

DEKF-Driven Few-Shot Class Imbalance Learning to Enhance Cyber Attack Detection

F. Saleh^{ID}, K. Dadashtabar Ahmadi*, M.A. Keyvanrad,^{ID}

*Assistant Professor, Imam Hossein University (AS), Tehran, Iran

Received:2024 /05/23, Revised: 2024/09/06, Accepted: 2024/10/07, Published: 2024/10/22

DOR: <https://dor.isc.ac/dor/20.1001.1.23224347.1403.12.3.10.2>

ABSTRACT

The escalating use of networks and the internet has led to a surge in cyber threats, making it imperative to develop sophisticated intrusion detection systems (IDS) capable of safeguarding against these malicious intrusions. While machine learning techniques have been extensively employed to enhance IDS, challenges persist, notably in handling imbalanced datasets and rare attack detection such as R2L and U2R due to the small number of their samples in the training dataset. Imbalanced datasets, a common challenge in IDS evaluation, often skew toward majority classes, hindering the detection of minority class attacks. Existing machine learning classifiers, primarily accuracy-driven, struggle to excel at identifying rare attacks, which are often more catastrophic. Moreover, overlapping classes complicate feature selection, further impeding accurate detection. To tackle these challenges, this article proposes a solution rooted in Few-Shot Learning, particularly MAML. Traditional MAML has limitations, including slow convergence and computational demands. To enhance MAML's performance, the article introduces the Node Decoupled Extended Kalman Filter (NDEKF) as an alternative to gradient descent in the inner loop. NDEKF optimizes MAML training, offering faster convergence and improved generalization. The DEKF (Decoupled Extended Kalman Filter) variant simplifies calculations, making it suitable for deep neural networks. The combination of MAML and NDEKF, termed NDEKF-based MAML, is applied to address the imbalanced data problem in IDS. The proposed approach is evaluated on the NSL-KDD dataset, demonstrating its potential to improve rare attack detection in intrusion detection systems. By adopting this approach, we achieved improved convergence speed, enhanced ability to generalize, and higher accuracy compared to the original MAML algorithm when dealing with a sparse and unstable dataset such as NSL-KDD. Particularly, our framework demonstrated significant advancements in accurately detecting rare U2R and R2L attacks. The accuracy rates for R2L and U2R attacks using our proposed framework surpassed those of the original MAML, increasing from 61% to 75% and from 51% to 66%, respectively, even with a reduced number of training epochs.

Keywords: Intrusion detection system, Class Imbalance Learning, Few Shot Learning MAML, Decoupled Extended Kalman Filter

* Corresponding Author Email: dadashtabar@mut.ac.ir

آموزش در داده‌های نامتوازن با استفاده از روش آموزش با داده محدود مبتنی بر فیلتر DEKF

برای بهبود تشخیص نفوذ حملات سایبری

فتاه صالح^۱، کوروش داداش تبار احمدی^{۲*} محمدعلی کیوان‌راد^۳

۱- دانشجوی دکتری ۲- استادیار، ۳- استادیار، دانشگاه صنعتی مالک‌اشتر، تهران، ایران

(دریافت: ۱۴۰۳/۰۳/۰۳، بازنگری: ۱۴۰۳/۰۶/۱۶، پذیرش: ۱۴۰۳/۰۷/۱۶، انتشار: ۱۴۰۳/۰۷/۰۱)

DOR: <https://dor.isc.ac/dor/20.1001.1.23224347.1403.12.3.10.2>

* این مقاله یک مقاله با دسترسی آزاد است که تحت شرایط و ضوابط مجوز Creative Commons Attribution (CC BY) توزیع شده است.



نویسندگان ©

ناشر: دانشگاه جامع امام حسین (ع)

چکیده

در حالی که تکنیک‌های یادگیری ماشین به‌منظور بهبود سیستم تشخیص نفوذ حملات سایبری به شکلی گسترده مورد استفاده قرار می‌گیرند، چالش‌های متعدد، به‌ویژه در زمینه مدیریت مجموعه داده‌های نامتوازن و تشخیص حملات نادر مانند $R2L^1$ و $U2R^2$ به دلیل ناکافی بودن تعداد نمونه‌های آنها در مجموعه داده‌های آموزشی، به قوت خود باقی هستند. مجموعه داده‌های نامتوازن، چالشی رایج و همیشگی در هنگام ارزیابی عملکرد سیستم تشخیص نفوذ بوده و اغلب منجر به انحراف به سمت کلاس اکثریت می‌شوند؛ این عامل به نوبه خود، مانع از شناسایی حملات کلاس‌های اقلیت خواهد شد. طبقه‌بندی‌کننده‌های نیز یادگیری ماشین که عمدتاً مبتنی بر دقت هستند، قادر به شناسایی حملات سایبری نادر نخواهند بود. به‌علاوه، همپوشانی کلاس‌ها انتخاب ویژگی را پیچیده‌تر کرده و مانع تشخیص دقیق نفوذ خواهد شود. ما در این مقاله، برای مقابله با این چالش‌ها، راهکاری ارائه می‌کنیم که منشأ آن در آموزش با داده محدود، به‌ویژه یادگیری متامدل/آگنوستیک $MAML^3$ نهفته است. الگوریتم $MAML$ سنتی محدودیت‌هایی دارد که از آن جمله می‌توان به همگرایی کند و الزامات محاسباتی اشاره کرد. به‌منظور ارتقای عملکرد $MAML$ ، این مقاله فیلتر کالمن گسترش‌یافته جداسده از گره $NDEKF^4$ را به‌عنوان جایگزینی برای گرادینان نزولی در حلقه داخلی معرفی می‌کند. الگوریتم $NDEKF$ باعث بهینه‌سازی آموزش $MAML$ ، تسریع همگرایی و بهبود تعمیم خواهد شد. فیلتر فوق محاسبات را ساده‌تر و آن را برای شبکه‌های عصبی عمیق بهینه‌سازی می‌کند. برای حل معضل داده‌های نامتوازن در سیستم تشخیص نفوذ از ترکیب دو الگوریتم $MAML$ و $NDEKF$ تحت عنوان $MAML-NDEKF$ استفاده شده است. این رویکرد پیشنهادی، بر روی مجموعه داده $NSL-KDD$ ارزیابی می‌شود. بعد از اعمال این رویکرد، همگرایی به‌سرعت بهبود می‌یابد، قابلیت تعمیم افزایش پیدا می‌کند و در مقایسه با الگوریتم اصلی $MAML$ ، هنگام رویارویی با مجموعه داده پراکنده و ناپایداری مانند $NSL-KDD$ دقت بالاتری حاصل می‌گردد. چارچوب پیشنهادی ما به طور خاص، پیشرفت‌های قابل توجهی در زمینه تشخیص دقیق حملات $R2L$ و $U2R$ نشان داده است. نرخ دقت تشخیص حملات $R2L$ و $U2R$ در این رویکرد علی‌رغم کاهش تعداد دوره‌های آموزش، بیشتر از $MAML$ اصلی بوده و به ترتیب از ۶۱٪ به ۷۵٪ و از ۵۱٪ به ۶۶٪ افزایش یافته است.

کلیدواژه‌ها: سیستم تشخیص نفوذ حملات سایبری، یادگیری نامتوازن کلاس، آموزش با داده محدود $MAML$ ، فیلتر کالمن $NDEKF$

* رایانامه نویسنده مسئول: dadashtabar@mut.ac.ir

¹ Root to Local² User to Root³ Model-Agnostic Meta-Learning⁴ Node Decoupled Extended Kalman Filter

۱. مقدمه

طی چند سال گذشته، استفاده سازمان‌ها و افراد از شبکه‌ها و اینترنت به‌منظور برقرار ارتباط و انتقال داده‌ها افزایش چشمگیری داشته است، به‌ویژه از وقتی شیوه‌های کاری با افزایش دورکاری کارمندان دستخوش تغییرات بسیار شد. اما متأسفانه هم‌زمان با این تحول، ریسک حملات سایبری و نفوذهای مخرب نیز افزایش یافته است.

گزارش‌های اخیر مرتبط با جرایم سایبری نشان می‌دهد که این نوع حملات در سطح جهان در چند سال گذشته به‌شدت افزایش داشته است. به همین سبب، سازمان‌ها مجبور می‌شوند برای بازیابی اطلاعات، رفع آسیبی که به نام تجاری آنها وارد می‌شود، هزینه‌های عملیاتی، بیمه و غیره میلیون‌ها دلار هزینه کنند.

برای جلوگیری از بروز چنین خطراتی، لازم است سیستم‌های دفاعی پیشرفته و هوشمندی باهدف پیشگیری، شناسایی و پاسخ به این حملات ابداع شوند. در اینجا است که نقش مهم سیستم‌های تشخیص نفوذ (IDS)^۱ مطرح می‌شود، چون این سیستم‌ها برای یافتن نشانه‌هایی از فعالیت‌های مخرب موجود، کل ترافیک شبکه را پایش می‌کنند.

چندین دهه است که فن‌ها و روش‌های یادگیری ماشین متعددی برای توسعه این نوع سیستم‌ها مورداستفاده قرار می‌گیرند [۱، ۲]. اما به دلیل وقوع مکرر حملات مخرب به سیستم‌ها و شبکه‌های اطلاعاتی، این موضوع همچنان در دست بررسی و تحقیق است. گرچه اکثر این فن‌ها در زمینه بهبود دقت کلی سیستم‌های تشخیص نفوذ (IDS) موفقیت قابل‌ملاحظه‌ای کسب کرده‌اند، در برخی از حملات نادر یا جدید، میزان دقت آنها چندان رضایت‌بخش نیست و علت آن این است که در اکثر الگوریتم‌های یادگیری ماشین، مجموعه داده‌های توزیع متوازن یا هزینه‌های خطای برابر در نظر گرفته می‌شود [۳]. بنابراین، سوگیری^۲ مجموعه داده‌های نامتوازن حاصله معمولاً به سمت کلاس‌های غالب (اکثریت) و برخلاف کلاس‌های نادر (اقلیت) خواهد بود که به نوبه خود، عملکرد سیستم را تحت‌تأثیر قرار می‌دهد [۴]. این وضعیت "مشکل یادگیری نامتعادل (نامتوازن) یا به عبارت دیگر، "یادگیری از مجموعه داده‌های نامتوازن" نامیده می‌شود.

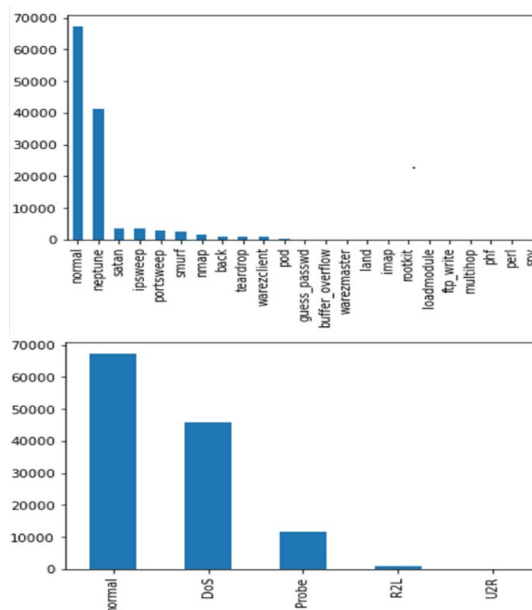
به تازگی، یادگیری نامتوازن توجه زیادی را در یادگیری ماشین به خود جلب کرده است، چراکه الگوریتم‌هایی که در حال حاضر پیاده‌سازی می‌شوند، به‌ویژه در خصوص کلاس‌های اقلیت، عملکرد رضایت‌بخشی نشان نمی‌دهند.

به چند دلیل باید چالش جدی تشخیص حملات نادر توسط سیستم‌های تشخیص نفوذ (IDS) را جدی گرفت که مهم‌ترین آنها عبارت‌اند از:

- بخش زیادی از پایگاه‌های داده‌ای که برای ارزیابی عملکرد این سیستم‌ها استفاده می‌شوند، دارای توزیع کلاس چوله (ناموزون) هستند [۵] و تعداد نمونه‌های آموزشی در یک کلاس، به‌شدت از تعداد نمونه‌های آموزشی در کلاس دیگر بیشتر است. به‌عنوان مثال، برخی از مهم‌ترین مجموعه‌های داده در زمینه سیستم‌های تشخیص نفوذ (IDS)، مانند KDD-CUP99، KDD، NSL، کیوتو یونیورسیتی^۳ و دادگان UNB ISCX 2012، نامتوازن هستند. دو دادگان KDD-CUP99 و NSL-KDD شامل دو کلاس حمله کوچک R2L و U2R هستند که تعداد رکوردهای آنها بسیار کمتر از کلاس‌های دیگر است (به شکل (۱) رجوع کنید).
- اکثر طبقه‌بندی‌کننده‌های یادگیری ماشین که در سیستم‌های تشخیص نفوذ (IDS) کاربرد دارند، دقت محور هستند. این طبقه‌بندی‌کننده‌ها در دستیابی به سطوح بالای دقت کلی موفق ظاهر شده‌اند، اما هنوز برای کلاس‌های اقلیت به موفقیت قابل‌قبولی نرسیده‌اند. مثلاً، یک مجموعه داده شامل یک کلاس خاص با ۹۸ نمونه و یک کلاس دیگر تنها با ۲ نمونه را در نظر بگیرید. اگر همه داده‌ها برای طبقه‌بندی‌کننده‌ای که از این مجموعه داده یاد می‌گیرد، توسط کلاس اول برچسب‌گذاری شوند، این طبقه‌بندی‌کننده همچنان به دقت ۹۸٪ دست خواهد یافت [۶].
- ریسک حملات کلاس‌های اقلیت می‌تواند به مراتب مخرب‌تر از کلاس‌های اکثریت باشد. هزینه طبقه‌بندی اشتباه این حملات بسیار زیاد است؛ بنابراین، مسئله رسیدن به درصد تشخیص بالا برای این کلاس‌ها برای تضمین قابلیت اطمینان، یکپارچگی و در دسترس بودن سیستم‌های شبکه، امری ضروری و چالش برانگیز است.
- همپوشانی کلاس‌ها در مرحله انتخاب ویژگی منجر به ایجاد چالش‌های مهمی می‌شود. زیرا ممکن است یک زیرمجموعه ویژگی مشخص نتواند نرخ تشخیص بالایی داشته باشد. تقریباً هدف کلیه تکنیک‌های انتخاب ویژگی اعم از الگوریتم پالایش یا فیلتر و پوشش، یافتن مجموعه ویژگی‌های بهینه برای مجموعه داده‌ها است. به‌طور کلی، نتایج کاملاً رضایت‌بخش هستند. هرچند، برای کلاس‌های اقلیت، نرخ تشخیص پایین اعمال شد. دو کلاس R2L و U2R در

³ Kyoto University¹ Intrusion detection systems² Bias

دادگان NSL-KDD این نوع مشکل همپوشانی را نشان می دهند [۴].



