

هوش مصنوعی در امنیت سایبری: فرصت‌ها، چالش‌ها و راهکارهای آینده‌نگر

عباس قدیم پور شبلو^۱، یاشار ابری^۱

^۱ دانشجوی کارشناسی ارشد مهندسی فناوری اطلاعات، دانشکده مهندسی دانشکدگان فارابی دانشگاه تهران،
ایران

{yasharabri,ghadimpourshablo}@ut.ac.ir

چکیده

در عصر دیجیتال، امنیت سایبری با رشد تهدیدات پیچیده به یکی از چالش‌های اساسی تبدیل شده است. این مقاله نقش هوش مصنوعی (AI) را در تقویت دفاع سایبری بررسی می‌کند. ابتدا به کاربردهای AI در شناسایی تهدیدات، تحلیل آسیب‌پذیری‌ها و بهبود پاسخگویی پرداخته و توانایی‌های آن در تحلیل داده‌های بزرگ و شناسایی الگوهای پیچیده را تحلیل می‌کند. چالش‌های مرتبط با AI از جمله عدم شفافیت، ملاحظات اخلاقی و خطرات مرتبط با هوش مصنوعی خصمانه نیز بررسی شده‌اند. در ادامه، کاربردهای عملی AI در سیستم‌هایی مانند Darktrace و Splunk که از خودکارسازی و تحلیل داده‌ها برای مقابله با تهدیدات سایبری استفاده می‌کنند، معرفی شده است. چالش‌های فنی و ریسک‌های مربوط به استفاده از AI، مانند تهدیدات ناشی از هوش مصنوعی خصمانه، در کنار استراتژی‌های بهینه‌سازی و تدوین چارچوب‌های قانونی و اخلاقی برای استفاده مسئولانه از این فناوری، نیز مورد بحث قرار گرفته‌اند. در نهایت، مقاله بر ضرورت مدیریت چالش‌ها و بهره‌برداری از پتانسیل AI در امنیت سایبری تأکید دارد.

کلمات کلیدی: هوش مصنوعی، امنیت سایبری، تشخیص تهدیدات سایبری، خودکارسازی امنیت سایبری، حملات سایبری مبتنی بر AI، یادگیری ماشینی، نوآوری‌های سایبری.

۱ مقدمه

در دنیای امروز، فناوری اطلاعات و ارتباطات به عنوان یکی از ستون‌های اساسی زندگی مدرن به شمار می‌آید. پیشرفت سریع در این حوزه به پیدایش چالش‌های جدیدی مانند امنیت سایبری منجر شده است [۳]. با گسترش روزافزون ارتباطات جهانی و انتقال داده‌ها به صورت آنلاین، تهدیدات سایبری بیش از هر زمان دیگری اهمیت پیدا کرده‌اند. این تهدیدات شامل حملات هکری، بدافزارها، فیشینگ و حملات توزیع‌شده منع سرویس (DDoS) هستند که می‌توانند خسارات جبران‌ناپذیری به سازمان‌ها و زیرساخت‌های دیجیتال وارد کنند [۲]. هدف این است که با بررسی دقیق و همه‌جانبه این موضوع، راهکارهایی ارائه شود که نه تنها

از تهدیدات سایبری پیشگیری کند، بلکه بتواند زمینه‌ساز بهره‌برداری بهینه از فناوری‌های نوین در جهت ایجاد یک فضای دیجیتال امن‌تر و قابل اعتمادتر باشد. در این میان، هوش مصنوعی (AI) به عنوان یک ابزار قدرتمند برای مقابله با این تهدیدات ظهور کرده است [۴]. هوش مصنوعی با استفاده از الگوریتم‌های یادگیری ماشینی و تحلیل داده‌های بزرگ، می‌تواند تهدیدات ناشناخته را شناسایی کرده و با سرعت و دقت بالا به آن‌ها پاسخ دهد [۵، ۱]. سیستم‌های امنیت سایبری مبتنی بر هوش مصنوعی می‌توانند به‌طور خودکار الگوهای غیرعادی را در شبکه‌ها شناسایی کرده و از وقوع حملات جلوگیری کنند [۲]. در خلاصه‌ای از مفاهیم هوش مصنوعی و امنیت سایبری می‌توان تعاریف هریک را این‌گونه تشریح کرد:

هوش مصنوعی (AI) : هوش مصنوعی شاخه‌ای از علوم کامپیوتر است که به توسعه سیستم‌هایی می‌پردازد که می‌توانند یادگیری، تحلیل و تصمیم‌گیری خودکار را انجام دهند. این فناوری در تحلیل داده‌های بزرگ و شناسایی الگوهای پیچیده بسیار مفید است. هوش مصنوعی در حال حاضر در زمینه‌های مختلفی از جمله پزشکی، حمل و نقل، و امنیت سایبری به کار گرفته می‌شود.

امنیت سایبری: امنیت سایبری به مجموعه‌ای از فرآیندها، فناوری‌ها و کنترل‌هایی اطلاق می‌شود که به منظور حفاظت از سیستم‌ها، شبکه‌ها و داده‌های حساس در برابر حملات سایبری، از جمله هک، بدافزار و حملات DDOS طراحی شده‌اند. هدف از امنیت سایبری، تضمین دسترسی‌پذیری، صحت و محرمانگی داده‌هاست.

