

Ethical Analysis of Terms and Conditions for the Use of AI-Based Products and Services in Iran

Reza Payandeh¹, Hasti Kanani²

1- Assistant Professor, Faculty of Governance, University of Tehran, Tehran, Iran.

(Corresponding Author: reza.payandeh@ut.ac.ir)

2- M.A. Student, Faculty of Governance, University of Tehran, Tehran, Iran.

Abstract

This study examines the "Terms and Conditions" (T&C) of AI-driven businesses, aiming to identify relevant ethical issues and evaluate their compliance with AI ethical frameworks. The term "Terms and Conditions" refers to a document that regulates the contractual relationship between service providers and users, typically encompassing privacy policies, terms of service, user and company rights and obligations, as well as data protection policies. This research employed a cross-sectional study design and qualitative content analysis to identify and analyze Iranian businesses operating in the AI sector. The sample was drawn from platforms such as Myket and Cafe Bazaar, complemented by Google searches, and data were collected from official websites and relevant documents. The collected data were coded and analyzed based on ethical principles. In the quantitative phase, scoring was conducted using a Likert scale, and businesses were ranked using the TOPSIS method. The findings revealed that the T&C policies often lack sufficient operational details to build user trust and fail to fully uphold principles such as the balance of power, accountability, and social responsiveness. This study highlights the necessity of revising these T&C policies, developing more comprehensive frameworks, and aligning them with international standards.

Keywords: Terms and Conditions, Business, Ethics, Artificial Intelligence, Cross-Sectional Study.

How to Cite this Paper:

Payandeh, R., Kanani, H. (2024). **Ethical Analysis of Terms and Conditions for the Use of AI-Based Products and Services in Iran**. *Journal of Science & Technology Policy*, 17(3), 27-39. {In Persian}.

doi: 10.22034/jstp.2025.11769.1813



تحلیل اخلاقی شرایط و مقررات استفاده از محصولات و خدمات شرکت‌های مبتنی بر هوش مصنوعی در ایران

رضا پاینده^۱، هستی کنعانی^۲

۱- استادیار دانشکده حکمرانی، دانشگاه تهران، تهران، ایران. (نویسنده عهده‌دار مکاتبات: reza.payandeh@ut.ac.ir)

۲- دانشجوی کارشناسی ارشد، دانشکده حکمرانی، دانشگاه تهران، تهران، ایران

چکیده

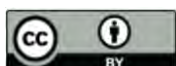
این مقاله به بررسی «شرایط و مقررات» کسب‌وکارهای مبتنی بر هوش مصنوعی می‌پردازد و با هدف شناسایی موارد اخلاقی مرتبط و ارزیابی میزان تطابق آن‌ها با چارچوب‌های اخلاقی هوش مصنوعی انجام شده است. عبارت «شرایط و مقررات» به سندی اشاره دارد که رابطه قراردادی میان ارائه‌دهنده خدمات و کاربران را تنظیم می‌کند و معمولاً شامل سیاست‌های حریم خصوصی، شرایط استفاده از خدمات، حقوق و تعهدات کاربران و شرکت و همچنین سیاست‌های حفاظت از داده‌ها است. این تحقیق با استفاده از روش مطالعه مقطعی و تحلیل محتوای کیفی، به شناسایی و تحلیل کسب‌وکارهای ایرانی فعال در حوزه هوش مصنوعی پرداخته است. نمونه‌گیری از پلتفرم‌های مایکت و کافه‌بازار و جستجو در گوگل انجام شده و داده‌ها از وبگاه‌های رسمی و مستندات مرتبط جمع‌آوری، کدگذاری و بر اساس اصول اخلاقی تحلیل شده‌اند. در بخش کمی، امتیازدهی با طیف لیکرت و رتبه‌بندی با روش تاپسیس انجام شد. یافته‌ها نشان داد که شرایط و مقررات اغلب فاقد جزئیات عملیاتی کافی برای جلب اعتماد کاربران هستند و اصولی همچون تعادل قدرت، مسئولیت‌پذیری و واکنش اجتماعی در آن‌ها به طور کامل رعایت نشده است. این مطالعه بر ضرورت بازنگری در این شرایط و مقررات، تدوین چارچوب‌های جامع‌تر و هماهنگی با استانداردهای بین‌المللی تأکید دارد.

کلیدواژه‌ها: شرایط و مقررات، کسب‌وکار، اخلاق، هوش مصنوعی، مطالعه مقطعی.

برای استنادات بعدی به این مقاله، قالب زیر به نویسندگان محترم مقالات پیشنهاد می‌شود:

پاینده، رضا. و کنعانی، هستی. (۱۴۰۳). تحلیل اخلاقی شرایط و مقررات استفاده از محصولات و خدمات شرکت‌های مبتنی بر هوش مصنوعی در ایران. سیاست علم و فناوری، ۱۷(۳)، ۲۷-۳۹.

doi: 10.22034/jstp.2025.11769.1813



۱- مقدمه

با رشد سریع هوش مصنوعی در حوزه‌های گوناگون، به‌ویژه در کسب‌وکارها، به‌کارگیری این فناوری در تعامل با مشتریان به یکی از چالش‌های اساسی در زمینه اخلاق داده‌ها تبدیل شده‌است. پیشرفت‌های چشمگیر این فناوری در تحلیل داده‌ها و تصمیم‌گیری‌های خودکار، نگرانی‌ها درباره نحوه استفاده از داده‌های شخصی، حفظ حریم خصوصی و تبعیض‌های الگوریتمی را افزایش داده‌است [۱، ۲]. در دهه‌های اخیر با پیشرفت‌های بیشتر در هوش مصنوعی و یادگیری ماشین، توجه به ملاحظات اخلاقی و حقوقی به یکی از موضوعات اصلی پژوهش‌های علمی و سیاست‌گذاری‌های جهانی تبدیل شد [۳، ۴].

استفاده از هوش مصنوعی فرصت‌های بی‌نظیری در ارائه خدمات بهتر به مشتریان فراهم می‌کند. به‌عنوان مثال، سیستم‌های هوش مصنوعی قادرند با تحلیل داده‌های مشتریان، پیشنهادهای شخصی‌سازی‌شده ارائه دهند تا تجربه مشتری را بهبود بخشند. همچنین، با استفاده از الگوریتم‌های پیشرفته، کسب‌وکارها می‌توانند نیازهای مشتریان را پیش‌بینی کرده و به شیوه‌ای سریع‌تر و کارآمدتر پاسخ دهند [۳، ۵]. با این حال، استفاده گسترده از هوش مصنوعی و داده‌های مشتریان مسائل اخلاقی جدی را نیز به همراه دارد. یکی از نگرانی‌های اصلی، «استعمار داده‌ها»^۱ است که به تلاش شرکت‌ها برای ادعای مالکیت و خصوصی‌سازی داده‌های تولیدشده توسط کاربران اشاره دارد. علاوه بر این، مسئله «شست‌وشوی اخلاقی»^۲ نیز مطرح شده که به معنای استفاده از اصول اخلاقی به‌صورت نمایشی است تا نگرانی‌های مربوط به استفاده از داده‌ها را کاهش دهد، بدون اینکه تغییرات معناداری در سیاست‌ها و رویه‌ها ایجاد شود [۶]. از دیگر مسائل مهم، «حباب الگوریتمی»^۳ است؛ وضعیتی که در آن الگوریتم‌های هوش مصنوعی تعصبات موجود در داده‌ها را بازتولید و تقویت می‌کنند و این امر می‌تواند به نابرابری‌های اجتماعی در زمینه‌هایی مانند استخدام، اعطای وام، یا قیمت‌گذاری شخصی‌سازی‌شده دامن بزند [۱، ۷]: در

مطالعه‌ای مرتبط با هوش مصنوعی مولد، نشان داده شده که این فناوری می‌تواند به تداوم الگوهای ناعادلانه‌ای منجر شود که نابرابری‌های اجتماعی موجود را تثبیت می‌کنند [۸]. همچنین، پژوهشی در حوزه بهداشت بر چالش‌های عدالت در استفاده از هوش مصنوعی تمرکز داشته و بر لزوم بازنگری الگوریتم‌ها برای جلوگیری از تبعیض تأکید کرده است [۹]. مسئله دیگر، «نظارت تهاجمی»^۴ است که به استفاده از سیستم‌های هوش مصنوعی برای جمع‌آوری گسترده داده‌ها و نظارت فراگیر اشاره دارد که می‌تواند به نقض حریم خصوصی و ایجاد حس عدم اعتماد در میان کاربران منجر شود [۱۰، ۱۱]: و ده‌ها مسئله اخلاقی دیگر.

