

The Application of the CatBoost Algorithm for Malware Category Detection

Mohsen Ghasemi*

Assistant Professor, Department of Computer Engineering, Larestan Branch, Islamic Azad university, Larestan, Iran. *Corresponding Author, Email: Mohsen.Ghasemi@iau.ac.ir

Abstract

Introduction: Android is the most widely used mobile operating system. Its popularity, open-source nature, and extensive app ecosystem have made it a target for malwares. While many studies focus on detecting Android malware, fewer have explored malware classification by attack types. This classification is essential for understanding malicious patterns. This paper proposes using the CatBoost algorithm to detect categories of Android malware, employing a hybrid analysis that combines static and dynamic features. The Kronodroid dataset, which contains 14 distinct malware category labels, serves as the benchmark for evaluation.

Method: A novel feature selection method called VFCE is introduced. Four feature selection techniques - variance, F-test, Chi-Square, and Extra Tree - were applied sequentially to the dataset. By selecting 50, 100, and 200 features using these methods along with VFCE, a total of 15 distinct feature sets were generated. The dataset was divided into training (70%) and testing (30%) subsets. I evaluated the performance of eight classification algorithms: CatBoost, Random Forest, Decision Tree, Support Vector Machine, Logistic Regression, Multi-Layer Perceptron, Bagging, and K-Nearest Neighbors across these feature sets.

Results: The VFCE feature selection method produced better results when selecting 100 features, especially in combination with the CatBoost algorithm. I compared the performance of various algorithms, including CatBoost, using the proposed feature selection method. They were assessed based on several metrics: accuracy, precision, recall, F-measure, false positive rate (FPR), root mean squared error (RMSE), training time, and testing time. CatBoost outperformed the other algorithms and previous studies on this dataset, achieving an accuracy of 93.28%, precision of 93.32%, recall of 93.28%, F-measure of 93.19%, and a rapid testing time of 0.07 seconds.

Discussion: This study examines the use of the CatBoost algorithm for Android malware category detection. The proposed VFCE feature selection method enhances both accuracy and speed. The CatBoost algorithm, in conjunction with the proposed feature selection method, improves accuracy, precision, recall, F-measure, RMSE, and testing time.

Keywords: Android Operating System, Malware Category, Machine Learning, CatBoost.

کاربرد الگوریتم CatBoost در شناسایی دسته بدافزار اندروید

دوره پنجم، پاییز ۱۴۰۳
شماره سوم، صص: ۱۷۱-۱۸۳

تاریخ دریافت: ۱۴۰۳/۰۶/۰۹
تاریخ پذیرش: ۱۴۰۳/۰۸/۲۴

محسن قاسمی*

استادیار، گروه کامپیوتر، واحد لارستان، دانشگاه آزاد اسلامی، لارستان، ایران. Mohsen.Ghasemi@iaau.ac.ir

چکیده: اندروید پرستفاده‌ترین سیستم‌عامل تلفن همراه است و بسیار مورد هدف بدافزارها قرار می‌گیرد. از الگوریتم‌های یادگیری ماشین برای تشخیص بدافزار اندروید، بسیار استفاده شده است. علاوه بر تشخیص مخرب بودن یک برنامه، شناسایی نوع حمله نیز مهم است. تعیین دسته بدافزار اندروید، گام مهمی برای درک دقیق تهدیدی است که در برابر آن آسیب‌پذیر هستیم. در این مقاله، از الگوریتم CatBoost و از رویکرد ترکیبی با ویژگی‌های ایستا و پویا برای تشخیص دسته بدافزار اندروید استفاده می‌شود. نتایج این الگوریتم و چند الگوریتم دیگر یادگیری ماشین، روی مجموعه داده Kronodroid با یکدیگر مقایسه می‌شوند. برای انتخاب ویژگی‌ها، چهار روش انتخاب ویژگی، به‌طور متوالی روی مجموعه داده اعمال می‌شود. برای ارزیابی مجموعه ویژگی پیشنهادی (VFCE)، براساس تعداد ویژگی و هر کدام از روش‌های انتخاب ویژگی، ۱۵ مجموعه ویژگی را ایجاد کرده و دقت الگوریتم‌ها، روی آن‌ها مقایسه می‌شوند. روش پیشنهادی VFCE با ۱۰۰ ویژگی، دقت بهتری دارد. سپس نتایج الگوریتم‌ها با معیارهای دقت، صحت، فراخوانی، معیار F، نرخ مثبت کاذب، ریشه میانگین مربعات خطا، زمان آموزش و زمان تست، مقایسه می‌شوند. همچنین نتایج با سایر کارهای مرتبط روی این مجموعه داده مقایسه می‌شوند. الگوریتم CatBoost با دقت ۹۳٫۲۸٪، صحت ۹۳٫۳۲٪، فراخوانی ۹۳٫۲۸٪، معیار F، ۹۳٫۱۹٪ و زمان تست ۰٫۰۷ ثانیه عملکرد بهتری از سایر الگوریتم‌ها دارد.

واژه‌های کلیدی: سیستم‌عامل اندروید، دسته بدافزار، یادگیری ماشین، CatBoost.

۱. مقدمه

بدافزار موبایل، نرم افزار مخربی است که به طور خاص برای هدف قرار دادن دستگاه های تلفن همراه مانند گوشی های هوشمند و تبلت ها طراحی شده است. بدافزار به هر نوع کد مخربی اطلاق می شود که بدون اطلاع یا رضایت کاربر، یکپارچگی و عملکرد سیستم تلفن همراه را تحت تأثیر قرار می دهد. انواع بدافزارها شامل باج افزار، تروجان، کرم، جاسوس افزار، روت کیت و بات نت است [۱].

اندروید یک سیستم عامل منبع باز برای دستگاه های تلفن همراه و یک پروژه متن باز به رهبری گوگل است. سیستم عامل اندروید براساس یک نسخه تغییر یافته لینوکس ساخته شده است و عمدتاً در دستگاه های موبایل استفاده می شود. اندروید، ۳،۳ بلیون کاربر در سال ۲۰۲۳ و ۷۱،۷۴٪ سهم بازار سیستم عامل های موبایل را دارد. ۹۶،۹٪ برنامه های آن مجانی است. در مارس ۲۰۲۳، ۲۶۷۳۲۹۲ برنامه در فروشگاه گوگل در دسترس بوده است [۲]. با سیاست گوگل مبنی بر هسته منبع باز، نویسندگان بدافزار به طور بالقوه می توانند دانش خوبی از پلت فرم تلفن همراه به دست آورند و مجرمان سایبری فرصت می یابند تا بدافزارهای تلفن همراه را ایجاد و منتشر کنند. علاوه بر این، با افزایش تعداد کاربرانی که برنامه ها را دانلود و نصب می کنند، احتمال نصب بدافزار موبایل نیز افزایش می یابد [۳].

دستگاه های اندرویدی به طور مداوم توسط بدافزارها مورد حمله قرار می گیرند و آلوده می شوند. در فصل اول سال ۲۰۲۴، بیش از ۱۰،۱ میلیون نمونه حمله موبایل به وسیله کاسپرا سکای گزارش شده است. تعداد نمونه های شناسایی شده بدافزار و نرم افزارهای ناخواسته اندروید در همین فصل، به ۳۸۹۱۷۸ بسته نصب شده رسید [۴]. بیش از ۹۸ درصد از حملات سایبری به دستگاه های تلفن همراه، مربوط به اندروید است [۵]. بدافزار اندرویدی می تواند باعث عملیاتی مانند دسترسی های غیرمجاز، سرقت اطلاعات، جاسوسی، ارسال پیامک غیرمجاز و دانلود اضافی شود [۶].

در دهه اخیر، حجم بسیاری از تحقیقات برای تشخیص بدافزارهای اندروید، از یادگیری ماشین استفاده می کنند. اکثریت قریب به اتفاق مطالعات تشخیص بدافزار اندروید مبتنی بر یادگیری ماشین، کارایی بالایی را گزارش می کنند [۷]. برای تشخیص بدافزار با استفاده از الگوریتم های یادگیری ماشین، برنامه ها باید از قبل پردازش شوند تا ویژگی هایی که رفتار آنها را به بهترین شکل توصیف می کنند، استخراج شوند. این کار با استفاده از تکنیک های تحلیل نرم افزاری پویا یا ایستا انجام می شود. در تجزیه و تحلیل پویا، رفتار یک برنامه شبیه سازی می شود. تجزیه و تحلیل ایستا بر اساس بررسی فایل های موجود در بسته برنامه است و نیاز به اجرای کد ندارد. این تکنیک ها مزایا و معایب خاص خود را دارند. از یک طرف، از طریق تحلیل پویا، دسترسی به کدهایی که در زمان اجرا بارگذاری و اجرا می شوند، امکان پذیر است. با این حال، موفقیت این تجزیه و تحلیل، از نظر پوشش کد، تا حد زیادی به مکانیسم شبیه سازی و عدم وجود تکنیک های فرار

از مجازی سازی، بستگی دارد [۸]. از سوی دیگر، تجزیه و تحلیل ایستا، قادر است تمام اطلاعات موجود در برنامه ها را ارزیابی کند، اما موفقیت آن در فقدان تکنیک های فرار مانند مبهم سازی یا بارگذاری پویای کد است. تحلیل دینامیکی، زمان و توان محاسباتی و منابع بیشتر نسبت به تحلیل ایستا مصرف می کند، به همین دلیل کمتر مورد توجه محققین قرار می گیرد [۷]. رویکرد ترکیبی، با هدف به دست آوردن مزایای هر دو رویکرد ایستا و پویا، آنها را با هم ترکیب می کند [۶]. در این تحقیق از رویکرد ترکیبی استفاده می شود.

تشخیص مخرب بودن یک برنامه به تنهایی کافی نیست. نوع حمله می تواند برای جلوگیری از آسیب مورد نظر بسیار مهم باشد تا برای تهدید، اقدامات مناسبی انجام داد [۹]. دسته بندی بدافزارها، بینش های عمیق تری در مورد رفتار و تکامل انواع مختلف بدافزار ارائه می دهد. انواع مختلف بدافزار، رفتارها و اهداف متفاوتی دارند. با دسته بندی بدافزارها، سیستم های امنیتی می توانند دفاع های خود را به طور موثرتری تنظیم کنند و حفاظت بهتری در برابر تهدیدات خاص ارائه دهند [۱۰]. بنابراین، ضروری است که راه حلی داشته باشیم که یک گام فراتر رفته و بدافزار را به خانواده یا دسته خاصی نسبت دهد [۱۱].

