

ادراک اساتید از چارچوب اخلاقی استفاده مسئولانه از هوش مصنوعی در پژوهش‌های آموزشی

امیرحسین کیزوری***

محمد محمودی بورنگ**

میثم غلام پور*

* استادیار گروه علوم تربیتی، دانشکده ادبیات و علوم انسانی، دانشگاه حکیم سبزواری، سبزوار، ایران (نویسنده مسئول) m.gholampour@Hsu.ac.ir

** دکتری مطالعات برنامه‌ریزی درسی، دانشگاه بیرجند، بیرجند، ایران m.mahmoodi@birjand.ac.ir

*** دانشیار گروه علوم تربیتی، دانشکده ادبیات و علوم انسانی، دانشگاه حکیم سبزواری، سبزوار، ایران akayzouri@hsu.ac.ir

تاریخ پذیرش مقاله: ۱۴۰۴/۱۲/۱

تاریخ شروع بررسی: ۱۴۰۴/۴/۲۳

تاریخ دریافت مقاله: ۱۴۰۴/۰۴/۱۱

چکیده

هدف پژوهش حاضر، تبیین ادراک اساتید از چارچوب اخلاقی استفاده از هوش مصنوعی در پژوهش‌های آموزشی است. این مطالعه با رویکرد کیفی و به روش پدیدار نگاری انجام شد. جامعه آماری، شامل اساتید و متخصصان در زمینه هوش مصنوعی و اخلاق پژوهش بودند و نمونه مدنظر با استفاده از روش نمونه‌گیری هدفمند و بر مبنای ملاک‌های از پیش تعیین‌شده انتخاب شدند. فرایند گردآوری داده‌ها با بهره‌گیری از مصاحبه‌های نیمه ساختاریافته تا دستیابی به اشباع نظری ادامه یافت. یافته‌ها در قالب هفت مضمون فراگیر، ۱۸ مضمون سازمان دهنده و ۵۶ مضمون پایه دسته‌بندی شد. ابعاد همچون، ابعاد اخلاقی در فرآیند پژوهش (مشمول بر مؤلفه‌هایی چون شفافیت در مدیریت داده، حفاظت حقوق فردی مشارکت‌کننده و شفافیت و محافظت اخلاقی در انتشار نتایج پژوهش)؛ معیارهای اخلاقی در بُعد تعهد (مشمول بر مؤلفه‌هایی چون تعهد حرفه‌ای محقق/محققین در به‌کارگیری هوش مصنوعی، تضمین عدالت و انسان‌محوری در تحلیل‌های هوش‌مبنا)؛ معیارهای اخلاقی ساختاری-زمینه‌ای مرتبط با هوش مصنوعی (مشمول بر محورهای چون تدوین و اجرای مقررات اخلاق پژوهی در کاربرد هوش مصنوعی، نهادینه‌سازی سواد اخلاقی و ترویج هنجارهای اخلاق پژوهی در استفاده از هوش مصنوعی و هم‌راستاسازی استانداردهای اخلاقی در سطوح سازمانی، ملی و بین‌المللی در زمینه کاربرد پژوهشی هوش مصنوعی)؛ حقوق مالکیت در کاربرد هوش مصنوعی و ابعاد اخلاق رفتاری در کاربرد هوش مصنوعی در پژوهش مورد سازمان‌دهی قرار گرفت. بر مبنای یافته‌ها می‌توان نتیجه گرفت توجه به ایجاد منشور اخلاقی در زمینه کاربرد هوش مصنوعی در پژوهش در ابعاد مختلف زمینه سازمان‌دهی استفاده از این ابزار در پژوهش‌ها می‌شود.

واژگان کلیدی: ابعاد اخلاقی، هوش مصنوعی، پژوهش، منشور اخلاقی

Professors' Perception of the Ethical Framework for Responsible Use of Artificial Intelligence in Academic Research

Meysam Gholampour*

Mohammad Mahmoodi**

Amir hossein kayzouri***

* Assistant Professor, Department of Educational Sciences, Faculty of Literature and Human Sciences, Hakim

Sabzevari University, Sabzevar, Iran. (Corresponding Author) m.gholampour@Hsu.ac.ir

** PhD in Curriculum Studies, University of Birjand, Birjand, Iran. m.mahmoodi@birjand.ac.ir

*** Associate Professor, Department of Educational Sciences, Faculty of Literature and Humanities, Hakim

Sabzevari University, Sabzevar, Iran akayzouri@hsu.ac.ir

Abstract

The aim of the present study is to explain the professors' perception of the ethical framework of using artificial intelligence in academic research. This study was conducted with a qualitative approach and phenomenographic method. The statistical population included university professors and experts in the field of artificial intelligence and research ethics, and the sample was selected using a purposive sampling method based on predetermined criteria. The data collection process continued using semi-structured interviews until theoretical saturation was achieved. The findings were categorized into seven overarching themes, 18 organizing themes, and 56 basic themes. Dimensions such as ethical dimensions in the research process (including components such as transparency in data management, protection of individual rights of participants, and transparency and ethical protection in publishing research results); ethical standards in the commitment dimension (including components such as professional commitment of the researcher/researchers in using artificial intelligence, ensuring justice, and human-centeredness in intelligence-based analyses); Structural-contextual ethical standards related to artificial intelligence (including areas such as the development and implementation of research ethics regulations in the application of artificial intelligence, institutionalization of ethical literacy and promotion of research ethics norms in the use of artificial intelligence, and alignment of ethical standards at organizational, national, and international levels in the field of research application of artificial intelligence); property rights in the application of artificial intelligence, and dimensions of behavioral ethics in the application of artificial intelligence in academic research were organized. Based on the findings, it can be concluded that attention should be paid to creating a charter of ethics in the application of artificial intelligence in academic research in various dimensions of organizing the use of this tool in university research.

Keywords: Ethical dimensions, artificial intelligence, research, code of ethics

هوش مصنوعی سریع‌ترین فناوری در حال رشد در جهان است و از ظرفیت عظیمی برای بازنویسی قوانین کل صنایع، ایجاد رشد اقتصادی قابل توجه و متحول کردن تمام زمینه‌های زندگی برخوردار است (پروانه و میرباقری، ۱۴۰۲). در چند سال اخیر نقش هوش مصنوعی در جنبه‌های مختلف زندگی بشر به‌طور فرایندهای پررنگ شده است و تقریباً بر روی تمامی عرصه‌ها از جمله زندگی روزمره، بازی و سرگرمی، علم و هنر و همچنین حوزه سلامت، پزشکی و بهداشت تأثیرات وسیعی گذاشته است (تای^۱، ۲۰۲۰). این فناوری‌ها به‌واسطه توانمندی‌هایی چون تحلیل کلان داده‌ها، کشف الگوها و پیش‌بینی روندها، پژوهشگران را قادر ساخته‌اند تا با دقت و سرعت بیشتری به نتایج تحقیق دست یابند (جوبین^۲ و همکاران، ۲۰۱۹). متخصصان پیش‌بینی می‌کنند که هوش مصنوعی موتور محرک اصلی انقلاب صنعتی چهارم خواهد بود (قریشی^۳، ۲۰۲۳).

هوش مصنوعی به‌عنوان یک حوزه مطالعاتی و فناوری، از دهه ۱۹۴۰ میلادی مطرح شده است و این فرآیند تاکنون توانسته در مراحل مختلفی را از جمله، پردازش زبان طبیعی و یادگیری ماشین، شبکه‌های عصبی و عصب‌شناسی محاسباتی، پردازش تصویر و تشخیص الگو را پیشرفت و... ارائه گردد (اخلاق پور، ۱۴۰۲). هوش مصنوعی به‌عنوان استفاده از ماشین‌های محاسباتی برای تقلید از قابلیت‌های ذاتی انسان، مانند انجام وظایف فیزیکی یا مکانیکی، تفکر و احساس تصور می‌شود (غلام‌پور و دهباشی، ۱۴۰۴). دیدگاه هوش مصنوعی چندگانه مبتنی بر این است که به‌جای اینکه هوش مصنوعی را یک ماشین متفکر بدانیم، می‌توان هوش مصنوعی را به‌گونه‌ای طراحی کرد که مانند انسان‌ها برای وظایف مختلف، هوش‌های چندگانه وجود داشته باشد. با توجه به زمینه‌هایی که هوش مصنوعی می‌تواند به آن‌ها رسیدگی کند، هوش مصنوعی مکانیکی، فکری و احساسی وجود دارد (هانگ^۴ و همکاران، ۲۰۲۱). باید در نظر داشت که هوش مصنوعی به‌عنوان ابزاری قدرتمند مشابه هیچ‌یک از نوآوری‌های سابق نمی‌باشد، زیرا به‌سرعت در حال تکامل بوده و هرروزه کاربردهای نوینی به آن اضافه می‌شود و می‌تواند علاوه بر تأثیرات مفید و یاری‌رسان خود، چالش‌هایی را نیز به وجود آورد که پیشگیری از وقوع و یا مدیریت این چالش‌ها یکی از مسئولیت‌های مهم در حیطه رعایت اخلاق در استفاده از آن است (گران‌هوت^۵ و همکاران، ۲۰۲۲).

اخلاق مفهومی پیچیده و درهم‌تنیده است. اخلاق را می‌توان به‌عنوان اصول اخلاقی حاکم بر رفتارها یا اعمال یک فرد یا یک گروه از افراد تعریف کرد (هانگ و همکاران، ۲۰۲۱). از جمله زمینه‌های مهم در حوزه اخلاق در دوره کنونی، اخلاق در هوش مصنوعی است. در زمینه هوش مصنوعی، اخلاقیات هوش مصنوعی، تعهدات اخلاقی و وظایف یک هوش مصنوعی و محققین آن را مشخص می‌کند. اخلاق در هوش مصنوعی بخشی از اخلاق فناوری پیشرفته است که بر روی ربات‌ها و سایرین به‌طور مصنوعی هوشمندانه تمرکز می‌کند. (سیائو و وانگ^۶، ۲۰۲۰). اخلاق در کاربست هوش مصنوعی مصادیق متعددی دارد و از جمله نمونه‌های کاربردی آن، استفاده از هوش مصنوعی در نوشتار و نگارش متون علمی است (کاکنا^۷ و همکاران، ۲۰۲۴).

اخلاق پژوهشی به مجموعه‌ای از اصول، ارزش‌ها و هنجارهایی گفته می‌شود که رفتار پژوهشگران را در تمام مراحل انجام تحقیق - از طراحی و گردآوری داده‌ها تا تحلیل، انتشار و کاربرد نتایج - هدایت می‌کند تا اصالت، صداقت، عدالت و احترام به حقوق انسان‌ها و جامعه علمی حفظ شود (کاکنا و همکاران، ۲۰۲۴). نوشتن مقالات علمی فرآیندی دقیق است که تحقیقات گسترده، استدلال سازمان‌یافته و بیان واضح را برای پیشبرد دانش ترکیب می‌کند (مالیک^۸ و همکاران، ۲۰۲۳). بسیاری از پژوهشگران در حال حاضر از هوش مصنوعی به‌عنوان دستیار تحقیق، برای سازمان‌دهی افکار، ارائه بازخورد، کمک به کد نویسی و خلاصه‌سازی

¹ Tai

² Jobin

³ Qureshi

⁴ Huang

⁵ Grunhut

⁶ Siau & Wang

⁷ Kacena

⁸ Malik

متون علمی استفاده می‌کنند. هوش مصنوعی و نگارش مقالات علمی چنان درهم تنیده‌اند که زمینه‌ساز یک همکاری تحول‌آفرین در پژوهش شده‌اند (التیما^۱ و همکاران، ۲۰۲۳). این ابزارهای قدرتمند مزایای قابل توجهی برای محققان در انجام جستجوها، بررسی مقالات، خلاصه‌سازی داده‌ها و نگارش متون فراهم می‌کنند. با این حال، نگرانی‌های جدی در مورد دقت محتوای تولیدشده توسط هوش مصنوعی و همچنین حفظ محرمانگی اطلاعات وجود دارد (اینام^۲ و همکاران، ۲۰۲۴). رعایت اخلاق در استفاده از هوش مصنوعی به‌ویژه هنگام انجام پژوهش‌های علمی اهمیتی دوچندان می‌یابد. پژوهشگران، به‌عنوان متولیان دانش، مسئول‌اند اطمینان حاصل کنند که فناوری‌های مبتنی بر هوش مصنوعی به‌گونه‌ای به کار گرفته شوند که همسو با اصول اخلاق پژوهشی باشد (اوسانیا^۳ و همکاران، ۲۰۲۴). علاوه بر این، پژوهشگران باید به دستورالعمل‌ها و چارچوب‌هایی اخلاقی پایبند باشند که استفاده مسئولانه از هوش مصنوعی را تضمین می‌کند. این چارچوب‌ها، به‌عنوان راهنمای عملی برای پژوهشگران، شامل اصولی همچون انصاف الگوریتمی، شفافیت فرایندها، حفظ حریم شخصی و پرهیز از سوگیری هستند (بجینو^۴، ۲۰۲۳). رعایت این چارچوب‌ها می‌تواند کمک کند نتایج پژوهش‌های مبتنی بر هوش مصنوعی قابل‌اعتمادتر باشد و منافع آن به‌صورت عادلانه میان همه گروه‌های ذینفع توزیع شود.

