

طراحی و اعتبارسنجی مدل پابندی به اخلاق نشر در پژوهش‌های مبتنی بر کاربرست هوش مصنوعی:

مطالعه‌ای در دانشگاه‌های علوم پزشکی ایران

سارا دخش

پژوهشگر پسادکتری گروه کتابداری و اطلاع‌رسانی پزشکی،

دانشکده علوم پیراپزشکی، دانشگاه علوم پزشکی تهران، تهران،

ایران.

شهناز خادمی‌زاده*

دکتری علم اطلاعات و دانش‌شناسی، استاد،

گروه علم اطلاعات و دانش‌شناسی، دانشگاه شهید چمران اهواز، اهواز،

ایران.

زینب محمدی

دکتری علم اطلاعات و دانش‌شناسی، استادیار،

گروه علم اطلاعات و دانش‌شناسی، دانشگاه شهید چمران اهواز، اهواز،

ایران

دریافت: ۱۴۰۵/۰۳/۱۹

پذیرش: ۱۴۰۵/۰۶/۱۸

مقاله برای اصلاح به مدت ۱۸ روز نزد پدیدآوران بوده است.

چکیده: امروزه هر چند پیشرفت‌های جدید مانند ظهور ابزارهای هوش مصنوعی منجر به استفاده از فناوری‌های بروز در مطالعات گوناگون شده است؛ اما علم و پژوهش نمی‌تواند فراتر از استانداردها و ارزش‌های تعیین شده قدم بردارد؛ چرا که همواره با اخلاق سروکار دارند. بدین منظور پژوهش حاضر با هدف ارائه مدل پابندی به اخلاق نشر در پژوهش‌های وابسته به هوش مصنوعی در دانشگاه‌های علوم پزشکی کشور انجام شده است.

مطالعه حاضر یک پژوهش آمیخته، با رویکرد و هدف اکتشافی و از نظر بعد نتایج پژوهش کاربردی است. در مرحله کیفی مصاحبه‌های نیمه‌ساختاریافته با ۲۰ کنشگر علمی نظام پژوهشی در دو سطح کلان (مدیران، سیاست‌گذاران و معاونان پژوهشی وزارت بهداشت و دانشگاه‌ها) و خرد (سرپرستان، داوران نشریات، اعضای هیئت تحریریه و پژوهشگران حوزه اخلاق پژوهش و هوش مصنوعی) انجام شد. لازم به ذکر است که نمونه‌گیری به صورت هدفمند و با راهبردهای موارد مطلوب و زنجیره‌ای صورت گرفت. داده‌ها با استفاده از تحلیل مضمونی در نرم‌افزار مکس

نشریه علمی (رتبه بین المللی)
پژوهشگاه علوم و فناوری اطلاعات ایران
شاپا (چاپی) ۸۲۲۳-۲۲۵۱
شاپا (الکترونیکی) ۸۲۳۱-۲۲۵۱
نمایه در SCOPUS و ISC
<http://jipm.irandoc.ac.ir>
دوره XX | شماره X | صص XX-XX
۱۳XX X

نوع مقاله: پژوهشی

به این مقاله به شکل زیر استناد کنید:

(دخش، خادمی‌زاده و محمدی، زودآیند)

در فهرست منابع:

دخش، سارا، خادمی‌زاده، شهناز و زینب محمدی.

زودآیند. طراحی و اعتبارسنجی مدل پابندی به

اخلاق نشر در پژوهش‌های مبتنی بر کاربرست هوش

مصنوعی: مطالعه‌ای در دانشگاه‌های علوم پزشکی

ایران. پژوهشنامه پردازش و مدیریت اطلاعات.

(دسترسی در <http://jipm.irandoc.ac.ir>)

روز/ماه/سال)

کیو. دی. ای. نسخه ۲۰ تحلیل شدند که منجر به استخراج ابعاد و مؤلفه‌های مدل پابیندی در چهار بعد فردی، سازمانی، محیطی و شایستگی‌های حرفه‌ای گردید. در مرحله کمی، مدل پیشنهادی در قالب پرسشنامه محقق ساخته با ۳۴ گویه طراحی و توسط ۱۲۴ نفر از خبرگان حوزه اخلاق پژوهش و هوش مصنوعی تکمیل شد. داده‌ها با رویکرد مدل‌سازی معادلات ساختاری به روش حداقل مربعات جزئی (PLS-SEM) و با استفاده از نرم‌افزار Smart PLS نسخه ۴ تحلیل و برازش مدل تأیید گردید.

نتایج پژوهش نشان داد که ابعاد اثرگذار در کاربست اخلاقی ابزارهای هوش مصنوعی در نشر نتایج پژوهش در ۴ بعد اصلی فردی (با ۵۲ کد تفسیری)، سازمانی (با ۱۴۵ کد تفسیری)، محیطی (با ۵۱ کد تفسیری) و شایستگی‌های حرفه‌ای (با ۷۶ کد تفسیری) قرار می‌گیرند. همچنین مشخص شد که سازه پابیندی به اخلاق نشر در پژوهش‌های علوم پزشکی وابسته به ابزارهای هوش مصنوعی، اثر مثبت و معناداری بر تمامی ابعاد فردی، سازمانی، محیطی و شایستگی‌های حرفه‌ای دارد. در این میان، بیشترین ضریب مسیر مربوط به بعد سازمانی با مقدار ۰/۹۵۵ و آماره t برابر با ۱۱۰/۲۳۱ است که نشان می‌دهد عوامل سازمانی، قوی‌ترین مؤلفه در تبیین پابیندی به اخلاق نشر در پژوهش‌های علوم پزشکی وابسته به ابزارهای هوش مصنوعی محسوب می‌شوند. پس از آن، ابعاد فردی (۰/۹۵۳)، محیطی (۰/۸۸۲) و شایستگی‌های حرفه‌ای (۰/۸۸۱) در رتبه‌های بعدی قرار دارند. با توجه به اینکه در تمامی روابط، مقادیر آماره t به مراتب بزرگ‌تر از ۱/۹۶ و مقادیر P برابر با صفر گزارش شده‌اند، می‌توان نتیجه گرفت که کلیه مسیرهای مدل در سطح اطمینان ۹۵ درصد معنادار هستند. این یافته‌ها نشان می‌دهد که مدل ساختاری پژوهش از قدرت تبیین بسیار بالایی برخوردار است.

با این تفاسیر، پابیندی به موازین اخلاق نشر در پژوهش‌های وابسته به ابزارهای هوش مصنوعی در دانشگاه‌های علوم پزشکی، پدیده‌ای چندوجهی و مجموعه‌ای از عوامل درهم‌تنیده است. لذا تنها در سایه چنین نگاه چندسطحی و هماهنگی است که می‌توان از ظرفیت‌های بی‌نظیر هوش مصنوعی در پیشبرد پژوهش‌های علوم پزشکی بهره برد و هم‌زمان، اصالت، صحت و امانت‌داری علمی را حفظ نمود. یافته‌های این پژوهش می‌تواند مبنای سیاست‌گذاری‌های اخلاقی، تدوین راهنماهای بالادستی و طراحی برنامه‌های توانمندسازی پژوهشگران در دانشگاه‌های علوم پزشکی قرار گیرد.

کلیدواژه‌ها: اخلاق پژوهش، اخلاق نشر، سوء رفتار علمی، مدل‌های زبانی بزرگ، هوش مصنوعی، پژوهشگران، علوم پزشکی

* پدیدآور رابط: شهناز خادمی‌زاده. s.khademi@scu.ac.ir

۱. مقدمه

پژوهش در رشد اقتصادی، اجتماعی، سیاسی و فرهنگی و ارتقای جایگاه بین‌المللی کشورها نقش به‌سزایی دارد و بی‌شک از مهم‌ترین مزیت‌های رقابتی جوامع در جهان پرشتاب کنونی به حساب می‌آید. اما توسعه این ظرفیت‌های پژوهشی به ویژه در عصر فناوری و رواج ابزارهای هوش مصنوعی، نیازمند تربیت پژوهشگران باانگیزه، تامین منابع مالی و از همه مهم‌تر مستلزم رعایت و پابندی به اصول و ارزش‌هایی است که زمینه پیشبرد پژوهش‌ها را فراهم آورد (دخش ۱۴۰۲). اخلاق نشر از موضوعات مورد توجه در مباحث اخلاق پژوهش است و با توجه به گسترش دانش بشری و رشد فناوری به یک نگرانی جهانی تبدیل شده است (Son, & Yun 2021). اخلاق نشر به عنوان یک رفتار حرفه‌ای تعریف می‌شود و بیانگر اصول و قواعد مرتبط با شفافیت و یکپارچگی در نشر نتایج پژوهش است و با مالکیت معنوی آثار ارتباط مستقیم دارد (Schroter et al. 2018). استانداردها، ضوابط، مقررات، راهنماها، بایدها و نبایدها و دستورالعمل‌های نشر نتایج علمی تحت عنوان موازین اخلاق نشر قابل تعریف است که دربرگیرنده مسئولیت‌های اخلاقی مهره‌های پژوهش مانند شرایط نویسندگی مقالات و آثار علمی، اصول اخلاقی در سردبیری و داوری نشریات علمی و مبارزه با تخلفات و سوءرفتارهای پژوهشی در گزارش و انتشار آثار علمی توسط مدیران و سیاست‌گذاران است (وزارت بهداشت، درمان و آموزش پزشکی، ۱۳۹۶). در این پژوهش موازین اخلاق نشر، به مسئولیت‌های اخلاق نشر کنشگران نظام پژوهشی دانشگاه‌های علوم پزشکی اشاره دارد که شامل نویسندگان، سردبیران، ناشران و داوران نشریات، همچنین سیاست‌گذاران یعنی مدیران و تصمیم‌گیران دانشگاه‌ها، سازمان‌ها و مراکز پژوهشی وابسته به وزارت بهداشت، درمان و آموزش پزشکی هستند و فعالیت‌های آن‌ها در چارچوب دستورالعمل‌های «کمیته بین‌المللی اخلاق نشر» و موازین ملی قرار می‌گیرد.

هوش مصنوعی از دهه ۱۹۵۰ وجود داشته است اما در چند سال اخیر بر جنبه‌های مختلف زندگی بشر به ویژه در علوم سلامت، بهداشت و پزشکی تأثیرات وسیعی گذاشته است؛ مانند مراقبت از بیمار، آموزش ارائه‌کنندگان خدمات بهداشت و درمان، ارائه خدمات پزشکی و تسهیل در انتشارات علمی این حوزه موضوعی (Peh & Saw 2023; Akinrinmade et al. 2023). ورود این فناوری در نظام پژوهشی، به ویژه با ظهور گپ‌افزارهایی مانند چت جی. پی. تی. موجب کاهش فرآیندهای زمان‌بر تولید و نشر آثار علمی شده و در مقابل با صرفه‌جویی در زمان، فرصت بیشتری را جهت تمرکز بر مسائل عمیق‌تر و اساسی‌تر فرآیند پژوهش ایجاد نموده تا زمینه پیشبرد نوآوری‌ها و اکتشافات علمی فراهم شود (بهمن‌آبادی ۱۴۰۲؛ خادمی‌زاده ۱۴۰۳).

«کمیته بین‌المللی اخلاق نشر»، پردازش زبان طبیعی و یادگیری ماشین را به عنوان زیرمجموعه‌های هوش مصنوعی می‌بیند و بیان دارد که هدف استفاده از هوش مصنوعی، به کارگیری یک مدل یادگیری عمیق است تا محصولی بهتر از آنچه ذهن انسان توسعه داده است را ایجاد کند (COPE 2021). بنابراین با توجه به اینکه امروزه کنشگران علمی، زندگی آموزشی، پژوهشی و حتی شخصی خود را در یک محیط وابسته به فناوری‌ها و ابزارهای هوش مصنوعی شکل می‌دهند؛ ضروری است که جایگاه و نقش این پلتفرم‌ها در نظام پژوهشی بررسی و مدیریت شود.

