

امنیت سایبری در عصر اینترنت اشیا با هوش مصنوعی

سیدمجید عارفی نژاد^۱

^۱ کارشناسی ارشد مهندسی برق-کنترل، دانشکده علوم و مهندسی دفاعی، دانشگاه افسری و تربیت پاسداری
امام حسین (ع)، تهران، ایران
sm.arefi@ihu.ac.ir

چکیده

با گسترش اینترنت اشیا، افزایش دستگاه‌های متصل چالش‌های قابل توجهی در زمینه امنیت سایبری ایجاد می‌کند. تکنیک‌های پیش‌بینی به طور فزاینده‌ای برای مقابله با این تهدیدات سایبری که به طور مداوم در حال تحول هستند، ضروری شده‌اند، زیرا اقدامات سنتی امنیت سایبری در برابر آن‌ها ناکارآمد نشان داده شده‌اند. این مقاله به بررسی دنیای تهدیدات سایبری می‌پردازد و به موضوعاتی مانند باج‌افزار، فیشینگ، بدافزار و حملات منع سرویس (DoS) می‌پردازد. هوش مصنوعی (AI) به عنوان یک نیروی تحول‌آفرین در زمینه امنیت سایبری ظاهر شده و نقش مهمی در حفاظت از زیرساخت‌های ملی ایفا می‌کند. این مقاله اهمیت هوش مصنوعی در حمایت از دفاع امنیت سایبری، از جمله سیستم‌های تشخیص نفوذ، امنیت شبکه و استفاده از عوامل هوشمند اشاره می‌کند. همچنین به اهمیت تکنیک‌های یادگیری ماشین و مدل‌سازی پیش‌بینی در پیش‌بینی و جلوگیری از حملات سایبری می‌پردازد. با وجود مزایای بالقوه امنیت سایبری مبتنی بر هوش مصنوعی، جدیت مشکلات مربوط به حریم خصوصی داده‌ها، مقیاس‌پذیری و همکاری انسان و ماشین را نمی‌توان نادیده گرفت. در محیط دیجیتالی فزاینده امروز، شرکت‌ها می‌توانند با پیاده‌سازی راه‌حل‌های امنیت سایبری مبتنی بر هوش مصنوعی، دفاع خود را در برابر حملات سایبری تقویت کرده و از دارایی‌های ارزشمند خود محافظت کنند.

کلمات کلیدی: اینترنت اشیا، امنیت سایبری، هوش مصنوعی.

۱ مقدمه

در چشم‌انداز معاصر، گسترش فناوری‌های دیجیتال بی‌تردید هر جنبه‌ای از جامعه انسانی را متحول کرده است. از اتصال بی‌وقفه‌ای که توسط رایانش ابری تسهیل شده تا شبکه پیچیده‌ای از دستگاه‌های متصل که اینترنت اشیا را تشکیل می‌دهند [۱]، وابستگی ما به زیرساخت‌های دیجیتال به جزء اساسی در زندگی مدرن تبدیل شده است. این فناوری‌ها سطوح بی‌سابقه‌ای از کارایی، بهره‌وری و ارتباطات را به ارمغان آورده‌اند و نحوه انجام کسب‌وکار، ارتباط با یکدیگر و تعامل با جهان اطراف را دگرگون کرده‌اند [۲]، [۳].

با گسترش سریع اینترنت اشیا و تهدید فزاینده حملات سایبری، نقش هوش مصنوعی در امنیت سایبری هرگز به این اندازه حیاتی نبوده است. این مقاله به بررسی انواع مختلف حملات سایبری و ماهیت در حال تحول آن‌ها می‌پردازد و نحوه استفاده از هوش مصنوعی برای پیش‌بینی و جلوگیری از این تهدیدات را مورد بحث قرار می‌دهد.

علاوه بر این، مدل‌سازی پیش‌بینی برای پیش‌بینی حملات سایبری و الگوریتم‌های یادگیری ماشین برای جلوگیری از حملات سایبری نیز بررسی می‌شود. با این حال، در میان امنیت سایبری مبتنی بر هوش مصنوعی، نیاز به توجه به چالش‌ها و محدودیت‌های هوش مصنوعی، مانند حریم خصوصی داده‌ها، مقیاس‌پذیری و همکاری انسان و ماشین وجود دارد. به طور خلاصه، این مقاله به حملات سایبری در حال تحول اشاره کرده و رویکرد پیشگیرانه‌ای را برای پیش‌بینی و جلوگیری از حملات سایبری با استفاده از مدل‌های هوش مصنوعی بررسی می‌کند.

۲ انواع حملات سایبری

حمله سایبری تلاشی عمدی از سوی یک فرد یا گروه برای نفوذ به اطلاعات سیستم به منظور اهداف اقتصادی یا سرقت داده‌ها از سیستم‌های هدف است [۴]. درک این حملات مختلف برای توسعه یک استراتژی پیش‌بینی مؤثر بسیار مهم است. انواع مختلف حملات سایبری به شرح زیر تعریف شده‌اند [۴]:

- حملات بدافزاری: انواع مختلفی از برنامه‌های خطرناک که برای دسترسی غیرقانونی به سیستم‌های کامپیوتری ایجاد شده‌اند، بدافزار نامیده می‌شوند که گاهی اوقات به آن‌ها نرم‌افزارهای مخرب نیز گفته می‌شود. اسب‌های تروجان، کرم‌ها، ویروس‌ها و بدافزارها چند نمونه از آن‌ها هستند. وب‌سایت‌های مخرب، نرم‌افزارهای هک‌شده و پیوست‌های ایمیل آلوده می‌توانند همگی بدافزار را گسترش دهند. بدافزار می‌تواند فایل‌ها را از بین ببرد، اطلاعات محرمانه را سرقت کند یا کنترل یک سیستم را به دست گرفته و از آن برای اهداف مخرب استفاده کند.
- حملات فیشینگ: حملات فیشینگ افراد را فریب می‌دهند تا اطلاعات حساس خود، مانند شماره‌های کارت اعتباری یا اطلاعات ورود، را فاش کنند و خود را به عنوان یک شرکت معتبر معرفی می‌کنند. مجرمین سایبری معمولاً بر ایجاد وب‌سایت‌های جعلی یا ارسال ایمیل‌های به ظاهر معتبر که واقعی به نظر می‌رسند، تکیه می‌کنند و هدفشان فریب دادن افراد بی‌خبر به این باور است که این منابع قابل اعتماد هستند. در حوزه مالی، حملات فیشینگ به طور خاص بر روی مشتریان یا کارکنان بانک‌ها متمرکز است. هدف اصلی این است که به طور مخفیانه به حساب‌ها یا اطلاعات مالی حساس آن‌ها دسترسی پیدا کنند بدون اینکه آن‌ها متوجه شوند.
- باج‌افزار: مهاجمان به هدف قراردادن افراد یا سازمان‌ها می‌پردازند تا از دسترسی آن‌ها به داده‌هایشان جلوگیری کنند مگر اینکه مبلغی به عنوان باج پرداخت کنند تا داده‌های خود را بازیابی کنند. قربانیان

ممکن است با تصمیمات دشواری در مورد اینکه آیا باج را پرداخت کنند یا سعی کنند داده‌های خود را از طریق روش‌های دیگر بازیابی کنند، مواجه شوند.