مطالعات زیادی برای بررسی اثربخشی مدل های مختلف یادگیری ماشین برای شناسایی بدافزار اندرویدی، انجام شده است. اکثر این مطالعات از مجموعه داده هایی استفاده می کنند که با استفاده از تجزیه و تحلیل ایستا یا پویا جمع آوری شده اند؛ اما استفاده از رویکردهای تجزیه و تحلیل ترکیبی به طور کامل مورد توجه قرار نگرفته است. در [۱۲]، با استفاده از رویکردهای تجزیه و تحلیل ترکیبی، یک مدل یادگیری ماشین نظارت شده برای تشخیص بدافزار اندروید و تشخیص دسته بدافزار اندروید، ارائه شد. از مجموعه داده Kronodroid که شامل ویژگی های ایستا و پویا است، استفاده کردند. این مجموعه داده شامل ۲۴۰ خانواده بدافزار است. با کمک چندین مخزن آنلاین، بدافزارها در ۱۵ گروه، دسته بندی شدند. با الگوریتم های Extra Tree و اطلاعات متقابل، تعداد ۳۰،۵۰ و ۱۰۰ ویژگی برتر را روی مجموعه داده عادی و نرمال شده انتخاب و ۱۲ نمونه داده ایجاد کردند. دقت الگوریتم های درخت تصمیم (DT)، جنگل تصادفی (RF)، K نزدیکترین همسایگی (KNN) و ماشین بردار پشتیبان (SVM) را روی این نمونه داده ها، با هم مقایسه نمودند. نمونه داده با ۵۰ ویژگی انتخاب شده با Extra Tree و الگوریتم RF، بالاترین دقت را در هر دو مورد شناسایی بدافزار و شناسایی دسته بدافزار به دست آورد. با تنظیم پارامترهای ۴ الگوریتم کلاس بندی مذکور، RF، دقت ۹۸،۰۳٪ برای تشخیص بدافزار و ۸۷،۵۶٪ برای تشخیص دسته بدافزار را به دست آورد.

در این تحقیق، از مجموعه داده و دسته بندی بدافزار ارائه شده در [۱۲] استفاده می شود. برای انتخاب ویژگی، چهار روش انتخاب ویژگی به صورت متوالی استفاده می شود. الگوریتم های CatBoost، پرسپترون چندلایه (MLP)، لجستیک رگرسیون (LR) و Bagging نیز به مدل یادگیری اضافه می شوند.

می‌تواند به‌طور خودکار الگوهای مخرب را با استفاده از دنباله فراخوانی-های متد خام در کد اسمبلی، تشخیص دهد. MalDozer می‌تواند خانواده بدافزار اندروید را مشخص کند.

در [۱۴]، یک روش تشخیص بدافزار اندروید بنام ICCDetector پیشنهاد شد. از ارتباطات و تعاملات بین مولفه‌های داخل برنامه (ICC) یا در سراسر مرزهای برنامه و اینتنت‌ها، به عنوان ویژگی‌هایی برای شناسایی بدافزار استفاده شد. از مجموعه داده شامل ۵۲۶۴ بدافزار و ۱۲۰۲۶ برنامه امن، ۱۲۱۶۲۱ ویژگی استخراج شد. با روش انتخاب ویژگی مبتنی بر همبستگی، ویژگی‌های اضافی حذف شدند. برای کلاسبندی از SVM و اعتبارسنجی متقابل ۱۰ فولد استفاده شد. ارزیابی روی این مجموعه داده، دقت ۹۷٫۴٪ و نرخ مثبت کاذب ۰٫۶۷٪ را نتیجه داد. همچنین، با توجه به ویژگی‌های مؤلفه‌های داخل برنامه، بدافزارهای شناسایی شده توسط ICCDetector، به پنج دسته بدافزار جدید کلاسبندی شدند.

در [۱۵]، یک مدل تشخیص بدافزار مبتنی بر یادگیری تصویر عمیق به نام DIDROID، برای شناسایی دسته بدافزار اندروید پیشنهاد شد. فرضیه این بود که تبدیل داده‌های ویژگی‌های ایستا به تصاویر، می‌تواند به شناسایی چندکلاسه برنامه‌های بدافزار اندروید کمک کند. نویسندگان، از مجموعه داده CCS-CICAndMal2020 با ۴۰۰ هزار برنامه (۲۰۰۰۰۰ بدافزار و ۲۰۰۰۰۰ برنامه امن) که خود در تولید مشارکت داشتند استفاده نمودند. مجموعه داده شامل ۱۹۱ خانواده بدافزار است که در ۱۲ دسته، برچسب‌گذاری شده‌اند. برای انتخاب ویژگی‌ها، از Extra Tree و برای کلاسبندی از شبکه عصبی کانولوشنال، استفاده شد. نتایج آزمایش‌ها روی داده‌های تست، نشان داد که ۹۳٫۳٪ (۳۱۹۰۱ نمونه) بدافزارها، به‌درستی در رده‌های مربوطه کلاسبندی شدند. ۶٫۷٪ (۲۲۹۹ نمونه)، اشتباه کلاسبندی شدند.

در [۱۶]، یک تکنیک تجزیه و تحلیل پویا مبتنی بر آنتروپی به نام EntropLyzer پیشنهاد شد. تجزیه و تحلیل پویا، در یک محیط شبیه-سازی شده روی مجموعه داده CCCS-CIC-AndMal2020 انجام شد که شامل ۱۴۷ خانواده و ۱۲ دسته بدافزار اندروید است. شش دسته از ویژگی‌ها برای کلاسبندی و شناسایی بدافزار اندروید استخراج شد. از آنتروپی شانن برای شناسایی تغییرات رفتاری در خانواده‌های بدافزار و رتبه‌بندی این ویژگی‌ها استفاده شد. از الگوریتم‌های مختلف یادگیری ماشین (بیز ساده، RF، DT)، برای کلاسبندی دسته و خانواده بدافزار استفاده شد. بهترین نتایج را DT با صحت ۹۸٫۴٪، فراخوانی و معیار F، ۰٫۹۸۳ به‌دست‌آورد.

یادگیری ماشین بانظارت، نیاز به تعداد قابل‌توجهی برنامه‌امن و بدافزار دارد که از قبل شناسایی و برچسب‌گذاری شده باشند. داده‌های برچسب‌گذاری شده در این زمینه، گران و به‌دست‌آوردن آن‌ها دشوار است. داده‌های بدون برچسب، فراوان و ارزان هستند. در [۱۷]، براساس یادگیری نیمه‌نظارت‌شده برای شبکه‌های عصبی عمیق، مجموعه‌ای از نمونه‌های برچسب‌گذاری شده و بدون برچسب را آموزش دادند. این

CatBoost یک کتابخانه منبع‌باز با کارایی بالا برای تقویت‌گرادیان در درخت‌های تصمیم است که یاندکس در سال ۲۰۱۷ ارائه کرد. این الگوریتم؛ در مقایسه با سایر الگوریتم‌های مبتنی بر درخت تصمیم افزایش‌گرادیان، می‌تواند ویژگی‌های کلاسبندی را به‌طور موثر و معقول پردازش کند و مشکلات سوگیری‌گرادیان و شیفت (توزیع خطا) پیش‌بینی را حل کند. در نتیجه بیش‌برازش را کاهش داده و دقت کلاسبندی و توانایی تعمیم مدل را بهبودمی‌بخشد [۱۳]. در ادامه، در بخش ۲، سابق و کارهای گذشته آورده شده است. در بخش ۳، متدولوژی، مجموعه داده و روش انتخاب ویژگی تشریح می‌شود. در بخش ۴، نتایج تجربی الگوریتم‌ها نمایش داده می‌شود و الگوریتم‌های پیاده‌سازی شده با یکدیگر و با نتایج تحقیق دیگران مقایسه می‌شوند. در پایان نیز نتیجه‌گیری و کارهای آینده بیان می‌شود.

۲. کارهای گذشته

تحقیقات بسیاری در زمینه تشخیص دسته و خانواده بدافزارهای اندروید با استفاده از یادگیری ماشین انجام شده است. در بخش قبل، تحقیق واحد و قدیر [۱۲] آورده شد. در ادامه چند مورد از آخرین این کارها آورده می‌شود.

در [۱۳]، چارچوب شناسایی سریع بدافزار اندروید (FAMD) معرفی شد. با استخراج دنباله‌های آپ‌کدها و مجوزها از نمونه برنامه‌ها، مجموعه ویژگی‌های اصلی ساخته می‌شوند. سپس آپ‌کدها با تکنیک N-Gram پیش‌پردازش می‌شوند و با الگوریتم فیلتر مبتنی بر همبستگی سریع بر اساس عدم قطعیت متقارن، ابعاد ویژگی‌ها کاهش می‌یابد. در این تحقیق، برای اولین بار از کلاسبند یادگیری ماشین CatBoost برای تشخیص بدافزار اندروید استفاده شد. نویسندگان نشان دادند که روش پیشنهادیشان از RF، KNN و XGBoost بهتر عمل می‌کند و زمان مصرف کمتری دارد. آن‌ها از دو مجموعه داده استفاده کردند. یک مجموعه داده شامل ۲۵۷۳۷ برنامه (۱۲۹۸۹ بدافزار و ۱۲۷۴۸ برنامه امن) را جمع‌آوری کردند و برای آزمایش‌ها از آن استفاده کردند. ۷۰٪ داده‌ها، برای آموزش و ۳۰٪ برای تست استفاده شد. برای مقایسه چارچوب پیشنهادیشان با دیگر کارها، از مجموعه داده Drebin استفاده کردند. علاوه بر شناسایی بدافزار، کلاسبندی خانواده بدافزار را نیز در این چارچوب موردتوجه قرار دادند. ۲۰ خانواده برتر بدافزار در این مجموعه داده را در نظر گرفتند و کلاسبندی خانواده بدافزار را با چارچوب ارائه شده، انجام دادند. الگوریتم CatBoost، بهترین عملکرد را داشت و دقت ۹۷٫۴٪ و ۹۷٫۳۸٪ را به ترتیب، برای شناسایی بدافزار و شناسایی خانواده بدافزار اندروید، به‌دست‌آورد.