مرور پیشینه پژوهش نشان می‌دهد که تلاش‌های متعددی در سطح بین‌المللی برای تدوین چارچوب‌های اخلاقی در کاربست هوش مصنوعی صورت گرفته است. به‌عنوان نمونه جوبین و همکاران (۲۰۱۹) با استفاده از مرور نظام‌مند اسناد و دستورالعمل‌های اخلاقی منتشرشده در سطح جهانی، بیش از ۸۰ چارچوب اخلاقی مرتبط با هوش مصنوعی را استخراج و مقایسه کردند. یافته‌های آنان نشان داد که اصولی چون شفافیت، عدالت، مسئولیت‌پذیری و حفظ حریم خصوصی، بیشترین تکرار را در این اسناد دارند؛ هرچند هم‌پوشانی‌ها زیاد و معیارهای عملیاتی سازی هنوز مبهم باقی‌مانده است. در همین راستا، فلورییدی^۵ و همکاران (۲۰۱۸) در قالب پروژه هوش مصنوعی برای مردم و با رویکرد تحلیلی-اجماعی میان ذی‌نفعان مختلف، چهار اصل بنیادین «خیررسانی، پرهیز از آسیب، عدالت و خودمختاری» را برای ایجاد جامعه‌ای اخلاق‌مدار در عصر هوش مصنوعی پیشنهاد کردند. همچنین، سازمان آموزشی، علمی و فرهنگی سازمان ملل متحد^۶ (۲۰۲۱) با رویکردی سیاست‌گذاری کلان و تطبیقی، توصیه‌نامه‌ای جهانی در زمینه اخلاق هوش مصنوعی منتشر نمود که ضمن تأکید بر حقوق بشر، بر ضرورت ایجاد سازوکارهای نظارتی و بومی‌سازی در نظام‌های آموزشی و پژوهشی کشورها تأکید داشت.

در فضای پژوهشی داخلی نیز، گام‌هایی در این زمینه برداشته شده است. به‌عنوان مثال، پروانه و میرباقری (۱۴۰۲) در یک مطالعه تطبیقی اسنادی، استانداردهای اخلاقی ملی ایران را با دستورالعمل‌های بین‌المللی مقایسه کردند و نتیجه گرفتند که هرچند برخی اصول مشترک مانند حفظ حریم خصوصی و عدالت موردتوجه است، اما شکاف قابل‌توجهی در حوزه ضمانت اجرایی و آموزش اخلاق پژوهش در استفاده از هوش مصنوعی وجود دارد. همچنین، اخلاق‌پور (۱۴۰۲) در پژوهشی کاربردی-تحلیلی نقش «سیستم‌های توصیه‌گر مبتنی بر هوش مصنوعی» را در تحولات آموزشی بررسی کرد و هشدار داد که فقدان چارچوب اخلاقی روشن می‌تواند منجر به سوگیری‌های جدی در تصمیم‌گیری‌های آموزشی شود.

علاوه بر این، پژوهش‌های نوینی به‌طور خاص به استفاده از هوش مصنوعی در نگارش علمی پرداخته‌اند. برای نمونه، التیما و همکاران (۲۰۲۳) در یک مقاله مرور نقادانه به بررسی نقش هوش مصنوعی در فرآیند نگارش علمی پرداختند و نشان دادند که این فناوری می‌تواند به بهبود ساختار مقالات و تسهیل فرآیند ویرایش کمک کند، اما درعین حال خطر «وابستگی بیش‌ازحد

1 Altmäe

2 Inam

3 Awosanya

4 Benichou

5 Floridi

6 United Nations Educational, Scientific and Cultural Organization (UNESCO)

پژوهشگر به ابزار» و «محو مرز خلاقیت انسانی» را نیز به همراه دارد. به همین ترتیب، اوسانیا و همکاران (۲۰۲۴) با رویکرد مطالعه مروری-کاربردی در حوزه پزشکی نشان دادند که استفاده از ابزارهای مولد هوش مصنوعی در نگارش مرورهای علمی، سرعت و کارایی تولید مقاله را افزایش می‌دهد، اما هم‌زمان ابهامات جدی در زمینه دقت علمی، سرقت ادبی و رعایت حقوق مالکیت معنوی ایجاد می‌کند.

با وجود گسترش پرشتاب کاربرد هوش مصنوعی در آموزش و پژوهش، مرور نظام‌مند پیشینه نشان می‌دهد که اغلب مطالعات موجود، ماهیتی سیاست‌گذاری، فناورانه یا نظری دارند و بر تدوین اصول کلی و دستورالعمل‌های اخلاقی جهانی متمرکز بوده‌اند، نه بر ادراک و تجربه زیسته اعضای هیئت علمی به‌عنوان کنشگران اصلی عرصه پژوهش‌های آموزشی. در چارچوب مطالعات ایرانی نیز بیشتر تحقیقات، ناظر به تحلیل اسناد یا بررسی تطبیقی و هنجاری بوده و از کاوش عمیق در چگونگی مواجهه استادان با چالش‌های اخلاقی واقعی هنگام استفاده از ابزارهای هوش مصنوعی غفلت کرده‌اند. این خلأ دانشی موجب شده درک جامعی از ابعاد درونی، زمینه‌ای و فرهنگی استفاده مسئولانه از هوش مصنوعی در بافت آموزش عالی ایران شکل نگیرد و رابطه میان «اصول اخلاقی تدوین‌شده» و «رفتار و ادراک واقعی استادان در میدان عمل» همچنان نامشخص بماند؛ بنابراین، ضرورت دارد پژوهشی کیفی با رویکرد پدیدارشناسی، از درون روایت‌های زیسته استادان، معنا و منطق نهفته در کنش اخلاقی آنان را در مواجهه با فناوری‌های هوش مصنوعی آشکار کند. در همین راستا، پژوهش به دنبال پاسخ به پرسش اصلی زیر است: اساتید دانشگاهی چه ادراکی از چارچوب‌های اخلاقی لازم برای استفاده مسئولانه از هوش مصنوعی در پژوهش‌های آموزشی دارند؟

روش پژوهش

رویکرد پژوهش حاضر کیفی و از نظر روش مبتنی بر مطالعات پدیدارنگاری با رویکرد توصیفی است. مطالعه پدیدارنگاری معنای تجارب زیسته افراد متعدد از یک مفهوم یا پدیده را توصیف می‌کند. هدف اساسی تحقیق پدیدارنگاری این است که جزئیات کاملاً توصیفی را در اختیار محقق قرار دهد تا مشخص سازد که این افراد چه درکی از پدیده موردنظر دارند (کرسول، ۱۳۹۶). جامعه آماری این پژوهش کلیه اساتید دانشگاه و متخصصان در زمینه هوش مصنوعی و اخلاق پژوهش بودند. نمونه‌گیری به‌صورت هدفمند - ملاک محور (ملاک ورود حداقل مرتبه علمی استادیار، داشتن تألیف و پژوهش در زمینه هوش مصنوعی و اخلاق پژوهشی بود همچنین معیار خروج عدم تمایل به مشارکت و مصاحبه و نداشتن تجربه مستقیم از کاربرد هوش مصنوعی) و تا رسیدن به اشباع نظری داده‌ها ادامه یافت. نمونه پژوهش حاضر شامل ۱۶ نفر از اساتید دانشگاه در گروه‌های علوم پایه، علوم انسانی و فنی و مهندسی بود. پس از انجام و تحلیل این تعداد مصاحبه، مشخص شد که دیگر داده‌های جدیدی از مصاحبه‌های بعدی به دست نمی‌آید، به عبارتی در مصاحبه‌های ۱۷ و ۱۸ کد جدیدی در تحلیل بیانات مشارکت‌کنندگان یافت نشد؛ به همین دلیل تعداد مصاحبه‌ها برای این پژوهش کافی و به مرحله اشباع نظری رسید. لازم به ذکر است در این مطالعه تنها اساتید دارای مدرک دکتری از گروه‌های مختلف شرکت کردند، اما در پژوهش‌های آینده می‌توان دانشجویان تحصیلات تکمیلی و مدیران پژوهشی دانشگاه‌ها را نیز به‌عنوان مشارکت‌کننده اضافه کنید. به‌منظور گردآوری داده‌ها، از مصاحبه‌های نیمه ساختاریافته استفاده شد. پروتکل مصاحبه شامل محورهای اصلی و پرسش‌های باز بود که بر مبنای مرور ادبیات و اهداف پژوهش طراحی گردید. نمونه‌ای از پرسش‌های محوری عبارت بودند از:

- تجربه شخصی شما از کاربری اخلاقی هوش مصنوعی در فعالیت‌های پژوهشی چگونه بوده است؟
- چه چالش‌های اخلاقی را در فرایند جمع‌آوری، تحلیل و انتشار داده‌های پژوهشی مبتنی بر هوش مصنوعی مشاهده کرده‌اید؟
- چه اصول یا چارچوب‌هایی می‌تواند به استفاده مسئولانه از هوش مصنوعی در پژوهش کمک کند؟
- به نظر شما دانشگاه‌ها و نهادهای علمی چه اقداماتی باید در زمینه اخلاق هوش مصنوعی در فرایند پژوهش انجام دهند؟
- به نظر شما مهم‌ترین ملاحظات اخلاقی در استفاده از هوش مصنوعی در پژوهش‌های دانشگاهی چیست؟

- در لحظه تردید یا تصمیم‌گیری اخلاقی چگونه عمل می‌کنید؟

هماهنگی اولیه از طریق ارسال دعوت‌نامه پیامکی/ایمیلی و سپس تماس تلفنی صورت گرفت. پس از اعلام رضایت آگاهانه، زمان و شیوه مصاحبه تعیین شد. هر مصاحبه به‌طور متوسط ۴۵ تا ۶۰ دقیقه به طول انجامید و در مکان‌های مورد توافق طرفین (دفتر کار استاد یا بستر آنلاین مانند اسکای روم/ادوبی کانکت یا تماس تلفنی) برگزار شد. تمامی مصاحبه‌ها با اجازه مشارکت‌کنندگان ضبط صوتی شد و بلافاصله پس از پایان، به‌صورت کامل پیاده‌سازی گردید. برای رعایت اصول اخلاقی، به هر یک از مشارکت‌کنندگان کدی اختصاص داده شد (م ۱ تا م ۱۶) و از ذکر نام و مشخصات فردی پرهیز گردید. همچنین، امکان انصراف از مصاحبه در هر مرحله برای مشارکت‌کنندگان فراهم بود. در ادامه ویژگی‌های مشارکت‌کنندگان در جدول ۱ بیان شده است.