در تحلیل مسائل اخلاقی و در مقام یافتن راهکار، برخی محققان به اهمیت «پاسخگویی الگوریتمی»^۵ تأکید کرده‌اند تا اطمینان حاصل شود که تصمیمات گرفته‌شده توسط الگوریتم‌های هوش مصنوعی قابل‌ردیابی هستند. این مسئله به‌ویژه در مواردی که هوش مصنوعی برای تصمیم‌گیری‌های مهم استفاده می‌شود، مانند تصمیم‌گیری‌های پزشکی یا حقوقی، اهمیت بیشتری پیدا می‌کند [۱۲].

ترکیب این مسائل و فرصت‌ها نشان می‌دهد که بهره‌گیری از هوش مصنوعی در کسب‌وکارها نیازمند رویکردی است متوازن که هم از مزایای این فناوری بهره‌برداری کند و هم ملاحظات اخلاقی را رعایت نماید. برخی پژوهش‌ها هشدار داده‌اند که بی‌توجهی به ملاحظات اخلاقی می‌تواند به تأثیرات منفی گسترده‌ای مانند «طرد اجتماعی»^۶ منجر شود، درحالی‌که تمرکز بیش از حد بر حمایت از حقوق فردی می‌تواند مزایای بالقوه هوش مصنوعی را محدود کند [۱].

هدف مقاله حاضر این است که با رویکردی ارزش‌یابانه، شرایط و مقررات^۷ تعدادی از کسب‌وکارهای مبتنی بر هوش مصنوعی در ایران را مورد بررسی قرار داده و میزان انطباق این شرایط با اصول اخلاقی در ادبیات علمی را بررسی کند. برخی پژوهش‌ها این موضوع را به صورت بین‌المللی بررسی کرده‌اند [۱۳] و البته پژوهش‌های ارزشیابانه‌ای در زمینه

^۴ Mass Surveillance

^۵ Algorithmic Accountability

^۶ Social Exclusion

^۷ Terms & Conditions

^۱ Data Colonialism

^۲ Ethics Washing

^۳ Algorithmic Bubble

اخلاقی شرایط و مقررات در کسب‌وکارهای مبتنی بر هوش مصنوعی پیشنهاد می‌شود. این اصول تلاش می‌کنند تا جامعیت بیشتری داشته و به نیازهای متنوع جامعه و کاربران در مواجهه با فناوری‌های نوین پاسخ دهند.

الف) مسئولیت‌پذیری و شفافیت^۱ به معنای اطمینان از این است که توسعه‌دهندگان مسئولیت عملکرد صحیح هوش مصنوعی را به عهده می‌گیرند و فرایندهای تصمیم‌گیری شفاف هستند. قوانین و شرایط استفاده باید به وضوح توضیح دهند که چه کسی مسئول هر تصمیم هوش مصنوعی است و چگونه این مسئولیت به کاربران توضیح داده می‌شود.

ب) قابلیت توضیح و تفسیر^۲ به معنای طراحی سیستم‌های هوش مصنوعی به گونه‌ای است که نحوه عملکرد آن‌ها برای کاربران قابل فهم و ردیابی باشد. قوانین و شرایط استفاده باید به وضوح توضیح دهند که تصمیمات هوش مصنوعی چگونه و بر چه اساسی گرفته می‌شوند که این می‌تواند شامل توضیحات فنی ساده و قابل فهم برای کاربران باشد.

ج) عدالت و برابری^۳ به معنای اطمینان از این است که هوش مصنوعی بدون تعصب و تبعیض عمل می‌کند و به عدالت اجتماعی کمک می‌کند. قوانین و شرایط استفاده باید اطمینان دهند که هوش مصنوعی به صورت عادلانه و بدون تبعیض عمل می‌کند که شامل بررسی منظم و مستقل الگوریتم‌ها برای شناسایی و حذف هرگونه تعصب است.

د) سودمندی^۴ به معنای استفاده از هوش مصنوعی به نفع رفاه عمومی جامعه و کیفیت زندگی افراد است. قوانین و شرایط استفاده باید نشان دهند که هوش مصنوعی به طور کلی به نفع کاربران است و تعهدات و اهداف اخلاقی خود در استفاده از این فناوری‌ها را به روشنی بیان می‌کند.

ه) عدم فریب‌کاری^۵ به معنای جلوگیری از هرگونه فریب، سردرگمی یا ارائه اطلاعات نادرست توسط هوش مصنوعی است. قوانین و شرایط استفاده باید به وضوح مشخص کنند که چگونه داده‌ها جمع‌آوری و استفاده می‌شوند و آیا این

توسعه هوش مصنوعی در ایران نیز انجام شده [۱۴]، ولی تاکنون پژوهشی از منظر اخلاق به توسعه هوش مصنوعی در کسب‌وکارهای ایرانی نپرداخته است. این پژوهش به دنبال پاسخ به این سؤال است که تا چه حد شرایط و مقررات استفاده از خدمات هوش مصنوعی که توسط شرکت‌های گوناگون تدوین شده، با اصول اخلاقی علمی در حوزه هوش مصنوعی مطابقت دارند؟

ساختار مقاله بدین شرح است: ابتدا مبانی نظری مرتبط با اصول اخلاقی در استفاده از هوش مصنوعی بحث می‌شود. بخش سوم به مرور و تحلیل پژوهش‌های مرتبط می‌پردازد، بخش چهارم به روش‌شناسی پژوهش اختصاص دارد و روش جمع‌آوری و تحلیل داده‌ها را در دو بخش کمی و کیفی شرح می‌دهد. در بخش پنجم، یافته‌های اصلی پژوهش ارائه و تحلیل می‌شوند که شامل شناسایی و تحلیل مصادیق موردنظر در شرایط و مقررات کسب‌وکارهای مبتنی بر هوش مصنوعی در ایران است. بخش ششم به بحث و تفسیر نتایج می‌پردازد و یافته‌های پژوهش در زمینه چارچوب‌های اخلاقی مقایسه می‌شوند. در نهایت، در بخش هفتم، نتیجه‌گیری و محدودیت‌های پژوهش ارائه می‌گردد.

۲ - مبانی نظری

در سال‌های اخیر، موضوع اخلاق در هوش مصنوعی در کسب‌وکارهای مبتنی بر این فناوری به یکی از موضوعات بحث‌برانگیز تبدیل شده است [۲، ۳]. مقالات از ابعاد متنوعی به این امر نگاه کرده و هریک به نحوی به تحلیل جنبه‌های اخلاقی هوش مصنوعی پرداخته‌اند. برخی بر مسئولیت‌پذیری و شفافیت در تصمیم‌گیری‌های هوش مصنوعی تأکید داشته‌اند [۱۵]، در حالی که دیگران بر عدالت، حفظ حریم خصوصی و پایداری تأکید کرده‌اند [۴، ۱۱]. در جدول ۱، مهمترین مقالات مرتبط با این موضوع فهرست شده‌اند. با وجود این تنوع دیدگاه‌ها، به نظر می‌رسد که همچنان نیاز به بازنویسی و بازتعریف اصول اخلاقی وجود دارد تا این اصول به شکلی جامع‌تر و کلی‌تر تدوین شوند [۵، ۱۶].

براین‌اساس، با ترکیب و تحلیل یافته‌های مقالات و بررسی اصول اخلاقی مطرح‌شده، اصول زیر به عنوان مبانی ارزیابی