۲ بررسی ادبیات و پیشینه پژوهشی

در دهه‌های اخیر، با رشد سریع فناوری‌های دیجیتال و توسعه گسترده اینترنت، تهدیدات سایبری به یکی از جدی‌ترین چالش‌های جهانی تبدیل شده است [۳]. این تهدیدات که در ابتدا بیشتر به حملات ساده و بدون پیچیدگی محدود می‌شدند، اکنون به سطحی پیچیده و هوشمند ارتقا یافته‌اند [۶]. همزمان با این تحولات، هوش مصنوعی (AI) به عنوان یکی از ابزارهای کلیدی در مبارزه با تهدیدات سایبری مطرح شده است [۴]. هوش مصنوعی با توانایی‌های منحصر به فرد خود، از جمله یادگیری ماشینی، پردازش زبان طبیعی و تحلیل داده‌های بزرگ، نقش مهمی در ارتقای امنیت سایبری ایفا می‌کند [۵].

۱.۲ نقش هوش مصنوعی در امنیت سایبری

نقش هوش مصنوعی در امنیت سایبری را می‌توان از چندین زاویه بررسی کرد. اولاً، AI به‌طور چشمگیری در بهبود فرآیندهای تشخیص تهدیدات مؤثر بوده است [۴]. در گذشته، سیستم‌های امنیتی عمدتاً به روش‌های تشخیص مبتنی بر امضا و قواعد از پیش تعیین شده متکی بودند. این روش‌ها با وجود کارایی نسبی، در برابر تهدیدات جدید و ناشناخته ناکارآمد بودند [۶]. با این حال، هوش مصنوعی با استفاده از الگوریتم‌های یادگیری عمیق و یادگیری ماشینی توانسته است این نقص‌ها را جبران کند [۱]. الگوریتم‌های هوش مصنوعی قادرند الگوهای پیچیده را در داده‌های بزرگ شناسایی کرده و تهدیدات جدید را به‌موقع تشخیص دهند [۴، ۳]. این

قابلیت‌ها نه تنها سرعت تشخیص را افزایش می‌دهند، بلکه دقت بیشتری را نیز به همراه دارند [۲]. در ادامه، جدول ۱ انواع حملات سایبری مبتنی بر هوش مصنوعی را نشان می‌دهد و الگوریتم‌های مرتبط با هر نوع حمله را معرفی می‌کند.

جدول ۱: طبقه‌بندی انواع حملات سایبری مبتنی بر هوش مصنوعی

نوع حمله	توضیحات	الگوریتم‌های معمول
فیشینگ	تلاش برای فریب کاربران به منظور افشای اطلاعات حساس	یادگیری نظارت شده
بدافزارها	نرم‌افزارهای مخربی که به سیستم‌ها آسیب می‌رسانند	یادگیری عمیق
حملات DDoS	حملات توزیع شده برای از کار انداختن خدمات آنلاین	یادگیری بدون نظارت

۲.۲ توسعه ابزارهای پیش‌بینی‌کننده مبتنی بر هوش مصنوعی

یکی از نوآوری‌های مهم هوش مصنوعی در امنیت سایبری، استفاده از مدل‌های پیش‌بینی‌کننده برای شناسایی تهدیدات سایبری قبل از وقوع آن‌ها است [۲]. این مدل‌ها با تجزیه و تحلیل الگوهای رفتاری شبکه و کاربران، می‌توانند فعالیت‌های مشکوک و ناهنجار را شناسایی کرده و هشدارهای لازم را برای جلوگیری از حملات صادر کنند [۹]. استفاده از این تکنیک می‌تواند زمان پاسخ‌دهی را کاهش داده و هزینه‌های امنیتی را به‌طور چشمگیری کاهش دهد [۱۰]. همین موضوع باعث شده است که پیش‌بینی مبتنی بر هوش مصنوعی به یکی از ارکان اصلی امنیت سایبری تبدیل شود [۶].

۳.۲ چالش‌ها و محدودیت‌های هوش مصنوعی در امنیت سایبری

با وجود تمام مزایای هوش مصنوعی، استفاده از این فناوری در امنیت سایبری با چالش‌ها و محدودیت‌های قابل توجهی همراه است [۷]. یکی از مهم‌ترین چالش‌ها، مسئله اعتماد به هوش مصنوعی است. سیستم‌های AI اغلب به عنوان «جعبه‌های سیاه» تلقی می‌شوند، زیرا فرآیند تصمیم‌گیری آن‌ها برای کاربران و مدیران سیستم‌ها شفاف نیست [۴، ۷]. این عدم شفافیت می‌تواند اعتماد به تصمیمات گرفته‌شده توسط این سیستم‌ها را کاهش دهد، به‌ویژه در شرایطی که این تصمیمات بر امنیت داده‌ها و زیرساخت‌های حساس تأثیر مستقیم دارند [۴].

۴.۲ خودکارسازی فرآیندهای دفاعی

یکی از توانایی‌های کلیدی هوش مصنوعی در امنیت سایبری، خودکارسازی فرآیندهای دفاعی است [۳]. سیستم‌های مبتنی بر هوش مصنوعی قادرند بطور خودکار تهدیدات را شناسایی کرده و بدون نیاز به دخالت انسان، اقدامات لازم برای مقابله با آن‌ها را انجام دهند [۹، ۱۰]. این قابلیت به‌ویژه در زمان بروز حملات بزرگ و پیچیده مانند حملات DDoS یا باج‌افزاری اهمیت ویژه‌ای پیدا می‌کند [۲].

۵.۲ ملاحظات اخلاقی و اجتماعی

ملاحظات اخلاقی یکی از مسائل مهمی است که در استفاده از هوش مصنوعی در امنیت سایبری باید مد نظر قرار گیرد. توسعه و استفاده از سیستم‌های هوش مصنوعی نیازمند رعایت اصول اخلاقی و احترام به حریم خصوصی افراد است [۴، ۷]. نگرانی‌هایی وجود دارد که استفاده از هوش مصنوعی می‌تواند به نقض حقوق حریم خصوصی و امنیت شخصی منجر شود [۷].

