

## نقش هوش مصنوعی در مدیریت مخاطرات امنیت داخلی با رویکرد حقوقی

نوید الهوتی<sup>۱\*</sup>، مهدی منصوری<sup>۲</sup>

### چکیده

پیشرفت‌های بی‌سابقه در حوزه هوش مصنوعی، تحولات شگرفی را در ابعاد مختلف زندگی بشر به همراه داشته است، به‌ویژه در زمینه امنیت داخلی که ازجمله مهم‌ترین عرصه‌ها برای استفاده از این فناوری به‌شمار می‌رود. این فناوری با بهره‌گیری از قابلیت‌های پیشرفته خود در شناسایی، تحلیل، پیش‌بینی و مقابله با مخاطرات، موجب افزایش چشمگیر توانمندی‌های نهادهای امنیتی شده است. به لطف هوش مصنوعی، سیستم‌های امنیتی قادر به پردازش حجم عظیمی از داده‌ها با سرعت و دقت بالا هستند که این امر به تشخیص سریع‌تر مخاطرات و واکنش مؤثرتر در برابر آن منجر می‌شود. با این حال، بهره‌گیری از هوش مصنوعی در حوزه امنیت داخلی، چالش‌های حقوقی قابل توجهی را نیز به همراه دارد. این چالش‌ها شامل مسائل مرتبط با حفظ حریم خصوصی افراد، تعیین مسئولیت در صورت بروز خطا یا سوءاستفاده، خطرات ناشی از تبعیض الگوریتمی و همچنین مخاطرات امنیت سایبری است. این مقاله به بررسی جامع این چالش‌ها می‌پردازد و تلاش می‌کند تا راهکارهایی عملی برای بهره‌برداری امن، کارآمد و مسئولانه از هوش مصنوعی در حوزه امنیت داخلی ارائه دهد، همچنین در این راستا تأکید ویژه‌ای بر ضرورت تدوین مقررات حقوقی مناسب و شفاف، آموزش و توسعه نیروی انسانی متخصص در زمینه هوش مصنوعی، و تقویت همکاری‌های بین‌المللی برای مقابله با مخاطرات جهانی مطرح می‌شود. از این رو، پژوهش حاضر تلاش می‌کند تا با ارائه بینشی جامع و عملی، به بهبود فرایندهای امنیتی و کاهش مخاطرات حقوقی و اخلاقی مرتبط با استفاده از هوش مصنوعی کمک کند و راه را برای بهره‌برداری مؤثر از این فناوری نوین در حوزه مدیریت مخاطرات در امنیت داخلی هموار سازد.

**کلیدواژه‌ها:** هوش مصنوعی، امنیت داخلی، مدیریت

---

۱. \* پژوهشگر دانا navidalh@chmail.i

۲. کارشناس علوم امنیتی m.۶۷۸۹@gmail.ir

## ۱. کلیات

۱-۱ - مقدمه و بیان مسئله

هوش مصنوعی به عنوان یکی از برجسته ترین دستاوردهای فناوری در عصر مدرن، تحولات گسترده ای را در ابعاد مختلف زندگی بشری به وجود آورده است. این فناوری نوین با ورود به حوزه های مختلف، به ویژه امنیت داخلی، می تواند نقش چشمگیری در بهبود و تسهیل فرایندها ایفا کند. در شرایط کنونی که مخاطرات امنیتی پیچیده تر، متنوع تر شده اند، نیاز به استفاده از ابزارهای پیشرفته برای پیش بینی و مقابله با این مخاطرات بیش از هر زمان دیگری احساس می شود. یکی از مهم ترین مخاطرات امنیتی در دنیای امروز، مخاطرات ترکیبی است که متشکل از روش های سخت و نرم برای ایجاد بی ثباتی و تهدید علیه امنیت داخلی کشورها می شود که مخاطرات مهمی تلقی می شوند. این مخاطرات به طور فزاینده ای از ابزارهای پیچیده ای همچون حملات سایبری، جنگ اطلاعاتی، به صورت تکنولوژی محور بهره می گیرند که می تواند مخاطرات جدی برای زیرساخت های کشورها به همراه داشته باشد. در این میان هوش مصنوعی با قابلیت های فراوان خود، به نهادهای امنیتی کمک می کند تا با تحلیل سریع و دقیق حجم عظیمی از داده ها، الگوهای پنهان را شناسایی و مخاطرات بالقوه را در مراحل ابتدایی کشف کنند. قدرت پردازش بی نظیر و الگوریتم های یادگیری ماشینی هوش مصنوعی، این امکان را فراهم کرده است که نهادهای مرتبط حوزه مدیریت مخاطرات بتوانند به سرعت و با دقت بالا به مخاطرات پاسخ دهند و از وقوع حملات جلوگیری کنند، اما در کنار این فرصت های بی بدیل، استفاده از هوش مصنوعی در امنیت داخلی با چالش های حقوقی پیچیده ای نیز همراه است، ماهیت خودکار و مبتنی بر الگوریتم های پیچیده هوش مصنوعی، مشکلات حقوقی جدی درخصوص مسئولیت پذیری ناشی از خطا، حفاظت از حریم خصوصی شهروندان و جلوگیری از بروز تبعیض های ناعادلانه ایجاد کرده است، برای مثال، اگر یک سیستم هوش مصنوعی به اشتباه فردی را به عنوان تهدید شناسایی کند، چالش های حقوقی متعددی درمورد مسئولیت پذیری مطرح می شود. همچنین نگرانی های زیادی درمورد نحوه استفاده از داده های شخصی شهروندان و تضمین حفظ حریم خصوصی آنان وجود دارد. علاوه بر این، الگوریتم های هوش مصنوعی ممکن است به طور ناخواسته باعث بروز تبعیض هایی شوند که ناشی از داده های نادرست یا جانبدارانه است. این چالش ها نه تنها نیازمند بررسی های دقیق حقوقی هستند، بلکه به ایجاد چارچوب های قانونی جدید نیاز دارند که بتوانند با سرعت پیشرفت فناوری همگام شوند. به همین دلیل، مطالعه و تحلیل دقیق فرصت ها و چالش های حقوقی استفاده از هوش مصنوعی در امنیت داخلی اهمیت بسیاری دارد. این مقاله با هدف بررسی این موضوع، تلاش می کند تا به تبیین فرصت ها و چالش های مرتبط با استفاده از هوش مصنوعی در امنیت داخلی بپردازد. با توجه به نقش برجسته هوش مصنوعی در عصر حاضر این پرسش مطرح می شود که آیا فرصت های بی بدیل هوش مصنوعی در امنیت داخلی می تواند بر چالش های حقوقی مرتبط با آن، از جمله مسئولیت پذیری، حفظ حریم خصوصی، و جلوگیری از تبعیض های ناعادلانه غلبه کند. این مقاله به دنبال بررسی جامع فرصت ها و چالش های حقوقی استفاده از هوش مصنوعی در حوزه امنیت داخلی است تا به این پرسش پاسخ دهد.

۱-۲ - اهمیت و ضرورت تحقیق

در دنیای کنونی که فناوری هوش مصنوعی به سرعت در حال گسترش است، کاربردهای آن در امنیت داخلی فرصت‌های بی‌نظیری برای افزایش کارآمدی، دقت و سرعت در شناسایی مخاطرات و مدیریت مخاطرات‌های امنیتی فراهم کرده است. فناوری‌های هوشمند، با امکان تحلیل داده‌های گسترده، شناسایی الگوهای تهدید و بهبود فرایندهای امنیتی، فرصت مهمی برای ارتقای امنیت داخلی کشور ایجاد می‌کنند. با این حال، کاربرد گسترده هوش مصنوعی همچنین نگرانی‌های جدی حقوقی را مطرح کرده است، از جمله مخاطرات مربوط به حریم خصوصی، احتمال سوءاستفاده از داده‌ها و پیش‌بینی‌پذیری تصمیمات امنیتی. افزون بر این، نبود قوانین جامع و مناسب برای کنترل و نظارت بر کاربردهای امنیتی هوش مصنوعی، چالش‌های جدیدی را به وجود آورده است که می‌تواند در آینده‌ای نزدیک امنیت داخلی را تحت تأثیر قرار دهد. این تحقیق می‌کوشد، با شناسایی فرصت‌ها و تهدیدهای هوش مصنوعی در حوزه امنیت داخلی، خلأهای قانونی و حقوقی موجود را بررسی کند و راهکارهایی برای تدوین چارچوب‌های قانونی منطبق با نیازهای روز ارائه دهد. با توجه به گسترش روزافزون فناوری‌های مبتنی بر هوش مصنوعی در امنیت داخلی، تحقیق حاضر اهمیت بسیاری دارد، زیرا می‌تواند به تدوین سیاست‌ها و مقرراتی کارآمد و جامع کمک کند که همگام با پیشرفت‌های فناوری، از حقوق شهروندان نیز محافظت کند و زمینه‌ای برای استفاده ایمن و مؤثر از این فناوری فراهم آورد.

