

Malicious behavior detection in grant-free access mechanisms for 5G cellular networks

Maryam Keyvani, Behrouz Shahgholi Ghahfarokhi*, Amin Jamali

¹Faculty of Computer Engineering, University of Isfahan, Isfahan, Iran
E-mails: maryam.keyvani1996@eng.ui.ac.ir; shahgholi@eng.ui.ac.ir; a.jamali@alumni.iut.ac.ir
* means corresponding author

Abstract

One of the services that 5G provides is URLLC. To meet the requirements of URLLC, Grant-Free access schemes have been proposed. In these schemes, the radio resources are utilized by User Equipments without any reservation. In these scenarios, we may see malicious behaviors such as the inclination of UEs to reduce their own latency with selfish behaviors or possibility of having misbehaving nodes that degrade the QoS of others. Therefore, detecting malicious users is important in these schemes. In this research, we seek to detect misbehaviors in K-repetition based slotted-ALOHA grant-free access scheme. In previous studies, the security of traditional MAC protocols such as ALOHA has been studied. However, the security of grant-free access based on ALOHA mechanism has not been investigated. Deep learning is considered for misbehavior detection and we employ a Long Short-Term Memory (LSTM) model to detect malicious UEs. The simulation results indicate that under stable conditions, the accuracy decreases as the number of nodes increases. Additionally, in the presence of a fixed number of nodes, as the number of malicious nodes increases, the accuracy decreases. Additionally, if malicious nodes behave properly at times and misbehave at other times, detecting such misbehavior is more difficult.

Keywords

Grant-free access, long short-term memory (LSTM), malicious node, ultra-reliable low-latency communications (URLLC), 5G cellular networks

Introduction

One of the services that is provided by 5G cellular technology is Ultra-Reliable Low-Latency Communications (URLLC). To meet URLLC requirements, grant-free (GF) access schemes have been proposed. In GF access methods, the use of radio resources is done directly without prior reservation, which is suitable for scenarios with a large number of devices such as the Internet of Things (IoT) [2]. But in these scenarios, the probability of users trying to reduce their delay with selfish behaviors, or the possibility of having malicious nodes that seek to decrease the quality of service (QoS) of users, increases. For this reason, it is very important to detect malicious or selfish behaviors in these systems.

Proposed Work and Methodology

The proposed method is presented for a grant-free access scheme based on slotted ALOHA with K repetition. The misbehaving nodes try to adjust the access probability of ALOHA and/or the value of K in a selfish or malicious manner. In proposed method, long short-term memory model (LSTM) is used to detect misbehavior of GF users in a 5G system. To this aim, the media access evidences are sequentially windowed, injected into LSTM network model, and the misbehavior/normal behavior of the node is extracted from it as the output. The model has been trained with the help of data extracted from a custom simulation using the mean square error function and Adam's optimizer. The simulation results indicate that under stable conditions, the accuracy decreases as the number of nodes increases. Additionally, in the presence of a fixed number of nodes, as the number of malicious nodes increases, the accuracy decreases. Moreover, if malicious nodes behave properly at times and misbehave at other times, detecting such misbehavior is more difficult.

Conclusion

In this research, the behavior learning approach with the help of deep neural networks employed to detect the misbehavior of devices in grant-free random access in 5G cellular networks. The long short-term memory (LSTM) used to learn the behavior of normal and misbehaving users in an ALOHA-based GF method with K repetitions mechanism. The simulation results show acceptable accuracy of proposed method in detecting misbehaviors. Of course, increasing the number of nodes, increasing the number of available frequency channels, increasing the number of misbehaving nodes, and on-off attack reduce the accuracy of the method. Using other learning methods and investigating more intelligent methods for detection of more complex misbehaviors are recommended as future works.

