



Using Artificial Intelligence to Deal with Computer Crimes and Improve Cyber Security Performance

Mohammad Hadi Kaveh^{*1}, Mohammad Barani²

1. M.A of Criminal Law and Criminology, Damghan Branch, Islamic Azad University, Damghan, Iran. (Corresponding Author)

2. Associate Professor, Department of Criminology, Amin University of Police Sciences, Tehran, Iran.

ARTICLE INFORMATION

Type of Article:

Original Research

Pages: 193-205

Corresponding Author's Info

ORCID: 0009-0005-6533-3839

TELL: +982335225041

Email: Hadikaveh895@yahoo.com

Article history:

Received: 03 Oct 2024

Revised: 06 Dec 2024

Accepted: 07 Feb 2025

Published online: 21 Mar 2025

Keywords:

Artificial Intelligence, Computer Crimes, Cyber Security, Phishing, Network Security.

ABSTRACT

Computer crimes are one of the important issues of criminal law, and therefore it is very important to deal with them. On the other hand, artificial intelligence has greatly changed the virtual space and technological innovations in the field of computers. The purpose of this article is to examine the question of how artificial intelligence is effective in dealing with computer crimes and improving cyber security performance. The findings indicate that considering human limitations and the fact that computer viruses and worms have become intelligent and flexible, the design of agents with intelligent sensors to deal with computer crimes can be considered. One of the fields that can be used to improve cyber security is artificial intelligence. Artificial intelligence is effective in dealing with computer crimes by helping to identify and detect crime and at the same time prevent it. Also, data security and network security through intrusion detection and botnet detection, e-mail security through spam detection and phishing attack detection, dealing with fake accounts, user data and information protection, authentication and detection of authentication fraud, program security and monitoring User activity is one of the most important measures of artificial intelligence in improving cyber security performance. Of course, artificial intelligence in improving cyber security faces some limitations, such as the lack of transparency in the decision-making process of artificial intelligence, the need for high data and the necessity of human participation and the interference of confusing variables.



This is an open access article under the CC BY license.

© 2025 The Authors.

How to Cite This Article: Kaveh, MH & Barani, M (2025). "Using Artificial Intelligence to Deal with Computer Crimes and Improve Cyber Security Performance". *Journal of Comparative Criminal Jurisprudence*, 5(1): 193-205.



انجمن علمی فقه‌پژای تطبیقی ایران

فصلنامه فقه‌پژای تطبیقی

www.jccj.ir



فصلنامه فقه‌پژای تطبیقی

دوره پنجم، شماره اول، بهار ۱۴۰۴

به کارگیری هوش مصنوعی برای مقابله با جرایم رایانه‌ای و بهبود عملکرد امنیت سایبری

محمدهادی کاوه^{۱*}، محمد بارانی^۲

۱. کارشناس ارشد گروه حقوق جزا و جرم‌شناسی، واحد دامغان، دانشگاه آزاد اسلامی، دامغان، ایران. (نویسنده مسؤول)

۲. دانشیار، گروه جرم‌شناسی، دانشگاه علوم انتظامی امین، تهران، ایران.

چکیده

جرایم رایانه‌ای، از موضوعات مهم حقوق کیفری بوده، لذا مقابله با آن از اهمیت به‌سزایی برخوردار است. از طرفی هوش مصنوعی به‌شدت فضای مجازی و نوآوری‌های فناوری در زمینه کامپیوتر را دستخوش تغییر کرده است. هدف مقاله حاضر بررسی این سؤال است که هوش مصنوعی در مقابله با جرایم رایانه‌ای و بهبود عملکرد امنیت سایبری چگونه تأثیرگذار است. یافته‌ها بر این امر دلالت دارد که با توجه به محدودیت‌های انسانی و این حقیقت که ویروس‌ها و کرم‌های رایانه‌ای هوشمند و انعطاف‌پذیر شده‌اند، طراحی عامل‌های با سنسور هوشمند برای مقابله با جرایم رایانه‌ای می‌تواند مورد توجه قرار گیرد. یکی از حوزه‌هایی که می‌توان با کمک آن امنیت سایبری را بالا برد، هوش مصنوعی است. هوش مصنوعی از طریق کمک به شناسایی و کشف جرم و در عین حال پیشگیری از جرم، در مقابله با جرایم رایانه‌ای تأثیرگذار است. همچنین امنیت داده و امنیت شبکه از طریق تشخیص تهاجم و شناسایی بات‌نت‌ها، امنیت ایمیل‌ها از طریق شناسایی هرزنامه و شناسایی حملات فیشینگ، مقابله با حساب‌های جعلی، حفاظت از اطلاعات و داده‌های کاربران، احراز هویت و تشخیص فریب احراز هویت، امنیت برنامه و نظارت بر فعالیت کاربران از مهم‌ترین اقدامات هوش مصنوعی در بهبود عملکرد امنیت سایبری است، البته هوش مصنوعی در بهبود امنیت سایبری با برخی محدودیت‌ها مانند عدم شفافیت در فرآیند تصمیم‌گیری هوش مصنوعی، نیاز به داده‌های بالا و ضرورت مشارکت انسانی و تداخل متغیرهای گنج‌کننده مواجه است.

اطلاعات مقاله

نوع مقاله: پژوهشی

صفحات: ۱۹۳-۲۰۵

اطلاعات نویسنده مسؤول

کد ارکید: ۳۸۳۹-۶۵۳۳-۰۰۵-۰۰۹-۰۰۰

تلفن: +۹۸۲۳۳۵۲۲۵۰۴۱

ایمیل: Hadikaveh895@yahoo.com

سابقه مقاله:

تاریخ دریافت: ۱۴۰۳/۰۷/۱۲

تاریخ ویرایش: ۱۴۰۳/۰۹/۱۶

تاریخ پذیرش: ۱۴۰۳/۱۱/۱۹

تاریخ انتشار: ۱۴۰۴/۰۱/۰۱

واژگان کلیدی:

هوش مصنوعی، جرایم رایانه‌ای، امنیت سایبری، فیشینگ، امنیت شبکه.



مقدمه

تکامل سریع فناوری‌های اطلاعات و ارتباطات، از جمله اینترنت، پیامدهای مثبتی را برای سازمان‌ها و زندگی اجتماعی به همراه داشته است، اما امنیت داده‌ها در فضای مجازی از چالش‌های مهمی است که باعث می‌شود که پیگیری جرایمی مانند سرقت مجازی و کلاهبرداری اینترنتی در فضای مجازی پیچیده‌تر شود (Calzavara et al, 2015: 30). نگران‌کننده‌تر آن که در فضای مجازی، همه‌جا اهمیت مرزهای جغرافیایی را از بین می‌برد و این امر همچنین منشأ فعالیت‌های مجرمانه را از هر قسمت از جهان فراهم می‌کند. از این رو، طیف گسترده‌ای از حملات سایبری سازمان‌ها را به طور فزاینده‌ای به چالش می‌کشد. این حملات با سطح بالایی از مهارت شناخته می‌شود که نیاز به استفاده از هوش مصنوعی یا عوامل هوشمند برای مبارزه با آن‌ها دارد (Kleinmann & Wool, 2016: 1-21). همان‌طور که جامعه ما بیشتر به هم متصل و از نظر فناوری پیشرفته‌تر می‌شود، نقش راه حل‌های امنیتی اهمیت بیشتری خواهد داشت. با این حال چالش ایمن‌سازی سیستم‌ها و جامعه ما که به این سیستم‌ها متکی است، با تهدید مواجه است. از این رو طراحی راه حل‌های امنیت سایبری کارآمدتر و مؤثرتر موضوعی است که همواره مورد علاقه است (Xiao et al, 2018; Guan & X.Ge, 2017: 48-59).

