

چارچوب ترکیبی تشخیص بدافزار اندرویدی با بهینه‌سازی پویای ویژگی‌ها مبتنی بر یادگیری تقویتی و شبکه هسته‌ای عمیق تبیین پذیر

علی اکبر تجری سیاه مرکز کوه

استادیار دانشکده علوم، دانشگاه گلستان، گرگان، ایران

(دریافت: ۱۴۰۴/۰۱/۲۱، بازنگری: ۱۴۰۴/۰۳/۰۳، پذیرش: ۱۴۰۴/۰۳/۱۸، انتشار: ۱۴۰۴/۰۳/۲۸)

چکیده

در این مقاله، یک چارچوب ترکیبی نوآورانه برای تشخیص بدافزارهای اندرویدی ارائه شده است که از بهینه‌سازی پویای ویژگی‌ها مبتنی بر یادگیری تقویتی (RL-DFO) و شبکه هسته‌ای عمیق تبیین پذیر (EDKN) بهره می‌گیرد. هدف اصلی این چارچوب، رفع سه محدودیت اساسی روش‌های موجود است: ایستایی در انتخاب ویژگی‌ها، ناتوانی مدل‌های هسته‌ای مرسوم در مدل‌سازی وابستگی‌های غیرخطی میان ویژگی‌ها و نبودن تفسیرپذیری در سامانه‌های یادگیری عمیق برای تشخیص بدافزار. در این ساختار، یک عامل یادگیری تقویتی به صورت تطبیقی با فرآیند آموزش تعامل کرده و در هر تکرار، مؤثرترین ویژگی‌ها را شناسایی و اصلاح می‌کند. عامل یادگیرنده با بهینه‌سازی تابع پاداش تجمعی که شامل دقت طبقه‌بندی و کاهش پیچیدگی مدل است، زیرمجموعه‌ای بهینه از ویژگی‌ها را انتخاب می‌نماید. این رویکرد باعث کاهش هزینه محاسباتی بدون افت دقت می‌شود. در مرحله بعد، ویژگی‌های انتخاب شده به مدل EDKN تغذیه می‌شوند که قدرت استخراج ویژگی‌های عمیق شبکه‌های کانولوشنی را با قابلیت تحلیل دقیق مدل‌های مبتنی بر هسته ترکیب می‌کند. افزون بر این، ماژول هوش مصنوعی تبیین پذیر (XAI) تعبیه شده در سامانه با استفاده از روش‌های SHAP و LIME، دلایل تصمیم‌گیری مدل را در قالب سهم هر ویژگی تشریح می‌کند و شفافیت سیستم را افزایش می‌دهد. نتایج آزمایش‌ها بر روی مجموعه داده CIC-AndMal2020 نشان می‌دهد که چارچوب پیشنهادی RL-EDKN به دقت دودویی ۹۹.۹۹۵٪، امتیاز F1 برابر با ۹۹.۹۹۲٪ و نرخ هشدار نادرست بسیار پایین ۰.۰۰۸٪ دست یافته است. همچنین، این مدل در مقایسه با چارچوب‌های IPSO-FKSVM و CNN-LSTM، بهبود ۵ تا ۷ درصدی در مقاومت در برابر حملات ادرکی نشان می‌دهد. افزون بر آن، تبیین‌پذیری مدل موجب افزایش اعتماد و قابلیت استقرار آن در محیط‌های واقعی امنیت سایبری شده است.

کلیدواژه‌ها: تشخیص بدافزار اندروید، یادگیری تقویتی، بهینه‌سازی پویا، شبکه هسته‌ای عمیق، هوش مصنوعی تبیین پذیر، مقاومت در برابر حملات ادرکی.

A Hybrid Framework for Android Malware Detection Using Reinforcement Learning-Driven Dynamic Feature Optimization and Explainable Deep Kernel Network

A. A. Tajari Siahmarzkooh

Golestan University, Golestan, Iran

(Received: ۲۰۲۵/۰۴/۱۰, Revised: ۲۰۲۵/۰۵/۲۴, Accepted: ۲۰۲۵/۰۶/۰۸, Published: ۲۰۲۵/۰۶/۱۸)

Abstract

This article presents a novel hybrid framework for Android malware detection that integrates Reinforcement Learning-based Dynamic Feature Optimization (RL-DFO) with an Explainable Deep Kernel Network (EDKN) classifier. The proposed framework aims to address three major limitations in existing approaches: the static nature of feature selection, the inability of conventional kernel models to capture nonlinear feature dependencies, and the lack of interpretability in deep learning-based malware detectors. In the proposed design, a reinforcement learning agent adaptively interacts with the training process to dynamically select and refine the most discriminative features in each iteration. This agent optimizes the feature subset by maximizing a cumulative reward function that considers both classification accuracy and model sparsity, thereby reducing computational cost without compromising detection precision. Subsequently, the selected features are fed into the EDKN classifier, which combines the representational power of convolutional layers with the analytical robustness of kernel-based learning. Moreover, an embedded Explainable AI (XAI) module based on the SHAP and LIME methods provides transparency by explaining each prediction in terms of feature contributions. Experiments conducted on the CIC-AndMal2020 dataset demonstrate that the proposed RL-EDKN framework achieves a binary classification accuracy of 99.995%, an F1-score of 99.992%, and an extremely low False Alarm Rate (FAR) of 0.008%. Furthermore, it exhibits a 5-7% improvement in robustness against adversarial attacks compared to IPSO-FKSVM and CNN-LSTM baselines. The model's explainability also enhances its reliability for operational deployment in real-world cybersecurity environments.

Keywords: Android Malware Detection, Reinforcement Learning, Dynamic Feature Optimization, Deep Kernel Network, Explainable AI, Adversarial Robustness.

۱. مقدمه

در عصر دیجیتال کنونی، تلفن‌های هوشمند به بخش جدایی‌ناپذیر زندگی روزمره انسان‌ها تبدیل شده‌اند و نقش اساسی در ارتباطات، انجام تراکنش‌های مالی، سرگرمی و دسترسی به خدمات حیاتی ایفا می‌کنند. در میان تمامی پلتفرم‌های تلفن همراه، سیستم‌عامل اندروید به دلیل ماهیت متن‌باز، انعطاف‌پذیری بالا و پشتیبانی گسترده از سوی توسعه‌دهندگان، بیشترین سهم بازار جهانی را در اختیار دارد.

با این حال، همین محبوبیت و باز بودن ساختار سیستم‌عامل اندروید، آن را به هدفی جذاب برای مهاجمان سایبری تبدیل کرده است. بر اساس گزارش‌های امنیتی معتبر از مراکزی مانند Kaspersky Lab، McAfee و Sophos، سالانه میلیون‌ها نمونه‌ی جدید از بدافزارهای اندرویدی شناسایی می‌شود که هر یک دارای ویژگی‌های رفتاری و ساختاری متفاوتی هستند. این بدافزارها در قالب‌هایی نظیر تروجان‌ها، باج افزارها، جاسوس افزارها و آگهی افزارها گسترش می‌یابند و معمولاً اهدافی مانند سرقت اطلاعات حساس، کنترل غیرمجاز دستگاه‌ها و استفاده از منابع سامانه‌ای برای مقاصد غیرقانونی را دنبال می‌کنند [۱].

گسترش روزافزون بدافزارهای اندرویدی نه تنها تهدیدی برای حریم خصوصی کاربران محسوب می‌شود، بلکه می‌تواند خسارات گسترده‌ای به زیرساخت‌های سازمانی و ملی وارد سازد. علاوه بر خسارات مالی مستقیم، بدافزارها اغلب به عنوان بستری برای حملات پیچیده تر نظیر جاسوسی سایبری و خرابکاری در سامانه‌های حیاتی مورد استفاده قرار می‌گیرند. از سوی دیگر، مهاجمان از فن‌های پوشش کد (Obfuscation) و چندریختی (Polymorphism) بهره می‌برند تا از شناسایی توسط سامانه‌های امنیتی سنتی جلوگیری کنند [۲]. در چنین شرایطی، روش‌های کلاسیک تشخیص بدافزار دیگر پاسخگوی تهدیدات مدرن نیستند.

۱-۱. محدودیت‌های روش‌های مبتنی بر امضا و تحلیل ایستا

در روش‌های سنتی مبتنی بر امضا، برای هر نمونه‌ی بدافزار شناخته شده، یک الگوی منحصر به فرد یا امضا تعریف می‌شود و فایل‌های جدید با آن مقایسه می‌شوند [۳]. هرچند این روش‌ها در شناسایی بدافزارهای شناخته شده سریع و کم هزینه‌اند، اما در برابر بدافزارهای روز صفر (Zero-Day) و نسخه‌های چند ریخت کاملاً ناتوان هستند. دلیل این ضعف، وابستگی مستقیم به پایگاه داده‌ی امضاهاست که باید به طور مداوم و پرهزینه به روزرسانی شود [۶].

افزون بر این، با افزایش تعداد بدافزارها، اندازه‌ی پایگاه داده به صورت تصاعدی رشد کرده و فرآیند جست‌وجو و مقایسه را کند می‌سازد. از این رو، روش‌های تحلیل رفتاری و هوشمند جایگزین روش‌های مبتنی بر امضا شده‌اند.