### ۳-۱- پرسش پژوهش

پرسش اصلی تحقیق این است که هوش مصنوعی چه فرصت‌ها، تهدیدها و چالش‌های حقوقی را در حوزه امنیت داخلی به وجود می‌آورد و چگونه می‌توان با استفاده از چارچوب‌های حقوقی موجود، این فناوری را به نحوی هدایت کرد که به امنیت داخلی آسیب نرساند و همزمان حقوق شهروندی را نیز حفظ کند. این پرسش به دنبال شناسایی نقاط قوت و فرصت‌هایی است که هوش مصنوعی در ایجاد سامانه‌های نظارتی، تحلیل داده‌های بزرگ و افزایش کارایی در امنیت داخلی به وجود می‌آورد. همچنین، هدف این پرسش تحلیل تهدیدها و چالش‌های حقوقی است که از به کارگیری این فناوری ممکن است مشکلات نظیر نقض حریم خصوصی، تبعیض‌های احتمالی در سیستم‌های هوش مصنوعی و تداخل آن با حقوق بشر و حقوق شهروندی پدید می‌آورد.

### ۴-۱- فرضیه پژوهش

فرضیه تحقیق این است که به کارگیری هوش مصنوعی در حوزه امنیت داخلی، علاوه بر ایجاد فرصت‌های جدید در پیشگیری و مقابله با مخاطرات، چالش‌های حقوقی پیچیده‌ای را نیز به همراه دارد که بدون تنظیم و تقویت قوانین مرتبط، ممکن است به کاهش امنیت داخلی و نقض حقوق شهروندی منجر شود. این فرضیه بیان می‌کند که هوش مصنوعی با فراهم آوردن قابلیت‌های پیشرفته‌ای مانند شناسایی مخاطرات بالقوه، تحلیل الگوهای مجرمانه و پیش‌بینی رفتارهای مخاطره‌آمیز، به بهبود کارایی و کارآمدی امنیت داخلی کمک می‌کند. با این حال، عدم وجود مقررات دقیق و شفاف درخصوص استفاده از هوش مصنوعی در این زمینه می‌تواند به سوءاستفاده‌ها، تبعیض‌ها و نقض حقوق شهروندان منجر شود. به همین دلیل، نیاز به چارچوب‌های حقوقی جامع و متناسب با تحولات تکنولوژیکی وجود دارد تا بتوان به شکلی متعادل از فرصت‌های هوش

مصنوعی بهره برد و تهدیدها و چالش‌های مرتبط را کنترل و کاهش داد. همچنین این تحقیق به بررسی راهکارهایی می‌پردازد که از طریق آن‌ها می‌توان فرصت‌های ایجادشده توسط هوش مصنوعی را به نحوی سامان داد که هم به امنیت داخلی کمک کند و هم حقوق شهروندان محفوظ بماند.

#### ۵-۱- پیشینه پژوهش

علی‌رغم مطالعات انجام گرفته موارد مشابهی مشاهده نشد، اما موارد ذیل تا حدودی نزدیک به موضوع هستند.

اونیل (۲۰۲۱) در مقاله سلاح‌های مخرب ریاضی چگونه داده‌های بزرگ نابرابری را افزایش می‌دهد و دموکراسی را تهدید می‌کند،<sup>۲</sup> به تبعات منفی استفاده نادرست از الگوریتم‌ها و هوش مصنوعی در حوزه‌های مختلف از جمله عدالت کیفری، استخدام و اعتبارسنجی می‌پردازد که چگونه الگوریتم‌های معیوب می‌توانند به تقویت تبعیض و نابرابری اجتماعی کمک کنند. این پژوهش به‌طور عمده بر مبنای مطالعه موردی و تحلیل کیفی است. از داده‌های موجود و مثال‌های واقعی استفاده می‌کند تا به تحلیل و ارزیابی پیامدهای منفی الگوریتم‌ها بپردازد و بررسی عمیق‌تری از نحوه عملکرد الگوریتم‌ها در زمینه‌های مختلف داشته باشد.

دانکن نایل<sup>۳</sup> و همکاران در مقاله (۲۰۲۲) بررسی هوش مصنوعی در امنیت سایبری؛<sup>۴</sup> به بررسی کاربرد یادگیری ماشین در تشخیص مخاطرات سایبری می‌پردازند. آن‌ها نشان می‌دهند که چگونه الگوریتم‌های یادگیری ماشین می‌توانند در شناسایی الگوهای غیرعادی در داده‌های شبکه و تشخیص حملات سایبری مؤثر باشند. این مقاله به جمع‌آوری و تحلیل مقالات پیشین در زمینه استفاده از یادگیری ماشین در امنیت سایبری با روش کیفی پرداخته و نقاط قوت و ضعف الگوریتم‌های مختلف هوش مصنوعی را شناسایی کرده است و روش‌هایی را برای شناسایی مخاطرات سایبری پیشنهاد می‌دهد.

#### ۶-۱- مفهوم‌شناسی

الف) هوش مصنوعی: مجموعه‌ای از روش‌ها و تکنیک‌های کامپیوتری است که به توسعه سیستم‌ها و برنامه‌های کامپیوتری می‌پردازد که قادرند وظایفی را انجام دهند (United Nations, ۲۰۲۳: ۳) و معمولاً نیازمند هوش انسانی مانند یادگیری، استدلال و حل مسئله است. هوش مصنوعی در صدد است تا بتواند به‌صورت خودکار و هوشمندانه مانند انسان‌ها اما با قدرت بیشتری از هوش انسانی عمل کند.

---

1 O'Neil

2 Destructive mathematical weapons: How big data increases inequality and threatens democracy

3Duncan Nyale

4 Exploring Artificial Intelligence in Cybersecurity

5 Artificial intelligence - AI

ب) امنیت داخلی<sup>۱</sup>: به مجموعه‌ای از اقداماتی گفته می‌شود که به منظور حفظ نظم و انسجام عمومی کشور، تضمین حقوق شهروندان و حفاظت از امنیت کشور در برابر مخاطرات داخلی اتخاذ شده‌اند، اقدامات قانونی که برای مقابله با مخاطرات و حفظ ثبات امنیتی به کار گرفته می‌شود (European Commission, ۲۰۲۰: ۵).

ج) حقوق امنیت<sup>۲</sup>: حقوق امنیت به مجموعه‌ای از قوانین و اصول حقوقی اطلاق می‌شود که هدف آن حفظ و تأمین امنیت فردی، اجتماعی، ملی و بین‌المللی است. این حقوق شامل حفاظت از جان، مال، آزادی‌ها و کرامت انسانی افراد در برابر مخاطرات داخلی و خارجی می‌شود (Supreme Court, ۲۰۲۴: ۹۰).

#### ۷-۱- مبانی نظری

##### ۷-۱-۱- نظریه امنیت شناختی

نظریه امنیت شناختی<sup>۳</sup> به محافظت از ادراکات، نگرش‌ها و باورهای افراد و جوامع در برابر مخاطرات شناختی می‌پردازد و با پیشرفت هوش مصنوعی، این مخاطرات شکل جدید و پیچیده‌تری پیدا کرده‌اند. در این نظریه، امنیت تنها به معنای حفاظت فیزیکی یا سایبری نیست (Zuboff, ۲۰۲۰: ۱۲۳)، بلکه شامل حفاظت از فرایندهای فکری و شناختی افراد و جلوگیری از دستکاری و گمراه‌سازی آنان نیز می‌شود. هدف اصلی امنیت شناختی مقابله با عملیات‌های اطلاعاتی، تبلیغات گمراه‌کننده و دیگر روش‌های نفوذ شناختی است که با استفاده از فناوری‌های هوش مصنوعی به مخاطبان هدف منتقل می‌شوند (خدیور و همکاران، ۱۴۰۳: ۳۴).