از طرف دیگر، باتوجه به حجم حملات سایبری، مداخله انسان در برابر این حملات برای دادن عکس‌العمل مناسب، کافی نیست، زیرا بیشتر این حملات، با استفاده از عامل‌های هوشمند مانند کرم‌ها و ویروس‌های کامپیوتری، صورت می‌گیرد، لذا مقابله با آن‌ها، نیاز به عامل‌های هوشمند دارد. این عامل‌های هوشمند وظیفه دارند: کل فرایند حمله را در زمان مناسب مدیریت کنند، نوع حمله را شناسایی کنند، کشف کنند که هدف حمله چه بخشی از شبکه می‌باشد، مناسب‌ترین پاسخ به حمله، چه می‌باشد و چگونه جلوی حملات مشابه آن را بگیرند. علاوه بر این، مشکل جرایم رایانه‌ای، یک معضل جهانی شده است. در گذشته نه‌چندان دور، فقط افراد تحصیل کرده که با کامپیوتر و اینترنت سروکار داشتند، در معرض این حملات بودند، اما امروزه با فراگیر شدن استفاده از کامپیوتر و اینترنت، هرکسی می‌تواند مورد حمله سایبری قرار بگیرد. متأسفانه الگوریتم‌ها و

برنامه‌های کامپیوتری که قبلاً برای شناسایی این حملات، مورد استفاده قرار می‌گرفتند، دیگر کارا نیستند، زیرا ویروس‌ها بسیار هوشمند شده‌اند و به محض ورود به یک سیستم، خود را با شرایط آن سیستم وفق می‌دهند و پیدا کردن آن‌ها ممکن نیست. ما نیاز به استفاده از روش‌های یادگیری ماشین و هوش مصنوعی برای مقابله با این حملات داریم که هم انعطاف‌پذیر هستند و هم به سرعت با شرایط هر محیطی خود را سازگار می‌کنند. براساس آن‌چه گفته شد، سؤال مقاله بدین شکل قابل طرح است که هوش مصنوعی چه نقشی می‌تواند در مقابله با جرایم رایانه‌ای و بهبود عملکرد امنیت سایبری ایفا نماید؟

این مقاله توصیفی - تحلیلی است و با استفاده از روش کتابخانه‌ای به بررسی سؤال مورد اشاره پرداخته است. به منظور بررسی سؤال مورد اشاره ابتدا نقش هوش مصنوعی در مقابله با جرایم رایانه‌ای بررسی شده و در ادامه از نقش آن در تحقق امنیت سایبری بحث می‌شود. در نهایت به محدودیت‌های هوش مصنوعی در بهبود امنیت سایبری پرداخته شده است.

۱- به کارگیری هوش مصنوعی برای مقابله با جرایم رایانه‌ای

توسعه سیستم‌های کامپیوتری، از یک طرف تأثیر مثبت قابل توجهی روی بهبود کیفیت زندگی بشر داشته، از طرف دیگر مشکلات جدیدی مانند جرایم پیشرفته را با خود به ارمغان آورده است. جرایم معمولی مانند دزدی و کلاهبرداری، با به کارگیری فناوری اطلاعات، شکل جدیدی به خود گرفته‌اند که به آن «جرایم رایانه‌ای» می‌گویند که با پیشرفت فناوری، این جرایم نیز هر روز پیشرفته‌تر می‌شوند. از طرف دیگر، روند جهانی‌سازی و کم‌رنگ شدن نقش مرز بین کشورها، باعث شده که کشف، رصد، پیگیری و ممانعت از رخداد آن‌ها، هر روز سخت‌تر شود (Gordon & Ford, 2006: 13-20). اکثر جرایم سایبری که امروزه صورت می‌گیرند، به سادگی نشان می‌دهند که جرایم، از فضای حقیقی به فضای مجازی، مهاجرت کرده‌اند. به عبارت دیگر جرایم قدیمی با شیوه جدیدی رخ می‌دهند (Brenner, 2010). در این قسمت به بررسی به کارگیری هوش مصنوعی برای مقابله با جرایم رایانه‌ای پرداخته می‌شود.

۱-۱- نقش هوش مصنوعی در کشف جرایم رایانه‌ای

۱-۲- نقش هوش مصنوعی در پیشگیری از جرایم رایانه‌ای
سیستم‌های ایمنی مصنوعی مانند سیستم ایمنی بیولوژیکی، برای نگهداشتن پایداری سیستم در یک محیط متغیر، طراحی شده‌اند. در سال ۲۰۰۷ یک سیستم مبتنی بر سیستم‌های ایمنی مصنوعی برای کشف هرزنامه‌ها، طراحی گردید (Sirisanyalak & Sornil, 2007: 3392-3398). یکی از روش‌های انتشار ویروس‌ها و کرم‌های کامپیوتری، استفاده از هرزنامه و سوءاستفاده از اعتماد کاربران به ایمیل‌ها می‌باشد. در پروژه دیگری، با بررسی مدل‌های مختلف سیستم‌های ایمنی مصنوعی تئوری خطر برای اولین بار مطرح گردید. در این پروژه از نقشه‌های خودسازمان‌ده برای پاسخگویی به خطرات در شبکه بی‌سیم استفاده گردید (Lebbe et al, 2007: 322-327).

در سال ۲۰۱۲ یک سیستم کشف نفوذ مبتنی بر الگوریتم ژنتیک قانون‌گرا، طراحی شد. این سیستم برای بهبود امنیت، اطمینان، تجمیع و در دسترس بودن منابع شبکه، ارائه گردید. سیستم ارائه‌شده، از مجموعه‌ای از قوانین طبقه‌بندی که از داده‌های شبکه به‌دست آمده‌اند، تشکیل شده است و از درجه اطمینان شبکه به‌عنوان تابع پرارزش به‌منظور ارزیابی هر قانون استفاده می‌کند (Ojugo et al, 2012: 1182-1194). از قوانین فازی به‌منظور طبقه‌بندی انواع حملات استفاده می‌شود و الگوریتم ژنتیک کمک می‌کند تا قوانین فازی مناسب پیدا شوند.

۲- به کارگیری هوش مصنوعی برای بهبود عملکرد امنیت سایبری

کاربردهای هوش مصنوعی برای امنیت سایبری به‌طور کلی موفقیت‌آمیز بوده است. در مقابله با جرایم اینترنتی کمک‌های قابل توجهی انجام شده است. عمدتاً موضوعات مرتبط با سیستم‌های تشخیص و پیشگیری از نفوذ، سیستم‌های جدید بهبودهایی را نسبت به سیستم‌های قبلی نشان داده‌اند (Tsai et al, 2009: 119-120).

فناوری‌های متداول پیشگیری امنیت سایبری از الگوریتم‌های ثابت و دستگاه‌های فیزیکی (مانند حسگرها و ردیاب‌ها) استفاده می‌کنند، بنابراین در مهار تهدیدهای جدید فضای مجازی بی‌اثر هستند. برای مثال، اولین نسل از سیستم‌های آنتی‌ویروس

هوش مصنوعی به‌عنوان یک علم جدید، برای اولین بار در سال ۱۹۵۶ ظاهر شد. هدف از آن، ایجاد ماشین‌های هوشمند و حل مسائلی بود که بدون استفاده از درجاتی از هوشمندی، قابل حل نبودند. به‌عنوان مثال برای مقابله با جرایم رایانه‌ای، ما به سیستمی نیاز داریم که تا حدودی هوشمند باشد و بتواند از روی مقدار زیادی اطلاعات، تصمیمات درستی بگیرد. در مبحث هوش مصنوعی، با عامل سروکار داریم که تعریف آن به‌صورت یک «موجود خودکار» که محیط اطراف خود را درک می‌کند، است. عامل دارای یک سیستم خود تصمیم‌گیرنده درونی است و عملی که انجام می‌دهد دارای تأثیر روی محیط اطراف است. هوش مصنوعی از تکنولوژی‌های متعدد برای حل مسائل استفاده می‌کند، مانند شبکه عصبی مصنوعی، منطق فازی محاسبات تکاملی، یادگیری ماشین، سیستم‌های ایمنی مصنوعی و کرم‌های مصنوعی. این فناوری‌ها توانایی تصمیم‌گیری انعطاف‌پذیر برای محیط‌های پویا مانند دایره امنیت و جرم را دارند. اکثر این روش‌ها، با الهام از طبیعت، از سیستم‌های بیولوژیکی به‌خوبی تقلید می‌کنند و مانند آن‌ها توانایی یادگیری، به‌خاطر سپردن، شناخت، طبقه‌بندی و پردازش اطلاعات را دارند.

