



فصلنامه علمی پژوهشی اخلاق پژوهی

سال نهم • شماره اول • بهار ۱۴۰۵

Quarterly Journal of Moral Studies
Vol. 9, No. 1, Spring 2026



ابهام اخلاقی در استفاده از هوش مصنوعی در پژوهش‌های برنامه درسی

صادق حامدی نسب*

doi: 10.22034/ethics.2026.52365.1863

چکیده

هدف این پژوهش بررسی ابهام اخلاقی در استفاده از هوش مصنوعی در پژوهش‌های برنامه درسی بود. پژوهش با رویکرد کیفی و روش پدیدارشناسانه انجام شد و جامعه پژوهش را کلیه اساتید برنامه درسی در سال تحصیلی ۱۴۰۴-۱۴۰۵ تشکیل دادند. نمونه‌گیری به صورت هدفمند و تا رسیدن به اشباع نظری ادامه یافت و در نهایت، ۱۴ نفر از اساتید در مصاحبه‌های نیمه‌ساختاریافته شرکت کردند. داده‌ها با روش تحلیل مضمون مورد بررسی قرار گرفت. نتایج تحلیل داده‌ها نشان داد که ابهام اخلاقی در استفاده از هوش مصنوعی، پدیده‌ای چند بعدی است که شامل ۹ مضمون فراگیر و ۲۹ مضمون سازمان‌دهنده می‌شود. این مضامین ناظر بر چالش‌هایی همچون اصالت علمی، مسئولیت‌پذیری پژوهشگر، حریم خصوصی داده‌ها، عدالت و شفافیت، تأثیر هوش مصنوعی بر تفکر و خلاقیت، ابهام در تفسیر داده‌ها، نبود سیاست‌ها و دستورالعمل‌های نهادی روشن، و ضعف در آموزش و توانمندسازی پژوهشگران و راهبردهای مواجهه اساتید با ابهام اخلاقی است. یافته‌ها نشان می‌دهد که فقدان چارچوب‌های اخلاقی شفاف و راهنمای نهادی، می‌تواند به کاهش اعتماد، کیفیت و اصالت پژوهش‌های مبتنی بر هوش مصنوعی منجر شود. بر این اساس، تدوین سیاست‌های اخلاقی روشن، تقویت بازبینی انسانی، کنترل کیفیت تحلیل‌ها و توانمندسازی پژوهشگران به عنوان راهبردهای کلیدی برای کاهش ابهام اخلاقی پیشنهاد می‌شود.

کلیدواژه‌ها

فناوری اطلاعات، ابهام اخلاقی، هوش مصنوعی، اخلاق هوش مصنوعی، اخلاق پژوهش.

* استادیار گروه آموزش علوم تربیتی، دانشگاه فرهنگیان، تهران، ایران. | hamedinasab@cfu.ac.ir

تاریخ دریافت: ۱۴۰۴/۱۰/۱۱ □ تاریخ تأیید: ۱۴۰۵/۰۱/۳۱

■ حامدی نسب، صادق. (۱۴۰۵). ابهام اخلاقی در استفاده از هوش مصنوعی در پژوهش‌های برنامه درسی. فصلنامه اخلاق پژوهی.

doi: 10.22034/ethics.2026.52365.1863 .۸۰-۵۱، (۳۰)۹

مقدمه

هوش مصنوعی^۱ به عنوان یکی از برجسته ترین فناوری های قرن بیست و یکم، تحولات بنیادی در حوزه های مختلف اجتماعی، اقتصادی، فرهنگی و به ویژه آموزشی ایجاد کرده است (Luckin et al., 2016; Zawacki-Richter et al., 2019, p. 1). در آموزش عالی، ابزارها و سامانه های مبتنی بر هوش مصنوعی نه تنها در مدیریت یادگیری و تولید محتوا به کار می روند، بلکه به طور گسترده ای در پژوهش های آموزشی نیز وارد شده اند و نقش تحلیل داده های بزرگ، تولید متون پژوهشی، و پشتیبانی از تصمیم گیری علمی را ایفا می کنند (Wiese et al., 2025, p. 100405). این پیشرفت ها، در کنار فرصت های آموزشی و پژوهشی، مجموعه ای از چالش های اخلاقی را نیز فراروی جامعه دانشگاهی نهاده است.

اخلاق پژوهش از منظر سنتی بر اصولی چون صداقت علمی، احترام به حقوق مشارکت کنندگان، اجتناب از تقلب و سرقت ادبی، و پاسخ گویی علمی تأکید دارد (Resnik, 2020, p. 232). این چارچوب ها بر مبنای این فرض بنا شده اند که پژوهشگر، به عنوان یک عامل انسانی، دارای قضاوت، نیت و مسئولیت اخلاقی است. با ورود هوش مصنوعی به فرایندهای پژوهشی، مرز میان تصمیمات انسان و عملکردهای ماشینی مبهم شده و پرسش هایی بنیادی درباره اصالت علمی، مسئولیت اخلاقی، انصاف و شفافیت در پژوهش ایجاد شده است (Floridi et al., 2018, pp. 689-707; Wiese et al., 2025, p. 100405). به طور خاص در پژوهش های برنامه درسی^۲، که می توانند بر جهت گیری های تربیتی و سیاست های آموزشی اثرگذار باشند، این چالش ها از اهمیت دوچندانی برخوردارند.

در این پژوهش، منظور از «پژوهش های برنامه درسی» بررسی نظام مند جنبه های طراحی، توسعه، اجرا و ارزشیابی برنامه های درسی است، از جمله اهداف و مبانی نظری، محتوای آموزشی، مواد و منابع آموزشی، روش های تدریس و طرحواره های ارزشیابی و تحلیل چگونگی تأثیر این عناصر بر عمل آموزشی، تصمیم گیری های مدرسه ای و سیاست گذاری آموزش (van den Akker, 2010, pp. 175-195). به عبارت دیگر، پژوهش های برنامه درسی هم زمان می توانند به صورت توصیفی (شناسایی آنچه هست)، تحلیلی-تبیینی (بررسی روابط و پیامدها) و طراحی عملیاتی (تولید و آزمون راهکارها / محصولات درسی) انجام شوند و بنابراین، حوزه ای

1. Artificial Intelligence (AI)
2. Curriculum research

میان رشته‌ای است که نیازمند تلفیق نظریه، پژوهش تجربی و دیدگاه‌های طراحی آموزشی است (Pinar, 2019, p. 244).

در سال‌های اخیر، مرورهای نظام‌مند، مطالعات تجربی و گزارش‌های سیاستی بین‌المللی به بررسی جنبه‌های مختلف اخلاق در کاربرد هوش مصنوعی در آموزش و پژوهش پرداخته است. مرورهای نظام‌مند پژوهش‌های منتشرشده درباره ابعاد اخلاقی کاربرد هوش مصنوعی در آموزش عالی و پژوهش علمی نشان می‌دهند که موضوعات محوری در حوزه اخلاق هوش مصنوعی آموزشی شامل حریم خصوصی داده‌ها، سوگیری الگوریتمی، شفافیت فرآیندهای تصمیم‌گیری، مسئولیت‌پذیری و انسجام با اصول عدالت آموزشی است (Wiese et al., 2025, p. 100405; Nature Commun. 2025). این مطالعات بر نیاز به چارچوب‌های اخلاقی، سیاست‌های نهادی و آموزش سواد اخلاقی برای دانشجویان و اساتید تأکید دارند تا بتوان بهره‌وری هوش مصنوعی را بدون خدشه‌دار شدن اصول اخلاق علمی ارتقا داد (ENAI et al., 2023; Lund & et al, 2025).

ابهام اخلاقی^۱، به‌عنوان یکی از مفاهیم کلیدی و مغفول در فلسفه اخلاق و تحقیقات اجتماعی، به وضعیت‌هایی اشاره دارد که در آن فرد نسبت به درستی یا نادرستی یک عمل دچار تردید و عدم قطعیت می‌شود (Jameton, 1984, p. 6; McConnell, 2014). در زمینه هوش مصنوعی، ابهام اخلاقی می‌تواند ناشی از عدم وجود سیاست‌های شفاف، سرعت بالای تحولات فناورانه، یا عدم تطابق میان انتظارات نهادی و عمل پژوهشی باشد. برای مثال، در پژوهش‌های برنامه درسی ممکن است یک استاد برای تحلیل داده‌ها از ابزارهای هوش مصنوعی استفاده کند، اما نتواند تعریف روشنی از نحوه گزارش یا تبیین روش اخلاقی خود ارائه دهد، یا نخواهد بداند چگونه باید نقش هوش مصنوعی را در تولید یافته‌ها در برابر داوران علمی توضیح دهد.

موضوع «ابهام اخلاقی» در پژوهش‌های برنامه درسی در پیوند تنگاتنگ با مفهوم گسترده‌تر یکپارچگی علمی قرار دارد. یکپارچگی علمی ناظر بر پایبندی به ارزش‌های صداقت، امانت‌داری، شفافیت و مسئولیت‌پذیری در همه مراحل تولید و گزارش دانش است (Chen et al., 2024, p. 38811). با ورود فناوری‌های هوش مصنوعی به فرایندهای پژوهشی، مرز میان تصمیم‌گیری انسانی و خروجی‌های ماشینی مبهم شده و پرسش‌هایی جدید درباره اصالت علمی، نسبت‌دهی نویسنده‌گی و مسئولیت اخلاقی پدید آمده است (Chen et al., 2024, p. 38811). مطالعات تجربی اخیر نشان می‌دهد که دانشجویان و پژوهشگران درک‌های متفاوت و گاه



متناقضی از مصادیق «رفتار اخلاقی» هنگام به کارگیری هوش مصنوعی دارند. برای مثال، Tan و Maravilla (2024) گزارش کردند که سهم قابل توجهی از دانشجویان رویکردهایی را نسبت به استفاده نامناسب از ابزارهای تولید محتوا اتخاذ می کنند که می تواند مرز میان استفاده مشروع و تقلب را مبهم سازد (Tan & Maravilla, 2024). همچنین Lund و همکاران (۲۰۲۵) نشان داده اند که باورهای اخلاقی فردی و نه صرفاً سیاست های نهادی، در تعیین رفتار دانشجویان نسبت به تولید متن با کمک هوش مصنوعی نقش مهمی دارند (Lund et al., 2025, pp. 1545-1565). این یافته ها همگی نشانه ای از «ابهام اخلاقی» هستند؛ یعنی وضعیتی که در آن کنشگران پژوهشی (اساتید و دانشجویان) درباره درست / نادرست رفتارهای مرتبط با هوش مصنوعی مطمئن نیستند.

