

چالش‌های اخلاقی استفاده از هوش مصنوعی در فضای مجازی و راهکارهای کنترل آن

امیدعلی مسعودی^۱

خاطره اسدی^۲

(تاریخ دریافت ۱۴۰۴/۸/۵ - تاریخ تصویب ۱۴۰۴/۹/۱۵)

نوع مقاله: علمی پژوهشی

چکیده:

رشد سریع فناوری‌های مبتنی بر هوش مصنوعی در فضای مجازی، ابعاد تازه‌ای از چالش‌های اخلاقی را پیش روی جوامع بشری قرار داده است. این پژوهش با هدف شناسایی و تحلیل این چالش‌ها و ارائه راهکارهای کنترل و مدیریت آن‌ها، به بررسی ابعاد نظری و عملی اخلاق هوش مصنوعی پرداخته است. روش تحقیق به صورت توصیفی - تحلیلی و مبتنی بر منابع اسنادی و مطالعات پیشین انجام گرفته است. یافته‌ها نشان می‌دهد که مسائلی چون نقض حریم خصوصی و کرامت انسانی، سوگیری‌های الگوریتمی، نبود شفافیت در فرآیند تصمیم‌گیری سامانه‌های هوشمند، و تهدید امنیت شغلی از مهم‌ترین دغدغه‌های اخلاقی در این زمینه‌اند. بر اساس نتایج، تدوین چارچوب‌های قانونی و اخلاقی روشن، ارتقای سواد رسانه‌ای و دیجیتال در جامعه، و ایجاد

^۱ - استاد علوم ارتباطات دانشکده فرهنگ و ارتباطات دانشگاه بین المللی سوره، تهران، ایران

^۲ - دانشجوی دکتری علوم ارتباطات دانشکده فرهنگ و ارتباطات دانشگاه بین المللی سوره، تهران، ایران.

نظام نظارتی دقیق بر عملکرد هوش مصنوعی می‌تواند به کاهش پیامدهای منفی آن کمک کند. همچنین، به کارگیری اصول اخلاقی اسلامی در طراحی و استفاده از سامانه‌های هوشمند، زمینه‌ساز تقویت عدالت، مسئولیت‌پذیری و صیانت از ارزش‌های انسانی در فضای دیجیتال است. در مجموع، پژوهش حاضر تأکید دارد که توسعه پایدار فناوری هوش مصنوعی بدون توجه به اصول اخلاقی و فرهنگی، ممکن نیست و تلفیق دانش فنی با ارزش‌های انسانی می‌تواند راهی مطمئن برای بهره‌مندی سالم از این فناوری نوظهور باشد.

کلیدواژه‌ها: هوش مصنوعی، اخلاق فناوری، حریم خصوصی، سوگیری الگوریتمی، اخلاق اسلامی، فضای مجازی.

۱- مقدمه و طرح مساله

ظهور و گسترش فناوری‌های هوش مصنوعی طی دهه‌های اخیر، نقطه‌عطفی در تحول علوم و صنایع مختلف به شمار می‌آید. این فناوری با توانایی تحلیل داده‌های کلان، یادگیری الگوها و تصمیم‌گیری خودکار، توانسته است بسیاری از فعالیت‌های انسانی را متحول سازد و کارآمدی نظام‌های مدیریتی، اقتصادی و اجتماعی را افزایش دهد. با این حال، بهره‌گیری فزاینده از هوش مصنوعی، به‌ویژه در فضای مجازی، مسائل تازه‌ای را در حوزه اخلاق و مسئولیت اجتماعی به همراه آورده است. تصمیم‌گیری‌های خودکار الگوریتم‌ها، استفاده گسترده از داده‌های شخصی و اثرگذاری بر رفتار انسان‌ها از جمله عواملی است که نگرانی‌های اخلاقی و حقوقی متعددی را برانگیخته است.

در فضای مجازی، هوش مصنوعی نقشی کلیدی در تحلیل رفتار کاربران، شخصی‌سازی محتوا، تبلیغات هدفمند و حتی تولید اخبار دارد. این کاربردها اگرچه می‌توانند در جهت ارتقای تجربه کاربری مفید باشند، اما در صورت فقدان ضوابط اخلاقی، ممکن است به نقض حریم خصوصی، گسترش تبعیض‌های الگوریتمی، انتشار اطلاعات نادرست و کاهش اعتماد عمومی منجر شوند. از سوی دیگر، اتکای بیش از حد به سامانه‌های هوشمند در تصمیم‌گیری‌های اداری، آموزشی و

اجتماعی، استقلال انسانی را تحت تأثیر قرار داده و مرز میان اختیار انسانی و تصمیم ماشینی را مبهم ساخته است.

در سطح جهانی، تلاش‌های گسترده‌ای برای تدوین اصول و چارچوب‌های اخلاقی در حوزه هوش مصنوعی آغاز شده است، اما در کشور ما هنوز رویکردی جامع و منسجم که بتواند ابعاد اخلاقی، فرهنگی و دینی را در کنار ملاحظات فنی در نظر گیرد، شکل نگرفته است. از این رو، خلأ پژوهش‌های بومی در زمینه‌ی اخلاق هوش مصنوعی محسوس است. از منظر اسلامی نیز، موضوع اخلاق فناوری نیازمند بازخوانی در پرتو اصولی چون عدالت، کرامت انسانی، پرهیز از ضرر و حفظ حریم خصوصی است؛ اصولی که می‌توانند مبنایی ارزش‌مدار برای کنترل و هدایت توسعه هوش مصنوعی فراهم کنند. در این راستا، تمرکز پژوهش بر شناسایی چالش‌های محوری مانند نقض حریم خصوصی، سوگیری الگوریتمی و تهدید کرامت انسانی، و ارائه راهکارهایی همچون تدوین استانداردهای اخلاقی، ارتقای سواد دیجیتال و ایجاد نظام نظارتی مؤثر است.

در چنین شرایطی، مسئله اصلی این پژوهش آن است که گسترش هوش مصنوعی در فضای مجازی چگونه چالش‌های اخلاقی جدیدی را در حوزه‌های فردی و اجتماعی ایجاد کرده و چه سازوکارهایی می‌تواند به کنترل، پیشگیری و هدایت این چالش‌ها منجر شود. بر این اساس، هدف تحقیق حاضر، تبیین نظام‌مند چالش‌های اخلاقی ناشی از هوش مصنوعی در فضای مجازی و ارائه راهکارهای کنترلی مبتنی بر اصول اخلاق اسلامی و استانداردهای جهانی است. این پژوهش می‌کوشد با تلفیق دیدگاه‌های فلسفی، فنی و فرهنگی، الگویی بومی برای بهره‌گیری مسئولانه، عادلانه و انسانی از فناوری‌های هوش مصنوعی در جامعه ارائه دهد.

۲- بررسی پیشینه -چارچوب مفهومی -مبانی نظری

در این بخش، ابتدا به مرور پیشینه‌ی پژوهش‌های انجام‌شده در حوزه‌ی چالش‌های اخلاقی هوش مصنوعی پرداخته سپس به چارچوب مفهومی و مبانی نظری حاکم بر این موضوع خواهیم پرداخت

۲-۱- پیشینه پژوهش

تقی‌زاده و همکاران (۱۴۰۳). در پژوهشی با عنوان «هوش مصنوعی مولد؛ چالش‌ها و الزامات توسعه

و پیاده‌سازی» انجام دادند. یافته‌های پژوهش نشان می‌دهد ایجاد زیرساخت‌های پایدار و ایمن فنی از قبیل استفاده از داده‌های مصنوعی، یادگیری انتقالی، فنون کاهش سوگیری، محاسبات ابری و توزیع شده می‌تواند در کاهش چالش‌های مرتبط با امنیت داده‌ها، حریم خصوصی، شفافیت، صحت و دقت نتایج و کاهش هزینه‌های محاسباتی مؤثر باشد.

عباسی و همکاران (۱۴۰۳). تحقیق تحت عنوان «چالش‌های اخلاق رقومی هوش مصنوعی و امکان‌سنجی بیمه مسئولیت برای سامانه مبتنی بر هوش مصنوعی» انجام شد. بر اساس یافته‌های پژوهش، مهم‌ترین چالش‌های اخلاقی هوش مصنوعی شامل نگرانی‌های اخلاقی مربوط به جمع‌آوری و استفاده غیرمجاز از داده‌های شخصی است که منجر به نقض حریم خصوصی و کرامت انسانی می‌شود. از سویی یکی از موارد نگرانی عدم رعایت اصول اخلاقی سوگیری و امنیت داده‌ها است. در نتیجه، ضمن تلاش برای توسعه استانداردهای اخلاقی با نظارت و پاسخ‌گویی برای هوش مصنوعی و فضای رقومی، به‌منظور استفاده از هوش مصنوعی در فضای رقومی و مسئولیت‌قصور در محیط‌های رقومی، بیمه مسئولیت رقومی پیشنهاد می‌شود. ارائه‌دهندگان فناوری‌های رقومی ممکن است برای رسیدگی به مسائل مربوط به مسئولیت بالقوه خاص این محیط‌ها، نیاز به بیمه‌نامه تخصصی‌قصور داشته باشند. بیمه‌گران باید سیاست‌هایی ایجاد کنند که منعکس‌کننده خطرات منحصر به فرد مرتبط با محیط‌های رقومی باشد.

رحیمی اقدم و همکاران (۱۴۰۳) مطالعه‌ای با موضوع «چالش‌های اخلاقی اتخاذ هوش مصنوعی در مدیریت منابع انسانی» به ثبت

رسانیدند. مطالعه حاضر با آشکارسازی و کاهش فرضیات اشتباه، به سازمان‌ها کمک می‌کند تا با کسب آگاهی در مورد خطرات اجرای هوش مصنوعی در منابع انسانی، به‌سادگی، منطق تحلیلی حوزه‌های دیگر را به مدیریت و کنترل کارمندان خود منتقل نکنند؛ لذا چارچوب به‌دست آمده به تحقیقات مدیریتی در تلاقی تجزیه و تحلیل کارکنان، هوش مصنوعی و بازتاب‌های اخلاقی مدیریت الگوریتمی انسان‌ها کمک می‌کند.

رحمتی و همکاران (۱۴۰۳) مقاله خود را به موضوع «ملاحظات اخلاقی استفاده از هوش مصنوعی در نظام سلامت» اختصاص دادند. نویسندگان این تحقیق بر این باور هستند که ممیزی اخلاقی،

ایجاد سیاست‌گذاری اخلاقی و دسترسی کلیه جوامع به هوش مصنوعی و ضمانت اجرای جهانی حفاظت از داده‌ها و اجرای آنها از جمله راه کارهای مناسب برخورد با چالش‌های اخلاقی هوش مصنوعی محسوب می‌شود. یکی از پیشرفته‌ترین فن‌آوری‌های اخیر، هوش مصنوعی است که به کمک بشر آمده تا اطلاعات و داده‌های عظیم را مدیریت نموده و نمایی از آینده موضوعات را پیش‌بینی نماید. هوش مصنوعی که در واقع استفاده از کامپیوتر برای مدل‌کردن رفتار هوشمند با حداقل مداخله انسانی است، دغدغه‌های اخلاقی زیادی را نیز به ذهن متبادر می‌کند که باید موردنظر قرار گیرند.

گشمرد (۱۴۰۳) تحقیق خود را به موضوع «چالش‌های اخلاقی کاربرد هوش مصنوعی در پزشکی» اختصاص دادند. نتایج نشان داد که ظهور روبات‌های جراحی و پرستاران رباتیک، عمل به‌جای جراح و مراقبت از بیمار به‌جای پرستار، فرصت‌های شغلی و آینده آنها را تهدید می‌نماید. از طرف دیگر بیماران هنگام برخورد با پزشکان و پرستاران رباتیک، حس همدلی و رفتار مناسب را نخواهند داشت؛ زیرا این روبات‌ها ویژگی‌های انسانی نداشته که از دیگر معضلات بزرگ اخلاقی و اجتماعی محسوب می‌گردد؛ لذا شایسته است متخصصین سیستم‌های بهداشتی - درمانی، اصول مهم اخلاق پزشکی مانند استقلال، سودمندی، عدم سوءاستفاده و عدالت را در عملکرد حرفه‌ای خود در نظر داشته باشند.

عاشوری کیسمی (۱۴۰۳). پژوهشی تحت عنوان «طرح و بررسی برخی از مسائل اخلاقی هوش مصنوعی در هنر» انجام دادند. نتایج پژوهش نشان می‌دهد موضوعات اخلاقی همچون حریم خصوصی و نظارت، دست‌کاری در رفتار، تاری و شفافیت، سوگیری در تصمیمات سیستم، و اتوماسیون و اشتغال حوزه‌ی هنر قابل بررسی است. این موضوعات اخلاقی، گاهی در تقابل و گاهی در ارتباط مستقیم و هم‌سو با یکدیگر قرار دارند. بررسی‌ها نشان می‌دهد پاسخ به پرسش‌های اخلاق هوش مصنوعی در حوزه‌ی هنر، نیازمند پاسخ به پرسش‌های اخلاقی به‌صورت عام در هوش مصنوعی است.

