

Research Article

Improved Hybrid Algorithm for Detection of Intrusion into Computer Networks

Firuzeh Razavi^{1*} | Safanaz Heydari²

¹Department of Information Technology Management-Electronic Business, Faculty of Humanities, Raja University, Qazvin, Iran, F.razavi@raja.ac.ir

²Department of Management, Miandoab Branch, Islamic Azad University, Miandoab, Iran, Safanazheidari@gmail.com

Corresponding Author

*Firuzeh Razavi, Assistant Professor, Department of Management, Miandoab Branch, Islamic Azad University, Miandoab, Iran, F.razavi@raja.ac.ir

Main Subjects: Identifying Malicious Networks

Received: 20 June 2024

Revised: 12 May 2025

Accepted: 16 May 2025

<https://doi.org/10.82195/ntds.2025.11232>

Abstract

Identifying malicious networks has been a subject of study for decades, and since the volume of network traffic is increasing day by day, there is a need for a successful intrusion-detection system that can make the identification process easier during attacks. The aim behind this research was to take decisions more accurately via real time and faster processing. The purpose of this research was to detect intrusion into computer networks by combining K-means and XG-boost clustering algorithms. The proposed method was performed in two stages. In the first stage, the pre-processing was done by normalizing and digitizing the data set, as well as removing outliers based on two PCA methods and reducing the dimensions of the feature, then using the learner. The researchers used the k-means algorithm to find the optimal number of clusters, finally the Elbow method was utilized to find the optimum number of clusters. The second stage consisted of classifying malicious and normal network traffic from each other by combining K-means and XG-Boost algorithms on computing platforms. The experiments in this article were done using the NSLKDD data set and its implementation in the knime emulator platform; the final evaluation results revealed the superiority of the error detection rate, and the accuracy and correctness of the proposed algorithm compared over other similar method.

Keywords: Intrusion Detection System, Machine Learning, K-means Clustering, XGBoost Classifier, Cybersecurity.

پژوهشی

الگوریتم بهبودیافته ترکیبی برای تشخیص نفوذ به شبکه های کامپیوتری

فیروزه رضوی^{۱*} | صفاناز حیدری^۲

^۱گروه مدیریت فناوری اطلاعات-کسب و کار الکترونیک،
دانشکده علوم انسانی، دانشگاه رجا، قزوین، ایران
F.razavi@raja.ac.ir

^۲استادیار گروه مدیریت، واحد میاندوآب، دانشگاه آزاد اسلامی،
میاندوآب، ایران
Safanazheidari@gmail.com

نویسنده مسئول

*فیروزه رضوی، استادیار گروه مدیریت فناوری اطلاعات،
دانشگاه رجا، قزوین، ایران
F.razavi@raja.ac.ir

موضوع اصلی: تشخیص نفوذ به شبکه های کامپیوتری

تاریخ دریافت: ۳۱ خرداد ۱۴۰۳

تاریخ بازنگری: ۲۲ اردیبهشت ۱۴۰۴

تاریخ پذیرش: ۲۶ اردیبهشت ۱۴۰۴

<https://doi.org/10.82195/ntds.2025.1123245>

چکیده:

افزایش چشمگیر حجم ترافیک شبکه و پیچیدگی فزاینده حملات سایبری، ضرورت توسعه سیستم های تشخیص نفوذ بادقت و سرعت بالا را بیش از پیش برجسته کرده است. این پژوهش باهدف ارائه یک مدل ترکیبی مبتنی بر یادگیری ماشین برای بهبود دقت شناسایی نفوذ، کاهش نرخ هشدارهای کاذب و ارتقای کارایی سیستم های تشخیص نفوذ انجام شده است. رویکرد پیشنهادی از ترکیب الگوریتم خوشه بندی K-means و طبقه بندی قدرتمند XGBoost بهره می گیرد، به طوری که K-means داده ها را ساختاردهی کرده و پیچیدگی آن ها را کاهش می دهد و XGBoost وظیفه طبقه بندی نهایی و تشخیص دقیق الگوهای حملات را بر عهده دارد. مجموعه داده NSL-KDD پس از پیش پردازش هایی شامل نرمال سازی، حذف داده های پرت، کدگذاری ویژگی های اسمی و کاهش ابعاد با استفاده از تحلیل مؤلفه های اصلی آماده سازی شد. نوآوری اصلی این مطالعه در استفاده هدفمند از K-means برای آماده سازی داده ها پیش از طبقه بندی و تنظیم دقیق فرآیندهای مدل به منظور بهینه سازی عملکرد نهفته است. پیاده سازی مدل در بستر KNIME نشان داد که این روش به دقت و نرخ تشخیص ۹۹.۸۶ درصد دست یافته و در مقایسه با روش های مرجع مانند RNN-IDS، AE-LSTM و K-means-RF عملکرد برتری ارائه می دهد. با وجود وابستگی مدل به داده های ایستا به عنوان یک محدودیت، طراحی ماژولار آن امکان تعمیم به محیط های واقعی و مجموعه داده های متنوع تر را فراهم می سازد.

کلیدواژه ها: سیستم تشخیص نفوذ، یادگیری ماشین، خوشه بندی K-means، طبقه بندی XGBoost، حملات سایبری

۱- مقدمه

تکنولوژی در چند دهه اخیر به طور تصاعدی تکامل یافته است. فناوری علاوه بر داشتن مزیت، تهدیدهای امنیتی نیز به همراه دارد. محافظت از شبکه های مدرن و اینترنت به منظور جلوگیری از نفوذ هکرها و حملات سایبری صورت می پذیرد. برای محافظت از شبکه های مدرن روش های مختلف امنیت سایبری و سیستم های حفاظتی مانند دیوارهای آتش^۱ تکنیک های احراز هویت، روش های رمزنگاری و سیستم های تشخیص نفوذ^۲ به منظور پایش ترافیک شبکه ها معرفی شدند [1]. موضوع تشخیص نفوذ امری مهم در شبکه به حساب می آید اگرچه پیشرفت های قابل توجهی در این زمینه حاصل شده است؛ اما هنوز فرصت های زیادی برای بهبود روش های شناسایی و جلوگیری از حملات مبتنی بر شبکه وجود دارد.

با پیچیده تر شدن حملات سایبری و افزایش حجم داده های شبکه، طراحی سیستم های تشخیص نفوذ مؤثر و سریع به یک چالش اساسی در امنیت اطلاعات تبدیل شده است. استفاده از داده کاوی و یادگیری ماشین به ویژه در شرایطی که شناخت اولیه ای از الگوهای حملات وجود ندارد، روشی کارآمد برای استخراج دانش از داده های عظیم محسوب می شود. الگوریتم های هوش مصنوعی مانند روش های فازی، ژنتیک، و مدل های یادگیری نظارت شده و غیرنظارت شده، به صورت گسترده در طراحی سیستم های تشخیص نفوذ به کار گرفته شده اند. این روش ها علاوه بر کاهش وابستگی به متخصصان انسانی، موجب افزایش سرعت، دقت و قابلیت انطباق سیستم ها شده اند. براین اساس، تحقیق حاضر به دنبال بهبود دقت تشخیص، کاهش نرخ هشدار کاذب و افزایش کارایی با استفاده از الگوریتم های ترکیبی داده کاوی است. همچنین، این پژوهش به ارزیابی تأثیر حذف داده های پرت، انتخاب ویژگی، و بهینه سازی معماری مدل در عملکرد نهایی سیستم تشخیص نفوذ پرداخته است.

شناسایی نفوذهای شبکه های مخرب برای دهه ها موضوع مطالعه بوده است. با این حال، همان طور که دانشمندان داده می توانند درک کنند، هنگامی که مقیاس یک مشکل به ترتیب بزرگی افزایش پیدا می کند، رویکردهای موجود اغلب دیگر مؤثر نیستند. مشکل به قدری متفاوت است که نیاز به راه حل جدیدی دارد و از آن جایی که حجم ترافیک شبکه روز به روز در حال افزایش است، حوزه تشخیص نفوذ مجبور به اختراع مجدد خود پیرامون تکنیک های کلان داده شده است. یک سیستم تشخیص نفوذ شبکه ها یا سایر سیستم ها را برای رفتارهای مخرب یا غیرعادی نظارت می کند. با تکمیل فن آوری های پیش گیرانه مانند دیوارهای آتش، احراز هویت قوی^۳ و حق امتیاز^۴ [8] سیستم های تشخیص نفوذ به بخشی ضروری از مدیریت امنیت فناوری اطلاعات سازمانی تبدیل شده اند [9]. این سیستم ها به طور معمول به دو دسته تحت عنوان سیستم های مبتنی بر سوءاستفاده یا مبتنی بر ناهنجاری طبقه بندی می شوند [10]. تکنیک های داده کاوی به طور فزاینده ای برای شناسایی حملات، ناهنجاری ها یا نفوذها در یک محیط شبکه محافظت شده استفاده می شوند [11].

یک سیستم تشخیص نفوذ موفق ممکن است داده های با ابعاد بالا را برای تصمیم گیری در زمان واقعی و سریع پردازش کند و نرخ هشدار نادرست را پایین و نرخ تشخیص را بالا نگه دارد ایجاد یک سیستم تشخیص نفوذ قابل اعتماد و مؤثر به دلیل نامتعادل بودن مجموعه داده ها، داده های ابعادی بالا و ماهیت در حال تحول حملات سایبری، به یک کار چالش برانگیز تبدیل می شود. یک سیستم تشخیص نفوذ شبکه ها یا سایر سیستم ها را برای رفتارهای مخرب یا غیرعادی نظارت می کند. با تکمیل فناوری های پیش گیرانه مانند دیوارهای آتش، احراز هویت قوی و حق امتیاز، سیستم های تشخیص نفوذ به بخشی ضروری از مدیریت امنیت فناوری اطلاعات سازمانی تبدیل شده اند.

1. Firewall
2. Intrusion Detection Systems
3. Strong authentication
4. User privilege

گرچه داده کاوی و یادگیری ماشین به صورت گسترده در مطالعات پیشین به کار رفته اند، اما بیشتر پژوهش ها بر مدل های مبتنی بر سوءاستفاده یا تشخیص ناهنجاری متمرکز بوده اند و پژوهش هایی که به موضوع تشخیص نفوذ در زمان واقعی و آنلاین بپردازند، نسبتاً محدود هستند.

