

دریافت مقاله: ۱۴۰۳/۰۷/۰۹

پذیرش مقاله: ۱۴۰۳/۰۸/۰۵

ارائه یک مدل یادگیری عمیق برای سیستم تشخیص نفوذ با استفاده از لایه هم آمیخت و روش تحلیل مولفه‌های اصلی

سید محمد جوادی مقدم*

دانشکده مهندسی، دانشگاه بزرگمهر قائنات، قاین، ایران

پست الکترونیکی: smjavadim@buqaen.ac.ir

حسین غلامعلی نژاد

دانشکده مهندسی، دانشگاه بزرگمهر قائنات، قاین، ایران

پست الکترونیکی: h_gholamalinejad@buqaen.ac.ir

طه‌ورا رمضانی مقدم

دانشکده مهندسی، دانشگاه بزرگمهر قائنات، قاین، ایران

پست الکترونیکی: tahooraramezani79@gmail.com

چکیده

امنیت و پایداری شبکه‌ها نقش اساسی در دنیای ارتباطات دارد. سیستم تشخیص نفوذ یک عنصر حیاتی در شناسایی و پیشگیری از انواع تهدیدات سایبری است و از عناصر ضروری امنیت محسوب می‌شود. یکی از مشکلات روش‌های سنتی تشخیص نفوذ، وابستگی آنها به الگوهای موجود و ناتوانی در شناسایی ناهنجاری‌ها است. شبکه‌های عصبی هم آمیخت^۱ می‌توانند الگوهای جدید و مبهم را شناسایی کرده و تعداد هشدارهای کاذب را به حداقل برسانند. این مقاله یک رویکرد جدید، بر اساس لایه هم آمیخت یک بعدی و لایه حذف تصادفی با روش تحلیل مؤلفه اصلی، برای حل مشکل و کمک به سیستم تشخیص نفوذ برای دستیابی به نرخ تشخیص بالاتر ارائه می‌کند. روش پیشنهادی شامل یک شبکه عمیق با شش لایه

هم آمیخت است که از تابع فعال سازی MISH بعد از هر لایه هم آمیخت استفاده می‌کند. تحلیل مولفه‌های اصلی برای کاهش ابعاد ویژگی‌های ورودی و استخراج ویژگی‌های ضروری استفاده شده است. ساختار پیشنهادی با استفاده از مجموعه داده CIC-IDS2017 تجزیه و تحلیل شده است. نتایج آزمایش‌ها عملکرد بهتری را از نظر زمان تشخیص و صحت، دقت، فراخوانی و کاپاکوهن در روش پیشنهادی نسبت به ساختارها و روش‌های مشابه نشان می‌دهد. زمان تشخیص روش پیشنهادی هشت میلی ثانیه بوده است. واژه‌های کلیدی: تشخیص نفوذ، الگوریتم تحلیل مولفه اصلی، یادگیری عمیق، الگوریتم MISH، لایه هم آمیخت.

۱. مقدمه

امروزه، مسائل امنیتی نقش اساسی در ارتباطات ایفا

* نویسنده مسئول

1- Convolutional neural network

مدل‌های یادگیری ماشین و یادگیری عمیق مانند فرکانس رادیویی^۲، ماشین بردار پشتیبانی، شبکه باور عمیق^۳ و حافظه کوتاه مدت طولانی^۴ ارزیابی کرده است. بالاترین مقدار دقت شبکه هم‌آمیخت پیشنهادی، با مقدار ۸۰/۱۳ درصد، در مجموعه داده NSL-KDD بود.

تعدادی از محققین روش‌هایی مبتنی بر شبکه‌های عصبی هم‌آمیخت و شبکه‌های حافظه کوتاه مدت طولانی برای تشخیص نفوذ در شبکه‌های کامپیوتری پیشنهاد کرده‌اند. خوزه [۱۱] با استفاده از مجموعه داده CIC-IDS 2017، مقدار دقت تشخیص را ۹۰/۶۱ درصد در شبکه هم‌آمیخت پیشنهادی و ۹۷/۶۷ درصد در شبکه حافظه کوتاه مدت طولانی گزارش کرد. به همین ترتیب، حمزه [۱۲] رویکردی را پیشنهاد کرد که به نتیجه دقت ۹۳ درصد در مجموعه داده CSE-CIC-IDS2018 دست یافت. این مدل می‌تواند داده‌های تصویر مخرب و ترافیک شبکه متوالی را شناسایی کند و میانگین دقت تشخیص نفوذ ۹۷/۶۳ درصد دارد. لین [۱۳] یک مدل تشخیص نفوذ پویا را با استفاده از حافظه کوتاه مدت طولانی، روش توجه‌ه و تکنیک نمونه‌زایی اقلیت ترکیبی^۶ ارائه کرد. معماری حافظه کوتاه مدت طولانی بر چالش‌های روزمره شبکه‌های عصبی بازگشتی^۷، مانند ناپدید شدن و انفجار گرادیان غلبه می‌کند. این مدل در مجموعه داده‌های CSE-CIC-IDS2018 به دقت ۹۶/۲ درصد دست یافت. به همین ترتیب، کیم [۱۴] به دقت ۹۸/۷۹ درصد در مجموعه داده KDD Cup 99 دست یافت.