۶.۲ مثال‌های عملی و نتایج موفق

در پایان این مرور ادبیات، اشاره به مثال‌های عملی و موفقیت‌های حاصل از به‌کارگیری هوش مصنوعی در امنیت سایبری می‌تواند به درک بهتر اهمیت این فناوری کمک کند [۶]. به عنوان مثال، سیستم‌های مبتنی بر AI مانند Darktrace و Splunk، که در تشخیص تهدیدات سایبری موفقیت‌های قابل توجهی کسب کرده‌اند، نمونه‌های موفق از کاربرد هوش مصنوعی در امنیت سایبری هستند [۹، ۱۲].

۳ کاربردهای هوش مصنوعی در امنیت سایبری

با گسترش روزافزون فناوری‌های دیجیتال و ارتباطات جهانی، پیچیدگی و شدت تهدیدات سایبری نیز به‌طور قابل‌ملاحظه‌ای افزایش یافته است [۳]. سازمان‌ها برای مقابله با این تهدیدات و حفظ امنیت اطلاعات و زیرساخت‌های حیاتی خود، به فناوری‌های پیشرفته‌ای همچون هوش مصنوعی (AI) روی آورده‌اند [۴]. با قابلیت‌های منحصر به فرد خود در تحلیل داده‌های بزرگ، یادگیری از الگوهای پیچیده و خودکارسازی فرآیندها، به یکی از ابزارهای اصلی در مقابله با تهدیدات سایبری تبدیل شده است [۵].

۱.۳ تشخیص تهدیدات و تحلیل آسیب‌پذیری‌ها

یکی از کاربردهای اصلی هوش مصنوعی در امنیت سایبری، تشخیص تهدیدات و تحلیل آسیب‌پذیری‌ها است [۳، ۴]. سیستم‌های سنتی تشخیص تهدیدات معمولاً بر اساس قواعد از پیش تعیین‌شده عمل می‌کردند که این امر آن‌ها را در برابر تهدیدات جدید و پیچیده آسیب‌پذیر می‌ساخت [۷]. در مقابل، سیستم‌های مبتنی بر هوش مصنوعی قادرند از الگوریتم‌های یادگیری ماشینی و یادگیری عمیق برای تحلیل داده‌های بزرگ و شناسایی الگوهای ناشناخته استفاده کنند [۵، ۱].

به‌عنوان مثال، سیستم‌های مبتنی بر هوش مصنوعی می‌توانند تهدیدات نوظهور، را با دقت و سرعت بیشتری شناسایی کنند و واکنش مناسب را در برابر آن‌ها انجام دهند. این قابلیت به سیستم‌های امنیت سایبری امکان می‌دهد که به‌طور مداوم از حملات جدید پیشگیری کرده و نقاط ضعف را شناسایی کنند.

۲.۳ پاسخگویی به حوادث و کاهش آسیب‌ها

روند پاسخگویی به حوادث سایبری معمولاً پیچیده و زمان‌بر است [۴]. تأخیر در پاسخ به یک حمله سایبری می‌تواند خسارات جبران‌ناپذیری به همراه داشته باشد [۵، ۲]. هوش مصنوعی با خودکارسازی فرآیندها و

تحلیل سریع داده‌ها، به‌طور قابل توجهی زمان پاسخگویی به حوادث را کاهش می‌دهد [۹]. این توانایی به ویژه در حملات پیچیده‌ای مانند DDoS یا حملات باج‌افزاری بسیار حیاتی است [۱۲].

۳.۳ مثال های کاربردی

برای درک بهتر نقش هوش مصنوعی در امنیت سایبری، بررسی برخی از مثال‌های عملی و کاربردی می‌تواند مفید باشد [۶، ۱۲]. این مثال‌ها توانایی‌های بالقوه هوش مصنوعی در تشخیص تهدیدات، پاسخ به حملات و خودکارسازی فرآیندهای امنیتی را نشان می‌دهند [۶].

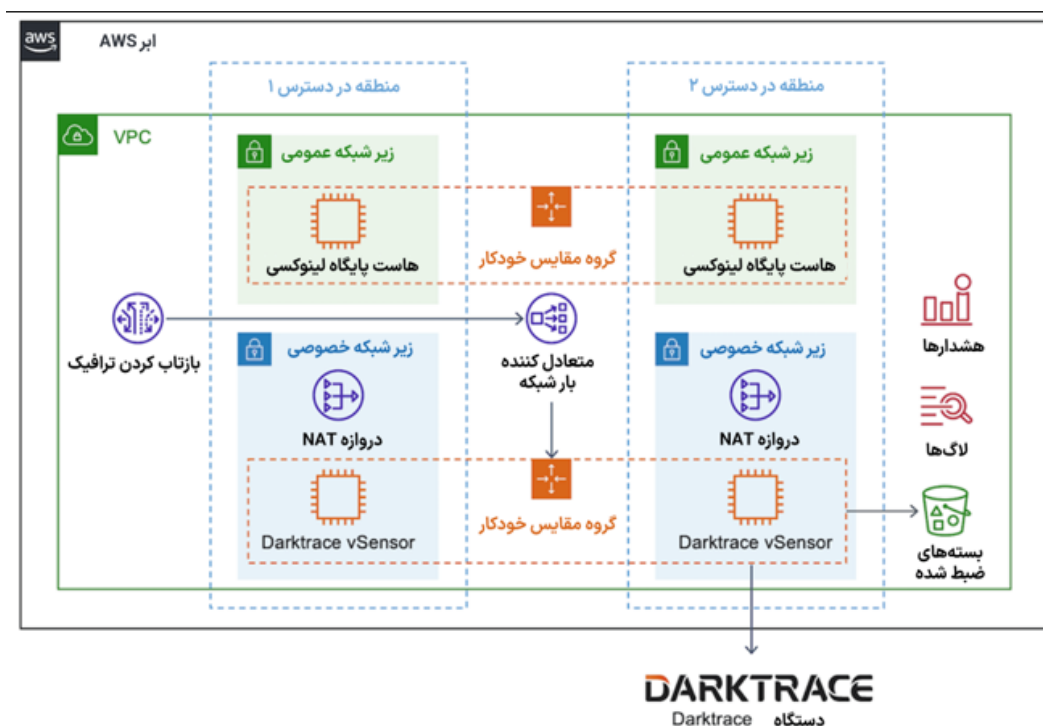
Darktrace. یکی از برجسته ترین مثال‌ها در زمینه استفاده از هوش مصنوعی برای امنیت سایبری، سیستم Darktrace است. این سیستم از الگوریتم‌های یادگیری ماشینی برای شناسایی تهدیدات سایبری و ناهنجاری‌های رفتاری در شبکه‌های سازمانی استفاده می‌کند. Darktrace قادر است به طور خودکار فعالیت‌های مشکوک را شناسایی کرده و در زمان واقعی هشدارهایی را به تیم‌های امنیتی ارسال کند.