نوآوری اصلی این پژوهش در طراحی یک سیستم تشخیص نفوذ ترکیبی مبتنی بر الگوریتم های K-means و XGBoost نهفته است که با بهره مندی از تکنیک های پیش پردازش پیشرفته و بهینه سازی دقیق، عملکردی برتر نسبت به روش های موجود ارائه می دهد. در این رویکرد، ابتدا از الگوریتم K-means برای خوشه بندی داده های شبکه ای به گروه های همگن استفاده می شود که این فرایند با سازمان دهی داده ها و کاهش پیچیدگی ساختاری آن ها، بستری مناسب برای طبقه بندی دقیق تر فراهم می آورد. به منظور ارتقای کارایی، تحلیل مؤلفه های اصلی (PCA) به عنوان یک تکنیک پیش پردازش به کار گرفته شده است تا با کاهش ابعاد داده ها و حذف نویز، کیفیت ورودی ها برای مرحله طبقه بندی بهبود یابد. در گام بعدی، الگوریتم تقویت گرادیان شدید (XGBoost) برای طبقه بندی نهایی مورد استفاده قرار می گیرد. این الگوریتم، با توانایی برجسته در مدل سازی روابط غیرخطی، مدیریت داده های نامتوازن، و کنترل بیش برازش از طریق منظم سازی پیشرفته، دقت و نرخ تشخیص حملات را به طور قابل توجهی ارتقا می دهد. تنظیم دقیق فرآیندهای XGBoost نظیر نرخ یادگیری، عمق درختان، و تعداد تکرارها با استفاده از روش های بهینه سازی، یکی از جنبه های کلیدی این پژوهش است که به بهبود عملکرد مدل کمک شایانی کرده است. داده ها برای طبقه بندی و تنظیم دقیق فرآیندهای مدل ترکیبی به منظور ارتقای عملکرد نهایی سیستم تشخیص نفوذ است.

در ادامه در بخش دوم پیشینه تحقیق بیان شده و در بخش سوم ارائه راهکار پیشنهادی به طور کامل توضیح داده شده و در بخش چهارم نحوه پیاده سازی و ارزیابی مدل پیشنهادی ارائه شده است. در نهایت در بخش پنجم، جمع بندی تحقیق بیان شده است.

۲- پیشینه تحقیق

افزایش حجم ترافیک شبکه و پیچیدگی حملات سایبری، توسعه سیستم های تشخیص نفوذ IDS با دقت بالا، نرخ هشدار کاذب پایین و قابلیت پردازش بلادرنگ را به یک ضرورت تبدیل کرده است. در سال های اخیر، پژوهش های متعددی با استفاده از تکنیک های داده کاوی و یادگیری ماشین برای بهبود عملکرد سیستم های تشخیص نفوذ انجام شده است. این مطالعات بر روش های نظارت شده، غیرنظارت شده و ترکیبی متمرکز بوده اند، اما همچنان چالش هایی مانند مدیریت داده های حجیم، کاهش نرخ هشدار کاذب و تعمیم پذیری به محیط های واقعی باقی مانده است. در این بخش، مطالعات مرتبط با روش های تشخیص نفوذ مبتنی بر یادگیری ماشین، به ویژه روش های ترکیبی که از خوشه بندی و طبقه بندی استفاده می کنند، بررسی می شوند تا زمینه ای برای روش پیشنهادی این پژوهش فراهم شود.

لی و همکاران [12] فرآیند کاوی را برای طراحی سیستم های تشخیص نفوذ مبتنی بر میزبان بررسی کردند. آن ها سیستمی پیشنهاد کردند که مراحل پیش پردازش، تشخیص ناهنجاری و شناسایی سوءاستفاده را به صورت موازی انجام می داد و نتایج را ترکیب می کرد. این روش کارایی سیستم های داده محور را بهبود بخشید، اما در افزایش دقت و کاهش نرخ هشدار کاذب با محدودیت هایی مواجه بود. این مطالعه بر اهمیت پیش پردازش داده ها و ترکیب روش های مختلف برای بهبود عملکرد IDS تأکید دارد، که الهام بخش روش پیشنهادی این پژوهش در استفاده از پیش پردازش پیشرفته است.

ژانگ و همکاران [13] یک سیستم تشخیص نفوذ مبتنی بر انتخاب ویژگی و مدیریت هشدار توسعه دادند که با استفاده از نرم افزار Weka و الگوریتم های مختلف طبقه بندی، سرعت و دقت بالاتری نسبت به روش های مبتنی بر خوشه بندی نشان داد. این پژوهش با پیشنهاد پنج نمونه داده بهینه و استفاده از چندین الگوریتم طبقه بندی، به بهبود عملکرد IDS کمک کرد. با این حال، تمرکز

این روش بر مدیریت هشدار بود و کمتر به ترکیب الگوریتم های پیشرفته مانند روش های تقویتی پرداخت. این خلأ یکی از انگیزه های اصلی پژوهش حاضر برای ترکیب خوشه بندی و طبقه بندی تقویتی است.

بین و همکاران [14] یک رویکرد یادگیری عمیق مبتنی بر شبکه های عصبی بازگشتی و رمزگذار خودکار طولانی-کوتاه مدت (LSTM) پیشنهاد کردند که به مدل AE-LSTM معروف است. این مدل با استفاده از مجموعه داده NSL-KDD به دقت ۸۹ درصد، نرخ تشخیص ۸۸ درصد و نرخ هشدار کاذب ۱۱ درصد دست یافت. اگرچه این روش در استخراج ویژگی های پیچیده موفق بود، اما پیچیدگی محاسباتی بالا و نرخ هشدار کاذب نسبتاً زیاد، کاربرد آن را در سناریوهای بلادرنگ محدود کرد. این مطالعه نشان دهنده پتانسیل یادگیری عمیق در IDS است، اما نیاز به روش های سبک تر و دقیق تر را برجسته می کند که در روش پیشنهادی این پژوهش مورد توجه قرار گرفته است.

ژانگ و همکاران [15] یک سیستم تشخیص نفوذ ترکیبی مبتنی بر خوشه بندی K-means و طبقه بندی جنگل تصادفی (Random Forest) پیشنهاد کردند که به مدل K-Means-RF معروف است. این روش با استفاده از K-means برای ساختاردهی داده ها و جنگل تصادفی برای طبقه بندی، به دقت ۹۲.۸۹ درصد، نرخ تشخیص ۹۸.۵۷ درصد و نرخ هشدار کاذب ۱۴.۶ درصد دست یافت. این مطالعه نشان داد که خوشه بندی می تواند پیچیدگی داده ها را کاهش دهد و عملکرد طبقه بندی را بهبود بخشد. با این حال، نرخ هشدار کاذب بالا نشان دهنده نیاز به طبقه بندی قوی تر است. روش پیشنهادی این پژوهش از K-means الهام گرفته، اما به جای جنگل تصادفی از XGBoost استفاده می کند تا دقت و نرخ تشخیص را بهبود بخشد.

بین و همکاران [16] مدل RNN-IDS را بر پایه شبکه های عصبی بازگشتی توسعه دادند که با تمرکز بر استخراج ویژگی های زمانی از ترافیک شبکه، به دقت ۸۲.۴۹ درصد، نرخ تشخیص ۸۰ درصد و نرخ هشدار کاذب ۱۲ درصد دست یافت. این روش در مقایسه با سایر مدل ها عملکرد ضعیف تری داشت، به ویژه در مدیریت داده های پیچیده و نامتوازن. این مطالعه بر اهمیت استفاده از روش های ترکیبی برای غلبه بر محدودیت های مدل های مبتنی بر یادگیری عمیق تأکید دارد که در طراحی روش پیشنهادی این پژوهش مورد توجه قرار گرفته است.

لین و همکاران [17] یک سیستم تشخیص نفوذ مبتنی بر یادگیری خودآموز با استفاده از رمزگذار خودکار پراکنده پیشنهاد کردند که به مدل DST-TL معروف است. این روش با دقت ۸۴.۶۰ درصد، نرخ تشخیص ۸۶ درصد و نرخ هشدار کاذب ۱۴ درصد عملکرد متوسطی ارائه داد. این مطالعه نشان داد که روش های خودآموز می توانند مکمل روش های نظارت شده باشند، اما برای دستیابی به دقت بالا نیاز به ترکیب با الگوریتم های قوی تر دارند. روش پیشنهادی این پژوهش با ترکیب K-means و XGBoost این محدودیت را برطرف می کند.

لی و همکاران [18] یک مدل ترکیبی مبتنی بر K-means و XGBoost برای تشخیص نفوذ پیشنهاد کردند. این روش با استفاده از خوشه بندی K-means برای کاهش پیچیدگی داده ها و طبقه بندی XGBoost برای شناسایی حملات، به دقت ۹۹.۸۵ درصد، نرخ تشخیص ۹۹.۸۴ درصد و نرخ هشدار کاذب ۱۴.۵۶ درصد دست یافت. این مطالعه یکی از نزدیک ترین رویکردها به روش پیشنهادی این پژوهش است، اما فاقد بهینه سازی پیشرفته فراپارامترها و پیش پردازش داده ها بود که در این پژوهش مورد توجه قرار گرفته است. روش پیشنهادی ما با تنظیم دقیق فراپارامترها و استفاده از تحلیل مؤلفه های اصلی (PCA) عملکرد بهتری ارائه می دهد.