یکی از مهمترین نگرانی‌ها در مورد سیستم‌های تشخیص نفوذ، پایداری آنها در برابر تغییرات سریع در شبکه‌ها و سخت‌افزار است. شون [۱۵] مدل جدیدی را بر اساس رمزگذار خودکار عمیق غیرمتقارن و طبقه‌بندی جنگل تصادفی پیشنهاد کرد. این روش با دو مجموعه

می‌کنند. چون جریان داده در این محیط به سرعت رخ می‌دهد [۱]. سیستم‌های تشخیص نفوذ به عنوان عناصر حیاتی در شناسایی و پیشگیری از تهدیدات سایبری مختلف معرفی شده‌اند و از عناصر ضروری امنیت محسوب می‌شوند [۲]. با پیشرفت‌های اخیر در یادگیری ماشین، به‌ویژه شبکه‌های عصبی هم‌آمیخت، امید جدیدی برای شناسایی دقیق‌تر و کارآمدتر حملات سایبری وجود دارد [۳].

شبکه‌های عصبی هم‌آمیخت به دلیل توانایی منحصر به فرد خود در استخراج و تجزیه و تحلیل ویژگی‌های داده، تحقیقات بینایی ماشین و سایر حوزه‌ها را متحول کرده‌اند. این شبکه‌ها به عنوان راه‌حل‌های پیشگام برای شناسایی الگوهای مخرب، زمینه جدیدی را در امنیت سایبری ایجاد می‌کنند [۴]. شبکه‌های عصبی هم‌آمیخت با قابلیت‌های خود برای تجزیه و تحلیل دقیق و سریع پدیده‌های پیچیده، به ابزاری مفیدی برای تشخیص نفوذ دقیق‌تر تبدیل شده‌اند. یکی از مشکلات روش‌های سنتی تشخیص نفوذ، وابستگی آنها به الگوهای موجود و ناتوانی در شناسایی ناهنجاری‌ها است. شبکه‌های عصبی هم‌آمیخت می‌توانند الگوهای جدید و مبهم را شناسایی کرده و تعداد هشدارهای کاذب را به حداقل برسانند. این شبکه‌ها امکان بهبود فعالیت‌های مشکوک را فراهم می‌کنند و تجزیه و تحلیل دقیق‌تری ارائه می‌دهند. این تحولات به معنای افزایش حفاظت از سازمان‌ها و افراد در برابر تهدیدات سایبری است [۵].

بسیاری از محققان ادعا کرده‌اند که شبکه‌های عصبی مصنوعی می‌توانند عملکرد سیستم تشخیص نفوذ را در مقایسه با روش‌های مرسوم افزایش دهند [۶-۸]. کیم [۹] یک مدل مبتنی بر شبکه عصبی عمیق ارائه کرد که با استفاده از مجموعه داده KDD CUP 99 به دقت ۹۹ درصد در این مجموعه داده دست یافت. دینگ و ژای [۱۰] یک مدل تشخیص نفوذ مبتنی بر شبکه هم‌آمیخت را پیشنهاد کردند. آنها بر روی چند طبقه‌بندی با مجموعه داده NSL-KDD تمرکز کردند. این تحقیق عملکرد مدل پیشنهادی را با سایر

2-Radio frequency

3-Deep belief network

4-Long short-term memory (LSTM)

5-Attention

6-Synthetic Minority Oversampling Technique (SMOTE)

7-RRN

۲. روش پیشنهادی

مدل پیشنهادی یک ساختار عمیق است که شامل لایه‌های هم‌آمیخت یک‌بعدی، نرمال‌سازی دسته‌ای^{۱۰}، و حذف تصادفی^{۱۱} می‌شود و از تابع فعال‌سازی Mish استفاده می‌کند. الگوریتم تحلیل مولفه اصلی ابتدا برای کاهش تعداد ویژگی‌ها روی داده‌های ورودی اعمال می‌شود. در زیر جزئیات هر لایه آورده شده است.