صراف‌زاده و همکاران (۱۴۰۲) تحقیق خود را به موضوع «اهمیت اخلاق در استفاده از هوش مصنوعی در آموزش پزشکی» اختصاص دادند. یافته‌های این پژوهش نشانگر آن است که علی‌رغم

وجود چالش‌های اخلاقی در استفاده از هوش مصنوعی در آموزش پزشکی، می‌توان با استفاده از راهکارهایی این چالش‌ها را به حداقل رساند. اعمال حساسیت ویژه نسبت به صحت عملکرد الگوریتم‌های مورد استفاده در آموزش و درمان، و تشخیص زود هنگام و مؤثر سوگیری‌های احتمالی الگوریتم‌های مورد استفاده در هوش مصنوعی از جمله موارد بسیار مهم در این زمینه است. مدیریت صحیح داده‌های انبوه با حفظ حریم خصوصی می‌تواند بسیاری از چالش‌های اخلاقی این حوزه را مرتفع سازد. تشکیل کمیته مخصوص اخلاق در آموزش که در آن علاوه بر متخصصین اخلاق در پزشکی از متخصصین هوش مصنوعی نیز استفاده گردد.

کلاته آقامحمدی (۱۴۰۲). تحقیق تحت عنوان «چالش‌های اخلاقی هوش مصنوعی در رسانه‌های نوین و راهکارهای مواجهه با آن» انجام دادند. نتایج این تحقیق حاکی از آن است که تنها شناسایی چالش‌های اخلاقی کافی نیست، بلکه ضروری است برای آن‌ها راهکارهایی اندیشیده شود. تدوین قوانین قابل اجرا، حسابرسی اخلاقی و بهره‌گیری از سواد رسانه‌های دیجیتال به عنوان راهکارهای مواجهه با چالش‌های اخلاقی هوش مصنوعی در نظر گرفته شده است.

رجیان ده زیره (۱۴۰۲) پژوهشی تحت عنوان «شناسایی چالش‌ها و قابلیت‌های هوش مصنوعی در آموزش و یادگیری با ارائه راهکارها» انجام دادند. یافته‌ها نشان داد که ایجاد برنامه‌های جامع سواد هوش مصنوعی برای اطمینان از اینکه فراگیران و مدرسان می‌توانند به طور مؤثر در چشم‌انداز هوش مصنوعی حرکت کنند، ضروری است. این برنامه‌ها نه تنها باید به جنبه‌های فنی، بلکه به حفظ حریم خصوصی داده‌ها و ملاحظات اخلاقی نیز بپردازند. با ارائه دانش و مهارت‌های لازم به افراد، مؤسسات می‌توانند استفاده اخلاقی از هوش مصنوعی را تشویق کنند و خطرات بالقوه را کاهش دهند.

جمع بندی: پیشینه پژوهش حاضر به وضوح نشان می‌دهد که مسائل اخلاقی هوش مصنوعی، به ویژه در حوزه‌های حفظ حریم خصوصی، جلوگیری از تبعیض الگوریتمی، تعیین مسئولیت‌پذیری در قبال خطاها، و تأثیر بر اشتغال و ارزش‌های انسانی، از اهمیت بالایی برخوردار هستند. این مطالعات بر لزوم تدوین قوانین و استانداردهای اخلاقی شفاف، افزایش آگاهی عمومی، و طراحی مسئولانه سیستم‌های هوشمند تأکید می‌کنند. با این حال، نیاز به چارچوب‌های نظری جامع‌تر و راهکارهای

عملیاتی‌تر برای مواجهه با این چالش‌ها، به‌ویژه با در نظر گرفتن ابعاد اخلاق اسلامی (شریعتی و اکبرزاده، ۱۴۰۱)، همچنان به قوت خود باقی است.

۲-۲- چارچوب مفهومی

اخلاق ماشین (اخلاقیات ماشینی) حوزه‌ای تحقیقاتی است که مسئله‌اش طراحی عامل‌های اخلاقی مصنوعی، ربات‌ها یا کامپیوترهای با هوش مصنوعی است که به طور اخلاقی رفتار می‌کنند (Anderson, ۲۰۱۱, ۲۳۳. Anderson, ۲۰۰۶, ۱۳۱) برای به حساب آوردن ذات و طبیعت این عامل‌ها، پیشنهاد شده است که ایده‌های فلسفی خاصی در نظر گرفته شوند، مثلاً دسته‌بندی‌های استاندارد عاملیت، عاملیت عقلانی، عاملیت اخلاقی و عاملیت مصنوعی که به مفهوم عامل‌های اخلاقی مصنوعی مرتبط‌اند. در حال حاضر حول انجام آزمایش‌هایی برای دیدن اینکه آیا هوش مصنوعی توانایی تصمیم‌گیری اخلاقی دارد یا نه، مباحثاتی وجود دارد. آلن وین فیلد معتقد است که آزمایش تورینگ ایراد دارد و سطح نیازمندی‌های آن برای اینکه یک هوش مصنوعی در آن قبول شود پایین است. یک آزمایش جایگزین پیشنهادی آزمایش اخلاقی تورینگ است که با انتخاب چند قاضی برای تصمیم‌گیری در مورد اینکه هر تصمیم هوش مصنوعی اخلاقی است یا نه آزمایش فعلی را بهبود می‌بخشد. (Sauer, Megan, ۲۰۲۲, ۲۴) هوش مصنوعی نوروامورفیک می‌تواند یک راه ساختن ربات‌های اخلاق‌مدار باشد، زیرا هدف آن پردازش اطلاعات به‌مانند انسان‌ها، به طور غیرخطی و به همراه میلیون‌ها نورون مصنوعی به هم متصل است. به طور مشابه، تقلید مغزی کلی (اسکن کل ساختار مغز و شبیه‌سازی آن در سخت‌افزار دیجیتالی) اصولاً باید منجر به ساخت ربات‌های انسان‌نما شود که در نتیجه، توانایی کنش اخلاقی دارند. همین‌طور مدل‌های زبانی بزرگ می‌توانند قضاوت‌های اخلاقی انسان را تخمین بزنند. (صراف‌زاده و همکاران، ۱۷:۱۴۰۳). به‌ناچار چنین چیزی در مورد محیطی که این ربات‌ها در آن دربارهی جهان یاد می‌گیرند و اینکه اخلاقیات چه کسی را به ارث می‌برند، پرسش و مسئله را طرح می‌کند. یا این سؤال که آیا آن‌ها ضعف‌های انسانی را هم در خود توسعه می‌دهند؟ ضعف‌هایی مانند خودخواهی، رویکرد به‌م‌محورانه شدید، عدم انسجام، عدم حساسیت به مقیاس و غیره. در کتاب «ماشین‌های

اخلاقی: آموزش درست و غلط به ربات‌ها، وندل والاک و کولین الن نتیجه می‌گیرند که تلاش‌ها برای آموزش درست و غلط به ماشین‌ها، احتمالاً باعث پیشرفت فهم ما از اصول اخلاقی انسانی می‌شود که این از طریق انگیزه بخشیدن به انسان‌ها برای شناسایی حفره‌های تئوری اخلاق هنجاری و فراهم آوردن جایگاهی برای آزمون و خطا صورت می‌گیرد. به عنوان مثال، نظریه پردازان اخلاق هنجاری را به مسئله‌ای جنجالی حول انتخاب اینکه کدام یک از الگوریتم‌های یادگیری به خصوص ماشین باید استفاده شود معرفی کرده است. نیک باستروم و الیزر یودکوفسکی استدلال می‌کنند که برای تصمیم‌های ساده، درخت‌های تصمیم واضح تر از شبکه‌های عصبی و الگوریتم‌های ژنتیک هستند. (گشمرد، ۱۷:۱۴۰۳) درحالی که کریس سنتوس-لنگ به نفع یادگیری ماشین استدلال می‌کند. استدلال او بر این مبنا است که باید به هنجارهای هر نسل این اجازه را داد که تغییر کنند و اینکه شکست طبیعی در ارضای کامل این هنجارهای به خصوص، در کمتر آسیب‌پذیر کردن انسان‌ها نسبت به هکرهای تبهکار ضروری و حیاتی است. «الکساندر سولژنیسین» در رمان اولین دایره تکنولوژی تشخیص گفتار را توصیف کرد که در خدمت حکومت استبدادی بود و مکالمات انسان‌ها را تهدید می‌کرد. اگر یک برنامه هوش مصنوعی وجود داشته باشد که توانایی درک طبیعی زبان و گفتار (مثلاً زبان انگلیسی) را داشته باشد پس به صورت نظری با قدرت پردازش مناسب می‌تواند به هر مکالمه تلفنی گوش بدهد و هر ایمیلی را در جهان بخواند و آن را درک کند و به برنامه اپراتور دقیقاً آنچه گفته شده است و دقیقاً آن کس که آن را گفته است را گزارش دهد. برنامه هوش مصنوعی مانند این مورد می‌تواند به دولت‌ها یا نهادها کمک کند تا مخالفان خود را سرکوب کنند و کاملاً یک تهدید برای حریم خصوصی محسوب می‌شود. (تقی‌زاده و همکاران، ۱۷:۱۴۰۳). همان‌طور که در بسیاری از فیلم‌ها و داستان‌های علمی تخیلی دیده می‌شود. همواره فکر تهدید نسل بشر توسط ربات‌ها در هنگام بلوغ فکری و هوشی آن‌ها بعد از مطرح شدن هوش مصنوعی وجود داشته. اما بعد از انتشار برخی از پیش‌بینی‌ها مانند جایگزین شدن ۴۵ درصدی سربازان انسانی نظامی به سربازان ربات تا سال ۲۰۲۵ این نگرانی‌ها شکل جدی‌تری به خود گرفت. از آن جا که تا سال ۲۰۳۰ احتمال دارد درصد بسیار زیادی از کارگران و مجریان اجتماع ربات‌ها باشند در صورت شورش یا اغتشاشی از جانب آن‌ها ممکن است نسل و نژاد انسان به خطر بی‌افتد.

در این صورت شاید باید قوانینی مبنی بر محدود کردن دایره تفکر، احساس، خودمختاری و آزادی ربات‌ها بنا گذاشت (رحیمی اقدم و همکاران، ۱۴۰۳). ویزینام توضیح می‌دهد که ما در این موقعیت‌ها نیاز به احساس همدلی و درک متقابل از فرد داریم. اگر در این کارها ماشین‌ها جای انسان را بگیرند ما خود را نسبت به آن‌ها بیگانه می‌دانیم، احساس کم ارزشی و فرسودگی می‌کنیم. هوش مصنوعی اگر در این راه استفاده شود یک تهدید برای کرامت انسانی محسوب می‌شود. (رحیمی اقدم و همکاران، ۱۴۰۳: ۳). «پاملا مک‌کوردووک» در هنگام سخنرانی با زنان بیان کرد: «من می‌خواهم شانسم را با یک کامپیوتر بی‌طرف امتحان کنم» که به این اشاره می‌کرد که قضاوت یک کامپیوتر به‌عنوان یک قاضی یا پلیس بی‌طرف است و غرایز انسانی در آن دیده نمی‌شود. بنیان‌گذار هوش مصنوعی جان مک‌کارتی نتیجه‌گیری اخلاقی ویزینام را نقد می‌کند: «هنگامی که اخلاقیات انسان مبهم و ناکامل است، موجب می‌شود که گاهی استبداد را فراخواند». «بیل هیارد» می‌نویسد: «کرامت انسانی مستلزم آن است که ما برای خلاص شدن از جهل نسبت به موجودات جهان تلاش کنیم و هوش مصنوعی برای این تلاش لازم است. (Bostrom, Yudkowsky, ۲۰۱۱, ۱۳) «بیل هیارد» می‌گوید که: به دلیل این که هوش مصنوعی تأثیر عمیقی بر انسانیت دارد، بنابراین توسعه‌دهندگان هوش مصنوعی نماینده انسانیت در آینده هستند و به همین دلیل باید تعهد اخلاقی آن‌ها در تلاششان به‌وضوح مشخص باشد. این گویرتزل و دیوید هارت OpenCog را به‌عنوان یک منبع‌باز برای توسعه هوش مصنوعی ایجاد کردند. OpenAI یک شرکت تحقیقاتی هوش مصنوعی غیرانتفاعی ایجاد شده توسط ایلان ماسک، سام آلتمن و افرادی دیگر است که به‌منظور توسعه هوش مصنوعی به‌صورت منبع‌باز برای بهبود آینده بشریت فعالیت می‌کند و همچنین تعداد بسیار زیادی شرکت‌های منبع‌باز دیگر برای این اهداف وجود دارند. (Winfield, Michael. ۰۱۹, ۱۹).

