

Using Linear Discriminant Analysis for Risk Score Estimation of Unified Resource Locators (URLs)

M. Deypir^{1*}, K. Halili²

¹ Associate Professor, Shahid Sattari University of Aeronautical Science and Technology, Tehran, Iran

(*Correspondence: mdeypir@ssau.ac.ir)

² Assistant Professor, Shahid Sattari University of Aeronautical Science and Technology, Tehran, Iran

ARTICLE INFO

Article history:

Article Type: Research paper

Received: ۲۰ March ۲۰۲۰

Revised: ۲۶ April ۲۰۲۰

Accepted: ۰۸ June ۲۰۲۰

Available online: ۱۱ July ۲۰۲۰

*Correspondence:

mdeypir@ssau.ac.ir

Keywords:

Linear Discriminant Analysis

Malicious Urls

Security Risk

Textual Feature

ABSTRACT

Sending a malicious URL to a victim is the starting point for various types of malicious activities by attackers. Therefore, identifying malicious URLs plays a significant role in the security of Internet users. Recently, calculating the security risk of URLs instead of classifying them using machine learning-based models has received attention. This is because, on the one hand, it provides necessary warnings to users and, on the other hand, it does not have the problems of classification models. In this study, a new criterion based on linear discriminant analysis is proposed to estimate the security risk of URLs. In this criterion, previously known examples of normal and malicious URLs are mapped into a new space in which the security risk can be calculated more accurately. Although deep learning is not used in the proposed criterion and it requires little training data, it provides a realistic estimate of the security risk values of malicious and safe URLs. Experiments conducted on real datasets show that the proposed criterion outperforms the previously proposed criteria in terms of detection rate. The implementations carried out in this research, along with the datasets used in the experiments, are publicly available at <https://github.com/mdeypir/LRU>.

Cite this article: Khorasani, Yaghoub[®], Shafiee, Ebrahim[®], Shamsi, Alireza (۲۰۲۰). Using Linear Discriminant Analysis for Risk Score Estimation of Unified Resource Locators (URLs). Journal of Electronic and Cyber Defens. ۲۰۲۰; ۱۳(۲): ۱۱۱-۱۲۲.

DOI: <https://dor.isc.ac/dor/۲۰.۱۰۰.۱.۲۳۲۲۴۳۴۷.۱۴۰۴.۱۳.۲.۱۰.۲>

© Author(s) retain the copyright and full publishing rights

Publisher: Imam Hossein University.



تخمین خطر امنیتی نشانی‌های اینترنتی به کمک تحلیل تفکیکی خطی

محمود دی پیر^{۱*}، خداداد هلیلی^۲^۱ دانشیار، دانشگاه علوم و فنون هوایی شهید ستاری، تهران، ایران (نویسنده مسئول: mdeypir@gmail.com)^۲ استادیار، دانشگاه علوم و فنون هوایی شهید ستاری، تهران، ایران (kh.halili@ssau.ac.ir)

چکیده (استایل عنوان چکیده)	مشخصات مقاله
ارسال یک آدرس اینترنتی (URL) مخرب به قربانی نقطه شروع انواع مختلف فعالیت‌های مخرب توسط مهاجمان است؛ بنابراین شناسایی URL های مخرب نقش بسزایی در امنیت کاربران اینترنت دارد. اخیراً محاسبه ریسک امنیتی URL ها به جای طبقه‌بندی آن‌ها با استفاده از مدل‌های مبتنی یادگیری ماشینی مورد توجه قرار گرفته است. این امر به این دلیل است که از یک‌طرف هشدارهای لازم را به کاربران می‌دهد و از طرف دیگر مشکلات مدل‌های طبقه‌بندی را ندارد. در این مطالعه، معیار جدیدی بر اساس تحلیل تفکیکی خطی برای تخمین ریسک امنیتی URL ها ابداع شده است. در این معیار، نمونه‌های شناخته‌شده قبلی از URL های عادی و مخرب در فضای جدیدی نگاشت می‌شوند که در آن می‌توان ریسک امنیتی را با دقت بیشتری محاسبه کرد. اگرچه یادگیری عمیق در معیار پیشنهادی استفاده نمی‌شود و به داده‌های آموزشی کمی نیاز دارد، تخمین واقع‌بینانه‌ای برای مقادیر ریسک امنیتی URL های مخرب و ایمن ارائه می‌دهد. آزمایش‌های انجام‌شده بر روی مجموعه‌های داده واقعی نشان می‌دهد که معیار پیشنهادی از نظر میزان تشخیص نسبت به معیارهای ارائه‌شده قبلی برتری دارد. پیاده‌سازی‌های انجام‌شده در این پژوهش به همراه مجموعه داده‌های مورد استفاده در آزمایش‌های انجام‌شده در آدرس https://github.com/mdeypir/LRU در دسترس عموم قرار گرفته‌اند.	تاریخچه مقاله: نوع مقاله: علمی پژوهشی دریافت: ۱۴۰۴/۰۱/۱۰ بازنگری: ۱۴۰۴/۰۲/۰۶ پذیرش: ۱۴۰۴/۰۳/۱۸ ارائه آنلاین: ۱۴۰۴/۰۴/۲۰
	کلید واژه‌ها: تحلیل تفکیکی خطی URL های مخرب خطر امنیتی ویژگی متنی

استناد: دی پیر، محمود^۱، هلیلی، خداداد^۲. تخمین خطر امنیتی نشانی‌های اینترنتی به کمک تحلیل تفکیکی خطی. پدافند الکترونیک و سایبری. ۱۴۰۴؛ ۱۳ (۲): ۱۱۱-۱۲۲.

<https://dor.isc.ac/dor/۲۰.۱۰۰۱.۱.۲۳۲۲۴۳۴۷.۱۴۰۴.۱۳.۲.۱۰.۷>

© نویسنده(گان) حق نشر و حقوق کامل انتشار را برای خود محفوظ می‌دارند.



ناشر: دانشگاه جامع امام حسین(ع).

OPEN ACCESS

۱- مقدمه

URL^۱ مخرب، پیوندی است که باهدف تسهیل کلاهبرداری، حمله سایبری یا هر فعالیت بدخواهانه دیگری ایجاد شده است. وقتی کاربر روی یک URL مخرب کلیک می‌کند، می‌تواند باعث دانلود باج افزار، دسترسی به صفحات فیشینگ، ایمیل‌های فیشینگ یا سایر اشکال نفوذ سایبری شود. حجم قابل توجهی از مطالعات تحقیقاتی در زمینه شناسایی URLهای مخرب انجام شده است. بر اساس این مطالعات، ویژگی‌های مختلفی را می‌توان از آدرس وبسایت، محتوای آن‌ها، دامنه مرتبط، سرور و میزبان استخراج کرد. طیف وسیعی از مدل‌هایی که از تکنیک‌های داده‌کاوی یا یادگیری ماشین استفاده می‌کنند با استفاده از این ویژگی‌ها ایجاد شده‌اند [۱-۱۹]. این مدل‌های طبقه‌بندی برای تشخیص آدرس‌های مخرب از آدرس‌های مفید استفاده می‌شوند. این روش‌ها به‌طور کلی نرخ مشخصی از تشخیص یا خطای طبقه‌بندی را نشان می‌دهند. علاوه بر این، با درک اینکه چگونه این روش‌ها به‌عنوان یک جعبه سیاه یا یک جعبه سفید عمل می‌کنند، نفوذ گران می‌توانند رویکردهایی برای دور زدن این مدل‌ها ابداع کنند. در نتیجه، عملکرد مدل‌های تشخیص URL کاهش می‌یابد [۲]. برای رسیدگی به این چالش‌ها، رویکرد پیشنهادی به جای طبقه‌بندی آن‌ها به‌عنوان مخرب یا غیر مخرب، بر تخمین خطر امنیتی URLها تمرکز دارد. برای این منظور، معیاری را معرفی می‌کنیم. در این معیار، به جای برچسب کلاس باینری (مخرب یا مفید)، برای هر URL ورودی ناشناخته، یک مقدار ریسک امنیتی محاسبه می‌شود. بنابراین، هر چه میزان ریسک بالاتر باشد، احتمال مخرب بودن URL بیشتر است. بنابراین، با مشاهده ارزش ریسک یک URL ورودی، کاربر می‌تواند تصمیم بگیرد که آیا با آن تعامل داشته باشد یا خیر. ارزیابی‌های تجربی ما بر روی مجموعه داده‌های URL واقعی، اثربخشی معیار پیشنهادی را در شناسایی URLهای مخرب نسبت به معیارهای ارائه شده قبلی، نشان می‌دهد. علاوه بر این، معیار پیشنهادی به‌طور قابل توجهی از معیارهای خطر امنیتی مشابهی که برای تشخیص بدافزارها طراحی شده بودند، بهتر عمل می‌کند.

ادامه این مقاله به شرح زیر سازمان‌دهی شده است. بخش بعدی مروری بر تحقیقات گذشته در زمینه شناسایی URLهای مخرب و محاسبه خطر امنیتی ارائه می‌دهد. بخش سوم به تشریح مسئله تحقیق می‌پردازد. بخش چهارم، معیار پیشنهادی و الگوریتم مربوطه را توضیح می‌دهد. بخش پنجم برتری معیار پیشنهادی را از نظر نرخ تشخیص URL مخرب، در مقایسه با معیارهای قبلی، با استفاده از دو مجموعه داده دنیای واقعی نشان می‌دهد. در نهایت، بخش ششم، مقاله را جمع‌بندی و نتیجه‌گیری می‌کند.