۲-۱- تحلیل مولفه اصلی^{۱۲}

تحلیل مولفه اصلی یک تکنیک آماری است که معمولاً در یادگیری ماشینی و تجسم داده‌ها استفاده می‌شود. ابعاد مجموعه داده را کاهش می‌دهد و قابلیت تفسیر را افزایش می‌دهد و در عین حال از دست دادن اطلاعات را به حداقل می‌رساند. تحلیل مولفه اصلی شامل استانداردسازی داده‌ها، محاسبه ماتریس کوواریانس، و شناسایی اجزای اصلی است که بیشترین واریانس را به خود اختصاص می‌دهند. مقادیر ویژه و بردارهای ویژه مرتب‌سازی می‌شوند و ماتریس اصلی بر اساس این مقادیر تبدیل می‌شود تا مؤلفه‌های اصلی ناهمبسته به دست آید. این مؤلفه‌ها بیشترین واریانس را در داده‌ها به خود اختصاص می‌دهند و هر مؤلفه بعدی واریانس کمتری را به خود اختصاص می‌دهد.

۲-۲- لایه هم‌آمیخت یک بعدی

یادگیری ماشین و یادگیری عمیق به طور گسترده از لایه‌های Conv1D یا هم‌آمیخت یک بعدی برای پردازش داده‌های زمانی، پردازش زبان طبیعی و تجزیه و تحلیل سری‌های زمانی استفاده می‌کنند. برخلاف Conv2D که برای پردازش تصویر استفاده می‌شود، Conv1D فقط هم‌آمیختی را در یک بعد انجام می‌دهد و از هسته‌ای با وزن‌های قابل یادگیری استفاده می‌کند که روی داده‌های ورودی می‌لغزند، در هر مرحله یک ضرب داخلی انجام

داده مختلف KDDCUP'99 و NSL-KDD مورد ارزیابی قرار گرفت و دقت آن به ترتیب ۹۷/۸۵ درصد و ۸۵/۴۲ درصد بود. به همین ترتیب، هو [۱۶] از مدل رمزگذار خودکار متغیر شرطی log-cosh^۸ و شبکه عصبی کوتاه‌مدت دوطرفه و یک‌طرفه برای سیستم تشخیص نفوذ استفاده کرد. این روش می‌تواند نمونه‌های مجازی با برچسب‌های خاص ایجاد کند و ویژگی‌های حمله مهمتری را از داده‌های ترافیک شبکه استخراج کند. این روش نرخ تشخیص برخی از رده‌های حمله را افزایش می‌دهد و عملکرد NID بهبود یافته است. نتایج روی مجموعه داده NSL-KDD شامل دقت، امتیاز F1، خوانایی و نرخ خطای نادرست به ترتیب ۸۷/۳۰ درصد، ۸۷/۸۹، ۸۰/۸۹ درصد و ۴/۳۶ درصد بود. بین [۱۷] یک شبکه عصبی مکرر با مجموعه داده NSL-KDD پیشنهاد کرد و طبقه‌بندی‌های دودویی و چندگانه را انجام داد. علاوه بر این، روی [۱۸] روشی مبتنی بر شبکه‌های عصبی بازگشتی حافظه کوتاه‌مدت دوطرفه (BLSTM-RNN) برای تشخیص نفوذ دودویی پیشنهاد کرد. این مدل به دقت ۹۵ درصد دست‌یافته است. کو و همکاران [۱۹] مدلی بر اساس شبکه باور عمیق^۹ ارائه کرد. آنها دقت ۹۵/۲۵ درصد را در مجموعه داده NSL-KDD گزارش کردند.

در یک سیستم تشخیص نفوذ، دقت تشخیص و سرعت تشخیص دو پارامتر مهم هستند که برتری روش‌ها را تعیین می‌کنند. این مقاله یک شبکه هم‌آمیخت و یک الگوریتم تحلیل مولفه اصلی را ترکیب می‌کند. این روش زمان تشخیص را کاهش می‌دهد و عملکرد بهتری نسبت به روش‌های مشابه در سایر معیارها دارد. بقیه مقاله به شرح زیر است. ساختار روش پیشنهاد در بخش ۲ توضیح داده شده است. بخش سه نتایج ارزیابی روش پیشنهادی و دیگر روش‌های مشابه براساس معیارهای توضیح داده شده و بحث در مورد آنها را بیان می‌کند. در آخر بخش چهار نتیجه‌گیری مقاله را ارائه می‌کند.

10-Batch Normalization

11-Dropout

12-Principal Component Analysis (PCA)

8-Log-cosh conditional variational autoencoder (LCVAE)

9-DBN

است. این می‌تواند به کاهش مشکلاتی مانند ناپدید شدن گرادین‌ها^{۱۷} و نورون‌های مرده^{۱۸} کمک کند. شکل غیر یکنواخت آن می‌تواند ظرفیت مدل را بدون افزایش تعداد پارامترها یا عمق مدل افزایش دهد.

در حالی که تابع فعال‌ساز Mish با موفقیت در کارهای مختلف استفاده شده است، شایان ذکر است که استفاده از یک تابع نمایی در محاسبه آن، سطح بالاتری از پیچیدگی محاسباتی را در مقایسه با سایر توابع مانند ReLU معرفی می‌کند. این به طور بالقوه می‌تواند بر کارایی مدل تأثیر بگذارد.