جدول ۱. اطلاعات مشارکت‌کنندگان در پژوهش

ردیف	مدرک تحصیلی	رشته تحصیلی	تجربه استفاده از هوش مصنوعی	کد مصاحبه
۱	دکتری	تکنولوژی آموزشی	بله	م ۱
۲	دکتری	تحقیقات آموزشی	بله	م ۲
۳	دکتری	تکنولوژی آموزشی	بله	م ۳
۴	دکتری	عضو هیئت مدیره انجمن اخلاق در علوم و فناوری ایران و پژوهشگر اخلاق	بله	م ۴
۵	دکتری	عضو انجمن اخلاق زیستی	بله	م ۵
۶	دکتری	تکنولوژی آموزشی	بله	م ۶
۷	دکتری	روانشناسی	بله	م ۷
۸	دکتری	مدیریت آموزشی	بله	م ۸
۹	دکتری	ریاضی کاربردی	بله	م ۹
۱۰	دکتری	برنامه‌ریزی درسی	بله	م ۱۰
۱۱	دکتری	نرم‌افزار گرایش شبکه	بله	م ۱۱
۱۲	دکتری	شیمی	بله	م ۱۲
۱۳	دکتری	آمار	بله	م ۱۳
۱۴	دکتری	مدیریت آموزشی	بله	م ۱۴
۱۵	دکتری	فیزیک	بله	م ۱۵
۱۶	دکتری	مهندسی کامپیوتر	بله	م ۱۶

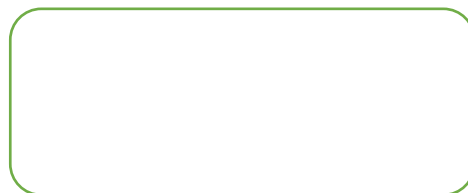
در پژوهش حاضر از مدل تحلیلی پدیدارنگاری اپوخه^۱ استفاده شد. پژوهشگر پیش از آغاز فرآیند تحلیل، پیش‌فرض‌ها و باورهای شخصی خود درباره‌ی اخلاق پژوهش و هوش مصنوعی را در یک دفتر یادداشت پژوهشی ثبت کرده تا از تأثیر آن‌ها بر تفسیر داده‌ها جلوگیری شود. فرآیند تحلیل داده‌ها در این پژوهش بر اساس الگوی شش‌مرحله‌ای تحلیل تم بازتابی براون و کلارک انجام شد. در گام نخست، پژوهشگر با غوطه‌وری در داده‌ها از طریق خواندن مکرر مصاحبه‌ها، به درک کلی از محتوای تجارب زیسته اساتید دست یافت. در گام دوم، اپوخه (برکتینگ) در قالب «واکاوی بازتابی» اجرا شد؛ بدین معنا که پژوهشگر هم‌زمان با خواندن داده‌ها، پیش‌فهم‌ها و سوگیری‌های شخصی خود را به‌صورت آگاهانه تعلیق کرد تا معانی از دل داده‌ها و روایت مشارکت‌کنندگان ظاهر شود. در گام سوم، کدهای اولیه^۲ بر اساس عبارات معنادار مرتبط با ادراک اخلاقی اساتید استخراج گردید. در گام چهارم، کدها در قالب زیرتم‌ها^۳ طبقه‌بندی و ارتباط‌های مفهومی میان آن‌ها بررسی شد. در گام پنجم، تم‌های

^۱ Epoche

^۲ initial codes

^۳ sub-themes

نهایی یا فراگیر^۱ با تأکید بر معانی بازتابی و زمینه‌ای شکل گرفت. نهایتاً در گام ششم، سنتز و بازنمایی نتایج انجام شد تا ابعاد مختلف «ادراک اساتید از استفاده مسئولانه از هوش مصنوعی در پژوهش دانشگاهی» آشکار شود. در کل فرایند، اپوخه به مثابه تمرین بازتابی مستمر رعایت شد تا فاصله‌ی تحلیلی میان پیش‌فهم پژوهشگر و تجارب واقعی مشارکت‌کنندگان حفظ شود. در ادامه نمونه‌ای عینی از فرایند کدگذاری ارائه شده است.



شکل ۱: فرایند کدگذاری داده‌ها

در زمینه روایی و پایایی از روش گوبا و لینکلن^۲ (۱۹۸۲) استفاده شد. آن‌ها چهار معیار «قابلیت اعتبار^۳، قابلیت انتقال^۴، قابلیت اتکا^۵ و قابلیت تأیید^۶» را به منظور ارزیابی دقت علمی پژوهش برشمردند. در زمینه اعتبار، از روش کنترل اعضا بهره گرفته شد، به این صورت که پس از تحلیل داده‌ها، خلاصه یافته‌ها برای تعدادی از مشارکت‌کنندگان ارسال شد و بازخوردهای آنان در فرایند تحلیل و نتایج پژوهش اعمال گردید. برای انتقال‌پذیری، توصیف غنی زمینه پژوهش، ویژگی‌های مشارکت‌کنندگان و نقل‌قول‌های مستقیم ارائه گردید تا امکان قضاوت درباره کاربرد نتایج در موقعیت‌های مشابه فراهم شود. در بعد اتکاپذیری، بخشی از مصاحبه‌ها توسط دو پژوهشگر مستقل کدگذاری شد و توافق میان کدگذاران بررسی گردید، ضمن اینکه تمامی مراحل پژوهش مستندسازی شد. نهایتاً، جهت تأییدپذیری، از روش اپوخه و بازاندیشی پژوهشگر استفاده شد و تحلیل‌ها توسط هم‌تایان علمی بازبینی گردید به این صورت که دو همکار دانشگاهی که در پژوهش مشارکت نداشتند، بخشی از متن مصاحبه‌ها و کدهای اولیه را به صورت مستقل مرور کردند تا از هم‌خوانی مفاهیم اطمینان حاصل شود. این اقدامات در مجموع موجب تقویت روایی و پایایی پژوهش شد. تحلیل داده‌ها بر اساس به صورت دستی در نرم افزار ورد^۷ با رنگ‌بندی مجزا برای هر واحد معنایی انجام شد.

یافته‌ها

الف: معیارهای اخلاقی در بُعد فرایند پژوهش

تحلیل دیدگاه مشارکت‌کنندگان در پژوهش نشان‌دهنده این مهم بود که یکی از ابعاد اخلاقی در زمینه استفاده از هوش مصنوعی در پژوهش‌های آموزشی، استفاده اخلاقی از هوش مصنوعی در فرایند پژوهش است در این زمینه سه زیر تم اصلی شامل شفافیت در مدیریت داده‌ها (مرحله جمع‌آوری داده)، حفاظت حقوق فردی مشارکت‌کنندگان (مرحله تحلیل داده) و شفافیت و محافظت اخلاقی در انتشار نتایج پژوهش (مرحله گزارش نهایی) در ادامه در جدول ۲ مضامین این مهم ارائه شده است.

¹ overarching themes

² Guba & Lincoln

³ Credibility

⁴ Transferability

⁵ Dependability

⁶ Conformability

⁷ Word

جدول ۲: معیارهای اخلاقی در بُعد فرایند پژوهش در کاربرد هوش مصنوعی در پژوهش‌های آموزشی

تم اصلی	زیر تم	نمونه بیانات مشارکت‌کنندگان (کدهای اولیه)
شفافیت در مدیریت داده‌ها (مرحله جمع‌آوری داده)	رضایت آگاهانه	محقق باید مشارکت‌کنندگان را از نحوه استفاده از اطلاعاتشان آگاه کند (مشارکت‌کننده، ۱۵).
		محقق باید چگونگی حفاظت از حریم خصوصی را در اختیار مشارکت‌کنندگان در پژوهش قرار دهد (مشارکت‌کننده، ۸).
		مشارکت‌کنندگان باید از حق انصراف از مشارکت در پژوهش برخوردار باشند (مشارکت‌کنندگان، ۱۴، ۳).
	وضوح رویه‌ها	مشارکت‌کنندگان باید از اهداف و ضروریات پژوهش مبتنی بر هوش مصنوعی آگاه باشند (مشارکت‌کننده، ۹).
		در پژوهش‌های مبتنی بر هوش مصنوعی باید از منابع مختلف جهت جمع‌آوری داده استفاده شود (مشارکت‌کننده، ۶).
		روش‌های جمع‌آوری داده‌ها، نحوه استفاده از داده‌ها و چگونگی حفاظت از داده‌ها در تحلیل‌های مبتنی بر هوش مصنوعی باید به‌صورت شفاف بیان شود (مشارکت‌کننده، ۱۱، ۱۳، ۱۵).
		شفافیت محدوده استفاده از داده‌ها برای مشارکت‌کنندگان در تحلیل‌های مبتنی بر مکانیزم‌های هوش مصنوعی (مشارکت‌کننده، ۴).
حفاظت حقوق فردی مشارکت‌کنندگان (مرحله تحلیل داده)	رمزگذاری اخلاقی داده‌ها	داده‌های پژوهش‌های مبتنی بر هوش مصنوعی باید به‌صورت اخلاقی و تک لایه تحلیل شوند و از آن سوءاستفاده نشود (مشارکت‌کنندگان، ۲ و ۱۲).
		مشارکت‌کنندگان باید از رویکرد رمزگذاری داده‌ها مطلع بوده و بر مبنای آن مشارکت خود را سازمان دهند (مشارکت‌کننده، ۷).
	کنترل دسترسی	در تحلیل داده‌های پژوهشی مبتنی بر مکانیزم‌ها و الگوریتم‌های هوش مصنوعی مشارکت‌کنندگان باید دسترسی به تحلیل‌ها داشته باشند (مشارکت‌کنندگان، ۱ و ۵).
		نظارت مستمر توسط اساتید باید بر فرآیند روش‌شناسی مرحله‌ای هوش مصنوعی در تحلیل وجود داشته باشد (مشارکت‌کننده، ۱۰).
	رفع تعصبات	تحلیل‌های مبتنی بر هوش مصنوعی چون بر مبنای تعصبات مقالات قبلی و برخی موارد نژادی است باید این مهم در نظر گرفته شود (مشارکت‌کنندگان، ۵ و ۱۶).
		باید تعصبات موجود در الگوریتم‌های تحلیل داده کیفی موجود در هوش مصنوعی در نظر گرفته شود (مشارکت‌کننده، ۳ و ۱۵).
		الگوریتم‌های هوش مصنوعی تعصبات موجود در داده‌های آموزشی خود را تقویت می‌کنند که باید تعدیل شوند (مشارکت‌کنندگان، ۴ و ۱۶).
نتایج پژوهش و محافظت اخلاقی در انتشار (مرحله گزارش نهایی)	حفاظت در برابر حملات سایبری	با توجه به مکانیزم هوش مصنوعی باید داده‌های مشارکت‌کنندگان در برابر حملات سایبری محافظت شود (مشارکت‌کنندگان، ۶، ۱۰).
		بهتره که به‌روزرسانی مداوم سرورهای هوش مصنوعی جهت جلوگیری از حملات سایبری بیگانه انجام بگیرد (مشارکت‌کننده، ۲).
	شفاف‌سازی الگوریتم‌ها و فرایند پژوهش	باید الگوریتم‌های تحلیلی هوش مصنوعی در پژوهش‌ها و انتشار آن دقیق مشخص شود (مشارکت‌کننده، ۸).
		متصل شدن سرور مجلات پژوهشی به سامانه هوش مصنوعی جهت پایش و بررسی مقالات پژوهش (مشارکت‌کنندگان، ۱ و ۱۲).
		باید در زمان انتشار یافته‌ها از رویکردهای پیشرفته ناشناس‌سازی استفاده شود (مشارکت‌کنندگان، ۵ و ۱۳).

باید محققان به این نکته توجه داشته باشند که هوش مصنوعی می‌تواند با ترکیب داده‌ها هویت افراد را فاش کند (مشارکت‌کننده، ۴ و ۹).	ناشناس‌سازی و محرمانه‌سازی داده‌ها
--	---

شفافیت در مدیریت داده: یکی از مراحل مهم در انجام پژوهش، مدیریت فرآیند جمع‌آوری داده‌ها است که باید موردتوجه قرار بگیرد. در مدیریت داده‌های در پژوهش، شفافیت یک اصل اخلاقی بنیادین است. این امر مستلزم آشکارسازی کامل روش‌های جمع‌آوری، پردازش و تحلیل داده‌ها، و همچنین الگوریتم‌های مورد استفاده است تا امکان بررسی و راستی آزمایی مستقل فراهم شود. تحلیل دیدگاه مشارکت‌کنندگان نشان‌دهنده این مهم بود که باید در پژوهش در این زمینه نشان‌دهنده دو محور اصلی رضایت آگاهانه در زمینه مشارکت در پژوهش و وضوح رویه‌ها در پژوهش در زمینه استفاده از هوش مصنوعی در پژوهش بود. در ادامه به نمونه‌هایی از بیانات مشارکت‌کنندگان در این زمینه اشاره می‌شود: "به‌طور کل در این دست از پژوهش‌ها محقق باید شفافیت لازم در زمینه محدود استفاده از داده‌ها برای مشارکت‌کنندگان در تحلیل‌های مبتنی بر مکانیزم‌های هوش مصنوعی را داشته باشد و اطمینان لازم را به مشارکت‌کنندگان بدهد" (م ۱۲) یا مشارکت‌کننده ۳ که بیان داشت «گر مشارکت‌کننده نماند داده‌هایش دقیقاً برای چه استفاده می‌شود، احساس ناامنی می‌کند» همچنین مشارکت‌کننده ۹ اشاره داشت که "در پژوهش‌های مبتنی بر هوش مصنوعی باید از همان ابتدا همه چیز روشن و شفاف باشد."