شکل (۱). توزیع کلاس های اصلی و کلاس های فرعی در دادگان NSL-KDD

تاکنون راهکارهای مختلفی مانند نمونه گیری مجدد^۱ و یادگیری حساس به هزینه^۲ برای حل مسئله یادگیری نامتوازن پیشنهاد شده است [۷]. علی رغم این که این راه حل ها در برخی موارد موفقیت آمیز بوده اند، مشکل بیش برآزش^۳ در نمونه برداری بیش از حد^۴ همچنان به چشم می خورد، در حالی که نمونه گیری undersampling به طور بالقوه اطلاعات مفید را حذف می کند. علاوه بر این، در این سیستم ها الگوریتم های یادگیری عمیق^۵ بسیاری پیاده سازی شده و عملکرد IDS به میزان قابل توجهی بهبود یافته است. با این حال، نتایج حاصل از آنها، از سطوح بالای دقت برای حملات نادر برخوردار نیستند. دلیل این امر آن است که این الگوریتم های یادگیری عمیق به مقادیر بسیار زیادی داده برای دستیابی به نسبت های تشخیص بزرگ نیاز دارند.

وقتی که صحبت از کاربرد در دنیای واقعی به میان باشد، نمی توان همیشه نمونه داده های کافی و مورد نیاز برای آموزش یک طبقه بندی کننده قوی با قابلیت شناسایی هرگونه حمله را به دست آورد. یکی از دلایل این وضعیت آن است که محیط فضای مجازی، محیطی پویا است. دلیل دیگر آن است که همیشه انواع جدیدی از حملات رخ می دهند. به علاوه اینکه، هیچ نهادی مایل نیست اطلاعات حساس خود را در معرض ریسک این گونه

حملات قرار دهد. جوامع امنیتی نیز در به دست آوردن نمونه های کافی از حملات در یک دوره زمانی کوتاه مدت، با مشکل مواجه هستند [۸]؛ بنابراین، این احتمال وجود دارد که به دلیل مسائل امنیتی یا حساسیت داده ها یا حتی هزینه زیاد یا هر دلیل دیگری، هنگام جمع آوری تعدادی از نمونه های برچسب زده شده، مشکلاتی رخ دهد که به نوبه خود منجر به ایجاد مسئله عدم توازن داده ها می شود. رفع این محدودیت ها در تشخیص نفوذ در دنیای واقعی، مستلزم استفاده از الگوریتم های کارآمدتری است که بتوانند مشکل داده های محدود برای حملات نادر را به منظور آموزش طبقه بندی کننده های نفوذ با قابلیت های تعمیم عالی، برطرف نمایند.

به غیر از روش های مرسوم که در بالا به آن اشاره شد، برای بهبود قابلیت تعمیم IDS در برابر حملات نادر یا حتی حملات جدید و حل چالش توزیع نامتوازن داده ها، ما از جدیدترین تکنیک در زمینه یادگیری عمیق یعنی آموزش با داده محدود (FSL)^۷ نیز استفاده خواهیم کرد. اگرچه مدت زیادی از عمر تکنیک FSL نمی گذرد، اما این روش در مواردی که از مشکل کمبود نمونه های برچسب گذاری شده مانند تشخیص تصویر و بازشناسی کاراکتر رنج می برند، به نتایج قابل توجهی دست یافته است [۹-۱۳]. علاوه بر این، شبکه های عصبی^۸ استاندارد قادر به یادگیری وظایف^۹ جدید نیستند. طی سال های اخیر، محققان با الهام از این که انسان ها چگونه می توانند مهارت های جدید را به سرعت یاد بگیرند یا فقط از طریق چند تجربه ساده، با آنها انطباق پیدا کنند، رویکردهای زیادی را برای افزودن قابلیت یادگیری از چند نمونه محدود، به شبکه های عصبی عمیق پیشنهاد کرده اند. یادگیری متا^{۱۰} یا «یادگیری برای یادگیری» یکی از رویکردهایی است که به تازگی به منظور مقابله با این نوع محدودیت در شبکه های عصبی استاندارد مورد استفاده قرار می گیرد. در فرایند یادگیری متا، روند یادگیری در دو سطح فرا یادگیری و استراتژی های انطباق^{۱۱} انجام می گردد. در سطح یادگیری متا، مدل مورد نظر، مجموعه پارامترهای اولیه مدل را در میان وظایف (task) مختلف یاد می گیرد. در طول استراتژی انطباق، این روند، بر مسئله کسب سریع دانش برای یادگیری پارامترهای مخصوص هر وظیفه تمرکز دارد [۱۴]. تاکنون تعداد زیادی از رویکردهای یادگیری متا برای حل مشکلات آموزش با داده محدود پیشنهاد شده است. این رویکردها اساساً بسته به نحوه انجام روند بهینه سازی فرایند آموزش، با یکدیگر فرق دارند. به طور کلی، رویکردهای مورد بحث را می توان به سه دسته اصلی

⁷ Few-Shot learning

⁸ Neural networks

⁹ Task

¹⁰ Meta-learning

¹¹ Adaptation strategy

¹ Resampling

² Cost sensitive learning

³ Overfitting

⁴ Oversampling

⁵ Deep learning

⁶ Cyberspace

تلاش‌های صورت گرفته برای بهبود قابلیت تعمیم و سرعت همگرایی MAML نسبت به الگوریتم اصلی MAML، اساساً به کاربرد مراحل SGD در حلقه درونی بستگی دارد که به نوبه خود از مشکلاتی که قبلاً اشاره کردیم رنج می‌برد.

ما برای غلبه بر این محدودیت‌ها در MAML، الگوریتم NDEKF-MAML را توسعه دادیم که در آن، به‌منظور بهینه‌سازی آموزش MAML، الگوریتم گرادینان نزولی در حلقه درونی با فیلتر کالمن گسترش‌یافته جداسده (NDEKF)^۸ جایگزین می‌شود [۱۹]. فیلتر کالمن گسترش‌یافته (EKF)^۹ به‌عنوان یک روش برآورد حالت در سیستم‌های غیرخطی، کاربرد فراوانی دارد. رویکرد فیلتر کالمن گسترش‌یافته جداسده برای آموزش شبکه عصبی، یکی از روش‌های مرتبه دوم^{۱۰} باتوجه به فرضیات ساده‌سازی^{۱۱} است [۲۰].

برای افزایش عملکرد همگرایی در مقایسه با الگوریتم پس انتشار، چندین الگوریتم یادگیری شبکه عصبی چندلایه با استفاده از EKF پیشنهاد شده است [۲۱-۲۳]. استفاده از EKF، برآورد حداقل واریانس وزن‌های اتصال^{۱۲} را در اختیار ما قرار می‌دهد به‌طوری که با سرعتی برابر مرتبه دوم همگرا می‌شود. با این حال، این روش نیز یک اشکال جدی در رابطه با پیچیدگی‌های محاسباتی دارد [۲۴].

اگر همه وزن‌های شبکه در الگوریتم گنجانده شوند، ماتریس‌های خروجی پیچیده شده و مدیریت آن دشوار خواهد بود. به همین دلیل، یک نسخه جدا از گره^{۱۳} از این الگوریتم برای ساده کردن فرایند محاسبات پیاده‌سازی شده است که عامل مهمی در شبکه‌های عصبی عمیق به شمار می‌رود.

دلیل اصلی این امر آن است که DEKF با تقریب‌زدن داده‌های مرتبه دوم در میان‌وزن‌ها، از EKF به دست می‌آید به‌طوری که این داده‌ها فقط به گروه‌های دویه‌دو ناسازگار تعلق داشته باشند. ممکن است در این‌بین، همبستگی بین تخمین وزن‌ها نادیده گرفته شود که نتیجه آن، یک ماتریس کوواریانس خطای پراکنده P_k خواهد بود. برای بهینه‌سازی روند آموزش شبکه‌های عصبی، الگوریتم DEKF [۲۴، ۲۵] با موفقیت پیاده‌سازی شده است. در مجموع، DEKF به مراحل آموزشی کمتری نیاز داشته و در مقایسه با الگوریتم‌های گرادینان نزولی، قابلیت تعمیم بهتری را ارائه می‌دهد.

ایده‌های اصلی در این مقاله بدین شرح هستند. اولین ایده آموزش یک مدل عمیق با داده محدود و با استفاده از روش

طبقه‌بندی کرد: رویکردهای مبتنی بر مدل^۱، مبتنی بر معیار^۲ و رویکردهای مبتنی بر بهینه‌سازی^۳.

برای مقابله با مشکل داده‌های نامتوازن در IDS و افزایش قابلیت تعمیم مدل در برابر حملات نادر، ما الگوریتم نویدبخش یادگیری متامدل - آگنوستیک (MAML) را ساخته‌ایم [۱۰] که یک رویکرد یادگیری متا مبتنی بر بهینه‌سازی ساده و درعین حال مؤثر است. الگوریتم MAML دو مرحله دو در دو را شامل می‌شود. مرحله داخلی برخی از مراحل گرادینان نزولی تصادفی (SGD)^۴ را برای هر وظیفه مجزا اجرا و مرحله خارجی، متا پارامتر را در تمام وظیفه‌های نمونه به‌روزرسانی می‌کند. با این حال، در مرحله یادگیری متا از الگوریتم MAML، از الگوریتم یادگیری استاندارد (گرادینان نزولی) در حلقه داخلی استفاده می‌شود. در این مرحله یک یا چند مرحله گرادینان نزولی کامل بر روی تمامی وزن‌ها اجرا می‌گردد. این امر باعث می‌شود این روش با مشکلات عمده‌ای روبرو گردد که مهم‌ترین آنها عبارت‌اند از:

- مقداردهی اولیه پارامترهای مدل مستلزم پس انتشار یا انتشار معکوس^۵ از طریق فرایند بهینه‌سازی داخلی با استفاده از چندین مرحله از گرادینان نزولی است. در نتیجه، فرایند یادگیری متا به مشتقات مرتبه بالاتر [۱۵] نیاز دارد که نتیجه آن همگرایی کند MAML و الزامات محاسباتی و حافظه بسیار است.
- نتیجه استفاده صرف از وزن‌های نهایی که در حلقه داخلی برای یادگیری حلقه بیرونی محاسبه می‌شوند، عملیات پس انتشار اضافی خواهد بود که می‌تواند منجر به مشکل محوشدگی^۶ یا انفجار^۷ گرادینان در MAML شود.
- از آنجایی که تمام وزن‌ها از طریق چند مرحله SGD و بر اساس چند نقطه داده (همان داده‌های پشتیبانی) به‌روزرسانی می‌شوند، ممکن است وضعیت بیش برآزش اتفاق بیفتد.

برای رفع این محدودیت‌ها در الگوریتم MAML و به‌منظور غلبه بر محدودیت‌های ناشی از تمایز از طریق حلقه درونی در داده‌های کوچک (فرایند تنظیم دقیق یا fine-tuning) که در آن پارامترهای شبکه با استفاده از چند مرحله SGD به‌روزرسانی می‌شوند، چندین راه حل پیشنهاد شده است [۱۵-۱۸]. هرچند، این محدودیت‌ها زمانی تأثیرگذارتر می‌شوند که مجموعه داده‌های آموزشی، همانند مجموعه داده‌های IDS پراکنده و ناپایدار باشند.

¹ Model-based

² Metric-based

³ Optimization-based approaches

⁴ Stochastic gradient descent

⁵ Backpropagation

⁶ Vanishing

⁷ Exploding

⁸ Decoupled Extended Kalman Filter

⁹ Extended Kalman filter

¹⁰ Second-order Methods

¹¹ Simplifying assumptions

¹² Connection-weights

¹³ Node-decoupled version

ساختار این مقاله در ادامه بدین شرح است. در بخش ۲ کارهای مرتبط معرفی می شوند. در بخش ۳ مدل جدید تشخیص نفوذ MAML-NDEKF ارائه شده و نحوه عملکرد آن به تفصیل بیان می شود. در بخش ۴ جزئیات و نتایج تجربی شرح داده می شوند. در نهایت در بخش ۵ نیز نتیجه گیری و مقایسه با کارهای تحقیقاتی دیگر ارائه می شود.

۲. کارهای مرتبط

بر اساس مقدمه فوق و روش پیشنهادی برای حل مشکل شناسایی حملات نادر ناشی از عدم توازن داده ها در IDS، از سه مرحله استفاده می شود. در مرحله اول روش هایی را مورد بحث قرار می گیرد که برای حل مشکل عدم توازن در مجموعه داده های IDS مورد استفاده قرار گرفته اند. در مرحله دوم، روش هایی را که از فیلتر کالمن در آموزش شبکه های عصبی استفاده کرده اند مورد نقد و بررسی قرار می گیرد. در آخرین مرحله نیز از آموزش با داده محدود بهره گرفتیم.

۲-۱. روش های رفع مشکل عدم توازن کلاس در IDS

علی رغم این واقعیت که فناوری های یادگیری ماشین مبتنی بر IDS قابلیت تشخیص و خودسازگاری قابل توجهی را برای جبران تغییرات محیط عملیاتی شبکه فراهم می سازد، ولی عملکرد آنها همچنان تحت تأثیر توزیع نامتوازن داده ها قرار دارد.