در سال ۲۰۰۶ یک سیستم چندعاملی برای کشف کرم‌های کامپیوتری در محیط اینترنت شهری، طراحی گردید. در این محیط کرم‌های زیادی منتشر می‌شوند که پهنای باند زیادی از شبکه را به هدر می‌دهند و باعث از کارافتادن مسیرپای‌ها می‌شوند (Gou et al, 2006: 259-265).

همچنین در سال ۲۰۰۶ سیستم دیگری براساس سیستم‌های چندعاملی طراحی شد که برای مقابله با حملات توزیعی DoS بسیار مناسب بود. این سیستم با استفاده از زبان پرولوگ پیاده‌سازی گردید و بدون دخالت انسان، عملکرد مناسبی داشت (Phillips et al, 2006). در سال ۲۰۰۷ چهارچوب یک سیستم مقابله با حملات سایبری طراحی شد که در آن تعداد از عامل‌های هوشمند در تعامل با یکدیگر هستند و به مرور زمان، باتوجه به شرایط شبکه و شدت و سختی حملات، ساختار، پیکربندی و رفتار خود را تغییر می‌دهند (Kotenko & Ulanov, 2007: 212-228).

بررسی به کارگیری هوش مصنوعی برای بهبود عملکرد امنیت سایبری پرداخته می‌شود.

۲-۱- امنیت داده و شبکه

سازمان‌ها برای به حداقل رساندن صدمات ناشی از حملات سایبری و جلوگیری از وقوع حوادث احتمالی و همچنین حفاظت از داده‌ها از فناوری‌های نوینی مانند هوش مصنوعی استفاده می‌کنند. ابزارهای هوش مصنوعی می‌توانند داده‌ها را مورد بررسی قرار داده و داده‌های ناهنجار را که از دید انسان پنهان مانده است، تشخیص دهند. یکی از کارکردهای هوش مصنوعی در زمینه امنیت داده شناسایی و جلوگیری از نشت داده می‌باشد، نشت داده، ممکن است از داخل و یا خارج از سازمان اتفاق افتد که در بسیاری از مواقع آسیب جدی به اطلاعات و دانش سازمان‌ها وارد می‌نماید، این امر می‌تواند به دلیل اتصال کاربران غیرمجاز به سیستم، سطح دسترسی‌های نادرست به اطلاعات و یا معماری اشتباه شبکه باشد، سیستم‌های مجهز به هوش مصنوعی با شناسایی نظارت و مراقبت از اطلاعات موجود در حافظه و همچنین مراقبت از داده‌هایی که در شبکه‌ها در حال جابه‌جایی هستند، امنیت داده‌ها را افزایش می‌دهد و در صورت وجود رفتار غیرعادی به‌طور خودکار هشدار داده و از ادامه آن رفتار جلوگیری می‌کند (Kowert, 2017: 181).

با گسترده‌شدن شبکه‌های کامپیوتری ایجاد امنیت در شبکه از اهمیت به‌سزایی برخوردار است. متأسفانه افراد سودجو با دسترسی به اطلاعات مهم مراکز خاص و یا اطلاعات افراد دیگر به قصد اعمال نفوذ فشار و یا ایجاد بی‌نظمی در سیستم‌ها، حمله به شبکه و سیستم‌ها را در پیش می‌گیرند. دو مدل یادگیری با ناظر و یادگیری بدون ناظر در حوزه هوش مصنوعی توانستند تأثیر به‌سزایی در تأمین امنیت شبکه بگذارند. مدل‌های یادگیری ماشین با ناظر براساس داده‌های برچسب‌دار، آموزش می‌بینند که مسائل طبقه‌بندی نیز در این گروه قرار می‌گیرند، در این مسائل داده‌های ورودی براساس الگوی نهفته در آن‌ها در یک دسته مشخص جای می‌گیرند، اما در یادگیری بدون ناظر به توسعه سیستم‌هایی پرداخته می‌شود که تنها با

برای شناسایی ویروس‌ها از طریق اسکن کردن بیت‌آن^۱ طراحی شدند. فرض اساسی این مفهوم این است که یک ویروس همان ساختار و الگوی بیت را در همه نمونه‌ها دارد. بنابراین این امضاها و الگوریتم‌ها ثابت هستند، اگرچه کاتالوگ امضاها به‌صورت روزانه به‌روز می‌شود (یا هر زمان که دستگاه به اینترنت متصل باشد)، پیچیدگی و انتشار منظم بدافزارهای گسترده این روش را بی‌اثر می‌کند. با این حال، معرفی روش‌های بدون امضا که قادر به شناسایی و کاهش حملات بدافزار با استفاده از روش‌های جدیدتر مانند تشخیص رفتاری و هوش مصنوعی هستند، مؤثرتر است (Barth et al, 2012: 482-493).

این نشان می‌دهد که پیشرفت در کاربردهای هوش مصنوعی امکان طراحی سیستم‌های نسبتاً کارآمد را فراهم کرده است که به‌طور خودکار فعالیت‌های مخرب را در فضای مجازی شناسایی و از آن جلوگیری می‌کند (Calzavara et al, 2015: 1-30). آن‌ها برای حمایت از روش‌های فناوری موجود به تصویب رسیده‌اند، زیرا استانداردها و سازوکارهای مؤثری را برای کنترل و جلوگیری بهتر از حملات سایبری فراهم می‌کنند (Regev et al, 2009: 1-40). علی‌رغم تمام مزایایی که هوش مصنوعی فراهم می‌کند، توسعه سریع آن باعث می‌شود تا محققان از به کارگیری کارآمدترین تکنیک و تأثیر آن بر امنیت فضای مجازی با مشکل مواجه شوند. هیچ ابهامی وجود ندارد که تصور عمومی در بین محققان امنیت اطلاعات و امنیت سایبری حاکی از آن است که هوش مصنوعی امنیت اطلاعات سازمانی را بهبود بخشیده است، اما باتوجه به دانش ما، این ادعاها حدس و گمان است و از نظر تجربی اثبات نشده است. بیشتر مطالعات موجود نشان داده‌اند که چگونه نوآوری آن‌ها از انتخاب روش‌های موجود بهتر عمل می‌کند یا نمونه‌ای از سیستم‌ها را بررسی کرده و عملکرد آن‌ها را در مقایسه با آن‌ها ارزیابی می‌کنند. در همه موارد، سطح تعصبات انتخاب نسبتاً زیاد است. بر این اساس، نیاز به یک ادبیات جمع‌بندی شده است که خلاصه‌ای از موضوعات، چالش‌ها و دستورالعمل‌های آینده تحقیق در این دامنه را ارائه دهد (Zhu & Tan, 2011: 202-206; Ali et al, 2011: 486-497). در این قسمت به

^۱ - Biton

شناسایی می‌نماید» (قربانپور ویشکاسوقی و میرعظیمی طالبوندانی، ۱۴۰۲: ۵).