حفاظت از حقوق فردی در مرحله تحلیل داده: تحلیل دیدگاه مشارکت‌کنندگان در پژوهش نشان‌دهنده این مهم بود که یکی از محورهای اصلی اخلاقیات در کاربرد هوش مصنوعی در فرآیند پژوهش حفاظت از حقوق فردی مشارکت‌کنندگان در تحلیل داده‌ها است. در تحلیل داده‌ها توسط هوش مصنوعی ناشناس‌سازی دقیق داده‌ها و استفاده محدود از اطلاعات شخصی برای جلوگیری از نقض حریم خصوصی و تبعیض است باید موردنظر قرار گیرد. در این زمینه به‌عنوان مثال مشارکت‌کننده ۴ اشاره داشت: "مهم است بدانیم چه کسی به داده‌ها دسترسی دارد و چطور رمزگذاری می‌شوند" به‌علاوه در این زمینه محورهای چون رمزگذاری اخلاقی داده‌ها، کنترل دسترسی و رفع تعصبات موردتوجه بود. در ادامه به نمونه‌هایی از بیانات مشارکت‌کنندگان اشاره می‌شود: "باید پژوهشگران بدانند که تحلیل‌های مبتنی بر هوش مصنوعی چون بر مبنای تعصبات مقالات قبلی و برخی موارد نژادی است باید این مهم در نظر گرفته شود و در تحلیل‌ها این موارد سعی شود با کاربردی الگوریتم‌های متفاوت و نظارت انسانی تعدیل شود"

شفافیت و محافظت اخلاقی در انتشار نتایج پژوهش: در انتشار نتایج پژوهش‌های مبتنی بر تحلیل داده توسط هوش مصنوعی، شفافیت و محافظت اخلاقی از اهمیت بالایی برخوردارند. از دیگر ابعاد اخلاقی که در زمینه کاربرد هوش مصنوعی در فرآیند تحلیل داده‌ها شفافیت و محافظت اخلاقی در مرحله انتشار داده‌ها است در این زمینه محورهایی چون حفاظت در برابر حملات سایبری، شفافیت سازی الگوریتم‌ها و فرآیند پژوهش و ناشناس سازی و محرمانه سازی داده‌ها موردتوجه بود. مشارکت‌کننده (۷) بیان داشت: "باید محققان به این نکته در انتشار یافته‌ها توجه داشته باشند که هوش مصنوعی می‌تواند با ترکیب داده‌ها هویت افراد را فاش کند. از این رو باید از تکنیک‌های ناشناس سازی و محرمانه سازی داده‌ها استفاده کنند" یا مشارکت‌کننده (۳) بیان داشت: "باید الگوریتم‌ها توضیح‌پذیر باشند تا کسی درک کند بر چه اساسی تصمیم گرفته‌اند"

ب: معیارهای اخلاقی بُعد تعهد

تحلیل داده‌های پژوهش نشان‌دهنده این مهم بود که یکی از ابعاد مهم معیارهای اخلاقی بُعد تعهد در به‌کارگیری هوش مصنوعی در پژوهش‌های آکادمیک است. در این زمینه تم اصلی پاسخ‌گویی اخلاقی و عدالت‌محور پژوهشگر در کاربردهای هوش مصنوعی مورد شناسایی قرار گرفت که دارای دو زیر تم تعهد حرفه‌ای محقق/محققین در به‌کارگیری هوش مصنوعی و تضمین عدالت و انسان‌محوری در تحلیل‌های هوش‌مبنا مورد شناسایی قرار گرفت. یافته‌های این مضمون در جدول ۳ ارائه شده است.

جدول ۳: معیارهای اخلاقی در بُعد تعهد نسبت به کاربرد هوش مصنوعی در پژوهش‌های آموزشی

تم اصلی	زیر تم	نمونه بیانات مشارکت‌کنندگان (کدهای اولیه)
پاسخ‌گویی اخلاقی و عدالت‌محور پژوهشگر در کار بست هوش مصنوعی	تعهد حرفه‌ای محقق/محققین در به‌کارگیری هوش مصنوعی	محقق مسئولیت هرگونه خطا در الگوریتم‌های هوش مصنوعی را بپذیرد (مشارکت‌کننده، ۱۶).
		محقق/محققین باید راهکارهایی برای گزارش خطا و رفع آن فراهم کنند (مشارکت‌کننده، ۲).
		محقق باید نحوه استفاده از هوش مصنوعی در پژوهش خود را به‌صورت کامل بیان کند (مشارکت‌کنندگان، ۵ و ۱۱).
		در پژوهش باید مسئول هرکدام از بخش‌های پژوهش در مرحله جمع‌آوری، تحلیل و گزارش نهایی داده‌ها مشخص شود (مشارکت‌کننده، ۹، ۱۱).
	تضمین عدالت و انسان‌محوری در تحلیل‌های هوش‌مبنا	در پژوهش‌ها باید از الگوریتم‌های هوش مصنوعی که ارائه نتایج را عادلانه‌تر و بدور از تعصب می‌کنند استفاده شود (مشارکت‌کننده، ۳).
		باید از الگوریتم‌های هوش مصنوعی توضیح‌پذیر در پژوهش‌های دانشگاهی استفاده کرد (مشارکت‌کننده، ۱۴).
		باید زمینه‌ای فراهم بشه تا خوانندگان در جریان الگوریتم‌های استفاده‌شده قرار بگیرند (مشارکت‌کنندگان، ۳ و ۱۰).
		در پژوهش باید زمینه‌های انسانی به‌عنوان مبنا در نظر گرفته شود تا عدالت در تحلیل رعایت شود (مشارکت‌کننده، ۷ و ۱۳).
		پژوهش‌هایی که از مدل‌های هوش مصنوعی پیچیده و داده‌های اختصاصی استفاده می‌کنند، باید به فکر قابلیت بازتولید و اعتبارسنجی نتایج توسط دیگران در آینده نیز باشند (مشارکت‌کننده، ۶ و ۱۵).
		تحلیل و دیدگاه تحلیلی باید در اختیار ذی‌نفعان قرار بگیرد تا قضاوت‌های اخلاقی صورت پذیرد (مشارکت‌کننده، ۴).

تعهد حرفه‌ای محقق/محققین در به‌کارگیری هوش مصنوعی: از دیگر محورهایی که در استفاده اخلاقی هوش مصنوعی در پژوهش‌های آموزشی موردتوجه مشارکت‌کنندگان بود، تعهد حرفه‌ای محقق در به‌کارگیری هوش مصنوعی در پژوهش بود. استفاده از هوش مصنوعی در نگارش و پژوهش برای محقق این مسئولیت را ایجاد می‌کند، که به موازین اخلاقی و مسئولیت‌پذیری خود متعهد باشد. در این زمینه مشارکت‌کننده (۱۱) بیان داشت: "در پژوهش باید مسئول هرکدام از بخش‌های پژوهش در مرحله جمع‌آوری، تحلیل و گزارش نهایی داده‌ها مشخص شود و مسئولیت تمام موارد و تحلیل‌های مبتنی بر هوش مصنوعی با محقق است". مشارکت‌کننده (۹) بیان داشت: "برای اینکه بتوانیم بحث اخلاق رو در پژوهش‌ها دقیق رعایت کنیم باید مسئول هرکدام از بخش‌های پژوهش در مرحله جمع‌آوری، تحلیل و گزارش نهایی داده‌ها مشخص باشه". همچنین مشارکت‌کننده (۶) بیان داشت: "به‌نظر من وقتی یه پژوهش از مدل‌های خیلی پیچیده هوش مصنوعی و داده‌های اختصاصی استفاده می‌کنه، مسئولیت پژوهشگر بیشتر می‌شه. چون اگه کسی دیگه نتونه اون نتایج رو بازتولید کنه یا بررسی کنه که واقعاً درست بوده یا نه، عملاً اعتبار اون کار زیر سؤال می‌ره. ما باید از حالا فکر کنیم که چطوری می‌شه داده‌هامون و روش‌هامون طوری طراحی بشه که در آینده هم بشه نتایج رو بازبینی و اعتبارسنجی کرد".

تضمین عدالت و انسان‌محوری در تحلیل‌های هوش‌مبنا: در تحلیل داده‌ها مبتنی بر هوش مصنوعی باید سعی شود سیستم‌ها و الگوریتم‌هایی به کارگرفته شود که منافع تمام گروه‌های اجتماعی در نظر بگیرد. در این زمینه مشارکت‌کننده (۳) بیان داشت: "در پژوهش‌ها باید از الگوریتم‌های هوش مصنوعی که ارائه نتایج را عادلانه‌تر و بدور از تعصب می‌کنند استفاده شود این مهم

در پژوهش‌های اجتماعی نمود خیلی زیادی دارد، خود من چند الگوریتم بر مبنای شبکه عصبی عمیق^۱ و پردازش زبان طبیعی^۲ رو بررسی کردم خیلی در زمینه تعصبات اجتماعی حساس هست و مبتنی بر اون است". همچنین مشارکت‌کننده (۱۴) بیان داشت: "به‌نظر من وقتی قراره تو پژوهش‌های آموزشی از هوش مصنوعی استفاده کنیم، بهتره سراغ مدل‌هایی بریم که بشه منطبق و تصمیم‌هاش رو توضیح داد. اینکه یه الگوریتم فقط جواب بده ولی نگه چجوری به اون نتیجه رسیده، از نظر علمی قابل قبول نیست. ما باید بدونیم پشت اون خروجی چی بوده و بتونیم به بقیه هم توضیحش بدیم."

معیارهای اخلاقی ساختاری-زمینه‌ای مرتبط با هوش مصنوعی در پژوهش

تحلیل دیدگاه مشارکت‌کنندگان در پژوهش نشان‌دهنده این مهم بود که برخی از عوامل در زمینه^۳ کاربرد اخلاقی هوش مصنوعی در پژوهش قابل سازمان‌دهی در بُعد ساختاری - زمینه‌ای است در این زمینه محورهایی چون تدوین مقررات اخلاق پژوهی در کاربرد هوش مصنوعی در پژوهش، نهادینه‌سازی سواد اخلاقی و ترویج هنجارهای اخلاق پژوهی در استفاده از هوش مصنوعی و هم‌راستاسازی استانداردهای اخلاقی در سطوح سازمانی، ملی و بین‌المللی در زمینه^۴ کاربرد پژوهشی هوش مصنوعی بود که در ادامه به آن پرداخته می‌شود.