۲-۴- حذف تصادفی^{۱۹}

شبکه‌های عصبی معمولاً از یک لایه حذف تصادفی برای جلوگیری از بیش‌برازش^{۲۰} استفاده می‌کنند. در این روش، نورون‌ها به طور تصادفی در حین آموزش حذف می‌شوند، بنابراین مدل داده‌های آموزشی را به خوبی یاد نمی‌گیرد و روی داده‌های دیده نشده عملکرد ضعیفی دارد. سهم آنها در نورون‌های پایین‌دست به طور موقت حذف می‌شود و به‌روزرسانی وزن اعمال نمی‌شود. انتخاب تصادفی نورون‌های حذف شده تضمین می‌کند که مدل به شدت به یک نورون وابسته نیست و آن را قوی‌تر می‌کند. در طول فرآیند آزمایش، هیچ نورونی حذف نمی‌شود، اما نرخ انصراف، مقادیر خروجی لایه را کاهش می‌دهد تا افزایش فعالیت در طول آموزش را در نظر بگیرد.

۲-۵- ساختار پیشنهادی

شکل ۱ ساختار پیشنهادی را نشان می‌دهد. این شبکه از شش لایه هم‌آمیخت یک بعدی تشکیل شده است. در هر لایه، تعداد هسته‌ها ۲۵۶، ۱۲۸، ۶۴، ۳۲، ۱۶ و ۸ است. پس از اعمال هم‌آمیختی، هر لایه از تابع فعال‌سازی MISH و سپس یک لایه حذف تصادفی با $p=0.1$ استفاده می‌کند.

می‌دهد و نتایج را برای تولید خروجی جمع می‌کند. پارامترهای ضروری Conv1D شامل تعداد فیلترها، اندازه هسته، گام‌ها^{۲۱} و لایه‌گذاری^{۲۲} است. لایه Conv1D برای تشخیص گفتار، تشخیص ناهنجاری در سری‌های زمانی و تجزیه و تحلیل توالی با شناسایی الگوهای ثابت برای ترجمه‌های آنی مناسب است.

در نظر بگیرید h_i^k نقشه ویژگی k ام در لایه h و $x_0 = x$ لایه ورودی است. فرض کنید w_i^k بردار وزن k امین فیلتر هم‌آمیختی و b_i بردار بایاس لایه i است. لایه هم‌آمیخت یک بعدی بسته داده ترافیک شبکه را به عنوان یک بردار ورودی $x = (x_1, x_2, x_3, \dots, x_k)$ دریافت می‌کند که x_k یک ویژگی است. f را به عنوان یک تابع فعال کننده در نظر بگیرید. در هر یک از لایه‌های هم‌آمیخت، نقشه ویژگی ورودی h_i^k با فیلترها یک نقشه ویژگی جدید تولید می‌کند

$$h_i^k = b_i^k + \sum_{k=1}^K w_i^{k-1} y_i^{k-1} \quad (۱)$$

$$y_i^k = f(h_i^k) \quad (۲)$$

۲-۳. تابع فعال‌ساز Mish

تابع فعال‌ساز Mish یک تابع فعال‌ساز نسبتاً جدید است که در یادگیری عمیق استفاده می‌شود. این تابع توسط دیگان‌تا میسرا^{۱۹} در سال ۲۰۱۹ پیشنهاد شد [۲۰]. تابع Mish یک تابع نرم^{۲۱} و قابل تمایز است که به صورت زیر تعریف می‌شود:

$$Mish(x) = x * \tanh(\text{softplus}(x)) \quad (۳)$$

که در آن $\text{softplus}(x)$ یک تقریب نرم برای تابع است و به صورت زیر تعریف می‌شود:

$$\text{softplus}(x) = \log(1 + \exp(x)) \quad (۴)$$

در برخی از مدل‌های یادگیری عمیق، تابع فعال‌ساز Mish نسبت به سایر توابع فعال‌ساز بهتر عمل کرده

17-Gradient disappearance

18-Dead neurons

19-Dropout

20-Overfitting

13-Stride

14-Padding

15-Diganta Misra

16-Smooth

آسانی قابل تشخیص است.

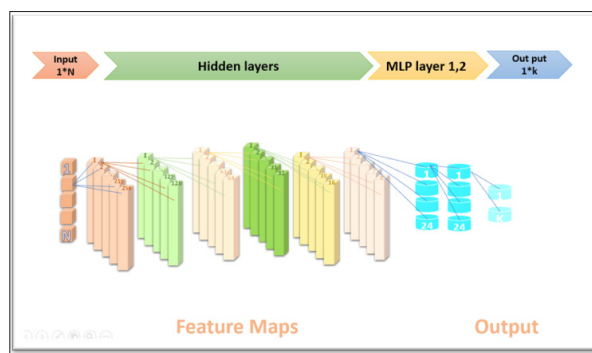
مجموعه داده CIC-IDS2017 یک مجموعه داده جامع برای سیستم‌های تشخیص نفوذ است که توسط موسسه امنیت سایبری کانادا توسعه یافته است [۲۲]. این مجموعه داده معیاری برای انواع مختلف حملات سایبری برای تسهیل و تشویق توسعه تکنیک‌ها و روش‌های جدید برای امنیت سایبری فراهم می‌کند.

