

مقاله پژوهشی: الگوی مفهومی کاربست حکمرانی هوش مصنوعی با رویکرد ارتقاء قدرت سایبری

حسن محمدی منفرد^۱، احمد رضا ترسلی^۲

تاریخ دریافت: ۱۴۰۳/۱۱/۰۳

تاریخ پذیرش: ۱۴۰۴/۰۳/۲۰

چکیده

عدم تنظیم‌گری هوش مصنوعی، می‌تواند منجر به ظهور قدرت سایبری مبتنی بر اراده و رفتار بازیگران غیردولتی و تشدید خلأ قانونی در این زمینه گردد. در چنین شرایطی ضمانت‌اجرائی به‌کارگیری فناوری هوش مصنوعی در همه عناصر قدرت ملی با تردیدهای امنیتی و مسئولیت‌پذیری مواجه گردیده، منجر به پیامدهای منفی خواهد شد؛ لذا، پژوهش حاضر باهدف دستیابی به الگوی مفهومی کاربست حکمرانی هوش مصنوعی انجام شده است. ازاین‌رو؛ پژوهش مذکور، از نظر هدف، کاربردی توسعه‌ای؛ از لحاظ رویکرد، کیفی و بر اساس روش فراترکیب انجام شده است. براین‌اساس، جامعه آماری پژوهش، شامل مجموعه مقالاتی است که در بازه زمانی سال‌های ۲۰۲۰ تا ۲۰۲۴ میلادی و با موضوعیت هوش مصنوعی، حکمرانی هوش مصنوعی، قدرت سایبری و ارتقاء قدرت سایبری به چاپ رسیده‌اند. نویسندگان با جستجو در میان پایگاه‌های علمی مطرح داخلی و خارجی ۷۱ مقاله را از میان ۳۷۶ مقاله برگزیدند. در این پژوهش بر اساس روش هفت‌مرحله‌ای فراترکیب، مقوله‌ها و مفاهیم استخراج شده در قالب ۱۴ مؤلفه (خودکارسازی و کارائی، مهارت شناختی، پیش‌بینی، تصمیم‌گیری، قابلیت‌های تهاجمی، دفاعی و امنیتی، قابلیت اعتماد، توضیح‌پذیری، مسئولیت‌پذیری، سازوکارها، فرایندها، هنجارهای اجتماعی، اخلاق‌مداری، تنظیم‌گری) و ۴ بعد (عوامل کارکردی، بارویکرد ارتقاء قدرت سایبری، عوامل سیستمی، عوامل سازمانی و عوامل محیطی و فراسازمانی) شناسایی شده است. در ادامه و با مراجعه به جامعه خبرگی تحقیق، آزمون کیفیت و اعتبارسنجی دسته‌بندی‌ها صورت گرفت که این مهم به تحصیل اعتبار خروجی نهایی پژوهش منجر شد. درنتیجه، با کاربست چنین مؤلفه‌هایی بتوان به یک الگوی عملگرها و استفاده اخلاقی و مسئولانه، ضمن به‌حداقل رساندن خطرات مرتبط با استفاده از هوش مصنوعی، از آن در راستای دستیابی به قدرت سایبری استفاده نمود.

کلیدواژه‌ها: هوش مصنوعی، قدرت سایبری، حکمرانی هوش مصنوعی، الگوی حکمرانی هوش مصنوعی

^۱ . استادیار دانشگاه عالی دفاع ملی (نویسنده مسئول) h. mohammadi@sndu.ac.ir

^۲ . دانش آموخته دکتری مدیریت راهبردی فضای سایبر دانشگاه عالی دفاع ملی

۱. مقدمه

پیشرفت‌های سریع در توسعه و گسترش فناوری‌های هوش مصنوعی در سال‌های اخیر توأم با ظهور بازیگران قدرتمند هوش مصنوعی در ابعاد مختلف عناصر قدرت، تأثیر و تغییرات شگرفی در معادله قدرت به وجود خواهد آورد (Dafie, 2018:34) و در نتیجه قدرت سایبری، به طور فزاینده‌ای به‌عنوان پیشران برتر در دستیابی به قدرت ملی برای هر دولت تبدیل می‌شود (Vuuren, et.al, 2016:1). این پدیده مستلزم حکمرانی مناسب هوش مصنوعی برای ارتقای قدرت سایبری است تا با شناسایی، هدایت مطلوب و انتخاب بهترین راه‌حل در توسعه و به‌کارگیری سامانه‌های پیشرفته هوش مصنوعی در ابعاد مختلف قدرت، منافع مردم و کشور تأمین گردد (Dafie, 2018:2).

پیشرفت شتابان فناوری و علم هوش مصنوعی و تأثیر تحول‌آفرین آن در طیف گسترده‌ای از مسائل، فرصت‌ها و چالش‌های جدیدی را برای سیاست‌گذاران و سایر ذی‌نفعان ایجاد می‌کند که به طور مستقیم می‌توانند بر حکمرانی و قدرت ملی تأثیر بگذارند (Schmitt, 2022:1). فقدان رسیدگی به این چالش‌ها، مسائل و دغدغه‌های فراوانی را برای اطمینان از چگونگی به‌کارگیری مناسب هوش مصنوعی برای دستیابی به قدرت سایبری و برخورداری از اثرات بازدارنده آن ایجاد خواهد کرد. دغدغه‌هایی که حاکی از فقدان تعریف هنجارها، سیاست‌ها و چارچوب‌های لازم برای اطمینان از توسعه و استفاده مفید از هوش مصنوعی است که خود از جمله مسائل مهم و اولویت‌دار در حکمرانی است. پیشرفت‌های اخیر هوش مصنوعی نشانگر ضرورت تدوین حکمرانی مناسب هوش مصنوعی مبتنی بر ویژگی‌های معتبر آن برای به حداکثر رساندن مزایای فناوری‌های هوش مصنوعی و کاهش خطرات آن‌ها است. در چنین شرایطی، ضمانت اجرایی به‌کارگیری فناوری هوش مصنوعی در همه عناصر قدرت ملی با تردیدهای اخلاقی، امنیتی (حریم خصوصی) و مسئولیت‌پذیری مواجه گردیده که خود مانع افزایش قدرت سایبری کشور و تحقق مزایای مورد انتظار از یک الگوی حکمرانی مناسب هوش مصنوعی برای ارتقای قدرت سایبری خواهد گردید (Ibid:3).

نهایتاً الگوی حکمرانی هوش مصنوعی به‌عنوان یک متغیر تعدیل‌کننده یک جنبه مهم برای اطمینان از شفافیت، توجه به اصول اخلاقی و قابل‌اعتماد بودن متغیر مستقل هوش مصنوعی است (Mäntymäki, et.al, 2023:2) و بایستی به نگرانی‌های شناختی، اخلاقی و امنیتی (Berente, et.al, 2021:4) توجه داشته باشد. الگوی حکمرانی، باید به طور خاص به خطرات و چالش‌های ناشی از شواهد غیرقطعی، غیرقابل‌درک و نادرست ایجاد شده توسط الگوریتم‌های یادگیری ماشینی که حکمرانی هوش مصنوعی را

از حکمرانی و مدیریت هر نوع سامانه فناوری اطلاعات متمایز می‌کند، رسیدگی کند (Martin, ۲۰۱۹:۸۴۳).

با این مقدمات، هدف نویسندگان در پژوهش حاضر، دستیابی به الگوی مفهومی کاربست حکمرانی هوش مصنوعی است. براین اساس سؤال اصلی پژوهش عبارت است از: «چگونه می‌توان به یک الگوی عمل‌گرا در حکمرانی هوش مصنوعی برای ارتقای قدرت سایبری دست یافت؟» و در ادامه، «مؤلفه‌های اصلی این الگو مبتنی بر چه مؤلفه‌هایی هستند؟» به بیانی شفاف‌تر، با تأکید بر چه مؤلفه‌هایی می‌توان به الگوی مورد نظر دست یافت؟

در خصوص ضرورت انجام تحقیق نیز می‌توان به چند نکته مهم اشاره داشت: نخست آنکه؛ همه الگوهای حکمرانی، برای هوش مصنوعی مناسب نیستند. به‌عنوان مثال، الگوی حکمرانی فناوری هسته‌ای برای هوش مصنوعی مناسب نیست؛ زیرا آنها دو نوع مختلف از فناوری هستند. از جمله تفاوت‌های مشهود در بین فناوری هسته‌ای و فناوری هوش مصنوعی می‌توان به واقعیت فیزیکی مواد در فناوری هسته‌ای و ماهیت دیجیتالی و نامشهود هوش مصنوعی و همچنین بازیگران اصلی و دولتی فناوری هسته‌ای و بازیگران غیردولتی و تجاری در حوزه هوش مصنوعی اشاره کرد که این تفاوت‌ها نقش مهمی در طراحی الگوی حکمرانی مناسب هر فناوری ایفا می‌نمایند (Stewart, 2023:2). دوم؛ نامناسب بودن الگوی حکمرانی می‌تواند منجر به تنظیم‌گری و قانون‌گذاری نامناسب هوش مصنوعی گردد و از آنجاکه بازیگران اصلی و فعال این فناوری غالباً در بخش‌های غیردولتی و تجاری مستقر هستند، تنظیم‌گری نامناسب هوش مصنوعی می‌تواند منجر به ظهور قدرت سایبری مبتنی بر اراده و رفتار بازیگران غیردولتی و تشدید خلأ قانونی در این زمینه گردد.