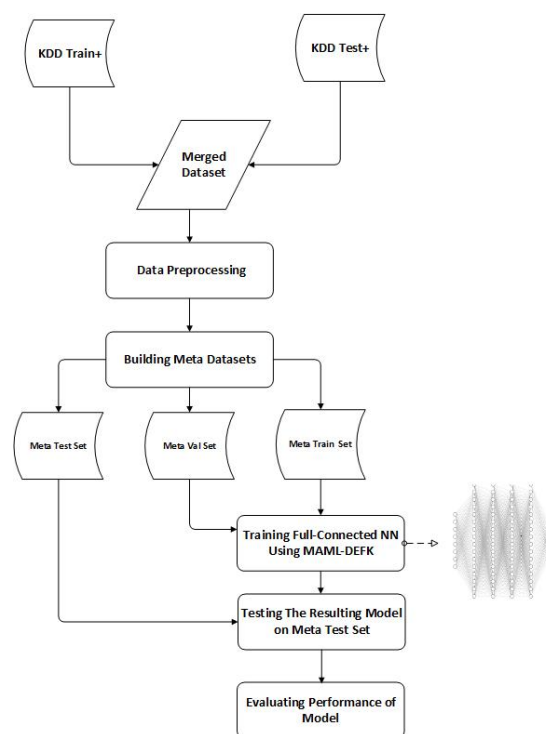
بسیاری از مطالعاتی که تکنیک های یادگیری ماشین را (خواه سطحی یا عمیق) مورد استفاده قرار داده اند (جدول (۱))، نتوانسته اند برای حملات نادر و حتی برای حملات سایبری ناشناخته یا مشاهده نشده ای که در مجموعه داده آموزشی وجود ندارند، به نرخ شناسایی رضایت بخشی دست پیدا کنند.

در [۲۶]، محققان یک مدل ترکیبی چند سطحی متشکل از ماشین بردار پشتیبان (SVM) و ماشین یادگیری حداکثری (ELM) را پیشنهاد کرده اند. این مدل با استفاده از مجموعه داده KDD Cup 1999 ارزیابی شده است. به علاوه، در این مدل، یک الگوریتم کی-میانگین^۴ اصلاح شده برای کاهش اندازه مجموعه داده های آموزشی و ایجاد توازن در مجموعه داده های آموزشی برای آموزش SVM و ELM اجرا می شود. با این حال، نرخ تشخیص برای کلاس های حملات نادر U2R و R2L، به ترتیب معادل ۲۱٫۹۳٪ و ۳۱٫۳۹٪ بود.

از سوی دیگر، در برخی از مطالعات حملات دارای فراوانی اندک بررسی نشده یا حتی به طور کلی کنار گذاشته شده اند. در پژوهش [۲۷]، محققان روشی را پیشنهاد کردند که خود رمزگذار

MAML برای یادگیری حملات سایبری و با چالش دسته های نامتوازن از طریق یادگیری تحت نظارت^۱ است. در ایده دوم از فیلتر کالمن برای بهینه سازی پارامترهای یک شبکه DNN نسبتاً بزرگ استفاده کردیم که از ۴ لایه پنهان هر کدام با اندازه ۳۲ نرون تشکیل شده بود. این کار در هیچ یک از مطالعات قبلی که از فیلتر کالمن برای بهینه سازی پارامترهای شبکه عصبی استفاده کرده اند، انجام نشده است و در همه کارهای گذشته یک لایه و نهایتاً با ۱۰ نرون استفاده شده است. ایده سوم استفاده از فیلتر NDEKF برای بهینه سازی پارامترهای شبکه DNN به جای گرادین کاهشی در حلقه داخلی MAML است. با به کارگیری چارچوب جدیدمان یعنی MAML-NDEKF توانستیم به بهبود سرعت همگرایی، افزایش قدرت تعمیم و افزایش دقت در مقایسه با الگوریتم اصلی MAML هنگام رسیدگی به مجموعه داده های نامتوازن مانند NSL-KDD دست یابیم.

روند اجرای کار در این روش در شکل (۲) آمده است. در بخش ۴ توضیح بخش های مختلف این شکل به طور کامل شرح داده شده است.



شکل (۲). روند کار MAML-NDEKF پیشنهادی

در نهایت میزان دقت برای این دو حمله اقلیت R2L و U2R با استفاده از چارچوب پیشنهادی ما حتی با کاهش تعداد دوره های آموزشی نیز از MAML اصلی پیشی گرفت و به ترتیب از ۶۱٪ به ۷۵٪ و از ۵۱٪ به ۶۶٪ افزایش یافت.

² Support vector machine

³ Extreme learning machine

⁴ K-means algorithm

¹ supervised learning

مدل پیشنهادی خود، به جز نرخ تشخیص، معیار دیگری را در نظر نگرفتند. این معیار به تنهایی در مواقع رویارویی با توزیع نامتوازن داده‌ها کافی نیست.

در [۲۹]، نویسندگان سه چارچوب مبتنی بر داده کاوی^۶ را پیشنهاد کردند که از الگوریتم جنگل‌های تصادفی (RF) در موارد کاربرد نابجا، ناهنجاری و تشخیص ترکیبی برای بهبود نرخ تشخیص IDS و شناسایی نفوذهای جدید بهره می‌گیرند. برای حل مشکل نفوذهای نامتوازن نیز از نفوذهای اقلیت نمونه‌گیری oversampling و هم از نفوذهای اکثریت نمونه‌گیری undersampling به عنوان یک مرحله پیش‌پردازش استفاده کردند. آنها با استفاده از این چارچوب‌ها در KDD Cup '99 ادعا کردند که مدل پیشنهادی در این پژوهش می‌تواند عملکرد کلی IDSهای فوق را بهبود ببخشد. این درحالی است که در گزارش نهایی آنها هیچ اشاره‌ای به نرخ شناسایی و عملکرد مدل در مورد هر حمله، نشده است. مشکل اصلی مطالعه مورد بحث این است که برخی از نفوذها با درجه تشابه بالا قابل شناسایی نبودند.

در پژوهش [۳۰]، نویسندگان علاوه بر الگوریتم‌های یادگیری متا vote و Random Committee، از شبکه بیزی^۷ و الگوریتم درخت تصادفی^۸ به عنوان طبقه‌بندی‌کننده اصلی استفاده کردند. مجموعه داده KDD Cup '99 نیز برای ارزیابی عملکرد معرفی شده است. اگرچه عملکرد مدل آنها در سطح خوب و قابل قبولی بود، اما میزان دقت آن برای حملات U2R بسیار پایین است.

در [۳۱] مدل تشخیص نفوذ جدیدی به نام ICVAE-DNN ارائه شد. برخلاف استفاده از خودکدگذار^۹ متعارف برای یادگیری ویژگی^{۱۰}، در این جا یک خودکدگذار متغیر شرطی بهبودیافته^{۱۱} (ICVAE) برای بیش‌نمونه‌برداری^{۱۲} از داده‌های آموزشی به‌منظور رسیدگی به مجموعه داده‌های نامتوازن به کار برده می‌شود. ICVAE نمونه‌های حمله جدیدی را برای متوازن‌سازی داده‌ها و افزایش تنوع تولید می‌کند. سپس از این نمونه‌ها برای مقارن‌دهی اولیه^{۱۳} وزن‌های یک شبکه عصبی عمیق (DNN) استفاده می‌شود. نتایج تجربی نشان می‌دهند که ICVAE-DNN در مقایسه با پیشرفته‌ترین الگوریتم‌ها و روش‌های بیش‌نمونه‌برداری متعارف در تشخیص حملات اقلیت مانند U2R و R2L دارای عملکرد بهتری است.

عمیق غیرمتقارن (NDAE)^۱ (یادگیری عمیق) را با دقت و سرعت الگوریتم جنگل تصادفی (RF)^۲ (یادگیری کم‌عمق) ترکیب می‌کند. آنها از NDAE برای یادگیری روابط پیچیده بین ویژگی‌های مختلف استفاده کردند. علاوه بر این، این روش از قابلیت استخراج ویژگی بالایی برخوردار است که به آن اجازه می‌دهد پارامترهای مدل را با اولویت‌بندی توصیفی‌ترین ویژگی‌ها^۳ اصلاح کند. آنها در ادامه از الگوریتم جنگل تصادفی به عنوان طبقه‌بندی‌کننده استفاده کردند. این الگوریتم در مقابل مجموعه داده‌های معیار NSL-KDD (همه مجموعه داده‌های آموزشی) و KDD Cup '99 (۱۰ درصد از زیر مجموعه برای آموزش) پیاده‌سازی و ارزیابی شده است. آنها مدل خود را بر اساس طبقه‌بندی ۵ کلاسه برای KDD Cup '99 و ۱۳ کلاسه برای NSL-KDD مورد ارزیابی قرار دادند، درحالی که الگوهای حمله سایبری با کمتر از ۲۰ مورد در مجموعه داده را برای همه تجربیات حذف کردند. در نهایت، نتایج مدل پیشنهادی و مدل DBN^۴ را با یکدیگر مقایسه کردند. آنها مدعی هستند که این مدل یادگیری عمیق، به ترتیب، ضریب دقت ۹۷.۸۵٪ و ۸۵.۴۲٪ را در طبقه‌بندی ۵ کلاسه برای KDD Cup '99 و NSL-KDD و دقت ۸۹.۲۲٪ را برای طبقه‌بندی ۱۳ کلاسه NSL-KDD ارائه می‌دهد.

بر اساس نتایج به‌دست آمده، این رویکردها در مقایسه با مدل DBN دقت بالاتری را ارائه کرده و مستلزم صرف زمان کمتری برای آموزش شبکه عصبی هستند.

علاوه بر این، به‌منظور حل مشکل عدم توازن در مجموعه داده‌های تشخیص نفوذ، روش‌های کارآمد متعددی با استفاده از تکنیک‌های نمونه‌گیری مجدد پیشنهاد شده است. در پژوهش [۲۸]، مؤلفین باتوجه به سطح داده‌ها، یک روش نمونه‌گیری موسوم به RA-SMOTE^۵ را برای حل مشکل عدم توازن کلاس پیشنهاد کردند. در این اثر، از سه نوع طبقه‌بندی‌کننده مختلف برای آزمایش اثربخشی الگوریتم استفاده شد که عبارت‌اند از: ماشین‌های بردار پشتیبان (SVM)، شبکه عصبی BP (یا BPNN) و جنگل‌های تصادفی (RF).

نتایج تجربی مربوط به مجموعه داده NSL-KDD نشان داد که الگوریتم پیشنهادی می‌تواند به شکلی کارآمد مشکل عدم توازن کلاس را حل کند و نسبت به سایر تکنیک‌های پیشرفته و نوین IDS موجود، نتایج به‌مراتب بهتری را ارائه دهد. در این شیوه، نرخ تشخیص حملات اقلیت U2R و R2L، به ترتیب برابر ۹۸.۶۳٪ و ۹۸.۸۷٪ بود. هرچند، محققان برای ارزیابی عملکرد

⁶ Data mining

⁷ Bayesian network

⁸ Random Tree

⁹ autoencoder

¹⁰ feature learning

¹¹ Improved Conditional Variational AutoEncoder

¹² oversampling

¹ Non-Symmetric Deep Auto-Encoder

² Random Forest

³ Descriptive features

⁴ Deep Belief Network

⁵ Region Adaptive Synthetic Minority Oversampling Technique (RA-SMOTE)

جدول (۱). خلاصه کارهای مرتبط که به موضوع داده‌های نامتوازن در سیستم‌های تشخیص نفوذ (IDS) می‌پردازد.

سال انتشار مقاله	الگوریتم‌های مورد استفاده برای رسیدگی به مسئله عدم توازن دسته‌ها	مدل دسته‌بندی کننده مورد استفاده	توضیحات	مجموعه داده‌های IDS	نتایج
[۲۶]، ۲۰۱۷	در این کار مطالعاتی، مجموعه داده‌های نامتوازن مورد بررسی قرار نگرفتند و صرفاً اندازه مجموعه داده‌ها با استفاده از الگوریتم کی - میانگین ^۱ اصلاح شده کاهش داده شد.	SVM ترکیبی چندسطحی و ELM (ماشین یادگیری افراطی) ^۲	نویسندگان مقاله ابتدا الگوریتم کی - میانگین اصلاح شده را برای کاهش اندازه مجموعه داده‌ها به کار بردند. سپس از یک مدل ترکیبی چندسطحی متشکل از ماشین بردار پشتیبان ^۳ (SVM) و ماشین یادگیری افراطی استفاده کردند.	KDD Cup 99 (دسته‌بندی ۵ دسته‌ای)	نرخ تشخیص برای دسته‌های نادر حملات U2R و R2L به ترتیب فقط ۲۱٫۹۳٪ و ۳۱٫۳۹٪ بود.
[۲۷]، ۲۰۱۸	حملات فرانکس پایین مورد بررسی قرار نگرفتند.	استخراج ویژگی با استفاده از روش عمیق RF و NDAE به عنوان دسته‌بندی کننده	نویسندگان مقاله، خود کدگذار عمیق نامتوازن ^۴ انباشته (NSAE) را برای یادگیری روابط پیچیده بین ویژگی‌های مختلف به کار بردند. سپس از جنگل تصادفی ^۵ به عنوان دسته‌بندی کننده استفاده کردند.	KDD Cup 99 (دسته‌بندی ۵ دسته‌ای) و NSL-KDD (دسته‌بندی ۵ دسته‌ای و ۱۳ دسته‌ای)	نتایج حاصل دارای دقت ۸۵٫۹۷٪ و ۸۵٫۴۲٪ در دسته‌بندی ۵ دسته‌ای به ترتیب برای KDD Cup 99 و NSL-KDD و دقت ۸۹٫۲۲٪ برای NSL-KDD ۱۳ دسته‌ای بودند.
[۲۹]، ۲۰۰۸	روش‌های بیش نمونه‌برداری و کم نمونه‌برداری	الگوریتم جنگل‌های تصادفی (RF) در تشخیص سوء استفاده، ناهنجاری و ترکیبی		KDD Cup 99	در این کار مطالعاتی اشاره نشد که نرخ تشخیص و عملکرد آن برای هر حمله چگونه است.
[۳۱]، ۲۰۱۹	خود کدگذار متغیر شرطی بهبود یافته (ICVAE) به عنوان روش استخراج ویژگی و بیش نمونه‌برداری	شبکه عصبی عمیق (DNN)	ICVAE برای یادگیری و بررسی نمایش‌های تُنک ^۶ بالقوه بین ویژگی‌ها و دسته‌های داده‌های شبکه به کار برده می‌شود. کدگشای آموزش دیده ICVAE، نمونه‌های حمله جدیدی را تولید می‌کند. کدگشای آموزش دیده ICVAE برای مقاردهی اولیه وزن لایه‌های پنهان DNN نیز مورد استفاده قرار گرفته است.	NSL-KDD (دسته‌بندی ۵ دسته‌ای) و UNSW-NB15 (دسته‌بندی ۱۰ دسته‌ای)	نویسندگان مقاله ادعا کردند که مقدار DR بالاتری برای U2R و R2L نسبت به مطالعات دیگری که ذکر کرده بودند، به دست آوردند. DR برای U2R و R2L در (KDDTest-21) به ترتیب ۱۱٫۰۰ و ۴۴٫۱۷ بود.