این مجموعه داده شامل هشت فایل مختلف است که حاوی پنج روز ترافیک عادی و داده‌های حمله از موسسه امنیت سایبری کانادا است. جدول ۱ شرح مختصری از تمام آن فایل‌ها را نشان می‌دهد.

جدول (۱): فایل‌های مجموعه داده CICIDS2017

نام فایل	فعالیت روز	حملات پیدا شده
Monday-WorkingHours.pcap_ISCX.csv	دوشنبه	Benign (Normal human activities)
Tuesday-WorkingHours.pcap_ISCX.csv	سه شنبه	Benign, FTP-Patator, SSH-Patator
WednesdayworkingHours.pcap_ISCX.csv	چهارشنبه	Benign, DoS GoldenEye, DoS Hulk, DoS Slowhttptest, DoS slowloris, Heartbleed
Thursday-WorkingHours-Morning-WebAttacks.pcap_ISCX.csv	پنج شنبه	Benign, Web Attack – Brute Force, Web Attack – Sql Injection, Web Attack – XSS
Thursday-WorkingHours-Afternoon-Infiltration.pcap_ISCX.csv	پنج شنبه	Benign, Infiltration
Friday-WorkingHours-Morning.pcap_ISCX.csv	جمعه	Benign, Bot
Friday-WorkingHours-AfternoonPortScan.pcap_ISCX.csv	جمعه	Benign, PortScan
Friday-WorkingHours-AfternoonDDos.pcap_ISCX.csv	جمعه	Benign, DDoS

جدول ۱ نشان می‌دهد که داده‌های پنجشنبه بعدازظهر



شکل (۱): ساختار پیشنهادی

علاوه بر این، ساختار از دو لایه کاملاً متصل^{۲۱} به ابعاد ۲۴ و در نهایت از لایه فعال‌کننده Softmax استفاده می‌کند

۳. نتایج و بحث

۳-۱- محیط ارزیابی

تمامی آزمایش‌ها در پایتون با استفاده از کتابخانه PyTorch روی پردازنده core i3-9100f با فرکانس ۳/۶ گیگاهرتز با ۱۲۸ گیگابایت RAM و پردازنده گرافیکی Nvidia GeForce 2080 Turbo انجام شد. برای اندازه‌گیری زمان تشخیص در همه آزمایش‌ها شرایط یکسانی حفظ شد و همه برنامه‌ها به جز VSCode غیرفعال بودند.

۳-۲- مجموعه داده

برای ارزیابی روش پیشنهادی از دو مجموعه داده استفاده شد. علاوه بر این از مجموعه داده برای ارزیابی بیشتر روش پیشنهادی استفاده شد. این مجموعه داده KDD 99 CUP که با هدف آزمون الگوریتم‌های تشخیص نفوذ گردآوری و طراحی شده است و با استفاده از حجم عظیم داده‌های گردآوری شده در پروژه DIDE^{۲۲} که با همکاری سازمان پروژه‌های تحقیقاتی پیشرفته دفاعی، وزارت دفاع ایالات متحده آمریکا و آزمایشگاه لینکلن دانشگاه MIT انجام شد، تهیه گردیده است. کلیه رکوردهای موجود در این مجموعه داده، توسط افراد خبره در حوزه امنیت اطلاعات برچسب‌گذاری شده است به گونه‌ای که تعلق هر رکورد به رده خاصی از حمله و یا عادی بودن رکورد به

21-Full-connected
22-DARPA Intrusion Detection Evaluation

۱۹۶۶	Bot
۱۵۰۷	Web Attack – Brute Force
۶۵۲	Web Attack – XSS
۳۶	Infiltration
۲۱	Web Attack – Sql Injection
۱۱	Heartbleed

۳-۳- معیارهای ارزیابی

معیارهای اندازه‌گیری زمان، صحت^{۲۵}، دقت^{۲۶}، فراخوانی^{۲۷} و کاپا کوهن^{۲۸} است که برای ارزیابی توانایی ساختار پیشنهادی در شناسایی نوع نفوذ استفاده شد.

- صحت

صحت تشخیص، نسبت رده‌های پیش‌بینی شده درست به کل داده‌ها است. معادله ۱ به صورت ریاضی بیانگر صحت است.