##### ۷-۱-۱-۱- کاربردهای عملیاتی نظریه امنیت شناختی<sup>۴</sup>

دولت‌ها از امنیت شناختی برای جلوگیری از تأثیرگذاری نیروهای خارجی بر افکار عمومی و امنیت داخلی استفاده می‌کنند. امنیت شناختی به دولت‌ها کمک می‌کند تا در مواقع مخاطرات‌های اجتماعی مانند اعتراضات و ناآرامی‌ها از طریق کنترل اطلاعات و مدیریت ادراکات عمومی، از تشدید مخاطرات جلوگیری کنند. کشورهای مختلف از تکنیک‌های شناختی و تبلیغات دیجیتال برای تأثیرگذاری بر افکار عمومی در کشورهای رقیب استفاده می‌کنند. امنیت شناختی می‌تواند به مقابله با این حملات و حفاظت از ذهنیت ملی کمک کند. نظریه امنیت شناختی، با توجه به توسعه هوش مصنوعی، به عنوان یک حوزه نوظهور و بسیار مهم در امنیت داخلی شناخته می‌شود (راسل، ۱۴۰۲: ۱۰۱). این نظریه بر اهمیت محافظت از افکار و باورهای جامعه در برابر مخاطرات ناشی از اطلاعات گمراه‌کننده تأکید دارد و چالش‌های قانونی و اخلاقی بسیاری را پیش روی دولت‌ها و نهادهای حقوقی قرار می‌دهد.

---

1. Internal security

2. Security Right

3. Cognitive Security

## ۱-۷-۱-۲- چالش‌ها و مسائل حقوقی نظریه امنیت شناختی

نظریه امنیت شناختی با چالش‌های حقوقی و اخلاقی پیچیده‌ای همراه است بسیاری از داده‌های شناختی و رفتاری افراد از طریق تحلیل‌های هوش مصنوعی به دست می‌آید که ممکن است به حریم خصوصی آن‌ها لطمه وارد کند. برای نمونه، جمع‌آوری اطلاعات شخصی از شبکه‌های اجتماعی بدون رضایت افراد، می‌تواند نقض حریم خصوصی تلقی شود. نظارت بر محتوای شناختی و اطلاعاتی که مردم دریافت می‌کنند، ممکن است به نقض آزادی بیان منجر شود. کنترل و مقابله با اطلاعات نادرست و تبلیغات گمراه‌کننده باید به گونه‌ای انجام شود که حقوق افراد برای دریافت و انتشار اطلاعات رعایت شود. اقدامات امنیت شناختی با استفاده از هوش مصنوعی، به دلیل توانایی دست‌کاری افکار و باورها، با چالش‌های اخلاقی مواجه است (میچل، ۱۴۰۳: ۱۲). هرگونه دخالت در فرایندهای فکری و شناختی افراد باید با اصول اخلاقی و رعایت حقوق شهروندی همراه باشد.

## ۱-۷-۲- نظریه امنیت ترکیبی

نظریه امنیت ترکیبی ۱ به مفهومی اشاره دارد که ترکیبی از استراتژی‌ها و ابزارهای مختلف امنیتی را برای مقابله با مخاطرات پیچیده و چندوجهی استفاده می‌کند. این نظریه با توجه به گسترش فناوری‌های نوین، از جمله هوش مصنوعی، اینترنت اشیا ۲ و مخاطرات سایبری، تلاش دارد تا یک رویکرد جامع و یکپارچه برای امنیت داخلی، با در نظر گرفتن رویکردهای ملی و جهانی و کمک به امنیت ملی ارائه دهد (Nyale & others, ۲۰۲۲: ۴۸۰). در این نظریه، به جای استفاده از یک روش امنیتی خاص، ترکیبی از رویکردهای فیزیکی، سایبری، شناختی و اطلاعاتی به کار گرفته می‌شود تا بتوان با مخاطرات متعدد و همزمان مقابله کرد.

### ۱-۷-۲-۱- کاربردهای عملیاتی نظریه امنیت ترکیبی

در مواقع مخاطرات داخلی یا جنگ‌های شناختی، استفاده از رویکردهای امنیتی ترکیبی به نیروهای امنیتی و نهادهای دولتی کمک می‌کند تا به طور همزمان به مخاطرات مختلف واکنش نشان دهند. برای مثال، در مخاطرات های اجتماعی، می‌توان از ابزارهای سایبری برای مدیریت اطلاعات از نیروی فیزیکی برای حفظ نظم استفاده کرد. استفاده از سیستم‌های امنیتی ترکیبی به ویژه در محیط‌های حساس مانند نیروگاه‌ها، سیستم‌های مالی و زیرساخت‌های حیاتی کمک می‌کند تا مخاطرات سایبری و فیزیکی به طور همزمان شناسایی و مدیریت شوند. استفاده از ابزارهای هوش مصنوعی در تجزیه و تحلیل محتوای دیجیتال و رسانه‌ها می‌تواند به شناسایی و مقابله با جنگ‌های شناختی و اطلاعاتی کمک کند و از انتشار اطلاعات نادرست و اخبار کذب جلوگیری کند (تودت، ۱۴۰۲: ۱۳۴).

## ۱-۷-۲-۲- چالش‌ها و مسائل حقوقی نظریه امنیت ترکیبی

یکی از بزرگ‌ترین چالش‌های این نظریه، احتمال نقض حقوق بشر و آزادی‌های فردی است. نظارت گسترده و استفاده از ابزارهای دیجیتال برای شناسایی مخاطرات می‌تواند به نقض حریم خصوصی و آزادی‌های شخصی منجر شود. همچنین، استفاده از سیستم‌های هوش مصنوعی برای تحلیل رفتارهای اجتماعی و سیاسی ممکن است به سرکوب آزادی بیان و ایجاد فضای نابرابری‌های اجتماعی بینجامد. در سامانه‌های امنیت ترکیبی، داده‌ها به‌طور مداوم جمع‌آوری، تجزیه و تحلیل می‌شوند. چالش بزرگ این است که چگونه این داده‌ها به‌طور صحیح و شفاف استفاده شوند تا مانع از سوءاستفاده و نقض حریم خصوصی شوند. قوانین و مقررات خاصی برای مدیریت داده‌ها و اطمینان از استفاده قانونی و اخلاقی از آن‌ها لازم است (حسینی مقدم، ۱۴۰۲: ۷۸).

## ۲- روش‌شناسی تحقیق

نوع تحقیق توصیفی - تحلیلی، روش تحقیق کتابخانه‌ای با استفاده از بررسی مقالات علمی، گزارش‌های پژوهشی، قوانین و مقررات مرتبط، و سایر منابع معتبر در زمینه هوش مصنوعی، حقوق و علوم امنیتی است تا فرصت‌ها و مخاطرات استفاده از هوش مصنوعی در تقویت امنیت داخلی و مخاطرات علیه امنیت داخلی را بررسی کند. همچنین این پژوهش درصدد است تا ضمن بررسی مسائل حقوقی و امنیتی مرتبط با هوش مصنوعی در زمینه امنیت داخلی کمک کند تا راهکارهایی برای سازگاری با چالش‌های آن در سطح ملی و بین‌المللی برای تقویت امنیت داخلی، دست یابد.

## ۳- تجزیه و تحلیل یافته‌های تحقیق

### ۳-۱- فرصت‌های هوش مصنوعی در امنیت داخلی

با پیشرفت سریع فناوری‌های هوش مصنوعی، فرصت‌های گسترده‌ای برای بهبود کارایی و اثربخشی سیستم‌های امنیت داخلی فراهم شده است. این فرصت‌ها نه تنها به ارتقای کارایی نیروهای امنیتی منجر می‌شود، بلکه می‌تواند به پیشگیری مؤثرتر از وقوع مخاطرات و کاهش هزینه‌های مرتبط با آن کمک کند.