^۱ k-means algorithm

^۲ Extreme learning machine

^۳ support vector machine

^۴ Non-Symmetric Deep Auto-Encoder

^۵ Random Forest

^۶ sparse representations

مشقتات درجه دوم نادیده گرفته می‌شوند. آنها دریافتند که رویکرد جدید هنگام اعمال بر روی مجموعه داده‌های ImageNet، عملکردی مشابه با MAML مرتبه بالاتر دارد. اگرچه پیاده‌سازی روش جدید، ساده‌تر و از نظر محاسباتی مقرون به صرفه‌تر از MAML اصلی است، این روند ساده‌سازی منجر به ازدست‌دادن اطلاعات گرادیان می‌شود.

یکی دیگر از الگوریتم‌های بهینه‌سازی یادگیری متا، Reptile نام دارد [۱۶] که اساساً الهام گرفته از رویکرد FOMAML است. این رویکرد شباهت بسیاری با MAML اصلی دارد، زیرا از طریق اجرای چند مرحله SGD روی تعداد کمی از نمونه‌ها در حلقه درونی، یک مقداردهی اولیه برای پارامترهای شبکه عصبی را می‌آموزد تا به سازگاری سریع‌تر در وظایف جدید و نرخ تعمیم قابل قبولی دست یابد. الگوریتم MAML از طریق فرایند "تنظیم دقیق" به مقادیر وزن اولیه پس‌انتشار می‌یابد، بنابراین محاسبه مشتقات دوم مورد نیاز است. در عوض، رویکرد Reptile گراف محاسباتی را باز نمی‌کند. بلکه بردار وزن مرکزی را در جهت وزن‌های تنظیم شده در طول حلقه بیرونی، به‌روزرسانی و از مشتقات دوم گرادیان‌ها اجتناب می‌کند. علی‌رغم محاسبات مقرون به صرفه و اجرای ساده Reptile در مقایسه با MAML اصلی بدون در نظر گرفتن مشتقات مرتبه دوم، انجام مراحل SGD در حلقه درونی آن باعث می‌شود کیفیت آن به واسطه مشکلات گرادیان اصلی مانند محوشدگی یا انفجار گرادیان، گیر کردن در بهینه محلی^۳ و بیش برآزش تهدید شود.

یکی دیگر از نمونه‌های بهبودیافته MAML، الگوریتم MAML++ نام دارد است. محققان در این پژوهش، راه‌حل‌های متعددی را برای مقابله با چالش‌های MAML و بهبود عملکرد تعمیم و سرعت همگرایی آن پیشنهاد کرده‌اند. یکی از راهکارهای پیشنهادی آنها برای غلبه بر مشکل تخریب گرادیان^۴ MAML و بهبود پایداری آن، روش بهینه‌سازی اتلاف چندمرحله‌ای (MSL)^۵ است. راهکار MSL به جای به حداقل رساندن تلفات مجموعه هدف پس از محاسبه تمام به‌روزرسانی‌های حلقه درونی، پس از هر بار به‌روزرسانی اتلاف در مجموعه پشتیبان در حلقه درونی، اتلاف در مجموعه هدف را به حداقل می‌رساند. یک راه‌حل دیگر که برای بهبود قابلیت تعمیم و سرعت همگرایی MAML پیشنهاد شده است، یادگیری خودکار نرخ یادگیری حلقه درونی است. [۱۷]

۳-۲. فیلتر کالمن برای آموزش شبکه‌های عصبی

برآورد حالت بهینه برای سیستم دینامیکی خطی با استفاده

علی‌رغم پیشرفت‌های قابل توجهی که از طریق روش‌های یادگیری عمیق برای ساختن سیستم‌های تشخیص نفوذ کارآمد و مؤثر حاصل شده است، هنگام اعمال بر روی داده‌های نامتوازن، عملکرد ضعیفی از خود نشان داده‌اند. همین امر لزوم بررسی بیشتر رویکردهای جدید برای رسیدگی به این موضوع را برجسته می‌کند. به عقیده ما، به جای جمع‌آوری نمونه‌های آموزشی بیشتر که در اکثر کاربردهای دنیای واقعی، عملاً کاری دشوار است، باید به فکر یافتن راهی برای یادگیری از چند نمونه یا حتی بدون نمونه بود؛ بنابراین، فکر کردیم که می‌توانیم با استفاده از الگوریتم آموزش با داده محدود و یادگیری متا باهدف ارتقای سطح عملکرد سیستم تشخیص نفوذ، به هدف نهایی خود در این تحقیق برسیم. اگرچه انجام کار پژوهشی در زمینه یادگیری آموزش با داده محدود، پیشینه زیادی ندارد و عمر این‌گونه تحقیقات به‌سختی به ۱۰ سال می‌رسد، آموزش با داده محدود در زمینه تشخیص کاراکتر^۱ [۳۲]، طبقه‌بندی تصویر [۳۳، ۳۴]، تشخیص رویدادهای ویدئویی [۳۵] و طبقه‌بندی متون به نتایج قابل توجهی دست‌یافته است. باین حال، تاکنون تنها در یک یا دو مطالعه از آموزش با داده محدود برای تشخیص نفوذ استفاده شده است.

۲-۲. آموزش با داده محدود با استفاده از یادگیری متامدل آگنوستیک (MAML)

الگوریتم MAML [۱۰] یک رویکرد یادگیری متا مبتنی بر بهینه‌سازی است که برای یادگیری پارامترهای اولیه مدل مورد استفاده قرار می‌گیرد به فرایند یادگیری در وظایف جدید از طریق بهینه‌سازی مبتنی بر گرادیان، سرعت می‌بخشد. علاوه بر این، MAML مدل اولیه f_θ به همراه متاپارامتر θ را به عنوان مدل پایه در نظر می‌گیرد. سپس پارامتر θ را می‌آموزد که هنگام استفاده به عنوان پارامترهای اولیه، امکان انطباق سریع‌تر با وظایف جدید را فراهم می‌کند. این الگوریتم از طریق فرایند انطباق، پس‌انتشار می‌یابد و آن را محدود به مدل آموزش با داده محدود می‌سازد، زیرا از نظر محاسباتی پرهزینه بوده و در معرض محوشدگی یا انفجار گرادیان قرار دارد.

برای رفع محدودیت‌های ناشی از تمایز از طریق حلقه داخلی در مجموعه داده‌های کوچک (فرایند تنظیم دقیق) در الگوریتم MAML تاکنون نمونه‌های بسیاری پیشنهاد شده است که در آنها پارامترهای شبکه با استفاده از چند مرحله SGD به‌روزرسانی می‌شوند. محققان MAML گونه‌ای از رویکرد پیشنهادی خود را تحت عنوان MAML مرتبه اول (FOMAML)^۲ معرفی کردند [۱۰]. در این رویکرد، هنگام تمایز از طریق گرادیان نزولی،

^۳ Stuck in local optima

^۴ Gradient degradation

^۵ Multi-Step Loss Optimization

^۱ Character recognition

^۲ First-Order MAML

طور خلاصه، DEKF علاوه بر قابلیت انتخاب پیچیدگی محاسبات، از کلیه ویژگی های GEKF نیز برخوردار است.

تاکنون محققان بسیاری تلاش کرده اند از یکی از انواع EKF برای برآورد وزن شبکه عصبی استفاده کنند. به عنوان مثال، در [۴۰]، برای تشخیص آسیب ساختاری روی یک پل بزرگراه، شبکه عصبی با استفاده از EKF آموزش داده شد. اما به دلیل دشوار بودن محاسبه ماتریس ژاکوبین^۷، از یک لایه پنهان تکین استفاده شد.

برای آموزش شبکه های عصبی عمیق با کمک EKF و به منظور غلبه بر چالش محاسبه ماتریس ژاکوبین، محققان در [۴۱] الگوریتم جدیدی را برای آموزش شبکه DBN با استفاده از EKF ابداع کردند. این الگوریتم بر اساس محاسبه مشتقات جزئی ماتریس به صورت لایه به لایه، به جای محاسبه مشتق شبکه عصبی در یک معادله واحد استوار است. چارچوب EKF-DBN پیشنهادی در سه نمونه کاربرد مورد آزمایش قرار گرفت. این چارچوب به لحاظ سرعت و دقت محاسبات برای مدل سازی تحلیلی پیش بینی کننده، عملکردی به مراتب بهتر از رویکردهای BP-DBN و RNN از خود نشان داد.

در پژوهش [۴۲]، برای آموزش پارامترهای شبکه LSTM از الگوریتم EKF مستقل استفاده شد. برای کاهش پیچیدگی محاسبه وزن ها، هر گره در مدل LSTM به عنوان یک زیرسیستم دینامیک مستقل در نظر گرفته می شود که در آن ماتریس کوواریانس^۸ حالت، به صورت بلوک مورب تقریب زده می شود در حالی که همبستگی بین وزن گره های مختلف مانند DEKF با ضمانت عملکرد، نادیده گرفته می شود.

با این حال، الگوریتم یادگیری جهانی فیلتر کالمن گسترش یافته (GEKF)^۹ به دلیل دقت و سرعت همگرایی بالایی که دارد، به شکلی گسترده به عنوان یک تکنیک مرتبه دوم، استفاده می شود [۱۹]. با این حال، پیچیدگی محاسباتی آن (یعنی $O(m^2)$ که در آن m تعداد وزن ها در شبکه عصبی است)، باعث شده است کاربرد این تکنیک برای آموزش شبکه های عمیق چالش برانگیز باشد، زیرا تعداد پارامترهایی که باید تخمین زده شوند بسیار زیاد است. به منظور به حداقل رساندن هزینه محاسباتی GEKF، به جای آن، الگوریتم یادگیری DEKF پیاده سازی شده است. الگوریتم DEKF وابستگی متقابل بین وزن هایی که به صورت دوبه دو به گره های ناسازگار تعلق دارند و در آن وزن ها در گره یا لایه گروهبندی می شوند را نادیده می گیرد [۴۳].

از فیلترینگ کالمن انجام می شود. نسخه های توسعه یافته الگوریتم کالمن را می توان با خطی سازی سیستم حول برآورد فعلی پارامترها برای سیستم های دینامیکی غیرخطی به کاربرد. اگرچه الگوریتم کالمن از نظر محاسباتی پیچیده است، اما پارامترها را به صورت بازگشتی به روز می کند که با تمام داده های ورودی قبلی سازگار است و معمولاً در چند تکرار همگرا می شود. استخراج معادلات کالمن در مراجع زیادی قابل دسترسی است [۳۶، ۳۷]. هنگامی که از فیلترینگ کالمن برای آموزش شبکه عصبی استفاده می شود، وزن های شبکه در واقع پارامترهایی هستند که به عنوان حالت های سیستم در نظر گرفته می شوند و برآورد این حالت ها با استفاده از اطلاعات اعمالی به ورودی شبکه انجام می شود. شرح مختصری از معادلات کالمنی که با استفاده از آن ها پارامترهای شبکه عصبی به صورت سازگار محاسبه می شوند، در بخش ۲-۳ ارائه می شود.

برای رفع اشکالات الگوریتم پس انتشار استاندارد (SBP)^۱، یعنی محوشدگی یا انفجار گرادیان، و بهبود عملکرد همگرایی شبکه های عصبی، اعم از نوع پیش خور^۲ [۳۸، ۳۹] یا بازگشتی^۳ [۲۴]، در تکرارها و نمونه داده های کمتر نسبت به SBP، چندین الگوریتم یادگیری مشتق شده از فیلتر کالمن گسترش یافته (EKF) مورد استفاده قرار گرفته است. با این حال، با افزایش پارامترهای شبکه عصبی و اندازه مجموعه داده آموزش، تعداد عملیات ماتریسی مورد نیاز برای EKF منجر به افزایش پیچیدگی های محاسباتی این رویکردها می شود. برای غلبه بر این محدودیت، یک رویکرد تازه با عنوان فیلتر کالمن گسترش یافته جداسده (DEFK) در [۲۵] ارائه شده است.

یکی از مزایای اصلی رویکرد DEKF این است که وزن های شبکه عصبی را می توان به گونه ای گروهبندی (جدا) کرد که عناصر خاصی از $P(n)$ مرتبط با وزن های چند گروه، قابل اغماض باشند.

علاوه بر این، اگر گروه های وزنی دوبه دو ناسازگار باشند، می توان $P(n)$ را به شکل یک ماتریس بلوک مورب^۴ بازآرایی کرد. اگر با تعداد زیادی از گروه ها مواجه باشیم، شاهد کاهش قابل توجهی در هزینه محاسباتی و الزامات ذخیره سازی خواهیم بود. همچنین، بسته به معیارهایی مانند گره^۵، لایه و یا زیرشبکه^۶، می توان گروه های وزنی را به شکلی منطقی انتخاب کرد تا زمان کلی آموزش به میزان قابل توجهی کاهش یابد. به

^۱ Standard backpropagation algorithm

^۲ Feedforward Neural Network

^۳ Recurrent Neural Network

^۴ Block diagonal

^۵ Node

^۶ Subnetwork

^۷ Jacobian matrix

^۸ Covariance matrix

^۹ Global extended Kalman filter

که در آن، α نرخ یادگیری، θ_i^b وزن‌های شبکه مبنا بعد از طی i مرحله به سمت وظیفه b و $L_{S_b}(f_{\theta_i^b})$ تلفات در مجموعه پشتیبان وظیفه b بعد از i مرحله به‌روزرسانی هستند.