۲. مبانی نظری و پیشینه‌شناسی تحقیق

۲-۱. پیشینه تحقیق

بررسی‌ها نشان می‌دهد در خصوص الگوی حکمرانی هوش مصنوعی از نگاه ارتقای قدرت سایبری، تحقیق چشمگیری صورت نگرفته و اکثریت منابع فارسی و لاتین در حوزه مفهومی و به‌صورت ژورنالیستی به جنبه‌هایی از آن پرداخته است؛ لذا پرداختن به موضوع حاضر را از حیث نو بودن در جایگاه مناسبی قرار می‌دهد. از مهم‌ترین پژوهش‌های صورت گرفته در خصوص موضوع مقاله، به موارد زیر می‌توان اشاره نمود:

(۱) قاسمی و همکاران (۱۴۰۲)، در تحقیقی با عنوان «الگوی ارزیابی قدرت آفند سایبری ارتش جمهوری اسلامی ایران» نشان دادند که نیروی انسانی، پیچیدگی سایبری، تسلیحات سایبری و آگاهی وضعیتی

سایبری مهم‌ترین مؤلفه‌ها در آفند سایبری هستند و در نهایت ۱۲ شاخص برای ارزیابی قدرت آفند سایبری و ارائه الگوی ارزیابی قدرت آفند سایبری ارتش جمهوری اسلامی ایران را ارائه نموده‌اند.

(۲) رمضان‌زاده و همکاران (۱۳۹۹) در تحقیقی با عنوان «ارائه مدل مفهومی ارزیابی قدرت سایبری نیروهای مسلح با تأکید بر بعد بازدارندگی سایبری» نشان دادند که مدل مفهومی دارای مؤلفه‌های پنج‌گانه: پشیمان‌کنندگی دشمن، استمرار عملیات، پاسخ به تهاجم، استحکام سازی و بازیابی است. همچنین، مؤلفه استمرار عملیات، از اولویت بالاتری نسبت به سایر مؤلفه‌ها برخوردار است. استفاده از نتیجه این تحقیق به‌وسیله معاونت فاوا می‌تواند برای شناسایی و اولویت‌بندی سرمایه‌های سایبری نیروهای مسلح به‌منظور ایجاد زمینه بازدارندگی سایبری در مقابل تهدیدات، مفید باشد.

(۳) صبا و پرتوریوس (۲۰۲۴)، در تحقیقی با عنوان «تأثیر سرمایه‌گذاری هوش مصنوعی بر رفاه انسان در کشورهای جی هفت: آیا نقش تعدیل‌کننده حکمرانی اهمیت دارد؟» نشان دادند که سیاست‌گذاران باید ادغام هوش مصنوعی در حکمرانی را برای تقویت رفاه انسانی در کوتاه‌مدت و بلندمدت در اولویت قرار دهند. با این حال، هنگام در نظر گرفتن تعامل هوش مصنوعی با حکمرانی نهادی، احتیاط لازم است، زیرا از رفاه انسانی پشتیبانی نمی‌کند. در حالی که هوش مصنوعی و کیفیت کلی حکمرانی امیدوارکننده به نظر می‌رسد، اجرای تدابیر امنیتی در برابر اثرات منفی بالقوه ناشی از تعامل بین هوش مصنوعی و حکمرانی نهادی بر رفاه انسانی ضروری است.

(۴) هادفیلد و کلارک (۲۰۲۳)، در تحقیقی با عنوان: «بازارهای نظارتی: آینده حکمرانی هوش مصنوعی» نشان دادند که فناوری‌های هوش مصنوعی ظرفیت بازنویسی ساختار بنیادی جوامع و اقتصادهای انسانی را دارند و امروزه به روش‌هایی این کار را انجام می‌دهند که از چارچوب‌های نظارتی که برای مدیریت جهان قرن بیستم وضع شده عبور می‌کنند. آن‌ها پیشنهاد کردند که زمان ظهور ایده‌های جدید جسورانه برای تنظیم‌گری هوش مصنوعی فرارسیده است و بازارهای نظارتی یکی از این ایده‌ها هستند که دولت‌ها می‌توانند و باید برای مقابله با چالش‌های حکمرانی اصلی هوش مصنوعی به کار گیرند.

(۵) بریکستد^۴ و همکاران (۲۰۲۳)، در تحقیقی با عنوان «حکمرانی هوش مصنوعی: موضوعات، شکاف‌های دانش و برنامه‌های آینده» این فرض را تأیید کردند که: حکمرانی هوش مصنوعی یک موضوع تحقیقاتی نوظهور با تعاریف صریح کمی است. علاوه بر این، بررسی نویسندگان چهار موضوع را در

^۱: Saba & Pretorius

^۲: G-7

^۳: Hadfield & Clark

4 : Birkstedt, et.al.

ادبیات حکمرانی هوش مصنوعی شناسایی کردند: فناوری، ذی‌نفعان و زمینه، مقررات و فرایندها. شکاف‌های دانشی آشکار شده عبارت بودند از درک محدود اجرای حکمرانی هوش مصنوعی، عدم توجه به زمینه حکمرانی هوش مصنوعی، اثربخشی نامشخص اصول و مقررات اخلاقی، و عملیاتی‌سازی ناکافی فرایندهای حکمرانی هوش مصنوعی.

۲-۲. مبانی نظری پژوهش

۲-۲-۱. قدرت سایبری

واژه قدرت به معنای توانایی تغییر رفتار افراد برای دستیابی به هدف‌ها یا پیشبرد منافع خود است. در گذشته، سرچشمه قدرت ناشی از مولفه‌هایی چون جمعیت و گستره سرزمینی، ذخائر و ثروت‌های معتبر و قدرت نظامی برآورد می‌گردید ولیکن منابع قدرت در طول زمان و با تأثیر گسترده فناوری‌های نوظهور و مخرب تغییر می‌کنند و هزینه‌های سیاسی و اجتماعی جدیدی بر کشور تحمیل می‌کند. تحولات عمده دهه‌های اخیر در حوزه دفاعی و امنیتی تابعی از تغییرات فناوری‌های نوین بوده که یکی از عوامل اصلی در شکل‌گیری و تقویت قدرت است (Cohen, et. al, 2007:6).

با رشد سریع فضای سایبری ظهور فناوری‌های نوین که سهولت دسترسی و گمنامی را برای بازیگران فعال فضای سایبر فراهم کرده‌اند، تقابل‌های نامتقارن میان بازیگران کوچک و قدرت‌های بزرگ، مفهوم قدرت سایبری را بیش‌ازپیش معرفی کرده است. مؤلفه اطلاعات که از ارکان فضای سایبری است، همیشه نقش مؤثری بر قدرت داشته است (Nye, Jr., 2010:5). به گفته جوزف نای^۱، قدرت مبتنی بر منابع اطلاعاتی، قدرت سایبری است. مک کانل نیز قدرت سایبری را توانایی به‌کارگیری فناوری اطلاعات و ارتباطات برای رشد اقتصادی، رفاه اجتماعی و امنیت تعریف کرده است (Vuuren, et.al, 2016:3)؛ بنابراین کارکرد قدرت سایبری، برای دستیابی به قدرت ملی و نهایتاً تأمین امنیت ملی است (Vuuren, et.al, 2019:1). به گفته کوئل: قدرت سایبری توانایی تولید مزیت‌های کلیدی و تأثیرگذار در فضای سایبری بر رویدادهای همه محیط‌های عملیاتی و ابزارهای قدرت است و اطلاعات، عنصر مهم قدرت سایبری است و به طور شگرفی بر عناصر قدرت ملی (سیاسی، اطلاعاتی، نظامی، اقتصادی)، اثر می‌گذارد. اطلاعات بین عناصر و ابزارهای قدرت ارتباط ایجاد کرده و آنها را به هم مرتبط می‌کند که موجب شکل‌گیری قدرت سایبری می‌گردد. داده‌ها در فضای سایبری، به‌عنوان مواد خام به اقتصاد جامعه کمک می‌کند. انواع ارتباطات در فضای سایبری مرزهای سنتی را تغییر و روابط جدید بین مردم شکل داده و با

^۱: Nye Jr, 2010

^۲: McConnell, 2012

تسهیل برقراری ارتباط بین مردم و دیگر دولت‌ها از فراسوی مرزها شکل نوینی از قدرت سیاسی ایجاد کرده است (Kuehl, 2009:12).

۲-۲-۲. هوش مصنوعی

هوش مصنوعی مبتنی بر شباهت ذهن انسان و قدرت یادگیری ماشین در بررسی و پردازش داده‌ها و تولید دانش بنا شده است و می‌توان از آن برای انتخاب اقدامی که باید انجام داد استفاده کرد. هوش مصنوعی قادر است عملکرد انسان در وظایف خاص مانند تصمیم‌گیری، تشخیص الگو، و پیش‌بینی را تقلید کند. هوش مصنوعی با درک ورودی‌های محیط و به کمک توابع مختلف، کارهای مختلفی چون تصمیم‌گیری، تحلیل ارتباطات پیچیده، تشخیص تصویر، پیش‌بینی جرم و... را انجام می‌دهد (Eriksson, 2020:3).

جینی رومتی^۱، مدیرعامل اسبق IBM، نقش هوش مصنوعی در مشارکت بین انسان‌ها و ماشین‌ها را «شناخت تقویت‌شده»^۲ توصیف کرده است. از نظر او، هوش مصنوعی با تقویت شناخت انسان به کارآمدی و اثربخشی انسان کمک می‌کند (Tontiwachwuthikul, et.al, 2020:1). هوش مصنوعی قادر است وظایف دستی و محاسباتی انسان را با دقت و سرعتی در مقیاس ماشین و حجمی بالا در میان کلان‌داده‌ها با الگوریتم‌های طبقه‌بندی و خوشه‌بندی اطلاعات، به صورت خودکار و در تمام طول ساعات روز انجام دهد. استفاده فراگیر هوش مصنوعی و یادگیری ماشین، موجب تولید قدرت و بهره‌وری بیشتری گردیده است. (Chakraborty, et. al, 2022:252). توانایی استفاده از مزیت‌های کاربردی سامانه‌های هوش مصنوعی در بهبود تصمیم‌گیری، طراحی مدل‌های کسب‌وکار و زیست‌بوم‌ها، از جمله ابتکارات دیجیتالی سال‌های پیشرو خواهد بود (Duan, et.al, 2019:1).