در سطح نهادی، مرورهای سیستماتیک و منابع تحلیلی پیشنهاد می کنند که کاهش ابهام نیازمند رویکردی چندجانبه است: تدوین سیاست های شفاف، آموزش اخلاقی هدفمند، و ایجاد سازوکارهای نظارت و گزارش دهی (Sozon et al., 2024, pp. 1-20). به ویژه، Sozon و همکاران (۲۰۲۴) بر ضرورت ترکیب ابزارهای پیشگیرانه (مانند چارچوب های اعلام استفاده از هوش مصنوعی در گزارش پژوهشی) و آموزش فعال در بستر مؤسسات آموزشی تأکید می کنند. به عبارتی، سیاست گذاری نهادی و آموزش توانمندساز برای اعضای هیئت علمی و دانشجویان می تواند ابهام اخلاقی را کاهش دهد (Sozon et al., 2024, pp. 1-20).

از منظر برنامه درسی، شواهد نشان می دهد که طراحی و بازآرایی ساختار آموزش می تواند نقش مداخله ای در کاهش تقلب مرتبط با هوش مصنوعی ایفا کند. به عنوان نمونه، Kldiashvili و همکاران (۲۰۲۵) در حوزه برنامه های پزشکی نشان دادند که روش های آموزشی مبتنی بر پروژه و تمرکز بر مهارت های پژوهشی، احتمال شباهت آثار دانشجویان و در نتیجه انگیزه سوءاستفاده از ابزارهای ماشینی را کاهش می دهد (Kldiashvili et al., 2025). همچنین مطالعات در حوزه های دیگر آموزشی نشان می دهند که آموزش معلمان و اساتید درباره چالش های اخلاقی هوش مصنوعی (از جمله شناسایی سوگیری های الگوریتمی، حفاظت از داده ها و قواعد نسبت دهی) پیش نیاز موفقیت هر اقدام برنامه ای - درسی است (Akgun & Greenhow, 2021, pp. 431-440).

تحقیقات کیفی نیز به تنوع دیدگاه های اعضای جامعه آموزشی درباره اخلاق هوش مصنوعی اشاره کرده اند. برای نمونه، برخی اساتید هوش مصنوعی را فرصتی برای افزایش کارایی آموزشی می دانند، در حالی که دیگران به فقدان آگاهی اخلاقی، عدم تطابق سیاست های نهادی و فقدان شفافیت در عملکرد الگوریتم ها اشاره کرده اند (Kamali et al., 2024, p. 20). همچنین مطالعات نشان می دهند دانشجویان نیز درک های متفاوتی از مصادیق «رفتار اخلاقی» هنگام استفاده از

هوش مصنوعی در نگارش آکادمیک دارند، به گونه‌ای که بسیاری از آنها مرز میان استفاده مشروع و تقلب را به روشنی نادیده می‌گیرند و این امر می‌تواند به تضعیف یکپارچگی علمی (یعنی پایبندی به صداقت، شفافیت، مسئولیت‌پذیری و ارجاع‌دهی صحیح در تولید و گزارش آثار علمی) منجر شود (Lund & et al, 2025, pp. 1545-1565).

با این وجود، به رغم رشد قابل توجه در پژوهش‌های بین‌المللی، بررسی‌های تجربی درباره ابهام اخلاقی ناشی از استفاده هوش مصنوعی در پژوهش‌های برنامه درسی بسیار محدود است. در ادبیات موجود، تأکید بیشتر بر مسائل عمومی مثل حریم خصوصی، سرقت علمی، یا سوگیری داده‌هاست، در حالی که تجربه زیسته پژوهشگران نسبت به این زمینه‌های اخلاقی پیچیده کمتر مورد توجه قرار گرفته است (Nature Commun. 2025; ENAI et al., 2023). این خلا نظری و تجربی، به ویژه در زمینه پژوهش‌های برنامه درسی که زیربنای جهت‌گیری‌های آموزشی را شکل می‌دهند، نیاز به تحقیق عمیق‌تر را بسیار ضروری می‌سازد.

در بافت آموزش عالی ایران، این مسأله پیچیدگی‌های خاص خود را دارد. برای نمونه، بسیاری از سندها و سرفصل‌های درسی فاقد سیاست‌ها و تعریف‌های روشن درباره سرقت ادبی، انتساب نویسندگی و نقش ابزارهای هوش مصنوعی در تولید متن هستند، به طوری که اساتید و دانشجویان راهنمای صریح و قابل استنادی برای گزارش یا نسبت‌دهی مشارکت‌های ماشینی ندارند (Nushi & Firoozkahi, 2017, pp. 1-18). علاوه بر این، محدودیت‌های ساختاری مانند تأثیر تحریم‌ها و دسترسی ناقص به فناوری، کتابخانه‌ها و خدمات بین‌المللی می‌تواند منجر به استفاده از نسخه‌های قدیمی‌تر ابزارها، راه‌حل‌های غیررسمی یا گذر از کانال‌های غیررسمی شود، وضعیتی که نه تنها ناهمگنی دسترسی را تشدید می‌کند، بلکه پیچیدگی‌های اخلاقی مربوط به اصالت و کیفیت خروجی‌های تولیدشده با هوش مصنوعی را افزایش می‌دهد (Madani, 2021). پژوهش‌های اخیر همچنین نشان می‌دهند که در میان پژوهشگران ایرانی، شیوع «رفتارهای پژوهشی مشکوک» (برای مثال، مسائل مربوط به نسبت‌دهی نویسندگی، افشای استفاده از هوش مصنوعی و مدیریت خطا / سوگیری) به صورت برجسته‌ای با نبود دستورالعمل‌های روشن و آموزش اخلاقی پیوند دارد؛ یعنی نبود راهنمایی نهادی مستقیم، خود یکی از عوامل بالا رفتن ابهام اخلاقی است (Farangi & Nejadghanbar, 2024, p. 103427). همین طور مطالعات مربوط به «سواد هوش مصنوعی» نشان می‌دهد سطح متفاوت آگاهی و تسلط میان دانشجویان و اعضای هیئت علمی می‌تواند مرزهای اخلاقی را مبهم سازد و باعث شود برداشت‌ها از «استفاده مشروع» و «تقلب» در عمل متفاوت باشد (Xiao et al., 2024, pp. 179-199) به خاطر همین ترکیب



محدودیت‌های فناوریانه (مثلاً ناشی از تحریم‌ها و زیرساخت)، خلأهای سیاست‌گذاری نهادی و تفاوت در سطح سواد هوش مصنوعی است که وضعیت ایران از بسیاری کشورها که سیاست‌های روشن یا دسترسی یکنواخت‌تری دارند، متمایز می‌شود؛ از این رو، تجربیات اساتید در مواجهه با ابهام‌های اخلاقی گوناگون و گاه متضاد است و نیاز به بررسی نظام‌مند و بومی دارد. پژوهش‌های تجربی در زمینه اخلاق هوش مصنوعی در آموزش نشان داده‌اند که کارکنان علمی (یعنی اعضای هیئت‌علمی و پژوهشگران دانشگاهی) از دیدگاه‌های متفاوتی درباره این موضوع برخوردارند. بررسی‌های کیفی اخیر با به‌کارگیری نظریه فعالیت (چارچوبی نظری برای تحلیل تعامل انسان، ابزار (مثلاً هوش مصنوعی)، قواعد و زمینه اجتماعی در فعالیت‌های سازمان‌یافته)، نشان می‌دهند که اساتید از دیدگاه‌های متنوع و گاه متناقضی درباره اهمیت، فهم و کاربرد اصول اخلاقی در سازمان‌دهی استفاده از هوش مصنوعی برخوردار هستند (Kamali et al., 2024, p. 20). این تنوع دیدگاه‌ها نشان می‌دهد که فهم «اخلاق هوش مصنوعی» در آموزش و پژوهش هنوز یک پدیده چندوجهی و پیچیده است که نیازمند بررسی عمیق‌تر است.

با توجه به اهمیت برنامه درسی به‌عنوان ستون اصلی نظام‌های آموزشی و نقش پژوهش در شکل‌دهی و بازتولید آن، درک عمیق و تجربی از چگونگی مواجهه پژوهشگران با ابهام اخلاقی در به‌کارگیری هوش مصنوعی در پژوهش‌های برنامه درسی ضروری است. این نیاز زمانی پررنگ‌تر می‌شود که بدانیم چارچوب‌های اخلاق پژوهش موجود عمدتاً برای شرایط پیش از هوش مصنوعی تدوین شده‌اند و توانایی پاسخ‌گویی به چالش‌های نوظهور را ندارند. هدف مشخص پژوهش حاضر، بررسی تجربه زیسته اساتید پژوهشگر در حوزه برنامه درسی درباره ابهام‌های اخلاقی هنگام به‌کارگیری ابزارهای هوش مصنوعی در فرایندهای پژوهشی (شامل طراحی مطالعه، گردآوری داده، تحلیل و گزارش‌دهی) است. پژوهش حاضر با استفاده از رویکرد پدیدارشناسی قصد دارد به این خلأ نظری و تجربی پاسخ دهد و تجربه زیسته اساتید را از ابهام‌های اخلاقی در استفاده از هوش مصنوعی در پژوهش‌های برنامه درسی تبیین کند. در نهایت، نتایج این پژوهش می‌تواند به غنای ادبیات نظری درباره اخلاق پژوهش در عصر هوش مصنوعی کمک کند و مبنای علمی و بومی برای تدوین سیاست‌ها و دستورالعمل‌های اخلاقی در آموزش عالی ایران و فراتر از آن فراهم آورد. بر این اساس، سؤال اصلی پژوهش این است: چگونه پژوهشگران برنامه درسی ابهام‌های اخلاقی مرتبط با به‌کارگیری هوش مصنوعی در فرایندهای پژوهشی خود را تجربه و معناگذاری می‌کنند؟

روش شناسی

با توجه به ماهیت تفسیری و تجربه محور ابهام اخلاقی، رویکردهای کمی توان تبیین عمق و تنوع این پدیده را به طور کامل ندارند؛ زیرا این رویکردها مبتنی بر کمی سازی مفاهیم و سنجش متغیرهای از پیش تعریف شده اند و در نتیجه نمی توانند معانی زمینه مند، روایت های شخصی و فرآیندهای شکل گیری ابهام اخلاقی را به صورت عمیق آشکار سازند. در این میان، پدیدارنگاری به عنوان یک رویکرد کیفی، امکان شناسایی شیوه های گوناگون فهم و ادراک اساتید از یک پدیده مشترک را فراهم می کند (Marton, 1981, pp. 177-200; Åkerlind, 2012, pp. 115-127). این رویکرد به جای تمرکز بر درست یا نادرست بودن کنش ها، به کشف الگوهای متفاوت فهم اخلاق پژوهشی در مواجهه با هوش مصنوعی می پردازد. بر این اساس این پژوهش یک مطالعه کیفی با رویکرد پدیدارنگارانه است.



۱. مشارکت کنندگان: جامعه پژوهش حاضر را اساتید رشته برنامه درسی در دانشگاه های ایران تشکیل می دهند که در سال تحصیلی ۱۴۰۵-۱۴۰۴ در حوزه آموزش عالی به تدریس، پژوهش و راهنمایی پایان نامه ها و رساله ها اشتغال داشته اند. با توجه به ماهیت کیفی و اکتشافی پژوهش و تمرکز آن بر فهم ها و تجربه های زیسته اساتید از ابهام های اخلاقی ناشی از کاربرد هوش مصنوعی در پژوهش های برنامه درسی، از روش نمونه گیری هدفمند مبتنی بر ملاک استفاده شد.