اگر فرض کنیم که اندازه گروه وظایف یا task batch برابر B است، آنگاه می‌توانیم یک تابع هدف متا را به این شرح تعریف کنیم:

$$L_{meta} = \sum_{b=1}^B L_{Q_b}(f_{\theta_N^b}(\theta_0)) \quad (2)$$

معادله (۱) نشان‌دهنده وابستگی مستقیم θ_N^b به θ_0 است.

معادله (۲) هدف متا را توصیف می‌کند، که در آن، کیفیت θ_0 به عنوان تابعی از کل تلفات ناشی از کاربرد θ_0 به عنوان مقداردهی اولیه برای همه وظایف است.

اکنون ما فرایند به‌روزرسانی حلقه بیرونی را به‌صورت بهینه‌سازی (در اینجا کمیته‌سازی) تابع هدف متا L_{meta} تعریف می‌کنیم، که به ما اجازه می‌دهد مقدار اولیه بهینه θ_0 را بدست آوریم.

پس از اینکه این مقدار به‌روزرسانی شد، مقدار اولیه بهینه حاصل برای متا پارامترهای θ_0 را می‌توان این‌گونه بیان کرد:

$$\theta_0 = \theta_0 - \beta \nabla_{\theta} \sum_{b=1}^B L_{Q_b}(f_{\theta_N^b}(\theta_0)) \quad (3)$$

که در آن، β نرخ یادگیری متا و L_{Q_b} تلفات در مجموعه هدف برای وظیفه b هستند. انتروپی متقاطع، تابع تلفات^۲ (یا هزینه) مشترک برای کلیه عبارتهای طبقه‌بندی است.

این الگوریتم از دوفاز بهینه‌سازی تشکیل شده است که در آن MAML در حلقه بیرونی، تلفات مورد انتظار هر وظیفه را پس از یک حلقه درونی یا فاز تطبیق، به حداقل می‌رساند و با چند مرحله گرادیان نزولی که در پارامترهای فعلی مدل مقداردهی شده‌اند، محاسبه می‌گردد.

اکنون می‌توانیم با اعمال تغییراتی اندک در برخی از اسامی اصطلاحات، مراحل مختلف این الگوریتم را به‌طور کلی به شکلی که در [۱۵] آمده است، این‌گونه بیان کنیم:

الگوریتم (۱). الگوریتم MAML
Require: α, β : step size hyperparameters
1: randomly initialize θ
2: while not done do
3: Sample batch of tasks $T_b = \{S_b, Q_b\} \sim p(T)$
4: for all S_b do
5: Evaluate $L_{S_b}(f_{\theta_i^b})$ with respect to K

² Loss function

مطالعات بسیاری، چه به‌صورت تئوری [۴۴] و چه به‌صورت عملی [۴۰، ۴۳، ۴۵] نشان داده‌اند که الگوریتم یادگیری DEKF به لحاظ پایداری و ویژگی‌های همگرایی، عملکردی مشابه با GEKF دارد؛ اما پیچیدگی محاسباتی آن کمتر است. در پژوهش حاضر، ما از فیلتر کالمن گسترش‌یافته از گره (NDEKF) برای محاسبه و به‌روزرسانی وزن‌های سریعی (پارامترهای شبکه) استفاده می‌کنیم که در حلقه درونی (حلقه تنظیم دقیق) الگوریتم MAML با استفاده از چند مرحله SGD محاسبه می‌شوند.

۳. روش پیشنهادی

در این بخش، ابتدا بررسی مختصری بر معادلات الگوریتم داده محدود MAML خواهیم داشت. سپس، سازوکار مورد استفاده خود برای کاربرد NDEKF به‌عنوان تخمین‌گر وزن‌های سریع در حلقه درونی MAML به‌جای تخمین آنها با استفاده از مراحل SGD را توضیح خواهیم داد.

۳-۱. یادگیری متامدل آگنوستیک (MAML)

یادگیری پارامترهای یک مدل برای یافتن بهترین پارامترهای اولیه، ایده اصلی MAML را تشکیل می‌دهد. به همین جهت، باوجود پارامترهای اولیه خوب، این مدل می‌تواند به‌سرعت وظایف جدید را با استفاده از یک بهینه‌ساز معروف، یعنی SGD، یاد بگیرد. اگر پارامترهای مدل در تکرار N بهینه باشند، تنها یک یا چند تکرار SGD مورد نیاز خواهد بود تا این پارامترها با θ_0 برای طبقه‌بندی وظایف جدید، برازش پیدا کند. به عنوان مثال، برازش پارامترهای θ با تکرار N از طریق به‌حداقل‌رساندن انتروپی متقاطع^۱ در مجموعه پشتیبان انجام می‌شود.

اگر شبکه عصبی f_{θ} (به عنوان تابعی از متاپارامترهای θ) مدل اصلی باشد، با شروع از مقدار اولیه $\theta_0 = \theta$ ، می‌خواهیم شبکه عصبی را به‌گونه‌ای آموزش دهیم که بتوانیم پس از چند مرحله نسبتاً اندک از به‌روزرسانی گرادیان N با استفاده از مجموعه پشتیبان S_b داده‌ها، در نهایت به وزن‌های سریع θ_N برسیم. این شبکه عملکرد خوبی در مجموعه هدف Q_b دارد، که در آن، b ایندکس یک مجموعه پشتیبانی خاص در گروهی از وظایف مجموعه پشتیبان می‌باشد. ما مجموعه مراحل به‌روزرسانی N را فرایند به‌روزرسانی حلقه درونی می‌نامیم.

پس از تکرار i که در آن $i = 0 \dots N - 1$ ، پارامترهای شبکه عصبی به‌روزرسانی شده این‌گونه تعریف می‌شوند:

$$\theta_{i+1}^b = \theta_i^b - \alpha \nabla_{\theta} L_{S_b}(f_{\theta_i^b}) \quad (1)$$

¹ Cross-entropy

مرحله از رویکردهای یادگیری متا مبتنی بر گرادیان مانند MAML اصلی، روی تعداد کمی از نمونه های داده اعمال شود. علاوه بر این، در بهینه سازی مبتنی بر گرادیان، مشتقات تابع هزینه باتوجه به وزن ها، فقط فاصله بین خروجی فعلی و هدف متناظر را در نظر می گیرند و اطلاعات تاریخیچه در به روزرسانی وزن هیچ کاربردی ندارند.

برای اولین بار، در پژوهش [۲۵] استفاده از الگوریتم NDEKF برای برآورد وزن شبکه عصبی پیشنهاد شد. NDEKF در مقایسه با الگوریتم EKF به میزان حافظه و زمان محاسبات کمتری نیاز دارد. محققان در آن پژوهش پیشنهاد کردند وزن ها با گره جدا شوند، به طوری که فقط وابستگی متقابل وزن هایی در نظر گرفته شود که به همان گره تغذیه می شوند. در نسخه "جداشده از گره"، از فیلتر کالمن گسترش یافته به طور مستقل بر روی هر نورون شبکه عصبی برای تخمین وزن های بهینه تغذیه شده روی آن نورون استفاده می شود. به این ترتیب، این الگوریتم فقط وابستگی های متقابل محلی را مد نظر قرار می دهد.

ما در اینجا، با اختصار معادلات توصیف کننده این فیلتر را بر اساس کار انجام شده در [۴۶] فهرست می کنیم. معادلات موجود در پژوهش آنها بر روی شبکه ای با عمق سه لایه و تعداد کمی نورون در هر لایه اعمال شده است به طوری که تعداد آنها از پنج لایه تجاوز نمی کند. باین وجود، با علم به این واقعیت که امروزه رایانه ها قادر به انجام محاسباتی هستند که تا پیش از این انجام آن ممکن نبود، می توان این معادلات را به هر شبکه ای با هر عمق و با هر تعداد نورونی تعمیم داد.

الگوریتم EKF، برای تکرار t ، بر مبنای نمودی از متغیر حالت به شکل زیر استوار است:

$$\theta_j(t+1) = \theta_j(t) \quad (5)$$

$$y_j(t) = h(\theta_j(t).u_j(t)) + v_j(t) \quad (6)$$

که در آن، M تعداد نورون ها در یک لایه خاص است (ما در اینجا نورون ها را بر اساس لایه ها گروه بندی می کنیم). به علاوه، z یک نورون خاص را تعریف می کند. $(1 \leq j \leq M)$ ؛ θ_j بردار وزن متناظر با نورون z ؛

u_j نشان دهنده یک بردار ورودی $1 \times n_u$ برای نورون j ام؛

y_i نشان دهنده بردار هدف برای نورون خروجی j ام در زمان t ؛ $(n_y \times 1)$

$h(\cdot)$ تابع غیر خطی $\theta(\cdot)$ و $u(\cdot)$ و v_j بردار نویز اندازه گیری هستند.

examples

6: Compute adapted parameters with N step of gradient descent:

$$\theta_{i+1}^b = \theta_i^b - \alpha \nabla_{\theta} \mathcal{L}_{s_b}(f_{\theta_i^b}) \quad (4)$$

7: end for

$$8: \text{Update } \theta_0 = \theta_0 - \beta \nabla_{\theta} \sum_{b=1}^B \mathcal{L}_{Q_b}(f_{\theta_0^b}(\theta_{i+1}))$$

9: end while

استفاده از بهینه ساز SGD حلقه درونی MAML برای یادگیری وزن های سریعی که امکان تطبیق سریع با یک وظیفه جدید را فراهم می کند، بسیار ساده است. هر چند، پس انتشار از طریق محاسبات گرادیان متشکل از حلقه درونی و به روزرسانی پارامترها در حلقه بیرونی، باعث شده است آموزش MAML به امری دشوار تبدیل شود و با مشکلاتی مثل محوشدگی یا انفجار گرادیان و یا گیر کردن در کمینه های محلی دست و پنجه نرم کند. به علاوه، همین ساختار پیچیده MAML باعث می شود انتخاب متا پارامترهای مناسب باتوجه به مجموعه داده ها و سیستم های مختلف منجر به افزایش ناپایداری آنها شود.

ایده کاهش وابستگی به عملیات گرادیان در داخل حلقه درونی با جایگزینی بهینه ساز SGD با بهینه ساز NDEKF از همین جا نشئت گرفت؛ بنابراین، ما وزن های سریع در حلقه درونی را با استفاده از NDEKF محاسبه می کنیم.

البته ایده استفاده از NDEKF برای آموزش شبکه های عصبی، اتفاق جدیدی نیست. همان طور که در کارهای انجام شده پیشین مشاهده کردیم، محاسبه پارامترهای شبکه با استفاده از EKF به لحاظ قابلیت تعمیم و سرعت، عملکرد به مراتب بهتری نسبت به SGD دارد، به خصوص زمانی که اندازه مجموعه داده آموزشی، مانند مورد ما، در حلقه درونی MAML کوچک باشد. هر چند، به دلیل پیچیدگی محاسبه ماتریس ژاکوبین در فیلتر کالمن، به ویژه با وجود اندازه بزرگ پارامترها در شبکه های عمیق، ماتریس های حاصل، به شدت غیر قابل حل خواهند شد. توصیه می شود برای کاهش پیچیدگی محاسبات و هزینه های محاسباتی، از یک نسخه "جداشده از گره" از این الگوریتم استفاده شود. ما نیز به همین دلیل الگوریتم NDEKF را به عنوان راه حلی برای برآورد وزن های سریع در حلقه درونی مورد استفاده قرار دادیم.

۳-۲. فرمول بندی MAML مبتنی بر NDEKF

بهینه سازی وزن های شبکه عصبی با استفاده از الگوریتم های گرادیان نزولی منجر به کند شدن سرعت همگرایی و ناپایداری آن شده و با مشکل بیش برآزش مواجه است، به ویژه زمانی که در دو

$$H_j(t) = \begin{bmatrix} \frac{\partial y_1}{\partial \theta_{1j}} & \frac{\partial y_1}{\partial \theta_{2j}} & \dots & \frac{\partial y_1}{\partial \theta_{n_j}} \\ \frac{\partial y_2}{\partial \theta_{1j}} & \frac{\partial y_2}{\partial \theta_{2j}} & \dots & \frac{\partial y_2}{\partial \theta_{n_j}} \\ \vdots & \vdots & \dots & \vdots \\ \frac{\partial y_{n_y}}{\partial \theta_{1j}} & \frac{\partial y_{n_y}}{\partial \theta_{2j}} & \dots & \frac{\partial y_{n_y}}{\partial \theta_{n_j}} \end{bmatrix} \quad (13)$$

ما با بهروزرسانی پارامترها در حلقه درونی با استفاده از چند مرحله DEKF به جای گرادیان نزولی، DEKF را با MAML ترکیب می‌کنیم. هرچند، بهروزرسانی پارامترها در حلقه بیرونی (بهینه‌سازی متا^۳) همانند همان MAML اصلی با استفاده از گرادیان نزولی باقی می‌ماند. برای انجام این کار، ما معادله (۴) در MAML اصلی را با معادلات (۹) تا (۱۲) جایگزین و تغییرات مربوطه را براین اساس اعمال می‌کنیم. الگوریتم MAML با استفاده از DEKF برای یک تکرار در حلقه درونی به شرح زیر خواهد بود (به شکل (۳) رجوع کنید):