سامانه‌های هوش مصنوعی با قابلیت‌های سازماندهی بسیاری از وظایف اجتماعی و تسهیل ارتباطات و تمایلات انسانی، فرصت‌ها و مزیت‌های مکمل مهارت‌های انسانی را برای مردم فراهم خواهند کرد تا با کاهش خطاهای ناشی از قضاوت‌های انسانی از یک سو و بهبود عملکرد الگوریتم‌های یادگیری از طریق خودآموزی، قدرت سایبری هوش مصنوعی در کشف الگوهای پنهان را فراهم آورند (Zekos, 2022:9).

۲-۲-۳. حکمرانی هوش مصنوعی

دولت‌ها همواره از روش‌های نوین حکمرانی که به رفع مشکلات مردم و تولید ثروت و رفاه اجتماعی و نهایتاً تولید اعتماد عمومی به دولت بینجامد استقبال می‌کنند. هوش مصنوعی که توانایی فوق‌العاده‌ای در پردازش سریع و دقیق کلان‌داده‌ها و ارائه بیش‌مورد نیاز برای تصمیم‌گیری آگاهانه دارد، ضمن ارائه

^۱: Ginni Rometty, IBM' CEO at that time (now Executive Chairperson of IBM)

^۲: augmented cognition

تسهیلات فراوان در اداره امور جامعه، خطرات جدی مختص خود را هم به همراه دارد که در جای خود بایستی مورد مطالعه و تنظیم‌گری قرار گیرد. هوش مصنوعی با کاهش هزینه‌های مرسوم اطلاعات و ارتباطات و به کمک الگوریتم‌های یادگیری ماشینی می‌تواند برای تعامل با افراد بیشتر و افزایش دقت و پاسخ به موقع، یادگیری منحصربه‌فردی داشته باشد. این توانایی، زندگی اقتصادی و اجتماعی را متحول کرده و پیامدهای دیجیتالی در کاهش یا افزایش قدرت خواهد داشت. پیشرفت‌های اخیر هوش مصنوعی نشانگر ضرورت تدوین حکمرانی مناسب هوش مصنوعی مبتنی بر ویژگی‌های معتبر آن برای به حداکثر رساندن مزایای فناوری‌های هوش مصنوعی و کاهش خطرات آن‌ها است.

حکمرانی هوش مصنوعی چارچوب قانونی برای اطمینان از توسعه فناوری‌های هوش مصنوعی مبتنی بر الزامات قانونی و اصول اخلاقی و ارزشی ذی‌نفعان است. پذیرش سریع ابزارها، سامانه‌ها و فناوری‌های هوش مصنوعی در ابعاد مختلف قدرت سایبری، نگرانی‌هایی را در مورد اخلاق، شفافیت و انطباق با سایر مقررات مانند مقررات حفاظت از داده‌های عمومی ایجاد می‌کند. بدون حکمرانی مناسب، سامانه‌های هوش مصنوعی می‌توانند خطراتی مانند تصمیم‌گیری مغرضانه، نقض حریم خصوصی و سوءاستفاده از داده‌ها و قدرت سایبری را به همراه داشته باشند. حکمرانی مناسب هوش مصنوعی به دنبال تسهیل استفاده سازنده از فناوری‌های هوش مصنوعی و درعین حال محافظت از حقوق کاربر و جلوگیری از آسیب است (Birkstedt, et.al, 2023:133-144).

۲-۲-۴. کاربرد هوش مصنوعی در ارتقاء قدرت سایبری

هوش مصنوعی این پتانسیل را دارد که به روش‌های مختلف بر قدرت سایبری تأثیر بگذارد و از این‌رو ارتباط نزدیکی با هم دارند. کاربرد هوش مصنوعی در ارتقای قدرت سایبری قابل توجه و چندوجهی است که در ذیل به اهم آن‌ها اشاره می‌گردد:

- **خودکارسازی و کارایی:** هوش مصنوعی قادر است وظایف دستی و محاسباتی انسان را با دقت و سرعتی در مقیاس ماشین و حجمی بالا در میان کلان‌داده‌ها با الگوریتم‌های مختلفی چون طبقه‌بندی و خوشه‌بندی اطلاعات، به صورت خودکار و در تمام طول ساعات روز انجام و کارایی و زمان پاسخگویی را بهبود بخشد. (Chakraborty, et.al, 2022:22).

- **افزایش مهارت شناختی:** فناوری‌های هوش مصنوعی چون پردازش زبان طبیعی و بینایی رایانه؛ نقش مهمی در توسعه مهارت‌های شناختی و دانش حاصل از تحلیل مجموعه‌های پیچیده داده دارد. هدف پردازش زبان طبیعی به عنوان شاخه‌ای از هوش مصنوعی که توانایی درک زبان نوشتاری و کلمات

^۱: Automation and Efficiency

گفتاری را همانند انسان به ماشین می‌دهد، استخراج دانش‌شناختی است (Schembri, et.al, 2023:1721). در بینایی رایانه نیز با درک و تفسیر محیط همانند بینایی چشم انسان، اطلاعات معنی‌داری از رسانه‌های تصویری (فیلم، عکس و سایر داده‌های بصری) استخراج و مبنای شناخت قرار می‌گیرد (Matsuzaka, et.al, 2023:2).

• **تجزیه و تحلیل پیش‌بینی:** هوش مصنوعی این پتانسیل را دارد که تحلیل اطلاعات پیش‌بینی را با ارائه پیش‌بینی‌های دقیق و کارآمدتر بهبود بخشد. الگوریتم‌های هوش مصنوعی می‌توانند حجم زیادی از داده‌ها را به سرعت و با دقت تجزیه و تحلیل کنند و الگوهایی را شناسایی کنند که ممکن است فوراً برای انسان آشکار نباشد. این ظرفیت مستلزم تأمین مجموعه داده‌های جامع‌تری است تا از فرایندهای تصمیم‌گیری آگاهانه‌تر پشتیبانی کند. هوش مصنوعی از انواع الگوریتم‌های یادگیری ماشینی برای تجزیه و تحلیل داده‌ها و پیش‌بینی نتایج آینده استفاده می‌کند (Cardenas-Canto, 2022:30). به عنوان مثال در الگوریتم یادگیری نظارت شده، رایانه یاد می‌گیرد که بر اساس داده‌های برچسب‌گذاری شده توسط انسان، تخمین بزند و در الگوریتم یادگیری بدون نظارت^۳ رایانه بدون برچسب‌زنی به داده‌ها، یاد می‌گیرد که الگوی پنهان در مجموعه داده را توسط الگوریتم‌ها کشف و رویدادی را پیش‌بینی کند (Chouhan, et.al, 2021:62).

• **بهبود تصمیم‌گیری:** یکی از مزایای اصلی هوش مصنوعی، توانایی آن در شناسایی و پاسخگویی در زمان واقعی است که نقش مهمی در ارتقای قدرت سایبری دارد. سامانه‌های مبتنی بر هوش مصنوعی می‌توانند با تحلیل حجم زیادی از داده‌ها و تشخیص ناهنجاری‌ها، الگوهای رفتاری و سایر شاخص‌های نظارتی، اطلاعاتی را برای تصمیم‌گیری بهتر در مورد نحوه واکنش به رویدادها ارائه نمایند. تصمیم‌گیری مبتنی بر هوش مصنوعی با استفاده از الگوریتم‌های هوش مصنوعی برای انتخاب از بین گزینه‌ها یا انجام اقدامات بر اساس داده‌ها، مدل‌ها و سایر ورودی‌ها انجام می‌شود. بر اساس تجزیه و تحلیل پیش‌بینی، الگوریتم‌های هوش مصنوعی می‌توانند توصیه‌هایی ایجاد کنند، مناسب‌ترین مسیر عمل را انتخاب کنند، یا حتی اقداماتی را به طور مستقل انجام دهند (Sánchez, 2022:160). هنگامی که قابلیت‌های هوش مصنوعی با قضاوت انسانی همراه شود، آگاهی و اشراف اطلاعاتی را افزایش می‌دهد، با تصمیم‌سازی و یا تصمیم‌گیری به موقع، فرایند تصمیم‌گیری دقیق‌تر و آگاهانه‌تر می‌شود (Schmidt, et.al, 2021:۱۱۱) که خود موجب افزایش قدرت سایبری در مواجهه با محیط‌های رقابتی و پیچیده می‌گردد.