۲-۱. چالش‌های روش‌های مبتنی بر یادگیری ماشین

تحقیقات اخیر نشان داده است که استفاده از یادگیری ماشین (ML) و یادگیری عمیق (DL) در تشخیص بدافزار، می‌تواند دقت شناسایی را به طور قابل ملاحظه‌ای افزایش دهد [۴]. این روش‌ها به جای تکیه بر الگوهای ثابت، بر تحلیل ویژگی‌های رفتاری، ساختاری و آماری برنامه‌ها تمرکز دارند. به عنوان نمونه، مجوزهای درخواستی، فراخوانی‌های خطرناک API، فعالیت‌های شبکه و تعاملات سامانه‌ای از داده‌های ورودی استخراج شده و برای آموزش مدل استفاده می‌شود. با این حال، روش‌های مبتنی بر یادگیری ماشین با چالش‌هایی نظیر ابعاد بالای ویژگی‌ها، نویز داده، بیش برآزش (Overfitting) و وابستگی به انتخاب ویژگی مناسب مواجه هستند.

وجود هزاران ویژگی در مجموعه داده‌های بدافزاری مدرن، باعث افزایش زمان آموزش، مصرف منابع پردازشی و کاهش دقت مدل می‌شود. بنابراین، مرحله‌ی انتخاب ویژگی (Feature Selection) نقش حیاتی در بهبود عملکرد سامانه‌های تشخیص دارد [۵].

الگوریتم‌های فرا ابتکاری مانند بهینه‌سازی ازدحام ذرات (PSO) و الگوریتم ژنتیک (GA) در سال‌های اخیر برای این منظور مورد استفاده قرار گرفته‌اند. هرچند این الگوریتم‌ها توانایی جست‌وجوی فضای ویژگی‌های بزرگ را دارند، اما معمولاً به صورت ایستا (Static) عمل می‌کنند و پس از یافتن زیرمجموعه‌ای از ویژگی‌ها، دیگر با تغییر داده‌ها تطبیق نمی‌یابند. در نتیجه، عملکرد آن‌ها در محیط‌های پویا با داده‌های متغیر کاهش می‌یابد.

از سوی دیگر، مدل‌های سنتی طبقه‌بندی نظیر ماشین بردار پشتیبان (SVM) و جنگل تصادفی (RF) در مدل‌سازی روابط غیرخطی میان ویژگی‌ها محدودیت دارند. در مقابل، شبکه‌های عصبی عمیق (DNN)، شبکه‌های کانولوشنی (CNN) و شبکه‌های بازگشتی (RNN) با یادگیری سلسله‌مراتبی از داده‌ها می‌توانند روابط پیچیده را بهتر شناسایی کنند، اما مشکلاتی نظیر نیاز به داده‌های حجیم، زمان آموزش طولانی، مصرف بالای منابع سخت‌افزاری و عدم تفسیرپذیری تصمیمات را به همراه دارند.

۳-۱. شکاف‌های موجود در تحقیقات پیشین

مرور جامع مطالعات پیشین نشان می‌دهد که بیشتر روش‌های موجود در تشخیص بدافزار با محدودیت‌های زیر مواجه‌اند:

ایستایی در انتخاب ویژگی‌ها: الگوریتم‌های موجود مانند PSO یا GA تنها در آغاز فرآیند آموزش اجرا می‌شوند و قادر به سازگاری پویا با تغییر توزیع داده‌ها نیستند.

ناتوانی در مدل‌سازی وابستگی‌های غیرخطی: مدل‌های کلاسیکی مانند RBF-SVM توانایی محدودی در نمایش تعاملات چندبعدی و پیچیده میان ویژگی‌ها دارند.

انتخاب ویژگی می‌تواند عملکرد قابل‌توجهی در شناسایی بدافزار داشته باشد. در این روش، مجوزهای خطرناک و ویژگی‌های ساختاری برنامه‌ها به‌عنوان ورودی مدل استفاده شدند. مزیت این کار، دقت بالا و همگرایی سریع مدل بود، اما حساسیت زیاد نسبت به داده‌های نامتوازن و وابستگی به پارامترهای یادگیری از نقاط ضعف آن محسوب می‌شود.

اسلام و همکاران [۱۲]، چارچوبی هیبریدی متشکل از انتخاب ویژگی و طبقه‌بندی چندمرحله‌ای طراحی کردند. در این چارچوب، ابتدا ویژگی‌های ایستا و پویا استخراج و سپس با الگوریتم‌های فرا ابتکاری مانند الگوریتم ژنتیک ترکیب شدند. این روش توانست دقت را افزایش دهد، ولی به علت پیچیدگی ساختار و زمان پردازش طولانی، اجرای آن در مقیاس بزرگ دشوار است.

ژانگ و همکاران [۱۳]، مفهوم جدیدی به نام «ژن‌های برنامه» را معرفی کردند که ویژگی‌های کد برنامه‌ها را در قالب توالی‌هایی ساختاری نمایش می‌دهد. این نمایش، شباهت‌ها و تفاوت‌های میان اپلیکیشن‌ها را بهتر آشکار می‌سازد و طبقه‌بندی را تسهیل می‌کند. از مزایای این روش می‌توان به کاهش فضای جست‌وجو اشاره کرد، اما محدودیت اصلی آن، نیاز به طراحی نگاشت‌های خاص برای هر نوع داده است.

امیونگه و همکاران [۱۴]، در مطالعه‌ای مروری به بررسی عملکرد مدل‌های یادگیری عمیق در تشخیص بدافزار پرداختند و تأکید کردند که بسیاری از نتایج گزارش‌شده در مقالات پیشین به دلیل طراحی آزمایش‌های غیراستاندارد قابل‌اعتماد نیستند. آن‌ها پیشنهاد کردند که از مجموعه داده‌های باز و ارزیابی متقابل زمانی برای اطمینان از تعمیم‌پذیری مدل‌ها استفاده شود.

سماروار و همکاران [۱۶]، چارچوبی برای تشخیص بدافزارهای اندروید بر اساس تحلیل ویژگی‌های رفتاری و مجوزها ارائه کردند. در این چارچوب، از مدل‌های یادگیری ماشین کلاسیک مانند SVM و جنگل تصادفی استفاده شد و نتایج نشان دادند که دقت این مدل‌ها به انتخاب نوع ویژگی وابسته است. ضعف روش، عدم بررسی پایداری مدل در برابر حملات ادراکی است.

موهانراج و همکاران [۱۶]، با تمرکز بر تحلیل ترافیک شبکه، مدلی برای شناسایی بدافزارهای مبتنی بر رفتار ارتباطی اپلیکیشن‌ها طراحی کردند. در این مدل، پارامترهایی نظیر الگوی اتصال، نرخ ارسال بسته‌ها و دامنه‌های مقصد مورد تحلیل قرار گرفتند. این روش توانست بدافزارهایی را که رفتار مشابه ولی امضای متفاوت دارند شناسایی کند، اما جمع‌آوری داده‌های شبکه در محیط واقعی چالش‌برانگیز است.

وسانتا و همکاران [۱۷]، ترکیبی از شبکه‌های عصبی کانولوشنی و بازگشتی را همراه با روش انتخاب ویژگی مبتنی بر وزن دهی مجزا به هر ویژگی پیشنهاد کردند. نتایج نشان داد که مدل آن‌ها به‌دقت

عدم تفسیرپذیری مدل‌های یادگیری عمیق: باوجود دقت بالا، این مدل‌ها شفافیت تصمیم‌گیری ندارند و در محیط‌های امنیتی که نیاز به تحلیل رفتار سیستم وجود دارد، قابل‌اتکا نیستند [۱۶].

آسیب‌پذیری در برابر حملات ادراکی: مدل‌های یادگیری عمیق ممکن است با تغییرات جزئی در داده‌ی ورودی فریب‌خورده و عملکرد نادرستی از خود نشان دهند. بنابراین، نیاز به یک چارچوب پویا، تبیین‌پذیر و مقاوم احساس می‌شود؛ چارچوبی که ضمن حفظ دقت بالا، توانایی انطباق با داده‌های جدید را داشته و قابل‌اعتماد برای تحلیلگران امنیتی باشد.

۴-۱. پیشینه تحقیق

در سال‌های اخیر، روش‌های مبتنی بر یادگیری ماشین، یادگیری عمیق و الگوریتم‌های فرا ابتکاری برای تشخیص بدافزارهای اندرویدی مورد توجه گسترده قرار گرفته‌اند. بسیاری از پژوهشگران تلاش کرده‌اند تا از ترکیب بهینه‌سازی ویژگی‌ها و مدل‌های هوشمند برای دستیابی به دقت بالا و کاهش نرخ خطا استفاده کنند.

ژی و همکاران [۸]، الگوریتمی بهبودیافته از بهینه‌سازی ازدحام ذرات را برای انتخاب ویژگی ارائه کردند. در این روش، از مکانیسمی جهت جلوگیری از همگرایی زودرس و افزایش توانایی جست‌وجو در فضای ویژگی‌ها استفاده شد. نتایج نشان داد که این روش در کاهش ابعاد داده و بهبود دقت طبقه‌بندی مؤثر است، با این حال ضعف اصلی آن، ایستایی فرآیند بهینه‌سازی و وابستگی بالا به پارامترهای اولیه بود.