امروزه ابزارهای هوش مصنوعی می‌توانند به بهترین شکل برای رشد آثار علمی به کار روند اما مهم است که محصول نهایی پژوهش، منعکس‌کننده فرآیندهای فکری نویسندگان باشد. زیرا این پلتفرم‌ها در کنار فرصت‌های بزرگی که برای تسریع فرآیند تولیدات علمی فراهم می‌کنند؛ می‌توانند چالش‌هایی را برای یکپارچگی پژوهش ایجاد کرده و نگرانی‌هایی را در ارتباط صحت، اعتبار و اصالت نتایج به همراه داشته باشند (Pearson 2024; Kim 2024). بنابراین کنشگران علمی از نویسندگان تا سردبیران و داوران مجلات در نظام پژوهشی سلامت، جهت حفظ کیفیت شواهد پژوهش می‌بایست با آگاهی از ساختار و قوانین مربوط به بهره‌گیری از هوش مصنوعی در بستر پژوهش، مانع از بروز پیامدهای منفی ناخواسته بکارگیری این فرآیندهای خودکار شوند. یکی از ملموس‌ترین این پیامدها و چالش‌های اخلاقی، سلب اعتبار آثار علمی به دلیل وجود تخلفات پژوهشی و یا استفاده ناصحیح و غیر اخلاقی از فناوری هوش مصنوعی در انتشارات علمی است.

این مساله به ویژه با ظهور فناوری‌های دردسترس هوش مصنوعی، می‌تواند در جریان پژوهش زنگ خطری برای نظام سلامت باشد و پیامدهای سنگینی برای نظام بهداشت و درمان و آموزش پزشکی کشور به همراه داشته باشد. چرا که نتایج و شواهد این گروه از پژوهش‌ها با زندگی و سلامت جامعه ارتباط تنگاتنگ دارد و از آنجا که هدف اصلی نظام سلامت، بهبود سطح سلامت مردم و پیشبرد عدالت بهداشتی در میان آن‌ها است؛ توجه به موازین اخلاقی نشر نتایج پژوهش‌ها و حفظ کیفیت مطالعات در عصر هوش مصنوعی، بی‌شک نقش به‌سزایی در تحقق اهداف نظام سلامت خواهد داشت.

اگرچه مطالعات و گزارش‌های علمی متعدد در سال‌های اخیر تلاش کرده‌اند تا ابعاد مختلف کاربرست هوش مصنوعی مولد را در فرآیندهای پژوهشی بررسی کرده و چالش‌های اخلاقی منفرد نظیر مسائل نویسندگی، سرقت علمی، ابهامات مالکیت فکری را شناسایی کنند (Bjelobaba et al.

بر توصیف چالش‌ها یا ارائه دستورالعمل‌های عمومی متمرکز مانده است. لذا با وجود درک نسبی از پیامدهای منفی این فناوری، یک شکاف دانشی عمده در زمینه ارائه الگویی یکپارچه و چندبعدی وجود دارد که بتواند عوامل مختلف را به طور هم‌زمان در مواجهه با ابزارهای هوش مصنوعی تبیین کند. خلاء ساختاری مذکور مانع از آن شده است که سیاست‌گذاران و نهادهای متولی بتوانند راهبردهای عملیاتی و منسجمی برای هدایت پژوهشگران به سمت استفاده مسئولانه از این ابزارها تدوین کنند.

این شکاف پژوهشی به‌ویژه در بافت دانشگاه‌های علوم پزشکی ایران که یافته‌های پژوهشی آن‌ها مستقیماً با فرآیندهای تصمیم‌گیری بالینی، ایمنی بیماران و سلامت جامعه در ارتباط است، حساسیتی مضاعف دارد؛ زیرا هرگونه خدشه به تمامیت علمی در این حوزه می‌تواند پیامدهای جبران‌ناپذیری بر نظام سلامت کشور داشته باشد. بر این اساس، ضرورت دارد تا با جلب مشارکت کنشگران کلیدی شامل متخصصان اخلاق پژوهش، سیاست‌گذاران نظام سلامت و پژوهشگران، چارچوبی منسجم جهت ترویج اخلاق‌مداری در نشر پژوهش‌های متأثر از هوش مصنوعی تدوین شود. از این‌رو، پژوهش حاضر با هدف طراحی و اعتبارسنجی مدل پابندی به اخلاق نشر در پژوهش‌های مبتنی بر کاربرست هوش مصنوعی در دانشگاه‌های علوم پزشکی ایران انجام شده است تا با تبیین ابعاد گوناگون این پدیده، راهنما و تسهیل‌گر مسیر کنشگران علمی در جهت کاهش فاصله میان وضع موجود و مطلوب باشد. در ادامه به سوال‌های پژوهشی همراستا با هدف مطالعه حاضر اشاره شده است:

(۱) ابعاد و مولفه‌های پابندی به اخلاق نشر در پژوهش‌های مبتنی بر کاربرست هوش مصنوعی در دانشگاه‌های علوم پزشکی ایران کدام‌اند؟

(۲) برازش مدل پابندی به اخلاق نشر در پژوهش‌های مبتنی بر کاربرست هوش مصنوعی در دانشگاه‌های علوم پزشکی ایران چگونه است؟

۲. پیشینه پژوهش

آموزش اخلاق پژوهش در طول دوران تحصیل، امری ضروری برای توسعه استانداردهای اخلاقی است. در پژوهشی از نیک‌روان فرد، خراسانی‌زاده و زنده‌دل (۲۰۱۷) وضعیت آموزش اخلاق

پژوهش در برنامه‌های درسی تحصیلات تکمیلی دانشجویان علوم پزشکی ایران مورد ارزیابی قرار گرفته است. نتایج پژوهش آنان نشان داد که در ۵۸ درصد برنامه‌های درسی، واحدی برای آموزش اخلاق موجود نبود و تنها ۱۷ برنامه دارای دوره اختصاصی برای اخلاق پژوهش بودند. در مطالعه‌ای از (Dorton et al. ۲۰۲۳) آینده‌نگاری برای هوش مصنوعی اخلاقی بررسی شده است. نتایج مطالعه نشان داد که با توجه به شیوع رفتارهای اضطرابی که هنگام ادغام هوش مصنوعی در سیستم‌های پیچیده اجتماعی و فنی اتفاق می‌افتد؛ عملی کردن این چارچوب‌ها و دستورالعمل‌های اخلاقی می‌تواند دشوار باشد. از نتایج این عدم عمل‌گرایی اخلاق در کاربست هوش مصنوعی، پدیده‌ای به نام دین اخلاقی است. استدلال می‌شود که اصول، مدل‌ها و ابزارهای تصمیم‌گیری طبیعت‌گرایانه برای مقابله با این چالش‌ها بسیار مناسب هستند. همچنین نتایج مطالعه (۲۰۲۴) Kasani et al نشان داد که به کارگیری ابزارهای هوش مصنوعی در پژوهش، هم پتانسیل‌های جذاب و هم چالش‌های اخلاقی را به ویژه در زمینه نویسندگی و نشر نتایج علمی به همراه دارد. در پژوهشی دیگر از Kim (۲۰۲۴) اشاره شده که چالش‌های اخلاقی متعددی در اثر استفاده از این ابزارها پیش روی کنشگران علمی است که مهم‌ترین آن‌ها عبارتند از: شباهت زیاد در خروجی‌های هوش مصنوعی، نقض قوانین کپی‌رایت، به خطر افتادن کیفیت آثار به دلیل استنادات ناکافی و یا عدم ذکر منابع، نگرانی در ارتباط با یکپارچگی و صحت پژوهش، نیاز به دستورالعمل‌های استاندارد نویسندگی برای هوش مصنوعی، انتساب نامناسب نویسندگی و اختلال بالقوه پژوهش‌های آینده به دلیل سلب اعتبارهای متعدد آثار. Shaw et al. (۲۰۲۴) نیز در گزارشی که از نشست جهانی اخلاق زیستی در پژوهش (برگزار شده از سوی سازمان جهانی بهداشت) مربوط به سال ۲۰۲۲ ارائه کرده‌اند؛ به چالش‌ها و استراتژی‌های ترویج اخلاق در پژوهش‌های مربوط به هوش مصنوعی در سلامت جهانی اشاره داشته‌اند. نتایج این نشست نشان داد که ۱. مناسب بودن سازوکار ابزارهای هوش مصنوعی، ۲. قابلیت انتقال نظام‌های هوش مصنوعی، ۳. مسئولیت‌پذیری تصمیم‌ها و نتایج هوش مصنوعی و ۴. رضایت فردی، از مواردی است که نیاز است در زمان استفاده از خدمات و ابزارهای هوش مصنوعی در سیاست، پژوهش و عمل مورد توجه قرار گیرد. Bjelobaba et al. (۲۰۲۵) در پژوهشی کیفی با تحلیل اسنادی، چالش‌های اخلاقی مرتبط با کاربرد هوش مصنوعی مولد در پژوهش را بررسی کردند. نتایج پژوهش پنج دسته چالش اصلی شامل شفافیت در استفاده از هوش مصنوعی، مسئولیت‌پذیری نویسندگان، توهّمات ارجاعی و اطلاعات نادرست، سوگیری الگوریتمی و مسائل حریم خصوصی داده‌ها را شناسایی کرد. همچنین پژوهشگران بر ضرورت افشای استفاده از هوش مصنوعی، نظارت انسانی بر خروجی‌ها و تدوین سیاست‌ها و آموزش‌های اخلاقی برای استفاده مسئولانه از این فناوری تأکید کردند.

همچنین در پژوهشی از Lin et al (۲۰۲۶) با عنوان «پیمایش یکپارچگی علمی در پژوهش‌های زیست‌پزشکی: تأثیر مدل‌های زبانی بزرگ بر رویه‌های جاری و جهت‌گیری‌های آینده»، سه حوزه اصلی تأثیرگذاری مدل‌های زبانی بزرگ شامل نگارش و ویرایش مقالات، تحلیل و تفسیر داده‌ها و تولید پیشینه و فرضیه‌پردازی مورد بررسی قرار گرفته و چالش‌های اخلاقی متناظر با هر حوزه شامل مسئله نویسندگی و انتساب نقش به هوش مصنوعی، خطر استنادات جعلی، سوگیری‌های الگوریتمی در تفسیر یافته‌های بالینی و تهدید حریم خصوصی داده‌های بیمارستانی و پژوهشی احصا شده است. نویسندگان در پایان چهار جهت‌گیری کلیدی برای آینده شامل: (۱) توسعه ابزارهای تشخیص محتوای تولیدشده توسط هوش مصنوعی با رویکرد تأیید انسانی، (۲) بازتعریف مفهوم نویسندگی و تدوین معیارهای شفاف برای نقش هوش مصنوعی در مقالات، (۳) طراحی دوره‌های آموزشی تخصصی اخلاق هوش مصنوعی برای پژوهشگران زیست‌پزشکی و (۴) ایجاد مکانیسم‌های پاسخگویی و نظارت مستمر بر فرآیندهای پژوهشی مبتنی بر هوش مصنوعی را پیشنهاد می‌دهند. در مقاله‌ای تحلیلی نیز (Delanghe 2026) نشان داد که اگرچه هوش مصنوعی می‌تواند فرایندهایی مانند تحلیل داده، نگارش علمی، تولید فرضیه و داوری هم‌تا را تسهیل کند، اما هم‌زمان چالش‌های اخلاقی مهمی را در حوزه نشر علمی ایجاد می‌کند. وی با اتکا به دستورالعمل‌های ناشران و شواهد موجود تأکید کرد که مسئولیت نهایی صحت، اصالت و تمامیت محتوای علمی همچنان بر عهده پژوهشگران انسانی است و استفاده از هوش مصنوعی باید در موارد فراتر از ویرایش زبانی و قالب‌بندی، به صورت شفاف افشا شود. در این مقاله همچنین مسائلی مانند مالکیت فکری، نابرابری در دسترسی به فناوری، سوءاستفاده در تولید مقالات و محدودیت‌های هوش مصنوعی در داوری هم‌تا مورد تأکید قرار گرفته و حفظ محرمانگی داده‌ها به عنوان مهم‌ترین دلیل برای محدودسازی استفاده از این ابزارها در فرایندهای نشر علمی برجسته شده است.

