

کاربرد روش‌های هوش مصنوعی در تحلیل رفتار کاربران و شناسایی تهدیدات سایبری در محیط‌های دیجیتال: رویکردی نوین در ارتقای امنیت سایبری

ندا مهدی‌زاده

دانشجو کاردانی مهندسی فناوری اطلاعات، ایران - تهران - دانشگاه جامع علمی کاربردی، مرکز جوادالائمه (ع)

چکیده

گسترش فناوری‌های دیجیتال و افزایش وابستگی سازمان‌ها به زیرساخت‌های اطلاعاتی، موجب افزایش پیچیدگی تهدیدات سایبری و ضرورت استفاده از راهکارهای هوشمند برای حفاظت از محیط‌های دیجیتال شده است. در این میان، هوش مصنوعی به عنوان یکی از فناوری‌های نوین، قابلیت‌های گسترده‌ای در تحلیل داده‌ها، شناسایی الگوهای رفتاری و تشخیص تهدیدات سایبری ارائه می‌دهد. هدف این پژوهش، بررسی کاربرد روش‌های هوش مصنوعی در تحلیل رفتار کاربران و شناسایی تهدیدات سایبری در محیط‌های دیجیتال و تبیین نقش این فناوری در ارتقای امنیت سایبری است. پژوهش حاضر از نظر هدف، کاربردی و از نظر روش اجرا، مروری-توصیفی است. داده‌های مورد نیاز از طریق بررسی مقالات علمی، کتاب‌ها و گزارش‌های تخصصی معتبر در حوزه هوش مصنوعی، امنیت سایبری، یادگیری ماشین و تحلیل رفتار کاربران گردآوری و تحلیل شده است. یافته‌های پژوهش نشان می‌دهد که الگوریتم‌های یادگیری ماشین، یادگیری عمیق و روش‌های تحلیل رفتار کاربران، توانایی قابل توجهی در شناسایی ناهنجاری‌ها، تشخیص حملات ناشناخته، کشف بدافزارها، شناسایی تهدیدات داخلی و پیش‌بینی رخدادهای امنیتی دارند. همچنین، ترکیب فناوری‌هایی مانند تحلیل رفتار کاربران و موجودیت‌ها، سامانه‌های تشخیص نفوذ، هوش تهدید و معماری Zero Trust با مدل‌های هوش مصنوعی، موجب افزایش دقت شناسایی تهدیدات، کاهش هشدارهای نادرست و بهبود سرعت واکنش امنیتی می‌شود. با وجود مزایای فراوان، چالش‌هایی مانند کیفیت داده‌های آموزشی، حملات خصمانه علیه مدل‌های هوشمند، تفسیرپذیری الگوریتم‌ها و حفظ حریم خصوصی همچنان از موانع توسعه این فناوری محسوب می‌شوند. نتایج این پژوهش نشان می‌دهد که استفاده هدفمند از هوش مصنوعی، در کنار راهبردهای امنیتی مناسب، می‌تواند نقش مؤثری در ایجاد سامانه‌های امنیت سایبری هوشمند، پیشگیرانه و مقاوم در برابر تهدیدات نوظهور ایفا کند.

واژگان کلیدی: هوش مصنوعی، امنیت سایبری، تحلیل رفتار کاربران، یادگیری ماشین، یادگیری عمیق، تشخیص ناهنجاری، UEBA، تهدیدات سایبری

۱- مقدمه و بیان مسئله

در عصر حاضر، تحول دیجیتال و توسعه شتابان فناوری‌های اطلاعات و ارتباطات، ساختار فعالیت‌های فردی، سازمانی، اقتصادی و حتی اجتماعی را به‌طور اساسی دگرگون کرده است. امروزه بخش قابل توجهی از تعاملات کاربران، ارائه خدمات، تبادل اطلاعات، انجام تراکنش‌های مالی، مدیریت کسب‌وکارها و فرآیندهای تصمیم‌گیری در بستر محیط‌های دیجیتال انجام می‌شود. گسترش استفاده از اینترنت، سامانه‌های تحت وب، شبکه‌های اجتماعی، رایانش ابری، اینترنت اشیا،

تجهیزات هوشمند و خدمات الکترونیکی موجب تولید حجم عظیمی از داده‌های دیجیتال شده است (Khan et al., 2022). که هر روز بر میزان آن افزوده می‌شود. این حجم گسترده از داده‌ها، علاوه بر ایجاد فرصت‌های ارزشمند برای توسعه خدمات هوشمند و شخصی‌سازی شده، زمینه مناسبی را برای تحلیل دقیق رفتار کاربران و استخراج الگوهای رفتاری فراهم کرده است.

داده‌های رفتاری کاربران، شامل اطلاعات مربوط به نحوه ورود به سامانه‌ها، زمان و مکان دسترسی، الگوهای کلیک، نوع فعالیت، میزان استفاده از خدمات، الگوهای جستجو، تراکنش‌های انجام شده و سایر تعاملات دیجیتال، منبع ارزشمندی برای شناخت رفتارهای عادی و غیرعادی محسوب می‌شوند. تحلیل صحیح این داده‌ها می‌تواند به شناسایی تغییرات رفتاری، پیش‌بینی رخدادهای مشکوک و کشف فعالیت‌هایی که احتمال وقوع حملات سایبری را افزایش می‌دهند، کمک کند. در سال‌های اخیر، سازمان‌ها و مراکز ارائه‌دهنده خدمات دیجیتال بیش از گذشته به اهمیت بهره‌گیری از این اطلاعات در تقویت امنیت سایبری پی برده‌اند.

از سوی دیگر، افزایش وابستگی جوامع به فناوری‌های دیجیتال، هم‌زمان موجب گسترش سطح حملات سایبری و پیچیده‌تر شدن روش‌های نفوذ شده است. مهاجمان سایبری با استفاده از تکنیک‌های پیشرفته‌ای مانند بدافزارهای هوشمند، حملات فیشینگ، باج‌افزارها، حملات انکار سرویس، سرقت هویت، نفوذ به حساب‌های کاربری و سوءاستفاده از آسیب‌پذیری‌های نرم‌افزاری، تهدیدات گسترده‌ای را علیه امنیت اطلاعات و حریم خصوصی کاربران ایجاد کرده‌اند. این حملات علاوه بر خسارات مالی، موجب کاهش اعتماد کاربران، اختلال در ارائه خدمات، افشای اطلاعات محرمانه و آسیب به اعتبار سازمان‌ها می‌شوند.

در چنین شرایطی، روش‌های سنتی امنیت سایبری که عمدتاً مبتنی بر قوانین از پیش تعریف‌شده و امضاهای شناخته‌شده هستند، در بسیاری از موارد توانایی لازم برای شناسایی تهدیدات جدید، حملات ناشناخته و رفتارهای پیچیده مهاجمان را ندارند. تغییر مداوم الگوهای حمله، افزایش حجم داده‌ها و سرعت بالای وقوع رخدادهای امنیتی، نیاز به استفاده از فناوری‌های هوشمند و روش‌های تحلیل پیشرفته را بیش از پیش آشکار ساخته است.

در این میان، هوش مصنوعی، به‌ویژه شاخه‌هایی مانند یادگیری ماشین، یادگیری عمیق و تشخیص ناهنجاری، به‌عنوان یکی از مهم‌ترین فناوری‌های نوین در حوزه امنیت سایبری مطرح شده است. این فناوری‌ها با تحلیل حجم گسترده‌ای از داده‌های رفتاری، شناسایی الگوهای پنهان، یادگیری رفتار عادی کاربران و تشخیص انحراف از این الگوها، امکان شناسایی سریع‌تر و دقیق‌تر تهدیدات سایبری را فراهم می‌کنند. (شان و همکاران^۱، ۲۰۱۸؛ گودفالو و همکاران^۲، ۲۰۱۶)

بهره‌گیری از سامانه‌های مبتنی بر هوش مصنوعی می‌تواند زمان تشخیص حملات را کاهش داده، نرخ هشدارهای اشتباه را کم کرده و توان سازمان‌ها را در مقابله با تهدیدات نوظهور افزایش دهد. (مؤسسه ملی استانداردها و فناوری آمریکا^۳، ۲۰۲۴؛ گزارش امنیت شرکت آی‌بی‌ام^۴، ۲۰۲۳)

با توجه به اهمیت روزافزون امنیت اطلاعات و افزایش پیچیدگی حملات سایبری، بررسی نقش روش‌های هوش مصنوعی در تحلیل رفتار کاربران و شناسایی تهدیدات سایبری، به یکی از موضوعات مهم و کاربردی پژوهش‌های حوزه امنیت سایبری تبدیل شده است. بوچزاک و گوون^۵ (۲۰۱۶) بیان می‌کنند که روش‌های مبتنی بر یادگیری ماشین می‌توانند با

¹ Shone et al

² Goodfellow et al

³ NIST

⁴ IBM Security

⁵ Buczak & Guven

تحلیل حجم بالای داده‌های امنیتی، دقت و سرعت شناسایی تهدیدات را نسبت به روش‌های سنتی افزایش دهند. همچنین، Apruzzese و همکاران (۲۰۲۳) تأکید کرده‌اند که هوش مصنوعی با بهره‌گیری از الگوریتم‌های پیشرفته، نقش مؤثری در شناسایی رفتارهای غیرعادی کاربران، کشف حملات ناشناخته و ارتقای سامانه‌های دفاع سایبری ایفا می‌کند. از این رو، پژوهش حاضر با تمرکز بر کاربرد روش‌های هوش مصنوعی در تحلیل رفتار کاربران، تلاش می‌کند ظرفیت این فناوری‌ها را در ارتقای امنیت محیط‌های دیجیتال، بهبود فرآیند تشخیص تهدیدات و افزایش کارایی سامانه‌های امنیتی مورد بررسی قرار دهد. انتظار می‌رود نتایج این پژوهش بتواند زمینه‌ساز طراحی سامانه‌های امنیتی هوشمند، پیشگیرانه و سازگار با تهدیدات نوظهور باشد و به سازمان‌ها در کاهش مخاطرات امنیتی و افزایش تاب‌آوری سایبری کمک کند.