#### ۳-۱-۱- تقویت توانایی‌های شناسایی و پیش‌بینی مخاطرات

هوش مصنوعی با پردازش حجم عظیمی از داده‌ها در زمان کوتاه قادر است تا الگوهای پیچیده مخاطرات علیه امنیت داخلی را شناسایی و پیش‌بینی کند، این امر به‌ویژه در شناسایی مخاطرات نوظهور و پیچیده بسیار مؤثر است. با استفاده از الگوریتم‌های یادگیری ماشین، هوش مصنوعی می‌تواند رفتارهای غیرعادی افراد و سیستم‌ها را شناسایی کند که ممکن است نشانه‌ای از یک تهدید باشد (Bostrom, ۲۰۲۰: ۱۲). هوش مصنوعی می‌تواند با تحلیل داده‌های مختلف، سیستم‌های هشدار زودهنگام را توسعه دهد که به دست‌اندرکاران امنیتی اجازه می‌دهد تا قبل از وقوع حادثه، اقدامات لازم را انجام دهند.

#### ۳-۱-۲- پردازش حجم عظیم داده‌ها در زمان واقعی

سیستم‌های مبتنی بر هوش مصنوعی قادرند حجم عظیمی از داده‌ها را از منابع مختلفی مانند شبکه‌های اجتماعی، تجمع و پردازش فعالیت‌های دوربین‌های امنیتی، سنسورها و پایگاه‌های داده جمع‌آوری و در زمان واقعی پردازش کنند. این امر به شناسایی مخاطرات الگوهای پنهان و غیرقابل تشخیص برای انسان در داده‌ها کمک شایانی می‌کند تا با تحلیل دقیق داده‌ها، هوش مصنوعی می‌تواند این مخاطرات را که با روش‌های کلاسیک و متداول شناسایی نمی‌شوند، شناسایی کند این شناسایی مخاطرات اعم از فعالیت‌های خرابکارانه، حملات سایبری، یا حتی نشانه‌های اولیه وقوع حادثه تروریستی باشند (Russell&Norving, ۲۰۲۰: ۸۹).

### ۳-۱-۳- تحلیل رفتار غیرعادی

ایجاد پروفایل‌های رفتاری که هوش مصنوعی می‌تواند با تحلیل داده‌های مربوط به رفتار افراد، پروفایل‌های رفتاری نرمال و غیرنرمال ایجاد کند (Howe&Elenberg ۲۰۲۰: ۳۴). این پروفایل‌ها می‌توانند شامل الگوهای شخصیتی، سبک زندگی، فعالیت‌های برخط، و تعاملات اجتماعی روزمره در شبکه‌های اجتماعی باشند. شناسایی انحرافات که با هرگونه انحراف قابل توجه از پروفایل‌های رفتاری نرمال می‌تواند به‌عنوان یک نشانه هشداردهنده تلقی شود و نیاز به بررسی بیشتر داشته باشد.

### ۳-۱-۴- پیش‌بینی مخاطرات

مدل‌سازی پیش‌بینی مخاطرات که با استفاده از الگوریتم‌های یادگیری ماشین، انجام می‌گیرد که هوش مصنوعی می‌تواند مدل‌های پیش‌بینی ایجاد کند که با مدل‌سازی پیش‌بینی احتمال وقوع یک تهدید خاص را بر اساس داده‌های موجود تخمین می‌زند (راسل، ۱۴۰۲: ۷۸). این مدل‌ها قادرند تا سناریوهای مختلفی را برای وقوع یک تهدید شبیه‌سازی کنند، بدین ترتیب به دست‌اندرکاران مربوطه با توسعه سناریوهای احتمالی کمک کنند تا برای مقابله با مخاطرات احتمالی آماده شوند.

### ۳-۱-۵- توسعه سیستم‌های هشدار زودهنگام

سیستم‌های هشدار زودهنگام مبتنی بر هوش مصنوعی، تحولی شگرف در حوزه امنیت و پیش‌بینی مخاطرات ایجاد کرده‌اند. شناسایی علائم اولیه تهدید توسط هوش مصنوعی به گونه‌ای است که می‌تواند علائم اولیه یک تهدید را شناسایی کند که ممکن است با تحلیل انسان‌محور قابل تشخیص نباشد. اعلان هشدار به موقع با شناسایی زودهنگام مخاطرات، به دست‌اندرکاران مربوطه می‌تواند با پیش‌بینی و اقدامات پیشگیرانه لازم را انجام دهند و از وقوع زیان‌های امنیتی جلوگیری کنند (COWLS, ۲۰۲۰: ۳۱).

### ۳-۱-۶- خودکارسازی فرایندهای امنیتی

کاهش دخالت انسانی در خودکارسازی بسیاری از فرایندهای امنیتی مانند تحلیل داده‌ها، شناسایی مخاطرات و پاسخ به آن‌ها می‌تواند با استفاده از هوش مصنوعی بدون دخالت انسانی و فقط با نظارت انسانی انجام شوند. افزایش کارایی با استفاده از خودکارسازی فرایندها به افزایش سرعت و دقت در انجام وظایف امنیتی منجر می‌شود (علیزاده، ۱۴۰۱: ۲۱۷) تا به

دست‌اندرکاران، تحلیل‌گران و کارشناسان حوزه امنیت داخلی اجازه می‌دهد تا بر روی وظایف پیچیده‌تر و راهبردی بهتر تمرکز کنند.

#### ۷-۱-۳- ارتقای کارآیی سیستم‌های نظارتی

الگوریتم‌های یادگیری عمیق<sup>۱</sup> قادرند چهره‌ها را با دقت بالا از میان جمعیت شناسایی کنند، حتی اگر فرد در تلاش برای مخفی کردن ویژگی‌های ظاهری خود باشد (Zarsky, ۲۰۲۰: ۲۵). این امر می‌تواند برای شناسایی مجرمان تحت تعقیب، افراد گمشده، یا کسانی که تهدید امنیتی به شمار می‌آیند، بسیار مؤثر باشد. تحلیل حرکات این فناوری می‌تواند حرکات و رفتارهای غیرعادی را از میان فعالیت‌های عادی تفکیک کرده و موارد مشکوک را در لحظه شناسایی کند. شناسایی افراد مشکوک از طریق تحلیل الگوهای حرکتی، حتی رفتارهای ناهنجار چهره و بدن می‌تواند رفتارهای مشکوک یا خلاف هنجار را در زمان واقعی تشخیص دهند و از طریق هشدار به نیروهای امنیتی، واکنش سریع‌تر و دقیق‌تری را ممکن سازند.

#### ۲-۳- مخاطرات ناشی از هوش مصنوعی در امنیت داخلی

هوش مصنوعی، با تمام قابلیت‌ها و توانایی‌هایی که به ارمغان می‌آورد، چالش‌ها و مخاطراتی نیز به همراه دارد که به‌طور مستقیم بر امنیت داخلی تأثیر می‌گذارد (King, ۲۰۲۳: ۱۲۵). همچنین می‌تواند به مخاطرات‌های جدی امنیتی تبدیل شوند که به آن می‌پردازیم.

##### ۱-۲-۳- خطرات اعتماد بیش از حد به هوش مصنوعی در امنیت داخلی

اتکای بدون کنترل انسانی به این فناوری می‌تواند به بروز خطاهایی منجر شود که پیامدهای خطرناکی داشته باشد. سیستم‌های هوش مصنوعی با وجود قدرت پردازشی بالا، به داده‌ها و الگوریتم‌های اولیه وابسته‌اند در صورت بروز خطا در داده‌های ورودی، مانند داده‌های نادرست یا ناقص و بروز مشکلات در مدل‌های یادگیری ماشین، ممکن است تصمیمات اشتباهی بگیرد این تصمیمات اشتباه در حوزه امنیت داخلی می‌تواند منجر به شناسایی نادرست مخاطرات، تعقیب افراد بی‌گناه یا حتی ناتوانی در شناسایی مخاطرات واقعی شود (Alpcan et al., ۲۰۲۰: ۴۳).

##### ۲-۲-۳- حملات هدفمند و شخصی‌سازی شده

هوش مصنوعی قادر است با تحلیل حجم عظیمی از داده‌ها، پروفایل‌های دقیق و شخصی‌سازی شده از افراد ایجاد کند. این پروفایل‌ها می‌توانند برای طراحی حملات هدفمند و مؤثر مورد استفاده قرار گیرند با انجام مهندسی اجتماعی پیشرفته توسط هوش مصنوعی، مهاجمان می‌توانند حملات مهندسی اجتماعی بسیار پیچیده و باورپذیری را طراحی کنند، این حملات اغلب با استفاده از اطلاعات شخصی افراد و ایجاد حس اعتماد در آن‌ها انجام می‌شود و هوش مصنوعی می‌تواند ایمیل‌ها و پیام‌های

فیشینگ را به گونه‌ای شخصی‌سازی کند که احتمال اغوا کردن و تحریک برای فریب به جهت کلیک کردن بر روی آن‌ها توسط قربانیان به شدت افزایش یابد (Tsamadosm&others, ۲۰۲۲:۲۲۰).