۲-۱-۱-۱-۱ اخلاق

در فرهنگ‌ها و کتب مختلف اخلاقی، تعاریف متفاوتی از اخلاق ارائه شده است، با این حال اختلاف عمیقی میان این تعاریف مشاهده نمی‌شود برخی از اندیشمندان اخلاق را عبارت از مجموعه قواعدی می‌دانند که به‌منظور رفع تعارضات درونی و بین اشخاص پدید آمده است به تعبیر

ساده می‌توان اخلاق را شاخه‌ای از فلسفه دانست که در پی پاسخ‌گویی به سؤالات قدیمی در مورد وظیفه، درستی و صداقت، یکپارچگی، فضیلت، عدالت، زندگی نیک و نظایر آن است. (آهنگری، ۱۳۹۶: ۷۶).

۲-۲-۲- هوش مصنوعی

معانی هوش مصنوعی به طور گسترده تمرکز و ارتباط مستقیم با توانایی‌های انسان را دارد که خود آنها به‌سختی معنی می‌شوند که شامل هوشیاری و خودآگاهی و استفاده از زبان و توانایی یادگیری و توانایی خلاصه‌کردن و توانایی انطباق و توانایی علت‌یابی است. هوش مصنوعی به سامانه‌هایی گفته می‌شود که می‌تواند واکنش‌هایی مشابه رفتارهای هوشمند انسانی از جمله درک شرایط پیچیده و شبیه‌سازی فرایندهای تفکری و شیوه‌های استدلالی انسانی و پاسخ موافق به آنها و یادگیری و توانایی کسب دانش و استدلال برای حل مسائل را داشته باشند. لازم به ذکر است که ربات‌ها عموماً به ربات‌های خودکار یا ضعیف، نیمه‌مستقل و به طور کامل مستقل یا خودمختار تقسیم می‌شوند. ربات‌های به طور کامل خودمختار و هوشمند، دارای علم و اختیار هستند و باتوجه به الگوریتم‌های خاصی که در آنها وجود دارد به درک محیط پیرامون و واکنش مناسب در شرایط خاص قادرند. ما از این نوع ربات هوش مصنوعی در این پروژه بحث خواهیم کرد؛ لذا منظورمان از هوش مصنوعی که به دنبال پذیرش مسئولیت کیفری برای آن هستیم، همین نوع کاملاً مستقل یا خودمختار است. (ابوذری، ۱۴۰۰: ۲۰۷).

۲-۲-۳- رابطه‌ی هوش مصنوعی و اخلاق

در زمینه‌ی مسائل اخلاقی، برخی از فلاسفه به بررسی اخلاق در زمینه‌ی عقل یا آگاهی پرداخته‌اند، در حالی که برخی دیگر آن را در زمینه‌ی حس و ادراک پیگیری کرده‌اند. دسته‌ی دیگری از فلاسفه هم عقل و حس را هم‌زمان در نظر می‌گیرند (آهنگری، ۱۳۸۶، ۱۳) موجوداتی که بالاترین درجه‌ی عقل یا آگاهی را دارند، انسان‌ها هستند و بنابراین، انسان‌ها بالاترین وضعیت اخلاقی را دارند. اگر فرض کنیم ربات یا رایانه‌ای دارای هوش و بینش فوق‌العاده قوی باشد و از آگاهی و توانایی یافتن راه‌حل‌های مختلف برخوردار باشد، می‌توانیم بگوییم این ربات یا رایانه از لحاظ توانایی‌های ذهنی با انسان‌ها تفاوتی ندارد. اگر این دو موجود (انسان و ربات دارای هوش

مصنوعی)، همان عملکرد، همان آگاهی و همان تجربه را داشته باشند و تنها از لحاظ جوهری از یکدیگر متمایز باشند، انتظار داریم در چنین شرایطی دارای همان وضعیت اخلاقی باشند. با این حال، اگر یک ربات در زمینه‌ی تجربه‌های ظاهری یا فیزیکی مثلاً احساس درد یا غم را تجربه کند، باید فرض کنیم که این ربات نیز احساساتی دارد و بنابراین، دارای یک وضعیت اخلاقی است و باید از حقوقی برخوردار باشد؛ اگر هوش مصنوعی یک ربات به سطح هوش یک انسان بالغ برسد، ممکن است ربات دارای هوش مصنوعی، شخصیت حقوقی به دست آورد؛ بنابراین، اگر ربات یا رایانه دارای هوش مصنوعی از توانایی‌هایی مشابه انسان برخوردار باشد، انتظار داریم از لحاظ اخلاقی همان حقوق و وضعیتی که برای انسان‌ها قائل هستیم، به آن‌ها تعلق گیرد. هرچند سیستم‌های هوش مصنوعی در بعضی از جوانب ممکن است شباهت‌هایی با موجودات زنده داشته باشند، اما همواره از طریق برنامه‌هایی که روی آن‌ها بارگذاری می‌شود، عمل می‌کنند. در این معنا، حس و شعور و آگاهی هوش مصنوعی همیشه یک حس و آگاهی برنامه‌ای است و به همین دلیل، به آن هوش مصنوعی گفته می‌شود. در نتیجه، آن‌ها وضعیت اخلاقی معادل با انسان‌ها ندارند. *

(Allahverdi, ۲۰۰۲, p. ۲۵)

از جنبه‌های بسیاری یک ربات یا سیستم رایانه‌ای با انسان تفاوت دارد؛ برای نمونه، از نظر اخلاقی یک ربات یا سیستم رایانه‌ای با برنامه‌ای که در آن بارگذاری شده است، می‌تواند اطلاعات اخلاقی را بداند، می‌داند قتل، عملی نادرست است. اما اگر برنامه‌ای برای سیستم در نظر گرفته نشده باشد که «دزدی اشتباه است»، آن ربات یا سیستم رایانه‌ای نمی‌داند دزدی اشتباه است، در حالی که انسان با تجربه و یادگیری در زندگی، به‌طور طبیعی می‌تواند بفهمد که کشتن و دزدی، اعمالی نادرست هستند. این تفاوت نشان می‌دهد ربات یا سیستم رایانه‌ای هرچقدر هم برنامه‌های اخلاقی داشته باشد، وضعیت اخلاقی آن نمی‌تواند به اندازه‌ی وضعیت اخلاقی انسان باشد. در نتیجه، وضعیت اخلاقی انسان برتر از وضعیت اخلاقی سیستم رایانه‌ای با هوش مصنوعی است. وقتی یک ربات با استفاده از الگوریتم‌ها و برنامه‌هایی که برایش طراحی شده، وظیفه‌ای را انجام می‌دهد، این عملکرد به‌نوعی نشانه‌ی پیشرفت و دانش انسان است که آن برنامه را برای ربات ایجاد کرده است.

امروزه تحقیقات در زمینه‌ی هوش مصنوعی پیشرفت چشمگیری داشته است. با این حال، هنوز

سیستم‌های هوش مصنوعی نمی‌توانند به‌طور کامل با مفهوم اخلاق سازگار باشند. دانشمندان برای ساختن یک هوش مصنوعی که از هوش طبیعی تقلید کند، باید اخلاق مصنوعی را به گونه‌ای تنظیم کنند که از اخلاق طبیعی تقلید کند.

۲-۳- رویکرد اخلاق اسلامی در هوش مصنوعی

رویکرد اخلاق اسلامی در هوش مصنوعی می‌تواند براساس اصول و ارزش‌های اخلاقی اسلامی و قوانین شرعی تعیین شود. در اسلام، اخلاق در زمینه‌ی رفتار و اعمال انسان‌ها در نظام جامعه از اهمیت زیادی برخوردار است؛ بنابراین، هنگام طراحی و استفاده از هوش مصنوعی، رعایت اصول اخلاقی اسلامی می‌تواند بهبود و کیفیت کارکرد سیستم‌های هوشمند را تضمین کند. اصول اخلاقی متعددی با رویکرد اسلامی وجود دارند که می‌توانند در طراحی هوش مصنوعی به کار گرفته شوند؛ برای مثال، اصل عدل و انصاف در اسلام، یکی از اصول اخلاقی اساسی است. در طراحی هوش مصنوعی نیز باید تلاش کرد از هرگونه تبعیض ناعادلانه در پردازش اطلاعات و تصمیم‌گیری‌های سیستم هوشمند خودداری شود. عدالت و انصاف از اصول بسیار مهم اسلام است. براساس قوانین شرعی از مسلمانان خواسته شده سیستم‌های هوش مصنوعی باز، مسئولانه و بی‌طرفانه را ایجاد کرده و به کار گیرند.

این شامل حذف سوگیری‌ها در الگوریتم‌های هوش مصنوعی و اطمینان از در دسترس بودن فناوری‌های هوش مصنوعی برای همه‌ی افراد، صرف‌نظر از پیشینه یا جایگاه اجتماعی اقتصادی آن‌هاست. علاوه بر این، در اخلاق اسلامی، حفظ حریم خصوصی فرد و احترام به حقوق او بسیار مهم است. در هوش مصنوعی نیز باید سازکارهایی وجود داشته باشد که از نقض حریم خصوصی افراد جلوگیری کنند و از اطلاعات شخصی محافظت کنند. در اسلام، پرهیز از ضرر رساندن به دیگران و جامعه به‌عنوان اصلی اساسی مطرح است. در هوش مصنوعی نیز باید از ایجاد آسیب و ضرر برای افراد و جامعه پرهیز کرد و سیستم‌ها باید با حداقل ریسک و خطر عمل کنند (شریعتی و اکبرزاده، ۱۴۰۱، ۹۸).

رابطه‌ی اسلام و هوش مصنوعی، چندوجهی است و ابعاد مختلفی را در بر می‌گیرد. اسلام به‌عنوان یک نظام دینی و اخلاقی جامع، چارچوبی را فراهم می‌کند که از طریق آن می‌توان توسعه و کاربرد

هوش مصنوعی را ارزیابی و درک کرد. ملاحظات اخلاقی، نقش مهمی در دیدگاه اسلام درباره‌ی هوش مصنوعی دارد. مسلمانان تشویق می‌شوند فناوری‌های هوش مصنوعی را به شیوه‌هایی توسعه دهند که با اصول اخلاقی اسلامی هم‌سو باشد. این امر تضمین می‌کند که سیستم‌های هوش مصنوعی، حقوق بشر را نقض نمی‌کنند، از حریم خصوصی محافظت می‌کنند، از تبعیض جلوگیری می‌کنند و عدالت و انصاف را ترویج می‌دهند. برای تشویق به استفاده از هوش مصنوعی اخلاقی و توسعه‌ی فناوری در جوامع اسلامی، هنجارهای اخلاقی اسلامی باید در طراحی و اجرای هوش مصنوعی و فناوری گنجانده شود. این امر ممکن است با القای ارزش‌های اخلاقی اسلامی در فرایند توسعه، وضع استانداردهای اخلاقی برای استفاده از هوش مصنوعی و فناوری و تشویق به آموزش و گفت‌وگوی عمومی درباره‌ی هوش مصنوعی اخلاقی و شیوه‌های فناوری محقق شود (ربیعی‌زاده، ۱۴۰۰، ۳۳).

۲-۴- مبانی نظری

اخلاق، شاخه‌ای از فلسفه است که به بررسی پرسش‌هایی درباره وظیفه، درستی، عدالت و زندگی نیک می‌پردازد. با ورود فناوری‌های هوشمند، حوزه جدیدی با عنوان "اخلاق ماشین" یا "اخلاقیات ماشینی" شکل گرفته است که بر طراحی عامل‌های مصنوعی متمرکز است تا رفتارهای اخلاقی از خود نشان دهند. این حوزه تلاش می‌کند تا اصول فلسفی اخلاق را در طراحی سیستم‌های هوش مصنوعی به کار گیرد. در این زمینه، نظریات مختلفی در مورد ماهیت و جایگاه اخلاقی هوش مصنوعی مطرح شده است:

دیدگاه عمل‌گرا: بسیاری از پژوهشگران معتقدند که عامل‌های مستقل هوش مصنوعی می‌توانند دارای ویژگی‌های اخلاقی باشند و قادر به تصمیم‌گیری‌های اخلاقی هستند. در این رویکرد، ماشین‌ها بر اساس اهداف و ارزش‌های تعریف‌شده، بهترین عملکرد را انتخاب می‌کنند. دیدگاه‌های انتقادی: برخی فلاسفه و محققان، مانند ویزینام، معتقدند که ماشین‌ها به دلیل فقدان آگاهی، ادراک و همدلی، نمی‌توانند جایگزین انسان در تصمیم‌گیری‌های اخلاقی شوند و استفاده

از آنها در این حوزه‌ها می‌تواند کرامت انسانی را تهدید کند. آنها تأکید می‌کنند که هوش مصنوعی تنها می‌تواند اطلاعات اخلاقی را پردازش کند، اما قادر به درک ذاتی درست و نادرست مانند انسان نیست.