۲- مبانی نظری و پیشینه پژوهش

۲-۱- هوش مصنوعی (Artificial Intelligence)

هوش مصنوعی به مجموعه‌ای از روش‌ها، الگوریتم‌ها و فناوری‌هایی اطلاق می‌شود که هدف آن‌ها ایجاد سیستم‌های رایانه‌ای قادر به انجام وظایفی است که معمولاً به هوش انسانی نیاز دارند؛ از جمله یادگیری، استدلال، تصمیم‌گیری، تشخیص الگوها و تحلیل داده‌ها. در حوزه امنیت سایبری، هوش مصنوعی با بهره‌گیری از قابلیت‌های پردازشی خود می‌تواند حجم عظیمی از داده‌های امنیتی را تحلیل کرده، الگوهای پنهان را شناسایی و در تشخیص و مقابله با تهدیدات نوظهور نقش مؤثری ایفا کند. (سارکر^۶، ۲۰۲۲؛ راسل و نورویگ^۷، ۲۰۲۱)

۲-۲- امنیت سایبری (Cybersecurity)

امنیت سایبری مجموعه‌ای از فناوری‌ها، فرآیندها و اقدامات مدیریتی است که با هدف حفاظت از سامانه‌های رایانه‌ای، شبکه‌ها، نرم‌افزارها و داده‌ها در برابر دسترسی‌های غیرمجاز، حملات مخرب و تهدیدات دیجیتال به کار گرفته می‌شود. با افزایش پیچیدگی حملات سایبری، رویکردهای سنتی امنیتی به تنهایی پاسخگوی نیازهای سازمان‌ها نیستند و استفاده از فناوری‌های هوشمند برای پیش‌بینی و شناسایی تهدیدات اهمیت بیشتری یافته است. (استالینگز^۸، ۲۰۱۹)

۲-۳- تحلیل رفتار کاربران (User Behavior Analytics - UBA)

تحلیل رفتار کاربران فرآیندی است که طی آن فعالیت‌ها، الگوهای استفاده و تعاملات کاربران با سیستم‌های دیجیتال بررسی می‌شود تا رفتارهای معمول از رفتارهای غیرعادی تفکیک شوند. این رویکرد با استفاده از تحلیل داده‌ها و الگوریتم‌های یادگیری ماشین، امکان شناسایی فعالیت‌های مشکوک مانند سوءاستفاده از حساب‌های کاربری یا دسترسی‌های غیرمجاز را فراهم می‌کند. (سامر و پاکسون^۹، ۲۰۱۰؛ وانگ و همکاران^{۱۰}، ۲۰۲۴؛ المنصوری و همکاران^{۱۱}، ۲۰۲۲)

۲-۴- یادگیری ماشین (Machine Learning)

یادگیری ماشین یکی از زیرشاخه‌های هوش مصنوعی است که به سیستم‌های رایانه‌ای امکان می‌دهد بدون برنامه‌نویسی مستقیم و صرفاً با استفاده از داده‌ها، الگوها را یاد گرفته و عملکرد خود را بهبود دهند. الگوریتم‌های یادگیری

⁶ Sarker

⁷ Russell & Norvig

⁸ Stallings

⁹ Sommer & Paxson

¹⁰ Wang et al

¹¹ Mansoori et al

ماشین در امنیت سایبری برای طبقه‌بندی حملات، تشخیص نفوذ، پیش‌بینی تهدیدات و شناسایی رفتارهای غیرعادی مورد استفاده قرار می‌گیرند. (میچل^{۱۲}، ۱۹۹۷)

۵-۲- یادگیری عمیق (Deep Learning)

یادگیری عمیق شاخه‌ای پیشرفته از یادگیری ماشین است که بر پایه شبکه‌های عصبی مصنوعی چندلایه عمل می‌کند و توانایی استخراج ویژگی‌های پیچیده از حجم زیادی از داده‌ها را دارد. این فناوری در سال‌های اخیر در زمینه‌هایی مانند تشخیص بدافزار، تحلیل ترافیک شبکه و شناسایی الگوهای پیچیده حملات سایبری کاربرد گسترده‌ای پیدا کرده است. (گودفالو و همکاران^{۱۳}، ۲۰۱۶)

۶-۲- تشخیص ناهنجاری (Anomaly Detection)

تشخیص ناهنجاری به فرآیندی گفته می‌شود که طی آن الگوها یا رخدادهایی که با رفتار معمول یک سیستم، شبکه یا کاربر تفاوت دارند شناسایی می‌شوند. در محیط‌های امنیت سایبری، تشخیص ناهنجاری یکی از روش‌های مهم برای کشف حملات ناشناخته، تهدیدات داخلی و فعالیت‌های غیرمعمول کاربران محسوب می‌شود. (چاندولا و همکاران^{۱۴}، ۲۰۰۹)

۷-۲- تحلیل رفتار کاربران و موجودیت‌ها (UEBA - User and Entity Behavior Analytics)

تحلیل رفتار کاربران و موجودیت‌ها رویکردی پیشرفته در امنیت سایبری است که علاوه بر رفتار کاربران، فعالیت سایر موجودیت‌ها مانند دستگاه‌ها، سرورها و برنامه‌های کاربردی را نیز بررسی می‌کند. UEBA با ایجاد الگوهای رفتاری معمول و استفاده از الگوریتم‌های هوشمند، می‌تواند انحراف از رفتار طبیعی را شناسایی کرده و احتمال وقوع تهدیدات امنیتی را افزایش دهد. (گارتنر^{۱۵}، ۲۰۱۸)

۸-۲- تهدیدات سایبری (Cyber Threats) تهدیدات سایبری به هرگونه اقدام یا فعالیت مخرب گفته می‌شود که با هدف آسیب‌رسانی، سرقت اطلاعات، اختلال در عملکرد سامانه‌ها یا دسترسی غیرمجاز به منابع دیجیتال انجام می‌شود. این تهدیدات شامل بدافزارها، حملات فیشینگ، نفوذ به شبکه، حملات مهندسی اجتماعی و تهدیدات داخلی هستند و شناسایی سریع آن‌ها یکی از چالش‌های اصلی امنیت اطلاعات محسوب می‌شود. (ویتمن و ماتورد^{۱۶}، ۲۰۲۱)

۹-۲- مبانی نظری

هوش مصنوعی (Artificial Intelligence) به مجموعه‌ای از فناوری‌ها، روش‌ها و الگوریتم‌های هوشمند گفته می‌شود که توانایی یادگیری، تحلیل داده‌ها، استدلال، تصمیم‌گیری و حل مسائل را بدون نیاز به برنامه‌ریزی صریح برای هر وضعیت خاص دارا هستند. هدف اصلی این فناوری، شبیه‌سازی بخشی از قابلیت‌های شناختی انسان در سامانه‌های رایانه‌ای است تا بتوانند با دریافت داده‌های جدید، الگوهای موجود را شناسایی کرده، از تجربیات گذشته بیاموزند و در شرایط مختلف بهترین تصمیم را اتخاذ کنند. در سال‌های اخیر، پیشرفت‌های چشمگیر در حوزه یادگیری ماشین (Machine Learning)، یادگیری عمیق (Deep Learning)، شبکه‌های عصبی مصنوعی و پردازش کلان‌داده‌ها، زمینه توسعه گسترده سامانه‌های هوشمند را در حوزه‌های مختلف از جمله پزشکی، صنعت، حمل‌ونقل، تجارت الکترونیک و به‌ویژه امنیت سایبری فراهم کرده است. امروزه بسیاری از سازمان‌ها برای افزایش سرعت شناسایی تهدیدات، کاهش خطاهای انسانی، بهبود فرآیند تصمیم‌گیری امنیتی و

¹² Mitchell

¹³ Goodfellow et al

¹⁴ Chandola et al

¹⁵ Gartner

¹⁶ Whitman & Mattord

مدیریت هوشمند رخدادهای امنیتی از سامانه‌های مبتنی بر هوش مصنوعی بهره می‌گیرند. (مؤسسه ملی استاندارد و فناوری آمریکا^{۱۷}، ۲۰۲۴؛ آژانس امنیت سایبری اتحادیه اروپا^{۱۸}، ۲۰۲۳)

رشد سریع فناوری‌های دیجیتال، توسعه اینترنت اشیا (IoT)، رایانش ابری، شبکه‌های نسل پنجم (5G) و افزایش تعاملات آنلاین موجب شده است حجم داده‌های تولیدشده در سازمان‌ها به‌طور چشمگیری افزایش یابد. این تحول، علاوه بر ایجاد فرصت‌های فراوان، سطح حملات سایبری را نیز گسترش داده است. مهاجمان سایبری با استفاده از روش‌های پیچیده و هوشمند، حملات خود را به‌صورت مستمر تغییر داده و از تکنیک‌هایی مانند بدافزارهای پیشرفته، حملات فیشینگ، باج‌افزارها، حملات روز صفر (Zero-Day Attacks) و تهدیدات مداوم پیشرفته (Advanced Persistent Threats یا APT) برای نفوذ به سامانه‌های اطلاعاتی استفاده می‌کنند. در چنین شرایطی، استفاده از روش‌های سنتی امنیتی به‌تنهایی پاسخگوی تهدیدات نوین نیست و به‌کارگیری فناوری‌های مبتنی بر هوش مصنوعی به یک ضرورت راهبردی تبدیل شده است. (مایکروسافت، گزارش دفاع دیجیتال مایکروسافت^{۱۹}، ۲۰۲۴)

در رویکردهای سنتی امنیت سایبری، شناسایی حملات عمدتاً بر اساس امضاهای شناخته‌شده، قوانین از پیش تعریف‌شده و الگوهای ثابت انجام می‌شود. اگرچه این روش‌ها در شناسایی تهدیدات شناخته‌شده عملکرد مناسبی دارند، اما در برابر حملات جدید، حملات روز صفر و الگوهای متغیر مهاجمان با محدودیت‌های جدی روبه‌رو هستند. از سوی دیگر، الگوریتم‌های یادگیری ماشین با تحلیل حجم عظیمی از داده‌های شبکه، گزارش‌های امنیتی، رخدادهای ثبت‌شده و رفتار کاربران، قادرند الگوهای غیرعادی را شناسایی کرده و احتمال وقوع حملات ناشناخته را با دقت بیشتری پیش‌بینی کنند. (بوچزاک و گوون، ۲۰۱۶؛ آی‌بی‌ام سکيوریتی، ۲۰۲۳) این الگوریتم‌ها با یادگیری مداوم از داده‌های جدید، به‌مرور زمان دقت خود را افزایش داده و می‌توانند با تغییر رفتار مهاجمان نیز سازگار شوند.

برای درک بهتر تفاوت میان رویکردهای سنتی و سامانه‌های مبتنی بر هوش مصنوعی در امنیت سایبری، مقایسه‌ای از مهم‌ترین ویژگی‌های این دو رویکرد در جدول ۱ ارائه شده است.

جدول ۱- مقایسه روش‌های سنتی و روش‌های مبتنی بر هوش مصنوعی در امنیت سایبری

کابرد	مزیت هوش مصنوعی	روش‌های مبتنی بر هوش مصنوعی	روش‌های سنتی	شاخص مقایسه
تشخیص تهدیدات	شناسایی الگوهای جدید	مبتنی بر یادگیری از داده‌ها	مبتنی بر امضا و قوانین ثابت	روش تشخیص
امنیت شبکه	کشف تهدیدات ناشناخته	بسیار مناسب	ضعیف	تشخیص حملات روز صفر
کنترل دسترسی	تشخیص رفتار غیرعادی	پیشرفته (UEBA)	محدود	تحلیل رفتار کاربران
مراکز عملیات امنیت (SOC)	پاسخ سریع‌تر	بسیار بالا و بلادرنگ	متوسط	سرعت تحلیل
مدیریت رخدادهای امنیتی	افزایش دقت	کمتر با آموزش مناسب	نسبتاً زیاد	نرخ هشدار نادرست

همان‌گونه که در جدول (۱) مشاهده می‌شود، سامانه‌های مبتنی بر هوش مصنوعی در بسیاری از شاخص‌ها، به ویژه در شناسایی تهدیدات ناشناخته، تحلیل رفتار کاربران و سرعت پاسخ‌گویی، عملکرد بهتری نسبت به روش‌های سنتی دارند.