در [۱۱]، براساس شبکه عصبی مصنوعی، یک چارچوب برای شناسایی بدافزار و شناسایی خانواده بدافزار بنام MalDozer، پیشنهاد شد. در این چارچوب، دنباله‌های خام فراخوانی‌های API از کد اسمبلی استخراج می‌شود. فراخوانی‌های API برداری شده با استفاده از شبکه عصبی کانولوشنال، آموزش داده می‌شوند. در طول آموزش، MalDozer

جدول ۱: تحقیقات اخیر برای تشخیص دسته بدافزار اندروید با استفاده از مجموعه داده های مختلف [۱۲]

مرجع	الگوریتم و سال	مجموعه داده و سال انتشار آن	بازه زمانی	رویکرد (ایستا، پویا، ترکیبی)	روش انتخاب ویژگی	شناسایی دسته بدافزار
[۲۰]	Deep ANN (2021)	CICANDMAL2019	-	ترکیبی	-	دقت رویکرد ایستا: ۹۲,۵٪ دقت رویکرد پویا: ۸۰,۳٪
[۲۱]	Machine learning (2018)	Updroid	۲۰۱۴-۲۰۱۸	ترکیبی	-	دقت: ۹۶,۳۷٪
[۱۷]	Deep neural network with pseudolabel (2020)	CICMALDroid2020	۲۰۱۸-۲۰۱۷	پویا	-	Fmeasure: ۴ دسته بدافزار با 97.8%
[۱۶]	Machine learning (2021)	CCS-CICAndMal2020	-	پویا	Shannon Entropy	۱۲ دسته با صحت: ۹۸,۴٪
[۲۲]	Machine learning (2017)	MacAfee, Drebin, Praguard	-	ایستا	Mean decrease impurity (MDI)	دقت: ۹۹,۲۶٪
[۲۳]	Machine learning (2018)	Drebin, Genome VirusShare	۲۰۰۹-۲۰۱۷	پویا	-	Fmeasure: ۹ دسته بدافزار با 97.8%
[۱۸]	Deep neural network with pseudo label stack auto encoder (2022)	CICMALDroid2020	۲۰۱۷-۲۰۱۸	ترکیبی	Variance, Mutual information	۵ دسته با صحت: ۹۸,۴٪
[۱۲]	Machine learning (random forest)	Subset of 2020 KronoDroid	۲۰۰۸-۲۰۲۰	ترکیبی	Extra tree, Mutual information	۱۵ دسته با دقت ۸۷,۶٪

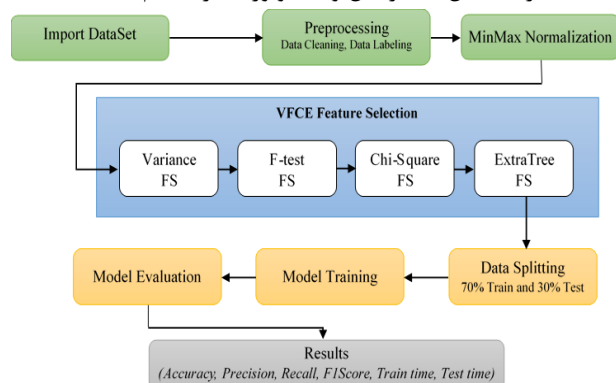
ویژگی ها با رویکرد Chi-Squared انجام شد و ۲۰ ویژگی از ۴۸۹ ویژگی انتخاب شدند. نتایج نشان داد که با زیرمجموعه کوچکی از ویژگی ها، می توان بخوبی برنامه های غیرمبهم را کلاس بندی کرد، ولی پس از اعمال مبهم سازی، کلاس بندی تغییر شدیدی دارد که نشان دهنده اهمیت نسبی تکنیک های مبهم سازی است.

در جدول ۱، تحقیقات اخیر برای تشخیص دسته بدافزار اندروید با استفاده از مجموعه داده های مختلف، آورده شده است.

۳. متدولوژی

۳.۱. توصیف محیط

تمام آزمایش ها با کامپیوتر در برنامه پایتون ۱۲,۳,۱ و در محیط ویندوز ۱۰,۶۴ بیتی انجام شده است. از نظر سخت افزاری از Intel Core i7-7700HQ 2.8 GHz Processor with 16 GB RAM استفاده شده است. در شکل ۱، مراحل و متدولوژی، ترسیم شده است.



شکل ۱: متدولوژی برای تشخیص دسته بدافزار اندروید

چارچوب به عنوان ورودی، رفتارهای مشاهده شده پویا را می گیرد تا با استفاده از یک تکنیک نیمه نظارت شده، کلاس بندی دسته بدافزار را فعال کند. نویسندگان، مجموعه داده CICMALDroid2020، را از نمونه برنامه های چند منبع، در فاصله دسامبر ۲۰۱۷ تا دسامبر ۲۰۱۸، ایجاد کردند. داده ها شامل برنامه های امن و چهار دسته بدافزار هستند. این مجموعه داده شامل ویژگی های ایستا و پویاست. آن ها مدل پیشنهادیشان (بنام PLDNN) را با روش یادگیری ماشین نیمه نظارتی انتشار برچسب و سایر الگوریتم های رایج یادگیری ماشین مقایسه کردند. نتایج نشان داد که این مدل می تواند، برنامه های اندرویدی را با توجه به دسته بدافزار، با معیار F، ۹۷,۸۴٪ و نرخ مثبت کاذب، ۲,۷۶٪ کلاس بندی کند که بهتر از روش انتشار برچسب است. در ادامه، در [۱۸]، با ترکیب تحلیل ایستا و پویا، روش نیمه نظارت شده بنام PLACE برای کلاس بندی دسته بدافزار پیشنهاد شد. نتایج نشان داد که مدل پیشنهادی، دقت ۹۸,۲۸٪ و نرخ مثبت کاذب ۱,۱۶٪ را به دست آورد. همچنین نتایج نشان دادند که با ۱٪ داده های برچسب گذاری شده، مدلی پیشنهادی دقت ۹۵,۱۹٪ را دارد که عملکرد خوبی است.

بدافزارهای هوشمند، برای پنهان کردن عملکرد خود و فرار از موتورهای ضد بدافزار، از تکنیک های مبهم سازی استفاده می کنند. یک رویکرد مبهم سازی می تواند نسخه های بدافزاری تولید کند که می تواند از استراتژی های تشخیص بدافزار، فرار کند و دقت تشخیص را به طور چشمگیری کاهش دهد. برنامه های امن نیز برای محافظت از حریم خصوصی کد یا مالکیت معنوی، از مبهم سازی استفاده می کنند. در [۱۹]، رویکردی برای تشخیص تکنیک های مبهم سازی بدافزارها با یادگیری عمیق و با استفاده از مکانیزم رای گیری جمعی پیشنهاد شد. از مجموعه داده Kronodroid که برای ارزیابی استفاده شد، انتخاب

۲.۳. مجموعه داده

در این مقاله، از مجموعه داده Kronodroid استفاده شده است [۲۴]. این مجموعه داده شامل ۷۸۱۳۷ نمونه (۴۱۳۸۲ بدافزار و ۳۶۷۵۵ برنامه امن) از سال ۲۰۰۸ تا ۲۰۲۰ است. نمونه‌ها دارای ۲۰۰ ویژگی ایستا و ۲۸۹ ویژگی پویا هستند. ویژگی‌های پویا با اجرای بدافزار بر روی دستگاه واقعی جمع‌آوری شده‌اند. در این مجموعه داده علاوه بر برچسب بدافزار یا امن بودن، خانواده بدافزارها نیز مشخص شده است.

در اولین مرحله، ۹۳۸۹۴ برنامه‌های اندروید را از ۷ مجموعه داده و مخزن مختلف جمع‌آوری کردند. ۵۴۸۳۴ بدافزارها را از AMD، Drebin، VirusTotal و VirusShare و ۳۹۰۶۰ برنامه‌های امن را از MARVIN، F-droid و APKMirror دانلود کردند. سپس تحلیل پویا را با دو پلتفرم مختلف یعنی شبیه‌ساز و دستگاه واقعی انجام دادند. فراخوانی‌های سیستمی برای ویژگی‌های پویا استفاده می‌شوند. سیستم‌عامل اندروید، درخواست‌های سرویس‌ها و منابعی را مانند کنترل فرآیند، مدیریت فایل، مدیریت دستگاه و ارتباطات، مدیریت می‌کند. فراخوان‌های سیستمی در این مجموعه داده، با استفاده از یک اسکریپت خودکار با استفاده از دستورات خط فرمان ADB انجام شد. این اسکریپت؛ نصب، اجرا، نظارت، گزارش‌گیری و حذف برنامه‌های جمع‌آوری شده را انجام می‌دهد. برای فراخوانی‌های سیستمی، از ابزار monkey استفاده کردند. پس از نصب، هر برنامه ۶۰ ثانیه اجرا می‌شود. در زمان اجرا، برنامه بررسی می‌شود و فراخوانی‌های سیستمی صادر شده بوسیله ابزار strace ثبت می‌شود. فایل log ایجاد شده، با استفاده از ADB، ذخیره می‌شود. پس از تحلیل پویا، دو مرحله متوالی تحلیل ایستا انجام می‌شود. در مرحله اول، فایل برنامه (apk) و مانیفست آن بررسی و اطلاعات لازم استخراج می‌شوند. فایل مانیفست شامل داده‌هایی از مجوزهای امنیتی (۱۶۶ درخواست مجوز)، سرویس‌ها، فعالیت‌ها، ویژگی‌های سخت‌افزاری و نرم‌افزاری مورد نیاز برای دسترسی به منابع است. ویژگی‌های ایستا مورد نظر از فایل مانیفست با استفاده از ابزارهای AndroGuard، Android Asset Packaging Tool (aapt) استخراج می‌شوند. با استفاده از ابزار ExifTool، متاداده مانند نام فایل، زمان آخرین تغییر و زمان اولین تغییر واکشی می‌شود. مرحله دوم تحلیل ایستا با استفاده از آنتی‌ویروس VirusTotal انجام می‌گیرد. یک سرویس پیمایشگر آنتی‌ویروس است که برنامه را می‌گیرد و براساس نتایج ۵۰ آنتی‌ویروس، گزارش تشخیص بدافزار، خانواده بدافزار، نرخ تشخیص بدافزار و ... را می‌دهد. پس از تحلیل پویا و ایستا، داده‌های به‌دست‌آمده برای هر برنامه پردازش شده و ویژگی‌ها استخراج می‌شوند. در کل ۴۸۹ ویژگی (۲۰۰ ویژگی ایستا و ۲۸۹ ویژگی پویا) برای هر برنامه به‌دست می‌آید [۲۴].