الگوریتم (۳). الگوریتم پیشنهادی MAML- NDEKF
Require: $p(T)$: distribution over tasks Require: β: step size hyperparameter Require: μ_p, μ_q, μ_r: Kalman filter hyperparameters Require: set $P_i(0) = \left(\frac{1}{\mu_p}\right)I$, $Q = \mu_q I$, $R = \mu_r I$ 1: randomly initialize θ where $\theta_0 = \theta$ 2: while not done do 3: Sample batch of tasks $T_b = \{S_b, Q_b\} \sim p(T)$ 4: for all S_b do 5: Compute adapted parameters with N step of DEKF: 6: for depth j of neural network $\hat{\theta}_j(t+1) = \theta_j(t) + K_j(t) \cdot \varepsilon$ $K_j(t) = P_j(t) \cdot H_j^T(t) \cdot A(t)$ $A(t) = \left[R(t) + \sum_{j=1}^M H_j(t) \cdot P_j(t) \cdot H_j^T(t) \right]^{-1}$ $P_j(t+1) = P_j(t) - K_j(t) \cdot H_j(t) \cdot P_j(t) + Q(t)$ 7: end for 8: set $\hat{\theta}_i = [\hat{\theta}_1^i, \hat{\theta}_2^i, \dots, \hat{\theta}_J^i]$ where $0 \leq j \leq J$, J is a depth of neural network 10: end for 11: Update $\theta_0 = \theta_0 - \beta \nabla_{\theta} \sum_{b=1}^B \mathcal{L}_{Q_b} \left(f_{\theta_0^b}(\hat{\theta}_i) \right)$ 12: end while

خطای میانگین مربعات (MSE)^۱ متناظر با گره ز معمولاً با معادله زیر محاسبه می‌شود:

$$\varepsilon = \frac{1}{2} \sum_{j=1}^{n_y} (y_j(t) - \hat{y}_j(t))^2 \quad (7)$$

در مورد وظایف طبقه‌بندی، می‌توان از خطای آنتروپی متقاطع بین خروجی‌های مطلوب و واقعی نیز استفاده کرد:

$$\varepsilon = \frac{1}{2} \sum_{j=1}^{n_y} (y_j(t) - \hat{y}_j(t))^2 \quad (8)$$

معادلات بهروزرسانی وزن از الگوریتم بازگشتی کالمن که تابع تلفات (معادله ۵ یا ۶) را به حداقل می‌رساند را این‌گونه نوشت:

$$\hat{\theta}_j(t+1) = \theta_j(t) + K_j(t) \cdot \varepsilon \quad (9)$$

$$K_j(t) = P_j(t) \cdot H_j^T(t) \cdot A(t) \quad (10)$$

$$A(t) = \left[R(t) + \sum_{j=1}^M H_j(t) \cdot P_j(t) \cdot H_j^T(t) \right]^{-1} \quad (11)$$

$$P_j(t+1) = P_j(t) - K_j(t) \cdot H_j(t) \cdot P_j(t) + Q(t) \quad (12)$$

$$P_j(0) = \left(\frac{1}{\mu_p}\right)I, Q = \mu_q I, R = \mu_r I$$

به‌طوری‌که n_j نشان‌دهنده r وزن‌های وارده بر هر نورون j و n_y بعد بردار خروجی،

$$\hat{x}_j(\cdot) \text{ مقدار تخمینی } x_j(\cdot);$$

$P_j(\cdot)$ ماتریس کوواریانس خطای شرطی تقریبی $n_{x_j} \times n_{x_j}$ در $\theta_j(\cdot)$ ؛

$R(\cdot)$ ماتریس کوواریانس نویز اندازه‌گیری $(n_y \times n_y)$ موثر بر نرخ همگرایی؛

$$K_j(\cdot) \text{ ماتریس بهره }^2 \text{ فیلتر کالمن } (n_{x_j} \times n_y)$$

$$\varepsilon \text{ خطای بین خروجی‌های مطلوب و واقعی } 1 \times n_y$$

$Q(\cdot)$ ماتریس کوواریانس مورب نویز فرایند $n_{x_j} \times n_{x_j}$ $A(t)$ ماتریسی که یک بار برای هر تکرار محاسبه می‌شود و در همه بهروزرسانی‌های گره‌های جداسده استفاده می‌شود؛

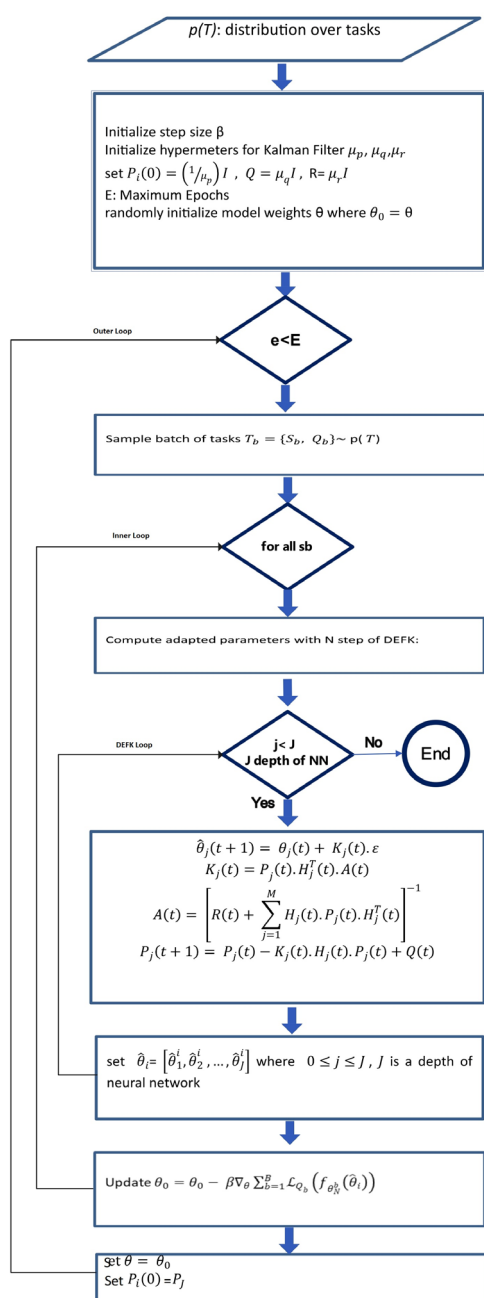
μ_p, μ_q, μ_r هایپرپارامترهای فیلتر که با استفاده از روش آزمون و خطا مقدار دهی اولیه می‌شوند؛

$H_j(\cdot)$ و ماتریس ژاکوبین $n_y \times n_{x_j}$ هستند. این ماتریس همه مشتقات نسبی تابع بردار تعریف‌کننده خروجی y شبکه عصبی متصل به هر وزن وارده بر نورون j را در برمی‌گیرد.

³ Meta-optimization

¹ Mean-squared error

² Gain matrix



شکل (۳). نمودار جریان الگوریتم پیشنهادی

۴. آزمایش ها و نتایج

ما در این بخش، عملکرد رویکرد یادگیری عمیق داده محدود پیشنهادی خود را به لحاظ تشخیص نفوذ بهبود یافته مورد ارزیابی قرار می دهیم و سپس آن را با رویکرد اصلی MAML مقایسه می کنیم. برای این منظور، معیار تشخیص نفوذ پرکاربرد NSL-KDD را پیاده سازی خواهیم کرد.

۴-۱. دادگان NSL-KDD

بررسی الگوریتم بالا مشخص می کند که اجرای الگوریتم DEKF نسبت به کاربرد الگوریتم گرادیان نزولی مستلزم صرف وقت بیشتری است. دلیل این امر آن است که در هر مرحله زمانی، علاوه بر محاسبه مشتقات، محاسبات اضافی بیشتری مانند محاسبه معکوس ماتریس مربع^۱ در معادله (۱۱) وجود دارد.

ما به دلیل پیچیدگی، این الگوریتم را در شبکه های غیر کانولوشن (غیر پیچشی)^۲ اعمال می کنیم. از آنجایی که ما این الگوریتم را در دادگان NSL-KDD اعمال می کنیم که یک پایگاه داده تصویری نیست، عملکرد آن تحت تأثیر قرار نخواهد گرفت. برای این منظور از شبکه ای با ۴ لایه مخفی با اندازه یکسان ۳۲ استفاده می کنیم که هر یک الگوریتم batch normalization و غیر خطی های ReLU به علاوه، یک لایه خطی و softmax را شامل می شوند. ما در مدل خود از آنتروپی متقاطع استفاده کردیم.

^۱ Square matrix

^۲ Non-convolutional network

مجموعه داده NSL-KDD از ۴۱ ویژگی، شامل ۳ ویژگی تفکیکی یا رسته‌ای^۶ و ۳۸ ویژگی مقادیر پیوسته^۷ تشکیل شده است.

انتقال داده: در یادگیری ماشین، نمی‌توان به طور مستقیم از ویژگی‌های رسته‌ای (تفکیکی) استفاده کرد. بلکه برای تبدیل آنها به چند روش رمزگذاری نیاز داریم. روش کدبندی وان‌هات^۸ یکی از متداول‌ترین تکنیک‌های مورد استفاده برای نام‌گذاری ویژگی‌های رسته‌ای است. هر ویژگی از نوع کاراکتر به یک بردار باینری^۹ تبدیل می‌شود که در آن هر رسته متناظر با عنوان ۱ یا ۰ برچسب‌گذاری می‌شود. برای مثال، ویژگی protocol_type در مجموع از سه مشخصه tcp، udp و icmp برخوردار است. با استفاده از روش وان‌هات، ویژگی tcp به صورت (1, 0, 0)، udp به صورت (0, 1, 0) و icmp به صورت (0, 0, 1) کدبندی می‌شود. به‌طور کلی، سه ویژگی از نوع کاراکتر: protocol_type، service و flag، در مقادیر باینری ۸۴ بعدی نگاشت می‌شوند. در غیر این صورت، مقدار ویژگی num_outbound_cmds برابر صفر است و به همین جهت، این ویژگی حذف می‌شود؛ بنابراین، NSL-KDD ۴۱ بُعدی اصلی را می‌توان به یک مجموعه داده ۱۲۱ بُعدی جدید تبدیل کرد.

نرمال‌سازی داده‌ها^{۱۰}: دادگان تبدیل‌شده NSL-KDD از ویژگی‌های ۱۲۱ بُعدی با چندین تفاوت در میان آنها برخوردار است. بنابراین، ما از تکنیک نرمال‌سازی داده‌ها برای به حداقل رساندن این تفاوت‌ها به منظور بهبود عملکرد سیستم استفاده می‌کنیم. در این مقاله، روش نرمال‌سازی min-max (معادله (۱۴)) برای کاهش تفاوت‌ها در همه ابعاد و نرمال کردن مقادیر در بازه مشخص [۰، ۱] مورد استفاده قرار گرفته است.

$$Normalized_x = \frac{x - x_{min}}{x_{max} - x_{min}} \quad (14)$$

۳-۴. تقسیم دادگان NSL-KDD برای یادگیری متا

ما مدل‌های داده محدود^{۱۱} خود را بر روی توزیع وظیفه FSL ($p(T)$) ارزیابی می‌کنیم. هر وظیفه FSL در توزیع به‌خودی‌خود، به‌عنوان یک مجموعه داده متشکل از دو مجموعه در نظر گرفته می‌شود که مورد اول، مجموعه support و مورد دوم، مجموعه query نامیده می‌شود. مجموعه support برای دربرگرفتن N کلاس و K نمونه از هر یک به کار می‌رود، درحالی‌که مجموعه query، تعمیم مدل به نمونه‌های دیده نشده را حفظ می‌کند؛ بنابراین، مجموعه‌ای متفاوت از نمونه‌های گرفته‌شده از همان N

دادگان NSL-KDD [۴۷] یک معیار کارآمد برای تشخیص نفوذ مبتنی بر یادگیری ماشین است. هرچند این دادگان نیز فاقد ایراد نبوده و با مشکل عدم توازن کلاس مواجه است. به‌عنوان مثال، در دادگان آموزشی NSL-KDD، تنها ۰.۰۴٪ از نمونه‌ها به نوع حمله U2R تعلق دارند که باعث می‌شود به‌شدت کمتر از حد نشان داده شود. این وضعیت در مورد انواع حمله R2L و حمله پروب شبکه^۱ نیز صادق است، درحالی‌که اکثر سوابق حمله، نشان‌دهنده نوع حمله DDOS هستند.

دسته‌های موجود در مجموعه داده NSL-KDD به پنج دسته اصلی گروه‌بندی می‌شوند:

۱. عادی (Normal)
۲. انکار سرویس^۲ (DoS): این حمله با جلوگیری از درخواست‌های قانونی برای دسترسی به منابع شبکه از طریق مصرف پهنای باند یا بارگذاری اضافی بر روی منابع محاسباتی^۳ انجام می‌شود.
۳. پویشی (Prob): حمله پویشی، دسته‌ای از حملاتی است که در آن مهاجم قبل از شروع حمله، شبکه را بررسی می‌کند تا اطلاعات سیستم هدف را جمع‌آوری کند.
۴. کاربر به ریشه (U2R): در این حالت، مهاجم با دسترسی به یک حساب کاربری عادی در سیستم اقدام می‌کند و می‌تواند از آسیب‌پذیری‌های سیستم برای دسترسی ریشه-ای^۴ به سیستم استفاده کند.
۵. ریشه‌ای به محلی (R2L): مهاجمی که حساب کاربری در یک ماشین راه دور^۵ ندارد، بسته‌ای را از طریق شبکه به آن ماشین ارسال می‌کند و از آسیب‌پذیری‌های سیستم برای دسترسی محلی به‌عنوان کاربر آن ماشین استفاده می‌کند.

همین واقعیت، باعث شده است تشخیص انواع حملات کمتر نمایش داده‌شده برای طبقه‌بندی‌کننده‌های ML معمولی دشوار شود نوبه خود، منجر به کاهش دقت آن می‌شود. ما برای بهبود دقت در تشخیص انواع حملات نادر، به‌غیر از رویکردهای مرسوم رفع مشکل عدم توازن، الگوریتم پرتفردار داده محدود یادگیری متای MAML را به کار گرفتیم و آن را با استفاده از الگوریتم DEKF بهینه‌سازی می‌کنیم. اما پیش از آموزش مدل خود با استفاده از این مجموعه داده، مراحل مختلف استاندارد پیش‌پردازش را انجام داده، سپس آن را برای وظایف یادگیری متا تقسیم می‌کنیم.

۲-۴. پردازش NSL-KDD

⁶ Categorical features

⁷ Continuous value features

⁸ One-hot-encoding

⁹ Binary vector

¹⁰ Data Normalization

¹¹ Few-Shot

¹ Probe attack

² Denial of Service

³ computational resources

⁴ root access

⁵ remote machine

متای خود استفاده خواهیم کرد.