۲-۱-۲- شناسایی باتنت‌ها

یکی از رایج‌ترین مشکلات در تشخیص ناهنجاری، شبکه مربوط به باتنت‌هاست. باتوجه به خطر چنین شبکه‌های مخفی شناسایی باتنت‌ها جهت جلوگیری از فرسودگی شبکه و مصرف بیهوده منابع محاسباتی و برای جلوگیری از انتشار اطلاعات حساس نشت داده‌ها به خارج از شرکت، ازجمله فعالیت‌های ضروری جهت برقراری امنیت داده‌ها و اطلاعات در سطح شبکه است، باتنت‌ها از مجموعه‌ای از کامپیوترها و یا سیستم‌های آلوده تشکیل شده‌اند که توسط یک بدافزار هدایت می‌شوند. افرادی که بدافزارها را ایجاد کرده‌اند می‌توانند به‌جای این که به یک کامپیوتر آلوده به‌صورت مستقیم متصل شوند، از باتنت‌ها برای مدیریت خودکار حجم زیادی از این کامپیوترهای آلوده استفاده کنند که از طریق یک کانال فرمان و کنترل C&C به یکدیگر وصل شده‌اند، یکی از مشکلات کشف باتنت‌ها شباهت زیاد ترافیک باتنت با ترافیک واقعی شبکه است، بنابراین به‌راحتی و با استفاده از روش‌های سنتی نمی‌توان آن‌ها را شناسایی نمود. الگوریتم‌های یادگیری ماشین در دو نوع با ناظر و بدون ناظر یکی از ابزارهای هوش مصنوعی در شناسایی باتنت‌ها هستند. شناسایی باتنت‌ها نمونه‌ای از مسائل طبقه‌بندی است که با هدف تعیین کلاس یک بسته یا توالی بسته‌ها به ترافیک باتنت‌ها و یا ترافیک معمولی مورد بررسی قرار می‌گیرند. از طرف دیگر جهت گروه‌بندی ترافیک باتوجه به ویژگی‌های مشابه و شناسایی ترافیک‌های مشکوک می‌توان از مدل‌های یادگیری ماشین بدون ناظر استفاده کرد (Bakhshi & Veisi, 2019).

۲-۲- امنیت ایمیل

تهدیدات امنیتی از ایمیل نیز به‌عنوان ابزاری برای حمله استفاده می‌کند. از آنجا که میزان ترافیک منتقل‌شده از این طریق به‌سیار زیاد است، استفاده از روش‌های تشخیص خودکار که از الگوریتم‌های یادگیری ماشین بهره می‌برند، امری ضروری در شناسایی هرزنامه و حملات فیشینگ است، کاربردهای هوش مصنوعی در امنیت ایمیل در ذیل شرح داده

یافتن شباهت بین داده‌های بدون برچسب آن‌ها را گروه‌بندی می‌نماید، باتوجه به این که داده‌ها برای این مدل‌ها همگی یکسان به‌نظر می‌رسند بدون برچسب، هدف این نوع از الگوریتم‌ها درک رابطه بین داده‌ها و درنهایت گروه‌بندی آن‌هاست، خوشه‌بندی، یکی از نمونه‌های رایج این نوع از مدل‌هاست. در ادامه کاربردهای هوش مصنوعی در امنیت شبکه شرح داده می‌شود (Bowman & Huang, 2019).

۲-۱-۲- تشخیص تهاجم

سیستم تشخیص تهاجم و (یا) نفوذ یک ابزار مؤثر جهت شناسایی و تشخیص هرگونه استفاده غیرمجاز سوءاستفاده و یا آسیب‌رسانی در شبکه را برعهده دارد. «برای ایجاد امنیت کامل در یک سیستم کامپیوتری، علاوه بر دیوارهای آتش و دیگر تجهیزات جلوگیری از نفوذ سیستم‌های دیگری نیز نیاز است تا بتوان در صورت عبور نفوذگر از دیواره آتش آنتی‌ویروس و سایر تجهیزات امنیتی آن‌ها را تشخیص داده و چاره‌ای برای مقابله با آن تعبیه نمود. به‌دلیل ماهیت غیرالگوریتمی روش‌های نفوذ در شبکه‌های کامپیوتری راهکارهای مطرح‌شده برای مقابله با ناهنجاری‌ها نیز باید دارای ماهیت غیرالگوریتمی باشد. هوش مصنوعی به‌ویژه شبکه عصبی که جزء سیستم‌های یادگیرنده محسوب می‌شود، توانسته در زمینه شناسایی و جلوگیری از این ناهنجاری‌ها تأثیر به‌سزایی بگذارد. به‌طور سنتی، فعالیت شناسایی نفوذ از طریق معرفی سیستم‌هایی معروف به سیستم‌های تشخیص نفوذ مدیریت می‌شود. این سیستم‌ها به دو دسته سیستم تشخیص نفوذ تحت شبکه و سیستم‌های تشخیص نفوذ میزبان تقسیم می‌شوند که براساس عواملی، از قبیل نوع و تعداد فرآیندهای در حال اجرا، رفتار کاربران ماژول‌های سیستم عامل در هنگام راه‌اندازی سیستم و یا الگوی مصرف ترافیک به نظارت و یا کنترل شبکه‌های داخلی و یا خارجی می‌پردازند. با معرفی تکنیک‌های هوش مصنوعی در زمینه امنیت سایبری، نوع سومی تحت عنوان سیستم تشخیص نفوذ مبتنی بر ناهنجاری مطرح گردید، این دسته با استفاده از الگوریتم‌های یادگیری با ناظر و یا بدون ناظر، یادگیری تقویتی و حتی الگوریتم‌های خوشه‌بندی حوزه داده‌کاوی رویدادهایی را که خارج از رفتار عادی سیستم اتفاق می‌افتد به‌عنوان ناهنجاری

رویکرد پیشگیرانه برای حفاظت حساب و اطلاعات کاربران می‌تواند متمر ثمر واقع شود. در اینجا است که هوش مصنوعی با استفاده از روش‌های مختلف داده‌کاوی و یادگیری ماشین جهت بهره‌برداری از داده‌های ساختاریافته و یا غیرساختاری استخراج‌شده از منابع ناهمگون سازمان به‌کار می‌آید، هوش مصنوعی با شروع تجزیه و تحلیل داده‌های گذشته قادر به نشان‌دادن الگوهای نهفته برآورد رفتارهای آینده کاربران و شناسایی به‌موقع تلاش‌های احتمالی برای کلاهبرداری است (Adekunle et al, 2019). مقابله با حساب کاربری جعلی، از جمله فعالیت‌های مستلزم نظارت است، زیرا گسترش حساب‌های جعلی باعث فراهم‌نمودن بستری ناامن برای انجام فعالیت‌های نامناسب، از جمله فریب کاربران قانونی و استفاده از حساب کاربری آن‌ها می‌گردد، نظارت بر حساب کاربری شبکه‌های اجتماعی مانند فیسبوک و توئیتر که روزانه تعداد زیادی حساب کاربری ایجاد می‌شود، کاری دشوار و زمان‌بر است که با گسترش کاربردهای هوش مصنوعی در امنیت، این فعالیت با دقت بالا و زمان کم قابل انجام است. به‌عنوان نمونه، فیسبوک با به‌کارگیری سیستم جدیدی مبتنی بر یادگیری ماشین موسوم به «Deep Entity Classification» به شناسایی حساب‌های جعلی می‌پردازد و از فعالیت آن‌ها جلوگیری می‌کند. این الگوریتم ۲۰۰۰۰ ویژگی متعلق به هر حساب و افراد وابسته به آن حساب را در نظر می‌گیرد و اصل بودن آن را بررسی می‌نماید. انجام این عمل به‌صورت دستی توسط انسان کاری تقریباً غیرممکن است (Adekunle et al, 2019).