نسخه اصلی مجموعه داده Kronodroid، شامل برچسب دسته بدافزار نیست. در [۱۲]، با کمک چندین مخزن آنلاین، ۲۴۰ خانواده بدافزار این مجموعه داده را در ۱۵ دسته برچسب گذاری کردند. در مرحله برچسب گذاری داده، هر بدافزار را با نام دسته بدافزار

برچسب‌گذاری کردند. برچسب دسته بدافزار از مخازن آنلاین VirusTotal، F-Secure و FortiGuard به‌دست آمده است. نسخه اصلاح شده و بهبود یافته نهایی مجموعه داده در Github در دسترس قرار گرفته است [۲۵]. ویژگی‌های غیرعددی sha256 و HighestModDate و EarliestModDate که تأثیری در خروجی ندارند، حذف شده‌اند. ۴۷۹ ویژگی باقی مانده است. مقادیر عددی با تکنیک MinMax نرمال می‌شوند.

در جدول ۲، تعداد بدافزار در هر دسته نمایش داده شده است. بدافزارهایی که دسته آن‌ها شناسایی نشده است در Blank-Category قرار گرفته‌اند. به هر دسته، عددی صحیح اختصاص یافته است که جانشین نام آن‌ها در مجموعه داده شده است. دسته برنامه‌های امن (Benign) با عدد ۱۶، مشخص شده‌اند.

جدول ۲: تعداد نمونه هر دسته بدافزار

ردیف	دسته بدافزار	تعداد	مقدار عددی متناظر
۱	Adware	۱۱۱۸۳	۱
۲	Trojan	۷۵۴۱	۲
۳	Trojan-SMS	۶۳۴۷	۳
۴	Riskware	۵۳۵۸	۴
۵	Trojan-spy	۴۳۳۶	۵
۶	Ransomware	۲۰۲۹	۶
۷	Trojan-banker	۱۶۸۱	۷
۸	Trojan-Backdoor	۲۲	۸
۹	Scareware	۴۲	۹
۱۰	PUA	۸۸۱	۱۰
۱۱	File-infecter	۴۴۲	۱۱
۱۲	Spyware	۲۶۰	۱۲
۱۳	Trojan-dropper	۴۶	۱۳
۱۴	Trojan/Riskware	۱۰۴۸	۱۴
۱۵	Blank-category	۱۶۶	۱۵
۱۶	Benign	۳۶۷۵۵	۱۶

هر بدافزار با توجه به سه مخزن، دسته‌بندی شده است. یک چالش، عدم توافق آنتی‌ویروس‌ها و برنامه‌های امنیتی در مورد خانواده و دسته بدافزارهاست. برخی بدافزارها، به وسیله برنامه‌های امنیتی در خانواده‌های مختلف، برچسب‌گذاری می‌شوند. رأی‌گیری بین نتایج چند برنامه امنیتی یا اولویت‌دهی به نتایج آنتی‌ویروس‌ها را می‌توان به کار برد. مثلاً اگر یک بدافزار به وسیله VirusTotal و F-Secure در دسته Riskware و به وسیله FortiGuard در دسته Spyware، قرار گیرد، با مکانیزم رأی‌گیری، در دسته Riskware قرار می‌گیرد.

چالش دیگری که در افزایش تعداد دسته‌بندی‌ها باید در نظر گرفت، نامتوازن شدن مجموعه داده است. یعنی تعداد بدافزارها در دسته‌های مختلف، تفاوت زیاد داشته باشند. در دنیای واقعی نیز انواع بدافزار متوازن نیستند و برخی انواع بدافزار بیشتر از انواع دیگر هستند. در صورتی که بخواهیم مجموعه داده را متوازن کنیم، باید تعدادی از بدافزارها را حذف کنیم. در صورت حذف، تعداد نمونه‌های مجموعه داده

کاهش می‌یابد. در این مجموعه داده نیز مشکل نامتوازن بودن تا حدی وجود دارد. به همین دلیل در محاسبه معیارهای ارزیابی از میانگین وزنی استفاده شده است. در جدول ۱، تحقیقات مشابه روی مجموعه داده‌های مختلف آورده شده است. در این تحقیقات نیز تعداد ۴، ۹، ۵ و ۱۲ دسته در نظر گرفته شده است. این مجموعه داده، ۱۵ دسته بدافزار دارد که از بقیه بیشتر است.

برای تنوع بیشتر، می‌توان از نمونه برنامه‌های بیشتر و مخازن دیگر نیز برای دسته‌بندی استفاده کرد. همچنین با توجه به تکامل بدافزارها، خانواده و دسته آن‌ها باید در بازه‌های کم زمانی، به‌روزرسانی شوند.

۳. انتخاب ویژگی

از آنجایی که تعداد زیادی ویژگی در مجموعه داده وجود دارد، یافتن مرتبط‌ترین ویژگی‌ها، نقشی حیاتی در یادگیری ماشین دارد و به آموزش دقیق و مؤثر مدل کمک می‌کند. هدف اصلی فرآیند انتخاب ویژگی، انتخاب مرتبط‌ترین ویژگی‌ها و حذف ویژگی‌های نامربوط و زائد است. انتخاب ویژگی، باعث می‌شود که مدل، سریع‌تر و با حافظه کمتر آموزش ببیند [۲۶]. در اینجا، چهار روش انتخاب ویژگی، به‌صورت متوالی روی مجموعه داده اعمال می‌شود. در شکل ۲، شبه کد الگوریتم انتخاب ویژگی آمده است. هدف این است که زنجیره‌ای از روش‌های انتخاب ویژگی استفاده شود که در کنار کاهش ابعاد، دقت و سرعت خوبی نیز داشته باشد. همچنین نشان دهیم که روش زنجیره‌ای متشکل از چند روش انتخاب ویژگی، نتایج بهتری از هر کدام از روش‌ها به‌تنهایی دارد.

در کنار این چهار روش، PCA، mutual information و Logistic regression نیز برای انتخاب ویژگی بررسی شدند. از بین این روش‌ها، براساس دقت و سرعت، چهار روش واریانس، F-test، Chi-squared و Extra Tree انتخاب شدند. ترتیب اعمال هر کدام از این چهار روش نیز مهم است. انجام به‌ترتیب واریانس، F-test، Chi-squared و Extra Tree، نتایج بهتری داشتند.

روش واریانس روی مجموعه داده پس از نرمال‌سازی، انجام شده است که بیشتر از ۱۸۰ ویژگی را حذف می‌کند. ابتدا واریانس هر ویژگی محاسبه می‌شود و سپس ویژگی‌هایی که واریانس بالاتری دارند، انتخاب می‌شوند. این کار باعث می‌شود که ویژگی‌های کم‌اهمیت و با واریانس پایین حذف شوند و مدل یادگیری ماشین با ویژگی‌های مهم‌تر و اطلاعات بیشتر آموزش داده شود. روش F-test، امکان مقایسه چند گروه را به‌طور هم‌زمان فراهم می‌کند و در سناریوهایی با بیش از دو گروه کارآمد است. در اینجا نیز چند کلاس (گروه) از بدافزارها را داریم. روش F-test، سرعت بالایی دارد. روش Extra Tree در کارهای دیگر نیز نتایج خوبی داشته‌اند. در [۱۲]، Extra Tree و mutual information روی مجموعه داده Kronodroid برای تشخیص دسته بدافزار، با هم مقایسه شدند که Extra Tree نتایج بهتری داشت. در [۲۷]، سه روش انتخاب ویژگی F-test، chi-square، mutual information انجام شد که chi-square نتیجه بهتری داشت. F-test در رتبه بعدی قرار گرفت.

در کارهای دیگر مانند [۱۹، ۲۸] نیز، از chi-square استفاده شده است. نام این روش را VFCE می‌نامیم که از قراردادن حرف اول هر کدام از چهار روش انتخاب ویژگی به‌دست آمده است. ابتدا ویژگی‌هایی را که واریانس صفر یا نزدیک به صفر دارند، حذف می‌کنیم. ۳۰۰ ویژگی باقی می‌ماند.

تابع SelectKBest در کتابخانه Scikit-Learn پایتون، k ویژگی با بالاترین امتیاز را نگه می‌دارد و بقیه را حذف می‌کند. با این تابع، ابتدا ۲۵۰ ویژگی از ۳۰۰ ویژگی مرحله قبل که امتیاز F-test بیشتر دارند، انتخاب می‌کنیم. در واقع، F-value بین هر ویژگی و برچسب (کلاس) را محاسبه می‌شود و ۲۵۰ ویژگی با مقدار بالاتر، انتخاب می‌شود. مجدداً با Chi-Square، ۲۲۰ ویژگی از ۲۵۰ ویژگی مرحله قبل را انتخاب می‌کنیم. Chi-Square روشی آماری برای آزمون استقلال دو متغیر است. در انتخاب ویژگی، Chi-Square بررسی می‌کند که یک ویژگی به چه میزان مستقل از برچسب است.

VFCE_Feature_Selection (X, Y)

```
X_Var = VarianceThreshold(threshold=0).fit_transform(X)
X_Ftest = SelectKBest(f_classif, k=250).fit_transform(X_Var, Y)
X_Chi = SelectKBest(chi2, k=220).fit_transform(X_Ftest, Y)
Model = ExtraTrees(max_features=2).fit(X_Chi, Y)
X_new = TopFeatures(Model.feature_importances_, 100)
Return X_new
```

شکل ۲: شبه کد انتخاب ویژگی VFCE

از کلاسبندی Extra Tree، برای انتخاب نهایی ویژگی‌ها استفاده می‌شود. Extra Tree یک نوع تکنیک یادگیری تجمعی است که چندین درخت تصمیم را به‌صورت موازی ایجاد می‌کند. هر درخت با استفاده از یک زیرمجموعه تصادفی از ویژگی‌ها و نقاط داده، ساخته می‌شود. این کلاسبند، نتایج چندین درخت تصمیم بدون همبستگی در یک جنگل را جمع‌آوری می‌کند تا نتیجه کلاسبندی آن را به‌دست

جدول ۳: مجموعه ویژگی‌های ایجاد شده

الگوریتم انتخاب ویژگی	نام مجموعه ویژگی	تعداد ویژگی‌ها
Variance	Variance 50	۵۰
	Variance 100	۱۰۰
	Variance 200	۲۰۰
F_Test	F_Test 50	۵۰
	F_Test 100	۱۰۰
	F_Test 200	۲۰۰
Chi-Square	Chi-Square 50	۵۰
	Chi-Square 100	۱۰۰
	Chi-Square 200	۲۰۰
Extra Tree	Extra Tree 50	۵۰
	Extra Tree 100	۱۰۰
	Extra Tree 200	۲۰۰
Combination Variance, F-Test, Chi-Square and Extraree Algorithms sequentially	VFCE 50	۵۰
	VFCE 100	۱۰۰
	VFCE 200	۲۰۰

آورد [۱۵]. در طول ساخت جنگل، به‌طور ضمنی ویژگی‌ها را بر اساس سهم آن‌ها در کاهش ناخالصی و پیش‌بینی‌های دقیق، رتبه‌بندی می‌کند. در اینجا تعداد مورد نظر (۵۰، ۱۰۰ و ۲۰۰) ویژگی بهتر را انتخاب می‌کنیم.

