گرو **استفاده هوشمندانه، مسئولانه و شفاف** از هوش مصنوعی است. این فناوری باید در نقش یک **دستیار پژوهشی** قدرتمند باقی بماند که مراحل مکانیکی و زمان بر را تسهیل می کند، نه یک **جانشین** برای تفکر نقاد، خلق ایده های اصیل، تحلیل عمیق و قضاوت تخصصی پژوهشگر. پژوهشگر باید همواره کنترل، نظارت و مسئولیت نهایی محتوا، تفسیر نتایج و رعایت اصول اخلاقی را بر عهده داشته باشد. تنها از این طریق است که می توان از قابلیت های شگفت انگیز هوش مصنوعی برای

در مقابل، **استفاده های غیرمجاز** مواردی است که جانشین قضاوت و خلاقیت انسانی می شود. این موارد شامل تولید کل مقاله یا بخش های اصلی (مانند چکیده، بحث و نتیجه گیری)، جعل داده ها یا منابع، ایجاد تصاویر یا نتایج ساختگی و استفاده از آن به صورت عدم افشاسازی است. به طور خلاصه، می توان از هوش مصنوعی برای تسریع کار استفاده نمود، اما نه برای تفکر، تحلیل عمیق و خلق دانش اصیل؛ این امور باید همواره بر عهده پژوهشگر باقی بماند.

بر اساس دستورالعمل **بنیاد ملی علم ایران** [۱۰]، استفاده از ابزارهای هوش مصنوعی در پژوهش مجاز است، اما با شرایط و محدودیت های دقیق و به عنوان یک ابزار کمکی، نه جایگزینی برای پژوهشگر. اصول کلیدی آن رعایت اخلاق علمی، شفافیت کامل، مسئولیت پذیری پژوهشگر و حفظ حریم خصوصی است. در این دستورالعمل هم آمده است که استفاده از هوش مصنوعی برای ایده پردازی، ترجمه، ویرایش، تحلیل داده، تولید کد و بهبود نگارش مجاز است. با این حال، پژوهشگر موظف است به صراحت در بخش روش شناسی یا بیانیه اخلاق پروژه، نام ابزار، نسخه و تاریخ استفاده را ذکر کرده و تمام خروجی ها را بازبینی و اعتبارسنجی علمی نماید. همچنین، جایگزینی کامل نقش پژوهشگر، تولید داده یا منابع ساختگی و استفاده از داده های حساس در ابزارهای عمومی ممنوع است. در نهایت، مسئولیت نهایی محتوا، تفسیر نتایج و رعایت اخلاق بر عهده پژوهشگر و نویسنده است و هوش مصنوعی هرگز نمی تواند مسئولیت اخلاقی یا خلاقیت انسانی را بر عهده گیرد.

۱۲. جمع بندی و نتیجه گیری

هوش مصنوعی به ویژه در قالب مدل های زبانی بزرگ، به یک همکار قدرتمند در فرآیند پژوهش و نگارش علمی تبدیل شده است. مزایای آن از جمله افزایش کارایی، سرعت بخشیدن به مراحل مقدماتی پژوهش، تحلیل داده های کلان، کاهش خطاهای تکراری و ارائه کمک های شخصی سازی شده در نگارش، غیرقابل انکار است. این فناوری می تواند با آزادسازی زمان پژوهشگران از وظایف مکانیکی، فضای بیشتری برای تفکر

[5]. OECD. (2019). Recommendation of the Council on Artificial Intelligence (OECD Legal Instruments). Organisation for Economic Co-operation and Development
Retrieved from: <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>
OR:
OECD. (2019). OECD Principles on Artificial Intelligence. Organisation for Economic Co-operation and Development.

Retrieved from: (<https://oecd.ai/en/ai-principles>)

[6]. European Parliament and Council. (2016). Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation). *Official Journal of the European Union*, L 119, 04/05/2016, P. 1–88. Retrieved from: (<https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:32016R0679>)

[7]. Smolčić, V. S. (2013). Salami publication: definitions and examples. *Biochemia Medica*, 23(3), 237–241. (doi:10.11613/BM.2013.030)

[8]. Committee on Publication Ethics (COPE) (2025). Retrieved from (<https://publicationethics.org>)

[۹]. وزارت علوم، تحقیقات و فناوری جمهوری اسلامی ایران (۲۵ آبان ماه ۱۴۰۴). *راهنمای استفاده از ابزارهای هوش مصنوعی در پژوهش*، تهیه شده توسط مؤسسه تحقیقات سیاست علمی کشور

[۱۰]. بنیاد علم ایران (۱۴۰۴). *دستورالعمل استفاده از ابزارهای هوش مصنوعی در پژوهش* (www.insf.org)

پیشبرد علم بهره برد، درحالی که اصالت، اعتماد و یکتایی پژوهش علمی که سنگ بنای پیشرفت دانش است، محفوظ و حتی تقویت می‌شود. همکاری مستمر نهادهای قانون‌گذار، ناشران، مؤسسات علمی و خود پژوهشگران برای به‌روزرسانی چارچوب‌های اخلاقی و حرفه‌ای، ضرورتی اجتناب‌ناپذیر در این مسیر پویا و در حال تکامل است.

Reference

منابع

[1]. European Parliament and Council (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union legislative acts. *Official Journal of the European Union*, L 2024/1689. Retrieved from EUR-Lex website: (<https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>)

[2]. Executive Order 14110 of October 30, 2023. Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence. Federal Register: *The Daily Journal of the United States Government*, Vol. 88, No. 211. (November 1, 2023).

Retrieved from: (<https://www.govinfo.gov/content/pkg/FR-2023-11-01/pdf/2023-24283.pdf>)

[3]. UK Department for Science, Innovation and Technology (DSIT). (2023, November 1). The Bletchley Declaration by Countries Attending the AI Safety Summit, 1-2 November 2023.

Retrieved from: (<https://www.gov.uk/government/publications/ai-safety-summit-2023-the-bletchley-declaration>)

[۴]. منتظری، حسام، غفاری، احمدرضا؛ موسوی موحدی، علی اکبر (۱۴۰۲). کاربردهای چت جی پی تی در تحقیقات علمی و ملاحظات اخلاقی استفاده از آن، نشریه نشا علم، ۱۳، ۱، ۱۴۸–۱۵۶.