جدول ۴: معیارهای اخلاقی در بُعد ساختاری-زمینه‌ای در کاربرد هوش مصنوعی در پژوهش‌های آموزشی

تم اصلی	زیر تم	نمونه بیانات مشارکت‌کنندگان(کدهای اولیه)
نهادینه‌سازی نظام‌مند اخلاق پژوهی در کاربست هوش مصنوعی در پژوهش	تدوین و اجرای مقررات اخلاق پژوهی در کاربرد هوش مصنوعی	باید در مؤسسات آموزشی-پژوهشی رویه‌های برخورد با تخلفات اخلاقی مرتبط با استفاده غیراخلاقی و غیراصولی از هوش مصنوعی مشخص شود(مشارکت‌کننده، ۷).
		باید در مؤسسات آموزشی - پژوهشی زمینه‌های مجاز و غیرمجاز استفاده از هوش مصنوعی مشخص بشه و براش برنامه‌ریزی صورت بگیره(مشارکت‌کننده، ۱۲)
		باید چارچوب قانونی در زمینه استفاده از هوش مصنوعی در مؤسسات آموزشی ایجاد بشه (مشارکت‌کنندگان، ۳ و ۱۵).
		زمینه‌های اخلاقی استفاده از هوش مصنوعی در پژوهش باید مبتنی بر معیارهای درست مشخص شود و وجهه قانونی بهش داده بشه (مشارکت‌کننده، ۱۳).
	نهادینه‌سازی سواد اخلاقی و ترویج هنجارهای اخلاق پژوهی در استفاده از هوش مصنوعی	باید هوش مصنوعی رو به‌عنوان یه ابزار مفید پژوهشی بپذیریم و بتوانیم درست از امکاناتش در محیط‌های دانشگاهی استفاده کنیم(مشارکت‌کننده، ۹).
		باید به افزایش آگاهی عمومی از ملاحظات اخلاقی کاربرد هوش مصنوعی در پژوهش استفاده کنیم(مشارکت‌کننده، ۳).
		آموزش محققان در زمینه ملاحظات اخلاقی کاربرد هوش مصنوعی در پژوهش موردنیاز است(مشارکت‌کننده، ۱۶).
		توجه به فرهنگ نوآوری در پژوهش‌های نوین مبتنی بر هوش مصنوعی در آموزش و برنامه‌های آموزشی دانشگاهی لازم است(مشارکت‌کننده، ۴).
	هم‌راستاسازی استانداردهای اخلاقی در سطوح سازمانی، ملی و بین‌المللی در زمینه کاربرد پژوهشی هوش مصنوعی	تلاش برای هماهنگ‌سازی سیاست‌های اخلاقی در زمینه هوش مصنوعی در سطح ملی و بین‌المللی باید از سمت نهادهای مجری صورت بگیره(مشارکت‌کنندگان، ۹ و ۱۱).
		باید معیارهای اخلاقی جهانی در زمینه کاربرد هوش مصنوعی ایجاد بشه و اون‌ها رو همه کشورها قبول کنند تا از سردرگمی و ناسازگاری استانداردهای آموزشی جلوگیری شود(مشارکت‌کننده، ۲).
		ایجاد همسویی قوانین اخلاقی استانداردهای کاربرد هوش مصنوعی در سطح بین‌المللی به ترویج پژوهش‌های بین‌المللی در این زمینه کمک می‌کند(مشارکت‌کنندگان، ۱۰ و ۱۲).

¹ Deep Neural Networks

² Natural Language Processing - NLP

تدوین و اجرای آیین‌نامه‌های اخلاق پژوهی در کاربرد هوش مصنوعی: از محورهای اصلی شناسایی شده در بُعد ساختاری-زمینه‌ای تدوین و اجرای آیین‌نامه‌های اخلاق پژوهی در کاربرد هوش مصنوعی در پژوهش‌ها است. مشارکت‌کنندگان بیان داشتند که تدوین آیین‌نامه اخلاق پژوهی در کاربرد هوش مصنوعی در پژوهش‌های دانشگاهی منجر به ارزیابی دقیق و مستمر ریسک‌های بالقوه، به‌ویژه در زمینه‌های حساس مانند سلامت و عدالت اجتماعی، می‌شود. در این زمینه مشارکت‌کننده (۱۳) بیان داشت "زمینه‌های اخلاقی/استفاده از هوش مصنوعی در پژوهش باید مبتنی بر معیارهای درست مشخص شود و وجهه قانونی بهش داده بشه تا زمینه‌های بروز تخلف گرفته شود و از سردرگمی جلوگیری بشه". همچنین مشارکت‌کننده (۱۵) بیان داشت: "ببینید، مسئله‌ی هوش مصنوعی و اخلاقیش دیگه صرفاً یه بحث آکادمیک یا فلسفی نیست. وقتی این ابزارها دارن میان داخل فرآیندهای آموزشی و پژوهشی ما، باید یه چارچوب قانونی و دستورالعمل اجرایی برای مؤسسات آموزشی داشته باشیم. نمی‌شه همین‌طور ره‌اش کرد. ما نیاز داریم یه مقررات داخلی تعریف بشه که مشخص کنه چطور باید از داده‌های دانشجوها استفاده بشه، یا خط قرمزها در مورد مالکیت فکری و تقلب چی هستن. الان همه‌چیز تو فضایی از ابهام و سلیقه‌ای پیش می‌ره که به صلاح دانشگاه نیست."

نهادینه‌سازی سواد اخلاقی و ترویج هنجارهای اخلاق پژوهی در استفاده از هوش مصنوعی: از دیگر محورهایی که مورد توجه مشارکت‌کنندگان در پژوهش بود نهادینه‌سازی سواد اخلاقی در استفاده از هوش مصنوعی در پژوهش بود. به اعتقاد محققین نمی‌توان در زمینه استفاده از هوش مصنوعی در فضای آکادمیک منفعل بود و از طرفی در برابر آن جبهه مخالف گرفت، چرا که هوش مصنوعی زمینه تحول در کلیه امور را فراهم آورده و خود به‌عنوان ابزاری برای پیشرفت است. با افزایش سواد استفاده از این ابزار می‌توان ضمن جلوگیری از نااخلاقی زمینه تحول مثبت در پژوهش را فراهم آورد. نهادینه‌سازی سواد اخلاقی در کاربرد هوش مصنوعی در پژوهش مستلزم تدوین چارچوب‌هایی است که پژوهشگران را به شناسایی، تحلیل و پیش‌بینی پیامدهای اخلاقی فناوری‌های هوش مصنوعی ملزم کند. مطالعات اخیر نشان می‌دهند که ادغام آموزش سواد اخلاقی در طراحی و اجرای پروژه‌های تحقیقاتی هوش مصنوعی، موجب افزایش مسئولیت‌پذیری و بهبود نتایج پژوهشی می‌شود. در این راستا، مشارکت‌کننده (۴) بیان داشت: "اساتید و دانشجویان باید آموزش ببینند تا بدانند استفاده درست از هوش مصنوعی چیست؟" همچنین مشارکت‌کننده (۱۶) بیان داشت: "واقعیت اینه که خیلی از پژوهشگرها هنوز آموزش ندیدن که بدونن استفاده از هوش مصنوعی تو کارهای پژوهشی چه ملاحظات اخلاقی‌ای داره. مثلاً ممکنه یه نفر بدون نیت بد از یه ابزار هوش مصنوعی استفاده کنه، ولی ندونه داده‌هایی که وارد می‌کنه یا خروجی‌هایی که می‌گیره ممکنه از نظر حریم خصوصی یا دقت علمی مشکل‌زا باشه. لازمه قبل از اینکه استفاده از این فناوری‌ها تو پژوهش عادی بشه، برای اساتید و دانشجوها دوره‌ها و آموزش‌های مشخصی در این زمینه برگزار بشه"

هم‌راستاسازی استانداردهای اخلاقی در سطوح سازمانی، ملی و بین‌المللی: یکی دیگر از محورهایی که مد نظر مشارکت‌کنندگان در پژوهش بود هم‌راستاسازی استانداردهای اخلاقی در سطوح سازمانی، ملی و بین‌المللی است تا در این زمینه وحدت رویه جهانی ایجاد شود و پژوهشگران به‌صورت عادلانه از این امکان استفاده کنند. در این زمینه همان‌طور که مشارکت‌کنندگان بیان داشتند، باید توجه داشت که هماهنگی در تدوین و اجرای اصول اخلاقی می‌تواند خطرات را کاهش داده و اعتماد عمومی نسبت به پژوهش‌های مبتنی بر هوش مصنوعی را افزایش دهد. مشارکت‌کننده (۲) بیان داشت: "باید معیارهای اخلاقی جهانی در زمینه کاربرد هوش مصنوعی ایجاد بشه و اون‌ها رو همه کشورها قبول کنند تا از سردرگمی و ناسازگاری استانداردهای آموزشی جلوگیری شود و همه بتوانند به‌صورت عادلانه از این ابزار استفاده کنند و مجلات هم رویه غالب و مشخصی رو دنبال کنند اما متأسفانه در این زمینه‌ها کشور ما همیشه تعلل می‌کنه و چند سال از روندها جهانی عقب می‌افته."

حقوق مالکیت در کاربرد هوش مصنوعی در پژوهش

تحلیل دیدگاه مشارکت‌کنندگان در پژوهش نشان‌دهنده این مهم بود که یکی از ابعاد مهم اخلاقی در کاربرد هوش مصنوعی در پژوهش‌های آموزشی حفظ حقوق مالکیت در دسترسی به نتایج پژوهش‌های مبتنی بر هوش مصنوعی است. در این زمینه سه زیر تم شامل احترام به مالکیت معنوی پژوهش‌ها، دسترسی‌پذیری هوش مصنوعی و رعایت منافع مشروع حامیان علمی مورد شناسایی قرار گرفت. یافته‌های این بُعد در جدول ۴ ارائه شده است.

جدول ۵: معیارهای اخلاقی در بُعد حقوق مالکیت در کاربرد هوش مصنوعی در پژوهش‌های آموزشی

تم اصلی	زیر تم	نمونه بیانات مشارکت‌کنندگان (کدهای اولیه)
حفظ حقوق مالکیت در دسترسی به نتایج پژوهش‌های مبتنی بر هوش مصنوعی	احترام به مالکیت معنوی پژوهش‌ها	در پژوهش باید به صورت واضح و صریح مالک داده‌ها، الگوریتم‌ها و نتایج هوش مصنوعی را بیان کنیم (مشارکت‌کننده، ۱).
		باید قرار داده‌های حقوق مالکیت معنوی پژوهش را بر مبنای موازین اخلاقی تنظیم کنیم تا مشکلی در این زمینه پیش نیاید (مشارکت‌کنندگان، ۳ و ۱۴).
		باید به این نکته توجه داشته باشیم که انسان مالک نهایی یافته‌ها پژوهش است و مسئولیت نهایی پژوهش با محقق/محققین است (مشارکت‌کننده، ۷).
	دسترسی‌پذیری هوش مصنوعی	باید تا حد امکان دسترسی‌پذیری باز برای نتایج تحقیقات هوش مصنوعی فراهم گردد (مشارکت‌کننده، ۹ و ۱۱).
		باید مؤسسات آموزشی - پژوهشی دسترسی آزاد همه افراد به نتایج پژوهش‌ها صرف‌نظر از جایگاه، موقعیت جغرافیایی و ملیت را فراهم آورند (مشارکت‌کننده، ۱۴).
	رعایت منافع مشروع حامیان علمی	اطمینان از اینکه هوش مصنوعی و الگوریتم‌های آن در راستای منافع حامیان و سیاست‌گذاران پژوهش است (مشارکت‌کننده، ۳).
		پژوهشگران باید از بسترهای داده‌ای و پلتفرم‌های تحلیل هوش مصنوعی که استانداردهای اخلاقی و امنیتی مشخص دارند استفاده کنند (مشارکت‌کننده، ۵).

احترام به مالکیت معنوی پژوهش‌ها: یکی از محورهایی که در زمینه مؤلفه‌های اخلاقی کاربرد هوش مصنوعی در پژوهش‌ها مورد توجه مشارکت‌کنندگان در پژوهش بود احترام به مالکیت معنوی پژوهش بود. در این زمینه همان‌طور که مشارکت‌کنندگان بیان داشتند، باید توجه داشت که رعایت مالکیت معنوی در پژوهش‌های مبتنی بر هوش مصنوعی ضمن حفظ اعتبار و حقوق پژوهشگران منجر به ایجاد انگیزه تولید علم و نوآوری می‌شود. در این راستا مشارکت‌کننده (۱۱) بیان داشت: "باید قرار داده‌های حقوق مالکیت معنوی پژوهش را بر مبنای موازین اخلاقی تنظیم کنیم تا مشکلی در این زمینه پیش نیاید. تنظیم این قوانین موجب شفافیت و عدم تعارض در مراحل مختلف پژوهش شده و از سویی منجر به بهبود دیدگاه محققان در زمینه حفظ آثار علمی آنان می‌شود." یا مشارکت‌کننده (۲) بیان داشت: "باید از اول مشخص شود مالک داده و نتایج چه کسی است"

دسترسی‌پذیری: بررسی دیدگاه مشارکت‌کنندگان نشان‌دهنده این مهم بود که در دسترسی‌پذیری پژوهش‌های مبتنی بر هوش مصنوعی نیز باید معیارهای اخلاقی رعایت شده و دسترسی‌پذیری باز در این زمینه تعریف شود. در این زمینه همان‌طور که مشارکت‌کنندگان بیان داشتند، باید توجه داشت که دسترسی‌پذیری در پژوهش‌های مبتنی بر هوش مصنوعی منجر به امکان بازتولید نتایج و اعتبار علمی آن‌ها تضمین شود. پدر این زمینه مشارکت‌کننده (۱۴) بیان داشت: "باید دانشگاه‌ها دسترسی آزاد همه افراد به نتایج پژوهش‌ها صرف‌نظر از جایگاه، موقعیت جغرافیایی و ملیت را فراهم آورند. تا زمینه انتشار دانش و استفاده مناسب از یافته‌های علمی فراهم شود در زمینه هوش مصنوعی و مطالعات منطبق با آن نیز این مهم منجر به اعتباریابی یافته‌های پژوهش‌ها می‌شود." همچنین مشارکت‌کننده (۱۴) بیان داشت: "دانشگاه‌ها باید دسترسی آزاد همه افراد به نتایج پژوهش را فراهم کنند."