- حملات منع سرویس (DoS): حملات DoS سعی دارند با بارگذاری شبکه‌ها، سیستم‌ها یا نرم‌افزارهای هدف با ترافیک مخرب، در آن‌ها اختلال ایجاد کنند. این حملات معمولاً به سمت سازمان‌ها هدف‌گذاری می‌شوند تا به شهرت آن‌ها آسیب بزنند، خسارت مالی به بارآورند و در کسب‌وکار آن‌ها اختلال ایجاد کنند. حمله منع سرویس توزیع شده (DDoS) به دلیل نیاز به همکاری چندین دستگاه هک‌شده برای متوقف کردن حمله، ثقل در آن دشوارتر است.

۳ نقش هوش مصنوعی در امنیت سایبری

هوش مصنوعی یک جزء اساسی در حوزه امنیت سایبری می‌باشد؛ زیرا استراتژی‌های دفاعی امنیت سایبری را قوی‌تر می‌کند [۵]. هوش مصنوعی با تحلیل حجم وسیعی از داده‌ها به پیش‌بینی و جلوگیری از حملات سایبری قبل از وقوع آن‌ها کمک می‌کند. هوش مصنوعی دارای توانایی منحصربه‌فردی برای تجزیه و تحلیل حجم زیادی از داده‌ها با سرعتی بی‌سابقه است، که به مراتب از توانایی تحلیل‌گران انسانی پیشی می‌گیرد [۶]. این توانایی در الگوریتم‌های پیشرفته و قدرت محاسباتی که سیستم‌های هوش مصنوعی را پشتیبانی می‌کند، ریشه دارد و به آن‌ها این امکان را می‌دهد تا با کارایی فوق‌العاده‌ای داده‌های عظیم متشکل از منابع مختلف اطلاعات امنیتی را بررسی کنند. در اینجا توضیح مختصری از نحوه تجزیه و تحلیل داده‌ها توسط هوش مصنوعی برای شناسایی ناهنجاری‌ها و تهدیدات احتمالی سریع‌تر از انسان آورده شده است:

- الگوریتم‌های پیشرفته: سیستم‌های هوش مصنوعی از الگوریتم‌های پیچیده‌ای مانند یادگیری ماشین (ML) و یادگیری عمیق (DL) برای پردازش و تجزیه و تحلیل حجم زیادی از داده‌ها استفاده می‌کنند. این الگوریتم‌ها طراحی شده‌اند تا از الگوهای داده‌ها بیاموزند، به اطلاعات جدید تطبیق پیدا کنند و بر اساس بینش‌های حاصل از داده‌ها پیش‌بینی یا تصمیم‌گیری کنند [۷].
- پردازش موازی: سیستم‌های هوش مصنوعی قادر به پردازش موازی هستند که به آن‌ها امکان می‌دهد چندین جریان داده را به طور هم‌زمان تجزیه و تحلیل کنند. برخلاف تحلیل‌گران انسانی که محدود به ظرفیت شناختی و دامنه توجه خود هستند، هوش مصنوعی می‌تواند مجموعه‌های عظیم داده را به طور هم‌زمان پردازش کند و به طور قابل‌توجهی زمان لازم برای تحلیل را کاهش دهد [۸].
- تحلیل بلادرنگ: سیستم‌های هوش مصنوعی قادر به انجام تحلیل بلادرنگ بر روی داده‌های در حال جریان مانند گزارش‌های ترافیک شبکه و گزارش‌های رویدادهای سیستم به‌محض تولید آن‌ها هستند [۹]. این امر به سازمان‌ها امکان می‌دهد تا تهدیدات امنیتی را به‌صورت بلادرنگ شناسایی و به آن‌ها پاسخ دهند و تاثیر احتمالی حملات سایبری را به حداقل برسانند.

- تشخیص الگو: هوش مصنوعی در شناسایی الگوها، ناهنجاری‌ها و انحرافات در داده‌ها برتری دارد. با تجزیه و تحلیل داده‌های تاریخی و یادگیری از رویدادهای گذشته، هوش مصنوعی می‌تواند رفتارها یا فعالیت‌های غیرعادی که ممکن است نشانگر تهدید امنیتی باشد را شناسایی کند [۱۰]. این قابلیت تشخیص الگو به هوش مصنوعی امکان می‌دهد تا رویدادهای مشکوک را به سرعت و دقت شناسایی کند.
- اتوماسیون: اتوماسیون‌های مبتنی بر هوش مصنوعی فرآیند تجزیه و تحلیل داده‌ها را ساده می‌کند و نیاز به دخالت دستی در هر مرحله را از بین می‌برد. پس از آموزش، سیستم‌های هوش مصنوعی می‌توانند به طور خودکار داده‌ها را تجزیه و تحلیل کرده، ناهنجاری‌ها را شناسایی و بر اساس معیارهای از پیش تعریف شده هشدارها یا واکنش‌ها را فعال کنند [۱۱]. این اتوماسیون‌ها به طور قابل توجهی سرعت تشخیص و پاسخ به تهدیدات را افزایش می‌دهد و سازمان‌ها را قادر می‌سازد تا به سرعت به تهدیدات سایبری نوظهور واکنش نشان دهند.
- مقیاس‌پذیری: سیستم‌های هوش مصنوعی بسیار مقیاس‌پذیر هستند و قادر به مدیریت مجموعه‌های داده‌ای بزرگ‌تر و پیچیده‌تر بدون کاهش عملکرد هستند [۱۲]. هوش مصنوعی می‌تواند با توزیع محاسبات در بین چندین پردازنده یا نود، به صورت افقی مقیاس‌پذیر باشد، و حتی با افزایش حجم داده‌ها نیز تجزیه و تحلیل پیوسته و کارآمد را تضمین کند.
- برای جلوگیری از حملات سایبری، نیاز به یک رویکرد جامع و چندبعدی است که چندین تدبیر امنیتی را در سطوح مختلف زیرساختی درون یک سازمان ادغام کند. در زیر، تحلیل دقیق‌تری از هر روش برای جلوگیری از حملات سایبری ارائه شده است:
- تدابیر امنیت شبکه: چندین تدبیر امنیتی برای محافظت از حاشیه شبکه یک سازمان و جلوگیری از دسترسی غیرمجاز به سیستم‌ها و داده‌های داخلی لازم است. از جمله این تدابیر عبارت‌اند از:
 - دیوارهای آتش: دیوارهای آتش برای ارائه یک بازدارنده مؤثر در برابر تلاش‌های هکرها برای دسترسی غیرقانونی به یک کامپیوتر در زمان اتصال آن به اینترنت یا سایر اتصالات شبکه ایجاد شده‌اند [۳].
 - رمزنگاری: داده‌های رمزنگاری شده غیرقابل فهم هستند و برای رمزگشایی به کلید صحیح نیاز دارند. رمزگشایی، رمزنگاری و حل معادلات ریاضی پیچیده، مانند تجزیه اعداد اول بزرگ، زمان و منابع زیادی می‌طلبد. در رمزنگاری متقارن برای کدگذاری و رمزگشایی پیام‌ها از کلیدهای مشابه با سطوح امنیتی قابل مقایسه استفاده می‌شود [۳].
- تشخیص نفوذ: با مطالعه الگوهای ترافیک شبکه و شناسایی فعالیت‌های غیرمعمول که نشان دهنده یک حمله احتمالی هستند، تکنیک‌های هوش مصنوعی (AI) مانند یادگیری ماشین و یادگیری عمیق