۱-۹-۲- تحلیل رفتار کاربران و موجودیت‌ها (UEBA)

¹⁷ NIST

¹⁸ ENISA

¹⁹ Microsoft Digital Defense Report

یکی از مهم‌ترین کاربردهای هوش مصنوعی در امنیت سایبری، تحلیل رفتار کاربران (User Behavior Analysis) و تحلیل رفتار کاربران و موجودیت‌ها (User and Entity Behavior Analytics - UEBA) است. (وانگ و همکاران^{۲۰}، ۲۰۲۴)

این رویکرد بر این اصل استوار است که هر کاربر یا هر موجودیت در یک سامانه اطلاعاتی دارای الگوی رفتاری مشخصی است. با ایجاد یک خط مبنا (Baseline) از رفتار عادی کاربران، سامانه‌های هوشمند می‌توانند هرگونه انحراف از الگوی معمول را شناسایی کرده و احتمال وقوع تهدیدات داخلی، سوءاستفاده از حساب‌های کاربری، سرقت اطلاعات یا نفوذ مهاجمان را تشخیص دهند. استفاده از UEBA به‌ویژه در سازمان‌های بزرگ که هزاران کاربر و دستگاه به‌طور هم‌زمان در حال فعالیت هستند، موجب افزایش دقت شناسایی تهدیدات و کاهش هشدارهای کاذب می‌شود. (المنصوری و همکاران^{۲۱}، ۲۰۲۲؛ میکروسافت، گزارش دفاع دیجیتال مایکروسافت^{۲۲}، ۲۰۲۴)

یکی دیگر از قابلیت‌های مهم هوش مصنوعی، تحلیل بلادرنگ (Real-Time Analysis) داده‌های امنیتی است. سامانه‌های هوشمند می‌توانند با بررسی مداوم ترافیک شبکه، رفتار کاربران، فعالیت تجهیزات، درخواست‌های دسترسی و تغییرات غیرمعمول در الگوهای ارتباطی، نشانه‌های اولیه حملات سایبری را پیش از وقوع خسارت شناسایی کنند. این ویژگی باعث می‌شود مراکز عملیات امنیت (SOC) بتوانند در کوتاه‌ترین زمان ممکن نسبت به رخداد‌های امنیتی واکنش نشان داده و از گسترش حملات جلوگیری کنند. همچنین، تحلیل بلادرنگ داده‌ها امکان اولویت‌بندی هشدارهای امنیتی، کاهش حجم هشدارهای غیرضروری و تخصیص بهینه منابع انسانی را فراهم می‌سازد. (آی‌بی‌ام سکیوریتی^{۲۳}، ۲۰۲۳) (آژانس امنیت سایبری اتحادیه اروپا^{۲۴}، ۲۰۲۳)



شکل ۱- نمودار چرخه شناسایی تهدید

شکل (۱)، چرخه کلی شناسایی و پاسخ‌گویی به تهدیدات سایبری مبتنی بر هوش مصنوعی را نشان می‌دهد. در این فرآیند، داده‌های محیط دیجیتال جمع‌آوری شده، پس از تحلیل توسط الگوریتم‌های هوشمند، ناهنجاری‌ها و تهدیدات احتمالی شناسایی شده و اقدامات پاسخ مناسب انجام می‌شود.

برای مقایسه قابلیت‌ها و کاربردهای مهم‌ترین الگوریتم‌های هوش مصنوعی در امنیت سایبری، اطلاعات مربوطه در جدول (۲) ارائه شده است.

جدول ۲- مقایسه مهم‌ترین الگوریتم‌های هوش مصنوعی مورد استفاده در امنیت سایبری

²⁰ Wang et al

²¹ Mansoori et al

²² Microsoft Digital Defense Report

²³ IBM Security

²⁴ ENISA

نمونه استفاده	محدودیت	مزیت	مهم ترین کاربرد	الگوریتم
IDS	مصرف حافظه بیشتر	دقت بالا	تشخیص نفوذ	Random Forest
Malware Detection	تنظیمات پیچیده	سرعت و دقت بالا	تشخیص بدافزار	XGBOOST
تحلیل ترافیک شبکه	نیاز به داده زیاد	استخراج ویژگی قوی	تشخیص الگوهای حمله	CNN
UEBA	آموزش زمان بر	توانایی مدل سازی وابستگی های زمانی	تحلیل رفتار کاربران	LSTM
Threat Intelligence	نیاز به پردازش بالا	یادگیری روابط پیچیده	شناسایی تهدیدات پیچیده	Transformer

همان گونه که در جدول (۲) مشاهده می شود، هر یک از الگوریتم های هوش مصنوعی با توجه به ویژگی های خود، در بخش های مختلف امنیت سایبری از جمله تشخیص نفوذ، تحلیل رفتار کاربران، شناسایی بدافزارها و پیش بینی تهدیدات کاربرد دارند و انتخاب الگوریتم مناسب به نوع داده ها و هدف سامانه امنیتی بستگی دارد.

در سال های اخیر، ترکیب فناوری های هوش مصنوعی با معماری Zero Trust، سامانه های تشخیص و جلوگیری از نفوذ (IDS/IPS)، هوش تهدید (Threat Intelligence) و سامانه های مدیریت اطلاعات و رخدادهای امنیتی (SIEM) به یکی از مهم ترین رویکردهای نوین در امنیت سایبری تبدیل شده است. این همگرایی موجب می شود داده های امنیتی از منابع مختلف به صورت یکپارچه تحلیل شده و سامانه بتواند با دقت بیشتری فعالیت های مشکوک را شناسایی کند. همچنین استفاده از مدل های پیش بینی کننده مبتنی بر هوش مصنوعی، امکان پیش بینی روند حملات آینده و اتخاذ اقدامات پیشگیرانه را برای سازمان ها فراهم می کند. (مؤسسه ملی استاندارد و فناوری آمریکا، ۲۰۲۳؛ شرکت گوگل کلاود، پیش بینی امنیت سایبری، ۲۰۲۵)

با وجود مزایای فراوان، استفاده از هوش مصنوعی در امنیت سایبری با چالش هایی نیز همراه است. کیفیت داده های آموزشی، وجود داده های نامتوازن، احتمال ایجاد هشدارهای نادرست، تغییر مداوم الگوهای حمله، قابلیت توضیح پذیری مدل های هوشمند، هزینه های پیاده سازی و نگهداری، حملات خصمانه علیه مدل های یادگیری ماشین (Adversarial Attacks) و نگرانی های مربوط به حفظ حریم خصوصی از مهم ترین مسائل مطرح در پژوهش های اخیر محسوب می شوند. به همین دلیل، پژوهشگران علاوه بر توسعه مدل های دقیق تر، بر افزایش شفافیت، اعتمادپذیری، مقاومت در برابر حملات و طراحی چارچوب های اخلاقی برای استفاده از هوش مصنوعی نیز تمرکز کرده اند. (آژانس امنیت سایبری اتحادیه اروپا^{۲۵}، ۲۰۲۳؛ مؤسسه ملی استاندارد و فناوری آمریکا، ۲۰۲۴)

در مجموع، بررسی مبانی نظری پژوهش نشان می دهد که هوش مصنوعی با بهره گیری از الگوریتم های پیشرفته یادگیری ماشین، یادگیری عمیق و تحلیل هوشمند داده ها، ظرفیت قابل توجهی برای ارتقای امنیت سایبری دارد. با این حال، موفقیت این فناوری در گرو استفاده از داده های باکیفیت، به روزرسانی مستمر مدل ها، رعایت اصول امنیت اطلاعات و ترکیب فناوری های هوشمند با راهبردهای مدیریتی و امنیتی مناسب خواهد بود.

۱۰-۲- پیشینه پژوهش

یوسفی و همکاران (۱۴۰۱) در پژوهشی با عنوان «شناسایی حملات سایبری و تهدیدات پیشرفته با استفاده از الگوریتم های یادگیری ماشین در شبکه های سازمانی» که با رویکرد کاربردی و روش توصیفی-تحلیلی انجام شده بود، در

جامعه آماری متخصصان امنیت فناوری اطلاعات به تعداد ۲۸۰ نفر و نمونه آماری ۱۶۲ نفر که به روش نمونه‌گیری تصادفی ساده انتخاب شدند، به این نتیجه رسیدند که تلفیق رفتارسنجی کاربران (UBA) با مدل‌های یادگیری نظارت‌نشده (مانند جنگل تصادفی و شبکه‌های عصبی عمیق) می‌تواند نرخ تشخیص مثبت کاذب در تهدیدات داخلی (Insider Threats) را تا کاهش قابل توجهی در نرخ هشدارهای کاذب ایجاد کند و سرعت پاسخ‌دهی به حوادث امنیتی را به طور چشمگیری افزایش دهد.

رضایی و احمدی (۱۴۰۲) در پژوهشی با عنوان «تحلیل رفتار کاربران در بانکداری الکترونیک به منظور پیشگیری از کلاهبرداری با رویکرد هوش مصنوعی» که با روش توسعه‌ای-کاربردی انجام شده بود، در جامعه آماری مشتریان فعال درگاه‌های بانکی به تعداد ۵۰۰ نفر و نمونه آماری ۲۱۷ نفر به این نتیجه رسیدند که مدل‌سازی رفتار تراکنشی کاربران بر پایه ویژگی‌های زمانی و مکانی و استفاده از الگوریتم‌های خوشه‌بندی، شناسایی الگوهای غیرعادی (Anomaly Detection) را در لحظه ممکن ساخته و امنیت تراکنش‌های دیجیتال را در بستر بانکداری مدرن تا ۹۲ درصد ارتقا می‌بخشد.

محمدی و عباسی (۱۴۰۳) در پژوهشی با عنوان «نقش سیستم‌های نوین تشخیص نفوذ مبتنی بر یادگیری عمیق در امنیت شبکه‌های داده صنعتی» که با رویکرد شبه‌تجربی و تحلیل داده‌های شبیه‌سازی شده انجام شده بود، در جامعه آماری کارشناسان امنیت سیستم‌های کنترل صنعتی به تعداد ۱۵۰ نفر و نمونه آماری ۱۰۸ نفر به این نتیجه رسیدند که الگوریتم‌های حافظه طولانی کوتاه‌مدت (LSTM) با تحلیل مداوم رفتار ترافیک شبکه و مقایسه آن با رفتارهای بهنجار ثبت‌شده، توانایی بالایی در شناسایی حملات روز صفر (Zero-day) و تهدیدات پیشرفته مستمر (APT) دارند و بستری امن‌تر برای انتقال داده‌ها فراهم می‌کنند.

المنصوری و همکاران^{۲۶} (۲۰۲۲) در پژوهشی با عنوان «تحلیل رفتار کاربران و موجودیت‌ها (UEBA) با استفاده از تکنیک‌های یادگیری ماشین برای شناسایی تهدیدات سایبری» که با رویکرد کاربردی و روش کمی انجام شد، در جامعه آماری متخصصان فناوری اطلاعات به تعداد ۴۰۰ نفر و نمونه آماری ۱۹۶ نفر، به این نتیجه رسیدند که ترکیب طبقه‌بندهای یادگیری نظارت‌شده به‌ویژه الگوریتم (XG Boost) با داده‌های لاگ فعالیت کاربران، به تیم‌های امنیتی امکان می‌دهد سرعت اعتبارنامه‌ها (Credential Theft) و حرکت جانبی مهاجم در شبکه (Lateral Movement) را در شبکه‌های سازمانی با دقتی بیش از ۹۴.۵٪ شناسایی کنند.

ژانگ و وانگ^{۲۷} (۲۰۲۳) در پژوهشی با عنوان «چارچوب هوشمند برای شناسایی تهدیدات داخلی مبتنی بر یادگیری عمیق و تحلیل رفتاری در محیط‌های دیجیتال» که با روش‌شناسی تجربی و طراحی شبیه‌سازی انجام شد، در جامعه آماری خبرگان امنیت سایبری و مدیران سامانه‌ها به تعداد ۳۱۰ نفر و نمونه آماری ۱۷۲ نفر، دریافتند که مدل‌های یادگیری عمیق مبتنی بر توالی مانند شبکه‌های عصبی بازگشتی (RNN) در صورت‌بندی و مشخصه‌یابی انحرافات در فرمان‌های کاربر و الگوهای دسترسی به فایل بسیار اثربخش هستند؛ به‌گونه‌ای که می‌توانند با شناسایی زودهنگام الگوهای مشکوک، از وقوع نشت داده (Data Exfiltration) پیش از ایجاد خسارت جدی جلوگیری کنند.