۲-۴- احراز هویت و تشخیص فریب احراز هویت

احراز هویت فرآیندی است که طی آن درستی هویت یک فرد شناسایی و تأیید می‌گردد. بنابراین زمانی که کاربر بخواهد وارد سیستم شده و یا به منبعی دسترسی پیدا کند، ابتدا باید خود را اثبات نماید (Chang et al, 2016). احراز هویت به‌صورت روش‌های متفاوتی، از قبیل سؤال‌های، امنیتی رمزهای عبور توکن‌ها، دستگاه‌های فیزیکی و ویژگی‌های بیومتریک انجام می‌گیرد. در حال حاضر احراز هویت از طریق ویژگی‌های بیومتریک بسیار مطرح شده است. منظور از ویژگی‌های بیومتریک در احراز هویت، ویژگی‌های فیزیکی منحصر به فرد

شده است (Bakhshi & Veisi, 2019). امنیت ایمیل‌ها از دو طریق شناسایی هرزنامه و شناسایی حملات فیشینگ انجام می‌شود. «هرزنامه یا اسپم به پیام الکترونیکی اطلاق می‌شود که بدون درخواست گیرنده و برای افراد بی‌شمار فرستاده می‌شود. هرزنامه به‌نوعی سوءاستفاده از سامانه انتقال پیام است. جهت شناسایی هرزنامه، استراتژی‌های هوش مصنوعی متفاوتی وجود دارد که یکی از رایج‌ترین و ساده‌ترین آن‌ها شبکه‌های عصبی است. از قابلیت‌های دیگری همچون ماشین‌های بردار پشتیبان (به‌ویژه برای اسپم‌های تصویری)، شبکه‌های بیز و همچنین استفاده از فناوری‌های پردازش زبان طبیعی می‌توان در این زمینه استفاده نمود» (کشاورز و حسینی، ۱۴۰۲: ۴).

۲-۳- حفاظت از حساب‌ها و اطلاعات کاربران و مقابله با حساب‌های جعلی

«حفاظت از حساب‌های کاربری برای هر سازمانی به‌خصوص در شرایطی که دارای مسئولیت‌های قانونی در قبال اشخاص ثالث است، یک مأموریت حساس به‌شمار می‌رود. این امر با گسترش حساب‌های جعلی جهت سرقت اطلاعات محرمانه و حساس و یا ایجاد تشنج و بیان مطالب محرک و توهین‌آمیز در فضای سایبری بحرانی‌تر شده است. یکی از نقاط ضعف در محافظت از حساب کاربران، محافظت ضعیف از رمز عبور آن‌هاست، هرچند امروزه در راستای افزایش امنیت حساب‌های کاربران، راهکارهایی، از جمله ارسال پیام و یا ایمیل به کاربر در حین اتصال سیستم دیگری به حساب مربوطه مطرح شده است. با این حال، این فعالیت‌ها، واکنشی هستند که با شناسایی دسترسی‌های غیرمجاز، سیستم در قالب هشدار، واکنشی را نشان داده و باعث بسته و یا معلق‌شدن حساب کاربر می‌گردد. این سیستم‌های هشدار واکنشی، معمولاً با مجموعه‌ای از محرک‌های پیش‌فرض و مرتبط با رویدادها فعال می‌شوند که برای تمامی کاربران یکسان هستند. به‌عبارت دیگر، این سیستم‌ها در شناسایی رفتار کاربران تلاشی نکرده تا بتوانند براساس الگوی رفتاری هر فرد عمل نمایند. علاوه بر این، سیستم‌های واکنشی، آینده را مشابه با گذشته در نظر می‌گیرند و توانایی سازگاری سریع با تغییرات را ندارند» (قربانپور و ویشکاسوقی و میرعظیمی طبالوندانی، ۱۴۰۲: ۶). اتخاذ یک

مکانیسم احراز هویت امنیتی شبکه ارتباطی تلفن همراه را براساس رگرسیون بردار پشتیبانی (SVR) ساختند، اما داده‌های کمتری در شبیه‌سازی استفاده شد. چانگ و همکاران از ماشین بردار پشتیبانی یک کلاس (SVM یک کلاس) برای تشخیص الگوی پویایی ضربه‌زدن به کلید استفاده کرد و این الگو به دلیل هوش مصنوعی توجه گسترده‌ای را به خود جلب کرد (Ali et al, 2017: 185-190). لو و همکاران از شبکه عصبی کانولوشن (CNN) یادگیری تقویتی و یادگیری انتقال برای ساخت یک طرح احراز هویت فیزیکی استفاده کردند. هدف آن محاسبات لبه موبایل بود و برای مقاومت در برابر حملات لبه‌های سرکش استفاده شد.

با پیشرفت تکنولوژی علاوه‌بر ایجاد راهکارهایی جهت تأمین امنیت بیشتر در فضای دیجیتال می‌توان ادعا نمود که به همان میزان حملات و تقلب‌های سایبری نیز پیچیده‌تر و پیشرفته‌تر شده است. جعل در تصاویر در حوزه Anti Spoofing شامل موارد متفاوتی، ازجمله جعل در صحبت کردن و عدم رعایت الگوی مورد نظر در صحبت می‌باشد. با ظهور فناوری دیپ فیک و مدل‌هایی مانند مدل‌های مولد تخصصی می‌توان تصاویری از افراد ایجاد کرد که هرگز حضور فیزیکی نداشته‌اند و سیستم‌های مربوط به احراز هویت را فریب داد. در این راستا، فرآیندی تحت عنوان صحت‌سنجی یا تشخیص زنده‌بودن مطرح شده که عبارت است از مجموعه‌ای از عملیات که تصاویر ویدیویی را مورد پردازش قرار داده و تشخیص می‌دهد که ویدیوی دریافت‌شده از فرد جعلی نبوده و مورد دستکاری عامدانه قرار نگرفته است. این مورد یکی از جالب‌ترین کاربردهای هوش مصنوعی در امنیت سایبری است که به‌نوعی این فناوری را در مقابل خود قرار می‌دهد (Adekunle et al, 2019).

۲-۵- امنیت برنامه

تأمین امنیت برنامه یکی دیگر از کاربردهای هوش مصنوعی در امنیت سایبری شامل ملاحظات امنیتی برای مواردی، از قبیل برنامه‌های وب، برنامه‌های کاربردی کلاینت و موبایل است که در طول، توسعه طراحی برنامه و پس از استقرار آن اتفاق می‌افتد. به‌طور کلی برنامه‌ها در همه مکان‌ها قابلیت اجرا شدن

مانند عنبیه چشم، چهره، اثر انگشت و صداست که از طریق آن می‌توان افراد را شناسایی و ردیابی نمود، هوش مصنوعی با استفاده از تکنولوژی‌های بینایی ماشین و تشخیص گفتار، تحول شگرفی در این زمینه ایجاد کرده است که می‌تواند با تلفیق با علم داده‌کاوی الگوهای رفتاری مانند نحوه تایپ مطالب و دست خط و یا الگوی امضا کردن را نیز شناسایی نماید و از آن برای صحت‌گذاری هویت افراد بهره‌برد. علاوه‌بر این، از قابلیت‌های دیگر سیستم‌های هوشمند می‌توان به موارد زیر اشاره کرد (Adekunle et al, 2019).