۲- پیشینه تحقیق

ایجاد و توزیع URLهای مخرب یک تاکتیک پرکاربرد است که توسط مهاجمان برای نفوذ به سیستم قربانیان، استفاده می‌شود. چندین سال از استفاده اولیه از این روش می‌گذرد و رویکردهای مختلفی برای رفع آن ارائه شده است، اما همچنان یکی از دغدغه‌های اصلی در حوزه امنیت سایبری است. این به دلیل استفاده مهاجمان از تکنیک‌های جدید برای غلبه بر این رویکردها است. تحقیقات گسترده‌ای برای شناسایی URLهای مخرب با موفقیت نسبی انجام شده است. به‌طور کلی، تکنیک‌های شناسایی لینک‌های مخرب را می‌توان به سه دسته، طبقه‌بندی کرد: روش‌های مبتنی بر لیست سیاه [۳]، روش‌های مبتنی بر قانون [۴]، و مهم‌تر از همه، روش‌های مبتنی بر یادگیری ماشینی و یادگیری عمیق [۵]. با توجه به محدودیت‌های دو گروه اول، یادگیری ماشینی و یادگیری عمیق توجه بیشتری را به خود جلب کرده است [۶، ۷]. برای یادگیری ماشین و روش‌های مبتنی بر یادگیری عمیق، ویژگی‌های استخراج شده از URL، شامل دامنه مربوطه به آن و میزبان‌های ارجاع شده، طول URL از نظر تعداد کاراکترها، تعداد کاراکترهای متوالی، تعداد آدرس‌های IP مربوط به هر URL، رتبه الکسا آن، طول عمر URL و زمان انقضای مجوز SSL سایت مربوطه، هستند. پس از استخراج ویژگی‌های لازم، نمونه‌های متعلق به هر گروه از URLهای مخرب یا مفید، با مقادیر مربوط به این ویژگی‌ها، برای ساخت مدل استفاده می‌شوند. سپس مدل ساخته شده برای شناسایی برچسب URLهای ناشناخته استفاده می‌شود. این مدل‌ها معمولاً طبقه‌بندی دو کلاسه هستند که می‌توانند هر URL ورودی را به‌عنوان نرمال (مفید) یا مخرب دسته‌بندی کنند. مدل‌های یادگیری ماشین مختلفی برای حل این مسئله پیشنهاد شده‌اند که موفق‌ترین آن‌ها که دقت‌های بالایی را نشان می‌دهند، درخت‌های تصمیم [۸]، رگرسیون لجستیک [۹]، تقویت گرادیان [۱۰]، ماشین‌های بردار پشتیبان [۱۱]، جنگل‌های تصادفی [۱۲]، و بیز ساده [۱۳] هستند. روش‌های مبتنی بر یادگیری عمیق نیز عملکرد خوبی در این زمینه نشان داده‌اند [۱۴، ۱۵، ۱۶]. برخی از محققان با استفاده از مدل‌های طبقه‌بندی گروهی و ترکیب طبقه‌بندی به عملکرد بهتری دست‌یافته‌اند [۱]. در [۱۶]، یک روش ترکیبی مبتنی بر هفت نوع طبقه‌بندی و یک تکنیک خوشه‌بندی برای فیلتر کردن URLهای فیشینگ ابداع شد که تصمیم‌نهایی با استفاده از تکنیک رأی‌گیری در میان کمیته‌ای از طبقه‌بندی‌کننده‌ها گرفته می‌شد. علاوه بر این، انواع مختلفی از درختان تصمیم، مانند C۴.۵، ADTree، CART، Random Tree و Random Forest، در یک ساختار یادگیری جمعی در [۱۷] به همراه استفاده از ویژگی‌های جدید برای افزایش دقت طبقه‌بندی، ارائه گردید. اگرچه محققان تلاش‌های قابل توجهی برای توسعه مدل‌های کارآمد برای شناسایی URLهای مخرب انجام داده‌اند،

^۱ Unified Resource Locator

اما هیچ تکنیک کاملی وجود ندارد که بتواند این مسئله را به طور کامل حل کند. درواقع، مدل‌ها مستعد خطاهای مثبت کاذب و منفی کاذب در طبقه‌بندی URL ها هستند [۱۸]. شناسایی نشانی‌های اینترنتی به طور غیرمستقیم بر امنیت وب تأثیرگذار است، زیرا شروع دسترسی به یک وبسایت، وارد کردن و استفاده از آدرس URL آن است که خود باید آدرس معتبری باشد وگرنه می‌تواند خود آغازی بر یک حمله سایبری باشد. اگرچه شناسایی URL های مخرب نقش تعیین کننده‌ای در امنیت وب‌گردی برای کاربران اینترنت دارد؛ اما امنیت وب به طور خاص مورد مطالعه بسیاری از محققین قرار گرفته است و روش‌های زیادی تاکنون برای مقابله با این حملات ارائه شده است [۲۰، ۲۱، ۲۲، ۱۹، ۲۰]. بنابراین در اینجا به مرور برخی از تحقیقات در این زمینه نیز می‌پردازیم. شهیدی نژاد و همکارانش [۱۹] به کمک مدل دسته‌بندی درخت تصمیم توانستند حملات تزریق SQL را در زمان اجرا تشخیص دهند. در [۲۰] روشی مبتنی بر امضا به منظور شناسایی حملات اجرای کد با استفاده از سامانه تشخیص نفوذ وب در زبان PHP ایجاد شد. یادگاری و همکاران [۲۱] توانستند به کمک مدل بردار پشتیبان و آنتروپی مربوط به هر آی پی در پنجره‌های زمانی، حملات منع خدمت وب را تشخیص دهند. در [۲۲] یک سازوکار دفاعی پویا با قابلیت سفارشی‌سازی برای شناسایی ربات‌های وب بدخواه مشارکت‌کننده در حملات سایبری، با استفاده از تحلیل رفتار مرورگری آن‌ها، ارائه شده است.

آدرس‌های URL مختص وبسایت‌ها نیستند و در هر سیستم مبتنی بر شبکه و اینترنت برای دستیابی به سرورها و منابع می‌توانند استفاده شوند. ارزیابی خطرات امنیتی برای کلیه سیستم‌های فیزیکی سایبری مانند اینترنت اشیا، نیز مورد بررسی پژوهشگران امنیت سایبری قرار گرفته است [۲۳]. با این حال، یکی از حوزه‌هایی که ارتباط زیادی با پژوهش ما دارد، محاسبه خطرات امنیتی برای برنامه‌های کاربردی تلفن همراه است [۲۴]. مطالعات متعددی باهدف ارائه معیارهای ریسک موثر برای تشخیص بدافزار اندروید از برنامه‌های معمولی انجام شده است [۲۵، ۲۶، ۲۷، ۲۸]. این معیارها در درجه اول بر اساس مجوزهای استفاده شده یا عملکرد برنامه‌های Android است. به عنوان مثال، گیتس و همکاران. تعدادی از معیارهای ریسک مبتنی بر مجوز [۲۴] را معرفی کرد که RSS (سیستم امتیاز ریسک) یکی از قابل توجه‌ترین آن‌هاست. در [۲۵] روشی بر اساس بهره اطلاعاتی به منظور محاسبه خطر امنیتی بدافزارهای اندروید ارائه شد. دی پیر و همکاران [۲۶] توانستند به کمک روش نزدیک‌ترین همسایه، رویکردی را برای شناسایی حملات و برنامه‌های مخرب اندروید ارائه دهند. در این روش از فاصله همینگ به منظور محاسبه نزدیک‌ترین همسایه استفاده می‌شد. در [۲۷]، یک معیار ریسک با استفاده از نمونه‌های قبلی برنامه‌های مخرب و غیر مخرب

اندروید پیشنهاد شد. در این معیار کل فضای ویژگی بدافزارها و برنامه‌های مفید به منظور محاسبه دقیق تر خطر امنیتی مورد استفاده قرار می‌گرفت. علاوه بر این، یک اندازه‌گیری امتیاز ریسک مبتنی بر آنتروپی در [۲۸] ابداع شد که از کسب اطلاعات مجوزها و همچنین ویژگی‌های دیگر برای تشخیص برنامه‌های پرخطر اندروید از برنامه‌های معمولی، استفاده می‌کند.

محاسبه امتیاز ریسک برای URL ها مستلزم استخراج ویژگی‌های مناسب برای به دست آوردن مقادیر ریسک منطقی و دقیق است. همان طور که گفته شد، از ویژگی‌های URL مختلف برای ساخت مدل طبقه‌بندی استفاده شده است. ما از این ویژگی‌ها برای محاسبه ریسک‌های امنیتی واقع بینانه استفاده کرده ایم. بسیاری از روش‌ها URL های آلوده را بر اساس ویژگی‌های متنی موجود در URL ها شناسایی می‌کنند و از تفاوت‌های متن URL های مخرب در مقایسه با URL های معمولی استفاده می‌کنند [۲۹]. برخی از مطالعات همچنین ویژگی‌های دامنه یا سرور مانند IP، SSL و اطلاعات Whois را علاوه بر ویژگی‌های متنی URL ها در نظر می‌گیرند [۳۰]. حجاج و همکاران [۲] سعی کرده‌اند ویژگی‌های قوی برای شناسایی URL های مخرب ارائه دهند، به این معنی که تغییر آن‌ها برای متجاوزان دشوار یا غیرممکن است تا نتوانند از مدل‌های یادگیری ماشین یا سایر راه‌حل‌های امنیتی عبور کنند. هامون و همکاران [۳۱] مجموعه‌ای از ویژگی‌های متنی را با استفاده از تحلیل متنی برای دسته‌بندی URL ها به چهار کلاس ارائه می‌کنند: هرزنامه، فیشینگ، بدافزار و تغییر ظاهر^۱. نویسندگان در [۳۲] یک روش مهندسی ویژگی را برای بهبود دقت طبقه‌بندی با استفاده از مدل‌های محبوبی مانند k-نزدیک‌ترین همسایه، ماشین بردار پشتیبان و پرسپترون چندلایه، ارائه کردند. برخی از مطالعات بر روی ویژگی‌های استخراج شده از سوابق DNS^۲، وب، میزبان‌ها و دامنه‌ها برای دستیابی به دقت طبقه‌بندی بهتر، تمرکز می‌کنند [۳۳، ۳۴]. اخیراً چندین روش تخصصی برای تشخیص URL مخرب معرفی شده است [۳۵-۴۰]. این روش‌ها شامل یک الگوریتم بهینه‌سازی [۳۵]، یک مدل مفصل عصبی موازی [۳۶]، یک رویکرد XGBoost حساس به هزینه برای داده‌های نامتعادل [۳۷]، یک طبقه‌بندی کننده انجمنی [۳۸]، یک شبکه کانولوشن برگشتی چندلایه بهبود یافته [۳۹] و یک تکنیک ترکیب یادگیری عمیق و فیلتر بلوم [۴۰] هستند. درحالی‌که خدمات آنلاین تجاری یا محصولاتی وجود دارند که سعی می‌کنند URL های مخرب را با استفاده از لیست سیاه یا یادگیری ماشینی شناسایی کنند، آن‌ها معمولاً مقدار ریسک را محاسبه نمی‌کنند که نشان دهنده میزان مشکوک بودن باشد. به عنوان مثال می‌توان به سرویس مستقل جداسازی تهدید ایمیل [۴۱] و ریسک وب گوگل [۴۲] اشاره کرد که سطوح خطر گسترده‌ای را بر اساس نمایه‌های سوابق URL یا فهرست‌های