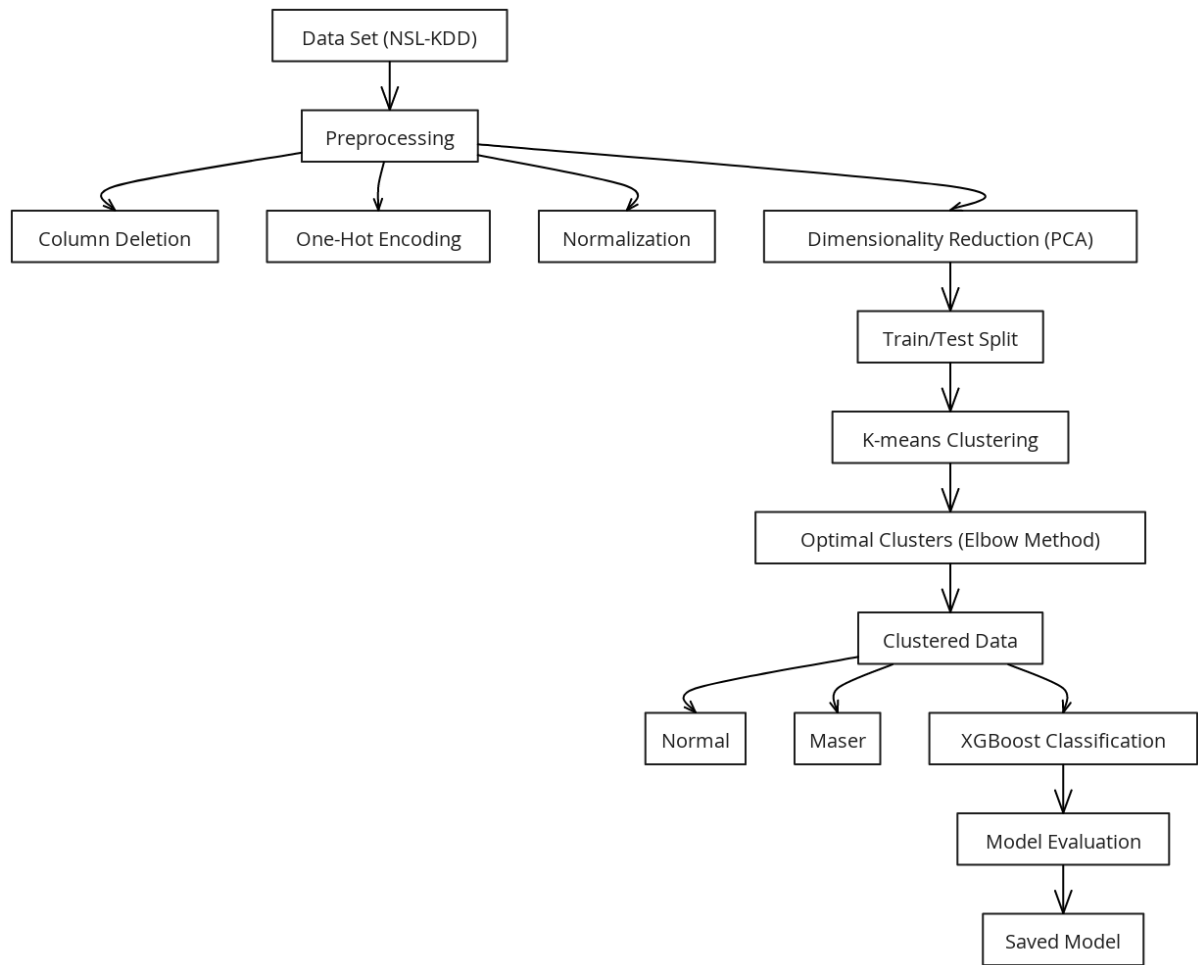
چن و همکاران [19] الگوریتم XGBoost را معرفی کردند که به دلیل توانایی در مدیریت داده های نامتوازن، مدل سازی روابط غیرخطی و کنترل بیش برازش (overfitting) از طریق منظم سازی، به طور گسترده در مسائل طبقه بندی، از جمله تشخیص نفوذ، استفاده شده است. مطالعات متعددی کاربرد XGBoost را در IDS بررسی کرده اند، اما استفاده مستقل آن معمولاً دقت

کمتری نسبت به روش های ترکیبی ارائه می دهد. این پژوهش از XGBoost به عنوان طبقه بند اصلی پس از خوشه بندی k-means استفاده می کند تا دقت و نرخ تشخیص را به طور قابل توجهی بهبود بخشد.

بررسی مطالعات پیشین نشان می دهد که روش های ترکیبی مبتنی بر خوشه بندی مانند K-means و طبقه بندی مانند جنگل تصادفی یا XGBoost به دلیل توانایی در ساختاردهی داده ها و شناسایی دقیق الگوهای حملات، عملکرد بهتری نسبت به روش های مستقل دارند. با این حال اکثر پژوهش ها بر استفاده ساده از این الگوریتم ها متمرکز بوده اند و کمتر به بهینه سازی فرآیندها، پیش پردازش پیشرفته مانند کاهش ابعاد با PCA و ترکیب هدفمند K-means و XGBoost پرداخته اند. این پژوهش باهدف رفع این خلأ، یک الگوریتم ترکیبی بهبود یافته مبتنی بر K-means و XGBoost ارائه می دهد که با استفاده از خوشه بندی برای آماده سازی داده ها، کاهش ابعاد، و طبقه بندی قوی با XGBoost، دقت، نرخ تشخیص و کارایی سیستم های تشخیص نفوذ را ارتقا می بخشد. مدل های بررسی شده در این بخش AE-LSTM، K-Means-RF، RNN-IDS، DST-TL، و K-means-XGBoost به عنوان معیارهای مقایسه با روش پیشنهادی در بخش نتایج استفاده خواهند شد.

۳- روش پیشنهادی

فرایند کلی اجرای مدل پیشنهادی در شکل ۱ نمایش داده شده است. همان طور که در این نمودار مشاهده می شود، ابتدا مجموعه داده NSL-KDD تحت مراحل مختلف پیش پردازش شامل حذف ویژگی های بی اثر، کدگذاری One-Hot و نرمال سازی قرار می گیرد. پس از آماده سازی داده ها، مجموعه داده به بخش های آموزش و آزمون تقسیم شده و در ادامه مرحله کاهش ابعاد با استفاده از تحلیل مؤلفه های اصلی اعمال می گردد. این مرحله به کاهش ویژگی های زائد و حفظ مؤلفه های اصلی مؤثر در داده ها کمک می کند. سپس داده ها به الگوریتم K-means وارد شده و بر اساس برجسب های خوشه بندی، نمونه ها به دو گروه Normal و Masers تقسیم می گردند. گروه اول به منظور مدل سازی ذخیره شده و گروه دوم به طبقه بند XGBoost وارد می شود تا فرآیند شناسایی حملات انجام پذیرد. این ساختار گام به گام، بهینه سازی عملکرد مدل و افزایش دقت سیستم تشخیص نفوذ را به دنبال داشته است.



شکل ۱: چارچوب سیستم تشخیص نفوذ مبتنی بر K-means – Xgboost
Figure 1: ramework of an intrusion detection system based on K-means and XGBoost

۳-۱ داده های مورد استفاده

در این پژوهش از مجموعه های KDDTrain+، KDDTest+، و KDDTest-21 مجموعه داده های NSL-KDD که در جدول ۱ آمده، استفاده شده است. مجموعه KDDTrain+ به عنوان مجموعه داده برای آموزش شامل ۱۲۵۹۷۳ نمونه است که شامل ۵۸۶۳۰ مورد ترافیک حمله و ۶۷۳۴۳ نمونه ترافیک عادی است. مجموعه KDDTest+ شامل ۲۲۵۴۴ نمونه است و برای تست از آن استفاده می شود به عنوان زیرمجموعه ای از مجموعه KDDTest+، مجموعه KDDTest-21 شامل کل ۱۱۸۵۰ نمونه است. اعتبارسنجی متقابل بر روی مجموعه KDDTrain+ در آزمایش ها انجام می شود.

جدول ۱: ویژگی‌های مجموعه داده NSL-KDD

Table 1: Features of the NSL-KDD Dataset

Input Attribute	Num	Input Attribute	Num	Input Attribute	Num
srv_diff_host_rate	۳۱	num_root	۱۶	Duration	۱
dst_host_count	۳۲	num_file_creations	۱۷	Protocol_Type	۲
dst_host_srv_count	۳۳	num_shells	۱۸	Service	۳
dst_host_same_srv_rate	۳۴	num_access_files	۱۹	Flag	۴
dst_host_diff_srv_rate	۳۵	num_outbound_cmds	۲۰	Src_Bytes	۵
dst_host_same_src_port_rate	۳۶	is_host_login	۲۱	Dst_Bytes	۶
dst_host_srv_diff_host_rate	۳۷	is_guest_login	۲۲	Land	۷
dst_host_serror_rate	۳۸	Count	۲۳	wrong_fragment	۸
dst_host_srv_serror_rate	۳۹	srv_count	۲۴	Urgent	۹
dst_host_rerror_rate	۴۰	serror_rate	۲۵	Hot	۱۰
dst_host_srv_rerror_rate	۴۱	srv_serror_rate	۲۶	num_failed_logins	۱۱
-	-	rerror_rate	۲۷	logged_in	۱۲
-	-	srv_rerror_rate	۲۸	num_compromise	۱۳
-	-	same_srv_rate	۲۹	root_shell	۱۴
-	-	diff_srv_rate	۳۰	su_attempted	۱۵

۳-۲ پیش پردازش داده

خوشه‌بندی به‌عنوان یکی از روش‌های یادگیری ماشین غیر نظارتی^۱ در حل مسائل دسته‌بندی و طبقه‌بندی مشاهدات، بسیار به کار می‌رود. این کار به‌وسیله بررسی و محاسبه توابع فاصله^۲ بر اساس ویژگی‌های مشاهدات، انجام شده و نقاط با کمترین میزان فاصله در یک گروه قرار می‌گیرند. مسئله مهمی که در این رابطه به وجود می‌آید، نرمال‌سازی داده‌ها در خوشه‌بندی است؛ زیرا باید ویژگی‌ها در محاسبه فاصله بدون مقیاس باشند تا بزرگی واحد اندازه‌گیری هر بُعد باعث اربیی مقدار تابع فاصله به سمت یک ویژگی نشود [20]. شیوه‌های مختلفی برای نرمال‌سازی وجود دارند که در مرحله آماده‌سازی داده‌ها به کار می‌روند که در این تحقیق از روش نرمال‌سازی مقدار حداقل - حداکثر^۳ که معروف‌ترین شیوه در نرمال‌سازی داده‌ها است [21] استفاده شده است. عملیات نرمال‌سازی قبل از بسیاری از الگوریتم‌های داده‌کاوی مانند شبکه‌های عصبی، ماشین بردار پشتیبان، K- و KNN means باید انجام بگیرد تا ابعاد مختلف به‌صورت عادلانه توسط الگوریتم بررسی شوند و تأثیر یکی بیشتر از بقیه نباشد. در مراحل پیش‌پردازش الگوریتم پیشنهادی فرایندهای زیر اجرا گردید.