در مطالعه‌ای مروری از بهمن‌آبادی (۱۴۰۲) تصویری کلی از کاربردهای هوش مصنوعی در پژوهش و برخی مسائل مرتبط با آن ارائه شده است. نتایج این بررسی نشان داد که از این فناوری برای فعالیت‌های رایجی در پژوهش مانند تهیه پیشینه پژوهشی، خلاصه‌سازی متون، ترجمه و ویرایش متن به صورت گسترده‌تری استفاده می‌شود. با این حال استفاده از هوش مصنوعی در پژوهش، چالش‌ها و خطرات خاص خود را دارد که سرقت علمی، خطا یا سوگیری در متن، دستکاری یا تحریف سوابق علمی، مالکیت معنوی اثر تولیدی و خطرات شغلی از جمله این موارد هستند. همچنین در مطالعه مرتضوی و زارعی (۱۴۰۳) نیز با عنوان «چالش‌های اخلاقی استفاده از

هوش مصنوعی در پژوهش‌های علمی» ضمن تأکید بر اصول پنج‌گانه اخلاق پژوهش شامل صداقت و امانت‌داری، بی‌طرفی و حقیقت‌جویی، قانون‌مداری و خودسامانی، مسئولیت‌پذیری و شفافیت و پرهیز از پیش‌داوری در نتایج، اشاره داشتند که مدل‌های هوش مصنوعی به دلیل ماهیت مبتنی بر داده‌های آموزشی، هرگز نمی‌توانند کاملاً بی‌طرف باشند و پژوهشگران باید علاوه بر آگاهی از این سوگیری‌ها، استفاده از هوش مصنوعی را به صورت شفاف افشا کنند و هرگز مسئولیت نهایی نتایج پژوهشی را به هوش مصنوعی واگذار نکنند. در پژوهشی دیگر از دختش و همکاران (۱۴۰۳) عوامل اثرگذار بر پایبندی به موازین اخلاق نشر در نظام پژوهشی وزارت عتف به روش تحلیل لایه‌ای علی بررسی شده است. در جریان مصاحبه با کنشگران علمی به مقوله «رصد نظام پژوهشی» از لایه لیتانی/عینی، مقوله «شفاف‌سازی قوانین و مقررات به لحاظ اجرایی» از لایه علل اجتماعی، مقوله «جهت‌دهی به گفتمان‌های اجتماعی حاکم با فرهنگ‌سازی» از لایه جهان‌بینی و مقوله «همبستگی آموزش و پژوهش» از لایه اسطوره و استعاره بیشترین تأکید شده بود. به علاوه شریف‌زاده (۱۴۰۴) در مطالعه خود با عنوان «تحلیل و بررسی مزایا و چالش‌های اخلاقی مدل‌های زبانی بزرگ در سپهر پژوهش»، در کنار ارائه مزایایی چون رفع موانع زبانی برای نویسندگان غیر انگلیسی‌زبان، کاهش هزینه‌ها، صرفه‌جویی در زمان و افزایش کیفیت نگارش، به دسته‌بندی جامعی از چالش‌های اخلاقی نیز پرداخته است. یافته‌ها نشان می‌دهد که چالش‌های این فناوری طیف گسترده‌ای از مسائل بنیادین همچون «توهم»، «ابهام در مسئولیت‌پذیری و مؤلف بودن»، «مشابه‌نویسی و سرقت ادبی»، «سوگیری»، «امنیت و حریم خصوصی»، «شفافیت و توضیح‌پذیری»، تا پیامدهای کلان‌تری چون «مسائل زیست‌محیطی» و «بیکاری و بهره‌کشی» را دربر می‌گیرد.

مرور پیشینه‌ها نشان می‌دهد که با گسترش سریع کاربرد هوش مصنوعی در فرایندهای پژوهش و نشر علمی، چالش‌های اخلاقی متعددی نیز پدیدار شده است؛ از جمله ابهام در نویسندگی و مسئولیت، سوگیری الگوریتمی، نقض محرمانگی داده‌ها و مسائل مربوط به مالکیت فکری. در عین حال، بیشتر مطالعات بر ضرورت شفافیت در اعلام استفاده از هوش مصنوعی، بازنگری در سیاست‌ها و حفظ نظارت انسانی تأکید کرده‌اند. با این حال، ادبیات موجود عمدتاً پراکنده و توصیفی است و کمتر به ارائه چارچوبی منسجم و بومی برای هدایت عمل پرداخته است. از این‌رو، طراحی یک مدل جامع و چندبعدی برای پایبندی به اخلاق نشر در پژوهش‌های مبتنی بر هوش مصنوعی، به‌ویژه در بستر دانشگاه‌های علوم پزشکی ایران، یک ضرورت پژوهشی محسوب می‌شود.

۳. روش پژوهش

مطالعه حاضر یک پژوهش آمیخته، با رویکرد و هدف اکتشافی و از نظر بعد نتایج پژوهش کاربردی است. از منظر زمان‌بندی، مراحل پژوهش به صورت پی‌درپی انجام شد؛ به گونه‌ای که ابتدا داده‌های کیفی جمع‌آوری و تحلیل گردید و سپس فاز کمی اجرا شد. از نظر وزن و اهمیت، اولویت اصلی به فاز کیفی اختصاص داشت تا ابعاد مدل بومی استخراج شود. ادغام داده‌ها در مرحله انتقال از فاز کیفی به فاز کمی و از طریق اتصال صورت گرفت؛ بدین معنا که کدهای حاصل از تحلیل کیفی، مبنای طراحی گویه‌های پرسشنامه کمی برای اعتبارسنجی مدل قرار گرفتند.

در گام اول یعنی بخش کیفی مطالعه، جهت ارائه مدل پایداری به اخلاق نشر در بستر استفاده از هوش مصنوعی در فرآیند پژوهش، از مصاحبه نیمه‌ساختاریافته استفاده شد. جامعه پژوهش در این مرحله کنشگران علمی نظام پژوهشی وزارت بهداشت در دو سطح کلان و خرد بودند. در سطح کلان، مدیران، سیاست‌گذاران و معاونان پژوهشی وزارت بهداشت و دانشگاه‌ها مشارکت داشتند و در سطح خرد، سردبیران، داوران نشریات، اعضای هیئت تحریریه و پژوهشگران حوزه اخلاق پژوهش و هوش مصنوعی همکاری نمودند. نمونه‌گیری به صورت هدفمند و با راهبردهای موارد مطلوب و زنجیره‌ای انجام شد و در نهایت ۲۰ مشارکت‌کننده تا رسیدن به اشباع نظری وارد مطالعه شدند (جدول ۱). لازم به ذکر است که پس از اتمام مصاحبه‌ها با مشارکت‌کنندگان و حصول اشباع نظری، به منظور افزایش جامعیت مدل، از ChatGPT نسخه ۵ صرفاً به عنوان ابزار کمکی در تحلیل تکمیلی و تقابلی استفاده شد. بدین صورت که مضامین و مقوله‌های استخراج‌شده از داده‌های انسانی، در یک مرحله بازبینی ثانویه، برای شناسایی احتمالی خلأهای مفهومی و بررسی پوشش ابعاد نوپدید در حوزه اخلاق نشر در بستر هوش مصنوعی، در معرض ارزیابی انتقادی این ابزار قرار گرفت.

جدول ۱. کنشگران مشارکت‌کننده در پژوهش (مصاحبه‌شوندگان)

سطوح کنشگران	سمت و مسئولیت	معیار ورود	فراوانی
--------------	---------------	------------	---------

۸	- مدیران و سیاست‌گذاران واحدهای معاونت پژوهشی و فناوری وابسته به وزارت بهداشت، درمان و آموزش پزشکی؛ - معاونین و مدیران پژوهشی و فناوری دانشگاه‌ها و مراکز تحقیقاتی وابسته به وزارت بهداشت، درمان و آموزش پزشکی. - آگاهی و شناخت کافی نسبت به مباحث اخلاق پژوهش، اخلاق نشر، سوءرفتار علمی و هوش مصنوعی در پژوهش.	سطح کلان
۱۲	- حداقل سه سال سابقه اجرایی در مجلات و انتشارات دانشگاه‌ها و مراکز پژوهشی وابسته به وزارت بهداشت، درمان و آموزش پزشکی؛ - حداقل سه تالیف و یا گزارش علمی و خبری در حوزه اخلاق پژوهش، اخلاق نشر، سوءرفتار علمی و هوش مصنوعی در پژوهش؛ - آگاهی و شناخت کافی نسبت به مباحث اخلاق پژوهش، اخلاق نشر، سوءرفتار علمی و هوش مصنوعی در پژوهش.	سطح خرد
۲۰		مجموع

گردآوری داده‌های کیفی از طریق مصاحبه‌های نیمه‌ساختاریافته صورت گرفت. ابزار گردآوری داده‌ها در مرحله کیفی بسته مصاحبه، شامل فرم رضایت آگاهانه و راهنمای مصاحبه نیمه‌ساختاریافته بود که پیش از اجرا در اختیار مشارکت‌کنندگان قرار گرفت. فرم راهنمای مصاحبه، دارای بخش مقدماتی برای معرفی هدف پژوهش و سوال‌های مصاحبه (۱۳ سوال) بود. ذکر این نکته ضروری است که در جریان مصاحبه تعدادی از سوالات تجمیع و یا حتی در مواردی افزایش یافتند. لذا متناسب با سطح دانش و جنس نظرات مصاحبه‌شوندگان تغییراتی در ترتیب و تعداد سوالات انجام شد تا در نهایت مرتبط‌ترین پاسخ‌ها دریافت شود.

به طور کلی گردآوری داده‌های مصاحبه به دو شکل حضوری و مجازی، در بازه زمانی ۱۳ هفته (اوایل مهر ماه ۱۴۰۴ تا اواخر آذر ماه ۱۴۰۴)، انجام شد و مدت زمان اجرای آن‌ها بین ۲۰ تا ۴۷ دقیقه متغیر بود. بدین صورت که با توجه به شیوه پیشنهادی هر یک از مصاحبه‌شوندگان و نوع مشغله کاری آنان، فرآیند مصاحبه اجرا شد. اکثریت مصاحبه‌ها (۱۴ مصاحبه)، با کنشگران به صورت حضوری و در محل کار آنان ضبط شد و دیگر مصاحبه‌ها که امکان شرکت در جلسه حضوری نداشتند؛ به صورت آنلاین در بستر گوگل میت انجام شد. لازم به ذکر است که مصاحبه‌ها با رضایت مصاحبه‌شوندگان ضبط و در حین مکالمه از تکنیک یادداشت‌برداری استفاده شد. پس از ضبط و پیاده‌سازی کامل مصاحبه‌ها، داده‌ها در نرم‌افزار MAXQDA نسخه ۲۰ وارد گردید. تحلیل داده‌های کیفی به روش تحلیل محتوای کیفی و با رویکرد استقرایی انجام شد. در این رویکرد، کدها و مقوله‌ها بدون اعمال فرضیات پیش‌بینی‌شده و مستقیماً از متن مصاحبه‌ها استخراج شدند. فرآیند تحلیل در سه مرحله با استفاده از تحلیل مضمونی سه‌مرحله‌ای شامل کدگذاری توصیفی، تفسیری و استخراج مضامین جامع صورت گرفت.

به‌منظور اعتبارسنجی ابزار گردآوری داده‌ها در مرحله کیفی (راهنمای مصاحبه)، از چارچوب پالایش پروتکل مصاحبه استفاده شد که شامل ارزیابی انطباق پرسش‌ها با اهداف پژوهش، دریافت بازخورد خبرگان و اجرای آزمایشی مصاحبه بود (Castillo-Montoya 2016)؛ این فرآیند به افزایش اعتبارپذیری ابزار پژوهش کمک نمود. همچنین، جهت سنجش و تضمین ثبات‌پذیری فرایند تحلیل و کدگذاری، از روش پایایی بین کدگذاران (توسط دو کدگذار مستقل) استفاده گردید که میزان توافق بین آن‌ها ۸۴ درصد محاسبه شد و نشان‌دهنده هماهنگی و ثبات مناسب در استخراج کدها است. همچنین اختلاف‌نظرهای باقی‌مانده از طریق بحث و توافق نهایی میان کدگذاران برطرف شد و نسخه نهایی کدبوک تثبیت گردید.