^۱ Defacement^۲ Domain Name Service

نامشخص است؛ بنابراین، ارائه مقدار ریسک امنیتی به جای طبقه‌بندی دقیق می‌تواند رویکرد منطقی و مؤثرتری برای تصمیم‌گیری صحیح باشد. مفهوم محاسبه ریسک امنیتی قبلاً در زمینه‌های دیگری نیز مانند شناسایی برنامه‌های مخرب مورد مطالعه قرار گرفته است و معیارهای مختلفی نیز ارائه شده است [۲۶-۲۸]. در این مطالعه، ما بر روی مشکل تخمین مقدار ریسک امنیتی برای URL‌های ناآشنا برای مقابله با چالش‌های فوق تمرکز می‌کنیم. به طور رسمی می‌توان گفت که یک نشانی اینترنتی یا URL که با u مشخص شده است را می‌توان با یک بردار ویژگی نشان دهیم، به طوری که خواهیم داشت:

$$u = \{f_1, f_2, \dots, f_n\} \quad (1)$$

این ویژگی‌ها می‌توانند ویژگی‌های متنی (واژه‌ای) درون URL، یا ویژگی‌های دامنه، میزبان و سروری باشند که URL به آن اشاره می‌کند، یا ترکیبی از انواع ویژگی‌های مختلف. هدف، ارائه معیاری است که سطح خطر امنیتی آدرس u را بر اساس اطلاعاتی که قبلاً در مورد URL‌های مخرب و غیر مخرب شناخته شده داریم، محاسبه می‌کند؛ بنابراین، ما به دنبال یافتن تابعی به شکل $R(u)$ هستیم که بردار ویژگی u را می‌گیرد و یک مقدار عددی مثبت پیوسته که به عنوان امتیاز ریسک امنیتی u شناخته می‌شود، برمی‌گرداند. ارزش ریسک محاسبه شده باید برای URL‌های شناخته شده قبلی معقول باشد. به عبارت دیگر، مقدار ریسک به دست آمده باید برای یک URL مخرب نسبتاً بزرگ‌تر از یک URL نرمال (مفید) باشد.

۴- روش پیشنهادی

روش ما برای به دست آوردن یک معیار اندازه‌گیری مقدار ریسک برای URL‌ها شامل پنج مرحله اصلی از جمله جمع‌آوری و پیش‌پردازش داده‌ها، مهندسی ویژگی، آموزش مدل، مکانیسم امتیاز ریسک و پیاده‌سازی سیستم است. در جمع‌آوری داده‌ها و پیش‌پردازش، مجموعه داده‌ای جامع از URL‌های مخرب و نرمال قبلاً شناخته شده را جمع‌آوری می‌کنیم. ابتدا داده‌ها را برای استخراج ویژگی‌ها و برجسته‌های مربوطه پیش‌پردازش می‌کنیم. در مهندسی ویژگی، تکنیک‌هایی برای استخراج و نمایش ویژگی‌ها از URL‌ها به کار می‌رود. در مرحله سوم برای ساخت و آموزش مدل، یک مدل ریسک مؤثر را بر روی مجموعه داده برجسته‌گذاری شده، آموزش می‌دهیم. در مکانیسم امتیازدهی ریسک، یک متریک یا مکانیزم امتیازدهی ریسک بر اساس ویژگی‌های استخراج شده از نمونه‌های قبلاً شناخته شده URL‌ها، توسعه می‌یابد. در آخرین مرحله برای پیاده‌سازی و ارزیابی سیستم، ما یک سیستم خودکار را پیاده‌سازی می‌کنیم که مدل آموزش دیده را برای تجزیه و تحلیل URL‌های جدید و برآورد میزان ریسک امنیتی آن‌ها به کار می‌برد.

سیاه اختصاص می‌دهند. مشکل محاسبه ریسک امنیتی URL‌ها برای اولین بار در [۴۳] مورد بررسی قرار گرفت، جایی که دو معیار برای تخمین ریسک URL‌های ناشناخته پیشنهاد شد. یکی از این معیارها مبتنی بر ویژگی و دیگری مبتنی بر نمونه بود که کارایی قابل قبولی در تخصیص مقدار خطر امنیتی به URL‌ها داشتند.

۳- بیان مسئله

همان‌طور که گفته شد روش‌های مختلفی برای طبقه‌بندی URL‌ها به دسته‌های مخرب و غیر مخرب ارائه شده است. با این حال، این روش‌ها اغلب از کاستی‌هایی رنج می‌برند. موفق‌ترین رویکردهای طبقه‌بندی، مانند ^۱SVM، Random Forest، درختان تصمیم، Naïve Bayes و مدل‌های مبتنی بر یادگیری عمیق، علی‌رغم استفاده از طیف گسترده‌ای از ویژگی‌های متنی و مبتنی بر میزبان، هنوز هم خطاهای طبقه‌بندی قابل توجهی در مواجهه با داده‌های واقعی دارند. علاوه بر این، نفوذ گران می‌توانند عملکرد مدل‌ها را از طریق تجزیه و تحلیل جعبه سفید یا جعبه سیاه، شناسایی کنند و سپس الگوهای URL خود را برای دورزن تشخیص آن‌ها، تغییر دهند. از سوی دیگر، درحالی که ایجاد یک لیست سیاه، روشی ساده و آسان برای پیاده‌سازی است، اما راه‌حل مؤثری نیست. زیرا URL‌های مخرب جدید روزانه ظاهر می‌شوند که در لیست وجود ندارند. این به دلیل مشکل در نگهداری فهرست جامع و به روز از URL‌های مخرب است که منجر به مشکل منفی‌های نادرست بالا می‌شود. علاوه بر این، بسیاری از مهاجمان تغییرات جزئی در URL اصلی ایجاد می‌کنند، مانند جایگزینی دامنه سطح بالا (TLD^۲)، استفاده از معادل‌سازی آدرس IP، تغییر ساختار دایرکتوری، جایگزینی رشته‌های پرس‌وجو یا استفاده از معادل نام تجاری که می‌تواند باعث ایجاد URL مخرب ناشناخته شود. علاوه بر این، استفاده از مجموعه‌ای از قوانین در فایروال‌ها برای فیلتر کردن URL‌های مخرب نیز مؤثر نیست، زیرا نفوذ گران می‌توانند این قوانین را شناسایی و دور بزنند. مهم‌تر از همه، در همه انواع روش‌های تشخیص URL مخرب، یک مشکل اساسی وجود دارد: یک URL ممکن است به شدت مخرب یا غیر مخرب نباشد، یا ممکن است تا حدی مخرب باشد. به عنوان مثال، ممکن است یک URL وبسایت تبلیغاتی توسط برخی از کاربران مخرب تلقی نشود، اما ممکن است همچنان نامطلوب تلقی شود؛ بنابراین، انواع مختلف URL‌های نامطلوب، از جمله هرزنامه، فیشینگ، بدافزار و حمله تغییر ظاهر، همه به یک اندازه مخرب نیستند. درحالی که یک دسته‌بندی URL‌های مخرب بر اساس نوع حمله آن‌ها در [۳۱] معرفی شد، میزان خطر مرتبط با هر نوع، هنوز برای کاربر