²⁶ Al-Mansoori et al

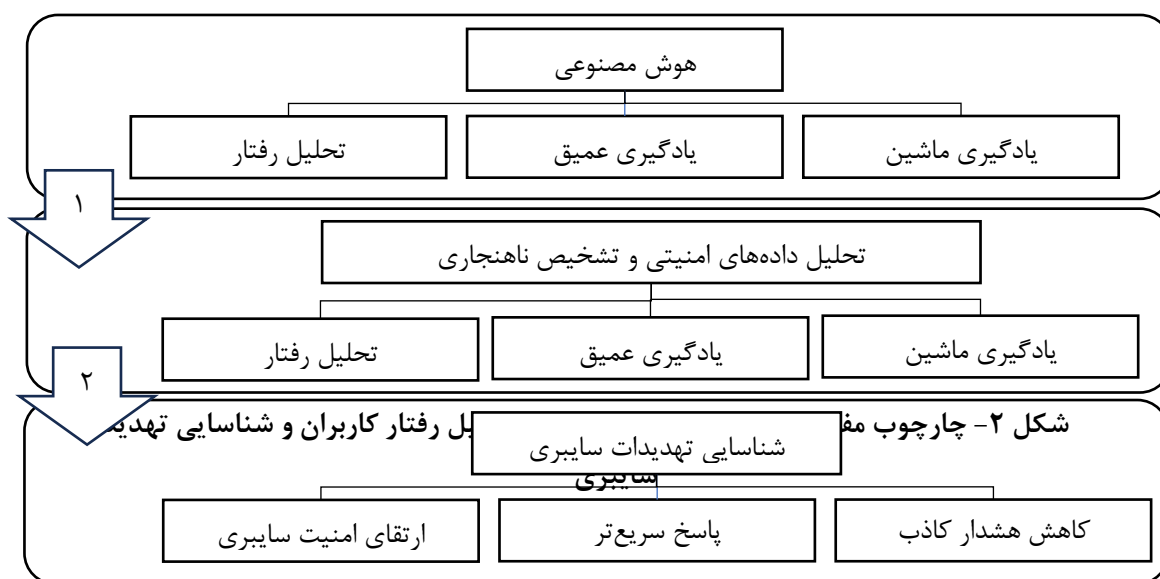
²⁷ Zhang & Wang

تامپسون و دیویس^{۲۸} (۲۰۲۴) در پژوهشی با عنوان «به کارگیری هوش مصنوعی برای تهدید یابی پیش‌دستانه و پروفایل سازی کاربران در رایانش ابری» که با رویکرد اکتشافی و روش تحلیلی انجام گرفت، در جامعه آماری کاربران و مدیران زیرساخت ابری به تعداد ۴۵۰ نفر و نمونه آماری ۲۰۷ نفر نشان دادند که پروفایل سازی پویای رفتاری در کنار عامل‌های یادگیری تقویتی (Reinforcement Learning Agents)، امکان انطباق بلادرنگ سطح دسترسی‌ها را فراهم می‌کند و در نتیجه، سطح حمله (Attack Surface) را در محیط‌های دیجیتال مبتنی بر ابر به طور معناداری کاهش می‌دهد.

۳- شکاف پژوهشی

بررسی مطالعات انجام‌شده نشان می‌دهد که پژوهش‌های پیشین نقش مؤثر فناوری‌های هوش مصنوعی، به‌ویژه الگوریتم‌های یادگیری ماشین و یادگیری عمیق، را در حوزه‌هایی مانند تشخیص نفوذ، تحلیل رفتار کاربران، شناسایی ناهنجاری‌ها و مقابله با تهدیدات سایبری مورد تأیید قرار داده‌اند. با این حال، بسیاری از مطالعات موجود بر یک الگوریتم یا یک حوزه خاص از امنیت سایبری تمرکز داشته‌اند و کمتر پژوهشی به بررسی جامع ارتباط میان تحلیل رفتار کاربران، فناوری‌های هوش مصنوعی، سامانه‌های UEBA و فرآیند شناسایی تهدیدات نوظهور در محیط‌های دیجیتال پرداخته است. بنابراین، پژوهش حاضر با تمرکز بر کاربرد روش‌های هوش مصنوعی در تحلیل رفتار کاربران و شناسایی تهدیدات سایبری، تلاش می‌کند با ارائه دیدگاهی یکپارچه، ظرفیت‌ها، مزایا و چالش‌های استفاده از این فناوری‌ها را در ارتقای امنیت سایبری بررسی کند.

با توجه به مبانی نظری و یافته‌های حاصل از مطالعات پیشین، چارچوب مفهومی پژوهش حاضر با هدف تبیین ارتباط میان داده‌های محیط دیجیتال، فناوری‌های هوش مصنوعی، تحلیل رفتار کاربران و فرآیند شناسایی تهدیدات سایبری ارائه شده است. شکل (۲) چارچوب مفهومی پیشنهادی پژوهش را نشان می‌دهد.



همان گونه که در شکل (۲) مشاهده می شود، هوش مصنوعی با بهره گیری از الگوریتم های یادگیری ماشین، یادگیری عمیق و تحلیل رفتار کاربران و موجودیت ها، امکان تشخیص ناهنجاری ها، شناسایی تهدیدات سایبری و ارتقای سطح امنیت دیجیتال را فراهم می سازد.

۴- روش تحقیق

پژوهش حاضر از نظر هدف، از نوع کاربردی و از نظر روش اجرا، مروری-توصیفی است. هدف اصلی این پژوهش، بررسی نقش هوش مصنوعی در تحلیل رفتار کاربران و شناسایی تهدیدات سایبری در محیط های دیجیتال و همچنین تبیین ظرفیت فناوری های هوشمند در ارتقای امنیت سایبری است. از آنجا که این حوزه با سرعت زیادی در حال تحول بوده و فناوری ها و روش های نوین به طور مستمر توسعه می یابند، استفاده از روش مروری-توصیفی امکان بررسی، مقایسه و تحلیل جامع یافته های پژوهش های معتبر را فراهم می کند. بر همین اساس، مطالعات، مقالات و گزارش های علمی منتشر شده در منابع معتبر بین المللی به صورت نظام مند مورد بررسی قرار گرفت تا تصویری جامع از وضعیت موجود و روندهای نوین این حوزه ارائه شود. (مؤسسه ملی استانداردها و فناوری آمریکا، ۲۰۲۴)

گردآوری اطلاعات در این پژوهش به روش مطالعه کتابخانه ای انجام شده است. بدین منظور، مقالات علمی، کتاب ها، گزارش های تخصصی و اسناد پژوهشی منتشر شده در پایگاه های معتبر بین المللی از جمله ScienceDirect، IEEE Xplore، SpringerLink، ACM Digital Library و MDPI مورد جست و جو و بررسی قرار گرفتند. انتخاب این پایگاه ها به دلیل اعتبار علمی، انتشار پژوهش های به روز و دسترسی به مقالات تخصصی در حوزه های هوش مصنوعی، امنیت سایبری، یادگیری ماشین و تحلیل رفتار کاربران صورت گرفته است. (گوگل کلود، ۲۰۲۵)

به منظور افزایش اعتبار علمی پژوهش و بهره گیری از جدیدترین استانداردها و گزارش های تخصصی، علاوه بر مقالات علمی، از گزارش های رسمی منتشر شده توسط سازمان ها و نهادهای معتبر بین المللی نیز استفاده شد. در این راستا، گزارش های مؤسسه ملی استاندارد و فناوری آمریکا (NIST)، آژانس امنیت سایبری اتحادیه اروپا (ENISA)، IBM Security، Microsoft Digital Defense Report و Google Cloud Cybersecurity Forecast به عنوان منابع اصلی تحلیل انتخاب شدند؛ زیرا این گزارش ها بر پایه داده های واقعی، رخدادهای امنیتی و تحلیل های تخصصی تهیه شده و از معتبرترین منابع در حوزه امنیت سایبری به شمار می روند. (مؤسسه ملی استاندارد و فناوری آمریکا، ۲۰۲۴)

برای حفظ به روز بودن یافته های پژوهش، تمرکز اصلی بر منابع منتشر شده در فاصله زمانی ۲۰۲۱ تا ۲۰۲۵ بوده است؛ با این حال، برای تبیین مفاهیم پایه و مبانی نظری، از برخی منابع بنیادین و معتبر حوزه هوش مصنوعی و امنیت سایبری نیز استفاده شده است. پس از گردآوری منابع، فرآیند غربالگری بر اساس معیارهایی مانند ارتباط مستقیم با موضوع پژوهش، اعتبار علمی ناشر یا مجله، کیفیت روش شناسی، میزان استنادپذیری، تازگی یافته ها و انطباق با اهداف پژوهش انجام شد. در نتیجه، منابعی که از نظر محتوایی بیشترین ارتباط را با کاربرد هوش مصنوعی در امنیت سایبری و تحلیل رفتار کاربران داشتند، برای تحلیل نهایی انتخاب شدند. (گوگل کلود، ۲۰۲۵)

به منظور افزایش دقت و اعتبار نتایج، فرآیند بررسی منابع در چند مرحله انجام شد. در مرحله نخست، منابع مرتبط با کلیدواژه هایی نظیر «هوش مصنوعی»، «امنیت سایبری»، «تحلیل رفتار کاربران»، «یادگیری ماشین»، «یادگیری عمیق»، «تشخیص ناهنجاری» و «UEBA» شناسایی شدند. در مرحله بعد، عنوان، چکیده و محتوای مقالات از نظر میزان ارتباط با

اهداف پژوهش مورد ارزیابی قرار گرفت و منابعی که فاقد ارتباط مستقیم با موضوع، دارای اطلاعات تکراری یا فاقد اعتبار علمی کافی بودند، کنار گذاشته شدند. در نهایت، منابع منتخب بر اساس کیفیت علمی، به روز بودن، استانداردپذیری و میزان انطباق با اهداف پژوهش برای تحلیل نهایی مورد استفاده قرار گرفتند. این فرآیند سبب شد یافته‌های پژوهش بر پایه معتبرترین مطالعات و گزارش‌های تخصصی موجود در حوزه هوش مصنوعی و امنیت سایبری تدوین شود.

در مرحله تحلیل داده‌ها، اطلاعات استخراج‌شده از مطالعات منتخب با استفاده از روش تحلیل محتوای کیفی بررسی و طبقه‌بندی شد. در این مرحله، یافته‌های پژوهش‌ها بر اساس محورهای شامل کاربرد الگوریتم‌های هوش مصنوعی در تحلیل رفتار کاربران، روش‌های تشخیص ناهنجاری، شناسایی تهدیدات سایبری، کاربرد یادگیری ماشین و یادگیری عمیق، فناوری تحلیل رفتار کاربران و موجودیت‌ها (UEBA)، معماری Zero Trust، سامانه‌های تشخیص نفوذ، مزایا، محدودیت‌ها و چالش‌های پیاده‌سازی سامانه‌های هوشمند امنیتی دسته‌بندی و با یکدیگر مقایسه شدند. این شیوه تحلیل امکان شناسایی نقاط اشتراک، تفاوت‌ها، مزایا و کاستی‌های مطالعات پیشین را فراهم ساخت و زمینه استخراج نتایج جامع‌تر را مهیا کرد (آی بی ام سکورتی، ۲۰۲۳).