¹ Accountability and Transparency

² Explainability and Interpretability

³ Justice and Fairness

⁴ Beneficence

⁵ Non-maleficence

جدول ۱) پژوهش‌های مرتبط با اصول اخلاقی هوش مصنوعی

نویسندگان	سال	تمرکز اصلی مقاله	اصول اخلاقی مورد اشاره
کامیلری [۱۵]	۲۰۲۴	حکمرانی هوش مصنوعی و مسئولیت اجتماعی	مسئولیت‌پذیری و شفافیت، قابلیت توضیح و تفسیر، برابری، حریم خصوصی و امنیت کاربران، امنیت و استحکام، پیشگیری از تبعیض الگوریتمی، پایداری
جی‌آرمالئو و همکاران [۱۶]	۲۰۲۴	بررسی نظام‌مند ادبیات اخلاقی در هوش مصنوعی	خیرخواهی، عدم ضرر، استقلال، عدالت، قابلیت توضیح
رابرتز و همکاران [۱]	۲۰۲۳	اخلاق و حکمرانی هوش مصنوعی در چین و اتحادیه اروپا	انسانیت، حریم خصوصی و حکمرانی داده‌ها، تنوع و عدم تبعیض، رفاه اجتماعی و محیطی، مسئولیت‌پذیری، شفافیت، استحکام فنی و ایمنی
بانکینز و فرموسا [۱۷]	۲۰۲۳	مفاهیم اخلاقی هوش مصنوعی برای کار معنادار	خیرخواهی، عدم ضرر، استقلال، عدالت، قابلیت توضیح
هرمان [۵]	۲۰۲۲	استفاده اخلاقی از هوش مصنوعی در بازاریابی	نیکوکاری، عدم آسیب‌رسانی، استقلال، عدالت، قابلیت توضیح
اشوک و همکاران [۳]	۲۰۲۲	چارچوب اخلاقی برای هوش مصنوعی و فناوری‌های دیجیتال	فهم‌پذیری، مسئولیت‌پذیری، مساوات، ارتقا رفاه، همبستگی، استقلال، کرامت و رفاه، امنیت، پایداری، حریم خصوصی، امنیت
دولگانوا [۱۳]	۲۰۲۱	بهبود تجربه مشتری با استفاده از هوش مصنوعی و با پیروی از اصول اخلاقی	مساوات، امنیت و بی‌ضرری، حریم خصوصی و محرمانگی، شفافیت و قابلیت توضیح، مسئولیت‌پذیری و پاسخگویی، استقلال و رضایت، پیشرفت و توسعه اخلاقی
ژانگ و همکاران [۱۱]	۲۰۲۱	اخلاق و حریم خصوصی هوش مصنوعی: درک از طریق بیبلیومتریک	برابری، حریم خصوصی داده‌ها، امنیت سایبری، اخلاق ماشین، عدم تبعیض
وزنیچ-آلویویچ و همکاران [۱۰]	۲۰۲۰	تأثیرات اجتماعی و اخلاقی هوش مصنوعی و چارچوب‌های سیاست‌گذاری اروپا	استقلال، کرامت انسانی، حریم خصوصی و حفاظت از داده‌ها، برابری، زندگی خوب و تنوع، مسئولیت‌پذیری و پاسخگویی، شفافیت و حسابرسی، دموکراسی و اعتماد
لو پیانو [۱۸]	۲۰۲۰	اصول اخلاقی در یادگیری ماشین و هوش مصنوعی	برابری، دقت، پاسخگویی، شفافیت، قابلیت توضیح، حریم خصوصی، امنیت، عدم ضرر
گرلیک و لیوزو [۴]	۲۰۲۰	ملاحظات اخلاقی و قانونی در تصمیم‌گیری الگوریتمی و قیمت‌گذاری شخصی‌سازی‌شده	مقابله با تقلب، برابری، عدالت اجتماعی، حریم خصوصی داده‌ها، عدم تبعیض
ژوبین و همکاران [۲]	۲۰۱۹	بررسی جهانی دستورالعمل‌های اخلاقی هوش مصنوعی	شفافیت، عدالت، برابری، عدم ضرر، مسئولیت‌پذیری، حریم خصوصی، نیکوکاری، آزادی و خودمختاری، اعتماد، پایداری، کرامت، همبستگی

قوانین و شرایط استفاده باید به مسائل محیط‌زیستی و استفاده پایدار از منابع توجه داشته باشند و راهکارهایی برای کاهش مصرف انرژی و کاهش ضایعات دیجیتال، از جمله داده‌های زائد و پردازش‌های غیرضروری ارائه دهند.

ح) **کرامت انسانی**^۳ به معنای احترام به ارزش ذاتی انسان و جلوگیری از نقض حقوق بشر است که فراتر از نقش او به‌عنوان مشتری یا کاربر قرار دارد. این اصل بر ضرورت همدلی، عدم نگاه ابزاری به انسان و توجه به شأن و اهمیت او تأکید دارد. قوانین و شرایط استفاده باید جایگاه انسان را به‌عنوان موجودی دارای شعور و احترام، در مرکز توجه قرار دهند و از هرگونه احساس بی‌اهمیتی، بهره‌برداری صرف از

اطلاعات به‌درستی به کاربران ارائه می‌شود. اطلاع‌رسانی دقیق درباره نحوه عملکرد هوش مصنوعی و پرهیز از بزرگ‌نمایی قابلیت‌ها از الزامات این اصل است.

و) **عدم مزاحمت**^۱ به معنای جلوگیری از هرگونه دخالت یا نفوذ ناخواسته در تعامل کاربران با سیستم‌های هوش مصنوعی است. و بر اهمیت احترام به مرزهای کاربران، اعم از فیزیکی، روانی یا اطلاعاتی تأکید دارد. این اصل همچنین بر طراحی هوش مصنوعی به‌گونه‌ای متمرکز است که تجربه کاربری بدون فشار، ابهام، یا مزاحمت تضمین شود.

ز) **پایداری**^۲ به معنای تضمین استفاده مسئولانه از منابع و حفاظت از محیط‌زیست در استفاده از هوش مصنوعی است.

^۱ Non-invasiveness

^۲ Sustainability

^۳ Dignity

از استفاده از داده‌هایشان منصفانه است، در حالی که نظریه تعادل قدرت-مسئولیت بر مسئولیت‌پذیری سازمان‌ها در مواجهه با قدرتشان در جمع‌آوری و پردازش داده‌ها تأکید دارد. همچنین، نظریه واکنش^۶ [۲۷] و نظریه تصمیم‌گیری رفتاری^۷ [۲۸، ۲۹] بر نقش کنترل و آزادی انتخاب کاربران تأکید می‌کنند و نشان می‌دهند که حس کنترل بر داده‌ها و شفافیت در سیاست‌ها می‌تواند نگرانی‌های حریم خصوصی را کاهش دهد و رفتار اشتراک‌گذاری داده را تسهیل کند.

۳ - پیشینه پژوهش

تحلیل اخلاقی هوش مصنوعی به یکی از موضوعات کلیدی در پژوهش‌های اخیر کسب‌وکارها تبدیل شده است. مطالعات انجام‌شده در این زمینه نشان می‌دهند که بسیاری از کسب‌وکارها در عمل با چالش‌هایی مانند شکاف میان اصول نظری و اجرایی، تعصب الگوریتمی و نبود شفافیت کافی مواجه هستند. این پژوهش‌ها از روش‌های گوناگون برای بررسی رعایت اصول اخلاقی استفاده کرده‌اند؛ برخی به تحلیل دقیق عملکرد شرکت‌ها پرداخته‌اند، درحالی‌که برخی دیگر تمرکز خود را بر دستورالعمل‌ها و چارچوب‌های نظری قرار داده‌اند.

مطالعه واکوری و همکاران (۲۰۲۲) با تمرکز بر ۳۹ شرکت نرم‌افزاری که در زمینه هوش مصنوعی فعالیت می‌کنند، نشان داد که اصولی همچون رفاه اجتماعی، تنوع و عدالت اجتماعی به‌ندرت رعایت می‌شوند. این پژوهش از روش تحلیل شکاف برای بررسی فاصله میان اصول نظری و عملکرد واقعی استفاده کرده است. نتایج نشان داد که بسیاری از شرکت‌ها، حتی با وجود دستورالعمل‌های بین‌المللی نظیر «دستورالعمل‌های اخلاقی برای هوش مصنوعی قابل اعتماد»، در اجرای عملی این اصول با چالش‌هایی مواجه هستند [۳۰]. در مطالعه کورنا و همکاران (۲۰۲۳)، تعداد ۲۰۰ دستورالعمل و گزاره اخلاقی مرتبط با هوش مصنوعی در سطح جهانی مورد تحلیل قرار گرفت. این پژوهش نشان داد که اصولی مانند شفافیت، عدالت و حفاظت از حریم خصوصی در این

کاربران، بی‌توجهی به اختیار و حقوق آن‌ها یا تبدیل افراد به ابزار سودآوری جلوگیری کنند.

ط) حریم خصوصی و حفاظت از داده‌ها^۱ به معنای حفاظت از داده‌های شخصی کاربران و جلوگیری از سوءاستفاده از آن‌ها است. قوانین و شرایط استفاده باید تضمین کنند که داده‌های شخصی به‌درستی محافظت شده و از آن‌ها سوءاستفاده نمی‌شود. سیاست‌های شفاف درباره جمع‌آوری، ذخیره‌سازی، و استفاده از داده‌ها و همچنین تدابیری برای جلوگیری از دسترسی غیرمجاز به این داده‌ها از نکات مورد تأکید این اصل است.

اصول اخلاقی ارائه‌شده پایه‌ای مهم برای ارزیابی شرایط و مقررات در کسب‌وکارهای مبتنی بر هوش مصنوعی فراهم می‌کنند. با این حال، برای تکمیل این تحلیل، بهره‌گیری از نظریات علمی می‌تواند ابعاد اخلاق در تعاملات میان کاربران و پلتفرم‌ها را روشن‌تر سازد. نظریات کلیدی در این حوزه هر یک بر جنبه‌هایی مانند عدالت، شفافیت، مسئولیت‌پذیری و کنترل تأکید دارند و راهبردهای مفهومی مفیدی برای بررسی پیچیدگی‌های اخلاقی در استفاده از هوش مصنوعی ارائه می‌دهند.