رعایت منافع مشروع حامیان علمی: یکی از محورهای که مورد توجه مشارکت کنندگان در پژوهش بود رعایت منافع مشروع حامیان علمی پژوهش‌ها بود که باید مورد توجه قرار گیرد. رعایت منافع حامیان نه تنها تضمین کننده تداوم حمایت‌ها و سرمایه گذاری در این حوزه است، بلکه موجب افزایش اعتماد و شفافیت در فرآیند پژوهش می‌شود. مطالعات معتبر تأکید دارند که رعایت حقوق و منافع مشروع حامیان، به‌ویژه در زمینه انتشار نتایج و استفاده منصفانه از داده‌ها و منابع، برای حفظ مشروعیت و پایداری تحقیق در حوزه هوش مصنوعی حیاتی است. در این زمینه مشارکت کننده (۵) بیان داشت "اطمینان از اینکه هوش مصنوعی و الگوریتم‌های آن در راستای منافع حامیان و سیاست‌گذاران پژوهش است باعث می‌شود حمایت‌ها از این رویکرد پژوهشی افزایش یافته و اختلالی در امر پژوهش‌های مبتنی بر فناوری‌های نوین ایجاد نشود". همچنین مشارکت کننده (۳) بین داشت: "به نظر من، به دغدغه مهم اینه که استفاده از هوش مصنوعی و الگوریتم‌هاش باید واقعاً در راستای منافع مؤسسه و هم همون سیاست‌هایی باشه که پژوهش رو حمایت می‌کنه، نه اینکه فقط ابزار دست افراد یا شرکت‌های خاص بشه. ما باید مطمئن باشیم که این فناوری‌ها به نفع کل مجموعه عمل می‌کنن، نه فقط به عده خاص یا ذی‌نفع بیرونی. سیاست‌گذارها و مدیران دانشگاهی باید دقت کنن که وقتی از هوش مصنوعی تو پژوهش‌ها مون استفاده می‌شه، جهت‌گیری و تصمیماتش کاملاً هم‌راستای اهداف کلان و منافع واقعی مجموعه دانشگاه باشه، نه صرفاً منافع موقت یا شخصی".

ابعاد اخلاق رفتاری در کاربرد هوش مصنوعی در پژوهش

منظور از بعد اخلاق رفتاری، همان رفتار انسان‌ها است. بنابراین بعد اخلاق رفتاری؛ رفتار و روابط انسانی، ارتباطات و الگوهای خاص به‌هم‌پیوسته می‌باشند. معیارهای اخلاقی عوامل رفتاری در کاربرد هوش مصنوعی در پژوهش‌ها در دو محور اصلی مشتمل بر حفظ و تقویت روابط انسانی معنادار در پژوهش و آموزش با هوش مصنوعی و پیشگیری از تمرکز قدرت و نابرابری‌های فناورانه مبتنی بر هوش مصنوعی مورد سازمان‌دهی قرار گرفت. یافته‌های این بعد در جدول ۶ گزارش شده است.

جدول ۶: ابعاد اخلاق رفتاری در کاربرد هوش مصنوعی در پژوهش

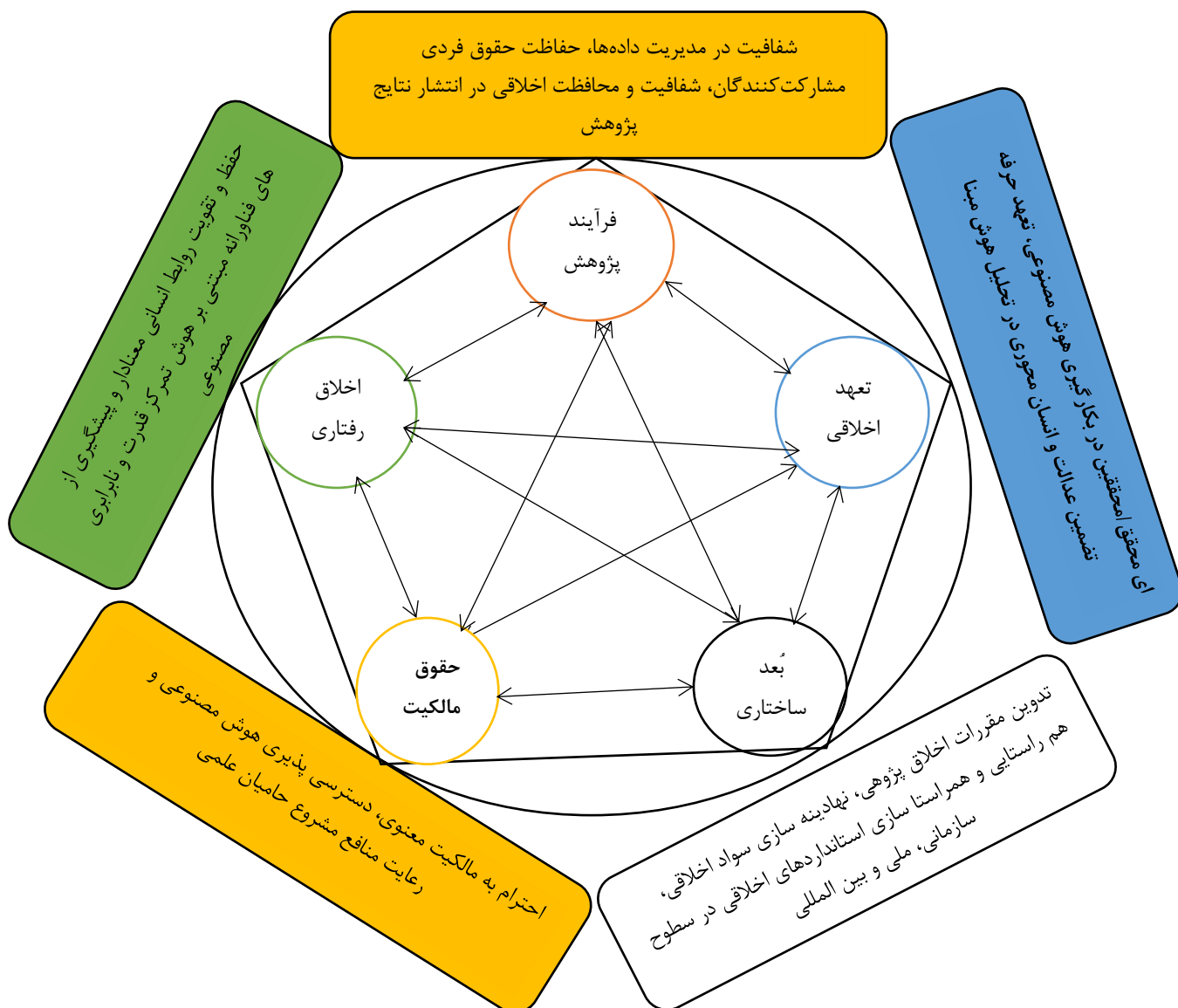
تم اصلی	زیر تم	نمونه بیانات مشارکت کنندگان (کدهای اولیه)
توجه به تأثیر هوش مصنوعی بر ساختار روابط انسانی، پویایی قدرت و فرهنگ همکاری در اکوسیستم‌های پژوهشی و آموزشی	حفظ و تقویت روابط انسانی معنادار در پژوهش و آموزش باهوش مصنوعی	استفاده از هوش مصنوعی در تصمیم‌گیری‌های پژوهشی (مانند تخصیص بودجه یا پذیرش دانشجو) نباید منجر به کاهش تعاملات انسانی و قضاوت‌های مبتنی بر شناخت عمیق افراد شود(مشارکت‌کننده، ۸).
		نگرانی وجود دارد که تمرکز بیش‌ازحد بر تحلیل‌های داده محور هوش مصنوعی، ارزش تجربیات کیفی و دیدگاه‌های انسانی منحصربه‌فرد اساتید و پژوهشگران را در فرآیندها به حاشیه براند(مشارکت‌کننده، ۵)
		باید مراقب بود که ابزارهای هوش مصنوعی به ابزاری برای کنترل یا نظارت بیش‌ازحد بر پژوهشگران جوان تبدیل نشوند و فضای اعتماد و استقلال فکری حفظ شود(مشارکت‌کننده، ۱۱ و ۱۴).
	پیشگیری از تمرکز قدرت و نابرابری‌های فناورانه مبتنی بر هوش مصنوعی	هوش مصنوعی نباید باعث ایجاد یا تشدید سلسله‌مراتب قدرت جدیدی در دانشگاه‌ها شود که در آن، متخصصان هوش مصنوعی بر سایر رشته‌ها یا رویکردهای پژوهشی تسلط نامتناسبی پیدا کنند(مشارکت‌کننده، ۳).
		نگرانی در مورد "جعبه سیاه" بودن برخی الگوریتم‌های هوش مصنوعی وجود دارد که می‌تواند منجر به تصمیم‌گیری‌های غیر شفاف و غیرقابل توضیح در مورد سرنوشت علمی افراد شود(مشارکت‌کننده، ۱).
		باید پژوهشگران را توانمند ساخت تا به‌جای پذیرش منفعلانه خروجی‌های هوش مصنوعی، بتوانند به شکلی انتقادی با آن تعامل کرده و در صورت لزوم آن را به چالش بکشند(مشارکت‌کننده، ۷).

		مؤسسات آموزشی - پژوهشی باید فضایی برای گفتگوی بین‌رشته‌ای در مورد تأثیر هوش مصنوعی بر فرهنگ همکاری، رقابت و مالکیت ایده‌ها در پژوهش ایجاد کنند و این قابلیت برای پژوهش در اختیار عده‌ای خاص نباشد (مشارکت‌کننده، ۱۰). نباید اساتیدی که پژوهش‌های مطلوب با استفاده از هوش مصنوعی می‌توانند تولید کنند بر سایر اساتید یا پژوهشگرانی که در بومیان دیجیتال نیستند برتری احساس کنند بلکه باید از بستر پژوهش برای ارتقا جمعی استفاده نمایند (مشارکت‌کننده، ۶).
--	--	--

حفظ و تقویت روابط انسانی معنادار در پژوهش و آموزش باهوش مصنوعی: از جمله چالش‌ها و ترس‌هایی که در بیانات مشارکت‌کنندگان در پژوهش مشهود بوده ترس از کاهش تعاملات انسانی در زمینه ظهور گسترده هوش مصنوعی به‌عنوان ابزار کمکی اصلی در پژوهش است. در این زمینه مشارکت‌کنندگان تأکید داشتند که تجربیات انسانی و دیدگاه‌های منحصر به فرد انسانی باید مدنظر قرار گرفته و از سویی نظارت نهایی بر عهده انسان است. در این زمینه مشارکت‌کننده (۵) بیان کرد: "این نگرانی وجود دارد که تمرکز بیش از حد بر تحلیل‌های داده محور هوش مصنوعی، ارزش تجربیات کیفی و دیدگاه‌های انسانی منحصر به فرد اساتید و پژوهشگران را در فرآیندها به حاشیه براند؛ این گنجینه یافته‌های نوین است که باید همیشه مدنظر قرار بگیرد چرا که ماشین نمی‌تواند عواطف و تجربیات ناب انسانی را تحلیل و تفسیر صد درصد کند." هم‌چنین مشارکت‌کننده (۱۲) بیان داشت: "تصمیم‌گیری پژوهشی نباید فقط بر داده‌های ماشین تکیه کند، قضاوت انسانی هم لازم است."