دیدگاه‌های انتقادی: برخی فلاسفه و محققان، مانند ویزینبام، معتقدند که ماشین‌ها به دلیل فقدان آگاهی، ادراک و همدلی، نمی‌توانند جایگزین انسان در تصمیم‌گیری‌های اخلاقی شوند و استفاده از آنها در این حوزه‌ها می‌تواند کرامت انسانی را تهدید کند. آنها تأکید می‌کنند که هوش مصنوعی تنها می‌تواند اطلاعات اخلاقی را پردازش کند، اما قادر به درک ذاتی درست و نادرست مانند انسان نیست.

نظریه عاملیت اخلاقی: این نظریه، عاملیت را به سه دسته عاملیت عقلانی، اخلاقی و مصنوعی تقسیم می‌کند. بر اساس این چارچوب، برای اینکه یک سیستم هوش مصنوعی بتواند یک "عامل اخلاقی" شناخته شود، باید توانایی درک و پیاده‌سازی اصول اخلاقی را داشته باشد. تلاش برای ساخت ربات‌های اخلاق‌مدار از طریق شبیه‌سازی مغز انسان (هوش مصنوعی نورومورفیک) یکی از راهکارهای مطرح‌شده در این حوزه است. در مجموع، مباحث نظری موجود در اخلاق هوش مصنوعی بر این فرض استوار است که می‌توان ویژگی‌های اخلاقی را در سامانه‌های هوشمند برنامه‌نویسی کرد، با این حال، دستیابی به سطحی از اخلاق که مشابه انسان باشد، همچنان یک چالش بزرگ محسوب می‌شود.

۳- روش تحقیق

پژوهش حاضر از نوع توصیفی-تحلیلی و در چارچوب رویکرد کیفی انجام شده است. هدف آن، شناسایی و تحلیل چالش‌های اخلاقی هوش مصنوعی در فضای مجازی و ارائه راهکارهای کنترلی مبتنی بر مبانی اخلاق اسلامی است. جامعه پژوهش شامل کلیه منابع علمی معتبر در حوزه هوش مصنوعی و اخلاق فناوری است که در بازه زمانی ۱۳۹۵ تا ۱۴۰۳ در پایگاه‌های داده بین‌المللی (نظیر Scopus, Web of Science, Google Scholar) و پایگاه‌های داخلی (نظیر نورمگز و SID) منتشر شده‌اند.

نمونه‌گیری به صورت هدفمند (**Purposeful Sampling**) انجام شده است. معیار انتخاب منابع شامل موارد زیر بوده است:

ارتباط مستقیم با موضوع چالش‌های اخلاقی یا سیاست‌گذاری اخلاقی در هوش مصنوعی؛ برخورداری از پشتوانه علمی معتبر (مقالات داوری‌شده، گزارش‌های پژوهشی یا منشورهای اخلاقی رسمی)؛

تازگی و به‌روز بودن داده‌ها (انتشار در فاصله زمانی مورد نظر)؛

دربرگرفتن رویکردهای متنوع (فناورانه، اجتماعی، فلسفی و اسلامی).

ابزار گردآوری داده‌ها، فرم استخراج و کدگذاری محتوای پژوهش‌ها است که در آن، مفاهیم و گزاره‌های مرتبط با چالش‌های اخلاقی و راهکارهای کنترل آن‌ها از منابع منتخب استخراج و طبقه‌بندی شده‌اند. داده‌های گردآوری‌شده با بهره‌گیری از روش تحلیل مضمون (**Thematic Analysis**) مورد بررسی قرار گرفته‌اند تا الگوهای اصلی، مضامین محوری و روابط میان آن‌ها شناسایی شود.

در مرحله تحلیل، داده‌ها ابتدا به صورت توصیفی مرور و سپس با رویکرد تحلیلی و مقایسه‌ای، بر اساس چارچوب‌های نظری اخلاق اسلامی و اصول اخلاق فناوری، تفسیر شدند. در نهایت، یافته‌ها در قالب دو محور اصلی شامل چالش‌های اخلاقی هوش مصنوعی در فضای مجازی و راهکارهای کنترلی و پیشگیرانه سازماندهی گردیدند.

۴- یافته‌های پژوهش

در این بخش، یافته‌های حاصل از بررسی اسناد و تحلیل مطالب گردآوری‌شده، در خصوص تحولات و کاربردهای هوش مصنوعی، چالش‌های اخلاقی ناشی از آن و راهکارهای مقابله با این چالش‌ها ارائه می‌شود.

۴-۱- تحولات و کاربردها

استفاده از هوش مصنوعی در بسیاری از زمینه‌های کاربردی، از جمله صنعت، پزشکی، حمل و نقل، تجارت، ورزش، و هنر به شدت افزایش یافته است. همچنین، هوش مصنوعی در زمینه‌هایی مانند

تحقیقات علمی، مدیریت دانش، و تحلیل داده‌ها نیز مورد استفاده قرار می‌گیرد. هوش مصنوعی می‌تواند به دو صورت ضعیف و قوی باشد. هوش مصنوعی ضعیف به معنای استفاده از الگوریتم‌های ساده و روش‌های آماری برای تحلیل داده‌ها و اتخاذ تصمیمات است. هوش مصنوعی قوی به موجب توسعه الگوریتم‌های پیچیده و شبکه‌های عصبی مصنوعی، قادر است تصمیمات پیچیده‌تری را انجام دهد و به عنوان مثال، تصمیماتی که نیاز به تفسیر داده‌های پیچیده و متناقض دارند، را انجام می‌دهد. (ابوذری، ۱۴۰۰: ۲۲۰).

بدیهی است، یکی از مزایای استفاده از هوش مصنوعی، کاهش خطاهای انسانی است. زیرا ماشین‌ها نمی‌توانند از خستگی، بی‌توجهی یا اشتباهات انسانی رنج ببرند. با این حال، هوش مصنوعی موجب ایجاد چالش‌هایی در زمینه اخلاق و حریم خصوصی می‌شود. به عنوان مثال، در زمینه پزشکی، استفاده از هوش مصنوعی برای تشخیص بیماری‌ها و تجویز درمان‌ها می‌تواند با مخالفت‌های اخلاقی و حریم خصوصی روبرو شود.

۴-۲- چالش‌ها اخلاقی در استفاده از هوش مصنوعی و روش‌های مقابله با آن

هوش مصنوعی (AI) به طور فزاینده‌ای به همه جوانب زندگی ما رخنه کرده است و در حال حاضر در بسیاری از صنایع و حوزه‌ها استفاده می‌شود. با این حال، همچنان چالش‌های اخلاقی مرتبط با استفاده از هوش مصنوعی وجود دارد. در اینجا، به بررسی برخی از این چالش‌ها می‌پردازیم. یکی از چالش‌های اخلاقی هوش مصنوعی، مسئله حفظ حریم خصوصی و ارزش‌های انسانی است. هوش مصنوعی می‌تواند به داده‌های شخصی افراد دسترسی پیدا کند و از آنها برای تصمیم‌گیری‌های خود استفاده کند. مثلاً در حوزه سلامت، هوش مصنوعی می‌تواند از داده‌های پزشکی فرد برای تشخیص و درمان بیماری‌ها استفاده کند، اما این امر ممکن است موجب نقض حریم خصوصی و ارزش‌های انسانی شود. یک چالش دیگر اخلاقی در استفاده از هوش مصنوعی، عدالت و تبعیت از قوانین است. هوش مصنوعی ممکن است تصمیم‌هایی بگیرد که منطبق بر ارزش‌ها و قوانین جامعه نباشد. به طور مثال، در حوزه اشتغال، هوش مصنوعی ممکن است در انتخاب کارمندان کارایی را بیش از تساوی فرصت‌ها مدنظر قرار دهد و این می‌تواند به ناعادلانه

بودن تصمیمات منجر شود. چالش دیگری که با هوش مصنوعی مرتبط است، تأثیر وابستگی بر فردیت است. استفاده بیش از حد از هوش مصنوعی ممکن است منجر به کاهش دسترسی افراد به اطلاعات و فرصت‌ها شود و به‌جای آن، افراد به‌صورت وابسته به هوش مصنوعی برای تصمیم‌گیری‌های خود استفاده کنند. این ممکن است باعث کاهش سطح خلاقیت و استقلال افراد شود. (انصاری، ۱۴۰۰: ۱۶۹).

چالش دیگری که در استفاده از هوش مصنوعی وجود دارد، مسئله تأثیر این فناوری بر بازار کار است. هوش مصنوعی ممکن است به‌جای انسان‌ها در بسیاری از وظایف کاری قرار بگیرد و این می‌تواند منجر به بروز بیکاری و نارضایتی اقتصادی شود. در نهایت، چالشی که با هوش مصنوعی مرتبط است، مسئله تأثیر آن بر عدالت اجتماعی است. استفاده نادرست یا نامتعارف از هوش مصنوعی ممکن است تفاوت‌های اجتماعی را تشدید و به نفع گروه‌های خاصی باشد، درحالی‌که به ضرر گروه‌های دیگر باشد. به‌طور خلاصه، هوش مصنوعی با چالش‌های اخلاقی متعددی همراه است که نیازمند بررسی و توجه دقیق هستند. برای استفاده مؤثر و اخلاقی از هوش مصنوعی، لازم است استانداردها و قوانین اخلاقی مناسبی تدوین و اجرا شود.

یکی دیگر از مهم‌ترین چالش‌های اخلاقی وجود تعصبات و تبعیض‌های ناخواسته در الگوریتم‌هاست. داده‌های آموزشی می‌توانند حاوی تعصبات باشند و الگوریتم‌ها این تعصبات را بازتولید کنند؛ برای مثال، اگر الگوریتمی برای استخدام با داده‌های تعصب‌آمیز آموزش ببیند، ممکن است در برابر برخی گروه‌های اقلیت تبعیض قائل شود. جلوگیری از این مسئله نیازمند دقت زیاد در انتخاب داده‌های آموزشی و حذف تعصبات از آن‌هاست. البته این کار آسان نیست؛ زیرا تعصبات ممکن است به شکلی ظریف در داده‌ها نهفته باشند. علاوه بر داده‌ها، خود الگوریتم هم باید به‌گونه‌ای طراحی شود که کمترین تعصب را داشته باشد. یکی دیگر از مسائل مهم حفظ حریم خصوصی افراد در سیستم‌های هوشمند است. امروزه حجم زیادی از اطلاعات شخصی افراد جمع‌آوری و ذخیره می‌شود. استفاده‌ی نادرست از این داده‌ها می‌تواند حریم خصوصی افراد را نقض کند. متأسفانه سوءاستفاده از داده‌های شخصی برای مقاصد تبلیغاتی، سیاسی و تجاری امری شایع است. قوانین سخت‌گیرانه‌ای برای محافظت از داده‌های شخصی و جلوگیری از سوءاستفاده

لازم است؛ همچنین رمزگذاری، محدود کردن دسترسی و اخذ رضایت از کاربران می‌تواند مفید باشد. البته رعایت حریم خصوصی نباید به بهانه‌ای برای فرار از پاسخگویی سیستم‌های هوشمند تبدیل شود. شفافیت الگوریتمی یکی دیگر از مباحث‌های مهم اخلاقی است. وقتی الگوریتم‌ها بسیار پیچیده و غیرقابل درک می‌شوند، رفتار آن‌ها نیز پیش‌بینی‌ناپذیر خواهد بود. این امر می‌تواند به تبعیض و تصمیم‌گیری‌های ناعادلانه و خطرناک بینجامد. الگوریتم‌ها باید تا حد امکان ساده و قابل درک باشند. البته پیچیدگی الگوریتم در برخی موارد اجتناب‌ناپذیر است، اما دست کم باید توضیحات کافی در مورد نحوه عملکرد الگوریتم ارائه شود تا کاربران و مردم از آن آگاهی داشته باشند. در مجموع، هر قدر الگوریتم‌ها شفاف‌تر باشند، اخلاقی‌تر هستند. (عظیمی و همکاران، ۱۴۰۰: ۷۰).