Amazon Web Services. آمازون وب سرویس (AWS) از فناوری‌های مختلف امنیت سایبری از جمله Darktrace برای نظارت و تحلیل ترافیک شبکه استفاده می‌کند که دیاگرام ساختار معماری این شبکه در شکل ۱ نمایش داده شده است. در این معماری، مکانیزم‌های مختلفی برای مدیریت و امنیت داده‌ها پیاده‌سازی شده است. در این ساختار، یک ابر مجازی خصوصی (VPC (Virtual Private Cloud) پیاده‌سازی شده که به سازمان اجازه می‌دهد یک شبکه خصوصی درون (AWS) ایجاد کرده و منابع و سرورهای خود را در آن مستقر کند. VPC به عنوان یک بستر اساسی، شبکه‌های عمومی و خصوصی را درون خود جای داده است. معماری در دو منطقه در دسترس (Availability Zones) متفاوت توزیع شده است که به معنای بهره‌گیری از زیرساخت چندمنطقه‌ای برای افزایش دسترسی پذیری و افزونگی است. این کار به منظور اطمینان از عدم قطع سرویس در صورت خرابی در یک منطقه است و به نوعی پشتیبان یکدیگر محسوب می‌شوند. در هر منطقه، دو زیرشبکه متفاوت وجود دارد: زیرشبکه عمومی که میزبان پایگاه‌های داده لینوکسی است که به شبکه عمومی متصل هستند و ممکن است نیاز به ارتباطات مستقیم با اینترنت یا سایر منابع خارجی داشته باشند. زیرشبکه خصوصی که در این بخش، اجزای حساس‌تر و مهم‌تر از جمله دروازه‌های NAT و Darktrace vSensor قرار دارند که مسئول نظارت و بررسی ترافیک شبکه هستند. دروازه‌های NAT به منابع موجود در زیرشبکه خصوصی امکان می‌دهند که به اینترنت دسترسی داشته باشند بدون اینکه مستقیماً در دسترس اینترنت باشند. گروه مقایس خودکار از سیستم به طور خودکار با توجه به نیاز شبکه، منابع بیشتری (مانند سرورها) را اضافه یا حذف می‌کند تا توان شبکه برای پردازش ترافیک‌های ورودی و خروجی بهینه شود. متعادل‌کننده بار هم نقش توزیع ترافیک بین منابع مختلف (مانند هاست‌های پایگاه لینوکسی) را بر عهده دارد تا منابع به صورت متعادل استفاده شوند و هیچ بخشی از سیستم تحت فشار زیاد قرار نگیرد.

در زیرشبکه‌های خصوصی، Darktrace vSensor نقش حیاتی در نظارت و تجزیه و تحلیل ترافیک شبکه

دارد. این ابزار پیشرفته ترافیک شبکه را بررسی می‌کند و به دنبال تهدیدات و حملات سایبری در زمان واقعی است. این بخش همچنین داده‌های بسته‌های ضبط‌شده را تحلیل کرده و اطلاعات لازم را به دستگاه‌های Darktrace ارسال می‌کند. Darktrace با تحلیل ترافیک شبکه و تشخیص هرگونه رفتار غیرعادی یا تهدید بالقوه، هشدارهای لازم را تولید کرده و اطلاعات مربوطه را لاگ می‌کند. بسته‌های ثبت شده نشان می‌دهد که بسته‌های مشکوک یا مهم به‌طور خاص ثبت شده و برای تحلیل‌های بیشتر ذخیره می‌شوند. بازتاب کردن ترافیک یا Traffic Mirroring قابلیتی است برای ارسال یک کپی از ترافیک شبکه به Darktrace vSensor که جهت تحلیل‌های بیشتر است. بازتاب کردن ترافیک اجازه می‌دهد که بدون تأثیرگذاری بر عملکرد اصلی سیستم، تحلیل و بررسی دقیقی از جریان داده‌ها در سطح شبکه صورت گیرد.