در ادامه جهت افزایش اعتبارپذیری یافته‌ها، کدها و مضامین استخراج‌شده به همراه متن پیاده‌سازی‌شده مصاحبه برای دو نفر از مشارکت‌کنندگان ارسال شد تا میزان انطباق آن‌ها با دیدگاه‌های خود را بررسی کنند. پس از دریافت بازخورد، پیشنهادهای مرتبط در جلسه تیم پژوهشی بررسی و موارد مورد توافق در تحلیل نهایی اعمال شد.

در فاز کمی پژوهش حاضر، به‌منظور آزمون مدل استخراج‌شده از مرحله کیفی، از روش توصیفی-تحلیلی استفاده شد. جامعه آماری مورد مطالعه در این مرحله متخصصان و پژوهشگران حوزه اخلاق پژوهش و هوش مصنوعی بودند که شناسایی آنان از طریق پایگاه داده «سامانه خبره‌یاب وزارت بهداشت، درمان و آموزش پزشکی» صورت گرفت. همراستا با ماهیت تخصصی موضوع و محدودیت در تعداد خبرگان واجد شرایط، از روش «نمونه‌گیری در دسترس» استفاده شد. فایل الکترونیکی پرسشنامه از طریق رایانامه برای ۱۸۰ نفر از خبرگان ارسال شد. در نهایت، ۱۲۴ پرسشنامه کامل و قابل‌تحلیل دریافت گردید که نرخ پاسخ‌دهی ۶۸/۸۰ درصد را نشان می‌دهد. همچنین به دلیل استفاده از روش مدل‌سازی معادلات ساختاری حداقل مربعات جزئی، کفایت حجم نمونه صرفاً بر مبنای نسبت تعداد مشارکت‌کنندگان به کل گویه‌ها ارزیابی نشد؛ بلکه با لحاظ پیچیدگی مدل ساختاری، تعداد سازه‌ها و روابط پیش‌بینی‌کننده هر سازه درون‌زا بررسی گردید. در این رویکرد، حجم نمونه موردنیاز تابع تعداد روابط ورودی به سازه درون‌زای دارای بیشترین تعداد پیش‌بین و توان آماری مورد انتظار است. بر این اساس، تعداد ۱۲۴ پرسشنامه کامل برای برآورد و آزمون مدل پژوهش کفایت داشت.

ابزار سنجش در این مرحله، پرسشنامه‌ای محقق‌ساخته بود که بر پایه طرح آمیخته اکتشافی متوالی و بر اساس نتایج حاصل از تحلیل مصاحبه‌های فاز کیفی (کدگذاری استقرایی مصاحبه‌ها) طراحی گردید. این ابزار مشتمل بر ۳۴ گویه در قالب ۴ بعد اصلی شامل: «فردی» (۴ مقوله و ۱۰ زیرمقوله)، «سازمانی» (۴ مقوله و ۱۲ زیرمقوله)، «محیطی» (۲ مقوله و ۶ زیرمقوله) و «شایستگی‌های حرفه‌ای» (۲ مقوله و ۶ زیرمقوله) تنظیم شد. سنجش متغیرها بر اساس طیف پنج‌درجه‌ای لیکرت (از ۱=خیلی کم تا ۵=خیلی زیاد) صورت پذیرفت. انتخاب طیف ۵ گزینه‌ای به دلیل کاهش بار شناختی پاسخ‌دهندگان (خبرگان)، سهولت و سرعت در پاسخ‌دهی، و ایجاد تعادل بهینه میان حساسیت اندازه‌گیری و پایایی ابزار انجام شد. بعلاوه جهت تأیید روایی صوری و محتوایی، ابزار مذکور توسط ۸ نفر از متخصصان حوزه موضوعی بازبینی و پس از اعمال اصلاحات ویرایشی، شفافیت و صراحت گویه‌ها احراز گردید.

تحلیل داده‌ها در این مرحله به کمک نرم‌افزار SmartPLS نسخه ۴ انجام شد. با توجه به ماهیت پژوهش، تمرکز اصلی بر طراحی و ارزیابی ویژگی‌های روان‌سنجی مدل بود؛ از این‌رو تحلیل‌ها در دو سطح توصیفی و استنباطی و در چارچوب ارزیابی مدل اندازه‌گیری سلسله‌مراتبی در سطوح مرتبه اول و دوم صورت گرفت. برای سنجش پایایی و روایی همگرا، از شاخص‌های آلفای کرونباخ، پایایی ترکیبی (CR) و میانگین واریانس استخراج‌شده (AVE) استفاده شد. نتایج نشان داد مقادیر آلفای کرونباخ و CR برای تمامی مؤلفه‌ها بالاتر از حد آستانه ۰/۷۰ بوده و بیانگر همسانی درونی مطلوب ابزار است. همچنین، مقادیر AVE برای همه سازه‌ها بیش از ۰/۵۰ به دست آمد که نشان‌دهنده روایی همگرای مناسب و تبیین قابل قبول واریانس گویه‌ها توسط سازه‌های مکنون مربوطه است.

۴. یافته‌ها

۴-۱. ابعاد و مؤلفه‌های پایداری به اخلاق نشر در پژوهش‌های مبتنی بر هوش مصنوعی

در مرحله اول و با اتخاذ رویکردی استقرایی، تحلیل متن مصاحبه‌ها به استخراج ۶۷۲ کد اولیه منجر شد. در مرحله بعد و طی فرآیند کدگذاری تفسیری و محوری، جهت کاهش و تلخیص داده‌ها، کدهای اولیه هم‌پوشان و مشابه ادغام گردیدند که حاصل آن شکل‌گیری ۳۲۴ کد مفهومی متمایز بود. در گام بعدی، این کدهای تفسیری بر اساس شباهت‌های مفهومی و سنخیت معنایی سازماندهی مجدد شدند و ۳۴ مقوله فرعی حاصل بود. در نهایت، از طریق کدگذاری انتخابی و انتزاع مفاهیم در سطح کلان، ابعاد اصلی پژوهش که نشان‌دهنده عوامل اثرگذار بر کاربرست اخلاقی ابزارهای هوش مصنوعی در نشر نتایج پژوهش هستند، مستقیماً از طریق تحلیل مصاحبه‌ها در ۴ بعد اصلی فردی (با ۵۲ کد تفسیری)، سازمانی (با ۱۴۵ کد تفسیری)، محیطی (با ۵۱ کد تفسیری) و شایستگی‌های حرفه‌ای (با ۷۶ کد تفسیری) صورت‌بندی و نام‌گذاری شدند. در ادامه به معرفی این ابعاد و مقوله‌ها پرداخته می‌شود.

بعد فردی، بیانگر نقش تعیین‌کننده ویژگی‌های فردی پژوهشگر در حفظ قضاوت انسانی، تصمیم‌گیری اخلاقی و استفاده مسئولانه از ابزارهای هوش مصنوعی هستند. از میان مقوله‌های فرعی، «پذیرش مسئولیت متون تولید شده با هوش مصنوعی» با فراوانی ۱۰، «ارزیابی دقیق خروجی هوش مصنوعی پیش از پذیرش» و «رعایت شفافیت در کاربرست هوش مصنوعی» هر یک با فراوانی ۷، بیشترین تکرار را در میان مشارکت‌کنندگان داشتند که نشان‌دهنده تأکید بالای آنان بر مسئولیت‌پذیری و شفافیت فردی در استفاده اخلاقی از ابزارهای هوش مصنوعی است (جدول ۲).

جدول ۲. مولفه‌های فردی اثرگذار بر پایداری به اخلاق نشر در پژوهش‌های مبتنی بر هوش مصنوعی

ابعاد	مقوله‌های اصلی	مقوله‌های فرعی	فراوانی
۶۰	مسئولیت‌پذیری اخلاقی	پذیرش مسئولیت متون تولید شده با AI؛ ۱۰	
		ارزیابی دقیق خروجی AI قبل از پذیرش؛ ۷	
		حفظ نگرش انسانی در تحلیل و تفسیر نتایج؛ ۶	
	آگاهی و نگرش اخلاقی نسبت به ابزارهای هوش مصنوعی	آگاهی نسبت به کدها و موازین اخلاقی AI؛ ۶	
		توجه به خطوط قرمز اخلاقی و قوانین استاندارد؛ ۴	
	خودکنترلی ^۱ و رعایت استانداردهای اخلاقی	استفاده آگاهانه و هدفمند از ابزارهای هوش مصنوعی؛ ۵	
		مدیریت زمان و انتشار علمی با رعایت موازین اخلاق؛ ۳	
		پیشگیری از هرگونه سوءاستفاده احتمالی؛ ۲	
	تعهد شخصی به شفافیت و صداقت علمی	رعایت شفافیت در کاربری AI؛ ۷	
		رعایت حقوق مالکیت فکری؛ ۲	

همچنین نتایج تحلیل داده‌ها در بعد سازمانی (جدول ۳) نشان داد که مقوله‌های «برگزاری آموزش‌ها و کارگاه‌های ساختارمند کاربری اخلاقی هوش مصنوعی» با فراوانی ۲۰، «تدوین و به‌روزرسانی آیین‌نامه‌ها و دستورالعمل‌های اخلاقی هوش مصنوعی» با فراوانی ۱۷ و «نظارت مستمر بر اجرای پژوهش و فرآیند نشر» با فراوانی ۱۵ بیشترین تکرار را در داده‌ها داشتند که نشان‌دهنده تأکید مشارکت‌کنندگان بر نقش محوری آموزش نهادی، سیاست‌گذاری شفاف و نظارت سازمانی در استفاده اخلاقی از هوش مصنوعی در فرآیند نشر علمی است.

جدول ۳. مولفه‌های سازمانی اثرگذار بر پایداری به اخلاق نشر در پژوهش‌های مبتنی بر هوش مصنوعی

¹ Self-control

ابعاد	مقوله‌های اصلی	مقوله‌های فرعی	فراوانی
سازمانی	سیاست‌گذاری و راهبری نهادی	تدوین و به‌روزرسانی آیین‌نامه‌ها و دستورالعمل‌های اخلاقی AI؛	۱۷
		شفاف‌سازی الزامات استفاده از AI در پژوهش و نشر؛	۱۴
		هماهنگی نهادی بین دانشگاه‌ها و مجلات؛	۱۰
	زیرساخت‌های آموزشی و توانمندسازی سازمانی	برگزاری آموزش‌ها و کارگاه‌های ساختارمند کاربرست اخلاقی AI؛	۲۰
		توسعه زیرساخت‌های پژوهشی و دسترسی کنترل‌شده به ابزارهای AI؛	۱۲
		الزام به دریافت گواهی اخلاق از دوره‌های آموزشی پیش از اجرای پژوهش؛	۴
	نقش کمیته‌های اخلاق و نظام‌های نظارتی	نظارت مستمر بر اجرای پژوهش و فرآیند نشر؛	۱۵
		داوری اخلاقی پروپوزال‌ها با تمرکز بر نحوه کاربرست AI؛	۱۲
		نقش آموزشی و مشاوره‌ای کمیته‌های اخلاق؛	۵
	فرهنگ سازمانی	ترویج فرهنگ شفافیت، گزارش‌دهی و مسئولیت‌پذیری در پژوهش؛	۱۴
		طراحی سازوکارهای الزام‌آور اخلاق‌محور در نشر علمی؛	۱۳
		ایجاد نظام‌های تشویقی و الگوسازی نهادی؛	۹

بعد محیطی شامل عوامل کلان فناورانه، ساختاری و نهادی است که می‌توانند محیط پژوهش را در کاربرست اخلاقی ابزارهای هوش مصنوعی شکل دهند. مطابق نتایج جدول ۴، بررسی مقوله‌های فرعی با بیشترین فراوانی نشان داد که «آموزش داوران و سردبیران برای تشخیص و

آشنایی با کدهای اخلاقی «AI (فراوانی ۱۲)، «مدیریت هم‌زمان تحولات فناوری و تدوین
هنجارهای اخلاقی» (فراوانی ۱۰) و «کنترل فشار نظام ارزیابی پژوهش و کاهش کمیت‌گرایی»
(فراوانی ۱۰) از مهم‌ترین عوامل محیطی اثرگذار در انتشار اخلاقی نتایج پژوهش هستند.