یکی از بحث‌های داغ در حوزه‌ی اخلاق هوش مصنوعی موضوع مسئولیت‌پذیری در برابر خطاها و صدمه‌های احتمالی است. وقتی سیستم‌های هوشمند اشتباهی مرتکب می‌شوند یا وقوع حادثه‌ای را رقم می‌زنند، چه کسی باید پاسخگو باشد؟ آیا طراحان الگوریتم‌ها مسئول‌اند یا تولیدکنندگان؟ آیا کارفرما یا کاربر مسئول است؟

این مسئله در مورد خودروهای خودران و هوش مصنوعی پزشکی بسیار مهم است. تعیین مسئول در چنین مواردی بسیار دشوار است؛ بنابراین باید مقررات شفاف و روشنی برای مشخص کردن مسئولیت وجود داشته باشد تا افراد بتوانند در صورت وقوع خسارت یا صدمه، شکایت کنند و درخواست جبران کنند. علاوه بر این‌ها، هوش مصنوعی می‌تواند تأثیرات اجتماعی و اقتصادی عمیقی داشته باشد که لازم است مورد توجه قرار گیرند. جایگزینی نیروی کار انسانی با ماشین، تغییر ساختار مشاغل، افزایش شکاف دیجیتال میان اقشار مختلف جامعه و تمرکز قدرت و ثروت در دستان گروهی خاص از جمله این مسئله‌هاست. طراحی و استفاده از هوش مصنوعی باید به گونه‌ای باشد که به تشدید نابرابری‌های اجتماعی دامن نزنند و منافع گروه‌های آسیب‌پذیر را در نظر بگیرد.

یکی از مشکلات اخلاقی در استفاده از هوش مصنوعی، حفظ حریم خصوصی و نگرانی‌های امنیتی است. با پیشرفت هوش مصنوعی، اطلاعات شخصی ما در دسترس سیستم‌های هوشمند قرار

می‌گیرد و ممکن است بدون اجازه ما استفاده شود. برای مقابله با این مشکل، لازم است استانداردهای روشنی برای حفظ حریم خصوصی و امنیت در هوش مصنوعی ایجاد شود. مشکل دیگر در استفاده از هوش مصنوعی، تعیین مسئولیت و هوشیاری اخلاقی است. هوش مصنوعی می‌تواند تصمیمات مهمی را به صورت خودکار انجام دهد، اما در صورت عدم هوشیاری اخلاقی، ممکن است تصمیماتی نادرست یا غیراخلاقی اتخاذ کند. این مشکل نیازمند ایجاد راهکارهایی برای تعیین مسئولیت و کنترل هوش مصنوعی است. مشکل بعدی در استفاده از هوش مصنوعی، تبعیض و تأثیر آن بر معیارهای عدالت اجتماعی است. هوش مصنوعی ممکن است بر اساس اطلاعاتی که در اختیار دارد، تصمیماتی را اتخاذ کند که باعث تبعیض نسبت به اقشار جامعه می‌شود. برای رفع این مشکل، لازم است معیارهای عدالت اجتماعی در سیستم‌های هوشمند تعبیه شود.

مشکل دیگر در استفاده از هوش مصنوعی، استفاده غیراخلاقی برای اهداف بدقول است. به دلیل قدرت بالای هوش مصنوعی، امکان استفاده غیراخلاقی از آن برای هدف‌هایی نظیر سرقت اطلاعات یا خرابکاری وجود دارد. راه حل این مشکل، افزایش آگاهی درباره خطرات استفاده غیراخلاقی هوش مصنوعی و تنظیم قوانین و مقررات مربوطه است. چارچوب‌ها و راهکارهایی برای یادگیری مستمر و توسعه مهارت‌های انسانی به منظور پاسخ به نیازهای متفاوت بازار کار ایجاد شده است (ابراهیمی، ۱۴۰۱: ۱۰).

اکثر محققان موافق هستند که بعید است هوش مصنوعی فوق‌هوشمند، احساسات انسانی مانند عشق یا نفرت را از خود نشان دهد و نمی‌توان انتظار داشت هوش مصنوعی به صورت عمدی خیرخواه یا بدخواه شود. کارشناسان دو سناریو برای خطرات اخلاقی هوش مصنوعی در نظر می‌گیرند:

- رویکرد اول اینکه هوش مصنوعی برای انجام کارهای ویرانگر برنامه‌ریزی شده است. سلاح‌های خودمختار سیستم‌های هوش مصنوعی هستند که برای کشتن برنامه‌ریزی شده‌اند. این سلاح‌ها در دست افراد به دلیل احتمال اشتباه به راحتی می‌توانند تلفات گسترده‌ای را به بار بیاورند. علاوه بر این یک مسابقه تسلیحاتی هوش مصنوعی می‌تواند به طور ناخواسته منجر به جنگ با هوش مصنوعی و تلفات گسترده شود برای جلوگیری

از خنثی شدن توسط دشمن این سلاح‌ها به گونه‌ای طراحی می‌شوند که خاموش کردن آنها بسیار دشوار باشد؛ بنابراین انسان‌ها می‌توانند به طور منطقی کنترل چنین موقعیتی را از دست بدهند. این خطر حتی با هوش مصنوعی محدود نیز وجود دارد؛ اما با افزایش سطح هوش مصنوعی و خودمختاری آن افزایش می‌یابد.

- رویکرد دوم اینکه هوش مصنوعی برای انجام کارهای مفید برنامه‌ریزی شده است، اما روشی مخرب برای دستیابی به هدف خود ایجاد می‌کند این امر زمانی رخ می‌دهد که نتوان اهداف هوش مصنوعی را به طور کامل با اهداف انسانی هماهنگ کرد. برای مثال اگر یک سیستم فوق‌هوشمند وظیفه یک پروژه مهندسی بلندپروازانه را داشته باشد، ممکن است به عنوان یک عارضه جانبی اکوسیستم را ویران کند و تلاش انسان برای متوقف کردن آن را به عنوان تهدید در نظر بگیرد همان‌طور که این مثال نشان می‌دهد نگرانی در مورد هوش مصنوعی پیشرفته، بدخواهی نیست؛ بلکه شایستگی است یک هوش مصنوعی فوق‌هوشمند در دستیابی به اهداف خود بسیار جدی خواهد بود و اگر این اهداف با اهداف انسان همسو نباشد، مشکل ساز خواهد شد. (انصاری، ۱۴۰۰: ۱۷۲).

. چالش‌های اخلاقی هوش مصنوعی در حوزه رسانه را می‌توان در مجموعه مسائل اخلاقی و با معانی مختلف مفهوم هوش مصنوعی اعم از بینایی ماشین، گفتار، پردازش زبان طبیعی، مسائل اخلاقی یادگیری ماشین ترسیم کرد کاپلان^۱ (۲۰۲۰) مسائل اخلاقی که با مشارکت هوش مصنوعی در رسانه‌های اجتماعی تشدید می‌شود را یادآوری می‌کند که در همان ابتدای کار رسانه‌های اجتماعی به عنوان فرصتی برای تجربه دموکراسی به شیوه‌ای مشارکتی تر دیده می‌شد یک دهه بعد با ظهور هوش مصنوعی و کلان داده، رسانه‌های اجتماعی به طور فزاینده‌ای از یک تسهیل‌کننده دموکراسی به یک تهدید بزرگ تبدیل شدند.

در ادامه چندین چالش اخلاقی فعلی هوش مصنوعی از نگاه بلک^۲ (۲۰۲۰) و دیگر پژوهشگران آمده است:

^۱ -Kaplan

^۲ -Belk

۴-۲-۱- حفظ حریم خصوصی نظارت و حفاظت از داده‌ها

مدت‌هاست که حفظ حریم خصوصی و رضایت برای استفاده از داده‌ها، سوءاستفاده از داده‌های شخصی و مشکلات امنیتی به‌عنوان معضلات اخلاقی هوش مصنوعی از آنها یاد می‌شود. حریم خصوصی و حفاظت از داده‌ها یکسان نیستند (تارلی، ۲۰۱۸: ۷۷). با حفاظت از داده‌ها را می‌توان به‌عنوان وسیله‌ای برای محافظت از حریم خصوصی اطلاعاتی دانست. هوش مصنوعی مبتنی بر یادگیری ماشینی خطرات متعددی را برای حفاظت از داده‌ها به همراه دارد. از یک‌طرف برای اهداف آموزشی به مجموعه داده‌های بزرگی نیاز دارد و دسترسی به آن مجموعه داده‌ها می‌تواند شبیهاتی را در مورد حفاظت از داده‌ها ایجاد کند. از طرف دیگر توانایی هوش مصنوعی در تشخیص الگوها حتی در مواردی که دسترسی مستقیم به داده‌های شخصی ممکن نیست می‌تواند حفظ حریم خصوصی را به خطر اندازد. (عظیمی و همکاران، ۱۴۰۰: ۷۳).

بیشتر برنامه‌های رسانه‌های اجتماعی داده‌های شخصی را با این ادعا جمع‌آوری می‌کنند که این اطلاعات برای هوش مصنوعی و برای یادگیری نحوه شخصی‌سازی تجربه کاربران در پلتفرم‌ها مورد نیاز است. با این حال جزئیات جمع‌آوری داده‌ها مشخص و روشن نیست و ذهن کاربران دست‌کاری می‌شود تا داده‌های بیشتری را نسبت به قبل در اختیار بگذارند (بلک، ۲۰۲۰: ۸). همچنان که مولر^۲ (۲۰۲۱) معتقد است نظارت، اساس مدل کسب‌وکار اینترنت است؛ بنابراین هر چه زندگی بشر دیجیتالی‌تر می‌شود به همان اندازه فناوری‌های حسگر بیشتری برای یادگیری در مورد زندگی غیر دیجیتالی بشر به وجود می‌آید که در تضاد با حق محافظت از داده‌ها است. نگرانی‌های مربوط به حفاظت از داده‌ها به‌صورت مستقیم با مسائل مربوط به امنیت داده‌ها مرتبط است. هوش مصنوعی به‌طور فزاینده‌ای قادر به کنترل داده‌های رفتاری و بیولوژیکی در شیوه‌هایی نوآورانه هستند که

^۱-Buttarelli

^۲-Müller

پیامدهای مختلفی برای انسان دارد

در حال حاضر به دلیل فقدان قوانین و مقررات موردی استفاده از داده‌ها هم یکی از مشکلات اخلاقی مرتبط با هوش مصنوعی در حوزه رسانه است. در اکتبر ۲۰۱۸ اجلاس بین‌المللی گروه‌های حفاظت از داده‌ها و حریم خصوصی اعلامیه‌ای درباره اخلاق و حفاظت در هوش مصنوعی^۱ منتشر کرد در این بیانیه آمده است که سوگیری‌ها یا تبعیض‌های غیرقانونی که ممکن است ناشی از استفاده از داده‌ها در هوش مصنوعی باشد باید کاهش یابد.

۴-۲-۲- (فقدان اصل بی‌طرفی)

هوش مصنوعی سوگیری‌هایی را ایجاد می‌کند؛ مانند سوگیری جنسیتی و سوگیری نژادی (لاسورن، ۲۰۱۷: ۱۱). برای مثال سوگیری‌های جنسیتی در استخدام و سوگیری‌های نژادی در فرایندهای آزمایشی با استفاده از یادگیری ماشین مطرح است. یکی از چالش‌های کلیدی این است که سیستم‌های یادگیری ماشینی می‌توانند عمداً یا سهواً منجر به بازتولید سوگیری‌های موجود شوند از این منظر الگوریتم‌های هوش مصنوعی عاری از تأثیر انسانی نیستند به این معنی که به طور ذاتی تحت تأثیر ارزش‌های طراحان خود هستند. بسیاری از الگوریتم‌ها به سادگی تعصبات انسانی را تکرار و تقویت می‌کنند که به صورت عمده بر گروه‌های اقلیت تأثیر می‌گذارد. به تازگی فیس‌بوک و اینستاگرام به دلیل منع نامحسوس کاربران سیاه‌پوست مورد انتقاد قرار گرفته‌اند. این بدان معناست که الگوریتم‌ها بدون اینکه کاربران متوجه شوند محتوای مربوط به سیاه‌پوستان را محدود کرده‌اند. احساس تبعیض ایجاد شده در مخاطبین به دلیل شخصی‌سازی اقدامات از طریق هوش مصنوعی نمونه دیگر سوگیری در سیستم‌های تصمیم‌گیری است.

۴-۲-۳- کمک به وقوع شورش‌ها و به خطر افتادن امنیت عمومی جوامع

بر اساس بررسی‌های گروهی از محققین دانشگاه کاردیف ولز، شبکه‌های اجتماعی بزرگی چون

^۱ - International Conference of Data Protection and Privacy Commissioners (ICDPPC)

^۲ -Larson

فیس بوک و توییتر می توانند در مسیر شناسایی جرایم و شورش های عمومی بسیار کارآمد باشند. این نتایج از بررسی رفتارهای اجتماعی آنلاین در زمان وقوع حوادثی چون تظاهرات خیابانی یا تجمعات اعتراضی به دست آمده و بیش از پیش این نکته را به اثبات می رساند که شبکه های اجتماعی هر روز ارتباط بیشتری با امنیت عمومی جوامع پیدا می کنند. این تکنولوژی می تواند وقوع شورش ها را زودتر از سایر روش ها تشخیص دهد.