شناسایی بدرفتاری در روش‌های دسترسی به رسانه بدون تخصیص در شبکه‌های سلولی نسل ۵

مریم کیوانی

دانشجوی کارشناسی ارشد، دانشکده مهندسی کامپیوتر، دانشگاه اصفهان، اصفهان، ایران

بهروز شاهقلی قهفرخی

دانشیار، دانشکده مهندسی کامپیوتر، دانشگاه اصفهان، اصفهان، ایران

امین جمالی

محقق پسادکتری، دانشکده مهندسی کامپیوتر، دانشگاه اصفهان، اصفهان، ایران

چکیده

برای برآورده کردن الزامات ارتباطات فوق‌العاده مطمئن با تاخیر کم در 5G، طرح‌های دسترسی بدون تخصیص (GF) پیشنهاد شده است. در این طرح‌ها، احتمال وجود داشتن گره‌های بدرفتار وجود دارد. به همین دلیل تشخیص بدرفتاری در این طرح‌ها اهمیت دارد. در این پژوهش به دنبال تشخیص بدرفتاری در روش دسترسی بدون تخصیص ALOHA شکاف‌دار مبتنی بر K بار ارسال مکرر هستیم. در پژوهش‌های پیشین، امنیت MAC‌های سنتی مانند ALOHA بررسی شده است؛ اما امنیت دسترسی بدون تخصیص در شبکه‌های نسل ۵ تا کنون بررسی نشده است. در این تحقیق، از یادگیری عمیق مبتنی بر مدل حافظه کوتاه مدت طولانی (LSTM) برای تشخیص بدرفتاری استفاده می‌شود و با تزریق دنباله‌های پنجره‌بندی شده از سوابق دسترسی به کانال به مدل آموزش داده شده، وضعیت یک دستگاه از نظر بدرفتار بودن بررسی می‌شود. مدل بدرفتاری که در این پژوهش برای گره‌های بدرفتار در نظر گرفته شده است، شامل تغییر در احتمال ارسال و همچنین تغییر در تعداد ارسال مکرر پروتکل است. نتایج حاصل از شبیه‌سازی حاکی از قابل قبول بودن دقت راهکار پیشنهادی در شناسایی دستگاه‌های بدرفتار است. دقت روش پیشنهادی در مقایسه با یک روش پایه آماری به طور متوسط، ۳۸٪ بهبود یافته است که این بهبود در ازای هزینه تعلیم مدل یادگیرنده است. همچنین، نتایج نشان می‌دهد که تحت شرایط ثابت محیط، با افزایش تعداد گره‌ها یا افزایش تعداد گره‌های بدرفتار، دقت تشخیص بدرفتاری کاهش می‌یابد. همچنین مشاهده شد که اگر بدرفتاری به صورت خاموش-روشن باشد، تشخیص آن دشوارتر است.

کلمات کلیدی

دسترسی تصادفی بدون تخصیص (GF)، حافظه کوتاه مدت طولانی (LSTM)، گره‌های بدرفتار، ارتباطات فوق‌العاده مطمئن با تاخیر کم (URLLC)، نسل پنجم شبکه‌های سلولی (5G).

نام نویسنده مسئول: بهروز شاهقلی قهفرخی

ایمیل نویسنده مسئول: shahgholi@eng.ui.ac.ir

تاریخ ارسال مقاله: ۱۴۰۲/۰۸/۰۲

تاریخ(های) اصلاح مقاله: ۱۴۰۲/۱۱/۱۶

تاریخ پذیرش مقاله: ۱۴۰۳/۰۲/۰۱

۱- مقدمه

نیاز پشتیبانی می‌کند. نسل پنجم ارتباطات سلولی از امواج میلی‌متری و سلول-های کوچک استفاده می‌کند تا تعداد دستگاه‌های زیاد را پوشش دهد [۱]. در عین فراهم کردن دسترسی انبوه، ارتباطاتی نیز مطرح می‌شوند که نیاز به تاخیر بسیار کم و قابلیت اطمینان بسیار زیاد دارند. یکی از خدماتی که 5G ارائه می‌دهد، ارتباطات فوق‌العاده مطمئن با تاخیر کم (URLLC) است. پشتیبانی از URLLC برای کاربردهای مختلف اینترنت اشیا نظیر سلامت هوشمند، وسایل نقلیه خودران، اتوماسیون کارخانه و شبکه‌های هوشمند انرژی مورد نیاز است. فناوری LTE کنونی به‌سختی می‌تواند نیازهای URLLC را برآورده کند. LTE فعلی از یک روش دسترسی تصادفی مبتنی بر تخصیص (GB)^۲ در لایه MAC استفاده می‌کند که مسبب بخشی از تاخیر است. طرح‌های دسترسی