این ساختار یک معماری چند لایه، مقاوم و تطبیق‌پذیر برای حفاظت از داده‌ها و نظارت بر شبکه در فضای ابری (AWS) ارائه می‌دهد. استفاده از زیرشبکه‌های عمومی و خصوصی، گروه‌های مقایس خودکار، و متعادل‌کننده بار، انعطاف‌پذیری و پایداری سیستم را افزایش می‌دهد. از سوی دیگر، پیاده‌سازی Darktrace برای نظارت هوشمندانه بر ترافیک شبکه، امکان تشخیص سریع و دقیق تهدیدات امنیتی را فراهم می‌کند و به این ترتیب یک لایه امنیتی مهم برای سازمان‌ها ایجاد می‌کند.



شکل ۱: معماری Darktrace vSensor در وب سرویس آمازون

SPLUNK. Splunk یک پلتفرم تحلیل داده است که به سازمان‌ها کمک می‌کند تا داده‌های تولیدشده

در زیرساخت‌های فناوری اطلاعات خود را جمع‌آوری و تجزیه و تحلیل کنند. Splunk با استفاده از هوش مصنوعی و یادگیری ماشینی، قادر به شناسایی الگوهای مشکوک و تشخیص تهدیدات سایبری در زمان واقعی است.

CYLANCE. Cylance یک شرکت امنیت سایبری است که از هوش مصنوعی برای پیشگیری از تهدیدات استفاده می‌کند. محصولات Cylance با بهره‌گیری از الگوریتم‌های یادگیری ماشینی، تهدیدات سایبری را پیش از وقوع شناسایی و خنثی می‌کنند. این تکنیک به سازمان‌ها کمک می‌کند تا از حملات پیچیده و ناشناخته جلوگیری کنند.

DEEP INSTINCT. Deep Instinct یکی دیگر از سیستم‌های پیشرفته است که از یادگیری عمیق (Deep Learning) برای شناسایی و جلوگیری از تهدیدات سایبری استفاده می‌کند. این سیستم قادر است تهدیدات ناشناخته را شناسایی کرده و به طور خودکار اقدام به مقابله با آن‌ها کند. Deep Instinct با استفاده از الگوریتم‌های پیشرفته، توانسته است توانایی‌های بی‌نظیری در محافظت از زیرساخت‌های دیجیتال ایجاد کند.

جدول ۲: انواع استراتژی‌های مقابله با حملات سایبری

استراتژی	توضیحات	مثال
تشخیص ناهنجاری‌ها	شناسایی رفتارهای غیرعادی در سیستم‌ها	Darktrace
تحلیل پیش‌بینی کننده	پیش‌بینی حملات با استفاده از داده‌های گذشته	Splunk
پاسخ خودکار	اقدامات خودکار برای خنثی‌سازی حملات	سیستم‌های SIEM

۴ چالش‌ها و ریسک‌های پیش رو

هوش مصنوعی به عنوان یک فناوری انقلابی در حوزه امنیت سایبری، نه تنها فرصت‌های بزرگی را فراهم می‌آورد، بلکه چالش‌ها و ریسک‌های قابل توجهی را نیز به همراه دارد [۴]. این چالش‌ها از پیچیدگی‌های فنی و عملیاتی گرفته تا ملاحظات اخلاقی و قانونی را دربرمی‌گیرد و به مدیریت دقیق و همه‌جانبه نیازمند است. در ادامه به بررسی چالش‌های کلیدی و ریسک‌های مرتبط با استفاده از هوش مصنوعی در امنیت سایبری پرداخته می‌شود.

۱.۴ پیچیدگی‌های فنی و عملیاتی

یکی از بزرگترین چالش‌های استفاده از هوش مصنوعی در امنیت سایبری، پیچیدگی فنی و عملیاتی آن است [۴، ۳]. توسعه و پیاده‌سازی سیستم‌های AI نیازمند دانش تخصصی و منابع قابل توجهی است [۷]. بدون داده‌های مناسب، این سیستم‌ها دچار خطاهای جدی می‌شوند و کارایی لازم را نخواهند داشت [۵]. بدون

جدول ۳: چالش‌های اصلی اعتماد به سیستم‌های هوش مصنوعی و پیامدهای آن

پیامدها	توضیحات	چالش
اعتماد کاربران	عدم امکان توضیح تصمیمات توسط هوش مصنوعی کاهش	عدم شفافیت الگوریتم
ریسک خطاهای ناشی از داده‌های نادرست	نیاز به داده‌های بزرگ و متنوع	وابستگی به داده‌ها
محدودیت در ارزیابی خطرات	دشواری در درک فرآیند تصمیم‌گیری	پیچیدگی تصمیم‌گیری

داده‌های مناسب، این سیستم‌ها دچار خطاهای جدی می‌شوند و کارایی لازم را نخواهند داشت. به همین دلیل، دستیابی به زیرساخت‌های قوی و تأمین منابع مناسب، یکی از چالش‌های اساسی در این حوزه است.

۲.۴ مسئله اعتماد و شفافیت

یکی دیگر از چالش‌های مهم در استفاده از هوش مصنوعی در امنیت سایبری، اعتماد به تصمیمات AI است. بسیاری از الگوریتم‌های هوش مصنوعی به عنوان «جعبه‌های سیاه» شناخته می‌شوند [۴]. این عدم شفافیت می‌تواند باعث کاهش اعتماد به سیستم‌های امنیتی مبتنی بر AI شود [۷].