می‌توانند سیستم‌های تشخیص نفوذ (IDS) را تقویت کنند [۱۳]. با آموزش از داده‌های گذشته، های IDS مبتنی بر هوش مصنوعی می‌توانند انحراف سیستم، شبکه یا خدمات را از رفتار طبیعی شناسایی کنند و بدین ترتیب امکان شناسایی زود هنگام تهدیدات را فراهم آورند. این سیستم‌ها می‌توانند قابلیت‌های پیش‌بینی خود را با تنظیم دینامیک قابلیت‌های تشخیص به استراتژی‌های حمله در حال تغییر، تکامل دهند.

- عامل‌های هوشمند: موجودات مستقل تولید شده توسط کامپیوتر که به عنوان عامل‌های هوشمند شناخته می‌شوند، با یکدیگر تعامل دارند تا اطلاعات را تبادل کرده و با هم همکاری کنند تا اقدامات مناسب را در صورت بروز شرایط غیرمنتظره سازماندهی و انجام دهند. فناوری عامل هوشمند به دلیل ماهیت همکاری، تحرک و سازگاری در شرایطی که استفاده می‌کنند، برای جلوگیری از حملات سایبری مناسب است [۱۴].

در نتیجه، تکنیک‌های هوش مصنوعی استراتژی‌های امنیت سایبری مانند سیستم‌های تشخیص نفوذ (IDS)، عامل‌های هوشمند و تدابیر امنیت شبکه را تقویت می‌کنند. سازمان‌ها می‌توانند با استفاده از فناوری‌های مبتنی بر هوش مصنوعی مانند یادگیری ماشین و یادگیری عمیق، پاسخ خود را به حملات سایبری بهبود بخشند. این امر موجب تقویت موضع امنیتی کلی آن‌ها در برابر تهدیدات در حال تغییر امروزی خواهد شد.

۴ مدل سازی برای پیش‌بینی حملات سایبری

یک جنبه حیاتی از هوش مصنوعی (AI) مدل سازی پیش‌بینی است که الگوها و ناهنجاری‌ها در داده‌ها را برای پیش‌بینی تهدیدات و حملات سایبری ممکن تحلیل می‌کند. این یک روش ترکیب داده‌ها و محاسبات جهت پیش‌بینی ریسک‌های آینده قبل از وقوع آنهاست [۱۵]. مدل سازی پیش‌بینی حملات سایبری را مشابه پیش‌بینی طوفان‌ها در پیش‌بینی آب‌وهوا انجام می‌دهد. یادگیری Federated از یادگیری ماشین توزیع شده سنتی متفاوت است، زیرا به چندین گره محاسباتی در یک شبکه اجازه می‌دهد که داده‌های آموزشی را به اشتراک بگذارند و به طور مشترک یک مدل یادگیری ماشین را آموزش دهند. این زمان مورد نیاز برای بازآموزی مدل‌ها را هنگام اعمال تغییرات کاهش می‌دهد و به روزرسانی گره‌ها در حین آموزش را ممکن می‌سازد [۱۶]. FedAvg، یکی از سه الگوریتم یادگیری Federated، در مقابل یک استراتژی یادگیری مرکزی با استفاده از مجموعه داده MNIST آزمایش شد. نتایج نشان داد که FedAvg بالاترین دقت کلی را داشت، اما زمانی که داده‌ها غیرمستقل و به طور هم‌زمان توزیع نشده بودند، تکنیک مرکزی بهتر عمل کرد [۱۷]. به دلیل اینکه یادگیری FL اجازه می‌دهد که از منابع داده بیشتری یاد گرفته شود بدون اینکه حریم خصوصی قربانی شود، مدل‌های پیش‌بینی ممکن است در شناسایی و جلوگیری از تهدیدات سایبری نوآورانه‌تر و موفق‌تر شوند [۱۸]. فرایند مدل سازی پیش‌بینی معمولاً شامل چندین مرحله است، از جمله:

- جمع‌آوری داده‌ها: جمع‌آوری اطلاعات مرتبط از منابع مختلفی از جمله حسگرها، پایگاه‌های داده و

وبسایت‌ها.

- پیش‌پردازش داده‌ها: پاک‌سازی و آماده‌سازی داده‌ها برای تحلیل. این شامل حل مقادیر گمشده، حذف ناهنجاری‌ها و تغییر متغیرها به‌عنوان ضرورت است.
- انتخاب ویژگی: انتخاب مهم‌ترین ویژگی‌ها یا متغیرهایی که احتمالاً بر نتیجه پیش‌بینی تأثیر می‌گذارند.
- آموزش مدل: شامل تنظیم پارامترهای مدل انتخاب شده برای کاهش خطاهای پیش‌بینی با استفاده از داده‌های تاریخی است.
- ارزیابی مدل دقت و قابلیت اعتماد: مدل آموزش‌دیده را با استفاده از روش‌های اعتبارسنجی ارزیابی می‌کند.
- پیش‌بینی: پیش‌بینی نتایج آینده با استفاده از مدل آموزش‌دیده بر روی داده‌های جدید یا آزمایش نشده.