برای ارزیابی روش ترکیبی انتخاب ویژگی، مجموعه ویژگی‌های هر کدام از روش‌های انتخاب ویژگی نیز، جداگانه روی مجموعه داده اعمال می‌شود. مجموعه ویژگی‌های متفاوت با ۲۰، ۳۰، ۵۰، ۱۰۰ و ۲۰۰ ویژگی برتر امتحان شدند. در جدول ۳، ۱۵ مجموعه ویژگی متفاوت که نتایج کلاسیک‌تری داشتند، انتخاب شدند.

۳.۴. الگوریتم کلاسیک CatBoost

CatBoost یک کتابخانه منبع‌باز با کارایی بالا برای تقویت گرادیان در درخت‌های تصمیم است که توسط یاندکس در سال ۲۰۱۷ ارائه شد. CatBoost الگوریتم بهبودیافته درخت تصمیم افزایش گرادیان (GBDT) است [۱۳].

در الگوریتم‌های بوستینگ کلاسیک، گرادیان‌های مورد استفاده در هر مرحله، با استفاده از همان مجموعه داده - که مدل فعلی بر روی آن‌ها ساخته شده است - تخمین زده می‌شوند و برای آموزش و به‌دست‌آوردن کلاسیک پایه، استفاده می‌شوند. این عمل، باعث سوگیری گرادیان، توزیع آن و بیش‌برازش می‌شود. به جای آن، CatBoost از مفهوم تقویت مرتب (Ordered Boosting) استفاده می‌کند. فرض کنید مجموعه داده $D = \{(X_i, y_i)\}_{i=1..n}$ با n نمونه، مشاهده شده است؛ به‌طوری که $X_i = (x_{i,1}, x_{i,2}, \dots, x_{i,m})$ برداری از m ویژگی است. ویژگی‌ها می‌توانند، عددی یا غیر عددی باشند. $y_i \in \mathbb{R}$ مقدار برچسب است. از D برای آموزش استفاده می‌نماییم. تعدادی جایگشت تصادفی از داده آموزش ایجاد می‌شود. فرض کنید که $\sigma = (\sigma_1, \dots, \sigma_n)$ یک جایگشت تصادفی روی نمونه‌های آموزش است و n مدل پشتیبان مختلف $s_1, \dots, s_n$ را داریم، طوری که مدل s_i فقط با استفاده از i نمونه اول جایگشت، آموزش می‌بیند. در هر مرحله، برای به‌دست‌آوردن باقیمانده برای نمونه z ام، از مدل s_{j-1} استفاده می‌شود [۲۹].

در CatBoost هر مقدار غیر عددی با میانگین برچسب همان مقدار جایگزین می‌شود. یعنی اگر $\sigma = (\sigma_1, \dots, \sigma_n)$ یک جایگشت تصادفی روی داده آموزش باشد، $x_{\sigma_j,k}$ با رابطه (۱)، جایگزین می‌شود.

$$x_{\sigma_j,k} = \frac{\sum_{j=1}^{i-1} [x_{\sigma_j,k} = x_{\sigma_i,k}] y_{\sigma_j} + a.p}{\sum_{j=1}^{q-1} [x_{\sigma_j,k} = x_{\sigma_i,k}] + a} \quad (1)$$

که منظور از $[.]$ ، آیورسون برکت است. یعنی اگر $x_{\sigma_j,k} = x_{\sigma_i,k}$ باشد، $[x_{\sigma_j,k} = x_{\sigma_i,k}] = 1$ و در غیر این صورت صفر است. p مقدار پیش‌بینی و معمولاً برابر با مقدار میانگین برچسب برای رگرسیون و احتمال پیش‌بینی کلاس مثبت، برای کلاسیک دودویی است. a مقداری بزرگتر از صفر است. با جایگذاری مقادیر غیر عددی با مقادیر رابطه بالا، ممکن است اطلاعات را از دست بدهیم. ترکیب ویژگی‌ها این

مشکل را کم می‌کند و ویژگی قویتری را ایجاد می‌کند. ترکیب همه ویژگی‌های غیر عددی در مجموعه داده، با توجه به رشد نمایی تعداد ترکیبات، مطلوب نیست. CatBoost ترکیب ویژگی‌ها را هنگام ساختن یک انشعاب جدید در درخت فعلی مورد نظر قرار می‌دهد. در اولین قسمت ساخت درخت، ترکیب ویژگی‌ها انجام نمی‌شود ولی برای ایجاد قسمت‌های بعدی درخت، تمام ویژگی‌های غیر عددی و ترکیبات موجود در درخت جاری را با تمام ویژگی‌های غیر عددی مجموعه داده، ترکیب می‌کند. مقادیر ترکیب ویژگی‌ها را طبق رابطه (۱) به مقدار عددی تبدیل می‌کند. همچنین CatBoost، ترکیبات ویژگی‌های عددی و غیر عددی را نیز تولید می‌کند. ویژگی‌های غیر عددی را مانند ویژگی‌های عددی به‌خوبی مدیریت می‌کند و عملکرد بهتری از الگوریتم‌های مبتنی بر تقویت گرادیان مانند XGBoost، LightGBM و H2O دارد [۳۰].

۳.۵. معیارهای ارزیابی

الگوریتم‌های کلاسیک، بر اساس معیارهای عمومی ارزیابی داده کاوی یعنی دقت، صحت، فراخوانی، معیار F، نرخ مثبت کاذب (FPR) و ریشه میانگین مربعات خطا (RMSE)؛ که به‌ترتیب بر اساس روابط (۲) تا (۷) به‌دست می‌آیند و همچنین بر اساس زمان آموزش و زمان تست با یکدیگر مقایسه می‌شوند.

در این کلاسیک چند کلاسی، معیارها را برای هر کلاس، به‌دست آورده و میانگین وزنی هر معیار محاسبه می‌شود. برای هر کلاس c_i :

TP: تعداد نمونه‌هایی است که به‌درستی، بعنوان کلاس c_i کلاسیک شده‌اند.

TN: تعداد نمونه‌هایی است که به‌درستی، بعنوان کلاسی غیر از c_i کلاسیک شده‌اند.

FP: تعداد نمونه‌هایی است که به‌اشتباه، بعنوان کلاس c_i کلاسیک شده‌اند.

FN: تعداد نمونه‌هایی است که به‌اشتباه، بعنوان کلاسی غیر از c_i کلاسیک شده‌اند.

$$Accuracy = \frac{TP + TN}{TP + TN + FN + FP} \quad (2)$$

$$Precision = \frac{TP}{TP + FP} \quad (3)$$

$$Recall = \frac{TP}{TP + FN} \quad (4)$$

$$F_measure = \frac{2 * Precision * Recall}{Precision + Recall} \quad (5)$$

$$FPR = \frac{FP}{FP + TN} \quad (6)$$

$$RMSE = \sqrt{1 - r^2} \times sd \quad (7)$$

r خطای پیش‌بینی و sd انحراف معیار است. خطای بین داده آموزش و تست را خطای پیش‌بینی می‌گویند.

۳.۶. تنظیم پارامترها

در این تحقیق برای مقایسه کارایی الگوریتم CatBoost روی مجموعه داده Kronodroid، از الگوریتم‌های کلاسیک RF، SVM، DT، KNN، MLP و LR نیز استفاده می‌شوند.

در کتابخانه Scikit-Learn، هایپر پارامترهای پیش‌فرضی برای هر مدل یادگیری در نظر گرفته‌اند. معمولاً این مقادیر، برای مسأله مورد نظر بهینه نیستند. تعیین بهترین هایپر پارامترها معمولاً کار غیرممکنی است ولی با آزمون و خطا می‌توان به مقادیر مناسبی دست یافت. برای این منظور، به نتایج تجربی نیاز است. تعداد زیادی از ترکیبات هایپر پارامترها، مورد آزمون قرار می‌گیرد و کارایی هر مدل به دست آمده، ارزیابی می‌شود. جستجوی گرید، یک الگوریتم بهینه‌سازی است که بهترین پارامترها را از فهرستی از گزینه‌های پارامتری، انتخاب می‌کند. این تابع روش آزمون و خطا را خودکار می‌کند.

GridSearchCV یک کتابخانه پایتون است که فرآیند انتخاب بهترین پارامترها برای یک الگوریتم یادگیری ماشین را تسهیل می‌کند. این تابع، در زمان آموزش یک مجموعه داده، همه ترکیبات ممکن از پارامترها را ارزیابی می‌کند و بهترین ترکیب را در نظر می‌گیرد. در [۱۲]، بهترین پارامترهای الگوریتم‌های DT، SVM، KNN و RF با این تابع مشخص شده‌اند. در اینجا نیز برای این الگوریتم‌ها، از همان مقادیر پارامترها استفاده می‌شود. مقادیر پارامترهای الگوریتم‌های CatBoost، LR، Bagging و MLP نیز با این تابع مشخص می‌شوند. در جدول ۴، مقادیر پارامترهای الگوریتم‌های استفاده شده، نمایش داده شده است.

برای بررسی دقیق‌تر تأثیر پارامترها، در جدول ۵، تأثیر مقادیر پارامترهای Loss_function، Eval_metric و Learning_rate در الگوریتم CatBoost آمده است. نتایج ۸ معیار ارزیابی با ۱۱ ترکیب از

جدول ۴: تنظیم پارامترهای الگوریتم‌ها

مقدار پارامترها	الگوریتم
loss_function = 'MultiClass', eval_metric = 'AUC', learning_rate = 0.3, verbose = True	CatBoost
bootstrap = True, max_depth = 300, max_features = 'log2'	RF
Criterion = 'gini', max_depth = 19	DT
C = 20, kernel = 'rbf', probability = True	SVM
n_neighbors = 1	KNN
max_iter = 100, activation = 'relu', solver = 'adam', learning_rate = 'adaptive'	MLP
C = 0.1, penalty = None, solver = 'saga', multi_class = 'multinomial'	LR
max_features: 0.5, max_samples: 0.5, n_estimators: 500	Bagging

مقادیر پارامترهای الگوریتم CatBoost نمایش داده شده است. هر کدام از پارامترها با لیستی از مقادیر مختلف، به تابع GridSearchCV داده شده است تا الگوریتم CatBoost را با ترکیب‌های مختلف پارامترهای داده شده، اجرا کند و بهترین ترکیب مقادیر پارامترها را انتخاب کند. از بین ترکیب‌های مختلف پارامترهای، ۱۱ ترکیب در این جدول نمایش داده شده است. بهترین نتایج بولد و دومین بهترین نتایج با رنگ قرمز مشخص شده‌اند. بهترین ترکیب در سطر آخر جدول آمده است که همان ترکیبی است که به وسیله GridSearchCV انتخاب شده است.