۵۷

انجام

اخلاقی در استفاده از هوش مصنوعی در پژوهش های برنامه درسی

ملاک های ورود به پژوهش

۱. سابقه تدریس یا پژوهش در حوزه برنامه درسی؛
۲. تجربه استفاده مستقیم یا مواجهه حرفه ای با ابزارهای مبتنی بر هوش مصنوعی (نظیر ابزارهای تحلیل داده، نگارش علمی یا تولید محتوا) در فعالیتهای پژوهشی؛
۳. مشارکت در راهنمایی یا داوری پژوهش های دانشگاهی مرتبط با برنامه درسی؛
۴. تمایل آگاهانه به مشارکت در مصاحبه و بیان دیدگاه های اخلاقی بود.

فرآیند گردآوری داده ها از طریق مصاحبه های نیمه ساختاریافته عمیق انجام شد و نمونه گیری تا دستیابی به اشباع نظری ادامه یافت. در نهایت، به صورت فرضی با ۱۴ نفر از اساتید برنامه درسی (۶ نفر مرد و ۸ نفر زن) مصاحبه انجام شد. تنوع در مرتبه علمی، سابقه تدریس و تجربه پژوهشی اساتید، با هدف دستیابی به طیفی از برداشتها و تفسیرهای متفاوت نسبت به



۵۸

فصلنامه علمی - پژوهشی اخلاق پژوهی

سال نهم

شماره اول

بهار ۱۴۰۵

ابهام‌های اخلاقی استفاده از هوش مصنوعی در پژوهش‌های برنامه درسی مدنظر قرار گرفت. مشخصات کلی مشارکت‌کنندگان به صورت کدگذاری شده در جدول ۱ ارائه شده است.

جدول ۱. مشخصات اطلاع‌رسان‌ها

کد	جنسیت	سابقه تدریس (سال)	مرتبه علمی
۱	مرد	۷	استادیار
۲	زن	۱۵	دانشیار
۳	زن	۶	استادیار
۴	مرد	۱۸	دانشیار
۵	زن	۲۵	استاد تمام
۶	مرد	۹	استادیار
۷	زن	۱۴	دانشیار
۸	مرد	۲۷	استاد تمام
۹	زن	۱۶	دانشیار
۱۰	زن	۸	استادیار
۱۱	مرد	۱۹	دانشیار
۱۲	زن	۳۰	استاد تمام
۱۳	زن	۵	استادیار
۱۴	مرد	۱۲	دانشیار

۲. روش تحلیل: داده‌های حاصل از این پژوهش با استفاده از روش تحلیل مضمون^۱ و نرم‌افزار MAXQDA10 تحلیل شد. تحلیل مضمون یکی از روش‌های رایج در تحقیقات کیفی است که به شناسایی، تحلیل و تفسیر الگوهای معنایی (موضوعات) در داده‌ها می‌پردازد. این روش، علاوه بر سادگی و کارآمدی، مهارت‌های اساسی مورد نیاز برای بسیاری از تحلیل‌های کیفی را فراهم می‌کند و یکی از مهارت‌های مشترک در تحلیل‌های کیفی محسوب می‌شود (عابدی جعفری،

1. Thematic Analysis

تسلیمی، فقیهی و شیخزاده، ۱۳۹۰، ص ۱۹۸-۱۵۱).

یکی از ابزارهای پیشرفته در تحلیل مضمون، شبکه مضامین است که آتراید-استیرلینگ آن را توسعه داده است (Attride-Stirling, 2001, pp. 385-405). این روش نقشه‌ای شبیه تارنما ارائه می‌دهد که به‌عنوان اصل سازمان‌دهنده و روش نمایش مضامین عمل می‌کند.

در این چارچوب، داده‌ها ابتدا به مضامین پایه (کدها و نکات کلیدی متن) تقسیم می‌شوند، سپس این مضامین به‌صورت مضامین سازمان‌دهنده ترکیب و تلخیص می‌شوند و در نهایت به مضامین فراگیر (برگیرنده اصول حاکم بر متن به‌مثابه کل) تبدیل می‌شوند. نقشه‌های شبکه تارنما امکان نمایش روابط میان این سه سطح را فراهم می‌آورند و مضامین برجسته هر سطح را نشان می‌دهند (عابدی جعفری و همکاران، ۱۳۹۰، ص ۱۹۸-۱۵۱).

تحلیل مضمون به پژوهشگران این امکان را می‌دهد که با یک رویکرد انعطاف‌پذیر، داده‌ها را به‌طور سیستماتیک بررسی و الگوهای معنادار را شناسایی کنند. این روش به‌ویژه در علوم اجتماعی و روان‌شناسی کاربرد گسترده‌ای دارد و کمک می‌کند تا درک عمیق‌تری از تجربیات و دیدگاه‌های شرکت‌کنندگان به‌دست آید (Clarke, 2017, pp. 241-260).

لازم به ذکر است که در این پژوهش، پدیدارشناسی به‌عنوان چارچوب نظری-روش‌شناختی و تحلیل مضمون به‌عنوان روش عملیاتی و تحلیلی استفاده شد. به عبارت دیگر، تحلیل مضمون در این مطالعه ابزار اصلی استخراج و سازمان‌دهی داده‌ها بود، اما تمام مراحل کدگذاری و تشکیل مضامین تحت گرایش پدیدارشناسانه انجام شد؛ یعنی تمرکز بر «معانی زیسته»، «ساختارهای تجربه» و «جوهر پدیدار» بود، نه صرفاً طبقه‌بندی توصیفی یا کمی رخدادها. این ترکیب، امکان بهره‌مندی همزمان از مزایای سازماندهی داده‌ها توسط تحلیل مضمون و کشف عمیق معناهای زیسته توسط پدیدارشناسی را فراهم می‌کند.

مراحل تحلیل شامل آشنایی با داده‌ها، تولید کدها، شناسایی و مرور موضوعات، و تعریف و نام‌گذاری مضامین است (Braun & Clarke, 2013, pp. 204-207). در تمامی مراحل، توجه ویژه‌ای به حفظ «جوهر پدیدار» و ارائه توصیف غنی از تجربه‌ها صورت پذیرفت تا مضامین نه صرفاً توصیفی، بلکه بازتاب‌دهنده معناهای زیسته باشند. در نهایت، شبکه مضامین مربوطه ترسیم شد و ارتباط میان مضامین پایه، سازمان‌دهنده و فراگیر مشخص گردید.

برای افزایش پایایی داده‌ها از روش توافق بین کدگذاران و برای ارتقای روایی یافته‌ها از بازنگری توسط متخصصان استفاده شد (Cernasev & Axon, 2023, pp. 751-755). این رویکرد، ضمن حفظ اصالت و عمق تجربیات اساتید، امکان ارائه نتایج قابل اعتماد و نظام‌مند را فراهم ساخت.





۶۰

فصلنامه علمی - پژوهشی اخلاق پژوهی | سال نهم | شماره اول | بهار ۱۴۰۵

۳. ملاحظات اخلاقی: در این پژوهش، تمامی ملاحظات اخلاقی مرتبط با اجرای مصاحبه‌های کیفی رعایت شده است. پیش از جمع‌آوری داده‌ها، رضایت آگاهانه از تمامی مشارکت‌کنندگان دریافت شد و اهداف پژوهش، روش‌ها، و نحوه استفاده از داده‌ها به آنها توضیح داده شد. محرمانگی اطلاعات از طریق استفاده از نام‌های مستعار و حذف داده‌های شناسایی شده تضمین گردید و مشارکت در هر مرحله از پژوهش کاملاً داوطلبانه بود.

یافته‌ها

تحلیل داده‌های مصاحبه با اساتید برنامه درسی ۹ مضمون فراگیر و ۲۹ مضمون سازمان‌دهنده مرتبط با ابهام اخلاقی در استفاده از هوش مصنوعی در پژوهش‌های برنامه درسی را نشان داد. مضامین فراگیر شامل: ابهام در اصالت علمی، ابهام در مسئولیت‌پذیری پژوهشگر، ابهام در حریم خصوصی داده‌ها، ابهام در عدالت و شفافیت، ابهام در تأثیر بر تفکر و خلاقیت، ابهام در تفسیر داده‌ها، ابهام در سیاست‌ها و دستورالعمل‌های نهادی، ابهام در آموزش و توانمندسازی پژوهشگران و راهبردهای مواجهه اساتید با ابهام اخلاقی بود. جدول زیر مضامین سازمان‌دهنده و فراگیر استخراج شده را نشان می‌دهد:

جدول ۲. مضامین فراگیر، سازمان‌دهنده و پایه ابهام اخلاقی
در استفاده از هوش مصنوعی در پژوهش‌های برنامه درسی

مضمون فراگیر	مضمون سازمان‌دهنده	مضامین پایه	کد مصاحبه‌شونده
ابهام در اصالت علمی	نقش هوش مصنوعی در تولید محتوا	ترس از کاهش اصالت علمی، تردید در مالکیت فکری، اعتماد به الگوریتم، اطمینان از اعتبار داده‌ها، تأثیر بر نوآوری پژوهشی، کنترل کیفیت محتوا، تفسیر نتایج توسط انسان، ارزیابی صحت منابع، تطبیق با اهداف پژوهش	1, 4, 9, 12
	اعتبار نتایج	دقت الگوریتم‌ها، نیاز به بازبینی انسانی،	1, 3, 7

مضمون فراگیر	مضمون سازمان‌دهنده	مضامین پایه	کد مصاحبه‌شونده
	تحلیل داده‌ها	اجتناب از خطاهای ماشینی، صحت داده‌های ورودی، انطباق با اهداف پژوهش، بررسی ناسازگاری داده‌ها، شفافیت روش تحلیل، مستندسازی مراحل، مقایسه با استانداردهای علمی، تفسیر نتایج	
	گزارش شفاف ابزارها	ارائه گزارش دقیق درباره نقش هوش مصنوعی، شفاف‌سازی مراحل تحلیل، توضیح محدودیت‌ها، ذکر الگوریتم‌های استفاده شده، یادداشت‌برداری از فرایند تصمیم‌گیری، توصیف تأثیر ابزارها، مستندسازی اصلاحات، اطلاع‌رسانی به تیم، ارائه مستندات تکمیلی	2, 4, 9, 13
ابهام در مسئولیت‌پذیری پژوهشگر	تفکیک نقش انسان و ماشین	تعیین حدود تصمیم‌گیری انسانی، شناسایی مسئولیت خطا، استانداردسازی تصمیمات پژوهشی، مستندسازی مداخله انسانی، تفکیک تحلیل ماشینی و انسانی، توضیح نقش پژوهشگر، رعایت اصول اخلاقی، پاسخگویی به پیامدها	2, 6, 9, 14
	پاسخگویی به خطاها	شناسایی مسئول خطاهای تحلیلی، بازخورد به تیم، اصلاح نتایج، اطلاع‌رسانی به ناظر، پیگیری پیامدها، برنامه‌ریزی پیشگیرانه، مستندسازی اقدامات، تدوین دستورالعمل اصلاحی، شفاف‌سازی اقدامات، بررسی پیامدهای اخلاقی	1, 5, 8, 10
ابهام در حریم خصوصی	حفاظت از داده‌های	محافظت از اطلاعات شخصی، رعایت محرمانگی، رضایت آگاهانه، نگهداری	2, 3, 10, 12, 14