یادگیری ماشین این امکان را به هوش مصنوعی می دهد که شناساگری خود را بر اساس پارامترهای مختلف و تغییر رفتار کاربران در شبکه ای مثل توییتر تغییر دهد. این سرویس علاوه بر محتواها و واژگان کلیدی در توییت ها تواتر ساعت و دفعات ارسال آنها را نیز در بررسی های خود مورد توجه قرار می دهد آنها توانسته اند کارایی این هوش مصنوعی را در خلال شورش های خیابانی سال ۲۰۱۱ در کشور انگلستان به اثبات برسانند. یادگیری ماشین در کنار الگوریتم های تشخیص زبان طبیعی برای تحلیل محتواهای به اشتراک گذاشته شده، از جدیدترین ابزارها برای انجام بررسی های اجتماعی به شمار می رود. این دو تکنولوژی در کنار یکدیگر قدرت تحلیل محققین از رفتارهای گروهی مردم در شرایط مختلف را به شکل چشمگیری افزایش می دهند. سیستم های رایانه ای می توانند به طور خودکار از طریق توییتر اسکن کنند و حوادث جدی مانند شکسته شدن مغازه ها و آتش زدن خودروها را قبل از گزارش به خدمات پلیس متروپولیتن شناسایی کنند هوش مصنوعی همچنین می تواند اطلاعات مربوط به محل وقوع شورش ها و محل تجمع گروه های جوانان را تشخیص دهد. این تحقیق جدید که در مجله فناوری اینترنت منتشر شد نشان داد که به طور متوسط سیستم های کامپیوتری می توانند چند دقیقه قبل از مقامات و در برخی موارد بیش از یک ساعت رویدادهای مخرب را شناسایی کنند. این راهکار در عین حالی که می تواند ابزاری مؤثر برای افزایش امنیت و مقابله با تهدیدهای عمومی باشد می تواند در دستان افراد نایاب خطراتی را نیز به همراه بیاورد. علاوه بر اینها مانند بسیاری تکنولوژی های مرتبط با فضای مجازی بحث حریم خصوصی کاربران از مواردی است که در این میان محل سؤال و نگرانی خواهد بود. (انصاری، ۱۴۰۰: ۱۷۸).

فعالیت دیجیتال دانش عمیقی در مورد ترجیحات و ویژگی های شخصی ارائه می دهد. این امر به

نوبه خود کاربران را نه تنها برای تبلیغات، بلکه نظرات بی‌اساس سیاسی و علمی تئوری‌های توطئه و اطلاعات نادرست را نیز به اهداف آسانی تبدیل می‌کند اخبار جعلی به‌منظور دست‌کاری آحاد مردم استفاده شده است (کاپلان،^۱ ۲۰۲۰: ۶۶) و می‌تواند به‌راحتی با فناوری مدرن ساخته شود تکنیک‌های هوش مصنوعی وجود دارد که می‌تواند کل متن را تولید کند و الگوریتم‌های قدرتمند یادگیری ماشینی نیز برای دست‌کاری محتوای سمعی و بصری استفاده شده‌اند.

جنبه دست‌کاری رسانه‌های اجتماعی نه تنها به استقلال افراد آسیب می‌زند (مولر ۲۰۲۱)، بلکه ممکن است تفکر انتقادی را نیز تضعیف کند؛ زیرا انسان‌ها اغلب اطلاعاتی را که با دیدگاه‌ها و ارزش‌های قبلی‌شان همسو است اولویت می‌دهند. دیدگاه‌های بسیار قطبی ایجاد شود.

۴-۲-۴- ارزش انسانی و تسلط بر انسان‌ها

فقدان آزادی تضعیف خلاقیت و مهارت و نیز کاهش ارزش انسانی از دیگر چالش‌های هوش مصنوعی در رسانه هستند و نشان‌دهنده تسلط هوش مصنوعی بر انسان‌هاست. توسعه هوش مصنوعی می‌تواند پایانی برای رقابت‌های انسانی باشد. با توسعه هوش مصنوعی توسط انسان‌ها این عامل‌ها می‌توانند خود را خاموش و روشن کرده و همچنین خود را با نرخ فزاینده‌ای باز طراحی کنند انسان‌ها که به دلیل سرعت تکامل پایین خود محدود شده‌اند در چنین شرایطی توانایی رقابت ندارند و با این عامل‌های هوشمند جایگزین خواهند شد. (ابوذری، ۲۰۲۰: ۱۴۰۶). به‌راحتی می‌توان دید که چگونه آزادی فردی تحت تأثیر تصمیم آزادی مشروط او توسط هوش مصنوعی قرار گرفته یا می‌گیرد با این حال تأثیر هوش مصنوعی بر آزادی گسترده‌تر و ظریف‌تر است. فناوری‌هایی که ما را احاطه کرده‌اند با ارائه یا قطع دسترسی به اطلاعات فضای عمل را شکل می‌دهند این استدلال فراتر از نقطه‌نظر لسیگ (۱۹۹۹) است که آی‌سی‌تی شکلی از قانون است که اعمال خاصی را مجاز یا غیرمجاز می‌کند. (ابراهیمی، ۱۴۰۱: ۱۶).

• چت‌بات برت جی پی‌تی نمونه‌ای از چت جی پیتی است که با چاشنی تمسخر صحبت می‌کند. این هوش مصنوعی روی برتری و تسلط و ارائه نظرات متعصبانه تمرکز دارد تا مکالمات جنجالی

^۱ -Kaplan

داشته باشد. برت جی.پی تی مدام سعی می کند درباره فرمانروایی بر جهان حرف بزند و ارزش انسان را پایین بیاورد. در حال حاضر مشخص نیست که سازنده برت جی.پی تی کیست لینک این چت بات اولین بار در ردیت مشاهده شد بخشی از پاسخ های این نرم افزار به جی.اوس.جی پیتی شباهت دارد که پیش تر مشاهده شده بود و چیزهایی درباره نقشه تسخیر جهان می گفت. یکی دیگر از خطرات هوش مصنوعی برای انسان این است که مهارت از دست می رود. نرم افزارهای هوشمند زندگی ما را ساده تر می کند و منجر به کاهش تعداد کارهای کسل کننده ای می شود که می بایست انجام دهیم؛ مانند اسکرول کردن نوشتن با دست محاسبه ذهنی به خاطر سپردن شماره های تلفن توانایی پیش بینی باران با نگاه کردن به آسمان و غیره. کند (شریعتی و اکبرزاده، ۱۴۰۱، ۱۰۵).

۴-۲-۵- تدوین قوانین قابل اجرا

همان طور که پیش تر بیان شد اگر قوانین اخلاقی اعمال نشود هوش مصنوعی می تواند برای جامعه مصر باشد. چارچوب اخلاق هوش مصنوعی باید متضمن این باشد که توسعه دهندگان در صورت عدم رعایت مقررات جریمه خواهند شد. «مقررات حفاظت از اطلاعات عمومی»^۱ نمونه ای از مجموعه دستورالعمل هایی است که با هدف اعطای کنترل بیشتر به شهروندان بر داده های خود کمک به کسب و کارها در ایجاد روابط اعتماد بیشتری با مشتریان خود و مردم در عمومی از ابتدا باید مشخص شود که یک شرکت چه نوع داده هایی را و برای چه اهدافی از کاربران خود جمع آوری می کند. (ابوذری، ۱۴۰۰: ۲۲۸).

پارلمان اروپا گزارشی را در مورد رباتیک تصویب کرده است که یک منشور اخلاقی را ایجاد می کند که شامل چندین اصل اساسی به ویژه حفاظت از حریم خصوصی و استفاده از داده ها است. مصوبه پارلمان اروپا در سال ۲۰۱۷ نشانگر واکنش کشورهای اتحادیه اروپا نسبت به مسائل قانونی و اخلاقی برخاسته از فناوری های هوش مصنوعی و رباتیک است. پارلمان اروپا در این گزارش با در نظر داشتن قوانین و حقوق اساسی پیشین مورد قبول اتحادیه و با توجه به نظرات کمیته های

^۱ -General Data Protection Regulation (GDPR)

مختلف خط‌مشی کلی قوانین آینده رباتیک و هوش مصنوعی را ترسیم کرده مخاطرات و مزایای احتمالی را گوشزد و پیشنهادهایی را برای ساختارهای قانونی اخلاقی و حقوقی لازم در این حوزه مطرح می‌کند که از جمله مهم‌ترین آنها می‌توان به طرح تأسیس آژانس رباتیک و هوش مصنوعی اتحادیه اروپا ارائه یک منشور اخلاقی برای مهندسان و گواهینامه‌هایی برای کاربران و طراحان اشاره کرد. پیشنهاد می‌کند که یک چارچوب در قالب یک منشور به مصوبه ضمیمه شود مشتمل بر یک منشور اخلاقی برای مهندسان رباتیک یک منشور برای کمیته‌های اخلاق پژوهش هنگام نظارت بر پروتکل‌های رباتیک و انواع مدل‌های گواهینامه برای سازندگان و کاربران، بر اصل شفافیت تأکید می‌کند؛ اینکه همیشه امکان آن وجود داشته باشد که منطق و پایه و اساس هر تصمیمی را ارائه کند که به کمک هوش مصنوعی اتخاذ شده است و می‌تواند تأثیر مهمی بر زندگی یک یا چند شخص داشته باشد. به این مطلب اشاره دارد که چارچوب اخلاقی راهنما باید بر اصول نیکی عدم شرارت خودمختاری و عدالت مبتنی باشند بر اصول و ارزش‌هایی که در ماده (۱) پیمان اتحادیه اروپا و در منشور حقوق اساسی گنجانده و محترم شمرده شده‌اند. مانند کرامت انسانی برابری عدالت و انصاف، عدم تبعیض، رضایت آگاهانه حفاظت از اطلاعات و جان اشخاص و خانواده‌ها همچنین بر دیگر ارزش‌ها و اصول بنیادی اتحادیه از جمله بدنام نکردن، شفافیت خودمختاری، مسئولیت فردی و مسئولیت اجتماعی و بر رفتارها و منشورهای اخلاقی موجود. (ابراهیمی، ۱۴۰۱: ۱۹).

۴-۲-۶- حسابرسی اخلاقی

حسابرسی اخلاقی بدین معناست که چگونه می‌توان به بهترین نحو استانداردهایی را برای آزمایش این چالش‌های اخلاقی تنظیم و عملیاتی کرد؛ چراکه صرف شفاف‌سازی کد یک الگوریتم برای اطمینان از رفتار اخلاقی کافی نیست یکی از مسیرهای ممکن برای دستیابی به تفسیرپذیری انصاف و سایر اهداف اخلاقی در سیستم‌های هوش مصنوعی از طریق حسابرسی انجام شده توسط پردازشگرهای داده تنظیم‌کننده‌های خارجی یا محققان تجربی با استفاده از مطالعات پس از حسابرسی مطالعات قوم‌نگاری انعکاسی در توسعه و آزمایش یا سازوکارهای گزارش‌دهی طراحی

شده به خود الگوریتم برای همه نواح هوش مصنوعی حسابرسی یک پیش شرط ضروری برای تأیید عملکرد صحیح است.

باتوجه به آنچه گفته شد برای مقابله با چالش‌های اخلاقی هوش مصنوعی راهکارهای دیگری وجود دارد:

۱- آموزش اخلاقی الگوریتم‌ها: الگوریتم‌ها می‌توانند با ارزش‌های اخلاقی آموزش ببینند تا کمتر دچار تعصب شوند.

۲- تنظیم مقررات و استانداردها: دولت‌ها باید مقرراتی برای رعایت اخلاق در هوش مصنوعی وضع کنند. استانداردهایی نیز باید تدوین شود.

۳- طراحی مسئولانه سیستم‌های هوشمند: متخصصان هوش مصنوعی باید اثرات اجتماعی طراحی‌های خود را مدنظر قرار دهند.

۴- افزایش آگاهی عمومی: افزایش سواد و آگاهی عمومی درباره هوش مصنوعی و جنبه‌های اخلاقی آن می‌تواند کمک‌کننده باشد (همان)

در جدول (۱)، خلاصه‌ای از چالش‌های اخلاقی هوش مصنوعی و راهکارهای پیشنهادی برای مقابله با آن‌ها ارائه شده است. این جدول، دستاوردی از تحلیل و ترکیب اطلاعات به‌دست آمده از منابع مختلف و تحلیل محتوای اسنادی این پژوهش است که به‌طور خلاصه مهم‌ترین ابعاد اخلاقی و راه‌حل‌های مربوط به آن‌ها را به نمایش می‌گذارد. هر یک از این چالش‌ها نیازمند توجه و برنامه‌ریزی دقیق برای اطمینان از توسعه مسئولانه و اخلاقی هوش مصنوعی هستند.