۳.۴ تهدیدات ناشی از هوش مصنوعی خصمانه

علاوه بر فرصت‌ها، تهدیدات ناشی از هوش مصنوعی خصمانه یکی دیگر از چالش‌های مهم در امنیت سایبری است. مهاجمان سایبری می‌توانند از AI برای تقویت حملات خود استفاده کنند [۹]. آن‌ها می‌توانند الگوریتم‌هایی بسازند که به‌طور خودکار آسیب‌پذیری‌ها را شناسایی کرده و به شیوه‌های غیرقابل پیش‌بینی به سیستم‌ها نفوذ کنند [۱۳]. این نوع حملات می‌توانند بسیار پیچیده‌تر از حملات سنتی باشند و سیستم‌های دفاعی معمولی قادر به شناسایی آن‌ها نباشند. برای مقابله با این تهدیدات، نیاز به دفاع‌های تطبیقی پیشرفته است که توانایی مقابله با حملات مبتنی بر هوش مصنوعی را داشته باشند و به‌صورت پویا در برابر تهدیدات جدید عمل کنند.

۴.۴ ملاحظات اخلاقی و قانونی

ملاحظات اخلاقی و قانونی پیرامون استفاده از هوش مصنوعی در امنیت سایبری نیز از اهمیت ویژه‌ای برخوردارند [۴، ۷]. یکی از دغدغه‌های اصلی این است که استفاده گسترده از AI ممکن است به نقض حریم خصوصی و حقوق فردی منجر شود [۱۴]. همچنین، مسئولیت‌پذیری در قبال تصمیمات خودکار گرفته شده توسط این سیستم‌ها یک موضوع مهم قانونی است. باید قوانین و مقرراتی تدوین شود که اطمینان حاصل کنند استفاده از هوش مصنوعی به شیوه‌ای اخلاقی و مسئولانه صورت می‌گیرد.

جدول ۴: ملاحظات کلیدی در زمینه اخلاق و قانون مرتبط با هوش مصنوعی امنیت سایبری

موضوع	توضیحات	اقدامات پیشنهادی
حریم خصوصی	احتمال نقض حقوق حریم خصوصی توسط هوش مصنوعی	توسعه قوانین محافظتی سختگیرانه تر
مسئولیت پذیری	تعیین مسئولیت در صورت وقوع خطا توسط هوش مصنوعی	تدوین چارچوب های قانونی روشن
شفافیت	نیاز به توضیح تصمیمات هوش مصنوعی	افزایش شفافیت از طریق توسعه الگوریتم های قابل توضیح

جدول ۵: توضیح پذیری و ارزیابی مدل ها مرتبط با هوش مصنوعی امنیت سایبری

استراتژی	توضیحات	مزایا
توسعه الگوریتم های قابل توضیح	ایجاد الگوریتم هایی که فرآیند تصمیم گیری را توضیح دهند	افزایش اعتماد کاربران
ارزیابی منظم مدل ها	بررسی دوره ای مدل های هوش مصنوعی برای اطمینان از عملکرد صحیح	کاهش ریسک خطا

۵.۴ آینده نگری و مدیریت ریسک

با توجه به پیشرفت های سریع در حوزه هوش مصنوعی، لازم است که رویکردهای آینده نگرانه ای برای مدیریت ریسک ها و چالش های مرتبط با این فناوری اتخاذ شود [۱۳]. این رویکردها باید شامل ارزیابی مداوم تهدیدات نوظهور، توسعه روش های جدید برای مقابله با این تهدیدات و همکاری های بین المللی در تدوین قوانین و استانداردهای مرتبط باشد [۱۴]. هوش مصنوعی به سرعت در حال پیشرفت است و سیستم های امنیت سایبری باید همگام با این تحولات پیشرفت کنند [۴].

۵ راهکارها و استراتژی های پیشنهادی

با توجه به چالش ها و ریسک هایی که در بخش قبلی بررسی شدند، ضروری است که سازمان ها و پژوهشگران به دنبال توسعه راهکارها و استراتژی های مؤثر برای بهینه سازی استفاده از هوش مصنوعی در امنیت سایبری باشند. این بخش به بررسی راهکارهای عملی و استراتژی های پیشنهادی برای مقابله با تهدیدات سایبری و بهبود امنیت دیجیتال با استفاده از AI می پردازد.

۱.۵ بهبود شفافیت و اعتماد

یکی از راهکارهای کلیدی برای مقابله با چالش های اعتماد به AI، افزایش شفافیت در فرآیندهای تصمیم گیری این سیستم ها است [۷]. توسعه الگوریتم های قابل توضیح (Explainable AI) می تواند به درک بهتر فرآیندهای تصمیم گیری هوش مصنوعی و افزایش اعتماد کاربران کمک کند [۶].

جدول ۶: راهکارهای پیشنهادی برای بهبود شفافیت و اعتماد در سیستم‌های هوش مصنوعی

مزایا	توضیحات	استراتژی
افزایش اعتماد کاربران	ایجاد الگوریتم‌هایی که فرآیند تصمیم‌گیری را توضیح دهند	توسعه الگوریتم‌های قابل توضیح
کاهش ریسک خطا	بررسی دوره‌ای مدل‌های هوش مصنوعی برای اطمینان از عملکرد صحیح	ارزیابی منظم مدل‌ها
هماهنگی در سطح جهانی	همکاری بین‌المللی برای ایجاد استانداردهای شفافیت	مشارکت بین‌المللی در توسعه استانداردها

۲.۵ تقویت دفاع سایبری

برای مقابله با تهدیدات سایبری پیشرفته، سازمان‌ها باید به تقویت دفاع سایبری خود با استفاده از هوش مصنوعی بپردازند [۱۲]. این شامل توسعه سیستم‌های پیش‌بینی و تشخیص تهدیدات، خودکارسازی پاسخ به حوادث و ایجاد مکانیزم‌های دفاعی تطبیقی می‌شود [۱۱]. سیستم‌های AI تطبیقی قادرند با توجه به تغییرات محیط دیجیتال و نوع تهدیدات، خود را به‌روزرسانی و تنظیم کنند [۴].