جدول ۴. مولفه‌های محیطی اثرگذار بر پایداری به اخلاق نشر در پژوهش‌های مبتنی بر هوش مصنوعی

ابعاد	مقوله‌های اصلی	مقوله‌های فرعی	فراوانی
فناوری	مدیریت فشار و تحولات ساختاری و کلان فناوری	• مدیریت هم‌زمان تحولات فناوری و تدوین هنجارهای اخلاقی؛	۱۰
		• کنترل فشار نظام ارزیابی پژوهش و کاهش کمیت‌گرایی؛	۱۰
		• کاهش ابهام و تعارض در استانداردهای ملی و بین‌المللی؛	۵
	نقش مجلات و ناشران در اخلاق نشر در هوش مصنوعی	• تدوین دستورالعمل‌ها و استانداردهای اخلاقی از سوی مجلات؛	۹
		• آموزش داوران و سردبیران برای تشخیص و آشنایی با کدهای اخلاقی AI؛	۱۲
		• الزام به شفاف‌سازی منابع و دسترسی به داده‌ها؛	۵

بعد شایستگی‌های حرفه‌ای نیز بر توسعه مهارت‌ها و دانش پژوهشگران برای استفاده مسئولانه و
اخلاق‌محور از AI تمرکز دارد. سه مقوله فرعی با بیشترین فراوانی شامل «کسب سواد اخلاقی و
سواد دیجیتال» (۱۹)، «تفکر انتقادی در ارزیابی خروجی AI» (۱۶) و «رعایت بایدها و نبایدها در
کاربست اخلاقی AI» (۱۴) هستند. این نتایج نشان می‌دهد که پژوهشگران برای اطمینان از
شفافیت، دقت و اصالت داده‌ها و نتایج تولیدشده توسط AI نیازمند آموزش اخلاقی مستمر، توانایی
تحلیل انتقادی خروجی‌ها و آشنایی کامل با چارچوب‌های اخلاقی مرتبط با فناوری‌های هوشمند
هستند (جدول ۵).

جدول ۵. مولفه‌های شایستگی‌های حرفه‌ای اثرگذار بر پایداری به اخلاق نشر در پژوهش‌های مبتنی بر هوش مصنوعی

ابعاد	مقوله‌های اصلی	مقوله‌های فرعی	فراوانی
-------	----------------	----------------	---------

- | | | |
|--|--|----|
| توانمندسازی حرفه‌ای و چارچوب‌سازی | • کسب سواد اخلاقی و سواد دیجیتال؛ | ۱۹ |
| | • تفکر انتقادی در ارزیابی خروجی AI؛ | ۱۶ |
| | • توجه به سواد هوش مصنوعی با یادگیری مستمر و مطابق با تحولات فناوری؛ | ۷ |
| رشد مهارت‌های حرفه‌ای-اخلاقی در کاربرست هوش مصنوعی | • رعایت باید‌ها و نبایدها در کاربرست اخلاقی AI؛ | ۱۴ |
| | • آشنایی با ابزارهای تخصصی AI در چارچوب‌های اخلاقی؛ | ۱۳ |
| | • پیشگیری از تخلفات با کسب مهارت‌های اخلاقی؛ | ۷ |

۴-۲. اعتبارسنجی مدل پابندی به اخلاق نشر در پژوهش‌های مبتنی بر هوش مصنوعی

در این مطالعه برای تجزیه و تحلیل داده‌ها و اعتبارسنجی مدل «پابندی به اخلاق نشر در پژوهش‌های مبتنی بر هوش مصنوعی»، از روش مدل‌سازی معادلات ساختاری با رویکرد حداقل مربعات جزئی (PLS-SEM) و نرم‌افزار SmartPLS نسخه ۴ استفاده شد. در جهت ساخت مفاهیم انتزاعی، پژوهش کنونی که به طراحی و بررسی ویژگی‌های روان‌سنجی مدل می‌پردازد، تحلیل داده‌های خود را در دو سطح توصیفی و استنباطی انجام داده است. در سطح استنباطی نیز ارزیابی مدل اندازه‌گیری، هم به صورت سلسله‌مراتب مرتبه اول و هم مرتبه دوم، مد نظر قرار گرفته است.

پیش از ارزیابی مدل، ویژگی‌های جمعیت‌شناختی ۱۲۴ نمونه مورد مطالعه، بررسی شد. یافته‌ها نشان داد که ۵۳/۲ درصد از شرکت‌کنندگان را زنان و ۴۶/۸ درصد را مردان تشکیل می‌دهند. از نظر مرتبه علمی، اکثریت نمونه را استادیاران (۵۴درصد) تشکیل داده و پس از آن دانشیاران (۳۲/۳ درصد) و استادان تمام (۱۳/۷ درصد) قرار داشتند. حضور بیش از ۴۶ درصدی پژوهشگران با مرتبه دانشجویی و بالاتر، حاکی از تخصص و آگاهی بالای پاسخ‌دهندگان نسبت به موضوعات اخلاق در نشر و فناوری‌های هوش مصنوعی است.

در گام بعدی جهت بررسی برازش مدل اندازه‌گیری، از سه معیار پایایی، روایی همگرا و روایی واگرا استفاده شد که نتایج آن در جدول ۶ گزارش شده است.

جدول ۶. شاخص‌های ضریب آلفای کروناخ، ضریب پایایی ترکیبی و میانگین واریانس استخراج شده

ابعاد	مولفه‌ها	آلفای کروناخ	پایایی ترکیبی	میانگین واریانس استخراج شده (AVE)	نتیجه
فردی	۱. مسئولیت‌پذیری اخلاقی	۰/۸۵۷	۰/۹۱۳	۰/۷۷۸	مطلوب
	۲. آگاهی و نگرش اخلاقی نسبت به ابزارهای هوش مصنوعی	۰/۷۰۸	۰/۸۷۲	۰/۷۷۳	مطلوب
	۳. خودکنترلی و رعایت استانداردهای اخلاقی	۰/۷۰۰	۰/۸۳۴	۰/۶۲۸	مطلوب
	۴. تعهد شخصی به شفافیت و صداقت علمی	۰/۷۶۸	۰/۸۹۶	۰/۸۱۱	مطلوب
سازمانی	۱. سیاست‌گذاری و راهبری نهادی	۰/۸۴۳	۰/۹۰۵	۰/۷۶۱	مطلوب
	۲. زیرساخت‌های آموزشی و توانمندسازی سازمانی	۰/۷۱۲	۰/۸۳۹	۰/۶۳۵	مطلوب
	۳. نقش کمیته‌های اخلاق و نظام‌های نظارتی	۰/۸۳۰	۰/۸۹۸	۰/۷۴۶	مطلوب
	۴. فرهنگ سازمانی	۰/۷۲۹	۰/۸۴۸	۰/۶۵۲	مطلوب
محیطی	۱. مدیریت فشار و تحولات ساختاری و کلان فناوری	۰/۷۳۲	۰/۸۴۹	۰/۶۵۲	مطلوب
	۲. نقش مجلات و ناشران در اخلاق نشر در هوش مصنوعی	۰/۸۱۲	۰/۸۸۹	۰/۷۲۸	مطلوب
شیاستی‌های حرفه‌ای	۱. توانمندسازی حرفه‌ای و چارچوب‌سازی	۰/۸۴۶	۰/۹۰۷	۰/۷۶۵	مطلوب
	۲. رشد مهارت‌های حرفه‌ای-اخلاقی در کاربست هوش مصنوعی	۰/۷۳۳	۰/۸۵۰	۰/۶۵۵	مطلوب

جهت بررسی تمایز مفهومی سازه‌ها، از معیار فورنل و لارکر^۱ استفاده شد. طبق این آزمون، روایی واگرا زمانی تأیید می‌شود که جذر AVE هر سازه (اعداد مندرج بر روی قطر اصلی ماتریس) از ضرایب همبستگی آن سازه با سایر متغیرهای مدل بیشتر باشد. مطابق با یافته‌های مندرج در جدول ۷، تمامی مقادیر روی قطر اصلی از مقادیر زیرین خود بزرگتر هستند؛ به عنوان مثال، جذر AVE برای سازه «آگاهی و نگرش اخلاقی نسبت به ابزارهای هوش مصنوعی» برابر با ۰/۸۷۹ است که از بیشترین همبستگی آن با سایر سازه‌ها (۰/۷۷۱) با سازه سیاست‌گذاری و راهبری نهایی بالاتر می‌باشد. این نتایج، استقلال مفاهیم و عدم همپوشانی نامطلوب میان ابعاد پژوهش را تأیید می‌کند.

جدول ۷. آزمون فورنل و لارکر مدل پیشنهادی

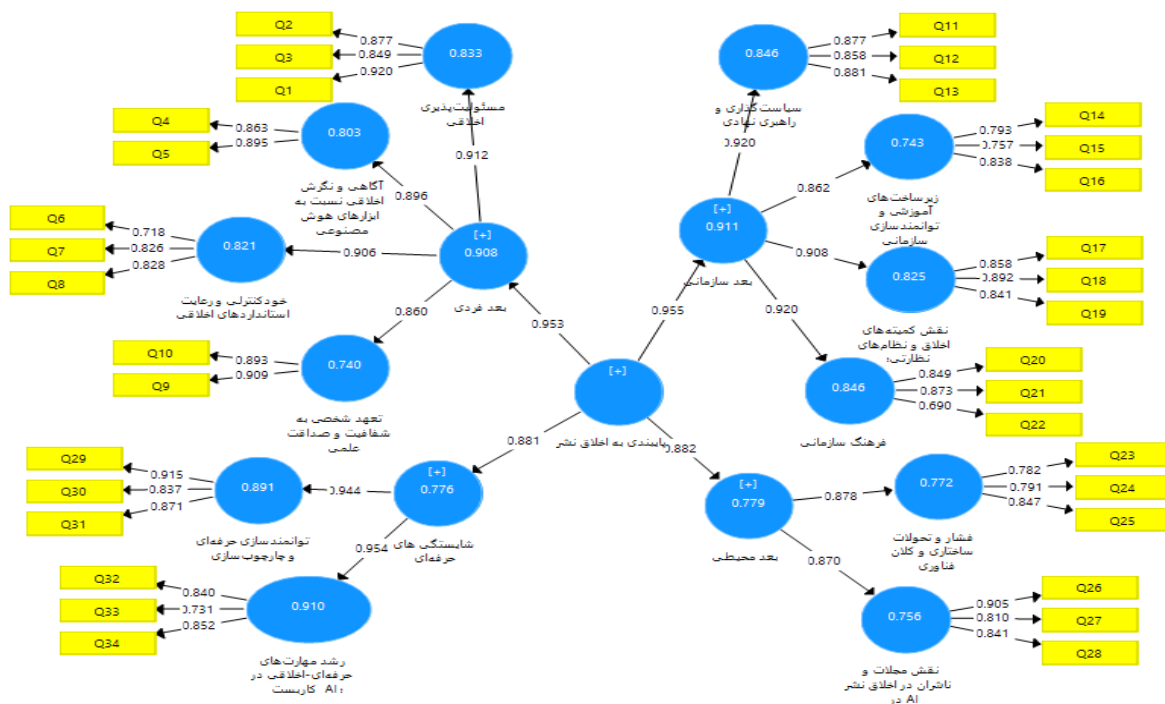
مؤلفه‌ها	آگاهی و نگرش اخلاقی نسبت به ابزارهای AI	تعهد شخصی به شفافیت و صداقت علمی	توانمندسازی حرفه‌ای و چارچوب‌سازی	خودکنترلی و رعایت استانداردهای اخلاقی	رشد مهارت‌های حرفه‌ای - اخلاقی در کاربرست AI	زیرساخت‌های آموزشی و توانمندسازی سازمانی	سیاست‌گذاری و راهبری نهادی	فرهنگ سازمانی	مدیریت فشار و تحولات ساختاری و کلان فناوری	مسئولیت‌پذیری اخلاقی	نقش مجلات و ناشران در اخلاق نشر در AI	نقش کمیته‌های اخلاق و نظام‌های فشارتی
آگاهی و نگرش اخلاقی نسبت به ابزارهای AI	۰/۸۷۹											
تعهد شخصی به شفافیت و صداقت علمی	۰/۷۱۷	۰/۹۰۱										
توانمندسازی حرفه‌ای و چارچوب‌سازی	۰/۶۳۹	۰/۶۶۹	۰/۸۷۵									
خودکنترلی و رعایت استانداردهای اخلاقی	۰/۷۵۲	۰/۷۰۲	۰/۶۷۶	۰/۷۹۲								
رشد مهارت‌های حرفه‌ای -	۰/۷۱۸	۰/۷۱۴	۰/۸۰۷	۰/۷۵۷	۰/۸۰۹							