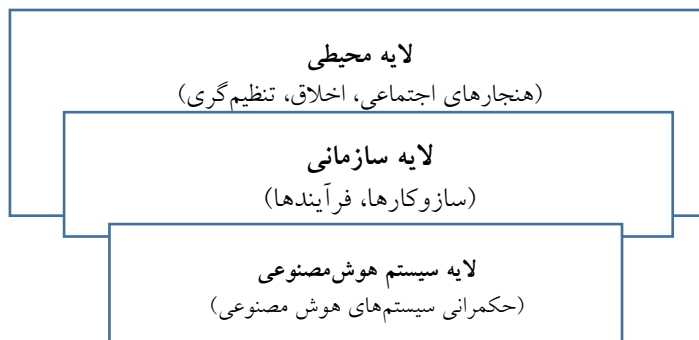
۳ supervised learning
 ۴ Estimation
 ۵ unsupervised learning

• **بهبود قابلیت‌های تهاجمی / امنیتی / دفاعی:** از آنجاکه افزایش امنیت سایبری بر تمایل دولت به ارتقای قدرت سایبری تأثیر مثبتی دارد (Arslan, 2023:5)، هوش مصنوعی با تقویت تیم‌های امنیتی از طریق داده‌ها و یادگیری ماشین به بهبود امنیت سایبری کمک می‌نماید. تا آن‌ها بتوانند سریع‌تر از حمله مهاجمان سایبری؛ پاسخ دهند و حتی حرکات آنها را پیش‌بینی کرده و از قبل اقدام کنند. توانایی هوش مصنوعی برای یادگیری و شناسایی الگوهای جدید به‌صورت تطبیقی می‌تواند تشخیص، مهار و واکنش را تسریع کند و بار تحلیل‌گران امنیتی را کاهش دهد و به آنها اجازه دهد فعال‌تر باشند (Mijwil, et.al, 2023:92). هوش مصنوعی می‌تواند نظارت در زمان واقعی و به شکل مستمر را که برای امنیت سایبری مدرن ضروری است، فراهم کند. ابزارهای امنیت سایبری مبتنی بر هوش مصنوعی و تحلیل‌های داده‌محور برای شناسایی حملات در زمان واقعی طراحی شده‌اند و می‌توانند فرایند واکنش به حادثه را خودکار کنند. آنها همچنین می‌توانند به کارشناسان حفاظت شخصی در شناسایی تهدیدها و روندهای نوظهور کمک و آنها را قادر به انجام اقدامات پیشگیرانه نمایند، چراکه تشخیص سریع و بموقع، هرگونه وقوع اختلال را به حداقل می‌رساند (Kaur, et.al, 2023:13). هوش مصنوعی می‌تواند برای افزایش قدرت سایبری کشور با بهبود قابلیت‌های دفاعی و تهاجمی در فضای سایبری استفاده شود. به‌عنوان مثال، هوش مصنوعی می‌تواند برای بهبود تشخیص و پاسخ به حملات سایبری، و همچنین برای توسعه قابلیت‌های سایبری تهاجمی پیچیده‌تر استفاده شود. (Bonfanti, 2022:65). هوش مصنوعی می‌تواند برای افزایش سرعت، دقت، برد و مرگ‌آوری تسلیحات سایبری استفاده شود (Johnson, 2020:22). توانایی‌های یادگیری و تصمیم‌گیری هوش مصنوعی مبتنی بر شرایط موجب سفارشی شدن و شبیه به انسان شدن حملات گردیده است. هوش مصنوعی تهاجمی می‌تواند با اطلاع از محیط خود، خود را تغییر دهد و با حداقل شناسایی، به طور ماهرانه سامانه‌ها را در معرض خطر قرار دهد (Guembe, et.al, 2022:31).

۲-۲-۵. چارچوب مفهومی الگوی حکمرانی هوش مصنوعی

یکی از راه‌های ممکن برای طراحی یک الگوی حکمرانی هوش مصنوعی پیروی از یک چارچوب لایه‌ای است که از سه سطح: محیطی، سازمانی و سامانه هوش مصنوعی تشکیل شده است. سطح محیطی به عوامل خارجی که چشم‌انداز حکمرانی هوش مصنوعی را شکل می‌دهند، مانند قوانین، تنظیم‌گری، هنجارهای اجتماعی و اخلاقی اشاره دارد. سطح سازمانی به سیاست‌ها، فرایندها و سازوکارهای داخلی

اشاره دارد که فعالیت‌های هوش مصنوعی را در یک سازمان هدایت می‌کند. سطح سامانه هوش مصنوعی به ویژگی‌های خاص هر سامانه هوش مصنوعی، مانند توضیح‌پذیری^۱، قابلیت اعتماد^۲، شفافیت^۳ و مسئولیت‌پذیری^۴ (پاسخگویی) آن به شرح شکل (۱) اشاره دارد.



شکل ۱: الگوی لایه‌ای حکمرانی فناوری هوش مصنوعی

با پیروی از این چارچوب لایه‌ای، یک سازمان می‌تواند یک الگوی حکمرانی هوش مصنوعی طراحی کند که جامع، منسجم و در سطوح مختلف سازگار باشد. علاوه بر این، یک سازمان می‌تواند الگوی حکمرانی هوش مصنوعی خود را متناسب با نیازها و اهداف خاص خود تنظیم کند و همچنین آن را با شرایط و الزامات در حال تغییر تطبیق دهد. یک الگوی راهبردی مؤثر هوش مصنوعی می‌تواند به سازمان کمک کند تا به نتایج مطلوب خود از سامانه‌های هوش مصنوعی دست یابد و درعین حال خطرات بالقوه را به حداقل برساند و مزایای اجتماعی را به حداکثر برساند.

۳. روش‌شناسی تحقیق

پژوهش حاضر از نظر هدف، توسعه‌ای - کاربردی و بر اساس ماهیت داده‌ها و سبک تحلیل، جزء پژوهش‌های کیفی به شمار می‌رود که در آن داده‌های پژوهش با استفاده از روش فراترکیب^۵

۱: Explainability

۲: مفهوم قابل اعتماد بودن در اطمینان از اینکه یک مدل یادگیری در مواجهه با یک مشکل معین، همانطور که در نظر گرفته شده عمل می‌کند یا خیر تعریف می‌گردد. (Arrieta, et.al., 2020:8).

۳: Transparency

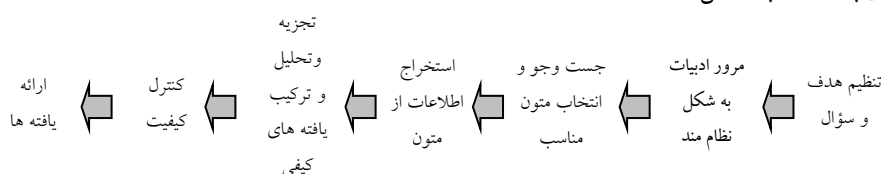
۴: Accountability

۵: Meta-Synthesis

جمع‌آوری و تحلیل شده است. این روش یک فراتحلیل کیفی از مفاهیم و نتایج مطالعات پیشین به شیوه کدگذاری متداول در پژوهش‌های کیفی است (Shojae Farahabadi et al., 2024).
براین اساس، جامعه آماری در این پژوهش شامل تمامی پژوهش‌های موجود در پایگاه‌های اطلاعاتی داخلی و خارجی است که از نظر محتوا با کلیدواژه‌های موضوع مورد مطالعه در این پژوهش ارتباط نزدیکی دارند. برای جستجوی پژوهش‌های منتشرشده، کلیدواژه‌ها در بازه زمانی سال‌های ۲۰۲۰ تا ۲۰۲۴ بررسی شد. درنهایت از میان ۳۷۶ پژوهش شناسایی شده، ۹۴ پژوهش به‌عنوان نمونه آماری انتخاب شدند. همچنین، به‌منظور بررسی روایی از نظر متخصصان و صاحب‌نظران امر در زمینه موضوع مورد مطالعه پژوهش، استفاده گردید و برای سنجش پایایی نیز شاخص کاپا به کار رفت.

۴. یافته‌ها و تجزیه و تحلیل داده‌ها

روش فراترکیب به سه شکل الگوی سه مرحله‌ای نوبلت و هیر^۲ (۱۹۹۸)، الگوی شش مرحله‌ای والش و داون^۳ (۲۰۰۵) و الگوی هفت مرحله‌ای سندلوسکی و باروسو^۴ (۲۰۰۷) قابل اجرا می‌باشد. الگوی سندلوسکی و باروسو (۲۰۰۷) به دلیل اینکه در پژوهش‌های فراترکیب بیشترین کاربرد را دارد، در این پژوهش نیز از این الگو استفاده شده است. بر اساس این الگو، مراحل روش فراترکیب در قالب شکل (۲) است.



شکل (۲) گام‌های متوالی روش فراترکیب (Sandelowski & Barroso, 2003)

مرحله اول: مشخص کردن سؤالات پژوهش: باتوجه به هدف، باید به شاخص‌های پژوهش، شامل چه جامعه‌ای، چه چیزی، چه زمانی و چگونگی روش پاسخ داده شود که براین اساس، سؤالات تنظیم شده است.

جدول ۱. سؤالات پژوهش

شاخص	سؤالات	پاسخ
What (چه چیزی؟)	الگوی حکمرانی هوش مصنوعی در راستای ارتقای قدرت سایبری	شناسایی مفاهیم و تعریف هوش مصنوعی، حکمرانی هوش مصنوعی و قدرت سایبری
Who (چه کسی؟ جامعه مطالعه)	جامعه مورد مطالعه برای شناسایی الگوی مورد نظر	جامعه پژوهش، مطالعات انجام شده در حوزه مورد مطالعه است که در مجلات معتبر علمی داخلی و خارجی منتشر شده‌اند.
When (چه زمانی؟ محدودیت و چارچوب زمانی مطالعه)	این منابع از چه دوره زمانی جستجو و بررسی شده‌اند؟	مطالعات در حوزه حکمرانی هوش مصنوعی و قدرت سایبری از سال ۲۰۲۰ تا ۲۰۲۴ مورد بررسی قرار گرفته است.
How (چگونه؟)	چه روشی برای فراهم نمودن این مطالعات استفاده شده است؟	در این پژوهش مرور ادبیات و تحلیل اسنادی مورد استفاده قرار گرفته است.

مرحله دوم: بررسی ادبیات موضوع به صورت نظام‌مند

در این مرحله پژوهشگر به طور نظام‌مند به جستجوی مقاله‌های منتشر شده در مجله‌های مختلف می‌پردازد. کلیدواژه‌های اصلی پژوهش یعنی هوش مصنوعی، حکمرانی هوش مصنوعی و قدرت سایبری در پایگاه‌های علمی مطرح داخلی و خارجی مانند نورمگز، مگیران، ساینس دایرکت، ریسرچ گیت، آرسکیو^۳ و اسپرینگر^۴ مورد جستجو واقع شدند. نتیجه جست‌وجو فهرستی از اسناد گوناگون، شامل ۷۲۵ مقاله بود. معیارهای پذیرش یا عدم‌پذیرش اسناد علمی مطابق جدول ۲ تعیین شد.