مضمون فراگیر	مضمون سازمان‌دهنده	مضامین پایه	کد مصاحبه‌شونده
داده‌ها	دانشجویان	امن داده‌ها، دسترسی محدود، حذف داده‌های غیرضروری، مدیریت مجوزها، بررسی خطرات احتمالی	
امنیت داده‌ها		رمزگذاری داده‌ها، جلوگیری از دسترسی غیرمجاز، مدیریت دسترسی‌ها، نگهداری پشتیبان، بررسی آسیب‌پذیری‌ها، به‌روزرسانی سیستم امنیتی، آموزش تیم پژوهشی، نظارت مستمر، ارزیابی ابزارهای امنیتی	2, 5, 10, 14
استفاده محدود		محدود کردن داده‌ها به اهداف پژوهشی، حذف داده‌های غیرضروری، تعیین محدودیت زمانی، کنترل استفاده ثانویه، بررسی انطباق با قوانین، مدیریت ریسک افشا	1, 3, 7, 11
ابهام در عدالت و شفافیت	سوگیری الگوریتم‌ها	شناسایی سوگیری‌های پنهان، اصلاح نابرابری‌ها، ارزیابی عادلانه نتایج، بررسی تبعیض‌های آموزشی، تحلیل تأثیر داده‌های پیشین، تطبیق الگوریتم با اهداف پژوهش، بازبینی مداوم، گزارش شفاف سوگیری‌ها	1, 3, 9, 10
	دسترسی برابر به ابزارها و زیرساخت‌های هوش مصنوعی	اطمینان از دسترسی تمام دانشجویان، فراهم‌سازی منابع برابر، کاهش نابرابری‌های فناوری، تضمین فرصت برابر در پژوهش، توجه به شرایط ویژه دانشجویان، مستندسازی دسترسی‌ها، اطلاع‌رسانی شفاف	4, 9, 11, 13

مضمون فراگیر	مضمون سازمان‌دهنده	مضامین پایه	کد مصاحبه‌شونده
	شفافیت تصمیمات	توضیح دلایل انتخاب ابزار، معیارهای استفاده از هوش مصنوعی، اعلام محدودیت‌ها، ثبت مراحل تصمیم‌گیری، ارائه مستندات تحلیلی، شفاف‌سازی تغییرات در روش، ارزیابی بازخوردها، تدوین راهنمای تصمیم‌گیری	1, 6, 12, 14
	وابستگی بیش از حد به هوش مصنوعی	کاهش تفکر انتقادی، کاهش خلاقیت، تأثیر بر توانایی حل مسئله، اتکا به نتایج ماشینی، کاهش استقلال فکری پژوهشگران، نادیده گرفتن مهارت‌های سنتی، ضعف در ارزیابی کیفیت محتوا، کاهش نوآوری، محدودیت ایده‌های جدید	1, 6, 9, 12, 14
ابهام در تأثیر بر تفکر و خلاقیت	محدودیت خلاقیت در تحلیل	کاهش تنوع تحلیل، تکیه صرف بر نتایج الگوریتمی، حذف ایده‌های جایگزین، محدودیت در روش‌های خلاقانه، کاهش توانایی فرضیه‌سازی، تضعیف نوآوری در پژوهش، کاهش استقلال تصمیم‌گیری، محدودیت تفکر باز، کاهش توانایی طرح سؤال پژوهشی	3, 7, 9, 14
	کاهش مهارت پژوهشی	کم‌توجهی به مهارت‌های سنتی پژوهش، ناتوانی در ارزیابی کیفی داده‌ها، ضعف در نگارش علمی، اتکا به نتایج ماشینی، کاهش تجربه عملی، تضعیف تحلیل انتقادی، کاهش توانایی حل مسائل پیچیده	2, 5, 10, 13





۶۴

فصلنامه علمی - پژوهشی اخلاق پژوهی | سال نهم | شماره اول | بهار ۱۴۰۵

مضمون فراگیر	مضمون سازمان دهنده	مضامین پایه	کد مصاحبه شونده
ابهام در تفسیر داده‌ها	تفاوت تفسیری انسان و ماشین	ناسازگاری بین تحلیل انسانی و ماشینی، تردید در تفسیر نتایج، مشکل در نتیجه‌گیری، اختلاف در ارزیابی داده‌ها، پیچیدگی در تطبیق با هدف پژوهش، نیاز به بازیابی انسانی، عدم قطعیت در تصمیم‌گیری	1, 7, 10, 11
	انعطاف‌پذیری در روش	تغییر روش‌ها بر اساس نتایج الگوریتم، انطباق با شرایط کلاس، اصلاح روش‌ها هنگام مواجهه با خطا، تطبیق با اهداف پژوهش، توجه به بازخوردها، هماهنگی با تیم پژوهشی	6, 9, 13, 14
	بازیابی چند مرحله‌ای	بررسی مجدد داده‌ها، اصلاح تحلیل‌ها، تطبیق با هدف پژوهش، شناسایی خطاهای تحلیلی، ثبت تغییرات، تطبیق نتایج با استانداردها، کنترل کیفیت خروجی‌ها	1, 4, 12, 14
ابهام در سیاست‌ها و دستورالعمل‌های نهادی	نبود دستورالعمل شفاف	فقدان راهنماهای نهادی، نبود چارچوب اخلاقی، تردید در تصمیم‌گیری‌ها، عدم وجود استانداردهای مشخص، نبود سیاست‌های متناسب با آموزش عالی، پیچیدگی در تطبیق با قوانین، فقدان نمونه‌های عملی	1, 9, 11, 12, 14
	نظارت و کنترل	کمبود نظارت نهادها، نبود ارزیابی‌های اخلاقی، نیاز به کمیته‌های اخلاق، بررسی عملکرد پژوهشگران، کنترل استفاده از ابزارها، ارزیابی منظم پروژه‌ها، گزارش‌دهی شفاف	2, 5, 7, 10



مضمون فراگیر	مضمون سازمان‌دهنده	مضامین پایه	کد مصاحبه‌شونده
	راهنمایی بومی	نبود دستورالعمل‌های متناسب با فرهنگ محلی، نیاز به سیاست‌های داخلی، تطبیق با ارزش‌های آموزشی، استفاده از تجربه‌های بومی، تدوین راهنماهای اخلاقی ملی، مشاوره محلی برای پژوهشگران، تطبیق با محیط دانشگاهی، ارائه دستورالعمل عملی	3, 6, 9, 13
ابهام در آموزش و توانمندسازی پژوهشگران	آموزش و فرهنگ‌سازی	آموزش اساتید درباره اخلاق هوش مصنوعی، توانمندسازی پژوهشگران، آشنایی با پیامدهای اخلاقی به‌کارگیری هوش مصنوعی در پژوهش، ارتقای آگاهی پژوهشگران، ایجاد حساسیت نسبت به چالش‌های AI، برگزاری کارگاه‌های تخصصی، تدوین راهنمای کاربردی	1, 3, 7, 11
	کارگاه‌ها و بازآموزی	ارائه کارگاه‌های عملی، شبیه‌سازی چالش‌ها، تمرین تصمیم‌گیری اخلاقی، مطالعه موردی، یادگیری تعاملی، ارائه مثال‌های واقعی، بازخورد مستمر، تمرین در محیط شبیه‌سازی شده	2, 4, 9, 12
	حمایت و مشاوره	راهنمایی فردی، مشاوره اخلاقی، ایجاد شبکه پشتیبانی، بررسی چالش‌های پژوهشگران، ارائه راهکارهای عملی، پاسخ به سوالات پیچیده، پشتیبانی در تصمیم‌گیری، نظارت مستمر	1, 6, 10, 13
راهنمادهای مواجهه اساتید	راهنبرد شفاف‌سازی و	اعلام صریح استفاده از ابزارهای هوش مصنوعی، ذکر نام و نسخه ابزار، توضیح	2, 5, 8, 11, 12, 13

مضمون فراگیر	مضمون سازمان دهنده	مضامین پایه	کد مصاحبه شونده
با ابهام اخلاقی	مستندسازی	نقش انسان در تحلیل، تفکیک سهم پژوهشگر و ماشین، مستندسازی مراحل پردازش داده، ثبت پرامپت های کلیدی، گزارش محدودیت های الگوریتم، شفاف سازی معیارهای تفسیر داده، نگهداری گزارش تصمیم گیری پژوهشی، ثبت اصلاحات انسانی، مستندسازی خطاهای احتمالی ابزار، ارائه پیوست روش شناسی، ایجاد قابلیت پیگیری تحلیل ها، پاسخ گویی در برابر داوران و کمیته های علمی	
	توانمندسازی و بومی سازی سیاست ها	تدوین دستورالعمل متناسب با بافت دانشگاهی ایران، انطباق با قوانین ملی، تعریف چارچوب اعلام استفاده از هوش مصنوعی، تعیین مسئولیت پژوهشگر، تنظیم آیین نامه داخلی گروه آموزشی، ارتقای سواد هوش مصنوعی اساتید، آموزش تحلیل انتقادی خروجی ها، آشنایی با سوگیری الگوریتمی، برگزاری کارگاه های تخصصی، ایجاد کمیته مشاوره اخلاقی، ارائه راهنمایی موردی در موقعیت های پیچیده	3, 4, 7, 9, 12, 14

۱. ابهام در اصالت علمی

یافته ها نشان داد که یکی از مؤلفه های مهم ابهام اخلاقی اساتید، ابهام در اصالت علمی است که شامل سه مضمون سازمان دهنده اصلی است: نقش هوش مصنوعی در تولید محتوا، اعتبار نتایج تحلیل داده ها و گزارش شفاف ابزارها. اساتید تجربه کرده اند که استفاده از هوش مصنوعی در تولید محتوا پرسش های اخلاقی و حرفه ای متعددی ایجاد می کند.

این ابهام شامل ترس از کاهش اصالت علمی، تردید در مالکیت فکری، اعتماد به الگوریتم و تأثیر منفی آن بر نوآوری پژوهشی است. برای مدیریت این چالش، اساتید اقدام به بازبینی انسانی، کنترل کیفیت محتوا و تفسیر نتایج توسط انسان می‌کنند و مطمئن می‌شوند که نتایج تولیدشده با اهداف پژوهش مطابقت دارد. به‌عنوان مثال، مصاحبه‌شونده شماره ۹ بیان کرد: «زمانی که از چت‌بات برای تولید محتوای پژوهشی استفاده می‌کنم، ابتدا خودم تمام متن را مرور می‌کنم، اصلاحات لازم را انجام می‌دهم و مطمئن می‌شوم که هیچ بخش مهمی از خلایقیت یا دیدگاه شخصی پژوهش از بین نرفته باشد. این کار باعث می‌شود دانشجو و خودم مطمئن باشیم که اصالت علمی حفظ شده است».