دسترسی کاربر به‌عنوان اولین خط دفاعی امنیت سایبری، این سیستم باید مدیریت احراز هویت دسترسی کاربر را تقویت کند، انواع رفتارهای استتاری را با دقت شناسایی کند و تشخیص اشیای غیرقانونی یا مخرب را شناسایی کند. قبل از عملیات سیستم باید از احراز هویت کاربران اطمینان حاصل کند، در عین حال داده‌های کاربر باید محرمانه باشد تا از سایر رویدادهای خطرناک مانند جمع‌آوری مخرب اطلاعات کاربر جلوگیری شود. در فرآیند احراز هویت فعلی بر افزودن ویژگی‌های دیگری برای افزایش منحصربه‌فرد بودن فرآیند تطبیق رمز عبور تمرکز می‌کند تا احتمال عبور دیگران به‌عنوان کاربران قانونی را به حداقل برساند.

حالت نحوه تطبیق رمزهای عبور و افزودن سایر مشخصات کاربر برای اطمینان از امنیت احراز هویت دوگانه چالشی است که باید در احراز هویت حالت حل شود. به‌عنوان مثال، دستگاه‌های خودپرداز فعلی فقط از کدهای پین برای تأیید هویت استفاده می‌کنند که به‌تنهایی امنیت احراز هویت را تضمین نمی‌کند (Adekunle et al, 2019). با توجه به کاستی‌های احراز هویت تک‌مرحله‌ای فناوری احراز هویت چندگانه مانند شوفان در نظر گرفته شده است، برای رسیدن به این هدف از جنگل تصادفی استفاده کرد. کورکماز نه‌تنها تطبیق رمز عبور را در سیستم احراز هویت رمز عبور انجام داد، بلکه صفحه کلید کاربر را با استفاده از برخی سبک‌ها از طریق شبکه عصبی آموزش داد، این سبک‌ها شامل سرعت تایپ کاربر و سبک تایپ، ترکیب کلیدها و سایر جنبه‌ها بود. وانگ و فانگ یک تابع هسته با عملکردهای جهانی و محلی طراحی کردند و یک

دارند و دائماً در حال تغییر هستند. به همین دلیل تأمین امنیت آن‌ها دشوار است. هدف از امنیت برنامه جلوگیری از سرقت یا ربودن داده‌ها یا کدهای موجود در برنامه است (Singh et al., 2017).

نقض امنیت سایبری در لایه برنامه می‌تواند به دلایل مختلف، از قبیل کدنویسی، ضعیف، انجام کم تست و رشد سریع و چشم گیر تکنولوژی‌های جدید مانند سرورهای ابری باشد که این لایه را در مقابل حمله مهاجمان آسیب‌پذیر کرده است. با این حال، حوزه‌های مختلف هوش مصنوعی، از جمله سیستم‌های خبره و یادگیری ماشین توانسته‌اند در بهبود تجزیه و تحلیل و پیش‌بینی حملات امنیتی شناسایی نقاط آسیب‌پذیر برنامه و ارائه گزینه‌های اصلاحی جهت رفع عیوب کدگذاری تأثیرات قابل توجهی بگذارند. به عنوان نمونه، میکروسافت با تکیه بر هوش مصنوعی ابزاری ارائه کرده که به توسعه‌دهندگان کمک می‌کند ایرادهای کدهایشان را با دقت ۹۹ درصد شناسایی کنند (Bakhshi & Veisi, 2019). این کار باعث می‌شود بخش بزرگی از امنیت برنامه در همان ابتدا و در مرحله طراحی و توسعه، تأمین گردد.

۳- محدودیت‌های هوش مصنوعی در بهبود امنیت سایبری

مانند هر چیز دیگری، استفاده از هوش مصنوعی در زمینه امنیت سایبری دارای معایبی بوده و با محدودیت‌هایی همراه است. به منظور ایجاد و حفظ یک سیستم هوش مصنوعی، سازمان‌ها به منابع و سرمایه‌گذاری‌های بیشتری نیاز دارند. علاوه بر این از آنجایی که سیستم‌های هوش مصنوعی با استفاده از مجموع داده‌ها آموزش می‌بینند، باید مجموعه‌های متمایز زیادی از کدهای بدافزار، کدهای غیرمخرب و ناهنجاری‌ها را به دست آورد. دستیابی به این همه مجموعه داده‌ها زمان بر است و نیاز به سرمایه‌گذاری‌هایی دارد که اکثر سازمان‌ها قادر به پرداخت آن نیستند. بدون حجم عظیمی از داده‌ها و رویدادها، سیستم‌های هوش مصنوعی می‌توانند نتایج نادرست و یا مثبت کاذب ارائه دهند و دریافت داده‌های نادرست از منابع غیرقابل اعتماد حتی می‌تواند نتیجه معکوس داشته باشد. نکته منفی دیگر این است که مجرمان سایبری می‌توانند از هوش مصنوعی برای تجزیه و تحلیل بدافزار خود و انجام حملات

پیشرفته‌تر استفاده کنند که ما را به نقطه بعدی می‌رساند. مشکل تضمین امنیت اطلاعات محرمانه یکی از کلیدهای همه موضوعات اقتصاد دیجیتال، از جمله مشکل امنیت سایبری با استفاده از هوش مصنوعی است. جامعه جهانی نگران استفاده از هوش مصنوعی برای مقاصد جنایی است (بهرامی‌زاده و صادق‌زاده، ۱۴۰۱: ۱). از طرفی آیا هوش مصنوعی می‌تواند تمام رویدادهای نامطمئن را شناسایی کند؟ پاسخ قطعاً منفی است. این فناوری جدید به عنوان یک «شمشیر دولبه» علاوه بر مزیت‌های فراوان دارای کاستی‌های خاص خود نیز می‌باشد، این بخش عواملی را مورد بحث قرار می‌دهد که مدل هوش مصنوعی را در زمینه امنیت سایبری ناصداق می‌کند.

۳-۱- تداخل داده‌های گیج‌کننده

چه مقدار تداخل می‌تواند هوش مصنوعی را فریب دهد؟ شاید یک پیکسل کافی باشد. آزمایش سو و همکاران در سال ۲۰۱۹ نشان داد که تنها تغییر یک پیکسل در تصویر می‌تواند منجر به طبقه‌بندی اشتباه شبکه عصبی شود. همان‌طور که از این مثال مشاهده می‌شود، زمانی که داده‌ها «آلوده شوند»، فرصتی برای تقلب در سیستم هوش مصنوعی وجود دارد که منجر به وضعیت ناامن شبکه می‌شود.

۳-۲- عدم شفافیت در فرآیند تصمیم‌گیری هوش مصنوعی

در فرآیند تصمیم‌گیری هوش مصنوعی همه شرکت‌کنندگان از جمله برنامه‌نویسان نمی‌دانند که چرا و چگونه مدل هوش مصنوعی نتایج نهایی تصمیم‌گیری را ارائه می‌دهد، یعنی فرآیند تصمیم‌گیری هوش مصنوعی فاقد شفافیت است. مدل هوش مصنوعی شبیه جعبه سیاه است. در فرآیند ایجاد و خودسازی می‌تواند پیچیدگی و تنظیم خودکار پارامترها را بدون دخالت بیش از حد کارکنان محقق کند و در نتیجه موجب صرفه‌جویی در منابع انسانی شود. با این وجود، در عین حال مشکل این است که فرآیند تصمیم‌گیری آن دشوار است که به وضوح توضیح داده شود، اگرچه مدل هوش مصنوعی می‌تواند به دقت بالایی دست یابد، اما همه آزمایش‌ها در مجموعه آزمایشی پیاده سازی می‌شوند. بنابراین هنگام مواجهه با رویدادهای ناشناخته، این که آیا مدل هوش مصنوعی می‌تواند به چنین دقت بالایی دست یابد، باید تأیید شود. هنگامی که به نتایج تصمیم‌گیری

هوش مصنوعی نقش مهمی در پیشگیری و تشخیص رفتارهای پرخطر شبکه ایفا می‌کند، اما عوامل متعددی وجود دارد که در قضاوت صحیح آن اختلال ایجاد کنند، هوش مصنوعی برای کمک به متخصصان امنیتی در این زمینه است، نه جایگزین کردن آن‌ها. بنابراین همچنان لازم است متخصصان مربوطه مداخله کنند و از دانش شبکه مربوطه برای قضاوت حرفه‌ای در مورد فرم فعلی شبکه استفاده کنند.