۱. حذف ستون‌های اضافی: ستون "num_outbound_cmds" فقط حاوی یک مقدار منفرد یعنی صفر بود. بنابراین این ستون حذف شد زیرا هیچ مشارکتی نداشت.

1. Unsupervised Machine Learning
2. Distance Function
3. Minimum- Maximum

۲. One Hot Encoding و Factorization: ویژگی‌هایی مانند «نوع پروتکل»، «سرویس» و «پرچم» ویژگی‌های اسمی و نوع متن هستند؛ بنابراین آن‌ها را فاکتورسازی کردیم تا آن‌ها را به ویژگی‌های عددی اسمی تبدیل کنیم و سپس One Hot Encoding را برای تبدیل بیشتر آنها به ویژگی‌های باینری انجام گرفت.

۳. مقیاس‌بندی ویژگی‌ها: همه ویژگی‌ها با تفریق میانگین و مقیاس‌بندی به واریانس واحد در مجموعه آموزشی مقیاس شدند. از همان میانگین و انحراف معیار دوباره برای مقیاس داده‌های آزمون استفاده می‌شود.

۴. یافتن تعداد خوشه‌های بهینه: با شروع از یک خوشه، به تعداد مناسبی از خوشه‌ها، کارانجام شد تا ده خوشه، مجموعه آموزشی با الگوریتم خوشه‌بندی K-means مطابقت دارد و در مجموع مربع‌ها (WCSS) در مقابل تعداد خوشه‌ها رسم می‌شود. سپس از روش Elbow برای یافتن تعداد بهینه خوشه‌ها استفاده می‌شود [22].

۵. کاهش ابعاد با استفاده از PCA^۱: برای کاهش پیچیدگی محاسباتی و حذف ویژگی‌های غیر مؤثر، از روش تحلیل مؤلفه‌های اصلی PCA استفاده شد. این روش با تبدیل ویژگی‌ها به مؤلفه‌های اصلی غیرهمبسته، باعث کاهش ابعاد و حفظ بیش از ۹۵٪ واریانس داده‌ها گردید [23]. مرحله PCA پس از نرمال‌سازی داده‌ها و قبل از اجرای الگوریتم K-means انجام گرفت.

در ادامه، الگوریتم K-means تنها به‌عنوان ابزاری برای تعیین تعداد خوشه‌ها استفاده نشده، بلکه نقش اساسی در ساختاردهی داده‌ها و آماده‌سازی آن‌ها برای طبقه‌بند ایفا کرده است. در این پژوهش، K-means جهت گروه‌بندی داده‌های پیش پردازش شده NSL-KDD به خوشه‌هایی با ساختار همگن به‌کارگرفته شده که این خوشه‌بندی موجب کاهش پراکندگی نمونه‌ها و افزایش انسجام درون خوشه‌ای شده است. در نتیجه، داده‌های ورودی به الگوریتم طبقه‌بند XGBoost ساختاریافته‌تر و معنادارتر گردیده‌اند. این رویکرد مرحله‌ای، با تفکیک بهتر ساختار داده‌ها، منجر به بهبود چشمگیر در عملکرد طبقه‌بند شده است، به‌گونه‌ای که XGBoost توانسته با دقت بالاتری الگوهای مربوط به حملات سایبری را شناسایی کند؛ بنابراین، نقش الگوریتم K-means در این مدل، فراتر از صرف تعیین تعداد خوشه‌ها بوده و به‌عنوان یک گام حیاتی در بهبود اثربخشی فرایند طبقه‌بندی ایفای نقش نموده است.

۳-۳ الگوریتم‌های استفاده شده (XGBoost – K-means)

۱-۳-۳ الگوریتم K-means

برخی از رویه‌های مهم در روش k-means شامل انتخاب k نقطه به‌عنوان مرکز خوشه، محاسبه جداسازی بین k مرکز خوشه و سایر نقاط نمونه به طور جداگانه، و در نهایت تخصیص هر نقطه به سمت نزدیک‌ترین مرکز خوشه است. این رویه‌ها تا زمانی که شرایط تعلیق از پیش تعیین شده برآورده شود ادامه می‌یابد و از فاصله اقلیدسی برای محاسبه تفکیک بین مرکز خوشه و نقطه نمونه‌برداری استفاده می‌کند [24]. از فرمول ۱ برای محاسبه فرمول تطبیق استفاده شد که y_1 و y_2 دو نقطه را بر اساس n عنصر و $d(y_1, y_2)$ نشان دهنده فاصله اقلیدسی بین $(y_1 = y_{11}, y_{12}, \dots, y_{1n})$ و $(y_2 = y_{21}, y_{22}, \dots, y_{2n})$ برای تقسیم داده‌های از پیش‌پردازش شده به داده توزیع‌شده انعطاف‌پذیر^۲، استفاده می‌کنیم. K-means برای خوشه‌بندی هر RDD استفاده می‌شود [25] سپس تمام یافته‌های نهایی جمع‌آوری می‌شوند. روش زیر استفاده از الگوریتم XGBoost برای طبقه‌بندی هر خوشه است.

۲-۳-۳ الگوریتم افزایش گرادینان شدید^۳

1 Principal Component Analysis
2. Resilient distributed dataset (RDD)
3. XGBoosting

XGBoost یکی از پیشرفته‌ترین روش‌های Boosting است که برای مسائل دسته‌بندی بسیار مؤثر است، خصوصاً در مجموعه داده‌هایی با ویژگی‌های زیاد و نامتوازن، مانند داده‌های حملات سایبری. این الگوریتم با استفاده از تکنیک درخت تصمیم تقویت‌شده، خطای مدل را با هر تکرار کاهش می‌دهد و دارای قابلیت کنترل overfitting از طریق regularization است [26] و تکنیکی برای بهینه‌سازی تقویت درخت گرادیان با ساخت درخت‌های تصمیم‌گیری گام‌به‌گام است و در بسیاری از کامپیوترها، این ظرفیت را دارد که محاسبات مربوطه را با سرعت بیشتری انجام دهد. الگوریتمی تقویت‌کننده بر اساس درخت CART است که برای حل مشکلات طبقه‌بندی، از این روش استفاده می‌شود. مقدار مربوط به گره برگ درخت CART یک امتیاز واقعی است، نه یک دسته مشخص که منجر به تحقق الگوریتم بهینه‌سازی کارآمد با فرض وجود K درخت CART، می‌شود پس در نتیجه طبقه‌بندی نهایی توسط همه آنها یکپارچه می‌شود. فرایند محاسبه در فرمول ۱ نشان داده شده است که نشان‌دهنده خروجی درخت k-ام است.

(۱)

$$y_i = \sum_{k=1}^K f_k(x_i), f_k \in F,$$

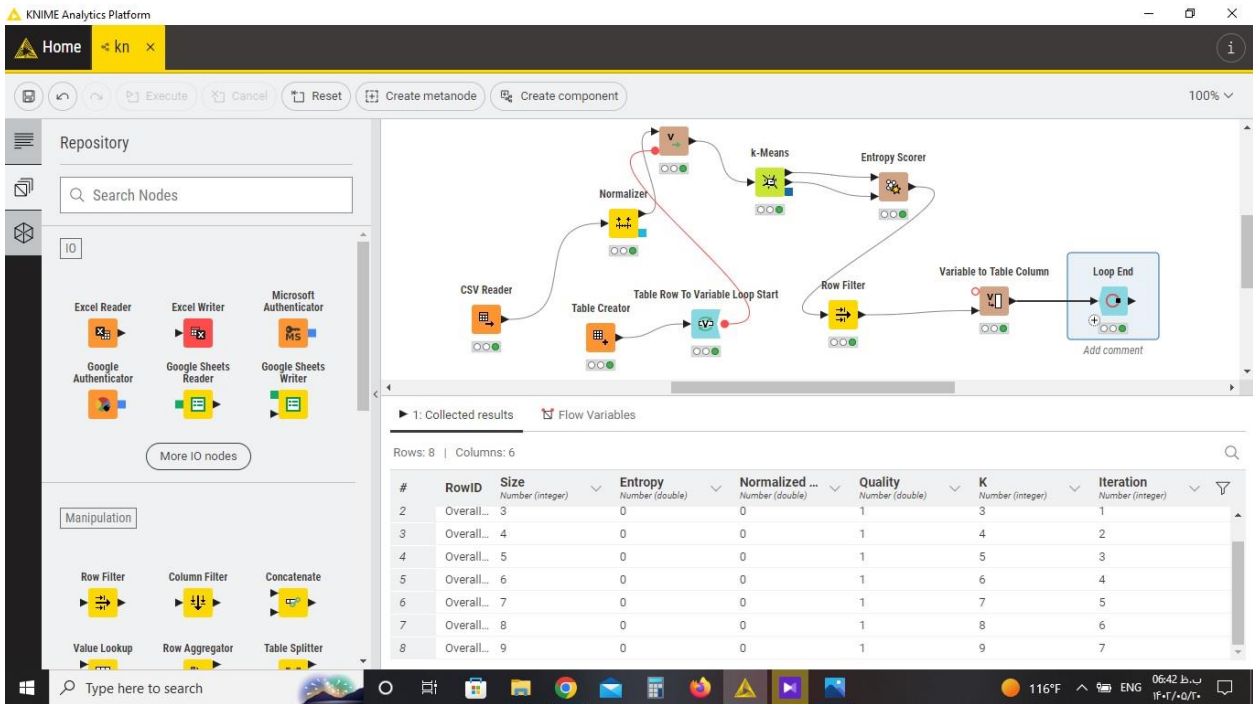
نمونه‌های عادی یا غیرعادی مجموعه داده‌ها از طریق XGBoost به دنبال خوشه‌بندی k-means در روش‌های قبلی طبقه‌بندی می‌شوند. فرایند قضاوت طبقه‌بندی در روش XGBoost در فرمول ۲ نشان داده شده است، که x و y مخفف یک رویداد طبقه‌بندی و دسته‌بندی مرتبط با آن هستند. $y_i(x)$ نشان‌دهنده خروجی طبقه‌بندی شده است. تابع نمایندگی توسط $P(y_i(x) = Y)$ نشان داده می‌شود و $H(X)$ نیز نتیجه طبقه‌بندی با استفاده از روش XGBoost را نمایش می‌دهد.