در تحلیل شرایط و مقررات و اخلاق هوش مصنوعی، شش نظریه کلیدی چارچوبی جامع برای درک روابط میان کاربران و پلتفرم‌ها ارائه می‌دهند. نظریه قرارداد اجتماعی^۲ [۱۹، ۲۰] و نظریه تبادل اجتماعی^۳ [۲۱] هر دو بر ایجاد ارزش متقابل در تعاملات تأکید دارند؛ به این معنا که کاربران زمانی اطلاعات شخصی خود را به اشتراک می‌گذارند که مزایای ملموس مانند خدمات شخصی‌سازی‌شده دریافت کنند. این نظریات بر اهمیت ارائه ارزش افزوده توسط پلتفرم‌ها برای تقویت اعتماد و مشارکت پایدار تأکید دارند.

از سوی دیگر، نظریه عدالت^۴ [۲۲-۲۴] و نظریه تعادل قدرت-مسئولیت^۵ [۲۵، ۲۶] به‌طور متداخل بر شفافیت و انصاف در استفاده از داده‌ها تمرکز دارند. عدالت رویه‌ای و توزیعی به کاربران اطمینان می‌دهد که فرایندها و نتایج حاصل

^۱ Data Privacy and Protection

^۲ Social Contract Theory

^۳ Social Exchange Theory

^۴ Justice Theory

^۵ Power-Responsibility Equilibrium Theory

^۶ Reaction Theory

^۷ Behavioral Decision Theory

پژوهش واکوری در شناسایی شکاف‌های اجرایی مؤثرتر بوده‌اند، در حالی که پژوهش‌های کلان‌نگر نظیر کورنا و بتال به اصول کلی و چالش‌های سیستماتیک توجه داشته‌اند.

۴- روش‌شناسی

این تحقیق با هدف تحلیل اخلاقی شرایط و مقررات کسب‌وکارهای مبتنی بر هوش مصنوعی در ایران، از راهبرد پژوهش ارزشیابی استفاده کرده که به‌عنوان رویکردی ساختارمند برای تحلیل میزان تطابق شرایط و مقررات با چارچوب‌های اخلاقی مرتبط با هوش مصنوعی به کار گرفته شده است [۳۵]. پژوهش در سال ۱۴۰۳ اجرا شده و طی آن داده‌ها در دو مرحله کمی و کیفی بررسی شدند.

برای انتخاب نمونه‌ها از روش نمونه‌گیری هدفمند استفاده شد. در این پژوهش، ۳۰ کسب‌وکار بررسی شدند که به طور خاص از هوش مصنوعی در تعامل با مشتریان خود بهره می‌برند. این کسب‌وکارها یا در بخش «درباره ما» وبگاه رسمی خود به استفاده از هوش مصنوعی اشاره کرده‌اند، یا نام آن‌ها حداقل در سه فهرست موجود از کسب‌وکارهای فعال در حوزه هوش مصنوعی ذکر شده تا از صحت انتخاب آن‌ها اطمینان حاصل شود.

شرایط و مقررات این کسب‌وکارها از منابع رسمی مانند وبگاه‌ها و مستندات منتشرشده جمع‌آوری شده تا داده‌های استفاده‌شده معتبر و قابل‌استناد باشند. برای ارزشیابی تطابق شرایط و مقررات با اصول اخلاقی، از طیف لیکرت پنج‌تایی استفاده شد. هر اصل اخلاقی به‌صورت سؤالات یا شاخص‌هایی طراحی شد و متن شرایط و مقررات بررسی‌شده برای هر شاخص امتیازدهی گردید. امتیازات به دست آمده برای هر کسب‌وکار تجمیع و وارد فرایند رتبه‌بندی شدند. برای رتبه‌بندی نهایی، از روش تاپسیس استفاده شد. در جدول ارزشیابی، نتایج به‌صورت رنگی ارائه شده‌اند. در این جدول از طیف رنگی سبز برای نشان دادن سطوح مختلف عملکرد مثبت استفاده شده است، به‌گونه‌ای که رنگ سبز پررنگ نشان‌دهنده بالاترین میزان تطابق و سبز کم‌رنگ نمایانگر تطابق اندک است. رنگ قرمز نیز به‌عنوان پایین‌ترین میزان تطابق و عملکرد ضعیف‌تر در نظر گرفته شده است.

دستورالعمل‌ها به طور گسترده‌ای تأکید شده‌اند. اما این دستورالعمل‌ها در بیشتر موارد فاقد جزئیات کاربردی برای مشتریان هستند. این مطالعه، علی‌رغم مقیاس گسترده خود، به تعداد مشخصی از شرکت‌ها نپرداخته و بیشتر بر تحلیل اصول و دستورالعمل‌ها متمرکز بوده است [۳۱].

مطالعه آرورا و تاتا (۲۰۲۴) به بررسی کاربرد هوش مصنوعی در تحلیل داده‌های بزرگ پرداخته و مسائل اخلاقی نظیر تعصب الگوریتمی، ذخیره‌سازی داده‌ها و نظارت تهاجمی را موردتوجه قرار داده است [۳۲]. در مقایسه با این پژوهش، مطالعه جین [۳۳] به‌صورت فنی‌تر به بررسی کاربرد هوش مصنوعی در شخصی‌سازی خدمات مشتری پرداخت و نشان داد که علی‌رغم استفاده گسترده از این فناوری، مشکلات مرتبط با حریم خصوصی و اخلاق همچنان چالش‌برانگیز است. این مطالعه بر استفاده از تکنیک‌هایی مانند رمزنگاری و یادگیری فدرال^۱ برای حفاظت از داده‌های کاربران تأکید کرده است.

مطالعه بتال و گلدنفین [۳۴] با رویکردی کل‌نگر به بررسی تعارض میان رویکردهای لیبرال در تنظیم قوانین و واقعیت‌های تکنو-اقتصادی و تکنو-سیاسی هوش مصنوعی پرداخت. این پژوهش بر اهمیت واسطه‌های داده و تعاونی‌های توجه^۲ به‌عنوان مدل‌هایی نو برای حفظ خودمختاری کاربران تأکید کرد. تعاونی‌های توجه به مدل‌هایی اشاره دارند که در آن کاربران به‌صورت جمعی درباره نحوه استفاده از داده‌ها و توجه خود در پلتفرم‌های دیجیتال تصمیم‌گیری می‌کنند و با جایگزینی مدل‌های اقتصادی تبلیغ‌محور، شفافیت و خودمختاری بیشتری برای کاربران فراهم می‌آورند. برخلاف مطالعات واکوری و کورنا که بر تحلیل‌های عملیاتی و اجرایی تمرکز داشتند، این پژوهش ابعاد سیستمی و ساختاری مرتبط با هوش مصنوعی را مورد بررسی قرار داد. به‌طور کلی، مطالعات عملیاتی مانند

^۱ یادگیری فدرال (Federated Learning) روشی در یادگیری ماشین است که مدل‌های هوش مصنوعی را به‌صورت غیرمتمرکز و بدون نیاز به ارسال داده‌های خام کاربران به سرور مرکزی آموزش می‌دهد. در این روش، تنها وزن‌های مدل یا به‌روزرسانی‌های الگوریتم به اشتراک گذاشته می‌شود، که به حفظ حریم خصوصی کاربران کمک می‌کند.

^۲ Attention Cooperatives

جی پی تی به عنوان یک نمونه، رویکردی حداقلی اتخاذ کرده و مسئولیت تمامی محتواهای تولیدشده را به کاربران منتقل کرده است. این اپلیکیشن صراحتاً اعلام می‌کند که «مسئولیت محتواهای تولید شده با وان جی پی تی بر عهده کاربران می‌باشد».

گپ جی پی تی تأکید می‌کند که خروجی‌ها توسط یک سیستم هوشمند تولید شده و نباید به انسان منتسب شوند. در این سیاست آمده است: «داده‌های خروجی به وسیله سیستم هوشمند غیرانسانی تولید می‌شود و در نتیجه کاربران نباید داده‌های خروجی را منتسب به انسان بدانند». این شفافیت در مورد ماهیت داده‌ها قابل تقدیر است، اما عدم مسئولیت‌پذیری در قبال پیامدهای ناشی از استفاده نادرست از داده‌ها نقطه ضعف مهمی محسوب می‌شود.

روبو مسئولیت محتوا را به طور کامل از خود سلب و اعلام کرده که «روبو هیچ‌گونه مسئولیتی در خصوص محتوای سؤالات و پاسخ‌های این سرویس ندارد». این سیاست با اشاره به وابستگی روبو به سرویس‌هایی مانند چت‌جی‌پی‌تی، مسئولیت را به طرف‌های ثالث واگذار کرده است.