جدول ۲. معیارهای پذیرش و عدم‌پذیرش اسناد علمی در گام دوم

معیارها	معیار پذیرش	معیار عدم‌پذیرش
زمان تحقیقات	تحقیقات منتشر شده در خلال سال‌های ۲۰۲۰ تا ۲۰۲۴ میلادی	تحقیقات غیر از این تاریخ
اعتبار اسناد	اسناد علمی چاپ شده در نشریات و پایگاه‌های اطلاعاتی معتبر داخلی و بین‌المللی	مصاحبه‌ها و نظرات شخصی در پایگاه‌های اطلاعاتی خبری غیرعلمی

^۱ Science Direct^۲ ResearchGate^۳ arxiv^۴ Springer

دسترسی به اسناد	دسترسی به متن کامل مقاله چاپ شده	عدم دسترسی به متن کامل مقاله
موضوع اسناد	کلیدواژه‌های اصلی پژوهش یعنی هوش مصنوعی، حکمرانی هوش مصنوعی و قدرت سایبری	غیر از موارد اشاره شده

مرحله سوم: بررسی و گزینش مقالات مناسب

هدف این مرحله، پالایش مقالاتی بود که بر اساس عنوان، چکیده و محتوا، تناسب چندانی باهدف پژوهش نداشتند. سرانجام، باتوجه به هدف پژوهش ۳۷۶ مورد با حوزه مورد بررسی مطابقت داشت و وارد فرایند ارزیابی گردید. در ادامه برای انتخاب مطالعات مناسب، بر اساس الگوریتم ارزیابی حیاتی (قاسمی و رعیت‌پیشه، ۱۳۹۴)، به شرح جدول (۳) در چند مرحله مورد ارزیابی قرار گرفتند و تعدادی از آنها به علت عدم تطابق باهدف پژوهش حذف شدند.

جدول ۳. مراحل پالایش منابع مورد استفاده بر اساس روش ارزیابی حیاتی

مرحله ۱	تعداد منابع یافت شده	۳۷۶
	تعداد منابع رده شده به علت عنوان	۷۳
مرحله ۲	تعداد منابع غربال شده بر اساس عنوان	۳۰۳
	تعداد منابع رده شده بر اساس چکیده	۶۶
مرحله ۳	تعداد منابع غربال شده بر اساس چکیده	۱۳۷
	تعداد منابع رد شده بر اساس محتوا	۲۳
مرحله ۴	تعداد منابع غربال شده بر اساس محتوا	۹۴

براین اساس نهایتاً ۷۱ مقاله به شرح جدول (۴) برای مطالعه عمیق گزینش شدند.

جدول ۴. تعداد مقالات منتخب از پایگاه‌های علمی

نام پایگاه علمی	نورمگز	مگیران	ساینس دایرکت	ریسرچ گیت	آرکیو	اسپرینگر	مجموع
تعداد	۱۲	۸	۱۳	۱۱	۱۵	۱۲	۷۱

مرحله چهارم استخراج اطلاعات مقالات

در مرحله چهارم، بعد از شناسایی پژوهش‌های مورد نظر کلیه متون این مقالات به‌عنوان یک داده برای دستیابی به یافته‌ها و پاسخگویی به سؤال پژوهش در نظر گرفته شد.

در این مرحله، پس از انتخاب محتوای مورد نظر برای تحلیل از بین متون تحقیق، و به دلیل کیفی بودن داده‌ها، از روش کدگذاری باز استفاده شد. این نوع کدگذاری دقیقاً با مرحله اول کدگذاری‌ها در پژوهش‌هایی که از روش نظریه داده‌بنیاد استفاده می‌کنند، مشابه است. در این شیوه کدگذاری، کدها از متن مقاله استخراج می‌شود (کدگذاری مرتبه اول) و سپس بر روی این کدهای استخراج شده مجدداً کدگذاری دیگری صورت می‌گیرد که مفاهیم را شکل می‌دهد (کدگذاری مرتبه دوم) و در نهایت بر روی مفاهیم نیز کدگذاری دیگری صورت می‌گیرد تا مقوله‌ها حاصل شود (متن - کد - مفهوم - مقوله).

مرحله پنجم: تجزیه و تحلیل و تلفیق یافته‌های کیفی

این مرحله که شاید حساس‌ترین مرحله فراترکیب باشد، باید با دقت خاصی انجام پذیرد. یافته‌های این گام مبنایی برای الگوی نهایی پژوهش به شمار می‌رود و باید در ترکیب آنها دقت داشت.

در این پژوهش داده‌ها از مجموع ۷۱ مقاله منتخب به شیوه کدگذاری گردآوری گردید. فرایند شناسایی کدها به صورت رفت و برگشتی بود، به این معنا که ابتدا با بررسی ادبیات موضوع، مفاهیم اولیه و کلی استخراج شدند. سپس با تحلیل متن مقاله‌ها و ایجاد مفاهیم جدید و جزئی‌تر، بار دیگر به ادبیات مراجعه شد تا معادل کدهای مطرح شده در مقاله‌ها، در ادبیات نیز جستجو شود. در این مرحله تعداد ۱۹۶ کد شناسایی گردید. در مرحله بعدی به تحلیل، ترکیب و تلفیق کدها در قالب مؤلفه‌ها پرداخته شد. در این گام کدهای شناسایی شده بر اساس میزان تشابه مفهومی دسته‌بندی و ترکیب شدند. در این مرحله کدهای استخراج شده در قالب ۱۴ مؤلفه و مؤلفه‌های شناسایی شده در سطح بالاتری در قالب ۴ بعد طبقه‌بندی شدند. نتایج آن در جدول (۵) آمده است.

جدول ۵. مؤلفه‌ها و ابعاد شناسایی شده در روش فراترکیب

مفهوم	بعد	مؤلفه	کد استخراجی	منبع
کارشناس حکمرانی هوش مصنوعی در ارتقاء قدرت هوش مصنوعی سیاسی	کارکردی اول	خودکارسازی	هوش مصنوعی قادر است وظایف دستی و محاسباتی انسان را با دقت و سرعتی در مقیاس ماشین و حجمی بالا در میان کلان‌داده‌ها با الگوریتم‌های مختلفی چون طبقه‌بندی و خوشه‌بندی اطلاعات، به صورت خودکار و در تمام طول ساعات روز انجام و	چاکرابورتی و همکاران، ۲۰۲۲: ۲۵۲. بدریچ و همکاران، ۲۰۲۳: ۱. کاردانانس-کانتو،

مفهوم	بعد	مؤلفه	کد استخراجی	منبع
			کارایی و زمان پاسخگویی را بهبود بخشد. هوش مصنوعی خودکار، برای تفکر و استدلالی فراتر از رگرسیون داده‌ها از مغز انسان و هوش شناختی و انعکاسی او تقلید کرده و یاد می‌گیرد. هوش مصنوعی خودکار نه تنها دانش انسان را گسترش داده؛ بلکه آن را با سرعت و وسعت بی‌سابقه‌ای افزایش می‌دهد.	۳۰: ۲۰۲۲ چوهان و همکاران، ۲۰۲۱: ۶۲ شمبری و همکاران، ۲۰۲۳: ۱۷۲۱ ماتسوزاکا و همکاران، ۲۰۲۳: ۲
			هوش مصنوعی این پتانسیل را دارد که تحلیل اطلاعات پیش‌بینی را با ارائه پیش‌بینی‌های دقیق و کارآمدتر بهبود بخشد. الگوریتم‌های هوش مصنوعی می‌توانند حجم زیادی از داده‌ها را به سرعت و با دقت تجزیه و تحلیل کنند و الگوهایی را شناسایی کنند که ممکن است فوراً برای انسان آشکار نباشد. هوش مصنوعی از انواع الگوریتم‌های یادگیری ماشینی برای تجزیه و تحلیل داده‌ها و پیش‌بینی نتایج آینده استفاده می‌کند. در الگوریتم یادگیری نظارت شده، رایانه یاد می‌گیرد که بر اساس داده‌های برچسب گذاری شده توسط انسان، تخمین بزند و در الگوریتم یادگیری بدون نظارت رایانه بدون برچسب زنی به داده‌ها، یاد می‌گیرد که الگوی پنهان در مجموعه داده را توسط الگوریتم‌ها کشف و رویدادی را پیش‌بینی کند.	سانچز، ۲۰۲۲: ۱۶۰ اشمیت و همکاران، ۲۰۲۱: ۱۱۱ میجویل و همکاران، ۲۰۲۳: ۹۲ کاور و همکاران، ۲۰۲۳: ۱۳ بونفانتی، ۲۰۲۲: ۶۵ جانسون، ۲۰۲۰: ۲۲ گومبه و همکاران، ۲۰۲۲: ۳۱
		شناختی / ادراکی	فناوری‌های هوش مصنوعی چون پردازش زبان طبیعی و بینایی رایانه؛ نقش مهمی در	