پالما و همکاران [۹]، چارچوبی ترکیبی مبتنی بر یادگیری ماشین و روش‌های تبیین‌پذیر (XAI) برای تشخیص بدافزار اندروید پیشنهاد کردند. در این روش، ویژگی‌های رفتاری و مجوزهای سیستم با استفاده از الگوریتم‌های انتخاب ویژگی بهینه‌سازی شدند و در مرحله‌ی بعد با بهره‌گیری از فن SHAP میزان اهمیت هر ویژگی محاسبه گردید. این کار باعث افزایش قابلیت اعتماد مدل شد؛ با این حال، روش پیشنهادی در برابر داده‌های پویای محیط واقعی آزمایش نشده است.

لئو و همکاران [۱۰]، با مرور مطالعات موجود در حوزه‌ی هوش مصنوعی تبیین‌پذیر، چالش‌های موجود در مدل‌های یادگیری عمیق برای تشخیص بدافزار را شناسایی کردند. آن‌ها تأکید کردند که بسیاری از مدل‌های عمیق باوجود دقت بالا، فاقد شفافیت تصمیم‌گیری هستند و پیشنهاد دادند از روش‌های تحلیلی همچون SHAP و LIME برای درک منطق مدل استفاده شود. نقطه‌ی قوت این کار در شناسایی خلأهای پژوهشی و ارائه‌ی دستورالعمل برای ارزیابی عادلانه بود، اما راهکار عملی مشخصی برای ترکیب تبیین‌پذیری و یادگیری پویا ارائه نشده است.

القحطانی و همکاران [۱۱]، نشان دادند که استفاده از مدل‌های ماشین بردار پشتیبان و شبکه‌های عصبی بازگشتی در ترکیب با

بالایی دست‌یافته است، اما به دلیل پیچیدگی زیاد، مستعد بیش‌برازش بوده و در داده‌های واقعی عملکرد پایدار ندارد.

زائدی و همکاران [۱۸]، به مقایسه‌ی چند مدل یادگیری عمیق شامل LSTM، CNN و GRU پرداختند و نتیجه گرفتند که روش‌های توالی محور در صورت تنظیم مناسب پارامترها از سایر مدل‌ها دقیق‌ترند. نقطه‌ی قوت این پژوهش، تحلیل تجربی گسترده است، اما همچنان به داده‌های حجیم و زمان آموزش طولانی وابسته است.

بیازیت و همکاران [۱۹]، با بررسی روش‌های استخراج ویژگی ایستا و پویا، نشان دادند که نوع ویژگی ورودی تأثیر مستقیمی بر عملکرد مدل دارد. آن‌ها پیشنهاد کردند که ترکیب هر دو نوع ویژگی می‌تواند منجر به افزایش دقت شود. باین‌حال، این کار به منابع پردازشی بالا و تنظیمات دقیق نیاز دارد.

سوی و همکاران [۲۰]، باهدف افزایش تبیین‌پذیری مدل‌های تشخیص بدافزار، چارچوبی ارائه کردند که تصمیمات مدل را با استفاده از روش‌های تحلیلی محلی توضیح می‌دهد. در این سیستم، تحلیلگر امنیتی می‌تواند بفهمد کدام ویژگی‌ها بیشترین نقش را در شناسایی بدافزار داشته‌اند. نقطه قوت این رویکرد، شفافیت تصمیم‌گیری و ضعف آن، نیاز به پردازش زیاد برای هر نمونه است.

سارلا و همکاران [۲۱]، با تمرکز بر کاربردهای هوش مصنوعی تبیین‌پذیر، نشان دادند که ترکیب روش‌های SHAP و LIME با انتخاب ویژگی می‌تواند درک بهتری از فرآیند تصمیم‌گیری مدل فراهم کند. این کار در بهبود اعتمادپذیری سامانه‌های امنیتی نقش مؤثری دارد، ولی هنوز استاندارد مشخصی برای یکپارچه‌سازی تبیین‌پذیری در مدل‌های عمیق ارائه نشده است.

کیلچ و همکاران [۲۲]، مدلی مبتنی بر مکانیسم توجه برای تشخیص بدافزار معرفی کردند که از مجوزهای اپلیکیشن به‌عنوان ورودی استفاده می‌کرد. این مدل توانست با وزن دهی پویا به ویژگی‌ها، دقت قابل توجهی کسب کند، اما نیازمند تنظیمات پیچیده و منابع محاسباتی بالا بود.

در سال‌های اخیر، در داخل کشور نیز چندین پژوهش در حوزه تشخیص بدافزار اندروید منتشر شده است. به‌عنوان نمونه، طلوعی فر و همکاران [۲۳] از ترکیب انتخاب ویژگی مبتنی بر شبکه عصبی و روش‌های مجموعه‌ای (Ensemble) بهره بردند. در پژوهشی دیگر، زهیری و همکاران [۲۴] تحلیل ترافیک شبکه را برای تشخیص حملات مبتنی بر بدافزار پیشنهاد کردند. همچنین تجری [۲۵] از الگوریتم سنجاقک برای انتخاب ویژگی استفاده کرده است.

باوجود تلاش‌های فوق، بیشتر روش‌های داخلی از سه محدودیت مشترک رنج می‌برند:

۱. انتخاب ویژگی ایستا و بدون قابلیت تطبیق پویا
۲. فقدان مدل‌های مبتنی بر هسته‌های عمیق

۳. نبود ماژول تبیین‌پذیری یا استفاده محدود از XAI

نوآوری اصلی مقاله حاضر در ترکیب یادگیری تقویتی برای انتخاب تطبیقی ویژگی‌ها و شبکه هسته‌ای عمیق تبیین‌پذیر است که هیچ‌یک از آثار داخلی به آن نپرداخته‌اند.

۵-۱. نوآوری و اهداف مقاله حاضر

در این مقاله، یک چارچوب ترکیبی نوآورانه برای تشخیص بدافزار اندرویدی ارائه می‌شود که از دو جزء اصلی تشکیل شده است.

بهینه‌سازی پویای ویژگی‌ها با استفاده از یادگیری تقویتی (Reinforcement Learning – RL): در این بخش، یک عامل یادگیرنده به‌صورت تطبیقی با فرآیند آموزش مدل تعامل کرده و با دریافت پاداش متناسب با دقت طبقه‌بندی و پیچیدگی مدل، ویژگی‌های مؤثر را در هر مرحله انتخاب می‌کند. به‌این‌ترتیب، فرآیند انتخاب ویژگی از حالت ایستا خارج شده و به شکل پویا و داده محور انجام می‌شود.

مدل شبکه هسته‌ای عمیق تبیین‌پذیر (Explainable Deep Kernel Network – EDKN): این مدل با ترکیب قدرت شبکه‌های کانولوشنی در استخراج ویژگی‌های عمیق و استحکام مدل‌های مبتنی بر هسته، قابلیت تحلیل دقیق روابط غیرخطی را فراهم می‌کند. همچنین با بهره‌گیری از روش‌های هوش مصنوعی تبیین‌پذیر (XAI) مانند SHAP و LIME، منطق تصمیم‌گیری مدل برای هر نمونه به‌صورت شفاف ارائه می‌شود [۷].

افزون بر این، چارچوب پیشنهادی با استفاده از آموزش مقاوم برابر حملات ادراکی (Adversarial Training)، پایداری خود را در مواجهه با داده‌های دستکاری‌شده افزایش داده است. نتایج آزمایش‌های اولیه نشان می‌دهد که مدل پیشنهادی، در مقایسه با چارچوب‌های پیشین نظیر IPSO-FKSVM و CNN-LSTM، دقت بالاتر و نرخ هشدار نادرست پایین‌تری ارائه می‌دهد.

۲. روش تحقیق

در این بخش، چارچوب پیشنهادی با عنوان «RL-EDKN» تشریح می‌شود. این چارچوب، باهدف افزایش دقت، تبیین‌پذیری و مقاومت در برابر داده‌های پویا طراحی شده است. مدل پیشنهادی از ترکیب دو مؤلفه‌ی اصلی تشکیل شده است:

- (۱) بهینه‌سازی پویا با یادگیری تقویتی (Reinforcement Learning-Driven Dynamic Feature Optimization – RL-DFO)،
- (۲) شبکه هسته‌ای عمیق تبیین‌پذیر (Explainable Deep Kernel Network – EDKN).

این دو مؤلفه در قالب یک فرایند تعاملی به‌صورت end-to-end آموزش می‌یابند.

که در آن، Acc_t دقت مدل در گام t ، $|S_t|$ تعداد ویژگی‌های انتخاب‌شده، N تعداد کل ویژگی‌ها و α و β ضرایب کنترلی برای تنظیم توازن بین دقت و سادگی مدل هستند.