بسیاری از صنایع از مدل‌سازی پیش‌بینی برای به‌دست‌آوردن بینش‌ها، اتخاذ تصمیمات قابل‌دفاع و بهینه‌سازی جریان‌های کاری استفاده می‌کنند. سازمان‌ها می‌توانند با استفاده از روش‌های هوش مصنوعی و داده‌های تاریخی، فرصت‌ها را شناسایی کنند، ریسک‌ها را کاهش دهند، روندهای آینده را پیش‌بینی کنند و عملکرد را بهبود بخشند.

۵ یادگیری ماشین الگوریتم‌ها برای پیشگیری از حملات سایبری

الگوریتم‌های یادگیری ماشین در پیشگیری از حملات سایبری با شناسایی الگوها در ترافیک شبکه و فعالیت‌های سیستمی ضروری هستند. الگوریتم‌های یادگیری ماشین که همراه با روش‌های سنتی امنیت سایبری به کار گرفته می‌شوند، یک سیستم امنیتی چندلایه در برابر تهدیدات سایبری فراهم می‌کنند. با پیش‌بینی نوع حملاتی که احتمال وقوع آنها وجود دارد، مدل‌های جدید یادگیری عمیق مانند LSTM، RNN و MLP می‌توانند به کاهش دشواری فزاینده حملات سایبری کمک کنند. نتایج امیدوارکننده‌ای از اعتبارسنجی مجموعه داده CTF به‌دست آمده است، به‌ویژه زمانی که یک مدل LSTM به معیار F بالای ۹۳ درصد دست یابد [۱۹]. امنیت ورزشگاه‌ها می‌تواند با استفاده از سیستم سایبر-فیزیکی مبتنی بر هوش مصنوعی (CPS-AI) که بر پیش‌بینی حملات سایبری و ناهنجاری‌های شبکه تمرکز دارد، بهبود یابد. مدل پیشنهادی CPS-AI که داده‌های جمع‌آوری شده را تحلیل می‌کند، نسبت‌های پیش‌بینی و دقت خوبی را به دست می‌آورد و وعده بهبود اقدامات امنیت سایبری در محیط‌های رویداد ورزشی را می‌دهد [۲۰].

رگرسیون لجستیک و مدل‌های زبان بزرگ، مدل‌های یادگیری ماشین هستند که برای شناسایی و پیشگیری از حملات سایبری مورد استفاده قرار می‌گیرند [۲۱].

رگرسیون لجستیک یک الگوریتم طبقه‌بندی داده‌ها است که بر اساس ویژگی‌های ورودی مدل‌سازی می‌کند و احتمال وقوع یک رویداد را محاسبه می‌کند. این الگوریتم در صنایع مختلفی مانند بهداشت و درمان، مالی، بازاریابی و امنیت سایبری مورد استفاده قرار می‌گیرد. در امنیت سایبری، رگرسیون لجستیک برای طبقه‌بندی رویدادهای غیرمعمول سیستم یا ترافیک شبکه به عنوان خوب یا مخرب به کار می‌رود [۲۲]. این طبقه‌بندی بر اساس چندین ویژگی مانند نوع پروتکل، آدرس IP و اندازه بسته انجام می‌شود. داده‌های لاگ سیستمی برای بررسی احتمال وقوع یک رویداد امنیتی، مانند نشت داده‌ها یا نفوذ به کار می‌رود. با تحلیل ویژگی‌های ایمیل، می‌توان به شناسایی فیشینگ پرداخت و به‌طور کلی، رگرسیون لجستیک یک الگوریتم قابل تفسیر و قابل تطبیق است [۲۲].

مدل‌های بزرگ زبانی می‌توانند ابزارهای ارزشمندی در امنیت سایبری باشند. پیاده‌سازی مدل‌های بزرگ زبانی در امنیت سایبری تصمیم‌گیری، روش‌های تشخیص خودکار و تخصص انسانی را تقویت می‌کند که به پیشگیری از تهدیدات سایبری کمک می‌کند [۲۳]. مدل‌های بزرگ زبانی داده‌ها را از منابع مختلفی، از جمله انجمن‌ها و گزارش‌های تهدید، تحلیل می‌کنند و داده‌های مرتبط را از این گزارش‌ها استخراج می‌کنند، مانند بدافزار و آسیب‌پذیری‌ها. این ویژگی اطلاعات تهدید مدل‌های بزرگ زبانی در تصمیم‌گیری کمک می‌کند. فیشینگ نیز می‌تواند با یکپارچه‌سازی مدل‌های بزرگ زبانی در روش‌های امنیت سایبری، مانند تشخیص تهدید مبتنی بر متن، جایی که آدرس‌های وب مشکوک، ایمیل‌های فیشینگ و بدافزار شناسایی می‌شوند، اجتناب شود. مدل‌های بزرگ زبانی چنین الگوهای مشکوکی را شناسایی کرده و برای بررسی‌های بیشتر علامت‌گذاری می‌کنند. با این حال، ادغام چنین مدل‌هایی در صنایع مختلفی مانند بهداشت و درمان و مالی می‌تواند به پیشگیری از نشت داده‌ها و تهدیدات سایبری کمک کند [۲۳].

خلاصه‌ای از درصد موفقیت روش‌های یادگیری ماشین در امنیت سایبری در جدول ۱ مورد بررسی قرار گرفته است.

۱.۵ چالش‌ها و محدودیت‌های هوش مصنوعی

۱.۱.۵ سازگاری با تهدیدات در حال تحول

سازمان‌ها باید با مسائلی که امنیت سایبری هنوز با آن‌ها مواجه است، علی‌رغم پیشرفت در استراتژی‌های پیشگیرانه، مقابله کنند. حمله‌کنندگان دائماً روش‌های جدیدی برای دورزدن سیستم‌های امنیتی پیدا می‌کنند که به این معناست که ریسک‌های سایبری همواره در حال تغییر هستند. سازمان‌ها باید به طور مداوم رویه‌های امنیتی خود را به‌روز کنند تا با تهدیدات در حال افزایش به‌خوبی مقابله کنند که این امر مستلزم حفظ اطلاعات مربوط به تهدیدات و سرمایه‌گذاری در راه‌حل‌های امنیتی پیشرفته است درحالی‌که فناوری پیشرفت می‌کند و حمله‌کنندگان ماهرتر می‌شوند.