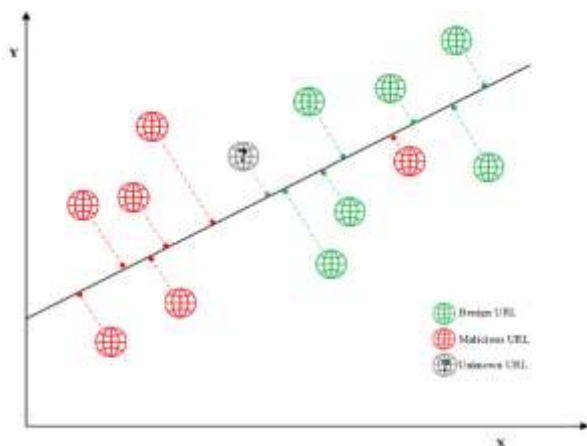
^۱ Support Vector Machine

^۲ Top Level Domain

تفکیک خطی URLهای مخرب و ایمن موجود به دست می‌آید. تصویر نمونه‌های قبلی به فضای جدید به این دلیل است که در فضای جدید این دو نوع نمونه از هم متمایزتر هستند و می‌توان ارزش ریسک URL ناشناخته را با دقت بیشتری محاسبه کرد. به‌طور کلی، خطر امنیتی یک نمونه ناشناخته با فاصله آن از نمونه‌های URL مخرب رابطه معکوس دارد. زیرا هرچه فاصله با نمونه‌های مخرب کمتر باشد، خطر بالاتری دارد. از طرفی ریسک امنیتی این نمونه ارتباط مستقیمی با فاصله با نمونه‌های ایمن دارد، یعنی هر چه فاصله با نمونه‌های ایمن کمتر باشد، ریسک کمتری دارد. این ایده هم در فضای اصلی و هم در فضای پیش‌بینی‌شده، معتبر است. برای محاسبه ریسک در فضای جدید که تفکیک‌پذیری بیشتری دارد، پارامترهای موردنیاز یعنی فواصل، به‌دست‌آمده و در نهایت مقادیر به‌دست‌آمده با یک معادله ساده به‌صورت زیر برای محاسبه خطر امنیتی، به کار می‌روند:

$$Risk(u) = \frac{ds}{dm} \quad (5)$$

معادله (۵) خطر امنیتی یک URL ناشناخته u را تخمین می‌زند. در این فرمول، ds و dm به ترتیب فاصله URLهای نگاشت شده u تا میانگین URLهای ایمن و میانگین URLهای مخرب در این فضای جدید هستند.



شکل (۱): نحوه محاسبه امتیاز ریسک یک URL ناشناخته با استفاده از URLهای مفید و مخرب قبلی

شکل (۲) شبه کد الگوریتم پیشنهادی LRU (تحلیل تفکیکی خطی برای تخمین ریسک URLها) را نشان می‌دهد. این الگوریتم مجموعه URLهای شناخته‌شده قبلی (SUS) و URL ناشناخته u را به‌عنوان ورودی دریافت می‌کند. هدف تخمین ریسک امنیتی URL ناشناخته u است. جدول (۱) نمادهای استفاده‌شده در الگوریتم را نشان می‌دهد.

در خط ۱، داده‌های ورودی، نرمال شده‌اند. در خطوط ۲ و ۳ به ترتیب URLهای مخرب و ایمن از لیست کلی URLها استخراج می‌شوند. مقادیر میانگین این دو مجموعه بر این اساس در خطوط ۴ و ۵ محاسبه می‌شوند. در خطوط ۶ و ۷ به ترتیب ماتریس پراکندگی درون کلاسی، برای هر دو کلاس (مجموعه)

در این تحقیق از تحلیل تفکیک خطی^۱ به‌منظور توسعه مدلی برای تخمین امتیاز ریسک URLها استفاده‌شده است. این تکنیک که به‌اختصار بنام LDA شناخته می‌شود در دسته‌بندی‌های دو کلاسه یا به‌عنوان روشی برای کاهش بعد، کاربرد دارد. هدف از تکنیک LDA این است که ماتریس داده اصلی را در فضایی با ابعاد پایین‌تر نشان دهد. برای رسیدن به این هدف باید سه مرحله انجام شود [۴۴] اولین مرحله محاسبه تفکیک‌پذیری بین کلاس‌های مختلف (فاصله بین میانگین‌های کلاس‌های مختلف) است که به آن واریانس بین کلاس یا ماتریس بین کلاس می‌گویند. این واریانس با فرمول زیر محاسبه می‌شود:

$$S_W = \sum_{i=1}^I (\sum_{j=1}^{n_i} (x_j^{(i)} - m^{(i)})(x_j^{(i)} - m^{(i)})^T) \quad (2)$$

مرحله دوم محاسبه فاصله بین میانگین و نمونه‌های هر کلاس است که به آن واریانس درون کلاسی یا ماتریس درون کلاسی می‌گویند که به‌صورت زیر محاسبه می‌شود:

$$S_B = \sum_{i=1}^I n_i (m^{(i)} - m)(m^{(i)} - m)^T \quad (3)$$

در فرمول‌های فوق I ، x و m به ترتیب تعداد کلاس‌ها، تعداد نمونه‌های کلاس i ام، هر نمونه موجود و میانگین نمونه‌ها هستند. مرحله سوم ساخت فضا با ابعادی پایین‌تر است که واریانس بین کلاس را به حداکثر می‌رساند و هم‌زمان واریانس درون کلاس را به حداقل می‌رساند. به‌عبارت دیگر تابع هدف تحلیل تفکیکی خطی به‌صورت زیر است:

$$W_{opt} = \operatorname{argmax}_w \frac{W^T S_B W}{W^T S_W W} \quad (4)$$

البته در ما در اینجا از LDA به‌منظور دسته‌بندی یا کاهش بعد استفاده نمی‌کنیم؛ بلکه آن را برای تخمین خطر امنیتی نشانی‌های ناشناخته با استفاده از نمونه‌های مخرب و مفید قبلی به کار گرفته‌ایم. ایده اصلی ما در شکل (۱) نشان داده‌شده است. در این شکل برای سادگی، فضای دوبعدی نشان داده‌شده است اما در عمل می‌توان از بسیاری از ابعاد (ویژگی‌ها) استخراج‌شده از URLها استفاده کرد. همان‌طور که در شکل (۱) نشان داده‌شده است، با توجه به تجزیه و تحلیل تفکیک خطی، هر دو نمونه URL مخرب و مفید به فضای با بعد پایین‌تر نگاشت می‌شوند، به‌نحوی که در این فضا تفکیک‌پذیری افزایش می‌یابد. درواقع با توجه به فرمول‌های فوق، نگاشت نمونه‌ها به بعد پایین، سبب از دست رفتن اطلاعات مربوط به جداپذیری نمونه‌ها نمی‌شود و عمل کاهش بعد با افزایش تفکیک‌پذیری نمونه‌های مخرب و مفید همراه است. معیار ارائه‌شده در اینجا پارامترهای لازم برای تخمین ریسک در فضای جدید را محاسبه می‌کند. این پارامترها مربوط به فاصله URL ناشناخته تا نمونه‌های مخرب و ایمن (نرمال)، در فضای نگاشت شده، است. این معیار در نهایت با کمک پارامترهای به‌دست‌آمده، میزان ریسک را تخمین می‌زند. درواقع، فضای متمایزکننده جدید با استفاده از تجزیه و تحلیل

^۱Linear Discriminant Analysis (LDA)

جدول (۱): نمادهای مورد استفاده و معنای آن‌ها در الگوریتم LRU.

نماد	مفهوم
SUs	مجموعه URLهای شناخته‌شده موجود (مضر و ایمن)
MS	مجموعه URLهای مخرب
SS	مجموعه URLهای ایمن
μ	مقدار متوسط برای مجموعه URLهای مخرب
m_s	مقدار متوسط برای مجموعه URLهای امن
SW_m	ماتریس پراکندگی درون کلاسی برای URLهای مخرب
SW_s	ماتریس پراکندگی درون کلاسی برای URLهای ایمن
SW	جمع ماتریس‌های پراکندگی درون کلاسی
W	ماتریس پراکندگی بین کلاسی
Y_m	مجموعه‌ای از URLهای مخرب نگاشت شده
Y_s	مجموعه‌ای از URLهای ایمن نگاشت شده
y_u	نگاشت URL ورودی u
mY_m	میانگین مقدار URLهای مخرب نگاشت شده
mY_s	میانگین مقدار URLهای ایمن نگاشت شده
dm	فاصله بین URL نگاشت شده u و mY_m
ds	فاصله بین URL نگاشت شده u و mY_s

جدول (۲): ویژگی‌های مجموعه داده‌های مورد استفاده

مجموعه داده	Textual Kaggle	OpenIntel
تعداد کل URLها	۴۵۰۱۷۶	۹۳۲۰۷۸
تعداد URLهای مخرب	۱۰۴۴۳۸	۵۸۷۲۵۷
تعداد URLهای ایمن	۳۴۵۷۳۸	۳۴۴۸۲۱
تعداد ویژگی‌ها	۱۸	۶
توصیف	ویژگی‌های متن استخراج‌شده از URLهای مفید و مخرب موجود در سایت [۴۵ Kaggle]	تعداد بالایی از URLهای مخرب و مفید [۴۶]

ما طیف گسترده‌ای از ویژگی‌های متن (واژگانی) را با استفاده از یک اسکرپت پایتون برای اولین مجموعه داده‌ای که از وبسایت Kaggle به دست آمده است، پیش‌پردازش و استخراج کرده‌ایم. ویژگی‌های متن یا واژگانی رایج‌ترین ویژگی‌هایی هستند که برای رویکردهای مختلف شناسایی URL مخرب استفاده می‌شوند. پس از حذف اطلاعات استفاده‌نشده یا اضافی از فایل ورودی، اسکرپت نوشته‌شده فایل ورودی را که شامل تعداد زیادی URL مخرب و غیر مخرب است را می‌گیرد و پردازش متن را برای استخراج ویژگی‌ها انجام می‌دهد. این اسکرپت ابتدا داده‌ها را با حذف اطلاعات استفاده‌نشده؛ مانند شماره ردیف‌ها و فاصله‌های اضافی، پیش‌پردازش می‌کند. سپس، توکن‌سازی را برای تسهیل استخراج ویژگی انجام می‌دهد. در مرحله اول استخراج ویژگی، اسکرپت، طول URL کامل، طول نام میزبان، طول مسیر و طول اولین زیرشاخه را اندازه می‌گیرد.