### ۳-۲-۳- توسعهٔ بدافزارهای خودیادگیر

تکامل مداوم که بدافزارهای مبتنی بر هوش مصنوعی قادر به یادگیری از محیط اطراف خود و تطبیق با سیستم‌های امنیتی هستند، این بدافزارها می‌توانند به سرعت خود را تغییر دهند تا از تشخیص و حذف جلوگیری کنند. پرهیز از تشخیص به گونه‌ای که هوش مصنوعی می‌تواند به بدافزارها کمک کند تا از روش‌های تشخیص سنتی مانند امضاهای ویروس ۱ و سیستم‌های تشخیص نفوذ اجتناب کنند. حملات چندمرحله‌ای که توسط بدافزارهای هوشمند انجام می‌شود این قابلیت را دارند تا حملات چندمرحله‌ای پیچیده و فرینده را اجرا کنند که شناسایی و مقابله با آن‌ها بسیار پیچیده و دشوار خواهد بود (وانتر، ۲۰۲۱:۴۱).

### ۴-۳-۲- جعل عمیق و اطلاعات نادرست

جعل چهره و صدا توسط هوش مصنوعی می‌تواند برای ایجاد ویدئوهای جعلی بسیار واقع‌گرایانه از افراد استفاده شود. این ویدئوها می‌توانند برای ایجاد شایعات، تخریب چهره اشخاص، حتی تأثیرگذاری بر انتخابات‌ها مورد استفاده قرار گیرند (Brynjolfsson & McAfee, ۲۰۲۴: ۵۴). هوش مصنوعی قادر است تا متن‌های جعلی بسیار قانع‌کننده‌ای را تولید کند. این متن‌ها می‌توانند برای انتشار اخبار جعلی، ایجاد اختلاف، عملیات‌های روانی یا دست‌کاری افکار عمومی استفاده شوند، همچنین جعل حرفه‌ای و هوشمند اسناد که هوش مصنوعی برای طراحی جهت جعل پول، اسناد و مدارک قانونی مورد استفاده قرار می‌گیرد. این امر می‌تواند به ایجاد اقدامات علیه امنیت اقتصادی و تضررهای حقوقی در زمینه‌های مختلف قابل توجهی منجر شود.

### ۵-۳-۲- حمله به زیرساخت‌های حیاتی

سیستم‌های کنترل صنعتی هوش مصنوعی می‌تواند برای حمله به سیستم‌های کنترل صنعتی و ایجاد اختلال در عملکرد آن‌ها مورد استفاده قرار گیرد. این امر می‌تواند به وقوع حوادث بزرگ و خسارات جانی و مالی گسترده منجر شود، برنامه‌ریزی و طرح‌ریزی توسط هوش مصنوعی حمله به شبکه‌های برق می‌تواند منجر به قطع برق گسترده و ایجاد اختلال در زندگی روزمره شود (آدینه، ۲۰۲۱:۷۸). در سیستم‌های حمل و نقل نیز هوش مصنوعی می‌تواند برای حمله و ایجاد اختلال در

---

۱ Virus Signature: الگوی منحصر به فردی از ویروس که به وسیلهٔ برنامه‌های ضدویروس برای شناسایی آن‌ها مورد

استفاده قرار می‌گیرد، امضای ویروس علامتی است که ویروس‌ها بر روی برنامه‌ای که به آن‌ها حمله کردند می‌گذارند تا دیگر به آن حمله نکنند.

ترافیک مورد استفاده قرار گیرد و صرفاً حمل نقل زمینی را شامل نمی گردد بلکه شامل حمل و نقل ریلی، دریایی و هوایی نیز می شود.

### ۳-۳- چالش های حقوقی در استفاده از هوش مصنوعی در امنیت داخلی

استفاده گسترده از هوش مصنوعی در حوزه امنیت داخلی، علی رغم مزایای فراوان آن، چالش های حقوقی و پیچیده ای را به همراه دارد این چالش ها که ریشه در ماهیت این فناوری و تعامل آن با ارزش های انسانی دارند، نیازمند توجه جدی هستند (Meca, ۲۰۲۲: ۲۷).

#### ۳-۳-۱- نقض حریم خصوصی

هوش مصنوعی برای عملکرد مؤثر به حجم عظیمی از داده نیاز دارد. این امر موجب جمع آوری بی رویه اطلاعات شخصی افراد می شود که می تواند به حریم خصوصی آن ها آسیب جدی وارد کند. پایش مداوم با استفاده از سیستم های نظارتی مبتنی بر هوش مصنوعی برای پایش مداوم رفتار افراد، می تواند احساس ناامنی و کاهش اعتماد عمومی را به دنبال داشته باشد. نحوه جمع آوری، پردازش و استفاده از داده های شخصی اغلب برای افراد مشخص نیست که این امر می تواند به سوءاستفاده از اطلاعات شخصی و نقض حریم خصوصی منجر شود (O'Neil, ۲۰۲۱: ۲۳).

#### ۳-۳-۲- تبعیض و نابرابری

الگوریتم های هوش مصنوعی ممکن است حاوی تعصبات ناخودآگاه باشند که به تصمیم گیری های ناعادلانه و تبعیض آمیز منجر شود؛ مثلاً که الگوریتم هایی که برای ارزیابی ریسک جرائم استفاده می شوند، ممکن است نسبت به برخی گروه های اجتماعی تعصب نشان دهند، استفاده از هوش مصنوعی در حوزه امنیت داخلی ممکن است به تقویت نابرابری های موجود در جامعه دامن بزند و افرادی که دسترسی محدودی به فناوری دارند، ممکن است بیشتر در معرض نظارت و کنترل قرار گیرند (رستمی، ۱۴۰۱: ۴۷).

#### ۳-۳-۳- عدم تعیین مسئولیت پذیری

تعیین مسئولیت در صورت بروز خطا در صورتی که سیستم های مبتنی بر هوش مصنوعی به اشتباه عمل کنند و خساراتی را به بار آورند، تعیین مسئولیت بسیار پیچیده خواهد بود که آیا باید سازنده سیستم، کاربر یا خود سیستم مسئول شناخته شود، یا طراح و شرکت های ارائه دهنده هوش مصنوعی که این خدمات را می دهند (Dasgupta & others, ۲۰۲۰: ۸۶). شفافیت در تصمیم گیری که در بسیاری از موارد، تصمیم گیری های سیستم های هوش مصنوعی برای انسان قابل درک نیست. این عدم شفافیت می تواند به کاهش اعتماد به این سیستم ها منجر شود، همچنین سیاری از الگوریتم های هوش مصنوعی به دلیل پیچیدگی بسیار، قابل تفسیر توضیح نیستند و این عدم شفافیت، تعیین مسئولیت در صورت بروز خطا یا آسیب را دشوار می کند (وانتر، ۱۴۰۲: ۵۶).

#### ۴-۳-۳- تعارض امنیت داخلی در برابر حقوق فردی

ایجاد تعادل بین نیاز به امنیت ملی و حفظ حقوق فردی یکی از چالش‌های اصلی در استفاده از هوش مصنوعی است. ابزارهای نظارتی مبتنی بر هوش مصنوعی می‌توانند به دولت‌ها اجازه دهند تا بر شهروندان خود نظارت گسترده‌ای داشته باشند که این امر می‌تواند تهدیدی برای دموکراسی باشد (Albrecht & Binns, ۲۰۲۰: ۳۸). استفاده از هوش مصنوعی در حوزه امنیت داخلی می‌تواند به افزایش امنیت ملی کمک کند. برای مثال، می‌توان از آن برای پیش‌بینی و پیشگیری جرائم، شناسایی مخاطرات تروریستی و نظارت بر فعالیت‌های مجرمان استفاده کرد. از سوی دیگر، استفاده گسترده از هوش مصنوعی در نظارت بر افراد می‌تواند به کاهش آزادی‌های فردی و ایجاد جوّی از ترس و بدبینی منجر شود.