۲.۱.۵ حریم خصوصی داده‌ها

شرکت‌ها به طور عمده به الگوریتم‌های یادگیری ماشین و یادگیری عمیق اتکا دارند تا از داده‌های تأمین شده توسط دستگاه‌ها، بینش‌ها و الگوهای مفید استخراج کنند. انتقال داده‌های حساس به یک مکان مرکزی برای

جدول ۱: روش‌های یادگیری ماشین در امنیت سایبری

منبع	سال	مدل استفاده شده	نتیجه
[۲۴]	۲۰۱۷	تکنیک یادگیری ماشین (ML)	حمله DoS را با دقت ۹۹/۷٪ شناسایی کرد.
[۲۵]	۲۰۱۹	تکنیک یادگیری ماشین (ML)	یک مجموعه داده جدید به نام Bot-IoT معرفی و روش استخراج ویژگی پیشنهاد شد.
[۲۶]	۲۰۲۰	مدل امنیتی IntruDTree مبتنی بر یادگیری ماشین (ML)	تعداد ویژگی‌های استفاده شده با انتخاب ویژگی‌های مهم کاهش یافته است.
[۲۷]	۲۰۲۰	شبکه‌های یادگیری عمیق، درخت تصمیم، و ماشین بردار پشتیبان	نرخ دقت ۹۵/۳٪ تا ۹۹/۶۶٪ با استفاده از تکنیک‌های ذکر شده در مجموعه‌های داده مختلف به دست آمد.
[۲۸]	۲۰۲۰	تکنیک یادگیری ماشین (ML)	بر اساس نتایج، USML بهترین روش شناسایی حمله از نظر دقت، نرخ هشدار کاذب و سایر معیارها است.
[۲۹]	۲۰۲۱	SABADT مبتنی بر یادگیری ماشین (ML)	تعداد ویژگی‌های مورد استفاده کاهش یافت و نرخ دقت ۹۹/۱۷٪ در شناسایی حملات به دست آمد.
[۳۰]	۲۰۲۱	یک چارچوب جدید مبتنی بر یادگیری ماشین (ML)	ارائه برای مقابله با تهدیدات امنیتی شبکه نوظهور در SDN
[۳۱]	۲۰۲۲	معماری دو مؤلفه‌ای یادگیری ماشین (ML)	نرخ دقت ۹۹/۸۱٪ در مجموعه داده‌های Bot-IoT به دست آمد.
[۳۲]	۲۰۲۰	تکنیک LSTM	دقت ۹۶/۹۶٪
[۳۳]	۲۰۲۲	تکنیک یادگیری ماشین (ML)	نرخ دقت ۹۹٪ در مجموعه داده N-BaloT به دست آمد.
[۳۴]	۲۰۲۳	Mayfly مبتنی بر یادگیری ماشین (ML)	نتایج بهبود یافته‌ای در معیارهای مختلف به نمایش گذاشته شد.

آموزش، یک ریسک کلیدی مرتبط با این فرایند است که نیاز به آموزش منظم الگوریتم‌ها با داده‌های کلان جمع‌آوری شده از شرکت‌ها و مکان‌های مختلف دارد. این به‌خاطر این است که احتمال دارد کاربرانی غیرمجاز و هکرها بتوانند به داده‌های خصوصی شرکت‌ها دسترسی پیدا کنند [۱۸].

۳.۱.۵ مقیاس‌پذیری

برای آموزش مؤثر الگوریتم‌های هوش مصنوعی، به حجم زیادی از داده‌ها نیاز است. حجم داده‌های تولید شده در امنیت سایبری می‌تواند بسیار بزرگ باشد و از منابع مختلفی از جمله فعالیت‌های کاربری، رویدادهای سیستمی و سوابق ترافیک شبکه ناشی شود. مسائلی در مقیاس‌پذیری ممکن است به دلیل پیچیدگی روزافزون پردازش و تجزیه و تحلیل داده‌ها به علت افزایش مقدار آن‌ها به وجود آید [۱۸].

۴.۱.۵ همکاری انسان-ماشین

در حالی که هوش مصنوعی (AI) مزایای قابل توجهی در تقویت دفاع‌های سایبری ارائه می‌دهد، همچنین چالش‌های پیچیده‌ای را به‌ویژه در همکاری انسان-ماشین به همراه دارد. ادغام سیستم‌های AI در عملیات امنیت سایبری نیازمند تعادل بین خودکارسازی و نظارت انسانی است [۲۳]. این تعادل نه تنها برای بهره‌برداری از تمام پتانسیل هوش مصنوعی بلکه برای محافظت در برابر عواقب ناخواسته‌ای که ممکن است سیستم‌های خودکار به همراه داشته باشند، حیاتی است. همکاری مؤثر انسان-ماشین تضمین می‌کند که سیستم‌های AI به‌عنوان ابزارهای تقویتی برای حرفه‌ای‌های امنیت سایبری عمل کنند و قابلیت‌های تصمیم‌گیری آن‌ها را افزایش دهند نه اینکه آن‌ها را جایگزین کنند. علاوه بر این، نظارت انسانی در تفسیر و زمینه‌سازی هشدارهای تولید شده توسط AI ضروری است، چرا که در غیر این صورت ممکن است به مثبت‌های کاذب منجر شود یا جزئیات ظریف تهدیدات سایبری را نادیده بگیرد [۱۸]. با تقویت یک رابطه هم‌افزا بین هوش مصنوعی و تخصص انسانی، سازمان‌ها می‌توانند وضعیت سایبری خود را تطبیق‌پذیرتر، پاسخگوتر و مقاوم‌تر بسازند. این رویکرد همکاری همچنین راه‌هایی برای یادگیری و بهبود مستمر مدل‌های AI باز می‌کند، زیرا بازخورد انسانی می‌تواند در پالایش الگوریتم‌های AI مؤثر باشد و تضمین کند که آن‌ها در برابر منظر در حال تحول تهدیدات سایبری مؤثر باقی بمانند؛ بنابراین، توسعه چارچوب‌ها و شیوه‌هایی که همکاری انسان-ماشین را افزایش دهد، برای تحقق کامل مزایای هوش مصنوعی در امنیت سایبری در حالی که از خطرات بالقوه کاسته می‌شود، بسیار حائز اهمیت است [۲۱]، [۲۲]. در ادامه طبق مطالعه [۳۵] در جدول ۲ مزایا و معایب استفاده از هوش مصنوعی در امنیت سایبری بیان شده است.