در نهایت، با تلفیق یافته‌های پژوهش‌های منتخب، چارچوبی جامع برای تبیین نقش هوش مصنوعی در ارتقای امنیت سایبری، تحلیل رفتار کاربران و بهبود فرآیند شناسایی تهدیدات نوظهور ارائه شد. این چارچوب ضمن تبیین مزایا و قابلیت‌های فناوری‌های هوشمند، چالش‌ها و محدودیت‌های موجود را نیز مورد توجه قرار داده و می‌تواند به عنوان مبنایی برای انجام پژوهش‌های آینده و توسعه راهکارهای نوین امنیت سایبری مورد استفاده پژوهشگران، مدیران فناوری اطلاعات و متخصصان امنیت قرار گیرد (مؤسسه ملی استانداردها و فناوری آمریکا، ۲۰۲۴).

۵- یافته‌های پژوهش

بررسی مطالعات منتخب نشان می‌دهد که هوش مصنوعی در سال‌های اخیر نقش مهمی در شناسایی و مقابله با تهدیدات سایبری ایفا کرده است. کاربرد الگوریتم‌های یادگیری ماشین، یادگیری عمیق و تحلیل رفتار کاربران موجب افزایش دقت تشخیص حملات، کاهش هشدارهای نادرست و بهبود سرعت واکنش به رخدادها امنیتی شده است. تهدیداتی مانند بدافزارها، حملات فیشینگ، حملات روز صفر، باج‌افزارها و تهدیدات داخلی از مهم‌ترین مخاطراتی هستند که سامانه‌های مبتنی بر هوش مصنوعی در شناسایی آن‌ها عملکرد قابل توجهی از خود نشان داده‌اند. این فناوری با تحلیل حجم عظیمی از داده‌های شبکه و رفتاری، الگوهای پنهان را استخراج کرده و توانایی شناسایی تهدیدات ناشناخته را نیز دارد. مطالعات فارسی متعدد نیز بر این نکته تأکید دارند که هوش مصنوعی امنیت سایبری را از حالت سنتی به سمت امنیت هوشمند سوق داده و امکان شناسایی، خنثی‌سازی و حتی پیش‌بینی حملات را فراهم کرده است (کشاورز و حسینی، ۱۴۰۲؛ گلدی نجف‌آباد، ۱۴۰۴).

به‌منظور ارائه تصویری روشن از مهم‌ترین تهدیدات سایبری و نقش فناوری‌های مبتنی بر هوش مصنوعی در مقابله با آن‌ها، جدول (۳) تهیه شده است.

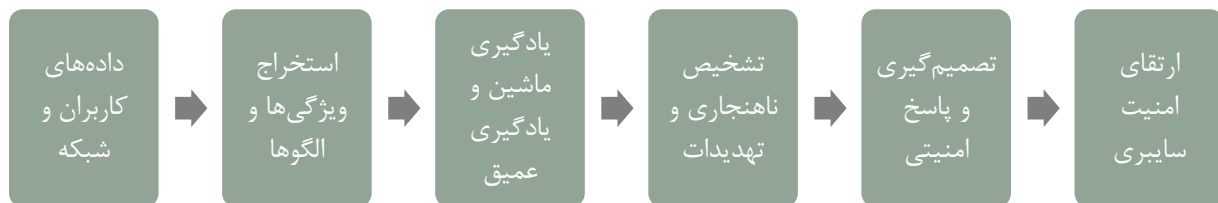
جدول ۳- مهم‌ترین تهدیدات سایبری و نقش هوش مصنوعی در شناسایی و مقابله با آن‌ها

فناوری یا روش مورد استفاده	نقش هوش مصنوعی در شناسایی	توضیح تهدید	نوع تهدید سایبری
یادگیری ماشین، یادگیری عمیق	تحلیل رفتار فایل ها و شناسایی الگوهای مشکوک	برنامه های مخربی که برای آسیب رسانی، سرقت اطلاعات یا کنترل سیستم ها طراحی می شوند.	بدافزار (Malware)
پردازش زبان طبیعی (NLP)، مدل های طبقه بندی	تشخیص پیام های جعلی و تحلیل محتوای مشکوک	تلاش برای فریب کاربران و سرقت اطلاعات حساس از طریق پیام ها و وبسایت های جعلی	فیشینگ (Phishing)
تشخیص ناهنجاری، یادگیری عمیق	کشف رفتارهای غیرعادی بدون نیاز به امضای قبلی	حملاتی که قبل از شناسایی و ارائه راهکار امنیتی رخ می دهند	حملات روز صفر (Zero-Day)
تحلیل رفتار، الگوریتم های ML	شناسایی فعالیت های غیر معمول در فایل ها و سیستم	رمزگذاری اطلاعات و درخواست باج برای بازگرداندن دسترسی	باج افزار (Ransomware)
تحلیل رفتار کاربران، UEBA	شناسایی انحراف از الگوی رفتاری معمول کاربران	سوءاستفاده کاربران مجاز از دسترسی های خود	تهدیدات داخلی (Threats Insider)

همان گونه که در جدول (۳) مشاهده می شود، هوش مصنوعی در مقابله با طیف گسترده ای از تهدیدات سایبری مؤثر است. یافته های حاصل از بررسی مطالعات نشان می دهد که میزان اثربخشی الگوریتم های هوش مصنوعی در شناسایی تهدیدات سایبری به عواملی مانند کیفیت داده های آموزشی، نوع الگوریتم، حجم و تنوع داده های ورودی، نحوه آموزش مدل ها و شرایط عملیاتی سامانه بستگی دارد. در میان الگوریتم های بررسی شده، مدل های مبتنی بر یادگیری عمیق به دلیل توانایی بالا در استخراج ویژگی های پیچیده از داده های حجیم و شناسایی روابط پنهان میان الگوهای رفتاری، عملکرد مطلوب تری در تشخیص حملات ناشناخته، حملات روز صفر و رفتارهای غیرعادی کاربران از خود نشان داده اند. همچنین، به کارگیری مدل های ترکیبی (Hybrid Models) که از تلفیق چند الگوریتم هوش مصنوعی بهره می برند، موجب افزایش دقت تشخیص، کاهش نرخ هشدارهای کاذب، بهبود سرعت واکنش به رخداد های امنیتی و ارتقای کارایی سامانه های دفاعی شده است. این یافته ها بیانگر آن است که استفاده از رویکردهای هوشمند می تواند بسیاری از محدودیت های روش های سنتی را برطرف کرده و سطح تاب آوری سایبری سازمان ها را در برابر تهدیدات نوظهور به طور چشمگیری افزایش دهد.

علاوه بر این، پژوهش های فارسی اخیر نیز تأیید می کنند که ترکیب یادگیری ماشین و یادگیری عمیق در تشخیص نفوذ، بدافزار و ناهنجاری های رفتاری، دقت و سرعت بالاتری نسبت به روش های مبتنی بر امضا ارائه می دهد و می تواند به کاهش زمان پاسخ گویی کمک کند (احمدزاده، ۱۴۰۴؛ بهرامی، ۱۴۰۴). مدل های ترکیبی چندلایه نیز با بهره گیری از نقاط قوت الگوریتم های مختلف، وظایف استخراج ویژگی، طبقه بندی و پیش بینی را به صورت مکمل انجام داده و سازگاری سامانه را با الگوهای جدید حملات افزایش می دهند. مطالعات علم سنجی نیز نشان دهنده رشد قابل توجه تولید علمی در حوزه هوش مصنوعی و امنیت سایبری و شکل گیری خوشه های پژوهشی متنوع در این زمینه هستند (اندایش و کیان راد، ۱۴۰۴).

به منظور جمع بندی یافته های پژوهش و نمایش ارتباط میان مهم ترین مؤلفه های مؤثر در شناسایی تهدیدات سایبری، چارچوب مفهومی نقش هوش مصنوعی در تحلیل رفتار کاربران و ارتقای امنیت سایبری در شکل (۳) ارائه شده است.



شکل ۳- چارچوب مفهومی نقش هوش مصنوعی در تحلیل رفتار کاربران و شناسایی تهدیدات سایبری

همان‌گونه که در شکل (۳) مشاهده می‌شود، سامانه‌های مبتنی بر هوش مصنوعی با تحلیل داده‌های کاربران و شبکه، استخراج الگوهای رفتاری، به‌کارگیری الگوریتم‌های یادگیری ماشین و یادگیری عمیق، تشخیص ناهنجاری‌ها و در نهایت تصمیم‌گیری امنیتی، نقش مهمی در شناسایی سریع تهدیدات سایبری و ارتقای سطح امنیت سازمان‌ها ایفا می‌کنند.

بررسی مطالعات منتشرشده در سال‌های اخیر نشان می‌دهد که هوش مصنوعی به یکی از مهم‌ترین فناوری‌های مورد استفاده در امنیت سایبری تبدیل شده است. الگوریتم‌های یادگیری ماشین و یادگیری عمیق با تحلیل حجم عظیمی از داده‌های شبکه، گزارش‌های امنیتی و الگوهای رفتاری کاربران، قادرند تهدیدات سایبری را با دقت و سرعت بیشتری نسبت به روش‌های سنتی شناسایی کنند. این قابلیت موجب شده است که بسیاری از سازمان‌ها از سامانه‌های مبتنی بر هوش مصنوعی برای افزایش سطح امنیت زیرساخت‌های اطلاعاتی خود استفاده کنند. پژوهش‌های داخلی نیز بر نقش دوگانه هوش مصنوعی تأکید کرده‌اند: از یک سو تقویت سیستم‌های دفاعی و پیش‌بینی تهدیدات، و از سوی دیگر ایجاد تهدیدات نوظهور و پیچیده‌تر که نیازمند سیاست‌گذاری هوشمندانه است (قدیم‌پور شبلو و ابری، ۱۴۰۳؛ روا، ۱۴۰۴).

یکی از مهم‌ترین کاربردهای هوش مصنوعی در امنیت سایبری، تحلیل رفتار کاربران (User Behavior Analytics - UBA) و تحلیل رفتار کاربر و موجودیت (UEBA) است. در این روش، الگوهای عادی فعالیت کاربران مانند زمان ورود به سامانه، موقعیت جغرافیایی، نوع دستگاه مورد استفاده، میزان دسترسی به اطلاعات و نحوه تعامل با منابع سازمانی به‌طور مستمر تحلیل می‌شود. هرگونه انحراف از الگوهای معمول می‌تواند به‌عنوان یک رفتار مشکوک شناسایی شده و پیش از وقوع حمله، هشدارهای لازم صادر شود. پژوهش‌های فارسی نشان داده‌اند که هوش مصنوعی با پردازش حجم وسیع داده‌ها و بدون نیاز به مداخله گسترده انسانی، ویژگی‌های رفتاری کاربران و موجودیت‌ها را شناسایی کرده و ناهنجاری‌ها را با سرعت بیشتری کشف می‌کند و قابلیت‌های مراکز عملیات امنیت را دوجندان می‌سازد (مقاله کاربست هوش مصنوعی در تجزیه و تحلیل رفتار کاربر و موجودیت، ۱۴۰۳).

یکی از مهم‌ترین مزایای تحلیل رفتار کاربران، توانایی آن در شناسایی تهدیدات داخلی (Insider Threats) است. برخلاف بسیاری از حملات خارجی که از طریق نفوذ به شبکه انجام می‌شوند، تهدیدات داخلی ممکن است توسط کاربران مجاز، کارکنان یا پیمانکارانی ایجاد شوند که به منابع سازمانی دسترسی قانونی دارند. سامانه‌های مبتنی بر هوش مصنوعی با ایجاد الگوی رفتاری عادی برای هر کاربر و مقایسه مداوم فعالیت‌های جاری با این الگو، می‌توانند تغییرات غیرعادی مانند افزایش ناگهانی حجم دانلود اطلاعات، دسترسی به فایل‌های غیرمرتبط با وظایف سازمانی، ورود از مکان‌های غیرمعمول یا استفاده از سامانه در ساعات غیرمعارف را شناسایی کرده و پیش از وقوع خسارت، هشدارهای لازم را صادر کنند. این قابلیت

نقش مهمی در کاهش مخاطرات ناشی از سوءاستفاده از دسترسی‌های مجاز و افزایش امنیت اطلاعات سازمان ایفا می‌کند. مطالعات داخلی نیز بر اهمیت UEBA در کشف تهدیدات پیشرفته و تهدیدات داخلی تأکید دارند.