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (5)$$

- فراخوانی

زمانی که مقدار منفی نادرست زیاد باشد، معیار فراخوانی مناسب خواهد بود. این معیار به صورت زیر تعریف می‌شود:

$$Recall = \frac{TP}{TP+FN} \quad (6)$$

- دقت

دقت معیاری برای نشان دادن صحت نتیجه مثبت است. حداکثر مقدار آن ۱ است. دقت بیشتر به معنای تعداد مثبت نادرست کمتر است. معادله ۳ این معیار را محاسبه می‌کند

$$Precision = \frac{TP}{TP+FP} \quad (7)$$

- کاپا کوهن

ضریب کاپا کوهن و تجزیه و تحلیل آماری بر اساس

و جمعه برای طبقه‌بندی دودویی مناسب هستند، در حالی که داده‌های صبح سه‌شنبه، چهارشنبه و پنج‌شنبه برای طراحی یک مدل تشخیص چند رده‌ای بهترین هستند. با این حال، ذکر این نکته ضروری است که یک مدل تشخیص ایده‌آل باید بتواند هر حمله‌ای را تشخیص دهد. بنابراین، همه داده‌های ترافیک روز در یک مجموعه داده ادغام شدند که منجر به ایجاد مجموعه داده‌ای با ۳۱۱۹۳۴۵ نمونه و ۸۳ ویژگی، از جمله ۱۵ برچسب رده (۱ برچسب عادی + ۱۴ برچسب حمله) شد. پس از حذف نمونه‌هایی با برچسب‌های رده از دست‌رفته ۲۳ و داده‌های از دست‌رفته ۲۴، مجموعه داده ترکیبی CICIDS2017 شامل ۲۸۳۰۵۴۰ نمونه بود. جداول ۲ و ۳ ویژگی‌های مجموعه داده ترکیبی و وقوع دقیق هر رده را نشان می‌دهد.

جدول (۲): ویژگی‌های کلی مجموعه داده CICIDS2017

CICIDS2018	نام مجموعه داده
چند کلاسه	نوع مجموعه داده
۲۰۱۷	سال تولید
۲۸۳۰۵۴۰	مجموع تعداد نمونه‌های مجزا
۸۳	تعداد ویژگی‌ها
۱۵	تعداد کلاس‌های مجزا

جدول (۳): تعداد نمونه رده‌ای مجموعه داده CICIDS2017

تعداد نمونه‌ها	برچسب‌های کلاس
۲۳۵۹۰۸۷	BENIGN
۲۳۱۰۷۲	DoS Hulk
۱۵۸۹۳۰	PortScan
۴۱۸۳۵	DDoS
۱۰۲۹۳	DoS GoldenEye
۷۹۳۸	FTP-Patator
۵۸۹۷	SSH-Patator
۵۷۹۶	DoS slowloris
۵۴۹۹	DoS Slowhttptest

23-Missing class lab

24-Missing data

25-Accuracy

26-Precision

27-Recall

28-Cohen's kappa

جدول (۶): مقایسه روش پیشنهادی با روش‌های مشابه با مجموعه داده KDD CUP 99

نام روش	صحت
[۹] DNN	۹۹
[۱۴] LSTM	۹۸/۷۹
روش پیشنهادی	۹۹/۳۲

۳-۵- بحث

جدول ۴ شامل یک نمای کلی از عملکرد مدل‌های مختلف در تشخیص الگو است. قابل ذکر است که مدل پیشنهادی با صحت ۹۸/۹۹ درصد، فراخوانی ۹۹/۸۵ درصد، دقت ۹۹/۹۶ درصد، و زمان تشخیص ۸ میلی ثانیه، برجسته است. این نتایج اعتماد به قابلیت‌های مدل را القا می‌کند. بر اساس اطلاعات ارائه شده در جدول ۵، مدل پیشنهادی دارای دقت بیش از ۹۹/۹۶ درصد و زمان تشخیص ۸ میلی ثانیه، مدل LSTM دارای دقت ۹۷/۶۷ درصد و زمان تشخیص ۲۵ میلی ثانیه و مدل CNN دارای یک دقت ۹۰/۶۱ درصد و زمان تشخیص ۵ میلی ثانیه می‌باشند. با توجه به این اطلاعات، مدل پیشنهادی از دقت بسیار بالا و زمان تشخیص نسبتاً کمتری نسبت به مدل‌های دیگر برخوردار است. این نشان‌دهنده بهبود عملکرد مدل پیشنهادی در مقایسه با مدل‌های CNN و LSTM است.