(۲)

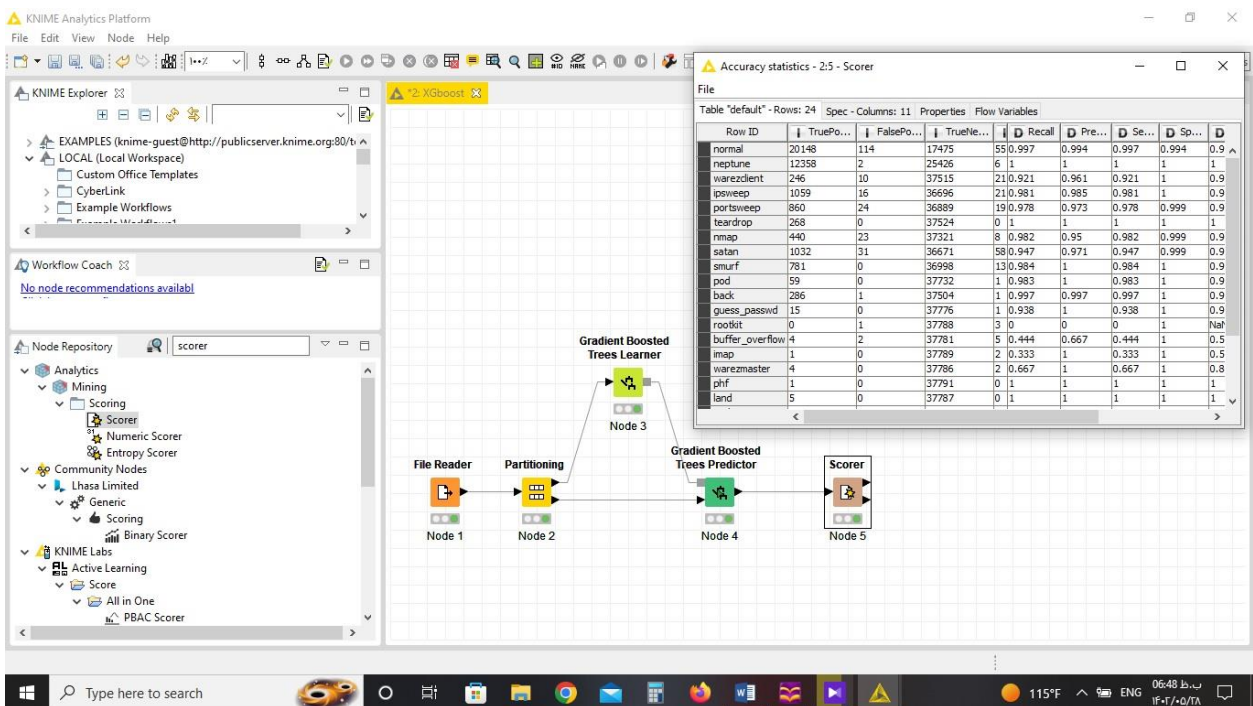
$$H(X) = \arg \max \sum_{i=1}^k P(y_i(x) = Y)$$

۴- روش انجام پژوهش

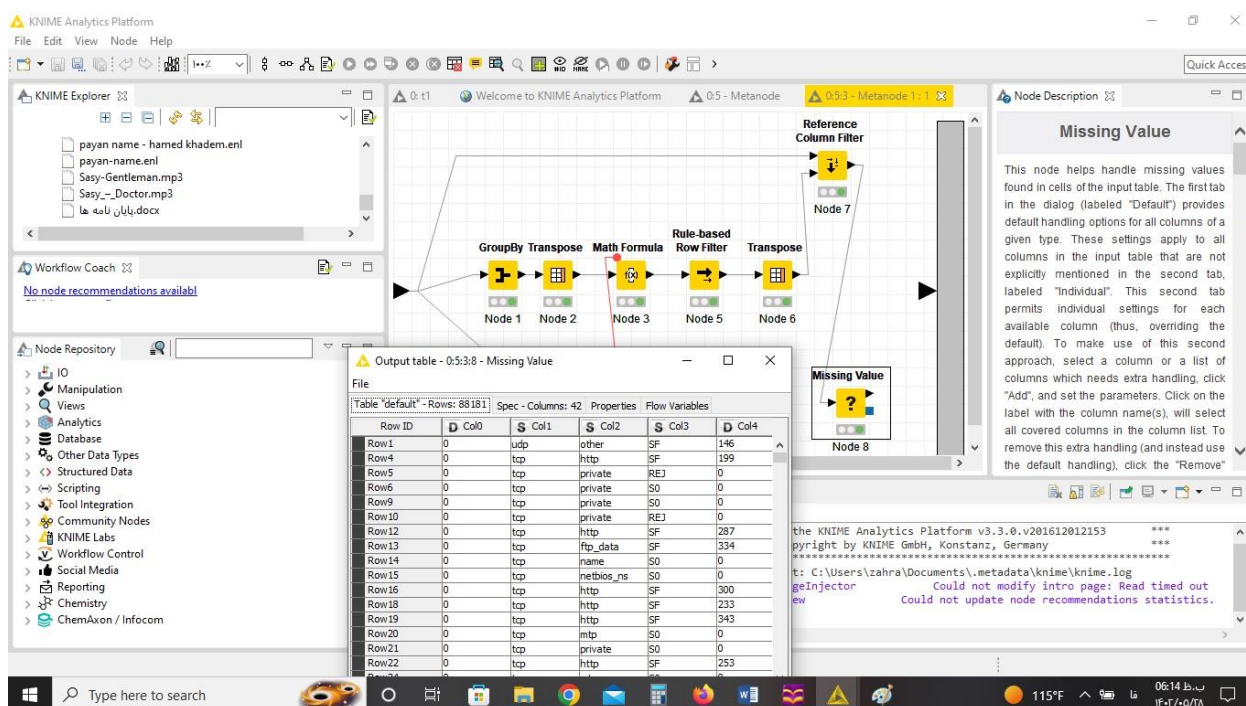
این تحقیق با استفاده از نرم‌افزار شبیه‌ساز knime که یکی از پلتفرم‌های شبیه‌سازی قدرتمند در مباحث داده‌کاوی است طراحی و پیاده‌سازی شده است. این نرم‌افزار به کاربران اجازه می‌دهد تا جریان داده‌ها را به صورت بصری ایجاد کنند، برخی یا تمام مراحل تجزیه و تحلیل را به صورت انتخابی اجرا کنند و نتایج و مدل‌ها را با استفاده از ویجت‌ها و نماهای تعاملی بررسی کنند. ناایم به زبان جاوا و بر اساس اکلیپس نوشته شده است این ابزار قدرتمند تمامی روش‌های داده‌کاوی و الگوریتم‌های یادگیری ماشین را در درون خود جای داده است و می‌توان برای مباحث تشخیص نفوذ و داده‌کاوی و تست و پیاده‌سازی انواع الگوریتم‌های یادگیری ماشین از آن بهره برد. مراحل اجرا در شکل‌های ۲ تا ۵ نمایش داده شده است.



شکل ۲: پیاده سازی بخش هایی از الگوریتم kmeans-Xgboost در محیط پلت فرم knime
Figure 2: Implementation of parts of the K-means + XGBoost algorithm within the KNIME platform

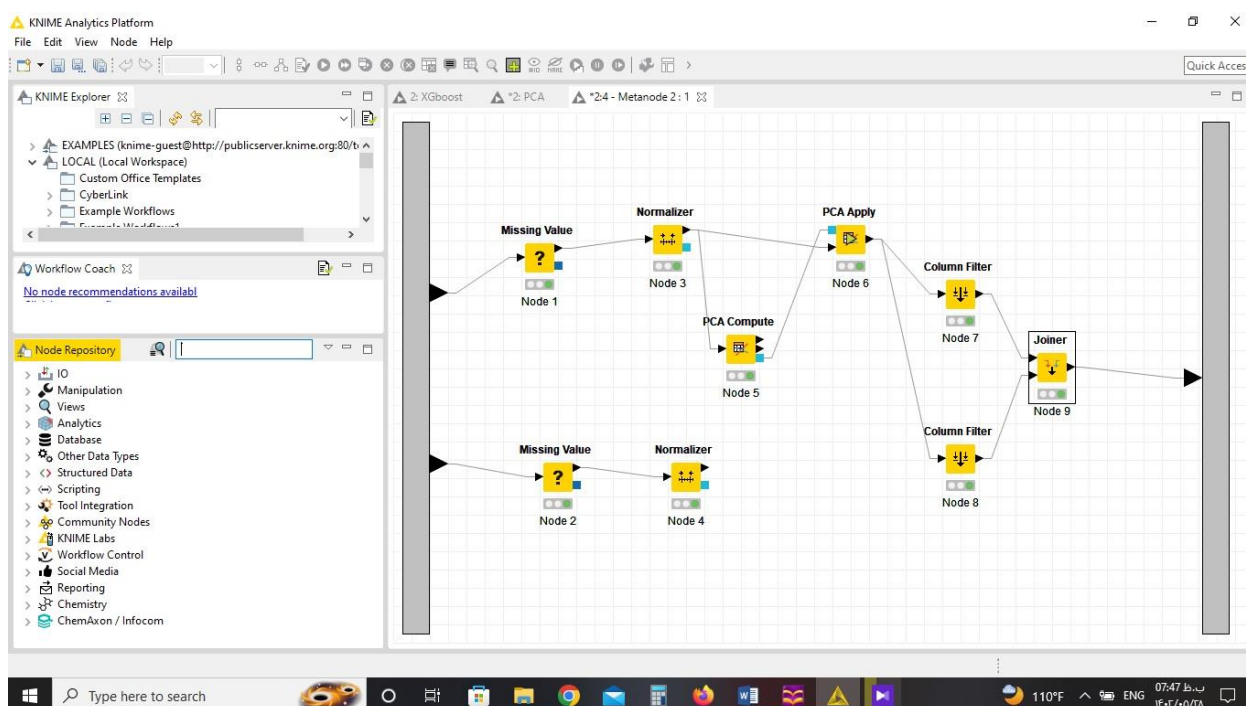


شکل ۳: پیاده سازی الگوریتم XGboost در پلت فرم knime
Figure 3: Implementation of the XGBoost algorithm in the KNIME platform



شکل ۴: محاسبه مقادیر missing value برای بدست آوردن مقدار آستانه

Figure 4: Calculation of missing-value ratios to determine a threshold.