۳.۷. آموزش مدل

داده‌ها به دو بخش، ۷۰٪ برای آموزش و ۳۰٪ برای تست تقسیم شده‌اند. در شکل ۳، شبه کد کلاسیک با تقسیم ۷۰٪ آمده است. مجموعه داده در قالب فایل اکسل است. هر سطر مربوط به ویژگی‌های یک برنامه است. در هر سطر، ستون آخر، کلاس آن برنامه و بقیه ستون‌ها، ویژگی‌های برنامه را مشخص می‌کنند. ماتریس X ۴۷۸ ستون اول داده‌ها را شامل می‌شود. ستون آخر، در بردار Y ، بعنوان برچسب داده‌ها قرار می‌گیرد. کلاسیک‌های مورد استفاده در لیست Classifiers

جدول ۵: تأثیر پارامترهای الگوریتم CatBoost در تشخیص دسته بدافزار اندروید

پارامترها	دقت	صحت	فراخوانی	معیار F	نرخ مثبت کاذب	ریشه میانگین مربعات خطا	زمان آموزش	زمان تست			
									Learning_rate	Eval_metric	loss_function
	۰.۹۲۸۲	۰.۹۲۸۱	۰.۹۲۸۲	۰.۹۲۷۳	۰.۳۶۲۸	۱.۳۹۹۴	۶۴۰.۸۶۱	۰.۴۵۳	۰.۱	AUC	MultiClass
	۰.۹۲۹۴	۰.۹۲۸۹	۰.۹۲۹۴	۰.۹۲۸۷	۰.۳۸۳۳	۱.۳۵۴۳	۳۷۳.۲۱	۰.۰۷۸	۰.۵	AUC	MultiClass
	۰.۸۸۱۶	۰.۸۸۲۳	۰.۸۸۱۶	۰.۸۸۱۸	۰.۱۸۴۹	۱.۸۹۵	۳۵۸.۸۴۷	۰.۰۶۷	۰.۸	AUC	MultiClass
	۰.۹۳۱۱	۰.۹۳۱۱	۰.۹۳۱۱	۰.۹۳۰۲	۰.۳۰۷۵	۱.۳۹۴	۳۸۱.۵۸	۰.۰۸۴	۰.۳	MCC	MultiClassOneVsAll
	۰.۹۲۵۵	۰.۹۲۵۵	۰.۹۲۵۵	۰.۹۲۴۶	۰.۳۵۷۵	۱.۴۲۱۸	۳۷۱.۶۳۲	۱.۸۷۵	۰.۱	AUC	MultiClassOneVsAll
	۰.۹۳۱۹	۰.۹۳۱۳	۰.۹۳۱۹	۰.۹۳۱۱	۰.۲۷۸۲	۱.۳۶۹۶	۴۰۱.۷۲۵	۰.۰۷۸	۰.۵	AUC	MultiClassOneVsAll
	۰.۹۳۱۲	۰.۹۳۰۹	۰.۹۳۱۲	۰.۹۳۰۴	۰.۲۴۷۹	۱.۳۷۷۷	۳۹۳.۵۲۴	۰.۰۹۴	۰.۳	MCC	MultiClass
	۰.۹۳۱۱	۰.۹۳۰۶	۰.۹۳۱۱	۰.۹۳۰۴	۰.۳۱۳۵	۱.۳۹۱۷	۳۷۲.۲۱	۰.۰۹۴	۰.۵	MCC	MultiClass
	۰.۹۳۲۷	۰.۹۳۲۹	۰.۹۳۲۷	۰.۹۳۲۲	۰.۳۹۴۱	۱.۳۵۹۷	۴۴۰.۷۷۵	۰.۱۰۹	۰.۳	AUC	MultiClassOneVsAll
	۰.۹۲۶۱	۰.۹۲۶	۰.۹۲۶۱	۰.۹۲۵	۰.۲۸۳۷	۱.۴۲۶	۳۷۳.۴۰۵	۰.۰۷۸	None	None	None
	۰.۹۳۲۸	۰.۹۳۳۲	۰.۹۳۲۸	۰.۹۳۱۹	۰.۳۶۱۳	۱.۳۶۶۲	۳۴۰.۲۵۳	۰.۰۷	۰.۳	AUC	MultiClass

قرار دارند. هر عنصر لیست، یک کلاس سبند است و همه کلاس‌بندهای لیست با حلقه For، آموزش و تست را انجام می‌دهند. با استفاده از تابع $train_test_split$ تقسیم داده‌ها به ۷۰٪ آموزش و ۳۰٪ X_train و X_test هستند. $Params$ پارامترهای کلاس‌بند $Classifier$ را مشخص می‌کند. $Model$ با تنظیم پارامترهای کلاس‌بند ایجاد می‌شود. آموزش با تابع fit انجام می‌شود. تابع $predict$ ، کلاس داده‌های تست را طبق مدل ایجادشده، پیش‌بینی می‌کند. تابع $time$ ، زمان فعلی را بر حسب ثانیه مشخص می‌کند. به کمک این تابع، زمان آموزش $Traintime$ و زمان تست $Testtime$ ، مشخص می‌شوند. تابع $CalcResults$ ، معیارهای ارزیابی را محاسبه و در $Results$ قرار می‌دهد.

۴. آزمایش‌ها و نتایج

۴.۱. نتایج

در جدول ۶، دقت الگوریتم‌های کلاس‌بندی پیاده‌سازی شده، روی ۱۵ مجموعه ویژگی انتخاب شده از مجموعه داده Kronodroid، نمایش داده شده است. بهترین نتیجه هر الگوریتم، بولد شده است. دومین بهترین نتایج با رنگ قرمز مشخص شده‌اند. بالاترین دقت را الگوریتم CatBoost با مجموعه ویژگی VFCE 100 دارد. دومین بهترین دقت نیز با الگوریتم CatBoost و مجموعه ویژگی Extra Tree 200 به دست آمده است. طبق نتایج الگوریتم‌ها در جدول ۶، مجموعه ویژگی VFCE 100 دقت بالاتری از سایر مجموعه ویژگی‌ها دارد. بالاترین دقت الگوریتم‌های RF، SVM، Bagging و CatBoost با مجموعه ویژگی VFCE 100 به دست آمده است. مجموعه ویژگی VFCE 200 در رتبه دوم قرار دارد. نتایج نشان می‌دهند که روش انتخاب ویژگی پیشنهادی VFCE، نتایج بهتری به دست می‌آورد. از مجموعه ویژگی VFCE 100 برای مقایسه نتایج الگوریتم‌ها براساس معیارهای ارزیابی، استفاده می‌شود.

جدول ۶: دقت کلاس‌بندی تشخیص دسته‌بند در مجموعه ویژگی‌های مختلف

روش انتخاب ویژگی	DT	KNN	SVM	LR	MLP	Bagging	RF	CatBoost
Variance 50	۰.۹۰۰۱	۰.۹۰۳۸	۰.۹۱۷۴	۰.۸۵۸۱	۰.۹۰۷۵	۰.۹۰۹	۰.۹۲۰۷	۰.۹۱۹۲
Variance 100	۰.۹۰۲۲	۰.۹۰۷۵	۰.۹۱۹۸	۰.۸۶۵۰	۰.۹۱۰۹	۰.۹۱۵۲	۰.۹۲۳۳	۰.۹۲۴۵
Variance 200	۰.۹۰۱۴	۰.۹۱۰۲	۰.۹۲۰۳	۰.۸۶۸۲	۰.۹۱۱۷	۰.۹۲۲۲	۰.۹۲۶۵	۰.۹۳۰۷
F_Test 50	۰.۹۰۵۸	۰.۹۰۹۹	۰.۹۰۹۶	۰.۸۵۲۲	۰.۹۰۱	۰.۹۲۲۱	۰.۹۳	۰.۹۲۸۷
F_Test 100	۰.۹۰۲۵	۰.۹۰۹۴	۰.۹۱۷۲	۰.۸۶۲۸	۰.۹۰۹۳	۰.۹۲۰۳	۰.۹۲۷۲	۰.۹۲۸۹
F_Test 200	۰.۹۰۴۵	۰.۹۱۰۶	۰.۹۲۰۶	۰.۸۶۸۳	۰.۹۱۳۳	۰.۹۱۳۳	۰.۹۲۳۷	۰.۹۳۱۹
Chi-Square 50	۰.۹۰۱۹	۰.۹۰۷۵	۰.۹۲۰۴	۰.۸۵۷۶	۰.۹۰۶۹	۰.۹۱۶۱	۰.۹۲۶۶	۰.۹۲۳
Chi-Square 100	۰.۹۰۳۴	۰.۹۰۹۴	۰.۹۲۰۱	۰.۸۶۴۸	۰.۹۱۱۹	۰.۹۲۲۷	۰.۹۲۴۸	۰.۹۲۴۹
Chi-Square 200	۰.۹۰۳۵	۰.۹۱۲	۰.۹۲۰۶	۰.۸۶۸۲	۰.۹۱۳۲	۰.۹۲۳۱	۰.۹۲۵۵	۰.۹۲۰۶
Extra Tree 50	۰.۹۰۱	۰.۹۰۳۹	۰.۹۱۳۱	۰.۸۵۵۳	۰.۹۰۳۱	۰.۹۲۱۶	۰.۹۳۰۷	۰.۹۳۰۶
Extra Tree 100	۰.۹۰۴	۰.۹۰۸۳	۰.۹۲۰۶	۰.۸۶۰۲	۰.۹۱۱۴	۰.۹۲۳۳	۰.۹۳	۰.۹۳۰۶
Extra Tree 200	۰.۹۰۱۸	۰.۹۱	۰.۹۲۰۴	۰.۸۶۸۴	۰.۹۱۳۴	۰.۹۲۴۴	۰.۹۲۵۴	۰.۹۳۲۱
VFCE 50	۰.۹۰۰۹	۰.۹۰۳۱	۰.۹۱۸	۰.۸۵۴۹	۰.۹۰۴۶	۰.۹۱۱۷	۰.۹۲۰۷	۰.۹۲۱۵
VFCE 100	۰.۹۰۳۶	۰.۹۱۱۴	۰.۹۲۰۸	۰.۸۶۲۳	۰.۹۱۱۷	۰.۹۲۸۷	۰.۹۳۰۹	۰.۹۳۲۸
VFCE 200	۰.۹۰۲۲	۰.۹۱۲	۰.۹۲۰۷	۰.۸۶۸۸	۰.۹۱۳۲	۰.۹۲۴۰	۰.۹۲۴۹	۰.۹۳۰۵