پیشگیری از تمرکز قدرت و نابرابری‌های فناورانه مبتنی بر هوش مصنوعی: از دیگر محورهایی که مورد توجه مشارکت‌کنندگان در پژوهش بود این مهم بود که باید توجه داشت که پژوهشگران و متخصصان پژوهش مبتنی بر هوش مصنوعی بر سایر پژوهشگران ارجعیت داشته باشند و کانون‌های قدرت جدیدی در فضای آکادمیک دانشگاهی ایجاد شود. در این زمینه مشارکت‌کننده (۶) بیان داشت: "نباید/اساتیدی که پژوهش‌های مطلوب با استفاده از هوش مصنوعی می‌توانند تولید کنند بر سایر اساتید یا پژوهشگرانی که در بومیان دیجیتال نیستند برتری احساس کنند بلکه باید از بستر پژوهش برای ارتقا جمعی استفاده نمایند" و مشارکت‌کننده (۳) بیان داشت: "اساتیدی که از هوش مصنوعی استفاده می‌کنند نباید احساس برتری بر دیگران داشته باشند."

در میان‌دهی این پنج بُعد را به صورت شبکه‌ای تعاملی نشان می‌دهد که در مرکز آن، مفهوم محوری پژوهش یعنی "استفاده مسئولانه از هوش مصنوعی در پژوهش" قرار دارد؛ مفهومی که از تلاقی تجربه‌های فردی، تعهدات حرفه‌ای و الزامات نهادی شکل می‌گیرد. در این مدل، تعامل متقابل بین ابعاد اخلاقی از جنس "تأثیر دوسویه" است؛ به گونه‌ای که ضعف در هر بعد، بر دیگر ابعاد نیز اثرگذار است. بر اساس روش پدیدارنگاری، این مدل نه یک نقشه علی، بلکه بیان ساختار تجربی پدیده در آگاهی مشارکت‌کنندگان است و جوهره آن را «اعتماد، شفافیت و عدالت اخلاقی در متن پژوهش دانشگاهی» تشکیل می‌دهد. در ادامه در شکل ۲ مدل مفهومی پژوهش ارائه شده است.



شکل ۲: چارچوب اخلاقی استفاده مسئولانه از هوش مصنوعی در پژوهش

بحث و نتیجه‌گیری

طی سال‌های اخیر، هوش مصنوعی به‌عنوان یک فناوری تحول‌آفرین در بسیاری از عرصه‌ها، از جمله آموزش و پژوهش، جایگاه ویژه‌ای پیدا کرده است. این فناوری نه‌تنها مسیر دستیابی به دانش را هموارتر ساخته، بلکه روندهای انجام پژوهش را به‌گونه‌ای تغییر داده است که بسیاری از پژوهشگران، هوش مصنوعی را به‌عنوان دستیار علمی و پژوهشی خود به کار می‌گیرند. این امر به‌ویژه در پژوهش‌های آموزشی اهمیت دوچندانی دارد، زیرا کیفیت و دقت فرایند تحقیق مستقیماً بر رشد و توسعه نظام آموزش و پرورش اثرگذار است. با این حال، استفاده از هوش مصنوعی بدون توجه به ملاحظات اخلاقی و چارچوب‌های استفاده مسئولانه، می‌تواند چالش‌هایی همچون سوگیری، عدم شفافیت، مشکلات دقت محتوا و نقض حریم خصوصی را به همراه داشته باشد. از این‌رو، پژوهش حاضر بر شناسایی و تبیین چارچوب‌های اخلاقی و رویکردهای کیفی مناسب برای استفاده مسئولانه از هوش مصنوعی در پژوهش‌های آموزشی متمرکز بوده است و در ادامه نتایج آن تبیین می‌شود.

نتایج نشان داد که در بُعد فرایند پژوهش، سه مؤلفه اصلی یعنی شفافیت در مدیریت داده‌ها، حفاظت از حقوق فردی مشارکت‌کنندگان، و شفافیت و محافظت اخلاقی در انتشار نتایج نقش تعیین‌کننده‌ای در استفاده مسئولانه از هوش مصنوعی دارند. این یافته بیانگر آن است که بدون شفاف‌سازی مراحل گردآوری و ذخیره‌سازی داده‌ها، رعایت حریم خصوصی و رضایت آگاهانه، و نیز ارائه نتایج به شکلی منصفانه و بدون سوگیری، اعتماد علمی به پژوهش‌های مبتنی بر هوش مصنوعی تضعیف خواهد شد. به عبارت دیگر، مشارکت‌کنندگان پژوهش تأکید کردند که این مؤلفه‌ها باید به عنوان پایه‌های اصلی هرگونه چارچوب اخلاقی بومی در نظر گرفته شوند؛ چراکه رعایت آن‌ها ضمن ارتقای کیفیت و اعتبار پژوهش، می‌تواند زمینه‌ساز استفاده پایدار و قابل‌اعتماد از هوش مصنوعی در آموزش عالی باشد. یافته‌های پژوهش در این بخش همسو با نتایج پژوهش اینام و همکاران (۲۰۲۴) و بچینو (۲۰۲۳) است که بر اهمیت پیاده‌سازی دستورالعمل‌ها و اصول اخلاقی جامع و شفاف، به‌منظور ارتقای مسئولیت‌پذیری و قابل‌اعتماد بودن پژوهش‌هایی که از هوش مصنوعی بهره می‌برند، تأکید دارند. بدین ترتیب، نهادینه کردن این مؤلفه‌ها در قالب چارچوب‌های اخلاقی بومی و کارآمد، نه تنها به بهبود کیفیت و اعتبار پژوهش کمک می‌کند، بلکه زمینه‌ساز استفاده مسئولانه و پایدار از هوش مصنوعی در حوزه‌های آموزش و پژوهش خواهد بود.

نتایج پژوهش نشان داد که در بُعد تعهد، مهم‌ترین مسئله برای استفاده اخلاقی از هوش مصنوعی، پاسخ‌گویی اخلاقی و عدالت‌محور پژوهشگر است. این به آن معناست که محققان نمی‌توانند صرفاً به قابلیت‌های فنی ابزارهای هوش مصنوعی تکیه کنند، بلکه باید مسئولیت کامل تصمیم‌ها و پیامدهای پژوهش را بر عهده بگیرند. تعهد حرفه‌ای، از پذیرش خطاهای احتمالی الگوریتم‌ها گرفته تا شفاف‌سازی شیوه استفاده از آن‌ها، نشان می‌دهد که نقش انسانی و اخلاقی پژوهشگر جایگزین‌ناپذیر است. در همین راستا، تضمین عدالت و انسان‌محوری در تحلیل‌های مبتنی بر هوش مصنوعی به معنای توجه به تأثیر این فناوری بر گروه‌های مختلف اجتماعی و پرهیز از بازتولید تبعیض‌ها و سوگیری‌هاست. به بیان دیگر، یافته‌های این پژوهش نشان می‌دهد که رعایت اصول اخلاقی در بُعد تعهد، نه یک انتخاب، بلکه پیش‌شرطی برای مشروعیت و اعتبار علمی پژوهش‌های دانشگاهی است. این برداشت با نتایج پژوهش‌های پیشین (گرانپوت و همکاران، ۲۰۲۲؛ التیما و همکاران، ۲۰۲۳) نیز همسوست که بر ضرورت تعهد اخلاقی و عدالت‌محور پژوهشگران تأکید کرده‌اند؛ اما آنچه در این مطالعه برجسته شد، تأکید مشارکت‌کنندگان بر پاسخ‌گویی فردی محقق به عنوان عامل اصلی اعتمادسازی و تضمین اصالت پژوهش بود.

یافته‌های پژوهش نشان داد که بعد ساختاری-زمینه‌ای، نقشی تعیین‌کننده در تضمین استفاده اخلاقی از هوش مصنوعی در دانشگاه‌ها دارد. مضمون فراگیر «نهادینه‌سازی نظام‌مند اخلاق پژوهی» بیانگر آن است که اصول اخلاقی نباید صرفاً در حد توصیه باقی بمانند، بلکه لازم است به صورت آیین‌نامه‌ها و سیاست‌های الزام‌آور در سطح دانشگاهی و ملی اجرا شوند. این نتایج نشان می‌دهد که بدون وجود مقررات روشن و کارآمد، استفاده از هوش مصنوعی در پژوهش می‌تواند به تصمیم‌گیری‌های سلیقه‌ای و حتی تخلفات اخلاقی منجر شود. همچنین، مشارکت‌کنندگان بر این نکته تأکید داشتند که ارتقای سواد اخلاقی اساتید و دانشجویان و آموزش عملی در این حوزه، کلید نهادینه‌سازی واقعی اخلاق پژوهش است. در کنار این، هماهنگی بین استانداردهای ملی و بین‌المللی می‌تواند از پراکندگی رویه‌ها جلوگیری کرده و اعتماد علمی را افزایش دهد. این تحلیل با نتایج پژوهش‌های تای (۲۰۲۰) و قریشی (۲۰۲۳) همسوست که بر ضرورت طراحی سازوکارهای جامع و تقویت ظرفیت‌های فرهنگی و نهادی تأکید داشته‌اند. با این حال، یافته‌های حاضر فراتر از این موضوع را نشان داد و بیانگر آن است که بدون تعهد عملی دانشگاه‌ها به تدوین مقررات بومی و آموزش اخلاقی مداوم، استفاده مسئولانه از هوش مصنوعی در پژوهش‌های آموزشی و دانشگاهی به سطح شعار محدود خواهد ماند.

معیارهای اخلاقی مرتبط با حقوق مالکیت در کاربست هوش مصنوعی در پژوهش سازوکارهایی اشاره دارد که حقوق مادی و معنوی ذینفعان را در بستر استفاده از فناوری‌های هوش مصنوعی حفظ می‌کند. یافته‌های این مطالعه نشان داد مضمون فراگیر «حفظ حقوق مالکیت در دسترسی به نتایج پژوهش‌های مبتنی بر هوش مصنوعی» به عنوان بُعدی کلیدی شناسایی شد و

مضامین سازمان‌دهنده آن شامل «احترام به مالکیت معنوی پژوهش‌ها»، «دسترسی‌پذیری هوش مصنوعی» و «رعایت منافع مشروع حامیان علمی» بود. این مضامین نشان می‌دهند که بهره‌برداری از هوش مصنوعی در پژوهش، بدون تضمین رعایت حقوق مؤلفان، پژوهشگران و نهادهای حامی، ناقص خواهد بود. احترام به مالکیت معنوی به‌عنوان یک اصل اساسی، شفافیت فرایندها و رعایت حق دسترسی منصفانه به نتایج تولیدشده را تسهیل می‌کند و بستر لازم را برای اشتراک‌گذاری مسئولانه دستاوردهای پژوهشی فراهم می‌سازد. همچنین، رعایت منافع مشروع حامیان مالی و علمی به‌عنوان بخشی از حقوق مالکیت کمک می‌کند تا روابط پایدار و اخلاقی میان همه ذینفعان حفظ شود. این نتایج همسو با پژوهش اسکیلک و ریمن^۱ (۲۰۲۵) است که به اهمیت نهادینه‌کردن ملاحظات حقوق مالکیت در توسعه و کاربست هوش مصنوعی در حوزه پژوهش اشاره داشته‌اند و بر ضرورت ایجاد چارچوب‌هایی منسجم و شفاف برای حمایت از این حقوق و بهبود اعتماد بین گروه‌های مختلف علمی تأکید دارند. بدین ترتیب، توجه به این بُعد اخلاقی و ساختار سازی مناسب در این زمینه، یکی از پیش‌شرط‌های اصلی استفاده مسئولانه و مؤثر از هوش مصنوعی در پژوهش‌های دانشگاهی محسوب می‌شود.