#### ۵-۳-۳- اخلاقیات و ارزش‌های انسانی

جهت حفظ کرامت انسانی استفاده از هوش مصنوعی در حوزه امنیت داخلی باید با حفظ کرامت انسانی سازگار باشد. برای مثال، استفاده از فناوری تشخیص چهره برای شناسایی افراد در مکان‌های عمومی می‌تواند احساس نظارت و کنترل را در افراد ایجاد کند (Cai & others, ۲۰۲۰: ۵۲). استفاده از هوش مصنوعی در برخی موارد ممکن است به کرامت انسانی آسیب برساند، برای مثال در مواردی که برای نظارت بر افراد و کنترل رفتار آن‌ها استفاده می‌شود و در موضوع عدالت تصمیم‌گیری‌های خودکار مبتنی بر هوش مصنوعی ممکن است با لطمه به عدالت، به تبعیض علیه گروه‌های خاص اجتماعی منجر شود. شفافیت بسیاری از الگوریتم‌های هوش مصنوعی به دلیل پیچیدگی بسیار زیاد، قابل توضیح و تفسیر نیستند که این امر می‌تواند باعث کاهش اعتماد عمومی به دلیل عدم درک آن توسط آحاد جامعه که ناشی از عدم تبیین است در این زمینه شود. در حوزه مسئولیت‌پذیری اجتماعی توسعه و استفاده از هوش مصنوعی باید با در نظر گرفتن مسئولیت‌های اجتماعی همراه باشد، همچنین نباید از هوش مصنوعی برای اهدافی غیر اخلاقی و غیر انسانی خودداری شود.

#### ۴-۳-۴- مقابله با چالش‌های حقوقی علیه امنیت داخلی

استفاده از هوش مصنوعی در این زمینه، چالش‌های حقوقی و اخلاقی متعددی را به همراه دارد که نیازمند بررسی دقیق و تدوین چارچوب‌های قانونی مناسب است.

#### ۱-۴-۳- چارچوب قانونی جامع

الگوریتم‌ها باید تا حد امکان شفاف باشند و نحوه تصمیم‌گیری آن‌ها قابل توضیح باشد. پاسخ‌گویی: در صورت بروز خطایا آسیب ناشی از تصمیم‌گیری‌های الگوریتمی، باید سازوکاری برای تعیین مسئولیت وجود داشته باشد. الگوریتم‌ها نباید به گونه‌ای طراحی شوند که به گروه‌های خاصی تبعیض قائل شوند. داده‌های شخصی باید به صورت محرمانه نگهداری شوند و تنها با رضایت فرد مورد نظر استفاده شوند (Tufekci, ۲۰۲۱: ۵۷). سیستم‌های هوش مصنوعی باید در برابر حملات سایبری ایمن باشند تعیین محدودیت‌هایی برای نوع داده‌هایی که می‌توان جمع‌آوری کرد و مدت زمان نگهداری آن‌ها مشخص شود. تعیین محدودیت‌هایی برای استفاده از داده‌ها ضروری است، مانند منع استفاده از داده‌ها برای اهداف غیر قانونی که مشخص

شوند. تعیین محدودیت‌هایی برای تصمیم‌گیری‌های خودکار که بر زندگی افراد تأثیر می‌گذارند. تعیین مراجع مستقل برای نظارت بر اجرای قوانین و مقررات مربوط به هوش مصنوعی بسیار حائز اهمیت است. انجام ارزیابی‌های مستمر از سیستم‌های هوش مصنوعی برای اطمینان از عملکرد صحیح و ایمن آن‌ها همواره باید مورد تأکید باشد. پاسخ‌گویی و ایجاد سازوکارهای ساده و کارآمد برای رسیدگی به شکایات مربوط به نقض قوانین و مقررات تعیین جریمه‌ها و مجازات‌های بازدارنده برای تخلفات جدی ضروری است (Hale, ۲۰۱۹: ۸۹).

#### ۲-۴-۳- سند ملی هوش مصنوعی و چالش‌های پیش رو

سند ملی هوش مصنوعی که به‌تازگی تصویب و ابلاغ شده، گامی مهم در جهت توسعه و ساماندهی حوزه هوش مصنوعی در کشور محسوب می‌شود. اما این سند نیز مانند هر سند دیگری، دارای نقاط قوت و ضعف است و برخی از چالش‌های موجود در حوزه تنظیم مقررات هوش مصنوعی را به‌طور کامل پوشش نمی‌دهد. سند ملی هوش مصنوعی به‌طور کلی به اصول و خط‌مشی‌های کلی در حوزه هوش مصنوعی پرداخته است، اما در برخی موارد به جزئیات کافی برای اجرای عملی این اصول ارائه نشده است و بیشتر بر جنبه‌های توسعه‌ای و توانمندسازی کشور در حوزه هوش مصنوعی تمرکز کرده و به‌اندازه کافی به چالش‌های حقوقی، اخلاقی و امنیتی ناشی از کاربرد هوش مصنوعی نپرداخته است (فرخی، ۱۴۰۲: ۶۷). هوش مصنوعی با سرعت بسیار بالایی در حال پیشرفت است و ممکن است مقررات موجود به سرعت منسوخ شوند. لذا نیاز به مکانیزم‌های به‌روزرسانی مداوم مقررات احساس می‌شود. برخی از مفاهیم کلیدی در حوزه هوش مصنوعی مانند هوش مصنوعی عمومی، یادگیری ماشین و شبکه‌های عصبی مصنوعی، در سند مذکور به صورت دقیق و روشن تعریف نشده‌اند که این امر می‌تواند به ابهام در تفسیر و اجرای مقررات منجر شود.

#### ۳-۴-۳- همکاری بین‌المللی

همکاری بین‌المللی در حوزه تنظیم مقررات هوش مصنوعی، کلیدی برای استفاده ایمن و مؤثر از این فناوری در سطح جهانی است، با ایجاد چارچوب‌های قانونی مناسب، توسعه استانداردهای مشترک و تقویت همکاری بین کشورها، می‌توان از مزایای هوش مصنوعی بهره‌مند شد و درعین حال از خطرات آن جلوگیری کرد. با ایجاد مجامع بین‌المللی برای تبادل اطلاعات و تجربیات در زمینه تنظیم مقررات هوش مصنوعی اقدام کرد. ضرورت تدوین استانداردهای مشترک، توسعه استانداردهای فنی و اخلاقی مشترک برای استفاده از هوش مصنوعی در سطح بین‌المللی لازم است همچنین تقویت همکاری بین‌المللی برای مقابله با مخاطرات ناشی از هوش مصنوعی همکاری بین‌المللی در مبارزه با جرائم این حوزه دارای اهمیت است (Binns, ۲۰۲۱: ۸۹).

#### ۴-۴-۳- توسعه نیروی انسانی متخصص

توسعه نیروی انسانی متخصص در حوزه هوش مصنوعی، یکی از اساسی‌ترین گام‌ها در جهت مقابله با مخاطرات سایبری و ارتقای امنیت داخلی است (Solove & Schwartz, ۲۰۲۰: ۶۷). این سرمایه‌گذاری بلندمدت، نه تنها در سطح ملی و کشوری

امکان می‌دهد تا در خط مقدم فناوری‌های نوین قرار گیرد، بلکه به تقویت توانایی‌های دفاعی و پاسخ‌گویی سریع به مخاطرات نیز کمک شایانی خواهد کرد.

#### ۴-۳-۵- گنجاندن مفاهیم هوش مصنوعی در برنامه‌های درسی

آشنایی با مفاهیم پایه هوش مصنوعی از سنین پایین، به نسل آینده کمک می‌کند تا به‌عنوان شهروندانی آگاه و متخصص در این حوزه نقش‌آفرینی کنند (Zarsky, ۲۰۲۰: ۱۱). گنجاندن این مفاهیم در برنامه‌های درسی مدارس و دانشگاه‌ها، به ایجاد جامعه‌ای دانش‌بنیان در این حوزه کمک کرده و زمینه‌ساز پرورش نسل جدیدی از متخصصان هوش مصنوعی خواهد شد.

#### ۴-۳-۶- ایجاد دوره‌های تخصصی و حمایت از تحقیقات

برگزاری دوره‌های تخصصی در سطوح مختلف، از کارشناسی تا دکتری، به تربیت نیروی انسانی متخصص و ماهر در زمینه‌های مختلف هوش مصنوعی منجر می‌شود (Regan, ۲۰۲۴: ۶۷). همچنین، حمایت از تحقیقات در این حوزه، به توسعه الگوریتم‌ها و کاربردهای جدید هوش مصنوعی در حوزه امنیت داخلی کمک کرده و باعث پیشرفت دانش در این زمینه خواهد شد.