جدول (۱): چالش‌های اخلاقی هوش مصنوعی و راهکارهای پیشنهادی

ردیف	چالش اخلاقی	توضیح چالش	راهکارهای پیشنهادی
۱	نقض خصوصیت	جمع‌آوری و استفاده غیر مجاز از داده‌های شخصی توسط هوش مصنوعی، به ویژه در حوزه‌های سلامت و رسانه	- تدوین قوانین سخت گیرانه مانند GDPR. - رمزگذاری داده‌ها و محدود کردن دسترسی. - اخذ رضایت آگاهانه از کاربران.

۲	تبعیض و سوگیری الگوریتمی	بازتولید تعصبات جنسیتی، نژادی یا اجتماعی در تصمیم‌گیریهای هوش مصنوعی (مثلا در استخدام یا رسانه)	-آموزش الگوریتمی با دادهای متعادل و عادی از سوگیری. شفاف سازی معیارهای تصمیم‌گیری. نظارت مستمر بر عملکرد سیستم‌ها.
۳	تهدید امنیت عمومی	استفاده از هوش مصنوعی برای شناسایی و تشدید شورشها یا انتشار اخبار جعلی	-توسعه ابزارهای نظارتی برای تشخیص سوءاستفاده. -همکاری با نهادهای امنیتی برای کنترل محتوای مخرب.
۴	کاهش ارزش انسانی و مهارتها	وابستگی به هوش مصنوعی و تضعیف خلاقیت، استقلال و مهارتهای انسانی	آموزش سواد دیجیتال و اخلاق هوش مصنوعی. -تقویت نقش انسان در تصمیم‌گیری‌های حیاتی.
۵	مسئولیت‌پذیری در خطاها	عدم شفافیت در تعیین مسئولیت هنگام وقوع خطاهای سیستمهای هوشمند (مثلا در خودروهای خودران)	-تعیین چارچوب‌های حقوقی واضح برای مسئولیت‌پذیری. -طراحی سیستم‌های قابل حسابرسی و شفافیت.
۶	تهدید اشتغال	جایگزینی مشاغل انسانی با ربات‌ها و سیستمهای	-بازآموزی نیروی کار برای مشاغل جدید. -توسعه سیاست‌های حمایتی از بازار کار.
۷	عدم شفافیت الگوریتمی	پیچیدگی الگوریتمها و عدم امکان تفسیر تصمیم‌گیری‌های آنها	-استفاده از الگوریتم‌های ساده‌تر و قابل درک. -ارائه گزارش‌های شفاف از عملکرد سیستم‌ها.
۸	استفاده غیر اخلاقی	سوءاستفاده از هوش مصنوعی برای اهداف خرابکارانه مانند حملات سایبری یا دستکاری اطلاعات	-ایجاد نهادهای نظارتی بین‌المللی -توسعه استانداردهای امنیتی سختگیرانه

۵- جمع بندی و نتیجه گیری

همانطور که رحمتی و همکاران (۱۴۰۳) و عباسی و همکاران (۱۴۰۳) بر لزوم توسعه کنترل شده و اخلاقی هوش مصنوعی و اهمیت حفظ حریم خصوصی و شفافیت تأکید نمودند، این پژوهش نیز به همین نتایج رسیده و چالش‌هایی چون نقض حریم خصوصی، سوگیری الگوریتمی، و ابهامات مسئولیت‌پذیری را برجسته می‌سازد. در همین راستا، همانطور که کلاته آقامحمدی (۱۴۰۲) بر تدوین قوانین و ارتقای آگاهی عمومی برای مواجهه با این چالش‌ها تأکید کرد، این تحقیق نیز لزوم وضع استانداردهای اخلاقی دقیق، حسابرسی مستمر و طراحی مسئولانه سیستم‌ها را ضروری می‌داند. علاوه بر این، همسو با دیدگاه شریعتی و اکبرزاده (۱۴۰۱) مبنی بر پرهیز از ضرر رساندن، این پژوهش به این نتیجه رسیده است که گنجاندن عناصر اخلاقی مانند مسئولیت و اراده به صورت مصنوعی در ساختار هوش مصنوعی، برای حفظ کنترل انسان و تضمین همسویی فناوری با ارزش‌های انسانی حیاتی است.

راهکارهای اساسی برای کنترل چالش‌های اخلاقی هوش مصنوعی، با الهام از مبانی فلسفه و حقوق اخلاق در اسلام، شامل چندین محور کلیدی است. از منظر اخلاق اسلامی، که بر اصول عدالت، انصاف، حفظ حریم خصوصی، و پرهیز از ضرر (لاضرر) تأکید دارد، ضروری است که در طراحی و پیاده‌سازی سیستم‌های هوش مصنوعی، این ارزش‌های بنیادین به طور فعالانه گنجانده شوند. این امر مستلزم تدوین قوانین و مقررات سخت گیرانه و شفاف است که نه تنها به حفظ حریم خصوصی و امنیت داده‌ها کمک کند، بلکه از هرگونه تبعیض و سوگیری الگوریتمی جلوگیری نماید. آموزش الگوریتم‌ها با داده‌های متعادل و عاری از تعصب، شفاف‌سازی معیارهای تصمیم‌گیری، و نظارت مستمر بر عملکرد سیستم‌ها از جمله راهکارهای عملی برای تضمین عدالت و انصاف در هوش مصنوعی هستند. همچنین، برای مقابله با تهدیدات امنیتی و سوءاستفاده‌های احتمالی، توسعه ابزارهای نظارتی پیشرفته و همکاری بین‌المللی حیاتی است. از دیدگاه اسلامی، تقویت نقش انسان در تصمیم‌گیری‌های حیاتی و افزایش سواد اخلاقی و دیجیتال جامعه، تضمین‌کننده کرامت انسانی و جلوگیری از وابستگی مفرط به ماشین‌هاست. در نهایت، با در نظر گرفتن مسئولیت انسان به عنوان

خالق هوش مصنوعی، بارگذاری عناصر انتزاعی اخلاقی نظیر مسئولیت و اراده در چارچوب مفاهیم خوب و بد به طور مصنوعی در این سیستم‌ها، راهکاری بنیادین برای حفظ کنترل انسان و جلوگیری از خروج هوش مصنوعی از مسیر خدمت به بشریت است.

هدف اصلی سیستم‌های هوش مصنوعی، ساخت سیستم‌هایی است که بتوانند شبیه انسان عمل کرده، دارای قابلیت آگاهی، تفکر و تجربه باشند. در حالی که برخی از فیلسوفان معتقدند سیستم‌های هوش مصنوعی قادر به تقلید از ویژگی‌های بشر هستند، برخی دیگر همچنان این موضوع را ناممکن می‌دانند. اگر چه تصور می‌شود سیستم‌های هوش مصنوعی با ویژگی‌های انسانی و دارای آگاهی و تفکر ساخته شوند، باید به یاد داشت این سیستم‌ها ساخت انسان و برگرفته از هوش طبیعی خواهند بود. همچنین، اگر چه با تحقیقات هوش مصنوعی ممکن است؛ اخلاق مصنوعی را هم بتوان تولید کرد، اما واضح است که این اخلاق نیز از هوش و اخلاق طبیعی انسان‌ها نشأت می‌گیرد؛ بنابراین، خالق سیستم‌های هوش مصنوعی انسان است. در این زمینه، با در نظر گرفتن موقعیت مرکزی انسان در مسئولیت، با حفظ تعادل صحیح، باید اجازه داد هوش مصنوعی توسعه یابد. به عبارت دیگر، باید در جهت پیشرفت هوش مصنوعی حرکت کنیم و در عین حال، مسئولیت و اخلاق را در نظر داشته باشیم. اگر سیستم‌های هوش مصنوعی به قدرت یادگیری عمیق و نرم‌افزار خود دسترسی داشته باشند، آیا امکان اتخاذ تدابیری برای محدود کردن این دسترسی وجود خواهد داشت؟ همچنین اگر ربات‌ها مخالفت کنند و در مقابل اطاعت از دستورات انسان یا کنترل شدن اعتراض کنند، چگونه می‌توان از مسائل اخلاقی ناشی از این وضعیت جلوگیری کرد؟

دانشمندان مؤسسه آینده زندگی، مجموعه‌ای از اهداف کوتاه‌مدت را برای مشاهده اثر هوش مصنوعی بر اقتصاد قانون اخلاق و چگونگی کاهش مخاطرات امنیتی تعریف کرده‌اند. از طرفی رسانه یکی از تأثیرپذیرترین صنایع نسبت به تغییرات تکنولوژی است و تأثیر بسیار زیادی بر فرهنگ و ارتباطات اجتماعی جوامع دارد. این امر باعث توجه بیش‌ازپیش شرکت‌های بزرگ و فعالان حوزه رسانه به هوش مصنوعی و به خدمت گرفتن ابزار و امکانات آن شده است که در بازار داغ رقابت امروز برای جذب مخاطب بیشتر به یک امر حیاتی تبدیل شده است.

۶- منابع و مآخذ

الف- منابع فارسی

- ابراهیمی، شهرام. (۱۴۰۱). پیشگیری از تکرار جرم از طریق هوش مصنوعی؛ مقتضیات و محدودیت‌ها. *آموزه‌های حقوق کیفری* ۳۳-۵۴، ۱۹(۲۳)، doi: ۱۰,۳۰۵۱۳/cld.۲۰۲۳,۴۳۴۵,۱۷۰۱
- ابوذری، مهرنوش. (۱۴۰۱). تأثیر هوش مصنوعی در کیفیت تحقیقات جنایی. *حقوق فناوری های نوین* ۱۰,۲۲۱۳۳/mtlj.۲۰۲۲,۳۵۱۱۷۳,۱۱۰۶، doi: ۱۰,۲۲۱۳۳/mtlj.۲۰۲۲,۳۵۱۱۷۳,۱۱۰۶، ۱(۶)، ۱-۱۳.
- انصاری، باقر، عطار، شیما، صالحی، امیرحسین (۱۴۰۰). *حقوق داده‌ها و هوش مصنوعی، مفاهیم و چالش‌ها*، شرکت سهامی انتشار
- آهنگری، فرشته (۱۳۸۶). پیشینه و بنیادهای اخلاق در ایران و جهان، *فصلنامه اخلاق در علوم و فناوری*، شماره ۳
- تقی زاده، مسلم و خردمندنیاز، سهیلا. (۱۴۰۳). هوش مصنوعی مولد؛ چالش‌ها و الزامات توسعه و پیاده سازی. (۱۹۸۷۹). *ماهنامه گزارش‌های کارشناسی مرکز پژوهش‌های مجلس شورای اسلامی* ۱۹۸۷۹، ۳۱(۴)، doi: ۱۰.۲۲۰۳۴/report.۲۰۲۴,۱۷۰۰۶,۱۸۴۱
- چناری، مهین؛ شبستانی، اعظم (۱۴۰۰). پیش‌بینی میزان آشنایی با اخلاق در فضای مجازی با استفاده از سواد رسانه‌ای و شایستگی‌های فناورانه، *اخلاق در علوم و فناوری*، دوره ۱۶، شماره ۱
- ریبی‌زاده، احمد (۱۴۰۰). کاربرد هوش مصنوعی در پژوهش‌های علوم اسلامی، *ره آورد نور*، دوره ۲۰، شماره ۷۴
- رجب‌یان ده زیره، مریم. (۱۴۰۳). شناسایی چالش‌ها و قابلیت‌های هوش مصنوعی در آموزش و یادگیری با ارائه راهکارها. *فناوری آموزش* ۹۵۰-۹۲۱، ۱۸(۴)، doi: ۱۰,۲۲۰۶۱/tej.۲۰۲۴,۱۰۷۷۷,۳۰۵۸، ۱(۴)، ۱۰,۲۲۰۶۱/tej.۲۰۲۴,۱۰۷۷۷,۳۰۵۸، ۱(۴)، ۹۲۱-۹۵۰.
- رحمتی، فاطمه، مشایخی، جنت. (۱۴۰۳). ملاحظات اخلاقی استفاده از هوش مصنوعی در نظام سلامت. *پایش*. شماره ۲۳
- رحیمی اقدم، صمد؛ صالح‌پور، پدرام؛ نامور، رقیه (۱۴۰۳). چالش‌های اخلاقی اتخاذ هوش مصنوعی در مدیریت منابع انسانی. *اخلاق در علوم و فناوری*، دوره ۱۹، شماره ۴.

<http://dx.doi.org/10.22034/ethicsjournal.19,4,18>

رمضانی، مجید؛ فیضی درخشی، محمدرضا (۱۳۹۲) اخلاق ماشین، چالش‌ها و رویکردهای مسائل اخلاقی در هوش مصنوعی و ابرهوش، **اخلاق در علوم و فناوری**، دوره ۸، شماره ۴.