امروزه زندگی الکترونیک برای انسان‌ها امری اجتناب‌ناپذیر است و اینترنت اشیا (IoT) یکی از ملزومات زندگی الکترونیک است. تصور آینده‌ی ما جامعه‌ای شبکه‌ای است که دسترسی بی‌حد و حصر به اطلاعات و به اشتراک‌گذاری داده‌ها دارد و در همه جا و هر زمان برای همه و همه‌چیز قابل دسترسی است. برای تحقق این تصویر، لازم است که فناوری‌های ارتباطی بی‌سیم برای تکامل مورد بررسی قرار گیرند. نسل پنجم شبکه تلفن همراه (5G)، هم اکنون جدیدترین نسل تجاری شده از سیستم‌های ارتباطات سیار سلولی است که برای تحقق نیازهای IoT مناسب است. 5G یک شبکه چندخدمته است که طیف گسترده‌ای از نیازمندی‌ها و درخواست‌ها را با مجموعه متنوعی از عملکردها و خدمات مورد

² Grant-based

¹ Ultra-Reliable Low-Latency Communications

کرده‌اند [۲۴، ۲۵]، اما امنیت روش ALOHA و شکافدار مبتنی بر K بار ارسال مکرر، از منظر بدرفتاری گره‌ها مورد توجه قرار نگرفته است.

با مرور کارهای پیشین، مشاهده شد که امنیت دسترسی بدون تخصیص در شبکه‌های سلولی به‌صورت محدود مورد توجه بوده است و امنیت این روش‌ها از منظر تشخیص و مقابله با بدرفتاری دستگاه‌ها بررسی نشده است. اگر چه در گذشته کارهایی در زمینه تشخیص و مقابله با بدرفتاری کاربران در روشهای دسترسی تصادفی کلاسیک نظیر ALOHA صورت پذیرفته است، اما روشهای دسترسی بدون تخصیص در شبکه‌های سلولی نسل ۵ و ۶ از طیف متنوع تری از تکنیکها نظیر ارسال مکرر نیز در کنار روشهای دسترسی تصادفی کلاسیک بهره می‌برند که شناسایی بدرفتاری را پیچیده‌تر می‌کند. در پژوهش حاضر سعی شده است با تمرکز بر نوآوری‌ها و مزایای پژوهش‌های پیشین، به تشخیص بدرفتاری دستگاه‌ها در دسترسی بدون تخصیص به رسانه بی‌سیم مبتنی بر ALOHA با K بار ارسال مکرر پرداخته شود. مدل بدرفتاری که در این پژوهش برای گره‌های بدرفتار در نظر گرفته شده شامل دستکاری در احتمال ارسال تنظیم شده توسط پروتکل و همچنین دستکاری در تعداد دفعات ارسال مکرر است که می‌تواند با مقاصد خودخواهانه یا بدخواهانه صورت پذیرد.

ادامه بخش‌های این مقاله به شرح زیر سازماندهی شده است. در بخش ۲ برخی مفاهیم و مقدمات دسترسی به رسانه مطرح می‌گردد. سپس در بخش ۳ مدل سیستم شرح داده می‌شود. در بخش ۴ روشی برای تشخیص بدرفتاری پیشنهاد می‌شود. سپس در بخش ۵ ارزیابی روش پیشنهادی مطرح می‌شود و در نهایت در بخش ۶ نتیجه‌گیری بیان می‌شود.

۲- روش‌های مدیریت دسترسی به رسانه در شبکه سلولی

در این بخش به معرفی روش‌های دسترسی به رسانه در شبکه‌های سلولی قدیم و جدید پرداخته می‌شود.