پژوهش‌های جدید نشان می‌دهند که الگوریتم‌هایی مانند XGBoost، Random Forest، Support Vector Machine (SVM)، شبکه‌های عصبی مصنوعی (ANN) و شبکه‌های عصبی عمیق (Deep Neural Networks) در تشخیص بدافزارها، حملات فیشینگ، نفوذ به شبکه و حملات روز صفر عملکرد مطلوبی داشته‌اند. این الگوریتم‌ها با یادگیری از داده‌های گذشته، الگوهای پنهان را استخراج کرده و توانایی شناسایی تهدیدات ناشناخته را نیز دارند. مطالعات فارسی نیز بر نقش یادگیری ماشین در حوزه‌هایی مانند تشخیص نفوذ، شناسایی بدافزار، کشف ناهنجاری و تشخیص فیشینگ تأکید کرده‌اند و چالش‌هایی مانند حملات خصمانه، کیفیت داده و تفسیرپذیری مدل‌ها را مورد توجه قرار داده‌اند (احمدزاده، ۱۴۰۴؛ عباس‌نژادورزی و همکاران، ۱۴۰۳).

یکی دیگر از یافته‌های مهم مطالعات اخیر، کاربرد یادگیری عمیق (Deep Learning) در تحلیل ترافیک شبکه است. مدل‌هایی مانند LSTM و CNN به دلیل توانایی بالا در یادگیری روابط پیچیده میان داده‌ها، در تشخیص رفتارهای غیرعادی و حملات پیشرفته سایبری دقت بالاتری نسبت به بسیاری از روش‌های کلاسیک نشان داده‌اند. این مدل‌ها به‌ویژه در محیط‌هایی که حجم داده بسیار زیاد است، عملکرد قابل توجهی دارند. بررسی‌های فارسی نیز نشان می‌دهد که شبکه‌های عصبی و روش‌های یادگیری عمیق دقت و سرعت تشخیص تهدیدات را به‌طور چشمگیری افزایش داده و در شناسایی الگوهای ناهنجار ترافیک شبکه و پیش‌بینی حملات موفق بوده‌اند (بهرامی، ۱۴۰۴؛ هاشمی شبیری، ۱۴۰۳).

بررسی مطالعات جدید همچنین نشان می‌دهد که استفاده از مدل‌های ترکیبی و چندلایه، عملکرد سامانه‌های امنیتی را نسبت به استفاده از یک الگوریتم منفرد بهبود می‌بخشد. در این رویکرد، الگوریتم‌های مختلف با بهره‌گیری از نقاط قوت یکدیگر، وظایفی مانند استخراج ویژگی‌ها، طبقه‌بندی تهدیدات، تحلیل رفتار کاربران و پیش‌بینی حملات را به‌صورت مکمل انجام می‌دهند. نتایج پژوهش‌ها بیانگر آن است که این مدل‌های ترکیبی علاوه بر افزایش دقت شناسایی تهدیدات، موجب کاهش هشدارهای کاذب، بهبود سرعت پردازش داده‌ها و افزایش قابلیت سازگاری سامانه با الگوهای جدید حملات سایبری می‌شوند. به همین دلیل، استفاده از معماری‌های ترکیبی به یکی از مهم‌ترین روندهای پژوهشی در حوزه امنیت سایبری مبتنی بر هوش مصنوعی تبدیل شده است. پژوهش‌های مروری فارسی نیز بر ضرورت مدل‌های مقاوم، توضیح‌پذیر و توزیع‌شده تأکید کرده‌اند که بتوانند در محیط‌های واقعی عملکرد دقیق ارائه دهند.

با وجود مزایای فراوان، استفاده از هوش مصنوعی در امنیت سایبری با چالش‌هایی نیز همراه است. کیفیت پایین داده‌های آموزشی، عدم توازن داده‌ها، احتمال تولید هشدارهای اشتباه (False Positive)، حملات خصمانه علیه مدل‌های یادگیری ماشین (Adversarial Attacks) و دشواری تفسیر تصمیمات برخی مدل‌های هوشمند، از مهم‌ترین محدودیت‌های این فناوری محسوب می‌شوند. پژوهشگران بر این باورند که توسعه مدل‌های توضیح‌پذیر (Explainable AI) و استفاده از داده‌های آموزشی باکیفیت، می‌تواند دقت و قابلیت اعتماد سامانه‌های امنیتی مبتنی بر هوش مصنوعی را افزایش دهد. مطالعات فارسی نیز بر ضرورت مدیریت چالش‌های فنی مانند مقیاس‌پذیری و تفسیرپذیری، و همچنین ملاحظات اخلاقی، حریم خصوصی و ضعف ساختارهای حقوقی تأکید کرده‌اند (روا، ۱۴۰۴؛ قدیم‌پور شبلو و ابری، ۱۴۰۳؛ مقاله تأثیر هوش مصنوعی مولد، ۱۴۰۳).

بررسی تطبیقی پژوهش‌های انجام‌شده نشان می‌دهد که استفاده از هوش مصنوعی در کنار سامانه‌های امنیتی سنتی، نسبت به استفاده مستقل از هر یک، نتایج مطلوب‌تری در شناسایی تهدیدات سایبری به همراه دارد. ترکیب فناوری‌هایی مانند هوش مصنوعی، تحلیل رفتار کاربران، سامانه‌های تشخیص نفوذ و اطلاعات تهدید (Threat Intelligence) موجب افزایش سرعت کشف حملات، کاهش زمان پاسخ‌گویی و بهبود تصمیم‌گیری امنیتی در سازمان‌ها می‌شود. در مجموع، یافته‌های این پژوهش نشان می‌دهد که هوش مصنوعی نقش مهمی در تحول امنیت سایبری ایفا می‌کند و با توسعه الگوریتم‌های یادگیری ماشین، یادگیری عمیق و تحلیل رفتار کاربران، می‌توان سامانه‌های امنیتی هوشمندتر، دقیق‌تر و مقاوم‌تری در برابر تهدیدات سایبری آینده طراحی و پیاده‌سازی کرد. مطالعات داخلی نیز بر لزوم تدوین سیاست‌ها و استانداردهای قانونی هوشمند برای تنظیم استفاده از هوش مصنوعی در این حوزه تأکید دارند.

۶- بحث و نتیجه‌گیری

تحول دیجیتال و افزایش وابستگی سازمان‌ها به زیرساخت‌های دیجیتال، موجب گسترش تهدیدات سایبری و پیچیده‌تر شدن روش‌های حمله شده است. یافته‌های این پژوهش نشان داد که روش‌های سنتی امنیت اطلاعات—که عمدتاً مبتنی بر امضاهای از پیش تعریف‌شده، قوانین ایستا و سامانه‌های تشخیص نفوذ مبتنی بر الگوهای شناخته‌شده هستند—به تنهایی توانایی کافی برای مقابله با تهدیدات نوظهور، حملات روز صفر (Zero-Day) و رفتارهای پیچیده مهاجمان را ندارند. در چنین فضایی، استفاده از فناوری‌های هوشمند، به‌ویژه هوش مصنوعی، به یک رویکرد ضروری و تحول‌آفرین در امنیت سایبری تبدیل شده است.

نتایج بررسی مطالعات نشان داد که هوش مصنوعی با بهره‌گیری از الگوریتم‌های یادگیری ماشین و یادگیری عمیق، قابلیت قابل‌توجهی در تحلیل رفتار کاربران، تشخیص ناهنجاری‌ها، شناسایی نفوذ، کشف بدافزارها و پیش‌بینی تهدیدات سایبری دارد. این فناوری با تحلیل حجم گسترده‌ای از داده‌های امنیتی (لاگ‌ها، ترافیک شبکه، رفتار کاربران و موجودیت‌ها) و استخراج الگوهای پنهان و غیرخطی، امکان شناسایی زود هنگام حملات، کاهش زمان واکنش و بهبود فرآیند تصمیم‌گیری امنیتی را فراهم می‌کند. برخلاف روش‌های سنتی که عمدتاً واکنشی عمل می‌کنند، مدل‌های هوشمند می‌توانند با ایجاد خط‌پایه‌های رفتاری پویا، انحرافات جزئی را شناسایی کرده و حتی تهدیدات ناشناخته را پیش‌بینی کنند.

همچنین نتایج نشان داد که استفاده از فناوری تحلیل رفتار کاربران و موجودیت‌ها (UEBA) در کنار سامانه‌های تشخیص و جلوگیری از نفوذ (IDS/IPS)، هوش تهدید (Threat Intelligence) و معماری Zero Trust، می‌تواند نقش مؤثری در افزایش تاب‌آوری سایبری سازمان‌ها داشته باشد. ترکیب این فناوری‌ها با مدل‌های هوش مصنوعی موجب افزایش دقت شناسایی تهدیدات ناشناخته، کاهش هشدارهای نادرست (False Positives) و بهبود عملکرد سامانه‌های امنیتی می‌شود. در این چارچوب، UEBA با تمرکز بر رفتارهای غیرعادی کاربران و موجودیت‌ها (مانند دسترسی‌های غیرمجاز، حرکت جانبی یا استخراج داده)، مکمل قدرتمندی برای IDS/IPS سنتی محسوب می‌شود و معماری Zero Trust نیز با اصل «هرگز اعتماد نکن، همیشه تأیید کن» زمینه را برای تصمیم‌گیری‌های مبتنی بر ریسک فراهم می‌سازد.

با وجود قابلیت‌های گسترده هوش مصنوعی، چالش‌هایی مانند کیفیت و تعادل داده‌های آموزشی، حملات خصمانه (Adversarial Attacks) علیه مدل‌های یادگیری ماشین، دشواری تفسیر تصمیمات مدل‌های جعبه سیاه (Black-Box) و مسائل مربوط به حفظ حریم خصوصی همچنان از موانع مهم توسعه این سامانه‌ها محسوب می‌شوند. کیفیت پایین داده‌ها

می تواند منجر به سوگیری مدل و افزایش هشدارهای کاذب شود؛ حملات خصمانه ممکن است با تغییرات جزئی در ورودی ها، مدل را فریب دهند؛ و عدم شفافیت تصمیمات، اعتماد تحلیل گران امنیتی و انطباق با مقررات را کاهش می دهد. بنابراین، توسعه مدل های توضیح پذیر (Explainable AI یا XAI)، مقاومت سازی الگوریتم ها در برابر حملات خصمانه، بهبود کیفیت و تنوع داده ها، استفاده از تکنیک های حفظ حریم خصوصی (مانند یادگیری فدرال) و رعایت اصول اخلاقی در استفاده از داده های رفتاری، از الزامات اساسی در آینده این حوزه خواهد بود.