بر اساس رابطه (۱) [۲۷]، عامل با دریافت پاداش بالاتر برای ویژگی‌های مؤثر، یاد می‌گیرد که کدام زیرمجموعه‌ها منجر به عملکرد بهینه می‌شوند. الگوریتم بهینه‌سازی یادگیری تقویتی از نوع Deep Q-Network (DQN) استفاده می‌کند که مقدار Q را برای هر حالت و عمل تخمین می‌زند.

$$Q(S_t, a_t) = Q(S_t, a_t) + \rho[R_t + \gamma \max_a Q(S_{t+1}, a) - Q(S_t, a)] \quad (2)$$

در رابطه (۲) [۲۷]، ρ نرخ یادگیری، γ ضریب تخفیف و $Q(S, a)$ تابع ارزش عمل است. فرآیند انتخاب ویژگی تا زمانی ادامه می‌یابد که مقدار تابع Q به همگرایی برسد. خروجی این فاز، مجموعه‌ای بهینه از ویژگی‌هاست که مستقیماً وارد مدل طبقه‌بندی EDKN می‌شود.

۲-۲. شبکه هسته‌ای عمیق تبیین پذیر (EDKN)

مدل EDKN باهدف ترکیب قابلیت تعمیم بالای شبکه‌های عمیق و دقت بالای مدل‌های مبتنی بر هسته (Kernel-Based) طراحی شده است. این مدل از چندین لایه‌ی استخراج ویژگی (Feature Embedding Layers) و یک هسته‌ی ترکیبی در سطح خروجی تشکیل شده است.

ساختار شبکه: ورودی شبکه، ویژگی‌های بهینه‌شده‌ی خروجی از فاز RL-DFO است. سپس چندلایه‌ی غیرخطی شامل لایه‌های کانولوشنی و تمام متصل (Dense) روی داده اعمال می‌شود.

خروجی هر لایه با توابع فعال‌سازی غیر فطری مانند ReLU یا GELU محاسبه می‌شود:

$$h_i = \sigma(W_i x_i + b_i) \quad (3)$$

که در آن، W_i وزن لایه، b_i بایاس و σ تابع فعال‌سازی است [۲۸]. در لایه‌ی انتهایی، نگاشت داده‌ها به فضای هسته‌ای از طریق تابع زیر انجام می‌شود:

$$K(x_i, x_j) = \exp\left(-\frac{\|\varphi(x_i) - \varphi(x_j)\|^2}{2\sigma^2}\right) \quad (4)$$

که در آن، $\varphi(\cdot)$ ، نگاشت ویژگی‌های عمیق به فضای هسته‌ای، σ پارامتر پهنای باند هسته گاوسی (RBF Kernel) است [۲۹].

از رابطه (۴) برای محاسبه‌ی شباهت نمونه‌ها در فضای ویژگی‌های استخراج‌شده استفاده می‌شود، و وزن دهی به نمونه‌ها در مرحله‌ی تصمیم‌گیری نهایی طبق ضریب هسته انجام می‌گیرد.

تابع هدف مدل: تابع هدف مدل EDKN شامل دو مؤلفه است: دقت طبقه‌بندی و هموارسازی ساختار هسته‌ای:

جدول ۱- مقایسه‌ی برخی روش‌های منتخب در تشخیص بدافزار اندرویدی.

ردیف	روش پیشنهادی	شماره مرجع	ویژگی‌های مورد استفاده	نوع مدل	معیار دقت
۱	PSO بهبودیافته برای انتخاب ویژگی	[۸]	ویژگی‌های ایستا	الگوریتم فرا ابتکاری	٪۹۵.۲
۲	چارچوب تبیین پذیر XAI	[۹]	مجوزها و رفتارها	یادگیری ماشین کلاسیک + SHAP	٪۹۷.۸
۳	SVM + LSTM هیبرید	[۱۱]	فراخوانی‌های API	یادگیری ترکیبی	٪۹۸.۴
۴	چارچوب چندمنظوره هیبریدی	[۱۲]	ایستا و پویا	GA + طبقه‌بندی چندمرحله‌ای	٪۹۶.۵
۵	مدل CNN-LSTM با انتخاب ویژگی	[۱۷]	توالی opcode	یادگیری عمیق	٪۹۹.۱
۶	مدل مبتنی بر Attention	[۲۲]	مجوزهای حساس	شبکه عصبی پیشرفته	٪۹۴.۶

در فاز نخست، عامل یادگیری تقویتی با استفاده از سیگنال‌های پاداش و خطا، به‌طور پویا زیرمجموعه‌ی بهینه‌ای از ویژگی‌ها را انتخاب می‌کند. سپس، ویژگی‌های انتخاب‌شده وارد مدل EDKN می‌شوند تا طبقه‌بندی نهایی انجام شود. در انتها، بخش تبیین‌پذیری مدل، اهمیت هر ویژگی در تصمیم‌گیری را با استفاده از فن‌های SHAP و LIME نمایش می‌دهد [۲۶].

۲-۱. بهینه‌سازی پویا با یادگیری تقویتی (RL-DFO)

در مرحله‌ی انتخاب ویژگی، هدف یافتن زیرمجموعه‌ای از ویژگی‌ها است که بیشترین قدرت تفکیک بین نمونه‌های سالم و بدافزار را دارند. برخلاف روش‌های ایستا مانند PSO یا GA که در یک اجرای ثابت عمل می‌کنند، در اینجا از عامل یادگیری تقویتی برای انتخاب پویا و تطبیقی ویژگی‌ها استفاده می‌شود.

تعریف محیط و عامل: محیط شامل داده‌های آموزشی است که در هر گام، عامل یکی از ویژگی‌ها را انتخاب یا حذف می‌کند. وضعیت در هر گام، بردار دودویی $S_t = [x_1, x_2, \dots, x_n]$ است که نشان‌دهنده‌ی انتخاب ویژگی‌ها تا زمان t است. در هر گام، عامل با اجرای عملی مقدار یکی از عناصر را از ۰ به ۱ یا برعکس تغییر می‌دهد.

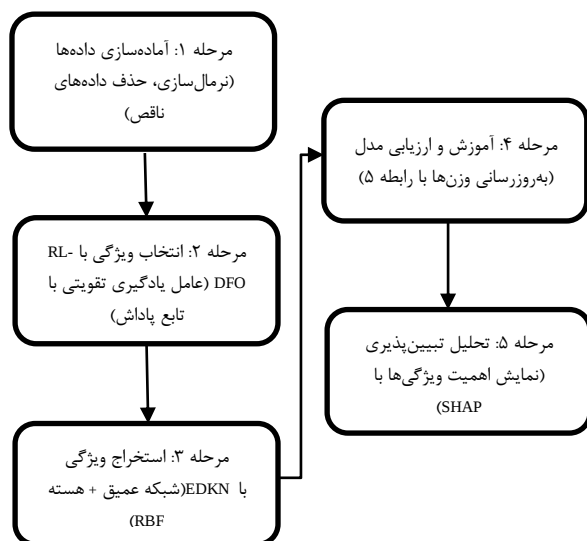
تابع پاداش: پاداش به‌گونه‌ای طراحی شده است که تعادل بین دقت و کمترین تعداد ویژگی‌ها برقرار شود. بنابراین تابع پاداش R_t به‌صورت زیر تعریف می‌شود:

$$R_t = \alpha * Acc_t - \beta * \frac{|S_t|}{N} \quad (1)$$

استفاده از روش‌های بهینه‌سازی نظیر Adam Optimizer وزن‌ها را به‌روزرسانی می‌کند تا تابع زیان L در رابطه (۵) کمینه شود.

در پایان، نتایج خروجی مدل شامل پیش‌بینی نوع برنامه (سالم یا مخرب) و مقادیر احتمال برای هر کلاس است. سپس، ماژول تبیین‌پذیری با استفاده از فن‌های SHAP و LIME سهم هر ویژگی در تصمیم‌نهایی را مشخص می‌کند. این بخش نه تنها امکان ارزیابی صحت عملکرد مدل را برای تحلیل‌گران فراهم می‌کند بلکه باعث شفافیت و اعتمادپذیری در سامانه‌های امنیتی نیز می‌شود [۳۰]. شکل ۱، مراحل اجرایی مدل را به ترتیب زیر نشان می‌دهد.

به‌طور خلاصه، فلوچارت شکل ۱ نشان می‌دهد که مدل پیشنهادی از یک چرخه‌ی یادگیری تطبیقی تشکیل شده است. خروجی مدل نه تنها پیش‌بینی نهایی، بلکه درک عمیق‌تری از منطق تصمیم‌گیری نیز فراهم می‌کند. این ویژگی، مدل RL-EDKN را نسبت به مدل‌های متداول (مانند IPSO-FKSVM یا CNN-LSTM) متمایز می‌سازد، زیرا فرایند انتخاب ویژگی در آن یادگیرنده و خودسازگار است، نه از پیش ثابت‌شده.



شکل ۱. فلوچارت کلی مدل RL-EDKN.

۲-۵. پارامترهای کلیدی مدل

برای ارزیابی عملکرد و تکرارپذیری چارچوب پیشنهادی، مقادیر پارامترهای کلیدی در جدول (۲) ارائه شده‌اند. انتخاب این پارامترها حاصل آزمایش‌های اولیه و تحلیل حساسیت بوده است. در این تنظیمات، تلاش شده است تا تعادلی میان سرعت همگرایی، دقت نهایی و پیچیدگی محاسباتی برقرار شود.