ابعاد اخلاق رفتاری در کاربرد هوش مصنوعی در پژوهش به‌عنوان بُعدی مهم و راهبردی، با مضمون فراگیر «توجه به تأثیر هوش مصنوعی بر ساختار روابط انسانی، پویایی قدرت و فرهنگ همکاری در اکوسیستم‌های پژوهشی و آموزشی» شناسایی شد. این بُعد اخلاقی بر این موضوع دلالت دارد که بهره‌گیری از هوش مصنوعی نباید منجر به تضعیف روابط و تعاملات انسانی در فرایندهای پژوهشی و آموزشی شود. بر این اساس، دو مضمون سازمان‌دهی «حفظ و تقویت روابط انسانی معنادار در پژوهش و آموزش با هوش مصنوعی» و «پیشگیری از تمرکز قدرت و نابرابری‌های فناورانه مبتنی بر هوش مصنوعی» استخراج گردید. در تبیین این نتایج باید بیان داشت که اگر سیاست‌گذاری و هدایت مناسب صورت نگیرد، هوش مصنوعی ممکن است به کاهش تعاملات اصیل میان استاد و دانشجو یا پژوهشگران منجر شود و جای روابط انسانی را با خروجی‌های مکانیکی پر کند. تحلیل داده‌ها نشان داد که مشارکت‌کنندگان بر استفاده از هوش مصنوعی به‌عنوان ابزاری تسهیلگر و مکمل تأکید دارند، نه جایگزین ارتباطات انسانی؛ زیرا تنها در این صورت می‌توان روح همکاری و همدلی را در پژوهش حفظ کرد. از سوی دیگر، نگرانی مهم دیگر به تمرکز قدرت فناورانه و ایجاد شکاف میان پژوهشگران «دیجیتال‌محور» و دیگران بازمی‌گردد. بدون سیاست‌گذاری دقیق، این خطر وجود دارد که گروهی محدود از متخصصان بر روند پژوهش مسلط شوند و فرصت‌های علمی نابرابر توزیع گردد. چنین نگرانی‌هایی همسو با پژوهش‌های اخلاق‌پور (۱۴۰۲) و کاکنا و همکاران (۲۰۲۴) است که بر ضرورت بازاندیشی در روابط پژوهشی و طراحی سازوکارهای عادلانه تأکید کرده‌اند. با این حال، یافته‌های حاضر نشان می‌دهد که در بافت آموزش عالی ایران، تدوین راهبردهایی برای حفظ تعاملات انسانی و پیشگیری از نابرابری فناورانه نه یک توصیه اختیاری، بلکه ضرورتی فوری برای اعتمادسازی و عدالت علمی است.

به طور خلاصه در بعد فرایند پژوهش، سه مؤلفه شفافیت داده‌ها، حفاظت از حقوق فردی و محافظت اخلاقی در انتشار نتایج، به‌عنوان پایه‌های اعتماد علمی شناسایی شد. در بعد تعهد، پیام کلیدی بر پاسخ‌گویی اخلاقی و عدالت‌محور پژوهشگران تأکید داشت. در بعد ساختاری-زمینه‌ای، نهادینه‌سازی اخلاق پژوهش و هماهنگی استانداردهای ملی و بین‌المللی ضروری دانسته شد. در بعد حقوق مالکیت، احترام به مالکیت معنوی و دسترسی منصفانه به نتایج پژوهش مطرح شد. و در بعد اخلاق رفتاری، ضرورت حفظ تعاملات انسانی و پیشگیری از تمرکز قدرت فناورانه برجسته گردید. پژوهش حاضر با محدودیت‌هایی چون تمرکز نمونه بر اساتید دانشگاه، بدون مشارکت مستقیم دیگر ذینفعان (چون دانشجویان تحصیلات تکمیلی و سایر پژوهشگران) و محدودیت جغرافیایی و فرهنگی، که تعمیم‌یافته‌ها به سایر کشورها را کاهش می‌دهد روبه‌رو بود. در ادامه بر اساس نتایج پژوهش و به تفکیک ابعاد آن، پیشنهادهای ارائه شده است.

¹ Schilke & Reimann

- با توجه به مضمون «رضایت آگاهانه» و «وضوح رویه‌ها»، دانشگاه‌ها باید فرم‌های استاندارد رضایت‌نامه مخصوص پژوهش‌های مبتنی بر هوش مصنوعی طراحی کنند که در آن محدوده استفاده از داده‌ها و نحوه پردازش شفاف ذکر شود.
 - بر اساس مضمون «تعهد حرفه‌ای»، پیشنهاد می‌شود ثبت گزارش استفاده از هوش مصنوعی در هر مرحله پژوهش (جمع‌آوری، تحلیل، نگارش) الزامی شود تا مسئولیت‌پذیری پژوهشگر مستند گردد.
 - متناسب با مضمون «تضمین عدالت»، کمیته‌های اخلاق دانشگاهی موظف شوند الگوریتم‌های مورد استفاده در پژوهش‌ها را از نظر سوگیری اجتماعی و فرهنگی ارزیابی کنند.
 - از مضمون «تدوین مقررات اخلاق پژوهی» برمی‌آید که وزارت علوم باید آیین‌نامه اخلاقی ویژه استفاده از هوش مصنوعی در پایان‌نامه‌ها و طرح‌های پژوهشی تدوین کند.
 - بر اساس مضمون «نهادینه‌سازی سواد اخلاقی»، پیشنهاد می‌شود کارگاه‌های مهارت‌محور اخلاق پژوهش و هوش مصنوعی برای دانشجویان تحصیلات تکمیلی اجباری شود.
 - در راستای مضمون «هم‌راستاسازی استانداردها»، لازم است کمیته مشترک بین دانشگاه‌ها و انجمن‌های علمی تشکیل گردد تا معیارهای ملی با چارچوب‌های بین‌المللی هماهنگ شود.
 - در پی مضمون «احترام به مالکیت معنوی»، هر دانشگاه باید سامانه ثبت مالکیت داده‌ها و نتایج مبتنی بر هوش مصنوعی راه‌اندازی کند تا حقوق محققان حفظ شود.
 - با توجه به مضمون «دسترسی‌پذیری»، پیشنهاد می‌شود نتایج پژوهش که با حمایت دولتی انجام می‌شوند، در قالب مخزن دسترسی آزاد در اختیار پژوهشگران قرار گیرد.
 - در راستای مضمون «پیشگیری از تمرکز قدرت»، پیشنهاد می‌شود دسترسی به ابزارهای هوش مصنوعی پژوهشی به‌صورت عادلانه در اختیار همه گروه‌ها و رشته‌ها قرار گیرد تا انحصار فناورانه شکل نگیرد.
- همچنین پیشنهادهای پژوهشی زیر برای پژوهشگران آینده پیشنهاد می‌شود:
- انجام مطالعات هم‌زمان بر دیدگاه‌های ذی‌نفعان مختلف (اساتید، دانشجویان، مدیران، متخصصان فناوری) برای ساخت چارچوب اخلاقی چندصدا .
 - طراحی مدل سنجش و پایش مسئولیت‌پذیری پژوهشگران در استفاده از هوش مصنوعی .
 - بررسی تأثیر استفاده اخلاقی از هوش مصنوعی بر کیفیت تعاملات انسانی در محیط‌های پژوهشی .
 - مطالعه مقایسه‌ای چارچوب‌های اخلاقی هوش مصنوعی ایران با کشورهای اسلامی و آسیایی .
 - تحلیل تجربیات عملی اجرای آیین‌نامه‌های اخلاقی در پروژه‌های دانشگاهی واقعی.

منابع

اخلاق‌پور، محمد. (۱۴۰۲). تأثیر «سیستم توصیه گر مبتنی بر هوش مصنوعی» در تحولات آموزشی. نشریه جامعه‌شناسی ارتباطات، ۳(۱۲)، ۱۲-۴۴. https://jsc.daneshpajooohan.ac.ir/article_734810.html?lang=fa

پروانه، مسعود، و میرباقری، سید محسن. (۱۴۰۲). مطالعه تطبیقی استانداردهای بین‌المللی و ملی در هوش مصنوعی. فصلنامه علمی مدیریت/استاندارد و کیفیت، ۱۳(۱۱)، ۱۶۹-۱۲۹. <https://doi.org/10.22034/jsqm.2023.397292.1492>

کرسول، جان دبلیو، و کلارک، ویکی پلانو. (۱۳۹۶). روش‌های پژوهش ترکیبی (ترجمه جاوید سرایی و علیرضا کیامنش، چاپ سوم). نشر آبیژ.

غلام‌پور، میثم، و دهباشی، اکرم. (۱۴۰۴). الگوی بهره‌گیری از هوش مصنوعی در ارزشیابی فرآیند طراحی و تدوین برنامه‌های درسی آموزش عالی: رویکرد نظریه داده بنیاد. نظریه و عمل در برنامه درسی، ۱۳(۲۵)، ۸۲-۶۱. <https://doi.org/10.22034/cstp.2025.546099.1101>

Altmäe, S., Sola-Leyva, A., & Salumets, A. (2023). Artificial intelligence in scientific writing: a friend or a foe?. *Reproductive BioMedicine Online*, 47(1), 3-19. <https://doi.org/10.1016/j.rbmo.2023.04.009>

Awosanya, O. D., Harris, A., Creecy, A., Qiao, X., Toepp, A. J., McCune, T., & Ozanne, M. V. (2024). The utility of AI in writing a scientific review article on the impacts of COVID-19 on musculoskeletal health. *Current Osteoporosis Reports*, 22(1), 146-151. <https://doi.org/10.1007/s11914-023-00855-x>

Benichou, L. (2023). The role of using ChatGPT AI in writing medical scientific articles. *Journal of stomatology, oral and maxillofacial surgery*, 124(5), Article 101456. <https://doi.org/10.1016/j.jormas.2023.101456>.

Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., & Vayena, E. (2018). AI4People—An ethical framework for a good AI society. *Minds and Machines*, 28(4), 689–707. <https://doi.org/10.1007/s11023-018-9482-5>

Grunhut, J., Marques, O., & Wyatt, A. T. (2022). Needs, challenges, and applications of artificial intelligence in medical education curriculum. *JMIR medical education*, 8(2), Article e35587. <https://doi.org/10.2196/35587>

Huang, J., Saleh, S., & Liu, Y. (2021). A review on artificial intelligence in education. *Academic Journal of Interdisciplinary Studies*, 10(3), Article 206. <https://doi.org/10.36941/ajis-2021-0077>

Inam, M., Sheikh, S., Minhas, A. M. K., Vaughan, E. M., Krittanawong, C., Samad, Z., Lavie, C. J., Khoja, A., D'Cruze, M., Slipczuk, L., Alarakhiya, F., Naseem, A., Haider, A. H., & Virani, S. S. (2024). A review of top cardiology and cardiovascular medicine journal guidelines regarding the use of generative artificial intelligence tools in scientific writing. *Current Problems in Cardiology*, 49(3), Article 102387. <https://doi.org/10.1016/j.cpcardi.2024.102387>

- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature machine intelligence*, 1(9), 389-399. <https://doi.org/10.1038/s42256-019-0088-2>
- Kacena, M. A., Plotkin, L. I., & Fehrenbacher, J. C. (2024). The use of artificial intelligence in writing scientific review articles. *Current Osteoporosis Reports*, 22(1), 115-121. <https://doi.org/10.1007/s11914-023-00852-0>
- Malik, A. R., Pratiwi, Y., Andajani, K., Numertayasa, I. W., Suharti, S., & Darwis, A. (2023). Exploring artificial intelligence in academic essay: Higher education student's perspective. *International Journal of Educational Research Open*, 5(3), Article 100296. <https://doi.org/10.1016/j.ijedro.2023.100296>
- Qureshi, B. (2023). Exploring the use of ChatGPT as a tool for learning and assessment in undergraduate computer science curriculum: Opportunities and challenges [Preprint]. *arXiv*. <https://doi.org/10.48550/arXiv.2304.11214>
- Schilke, O., & Reimann, M. (2025). The transparency dilemma: How AI disclosure erodes trust. *Organizational Behavior and Human Decision Processes*, 188, Article 104405. <https://doi.org/10.1016/j.obhdp.2025.104405>
- Siau, K., & Wang, W. (2020). Artificial intelligence (AI) ethics: ethics of AI and ethical AI. *Journal of Database Management (JDM)*, 31(2), 74-87. <https://doi.org/10.4018/JDM.2020040105>
- Tai, M. C. T. (2020). The impact of artificial intelligence on human society and bioethics. *Tzu chi medical journal*, 32(4), 339-343. https://doi.org/10.4103/tcmj.tcmj_71_20
- United Nations Educational, Scientific and Cultural Organization. (2021). *Recommendation on the ethics of artificial intelligence*. UNESCO Publishing. https://www.geoethics.org/files/ugd/5195a5_56fc787932864ee199e73208ee6784ec.pdf?index=true
- Guba, E. G. & Lincoln, Y.S. (1982). Epistemological and methodological bases of naturalistic inquiry. *Educational Communication and Technology Journal*, 30(4), 233-252. <https://link.springer.com/article/10.1007/BF02765185>