URLهای مخرب و ایمن به دست آمده است. پس از محاسبه ماتریس پراکندگی بین کلاسی در خط ۸، ماتریس تبدیل LDA در خط ۹ محاسبه می‌شود. در خطوط ۱۰ و ۱۱، مجموعه‌های نگاشت شده از URLهای مخرب و ایمن، توسط تبدیل LDA به دست می‌آیند. سپس در خط ۱۲ نگاشت URL ورودی ناشناخته u نیز توسط همین تبدیل محاسبه می‌شود. پس از آن، مقادیر میانگین دو مجموعه نگاشت شده به ترتیب در خطوط ۱۳ و ۱۴ محاسبه می‌شوند. در خطوط ۱۵ و ۱۶، فواصل نمونه ورودی u تا هر دو مقدار میانگین، محاسبه می‌شوند. این فواصل برای تخمین مقادیر ریسک URL ورودی ناشناخته u مورد نیاز هستند. پس از به دست آوردن پارامترهای مورد نیاز، در نهایت، در خط ۱۷، ریسک امنیتی URL ورودی u با استفاده از فرمول (۵) همان‌طور که قبلاً توضیح داده شد، محاسبه می‌شود.

۵- ارزیابی و مقایسه کارایی

ما الگوریتم مربوط به معیار پیشنهادی LRU خود را با استفاده از نرم‌افزار ۲۰۲۳ Matlab و زبان برنامه‌نویسی آن پیاده‌سازی کرده‌ایم. برای افزایش سرعت محاسبات، به جای محاسبه ریسک برای هر URL به صورت جداگانه طبق رویکرد نشان داده‌شده در شکل (۲)، ما محاسبات مبتنی بر ماتریس را در Matlab انجام می‌دهیم تا امتیاز ریسک مجموعه‌ای از URLهای ناشناخته را به طور هم‌زمان محاسبه کنیم. ما از دو مجموعه داده در دسترس عموم برای ارزیابی کارایی معیار LRU پیشنهادی در سناریوهای مختلف استفاده می‌کنیم. جدول (۲) خصوصیات کلیدی این مجموعه داده‌ها را خلاصه می‌کند. مجموعه داده اول فقط شامل ویژگی‌های متن (واژگانی) است ولی مجموعه داده دوم حاوی انواع ویژگی‌های واژگانی و مبتنی بر میزبان است.

Procedure LRU (SUs, u)

begin

1. $SUs = \text{normalize}(SUs)$;
2. $MS = \{\forall \text{ malicious } su \in SUs\}$;
3. $SS = \{\forall \text{ safe } su \in SUs\}$;
4. $\mu = \text{mean}(MS)$;
5. $m_s = \text{mean}(SS)$;
6. $SW_m = \sum_{x_i \in MS} (x_i - \mu) \times (x_i - \mu)^T$;
7. $SW_s = \sum_{x_i \in SS} (x_i - m_s) \times (x_i - m_s)^T$;
8. $SW = SW_m + SW_s$;
9. $W = SW^{-1} \times (\mu - m_s)$;
10. $Y_m = W^T \times MS$;
11. $Y_s = W^T \times SS$;
12. $y_u = W^T \times u$;
13. $mY_m = \text{mean}(Y_m)$;
14. $mY_s = \text{mean}(Y_s)$;
15. $dm = |y_u - mY_m|$;
16. $ds = |y_u - mY_s|$;
17. **return** $\frac{ds}{dm}$;

end

شکل (۲): الگوریتم LRU

مختلف درصد، بالاترین URLها از نظر امتیاز خطر امنیتی را از فهرست مرتب شده انتخاب می کنیم. پس از آن، برای فهرست مرتب شده هر متریک، مقدار درصد URLهای مخرب موجود در قسمت انتخاب شده را تعیین می کنیم. در نتیجه، درصد انتخاب هر لیست مرتب و درصد شناسایی شده از URLنای مخرب به ترتیب به عنوان نرخ هشدار و نرخ تشخیص نام گذاری می شوند. تعداد موارد مثبت کاذب و مثبت واقعی به ترتیب با میزان هشدار و تشخیص نسبت مستقیم دارند. اگرچه محدوده های ارزش ریسک محاسبه شده برای معیارهای مقایسه شده متفاوت است، اما می توان آن ها را با این روش به صورت عادلانه مقایسه کرد زیرا آستانه نرخ هشدار مطلق در اینجا استفاده نمی شود. در نتیجه، فرآیند ارزیابی ما مستقل از محدوده ارزش ریسک معیارها است. واضح است که با معیار ریسک قوی تر، تعداد بیشتری از URLهای مخرب در بالای لیست مرتب شده قرار می گیرند و بنابراین معیار، نرخ تشخیص بیشتری دارد. به عبارت دیگر، کاربر نهایی انتظار دارد که URLهایی با بالاترین امتیاز ریسک، مخرب باشند و ایمن نباشند. ما معیارهای ابداع شده و پنج معیار قبلی را با استفاده از دو مجموعه داده نشان داده شده در جدول (۲) باهم مقایسه می کنیم. شکل (۳) نتایج به دست آمده را برای هر دو مجموعه داده، در قالب نمودارهای ROC نشان می دهد. همان طور که در این شکل نشان داده شده است، معیار LRU پیشنهادی از نظر نرخ تشخیص بهتر از سایر معیارها است.

برای اولین مجموعه داده، شکاف عملکرد بین معیارهای پیشنهادی و موارد قبلی قابل توجه است. در واقع معیار پیشنهادی، هم در سطوح هشدار کم و هم سطوح هشدار بالاتر، دارای نرخ تشخیص بسیار بهتری است که نشان دهنده تخمین دقیق خطر امنیتی در این مجموعه داده است. استفاده از تحلیل تفکیکی خطی سبب می شود نمونه های مفید و مخرب قبلی جدا پذیری بیشتری پیدا کنند و مقدار واقع بینانه تری برای خطر امنیتی نمونه های ناشناخته به دست آید. در این مجموعه داده، FRU و NRU به ترتیب رتبه های دوم و سوم را دارند. از بین این دو، NRU پایداری بیشتری در سطوح هشدار مختلف نشان می دهد. در مجموعه داده دوم، پس از LRU پیشنهادی، NRU و RSS از نظر عملکرد به ترتیب در رتبه های دوم و سوم قرار دارند. همچنین با دقت در این منحنی ها می توان دریافت که روش پیشنهادی علاوه بر برتری دارای عملکرد یکنواخت تری در سطوح هشدار مختلف است. نتایج به دست آمده برای نرخ تشخیص، روی هر دو مجموعه داده، نشان می دهد روش پیشنهادی LRU برای نمونه های مفید مقدار خطر امنیتی کمتر و برای نمونه های مخرب خطر امنیتی بالاتری را محاسبه می کند.

برای ارزیابی جامع تر در همه مقادیر سطح هشدار، شکل (۴) سطح کل زیر منحنی (AUC) همه معیارها را مقایسه می کند.

همچنین کاراکترهای ویژه هر URL از جمله کاراکترهای خاصی مثل @, #, %, ., /, = را می شمارد. علاوه بر این، تعداد ارقام، حروف و دایرکتوری های فرعی را می شمارد. متعاقباً، نوع پروتکل (یعنی http یا https)، وجود عبارت «www»، حاوی آدرس IP بودن و استفاده از سرویس کوتاه کننده را تشخیص می دهد. URLهایی که از سرویس های کوتاه کننده استفاده می کنند، عبارات خاصی در متن خود دارند که در اسکریپت فهرست شده اند. به طور خلاصه، هجده ویژگی از این مجموعه داده استخراج شد. ما دومین مجموعه داده [۴۶] را از GitHub^۱، به دست آوردیم. این بزرگ ترین مجموعه داده مورد استفاده در آزمایش های ماست. برخلاف مورد اول، در این مجموعه داده تعداد URLهای مخرب بیشتر از تعداد URLهای مفید است. مجموعه داده های استفاده شده و پیاده سازی های ما برای معیار پیشنهادی و معیارهای ریسک قبلی، به صورت رایگان برای محققین در دسترس هستند.^۲

ما آزمایش هایی را روی مجموعه داده های اول و دوم انجام دادیم تا عملکرد تشخیص متریک LRU پیشنهادی را با دو معیار پیشنهادی قبلی، FRU و NRU، برای اندازه گیری امتیاز ریسک URL مقایسه کنیم [۴۳]. علاوه بر این، برای یک تجزیه و تحلیل مقایسه ای جامع تر، ما همچنین سه معیار دیگر را که در ابتدا برای اندازه گیری امتیاز ریسک برنامه اندروید پیشنهاد شده بود را در مقایسه ها استفاده کردیم. این معیارهای عبارت اند از: RSS [۲۶]، IRS [۲۷] و ERS [۲۸]. این سه معیار در این پژوهش برای محاسبه ریسک URL تطبیق داده شدند. تنها تفاوت این است که معیارهای امتیاز ریسک برنامه بر روی مقادیر ویژگی های دودویی کار می کنند، در حالی که در معیارهای پیشنهادی، مقادیر ویژگی برای URLها و دامنه ها می توانند به صورت پیوسته یا دودویی باشند. بنابراین، ما به سادگی این سه معیار را برای مطالعه تطبیقی، پیاده سازی و به کار گرفته شد. در مجموع می توانیم معیار پیشنهادی خود را با پنج معیار محاسبه ریسک قبلی مقایسه کنیم.