۲-۱- دسترسی چندگانه مبتنی بر تخصیص

در شبکه‌های سلولی نسل ۴ و ماقبل، انتقال با تخصیص منبع قابل انجام است؛ یعنی تجهیزات کاربر به شرط دریافت یک تخصیص زمان‌بندی شده (SG)^۷ از یک ایستگاه پایه خدمت‌دهنده^۸ (BS)، می‌توانند داده‌ی خود را انتقال دهند. مراحل دسترسی به رسانه GB به شرح زیر (نشان داده شده در شکل ۱) است [۲۶]:

۱. UE داده‌های جدیدی را از لایه بالا برای انتقال دریافت می‌کند.
۲. سپس منتظر فرصت مشخص شده برای ارسال درخواست زمان‌بندی (SR)^۹ می‌ماند.
۳. UE یک درخواست زمان‌بندی (SR) را به ایستگاه پایه می‌فرستد.
۴. در صورت موفقیت ارسال، ایستگاه پایه یک تخصیص زمان‌بندی شده را برای او ارسال می‌کند.
۵. UE داده‌ها را در منبع تخصیص یافته منتقل می‌کند.

۲-۲- دسترسی چندگانه بدون تخصیص

در این روش تجهیزات کاربر بدون نیاز به دریافت SG، داده‌های خود را مستقیماً روی رسانه مشترک انتقال می‌دهند؛ که یک راه حل امیدوارکننده برای خدمات‌های حساس به تأخیر است [۵]. تفاوت این روش با روش قبل در شکل ۱ نشان داده شده است. رویکردهای مختلفی برای دسترسی GF وجود دارد که برخی از آن‌ها در ادامه تشریح می‌شود.

بدون تخصیص (GF)^۳ به‌عنوان یک فناوری عملی و امیدوارکننده برای تأمین چنین شرایطی، پیشنهاد می‌شود [۲]. همچنین طرح‌های دسترسی بدون تخصیص همراه با طرح‌های انتقال غیرمتعامد، یک راه‌حل امیدوارکننده برای اینترنت اشیاء در 6G است [۳، ۴]. در این روش تجهیزات کاربر بدون نیاز به دریافت اجازه و رزرو منابع، داده‌های خود را در منابع رادیویی مشترک انتقال می‌دهند. طرح‌های دسترسی GF، تمام تأخیرهای ناشی از فرایند دست‌دهی^۴ موجود در روش دسترسی GB را حذف می‌کنند. همچنین مصرف انرژی تجهیزات کاربر (UE)^۵ را بهبود می‌بخشند، پیچیدگی آن‌ها را کاهش می‌دهند و سربار سیگنالینگ را در مقایسه با روش دسترسی مبتنی بر تخصیص کاهش می‌دهند [۵].

در دنیای پرسرعت سیستم‌های ارتباط سلولی نسل ۵ و ۶، مدیریت رفتار درست در دسترسی به رسانه یک چالش جدید است. در واقع، تخصیص کارآمد منابع رادیویی مشترک برای اطمینان از انتقال عادلانه و قابل اعتماد داده‌ها در میان کاربران متعدد بسیار مهم است. در اولین روش‌های دسترسی به رسانه غیر متمرکز، سیستم‌ها به شیوه‌ای ناهماهنگ عمل می‌کردند که منجر به استفاده ناکارآمد از منابع مشترک و نرخ برخورد بالا در میان گره‌های انتقال می‌شد. در دهه ۱۹۷۰، دانشگاه هاوایی پیشگام یک روش دستیابی به رسانه به نام ALOHA بود. این پروتکل دسترسی ناهماهنگ، به کاربران این امکان را می‌داد که هر زمان که اطلاعاتی برای ارسال داشتند، بدون نیاز به هماهنگی متمرکز، داده‌ها را منتقل کنند. با این حال، سادگی ALOHA با چالش‌هایی همراه بود، زیرا از برخوردهای مکرر رنج می‌برد که منجر به از دست دادن قابل توجه بسته و کاهش کارایی سیستم می‌شد. سپس برای غلبه بر محدودیت‌های ALOHA، مفهوم ALOHA شکافدار^۶ پیشنهاد شد. ALOHA شکافدار با همگام‌سازی گره‌های انتقال و تخصیص شکاف‌های زمانی برای انتقال، نرخ برخورد را به‌طور قابل‌توجهی کاهش داد و توان کلی سیستم را بهبود بخشید. در طول توسعه ALOHA و ALOHA شکافدار، محققان استراتژی‌های مختلفی را برای کاهش رفتارهای نادرست و رقابت‌ها در دسترسی به رسانه بررسی کردند. تکنیک‌هایی مانند عقب‌نشینی نمایی، تصادفی‌سازی و مکانیسم‌های تصدیق برای افزایش عدالت و بهینه‌سازی استفاده از منابع معرفی شدند. اما امروزه در شبکه‌های سلولی نسل ۵ و ۶، ترکیبی از این ایده‌ها در کنار روش‌هایی نظیر ارسال مکرر بسته، برای مدیریت دسترسی بدون تخصیص به رسانه استفاده می‌شود.