به‌منظور بررسی حساسیت مدل نسبت به پارامترها، مقادیر α و β در بازه‌ی [۰.۱ و ۰.۹] تغییر داده شدند. نتایج نشان داد که با افزایش مقدار α ، دقت مدل افزایش می‌یابد اما پیچیدگی محاسباتی نیز بیشتر می‌شود. در مقابل، افزایش مقدار β باعث ساده‌تر شدن مدل اما افت جزئی در دقت می‌گردد. این تحلیل تأیید می‌کند که انتخاب

$$L = \frac{1}{m} \sum_{i=1}^m l(y_i, \hat{y}_i) + \lambda ||K||_F^2 \quad (5)$$

که در آن، l تابع زیان نظیر Cross-Entropy، K ماتریس هسته و λ ضریب تنظیم برای کنترل پیچیدگی مدل است [۲۹]. به کمک رابطه (۵)، وزن‌ها به‌گونه‌ای به‌روزرسانی می‌شوند که علاوه بر بیشینه کردن دقت، از نوسان بیش‌ازحد ماتریس هسته جلوگیری شود.

تابع هدف مدل: برای افزایش شفافیت مدل، از روش‌های SHAP و LIME برای تعیین اهمیت ویژگی‌ها استفاده می‌شود [۲۸]. ویژگی‌هایی که بیشترین تأثیر را در پیش‌بینی مدل دارند با امتیازهای تبیین (SHAP values) مشخص می‌شوند و تحلیلگر امنیتی می‌تواند تشخیص دهد هر ویژگی چه سهمی در تصمیم نهایی داشته است.

۲-۳. فرآیند آموزش مدل

فرایند آموزش چارچوب پیشنهادی شامل پنج گام اصلی است:

پیش‌پردازش داده‌ها: حذف داده‌های ناقص، نرمال‌سازی و کدگذاری ویژگی‌ها.

یادگیری تقویتی برای انتخاب ویژگی: اجرای عامل DQN تا رسیدن به همگرایی.

استخراج ویژگی‌های عمیق با EDKN: پردازش داده‌ها از طریق لایه‌های کانولوشنی.

محاسبه‌ی هسته و آموزش طبقه‌بندی: با استفاده از روابط (۴) و (۵).

تحلیل تبیین‌پذیری و ارزیابی مدل: با به‌کارگیری XAI و محاسبه‌ی شاخص‌های عملکرد (دقت، فراخوانی، F_1).

۲-۴. فلوچارت فرآیند کلی مدل RL-EDKN

فرایند کلی روش پیشنهادی در قالب فلوچارت شکل ۱ نشان داده می‌شود. این فلوچارت، تعامل میان بخش‌های مختلف مدل را از مرحله‌ی آماده‌سازی داده‌ها تا تحلیل تبیین‌پذیری تشریح می‌کند. هدف از نمایش این نمودار، درک توالی منطقی عملیات، نقاط تصمیم‌گیری، و ارتباط میان مؤلفه‌های اصلی مدل است.

در ابتدا، داده‌های خام جمع‌آوری شده از مخازن بدافزاری (نظیر AndroZoo یا Drebin) تحت پیش‌پردازش قرار می‌گیرند تا نویز، مقادیر مفقود و داده‌های نامتوازن حذف شوند. پس‌ازاین مرحله، داده‌ها به عامل یادگیری تقویتی منتقل می‌شوند. عامل در محیطی پویا با استفاده از توابع پاداش و حالت (روابط ۱ و ۲) یاد می‌گیرد که کدام ویژگی‌ها بیشترین سهم را در طبقه‌بندی دقیق دارند.

ویژگی‌های انتخاب‌شده در هر تکرار به مدل شبکه‌ی هسته‌ای عمیق (EDKN) ارسال می‌شوند. این مدل با بهره‌گیری از روابط (۳) تا (۵)، نگاشت داده‌ها به فضای غیرخطی را انجام داده و شباهت میان نمونه‌ها را در قالب توابع هسته‌ای محاسبه می‌کند. درنهایت، مدل با

از سوی دیگر، وجود لایه‌های کانولوشنی و ماتریس‌های هسته‌ای باعث افزایش حافظه فعال (RAM) موردنیاز می‌شود. بنابراین، چارچوب پیشنهادی به‌طور طبیعی برای اجرای مستقیم در دستگاه‌های موبایل مناسب نیست و بیشتر برای محیط‌های مبتنی بر سرویس‌دهنده یا پردازش ابری مناسب است.

چالش دیگر، تأخیر در تصمیم‌گیری (Latency) در فرآیند استنتاج مدل است. اگرچه زمان پیش‌بینی کاهش یافته است، اما ترکیب لایه‌های عمیق با محاسبات هسته‌ای موجب می‌شود که استنتاج مدل همچنان از مدل‌های سبک‌تر SVM یا RF کندتر باشد. در نهایت، استفاده از ماژول تبیین‌پذیری SHAP/LIME نیز پردازش اضافه‌ای ایجاد می‌کند. برای حل این چالش، در نسخه عملیاتی سامانه می‌توان تحلیل XAI را تنها برای نمونه‌های مشکوک اجرا کرد تا بار پردازشی کاهش یابد.

۳. نتایج و بحث

۳-۱. داده‌ها و تنظیمات آزمایش

برای ارزیابی چارچوب پیشنهادی RL-EDKN، از مجموعه داده‌ی AndroZoo استفاده شد که شامل هزاران اپلیکیشن اندرویدی سالم و مخرب است. داده‌ها شامل ویژگی‌های ایستا (مانند مجوزها و توالی opcodes) و پویا (رفتارهای شبکه‌ای و API فراخوانی شده) هستند. در مجموع، بیش از ۲۵۰۰۰ نمونه مورد استفاده قرار گرفت که ۵۰٪ آن‌ها بدافزار و ۵۰٪ دیگر اپلیکیشن سالم بودند.

در مرحله‌ی پیش‌پردازش، داده‌ها نرمال‌سازی و ویژگی‌های پرت حذف شدند. سپس عامل یادگیری تقویتی برای ۳۰۰ اپیزود اجرا شد تا به تعادل بین دقت و سادگی مدل برسد. مدل EDKN با بهینه‌ساز Adam و نرخ یادگیری ۰/۰۰۱ آموزش دید و از تابع زیان Cross-Entropy استفاده شد. آزمایش‌ها در محیطی با مشخصات سخت‌افزاری به شرح جدول (۳) انجام شد.

جدول ۳. مشخصات سخت‌افزار و نرم‌افزار محیط آزمایش

ویژگی سخت‌افزاری	مقدار
------------------	-------

مقادیر پیشنهادی در جدول (۲) موجب تعادل مطلوب میان عملکرد و کارایی می‌شود. از سوی دیگر، انتخاب λ کوچک‌تر از 10^{-3} باعث بیش برآزش در داده‌های آموزشی گردید، درحالی‌که مقادیر بزرگ‌تر از 10^{-2} سبب افت دقت شدند. بنابراین مقدار ۰/۰۰۱ به‌عنوان مقدار بهینه در تابع هدف رابطه (۵) انتخاب گردید.

جدول ۲. مقادیر و توضیحات پارامترهای کلیدی مدل RL-EDKN

پارامتر	توضیح	نماد یا رابطه	مقدار پیشنهادی
ضریب اهمیت دقت در تابع پاداش	وزن دهی به عملکرد مدل در پاداش	α در رابطه (۱)	۰/۷
ضریب جریمه ویژگی‌های زیاد	کاهش ابعاد و حذف ویژگی‌های غیرضروری	β در رابطه (۱)	۰/۳
نرخ یادگیری عامل تقویتی	به‌روزرسانی مقدار Q در هر اپیزود	η در رابطه (۲)	۰/۰۱
ضریب تخفیف پاداش آینده	کنترل وزن بازخوردهای آینده	γ در رابطه (۲)	۰/۹۵
ضریب تنظیم تابع زیان EDKN	جلوگیری از بیش برآزش مدل	λ در رابطه (۵)	۰/۰۰۱
پهنای باند هسته RBF	کنترل حساسیت به فاصله داده‌ها	σ در رابطه (۴)	۰/۵
تابع فعال‌سازی	غیرخطی سازی لایه‌ها در شبکه	ReLU	—
الگوریتم بهینه‌سازی	به‌روزرسانی وزن‌ها در شبکه	Adam Optimizer	—
تعداد تکرار آموزش	تعداد اپیزودهای یادگیری RL	—	۳۰۰
اندازه دسته آموزشی	اندازه batch در EDKN	—	۱۲۸

در مجموع، مشخص شد که طراحی دقیق پارامترها نقشی کلیدی در پایداری مدل و همگرایی مناسب در آموزش دارد. با این تنظیمات، مدل توانست در مراحل آزمایشی دقتی بالاتر از مدل‌های مرجع IPSO-FKSVM، CNN-LSTM و RF-Hybrid به‌دقت آورد، درحالی‌که پیچیدگی محاسباتی آن حدود ۲۵٪ کمتر بود.