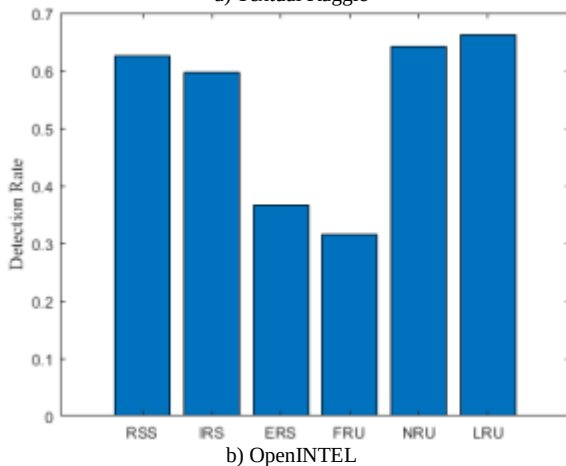
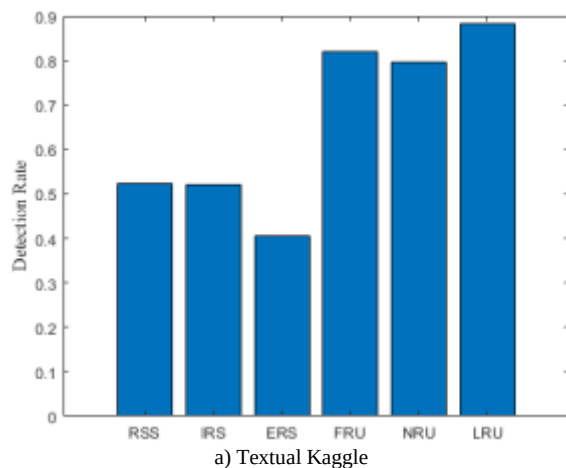
برای ارزیابی هر معیار ریسک، همه بردارهای URL مخرب و عادی را در یک لیست قرار می دهیم. ما از روش اعتبارسنجی متقابل ۱۰ بخشی^۳ برای ارزیابی استفاده می کنیم. به این صورت که در هر بخش، معیار امتیاز پیشنهادی را با استفاده از ۹۰٪ URLها به عنوان داده های آموزشی و ۱۰٪ باقی مانده به عنوان داده های آزمایشی به کار می بریم. با استفاده از هر معیار ریسک، مقادیر ریسک همه URLهای موجود در لیست آزمایشی را به طور جداگانه تخمین می زنیم و لیست را به ترتیب نزولی مقادیر ریسک امنیتی محاسبه شده، مرتب می کنیم. هر چه URLهای مخرب بیشتری در بالای لیست مرتب شده قرار گیرند، معیار خطر امنیتی، قوی تر است. برای اندازه گیری عملکرد برای مقادیر

^۱ <https://github.com/slrbl/malicious-urls-detection-with-autoencoder-neural-networks>

^۲ <https://github.com/mdeypir/LRU>

^۳ ۱۰-fold Cross Validation

$$Percision = \frac{TP}{TP + FP} \quad (5)$$



شکل (۴): مقایسه مساحت زیر منحنی (AUC)

یادآوری (Recall) بیانگر نسبت تعداد نشانی‌های مخربی است که با موفقیت شناسایی شدند به تعداد همه نشانی‌های مخرب که با موفقیت شناسایی شدند یا به‌طور ناموفق به‌عنوان نشانی بی‌ضرر تشخیص داده شدند:

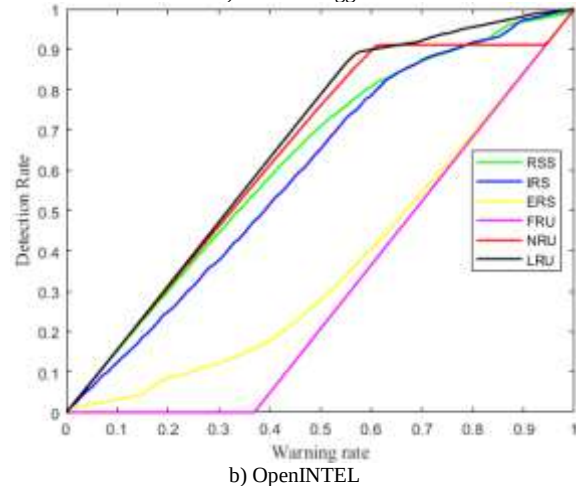
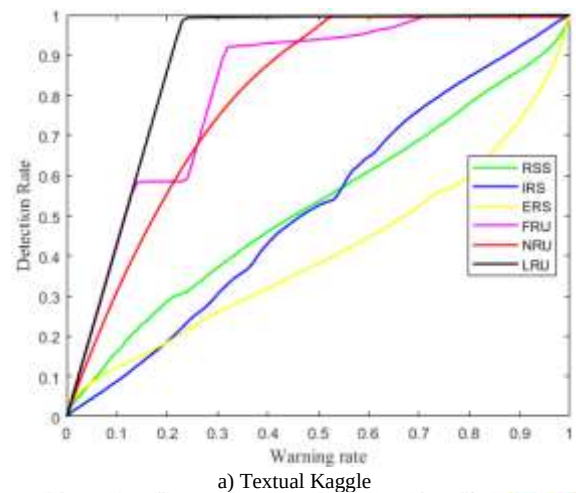
$$Recall = \frac{TP}{TP + FN} \quad (6)$$

نمره F_1 (F1-Score) میانگین هارمونیک از دقت و یادآوری است. این مقدار نشان می‌دهد که معیار پیشنهادی ریسک چقدر توانایی عمل در تمایز بین دو کلاس مخرب و بی‌خطر را دارد:

$$F_1\text{-Score} = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (7)$$

اگرچه فرمول‌های فوق اصولاً برای ارزیابی طبقه‌بندی‌ها به کار می‌روند، ما در اینجا از آن‌ها به‌منظور مقایسه توانایی معیارهای محاسبه خطر، استفاده کرده‌ایم. نتایج به‌دست‌آمده در جدول (۳) گزارش شده‌اند. همان‌طور که در این جدول مشاهده می‌شود، برای هر دو مجموعه داده، روش پیشنهادی LRU دقت، یادآوری و نمره F_1 خوبی دست می‌آورد. به‌ویژه در مجموعه داده اول برتری عددی روش پیشنهادی بسیار چشمگیر است. اما در مجموعه داده

به‌طور طبیعی، یک متریک با AUC نزدیک به یک، معیار خوبی به‌عنوان تخمین گر ریسک URL است.



شکل (۳): مقایسه نرخ تشخیص معیارهای امتیاز ریسک پیشنهادی و معیارهای قبلی توسط نمودار ROC

همان‌طور که در شکل (۴) نشان داده شده است، معیارهای LRU پیشنهادی دارای AUCهای بزرگ‌تری برای هر دو مجموعه داده است زیرا نرخ تشخیص بالاتری در مقادیر مختلف سطح هشدار دارد. در آزمایش بعدی، نمرات دقت، یادآوری و نمره F_1 تمام معیارهای محاسبه ریسک در هر دو مجموعه داده و برای همه معیارها محاسبه شده است. برای این منظور ابتدا باید مقادیر مثبت درست (True Positive)، منفی کاذب (False Negative) و مثبت کاذب (False Positive) محاسبه شوند. مثبت درست نشان‌دهنده URLهای مخربی است که با موفقیت شناسایی شده‌اند. مثبت کاذب نشان‌دهنده تعداد URLهای مفیدی است که به‌اشتباه به‌عنوان مخرب شناسایی شده‌اند. منفی‌های کاذب تعداد URLهای مخربی را نشان می‌دهد که به‌اشتباه به‌عنوان نشانی‌های مفید شناسایی شده‌اند. دقت (Precision) بیانگر نسبت تعداد نشانی‌های مخرب شناسایی شده درست به تعداد کل همه نشانی‌های مخرب شناسایی شده (درست یا غلط) است. معادله این نسبت به‌صورت زیر نشان داده شده است:

دوم روش پیشنهادی با اختلاف کمی رتبه دوم را به دست آورده است. دلیل کارایی خوب روش پیشنهادی این است که مقدار ارزش ریسک واقع بینانه تری را برای نمونه های ناشناخته تخمین می زند. بنابراین در مقایسه با سایر معیارهای اندازه گیری ریسک، مثبت های واقعی بالاتر، مثبت های کاذب کمتر، منفی های واقعی بیشتر و منفی های کاذب کمتری به دست می آیند.

جدول (۳): مقایسه دقت، یادآوری، و نمره F1

مجموعه داده	معیار	دقت	یادآوری	نمره F1
Textual Kaggle	RSS	۰.۳۰۸۰	۰.۳۰۷۹	۰.۳۰۸۰
	IRS	۰.۲۲۵۶	۰.۲۲۵۵	۰.۲۲۵۶
	ERS	۰.۲۰۹۸	۰.۲۰۹۷	۰.۲۰۹۹
	FRU	۰.۵۸۵۸	۰.۵۸۵۷	۰.۵۸۵۸
	NRU	۰.۶۲۳۲	۰.۶۲۳۱	۰.۶۲۳۲
	LRU	۰.۹۹۱۰	۰.۹۹۰۹	۰.۹۹۱۰
OpenINTEL	RSS	۰.۸۲۸۱	۰.۸۲۸۱	۰.۸۲۸۱
	IRS	۰.۸۲۴۷	۰.۸۲۴۶	۰.۸۲۴۷
	ERS	۰.۴۴۵۱	۰.۴۴۵۰	۰.۴۴۵۱
	FRU	۰.۴۱۲۸	۰.۴۱۲۷	۰.۴۱۲۸
	NRU	۰.۹۱۱۴	۰.۹۱۱۳	۰.۹۱۱۴
	LRU	۰.۹۰۶۱	۰.۹۰۶۰	۰.۹۰۶۱

بر اساس آزمایش ها، نتیجه می گیریم که معیار برنده، معیار LRU پیشنهادی است، زیرا بهترین عملکرد را در همه موقعیت ها با هر دو مجموعه داده و تقریباً همه نرخ های هشدار دارد. دلیل برتری معیار پیشنهادی نسبت به سایر معیارها را می توان در تحلیل تفکیک خطی جستجو کرد که منجر به نگاشت نمونه ها در فضای جدید می شود. در این فضا، نمونه های مخرب و ایمن، بیشتر از هم تفکیک شده و در نهایت برآورد دقیق تری از ریسک امنیتی به دست می آید.