همچنین، اساتید دریافتند که نتایج تولیدشده توسط هوش مصنوعی بدون بازبینی انسانی ممکن است نادرست یا ناقص باشد، بنابراین، دقت الگوریتم‌ها، بررسی صحت داده‌های ورودی، اجتناب از خطاهای ماشینی و مستندسازی مراحل تحلیل ضروری است. بررسی ناسازگاری داده‌ها، مقایسه با استانداردهای علمی و شفافیت روش تحلیل نیز از اقداماتی است که اعتبار نتایج را تضمین می‌کند.

علاوه بر این، شفافیت در استفاده از هوش مصنوعی و گزارش دقیق مراحل کار شامل توضیح محدودیت‌ها، ذکر الگوریتم‌های استفاده شده، مستندسازی اصلاحات و اطلاع‌رسانی به تیم پژوهشی، بخش مهمی از رعایت اصول اخلاقی پژوهش محسوب می‌شود. مصاحبه‌شونده شماره ۱ گفت: «وقتی داده‌ها توسط الگوریتم تحلیل می‌شوند، همیشه یک بار خودم تمام نتایج را بررسی می‌کنم و اگر ناسازگاری یا خطایی دیدم اصلاح می‌کنم. این کار باعث می‌شود اطمینان داشته باشم که نتایج معتبر و قابل استناد هستند». همچنین مصاحبه‌شونده شماره ۹ بیان کرد: «در گزارش نهایی، دقیقاً توضیح می‌دهم که از چه الگوریتمی استفاده شده، چه اصلاحاتی انجام داده‌ام و چگونه تصمیم‌گیری‌ها شکل گرفته‌اند. این کار باعث می‌شود هم تیم پژوهشی و هم داوران مقاله مطمئن شوند که استفاده از هوش مصنوعی شفاف و اخلاقی بوده است».

۲. ابهام در مسئولیت‌پذیری پژوهشگر

یافته‌ها نشان داد که یکی دیگر از مؤلفه‌های مهم ابهام اخلاقی اساتید، ابهام در مسئولیت‌پذیری پژوهشگر است که شامل دو مضمون سازمان‌دهنده اصلی است: تفکیک نقش انسان و ماشین و پاسخگویی به خطاها. اساتید تجربه کرده‌اند که با ورود هوش مصنوعی به فرآیند پژوهش، تمایز



بین تصمیمات انسانی و تصمیمات ماشینی، شناسایی مسئولیت خطا و استانداردسازی تصمیمات پژوهشی اهمیت ویژه‌ای پیدا می‌کند. آنها برای مدیریت این چالش، مرزهای دخالت انسان و ماشین را مشخص می‌کنند، مستندسازی مداخلات انسانی را انجام می‌دهند و نقش خود را در تحلیل و تفسیر داده‌ها توضیح می‌دهند تا اصول اخلاقی پژوهش رعایت شود.

به عنوان مثال، مصاحبه‌شونده شماره ۶ بیان کرد: «هرگاه الگوریتم داده‌ای را تحلیل می‌کند، من دقیقاً مشخص می‌کنم کدام بخش تصمیم‌گیری انسانی است و کدام بخش ماشینی. این تفکیک باعث می‌شود اگر خطایی رخ داد، مسئول آن روشن باشد و بتوان اقدامات اصلاحی انجام داد.» همچنین، اساتید تأکید کردند که پاسخگویی به خطاها ضروری است؛ فرآیند پاسخگویی شامل شناسایی مسئول خطاهای تحلیلی، بازخورد به تیم پژوهشی، اصلاح نتایج و اطلاع‌رسانی به ناظر پژوهش است. به عنوان مثال، مصاحبه‌شونده شماره ۸ گفت: «هر خطایی که در تحلیل داده‌ها پیش می‌آید، ابتدا مسئول آن را شناسایی می‌کنیم، به تیم اطلاع می‌دهیم و اصلاحات لازم انجام می‌شود. سپس مستندسازی می‌کنیم و بررسی می‌کنیم که پیامدهای اخلاقی آن چیست تا در آینده تکرار نشود». این اقدامات، تضمین می‌کند که مسئولیت‌پذیری پژوهشگر در استفاده از هوش مصنوعی شفاف و اخلاقی باقی بماند.

۳. ابهام در حریم خصوصی داده‌ها

یافته‌ها نشان داد که یکی دیگر از مؤلفه‌های ابهام اخلاقی اساتید، ابهام در حریم خصوصی داده‌ها است که شامل سه مضمون سازمان‌دهنده اصلی است: حفاظت از داده‌های دانشجویان، امنیت داده‌ها و استفاده محدود از داده‌های شخصی دانشجویان. اساتید تجربه کرده‌اند که ورود هوش مصنوعی به پژوهش‌های برنامه درسی حساسیت بیشتری نسبت به حریم خصوصی داده‌ها ایجاد کرده است. برای مدیریت این چالش، آنها اطلاعات شخصی دانشجویان را با رعایت محرمانگی، کسب رضایت آگاهانه، نگهداری امن داده‌ها و محدود کردن دسترسی محافظت می‌کنند و همچنین داده‌های غیرضروری را حذف می‌کنند تا ریسک افشا کاهش یابد.

در این راستا، مصاحبه‌شونده شماره ۱۲ بیان کرد: «اطلاعات شخصی دانشجویان را فقط برای اهداف پژوهشی جمع‌آوری و نگهداری می‌کنم و به هیچ وجه اجازه دسترسی غیرمجاز نمی‌دهم. حتی پس از پایان پژوهش، داده‌های غیرضروری حذف می‌شوند تا حریم خصوصی حفظ شود». علاوه بر این، اساتید بر امنیت داده‌ها تأکید دارند و اقداماتی مانند رمزگذاری

داده‌ها، مدیریت دسترسی‌ها، نگهداری پشتیبان و آموزش تیم پژوهشی برای پیشگیری از دسترسی غیرمجاز انجام می‌دهند. مصاحبه‌شونده شماره ۱۴ گفت: «تمام داده‌ها رمزگذاری می‌شوند و تیم پژوهشی آموزش دیده تا از دسترسی‌های غیرمجاز جلوگیری شود. هر مرحله از جمع‌آوری تا تحلیل داده‌ها تحت نظارت دقیق است». همچنین، استفاده از داده‌ها محدود به اهداف پژوهشی بوده و زمان‌بندی، کنترل استفاده ثانویه و انطباق با قوانین رعایت می‌شود تا اطمینان حاصل شود که حریم خصوصی دانشجویان در تمام مراحل پژوهش حفظ می‌شود.

۴. ابهام در عدالت و شفافیت

یافته‌ها نشان داد که یکی دیگر از مؤلفه‌های ابهام اخلاقی اساتید، ابهام در عدالت و شفافیت است که شامل سه مضمون سازمان‌دهنده اصلی: سوگیری الگوریتم‌ها، دسترسی برابر و شفافیت تصمیمات است. اساتید تجربه کرده‌اند که استفاده از هوش مصنوعی در پژوهش‌های برنامه درسی می‌تواند منجر به ایجاد یا تقویت نابرابری‌ها و سوگیری‌های پنهان شود. از این‌رو، آنها سوگیری‌های احتمالی الگوریتم‌ها را شناسایی، اصلاح نابرابری‌ها، ارزیابی عادلانه نتایج و بررسی تبعیض‌های آموزشی را به‌صورت مداوم انجام می‌دهند.

در این راستا، مصاحبه‌شونده شماره ۳ بیان کرد: «هر بار که الگوریتمی برای تحلیل داده‌ها استفاده می‌کنم، نتایج را بررسی می‌کنم تا مطمئن شوم که هیچ دانشجویی به دلیل داده‌های پیشین یا شرایط خاص خود، ناعادلانه تحت تأثیر قرار نگرفته باشد». علاوه بر این، اساتید بر دسترسی برابر تأکید دارند و اطمینان حاصل می‌کنند که تمام دانشجویان به منابع و فرصت‌های پژوهشی دسترسی یکسان دارند و مستندسازی و اطلاع‌رسانی شفاف در این زمینه انجام می‌شود. همچنین، در حوزه شفافیت تصمیمات، اساتید توضیح دلایل انتخاب ابزار، معیارهای استفاده از هوش مصنوعی، ثبت مراحل تصمیم‌گیری و ارائه مستندات تحلیلی را انجام می‌دهند تا فرآیند پژوهش شفاف و قابل پیگیری باشد. مصاحبه‌شونده شماره ۱۲ گفت: «تمام مراحل تصمیم‌گیری و تغییرات روش تحلیل را مستندسازی می‌کنم و برای تیم و داوران روشن می‌کنم که چرا از ابزار خاصی استفاده شد و چگونه تصمیمات اتخاذ شده است».

۵. ابهام در تأثیر بر تفکر و خلاقیت

یافته‌ها نشان داد که یکی دیگر از مؤلفه‌های ابهام اخلاقی اساتید، ابهام در تأثیر بر تفکر و





خلاقیت است که شامل سه مضمون سازمان‌دهنده اصلی: وابستگی بیش از حد به هوش مصنوعی، محدودیت خلاقیت در تحلیل، و کاهش مهارت پژوهشی می‌باشد. اساتید تجربه کرده‌اند که اتکای بیش از حد به ابزارهای هوش مصنوعی می‌تواند توانایی تفکر انتقادی، خلاقیت و استقلال فکری پژوهشگران را کاهش دهد و مهارت‌های سنتی پژوهش، مانند تحلیل داده‌ها و نگارش علمی، تحت تأثیر قرار گیرد. برای مقابله با این چالش، آنها همواره بازبینی انسانی، تلفیق روش‌های سنتی و ماشینی و تقویت مهارت‌های پژوهشی دانشجویان و خود پژوهشگر را در دستور کار قرار می‌دهند.

در این راستا، مصاحبه‌شونده شماره ۹ گفت: «گاهی دانشجویان و حتی خودم وسوسه می‌شویم که تمام تحلیل‌ها را به الگوریتم بسپاریم، اما همیشه تلاش می‌کنم خودم نیز تحلیل مستقل انجام دهم تا تفکر انتقادی و خلاقیت از بین نرود». همچنین مصاحبه‌شونده شماره ۵ بیان کرد: «هوش مصنوعی ابزار قدرتمندی است، اما اگر تمام مراحل پژوهش به آن واگذار شود، مهارت‌های پژوهشی مانند ارزیابی کیفی داده‌ها و نگارش علمی کم‌رنگ می‌شود؛ بنابراین من همیشه ترکیبی از تحلیل ماشینی و انسانی را به کار می‌برم تا استقلال فکری و توانایی حل مسئله حفظ شود».