۱. supervised learning

۲. Estimation

مفهوم	بعد	مؤلفه	کد استخراجی	منبع
		تصمیم‌گیری	توسعه مهارت‌های شناختی و دانش حاصل از تحلیل مجموعه‌های پیچیده داده دارد. هدف پردازش زبان طبیعی به عنوان شاخه‌ای از هوش مصنوعی که توانایی درک زبان نوشتاری و کلمات گفتاری را همانند انسان به ماشین می‌دهد، استخراج دانش شناختی است. در بینایی رایانه، با درک و تفسیر محیط همانند بینایی چشم انسان، اطلاعات معنی‌داری از رسانه‌های تصویری (فیلم، عکس و سایر داده‌های بصری) استخراج و مبنای شناخت قرار می‌گیرد.	
			تصمیم‌گیری مبتنی بر هوش مصنوعی با استفاده از الگوریتم‌های هوش مصنوعی برای انتخاب از بین گزینه‌ها یا انجام اقدامات بر اساس داده‌ها، مدل‌ها و سایر ورودی‌ها انجام می‌شود. بر اساس تجزیه و تحلیل، الگوریتم‌های هوش مصنوعی می‌توانند توصیه‌هایی ایجاد کنند، مناسب‌ترین مسیر عمل را انتخاب کنند، یا حتی اقداماتی را به‌طور مستقل انجام دهند هنگامی که قابلیت‌های هوش مصنوعی با قضاوت انسانی همراه شود، آگاهی و اشراف اطلاعاتی را افزایش می‌دهد، با تصمیم‌سازی و یا تصمیم‌گیری به‌موقع، فرآیند تصمیم‌گیری دقیق‌تر و آگاهانه‌تر می‌شود که خود موجب افزایش قدرت سایبری در مواجهه با محیط‌های رقابتی و پیچیده می‌گردد.	
		عملیاتی (تهاجمی/دفاعی / امنیتی)	هوش مصنوعی می‌تواند نظارت در زمان واقعی و به شکل مستمر را که برای امنیت سایبری مدرن ضروری است، فراهم کند. ابزارهای امنیت سایبری مبتنی بر	

مفهوم	بعد	مؤلفه	کد استخراجی	منبع
عوامل سازمانی	سیاست‌ها		هوش مصنوعی و تحلیل‌های داده‌محور برای شناسایی حملات در زمان واقعی طراحی شده‌اند و می‌توانند فرآیند واکنش به حادثه را خودکار کنند.	
			هوش مصنوعی می‌تواند برای بهبود تشخیص و پاسخ به حملات سایبری، و همچنین برای توسعه قابلیت‌های سایبری تهاجمی پیچیده‌تر استفاده شود.	
	فرآیندها		هوش مصنوعی می‌تواند برای افزایش سرعت، دقت، برد و مرگ‌آوری تسلیحات سایبری استفاده شود	
			هوش مصنوعی تهاجمی می‌تواند با اطلاع از محیط خود، خود را تغییر دهد و با حداقل شناسایی، به‌طور ماهرانه سیستم‌ها را در معرض خطر قرار دهد	
			تعریف سیاست‌ها و چارچوب‌های لازم برای اطمینان از توسعه و استفاده مفید از هوش مصنوعی از جمله مسائل مهم و اولویت دار در حکمرانی است.	بریکستد و همکاران، ۲۰۲۳: ۱.
			درک محدود اجرای حکمرانی هوش مصنوعی، عدم توجه به زمینه حکمرانی هوش مصنوعی، اثربخشی نامشخص اصول و مقررات اخلاقی، و عملیاتی سازی ناکافی	مانتیمایکی و همکاران، ۲۰۲۳: ۱
			فرآیندهای حکمرانی هوش مصنوعی از جمله خلاءهای دانشی نسبت به حکمرانی هوش مصنوعی است.	برنته و همکاران، ۲۰۲۱: ۴
			داده‌ها محرک اصلی توسعه چشم‌گیر هوش مصنوعی در عصر کنونی هستند. نقش کلان داده‌ها در فرآیندهای اقتصادی،	استوارت، ۲۰۲۳: ۲
				دافو، ۲۰۱۸: ۷
				ونکی و همکاران، ۲۰۲۲: ۱۷۹
				کاردنانس-کانتو،

مفهوم	بعد	مؤلفه	کد استخراجی	منبع
			اجتماعی، سیاسی بیانگر قابلیت‌های حکمرانی داده و هوش مصنوعی است الگوریتم‌های هوش مصنوعی می‌توانند حجم زیادی از داده‌ها را به سرعت و با دقت تجزیه و تحلیل کنند و الگوهایی را شناسایی کنند که ممکن است فوراً برای انسان آشکار نباشد. این ظرفیت از فرآیندهای تصمیم‌گیری آگاهانه‌تر پشتیبانی می‌کند. شفافیت و پاسخگویی فرآیندهای مبتنی بر هوش مصنوعی موجب کاهش نگرانی‌های اجتماعی، سیاسی، اقتصادی مردم از تصمیم‌گیری‌های مبتنی بر راه‌حل‌های هوشمند در عرصه‌های نقش‌آفرینی و ارتقاء قدرت سایبری می‌گردد.	۳۰:۲۰۲۲ بونی، ۲۴:۲۰۲۱
	الزامات و سازوکارها		حکمرانی مناسب هوش مصنوعی مستلزم آن است که ساختارهای سازمانی، نقش‌ها و مسئولیت‌ها، معیارهای عملکرد و پاسخگویی برای نتایج سیستم‌های هوش مصنوعی تعریف شوند. دانش طراحی و ایجاد چارچوب‌های حکمرانی هوش مصنوعی در سازمان‌ها در قالب نیازمندی‌های اساسی به تبیین اصول اخلاقی هوش مصنوعی در عمل کمک می‌کند حکمرانی مناسب هوش مصنوعی چارچوب قانونی برای اطمینان از توسعه فناوری‌های هوش مصنوعی مبتنی بر الزامات قانونی و اصول اخلاقی و ارزشی ذینفعان است. پذیرش سریع ابزارها، سیستم‌ها و فناوری‌های هوش مصنوعی در ابعاد مختلف قدرت	

مفهوم	بعد	مؤلفه	کد استخراجی	منبع
اعتماد پذیری	اعتماد پذیری		سایبری، نگرانی‌هایی را در مورد اخلاق، شفافیت و انطباق با سایر مقررات مانند مقررات حفاظت از داده‌های عمومی ایجاد می‌کند.	
			نگرانی‌هایی در چگونگی اعتماد به نتایج هوش مصنوعی و سوء استفاده احتمالی از آن ایجاد گردیده که خود بیانگر نیاز به الگوی حکمرانی مناسب برای اطمینان از استفاده اخلاقی و مسئولانه از آن می‌باشد. ضمانت اجرایی بکارگیری فناوری هوش مصنوعی در همه عناصر قدرت ملی با تردیدهای اخلاقی، امنیتی (حریم خصوصی) و مسئولیت‌پذیری مواجه می‌گردد. الگوی حکمرانی مناسب هوش مصنوعی بایستی به مسائل مربوط به سوء استفاده از هوش مصنوعی که اعتماد عمومی را کاهش داده و مانع استقرار مناسب هوش مصنوعی می‌شود بپردازد.	سعید و همکاران، ۲۰۲۳: ۱ لاتو و همکاران، ۲۰۲۲: ۲ حفیظ، ۲۰۲۲: ۲ بوچر، ۲۰۱۹: ۹۴ بریکستد و همکاران، ۲۰۲۳: ۱۳۳ راسکا، ۲۰۲۳: ۲۳۳ موکاندر و همکاران، ۲۰۲۱: ۲ دافو، ۲۰۱۸: ۷ استوارت، ۲۰۲۳: ۲ مانتیمای و همکاران، ۲۰۲۳: ۲
	توضیح‌پذیری	توضیح‌پذیری	هوش مصنوعی قابل توضیح (XAI) برای شفاف‌تر کردن پیچیدگی‌های مدل‌های هوش مصنوعی و در نتیجه پیشبرد پذیرش هوش مصنوعی در حوزه‌های حیاتی پیشنهاد شده	

مفهوم	بعد	مؤلفه	کد استخراجی	منبع
عوامل محیطی و فرا سازمانی	مسئولیت‌پذیری	مسئولیت‌پذیری	است. ابداع استراتژی‌های توضیح‌پذیری هوش مصنوعی است که تصمیم‌گیری‌ها را برای طرف‌های آسیب‌دیده توجیه می‌کند.	
			حکمرانی مناسب هوش مصنوعی مستلزم آن است که ساختارهای سازمانی، نقش‌ها و مسئولیت‌ها، معیارهای عملکرد و پاسخگویی برای نتایج سیستم‌های هوش مصنوعی تعریف شوند. استفاده مسئولانه از هوش مصنوعی ضمن شکل‌گیری اعتماد عمومی به سیستم‌ها و زیرساخت‌های آن با تأثیرات منفی احتمالی آن مبارزه می‌کند. ضمانت اجرائی بکارگیری فناوری هوش مصنوعی در همه عناصر قدرت ملی با تردیده‌های اخلاقی، امنیتی (حریم خصوصی) و مسئولیت‌پذیری مواجه می‌گردد.	
	هنجارسازی اجتماعی	هنجارسازی اجتماعی	رژیم نوپای هوش مصنوعی که ظهور می‌کند، چند مرکزی و تکه‌تکه است، اما در محیط عملکرد سازمان همکاری اقتصادی و توسعه با اقتدار معرفتی و قدرت تنظیم هنجارهای اجتماعی فعال می‌باشد. تعریف هنجارها و چارچوب‌های لازم برای اطمینان از توسعه و استفاده مفید از هوش مصنوعی از جمله مسائل مهم و اولویت دار در حکمرانی است.	اشمیت، ۲۰۲۲: ۱ استوارت، ۲۰۲۳: ۲ زانگ و همکاران، ۲۰۲۱: ۲۷ بریکستد و همکاران، ۲۰۲۳: ۱
			الگوی حکمرانی مناسب هوش مصنوعی بایستی به نگرانی‌های شناختی، اخلاقی و امنیتی توجه داشته باشد. الگوی حکمرانی مناسب هوش مصنوعی به عنوان یک متغیر تعدیل‌کننده یک جنبه مهم	مانتیمای و همکاران، ۲۰۲۳: ۱ بریکستد و همکاران، ۲۰۲۳: ۱