۳.۵ بهبود زیرساخت‌های داده‌ای

داده‌های با کیفیت و متنوع یکی از ارکان اصلی برای عملکرد موفقیت‌آمیز سیستم‌های هوش مصنوعی در امنیت سایبری است [۲]. سازمان‌ها باید زیرساخت‌های داده‌ای خود را بهبود دهند تا بتوانند از داده‌های بزرگ و پیچیده برای آموزش و بهینه‌سازی مدل‌های AI استفاده کنند [۵]. این داده‌ها باید به‌درستی مدیریت شده و به‌طور مداوم به‌روز شوند [۹].

۴.۵ رویکردهای قانونی و اخلاقی

برای بهره‌برداری مسئولانه از AI در امنیت سایبری، بهبود چارچوب‌های قانونی و اخلاقی ضروری است [۷]. این شامل تدوین قوانین دقیق برای حفاظت از حریم خصوصی، شفافیت در استفاده از داده‌ها و مسئولیت‌پذیری در قبال تصمیمات گرفته شده توسط سیستم‌های هوش مصنوعی می‌شود [۱۳، ۱۴]. این قوانین باید به‌گونه‌ای تدوین شوند که از حقوق کاربران محافظت کنند و در عین حال، نوآوری‌های فناوری را محدود نکنند [۱۲].

۵.۵ آینده‌نگری و تحقیق

پژوهش‌های آینده باید بر روی توسعه فناوری‌ها و رویکردهای جدید در امنیت سایبری با استفاده از هوش مصنوعی تمرکز کنند [۱۱]. این شامل تحقیق در زمینه‌های نوظهور مانند یادگیری ماشینی خصمانه، سیستم‌های هوش مصنوعی تطبیقی و تحلیل داده‌های پیچیده است [۱۰]. آینده امنیت سایبری به توانایی ما در پیش‌بینی تهدیدات و مقابله با آن‌ها بستگی دارد و هوش مصنوعی در این مسیر نقش کلیدی ایفا می‌کند.

۶ نتیجه گیری

در نهایت، بررسی جامع ادبیات و کاربردهای هوش مصنوعی در امنیت سایبری نشان می‌دهد که این فناوری نه تنها می‌تواند به بهبود امنیت دیجیتال کمک کند، بلکه پتانسیل زیادی برای مقابله با تهدیدات پیچیده سایبری دارد [۳، ۶]. AI با قابلیت‌های تحلیلی پیشرفته خود، توانایی پردازش و تفسیر حجم عظیمی از داده‌ها را در زمان واقعی فراهم می‌کند [۵]. این ویژگی به سازمان‌ها این امکان را می‌دهد که نه تنها تهدیدات شناخته‌شده را شناسایی کنند، بلکه در مواجهه با تهدیدات نوظهور و ناشناخته نیز واکنش سریع و مؤثری داشته باشند [۴، ۷].

هوش مصنوعی قادر است الگوهای پیچیده رفتاری را در داده‌های بزرگ شناسایی کند و به این ترتیب، تشخیص تهدیدات سایبری را به‌طور چشمگیری بهبود بخشد [۶]. با این حال، بهره‌برداری از این فناوری با چالش‌های جدی نیز همراه است. یکی از مهم‌ترین چالش‌ها، پیچیدگی فنی و عملیاتی هوش مصنوعی است که به منابع گسترده از نظر داده‌ها، فناوری و نیروی انسانی نیاز دارد [۷، ۱۰]. علاوه بر آن، عدم شفافیت در تصمیم‌گیری‌های مبتنی بر AI که به عنوان «جعبه سیاه» شناخته می‌شود، نگرانی‌های زیادی در مورد اعتماد به این سیستم‌ها به وجود آورده است [۴].

همچنین، تهدیدات ناشی از هوش مصنوعی خصمانه یکی از نگرانی‌های اصلی است. مهاجمان سایبری می‌توانند از AI برای تقویت حملات خود و ایجاد تهدیدات جدیدی که توسط سیستم‌های دفاعی سنتی قابل شناسایی نیستند، استفاده کنند [۹، ۱۱]. این حملات می‌توانند به شکل خودکار و بدون دخالت مستقیم انسان انجام شوند و چالش‌های جدی برای سیستم‌های امنیتی ایجاد کنند [۱۲].

ملاحظات اخلاقی و قانونی نیز یکی دیگر از جنبه‌های مهم استفاده از AI در امنیت سایبری است [۱۴]. از آنجایی که بسیاری از الگوریتم‌های هوش مصنوعی به داده‌های حساس دسترسی دارند، تدوین قوانین و مقررات مناسبی برای حفاظت از حریم خصوصی و حقوق افراد ضروری است [۱۳]. همچنین، باید تعیین شود که در صورت وقوع خطاهای ناشی از AI، مسئولیت بر عهده چه کسی خواهد بود و چگونه می‌توان پاسخگو بود [۱۲، ۱۴].

برای بهره‌برداری کامل و بهینه از هوش مصنوعی در امنیت سایبری، سازمان‌ها و پژوهشگران باید به دنبال توسعه راهکارهای عملی و استراتژی‌های مؤثر باشند [۱۱، ۱۰]. یکی از مهم‌ترین اقدامات، افزایش شفافیت و اعتماد در سیستم‌های AI است [۷]. توسعه الگوریتم‌های قابل توضیح (Explainable AI) می‌تواند به کاربران کمک کند تا فرآیندهای تصمیم‌گیری هوش مصنوعی را بهتر درک کرده و اعتماد بیشتری به این سیستم‌ها داشته باشند [۵].