۶. ابهام در تفسیر داده‌ها

یافته‌ها نشان داد که یکی دیگر از مؤلفه‌های ابهام اخلاقی اساتید، ابهام در تفسیر داده‌ها است که شامل سه مضمون سازمان‌دهنده اصلی: تفاوت تفسیری انسان و ماشین، انعطاف‌پذیری در روش و بازبینی چند مرحله‌ای می‌باشد. اساتید دریافتند که نتایج تولیدشده توسط هوش مصنوعی گاه با برداشت انسانی تفاوت دارد و این ناسازگاری می‌تواند منجر به تردید در تفسیر نتایج، مشکل در نتیجه‌گیری و اختلاف در ارزیابی داده‌ها شود. برای مدیریت این چالش، آنها روش‌های انعطاف‌پذیر طراحی می‌کنند که امکان اصلاح تحلیل‌ها بر اساس بازخورد و تغییر شرایط فراهم شود و از بازبینی چند مرحله‌ای برای تضمین اعتبار و انطباق نتایج با اهداف پژوهش استفاده می‌کنند.

در این راستا، مصاحبه‌شونده شماره ۱۰ بیان کرد: «گاهی نتایج الگوریتم با برداشت من تفاوت دارد؛ در این مواقع داده‌ها را دوباره بررسی و تحلیل انسانی انجام می‌دهم تا مطمئن شوم نتیجه‌گیری صحیح است». همچنین مصاحبه‌شونده شماره ۱۴ افزود: «روش تحلیل را به گونه‌ای طراحی کرده‌ام که انعطاف‌پذیر باشد؛ اگر خطایی دیدم، مراحل بازبینی چند مرحله‌ای را اجرا و با تیم هماهنگ می‌کنم تا نتایج همسو با اهداف پژوهش و استانداردهای علمی باشد».

۷. ابهام در سیاست‌ها و دستورالعمل‌های نهادی

یافته‌ها نشان داد که یکی دیگر از مؤلفه‌های ابهام اخلاقی اساتید، ابهام در سیاست‌ها و دستورالعمل‌های نهادی است که شامل سه مضمون سازمان‌دهنده اصلی: نبود دستورالعمل شفاف، نظارت و کنترل، و راهنمایی بومی (دستورالعمل‌ها و رویه‌های سازگار با مقتضیات حقوقی، فرهنگی، زبانی و زیرساختی ایران و یا هر مؤسسه دانشگاهی) می‌باشد. اساتید تجربه کرده‌اند که فقدان چارچوب‌های نهادی مشخص و راهنماهای اخلاقی باعث تردید در تصمیم‌گیری‌ها، پیچیدگی در تطبیق با قوانین و عدم وجود استانداردهای روشن می‌شود. همچنین کمبود نظارت و ارزیابی منظم، نبود کمیته‌های اخلاق و کنترل ناکافی بر استفاده از ابزارهای هوش مصنوعی، شرایطی ایجاد می‌کند که پژوهشگران در مواجهه با تصمیمات اخلاقی دچار ابهام می‌شوند. برای مثال: ۱) نبود رویه یکپارچه برای اعلام و گزارش به‌کارگیری ابزارهای هوش مصنوعی در مقالات و گزارش‌های پژوهشی (آیا باید «نام ابزار، نسخه و پرامپت» اعلام شود یا نه؟)، و ۲) نبود پروتکل واحد برای حفاظت از داده‌های دانشجویان هنگام بهره‌گیری از سرویس‌های ابری خارجی (چه نهاد / فردی مسئول تضمین امنیت و محرمانگی است؟). علاوه بر این، کمبود نظارت و ارزیابی منظم، نبود کمیته‌های اخلاق فعال یا کنترل ناکافی بر استفاده از ابزارهای هوش مصنوعی، شرایطی ایجاد می‌کند که پژوهشگران در مواجهه با تصمیمات اخلاقی دچار ابهام می‌شوند. در نتیجه، هم سطح سیاست‌گذاری ملی و هم ساختارهای نهادی دانشگاهی نیاز به شفاف‌سازی و هماهنگی دارند؛ از قبیل اسناد راهنمای ملی یا اصلاح آیین‌نامه‌های داخلی دانشگاه‌ها که نقش وزارت و نهادهای مربوطه را مشخص کند. در این راستا مصاحبه‌شونده شماره ۹ بیان کرد: «گاهی نمی‌دانیم کدام استانداردها یا دستورالعمل‌ها باید در استفاده از هوش مصنوعی رعایت شود، زیرا هیچ راهنمای نهادی شفاف وجود ندارد؛ این باعث می‌شود تصمیم‌گیری‌هایمان دچار تردید شود.» همچنین مصاحبه‌شونده شماره ۱۳ افزود: «ما به راهنمایی بومی و سیاست‌های داخلی نیاز داریم که با ارزش‌ها و فرهنگ دانشگاهی ما هماهنگ باشد؛ با چنین چارچوبی می‌توانیم پژوهش اخلاقی و مسئولانه انجام دهیم و از ابهامات جلوگیری کنیم».

۸. ابهام در آموزش و توانمندسازی پژوهشگران

یافته‌ها نشان داد که یکی دیگر از مؤلفه‌های ابهام اخلاقی اساتید، ابهام در آموزش و توانمندسازی



پژوهشگران است که شامل سه مضمون سازمان‌دهنده اصلی: آموزش و فرهنگ‌سازی، کارگاه‌ها و بازآموزی، و حمایت و مشاوره می‌باشد. اساتید تأکید کردند که نبود آموزش کافی در زمینه اخلاق هوش مصنوعی، فقدان آگاهی نسبت به پیامدهای اخلاقی و عدم توانمندسازی پژوهشگران، منجر به تردید و سردرگمی در تصمیم‌گیری‌های پژوهشی می‌شود. ارائه آموزش‌های تخصصی، کارگاه‌های عملی و شبیه‌سازی چالش‌های واقعی می‌تواند حساسیت و آگاهی پژوهشگران و ذی‌نفعان حوزه آموزش را نسبت به مسائل اخلاقی افزایش دهد و مهارت‌های لازم برای مواجهه با چالش‌های هوش مصنوعی را تقویت کند.

در این راستا مصاحبه‌شونده شماره ۱۱ بیان کرد: «کارگاه‌های عملی درباره اخلاق هوش مصنوعی باعث شد دیدگاه من نسبت به مسئولیت‌های پژوهشی شفاف‌تر شود و بتوانم تصمیم‌های دقیق‌تری در مواجهه با داده‌ها و نتایج ماشینی بگیرم». همچنین مصاحبه‌شونده شماره ۱۳ افزود: «داشتن شبکه مشاوره‌ای و پشتیبانی اخلاقی باعث می‌شود در مواجهه با موقعیت‌های دشوار، تصمیمات پژوهشی‌ام با اطمینان و مسئولیت‌پذیری اتخاذ شود».

۹. راهبردهای مواجهه اساتید با ابهام اخلاقی

یافته‌ها نشان داد که یکی از راهبردهای اصلی اساتید برای مدیریت ابهام‌های اخلاقی، راهبرد شفاف‌سازی و مستندسازی است که خود شامل چند مضمون مانند اعلام صریح و مفصل به‌کارگیری ابزارهای هوش مصنوعی، نگهداری اسناد تصمیم‌گیری پژوهشی و ثبت مداخلات انسانی، پیوست روش‌شناسی حاوی پرامپت‌ها، اصلاحات و محدودیت‌ها است. اساتید گزارش دادند که این اقدامات به‌عنوان ابزارهای پاسخ‌گویی و تضمین اصالت عمل می‌کنند؛ آنها معتقدند با شفاف‌سازی نقش انسان در تحلیل و ارائه مستندات فرایندی می‌توان ابهام در اصالت و مسئولیت‌پذیری را کاهش داد.

برای نمونه، اساتید اقدام به درج پیوست‌های روش‌شناسی در گزارش نهایی، ثبت کامل مراحل تحلیل و نگهداری سوابق اصلاحات انسانی می‌کنند تا امکان پیگیری و ارزیابی اخلاقی فراهم شود. به‌عنوان مثال، مصاحبه‌شونده شماره ۱۱ گفت: «در گزارش نهایی، دقیقاً توضیح می‌دهم که از چه الگوریتمی استفاده شده، چه اصلاحاتی انجام داده‌ام و چگونه تصمیم‌گیری‌ها شکل گرفته‌اند» و مصاحبه‌شونده شماره ۲ نیز اظهار کرد: «همیشه یک‌بار تمام نتایج تحلیل شده توسط الگوریتم را بازبینی می‌کنم تا اگر ناسازگاری دیدم، آن را اصلاح کنم».

یافته‌ها همچنین نشان داد که اساتید برای فراهم‌سازی راه‌حل‌های پایدارتر به راهبرد توانمندسازی و بومی‌سازی سیاست‌ها متوسل می‌شوند؛ این راهبرد شامل مضامینی مانند تدوین و اقتباس دستورالعمل‌های بومی و پروتکل‌های داخلی متناسب با مقتضیات حقوقی و زیرساختی ایران، ارتقای سواد و مهارت‌های عملی اساتید و دانشجویان از طریق کارگاه‌ها و آموزش‌های عملی، و ایجاد سازوکارهای حمایتی نهادی مانند کمیته‌های مشاوره اخلاقی یا شبکه‌های پشتیبانی است. اساتید گزارش دادند که آموزش‌های هدفمند و وجود راهنمای بومی سبب می‌شود تصمیم‌گیری‌های اخلاقی ملموس‌تر و قابل اتکا شوند و فشار تصمیم‌گیری فردی کاهش یابد. از این رو، بسیاری کارگاه‌های حضوری و شبیه‌سازی موارد واقعی را برگزار کرده، راهنمای ساده‌شده‌ای برای گزارش استفاده از هوش مصنوعی تدوین کرده و پیگیری ایجاد ساختارهای مشورتی را آغاز نموده‌اند. به‌عنوان مثال، مصاحبه‌شونده شماره ۳ بیان کرد: «کارگاه‌های عملی درباره اخلاق هوش مصنوعی باعث شد دیدگاه من نسبت به مسئولیت‌های پژوهشی شفاف‌تر شود و بتوانم تصمیم‌های دقیق‌تری در مواجهه با داده‌ها و نتایج ماشینی بگیرم» و مصاحبه‌شونده شماره ۹ بر ضرورت داشتن یک راهنمای بومی تأکید کرد: «ما به راهنمایی بومی نیاز داریم تا در مواجهه با موارد پیچیده، دستورالعمل‌های کاربردی داشته باشیم».