۶- بحث و نتیجه گیری

در این پژوهش، به جای طبقه بندی دقیق، تمرکز بر تخمین امتیاز ریسک امنیتی برای URL های ناشناخته بود. معیاری به نام LRU (تحلیل تفکیکی خطی برای تخمین ریسک URL ها) برای حل این مسئله پیشنهاد شد. برای محاسبه معیار بر روی داده های واقعی، الگوریتم مورد نیاز طراحی و پیاده سازی شد. این الگوریتم می تواند با محاسبه ریسک بالا برای URL های مخرب و اعلام هشدار به کاربر، از کلاهبرداری های اینترنتی پیشگیری کند. در صورتی که روش پیشنهادی به عنوان یک افزونه به مرورگرها اضافه شود و یا در شبکه های اجتماعی و سیستم های تبادل پیام به کار رود می تواند در اعلام هشدار به کاربران بسیار مؤثر باشد و جلوی بسیاری از نفوذها را بگیرد؛ زیرا بسیاری از روش های نفوذ با ارسال یک URL از سوی مهاجم شروع می شوند.

آزمایش های انجام شده، اثربخشی معیار پیشنهادی LRU را از نظر میزان تشخیص نشان می دهند. نتایج نشان می دهد که LRU می تواند مقادیر بالای ریسک را به URL های مخرب و مقادیر کم را به آدرس های غیرمخرب اختصاص دهد. این به دلیل مدل ریاضی بهتر آن است که از فضای ویژگی های نگاشت شده URL های مخرب و غیر مخرب قبلاً شناخته شده، برای بهبود تخمین ریسک URL های ناشناخته استفاده می کند. همان طور که نتایج آزمایش ها نشان داد روش پیشنهادی بدون استفاده از یادگیری عمیق، کارایی خوبی از خود نشان داده است. از طرفی با استفاده از داده های کمی می توان مدل ساده پیشنهادی را آموزش داد. به طور کلی روش تحلیل تفکیکی خطی (LDA) که روش پیشنهادی ما مبتنی بر آن است دارای سه مشکل اصلی است که این مشکلات عبارت اند از مشکل اندازه نمونه کوچک، حساسیت به نویز و نقاط پرت، و ناتوانی در برخورد با داده های کلاس چندوجهی (Multi Modal classes) که می تواند در شرایطی سبب کاهش کارایی آن شود. البته برای حل یا تعدیل این مشکلات راه حل هایی نیز وجود دارد که تقسیم کلاس ها به زیر کلاس های مختلف از جمله این راه حل هاست که تفکیک پذیری داده ها در ساختارهای پیچیده را بیشتر می کند [۴۶] روش های مبتنی بر یادگیری عمیق ممکن است کارایی بهتری داشته باشند؛ ولی علاوه بر نیاز به داده های زیاد، زمان آموزش بیشتری را نیاز خواهند داشت. در آینده قصد داریم محاسبه خطر امنیتی را با استفاده از یادگیری عمیق انجام دهیم و روشی را در این زمینه ارائه دهیم. همچنین استفاده از روش های متنوع تر استخراج ویژگی به منظور تخمین دقیق تر خطر امنیتی نشانی های اینترنتی از دیگر کارهای آینده است.

۷- مراجع

- [۱] P. S. Pakhare, S. Krishnan, N. N. Charniya, "Malicious url detection using machine learning and ensemble modeling," In Computer Networks, Big Data and IoT: Proceedings of ICCBI ۲۰۲۰, pp. ۸۳۹-۸۵۰, Springer Singapore, ۲۰۲۰, doi: https://doi.org/10.1007/978-981-16-0965-7_65
- [۲] C. Hajaj, N. Hason, and A. Dvir, "Less is more: Robust and novel features for malicious domain detection," Electronics, vol. ۱۱, no. ۶, p. ۹۶۹, ۲۰۲۲, <https://doi.org/10.3390/electronics11060969>.
- [۳] S. Kim, J. Kim, and B. B. Kang, "Malicious URL protection based on attackers' habitual behavioral analysis," Computers & Security, vol. ۷۷, pp. ۷۹-۸۰۶, ۲۰۱۸, <https://doi.org/10.1016/j.cose.2018.01.013>.
- [۴] A. S. Raja, G. Pradeepa, and N. Arulkumar, "Mudhr: Malicious URL detection using heuristic rules based

- Computer Technology and Transportation (ISCTT), pp. ۳۰۲-۳۲۰, IEEE, ۲۰۲۰, doi: 10.1109/ISCTT51595.2020.00060.
- [۱۴] R. Vinayakumar, K. P. Soman, and P. Poornachandran, "Evaluating deep learning approaches to characterize and classify malicious URL's," *Journal of Intelligent & Fuzzy Systems*, vol. ۳۴(۲), pp. ۱۳۳۳-۱۳۴۳, ۲۰۱۸, DOI: 10.۳۲۳۳/JIFS-۱۶۹۴۲۹.
- [۱۵] J. Yuan, Y. Liu, and L. Yu, "A novel approach for malicious url detection based on the joint model," *Security and Communication Networks*, p. ۴۹۱۷۰۱۶, ۲۰۲۱, <https://doi.org/10.1155/2021/4917016>.
- [۱۶] P. L. Indrasiri, M. N. Halgamuge, and A. Mohammad, "Robust Ensemble Machine Learning Model for Filtering Phishing URLs: Expandable Random Gradient Stacked Voting Classifier," (ERG-SVC). *IEEE Access*, vol. ۹, pp. ۱۵۰۱۴۲-۱۵۰۱۶۱, ۲۰۲۱, doi: 10.1109/access.2021.3124628.
- [۱۷] D. R. Patil, and J. B. Patil, "Malicious URLs detection using decision tree classifiers and majority voting technique," *Cybernetics and Information Technologies*, vol. ۱۸, no. ۱, pp. ۱۱-۲۹, ۲۰۱۸, doi: 10.۲۴۷۸/cait-2018-0002.
- [۱۸] D. K. Mondal, B. C. Singh, H. Hu, S. Biswas, Z. Alom, and M. A. Azim, "SeizeMaliciousURL: A novel learning approach to detect malicious URLs," *Journal of Information Security and Applications*, vol. ۶۲, ۱۰۲۹۶۷, ۲۰۲۱, <https://doi.org/10.1016/j.jisa.2021.102967>.
- [۱۹] A. Shahidinejad, M. Torabi, "Detection and Prevention of SQL Injection Attacks at Runtime Using Decision Tree Classification," *Electronic and Cyber Defense*, vol. ۸, no. ۴, pp. ۷۵-۹۳, ۲۰۲۱, doi: 10.1001.1.23222434.1399.4.7.3 (in persian)
- [۲۰] M. Maghale, M. Bagheri, "Detection of the Remote Code Execution Attacks Using the PHP Web Application Intrusion Detection System," *Electronic and Cyber Defense*, vol. ۱۰, no. ۲, pp. ۷۵-۸۵, ۲۰۲۲, doi: 10.1001.1.23222434.1401.102.7.3 (in persian)
- [۲۱] V. Yadegari, and A. R. Matinfar, "Detect Web Denial of Service Attacks Using Entropy and Support Vector Machine Algorithm," *Electronic and Cyber Defense*, vol. ۶, no. ۴, pp. ۷۹-۸۹, ۲۰۱۹, doi: 10.1001.1.23222434.1397.4.7.9 (in persian)
- [۲۲] M. Fathian, A. M. Abdollahi, and H. Dehghani, "Modeling Browsing Behavior Analysis for Malicious Robot Detection in Distributed Denial of Service Attacks," *Electronic and Cyber Defense*, vol. ۴, no. ۲,
- approach," In *AIP Conference Proceedings*, vol. ۲۳۹۳, no. ۱, p. ۰۲۰۱۷۶, AIP Publishing LLC, ۲۰۲۲, <https://doi.org/10.1063/5.0074077>
- [۲۳] R. Madhubala, N. Rajesh, L. Shaheetha, and N. Arulkumar, "Survey on Malicious URL Detection Techniques," In ۲۰۲۲ ۹th International Conference on Trends in Electronics and Informatics (ICOEI), pp. ۷۷۸-۷۸۱, IEEE, ۲۰۲۲, doi: 10.1109/ICOEI53556.2022.9777221.
- [۲۴] R. Vinayakumar, M. Alazab, K. P. Soman, P. Poornachandran, and S. Venkatraman, "Robust intelligent malware detection using deep learning," *IEEE Access*, vol. ۷, pp. ۴۶۷۱۷-۴۶۷۳۸, ۲۰۱۹, doi: 10.1109/ACCESS.2019.2906934.
- [۲۵] Y. Liang, Q. Wang, K. Xiong, X. Zheng, Z. Yu, and D. Zeng, "Robust Detection of Malicious URLs With Self-Paced Wide & Deep Learning," *IEEE Transactions on Dependable and Secure Computing*, vol. ۱۹, no. ۲, pp. ۷۱۷-۷۳۰, ۲۰۲۱, doi: 10.1109/TDSC.2021.3121388.
- [۲۶] R. Rakesh, S. Muthurajikumar, L. Sairamesh, M. Vijayalakmi, and A. Kannan, "Detection of URL based attacks using reduced feature set and modified C.F. ۵ algorithm," *Adv. Nat. Appl. Sci*, vol. ۹, pp. ۳۰۴-۳۱۱, ۲۰۱۵.
- [۲۷] F. A. Ghaleb, M. Alsaedi, F. Saeed, J. Ahmad, and M. Alasli, "Cyber Threat Intelligence-Based Malicious URL Detection Model Using Ensemble Learning," *Sensors*, vol. ۲۲, no. ۹, p. ۳۳۷۳, ۲۰۲۲, doi: 10.3390/s22093373.
- [۲۸] S. He, B. Li, H. Peng, J. Xin, and E. Zhang, "An effective cost-sensitive XGBoost method for malicious URLs detection in imbalanced dataset," *IEEE Access*, vol. ۹, pp. 93089-93096, ۲۰۲۱, doi: 10.1109/access.2021.3093094.
- [۲۹] R. Patgiri, H. Katari, R. Kumar, and D. Sharma, "Empirical study on malicious URL detection using machine learning," In *International Conference on Distributed Computing and Internet Technology*, pp. 380-388, Springer, Cham, ۲۰۱۹, https://doi.org/10.1007/978-3-030-53666-6_31.
- [۳۰] J. Chen, Z. Hu, and Z. Qian, "Research on malicious URL detection based on random forest," In ۲۰۲۲ 11th International Conference on Computer Research and Development (ICCRD), pp. 30-36, IEEE, ۲۰۲۲, January, doi: 10.1109/iccrd54409.2022.9730451.
- [۳۱] C. Ding, "Automatic detection of malicious urls using fine-tuned classification model," In ۲۰۲۰ ۵th International Conference on Information Science,