¹ Fornell & Larcker

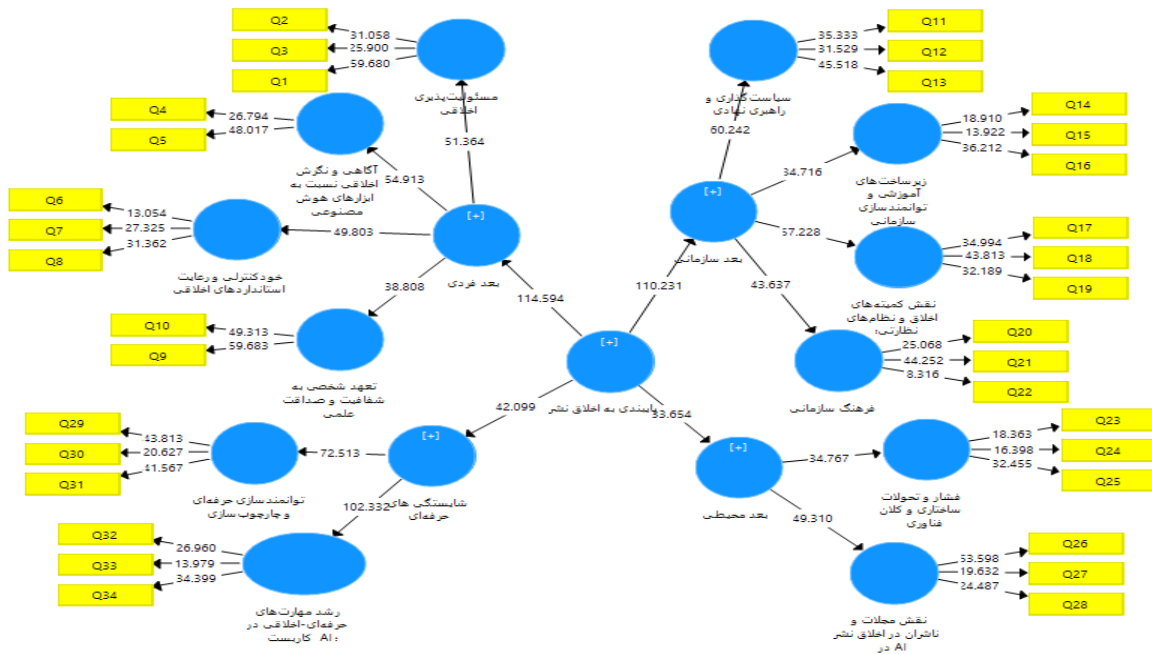
اخلاقی در کاربری AI											
زیرساخت‌های آموزشی و توانمندسازی سازمانی	۰/۶۱۹	۰/۵۵۳	۰/۵۷۴	۰/۷۰۱	۰/۶۷۱	۰/۷۹۷					
سیاست‌گذاری و راهبری نهادی	۰/۷۷۱	۰/۷۸۷	۰/۷۱۵	۰/۷۸۳	۰/۷۳۸	۰/۶۹۵	۰/۸۷۲				
فرهنگ سازمانی	۰/۷۱۹	۰/۷۲۱	۰/۷۰۹	۰/۷۲۰	۰/۷۲۵	۰/۷۰۹	۰/۸۳۸	۰/۸۰۸			
مدیریت فشار و تحولات ساختاری و کلان فناوری	۰/۶۰۳	۰/۵۷۴	۰/۶۵۴	۰/۵۸۶	۰/۶۴۱	۰/۵۵۳	۰/۷۰۹	۰/۷۲۳	۰/۸۰۷		
مسئولیت‌پذیری اخلاقی	۰/۷۵۹	۰/۷۴۲	۰/۶۳۳	۰/۷۴۱	۰/۷۶۲	۰/۵۷۷	۰/۷۶۰	۰/۷۲۰	۰/۶۷۲	۰/۸۸۲	
نقش مجلات و ناشران در اخلاق نشر در AI	۰/۶۷۷	۰/۶۹۸	۰/۶۱۵	۰/۵۷۶	۰/۵۸۶	۰/۵۲۰	۰/۶۲۶	۰/۶۷۳	۰/۵۲۹	۰/۶۲۷	۰/۸۵۳
نقش کمیته‌های اخلاق و نظام‌های نظارتی	۰/۶۹۵	۰/۶۸۹	۰/۵۸۸	۰/۷۱۰	۰/۶۴۲	۰/۷۰۷	۰/۷۸۸	۰/۷۸۵	۰/۵۴۰	۰/۶۹۶	۰/۶۰۶
											۰/۸۴۶

پس از تایید روایی و پایایی در سطح مولفه‌ها، ارزیابی مدل اندازه‌گیری سلسله‌مراتبی (مرتب دوم) به منظور اعتبارسنجی چگونگی بازنمایی ابعاد توسط سازه اصلی « پایبندی به اخلاق نشر در پژوهش‌های مبتنی بر هوش مصنوعی » انجام شد. شکل ۱، مدل اندازه‌گیری سلسله‌مراتبی را به همراه بارهای عاملی در حالت استاندارد نشان می‌دهد. شکل ۲ نیز بیانگر مقادیر آماره t مربوط به روابط اندازه‌گیری و بارهای عاملی است.

بررسی بارهای عاملی نشان می‌دهد که تمامی شاخص‌ها دارای بار عاملی بالاتر از ۰/۷ هستند. این امر بیانگر آن است که هر یک از مولفه‌ها از قدرت بازنمایی و تبیین مناسبی برخوردارند. بنابراین می‌توان نتیجه گرفت که شاخص‌های به کاررفته، از روایی همگرای مطلوبی در سنجش ابعاد پایبندی به اخلاق نشر برخوردارند (شکل ۱). همچنین نتایج حاصل از تحلیل بوت‌استرپ نشان می‌دهد که تمامی مقادیر آماره t بزرگ‌تر از ۱/۹۶ هستند که این موضوع حاکی از معناداری آماری بارهای عاملی در سطح اطمینان ۹۵ درصد است (شکل ۲).



شکل ۱. مدل اندازه گیری سلسله مراتبی سازه پایداری به اخلاق نشر در حالت استاندارد



شکل ۲. نتایج بوت استرپ مدل اندازه گیری سلسله مراتبی در حالت معناداری

پس از تایید پایایی و روایی ابعاد در مرتبه اول، ارزیابی مدل اندازه گیری سلسله مراتبی در مرتبه دوم به منظور اعتبارسنجی نحوه بازنمایی ابعاد چهارگانه (فردی، سازمانی، محیطی و شایستگی های حرفه ای) توسط سازه کلان «پایندی به اخلاق نشر در پژوهش های مبتنی بر هوش مصنوعی» انجام شد. با توجه به ماهیت توسعه ای و طراحی ابزار در این پژوهش، از مدل اندازه گیری سلسله مراتبی انعکاسی-انعکاسی^۱ استفاده شده است. در این ساختار، ابعاد مرتبه اول به عنوان شاخص های انعکاسی سازه کلان مرتبه دوم عمل می کنند.

جهت ارزیابی کیفیت این مدل اندازه گیری سلسله مراتبی، سه شاخص کلیدی بررسی شد. (۱) بارهای عاملی مرتبه دوم (ضرایب مسیر بیرونی)، که میزان اشتراک و قدرت ارتباط هر بعد را با سازه کلان نشان می دهد. (۲) ضریب تعیین (R^2)، که نشان می دهد چه میزان از واریانس هر یک از ابعاد مرتبه اول توسط

¹ Reflective-Reflective

سازه کلان مرتبه دوم تبیین و پوشش داده می شود. (۳) میانگین واریانس استخراج شده (AVE)، جهت اطمینان از حفظ روایی همگرا در سطح ساختار سلسله مراتبی. نتایج ارزیابی مدل اندازه گیری مرتبه دوم در جدول ۸ گزارش شده است.

جدول ۸. نتایج نهایی ارزیابی مدل اندازه گیری سلسله مراتبی (مرتبه دوم) پایداری به اخلاق نشر

ابعاد مدل	بارعاملی مرتبه دوم	آماره تی	مقدار معناداری	ضریب تعیین (R^2)	روایی همگرا (AVE)
پایداری به موازین اخلاق نشر - فردی	۰/۹۵۳	۱۱۴/۵۹۴	۰/۰۰۱	۰/۹۰۸	۰/۷۷۸
پایداری به موازین اخلاق نشر - سازمانی	۰/۹۵۵	۱۱۰/۲۳۱	۰/۰۰۱	۰/۹۱۲	۰/۷۶۱
پایداری به موازین اخلاق نشر - محیطی	۰/۸۸۲	۳۳/۶۵۴	۰/۰۰۱	۰/۷۷۸	۰/۶۵۲
پایداری به موازین اخلاق نشر - شایستگی های حرفه ای	۰/۸۸۱	۴۲/۰۹۹	۰/۰۰۱	۰/۷۷۶	۰/۷۶۵

همان طور که نتایج جدول ۸ نشان می دهد، تمامی بارهای عاملی مرتبه دوم بسیار بالا و بزرگتر از ملاک استاندارد ۰/۷۰ هستند که نشان دهنده همگرایی قوی ابعاد در بازنمایی سازه کلان پایداری به اخلاق نشر است. آماره های t به دست آمده از آزمون بوت استرپ همگی بزرگتر از ۱/۹۶ و در سطح $p < ۰/۰۰۱$ معنادار می باشند که اعتبار روابط اندازه گیری را تایید می کند.

مقادیر ضریب تعیین (R^2) برای ابعاد چهارگانه بین ۰/۷۷۵ تا ۰/۹۱۲ قرار دارد؛ این یافته ها نشان می دهند که سازه کلان «پایداری به اخلاق نشر»، سهم عمده و بسیار بالایی در تبیین واریانس ابعاد خود دارد. همچنین، مقادیر AVE برای تمامی ابعاد بالاتر از مرز پذیرش ۰/۵۰ تثبیت شده است که بقای روایی همگرا را در کل مدل اندازه گیری سلسله مراتبی اثبات می سازد.

علاوه بر شاخص های پایایی و روایی، برای بررسی میزان ناهمخوانی بین ماتریس همبستگی مشاهده شده و ماتریس بازتولید شده، شاخص ریشه میانگین مربعات باقیمانده استاندارد شده (SRMR) نیز بررسی شد. نتایج نشان داد مقدار SRMR در مدل اشباع شده برابر با ۰/۰۶۹ و در مدل تخمینی برابر با ۰/۰۸۳ است. مقدار ۰/۰۶۹ در مدل اشباع شده نشان دهنده برازش مناسب ساختار کلی اندازه گیری با

داده‌های مشاهده‌شده است. همچنین مقدار ۰/۰۸۳ در مدل تخمینی، اگرچه اندکی بالاتر از آستانه سخت‌گیرانه ۰/۰۸ قرار دارد، در چارچوب پژوهش‌های توسعه‌ای و اعتبارسنجی ابزار می‌تواند در محدوده قابل قبول تفسیر شود. با این حال، قضاوت نهایی درباره کفایت مدل صرفاً بر مبنای SRMR انجام نشد و شواهد اصلی اعتبار مدل بر پایه بارهای عاملی مرتبه دوم، ضرایب معناداری بوت‌استرپ، پایایی ترکیبی، میانگین واریانس استخراج‌شده و روایی واگرا استوار است.