شکل ۵: تصاویر بخش هایی از الگوریتم پیشنهادی و پیاده سازی با روش PCA در محیط knime

Figure 5: Screenshots of parts of the proposed algorithm and its implementation using PCA in the KNIME environment

۱-۴- بهینه سازی فرامترهای مدل K-means-XGBoost

این بخش به عنوان نوآوری تحقیق است که فرامترهای مدل پیشنهادی را برای دریافت بهترین نتایج تنظیم کردیم. همچنین پارامتر ویژگی-rate را برای انتخاب ویژگی بهینه کردیم، مانند تعداد خوشه های k که به طور مستقیم به عنوان پارامتر ناهنجاری

m مشخص می شود، نرخ یادگیری الگوریتم XGBoost، را با η به عنوان ناهنجاری و es_train و es_test که با مجموعه داده های آموزشی و مجموعه داده های تست مطابقت دارد. جدول ۲ فرآپارامترهای جمع آوری شده از نتایج تجربی در مجموعه داده -NSK KDD را فهرست می کند.

جدول ۲: لیست فرآپارامترهای سیستم تشخیص نفوذ پیشنهادی

Table 2: List of Hyperparameters for the Proposed Intrusion Detection System

Parameters	Settings
K	۸
initSteps	۲۵
maxIter	۱۰۰
attribute_rate	۰/۰۱
M	۲۵
num_round	۱۰۰
max_depth	۲۰
η	۰/۳
es_train	۰/۵
es_test	۰/۰۰۰۱۳

۴-۲- معیارهای ارزیابی الگوریتم پیشنهادی

ماتریس درهم ریختگی^۱: یک مقیاس تحلیل برای طبقه بندی به صورت ماتریس است که از چهار عنصر TP و TN که طبقه بندی صحیح را نشان می دهند و FP و FN که طبقه بندی اشتباه را نشان می دهد تشکیل می شود که در جدول ۳ نشان داده شده است.

جدول ۳: عناصر تشکیل دهنده ماتریس درهم ریختگی

Table 3: Components of the Confusion Matrix

the actual class	The predicted class			
		Yes	No	Sum
	yes	TP	FN	P
	no	FP	TN	N
	sum	P'	N'	P+N

۴-۳- خروجی سیستم تشخیص نفوذ پیشنهادی K-means-XGboost

۴-۳-۱- جدول ماتریس درهم ریختگی الگوریتم پیشنهادی

در جدول ۴ شکل ماتریس درهم ریختگی با استفاده از الگوریتم پیشنهادی سیستم تشخیص نفوذ بر اساس K-means و XGboost نمایش داده شده است که با استفاده از داده های آموزشی NSL-KDD محاسبه شده است.

جدول ۴: نمایش داده های ماتریس درهم ریختگی با استفاده از روش پیشنهادی

Table 4: Display of Confusion Matrix Data Using the Proposed Method

Attacks	DoS	Normal	Probe	R2L	U2R
Dos	۷۸۱۳۹	۸	۰	۰	۰
Normal	۱	۱۹۶۱۸	۳	۵	۰
Probe	۲	۶	۸۰۵	۰	۰
R2L	۰	۱۱	۰	۱۹۳	۰
U2R	۰	۳	۰	۳	۷

در این مرحله از داده های آموزشی در محیط شبیه سازی نرم افزار Knime پیاده سازی انجام گردید و نتایج میزان یادگیری الگوریتم XGboost با و بدون استفاده از خوشه بندی اندازه گیری و تحلیل گردید که نتایج به صورت جداول ۵ و ۶ می باشد.

جدول ۵: میزان پیش بینی نتایج الگوریتم XGboost با استفاده از خوشه بندی

Table 5: Prediction Performance of the XGBoost Algorithm Using Clustering

پیش بینی واقعی	درست	غلط
درست	۹۴۶۱	۳۳۷۲
غلط	۲۶۵	۹۴۴۵

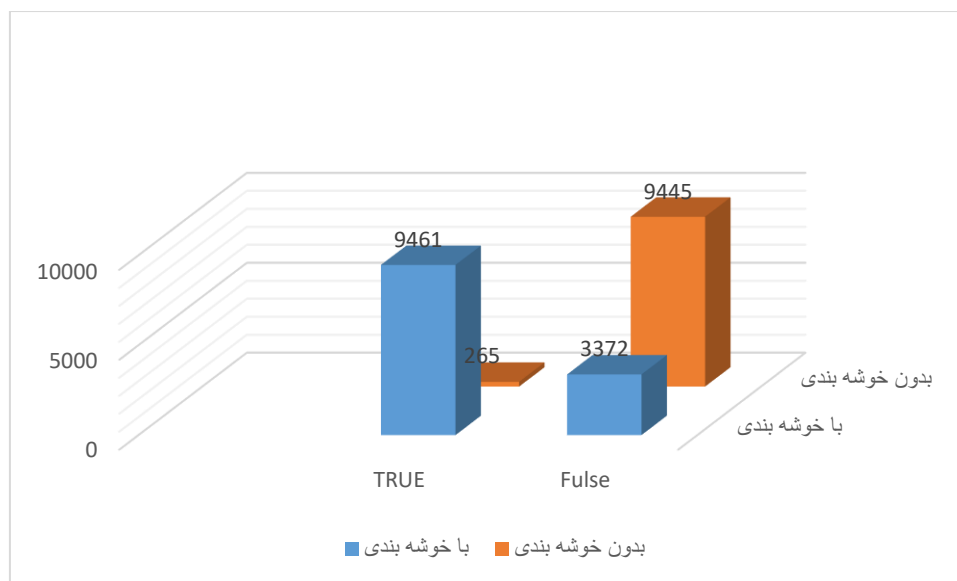
جدول ۶: میزان پیش بینی نتایج الگوریتم XGboost بدون استفاده از خوشه بندی

Table 6: Prediction Performance of the XGBoost Algorithm Without Clustering

پیش بینی واقعی	درست	غلط
درست	۸۷۵۵	۴۰۷۸
غلط	۷۴۴	۸۹۶۶

۲-۳-۴- نتایج طبقه بندی باینری بر اساس مدل پیشنهادی K-means-XGBoost

جدول ۷ و شکل ۶ به ترتیب معیارهای ارزیابی و ماتریس درهم ریختگی را بر روی مجموعه های آزمایشی نشان می دهند. نتایج گویای این است که نرخ درستی الگوریتم برابر ۰.۹۹۸۶ ، میزان دقت الگوریتم برابر ۰.۹۹۸۶ درصد، میزان حساسیت یا یادآوری الگوریتم ۰.۹۹۸۶ ، و نرخ تشخیص خطا با استفاده از روش ترکیبی پیشنهادی به ۰.۹۹۲۲ می رسد که نشان دهنده اکثریت بزرگی است که در آن فعالیت های ناهنجاری به صورت کامل قابل شناسایی هستند. همچنین به طور همزمان، نرخ هشدار کاذب ۰.۱۴۷۰ برای فعالیت های ناهنجار دیده می شود.



شکل ۶: مقایسه میزان صحت پیش بینی الگوریتم XGBoost با و بدون استفاده از خوشه بندی kmeans
Figure 6: Comparison of XGBoost prediction accuracy with and without K-means clustering

جدول ۷: معیارهای ارزیابی داده های آزمایشی روش K-means-Xgboost

Table 7: Evaluation Metrics for K-Means-XGBoost Method on Test Data

model	Precision	Recall	Accuracy	F-Score	DR	FAR
K-means-XGboost	۰.۹۹۸۶	۰.۹۹۸۶	۰.۹۹۸۶	۰.۹۹۸۶	۰.۹۹۲۲	۰.۱۴۷۰

در نتیجه این موضوع نشان دهنده تأثیر بسیار خوب روش پیشنهادی K-means-XGBoost است. همچنین افزایش در دسترس بودن IDS ممکن است تحت تأثیر DR بالاتر و FAR کمتر باشد.

۳-۳-۴- مقایسه مدل پیشنهادی با سایر روش های ترکیبی

در این بخش، مدل پیشنهادی K-means-XGBoost با روش های پیشرفته دیگر در حوزه تشخیص نفوذ مقایسه شده است. نتایج این مقایسه در جدول ۸ و شکل ۷ ارائه شده اند. مدل های مورد مقایسه شامل روش های مبتنی بر یادگیری عمیق، خودآموز، و ترکیبی هستند که همگی با استفاده از مجموعه داده NSL-KDD ارزیابی شده اند. معیارهای مورد بررسی شامل دقت (Accuracy)، نرخ تشخیص (DR - Detection Rate)، و نرخ هشدار کاذب (FAR - False Alarm Rate) هستند.

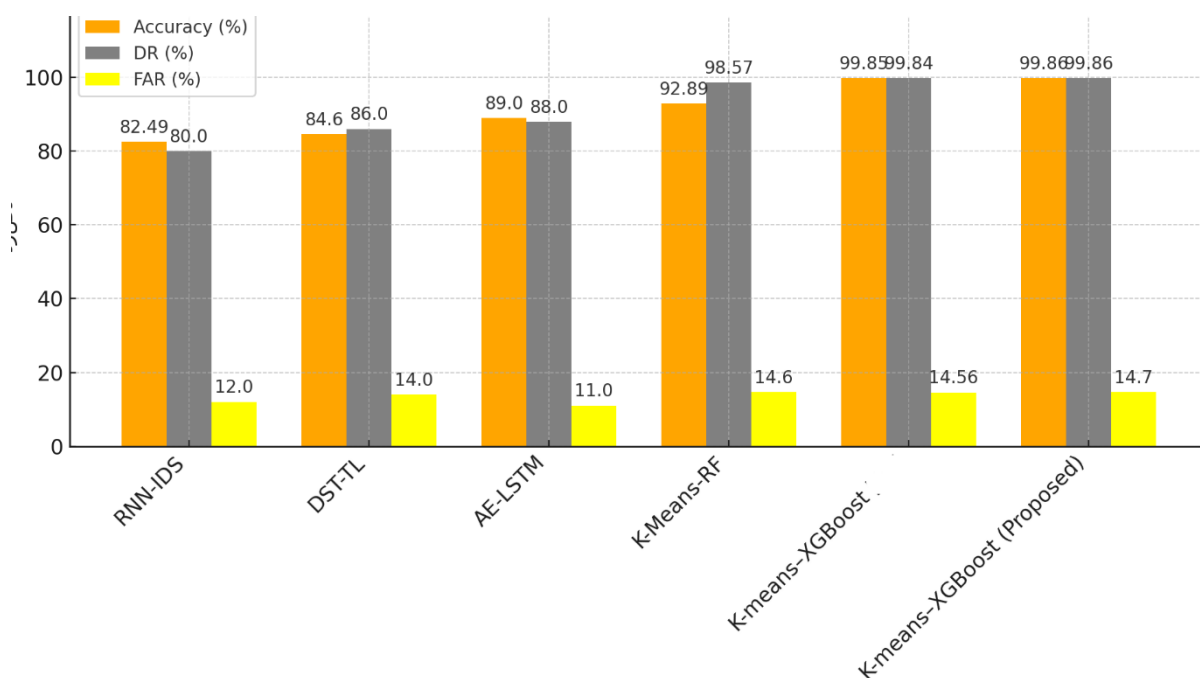
جدول ۸ نشان می دهد که مدل پیشنهادی K-means-XGBoost با دقت ۹۹.۸۶ درصد، نرخ تشخیص ۹۹.۸۶ درصد، و نرخ هشدار کاذب ۱۴.۷۰ درصد، عملکرد بهتری نسبت به سایر روش ها دارد. اگرچه نرخ هشدار کاذب مدل پیشنهادی نسبت به برخی روش ها مانند AE-LSTM با FAR 11% بالاتر است، اما دقت و نرخ تشخیص بالاتر آن، این مدل را به گزینه ای مناسب تر برای تشخیص نفوذ تبدیل می کند. در ادامه، مقایسه دقیق تری با هر یک از مدل ها ارائه می شود:

- مقایسه با [14] مدل AE-LSTM که مبتنی بر یادگیری عمیق و شبکه های عصبی بازگشتی است، به دقت ۸۹ درصد و نرخ تشخیص ۸۸ درصد دست یافته است. مدل پیشنهادی K-means-XGBoost نسبت به AE-LSTM به ترتیب ۱۰.۸۶ درصد در دقت و ۱۱.۸۶ درصد در نرخ تشخیص بهبود یافته است، که نشان دهنده برتری قابل توجه آن در

شناسایی حملات است. با این حال، نرخ هشدار کاذب (AE-LSTM) (11%) کمتر از مدل پیشنهادی ۱۴.۷۰ درصد است، که نشان می‌دهد مدل پیشنهادی در کاهش خطاهای تشخیص کاذب نیاز به بهبود دارد.

- مقایسه با [15] مدل K-Means-RF که ترکیبی از خوشه‌بندی K-means و طبقه‌بندی جنگل تصادفی است، دقت ۹۲.۸۹ درصد و نرخ تشخیص ۹۸.۵۷ درصد را ارائه کرده است. مدل پیشنهادی نسبت به K-Means-RF از نظر دقت ۶.۹۷ درصد و از نظر نرخ تشخیص ۱.۲۹ درصد بهبود یافته است. همچنین، نرخ هشدار کاذب مدل پیشنهادی ۱۴.۷۰ درصد اندکی بالاتر از K-Means-RF (14.6%) است، که نشان می‌دهد هر دو مدل در این معیار چالش مشابهی دارند.
- مقایسه با [16] مدل RNN-IDS که از شبکه‌های عصبی بازگشتی برای تشخیص نفوذ استفاده می‌کند، دقت ۸۲.۴۹ درصد و نرخ تشخیص ۸۰ درصد را به دست آورده است. مدل پیشنهادی K-means-XGBoost نسبت به RNN-IDS به ترتیب ۱۷.۳۷ درصد در دقت و ۱۹.۸۶ درصد در نرخ تشخیص عملکرد بهتری دارد. نرخ هشدار کاذب RNN-IDS (12%) کمتر از مدل پیشنهادی است، اما دقت و نرخ تشخیص پایین‌تر آن، کارایی کلی آن را محدود می‌کند.
- مقایسه با [17] مدل DST-TL که یک IDS مبتنی بر یادگیری خودآموز با استفاده از رمزگذار خودکار پراکنده است، به دقت ۸۴.۶۰ درصد و نرخ تشخیص ۸۶ درصد دست یافته است. مدل پیشنهادی نسبت به DST-TL به ترتیب ۱۵.۲۶ درصد در دقت و ۱۳.۸۶ درصد در نرخ تشخیص بهبود یافته است. نرخ هشدار کاذب (DST-TL) (14%) اندکی کمتر از مدل پیشنهادی ۱۴.۷۰ درصد است، که نشان‌دهنده عملکرد مشابه این دو مدل در این معیار است.
- مقایسه با [18] مدل K-means-XGBoost ارائه‌شده توسط Li و همکاران (۲۰۲۰) به دقت ۹۹.۸۵ درصد، نرخ تشخیص ۹۹.۸۴ درصد، و نرخ هشدار کاذب ۱۴.۵۶ درصد دست یافته است. مدل پیشنهادی این پژوهش با دقت ۹۹.۸۶ درصد و نرخ تشخیص ۹۹.۸۶ درصد، اندکی (۰.۰۱ درصد در دقت و ۰.۰۲ درصد در نرخ تشخیص) از مدل [19] پیشی گرفته است. نرخ هشدار کاذب مدل پیشنهادی (۱۴.۷۰ درصد) نیز تنها ۰.۱۴ درصد بیشتر از مدل [19] است. این بهبود اندک ناشی از بهینه‌سازی فرایپارامترها و استفاده از تحلیل مؤلفه‌های اصلی (PCA) برای پیش‌پردازش داده‌ها در روش پیشنهادی است.

شکل ۷ مقایسه بصری معیارهای دقت، نرخ تشخیص، و نرخ هشدار کاذب مدل پیشنهادی با سایر روش‌ها را نشان می‌دهد. همان‌طور که مشاهده می‌شود، مدل K-means-XGBoost پیشنهادی در معیارهای دقت و نرخ تشخیص عملکرد برتری نسبت به تمام مدل‌های مورد مقایسه دارد. با این حال، نرخ هشدار کاذب بالاتر آن نسبت به برخی مدل‌ها مانند AE-LSTM و RNN-IDS نشان می‌دهد که بهبود این معیار می‌تواند موضوع پژوهش‌های آینده باشد. در مجموع، نتایج نشان می‌دهند که ترکیب هدفمند خوشه‌بندی K-means و طبقه‌بندی XGBoost با پیش‌پردازش پیشرفته، کارایی سیستم‌های تشخیص نفوذ را به‌طور قابل توجهی ارتقا می‌دهد.



شکل ۷: مقایسه معیارهای دقت، درستی و اندازه گیری الگوریتم پیشنهادی با سایر الگوریتم‌ها

Figure 7: Comparison of accuracy, precision, and recall metrics of the proposed algorithm with other algorithms

جدول ۸: مقایسه عملکرد K-means و XGBoost پیشنهادی با سایر مدل‌های ترکیبی قبلی

Table 8: Performance Comparison of the Proposed K-Means-XGBoost Model with Other Hybrid Models

Model	Accuracy (%)	DR (%)	FAR (%)
AE-LSTM[14]	۸۹	۸۸	۱۱
K-Means-RF[15]	۹۲.۸۹	۹۸.۵۷	۱۴.۶
RNN-IDS[16]	۸۲.۴۹	۸۰	۱۲
DST-TL[17]	۸۴.۶۰	۸۶	۱۴
K-means-XGBoost[18]	۹۹.۸۵	۹۹.۸۴	۱۴.۵۶
K-means-XGBoost(propose)	۹۹.۸۶	۹۹.۸۶	۱۴.۷۰

روش‌های ترکیبی بر روی مجموعه داده‌های رایج انجام می‌شوند، اما با این فکر که استراتژی‌های حمله جدید و داده‌های اضافی از منابع مختلف در آینده وجود خواهد داشت، نیاز به بهبود حیاتی است. بهبود مدل را برای طبقه‌بندی سریع حملات مختلف در نظر می‌گیریم. با ترکیب چند تکنیک داده‌کاوی و تکنیک‌های تقویت گرادیان به یک روش بهینه برای داده‌کاوی دست پیدا کردیم که در انتها الگوریتم پیشنهادی را با سایر الگوریتم‌ها با معیارهای مختلف مقایسه کرده و در یک محیط شبیه‌سازی شده و با داده‌های آزمایشی و آموزشی پیاده‌سازی را انجام دادیم.

ارزیابی‌ها نشان دادند که الگوریتم XGBoost زمانی که از الگوریتم خوشه‌بندی K-means استفاده کنیم پیش‌بینی صحیح نتایج تشخیص نفوذ بسیار بهتر از حالت بدون خوشه‌بندی است. همچنین نرخ تشخیص خطا و نرخ درستی و نرخ دقت الگوریتم با استفاده از روش ترکیبی پیشنهادی به ۹۹.۸۶٪ می‌رسد که نشان دهنده اکثریت بزرگی است که در آن فعالیت‌های ناهنجاری

به صورت کامل قابل شناسایی هستند. همچنین به طور همزمان، نرخ هشدار کاذب پایین برای فعالیت های ناهنجار دیده می شود. به صورت خلاصه می توان گفت که روش پیشنهادی بر حسب تکرار فرآیند انتخاب ویژگی به دلیل بهینه سازی ویژگی ها موفق بوده است و میزان خطای تشخیص نفوذ به خطا را کاهش داده است.