جدول ۲: مزایا و معایب هوش مصنوعی [۳۵]

مزایا	معایب
پردازش حجم عظیمی از داده‌ها به وسیله هوش مصنوعی	جمع‌آوری حجم زیادی داده‌ها برای آموزش هوش مصنوعی می‌تواند منجر به مسائل حریم خصوصی و حفاظت می‌شود.
با راه‌حل‌های امنیت سایبری فعال هوش مصنوعی می‌توانند هر گونه تغییراتی را شناسایی کنند تا خطرات را از بین ببرند.	هکرها می‌توانند از هوش مصنوعی برای راه‌اندازی حملات پیچیده و در مقیاس بزرگ استفاده کنند.
هوش مصنوعی قابلیت خودکارسازی ایجاد الگوریتم‌ها برای شناسایی خطرات در امنیت سایبری	هوش مصنوعی می‌تواند به هکرها کمک کند تا به طور مؤثرتر نقاط آسیب‌پذیری‌ها را پیدا کرده و از آن‌ها بهره‌برداری کنند.
هوش مصنوعی با نظارت بر زیرساخت فناوری اطلاعات به منظور شناسایی موجودیت‌های مخرب و تلاش‌های نفوذ به شبکه کمک می‌کند.	روش‌های هوش مصنوعی می‌توانند توسط کشورهای سرکوبگر و دولت‌ها برای ردیابی مخالفان مورد استفاده قرار گیرند.
به محققان امنیت سایبری این امکان را می‌دهد که بر روی توسعه الگوریتم‌ها کار کنند یا تهدیدات نوظهور را بررسی کنند.	هوش مصنوعی می‌تواند برای نظارت بر حریم خصوصی شخصی، ردیابی و سایر نقض‌ها مورد سوء استفاده قرار گیرد.

۶ آینده و فرصت‌ها

با تحولاتی که هوش مصنوعی در امنیت سایبری به وجود می‌آورد، روندها و فرصت‌های جدید و هیجان‌انگیزی در حال ظهور هستند. این بخش به بررسی مسیرهای امیدوارکننده برای آینده می‌پردازد و به معرفی حوزه‌های کلیدی و چالش‌های احتمالی می‌پردازد.

۱.۶ روندهای نوظهور در امنیت سایبری مبتنی بر هوش مصنوعی

- یادگیری ماشینی مقابله‌ای: تکنیک‌های یادگیری ماشینی مقابله‌ای به عنوان یک حوزه مهم تحقیق در امنیت سایبری مبتنی بر هوش مصنوعی در حال ظهور هستند و بر توسعه دفاع‌های قوی در برابر حملات مقابله‌ای که الگوریتم‌های هوش مصنوعی را هدف قرار می‌دهند، متمرکز شده‌اند.
- امنیت بدون اعتماد (Zero Trust): پذیرش معماری‌های امنیت بدون اعتماد در حال افزایش است که به نیاز به حفاظت در برابر تهدیدات پیچیده سایبری و حملات داخلی پاسخ می‌دهد. فناوری‌های هوش مصنوعی در پیاده‌سازی و عملیاتی‌کردن اصول امنیت بدون اعتماد از طریق نظارت و تأیید مداوم رفتارهای کاربران و دستگاه‌ها نقشی حیاتی دارند.

۲.۶ فرصت‌های تحقیق و توسعه بیشتر

- هوش مصنوعی قابل توضیح: پیشرفت تحقیقات در زمینه هوش مصنوعی قابل توضیح برای بهبود شفافیت و قابلیت تفسیر سیستم‌های امنیت سایبری مبتنی بر هوش مصنوعی ضروری است. پژوهشگران با توسعه تکنیک‌هایی برای توضیح فرایندهای تصمیم‌گیری هوش مصنوعی به زبان قابل فهم، می‌توانند اعتماد و اطمینان در الگوریتم‌های هوش مصنوعی را در میان متخصصان امنیت و ذی‌نفعان افزایش دهند.
- هوش مصنوعی حفظ‌کننده حریم خصوصی: تحقیق در زمینه هوش مصنوعی حفظ‌کننده حریم خصوصی به توسعه تکنیک‌ها و روش‌شناسی‌هایی می‌پردازد که به الگوریتم‌های هوش مصنوعی اجازه می‌دهد تا محاسبات پیچیده را بدون نقض حریم خصوصی داده‌های حساس انجام دهند. با ادغام تکنیک‌های حفظ حریم خصوصی در تدابیر امنیتی مبتنی بر هوش مصنوعی، پژوهشگران می‌توانند به نگرانی‌های حریم خصوصی و الزامات انطباق بپردازند بدون اینکه عملکرد یا دقت را قربانی کنند.

۳.۶ پیش‌بینی‌ها برای آینده ادغام هوش مصنوعی و امنیت سایبری

- افزایش اتوماسیون: انتظار می‌رود که ادغام هوش مصنوعی در امنیت سایبری به افزایش اتوماسیون فرایندهای امنیتی منجر شود و به سازمان‌ها این امکان را بدهد که تهدیدات سایبری را به طور مؤثر و کارآمد شناسایی، پاسخ دهند و کاهش دهند.

• هوش تهدیدات مبتنی بر هوش مصنوعی: پلتفرم‌های هوش تهدیدات مبتنی بر هوش مصنوعی بینش‌های قابل‌اجرایی در مورد تهدیدات سایبری نوظهور به سازمان‌ها ارائه خواهند داد و به استراتژی‌های دفاعی پیشگیرانه و فعالیت‌های شکار تهدیدات کمک خواهند کرد. در نتیجه، پرداختن به ملاحظات اخلاقی و حریم خصوصی، بررسی روندهای نوظهور و بهره‌برداری از فرصت‌های تحقیق و توسعه بیشتر برای تحقق پتانسیل کامل هوش مصنوعی در امنیت سایبری و تضمین آینده دیجیتالی ایمن‌تر ضروری است.