بحث و نتیجه‌گیری

یافته‌های پژوهش نشان داد که ابهام اخلاقی اساتید در مواجهه با هوش مصنوعی یک پدیده چندبعدی است که ۹ مضمون فراگیر و ۲۹ مضمون سازمان‌دهنده می‌شود. این مضامین ناظر بر چالش‌هایی همچون اصالت علمی، مسئولیت‌پذیری پژوهشگر، حریم خصوصی داده‌ها، عدالت و شفافیت، تأثیر هوش مصنوعی بر تفکر و خلاقیت، ابهام در تفسیر داده‌ها، نبود سیاست‌ها و دستورالعمل‌های نهادی روشن، و ضعف در آموزش و توانمندسازی پژوهشگران و راهبردهای مواجهه اساتید با ابهام اخلاقی است. این ابعاد نمایانگر پیچیدگی‌های اخلاقی مرتبط با استفاده از فناوری‌های نوین در پژوهش و تحلیل داده‌ها هستند و ضرورت توجه به چارچوب‌های اخلاقی، آموزش و سیاست‌های نهادی را برجسته می‌کنند. در ادامه به تبیین هر یک از این مضامین پرداخته می‌شود:

یکی از مهم‌ترین ابعاد، ابهام در اصالت علمی بود. اساتید اظهار کردند که استفاده از هوش



مصنوعی در تولید و تحلیل داده‌ها می‌تواند منجر به کاهش کیفیت علمی و اعتبار نتایج شود، به ویژه زمانی که بازمینی انسانی کافی صورت نگیرد. برای کاهش این ابهام، آنها بر ضرورت بازمینی انسانی، کنترل کیفیت محتوا و تفسیر نتایج با رویکرد انسانی تأکید کردند. این یافته‌ها با مطالعات Floridi (2019) و Cowls (2019) و Jobin, Ienca, & Vayena (2019, pp. 389-399) همسو است که اهمیت حفظ اصالت علمی و صحت داده‌ها در محیط‌های دیجیتال را برجسته کرده‌اند. همچنین شفافیت در مستندسازی و گزارش‌گیری مراحل تحلیل، که توسط Mittelstadt (2019, pp. 501-507) به عنوان کلید اعتماد به نتایج هوش مصنوعی معرفی شده است، مورد توجه اساتید قرار گرفت و نشان داد که اصالت علمی بدون مشارکت فعال انسانی و کنترل مستمر تضمین نمی‌شود و چارچوب‌های اخلاقی روشن برای حمایت از پژوهشگران ضروری است.

بعد دیگر یافته‌ها مربوط به ابهام در مسئولیت‌پذیری پژوهشگر بود. اساتید نگرانی خود را از ابهام مرز میان تصمیم‌گیری انسانی و ماشینی و شناسایی مسئولیت در خطاهای تحلیلی بیان کردند. آنها بر نیاز به مستندسازی مداخله انسانی، پاسخگویی به خطاها و شفاف‌سازی اقدامات تأکید داشتند. مطالعات Bryson & Theodorou (2019, pp. 3-21) و Weller (2019, pp. 23-40) نیز بر ضرورت تفکیک مسئولیت‌های انسانی و ماشینی و ایجاد شفافیت در محیط‌های هوشمند تأکید دارند. علاوه بر این، ایجاد چارچوب‌های استاندارد برای تصمیم‌گیری پژوهشی و مدیریت پیامدهای اخلاقی به عنوان راهکارهای کلیدی برای کاهش ابهام مسئولیت‌پذیری مطرح شد، زیرا مسئولیت‌پذیری پژوهشگر باید در ترکیب با هوش مصنوعی به گونه‌ای تعریف شود که هم اعتبار علمی و هم اعتماد به نتایج حفظ شود.

ابعاد دیگر ابهام اخلاقی مربوط به حریم خصوصی داده‌ها بود. یافته‌ها نشان داد که اساتید در مورد حفاظت از داده‌های دانشجویان، رعایت محرمانگی و مدیریت دسترسی‌ها حساسیت دارند و محدود کردن استفاده از داده‌ها به اهداف پژوهشی و حذف داده‌های غیرضروری، به کاهش ریسک افشای اطلاعات کمک می‌کند. این نتیجه با مطالعات Taddeo & Floridi (2018, pp. 751-752) و Mittelstadt (2017, pp. 157-175) همسو است که اهمیت حفظ حریم خصوصی و امنیت داده‌ها در استفاده از هوش مصنوعی را تأکید می‌کنند. همچنین، مدیریت مجوزهای دسترسی و نظارت مستمر بر فرآیندهای داده‌ای موجب افزایش اعتماد پژوهشگران و دانشجویان به محیط پژوهشی می‌شود و پایه‌های اخلاقی پژوهش علمی را تقویت می‌کند.

بعد دیگر ابهام اخلاقی، عدالت و شفافیت است. یافته‌ها نشان داد که اساتید به شناسایی سوگیری‌های الگوریتمی، تضمین دسترسی برابر و شفافیت در تصمیم‌گیری‌ها توجه دارند و معتقدند که رعایت این اصول، اعتماد پژوهشگران و عدالت در دسترسی به فرصت‌ها را افزایش می‌دهد. این نتایج با پژوهش‌های O'Neil (2016) و Binns (2018, pp. 149-159) همسو است که ارزیابی منصفانه الگوریتم‌ها و کاهش تبعیض‌های پنهان را کلید بهره‌برداری اخلاقی از فناوری هوش مصنوعی معرفی کرده‌اند. همچنین، شفافیت در گزارش نتایج و اعلام محدودیت‌های تحلیل الگوریتمی، یکی از الزامات ایجاد اعتماد و افزایش پذیرش فناوری در محیط پژوهشی محسوب می‌شود.

پژوهش همچنین نشان داد که اساتید با چالش‌هایی مانند وابستگی بیش از حد به هوش مصنوعی، محدود شدن خلاقیت و کاهش مهارت‌های پژوهشی مواجه‌اند. برای مدیریت این مشکلات، اساتید بر نیاز به انعطاف‌پذیری در روش‌ها، بازبینی چندمرحله‌ای تحلیل‌ها و ایجاد سیاست‌های نهادی شفاف تأکید کردند. همچنین کمبود آموزش و توانمندسازی پژوهشگران نبود راهنمایی بومی و نهادی ضرورت تدوین برنامه‌های آموزشی و سیاست‌گذاری دقیق را برجسته می‌سازد. این یافته‌ها با پژوهش‌های Mittelstadt (2019, pp. 501-507)، Floridi et al. (2018) و Taddeo & Floridi (2018, pp. 751-752) همخوانی دارد که بر نقش آموزش و فرهنگ‌سازی در کاهش چالش‌های اخلاقی هوش مصنوعی تأکید دارند.

یافته‌ها نشان می‌دهد که برای کاهش ابهام اخلاقی در به‌کارگیری هوش مصنوعی در پژوهش‌های برنامه‌داری در ایران بایستی دو راهبرد مکمل هم‌زمان به‌کار گرفته شوند: نخست شفاف‌سازی و مستندسازی؛ اقداماتی مانند اعلام صریح ابزارها، درج پیوست روش‌شناسی، نگهداری «دفتر تصمیم» برای ثبت مداخلات انسانی و مستندسازی اصلاحات که به‌سرعت قابلیت پیگیری و پاسخ‌گویی را افزایش می‌دهد و از تضعیف اصالت علمی جلوگیری می‌کند. دوم توانمندسازی و بومی‌سازی سیاست‌ها که شامل ارتقای سواد هوش مصنوعی از طریق کارگاه‌های عملی و آموزش‌های هدفمند، تدوین راهنمای بومی سازگار با مقتضیات حقوقی و زیرساختی کشور، و ایجاد سازوکارهای مشورتی نهادی که با کاهش بار تصمیم‌گیری فردی و نهادینه‌سازی شیوه‌های اخلاقی، ظرفیت سازمانی را برای مدیریت موارد پیچیده تقویت می‌کند. این یافته‌ها با پژوهش‌های Farangi & Nejadghanbar (2023) و UNESCO (2023) همسو است. ترکیب فوری اقدامات شفاف‌ساز با برنامه‌ریزی بلندمدت توانمندسازی بومی، با توجه به



خلاً آیین‌نامه‌ای و ناهمگنی سطح سواد فناورانه در کانتکس ایران، مسیری عمل‌گرا برای تبدیل مواجهه‌های پراکنده به پاسخ‌های سازمان‌یافته و پایدار فراهم می‌آورد.

اگرچه ابعاد ابهام اخلاقی در استفاده از هوش مصنوعی در پژوهش‌های بین‌المللی پیش‌تر گزارش شده است، یافته‌های این پژوهش نشان می‌دهد که در بافت دانشگاه‌های ایران، این ابهام‌ها با تنش‌های مضاعفی همراه‌اند. از یک‌سو، سیاست‌های نهادی روشنی و مدون درباره استفاده از هوش مصنوعی هنوز تثبیت نشده و مسئولیت تصمیم‌گیری به سطح فردی اساتید واگذار شده است؛ از سوی دیگر، فشارهای ساختاری برای تولید علمی و انتشار مقاله، زمینه استفاده ابزاری از فناوری را تقویت می‌کند. این دو وضعیت، تعارضی عملی میان حفظ اصالت علمی و بهره‌گیری از ظرفیت‌های فناورانه ایجاد می‌کند. بدین ترتیب، ابهام اخلاقی صرفاً یک مسئله نظری نیست، بلکه به تجربه‌ای زیسته و تنش‌آفرین در کنش حرفه‌ای اساتید برنامه‌درسی تبدیل شده است.

یافته‌های این پژوهش اگرچه در سطح تجربه زیسته اساتید تحلیل شده‌اند، اما دلالت‌های آن فراتر از سطح فردی بوده و با ابعاد سیاسی-اجتماعی و حکمرانی دانشگاهی نیز پیوند می‌خورند. ابهام اخلاقی گزارش شده از سوی اساتید را می‌توان در چارچوب خلأهای سیاست‌گذاری و نبود دستورالعمل‌های شفاف در سطح نهادی و ملی نیز تفسیر کرد؛ به‌گونه‌ای که فقدان چارچوب‌های رسمی، بار تصمیم‌گیری اخلاقی را به سطح فردی منتقل کرده و مسئولیت را شخصی‌سازی می‌کند. در بافت ایران، که سیاست‌های فناوری، دسترسی به ابزارها، و محدودیت‌های زیرساختی تحت تأثیر شرایط کلان اقتصادی و سیاسی قرار دارند، تجربه ابهام اخلاقی صرفاً یک مسئله فردی نیست، بلکه بازتابی از وضعیت حکمرانی فناوری در آموزش عالی است. همچنین نابرابری در دسترسی به ابزارهای پیشرفته هوش مصنوعی و نبود استانداردهای یکپارچه ارزیابی، می‌تواند پیامدهایی برای عدالت پژوهشی و رقابت علمی ایجاد کند. از این منظر، اخلاق هوش مصنوعی نه تنها یک موضوع حرفه‌ای، بلکه بخشی از مسئله حکمرانی دانشگاهی و تنظیم‌گری علمی محسوب می‌شود که نیازمند سیاست‌های شفاف، سازوکارهای پاسخ‌گویی نهادی و گفت‌وگوی مستمر میان سیاست‌گذاران و کنشگران دانشگاهی است.