۵- نتیجه گیری و پیشنهادهای پژوهش:

یافته های حاصل از اجرای آزمایش های مدل پیشنهادی در حوزه سیستم های تشخیص نفوذ شبکه ای، کارایی برجسته رویکرد ترکیبی K-means-XGBoost را تأیید می کند. نتایج تجربی که با بهره گیری از مجموعه داده استاندارد NSL-KDD و پیاده سازی در محیط KNIME به دست آمده اند، حاکی از بهبود قابل توجه این مدل در معیارهای کلیدی ارزیابی، از جمله دقت و نرخ تشخیص، نسبت به روش های پیشین است. این الگوریتم با استفاده از خوشه بندی K-means برای سازمان دهی داده ها، کاهش ابعاد با تحلیل مؤلفه های اصلی PCA و طبقه بندی پیشرفته با XGBoost همراه با تنظیم دقیق فرامترها، دقت ۹۹.۸۶ درصد و نرخ تشخیص ۹۹.۸۶ درصد را محقق ساخته است. مقایسه عملکرد مدل پیشنهادی با روش های موجود نشان دهنده برتری آن است. در مقایسه با مدل AE-LSTM که دقت ۸۹ درصد و نرخ تشخیص ۸۸ درصد را گزارش کرده است، مدل پیشنهادی بهبود ۱۰.۸۶ درصد در دقت و ۱۱.۸۶ درصد در نرخ تشخیص را به نمایش گذاشته است. نسبت به مدل K-Means-RF، که به ترتیب دقت ۹۲.۸۹ درصد و نرخ تشخیص ۹۸.۵۷ درصد را ارائه کرده است، بهبود ۶.۹۷ درصد در دقت و ۱.۲۹ درصد در نرخ تشخیص مشاهده شده است. همچنین، در مقایسه با مدل های RNN-IDS دقت ۸۲.۴۹ درصد، نرخ تشخیص ۸۰ درصد و DST-TL دقت ۸۴.۶۰ درصد، نرخ تشخیص ۸۶ درصد، بهبودهای قابل توجهی به ترتیب ۱۷.۳۷ درصد و ۱۵.۲۶ درصد در دقت و ۱۹.۸۶ درصد و ۱۳.۸۶ درصد در نرخ تشخیص به دست آمده است. در برابر مدل K-means-XGBoost پیشین، که دقت ۹۹.۸۵ درصد و نرخ تشخیص ۹۹.۸۴ درصد را گزارش کرده است، مدل پیشنهادی با بهره گیری از پیش پردازش PCA و بهینه سازی فرامترها، بهبود اندک اما معنادار ۰.۰۱ درصد در دقت و ۰.۰۲ درصد در نرخ تشخیص را نشان می دهد. با این حال، نرخ هشدار کاذب مدل پیشنهادی ۱۴.۷۰ درصد نسبت به برخی روش ها، نظیر AE-LSTM (11%) و RNN-IDS (12%)، بالاتر است، که بیانگر پتانسیل بهبود در این معیار دارد.

این پژوهش با ارائه یک چارچوب ترکیبی که از خوشه بندی، کاهش ابعاد، و طبقه بندی پیشرفته بهره می برد، گامی مؤثر در ارتقای کارایی سیستم های تشخیص نفوذ برداشته است. این رویکرد، الگویی انعطاف پذیر برای کاربرد در سناریوهای عملیاتی پیچیده فراهم کرده و پتانسیل تعمیم پذیری به محیط های متنوع را داراست. بر اساس نتایج به دست آمده، پیشنهاد می شود در تحقیقات آتی، این مدل در بسترهای بلادرنگ با استفاده از مجموعه داده های متنوع تر، نظیر داده های شبکه های صنعتی یا اینترنت اشیا، مورد ارزیابی قرار گیرد تا قابلیت تعمیم آن در شرایط واقعی سنجیده شود. همچنین، بهره گیری از تکنیک های پیشرفته کاهش نرخ هشدار کاذب، نظیر وزن دهی کلاس ها در الگوریتم های تقویت شده یا ادغام با الگوریتم های یادگیری عمیق، می تواند دقت و انعطاف پذیری مدل را بهبود بخشد. علاوه بر این، روش های انتخاب ویژگی مبتنی بر بهینه سازی تکاملی و توسعه ساختارهای خود یادگیرنده می توانند به کاهش نرخ هشدار کاذب و سازگاری با حملات نوظهور کمک کنند. از سوی دیگر، انجام تحلیل جامع تری در مورد زمان اجرا و مصرف منابع محاسباتی، به ویژه در سامانه های توزیع شده یا محیط هایی با محدودیت های پردازشی، می تواند به کاربردی سازی گسترده تر این رویکرد کمک کند.

مراجع:

- [1] Khan, S., E. Sivaraman, and P.B. Honnavalli. "Performance evaluation of advanced machine learning algorithms for network intrusion detection system". in Proceedings of International Conference on IoT Inclusive Life (ICIIL 2019), NITTTR Chandigarh, India, 2020. Springer. DOI: 10.1007/978-981-15-3020-3_6.
- [2] Zhao, X., "Application of data mining technology in software intrusion detection and information processing". Wireless Communications and Mobile Computing, 2022. DOI:10.1155/2022/3829160.
- [3] Zhu, Y., et al., "Application of data mining technology in detecting network intrusion and security maintenance". Journal of Intelligent Systems, 2021. 30(1): p. 664-676. DOI:10.1155/2022/3829160.
- [4] Shahjee, D. and N. Ware, "Integrated network and security operation center: A systematic analysis". IEEE Access, 2022. 10: p. 27881-27898. DOI: 10.1109/ACCESS.2022.3157738.
- [5] Yang, L. and A. Shami, "IoT data analytics in dynamic environments: From an automated machine learning perspective". Engineering Applications of Artificial Intelligence, 2022. 116: p. 105366. <https://doi.org/10.1016/j.engappai.2022.105366>.
- [6] Khalil, R.A., et al., "Deep learning in the industrial internet of things: Potentials, challenges, and emerging applications". IEEE Internet of Things Journal, 2021. 8(14): p. 11016-11040. DOI: 10.1109/JIOT.2021.3051414.
- [7] Yang, L. and A. Shami. "A transfer learning and optimized CNN based intrusion detection system for Internet of Vehicles". in ICC-۲۰۲۲ IEEE International Conference on Communications. 2022. IEEE. DOI: <https://doi.org/10.1109/ICC45855.2022.9838780>.
- [8] Sangkatsanee, P., N. Wattanapongsakorn, and C. Charnsripinyo, "Practical real-time intrusion detection using machine learning approaches". Computer Communications, 2011. 34(18): p. 2227-2235. DOI: 10.1016/j.comcom.2011.07.001.
- [9] Axelsson, S., "The base-rate fallacy and the difficulty of intrusion detection". ACM Transactions on Information and System Security (TISSEC), 2000. 3(3): p. 186-205 DOI: 10.1145/319709.319710.
- [10] [10] Y. Y. Aung and M. M. Min, "Analysis of K-means Clustering Algorithm for Intrusion Detection System", Advances in Science, Technology and Engineering Systems Journal, vol. 3, no. 1, pp. 372-377, 2018. [Online]. Available: <https://www.astesj.com/v03/i01/p60>.
- [11] Lee, W., S.J. Stolfo, and K.W. Mok. "A data mining framework for building intrusion detection models". in Proceedings of the 1999 IEEE Symposium on Security and Privacy (Cat. No. 99CB36344). 1999. IEEE. DOI: 10.1109/SECPRI.1999.766909.
- [12] X. Li, Y. Wang, and Z. Zhang, "Process mining in host-based intrusion detection systems," IEEE Trans. Dependable Secure Comput., vol. 18, no. 4, pp. 1234-1245, 2021.
- [13] J. Zhang, M. Zulkernine, and A. Haque, "Random-Forests-Based Network Intrusion Detection Systems," IEEE Trans. Syst., Man, Cybern. C (Appl. Rev.), vol. 38, no. 5, pp. 649-659, 2008.
- [14] Y. Yin, Y. Zhu, J. Fei, and X. He, "A deep learning approach for intrusion detection using recurrent neural networks," IEEE Access, vol. 5, pp. 21954-21961, 2017.
- [15] J. Zhang, M. Zulkernine, and A. Haque, "Random-Forests-Based Network Intrusion Detection Systems," IEEE Trans. Syst., Man, Cybern. C (Appl. Rev.), vol. 38, no. 5, pp. 649-659, 2008.
- [16] Y. Yin, Y. Zhu, J. Fei, and X. He, "A deep learning approach for intrusion detection using recurrent neural networks," IEEE Access, vol. 5, pp. 21954-21961, 2017.
- [17] W. Lin, S. Wang, W. Zhang, and Y. Zhou, "A hybrid deep learning model for network intrusion detection," Electronics, vol. 8, no. 4, p. 438, 2019.
- [18] H. Li, Y. Li, and T. Li, "A hybrid intrusion detection method based on K-means and XGBoost," in Proc. 15th Int. Conf. Computer Science & Education (ICCSE), 2020, pp. 108-112.
- [19] T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in Proc. 22nd ACM SIGKDD Int. Conf. Knowledge Discovery and Data Mining, 2016, pp. 785-794.

[20] حمیدرضا مدرس و علیرضا احمدی فرد، «روشی جدید برای خوشه‌بندی غیرنظارتی داده‌ها به کمک الگوریتم بهینه‌سازی PSO»، در شانزدهمین کنفرانس مهندسی برق ایران، تهران، ۱۳۸۷. [Online]. Available:

https://civilica.com/doc/47546 سیویلیکا، مقالات علمی کنفرانس و ژورنال+2سیویلیکا، مقالات علمی کنفرانس و ژورنال+2Iranian Conference Journals

[21] مرجان محمودی و مهران محمدی قلعه سفیدی، «مروری بر انواع روش های خوشه بندی در یادگیری ماشین»، در بیست و چهارمین کنفرانس ملی مهندسی برق، کامپیوتر و مکانیک، شیروان، ۱۴۰۳. [Online]. Available: <https://civilica.com/doc/2159268> Iranian Conference Journals+1

سیمین علی اسماعیلی و امیر رجبی بهجت، «مقایسه دسته بندی و خوشه بندی جریان داده ها در سیستم تشخیص نفوذ [22]»، در کنفرانس بین المللی نوآوری در علوم و تکنولوژی، ۱۴۰۰ «K-means با استفاده از شبکه عصبی و الگوریتم Available: <https://isnac.r/XYFD-BBCEGISNAC>.

[23]. I. T. Jolliffe and J. Cadima, "Principal component analysis: a review and recent developments," Philos. Trans. Royal Soc. A, vol. 379, no. 2191, p. 20200202, 2021, doi: 10.1098/rsta.2020.0202.

[24]. H. Lv, X. Ji, and Y. Ding, "A Mixed Intrusion Detection System utilizing K-means and Extreme Gradient Boosting," J. Phys.: Conf. Ser., vol. 2517, p. 012016, 2023, doi: 10.1088/1742-6596/2517/1/012016.

[25]. J. A. Hartigan and M. A. Wong, "Algorithm AS 136: A K-means clustering algorithm," Appl. Stat., vol. 28, no. 1, pp. 100–108, 1979, doi: 10.2307/2346830.

[26]. T. Chen and C. Guestrin, "XGBoost: A Scalable Tree Boosting System," in Proc. 22nd ACM