ب) قابلیت توضیح و تفسیر: در این بخش، خدمات ارائه‌شده توسط دیجی مارک و مد‌آی تحلیل شده‌اند.

دیجی مارک خدمات خود را با استفاده از واسط برنامه‌نویسی شرکت‌های پیشرو ارائه می‌دهد. عبارت «خدمات خود را با بهره‌گیری از API‌های ارائه‌شده توسط شرکت‌های پیشرو» در شرایط و مقررات، فاقد شفافیت کافی در خصوص نوع واسط برنامه‌نویسی، قابلیت‌ها و نحوه تعامل این پلتفرم با آن‌ها است. مد‌آی بر گردآوری محتوا از اینترنت و اینستاگرام با استفاده از هوش مصنوعی تمرکز دارد. در سیاست این پلتفرم آمده است که «اکثر محتواهای موجود در مد‌آی توسط هوش مصنوعی از فروشگاه‌های اینترنتی و اینستاگرامی جمع‌آوری می‌شود که در مواردی احتمال اشتباه در آنها وجود دارد». این شفافیت در اشاره به احتمال وجود خطا در داده‌ها، یک گام مثبت در راستای ایجاد انتظارات واقعی برای کاربران است. با این حال، جزئیاتی درباره سازوکارهای گردآوری داده‌ها، چگونگی ارزیابی صحت اطلاعات و رویه‌های اصلاح اشتباهات احتمالی ارائه نشده است.

برای تحلیل عمیق‌تر، از روش تحلیل محتوای کیفی به منظور شناسایی مضامین و الگوهای مرتبط با اصول اخلاقی بهره گرفته شد. فرایند تحلیل شامل کدگذاری سیستماتیک داده‌ها بود که در آن مفاهیم کلیدی شناسایی و دسته‌بندی شدند. سپس، این مفاهیم با استفاده از چارچوب نظری تحقیق ارزشیابی شدند. برای اطمینان از دقت و انسجام تحلیل، دستورالعمل‌های مایلز و هوبرمن به کار گرفته شد. تمامی اسناد و متن‌های شرایط و مقررات به صورت کامل بررسی و دوباره خوانی شدند تا از جامعیت و پوشش تمام جوانب داده‌ها اطمینان حاصل شود [۳۶].

۵- یافته‌های پژوهش

۵-۱ تحلیل کمی

در بخش کمی، ابتدا با مطالعه دقیق شرایط و مقررات هر پلتفرم، به هر یک از اصول پژوهش امتیازی بر اساس طیف لیکرت اختصاص داده شد. پس از امتیازدهی، به کمک چهار نفر از متخصصان خط‌مشی‌گذاری تجاری برای وزن‌دهی به هریک از اصول کمک گرفته شد. سپس، با استفاده از روش تاپسیس ماتریس تصمیم‌گیری تشکیل شد و پس از آن نرمال‌سازی انجام گرفت. در مرحله بعد وزن هریک از اصول در پلتفرم‌ها محاسبه شد و ماتریس نرمال‌سازی شده و وزن‌دار شده به دست آمد. پس از تعیین ماتریس ایده‌آل مثبت و منفی، فاصله هر گزینه از ماتریس ایده‌آل مثبت و منفی محاسبه شد. در این داده‌ها تمامی مقادیر جنبه مثبت داشتند. در نهایت، ضریب نسبی نزدیکی به راه‌حل ایده‌آل محاسبه و در ستون آخر جدول اضافه شد و براین اساس اپلیکیشن‌ها رتبه‌بندی شدند (جدول ۲).

۵-۲ تحلیل کیفی

تحلیل کیفی شرایط و مقررات کسب‌وکارها ذیل اصول اخلاقی منتخب در این مقاله انجام شده است. مواردی که ذیل این اصول ذکر شده‌اند، صرفاً به عنوان نمونه هستند و هدف از آن‌ها توضیح بیشتر نحوه تحلیل است، نه پوشش تمامی مصادیق.

الف) مسئولیت‌پذیری و شفافیت: در میان کسب‌وکارها، وان

جدول ۲) رتبه‌بندی کسب‌وکارها بر اساس میزان انطباق شرایط و مقررات با اصول اخلاقی

اصول (وزن) پلتفرم‌ها	مسئولیت پذیری و شفافیت (۰.۱۵)	قابلیت توضیح و تفسیر (۰.۱۲)	عدالت و برابری (۰.۱۲)	سودمندی (۰.۱۳)	عدم فریب‌کاری (۰.۱)	عدم مزااحت (۰.۰۴)	پایداری (۰.۰۶)	کرامت انسانی (۰.۱۳)	حریم خصوصی و حفاظت از داده‌ها (۰.۱۵)	Pi
روینو	۱	۱	۱	۲	۵	۴	۱	۳	۵	0.603213684
دستیار	۱	۱	۱	۲	۴	۲	۱	۲	۵	0.522874248
زیرک	۵/۲	۱	۱	۲	۳	۱	۱	۲	۲	0.516773515
نوبین هاب	۱	۱	۱	۱	۵	۲	۱	۲	۵	0.482257011
گپ جی پی تی	۲	۱	۱	۱	۴	۱	۱	۲	۳	0.45935823
زیگپ	۲	۱	۱	۱	۴	۱	۱	۱	۵/۴	0.45176499
هوشک	۲	۱	۱	۱	۴	۳	۱	۱	۴	0.439862712
سهاب	۱	۱	۱	۱	۵/۴	۵/۳	۱	۲	۴	0.43911829
نویسه نگار	۱	۱	۱	۱	۵/۴	۵/۳	۱	۲	۴	0.43911829
مدآی	۱	۲	۱	۱	۵/۴	۲	۱	۱	۴	0.425576063
دیجی مارک	۱	۵/۱	۱	۱	۵	۲	۱	۱	۵/۴	0.424878348
دکتر نکست	۱	۱	۱	۱	۵	۵	۱	۱	۵/۴	0.41839488
فراشناسا	۱	۱	۱	۱	۵/۴	۳	۱	۱	۵/۴	0.395690577
Avalai	۱	۱	۱	۱	۴	۴	۱	۱	۵/۴	0.391203519
وان جی پی تی	۵/۲	۱	۱	۱	۵/۱	۱	۱	۱	۲	0.362770888
تاک بات	۱	۱	۱	۱	۱	۱	۱	۳	۱	0.354936618
آوانگار	۱	۱	۱	۱	۳	۳	۱	۱	۴	0.334246431
آواشو	۱	۱	۱	۱	۳	۳	۱	۱	۴	0.334246431
روبو	۲	۱	۱	۱	۱	۱	۱	۱	۵/۲	0.298648375
روبو وکیل	۱	۱	۱	۱	۵/۴	۵/۲	۱	۱	۲	0.265942597
پیشکاربات	۲	۱	۱	۱	۱	۱	۱	۱	۱	0.247833744
ویرا	۱	۱	۱	۱	۲	۱	۱	۱	۳	0.226094607
رخشای	۱	۱	۱	۱	۱	۱	۱	۱	۵/۲	0.165187313
جیبیت	۱	۱	۱	۱	۱	۲	۱	۱	۲	0.122160931
ایبو	۱	۱	۱	۱	۱	۱	۱	۱	۲	0.114157011
فیوناچی	۱	۱	۱	۱	۱	۱	۱	۱	۲	0.114157011

0.088354774	۱	۱	۱	۲	۲	۱	۱	۱	۱	بایتیکل
0.077005372	۱	۱	۱	۱	۲	۱	۱	۱	۱	آتنا
0	۱	۱	۱	۱	۱	۱	۱	۱	۱	بولوز
0	۱	۱	۱	۱	۱	۱	۱	۱	۱	ویزیسان

ج) عدالت و برابری علی‌رغم اهمیت بالا، در شرایط و مقررات هیچ یک از کسب‌وکارها ذکر نشده‌اند.