علاوه بر این، خودکارسازی دفاع سایبری می‌تواند به سازمان‌ها کمک کند تا با سرعت و دقت بیشتری به حملات سایبری پاسخ دهند و از آسیب‌های ناشی از تأخیر در واکنش جلوگیری کنند [۹]. بهبود زیرساخت‌های داده‌ای نیز اهمیت زیادی دارد، زیرا سیستم‌های AI برای عملکرد مؤثر خود به داده‌های بزرگ و با کیفیت نیاز دارند [۱۳].

همچنین، باید به تدوین چارچوب‌های قانونی و اخلاقی مناسب برای استفاده از AI در امنیت سایبری

توجه کرد تا از هرگونه سوءاستفاده یا نقض حقوق افراد جلوگیری شود [۱۴]. پژوهش‌های آینده باید به سمت توسعه رویکردهای جدید برای مقابله با تهدیدات سایبری پیش روند [۱۱]. یکی از این رویکردها می‌تواند ترکیب AI با فناوری‌های نوظهور دیگر مانند بلاک‌چین و اینترنت اشیا باشد تا یک لایه امنیتی قوی‌تر و جامع‌تر ایجاد شود [۹].

در نهایت، هوش مصنوعی می‌تواند به عنوان یک ابزار قدرتمند و مؤثر در مبارزه با تهدیدات سایبری عمل کند، به شرط آنکه با دقت و در نظر گرفتن تمامی جنبه‌های فنی، اخلاقی و قانونی مورد استفاده قرار گیرد [۱۰]. تنها با چنین رویکردی است که می‌توان از پتانسیل کامل این فناوری بهره‌برداری کرد و به ایجاد یک فضای دیجیتال امن‌تر و قابل اعتمادتر کمک نمود [۶، ۳]. آینده امنیت سایبری در گرو توانایی ما در بهره‌برداری هوشمندانه از AI است و در این راستا، همکاری‌های بین‌المللی، تحقیق و توسعه و تعهد به مسئولیت‌پذیری از اهمیت بسیار بالایی برخوردارند [۴].

مراجع

- [1] Bhardwaj, M.D., Alshehri, K., Kaushik, H.J., Alyamani, M., Kumar, "Secure framework against cyber-attacks on cyber-physical robotic systems", Journal of Electronic Imaging, 31(6), 2022, pp. 061802-061802.
- [2] Chithaluru, P., Fadi, A.T., Kumar, M., Stephan, T., "Computational intelligence inspired adaptive opportunistic clustering approach for industrial IoT networks", IEEE Internet of Things Journal, 2023. DOI: 10.1109/JIOT.2022.3231605.
- [3] Wiafe, I., Koranteng, F.N., Obeng, E.N., Assyne, N., Wiafe, A., Gulliver, S.R., "Artificial intelligence for cybersecurity: a systematic mapping of literature", IEEE Access, 8, 2020, pp. 146598-146612.
- [4] Zhang, Z., Ning, H., Shi, F., Farha, F., Xu, Y., Xu, J., Zhang, F., Choo, K.K.R., "Artificial intelligence in cyber security: research advances, challenges, and opportunities", Artificial Intelligence Review, 55, 2022, pp. 1029-1053.
- [5] Martínez Torres, J., Iglesias Comesaña, C., García-Nieto, P.J., "Machine learning techniques applied to cybersecurity", International Journal of Machine Learning and Cybernetics, 10(10), 2019, pp. 2823-2836.
- [6] Truong, T.C., Zelinka, I., Plucar, J., Candík, M., Sulc, V., "Artificial intelligence and cybersecurity: past, presence, and future", in Artificial intelligence and evolutionary computations in engineering systems, 2020, pp. 351-363.
- [7] Promyslov, V.G., Semenov, K.V., Shumov, A.S., "A clustering method of asset cybersecurity classification", IFAC-PapersOnLine, 52(13), 2019, pp. 928-933.
- [8] Millar, K., Cheng, A., Chew, H.G., Lim, C.C., "Operating system classification: a minimalist approach", in 2020 International Conference on Machine Learning and Cybernetics (ICMLC), 2020, pp. 143-150.
- [9] Aksoy, A., Gunes, M.H., "Automated IoT device identification using network traffic", in IEEE International Conference on Communications (ICC), 2019, pp. 1-7.

- [10] Sivanathan, A., Gharakheili, H.H., Loi, F., Radford, A., Wijenayake, C., Vishwanath, A., Sivaraman, V., "Classifying IoT devices in smart environments using network traffic characteristics", IEEE Transactions on Mobile Computing, 18(8), 2018, pp. 1745–1759.
- [11] Cvitić, I., Peraković, D., Perić, M., Gupta, B., "Ensemble machine learning approach for classification of IoT devices in smart home", International Journal of Machine Learning and Cybernetics, 12(11), 2021, pp. 3179–3202.
- [12] Cam, H., "Online detection and control of malware infected assets", in IEEE Military Communications Conference (MILCOM), 2017, pp. 701–706.
- [13] Kure, H.I., Islam, S., Ghazanfar, M., Raza, A., Pasha, M., "Asset criticality and risk prediction for an effective cybersecurity risk management of cyber-physical system", Neural Computing and Applications, 34(1), 2022, pp. 493–514.
- [14] Vega-Barbas, M., Villagrà, V.A., Monje, F., Riesco, R., Larriva-Novo, X., Berrocal, J., "Ontology-based system for dynamic risk management in administrative domains", Applied Sciences, 9(21), 2019, p. 4547.