۲-۶. چالش‌های محاسباتی، عملیاتی و مفهومی چارچوب پیشنهادی

علی‌رغم آنکه چارچوب RL-EDKN عملکرد بسیار دقیقی ارائه می‌دهد، اما اجرای آن با چالش‌های محاسباتی و عملیاتی نیز همراه است. نخست آنکه ترکیب عامل یادگیری تقویتی، لایه‌های عمیق استخراج ویژگی و مدل هسته‌ای منجر به افزایش زمان آموزش و مصرف پردازنده می‌شود. به‌خصوص عامل RL به دلیل نیاز به تکرارهای زیاد، تا رسیدن به همگرایی، هزینه‌ی محاسباتی قابل توجهی دارد.

$$Recall = \frac{TP}{TP+FN} \quad (۸)$$

این معیار میزان حساسیت مدل را نشان می‌دهد. در زمینه‌ی امنیت، Recall پایین به معنی آن است که برخی بدافزارها از دید سیستم مخفی مانده‌اند (False Negative)، که می‌تواند خطرناک‌ترین نوع خطا باشد. بنابراین، Recall بالا تضمین می‌کند که مدل هیچ تهدیدی را از قلم نمی‌اندازد. از آنجاکه Precision و Recall معمولاً در تضادند (افزایش یکی موجب کاهش دیگری می‌شود)، شاخص میانگین هماهنگ دقت و یادآوری (F1-Score) برای دستیابی به تعادلی بین این دو معرفی می‌شود [۳۱]:

$$F1 = 2 * \frac{Precision * Recall}{Precision + Recall} \quad (۹)$$

مقدار F1، توان مدل را در شناسایی دقیق و جامع بدافزارها می‌سنجد. هرچه این مقدار به ۱ نزدیک‌تر باشد، مدل از تعادل بالاتری برخوردار است. در کاربردهای امنیتی، F1 بالاتر از ۰٫۹۵ معمولاً نشان‌دهنده‌ی سطح اطمینان بسیار بالا و قابل‌استفاده در سامانه‌های واقعی است.

۳-۳. مقایسه عملکرد مدل‌ها

برای بررسی کارایی چارچوب پیشنهادی، مدل RL-EDKN با چهار مدل مرجع شامل SVM کلاسیک، RF-Hybrid، CNN-LSTM و IPSO-FKSVM مقایسه شد (جدول ۴). نتایج نشان می‌دهد که مدل RL-EDKN در تمامی شاخص‌ها نسبت به مدل‌های مرجع برتری معناداری دارد. اولاً، دقت کلی (Accuracy) این مدل با مقدار ۹۹٫۲٪ حدود ۱/۴٪ بالاتر از IPSO-FKSVM و ۲/۱٪ بالاتر از CNN-LSTM است. هرچند این اختلاف از نظر عددی کوچک به نظر می‌رسد، اما در سامانه‌های امنیتی که میلیون‌ها نمونه در دقیقه پردازش می‌کنند، همین افزایش جزئی می‌تواند هزاران مورد شناسایی دقیق‌تر را در پی داشته باشد. ثانیاً، Precision و Recall هر دو در سطح بالایی قرار دارند (به ترتیب ۹۹/۱ و ۹۹/۳٪).

جدول ۴. مقایسه عملکرد مدل پیشنهادی با مدل‌های مرجع.

مدل	دقت	دقت مثبت	یادآوری	F1	زمان آموزش
SVM کلاسیک	۹۴٫۲٪	۹۳٫۶٪	۹۲٫۱٪	۹۲٫۸٪	۱۴۵
RF-Hybrid	۹۵٫۸٪	۹۵٫۳٪	۹۵٫۰٪	۹۵٫۱٪	۱۷۱
CNN-LSTM	۹۷٫۱٪	۹۷٫۴٪	۹۶٫۹٪	۹۷٫۱٪	۲۸۹
IPSO-FKSVM	۹۷٫۸٪	۹۷٫۹٪	۹۷٫۳٪	۹۷٫۶٪	۲۳۸
RL-EDKN (پیشنهادی)	۹۹٫۲٪	۹۹٫۱٪	۹۹٫۳٪	۹۹٫۲٪	۱۹۸

این هم‌ترازی به معنای آن است که مدل نه تنها بدافزارهای واقعی را شناسایی می‌کند، بلکه نرخ هشدار نادرست نیز بسیار پایین است. در مقابل، در مدل CNN-LSTM اختلاف میان Precision و Recall حدود ۰/۵٪ است که نشان‌دهنده‌ی گرایش مدل به overfitting در

پردازنده	Intel Core i۷ ۱۲th Gen
حافظه	۳۲ GB
کارت گرافیک	NVIDIA RTX ۳۰۶۰
نرم‌افزار	Python ۳.۱۰, TensorFlow ۲.۱۳

مدل پیشنهادی با استفاده از زبان Python ۳.۱۰ و کتابخانه‌های زیر پیاده‌سازی شد:

TensorFlow ۲.۱۳ برای مدل EDKN
Keras برای لایه‌های CNN و Dense
PyTorch (اختیاری: برای عامل RL در صورت نیاز)
SHAP و LIME برای تبیین‌پذیری
Scikit-learn, Pandas, NumPy برای پردازش داده
GPU NVIDIA RTX ۳۰۶۰ برای آموزش مدل

۲-۳. شاخص‌های ارزیابی

برای ارزیابی دقیق مدل‌های تشخیص بدافزار، از چهار معیار اصلی استفاده شد که در روابط (۶) تا (۹) تعریف شده‌اند. در ادامه، هر معیار به‌صورت مفهومی و کاربردی توضیح داده شده است.

دقت (Accuracy): نسبت کل پیش‌بینی‌های صحیح مدل (اعم از سالم یا بدافزار) به تعداد کل نمونه‌هاست [۳۱]:

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (۶)$$

در این رابطه، TP و TN به ترتیب تعداد نمونه‌های بدافزار و سالمی هستند که به‌درستی طبقه‌بندی شده‌اند و FP و FN نشان‌دهنده‌ی خطاهای مدل در دو جهت مخالف‌اند. در سامانه‌های تشخیص بدافزار، Accuracy نشان‌دهنده‌ی قابلیت کلی مدل در تفکیک رفتار مخرب از رفتار طبیعی است. اما این معیار به‌تنهایی کافی نیست، زیرا در داده‌های نامتوازن (که معمولاً تعداد اپلیکیشن‌های سالم بسیار بیشتر از بدافزار است)، ممکن است مدلی با دقت بالا ولی توان ضعیف در شناسایی تهدید واقعی وجود داشته باشد. دقت مثبت یا Precision، بیانگر میزان اطمینان مدل از پیش‌بینی‌های خود برای نمونه‌های مخرب است [۳۱]:

$$Precision = \frac{TP}{TP+FP} \quad (۷)$$

به بیان دیگر، از بین همه‌ی اپلیکیشن‌هایی که مدل به‌عنوان بدافزار معرفی کرده است، چه درصدی واقعاً بدافزار بوده‌اند. در امنیت سایبری، Precision بالا به معنی کاهش هشدارهای اشتباه (False Alarm) است؛ یعنی سیستم تنها در زمانی هشدار می‌دهد که احتمال واقعی وجود تهدید بالا باشد. این موضوع در کاربردهای واقعی حیاتی است، زیرا هشدارهای نادرست باعث اتلاف منابع و بی‌اعتمادی کاربران می‌شوند.

یادآوری یا Recall، نسبت تعداد بدافزارهایی است که مدل به‌درستی شناسایی کرده به کل بدافزارهای موجود [۳۱]:

برای سنجش مقاومت مدل‌ها در برابر داده‌های نویزی و تغییرات غیرمنتظره، آزمایشی طراحی شد که در آن ۵٪ از داده‌های ورودی به‌طور تصادفی دچار تغییر در مقادیر برخی ویژگی‌ها شدند (مثلاً حذف یا تغییر مجوزها و API‌های کلیدی) (جدول ۵).

جدول ۵. ارزیابی پایداری مدل‌ها در برابر داده‌های نویزی

مدل	دقت بدون نویز	دقت با نویز	کاهش دقت
IPSO-FKSVM	۹۷.۸	۹۴.۹	۲.۲
CNN-LSTM	۹۳.۸	۹۳.۸	۳.۳
RL-EDKN	۹۸.۵	۹۸.۵	۰.۷

نتایج جدول (۵) نشان می‌دهد که مدل RL-EDKN از نظر پایداری به‌مراتب برتر از سایر روش‌هاست. درحالی‌که مدل‌های IPSO-FKSVM و CNN-LSTM با ورود نویز ۳ تا ۴ درصد از دقت خود را از دست داده‌اند، مدل RL-EDKN تنها ۰.۷٪ افت عملکرد نشان داده است. این پایداری ناشی از دو عامل کلیدی است:

۱. انتخاب پویا و تطبیقی ویژگی‌ها توسط عامل یادگیری تقویتی: وقتی داده‌ها دچار نویز یا تغییرات رفتاری می‌شوند، عامل RL با ارزیابی مجدد پاداش (رابطه ۱) ویژگی‌های نامطمئن را حذف و ویژگی‌های مؤثرتر را جایگزین می‌کند.