#### ۴-۳-۸- ارتقای مهارت‌های موجود

باتوجه به سرعت بالای تغییرات در حوزه هوش مصنوعی، ارتقای مهارت‌های متخصصان موجود از اهمیت بالایی برخوردار است. برگزاری کارگاه‌ها و دوره‌های آموزشی کوتاه‌مدت، به متخصصان امکان می‌دهد تا با آخرین دستاوردهای این حوزه آشنا شده و مهارت‌های خود را بروز کنند. همچنین، ایجاد فرصت‌هایی برای تبادل دانش و تجربیات بین متخصصان داخلی و خارجی، به گسترش شبکه‌های دانش و افزایش بهره‌وری در این حوزه کمک خواهد کرد (حسینی مقدم، ۱۴۰۲: ۸۸).

#### ۴-۳-۹- ایجاد مشوق‌های شغلی

پرداخت حقوق رقابتی و ایجاد مسیرهای شغلی مشخص و جذاب، به جذب و حفظ متخصصان هوش مصنوعی کمک کرده و انگیزه‌ای برای فعالیت در این حوزه خواهد بود. همچنین، حمایت از کارآفرینانی که در حوزه هوش مصنوعی فعالیت می‌کنند، به ایجاد اکوسیستم نوآوری در این حوزه کمک کرده و به توسعه محصولات و خدمات جدید منجر خواهد شد (راسل، ۱۴۰۲: ۱۸۷).

#### ۴-۳-۱۰- ترویج فرهنگ نوآوری

ایجاد محیط‌های کاری خلاق و حمایت از ایده‌های جدید، به پرورش روحیه نوآوری و خلاقیت در بین متخصصان هوش مصنوعی کمک کرده و به ارائه راهکارهای نوآورانه در این حوزه منجر خواهد شد و ایجاد محیط‌های کاری که به نوآوری و خلاقیت تشویق کنند و حمایت از ایده‌های جدید و ارائه حمایت‌های مالی و فنی برای اجرای آن‌ها باعث پیشرفت این حوزه خواهد شد (کسینجر و همکاران، ۱۴۰۳: ۹۰).

## ۴. نتیجه گیری و پیشنهادات

### ۴-۱. نتیجه گیری

در عصر دیجیتال، هوش مصنوعی به عنوان یک نیروی تحول آفرین در حوزه های مختلف، از جمله امنیت داخلی، شناخته می شود. کاربردهای گسترده این فناوری از جمله در نظارت هوشمند، پیش بینی مخاطرات، تحلیل داده های عظیم و خودکارسازی فرایندهای امنیتی، نشان دهنده قابلیت های بی بدیل آن در ارتقای سطح امنیت داخلی است. اما در کنار این فرصت های بزرگ، چالش های حقوقی و اخلاقی مهمی نیز ظهور کرده اند که نیازمند توجه جدی و تدوین راهکارهای مناسب هستند. یکی از بزرگ ترین فرصت های هوش مصنوعی، توانایی آن در پردازش و تحلیل سریع حجم عظیمی از داده ها است که می تواند به شناسایی مخاطرات بالقوه پیش از وقوع آنها کمک کند. با این حال، این قدرت تحلیل با خطرات جدی در زمینه حریم خصوصی و نظارت گسترده مواجه است. استفاده بی رویه از هوش مصنوعی بدون در نظر گرفتن حریم خصوصی افراد، می تواند به نقض گسترده حقوق فردی منجر شود که این امر می تواند اعتماد عمومی به نهادهای امنیتی را تضعیف کند. چالش دیگر، مسئله مسئولیت پذیری در تصمیمات گرفته شده توسط سیستم های هوش مصنوعی است. در مواردی که یک سیستم هوش مصنوعی تصمیم اشتباهی بگیرد که به ضرر افراد یا نقض حقوق آنها منجر شود، مشخص نیست که چه کسی باید پاسخگو باشد؛ این امر به پیچیدگی های قانونی منجر شده و نیازمند تدوین قوانین جدید برای تعیین مسئولیت های دقیق است. از سوی دیگر، هوش مصنوعی ممکن است به تقویت تبعیض ها و بی عدالتی ها منجر شود. الگوریتم هایی که بر اساس داده های جانب دارانه آموزش دیده اند، می توانند تصمیماتی اتخاذ کنند که به تبعیض های ناعادلانه منجر شوند. این موضوع نیازمند توجه دقیق به فرایندهای آموزش و توسعه الگوریتم ها است تا اطمینان حاصل شود که این فناوری به جای تقویت نابرابری ها، به عدالت و انصاف در جامعه کمک می کند. در مجموع، استفاده از هوش مصنوعی در حوزه امنیت داخلی با فرصت های بی نظیری همراه است که می تواند به ارتقای امنیت و کاهش مخاطرات کمک کند. با این حال، برای بهره گیری از این فرصت ها، لازم است که چالش های حقوقی آن به دقت مدیریت شوند.

### ۴-۲. پیشنهادات

با اتخاذ اقدامات مناسب می توان از هوش مصنوعی به عنوان ابزاری قدرتمند و در عین حال منصفانه در حفظ امنیت داخلی

استفاده کرد و پیشنهادها زیر در این راستا ارائه می شوند:

الف) تدوین چارچوب های قانونی شفاف: برای مدیریت چالش های حقوقی هوش مصنوعی در ابتدای امنیت داخلی، لازم است که قوانین و مقرراتی شفاف و جامع تدوین شود که استفاده از این فناوری را هدایت کند و از یک سو امنیت داخلی را حفظ و از سوی دیگر، حقوق شهروندی را تضمین کند.

ب) تقویت حاکمیت داده و حریم خصوصی: باید مکانیزم‌هایی برای حفاظت از حریم خصوصی در استفاده از هوش مصنوعی ایجاد شود. این مکانیزم‌ها می‌توانند شامل رمزنگاری داده‌ها، کنترل دسترسی‌های دقیق، و روش‌های نوین برای حفظ حریم خصوصی باشند که استفاده از داده‌ها را در سطحی ایمن و کنترل‌شده ممکن می‌سازد.

ج) ایجاد ساختارهای نظارتی و ارزیابی مستمر: یکی از راهکارهای کاهش تبعیض و بی‌عدالتی در هوش مصنوعی، تعیین یا ایجاد نهادهای نظارتی مستقل است که به ارزیابی و بررسی مداوم الگوریتم‌ها و سیستم‌های هوش مصنوعی بپردازند. این نهادها باید بتوانند الگوریتم‌های جانب‌دارانه را شناسایی و اصلاح کنند و از به‌کارگیری آن‌ها جلوگیری نمایند.

د) آموزش و ظرفیت‌سازی برای مسئولان و کاربران: ضروری است که نهادهای امنیتی و سایر کاربران هوش مصنوعی در حوزه امنیت داخلی، آموزش‌های لازم را در زمینه جنبه‌های حقوقی و اخلاقی هوش مصنوعی کسب کنند، این آموزش‌ها می‌تواند به درک بهتر پیچیدگی‌های فناوری و چالش‌های مرتبط با آن کمک و استفاده مسئولانه از آن را ترویج کند.

ه) همکاری‌های بین‌المللی: باتوجه به ماهیت جهانی مخاطرات امنیتی و فناوری‌های هوش مصنوعی باید همکاری‌های بین‌المللی برای تدوین استانداردها و قوانین جهانی ضروری است. این همکاری‌ها می‌توانند به تبادل دانش و تجربه بین کشورها کمک کند و به تدوین راهبردهای مؤثرتر برای استفاده از هوش مصنوعی در امنیت داخلی منجر شوند.

و) توسعه الگوریتم‌های شفاف و مسئولیت‌پذیر: یکی از چالش‌های مهم در استفاده از هوش مصنوعی، فقدان شفافیت در نحوه تصمیم‌گیری الگوریتم‌هاست. توسعه الگوریتم‌های مسئولیت‌پذیر که بتواند فرایند تصمیم‌گیری و قبول مسئولیت را به‌طور شفاف بیان کند، می‌تواند به مقابله با چالش‌های حقوقی و کاهش نگرانی‌های مرتبط با تصمیمات نادرست کمک کند.