در حال حاضر نوع جدیدی از هوش مصنوعی در حال توسعه است که انسان در حلقه نامیده می‌شود. در سال ۲۰۱۷، آژانس پروژه‌های تحقیقاتی پیشرفته دفاعی دارپا چالش رباتیک دارپا را طراحی کرد. در ژانویه ۲۰۱۹، دارپا پروژه هوش مصنوعی به نام KAIROS را برای پیاده‌سازی سیستمی منتشر کرد که می‌تواند رویدادها را شناسایی کرده و توجه انسان‌ها را به خود جلب کند. همچنین در ماه مه، راه‌اندازی پروژه ACE را برای توسعه قابلیت نبرد هوایی مشترک بین انسان و ماشین اعلام کرد در ماه نوامبر وزارت دفاع ایالات متحده گزارشی در مورد جنگنده‌های مکانیکی و ادغام انسان و ماشین دریافت کرد که در آن مبارزان سایبورگ برای جنگ‌های آینده ساخته می‌شدند. در زمینه نظامی نیز تحقیقات این فناوری جدید آغاز شده است که اهمیت آن را نیز منعکس می‌کند. تکنیک انسان در حلقه می‌تواند خرد انسانی و هوش ماشینی را ترکیب کند که روشی مهم برای درک مزایای مکمل انسان و ماشین است. هوش مصنوعی می‌تواند تعداد زیادی از داده‌ها را به سرعت پردازش کند و جلوه تشخیص خوبی برای صحنه‌های خاص دارد، اما ممکن است مختل شود و وضعیت جدید را به درستی قضاوت نکند، در مقایسه با ماشین‌ها، انعطاف‌پذیرتر هستند و می‌توانند در مواجهه با تغییرات جدید در شبکه سریع‌تر قضاوت کنند، اما همچنین به ماشین‌هایی برای ارائه نقش کمکی نیاز دارند، یادگیری ماشین، تعاملی که در هوش مصنوعی مورد استفاده قرار گرفت، در امنیت سایبری نیز تجسم یافته است (Santhanam et al, 2017: 209-230). افزودن این ایده به آگاهی، موقعیتی ضمن دستیابی به تجسم قابلیت اطمینان سیستم را نیز افزایش می‌دهد. بنابراین استفاده از هوش مصنوعی و انسان در حلقه در امنیت سایبری قابلیت‌های مدل‌ها را بیشتر خواهد کرد.

ارائه‌شده توسط مدل هوش مصنوعی اعتراضاتی وجود دارد، توضیح فرآیند تصمیم‌گیری دشوار است. بنابراین برخی از افراد نسبت به نتایج تصمیم‌گیری بدبین خواهند بود. این برای قضاوت سریع وضعیت شبکه مفید نخواهد بود و یا حتی باعث عواقب جبران‌ناپذیری خواهد شد. برخی از تیم‌های تحقیقاتی شروع به انجام تحقیقات عمیق در مورد این موضوع کرده‌اند (Schlegel et al, 2019: 419-420).

۳-۳- نیاز به داده‌های بالا

در حال حاضر، مدل‌های هوش مصنوعی اساساً به داده‌های زیادی برای تکمیل آموزش نیاز دارند. قبل از استفاده از داده‌ها، آن‌ها ممکن است عملیاتی را انجام دهند که عمده‌تاً شامل یک سری مراحل مانند کاهش نویز، داده‌ها، نرمال‌سازی، پرکردن مقادیر گم‌شده و ازدست‌رفته و ... است. اگر از روش نظارتی استفاده می‌شود، لازم است داده‌ها را به صورت دستی برچسب گذاری کرد. با این حال به دلیل ناهمگونی شدید فضای سایبری ساختارهای سایبری مختلف ممکن است رویدادهای پرخطر متفاوتی را ایجاد کنند و این رویدادها دارای ویژگی‌های ناگهانی هستند، بنابراین هر رویداد احتمالی پرخطر را نمی‌توان قبل از طراحی مدل‌ها تخمین زد و همچنین نمی‌توان این رویدادهای پرخطر را از قبل تحلیل و برچسب‌گذاری کرد. در همین حال مدل‌های هوش مصنوعی تقاضای بالایی برای داده‌ها دارند که ممکن است نتوانند به موقع بررسی کنند.

۳-۴- ضرورت مشارکت انسان با هوش مصنوعی

طراحی و اجرای هوش مصنوعی به‌ویژه شبکه‌های عصبی سعی در تقلید از مغز انسان دارد. هدف آن دستیابی به همان شیوه تفکر انسان با استفاده از نورون‌های متصل است. متأسفانه هوش مصنوعی برای دستیابی به یادگیری به تعداد زیادی نمونه نیاز دارد. بدون توانایی استدلال، مدل نهایی بعد از آموزش برای ما پیچیده است تا بفهمیم چگونه تصمیم می‌گیرد، اگرچه برخی تلاش‌ها برای یافتن توضیح هوش مصنوعی انجام شد، اما این کار هنوز در مرحله اولیه است، بنابراین برای دستیابی به استفاده کارآمد از هوش تکیه بر این ابزارها بدون مشارکت انسانی کافی نیست (Nunes et al, 2015: 965-944).

نتیجه‌گیری

استفاده از هوش مصنوعی در حوزه امنیت سایبری امیدوارکننده است و سیستم‌های تشخیص و پیشگیری از نفوذ در حال پیشرفت است. هوش مصنوعی باعث کاهش پیچیدگی محاسباتی و کاهش زمان آموزش مدل شده است. همچنین مشاهده شد که یک انحراف قابل توجهی در دامنه وجود دارد. علاوه بر این، محققان روی الگوریتم‌های کمتری تمرکز کرده‌اند و به همین ترتیب الگوریتم‌های جدید محبوب نیستند. این امر به‌عنوان یک چالش و همچنین فرصتی برای محققان به‌وجود می‌آید. اعتقاد بر این است که برنامه‌های هوش مصنوعی همچنان فرصت‌هایی برای امنیت سایبری ارائه می‌دهند. یکی از کارهایی که در این حوزه باید انجام گیرد، پیدا کردن تئوری «توابع تسهیلاتی گروهی» برای اجازه‌دادن به عامل‌ها برای تصمیم‌گیری گروهی است. استفاده از الگوریتم‌های یادگیری بدون نظارت برای ساخت سیستم‌های ترکیبی «کشف نفوذ غیرمعمول». ترکیب کلیه تکنیک‌های هوش مصنوعی برای ساخت ویروس‌یاب‌های جدید، پیشنهادات دیگر در این زمینه می‌باشند.