د) سودمندی: بررسی شرایط و مقررات اپلیکیشن‌های دستیار و زیرک نشان می‌دهد که این پلتفرم‌ها با اهدافی به دنبال ارائه سودمندی برای کاربران و جامعه هستند. دستیار در سیاست خود تأکید دارد که «همواره به اصل کاربردی‌بودن در طراحی و توسعه ابزارهای قرار داده شده در دستیار توجه کرده‌ایم و در تلاشیم تا زندگی دیجیتال کاربران را از هر زاویه بهبود ببخشیم». این جمله نشان‌دهنده تعهدی قابل‌تحسین به طراحی کاربرمحور و تمرکز بر بهبود زندگی دیجیتال کاربران است. با این حال، گستردگی مفهوم «زندگی دیجیتال» و عدم ارائه شاخص‌های قابل‌اندازه‌گیری برای این بهبود، تحلیل سودمندی این هدف را دشوار می‌کند. زیرک، بر نوآوری تأکید کرده و بیان می‌کند: «ما همیشه دغدغه این را داشته‌ایم که هم‌پای هم‌نوعان خود در سراسر جهان بتوانیم در سریع‌ترین زمان امکان استفاده از به‌روزترین فناوری‌ها را داشته باشیم و از آن برای خلق ارزش‌های جدید استفاده کنیم». اما، عبارت «خلق ارزش‌های جدید» نیازمند تعریف دقیق‌تر و ارائه مثال‌های مشخص است. نبود شفافیت در نوع ارزش‌های خلق‌شده و تأثیرات آن‌ها بر کاربران، این سیاست را بیشتر به یک شعار بازاریابی نزدیک می‌کند تا یک تعهد عملی.

ه) عدم فریب‌کاری: تحلیل سیاست‌های آتنا و دیجی مارک در زمینه جلوگیری از فریب‌کاری نشان می‌دهد که این دو پلتفرم تلاش کرده‌اند تا با شفاف‌سازی در مورد داده‌های جمع‌آوری‌شده و نحوه استفاده از آن‌ها، اعتماد کاربران را جلب کنند. آتنا در شرایط خود صراحتاً به نوع داده‌های جمع‌آوری‌شده اشاره کرده است: «اطلاعات حساب کاربری برای شناسایی شما و ارتقای تجربه کاربری»، «اطلاعات چت

برای بهبود عملکرد و توسعه سیستم در پاسخگویی» و «اطلاعات دستگاه برای بهینه‌سازی تطبیقی ارائه محتوا». این سطح از شفافیت در شناسایی داده‌ها و اهداف مرتبط با آن‌ها، یک گام مثبت در جهت جلوگیری از فریب‌کاری و افزایش اعتماد کاربران است. با این حال، سیاست آتنا فاقد جزئیات درباره نحوه محافظت از داده‌های حساس یا مکانیزم‌هایی برای پیشگیری از سوءاستفاده‌های احتمالی است. همچنین، استفاده از داده‌های چت برای «توسعه سیستم پاسخگویی» ممکن است نگرانی‌هایی در مورد حریم خصوصی کاربران ایجاد کند که نیازمند توضیح بیشتر درباره روش‌های مدیریت و حفاظت از این داده‌ها است. دیجی مارک بیان کرده اطلاعاتی را برای «بهبود خدمات و پشتیبانی مشتری»، «تحلیل و بهینه‌سازی تجربه کاربری» و «شناسایی، پیشگیری و رفع مشکلات فنی» جمع‌آوری می‌کند. این اهداف، هرچند شفاف هستند، اما ممکن است برای برخی کاربران، به‌ویژه در بخش‌هایی نظیر ارسال پیامک و ایمیل‌های دوره‌ای، نگرانی‌هایی درباره احتمال استفاده تبلیغاتی از داده‌ها ایجاد کنند.

و) عدم مزاحمت: آوانگار تصریح می‌کند که «صرفاً در صورتی و در زمانی، اطلاعات شخصی کاربران را درخواست می‌کند یا این اطلاعات را پردازش می‌کند که برای ارائه خدمات، یا ارائه بهتر خدمات، به ایشان ضروری باشد». این سیاست با وجود تمرکز بر کاهش تعاملات غیرضروری، فاقد جزئیات در مورد نحوه ارزیابی این ضرورت‌ها است. هوشک نیز تأکید می‌کند که اطلاعات کاربران «به‌اندازه‌ای که برای پایبندی به تعهدات قانونی ضروری است»، اما توضیحات بیشتری درباره نحوه مدیریت اطلاعات اضافی ارائه نشده است. در زمینه کنترل تبلیغات و اعلان‌ها، دکتر نکست حق ارسال پیامک، ایمیل و اعلان‌های مختلف را برای خود

محفوظ می‌داند، اما تأکید می‌کند که کاربران می‌توانند «ارتباط دهی از هر یک از طرق یاد شده را قطع نموده یا درخواست قطع آن ارسال نمایند». این شفافیت در بیان گزینه پیش‌فرض و ارائه ابزارهای کنترلی برای کاربران، گامی مثبت در کاهش مزاحمت‌های تبلیغاتی است.

ز) پایداری نیز همچون عدالت و برابری، در شرایط و مقررات هیچ یک از کسب‌وکارها ذکر نشده است.

ح) کرامت انسانی: کسب‌وکارها در شرایط و مقررات خود کمتر به این موضوع پرداخته‌اند. تاک‌بات در این زمینه اعلام می‌کند که «وبگاه اجازه نمی‌دهد که محتوای نژادپرستانه، تهدیدآمیز، توهین‌آمیز، جنسی، خشونت‌آمیز، تبلیغاتی، تقلبی و یا هرگونه محتوای غیرقانونی قرار داده شود». این رویکرد، حاکی از تعهد به کرامت انسانی و حفاظت از فضای عمومی پلتفرم است. با این حال، سیاست تاک‌بات نیز فاقد جزئیات عملیاتی برای نظارت مداوم و حذف محتوای ممنوعه است. در زمینه عدم تبعیض بین کاربران، تاک‌بات سیاست «مصرف منصفانه» را اعمال می‌کند و تصریح می‌کند که «هیچ اولویتی میان کاربران در نظر گرفته نمی‌شود». این سیاست تضمین می‌کند که همه کاربران، بدون در نظر گرفتن میزان خرید یا استفاده از خدمات، به‌صورت برابر با پلتفرم تعامل دارند.

ط) حریم خصوصی و حفاظت از داده‌ها در اپلیکیشن‌های مختلف مانند دستیار، نوین‌هاب، دکتر نکست و رویینو نشان‌دهنده تنوع در رویکردها و تعهدات آن‌ها نسبت به حفظ امنیت اطلاعات کاربران است. تضمین محرمانگی اطلاعات یکی از موضوعات کلیدی در سیاست‌های این اپلیکیشن‌ها است. دستیار با تأکید بر این که «امنیت و مراقبت از اطلاعات خط قرمز ماست»، تعهد خود به حفظ محرمانگی را نشان می‌دهد، اما فاقد جزئیاتی درباره رویه‌های اجرایی برای دستیابی به این هدف است. عدم استفاده غیرمجاز و افشاکری قانونی نیز به‌عنوان یکی از اصول کلیدی در این سیاست‌ها مطرح شده است. نوین‌هاب اعلام می‌کند که «جز به حکم قانون یا دستور مقام صالح قضایی، اطلاعات کاربر در اختیار هیچ فرد و یا سازمانی قرار داده نخواهد شد». این تعهد، قابل تحسین است، اما باید توضیحات بیشتری در مورد نوع داده‌ها و شرایط افشا مشخص کند.

در مورد به‌اشتراک‌گذاری داده‌های غیرشخصی، رویینو اعلام کرده که «اطلاعات غیرشخصی تجمعی ممکن است با همکاران تجاری به اشتراک گذاشته شود». اگرچه این داده‌ها کاربران را به طور فردی شناسایی نمی‌کنند، سیاست مذکور نیازمند شفاف‌سازی درباره نحوه تضمین محرمانگی این داده‌ها در فرایند اشتراک‌گذاری است. در زمینه معرفی اشخاص ثالث برای دسترسی به اطلاعات، دکتر نکست توضیح می‌دهد که ممکن است «بخشی ضروری از داده‌های شما را در اختیار خدمات‌دهندگان فنی ثالث قرار دهیم». این شفافیت، قدمی مثبت است، اما نیازمند توضیح بیشتری درباره نوع داده‌ها و شرایط استفاده آن‌ها توسط اشخاص ثالث است. از سوی دیگر، رویینو ضمن اشاره به «دسترسی اعضای کلیدی تیم و قراردادهای به اطلاعات کاربران»، تصریح کرده که این اشخاص به توافق‌نامه‌های محرمانگی متعهد هستند. این رویکرد، شفافیت بیشتری دارد، اما نحوه نظارت و تضمین اجرای این توافق‌نامه‌ها نامشخص است.

۶- بحث

یافته‌های این پژوهش نشان می‌دهند که بسیاری از کسب‌وکارهای ایرانی، علی‌رغم بهره‌گیری از فناوری‌های پیشرفته هوش مصنوعی، در تدوین شرایط و مقررات، ناکارآمد عمل کرده‌اند. نظریه‌های علمی، دریچه‌ای برای درک رفتار کاربران ارائه می‌دهند و به این موضوع کمک می‌کنند تا شرایط و مقررات متناسب با نیازها و ویژگی‌های هر پلتفرم بازطراحی شوند. در ادامه، چند مثال از این بازطراحی به‌عنوان پیشنهاد ارائه می‌شود.