در مجموع، هوش مصنوعی را می توان یکی از مهم ترین فناوری های تحول آفرین در حوزه امنیت سایبری دانست که با توانایی تحلیل هوشمند رفتار کاربران، شناسایی ناهنجاری ها و مقابله با تهدیدات پیچیده، نقش مهمی در حفاظت از زیرساخت های دیجیتال ایفا می کند. انتظار می رود با توسعه مدل های پیشرفته هوش مصنوعی، افزایش همکاری میان مراکز علمی، صنعتی و امنیتی، و ادغام عمیق تر این فناوری با معماری هایی مانند Zero Trust و سامانه های خودکار پاسخ به حوادث (SOAR)، سامانه های امنیت سایبری آینده از دقت، سرعت، خودکار سازی و قابلیت اطمینان بیشتری برخوردار شوند و توانایی مقابله مؤثرتری با تهدیدات سایبری در حال تحول داشته باشند. این تحول نه تنها نیازمند پیشرفت فنی است، بلکه مستلزم توجه به جنبه های سازمانی، حقوقی و اخلاقی نیز می باشد تا بتوان از مزایای هوش مصنوعی به صورت پایدار و مسئولانه بهره برداری کرد.

۷- مقایسه با تحقیقات پیشین

یافته های این پژوهش با نتایج مطالعات پیشین در حوزه کاربرد هوش مصنوعی در امنیت سایبری همخوانی قابل توجهی دارد و در عین حال، با تمرکز بر ترکیب چندلایه فناوری ها و چالش های عملی، ابعاد جدیدی را به ادبیات موجود اضافه می کند. در ادامه، این مقایسه به صورت نظام مند و بر اساس محورهای اصلی پژوهش (کاستی روش های سنتی، قابلیت های هوش مصنوعی، نقش UEBA و فناوری های مکمل، و چالش های پیش رو) ارائه می شود.

۱. همخوانی در نقد روش های سنتی امنیت اطلاعات

تقریباً تمام مطالعات مرور شده، چه فارسی و چه انگلیسی، بر این نکته اجماع دارند که روش های سنتی مبتنی بر امضا، قوانین ایستا و سامانه های تشخیص نفوذ کلاسیک، در برابر تهدیدات نوظهور، حملات روز صفر و رفتارهای پیچیده مهاجمان ناکافی هستند.

بهرامی (۱۴۰۴) در مرور جامع خود تأکید کرده است که روش های سنتی به دلیل ناتوانی در شناسایی حملات پیچیده، چندلایه و ناشناخته، کارایی خود را تا حد زیادی از دست داده اند. این یافته دقیقاً با نتیجه گیری پژوهش حاضر همسو است. مطالعات فارسی دیگر نیز، از جمله «نقش های یادگیری ماشین در ارتقای امنیت سایبری» (۱۴۰۵) و مقالات مرتبط با تشخیص نفوذ در اینترنت اشیا (۱۴۰۲)، به همین کاستی اشاره کرده اند و نشان داده اند که سامانه های مبتنی بر قوانین ثابت نمی توانند با سرعت و پیچیدگی حملات امروزی همگام شوند.

در ادبیات انگلیسی نیز، مرور دامنه وسیع میسرا و همکاران^{۲۹} (۲۰۲۶) و مطالعات مشابه، همین محدودیت را به عنوان نقطه شروع بحث خود قرار داده اند و بر ضرورت گذار به رویکردهای مبتنی بر یادگیری ماشین و یادگیری عمیق

²⁹ Mishra et al.

تأکید کرده‌اند. بنابراین، بخش مربوط به ناکارآمدی روش‌های سنتی در پژوهش حاضر، کاملاً با اجماع موجود در ادبیات همخوانی دارد.

۲. تأیید قابلیت‌های هوش مصنوعی در تشخیص تهدیدات

نتایج پژوهش حاضر مبنی بر توانایی بالای الگوریتم‌های یادگیری ماشین و یادگیری عمیق در تحلیل رفتار کاربران، تشخیص ناهنجاری، شناسایی نفوذ، کشف بدافزار و پیش‌بینی تهدیدات، با یافته‌های مطالعات پیشین کاملاً سازگار است.

بهرامی (۱۴۰۴) نشان داده است که مدل‌های مبتنی بر هوش مصنوعی با استخراج الگوهای غیرخطی و تطبیقی، دقت تشخیص را افزایش و نرخ هشدارهای کاذب را کاهش می‌دهند. مقالات فارسی مرتبط با تشخیص نفوذ در شبکه‌های اینترنت اشیا نیز به برتری روش‌های ترکیبی مبتنی بر ناهنجاری و یادگیری عمیق اشاره کرده‌اند. در مطالعات انگلیسی، Fuentes و همکاران (۲۰۲۵) با استفاده از Deep Autoencoders در چارچوب UEBA، قابلیت تشخیص تهدیدات پیچیده و تهدیدات داخلی را با دقت بالا و قابلیت توضیح‌پذیری نشان داده‌اند. مرورهای سیستماتیک دیگر نیز تأیید کرده‌اند که مدل‌های یادگیری عمیق (مانند LSTM، CNN و Autoencoder) در شناسایی انحرافات رفتاری و تهدیدات صفر-روز عملکرد بهتری نسبت به روش‌های سنتی دارند.

پژوهش حاضر این یافته‌ها را تأیید می‌کند و علاوه بر آن، بر اهمیت تحلیل حجم بالای داده‌های امنیتی و استخراج الگوهای پنهان تأکید دارد؛ موضوعی که در بیشتر مطالعات پیشین نیز به‌عنوان نقطه قوت اصلی هوش مصنوعی مطرح شده است.

۳. نقش ترکیبی UEBA، IDS/IPS، Threat Intelligence و معماری Zero Trust

یکی از نقاط تمایز و تقویت‌کننده پژوهش حاضر، تأکید بر ترکیب چند فناوری مکمل است. مطالعات پیشین اغلب بر یکی از این اجزا تمرکز کرده‌اند:

- برخی تحقیقات (مانند Fuentes et al., 2025) صرفاً بر UEBA و مدل‌های عمیق تمرکز داشته‌اند و نشان داده‌اند که تحلیل رفتار کاربران و موجودیت‌ها می‌تواند تهدیدات پنهان را شناسایی کند.

- مطالعات دیگر بیشتر به بهبود IDS/IPS با هوش مصنوعی پرداخته‌اند و کاهش نرخ هشدارهای کاذب و افزایش دقت را گزارش کرده‌اند.

- معماری Zero Trust نیز در ادبیات اخیر به‌عنوان چارچوب مکمل مطرح شده، اما کمتر به‌صورت یکپارچه با UEBA و Threat Intelligence بررسی شده است.

پژوهش حاضر با پیشنهاد ترکیب این چهار عنصر (UEBA + IDS/IPS + Threat Intelligence + Zero Trust) در کنار مدل‌های هوش مصنوعی، گامی فراتر از مطالعات تک‌محوره برداشته است. این رویکرد یکپارچه، دقت شناسایی تهدیدات ناشناخته را افزایش می‌دهد، هشدارهای نادرست را کاهش می‌دهد و تاب‌آوری کلی سازمان را تقویت می‌کند؛ موضوعی که در ادبیات موجود به‌صورت پراکنده مطرح شده اما کمتر به‌صورت منسجم تحلیل شده است.

۴. همخوانی در شناسایی چالش‌ها و مسیرهای آینده

چالش‌های شناسایی شده در این پژوهش کیفیت داده‌های آموزشی، حملات خصمانه علیه مدل‌ها، دشواری تفسیر تصمیمات (جعبه سیاه بودن مدل‌ها) و مسائل حریم خصوصی دقیقاً با اجماع ادبیات پیشین همخوانی دارد.

۸- پیشنهادها

با توجه به نتایج حاصل از این پژوهش و نقش روزافزون فناوری‌های هوش مصنوعی در ارتقای امنیت سایبری، پیشنهاد می‌شود سازمان‌ها به منظور افزایش سطح حفاظت از زیرساخت‌های دیجیتال، سامانه‌های مبتنی بر هوش مصنوعی را در کنار راهکارهای سنتی امنیت اطلاعات به کار گیرند. استفاده ترکیبی از روش‌های هوشمند و سنتی می‌تواند موجب بهبود قابلیت شناسایی تهدیدات ناشناخته، تحلیل دقیق‌تر رفتار کاربران، کاهش زمان واکنش به رخدادها و امنیتی و افزایش کارایی مراکز عملیات امنیت شود.

با توجه به توانایی بالای الگوریتم‌های یادگیری ماشین و یادگیری عمیق در پردازش حجم گسترده‌ای از داده‌های امنیتی، پیشنهاد می‌شود در طراحی سامانه‌های امنیتی هوشمند از مدل‌های پیشرفته‌ای مانند XGBoost, Random Forest, LSTM, CNN و Transformer استفاده شود. این مدل‌ها با قابلیت استخراج الگوهای پیچیده از داده‌ها، می‌توانند در حوزه‌هایی مانند تشخیص نفوذ، شناسایی بدافزارها، کشف حملات فیشینگ، تشخیص حملات روز صفر و پیش‌بینی رفتار مهاجمان عملکرد مؤثری داشته باشند.

یکی از عوامل کلیدی در موفقیت سامانه‌های مبتنی بر هوش مصنوعی، کیفیت و اعتبار داده‌های آموزشی است. بنابراین پیشنهاد می‌شود سازمان‌ها با ایجاد پایگاه‌های داده استاندارد، جمع‌آوری داده‌های متنوع و به‌روز، پاک‌سازی داده‌های نامعتبر و مدیریت صحیح چرخه عمر داده‌ها، زمینه آموزش دقیق‌تر و عملکرد قابل اعتمادتر مدل‌های هوشمند را فراهم کنند. همچنین توجه به مسئله عدم توازن داده‌ها و کاهش خطاهای ناشی از داده‌های ناقص می‌تواند نقش مهمی در افزایش دقت سامانه‌های تشخیص تهدید داشته باشد.

با توجه به افزایش حملات خصمانه علیه مدل‌های یادگیری ماشین، توصیه می‌شود پژوهشگران و توسعه‌دهندگان سامانه‌های امنیتی، علاوه بر افزایش دقت مدل‌ها، بر موضوعاتی مانند امنیت مدل‌های هوش مصنوعی، مقاومت در برابر حملات Adversarial Attacks و توسعه روش‌های هوش مصنوعی توضیح‌پذیر (Explainable AI) تمرکز کنند. استفاده از مدل‌های قابل تفسیر می‌تواند اعتماد متخصصان امنیتی به تصمیمات سامانه‌های هوشمند را افزایش داده و امکان بررسی علت صدور هشدارهای امنیتی را فراهم سازد.

همچنین پیشنهاد می‌شود سازمان‌ها فناوری تحلیل رفتار کاربران و موجودیت‌ها (User and Entity Behavior Analytics - UEBA) را در مراکز عملیات امنیت (Security Operations Center - SOC) به کار گیرند. این فناوری با ایجاد الگوهای رفتاری عادی برای کاربران، دستگاه‌ها و سایر موجودیت‌های شبکه، امکان شناسایی سریع‌تر رفتارهای غیرعادی، تهدیدات داخلی، حساب‌های کاربری آسیب‌دیده و فعالیت‌های مشکوک را فراهم می‌کند و می‌تواند مکمل مناسبی برای سامانه‌های سنتی تشخیص تهدید باشد.

از سوی دیگر، پیشنهاد می‌شود معماری امنیتی Zero Trust در کنار فناوری‌های هوش مصنوعی مورد توجه سازمان‌ها قرار گیرد. این رویکرد با اصل «اعتماد نکن، همیشه اعتبارسنجی کن» و بررسی مداوم هویت کاربران، وضعیت

دستگاه‌ها و درخواست‌های دسترسی، می‌تواند سطح مقاومت سازمان‌ها را در برابر نفوذهای سایبری افزایش دهد. ترکیب Zero Trust با تحلیل‌های هوشمند رفتاری، امکان شناسایی سریع‌تر فعالیت‌های غیرمجاز و کاهش احتمال سوءاستفاده از دسترسی‌های قانونی را فراهم خواهد کرد.