۵. نتیجه‌گیری

مطالعه حاضر با هدف طراحی و اعتبارسنجی مدل پابندی به اخلاق نشر در پژوهش‌های مبتنی بر کاربست هوش مصنوعی انجام شد. یافته‌های این پژوهش نشان داد که صورت‌بندی این پدیده، مستلزم نگاهی جامع و فراتر از رویکردهای تک‌عاملی است؛ به‌طوری‌که پابندی پژوهشگران به اخلاق نشر در مواجهه با ابزارهای هوش مصنوعی، سازه‌ای پویا و چندسطحی است که در قالب چهار بعد اصلی فردی، سازمانی، محیطی و شایستگی‌های حرفه‌ای تبیین می‌شود. در ادامه، جهت مشخص شدن سهم و جایگاه هر یک از این ابعاد، یافته‌های حاصل با ادبیات تجربی و مبانی نظری این حوزه مورد بحث و تفسیر قرار می‌گیرد.

بر مبنای نتایج پژوهش، «بعد سازمانی» به عنوان قوی‌ترین تبیین‌کننده پابندی به موازین اخلاق نشر در پژوهش‌های علوم پزشکی وابسته به ابزارهای هوش مصنوعی شناسایی شد. این یافته حاکی از آن است که بدون وجود زیرساخت‌های نهادی، آیین‌نامه‌های شفاف و سازوکارهای نظارتی کارآمد، نمی‌توان انتظار پابندی پایدار از سوی پژوهشگران داشت. دلیل اهمیت این بعد را می‌توان در ماهیت جمعی و ساختاریافته پژوهش در دانشگاه‌ها جستجو کرد. در میان مؤلفه‌های این بعد، «سیاست‌گذاری و راهبری نهادی» و «نقش کمیته‌های اخلاق و نظام‌های نظارتی» از اهمیت ویژه‌ای برخوردار بودند. تدوین و به‌روزرسانی آیین‌نامه‌های اخلاقی متناسب با تحولات فناوری، شفاف‌سازی الزامات استفاده از هوش مصنوعی و ایجاد هماهنگی بین دانشگاه‌ها و مجلات، بستر قانونی لازم را برای هدایت رفتار پژوهشگران فراهم می‌کند. از سوی دیگر، داوری اخلاقی دقیق پروپوزال‌ها با تمرکز بر نحوه کاربست هوش مصنوعی و نظارت مستمر بر اجرای پژوهش‌ها، ضمانت اجرایی این قوانین را تأمین می‌نماید. این یافته با پژوهش عباس‌زاده و همکاران (۱۳۹۵) همخوانی دارد. عبداللهی و همکاران (۱۳۹۷) نیز در الگوی ارتقای اخلاق پژوهش، بر اهمیت سیاست‌گذاری و راهبردهای سازمانی تأکید کرده‌اند. در سطح بین‌المللی، یافته‌های پژوهش Shaw

et al. (2024) در نشست جهانی اخلاق زیستی سازمان جهانی بهداشت، بر ضرورت روزآمدسازی مقررات مربوط به هوش مصنوعی در سلامت، تمرکز بر حاکمیت اخلاق در داده‌ها، ارزیابی تأثیر هوش مصنوعی و نظارت بر اقدامات صورت گرفته از سوی هوش مصنوعی تأکید دارد. این موارد همگی ذیل بعد سازمانی قرار می‌گیرند و اهمیت این بعد را در سطح جهانی نیز تأیید می‌کنند.

«بعد فردی» نیز با ضریب تأثیر بسیار نزدیک به بعد سازمانی، دومین تبیین‌کننده مهم پایداری به موازین اخلاق نشر در پژوهش‌های علوم پزشکی وابسته به ابزارهای هوش مصنوعی است. این یافته نشان می‌دهد که اگرچه زیرساخت‌های سازمانی لازم هستند، اما در نهایت این پژوهشگران هستند که در مقام عمل تصمیم‌گیرنده‌اند و باید مسئولیت اخلاقی کاربست هوش مصنوعی را بپذیرند. در میان مؤلفه‌های این بعد، «مسئولیت‌پذیری اخلاقی» و «تعهد شخصی به شفافیت و صداقت علمی» از بالاترین اهمیت برخوردار بودند. پذیرش مسئولیت متون تولیدشده با هوش مصنوعی، ارزیابی دقیق خروجی آن‌ها پیش از پذیرش و حفظ نگرش انسانی در تحلیل و تفسیر نتایج، از مصادیق بارز مسئولیت‌پذیری اخلاقی هستند. همچنین رعایت شفافیت در کاربست هوش مصنوعی و حقوق مالکیت فکری، نشان‌دهنده تعهد پژوهشگر به صداقت علمی است. این یافته‌ها با پژوهش Kim (2024) درباره چالش‌های اخلاقی چت‌جی‌بی‌تی همخوانی کامل دارد. Kasani et al (2024) نیز در مطالعه خود به این نتیجه رسیدند که کاربست هوش مصنوعی در پژوهش، هم پتانسیل‌های جذاب و هم چالش‌های اخلاقی به ویژه در زمینه نویسنده‌گی و نشر نتایج علمی به همراه دارد. آنان بر ضرورت کنترل و نظارت متخصصان انسانی بر فرآیندهای پژوهش تأکید کرده‌اند که با مؤلفه «حفظ نگرش انسانی در تحلیل و تفسیر نتایج» در پژوهش حاضر تطابق دارد.

بعد «شایستگی‌های حرفه‌ای» با وجود ضریب تأثیر مشابه بعد محیطی، از نظر ماهیت، وجهی زیربنایی دارد. این بعد به توانایی‌ها و مهارت‌هایی اشاره دارد که پژوهشگر را برای عمل اخلاقی در مواجهه با هوش مصنوعی مجهز می‌کند. دلیل اهمیت این بعد را می‌توان در پیچیدگی فزاینده فناوری‌های هوش مصنوعی و سرعت تحول آن‌ها جستجو کرد. در میان مؤلفه‌های این بعد، «کسب سواد اخلاقی و سواد دیجیتال» و «تفکر انتقادی در ارزیابی خروجی هوش مصنوعی» از اهمیت ویژه‌ای برخوردار بودند. همچنین «یادگیری مستمر و مطابق با تحولات فناوری» به عنوان یکی دیگر از مؤلفه‌های کلیدی شناسایی شد. این یافته‌ها با پژوهش نیک‌روان فرد و همکاران (۲۰۱۷) درباره وضعیت آموزش اخلاق پژوهش در برنامه‌های درسی تحصیلات تکمیلی علوم پزشکی ایران ارتباط نزدیکی دارد. پژوهش Schroter et al. (2018) نیز به

اهمیت کیفیت آموزش‌های اخلاقی اشاره کرده و نشان داده است که تنها تعداد محدودی از آموزش‌ها از نظر کیفیت در سطح عالی بودند. این یافته بر ضرورت طراحی برنامه‌های آموزشی ساختارمند، هدفمند و مبتنی بر نیازهای واقعی پژوهشگران تأکید می‌کند. صرف برگزاری کارگاه‌های یک‌روزه و سطحی، نمی‌تواند شایستگی‌های لازم برای مواجهه با چالش‌های اخلاقی هوش مصنوعی را ایجاد کند.

«بعد محیطی» با وجود ضریب تأثیر نسبتاً پایین‌تر، نقشی حیاتی در تکمیل مدل پابندی به موازین اخلاق نشر در پژوهش‌های علوم پزشکی وابسته به ابزارهای هوش مصنوعی ایفا می‌کند. این بعد به عواملی اشاره دارد که در سطح فرادانشگاهی و فراملی عمل می‌کنند و عمدتاً خارج از کنترل مستقیم دانشگاه‌ها و پژوهشگران هستند. دلیل اهمیت این بعد را می‌توان در وابستگی متقابل نظام‌های علمی در عصر جهانی شدن جستجو کرد. دانشگاه‌ها و پژوهشگران امروزی در شبکه‌ای از ارتباطات بین‌المللی فعالیت می‌کنند و استانداردهای نشر جهانی، سیاست‌های مجلات بین‌المللی و تحولات فناوری در سطح جهان، همگی بر رفتار آنان تأثیر می‌گذارند. در میان مؤلفه‌های این بعد، «نقش مجلات و ناشران در اخلاق نشر هوش مصنوعی» و «مدیریت فشار و تحولات ساختاری و کلان فناوری» از اهمیت بیشتری برخوردار بودند. تدوین دستورالعمل‌ها و استانداردهای اخلاقی از سوی مجلات، آموزش داوران و سردبیران برای آشنایی با کدهای اخلاقی هوش مصنوعی و الزام به شفاف‌سازی منابع، از جمله اقداماتی است که مجلات می‌توانند برای هدایت رفتار اخلاقی پژوهشگران انجام دهند. از سوی دیگر، مدیریت هم‌زمان تحولات فناوری و تدوین هنجارهای اخلاقی، کنترل فشار نظام ارزیابی پژوهش و کاهش کمیت‌گرایی، و کاهش ابهام و تعارض در استانداردهای ملی و بین‌المللی، چالش‌هایی هستند که در سطح محیطی بر پابندی اخلاقی تأثیر می‌گذارند. یافته‌های پژوهش Zhang et al. (2024) نشان داد که این موضوع در مجلات زیست‌پزشکی بیش از دیگر علوم مورد توجه قرار گرفته و ضرورت افزایش همکاری میان کشورها، نویسندگان و مؤسسات برای ارتقای اخلاق‌مداری در انتشارات علمی، از نتایج کلیدی این پژوهش بود. این یافته با مؤلفه «هماهنگی نهادی بین دانشگاه‌ها و مجلات» در پژوهش حاضر ارتباط دارد. لازم به ذکر است که این پژوهش با وجود تلاش برای دقت روش‌شناختی، با چند محدودیت همراه بود. در فاز کیفی و اجرای مصاحبه‌ها، دسترسی به برخی سیاست‌گذاران کلان و افراد کلیدی حوزه اخلاق پژوهش و نشر به دلیل محدودیت‌های زمانی و اداری با چالش‌هایی همراه بود. در فاز کمی نیز استفاده از پرسشنامه خودگزارش‌دهی می‌تواند با سوگیری مطلوبیت اجتماعی همراه باشد. لذا به دلیل محدود بودن جامعه مطالعه به دانشگاه‌های علوم پزشکی، تعمیم نتایج به سایر حوزه‌های علمی باید با احتیاط انجام شود.

با این تفاسیر، پابندی به موازین اخلاق نشر در پژوهش‌های وابسته به ابزارهای هوش مصنوعی در دانشگاه‌های علوم پزشکی، پدیده‌ای چندوجهی و مجموعه‌ای از عوامل درهم‌تنیده است. لذا دستیابی به نظام پژوهشی اخلاق‌مدار در عصر هوش مصنوعی، نیازمند عزمی ملی و همکاری هماهنگ در سطوح مختلف است. وزارت بهداشت، درمان و آموزش پزشکی به عنوان سیاست‌گذار کلان، می‌بایست با تدوین و ابلاغ آیین‌نامه‌های جامع و ایجاد سازوکارهای نظارتی کارآمد، بستر قانونی را فراهم آورد. دانشگاه‌ها به عنوان مجریان اصلی، موظف به ایجاد زیرساخت‌های آموزشی و توانمندسازی کمیته‌های اخلاق هستند. سردبیران مجلات نیز با وضع استانداردهای شفاف و آموزش داوران، نقش مهمی در هدایت رفتار اخلاقی پژوهشگران ایفا می‌کنند و در نهایت، پژوهشگران به عنوان کنشگران اصلی، باید با پذیرش مسئولیت و تعهد به شفافیت، نشان دهند که هوش مصنوعی را نه جایگزینی برای تفکر و تعهد انسانی، که ابزاری در خدمت حقیقت‌جویی علمی می‌دانند. بر این اساس، اتخاذ رویکردی چندسطحی و هماهنگ میان ارکان یادشده؛ پیش‌شرط بهره‌گیری بهینه از فناوری هوش مصنوعی در پژوهش‌های علوم پزشکی، هم‌زمان با حفظ اصالت، اعتبار و امانت‌داری علمی در خروجی‌های پژوهشی است. جهت توسعه این حوزه، پیشنهاد می‌شود در پژوهش‌های آینده، اثربخشی سیاست‌ها و دستورالعمل‌های اخلاقی موجود در مجلات و دانشگاه‌های علوم پزشکی در سطح ملی مورد ارزیابی تجربی قرار گیرد.