مفهوم	بعد	مؤلفه	کد استخراجی	منبع
			برای اطمینان از شفافیت، توجه به اصول اخلاقی و قابل اعتماد بودن متغیر مستقل هوش مصنوعی است. حکمرانی مناسب هوش مصنوعی می‌تواند با تدوین هنجارهای رفتاری و دستورالعمل‌های اخلاقی برای توسعه دهندگان، ایجاد سازوکارهایی برای ارزیابی تأثیر اجتماعی و اقتصادی هوش مصنوعی و ایجاد بازبینه‌های نظارتی، از استفاده ایمن و قابل اعتماد از هوش مصنوعی اطمینان حاصل نماید. با اجرای حکمرانی مناسب هوش مصنوعی، هوش مصنوعی سازمان‌ها را ارتقا می‌دهد و به آنها قدرت می‌دهد تا به جای کاهش سرعت، با اعتماد و چابکی کامل و متعهد به اخلاق کار کنند.	۱۳۳. همکاران، ۲۰۲۳: بریکستد و ۱۷۳. ماتیمای و همکاران، ۲۰۲۳: ۱۸ برننه و همکاران، ۲۰۲۱: ۴ استوارت، ۲۰۲۳: ۲ دافو، ۲۰۱۸: ۷ موکاندر و همکاران، ۲۰۲۱: ۲ راسکا، ۲۰۲۳: ۲۳۳ ویسمولر: ۲۰۲۳: ۲۹
			نامناسب بودن الگوی حکمرانی می‌تواند منجر به تنظیم‌گری و قانون‌گذاری نامناسب هوش مصنوعی گردد. از آنجا که بازیگران اصلی و فعال این فناوری غالباً در بخش‌های غیردولتی و تجاری مستقر هستند، تنظیم‌گری نامناسب هوش مصنوعی می‌تواند منجر به ظهور قدرت سایبری مبتنی بر اراده و رفتار بازیگران غیردولتی و تشدید خلاء قانونی در این زمینه گردد. پذیرش فناوری‌های هوش مصنوعی در ابعاد مختلف قدرت سایبری، نگرانی‌هایی را در مورد انطباق با سایر مقررات مانند مقررات حفاظت از داده‌های عمومی ایجاد می‌کند.	

مرحله ششم: واپایش کیفی یافته‌ها: در این پژوهش برای روایی توصیفی سعی شده است تا حد امکان، بیشترین تعداد مقاله‌های مرتبط شناسایی و گردآوری شود. برای روایی نظری سعی شده است تا پژوهش‌هایی مورد استفاده قرار گیرد که از اعتبار علمی بالایی به‌ویژه از نظر ارجاع مقالات علمی، برخوردار باشند. در نهایت به‌منظور روایی تفسیری، به این صورت عمل شد که از چهار نفر پژوهشگر به‌عنوان کدگذار و مفسر استفاده شد و در جلسات هماهنگی توافق نهایی در مورد کدهای مورد استفاده به دست آمد.

همچنین؛ برای سنجش پایایی کدهای استخراج شده از روش توافق بین دو کدگذار استفاده شد، به صورتی که علاوه بر پژوهشگر که اقدام به کدگذاری اولیه نموده است پژوهشگر دیگری نیز همان متنی که خود پژوهشگر کدگذاری نموده را بدون اطلاع از کدهای وی به طور جداگانه کدگذاری می‌کند. در صورتی که کدهای این دو پژوهشگر به هم نزدیک باشد نشان‌دهنده توافق بالای این دو کدگذار است. برای محاسبه ضریب توافق دو کدگذار از ضریب کاپا استفاده شد. مقدار ۰ محاسبه شد که در جدول (۶) نشان داده شده / ۰ عدد ۸۳۲ / در سطح معناداری ۰۰۰ شاخص با استفاده از نرم‌افزار SPSS است. استخراج کدها از پایایی مناسبی برخوردار بوده است.

جدول ۶. مقدار اندازه توافق بین دو کدگذار

مقدار	انحراف استاندارد	سطح معناداری
۰/۸۳۲	۰/۰۴۱	۰/۰۰۰
۱۹۶		
ضریب کاپای مورد توافق		
تعداد مورد معتبر		

مرحله هفتم: ارائه یافته‌ها: در نهایت در مرحله هفتم یافته‌های تحقیق بیان می‌شود. همان‌طور که در گام پنجم اشاره شد، در این پژوهش بر اساس مطالعات پیشین و کدهای استخراج شده، در نهایت ۱۴ مؤلفه و ۴ بعد شناسایی و آزمون کیفیت آنها تأیید شده است. در این گام الگوی مفهومی پژوهش به شرح شکل (۳) ارائه می‌شود.



شکل ۳. الگوی مفهومی کاربست حکمرانی فناوری هوش مصنوعی با رویکرد ارتقاء قدرت سایبری

۵. نتیجه‌گیری و پیشنهاد

۵-۱. نتیجه‌گیری

در این پژوهش با رویکرد فراترکیب یک‌طبقه بندی مناسب بر اساس پیشینه تحقیقات انجام شده ارائه گردید. بدین منظور ابتدا پیشینه پژوهش در این خصوص بررسی شد و سپس بر اساس رویکرد مذکور و مطالعه ۷۱ پژوهش معتبر در حوزه مورد مطالعه، به جمع‌بندی از نظریات مختلف و اشباع نظری پیرامون موضوع حکمرانی هوش مصنوعی، پرداخته شد. با توجه به یافته‌های این پژوهش، الگوی حکمرانی هوش مصنوعی برای ارتقای قدرت سایبری، شامل ۴ بعد عوامل فنی و قابلیت، عوامل سیستمی، عوامل سازمانی و عوامل محیطی و فراسازمانی است که متشکل از ۱۴ مؤلفه است. البته اعتبارسنجی دسته‌بندی‌ها، با مراجعه به جامعه خبرگی صورت گرفت که این مهم به تحصیل اعتبار خروجی نهایی پژوهش در قالب شکل (۳) منجر شد.

الگوی حکمرانی هوش مصنوعی به عنوان موضوع تحقیقاتی نوظهور برای ارتقای قدرت سایبری می‌تواند تضمین کند که هوش مصنوعی در جایی که از کلان داده‌ها و الگوریتم‌های یادگیری ماشین برای تصمیم‌گیری استفاده می‌شود، به صورت اخلاقی، قابل اعتماد و مسئولانه عمل می‌کند، نوآوری و بهره‌وری را در همه عناصر قدرت ارتقا می‌دهد، نگرانی‌های امنیتی و حریم خصوصی را برطرف

می‌کند، رویه‌های پاسخگویی را فراهم می‌کند، چارچوب‌های نظارتی ایجاد می‌کند، و مسائل مربوط به سوءاستفاده از هوش مصنوعی در ارتقای قدرت سایبری را برطرف می‌کند. حکمرانی مناسب هوش مصنوعی قادر به ساخت جامعه‌ای هوشمند و توانمند باهوش مصنوعی مسئول و پاسخگو است و خطرات ناشی از هوش مصنوعی برای جامعه بشری را کاهش می‌دهد. تأثیر مثبت هوش مصنوعی بر افزایش قدرت سایبری و اثرات آن بر قدرت ملی و امنیت ملی به‌عنوان یک راهبرد جدید به حکمرانی مناسب هوش مصنوعی کمک می‌کند.

باتوجه به اینکه نتایج این مطالعه، کمک شایان و درخور توجهی به ترویج گفتمان آن در جامعه، به‌ویژه جامعه علمی و اجرایی، خواهد کرد. نتایج فوق می‌تواند، بینشی ارزشمند به سیاست‌گذاران و تنظیم‌گران این حوزه که درصدد بهبود و ارتقا قدرت سایبری به واسطه حکمرانی هوش مصنوعی هستند، ارائه دهد. این پژوهش علاوه بر این، می‌تواند در سطوح اجرایی توسط مدیران سطوح میانی مورد استفاده قرار گیرد. در واقع نتایج این پژوهش به پژوهشگران و دست‌اندرکاران این حوزه کمک می‌کند تا بدانند، باید به چه ابعاد و مضامینی توجه کنند. چراکه، نتایج این پژوهش نشان داد از آنجاکه بازیگران اصلی و فعال این فناوری غالباً در بخش‌های غیردولتی و تجاری مستقر هستند، تنظیم‌گری نامناسب هوش مصنوعی می‌تواند منجر به ظهور قدرت سایبری مبتنی بر اراده و رفتار بازیگران غیردولتی و تشدید خلأ قانونی در این زمینه گردد. در چنین شرایطی ضمانت‌اجرائی به‌کارگیری فناوری هوش مصنوعی در همه عناصر قدرت ملی با تردیدهای اخلاقی، امنیتی (حریم خصوصی) و مسئولیت‌پذیری مواجه گردیده، منجر به پیامدهای منفی خواهد شد.

۲-۵. پیشنهادها

در نهایت باتوجه به نتایج به‌دست‌آمده از این پژوهش، پیشنهادهای زیر ارائه می‌گردد:

۱. به‌منظور کاربست الگوی مذکور در سطح سازمانی، پیشنهاد می‌شود که مسئولان مرتبط از تشکیل و راه‌اندازی میزهای پژوهشی مستقل باهدف ایجاد هماهنگی، ارائه اطلاعات و ارتباط با شبکه نخبگان، حمایت کنند تا از این طریق بتوان در سیاست‌گذاری‌های مرتبط با این حوزه از نظرات شبکه نخبگانی مذکور بهره برد.

۲. همچنین دوره‌های آموزشی جامع برای مدیران اجرایی در خصوص شناخت و تبیین ابعاد و مؤلفه‌های الگوی مذکور برگزار شود.