- [۳۲] T. Li, G. Kou, and Y. Peng, "Improving malicious URLs detection via feature engineering: Linear and nonlinear space transformation methods," *Information Systems*, vol. ۹۱, pp. ۱۰۱۴۹۴-۲۰۲۰, <https://doi.org/10.1016/j.is.2020.101494>.
- [۳۳] G. Palaniappan, S. Sangeetha, B. Rajendran, S. Goyal, and B. S. Bindhumadhava, "Malicious domain detection using machine learning on domain name features, host-based features and web-based features," *Procedia Computer Science*, vol. ۱۷۱, pp. ۶۵۴-۶۶۱, ۲۰۲۰, <https://doi.org/10.1016/j.procs.2020.04.071>.
- [۳۴] K. A. Messabi, M. Aldwairi, A. A. Yousif, A. Thoban, and F. Belqasmi, "Malware detection using dns records and domain name features," In *Proceedings of the 12nd International Conference on Future Networks and Distributed Systems*, pp. ۱-۷, ۲۰۱۸, doi: 10.1145/3231053.3231082.
- [۳۵] W. Bo, Z. B. Fang, L. X. Wei, Z. F. Cheng, Z. X. Hua, "Malicious URLs detection based on a novel optimization algorithm," *IEICE TRANSACTIONS on Information and Systems*, vol. 101(۴), pp. ۵۱۳-۵۱۶, ۲۰۲۱, doi: 10.1587/transinf.2020EDL114۷.
- [۳۶] J. Yuan, G. Chen, S. Tian, and X. Pei, "Malicious URL detection based on a parallel neural joint model," *IEEE Access*, vol. ۹, pp. ۹۴۶۴-۹۴۷۲, ۲۰۲۱, doi: 10.1109/access.2021.304962۵.
- [۳۷] S. He, B. Li, H. Peng, J. Xin, and E. Zhang, "An effective cost-sensitive XGBoost method for malicious URLs detection in imbalanced dataset," *IEEE Access*, vol. ۹, pp. ۹۳۰۸۹-۹۳۰۹۶, ۲۰۲۱, doi: 10.1109/access.2021.3093۰۹۴.
- [۳۸] S. Kumi, C. Lim, S. G. Lee, "Malicious url detection based on associative classification," *Entropy*, vol. ۲۳(۲), p. ۱۸۲, ۲۰۲۱, doi: 10.3390/e23020182.
- [۳۹] Z. Chen, Y. Liu, C. Chen, M. Lu, and X. Zhang, "Malicious url detection based on improved multilayer recurrent convolutional neural network model," *Security and Communication networks*, no. ۱, p. ۹۹۹۴۱۲۷, ۲۰۲۱, <https://doi.org/10.1155/2021/9994127>.
- [۴۰] R. Patgiri, A. Biswas, S. Nayak, "deepBF: Malicious URL detection using learned bloom filter and evolutionary deep learning," *Computer Communications*, vol. ۲۰۰, pp. ۳۰-۴۱, ۲۰۲۳, <https://doi.org/10.1016/j.comcom.2022.12.02۷>.
- [۴۱] Broadcom, "URL Risk Levels," <https://knowledge.broadcom.com/external/article/1۷۵۵۹/url-risk-levels.html>
- pp. ۱-۱۳, ۲۰۱۶, doi: 10.1001.2322434۷.139۵.۴.۲.1.۵ (in persian)
- [۴۲] X. Lyu, Y. Ding, S. H. Yang, "Safety and security risk assessment in cyber- physical systems," *IET Cyber- Physical Systems: Theory & Applications*, vol. ۴, no. ۳, pp. ۲۲۱-۲۳۲, ۲۰۱۹, <https://doi.org/10.1049/iet-cps.2018.506۸>.
- [۴۳] C. S. Gates, N. Li, H. Peng, B. Sarma, Y. Qi, R. Potharaju, C. Nita-Rotaru, and I. Molloy, "Generating summary risk scores for mobile applications," *IEEE Transactions on dependable and secure computing*, vol. ۱۱, no. ۳, pp. ۲۳۸-۲۵۱, ۲۰۱۴, doi: 10.1109/tdsc.2014.2302293.
- [۴۴] M. Deypir, "Estimating Security Risks of Android Apps Using Information Gain," *Electronic and Cyber Defense*, vol. ۵, no. ۱, pp. ۷۳-۸۳, ۲۰۱۷. (in persian)
- [۴۵] M. Deypir, "Presenting A Method Based on Nearest Neighbors and Hamming Distance in Order to Identify Malicious Applications," *Scientific Journal of Electronical & Cyber Defence*, vol. ۱۱, no. ۲, ۲۰۲۳, doi: 10.1001.2322434۷.1402.1.۲.۷.۰. (in persian)
- [۴۶] M. Deypir, and A. Horri, "Instance based security risk value estimation for Android applications," *Journal of information security and applications*, vol. ۴۰, pp. ۲۰-۳۰, ۲۰۱۸, <https://doi.org/10.1016/j.jisa.2018.02.002>.
- [۴۷] M. Deypir, "Entropy-based security risk measurement for Android mobile applications," *Soft Computing*, vol. ۲۳, no. ۱۶, pp. ۷۳۰۳-۷۳۱۹, ۲۰۱۹, <https://doi.org/10.1007/s00500-018-33۷۷-۵>.
- [۴۸] A. S. Raja, R. Vinodini, and A. Kavitha, "Lexical features based malicious URL detection using machine learning techniques," *Materials Today: Proceedings*, vol. ۴۷, pp. ۱۶۳-۱۶۶, ۲۰۲۱, <https://doi.org/10.1016/j.matpr.2021.04.041>.
- [۴۹] M. Kuyama, Y. Kakizaki, R. Sasaki, "Method for detecting a malicious domain by using whois and dns features," In *Proceedings of the Third International Conference on Digital Security and Forensics (DigitalSec2019)*, Kuala Lumpur, Malaysia, pp. ۶-۸, ۲۰۱۶.
- [۵۰] M. S. I. Mamun, M. A. Rathore, A. H. Lashkari, N. Stakhanova, and A. A. Ghorbani, "Detecting malicious urls using lexical analysis," In *International Conference on Network and System Security*, pp. ۴۶۷-۴۸۲. Springer, Cham, ۲۰۱۶, doi: 10.1007/978-۳-۳۱۹-۴۶۲۹۸-1_۳۰.

- large-scale active DNS measurements,” IEEE journal on selected areas in communications, vol. ۳۴, no. ۶, pp. ۱۸۷۷-۱۸۸۸, ۲۰۱۶, doi: ۱۰.۱۱۰۹/jsac.۲۰۱۶.۲۵۵۸۹۱۸.
- [۴۷] L. Qu, Y. Pei, “A Comprehensive Review on Discriminant Analysis for Addressing Challenges of Class-Level Limitations, Small Sample Size, and Robustness,” *Processes*, vol. ۱۲, no. ۷, p. ۱۳۸۲, ۲۰۲۴, <https://doi.org/10.3390/pr12071382>.
- [۴۸] Github, “Google Web Risk,” <https://github.com/google/webrisk>.
- [۴۹] M. Deypir, T. Zoughi, “Novel Security Metrics for Identifying Risky Unified Resource Locators (URLs),” *Iranian Journal of Science and Technology, Transactions of Electrical Engineering*, pp. ۱-۱۹, ۲۰۲۴, doi: ۱۰.۱۰۰۷/s۴۰۹۹۸-۰۲۳-۰۰۶۹۰-x.
- [۴۰] A. Tharwat, T. Gaber, A. Ibrahim, A. E. Hassanien, “Linear discriminant analysis: A detailed tutorial,” *AI communications*, vol. ۳۰, no. ۲, pp. ۱۶۹-۱۹۰, ۲۰۱۷, doi: ۱۰.۳۲۳۳/aic-۱۷۰۷۲۹.
- [۴۱] Kaggle, “Malicious URL Detection using MLP,” <https://www.kaggle.com/code/ashisharya-1/malicious-url-detection-using-mlp-99-accuracy/data?select=urldata.csv>
- [۴۲] R. van Rijswijk-Deij, M. Jonker, A. Sperotto, and A. Pras, “A high-performance, scalable infrastructure for