۴. نتیجه‌گیری

با توجه به رشد روزافزون فناوری اطلاعات، امنیت به یکی از مهمترین موضوعات تبدیل شده است. سیستم‌های تشخیص نفوذ نقش حیاتی در شبکه‌ها دارند. با این حال، سیستم‌های تشخیص نفوذ قدیمی و باروها^{۲۹} از امنیت سیستم‌های اطلاعاتی در برابر حملات جدید محافظت نمی‌کنند. این مطالعه یک الگوریتم تشخیص نفوذ جدید بر اساس لایه هم‌آمیخت یک بعدی با تابع فعال‌سازی MISH، لایه حذف تصادفی و روش کاهش تحلیل مولفه اصلی پیشنهاد می‌کند. مدل پیشنهادی از تابع فعال‌سازی Mish برای کاهش مشکلاتی مانند ناپدید شدن گرادیان‌ها،

آن، معیارهای عددی بین ۱- و ۱+ است و هر چه به ۱ نزدیک‌تر باشد، نشان‌دهنده وجود توافق متناسب و مستقیم است. اندازه‌های نزدیک به ۱- نشان‌دهنده تطابق معکوس و اندازه‌های نزدیک به صفر نشان‌دهنده عدم توافق است. این ضریب به صورت زیر محاسبه می‌شود:

$$Kappa = \frac{Pr(a) - Pr(e)}{1 - Pr(e)} \quad (۸)$$

به طوری که $Pr(a)$ توافق نسبی مشاهده شده بین ارزیاب‌ها است و $Pr(e)$ احتمال توافق احتمالی فرضی است.

۳-۴- نتایج

جدول ۴ نشان می‌دهد که روش پیشنهادی از نظر معیارهای طبقه‌بندی مشابه با شبکه Inception v2 رفتار می‌کند. با این حال، زمان تشخیص روش پیشنهادی تقریباً یک پنجم روش Inception v2 است.

جدول (۴): مقایسه ساختار پیشنهادی با ساختارهای عمیق رایج

نام روش	زمان تشخیص (میلی ثانیه)	کاپا کوهن	دقت	فراخوانی	صحت
Vgg19-BN	۵۶	۰/۹۹۵	۹۹/۳۵	۹۸/۵۸	۹۹/۹۵
Resnet50	۱۸	۰/۹۱۵	۹۷/۵۲	۹۶/۶۹	۹۷/۲۵
Inception_v2	۳۵	۰/۹۹۹	۹۹/۰۱	۹۹/۹۱	۹۹/۹۸
DarkNet	۲۱	۰/۸۹۵	۹۵/۶۸	۹۴/۶۵	۹۵/۵۶
روش پیشنهادی	۸	۰/۹۹۸	۹۸/۹۹	۹۹/۸۵	۹۹/۹۶

جدول ۵ روش پیشنهادی را با روش‌های مشابه مقایسه می‌کند.

جدول (۵): مقایسه روش پیشنهادی با روش‌های مشابه

نام روش	صحت	زمان تشخیص (میلی ثانیه)
[۲۳] CNN	۹۰/۶۱	۵
[۲۳] LSTM	۹۷/۶۷	۲۵
روش پیشنهادی	۹۹/۹۶	۸

جدول ۶ روش پیشنهادی را با روش‌های مشابه با مجموعه داده KDD CUP 99 مقایسه می‌کند.

- Intelligence and Applications, p. 2342003.
9. Kim, J., Shin, N., Jo, S. Y. & Kim, S. H. 2017. "Method of intrusion detection using deep neural network," in 2017 IEEE international conference on big data and smart computing (BigComp), IEEE, pp. 313-316.
 10. Ding Y. & Zhai, Y. 2018. "Intrusion detection system for NSL-KDD dataset using convolutional neural networks," in Proceedings of the 2018 2nd International conference on computer science and artificial intelligence, pp. 81-85.
 11. Jose J., & Jose, D. V. 2023. "Deep learning algorithms for intrusion detection systems in internet of things using CIC-IDS 2017 dataset," International Journal of Electrical and Computer Engineering (IJECE), vol. 13, no. 1, pp. 1134-1141.
 12. Erlambang, M. R., Hamzah, I. W. & Dewanta, F. 2023. "Machine Learning Approach for Intrusion Detection System to Mitigate Distributed Denial of Service Attack Based on Convolutional Neural Network Algorithm," eProceedings of Engineering, vol. 9, no. 6.
 13. Lin, P., Ye, K. & Xu, C.-Z. 2019. "Dynamic network anomaly detection system by using deep learning techniques," in Cloud Computing-CLOUD 2019: 12th International Conference, Held as Part of the Services Conference Federation, SCF 2019, San Diego, CA, USA, June 25-30, 2019, Proceedings 12, Springer, pp. 161-176.
 14. Kim, J., Kim, J., Thu, H. L. T. & Kim, H. 2016. "Long short term memory recurrent neural network classifier for intrusion detection," in 2016 international conference on platform technology and service (PlatCon), IEEE, pp. 1-5.
 15. Shone, N., Ngoc, T. N., Phai, V. D. & Shi, Q. 2018. "A deep learning approach to network intrusion detection," IEEE transactions on emerging topics in computational intelligence, vol. 2, no. 1, pp. 41-50.
 16. Hou, T., Xing, H., Liang, X., Su, X. & Wang, Z. 2023. "A Marine Hydrographic Station Networks Intrusion Detection Method Based on LCVAE and CNN-BiLSTM," Journal of Marine Science and Engineering, vol. 11, no. 1, p. 221.
 17. Yin, C., Zhu, Y., Fei, J. & He, X. 2017. A deep learning approach for intrusion detection using recurrent neural networks," Ieee Access, vol. 5, pp. 21954-21961.
 18. Roy, B. & Cheung, H. 2018. "A deep learning approach for intrusion detection in internet of things using bi-directional long short-term memory recurrent neural network," in 2018 28th international telecommunication networks and applications conference (ITNAC). IEEE, pp. 1-6.
 19. F. Qu, J. Zhang, Z. Shao, and S. Qi, "An intrusion detection model based on deep belief network," in Proceedings of the 2017 VI international conference on network, communication and computing, 2017, pp. 97-101.
 20. Misra, D. 2019. "Mish: A self regularized non-monotonic activation function," arXiv preprint arXiv:1908.08681.
 21. Kumar, S., et al. 2019. "A statistical analysis on KDD Cup'99 dataset for the network intrusion detection sys-