بر اساس نظریه قرارداد اجتماعی، شرایط و مقررات باید بر ایجاد ارزش متقابل میان کسب‌وکارها و کاربران تأکید کنند. این نظریه بر این باور است که کاربران تنها در صورتی داده‌های خود را به اشتراک می‌گذارند که منافع ملموس و مشخصی دریافت کنند [۱۹، ۲۰]. به‌عنوان مثال، سیاست‌های شفاف می‌توانند این جمله را شامل شوند: «ما اطلاعات شما را برای ارائه پیشنهادها و شخصی‌سازی شده، از جمله ارسال تخفیف‌های ویژه متناسب با علایق شما، اولویت در دسترسی به محصولات جدید و ارائه محتوای آموزشی رایگان استفاده

در نهایت، نظریه واکنش بر اهمیت حفظ آزادی انتخاب کاربران تأکید دارد [۲۷]. سیاست‌های طراحی شده می‌توانند عباراتی مانند این جمله را در بر داشته باشند: «شما می‌توانید هر زمان که بخواهید، اطلاعات خود را از سیستم ما حذف کنید. درخواست شما از طریق پنل کاربری یا ایمیل پشتیبانی ثبت شده و ظرف ۴۸ ساعت تأییدیه حذف اطلاعات از طریق ایمیل برای شما ارسال خواهد شد. در این فرآیند، تمامی اطلاعات شناسایی شده به‌صورت دائمی از پایگاه‌های داده ما پاک خواهد شد».

۷- نتیجه‌گیری

این پژوهش با هدف تحلیل شرایط و مقررات کسب‌وکارهای مبتنی بر هوش مصنوعی در ایران، به شناسایی و ارزیابی انطباق این سیاست‌ها با اصول اخلاقی پرداخته است. نتایج نشان داد که اگرچه برخی از کسب‌وکارها تلاش‌هایی برای بهبود شفافیت و حفظ حریم خصوصی انجام داده‌اند، اما کاستی‌های قابل‌توجهی در زمینه‌های شفافیت، عدالت، و مسئولیت‌پذیری وجود دارد. این سیاست‌ها غالباً فاقد جزئیات عملیاتی و سازکارهای نظارتی بوده و به‌جای ایجاد ارزش متقابل و اعتمادسازی، مسئولیت پیامدها را به کاربران منتقل کرده‌اند.

مطالعات پیشین، نظیر پژوهش واکوری و همکاران (۲۰۲۲) و کورنا و همکاران (۲۰۲۳)، نیز بر اهمیت شکاف میان اصول نظری و عملکرد واقعی تأکید دارند. یافته‌های این پژوهش نشان می‌دهد که سیاست‌های ایرانی در مقایسه با نمونه‌های بین‌المللی، بیشتر بر کلیات متمرکز بوده و کمتر به جزئیات فنی و تعهدات عملی پرداخته‌اند. این مقایسه نشان می‌دهد که رویکردهای نوآورانه نظیر یادگیری فدرال که در سطح بین‌المللی توصیه می‌شوند، در سیاست‌های داخلی نادیده گرفته شده‌اند.

این پژوهش با چند محدودیت روبه‌رو بوده است که ممکن است بر تعمیم‌پذیری و گستردگی نتایج تأثیر بگذارد. یکی از محدودیت‌های اصلی، محدودیت در دسترسی به داده‌ها بود. شرایط و مقررات بررسی شده عمدتاً از اسناد رسمی و منابع عمومی استخراج شده‌اند که ممکن است اطلاعات دقیق‌تری

می‌کنیم. این تخفیف‌ها از طریق پیامک و ایمیل به شما اطلاع‌رسانی خواهند شد.» یافته‌های این پژوهش نشان داد که چنین مواردی در سیاست‌های ایرانی به‌ندرت ذکر می‌شوند و این باعث کاهش اعتماد کاربران می‌شود.

نظریه عدالت، نشان می‌دهد که کاربران انتظار دارند فرایندهای جمع‌آوری و استفاده از داده‌ها عادلانه و شفاف باشد [۲۲-۲۴، ۳۷]. به‌عنوان مثال، سیاست‌های طراحی شده بر اساس این نظریه می‌توانند شامل عباراتی نظیر: «تمام داده‌های جمع‌آوری شده توسط الگوریتم‌های ما، هر سه ماه یک‌بار، توسط نهادهای مستقل نظیر سازمان ملی هوش مصنوعی بازبینی می‌شوند تا از عدم وجود تعصب و انطباق با قوانین اطمینان حاصل شود. در صورت مشاهده هرگونه تخلف یا تعصب، تیم فنی ما تغییرات لازم را طی دو هفته اعمال کرده و گزارش اصلاحات در بخش شفافیت وبسایت ما منتشر خواهد شد».

نظریه تعادل قدرت-مسئولیت تأکید دارد که قدرت گسترده کسب‌وکارها در جمع‌آوری و پردازش داده‌ها باید با مسئولیت‌پذیری اجتماعی متناسب باشد [۲۵، ۲۶]. به‌عنوان نمونه، شرایط و مقررات می‌توانند شامل چنین جمله‌ای باشند: «در صورت وقوع هرگونه نقض امنیت اطلاعات، شما ظرف ۲۴ ساعت از طریق ایمیل و پیامک مطلع خواهید شد. گزارشی شامل نوع داده‌های تحت تأثیر، علت نقض و اقداماتی که برای جلوگیری از تکرار آن انجام داده‌ایم، در بخش امنیت وبسایت ما منتشر خواهد شد. ما همچنین از فناوری‌های رمزنگاری پیشرفته مانند AES-256 برای حفاظت از داده‌ها استفاده می‌کنیم که مطابق با استانداردهای جهانی است».

نظریه تبادل اجتماعی بیان می‌کند که کاربران زمانی داده‌های خود را به اشتراک می‌گذارند که مزایا به‌وضوح از هزینه‌ها بیشتر باشند [۲۱]. در شرایط و مقررات می‌توان این‌گونه بیان کرد: «ما هر سه ماه یک گزارش شخصی‌سازی شده از الگوهای استفاده و خرید شما، شامل پیشنهادات ویژه برای کاهش هزینه‌های شما، ارائه خواهیم کرد. این گزارش از طریق پنل کاربری و ایمیل در دسترس شما قرار می‌گیرد».

Kogan Page. <https://www.koganpage.com/risk-compliance/data-ethics-9781398610279>

[7] Kriebitz, A., & Lütge, C. (2020). **Artificial intelligence and human rights: A business ethical assessment**. *Business and Human Rights Journal*, 5(1), 84–104. <https://doi.org/10.1017/bhj.2019.28>

[8] Ahmad, S., Ahmadi, A. M., Abedinzadeh, Z., Fatemi, F., Roknadini, F., Mirzaei, M., et al. (2024). **Generative artificial intelligence**. *Journal of Science and Technology Policy*, 17(0), 1–100. {In Persian}. https://jstp.nrisp.ac.ir/article_14031.html?lang=en

[9] Katirai, A. (2023). **The ethics of advancing artificial intelligence in healthcare: Analyzing ethical considerations for Japan's innovative AI hospital system**. *Frontiers in Public Health*, 11, Article 1142062.

<https://doi.org/10.3389/fpubh.2023.1142062>

[10] Vesnic-Alujevic, L., Nascimento, S., & Pólvara, A. (2020). **Societal and ethical impacts of artificial intelligence: Critical notes on European policy frameworks**. *Telecommunications Policy*, 44(6), Article 101961.

<https://doi.org/10.1016/j.telpol.2020.101961>

[11] Zhang, Y., Wu, M., Tian, G. Y., Zhang, G., & Lu, J. (2021). **Ethics and privacy of artificial intelligence: Understandings from bibliometrics**. *Knowledge-Based Systems*, 222, Article 106994. <https://doi.org/10.1016/j.knosys.2021.106994>

[12] Diakopoulos, N. (2015). **Algorithmic accountability: Journalistic investigation of computational power structures**. *Digital Journalism*, 3(3), 398–415.