۴-۴. تنظیم^۳ مجموعه متا

ما ۴۰ زیرکلاس موجود را به صورت تصادفی به نسبت ۱۰/۱۰/۲۰ به مجموعه داده های متا آموزش، اعتبارسنجی و آزمون تقسیم کردیم. هرچند، معیار تقسیم ما بر پایه گنجاندن نوع ترافیک نرمال در مجموعه متا آموزشی استوار است. علاوه بر این، ما زیرکلاس های مختلفی را از هر چهار حمله در هر مجموعه متا آموزش، اعتبارسنجی و آزمون انتخاب می کنیم و در عین حال، شرایطی را فراهم می سازیم که بین این مجموعه های متا هیچگونه همپوشانی ایجاد نشود. ما در آزمایش خود، به طور تصادفی ۵۰ نمونه از هر زیرکلاس را نمونه گیری کردیم.

لازم به ذکر است که ما رویکرد MAML و مدل KF-MAML پیشنهادی خود را در ۱۰ شات ۵ طرفه با اندازه دسته وظایف ۲۵ ارزیابی می کنیم.

۴-۵. تنظیم شبکه عصبی و هایپر پارامترها^۴

ما برای ارزیابی مدل های MAML و KF-MAML، از یک شبکه کاملاً متصل متشکل از ۴ لایه پنهان هر یک با اندازه ۳۲ استفاده می کنیم. هر لایه الگوریتم batch normalization (نرمال سازی دسته ای)، تابع فعال سازی ReLU برای غیرخطی ها و همچنین یک لایه خطی و یک softmax را شامل می شود.

مدل MAML با استفاده از بهینه ساز SGD با نرخ یادگیری ۰،۰۰۰۲، نرخ یادگیری درونی ۰،۰۱ برای فرایند انطباق و اندازه گام ۴ (یعنی تعداد گام های گرادیان) آموزش داده می شود.

مدل KF-MAML نیز با استفاده از بهینه ساز DEKF متشکل از ۴ مرحله برای تطبیق حلقه درونی، نرخ یادگیری متا ۰،۰۰۰۲ و رابطه $\mu_p = 100$ و $\mu_r = 1$ و $\mu_q = 10e^{-6}$ برای هایپر پارامترهای فیلتر کالمن آموزش داده می شود.

همان طور که در شکل (۴) نشان داده شده است، هنگام استفاده از مراحل DEKF در فرایند به روزرسانی حلقه درونی MAML به جای گرادیان نزولی در مجموعه داده IDS، پس از چند گام زمانی، به نتایج بسیار خوبی دست یافتیم. استفاده از ترکیب DEKF-MAML برای به روزرسانی وزن ها در حلقه درونی، باعث ارتقای سطح دقت مدل و افزایش سرعت همگرایی آن می شود.

کلاس مجموعه support خواهد بود. همانند یادگیری ماشین با نظارت^۱ رایج، ما باید مجموعه داده های خود را به مجموعه های آموزشی، اعتبارسنجی و آزمون تفکیک کنیم. در مورد یادگیری متا، تفاوت در این است که برای جلوگیری از همپوشانی کلاس ها و سازماندهی هر مجموعه در وظایف FSL (اپیزودهای مختلف)، به جای نمونه ها، رده ها به مجموعه های آموزش متا، اعتبارسنجی متا و آزمون متا تقسیم می شوند. برخلاف سایر حوزه ها مانند طبقه بندی تصویر که در آن معیارهای مختلف FSL مانند Mini-Image Net [۳۳] یا Omniglot [۳۲] وجود دارد، در مورد مجموعه داده NSL-KDD، هیچ مجموعه متای کلی برای آموزش با داده محدود نداریم؛ بنابراین، ما مجموعه متا NSL-KDD FSL خود را می سازیم. مجموعه داده NSL-KDD به نوبه خود از مجموعه داده های KddTrain+ و KddTest+ تشکیل شده است که محققان برای ساخت مدل های یادگیر ماشین و مقایسه عملکرد آنها، به طور گسترده از آنها استفاده می کنند. جدول (۲)، توزیع نمونه برای هر کلاس پایه NSL-Kdd در این مجموعه داده ها را نشان می دهد.

علی رغم وجود انواع مختلف حملات سایبری در داده های آموزشی، توزیع آنها چولگی^۲ زیادی دارد. گفتنی است که ترافیک نرمال شامل نیمی از تعداد نمونه ها در دو مجموعه داده است، درحالی که توزیع U2R و R2L بسیار اندک است. با این حال، TestKdd+ تعداد بیشتری از نمونه های این دو حمله را نسبت به TrainKdd+ شامل می شود، به علاوه اینکه زیر کلاس های جدید (جدول (۳)) این حملات که در TrainKdd+ وجود ندارند را نیز در برمی گیرد. همان طور که قبلاً اشاره کردیم، در مسائل FSL، مجموعه متا به داده های آموزشی زیادی نیاز ندارد و شامل وظایفی است که یک مجموعه support و یک مجموعه query را شامل می شود. به همین دلیل، ما برای نمونه گیری مجدد وظایف، TrainKdd+ و TrainKdd+ را با هم ادغام می کنیم و مجموعه متا IDS خود را می سازیم. از طرف دیگر، با این ترکیب می توانیم به بازمود خوبی از انواع حملات دست یافته و از بروز بیش برآزش جلوگیری کنیم.

هر نوع حمله مورد اشاره در جدول (۲)، زیر کلاس های مختلفی از تهدیدات را در برمی گیرد که در جدول (۳) نشان داده شده است. ما به منظور ارزیابی مدل های FSL باید داده های خود را بر اساس دسته ها به مجموعه های متای آموزش، اعتبارسنجی و آزمون تقسیم کنیم. علاوه بر ترافیک نرمال، ۳۹ زیر کلاس در داده های آموزشی ترکیبی خواهیم داشت که نتیجه آن ۴۰ زیر کلاس کلی است. بنابراین، ما از ۴۰ دسته برای ساخت مجموعه

³ Setup

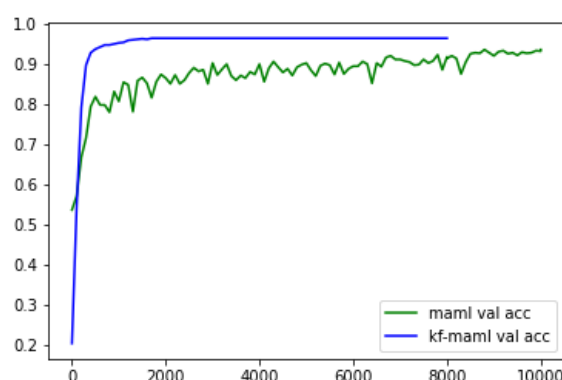
⁴ Hyper Parameters

¹ Supervised machine learning

² Skewed

بهینه‌سازی شبکه‌های دارای چندلایه پنهان با تعداد نسبتاً کمی گره انجام نشده بود. علاوه بر این، باتوجه به معماری شبکه کاملاً متصل ما با اندازه لایه‌های پنهان ۳۲، هنگام بهینه‌سازی MAML با استفاده از DEKF، سرعت همگرایی و عملکرد مدل به شکل قابل ملاحظه‌ای بهبود یافت. از سوی دیگر، رویکرد بهینه‌سازی ما در عمل غیرقابل اجرا شده و هنگام اعمال آن در یک شبکه کانولوشن یا زمانی که اندازه هر لایه پنهان ۶۴ باشد، عملکرد آن تنزل پیدا می‌کند. دلیل این امر الزامات محاسباتی زیاد برای محاسبه ماتریس‌ها در فیلتر کالمن است.

برای این که کارایی مدل IDS پیشنهادی نشان داده شود، عملکرد ۵ دسته آن با پیشرفته‌ترین IDS‌ها مقایسه شد. جدول (۵) نتایجی را که توسط مدل پیشنهادی به دست آمده است، در مقایسه با سایر مدل‌هایی که بر روی مجموعه داده NSL-KDD مقایسه شده‌اند، از نظر نرخ تشخیص هر دسته و دقت کلی ارائه می‌دهد.



شکل (۴). ارزیابی عملکرد مدل اصلی MAML و مدل KF-MAML پیشنهادی ما در مجموعه آزمون متای NSL-Kdd

از سوی دیگر، بر اساس جدول (۴) می‌بینیم که این میزان بهبودی، نه تنها در دقت کلی مدل، بلکه در دقت شناسایی دو حمله نادر R2L و R2U نیز اتفاق افتاده است.

در مطالعه حاضر، ما موفق شدیم فیلتر کالمن را برای بهینه‌سازی یک شبکه عصبی با اندازه نسبتاً بزرگ مورد استفاده قرار دهیم. این کار در مطالعات قبلی مربوط به کاربرد فیلترها در

جدول (۲). توزیع نمونه برای هر کلاس پایه NSL-Kdd در مجموعه داده‌های آموزش و آزمون

Dataset	Number of Samples					
	Total	Normal	DoS	Probe	U2R	R2L
TrainKdd+	125973	67343 (53%)	45927 (37%)	11656 (9.11%)	52 (0.04%)	955 (0.85%)
TestKdd+	22544	9711 (43%)	7458 (33%)	2421 (11%)	200 (0.9%)	2654 (12.1%)
Total NSL-Kdd	148517	77054 (52%)	53385 (36%)	14077 (9.5%)	252 (0.2%)	3609 (2.3%)

جدول (۳). زیرکلاس‌های هر حمله در مجموعه داده NSL-Kdd

Attack Types	DoS	Probe	U2R	L2R
Subclasses	apache2, back, land, neptune, mailbomb, pod, processtable, smarf, treadrop, udpstomp, worm	Ipsweep, mscan, nmap, portsweep, saint, satan	Buffer_overflow, loadmodule, peri, ps, rootkit, sqlattack, xterm	ftp_write, quess_passwd, httptunnel, imap, multihop, named, phf, sendmail, Snmpgetattack, spy, snmpguess, warezmaster, xlock, xsnoop
Total	11	5	7	15

جدول (۴). مقایسه دقت مدل‌های MAML و DEKF-MAML در مجموعه آزمون متای IDS

Approach	Accuracy	DoS	Probe	R2L	U2R	Normal
MAML	0.936	0.98	0.821	0.61	0.51	0.971
DEKF-MAML	0.953	0.951	0.85	0.75	0.662	0.962

جدول (۵). مقایسه بین روش‌های کارهای مرتبط و روش پیشنهادی با استفاده از NSL-KDD برای دسته‌بندی چنددسته‌ای

Approach	Normal	DoS	Probe	U2R	R2L	Accuracy
multi-level hybrid SVM and ELM [26]	98.13	99.54	87.22	21.93	31.39	95.75
S-NDAE [27]	97.73	94.58	72.97	2.70	3.82	85.42
Random forests (RF) algorithm with undersampling and	—	—	—	—	—	94.2

oversampling [29]						
ICVAE-DNN [31]	97.26	85.65	74.97	11.00	44.41	85.97
Original MAML	93.60	98.00	82.10	51.00	61.00	97.10
Proposed DEKF-MAML	96.20	95.10	85.00	66.200	75.00	95.30

جدا شده (DEKF) را در عملکرد درونی MAML ادغام کردیم تا جایگزین گرادیان نزولی سنتی شود. این رویکرد ترکیبی که MAML مبتنی بر DEKF نام دارد، با پیشرفت های قابل توجهی همراه شده و سرعت همگرایی و قابلیت های تعمیم مدل را به میزان زیادی بهبود داده است.

ما با انجام ارزیابی دقیق مجموعه داده NSL-KDD، اثربخشی رویکرد خود در زمینه بهبود تشخیص حملات نادر در سیستم های تشخیص نفوذ را به اثبات رساندیم. نتایج به دست آمده، بر کارایی این هم افزایی میان آموزش با داده محدود و فیلتر کالمن به عنوان ابزاری دقیق برای تقویت قابلیت های IDS، حتی در مواجهه با تهدیدات سایبری در حال تکامل تأکید دارند.

با استفاده از قابلیت های مدل MAML مبتنی بر DEKF، ما گام مهمی در راستای افزایش امنیت شبکه و محافظت از دارایی های اطلاعاتی مهم در برابر تهدیدات نوظهور و پیش بینی نشده برداشته ایم.

باتوجه به نتایج اخذ شده و مطالعات تحقیقات این حوزه، روش ها و راه حل های امیدوار کننده ای وجود دارد که در آینده می توان از آنها بهره جست.

یکی از موضوعات اساسی، بهینه سازی رویکرد DEKF-MAML برای مقیاس پذیری با تمرکز بر کاهش هزینه های محاسباتی و اطمینان از عملکرد درست و الگوریتم شبکه های بزرگ مقیاس است. توانایی سیستم تشخیص نفوذ ما در مواجهه با گستره روزافزون شبکه ها نیز عملکرد درستی خواهد داشت.

چالش های مقیاس پذیری و کارایی اختلالی در عملکرد سیستم ما نداشته، بلکه آن را در زیرساخت های شبکه پیچیده سازمان های مختلف قابل دسترس تر و کاربردی تر می کند.

۶. مراجع

- [1] Mazloum, H. Bigdeli, "An Optimized Compound Deep Neural Network Integrating With Feature Selection for Intrusion Detection System in Cyber Attacks," *Electronic and Cyber Defense*, vol. 10, no. 4, pp. 41-51, 2023 (In Persian). <https://dorl.net/dor/20.1001.1.23224347.1401.10.4.5.5>
- [2] M. Hassan Nataj Solhdar, "Investigation of a new ensemble method of intrusion detection system on different data sets," *Electronic and Cyber Defense*, vol. 10, no. 3, pp. 43-57, 2022 (In Persian).