قدردانی

پژوهشگران مراتب سپاس و قدردانی خود را از مشارک‌کنندگان در فازهای کیفی و کمی مطالعه، جهت همکاری در این پژوهش، اعلام می‌دارند. لازم به ذکر است که این مقاله حاصل طرح تحقیقاتی پسادکتری دانشگاه علوم پزشکی تهران به شماره ۱۴۰۳-۳-۱۰۲-۷۴۲۴۴ و کد اخلاق IR.TUMS.SPH.REC.1403.184 است. همچنین این اثر تحت حمایت مادی بنیاد ملی علم ایران (INSF) برگرفته شده از طرح شماره «۴۰۳۷۴۷۰» انجام شده است.

فهرست منابع

- بهمن آبادی، علیرضا. ۱۴۰۲. هوش مصنوعی و کاربردهای آن در فعالیتهای پژوهشی. علوم و فناوری اطلاعات کشاورزی، ۶ (۲)، ۳۳-۴۳.
- خادمی زاده، شهناز. ۱۴۰۳. هوش مصنوعی و رعایت اخلاق در نگارش تحقیقات علمی. اولین کنگره ملی پژوهش پاک.
- دخش، سارا. ۱۴۰۲. آینده نگاری اخلاق نشر در سیاست گذاری های نظام پژوهشی وزارت علوم، تحقیقات و فناوری. رساله دکتری. گروه علم اطلاعات و دانش شناسی، دانشکده علوم تربیتی و روانشناسی. دانشگاه شهید چمران اهواز.
- دخش، سارا؛ خادمی زاده، شهناز؛ فرج پهلوی، عبدالحسین؛ فرهادی راد، حمید. ۱۴۰۳. عوامل اثرگذار بر پایداری به موازین اخلاق نشر در نظام پژوهشی وزارت عتف. پژوهشنامه پردازش و مدیریت اطلاعات، ۳۹ (۴)، ۱۱۴۱-۱۱۷۰.
- شریف زاده، رحمان. ۱۴۰۴. تحلیل و بررسی مزایا و چالش های اخلاقی مدل های زبانی بزرگ در سپهر پژوهش. پژوهشنامه پردازش و مدیریت اطلاعات، ۴۰ (۴)، ۱۲۱۹-۱۲۵۰.
- عباس زاده، محمد؛ بنی فاطمه، حسین؛ علیزاده اقدم، محمدباقر؛ بوداقي، علی. ۱۳۹۵. عوامل زمینه ساز پایداری به اخلاق پژوهش در میان دانشجویان تحصیلات تکمیلی: مورد دانشگاه تبریز. پژوهش و برنامه ریزی در آموزش عالی، ۲۲ (۱)، ۷۵-۹۸.
- عبداللهی، حسین؛ طاهری، مرتضی؛ غیاثی، سعید؛ نعمتی، محمدعلی؛ بابایی محمدرضا. ۱۳۹۷. الگوی ارتقای اخلاق پژوهش؛ شرایط، بسترها، راهبردها و پیامدها. مدیریت اسلامی، ۲۶ (۴)، ۱۳۵-۱۵۷.
- مرتضوی شاهرودی، سید محمدعلی؛ زارعی، عیسی. ۱۴۰۳. اصول و چالش های اخلاقی استفاده از هوش مصنوعی در تحقیقات علمی. اخلاق پژوهی، ۷ (۴)، ۲۵-۵.
- وزارت بهداشت، درمان و آموزش پزشکی. ۱۳۹۶. راهنمای کشوری اخلاق در انتشار آثار پژوهشی. بانک اطلاعات نشریات علوم پزشکی کشور، دسترسی در ۱۴۰۵/۲/۱۵

https://ethics.research.ac.ir/docs/publication_guideline.pdf

References

- Akinrinmade, A. O., Adebile, T. M., Ezuma-Ebong, C., Bolaji, K., Ajufo, A., Adigun, A. O., ... & Okobi, O. E. 2023. Artificial intelligence in healthcare: perception and reality. Cureus, 15(9), e45594.

- Bjelobaba, S., Waddington, L., Perkins, M., Foltýnek, T., Bhattacharyya, S., & Weber-Wulff, D. 2025. Maintaining research integrity in the age of GenAI: an analysis of ethical challenges and recommendations to researchers. *International Journal for Educational Integrity*, 21(1), 18.
- Castillo-Montoya, M. 2016. Preparing for interview research: The interview protocol refinement framework. *Qualitative report*, 21(5), 811-831.
- Committee on Publication Ethics Council. 2021. COPE discussion document: Artificial intelligence (AI) in decision making (Online Retrieved 9/3/2025). <https://doi.org/10.24318/9kvAgmJ>
- Delanghe, J. R. 2026. The ethical aspects of AI in scientific publishing. *Ejifcc*, 37(1), 177.
- Dorton, S. L., Ministerio, L. M., Alaybek, B., & Bryant, D. J. 2023. Foresight for ethical AI. *Frontiers in Artificial Intelligence*, 6:1143907.
- Kasani, P. H., Cho, K. H., Jang, J. W., Yun, C. H. 2024. Influence of artificial intelligence and chatbots on research integrity and publication ethics. *Science Editing*, 11(1), 12-25.
- Kim, S. J. 2024. Research ethics and issues regarding the use of ChatGPT-like artificial intelligence platforms by authors and reviewers: a narrative review. *Science Editing*, 11(2), 96-106.
- Lin, A., Chen, Z., Jiang, A., Tang, B., Qi, C., Zhu, L., ... & Luo, P. 2026. Navigating academic integrity in biomedical research: the impact of large language models on current practices and future directions. *International Journal of Surgery*, 112(2), 4418-4433.
- Nikravanfard, N., Khorasanizadeh, F., & Zendehehdel, K. 2017. Research Ethics Education in Post-Graduate Medical Curricula in IR Iran. *Developing world bioethics*, 17(2), 77-83.
- Pearson, G. S. 2024. Artificial Intelligence and Publication Ethics. *Journal of the American Psychiatric Nurses Association*, 30(3), 453-455.
- Peh, W. C. G., & Saw, A. 2023. Artificial intelligence: Impact and challenges to authors, journals and medical publishing. *Malaysian Orthopaedic Journal*, 17(3), 1-4.
- Schroter, S., Roberts, J., Loder, E., Penzien, D. B., Mahadeo, S., & Houle, T. T. 2018. Biomedical authors' awareness of publication ethics: an international survey. *BMJ open*, 8(11), e021282.
- Shaw, J., Ali, J., Atuire, C. A., Cheah, P. Y., Español, A. G., Gichoya, J. W., ... & Vayena, E. 2024. Research ethics and artificial intelligence for global health: perspectives from the global forum on bioethics in research. *BMC Medical Ethics*, 25(1), 46.
- Son, W. J., & Yun, C. H. 2021. Consultation questions on publication ethics from 2016 to 2020 addressed by the Committee on Publication Ethics of the Korean Council of Science Editors. *Science Editing*, 8(1), 112-116.
- Zhang, M., Xu, J., Xu, C., Zheng, Q., Liu, M., Zhang, J., ... & Tian, J. 2024. Bibliometric review and mapping analysis of publication ethics research. *Ethics & Behavior*, 1-13.

Designing and Validating a Publication Ethics Compliance Model in Research Based on the Use of Artificial Intelligence: A Study in Iranian Universities of Medical Sciences

Sara Dakhesh

Postdoctoral researcher, Department of Medical Library and Information Sciences, School of Allied Medical Sciences, Tehran University of Medical Sciences, Tehran, Iran.; E-mail: saradakhesh@gmail.com

Shahnaz Khademizadeh*

(Corresponding author) PhD in Knowledge and Information Science; Professor; Department of Knowledge and Information Science; Shahid Chamran University of Ahvaz; Ahvaz, Iran; E-mail: s.khademi@scu.ac.ir

Zeinab Mohammadi

PhD in Knowledge and Information Science; Assistant Professor; Department of Knowledge and Information Science; Shahid Chamran University of Ahvaz; Ahvaz, Iran; E-mail: z-mohammadi@scu.ac.ir

Abstract

Although modern advancements, such as the emergence of artificial intelligence (AI) tools, have enabled the integration of cutting-edge technologies into various studies, scientific inquiry and research cannot bypass established standards and values, as they are inherently bound by ethical principles. This study aimed to develop a model for publication ethics compliance in AI-driven research within Iranian universities of medical sciences.

This study employed an exploratory sequential mixed-methods design and was applied in terms of its purpose. In the qualitative phase, semi-structured interviews were conducted with 20 academic actors in the research system at both macro levels (directors, policymakers, and research vice-chancellors of the Ministry of Health and universities) and micro levels (editors-in-chief, peer reviewers, editorial board members, and researchers in the fields of research ethics and AI). Purposive sampling was performed using critical case and snowball strategies. Qualitative data were analyzed using thematic analysis in MAXQDA (version 20), resulting in the extraction of compliance dimensions across four categories: individual, organizational, environmental, and professional competencies. In the quantitative

phase, the proposed model was operationalized into a researcher-developed questionnaire with 34 items and completed by 124 experts in research ethics and AI. The data were analyzed using partial least squares structural equation modeling (PLS-SEM) in SmartPLS (version 4) to evaluate the model fit.

The findings indicated that the factors influencing the ethical application of artificial intelligence tools in the publication of research results fall into four categories: individual (52 interpretive codes), organizational (145 interpretive codes), environmental (51 interpretive codes), and professional competencies (76 interpretive codes). The results also showed that the construct of adherence to publication ethics in AI-based medical research has a positive and significant effect on all four dimensions. Among these, the organizational dimension demonstrated the strongest effect, with a path coefficient of 0.955 and a t-value of 110.231, indicating that organizational factors play the most influential role in explaining adherence to publication ethics in AI-based medical research. This was followed by the individual (0.953), environmental (0.882), and professional competencies (0.881) dimensions. Since the t-values for all relationships were considerably higher than 1.96 and the p-values were reported as 0.000, all paths in the model were statistically significant at the 95% confidence level. These findings indicate that the structural model possesses very strong explanatory power.

Publication ethics compliance in AI-driven medical research is a multidimensional phenomenon shaped by interrelated factors. The findings of this study can serve as an operational framework for ethical policymaking, the formulation of national guidelines, and the design of researcher empowerment programs in universities of medical sciences.

Keywords: Research Ethics, Publication Ethics, Scientific Misconduct, Large Language Models, Artificial Intelligence, Researchers, Medical Sciences.

سارا دخش

متولد سال ۱۳۷۰، دارای پسادکتری دانشگاه علوم پزشکی تهران و دکتری علم اطلاعات و دانش‌شناسی از دانشگاه شهید چمران اهواز است. ایشان هم‌اکنون استادیار گروه علم اطلاعات و دانش‌شناسی دانشگاه یزد است. علم‌سنجی، ترجمان دانش، مدیریت دانش، اخلاق نشر، اخلاق در پژوهش و فناوری اطلاعات، هوش مصنوعی در پژوهش، آینده‌نگاری و سیاست‌گذاری علم از جمله علایق پژوهشی وی است.



شهناز خادمی‌زاده

متولد سال ۱۳۵۹، دارای مدرک تحصیلی دکتری در رشته علم اطلاعات و دانش‌شناسی از دانشگاه میسور هندوستان است. ایشان هم‌اکنون استاد گروه علم اطلاعات و دانش‌شناسی دانشگاه شهید چمران اهواز است. مدیریت دانش، سیستم‌های اطلاعاتی، داده‌کاوی، ارزیابی و سیاست‌گذاری پژوهش از جمله علایق پژوهشی وی است.



زینب محمدی

متولد سال ۱۳۶۱، دارای مدرک تحصیلی دکتری در رشته علم اطلاعات و دانش‌شناسی از دانشگاه شهید چمران اهواز است. ایشان هم‌اکنون استادیار گروه علم اطلاعات و دانش‌شناسی دانشگاه شهید چمران اهواز است. مدیریت دانش، سواد اطلاعاتی، مدیریت کتابخانه‌های عمومی و کاربرد هوش مصنوعی در پژوهش از جمله علایق پژوهشی وی است.