نورون‌های مرده و حذف‌های تصادفی استفاده می‌کند تا از بیش‌برازش جلوگیری شود. الگوریتم تحلیل مولفه اصلی ابتدا روی داده‌های ورودی اعمال می‌شود تا تعداد ویژگی‌های ورودی کاهش یابد. برای ارزیابی از مجموعه داده CIC-IDS2017 استفاده شد. نتایج ارزیابی صحت اعتبار و زمان تشخیص همگرایی خوبی را نشان می‌دهد. علاوه بر این نشان می‌دهد که دقت الگوریتم پیشنهادی مشابه ساختارهای عمیق قوی vgg16 و Inception-v3 است. با این حال، این مدل در زمان تشخیص بهترین است. زمان تشخیص مدل پیشنهادی ۸ میلی‌ثانیه بود که بهترین زمان در بین مدل‌های مشابه می‌باشد.

مراجع

1. Wenhua, Z., Qamar, F., Abdali, T.-A. N., Hassan, R., Jafri, S. T. A. & Nguyen, Q. N. 2023. "Blockchain technology: security issues, healthcare applications, challenges and future trends," Electronics, vol. 12, no. 3, p. 546.
2. Kumar, S., Selvi, M. & Kannan, A. 2023. "A comprehensive survey on machine learning-based intrusion detection systems for secure communication in internet of things," Computational Intelligence and Neuroscience: CIN, vol. 2023.
3. Abdulganiyu, O. H., Tchakoucht, T. A. & Saheed, Y. K. 2024. "Towards an efficient model for network intrusion detection system (IDS): systematic literature review," Wireless Networks, vol. 30, no. 1, pp. 453-482.
4. Binbusayyis, A. 2024. "Hybrid VGG19 and 2D-CNN for intrusion detection in the FOG-cloud environment," Expert Systems with Applications, vol. 238, p. 121758.
5. Saied, M., Guirguis, S. & Madbouly, M. 2024. "Review of artificial intelligence for enhancing intrusion detection in the internet of things," Engineering Applications of Artificial Intelligence, vol. 127, p. 107231.
6. Srivastava, A., Addimulam, S. C., Basu, M. T., Sindhuri, B. P. & Maurya, R. K. 2024. "Network Intrusion Detection System (NIDS) for WSN using Particle Swarm Optimization based Artificial Neural Network," International Journal of Intelligent Systems and Applications in Engineering, vol. 12, no. 15s, pp. 143-150.
7. Najafi Mohsenabad H. & Tut, M. A. 2024. "Optimizing Cybersecurity Attack Detection in Computer Networks: A Comparative Analysis of Bio-Inspired Optimization Algorithms Using the CSE-CIC-IDS 2018 Dataset," Applied Sciences, vol. 14, no. 3, p. 1044, 2024.
8. Jiao, M. 2024. "Application of Random Forest Algorithm in Network Intrusion Detection of Government Affairs Departments," International Journal of Computational

- tem.” Applied Soft Computing and Communication Networks: Proceedings of CAN, pp. 131-157.
22. Khan, Z. I., Afzal, M. M. & Shamsi, K. N. 2024. “A Comprehensive Study on CIC-IDS2017 Dataset for Intrusion Detection Systems,” International Research Journal on Advanced Engineering Hub (IRJAEH), vol. 2, no. 02, pp. 254-260.
23. Jose, J. & Jose, DV. 2023. Deep learning algorithms for intrusion detection systems in internet of things using CIC-IDS 2017 dataset. International Journal of Electrical and Computer Engineering (IJECE). Vol. 13, no. 01, pp.1134-41.