۷ نتیجه‌گیری

در نتیجه، پتانسیل قابل‌توجهی برای بهبود دفاع‌ها در برابر چشم‌انداز دائماً در حال تغییر تهدیدات سایبری از طریق ادغام هوش مصنوعی در امنیت سایبری وجود دارد، به‌ویژه از طریق مدل‌سازی پیش‌بینی و الگوریتم‌های یادگیری ماشین. اما همچنین چالش‌های جدیدی را به همراه دارد که نیاز به مطالعه و بهبود مداوم دارند. اثربخشی هوش مصنوعی در پیش‌بینی و جلوگیری از حملات سایبری اثبات شده است، اما توسعه‌های جدید در این حوزه می‌تواند روندهای امنیت سایبری را به طور بیشتری متحول کند. تحقیقات آینده باید بر روی حوزه‌های کلیدی تمرکز کند تا پتانسیل راه‌حل‌های امنیت سایبری مبتنی بر هوش مصنوعی به طور کامل تحقق یابد. با استفاده از آخرین پیشرفت‌ها در یادگیری ماشین و یادگیری عمیق، ممکن است روش‌های مدل‌سازی پیش‌بینی پیچیده‌تری توسعه یابند که امکان شناسایی و پاسخ به تهدیدات را با دقت بیشتری فراهم کند. تحقیق در مورد تحلیل‌های پیش‌بینی در زمان واقعی بسیار مهم است زیرا می‌تواند بینش‌های سریع ارائه دهد و به مقابله با خطرات نوظهور بپردازد. رسیدگی به محدودیت‌های مقیاس‌پذیری و حریم خصوصی داده‌ها که در آموزش مدل‌های هوش مصنوعی بر روی مجموعه داده‌های بزرگ وجود دارد، همچنان یک چالش حیاتی است. یادگیری Federated که به مدل‌ها اجازه می‌دهد از منابع داده غیرمتمرکز یاد بگیرند در حالی که ناشناسی را حفظ می‌کند، یک حوزه تحقیقاتی امیدوارکننده است. مطالعه محاسبات امن چندطرفه و استراتژی‌های حفظ حریم خصوصی تفاضلی می‌تواند قابلیت‌های حفظ حریم خصوصی سیستم‌های هوش مصنوعی در امنیت سایبری را بهبود بخشد.

موضوع دیگری که به‌شدت نیاز به تحقیق دارد، کاربرد اخلاقی هوش مصنوعی در امنیت سایبری است. اطمینان از اینکه مدل‌های پیش‌بینی منجر به افزایش تعصب نشده یا استانداردهای حریم خصوصی را نقض نکنند، بسیار حیاتی است. تحقیق در زمینه توسعه مدل‌های هوش مصنوعی شفاف و قابل توضیح می‌تواند به کاهش این نگرانی‌ها کمک کرده و اعتماد ذی‌نفعان را به انتخاب‌های مبتنی بر هوش مصنوعی افزایش دهد. ادغام هوش مصنوعی با فناوری‌های موجود در امنیت سایبری، مانند سیستم‌های تشخیص نفوذ و دروازه‌های برخط امن، فرصت‌های رشد جدیدی را ایجاد می‌کند. این ادغام بی‌دردسر ممکن است واکنش‌پذیری و سازگاری تدابیر امنیت سایبری را بهبود بخشد و منجر به دفاعی قوی‌تر در برابر حملات سایبری پیچیده شود.

علاوه بر این، ایجاد یک محیط همکاری برای دانشگاه‌ها و صنعت به‌منظور به‌اشتراک‌گذاری دانش، داده‌ها

و بهترین شیوه‌ها بسیار مهم است. چنین همکاری می‌تواند توسعه راه‌حل‌های موفق امنیت سایبری مبتنی بر هوش مصنوعی را تسریع کند و به شرکت‌ها در سراسر جهان کمک کند. با توجه به این ملاحظات، رشد مداوم تهدیدات سایبری نیاز به یک رویکرد پیشگیرانه در امنیت سایبری دارد. استفاده از هوش مصنوعی و یادگیری ماشین، همراه با تحقیق دقیق در مورد روش‌های پیشرفته و ملاحظات اخلاقی، می‌تواند به کسب‌وکارها کمک کند تا دفاع‌های خود را در برابر تهدیدات آینده افزایش دهند. مسیر پیشرو پیچیده و پر از چالش است، اما همچنین وعده یک دنیای دیجیتال امن‌تر را می‌دهد که توسط فناوری‌های پیشرفته محافظت می‌شود.

مراجع

- [1] Lacka, E., & Wong, T. C. Examining the impact of digital technologies on students' higher education outcomes: the case of the virtual learning environment and social media. *Studies in higher education*, 46(8), 1621-1634, 2021.
- [2] Urbinati, A., Chiaroni, D., Chiesa, V., & Frattini, F. The role of digital technologies in open innovation processes: an exploratory multiple case study analysis. *R&D Management*, 50(1), 136-160, 2020.
- [3] Sontan, A. D., & Samuel, S. V. The intersection of Artificial Intelligence and cybersecurity: Challenges and opportunities. *World Journal of Advanced Research and Reviews*, 21(2), 1720-1736, 2024.
- [4] A. Mtair, "Predictions of Cybersecurity Experts on Future Cyber-Attacks and Related Cybersecurity Measures", (IJACSA) *International Journal of Advanced Computer Science and Applications*, vol. 14, no. 2, 2023.
- [5] Nadella, G. S., & Gonaygunta, H. Enhancing Cybersecurity with Artificial Intelligence: Predictive Techniques and Challenges in the Age of IoT. *International Journal of Science and Engineering Applications*, 13(04), 30-33, 2024.
- [6] Jha, N., & Ramaprabhu, S. Thermal conductivity studies of metal dispersed multiwalled carbon nanotubes in water and ethylene glycol based nanofluids. *Journal of applied physics*, 106(8), 2009.
- [7] Javed, S. H., Ahmad, M. B., Asif, M., Almotiri, S. H., Masood, K., & Ghamdi, M. A. A. An intelligent system to detect advanced persistent threats in industrial internet of things (IIoT). *Electronics*, 11(5), 742, 2023.
- [8] Boukhalfa, A., Hmina, N., & Chaoni, H. Parallel processing using big data and machine learning techniques for intrusion detection. *IAES International Journal of Artificial Intelligence*, 9(3), 553, 2020.
- [9] Costin, A. T., Zinca, D., & Dobrota, V. A Real-Time Streaming System for Customized Network Traffic Capture. *Sensors*, 23(14), 6467, 2023.
- [10] Paolanti, M., & Frontoni, E. Multidisciplinary pattern recognition applications: A review. *Computer Science Review*, 37, 100276, 2020.