<https://doi.org/10.1080/21670811.2014.976411>

[13] Dolganova, O. I. (2021). **Improving customer experience with artificial intelligence by adhering to ethical principles**. *Business Informatics*, 15(2), 34–46. <https://doi.org/10.17323/2587-814X.2021.2.34.46>

[14] Kanani, F., Rasoulzadeh, P., Hafezi, R., & Ahangari, S. S. (2023). **Analysis of the artificial intelligence ecosystem in Iran and identifying institutional and functional gaps**. *Journal of Science and Technology Policy*, 16(2), 59–77. {In Persian}. https://jstp.nrisp.ac.ir/article_13993.html?lang=en

[15] Camilleri, M. A. (2024). **Artificial intelligence governance: Ethical considerations and implications for social responsibility**. *Expert Systems*, 41(7), Article e13406. <https://doi.org/10.1111/exsy.13406>

[16] Giarmoleo, F. V., Ferrero, I., Rocchi, M., & Pellegrini, M. M. (2024). **What ethics can say on artificial intelligence: Insights from a systematic literature review**. *Business and Society Review*, 129(2), 258–292. <https://doi.org/10.1111/basr.12336>

[17] Bankins, S., & Formosa, P. (2023). **The ethical implications of artificial intelligence (AI) for meaningful work**. *Journal of Business Ethics*, 185(4), 725–740. <https://doi.org/10.1007/s10551-023-05339-7>

[18] Lo Piano, S. (2020). **Ethical principles in machine learning and artificial intelligence: Cases from the field and possible ways forward**. *Humanities and Social Sciences Communications*, 7(1), 1–7. <https://doi.org/10.1057/s41599-020-0501-9>

در مورد نحوه اجرای این سیاست‌ها در عمل فراهم نکند. این امر باعث می‌شود تحلیل‌های ارائه‌شده بیشتر در سطح نظری باقی بماند. محدودیت دیگر، تمرکز بر پلتفرم‌های ایرانی بود که ممکن است نتایج را به زمینه‌های خاص فرهنگی و حقوقی محدود کند. اگرچه این تمرکز به درک عمیق‌تر از وضعیت ایران کمک کرده است، اما مقایسه با سیاست‌های مشابه در سایر کشورها می‌توانست دیدگاه جامع‌تری ارائه دهد. علاوه بر این، نبود مشارکت مستقیم کاربران یا نمایندگان پلتفرم‌ها در فرایند جمع‌آوری داده‌ها، باعث شد برخی از جنبه‌های کلیدی، مانند تأثیر واقعی سیاست‌ها بر رفتار کاربران یا چالش‌های عملیاتی، کمتر مورد توجه قرار گیرد. ترکیب این روش با روش‌های دیگر، مانند نظرسنجی از کاربران، می‌تواند به ارائه نتایج جامع‌تر کمک کند.

تعارض منافع

نویسندگان تعهد می‌کنند که هیچ تعارض منافی در این مقاله وجود نداشته‌است.

References

- [1] Roberts, H., Cowls, J., Hine, E., Morley, J., Wang, V., Taddeo, M., et al. (2023). **Governing artificial intelligence in China and the European Union: Comparing aims and promoting ethical outcomes**. *The Information Society*, 39(2), 79–97. <https://doi.org/10.1080/01972243.2022.2124565>
- [2] Jobin, A., Ienca, M., & Vayena, E. (2019). **The global landscape of AI ethics guidelines**. *Nature Machine Intelligence*, 1(9), 389–399. <https://doi.org/10.1038/s42256-019-0088-2>
- [3] Ashok, M., Madan, R., Joha, A., & Sivarajah, U. (2022). **Ethical framework for artificial intelligence and digital technologies**. *International Journal of Information Management*, 62, Article 102433. <https://doi.org/10.1016/j.ijinfomgt.2021.102433>
- [4] Gerlick, J. A., & Liozu, S. M. (2020). **Ethical and legal considerations of artificial intelligence and algorithmic decision-making in personalized pricing**. *Journal of Revenue and Pricing Management*, 19(2), 85–98. <https://doi.org/10.1057/s41272-019-00225-2>
- [5] Hermann, E. (2022). **Leveraging artificial intelligence in marketing for social good—An ethical perspective**. *Journal of Business Ethics*, 179(1), 43–61. <https://doi.org/10.1007/s10551-021-04843-y>
- [6] O'Keefe, K., & O'Brien, D. (2023). **Data ethics: Practical strategies for implementing ethical information management and governance** (2nd ed.).

- [30] Vakkuri, V., Kemell, K. K., Tolvanen, J., Jantunen, M., Halme, E., & Abrahamsson, P. (2022). **How do software companies deal with artificial intelligence ethics? A gap analysis.** *In Proceedings of the International Conference on Evaluation and Assessment in Software Engineering* (pp. 100–109). ACM. <https://doi.org/10.1145/3530019.3530030>
- [31] Corrêa, N. K., Galvão, C., Santos, J. W., Del Pino, C., Pinto, E. P., Barbosa, C., et al. (2023). **Worldwide AI ethics: A review of 200 guidelines and recommendations for AI governance.** *Patterns*, 4(10). <https://doi.org/10.1016/j.patter.2023.100857>
- [32] Arora, S., & Thota, S. R. (2024). **Ethical considerations and privacy in AI-driven big data analytics.** *International Research Journal of Engineering and Technology*, 11(5), 776–788. <https://www.irjet.net/volume11-issue5>
- [33] Gupta, D. G., & Jain, V. (2023). **Use of artificial intelligence with ethics and privacy for personalized customer services.** In J. N. Sheth, V. Jain, E. Mogaji, & A. Ambika (Eds.), *Artificial intelligence in customer service*. Palgrave Macmillan. https://doi.org/10.1007/978-3-031-33898-4_10
- [34] Benthall, S., & Goldenfein, J. (2021). **Artificial intelligence and the purpose of social systems.** *In Proceedings of the 2021 AAAI/ACM Conference on AI, Ethics, and Society* (pp. 3–12). ACM. <https://dl.acm.org/doi/10.1145/3461702.3462526>
- [35] Saunders, M., Lewis, P., & Thornhill, A. (2024). **Research methods for business students** (8th ed.). Pearson. https://www.pearson.com/nl/en_NL/higher-education/subject-catalogue/business-and-management/Research-methods-for-business-students-8e-saunders.html
- [36] Miles, M. B., & Huberman, A. M. (1994). **Qualitative data analysis: An expanded sourcebook.** Sage Publications. https://books.google.com/books?id=U4IU_-wJ5QEC
- [37] Barraket, J., & Loosemore, M. (2018). **Co-creating social value through cross-sector collaboration between social enterprises and the construction industry.** *Construction Management and Economics*, 36(7), 394–408. <https://doi.org/10.1080/01446193.2017.1416152>
- [19] Caudill, E. M., & Murphy, P. E. (2000). **Consumer online privacy: Legal and ethical issues.** *Journal of Public Policy & Marketing*, 19(1), 7–19. <https://doi.org/10.1509/jppm.19.1.7.16951>
- [20] Dunfee, T. W., Smith, N. C., & Ross, W. T. (1999). **Social contracts and marketing ethics.** *Journal of Marketing*, 63(3), 14–32. <https://doi.org/10.1177/002224299906300302>
- [21] Chellappa, R. K., & Sin, R. G. (2005). **Personalization versus privacy: An empirical examination of the online consumer's dilemma.** *Information Technology and Management*, 6, 181–202. <https://doi.org/10.1007/s10799-005-5879-y>
- [22] Rawls, J. (2001). **Justice as fairness: A restatement** (E. Kelly, Ed.). Harvard University Press. <https://www.hup.harvard.edu/books/9780674005112>
- [23] Ashworth, L., & Free, C. (2006). **Marketing dataveillance and digital privacy: Using theories of justice to understand consumers' online privacy concerns.** *Journal of Business Ethics*, 67, 107–123. <https://doi.org/10.1007/s10551-006-9007-7>
- [24] Vail, M. W., Earp, J. B., & Antón, A. I. (2008). **An empirical study of consumer perceptions and comprehension of website privacy policies.** *IEEE Transactions on Engineering Management*, 55(3), 442–454. <https://doi.org/10.1109/TEM.2008.922634>
- [25] Murphy, P. E., Laczniak, G. R., Bowie, N. E., & Klein, T. A. (2005). **Ethical marketing.** Pearson. https://epublications.marquette.edu/marq_fac-book/39/
- [26] Lwin, M., Wirtz, J., & Williams, J. D. (2007). **Consumer online privacy concerns and responses: A power-responsibility equilibrium perspective.** *Journal of the Academy of Marketing Science*, 35, 572–585. <https://doi.org/10.1007/s11747-006-0003-3>
- [27] Tucker, C. E. (2014). **Social networks, personalized advertising, and privacy controls.** *Journal of Marketing Research*, 51(5), 546–562. <https://doi.org/10.1509/jmr.10.0355>
- [28] Mothersbaugh, D. L., Foxx, W. K., Beatty, S. E., & Wang, S. (2012). **Disclosure antecedents in an online service context: The role of sensitivity of information.** *Journal of Service Research*, 15(1), 76–98. <https://doi.org/10.1177/1094670511424924>
- [29] Acquisti, A., John, L. K., & Loewenstein, G. (2013). **What is privacy worth?** *The Journal of Legal Studies*, 42(2), 249–274. <https://doi.org/10.1086/671754>