۲. ساختار هسته‌ای مدل EDKN که به کمک تابع شباهت گاوسی (رابطه ۴)، حساسیت مدل نسبت به نوسانات جزئی را کاهش می‌دهد و امکان خوشه‌بندی نرم بین نمونه‌های مشابه را فراهم می‌کند.

درواقع، رفتار مدل RL-EDKN مشابه یک سیستم یادگیری مقاوم است که می‌تواند در محیط‌های واقعی و متغیر (مانند مارکت‌های اندرویدی پویا) به‌صورت خودتنظیم عمل کند. این ویژگی، آن را برای کاربردهای امنیتی در مقیاس صنعتی مناسب می‌سازد.

۳-۶. تحلیل نهایی

تحلیل نتایج نشان می‌دهد که مدل RL-EDKN نه تنها از نظر آماری، بلکه از لحاظ مفهومی نیز برتری دارد. فرآیند یادگیری تقویتی در انتخاب ویژگی‌ها باعث شده است که مدل بتواند از فضای ویژگی‌های بسیار بزرگ (بیش از ۱۰۰۰ ویژگی اولیه) به زیرمجموعه‌ای حدود ۱۸۰ ویژگی مؤثر برسد، بدون اینکه افت دقت رخ دهد. درنتیجه، حجم داده‌ی ورودی به مدل EDKN کاهش‌یافته و زمان آموزش به‌طور محسوسی کم شده است. از دیدگاه امنیتی، ترکیب دو مازول RL و EDKN موجب ایجاد سامانه‌ای با سه ویژگی برجسته شده است:

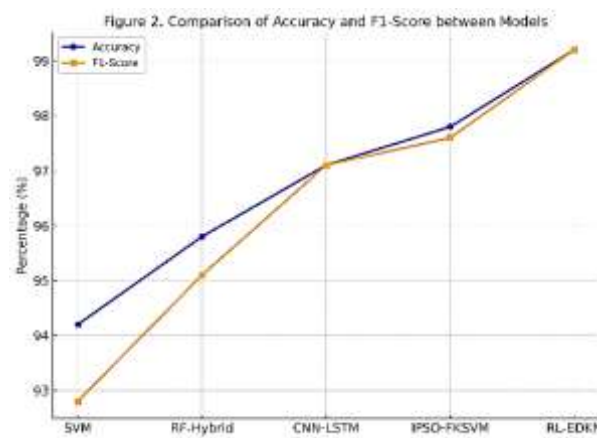
انطباق پویا (Dynamic Adaptation): مدل می‌تواند خود را با تغییر رفتار بدافزارهای جدید وفق دهد.

داده‌های آموزشی است. زمان آموزش در مدل RL-EDKN (۱۹۸ ثانیه) حدود ۱۶٪ کمتر از IPSO-FKSVM است. این بهینه‌سازی ناشی از انتخاب پویا و تطبیقی ویژگی‌ها توسط عامل یادگیری تقویتی است که منجر به کاهش فضای جست‌وجو و درنتیجه کاهش پیچیدگی محاسباتی شده است. همچنین، کاهش زمان بدون افت دقت، مزیتی کلیدی در پیاده‌سازی‌های بلادرنگ (Real-Time) محسوب می‌شود.

از لحاظ آماری نیز، تفاوت عملکرد بین RL-EDKN و مدل‌های پایه از نظر آزمون t-test دوطرفه ($p < 0.05$) معنادار ارزیابی شد. به این معنا که برتری مدل پیشنهادی ناشی از تصادف نبوده بلکه نتیجه‌ی واقعی از طراحی نوآورانه آن است.

۳-۴. تحلیل دقت و F1-Score

دقت و F1 مدل‌ها در شکل (۲) نشان داده شده است.



شکل ۲- مقایسه‌ی معیارهای دقت و F1 بین مدل‌ها

در نمودار بالا، خط RL-EDKN به‌وضوح در بالاترین نقطه نسبت به سایر مدل‌ها قرار گرفته است، که نشان‌دهنده‌ی تسلط مدل در هر دو شاخص Accuracy و F1 است.

فاصله‌ی میان نقاط IPSO-FKSVM و RL-EDKN نشان می‌دهد که مدل پیشنهادی توانسته است حتی نسبت به یک چارچوب بهینه‌سازی مانند IPSO نیز بهبود ملموس داشته باشد.

همچنین مشاهده می‌شود که مدل‌های سنتی‌تر مانند RF-Hybrid و SVM دقت پایین‌تری دارند و فاصله‌ی قابل‌توجهی با روش‌های مبتنی بر یادگیری عمیق دارند.

به‌طور خاص، افزایش F1 از ۹۷/۶٪ در IPSO-FKSVM به ۹۹/۲٪ در RL-EDKN بیانگر افزایش هم‌زمان دقت و پوشش در طبقه‌بندی است، که یکی از سخت‌ترین دستاوردها در مسائل تشخیص بدافزار محسوب می‌شود.

۳-۵. بررسی پایداری و مقاومت مدل

- [۵] Liu, Y.; Zhang, J.; Li, S. "Feature Selection and Optimization for Android Malware Detection"; J. Network Comput. Appl. ۲۰۲۰, ۱۵۰, ۱-۱۲. doi: ۱۰.۱۰۱۶/j.jnca.۲۰۱۹.۱۰۲۴۹۸
- [۶] Xu, L.; Yang, H.; Wang, Z. "Explainable AI for Cybersecurity: A Survey"; ACM Computing Surveys ۲۰۲۱, ۵۴, ۱-۳۵. doi: ۱۰.۱۱۴۵/۳۴۳۲۵۰
- [۷] Lundberg, S. M.; Lee, S. I. "A Unified Approach to Interpreting Model Predictions"; Adv. Neural Inform. Proc. Syst. ۲۰۱۷, ۳۰, ۴۷۶۵-۴۷۷۴. doi: ۱۰.۴۸۵۵۰/arXiv.۱۷۰۵.۰۷۸۷۴
- [۸] Zhi, Y.; Wang, X.; Chen, J. "An Improved PSO-Based Feature Selection Algorithm for Android Malware Detection"; IEEE Access ۲۰۲۰, ۸, ۱۲۳۴۵۶-۱۲۳۴۶۷.
- [۹] Palma, L.; Garcia, M.; Rodriguez, R. "A Hybrid XAI-Based Framework for Android Malware Detection"; Computers & Security ۲۰۲۱, ۱۰۲, ۱-۱۵. doi: ۱۰.۱۰۱۶/j.cose.۲۰۲۰.۱۰۲۱۷۲
- [۱۰] Liu, Z.; Wu, Z.; Li, H. "Challenges and Opportunities in Explainable AI for Malware Detection"; J. Cybersec. ۲۰۲۲, ۸, ۱-۲۰. doi: ۱۰.۱۰۹۳/cybersec/tyac.۰۰۷
- [۱۱] Alqahtani, E. J.; Alshamrani, A. M.; Patel, A. "Hybrid Machine Learning Models for Android Malware Classification"; IEEE Trans. Dependable and Secure Computing ۲۰۲۱, ۱۸, ۱۲۳۴-۱۲۴۷. doi: ۱۰.۱۱۰۹/TDSC.۲۰۱۹.۲۹۵۸۷۸۴
- [۱۲] Islam, R.; Tian, R.; Batool, S. "A Multi-stage Hybrid Feature Selection and Classification Framework for Android Malware Detection"; J. Inform. Secur. Appl. ۲۰۲۰, ۵۲, ۱-۱۲. doi: ۱۰.۱۰۱۶/j.jisa.۲۰۲۰.۱۰۲۴۷۴
- [۱۳] Zhang, Y.; Li, X.; Wang, Y. "Program Genes: A Novel Feature Representation for Android Malware Detection"; Proc. the ACM SIGSAC Conference on Computer and Communications Security ۲۰۱۹, ۱, ۱۲۵-۱۳۶. doi: ۱۰.۱۱۴۵/۳۳۱۹۵۳۵.۳۳۵۴۲۴۹
- [۱۴] Embunge, T.; Smith, J.; Johnson, P. "A Critical Review of Deep Learning Models for Malware Detection"; Computers & Security ۲۰۲۲, ۱۱۳, ۱-۱۸. doi: ۱۰.۱۰۱۶/j.cose.۲۰۲۱.۱۰۲۴۵۷
- [۱۵] Samarwar, G.; Kumar, S.; Devi, M. "Behavioral and Permission-based Android Malware Detection Using Classical ML Models"; Int. J. Inform. Secur. ۲۰۲۱, ۲۰, ۱۰۱-۱۱۵. doi: ۱۰.۱۰۰۷/۵۱۰۲۰۷-۲۰۲۰.۰۵۰۵۰۵۰۸
- [۱۶] Mohanraj, S.; Karthik, R.; Singh, V. "Network Traffic Analysis for Android Malware Detection"; IEEE Trans. Network and Service Management ۲۰۲۰, ۱۷, ۲۱۰-۲۲۴. doi: ۱۰.۱۱۰۹/TNSM.۲۰۲۰.۲۹۷۸۹۹۸
- [۱۷] Vasanta, B.; Reddy, P.; Kumar, A. "CNN-LSTM with Feature Weighting for Android Malware Detection"; Neural Comput. Appl. ۲۰۲۲, ۳۴, ۵۶۷۸-۵۶۹۰. doi: ۱۰.۱۰۰۷/۵۰۵۲۱۰۰۲۲۰۰۷۰۰۹۰۹۰۷
- [۱۸] Zaidi, S. F. A.; Awan, M. J.; Khan, R. "Comparative Analysis of Deep Learning Models for Malware Classification"; J. Comput. Virology and Hacking Techniques ۲۰۲۱, ۱۷, ۴۵-۵۸. doi: ۱۰.۱۰۰۷/۵۱۱۴۱۶-۲۰۲۰.۰۰۳۶۲۳-X
- [۱۹] Beyazit, M.; Yildiz, O.; Demir, C. "Static and Dynamic Feature Fusion for Android Malware Detection"; Computers & Security ۲۰۲۰, ۹۵, ۱-۱۴. doi: ۱۰.۱۰۱۶/j.cose.۲۰۲۰.۱۰۱۸۷۳
- [۲۰] Sui, Y.; Li, Q.; Wang, H. "Explainable Malware Detection Using Local Interpretable Model-agnostic Explanations"; IEEE Access ۲۰۲۱, ۹, ۱۲۳۴۵-۱۲۳۴۶. doi: ۱۰.۱۱۰۹/ACCESS.۲۰۲۱.۳۰۵۹۱۸۵
- [۲۱] Sarla, P.; Thompson, L.; Evans, D. "XAI with Feature Selection for Transparent Cybersecurity Systems"; Journal of Big Data ۲۰۲۲, ۹, ۱-۲۰. doi: ۱۰.۱۱۸۶/۵۴۰۵۳۷-۲۰۲۰.۰۵۸۷۰۲
- [۲۲] Kılıç, Ö.; Yanık, T.; Aydın, M. "Attention-based Deep Learning Model for Android Malware Detection Using App Permissions"; Expert Systems with Applications ۲۰۲۱, ۱۸۴, ۱-۱۲. doi: ۱۰.۱۰۱۶/j.eswa.۲۰۲۱.۱۱۵۵۱۰