در نهایت، با توجه به سرعت بالای پیشرفت فناوری‌های هوش مصنوعی، پیشنهاد می‌شود پژوهش‌های آینده بر توسعه مدل‌های ترکیبی هوش مصنوعی، استفاده از یادگیری فدرال (Federated Learning)، تحلیل داده‌های چندمنبعی، بهره‌گیری از هوش مصنوعی مولد (Generative AI) در امنیت سایبری و طراحی سامانه‌های خودسازگار تمرکز کنند. این رویکردها می‌توانند نقش مهمی در افزایش توانایی سامانه‌های امنیتی برای مقابله با تهدیدات پیچیده و نوظهور آینده ایفا کنند.

منابع

احمدزاده، ب. (۱۴۰۴). استفاده از یادگیری ماشین و یادگیری عمیق برای شناسایی تهدیدات و حملات سایبری: روندها و تکنیک‌ها. بیست و سومین کنفرانس ملی پژوهش‌های کاربردی در علوم برق، کامپیوتر و مهندسی پزشکی. [/https://civilica.com/doc/2287078](https://civilica.com/doc/2287078)

اندایش، س.، و کیان‌راد، ز. (۱۴۰۴). تحلیل علم‌سنجی هوش مصنوعی در امنیت سایبری. کنفرانس بین‌المللی هوش مصنوعی و فناوری‌های مرتبط. [/https://civilica.com/doc/2447729](https://civilica.com/doc/2447729)

بهرامی، س. (۱۴۰۴). هوش مصنوعی و امنیت وب در تشخیص حملات سایبری و حفاظت از حریم خصوصی. ششمین همایش بین‌المللی مهندسی کامپیوتر، برق و تکنولوژی. [/https://civilica.com/doc/2585676](https://civilica.com/doc/2585676)

جلالی، ش. (۱۴۰۴). هوش مصنوعی در امنیت سایبری: حفاظت از داده‌ها و سیستم‌ها. [/https://civilica.com/doc/2240847](https://civilica.com/doc/2240847)

روا، ب. (۱۴۰۴). هوش مصنوعی و امنیت سایبری. بیستمین کنفرانس ملی حقوق، علوم اجتماعی و انسانی، روانشناسی و مشاوره. [/https://civilica.com/doc/2306082](https://civilica.com/doc/2306082)

عباس‌نژادورزی، ر.، رهی، م.، و عباس‌نژادورزی، س. (۱۴۰۳). هوش مصنوعی در امنیت سایبری: مفاهیم، کاربردها، چالش‌ها، الگوریتم‌ها و ابزارها. انتشارات فناوری نوین.

قدیم‌پور شبلو، ع.، و ابری، ی. (۱۴۰۳). نقش هوش مصنوعی در امنیت سایبری: فرصت‌ها، چالش‌ها و راهکارهای آینده‌نگر. سومین کنفرانس فضای سایبر. [/https://civilica.com/doc/2384068](https://civilica.com/doc/2384068)

کشاوری، ز.، و حسینی، ح. (۱۴۰۲). هوش مصنوعی در امنیت سایبری (کاربردها، چالش‌ها و فرصت‌ها). ششمین همایش ملی فناوری‌های نوین در مهندسی برق، کامپیوتر و مکانیک ایران. [/https://civilica.com/doc/1744302](https://civilica.com/doc/1744302)

گلدی نجف‌آباد، م. (۱۴۰۴). هوش مصنوعی و امنیت سایبری. نهمین همایش بین‌المللی دانش و فناوری مهندسی برق، کامپیوتر و مکانیک ایران. [/https://civilica.com/doc/2299777](https://civilica.com/doc/2299777)

رضایی، م.، و احمدی، س. (۱۴۰۲). تحلیل رفتار کاربران در بانکداری الکترونیک به منظور پیشگیری از کلاهبرداری با رویکرد هوش مصنوعی. نشریه علمی امنیت اطلاعات و سیستم‌های دیجیتال، ۱۰(۲)، ۴۵-۵۸.

محمدی، ع.، و عباسی، ح. (۱۴۰۳). نقش سیستم‌های نوین تشخیص نفوذ مبتنی بر یادگیری عمیق در امنیت شبکه‌های داده صنعتی. پژوهش‌نامه پدافند غیرعامل و امنیت سایبری، ۱۲(۱)، ۸۹-۱۰۴.

یوسفی، ر.، مرادی، الف.، و کریمی، ف. (۱۴۰۱). شناسایی حملات سایبری و تهدیدات پیشرفته با استفاده از الگوریتم‌های یادگیری ماشین در شبکه‌های سازمانی. فصلنامه فناوری اطلاعات و ارتباطات انتظامی، ۹(۳۴)، ۱۲-۲۹.

هاشمی شبیری، ه. (۱۴۰۳). مطالعه مبانی هوش مصنوعی و بررسی راهکارهای ارتقای امنیت سایبری در شبکه مبتنی بر AI. هفتمین همایش ملی فناوری‌های نوین در مهندسی برق، کامپیوتر و مکانیک ایران. <https://civilica.com/doc/2050314>

(مقاله بدون نام نویسنده مشخص). (۱۴۰۳). کاربست هوش مصنوعی در تجزیه و تحلیل رفتار کاربر و موجودیت. <https://civilica.com/doc/2133590>

(مقاله بدون نام نویسنده مشخص). (۱۴۰۳). تأثیر هوش مصنوعی مولد در امنیت سایبری و حریم خصوصی. هفتمین کنفرانس بین‌المللی مهندسی برق، کامپیوتر، مکانیک و هوش مصنوعی. <https://civilica.com/doc/2046535>

Abdel-Basset, M., Hawash, H., Chakraborty, R. K., & Ryan, M. J. (2022). Deep learning for cyber security applications: A comprehensive review. *Journal of Information Security and Applications*, 70, 103361.

Al-Mansoori, S., Al-Qurran, R., & Abu-Nasser, A. (2022). User and Entity Behavior Analytics (UEBA) Using Machine Learning Techniques for Cyber Threat Detection. *Journal of Cybersecurity and Information Management*, 18(4), 112-127.

Alzahrani, A., Alghazzawi, D., Cheng, L., Alzahrani, S., & Al-Barakati, A. (2023). Deep learning approaches for intrusion detection systems: A systematic review. *Electronics*, 12(4), 890.

Apruzzese, G., Colajanni, M., Ferretti, L., Guido, A., & Marchetti, M. (2023). On the effectiveness of machine learning for cybersecurity. *Computers & Security*, 124, 102957.

Buczak, A. L., & Guven, E. (2016). A survey of data mining and machine learning methods for cyber security intrusion detection. *IEEE Communications Surveys & Tutorials*, 18(2), 1153-1176.

Chandola, V., Banerjee, A., & Kumar, V. (2009). Anomaly detection: A survey. *ACM Computing Surveys*, 41(3), 1-58.

ENISA. (2023). ENISA Threat Landscape 2023. European Union Agency for Cybersecurity.

Ferrag, M. A., Shu, L., Yang, X., Derhab, A., & Maglaras, L. (2021). Security and privacy for IoT-based smart healthcare systems using machine learning: A survey. *IEEE Internet of Things Journal*, 8(15), 11931-11951.

Gartner. (2018). Market Guide for User and Entity Behavior Analytics. Gartner Research.

Goodfellow, I., Bengio, Y., & Courville, A. (2016). Deep Learning. MIT Press.

- Google Cloud. (2025). Cybersecurity Forecast 2025. Google Cloud Security.
- IBM Security. (2023). Cost of a Data Breach Report 2023. IBM Corporation.
- Khraisat, A., & Alazab, A. (2021). A critical review of intrusion detection systems in the era of artificial intelligence. *Cybersecurity*, 4, 1–18.
- Khraisat, A., Gondal, I., Vamplew, P., & Kamruzzaman, J. (2023). Survey of intrusion detection systems: Techniques, datasets and challenges. *Cybersecurity*, 6(1), 1–29.
- Microsoft. (2024). Microsoft Digital Defense Report 2024. Microsoft Corporation.
- Mitchell, T. M. (1997). *Machine Learning*. McGraw-Hill.
- National Institute of Standards and Technology. (2023). Implementing a Zero Trust Architecture (SP 1800-35). NIST.
- National Institute of Standards and Technology. (2024). Adversarial Machine Learning: A Taxonomy and Terminology of Attacks and Mitigations (NIST AI 100-2e2024). NIST.
- Russell, S., & Norvig, P. (2021). *Artificial Intelligence: A Modern Approach* (4th ed.). Pearson.
- Sarker, I. H. (2022). AI-enabled cybersecurity: Challenges and opportunities. *SN Computer Science*, 3(3), 1–15.
- Sharma, A., & Rattan, D. (2023). Artificial intelligence in cybersecurity: A systematic literature review. *IEEE Access*, 11, 86545–86576.
- Sommer, R., & Paxson, V. (2010). Outside the closed world: On using machine learning for network intrusion detection. *IEEE Symposium on Security and Privacy*, 305–316.
- Stallings, W. (2019). *Effective Cybersecurity: A Guide to Using Best Practices and Standards*. Addison-Wesley.
- Thompson, K., & Davies, L. (2024). Harnessing Artificial Intelligence for Proactive Threat Hunting and User Profiling in Cloud Computing. *Computers & Security*, 136, 103521.
- Wang, W., Li, Y., Wang, X., Liu, J., & Zhang, Y. (2024). Artificial intelligence techniques for user behavior analytics in cybersecurity: A comprehensive review. *IEEE Access*, 12, 42851–42879.
- Whitman, M. E., & Mattord, H. J. (2021). *Principles of Information Security* (7th ed.). Cengage.
- Zhang, H., & Wang, Y. (2023). An Intelligent Framework for Insider Threat Detection Based on Deep Learning and Behavioral Analysis in Digital Environments. *IEEE Transactions on Information Forensics and Security*, 18, 1542–1555.

Abstract

The rapid growth of digital technologies and increasing dependence on information infrastructures have led to the emergence of more complex cyber threats and the need for intelligent approaches to protect digital environments. In this context, artificial intelligence has emerged as a powerful technology with significant capabilities in data analysis, behavioral pattern recognition, and cyber threat detection. The aim of this study is to investigate the applications of artificial intelligence methods in user behavior analysis and cyber threat detection within digital environments and to explain the role of these technologies in enhancing cybersecurity. This research is applied in terms of purpose and descriptive-review in terms of methodology. The required data

were collected and analyzed through the examination of scientific articles, books, and specialized reports related to artificial intelligence, cybersecurity, machine learning, and user behavior analytics .The findings indicate that machine learning algorithms, deep learning approaches, and user behavior analysis techniques have significant capabilities in anomaly detection, identifying unknown attacks, malware detection, insider threat recognition, and predicting security incidents. Furthermore, integrating artificial intelligence models with technologies such as User and Entity Behavior Analytics (UEBA), intrusion detection systems, threat intelligence, and Zero Trust architecture can improve threat detection accuracy, reduce false alarms, and enhance security response speed .Despite its advantages, the implementation of artificial intelligence-based cybersecurity solutions faces challenges, including training data quality, adversarial attacks against intelligent models, model interpretability, and privacy preservation. The results demonstrate that the effective utilization of artificial intelligence combined with appropriate security strategies can play a significant role in developing intelligent, proactive, and resilient cybersecurity systems against emerging threats.

Keywords: Artificial Intelligence, Cybersecurity, User Behavior Analysis, Machine Learning, Deep Learning, Anomaly Detection, UEBA, Cyber Threats.