به طور کلی نتایج این پژوهش نشان می‌دهد که ابهام اخلاقی اساتید در مواجهه با هوش مصنوعی یک پدیده چندبُعدی است که ترکیبی از چالش‌های اصالت علمی، مسئولیت‌پذیری، حریم خصوصی و عدالت را شامل می‌شود و این ابعاد نیازمند بازبینی انسانی، کنترل کیفیت،

مستندسازی و شفافیت در فرآیندهای پژوهشی هستند. استفاده از فناوری بدون چارچوب اخلاقی مناسب می‌تواند اعتماد و کیفیت پژوهش را کاهش دهد و موجب مخاطرات اخلاقی شود. محدودیت‌های پژوهش شامل نمونه کوچک اساتید و تمرکز بر دانشگاه‌های استان خراسان جنوبی بود که تعمیم‌پذیری یافته‌ها را محدود می‌کند. همچنین، ماهیت کیفی پژوهش اجازه تعمیم آماری نمی‌دهد، هرچند بینش‌های عمیق و کاربردی ارائه شده است. با این حال، پیشنهادها کاربردی شامل تدوین دستورالعمل‌های اخلاقی شفاف برای استفاده از هوش مصنوعی، آموزش و توانمندسازی پژوهشگران، ایجاد راهنمایی بومی و نهادی و تقویت شفافیت و پاسخگویی در فرآیندهای پژوهشی است. ادغام این راهبردها می‌تواند کیفیت تحلیل داده‌ها، اصالت علمی و اعتماد به نتایج پژوهشی را ارتقا دهد و به ایجاد فرهنگ پژوهشی اخلاق‌محور در محیط‌های دانشگاهی کمک کند.



فهرست منابع

- عابدی جعفری، حسن؛ و دیگران. (۱۳۹۰). تحلیل مضمون و شبکه مضامین: روشی ساده و کارآمد برای تبیین الگوهای موجود در داده‌های کیفی. اندیشه مدیریت راهبردی (اندیشه مدیریت)، ۵ (۲)، ۱۵۱-۱۹۸.
- Åkerlind, G. S. (2012). Variation and commonality in phenomenographic research methods. *Higher Education Research & Development*, 31(1), 115–127. doi: 10.1080/07294360.2011.642845
- Akgun, S., & Greenhow, C. (2021). Artificial intelligence in education: Addressing ethical challenges in K-12 settings. *AI and Ethics*, 2(3), 431–440. doi: 10.1007/s43681-021-00096-7
- Attridge-Stirling, J. (2001). Thematic networks: an analytic tool for qualitative research. *Qualitative Research*, 1(3), 385–405. doi: 10.1177/146879410100100307
- Binns, R. (2018). Fairness in machine learning: Lessons from political philosophy. *Proceedings of the 2018 Conference on Fairness, Accountability, and Transparency*, 149–159. <https://proceedings.mlr.press/v81/binns18a/binns18a.pdf>
- Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77-101. doi: 10.1191/1478088706qp063oa
- Braun, V., & Clarke, V. (2013). *Successful qualitative research: A practical guide for beginners*. SAGE Publications.
- Bryson, J., & Theodorou, A. (2019). How society can maintain human-centric AI.

- Ethics and Information Technology*, 21(1), 3–21. [10.1007/978-981-13-7725-9](https://doi.org/10.1007/978-981-13-7725-9)
- Cernasev, A., & Axon, D. R. (2023). Research and scholarly methods: Thematic analysis. *Jaccp: Journal Of The American College Of Clinical Pharmacy*, 6(7), 751–755. doi: [10.1002/jac5.1817](https://doi.org/10.1002/jac5.1817)
- Cernasev, A., & Axon, D. R. (2023). Thematic analysis in qualitative research: An overview. *JACCP Journal of the American College of Clinical Pharmacy*, 6(1), Article e1817. doi: [10.1002/jac5.1817](https://doi.org/10.1002/jac5.1817)
- Chen, Z., Chen, C., Yang, G., He, X., Chi, X., Zeng, Z., & Chen, X. (2024). Research integrity in the era of artificial intelligence: Challenges and responses. *Medicine*, 103(27), e38811. doi: [10.1097/MD.00000000000038811](https://doi.org/10.1097/MD.00000000000038811)
- Clarke, V. (2017). *Thematic analysis*. In C. Willig & W. Stainton-Rogers (Eds.), *The SAGE Handbook of Qualitative Research in Psychology* (pp. 241-260). SAGE Publications.
- ENAI (European Network for Academic Integrity) et al. (2023). Recommendations on the ethical use of Artificial Intelligence in Education. *International Journal for Educational Integrity*. <https://link.springer.com/article/10.1007/s40979-023-00133-4>
- Farangi, M. R., & Nejadghanbar, H. (2024). Investigating questionable research practices among Iranian applied linguists: Prevalence, severity, and the role of artificial intelligence tools. *System*, 125, 103427. doi: [10.1016/j.system.2024.103427](https://doi.org/10.1016/j.system.2024.103427)
- Floridi, L., & Cowls, J. (2019). A unified framework of five principles for AI in society. *Harvard Data Science Review*, 1(1). doi: [10.1162/99608f92.8cd550d1](https://doi.org/10.1162/99608f92.8cd550d1)
- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(9), 389–399. doi: [10.1038/s42256-019-0088-2](https://doi.org/10.1038/s42256-019-0088-2)
- Kamali, J., Alpat, M. F., & Bozkurt, A. (2024). AI ethics as a complex and multifaceted challenge: decoding educators' AI ethics alignment through the lens of activity theory. *International Journal of Educational Technology in Higher Education*. DOI: [10.1186/s41239-024-00496-9](https://doi.org/10.1186/s41239-024-00496-9)
- Kldiashvili, E., Mamiseishvili, A., & Zarnadze, M. (2025). Academic integrity within the medical curriculum in the age of generative artificial intelligence. *Health Science Reports*, 8(2), e70489. doi: [10.1002/hsr2.70489](https://doi.org/10.1002/hsr2.70489)
- Luckin, R., Holmes, W., Griffiths, M., & Forcier, L. B. (2016). *Intelligence unleashed: An argument for AI in education*. Pearson. file: [//C:/Users/Zomorrod/Downloads/PearsonIntelligenceUnleashedFINAL.pdf](https://www.pearson.com/content/dam/pearson/education/usa/IntelligenceUnleashedFINAL.pdf)
- Lund, B., Mannuru, N. R., Teel, Z. A., Lee, T. H., Ortega, N. J., Simmons, S., & Ward, E. (2025). Student perceptions of AI-assisted writing and academic integrity: Ethical concerns, academic misconduct, and use of generative AI in higher education. *AI and Ethics Education*, 1(1), 2. doi: [10.3390/aieduc1010002](https://doi.org/10.3390/aieduc1010002)





- Madani, K. (2021). Have international sanctions impacted Iran's environment? *World*, 2(2), 231–252. doi: [10.3390/world2020015](https://doi.org/10.3390/world2020015)
- Marton, F. (1981). Phenomenography—Describing conceptions of the world around us. *Instructional Science*, 10(2), 177–200. doi: [10.1007/BF00132516](https://doi.org/10.1007/BF00132516)
- McConnell, T. (2014). *Moral dilemmas*. Oxford University Press. <https://philpapers.org/rec/MCCMD>
- Mittelstadt, B. (2017). Ethics of the health-related internet of things: A narrative review. *Ethics and Information Technology*, 19(3), 157–175. doi: [10.1007/s10676-017-9426-4](https://doi.org/10.1007/s10676-017-9426-4)
- Mittelstadt, B. (2019). Principles alone cannot guarantee ethical AI. *Nature Machine Intelligence*, 1, 501–507. doi: [10.1038/s42256-019-0114-4](https://doi.org/10.1038/s42256-019-0114-4)
- Nushi, M., & Firoozkahi, A. H. (2017). Plagiarism policies in Iranian university TEFL teachers' syllabuses: An exploratory study. *International Journal for Educational Integrity*, 13, Article 12. doi: [10.1007/s40979-017-0023-4](https://doi.org/10.1007/s40979-017-0023-4)
- O'Neil, C. (2016). *Weapons of math destruction: How big data increases inequality and threatens democracy*. Crown. DOI: 10.5860/crl.78.3.403
- Pinar, W. F. (2019). *Curriculum theory and historical connections*. In Oxford Research Encyclopedia of Education. Oxford University Press. doi: [10.1093/acrefore/9780190264093.013.1034](https://doi.org/10.1093/acrefore/9780190264093.013.1034)
- Resnik, D. B. (2020). *The ethics of science: An introduction* (2nd ed.). Routledge.
- Sozon, M., Hamdan Mohammad Alkharabsheh, O., Chuan, S. B., & Lim, T.-S. (2024). Cheating and plagiarism in higher education institutions (HEIs): A literature review. *F1000Research*, 12, 1017. doi: [10.12688/f1000research.147140.2](https://doi.org/10.12688/f1000research.147140.2)
- Taddeo, M., & Floridi, L. (2018). How AI can be a force for good. *Science*, 361(6404), 751–752. doi: [10.1126/science.aat5991](https://doi.org/10.1126/science.aat5991)
- Tan, M. J. T., & Maravilla, N. M. A. (2024). Shaping integrity: Why generative artificial intelligence does not have to undermine education. *Frontiers in AI*. Advance online publication. doi: [10.3389/frai.2024.1471224](https://doi.org/10.3389/frai.2024.1471224)
- UNESCO. (2023). *Guidance for generative AI in education and research*. UNESCO. doi: [10.54675/EWZM9535](https://doi.org/10.54675/EWZM9535)
- van den Akker, J. (2010). Building bridges: How research may improve curriculum policies and classroom practices. In S. Nieveen & J. van den Akker (Eds.), *An introduction to educational design research* (pp. 25–39). Netherlands Institute for Curriculum Development (SLO). doi: [10.1007/978-94-6209-359-1_30](https://doi.org/10.1007/978-94-6209-359-1_30)
- Weller, A. (2019). Transparency: Motivations and challenges. *Explainable AI: Interpreting, Explaining and Visualizing Deep Learning*, 23–40. Springer. doi: [10.1007/978-3-030-28954-6_2](https://doi.org/10.1007/978-3-030-28954-6_2)
- Wiese, L. J., Patil, I., Schiff, D. S., & Magana, A. J. (2025). AI ethics education: A systematic literature review. *Computers and Education: Artificial Intelligence*, 8, 100405. doi: [10.1016/j.caeai.2025.100405](https://doi.org/10.1016/j.caeai.2025.100405)

- Williamson, B., & Eynon, R. (2020). Historical threads, missing links, and future directions in AI in education. *Learning, Media and Technology*, 45(3), 223–235 . doi: [10.1080/17439884.2020.1798995](https://doi.org/10.1080/17439884.2020.1798995)
- Xiao, J., Alibakhshi, G., Zamanpour, A., Zarei, M. A., Sherafat, S., & Behzadpoor, S.-F. (2024). How AI literacy affects students' educational attainment in online learning: Testing a structural equation model in higher education context. *The International Review of Research in Open and Distributed Learning*. doi: [10.19173/irrodl.v25i3.7720](https://doi.org/10.19173/irrodl.v25i3.7720)
- Zawacki-Richter, O., Marín, V. I., Bond, M., & Gouverneur, F. (2019). Systematic review of research on artificial intelligence applications in higher education. *International Journal of Educational Technology in Higher Education*, 16(39).