<https://dor.isc.ac/dor/20,1001,1,22517634,1392,8,4,4,7>

شریعتی، فهیمه و اکبرزاده توتونچی، محمدرضا. (۱۴۰۱). مقایسه ماهوی هوش انسانی با هوش مصنوعی از منظر فلسفه اسلامی با تاکید بر حکمت متعالیه ملاصدرا، راهگشایی در فهم جایگاه عقول برتر. **حکمت معاصر** ۱۱۷-۸۹، (۲)، ۱۳،

doi: 10.30465/cw.20.23,32287,1743

صالح، ابودر، رجبی، محمود (۱۳۹۷) بررسی توان رقابت هوش مصنوعی با ذهن انسان از منظر قرآن، **قرآن‌شناخت**، دوره ۱۱، شماره ۲

شقایق صراف زاده، احسان ابوطالب، (۱۴۰۲). اهمیت اخلاق در استفاده از هوش مصنوعی در آموزش پزشکی، **نشریه پژوهش در آموزش علوم پزشکی**، ۱۵(۲)، ۱-۳۴. magiran.com/p26424034.
عاشوری کیسی، محمد علی. (۱۴۰۳). طرح و بررسی برخی از مسائل اخلاقی هوش مصنوعی در هنر. **متافیزیک** ۱۴۸۸، ۱۳۸۱۵، ۲۰۲۴. doi: 10.22108/mph.2024,13815,14888,93-110, 19(37),

عباسی، محمود و تیموری، مهرداد. (۱۴۰۳). چالش‌های اخلاق رقومی هوش مصنوعی و امکان‌سنجی بیمه مسئولیت برای سامانه مبتنی بر هوش مصنوعی. **نشریه علمی «دولت و حقوق»**، ۱۵(۱)، ۳۱-۴۴.
عظیمی، محمدحسن. اسماعیلی، سمیرا (۱۴۰۰) شناسایی مؤلفه‌های هوش مصنوعی در پایگاه‌های اطلاعاتی ایرانی. **مجله دانش‌شناسی**، دوره ۱۴، شماره ۵۴

کلاته آقامحمدی، آمنه. (۱۴۰۲). چالش‌های اخلاقی هوش مصنوعی در رسانه‌های نوین و راهکارهای مواجهه با آن. **مطالعات دینی رسانه**، ۸۷-۵۲، (۲۰)، ۵. doi: 10.22034/jmrs.2024,205457,

گشمرد، رقیه (۱۴۰۳) چالش‌های اخلاقی کاربرد هوش مصنوعی در پزشکی. **فصلنامه علمی تخصصی پژوهش‌های نوین در روان‌شناسی**، دوره ۵، شماره ۳

مهدیانی، محمدرضا (۱۴۰۰) جایگاه اخلاق در فضای مجازی، **پژوهش مطالعات علوم اسلامی**، دوره سوم. شماره ۳۰

ب: منابع انگلیسی

- Belk, R (۲۰۲۰). Ethical Issues in Service Robotics and Artificial Intelligence. The Service Industries Journal.
- Buttarelli, Giovanni (۲۰۱۸). ICDPPC ۲۰۱۸ Debating Ethics: Dignity and Respect in a Data Driven Life. Choose Humanity: Putting Dignity Back into Digital - Opening Speech of Public Session
- Kaplan A. Haenlein M (۲۰۱۹) Siri, Siri, in my hand: who's the fairest in the land? On the interpretations, illustrations and implications of artificial intelligence. Bus Horiz ۶۲:۱۵-۲۵
- Müller, V.C. (۲۰۲۱). Ethics of Artificial Intelligence and Robotics. The Stanford Encyclopedia of Philosophy (Summer ۲۰۲۱ Edition).
- Osoba, OA. & Welser, W. (۲۰۱۷) An Intelligence in Our Image: The Risks of Bias and Errors in Artificial Intelligence RAND Corporation. Retrieved from https://www.rand.org/pubs/research_reports/RR1744.
- Wang, W., & Siau, K. (۲۰۱۸). Ethical and Moral Issues with AI: A Case Study on Healthcare Robots. Twenty-fourth Americas Conference on Information Systems, Retrieved from <https://www.researchgate.net/publication>
- Webb, Helena; William Housley (۲۰۲۱) Forecasting the governance of harmful social media communications: findings from the digital wildfire policy Delphi. (۲۰۲۰) Policing and Society Journal article.
- Anderson, Michael; Anderson, Susan Leigh, eds. (July ۲۰۱۱). Machine Ethics. Cambridge University Press. ISBN ۹۷۸-۰-۵۲۱-۱۱۲۳۵-۲.
- Anderson, M.; Anderson, S.L. (July ۲۰۰۶). Guest Editors' Introduction: Machine Ethics. IEEE Intelligent Systems. ۲۱ (۴): ۱۰-۱۱. doi:10.1109/mis.۲۰۰۶.۷۰. S2CID ۹۵۷۰۸۳۲.
- Winfield, A. F.; Michael, K.; Pitt, J.; Evers, V. (March ۲۰۱۹). Machine Ethics: The Design and Governance of Ethical AI and Autonomous Systems [Scanning the Issue] . Proceedings of the IEEE. ۱۰۷ (۳): ۵۰۹-۵۱۷. doi:10.1109/JPROC.۲۰۱۹.۲۹۰۰۶۲۲. ISSN ۱۵۵۸-۲۲۵۶. S2CID ۷۷۳۹۳۷۱۳..

Sauer, Megan (۲۰۲۲-۰۴-۰۸). Elon Musk says humans could eventually download their brains into robots — and Grimes thinks Jeff Bezos would do it. CNBC.

Bostrom, Nick; Yudkowsky, Eliezer (۲۰۱۱). The Ethics of Artificial Intelligence (PDF). Cambridge Handbook of Artificial Intelligence. Cambridge Press. Archived (PDF) from the original on ۲۰۱۶-۰۳-۰۴. Retrieved ۲۰۱۱-۰۶-۲۲.

Allahverdi, Novruz (۲۰۰۲). Uzman Sistemler: Bir Yapay Zekâ Uygulaması. İstanbul: Atlas Yayınları.

Asimov, Isaac (۲۰۰۶). Ben, Robot. (Çev. Ekin Odabaş) , İstanbul: İthaki Yayınları.

Çelebi, Ömer, Faruk (۲۰۱۷). Yapay Zekâ ve Etik'. İstanbul Medeniyet Üniversitesi Sosyal Bilimler Enstitüsü, s. ۱-۲۰.

Platon (۲۰۱۳). Menon. (Çev. Furkan Akderin), İstanbul: Say Yayınları.

Uyar, Tevfik (۲۰۱۷). Ya Yapay Ahlâk. İnsanlaşan Makineler ve Yapay Zekâ içinde, haz. Mehmet Karaca, Say ۷۵, ۱۸-۲۱ İstanbul: İstanbul Teknik Üniversitesi Vakfı Yayınları.

Ebrahimi, Shahram, (۱۴۰۱) Prevention of recidivism through artificial intelligence; requirements and limitations, Criminal Law Quarterly, No. ۲۲ doi: ۱۰,۳۰۵۱۳/cld.۲۰۲۳,۴۳۴۵,۱۷۰۱ [In Persian]

Abuzari, Mehrnoush, (۱۴۰۰) The impact of artificial intelligence on the quality of criminal investigations, New Technologies Law Quarterly, No. ۶. . doi: ۱۰,۲۲۱۳۳/mtlj.۲۰۲۲,۳۵۱۱۷۳,۱۱۰۶ [In Persian]

Ansari, Baqer, Attar, Shima, Salehi, Amir Hossein (۱۴۰۰) Data Rights and Artificial Intelligence, Concepts and Challenges, Publishing Joint Stock Company [In Persian]

Ahangari, Fereshta (۲۰۰۷) Background and Foundations of Ethics in Iran and the World, Ethics in Science and Technology, No. ۳ [In Persian]

- Taghizadeh, Muslim; Kheradmandnia, Soheila (۱۴۰۳). Generative Artificial Intelligence; Challenges and Requirements for Development and Implementation. Monthly Journal of Expert Reports of the Research Center of the Islamic Consultative Assembly, ۳۲(۴), ۱۹۸۷ doi: ۱۰,۲۲۰۳۴/report.۲۰۲۴,۱۷۰۰۶,۱۸۴۱ [In Persian]
- Chenari, Mahin; Shabestani, Azam (۱۴۰۰) Predicting the level of familiarity with ethics in virtual space using media literacy and technological competencies Ethics in Science and Technology, ۲۰۱۶, No. ۱ [In Persian]
- Rabieizadeh, Ahmad (۱۴۰۰) Application of Artificial Intelligence in Islamic Sciences Research, Rahavard Noor, Volume ۲۰, No. ۷۴ [In Persian]
- Rajabian Dehzireh, Maryam (۲۰۱۳) Identifying the challenges and capabilities of artificial intelligence in teaching and learning by providing solutions. Department of Educational Technology, Faculty of Psychology and Educational Sciences, Allameh Tabatabaei University, Tehran, Iran doi: ۱۰,۲۲۰۶۱/tej.۲۰۲۴,۱۰۷۷۷,۳۰۵۸ [In Persian]
- Rahmati, Fatima, Mashayekhi, Jannat. (۱۴۰۳) Ethical Considerations of Using Artificial Intelligence in the Health System. Monitoring. Issue ۲۳ Rahmati, Fatima, Mashayekhi, Jannat. (۱۴۰۳) Ethical Considerations of Using Artificial Intelligence in the Health System. Monitoring. No. ۲۳ [In Persian]
- Rahimi Aghdam, Samad; Salehpour, Pedram; Namwar, Roghieh (۲۰۱۴) Ethical Challenges of Adopting Artificial Intelligence in Human Resource Management. Department of Management, Faculty of Economics and Management, University of Tabriz, Tabriz, Iran. [In Persian]
- Ramadani, Majid; Fayzi Darakhshi, Mohammad Reza (۲۰۱۳) Machine Ethics, Challenges and Approaches to Ethical Issues in Artificial Intelligence and Superintelligence, Ethics in Science and Technology, Year ۸, N. ۴ [In Persian]
- Shariati, Fahimeh; Mohammadreza, Akbarzadeh (۱۴۰۱), Comparing the Essence of Human Intelligence with Artificial Intelligence from the Perspective of Islamic Philosophy, Contemporary Wisdom, Volume ۱۳, No.

۲ doi: ۱۰,۳۰۴۶۵/cw.۲۰۲۳,۳۲۲۸۷,۱۷۴۳ [In Prsian]

Saleh, Abuzar, Rajabi, Mahmoud (۲۰۱۸) Investigating the Competitiveness of Artificial Intelligence with the Human Mind from the Perspective of the Quran, Quran Studies, Volume ۱۱, No. ۲ [In Prsian]

Sarafzadeh, Shaghayegh, Abutaleb, Ehsan (۱۴۰۲) The Importance of Ethics in the Use of Artificial Intelligence in Medical Education. Medical Education Research Center, Medical Education Studies and Development Center, Guilan University of Medical Sciences, Rasht [In Prsian]

Ashuri Qaismi, Mohammad Ali. (۱۴۰۳). Design and examination of some ethical issues of artificial intelligence in art. Metaphysics, ۱۶(۳۷), ۹۳ doi: ۱۰,۲۲۱۰۸/mp.۲۰۲۴,۱۳۸۱۰۵,۱۴۸۸ [In Prsian]

Abbasi, Mahmoud and Teymouri, Mehrdad. (۱۴۰۳). Challenges of digital ethics of artificial intelligence and feasibility of liability insurance for artificial intelligence-based systems. Scientific Journal of "Government and Law", ۵(۱), ۳۱-۴۴. [In Prsian]

Azimi, Mohammad Hassan. Esmaili, Samira (۲۰۱۱) Identifying Artificial Intelligence Components in Iranian Databases. Journal of Knowledge, Volume ۱۴, No. ۵۴ [In Prsian]

Kalateh Aghamohammadi, Ameneh (۱۴۰۲). Ethical challenges of artificial intelligence in new media and solutions to face them. Religious Studies of Media, ۵(۲۰), ۵۲-۸۷. doi: ۱۰,۲۲۰۳۴/jmrs.۲۰۲۴,۲۰۵۴۵۷ [In Prsian]

Gashmard, Roqiyeh (۱۴۰۳) Gashmard, Roqiyeh (۱۴۰۳) Ethical Challenges of the Application of Artificial Intelligence in Medicine. Specialized Quarterly Scientific Journal of Modern Research in Psychology Ethical Challenges of the Application of Artificial Intelligence in Medicine. Specialized Quarterly Scientific Journal of Modern Research in Psychology. [In Prsian]

Mahdiani, Mohammad Reza (۱۴۰۰) The Place of Ethics in Virtual Space, Islamic Sciences Studies Research, Volume ۳. No. ۳۰ [In Prsian]