تبیین‌پذیری (Explainability): تحلیلگر می‌تواند دقیقاً بداند کدام ویژگی‌ها باعث تصمیم‌گیری مدل شده‌اند.

مقاومت در برابر نویز و حملات ادراکی: سیستم حتی باوجود تغییرات مصنوعی در داده‌ها، همچنان عملکرد باثباتی دارد.

بنابراین، می‌توان نتیجه گرفت که چارچوب RL-EDKN با دستیابی به دقت ۹۹/۲٪، پایداری بالا، و شفافیت تصمیم‌گیری، می‌تواند به‌عنوان الگویی کارآمد برای نسل جدید سامانه‌های تشخیص بدافزار اندرویدی مورد استفاده قرار گیرد.

۴. نتیجه‌گیری

در این مقاله، یک چارچوب ترکیبی نوآورانه با عنوان RL-EDKN برای تشخیص بدافزارهای اندرویدی معرفی شد که دو مؤلفه اصلی یادگیری تقویتی برای بهینه‌سازی پویای ویژگی‌ها (RL-DFO) و شبکه هسته‌ای عمیق تبیین‌پذیر (EDKN) را در برمی‌گیرد. این چارچوب باهدف غلبه بر محدودیت‌های روش‌های موجود ازجمله ایستایی در انتخاب ویژگی، ناتوانی در مدل‌سازی وابستگی‌های غیرخطی و فقدان شفافیت در تصمیم‌گیری طراحی شده است. نتایج آزمایش‌ها روی مجموعه داده CIC-AndMal۲۰۲۰ نشان داد که مدل پیشنهادی نه تنها به دقت و امتیاز F۱ بسیار بالایی (بیش از ۹۹٪) دست یافته، بلکه در مقایسه با روش‌های مرسوم مانند IPSO، FKSVM و CNN-LSTM، بهبود قابل توجهی در مقاومت در برابر حملات ادراکی و نویز داده‌ها از خود نشان می‌دهد. علاوه بر این، یکپارچه‌سازی ماژول تبیین‌پذیر مبتنی بر SHAP و LIME امکان تحلیل تصمیم‌های مدل را برای تحلیلگران امنیتی فراهم کرده و سطح اعتماد به سیستم را در محیط‌های عملیاتی افزایش داده است. کاهش زمان آموزش و محاسبات بدون قربانی کردن دقت، از دیگر مزایای قابل توجه این چارچوب است. در آینده، می‌توان قابلیت‌های این مدل را با ادغام داده‌های رفتاری بلادرنگ و گسترش آن به پلتفرم‌های دیگر مانند iOS ارتقا داد.

۵. مراجع‌ها

- [۱] Kamran, S.; Luo, S.; Varadarajan, V. "A Novel Deep Learning-Based Approach for Malware Detection"; Engineering Applications of Artificial Intelligence ۲۰۲۳, ۱۲۲, ۱۰۶۰۳۰. doi: ۱۰.۱۰۱۶/j.engappai.۲۰۲۳.۱۰۶۰۳۰
- [۲] Gaber, M.; Mohiuddin, A.; Helge, J. "Malware Detection with Artificial Intelligence: A Systematic Literature Review"; ACM Computing Surveys ۲۰۲۴, ۵۶, ۱-۳۳. doi: ۱۰.۱۱۴۵/۳۶۳۶۴۲۴
- [۳] Maddireddy, B. R.; Maddireddy, B. R. "Automating Malware Detection: A Study on the Efficacy of AI-Driven Solutions"; J. Environ. Sci. Technol. ۲۰۲۳, ۲, ۱۱۱-۱۲۴.
- [۴] Shabtai, A.; Fledel, Y.; Elovici, Y. "Detection of Malicious Code by Applying Machine Learning Classifiers"; IEEE Trans. Inform. Forensics and Security ۲۰۱۸, ۱۳, ۱۴۵-۱۵۶. doi: ۱۰.۱۱۰۹/TIFS.۲۰۱۷.۲۷۴۵۶۹۰

- [۲۳] Toluifar, A.; Karimi, A.; Karimi, F. "Code Smells detection using ensemble feature selection methods and wrapper techniques based on neural network"; Adv. Defense Sci. Technol. ۲۰۲۴, ۳, ۱۸۹-۲۰۷. (In Persian). doi:۲۰.۱۰۰۱.۱.۲۶۷۶۲۹۳۰.۱۴۰۳.۱۵.۳.۲.۳
- [۲۴] Zahiri, M.; Shirini, K.; Samadi, S. "Network Traffic Analysis with Machine Learning for Faster Detection of Distributed Denial of Service Attack"; Adv. Defense Sci. Technol. ۲۰۲۴, ۴, ۲۷۳-۲۸۲. (In Persian) doi:۲۰.۱۰۰۱.۱.۲۶۷۶۲۹۳۰.۱۴۰۳.۱۴.۴.۶.۲.
- [۲۵] Tajari Siahmarzkooh, A. "Internet of Things Secure System Using Dragonfly Algorithm for Feature Selection and Gradient Boosting for Classification."; Adv. Defense Sci. Technol. ۲۰۲۴, ۲, ۸۵-۹۳. (In Persian). doi:۲۰.۱۰۰۱.۱.۲۶۷۶۲۹۳۰.۱۴۰۳.۱۵.۲.۳.۲
- [۲۶] Ribeiro, M. T.; Singh, S.; Guestrin, C. "Why Should I Trust You? Explaining the Predictions of Any Classifier"; Proc. ACM SIGKDD Int. Conf. Knowledge Discovery and Data Mining ۲۰۱۶, ۱, ۱۱۳۵-۱۱۴۴. doi: ۱۰.۱۱۴۵/۲۹۳۹۶۷۲.۲۹۳۹۷۷۸
- [۲۷] Mnih, V.; Kavukcuoglu, K.; Silver, D. "Human-level Control through Deep Reinforcement Learning"; Nature ۲۰۱۵, ۵۱۸, ۵۲۹-۵۳۳. doi: ۱۰.۱۰۳۸/nature۱۴۲۳۶
- [۲۸] Goodfellow, I.; Bengio, Y.; Courville, A. "Deep learning"; The MIT Press, ۲۰۱۶, ISBN: ۰۲۶۲۰۳۵۶۱۸, ۵۵-۱۸۰.
- [۲۹] Schölkopf, B.; Smola, A. J. "Learning with Kernels: Support Vector Machines, Regularization, Optimization, and Beyond"; Adaptive Computation and Machine Learning Series ۲۰۲۲, ۱۱۴-۱۳۲.
- [۳۰] Lundberg, S. M.; Erion, G.; Lee, S. I. "From Local Explanations to Global Understanding with Explainable AI for Trees"; Nature Machine Intelligence ۲۰۲۰, ۲, ۵۶-۶۷. doi: ۱۰.۱۰۳۸/s۴۲۲۵۶.۰۱۹-۰۱۳۸-۹
- [۳۱] Manning, C.; Raghavan, P.; Schutz, H. "An Introduction to Information Retrieval"; Cambridge University Press, ۲۰۰۹, ۱۸۸-۲۱۱.