۳. به پژوهشگران علاقه‌مند در زمینه حکمرانی هوش مصنوعی و دغدغه‌مندان در حوزه قدرت سایبری پیشنهاد می‌شود که با انجام مطالعات تطبیقی و باتوجه‌به عملکرد موفق سایر کشورها در زمینه حکمرانی هوش مصنوعی، به بهینه‌کاوی و بومی‌سازی تجارب موفق در گام‌های اجرایی و البته با رویکرد ارتقاء قدرت سایبری اقدام کنند و شاخص‌های این الگوی را به طور شفاف و گام‌به‌گام، در این زمینه ارائه نمایند.

منابع

الف- فارسی

۱. رمضان‌زاده، مجید و همکاران (۱۳۹۹)، ارائه الگوی مفهومی ارزیابی قدرت سایبری نیروهای مسلح با تأکید بر بعد بازدارندگی، *فصلنامه مدیریت نظامی*، دوره بیستم، شماره ۷۸، صص ۹۲-۶۱
۲. قاسمی، علی و همکاران (۱۴۰۲)، الگوی ارزیابی قدرت آفند سایبری ارتش ج.ا.ایران، *فصلنامه علوم و فنون نظامی*، دوره ۱۹، شماره ۶۴، صص ۳۶-۶۱

ب- انگلیسی

الف) کتب انگلیسی

۱. Haunschild, J., Kaufhold, M. A., & Reuter, C. (2022). Cultural violence and fragmentation on social media: Interventions and countermeasures by humans and social bots. In *Cyber security politics* (pp. 48–63). Routledge.
۲. Sánchez-Marrè, M. (2022). *Intelligent decision support systems*. Springer International Publishing.
۳. Schembri, J., Gentile, R., & Galasso, C. (2023). Enhancing natural-hazard exposure modeling using natural language processing: A case-study for Maltese planning applications. *Procedia Structural Integrity*, 44.
۴. Voenekey, S., Kellmeyer, P., Mueller, O., & Burgard, W. (Eds.). (2022). *The Cambridge handbook of responsible artificial intelligence: Interdisciplinary perspectives*. Cambridge University Press.
۵. Agrawal, A., Gans, J., & Goldfarb, A. (2018). *Prediction machines: The simple economics of artificial intelligence*. Harvard Business Press.
۶. Chakraborty, S. P., Dashora, P., & Gupta, S. (2022). Application of machine learning in open government database. In *Impact of artificial intelligence on organizational transformation*. Wiley.
۷. Chouhan, A., Chutia, D., & Raju, P. L. N. (2021). Deep learning applications on very high-resolution aerial imagery. In *Artificial intelligence*. Chapman and Hall/CRC.
۸. Cohen, C., Nye Jr., J. S., & Armitage, R. L. (2007). *A smarter, more secure America*.
۹. Nye, J. S. (2010). *Cyber power*. Belfer Center for Science and International Affairs, Harvard Kennedy School.
۱۰. Raska, M., & Bitzinger, R. A. (Eds.). (2023). *The AI wave in defence innovation: Assessing military artificial intelligence strategies, capabilities, and trajectories*. Taylor & Francis.
۱۱. Sandelowski, M., & Barroso, J. (2007). *Handbook for synthesizing qualitative research*. Springer.
۱۲. Ventre, D. (2020). *Artificial intelligence, cybersecurity and cyber defense*. ISTE Ltd and John Wiley & Sons.
۱۳. Vuuren, J. V., Plint, G., Leenen, L., Zaaïman, J., Kadyamatimba, A., & Phahlahmohlaka, J. (2016). *Building blocks for national cyberpower*.

۱۴. Vuuren, J. V., & Leenen, L. (2019). Building blocks and measurement of national cyberpower. In M. Sarfraz (Ed.), *Developments in information security and cybernetic wars* (pp. xx-xx). IGI Global.
۱۵. Wiesmüller, S. (2023). *The relational governance of artificial intelligence: Forms and interactions*. Springer Nature.
16. Zekos, G. I. (2022). *Political, economic and legal effects of artificial intelligence: Governance, digital economy and society*. Springer Nature

ب) مقالات انگلیسی

۱. Arrieta, A. B., Díaz-Rodríguez, N., Del Ser, J., Bennetot, A., Tabik, S., Barbado, A., ... & Herrera, F. (2020). Explainable artificial intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information Fusion*, 58, 82–115.
۲. Arslan, A. S. (2023). Neorealist analysis of security dilemma in cyberspace: A quantitative study. *APSA Preprints*.
۳. Berente, N., Gu, B., Recker, J., & Santhanam, R. (2021). Managing artificial intelligence. *MIS Quarterly*, 45(3).
۴. Birkstedt, T., Minkinen, M., Tandon, A., & Mäntymäki, M. (2023). AI governance: Themes, knowledge gaps and future agendas. *Internet Research*, 33(7), 133–167.
۵. Bonfanti, M. E. (2022). Artificial intelligence and the offence-defence balance in cyber security. In *Cyber security: Socio-technological uncertainty and political fragmentation* (pp. 64–79). Routledge.
۶. Boni, M. (2021). Artificial intelligence and governance: Going beyond ethics. *European View*, 20(2).
۷. Butcher, J., & Beridze, I. (2019). What is the state of artificial intelligence governance globally? *The RUSI Journal*, 164(5–6), 88–96.
۸. Dafoe, A. (2018). AI governance: A research agenda. *Future of Humanity Institute*. Retrieved from <https://www.fhi.ox.ac.uk/>
۹. Duan, Y., Edwards, J. S., & Dwivedi, Y. K. (2019). Artificial intelligence for decision making in the era of big data: Evolution, challenges and research agenda. *International Journal of Information Management*, 48, 63–71.
۱۰. Eriksson, T., Bigi, A., & Bonera, M. (2020). Think with me, or think for me? On the future role of artificial intelligence in marketing strategy formulation. *The TQM Journal*, 32(4).
۱۱. Guembe, B., Azeta, A., Misra, S., Osamor, V. C., Fernandez-Sanz, L., & Pospelova, V. (2022). The emerging threat of AI-driven cyber-attacks: A review. *Applied Artificial Intelligence*, 36(1), 2037254.
۱۲. Johnson, J. (2020). Artificial intelligence: A threat to strategic stability. *Strategic Studies Quarterly*, 14(1), 16–39.
۱۳. Kaur, R., Gabrijelčič, D., & Klobučar, T. (2023). Artificial intelligence for cybersecurity: Literature review and future research directions. *Information Fusion*, 101804.
۱۴. Kuehl, D. T. (2009). From cyberspace to cyberpower: Defining the problem. In *Cyberpower and national security* (pp. 30–xx).
۱۵. Laato, S., Tiainen, M., Islam, A. K. M. N., & Mäntymäki, M. (2022). How to explain AI systems to end users: A systematic literature review and research agenda. *Internet Research*, 32(7).

۱۶. Mäntymäki, M., Minkkinen, M., Zimmer, M. P., Birkstedt, T., & Viljanen, M. (2023). Designing an AI governance framework: From research-based premises to meta-requirements. *ECIS 2023 Research Papers*.
۱۷. Martin, K. (2019). Ethical implications and accountability of algorithms. *Journal of Business Ethics*, 160, 835–850.
۱۸. Matsuzaka, Y., & Yashiro, R. (2023). AI-based computer vision techniques and expert systems. *AI*, 4(1).
۱۹. Mijwil, M., Salem, I. E., & Ismaeel, M. M. (2023). The significance of machine learning and deep learning techniques in cybersecurity: A comprehensive review. *Iraqi Journal for Computer Science and Mathematics*, 4(1), 87–101.
۲۰. Mökander, J., & Floridi, L. (2021). Ethics-based auditing to develop trustworthy AI. *Minds and Machines*, 31(2), 323–327.
۲۱. Nye, J. S. (2009). Get smart: Combining hard and soft power. *Foreign Affairs*.
۲۲. Schmitt, L. (2022). Mapping global AI governance: A nascent regime in a fragmented landscape. *AI and Ethics*, 2(2), 303–314.
۲۳. Soni, V. D. (2019). Role of artificial intelligence in combating cyber threats in banking. *International Engineering Journal for Research & Development*, 4(1), 7.
۲۴. Tontiwachwuthikul, P., Chan, C. W., Zeng, F., Liang, Z., Sema, T., & Chao, M. (2020). Recent progress and new developments of applications of AI, KBS, and ML in the petroleum industry. *Petroleum Journal*, 6(4), 319–320.
۲۵. Zhang, D., Pee, L. G., & Cui, L. (2021). Artificial intelligence in e-commerce fulfillment: A case study of resource orchestration at Alibaba's smart warehouse. *International Journal of Information Management*, 57, 102304.
۲۶. Schmidt, E., Work, B., Catz, S., Chien, S., Darby, C., Ford, K., ... & Matheny, J. (2021). *National security commission on artificial intelligence (AI)*. National Security Commission on AI.

ج) وبگاه‌ها

۱. Hafiz Sheikh Adnan, A. (2022). Developing an artificial intelligence governance framework. *ISACA*. <https://www.isaca.org/resources/news-and-trends/newsletters/atisaca/2022/volume-38/developing-an-artificial-intelligence-governance-framework>
۲. Stewart, I. J. (2023). Why the IAEA model may not be best for regulating artificial intelligence. *Bulletin of the Atomic Scientists*. <https://thebulletin.org/2023/06/why-the-iaea-model-may-not-be-best-for-regulating-artificial-intelligence/>

COPYRIGHTS

۲۰۲۵ by the authors. Published by The National Defense University. This article is an open-access article distributed under the terms and conditions of the Creative Commons Attribution 4.0 International (CC BY 4.0) <https://creativecommons.org/licenses/by/4.0>