- [11] Kaur, R., Gabrijelčič, D., & Klobučar, T. Artificial intelligence for cybersecurity: Literature review and future research directions. *Information Fusion*, 97, 101804, 2023.
- [12] Dini, P., & Saponara, S. Analysis, design, and comparison of machine-learning techniques for networking intrusion detection. *Designs*, 5(1), 9, 2021.
- [13] A. Bécue, I. Praça, and J. Gama, "Artificial intelligence, cyber-threats and Industry 4.0: challenges and opportunities", *Artificial Intelligence Review*, vol. 54, no. 5, pp. 3849–3886, Feb. 2021, DOI: <https://doi.org/10.1007/s10462-020-09942-2>.
- [14] S. Dilek, H. Cakır, and M. Aydın, "Applications of Artificial Intelligence Techniques to Combating Cyber Crimes: A Review", *International Journal of Artificial Intelligence & Applications*, vol. 6, no. 1, pp. 21–39, Jan. 2015, DOI: <https://doi.org/10.5121/ijaia.2015.6102>
- [15] Dr. A. M. Shamiulla, "Role of Artificial Intelligence in Cyber Security", *International Journal of Innovative Technology and Exploring Engineering*, vol. 9, no. 1, pp. 4628–4630, Nov. 2019, DOI: <https://doi.org/10.35940/ijitee.a6115.119119>.
- [16] H. Gonaygunta, D. Kumar, S. Maddini, and S. F. Rahman, "How can we make IoT Applications better with Federated Learning- A Review", *IJARCCCE*, vol. 12, no. 2, Feb. 2023, DOI: <https://doi.org/10.17148/ijarcce.2023.12213>.
- [17] A. Nilsson, S. Smith, G. Ulm, E. Gustavsson, and M. Jirstrand, "A performance evaluation of federated learning algorithms", Rennes, France: Association for Computing Machinery, 2018, pp. 1–8. DOI: <https://doi.org/10.1145/3286490.3286559>.
- [18] H. Gonaygunta, D. Kumar, S. Maddini, and S. F. Rahman, "How can we make IoT Applications better with Federated Learning- A Review", *IJARCCCE*, vol. 12, no. 2, Feb. 2023, DOI: <https://doi.org/10.17148/ijarcce.2023.12213>.
- [19] B. Fredj, A. Mihoub, M. Krichen, O. Cheikhrouhou, and A. Derhab, "CyberSecurity attack prediction: A deep learning approach", Merkez, Turkey: Association for Computing Machinery, 2021. DOI: <https://doi.org/10.1145/3433174.3433614>.
- [20] B. Wan, C. Xu, R. P. Mahapatra, and P. Selvaraj, "Understanding the Cyber-Physical System in International Stadiums for Security in the Network from Cyber-Attacks and Adversaries using AI", *Wireless Personal Communications*, vol. 127, no. 2, pp. 1207–1224, May 2021, DOI: <https://doi.org/10.1007/s11277-021-08573-2>.
- [21] H. Gonaygunta, "MACHINE LEARNING ALGORITHMS TO DETECT CYBER THREATS 1 Factors Influencing the Adoption of Machine Learning Algorithms to Detect Cyber Threats in the Banking Industry", 2023.
- [22] Hari Gonaygunta Using and Regression, "Article 6", *International Journal of Smart Sensors and Ad Hoc Networks*, vol. 3, no. 4, DOI: <https://doi.org/10.47893/IJSSAN.2023.1229>
- [23] K. Meduri, H. Gonaygunta, G. S. Nadella, P. P. Pawar, and D. Kumar, "Adaptive Intelligence: GPT-Powered Language Models for Dynamic Responses to Emerging Healthcare Challenges", *IJARCCCE*, vol. 13, no. 1, Jan. 2024, DOI: <https://doi.org/10.17148/ijarcce.2024.13114>.

- [24] Z. He, T. Zhang, and R. B. Lee, "Machine learning based DDoS attack detection from source side in cloud", in Proc. IEEE 4th Int. Conf. Cyber Secur. Cloud Comput. (CSCloud), Jun. 2017, pp. 114–120.
- [25] J. Alsamiri and K. Alsubhi, "Internet of Things cyber attacks detection using machine learning", Int. J. Adv. Comput. Sci. Appl., vol. 10, no. 12, pp. 627–634, 2019.
- [26] I. H. Sarker, Y. B. Abushark, F. Alsolami, and A. I. Khan, "IntruDTree: A machine learning based cyber security intrusion detection model", Symmetry, vol. 12, no. 5, p. 754, May 2020.
- [27] K. Shaukat, S. Luo, S. Chen, and D. Liu, "Cyber threat detection using machine learning techniques: A performance evaluation perspective", in Proc. Int. Conf. Cyber Warfare Secur. (ICCWS), Oct. 2020, pp. 1–6.
- [28] T. A. Tuan, H. V. Long, L. H. Son, R. Kumar, I. Priyadarshini, and N. T. K. Son, "Performance evaluation of botnet DDoS attack detection using machine learning", Evol. Intell., vol. 13, no. 2, pp. 283–294, Jun. 2020.
- [29] M. Ozkan-Okay, Ö. Aslan, R. Eryigit, and R. Samet, "SABADT: Hybrid intrusion detection approach for cyber attacks identification in WLAN", IEEE Access, vol. 9, pp. 157639–157653, 2021.
- [30] Z. A. El Houda, A. S. Hafid, and L. Khoukhi, "A novel machine learning framework for advanced attack detection using SDN", in Proc. IEEE Global Commun. Conf. (GLOBE-COM), Dec. 2021, pp. 1–6.
- [31] A. Mihoub, O. B. Fredj, O. Cheikhrouhou, A. Derhab, and M. Krichen, "Denial of service attack detection and mitigation for Internet of Things using looking-back-enabled machine learning techniques", Comput. Electr. Eng., vol. 98, Mar. 2022, Art. no. 107716.
- [32] A. Makkar and N. Kumar, "An efficient deep learning-based scheme for Web spam detection in IoT environment", Future Gener. Comput. Syst., vol. 108, pp. 467–487, Jul. 2020.
- [33] M. Waqas, K. Kumar, A. A. Laghari, U. Saeed, M. M. Rind, A. A. Shaikh, F. Hussain, A. Rai, and A. Q. Qazi, "Botnet attack detection in Internet of Things devices over cloud environment via machine learning", Concurrency Comput., Pract. Exper., vol. 34, no. 4, Feb. 2022, Art. no. e6662.
- [34] F. Alrowais, S. Althahabi, S. S. Alotaibi, A. Mohamed, M. A. Hamza, and R. Marzouk, "Automated machine learning enabled cyber security threat detection in Internet of Things environment", Comput. Syst. Sci. Eng., vol. 45, no. 1, pp. 687–700, 2023.
- [35] De Azambuja, A. J. G., Plesker, C., Schützer, K., Anderl, R., Schleich, B., & Almeida, V. R., "Artificial intelligence-based cyber security in the context of industry 4.0—a survey. Electronics, 12(8), 1920, 2023.
- [36] De Azambuja, A. J. G., Plesker, C., Schützer, K., Anderl, R., Schleich, B., & Almeida, V. R., "Artificial intelligence-based cyber security in the context of industry 4.0—a survey. Electronics, 12(8), 1920, 2023.