قرار دارند. هر عنصر لیست، یک کلاس سبند است و همه کلاس‌بندهای لیست با حلقه For، آموزش و تست را انجام می‌دهند. با استفاده از تابع $train_test_split$ تقسیم داده‌ها به ۷۰٪ آموزش و ۳۰٪

Inputs: *Dataset*, *Params*: params of each classifier,
Classifiers: list of classifiers

Output: *Results*: performance metrics (*Accuracy*,
Precision, *Recall*, *FMeasure*, *Traintime*, *Testtime*)

$X=Dataset [1:478]$

$Y=Dataset ['Category']$

$X=Feature_Selection (X, Y)$

$X_train, X_test, Y_train, Y_test = train_test_split (X,$
 $Y, train_size=0.70, shuffle=True)$

For Classifier in Classifiers:

$Model=Classifier (Params)$

$Tick = time ()$

$Model.fit (X_train, Y_train)$

$Traintime = time () - Tick$

$y_predict=Model.predict (Xtest)$

$Testtime = time () - Traintime$

$Results=CalcResults (Xtest, Ytest, y_predict, Model)$

شکل ۳: شبه کد الگوریتم‌های کلاس‌بندی با تقسیم ۷۰٪ تست، انجام می‌شود. پارامتر $Shuffle$ باعث درهم ریختگی داده‌ها به صورت تصادفی می‌شود تا تقسیم داده‌ها به طور متوالی انجام نگردد. قسمت داده برای آموزش، در X_train و برای تست در X_test قرار می‌گیرند. Y_train و Y_test ، به ترتیب، کلاس‌های متناظر با سطرهای

جدول ۷: نتایج کلاسنبدی تشخیص دسته بدافزار با مجموعه ویژگی VFCE 100

الگوریتم	دقت	صحت	فراخوانی	معیار F	نرخ مثبت کاذب	ریشه میانگین مربعات خطا	زمان آموزش	زمان تست
DT	۰,۹۰۳۶	۰,۹۰۳۶	۰,۹۰۳۶	۰,۹۰۳۵	۰,۰۱۵۸	۱,۷۲۶۴	۳,۰۰۸	۰,۰۳۱
KNN	۰,۹۱۱۴	۰,۹۱۰۷	۰,۹۱۱۴	۰,۹۱۱۱	۰,۰۳۳۵	۱,۵۶۴۸	۰,۰۷۸	۵,۸۵۹
SVM	۰,۹۲۰۸	۰,۹۲۰۱	۰,۹۲۰۸	۰,۹۱۹۸	۰,۰۳۴۷	۱,۵۲۳۳	۴۳۰,۴۶	۶۱,۰۲۷
LR	۰,۸۶۲۳	۰,۸۶۰۹	۰,۸۶۲۳	۰,۸۵۹۲	۰,۰۴۶۶	۱,۸۷۶۹	۵۱,۳۳۳	۰,۰۳۸
MLP	۰,۹۱۱۷	۰,۹۱۱۴	۰,۹۱۱۷	۰,۹۱۰۶	۰,۰۴۲۴	۱,۵۹۷۶	۷۰,۵۳۵	۰,۰۶۲
Bagging	۰,۹۲۸۷	۰,۹۲۸۴	۰,۹۲۸۷	۰,۹۲۶۷	۰,۰۶۲۶	۱,۴۱۶۲	۳۳۰,۴۱۶	۳۵,۹۶۱
RF	۰,۹۳۰۹	۰,۹۳۱۲	۰,۹۳۰۹	۰,۹۲۸۸	۰,۰۱۳	۱,۴۰۹۳	۱۸,۰۶	۰,۸۷۹
CatBoost	۰,۹۳۲۸	۰,۹۳۲۲	۰,۹۳۲۸	۰,۹۳۱۹	۰,۰۳۶۱۳	۱,۳۶۶۲	۳۴۰,۲۵۳	۰,۰۷

به ترتیب، بهترین نرخ مثبت کاذب را دارند. CatBoost بدترین نرخ مثبت کاذب را دارد.

نتایج کلی نشان می‌دهند که الگوریتم CatBoost با دقت ۰,۹۳,۲۸٪، صحت ۰,۹۳,۳۲٪، فراخوانی ۹۳,۲۸٪ و معیار F، ۰,۹۳,۱۹٪ و ریشه میانگین مربعات خطا ۱,۳۶۶۲ عملکرد بهتری از سایر الگوریتم‌ها دارد. نقطه ضعف آن در نرخ مثبت کاذب و زمان آموزش است.

۲.۴. مقایسه با دیگر کارها

با توجه به اینکه در این مقاله از مجموعه داده و همچنین دسته بدافزارهایی که در [۱۲] ایجاد شده است، استفاده می‌شود، نتایج به دست آمده با آن مقایسه می‌شود. مقاله مذکور در سال ۲۰۲۴ انتشار یافته است و تاکنون دسته بدافزارهای مشخص شده در این مقاله در کار دیگری بررسی نشده است. نتایج در جدول ۸، نمایش داده شده است. در [۱۲]، ۵۰ ویژگی برتر انتخاب شده با الگوریتم Extra Tree بهترین نتایج را داشت. برای مقایسه عادلانه، تعداد ویژگی‌های انتخاب شده را ۵۰ در نظر گرفته و نتایج الگوریتم‌ها روی مجموعه ویژگی VFCE 50 که بدتر از VFCE 100 هستند، در جدول آورده شده است. نرخ مثبت کاذب، ریشه میانگین مربعات خطا، زمان آموزش و زمان تست، در [۱۲] گزارش نشده است.

بهترین نتایج هر معیار ارزیابی، به ترتیب، بولد و قرمز شده است. در اینجا نیز، در همه معیارها، مدل پیشنهادی CatBoost بهترین نتایج را به دست آورده است. الگوریتم RF در رتبه دوم قرار دارد.

۵. نتیجه‌گیری

در این مقاله، تشخیص دسته بدافزار اندروید با الگوریتم CatBoost و به روش ترکیبی (ویژگی‌های ایستا و پویا) انجام شد. از دسته‌بندی بدافزارها در [۱۲] روی مجموعه داده Kronodroid استفاده شد. برای انتخاب ویژگی‌ها، روش VFCE ارائه شد. در این روش، چهار روش انتخاب ویژگی (وار یانس، F-test، Chi-Square و Extra Tree) به طور متوالی روی مجموعه داده اعمال می‌شود. براساس تعداد ۵۰، ۱۰۰ و ۲۰۰ ویژگی در هر کدام از چهار روش انتخاب ویژگی و روش پیشنهادی VFCE، ۱۵ مجموعه ویژگی را ایجاد کرده و دقت الگوریتم‌های

در جدول ۷، نتایج الگوریتم‌های کلاسنبدی برای تشخیص دسته بدافزار اندروید، براساس معیارهای ارزیابی دقت، صحت، فراخوانی، معیار F، نرخ مثبت کاذب، ریشه میانگین مربعات خطا، زمان آموزش و زمان تست نشان داده شده است. از مجموعه ویژگی VFCE 100 استفاده شده است. بهترین نتایج بولد شده‌اند. دومین بهترین نتایج با رنگ قرمز مشخص شده‌اند.

جدول ۸: مقایسه الگوریتم‌های پیاده‌سازی شده با نتایج دیگر

الگوریتم	دقت	صحت	فراخوانی	معیار F
DT [۱۲]	۰,۸۳۰۹	۰,۷۴۱۴	۰,۷۰۸۷	۰,۷۲۴۷
SVM [۱۲]	۰,۸۴۷۸	۰,۷۸۴	۰,۷۱۳۴	۰,۷۴۷
KNN [۱۲]	۰,۸۴۹۲	۰,۷۷۲۷	۰,۷۳۶۷	۰,۷۵۲
RF [۱۲]	۰,۸۷۵۶	۰,۹۰۱۴	۰,۷۴۲	۰,۸۱۳۹
Proposed DT	۰,۹۰۰۹	۰,۹۰۱	۰,۹۰۰۹	۰,۹۰۰۷
Proposed KNN	۰,۹۰۳۱	۰,۹۰۲۴	۰,۹۰۳۱	۰,۹۰۲۷
Proposed SVM	۰,۹۱۸	۰,۹۱۷۲	۰,۹۱۸	۰,۹۱۶۹
Proposed LR	۰,۸۵۴۹	۰,۸۵۱۴	۰,۸۵۴۹	۰,۸۵۰۴
Proposed MLP	۰,۹۰۴۶	۰,۹۰۶	۰,۹۰۴۶	۰,۹۰۴۳
Proposed Bagging	۰,۹۱۱۷	۰,۹۱۷۲	۰,۹۱۳۴	۰,۹۱۶۹
Proposed RF	۰,۹۲۰۷	۰,۹۲۰	۰,۹۲۰۷	۰,۹۱۹۴
Proposed CatBoost	۰,۹۲۱۵	۰,۹۲۱	۰,۹۲۱۵	۰,۹۲۰۸

الگوریتم‌های CatBoost و RF، به ترتیب؛ بهترین نتایج را در معیارهای دقت، صحت، فراخوانی، معیار F و ریشه میانگین مربعات خطا دارند. الگوریتم‌های KNN و DT، به ترتیب؛ کمترین زمان آموزش را دارند. DT و CatBoost، به ترتیب؛ کمترین زمان تست را دارند. در مجموع، CatBoost و RF، به ترتیب، رتبه اول و دوم را دارند. SVM و CatBoost بیشترین زمان آموزش را دارند. RF زمان آموزش بسیار کمتری از CatBoost دارد. یادآوری می‌شود که آموزش برای ۷۰٪ داده و تست برای ۳۰٪ دیگر داده‌ها، انجام شده است. با توجه به اینکه مدل پس از آموزش ایجاد می‌شود و پس از آن برای پیش‌بینی استفاده می‌شود، زمان تست اهمیت بیشتری دارد. زمان تست DT و CatBoost به ترتیب، ۰,۰۳۱ و ۰,۰۷ ثانیه هستند که زمان خوبی هستند. RF نیز زمان آموزش کمتر از یک ثانیه را به دست آورده است. RF و DT،