## منابع

۱. آدینه، رضا (۱۴۰۰). هوشمندسازی امنیت و مقابله با مخاطرات پیشرفته. تهران: مؤسسه فرهنگی هنری دیباگران.
۲. تودت، اندرو (۱۴۰۲). پرداختن به مخاطرات ترکیبی. ترجمه محمدحسین قربانی زواره. تهران: دافوس.
۳. حسینی مقدم، محمد. (۱۴۰۲). هوش مصنوعی و آینده پیشرفت علمی: گذار از علم عادی به علم پساعادی. مطالعات مدیریت کسب و کار هوشمند ۱۲(۴۵): ۷۱-۱۱۶. [https://ims.atu.ac.ir/article\\_۱۶۴۸۲.html](https://ims.atu.ac.ir/article_۱۶۴۸۲.html)
۴. حسینی مقدم، مهدی (۱۴۰۲). هوش مصنوعی و آینده آموزش دانشگاهی در ایران تهران، پژوهشکده مطالعات فرهنگی و اجتماعی.
۵. خدیور؛ آمنه، سیاه سرانی؛ حانیه و بصرای؛ رها، (۱۴۰۳). هوش مصنوعی برای مدیران، آتی نگر، تهران.
۶. راسل، استوارت (۱۴۰۲). هوش مصنوعی سازگاری با انسان و مسئله کنترل. ترجمه: فرهاد وداد، تهران، اندیشه مولانا.

۷. رستمی، محسن (۱۴۰۱). شناسایی و معرفی ظرفیت‌های کاربردی هوش مصنوعی در توسعه مضمون‌های راهبردی در سازمان‌های نظامی. فصلنامه علمی راهبرددفاعی، ۷۸، ۳۴-۷۴. [https://ds.sndu.ac.ir/article\\_2167.html](https://ds.sndu.ac.ir/article_2167.html)
۸. عظیم، علیزاده (۱۴۰۱). تبیین عوامل موثر در آینده پژوهی راهبردی امنیت داخلی، نشریه امنیت ملی، دوره ۱۴، شماره ۴۶ بهمن ۱۴۰۱ صفحه ۲۳۲-۲۰۳. [https://ns.sndu.ac.ir/article\\_2368.html](https://ns.sndu.ac.ir/article_2368.html)
۹. فرخی؛ نوید (۱۴۰۲). هوش مصنوعی رویارویی با انقلاب صنعتی چهارم، تهران، سبزان.
۱۰. کسینجر، هنری، اشمیت، اریک و هونتلوچر، دانیل (۱۴۰۳). عصر هوش مصنوعی و آینده ما انسان‌ها. ترجمه: سامان صفرزائی، تهران، پارسه.
۱۱. میچل، میلانی (۱۴۰۲). هوش مصنوعی شیوه نامه ای برای انسان متفکر. ترجمه: حمیدرضا کفاش، تهران، عابد..
۱۲. وانتر، دانیل (۱۴۰۲). هوش مصنوعی، امنیت سایبری و دفاع سایبری. ترجمه محمد رضا کریمی قهرودی، و ندا انعامی. تهران: مؤسسه آموزشی و تحقیقاتی صنایع دفاعی..

13. Albrecht, Jörn & Binns, Reuben (2020). Fairness and accountability in machine learning: A survey. *Proceedings of the 2019 Conference on Fairness, Accountability, and Transparency (1-22)*.

14. Alpcan, Tansu, Başar, Tamer, & Sastry, Shankar. (2020). Distributed control of smart grids via hierarchical game theory. In *2020 American Control Conference (ACC) (256-261) IEEE*. <https://doi.org/10.23919/ACC45564.2020.9147488>

15. Binns, Reuben. (2019). Fairness in machine learning: Lessons from political philosophy. In *Proceedings of the 2019 Conference on Fairness, Accountability, and Transparency (149-158)*. Association for Computing Machinery. <https://doi.org/10.48550/arXiv.1712.03586>

16. Bostrom, Nick (2020). *Superintelligence: Paths, dangers, strategies*. Oxford University Press.

17. Brynjolfsson, Erik & McAfee, Andrew (2024). *The second machine age: Work, progress, and prosperity in a time of brilliant technologies*. New York, W. W. Norton & Company.

18. Cai, Zhipeng & Xiong, Zuobin & Xu, Honghui & Wang, Peng & Li, Wei, & Pan, Yi (2021). Generative Adversarial Networks: A survey toward private and secure applications. *ACM Computing Surveys*, 54(6), 113-132. <https://doi.org/10.1145/345999>

19. Cows, Josh. (2020). A unified framework of digital ethics. *Mind*, 128(511), 563-598.

20. Dasgupta, Dipankar., Akhtar, Zahid, & Sen, Sajib . (2021). Machine learning in cybersecurity: A comprehensive survey. *Journal of Defense Modeling and Simulation*, 19(1), 57-106. <https://doi.org/10.1177/1548512920951275>

21. European Commission. (2020). *White paper on artificial intelligence: A European approach to excellence and trust*. [https://commission.europa.eu/publications/white-paper-artificial-intelligence-european-approach-excellence-and-trust\\_en](https://commission.europa.eu/publications/white-paper-artificial-intelligence-european-approach-excellence-and-trust_en)

- 22.Hale, Christopher (2019). *Fear of Crime: A Review of the Literature*. *International Review of Victimology*,4: 79-105.
- 23.Howe, Edmund G & Elenberg, Falcia (2020). *Ethical challenges posed by big data*. *Innovations in Clinical Neuroscience*, 17(10–12), 24–30. <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC7819582/>  
<http://arxiv.org/pdf/1908.09635>
- 24.King, Thomas C (2023). *The malicious use of AI-based deepfake techniques in cybercrime: Threats and countermeasures*. *Journal of Cybersecurity*, 9(1), 114-135. <https://doi.org/10.1093/cybsec/tyad015>
- 25.Mecaj, Stela. (2022). *Artificial intelligence and legal challenges*. *Revista Opinião Jurídica (Fortaleza)*, 20(34), 180–196 <https://periodicos.unichristus.edu.br/opiniaojuridica/article/view/4329>
- 26.Nyale, Duncan & Mwangi, Salome & Karume, Simon (2022). *A survey of contemporary and emerging research methods in information systems research*. *International Journal of Computer Applications Technology and Research*, 11(12), 478–482. IJCATR.  
<https://ijcat.com/archieve/volume11/issue12/ijcatr11121015.pdf>
- 27.O'Neil, Cathy. (2021). *Weapons of math destruction: How big data increases inequality and threatens democracy*. New York: Crown Publishing Group.
- 28.Regan, Priscilla (2024). *Legislating privacy: Technology, social values, and public policy*. New York: Springer.
- 29.Russell,Stuart J&Norvig,Peter (2020 *Artificial Intelligence A Modern Approach Third Edition*.Communications of the ACM,London,Pearson Education
- 3.Solove, Daniel & Schwartz, Paul (2020). *Information privacy law (5th ed.)*. New York: Wolters Kluwer.
- 31.Supreme Court : Expands Window to Challenge Federal Regulations Under (2024). *Jones Day Insights*. Retrieved from <https://www.jonesday.com/en/insights/2024/07/supreme-court-expands-window-to-challenge-federal-regulations-under-apa>
- 32.Tsamados, Andreas & Aggarwal, Nikita & Cows, Josh & Morley, Jessica & Roberts & Huw, Taddeo, Mariarosaria, & Floridi, Luciano. (2022). *The ethics of algorithms: Key problems and solutions*. *AI & Society*, 37(1), 215–230. <https://doi.org/10.1007/s00146-021-01154-8>
- 33.Tufekci, Zeynep (2021). *Twitter and tear gas: The power and fragility of networked protest*. New Haven: Yale University Press.
- 34.United Nations. (2023). *Secretary-General's remarks to the Security Council on Artificial Intelligence*. Retrieved from <https://www.un.org/sg/en/content/sg/speeches/2023-07-18/secretary-generals-remarks-the-security-council-artificial-intelligence>
- 35.Zarsky, Tal (2020). *Governing artificial intelligence: A roadmap for lawmakers*. *Berkeley Technology Law Journal*, 37(5), 67-78
- 36.Zuboff, Shoshana. (2020). *The age of surveillance capitalism: The fight for a human future at the new frontier of power*. New York: PublicAffairs.