روش پیشنهادی از نظر تشخیص نادرترین حملات U2R و R2L دارای افزایش دقت قابل توجهی است. همچنین می توان مشاهده نمود که باوجود افزایش دقت در این دو حمله، دقت بر روی کل مجموعه داده تغییر چندانی نداشته است. شایان ذکر است که مدل ما عملکرد قابل تحسینی در سایر دسته ها نیز دارد. باتوجه به اثر بسیار مخرب این دو حمله به زیرساخت سیستم ها، اهمیت این دو بسیار زیاد است، اما متأسفانه روش های قبلی نتوانستند به اندازه کافی به آن بپردازند و دقت مناسبی بر روی تشخیص این حملات خطرناک نداشته اند. البته روش ما نه فقط باعث افزایش دقت جدی در این دو حمله شده است، بلکه تقریباً در دسته های باقیمانده نیز دقت پیشین خود را حفظ کرده است. از طریق انتخاب دقیق هایپرپارامترها^۱ و بررسی ساختارهای جایگزین شبکه عصبی عمیق می توان به این هدف دست یافت. در این کار پژوهشی از یک شبکه عصبی نسبتاً ساده متشکل از چهار لایه پنهان (DNN) استفاده شد.

۵. نتیجه گیری

ما در این مقاله، رویکردی نوآورانه در این زمینه تشخیص نفوذ ارائه کردیم که در آن، آموزش داده محدود (FSL) به ویژه چارچوب یادگیری متامدل - آگنوستیک (MAML) با فیلتر کالمن گسترش یافته جدا شده (DEKF) برای بهبود قابلیت های سیستم های تشخیص نفوذ (IDS) با یکدیگر ادغام شده اند.

مدل های سنتی یادگیری ماشین که عمدتاً دقت کلی خوبی دارند، اغلب در تشخیص حملات نادر و غالباً فاجعه بار، ناکارآمد عمل کرده اند. علاوه بر این، وجود کلاس های هم پوشان، انتخاب ویژگی را پیچیده تر کرده و به این چالش می افزاید.

ما برای حل این چالش ها، راه حل جدیدی را معرفی کردیم که از سرعت بالای آموزش با داده محدود با تمرکز ویژه بر مدل MAML بهره می برد. ظرفیت MAML در سازگاری سریع با وظایف جدید، حتی باوجود داده های برچسب گذاری شده محدود، با ماهیت پویای تشخیص نفوذ انطباق کامل دارد.

ما در این مقاله، برای غلبه بر محدودیت های سرعت همگرایی و الزامات محاسباتی مدل MAML، فیلتر کالمن گسترش یافته

¹ hyperparameters

3.5.3

[19] S.S. Haykin, and S.S. Haykin, "Kalman filtering and neural networks," vol. 284: Wiley Online Library, 2001. <https://doi.org/10.1002/0471221546>

[20] B.C. Aissa, and C. Fatima, "Neural Networks Trained with Levenberg-Marquardt-Iterated Extended Kalman Filter for Mobile Robot Trajectory Tracking," Journal of Engineering Science & Technology Review, vol. 10, no. 4, 2017. <https://doi.org/10.25103/jestr.104.23>

[21] I. Yaesh, and N. Grinfeld, "Training without Gradients--A Filtering Approach," arXiv preprint arXiv:2010.04908, 2020. <https://doi.org/10.48550/arXiv.2010.04908>

[22] Y. Ollivier, "Online natural gradient as a Kalman filter," Electronic Journal of Statistics, vol. 12, no. 2: pp. 2930-2961, 2018. <https://doi.org/10.1214/18-EJS1468>

[23] L. Luttmann, and P. Mercorelli, "Comparison of Backpropagation and Kalman Filter-based Training for Neural Networks," 25th International Conference on System Theory, Control and Computing (ICSTCC), 2021. <https://doi.org/10.1109/ICSTCC52150.2021.9607274>

[24] J.A. Pérez-Ortiz, F.A. Gers, D. Eck, and J. Schmidhuber, "Kalman filters improve LSTM network performance in problems unsolvable by traditional recurrent nets," Neural Networks, vol. 16, no. 2: pp. 241-250, 2003. [https://doi.org/10.1016/S0893-6080\(02\)00219-8](https://doi.org/10.1016/S0893-6080(02)00219-8)

[25] G.V. Puskorius, and L.A. Feldkamp, "Decoupled extended Kalman filter training of feedforward layered networks," IJCNN-91-Seattle International Joint Conference on Neural Networks, 1991. <https://doi.org/10.1109/IJCNN.1991.155276>

[26] W.L. Al-Yaseen, Z.A. Othman, and M.Z.A. Nazri, "Multi-level hybrid support vector machine and extreme learning machine based on modified K-means for intrusion detection system," Expert Systems with Applications, vol. 67: pp. 296-303, 2017. <https://doi.org/10.1016/j.eswa.2016.09.041>

[27] N. Shone, T.N. Ngoc, V.D. Phai, and Q. Shi, "A deep learning approach to network intrusion detection," IEEE Transactions on Emerging Topics in Computational Intelligence, vol. 2, no. 1: pp. 41-50, 2018. <https://doi.org/10.1109/TETCI.2017.2772792>

[28] B. Yan, G. Han, M. Sun, and S. Ye, "A novel region adaptive SMOTE algorithm for intrusion detection on imbalanced problem," 3rd IEEE International Conference on Computer and Communications (ICCC), 2017. <https://doi.org/10.1109/CompComm.2017.8322749>

[29] J. Zhang, M. Zulkernine, and A. Haque, "Random-forests-based network intrusion detection systems," IEEE Transactions on Systems, Man, and Cybernetics, Part C (Applications and Reviews), vol. 38, no. 5: pp. 649-659, 2008. <https://doi.org/10.1109/TSMCC.2008.923876>

[30] Y. Wang, Y. Shen, and G. Zhang, "Research on intrusion detection model using ensemble learning methods," 7th IEEE International Conference on Software Engineering and Service Science (ICSESS), 2016. <https://doi.org/10.1109/ICSESS.2016.7883100>

[31] Y. Yang, K. Zheng, C. Wu, and Y. Yang, "Improving the classification effectiveness of intrusion detection by using improved conditional variational autoencoder and deep neural network," Sensors, vol. 19, no. 11, p.2528, 2019. <https://doi.org/10.3390/s19112528>

[32] B.M. Lake, R. Salakhutdinov, and J.B. Tenenbaum, "Human-level concept learning through probabilistic program induction," Science, vol. 350, no. 6266: pp. 1332-1338, 2015. <https://doi.org/10.1126/science.aab3050>

[Persian\). https://doi.org/10.1109/TKDE.2008.239](https://doi.org/10.1109/TKDE.2008.239)

[3] H. He, and E.A. Garcia, "Learning from imbalanced data," IEEE Transactions on knowledge and data engineering, vol. 21, no. 9. pp. 1263-1284. 2009. <https://doi.org/10.1109/TKDE.2008.239>

[4] K.-C. Khor, C.-Y. Ting, and S. Phon-Amnuaisuk, "The effectiveness of sampling methods for the imbalanced network intrusion detection data set," Recent Advances on Soft Computing and Data Mining, Springer. pp. 613-622, 2014. https://doi.org/10.1007/978-3-319-07692-8_58

[5] M. Ring, S. Wunderlich, D. Scheuring, D. Landes, and A. Hotho, "A survey of network-based intrusion detection data sets," Computers & Security, vol. 86. pp.147-167. 2019. <https://doi.org/10.1016/j.cose.2019.06.005>

[6] S. Visa, and A. Ralescu. "Issues in mining imbalanced data sets-a review paper," in Proceedings of the sixteen midwest artificial intelligence and cognitive science conference, vol. 2005, pp. 67-73. sn. 2005.

[7] Y. Liu, T. Gao, and H. Yang, "Selectnet: Learning to sample from the wild for imbalanced data training," Mathematical and Scientific Machine Learning. PMLR. Vol. 107. pp. 193-206. 2020.

[8] H. Liu, and B. Lang, "Machine learning and deep learning methods for intrusion detection systems: A survey," Applied Sciences, vol. 9, no. 20: p. 4396, 2019. <https://doi.org/10.3390/app9204396>

[9] G.S. Dhillon, P. Chaudhari, A. Ravichandran, and S. Soatto, "A baseline for few-shot image classification," in Proceedings of the International Conference on Learning Representations (ICLR), 2020.

[10] C. Finn, P. Abbeel, and S. Levine. "Model-agnostic meta-learning for fast adaptation of deep networks," in Proceedings of 34th International conference on machine learning. vol. 70. pp 1126-1135. 2017. <https://dl.acm.org/doi/10.5555/3305381.3305498>

[11] G. Koch, R. Zemel, and R. Salakhutdinov, "Siamese neural networks for one-shot image recognition," ICML deep learning workshop. 2015.

[12] H. Tanha, M. Abbasi, "Identify malicious traffic on IoT infrastructure using neural networks and deep learning," Electronic & Cyber Defence, vol. 11, no. 2, 2023 (In Persian). <https://dori.net/dor/20.1001.1.23224347.1402.11.2.1.4>

[13] A. A. Tajari Siahmarzkooh, "Intrusion Detection in Computer Networks Using Decision Tree and Feature Reduction," Electronic & Cyber Defence, vol. 9, no. 3, 2021 (In Persian). <https://dori.net/dor/20.1001.1.23224347.1400.9.3.8.9>

[14] S. Thrun, and L. Pratt, "Learning to learn: Introduction and overview," in Learning to learn, Springer. pp. 3-17, 1998. https://doi.org/10.1007/978-1-4615-5529-2_1

[15] A. Rajeswaran, C. Finn, S.M. Kakade, and S. Levine, "Meta-learning with implicit gradients," Advances in neural information processing systems, vol. 32, 2019.

[16] A. Nichol, J. Achiam, and J. Schulman, "On first-order meta-learning algorithms," arXiv preprint arXiv:1803.02999, 2018. <https://doi.org/10.48550/arXiv.1803.02999>

[17] A. Antoniou, H. Edwards, and A. Storkey, "How to train your MAML," in Proceedings of the International Conference on Learning Representations (ICLR) 2019.

[18] J. Chen, W. Yuan, S. Chen, Z. Hu, and P. Li, "Evo-MAML: Meta-Learning with Evolving Gradient," Electronics, vol. 12, no. 18: p. 3865, 2023. <https://doi.org/10.3390/electronics12183865>

- prediction modeling and efficient data assimilation," *Journal of Computing in Civil Engineering*, vol. 33, no. 3, 2019. [https://doi.org/10.1061/\(ASCE\)CP.1943-5487.0000835](https://doi.org/10.1061/(ASCE)CP.1943-5487.0000835)
- [42] N.M. Vural, S. Ergüt, and S.S. Kozat, "An efficient and effective second-order training algorithm for lstm-based adaptive learning," *IEEE Transactions on Signal Processing*, vol. 69 pp. 2541-2554, 2021. <https://doi.org/10.1109/TSP.2021.3071566>
- [43] M.C. Nechyba, and Y. Xu. "Cascade neural networks with node-decoupled extended Kalman filtering," *IEEE International Symposium on Computational Intelligence in Robotics and Automation CIRA'97. 'Towards New Computational Principles for Robotics and Automation'*, 1997. <https://doi.org/10.1109/CIRA.1997.613860>
- [44] N.M.Vural, F. Ilhan, and S.S. Kozat, "Stability of the Decoupled Extended Kalman Filter Learning Algorithm in LSTM-Based Online Learning," *arXiv preprint arXiv:1911.12258*, 2019. <https://doi.org/10.48550/arXiv.1911.12258>
- [45] S. Feng, X.Li, S. Zhang, Z. Jian, H. Duan, and Z. Wang, "A review: state estimation based on hybrid models of Kalman filter and neural network," *Systems Science & Control Engineering*, vol. 11, no. 1, 2023. <https://doi.org/10.1080/21642583.2023.2173682>
- [46] S. Murtuza, and S. Chorian, "Node decoupled extended Kalman filter based learning algorithm for neural networks," *9th IEEE International Symposium on Intelligent Control*, 1994. <https://doi.org/10.1109/ISIC.1994.367790>
- [47] M. Tavallae, E. Bagheri, W. Lu, and A.A. Ghorbani, "A detailed analysis of the KDD CUP 99 data set," *IEEE symposium on computational intelligence for security and defense applications*, 2009. <https://doi.org/10.1109/CISDA.2009.5356528>
- [33] O. Vinyals, C. Blundell, T. Lillicrap, and D. Wierstra, "Matching networks for one shot learning," *Advances in neural information processing systems*, 2016. <https://dl.acm.org/doi/10.5555/3157382.3157504>
- [34] J. Snell, K. Swersky, and R. Zemel. "Prototypical networks for few-shot learning," *Advances in neural information processing systems*, 2017. <https://dl.acm.org/doi/10.5555/3294996.3295163>
- [35] L.Y. Gui, Y.X. Wang, D. Ramanan, and J.M. Moura, "Few-shot human motion prediction via meta-learning," *Proceedings of the European Conference on Computer Vision (ECCV)*. 2018. https://doi.org/10.1007/978-3-030-01237-3_27
- [36] B.D. Anderson, and J.B. Moore, "Optimal filtering," *Courier Corporation*. 2012. <https://doi.org/10.1109/TSMC.1982.4308806>
- [37] A. Gelb (ed.) , "Applied optimal estimation," *MIT press*. 1974.
- [38] Y. Iiguni, H. Sakai, and H. Tokumaru, "A real-time learning algorithm for a multilayered neural network based on the extended Kalman filter," *IEEE Transactions on Signal processing*, vol. 40, no. 4: pp. 959-966, 1992. <https://doi.org/10.1109/78.127966>
- [39] Y. Shao, F.M. Dietrich, C. Nettelblad, and C. Zhang, "Training algorithm matters for the performance of neural network potential: A case study of Adam and the Kalman filter optimizers," *The Journal of Chemical Physics*, vol. 155, no. 20, 2021. <https://doi.org/10.1063/5.0070931>
- [40] C. Jin, S. Jang, X. Sun, J. Li, and R. Christenson, "Damage detection of a highway bridge under severe temperature changes using extended Kalman filter trained neural network," *Journal of Civil Structural Health Monitoring*, vol. 6, no. 3: pp. 545-560, 2016. <https://doi.org/10.1007/s13349-016-0173-8>
- [41] Q. Li, Z.Y. Wu, and A. Rahman, "Evolutionary deep learning with extended Kalman filter for effective