- [2] R. Shewale, 20 Android Statistics For 2024 (Market Share & Users). Retrieved December 27, 2023. Available: <https://www.demandsage.com/android-statistics>
- [3] M. Alazab, M. Alazab, A. Shalaginov, A. Mesleh, and A. Awajan, "Intelligent mobile malware detection using permission requests and API calls", *Future Generation Computer Systems*, vol. 107, pp. 509-521, 2020.
- [4] A. Kivva, Android malware and unwanted software statistics for Q1 2024, Securelist, 2024. Available: <https://securelist.com/it-threat-evolution-q1-2024-mobile-statistics/112750>.
- [5] A. Guerra-Manzanares, "Machine learning for android malware detection: mission accomplished? a comprehensive review of open challenges and future perspectives", *Computers & Security*, vol. 138, p. 103654, 2024.
- [6] R. A. Yunmar, S. S. Kusumawardani, and F. Mohsen, "Hybrid android malware detection: a review of heuristic-based approach", *IEEE Access*, vol. 12, pp. 41255-41286, 2024.
- [7] A. Muzaffar, H. R. Hassen, M. A. Lones, and H. Zantout, "An in-depth review of machine learning based Android malware detection", *Computers & Security*, vol. 121, p. 102833, 2022.
- [8] A. Afianian, S. Niksefat, B. Sadeghiyan, and D. Baptiste, "Malware dynamic analysis evasion techniques: a survey", *ACM Computing Surveys (CSUR)*, vol. 52, no. 6, pp. 1-28, 2020.
- [9] I. Rajabizadeh, and N. Modiri, "Strategic, tactical, technical, and operational requirements of executive organizations for developing cyber security vulnerability calculator", *Intelligent Multimedia Processing and Communication Systems (IMPCS)*, vol. 2, no. 3, pp. 13-23, 2021. Available: <https://sanad.iau.ir/fa/Article/903356>.
- [10] H. H. R. Manzil and S. Manohar Naik, "Android malware category detection using a novel feature vector-based machine learning model", *Cybersecurity*, vol. 6, pp. 1-11, 2023.
- [11] E. B. Karbab, M. Debbabi, A. Derhab, and D. Mouheb, "MalDozer: automatic framework for android malware detection using deep learning", *Digital Investigation*, 24, pp. S48-S59, 2018.
- [12] M. Waheed and S. Qadir, "Effective and efficient android malware detection and category classification using the enhanced kronodroid dataset", *Security and Communication Networks*, vol. 2024, no. 1, p. 7382302, 2024.
- [13] H. Bai, N. Xie, X. Di and Q. Ye, "Famd: a fast multifeature android malware detection framework, design, and implementation", in *IEEE Access*, vol. 8, pp. 194729-194740, 2020, doi: 10.1109/ACCESS.2020.3033026
- [14] K. Xu, Y. Li and R. H. Deng, "ICCDetector: icc-based malware detection on android", *IEEE Transactions on Information Forensics and Security*, vol. 11, no. 6, pp. 1252-1264, June 2016, doi: 10.1109/TIFS.2016.2523912.
- [15] A. Rahali, A. H. Lashkari, G. Kaur, L. Taheri, F. GAGNON, and F. Massicotte, "DIDroid: android malware classification and characterization using deep image learning", *Proceedings of the 2020 10th International Conference on Communication and Network Security*, Tokyo, Japan, 2021. Available: <https://doi.org/10.1145/3442520.3442522>.
- [16] D. S. Keyes, B. Li, G. Kaur, A. H. Lashkari, F. Gagnon and F. Massicotte, "EntropyLyzr: android malware classification and characterization using entropy analysis of dynamic characteristics", 2021 Reconciling Data
- مجموعه ویژگی‌ها، مقایسه می‌شوند. مجموعه داده به دو بخش ۷۰٪ برای آموزش و ۳۰٪ برای تست تقسیم می‌شود. نتایج الگوریتم‌های کلاسیک، روی مجموعه ویژگی VFCE 100 براساس معیارهای دقت، صحت، فراخوانی، معیار F، نرخ مثبت کاذب، ریشه میانگین مربعات خطا، زمان آموزش و زمان تست، مقایسه شدند. نتایج مقایسه الگوریتم‌های پیاده‌سازی شده، نشان می‌دهد که:
- روش انتخاب ویژگی پیشنهادی VFCE با تعداد ۱۰۰ ویژگی، دقت بهتری از چهار روش دیگر انتخاب ویژگی به دست آورد.
 - الگوریتم CatBoost، با دقت ۹۳٫۲۸٪، صحت ۹۳٫۳۲٪، فراخوانی ۹۳٫۲۸٪، معیار F، ۹۳٫۱۹٪ و ریشه میانگین مربعات خطا ۱٫۳۶۶۲، عملکرد بهتری از سایر الگوریتم‌ها دارد. الگوریتم RF در رتبه دوم این معیارها قرار دارد.
 - DT و CatBoost، با زمان تست کمتر از یک ثانیه، به ترتیب؛ کمترین زمان تست را دارند.
 - نقطه ضعف CatBoost در نرخ مثبت کاذب و زمان آموزش است.
- در [۱۲]، الگوریتم‌های RF، SVM، KNN، DT روی این مجموعه داده، پیاده‌سازی شده‌اند. مقایسه نتایج به دست آمده با نتایج [۱۲]، نشان می‌دهد که:
- الگوریتم CatBoost، عملکرد بهتری از دیگر الگوریتم‌ها دارد.
 - الگوریتم‌های RF، SVM، KNN و DT با روش انتخاب ویژگی پیشنهادی VFCE، نتایج بهتری از همین الگوریتم‌ها در [۱۲]، به دست می‌آورند؛ در نتیجه، روش انتخاب ویژگی پیشنهادی VFCE، موثر است.
- ## ۶. مسیر کارهای آتی و تحقیقات بعدی
- مجموعه داده Kronodroid، شامل خانواده بدافزارها نیز است. دسته بدافزار براساس خانواده بدافزار ایجاد شده است. در آینده می‌توان، تشخیص خانواده بدافزار را با CatBoost و روش انتخاب ویژگی VFCE، انجام داد. همچنین با بهبود الگوریتم CatBoost و ترکیب آن با الگوریتم‌های دیگر مانند DT، نرخ مثبت کاذب و زمان آموزش را بهبود داد.
- روش انتخاب ویژگی‌های پیشنهادی شده را در حوزه‌های دیگر مانند تشخیص بات‌نت نیز می‌توان، به کاربرد.
- با توجه به تکامل بدافزارها، خانواده و دسته آن‌ها باید در بازه‌های کم زمانی، به‌روزرسانی شود. می‌توان، در مجموعه داده Kronodroid، نمونه برنامه‌های بیشتری را اضافه کرد تا مجموعه داده، متوازن‌تر شود و همچنین از مخازن دیگر نیز برای دسته‌بندی بدافزارهای این مجموعه داده، استفاده کرد.
- ## References
- [1] M. Alazab, "Profiling and classifying the behavior of malicious codes", *Journal of Systems and Software*, vol. 100, pp. 91-102, 2015.

- Analytics, Automation, Privacy, and Security: A Big Data Challenge (RDAAPS), Hamilton, Canada, pp. 1-12, 2021.
- [17] S. MahdaviFar, A. F. Abdul Kadir, R. Fatemi, D. Alhadidi and A. A. Ghorbani, "Dynamic android malware category classification using semi-supervised deep learning", 2020 IEEE Intl Conf on Dependable, Autonomic and Secure Computing, Canada, 2020.
- [18] S. MahdaviFar, D. Alhadidi, and A. A. Ghorbani, "Effective and efficient hybrid android malware classification using pseudo-label stacked auto-encoder", *Journal of Network and Systems Management*, vol. 30, no. 1, p. 22, 2022.
- [19] S. Aurangzeb and M. Aleem, "Evaluation and classification of obfuscated android malware through deep learning using ensemble voting mechanism", *Scientific Reports*, 13, 3093, 2023.
- [20] S. I. Imtiaz, S. Rehman, A. R. Javed, Z. Jalil, X. Liu, and W. S. Alnumay, "DeepAMD: detection and identification of android malware using high-efficient deep artificial neural network", *Future Generation Computer Systems*, vol. 115, pp. 844-856, 2021.
- [21] K. Aktas and S. Sen, "Updroid: updated android malware and its familial classification", In *Secure IT Systems: 23rd Nordic Conference, NordSec 2018, Oslo, Norway, 2018, Proceedings 23* (pp. 352-368). Springer International Publishing, 2018.
- [22] G. Suarez-Tangil, S. K. Dash, M. Ahmadi, J. Kinder, G. Giacinto, and L. Cavallaro, "Droidsieve: fast and accurate classification of obfuscated android malware", In *Proceedings of the seventh ACM on conference on data and application security and privacy*, 2017.
- [23] H. Cai, N. Meng, B. Ryder, and D. Yao, "Droidcat: effective android malware detection and categorization via app-level profiling". *IEEE Transactions on Information Forensics and Security*, vol. 14, no. 6, pp. 1455-1470, 2018.
- [24] A. Guerra-Manzanares, H. Bahsi, and S. Nomm, "Kronodroid: time-based hybrid-featured dataset for effective android malware detection and characterization", *Computers & Security*, vol. 110, p. 102399, 2021.
- [25] M. Waheed and S. Qadir, Kronodroid improved dataset, 2022. Retrieved June 10, 2024 from https://github.com/semw/kronodroid_improved_hybrid_detection_v2.git.
- [26] Y. P. Khah, M. H. Shirvani, and H. Motameni, "Metaheuristic algorithms for feature selection in intrusion detection systems: a systematic review", *Intelligent Multimedia Processing and Communication Systems (IMPCS)*, vol.4, no.3, pp. 63-92, 2023. [Online]. Available: <http://sanad.iau.ir/fa/Article/1123897>.
- [27] C. Ding, N. Luktarhan, B. Lu, and W. Zhang, "A hybrid analysis-based approach to android malware family classification", *Entropy*, vol. 23, no. 8, p. 1009, 2021.
- [28] X. Wang, L. Zhang, K. Zhao, X. Ding, and M. Yu, M, "MFDroid: a stacking ensemble learning framework for Android malware detection", *Sensors*, vol. 22, no. 7, p. 2597, 2021.
- [29] L. Prokhorenkova, G. Gusev, A. Vorobev, A. V. Dorogush, and A. Gulin, "CatBoost: unbiased boosting with categorical features", *Advances in Neural Information Processing Systems*, vol. 31, pp. 6638-6648, 2018.
- [30] A. V. Dorogush, V. Ershov, and A. Gulin, "CatBoost: gradient boosting with categorical features support", *arXiv: 1810.11363*, 2018.