باتوجه به محدودیت‌های انسانی و این حقیقت که ویروس‌ها و کرم‌های رایانه‌ای هوشمند و انعطاف‌پذیر شده‌اند، طراحی عامل‌های با سنسور هوشمند بسیار ضروری می‌باشد. این عامل‌ها توانایی کشف، ارزیابی و پاسخ‌دهی به حملات سایبری در زمان مناسب را باید در خود داشته باشند. به‌کارگیری تکنیک‌های هوش مصنوعی در حوزه دفاع سایبری نیاز به برنامه‌ریزی و مطالعه و تحقیق دارد. پیش‌بینی در زمینه‌های مربوط به جرم از طریق مطالعه نظام‌مند روندها، آثار و عوارض آن‌ها را در آینده بررسی و تبیین کرده و براساس آن‌ها تمهیداتی را برای جلوگیری از رخدادها یا کاهش تبعات در نظر می‌گیرد. حال باتوجه به این که یکی از کاربردی‌ترین و تخصصی‌ترین استفاده‌های هوش مصنوعی، پیش‌بینی و تصمیم‌گیری در مقابل داده‌ها و قوانینی است که از قبل برای آن تعریف شده است، این علم می‌تواند در زمینه پیشگیری، به‌خصوص پیشگیری انتظامی از جرم کمک شایانی به نهادهایی کند که وظیفه پیشگیری انتظامی از جرم را برعهده دارند، ولی مشروط بر این که داده‌هایی که به این عامل وارد می‌شوند و همچنین

قوانینی که برای آن تعریف می‌شود، به‌صورت واقعی، دقیق و منطقی باشد تا بتوان شاهد خروجی کارآمد و مناسبی بود. هوش مصنوعی با استفاده از شبکه‌های عصبی مصنوعی، طراحی عامل یادگیر، منطق فازی، پردازش سیگنال نقش به‌سزایی را در پیشگیری انتظامی از جرم ایفا می‌کند و می‌توان با استفاده مناسب از خروجی این عامل‌ها و همچنین با دسته‌بندی و در اختیار قراردادن این اطلاعات به نهادهای ذی‌ربط گامی نو و اساسی در راستای پیشگیری انتظامی از جرایم برداشت.

ملاحظات اخلاقی: در این پژوهش تمامی ملاحظات اخلاقی رعایت گردیده است.

تعارض منافع: نگارش این مقاله، فاقد هرگونه تعارض منافی بوده است.

سهم نویسندگان: برابر.

تشکر و قدردانی: ابراز نشده است.

تأمین اعتبار پژوهش: این پژوهش بدون تأمین مالی انجام گرفته است.

منابع و مأخذ

الف. منابع فارسی

- بهرامی‌زاده، امیرحسین و صادق‌زاده، سینا (۱۴۰۱). «بررسی هوش مصنوعی و مشکلات تأمین امنیت سایبری». مازندران: شانزدهمین کنفرانس ملی علوم و مهندسی کامپیوتر و فناوری اطلاعات.

- قربانپور ویشکاسوقی، سینا و میرعظیمی طبالوندانی، سیداحمد (۱۴۰۲). «راهنمای جامع درباره کاربردهای هوش مصنوعی در امنیت سایبری». رشت: یازدهمین کنفرانس ملی ایده‌های نوین در علوم انسانی و مهندسی.

- کشاورز، زهرا و حسینی، حمیدرضا (۱۴۰۲). «هوش مصنوعی در امنیت سایبری (کاربردها، چالش‌ها و فرصت‌ها)». تهران: ششمین همایش ملی فناوری‌های نوین در مهندسی برق، کامپیوتر و مکانیک ایران.

ب. منابع انگلیسی

Transactions on Signal Information Processing over Networks, 4(1): 48-59.

- Kleinmann, A & Wool, A (2016). "Automatic Construction of Statechart-Based Anomaly Detection Models for MultiThreaded Industrial Control Systems". *Journal of International publication*, 8(4): 1-21.

- Kotenko, I & Ulanov, A (2007). "Multi-Agent Framework fo Simulation of Adaptive Cooperative Defense Against Internet Attacks". *Journal of International Workshop on Autonomous Intelligent Systems Agents and Data Mining*, 10(447): 221-228.

- Kowert, W (2017). "The foreseeability of human-artificial intelligence interactions". *Journal of Tex L Rev*, 96(1): 181.

- Lebbe, M. A; Agbinya, J. I; Chaczko, Z & Chiang, F (2007). "Self-Organized Classification of Dangers for Secure Wireless Mesh Networks". Australa: Australasian Telecommunication Networks and Applications Conference.

- Nunes DS; Zhang P & S.S. J(2015). "A survey on human-in-the-loop applications towards an internet of all". *Journal of IEEE Commun Surv Tutor*, 17(2): 944-965.

- Ojugo, A.A; Eboka, A.O; Okonta, O.E; Yoro, R.E & Aghware, F.O (2012). "Genetic Algorithm Rule-Based Intrusion Detection System (GAIDS)". *Journal of Emerging Trends in Computing and Information Sciences*, 3(8): 118-119.

- Phillips, L; Link, H; Smith, R & Weiland, L(2006). *Agent-Based Control of Distributed Infrastructure Resources*. USA: Published USA Department of Energy (Sandia Nati onal Laboratories).

- Regev, O (2009). "On lattices learning with errors random linear codes and cryptography". *Journal of J ACM*, 56(6): 1-43.

- Santhanam, GR & et al (2017). *Human-on-the-loop automation for detecting software side-channel vulnerabilities*. Beja: Publications Information systems security (Springer Cham).

- Adekunle, Y, et al(2019)." Holistic exploration of gaps vis-à-vis artificial intelligence in automated teller machine and internet banking". *International journal of applied information systems*, 2(25): 12-25.

- Bakhshi, B & Veisi, H (2019). "End to end fingerprint verification based on convolutional neural network in 2019". America: 27th Iranian conference on electrical engineering (ICEE).

- Barth, A; Rubinstein, P; Sundararajan, M; Mitchell, J. C; Song, D & Bartlett, P. L(2012). "A Learning-Based Approach to Reactive Security". *Journal of IEEE Trans Dependable Secur Comput*, 9(4): 482-490.

- Bowman, B & Huang, H.H(2019). "Securing malware cognitive systems against adversarial attacks in 2019". USA: 17th IEEE international conference on cognitive computing (ICCC).

- Brenner, S. W (2010). *Cybercrime Criminal Threats from Cyberspace*. 2th edition, USA: Greenwood publishing group Library of Congress Cataloging-in-Publication Data USA.

- Calzavara, S; Tolomei, G; Casini, A; Bugliesi, M and Orlando, S (2015). "A Supervised Learning Approach to Protect Client Authentication on the Web". *Journal of ACM Trans Web*, 9(1): 1-30.

- Chang, C; T Eude & Obando Carbajal, L.E(2016). "Biometric authentication by keystroke dynamics for remote evaluation with one-class classification". Canada: in Advances in Artificial Intelligence 29th Canadian Conference on Artificial Intelligence Canadian AI 2016.

- Gordon, S; Ford, R (2006). "On the definition and classification of cybercrime". *Journal in Computer Virology*, 2(1): 13-20.

- Gou, X; Jin, W; Zhao, D (2006). "Multi-agent system for worm detection and containment in metropolitan area networks". *Journal of Electronics*, 23(2): 259-265.

- Guan, Y & Ge, X (2017). "Distributed attack detection and secure estimation of networked cyber-physical systems against false data injection attacks and jamming attacks". *IEEE*

- Schlegel, U & et al (2019). “*Towards a rigorous evaluation of xai methods on time series*”. New York: international conference on computer vision workshop (ICCVW).
- Singh, K & et al (2017). “*Sparse proximity based robust fingerprint recognition*”. England: international conference on computing communication and automation (ICCCA).
- Sirisanyalak, B & Sornil, O (2007). “*An artificial immunity-based spam detection system*”. USA: IEEE Congress on Evolutionary Computation (CEC 2007).
- Tsai, C-F; Hsu, Y-F; Lin, C-Y and Lin & W-Y (2009). “Intrusion detection by machine learning A review”. *Journal of Expert Syst*, 6(13): 11994-12333.
- Xiao, R & et al (2018).” *Attacking network isolation in software-defined networks: New attacks and countermeasures in 2018*”. Japan: international conference on computer communication and networks (ICCCN).
- Zhu. Y & Tan, Y (2011). “A local-concentration-based feature extraction approach for spam _ltering”. *Journal of IEEE Trans Inf Forensics Secur*, 6(2): 486-490.