شناسایی بدرفتاری در دسترسی به رسانه مشترک به دلایلی ازجمله استفاده کارآمد از منابع، انصاف و برابری، و افزایش توان عملیاتی، از اهمیت بالایی برخوردار است. ازجمله مواردی که تشخیص بدرفتاری مصداق و اهمیت پیدا می‌کند، طرح‌های دسترسی GF در شبکه‌های 5G است. زیرا تعداد دستگاه‌ها بسیار زیاد است و تمایل کاربران به کاهش تأخیر خود با رفتارهای خودخواهانه و یا ایجاد اختلال برای سایرین، ریسک احتمال وجود داشتن گره‌های بدرفتار را بیشتر می‌کند. از این رو، در این پژوهش به دنبال ارائه راهکاری برای تشخیص بدرفتاری در روش‌های دسترسی بدون تخصیص هستیم. بررسی پیشینه پژوهش‌های انجام شده در زمینه دسترسی بدون تخصیص به رسانه در 5G [۵-۱۰] نشان می‌دهد که تا کنون تشخیص بدرفتاری مورد توجه مطالعات مذکور نبوده است. اگرچه امنیت MAC‌های سنتی مانند ALOHA و ALOHA شکافدار در کارهای پیشین [۱۱-۲۳] مورد توجه بوده است و پژوهش‌هایی نیز امنیت دسترسی بدون تخصیص را به‌صورت محدود بررسی

⁷ Scheduling Grant

⁸ Base Station

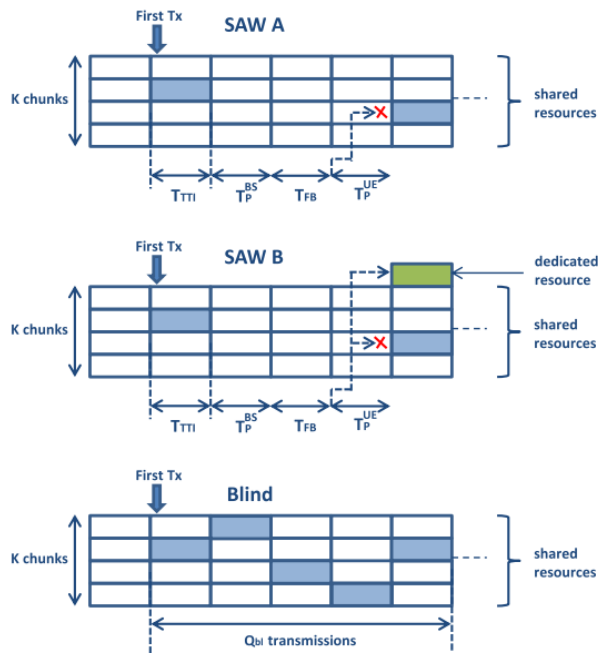
⁹ Scheduling Request

³ Grant Free

⁴ handshaking

⁵ User Equipment

⁶ Slotted



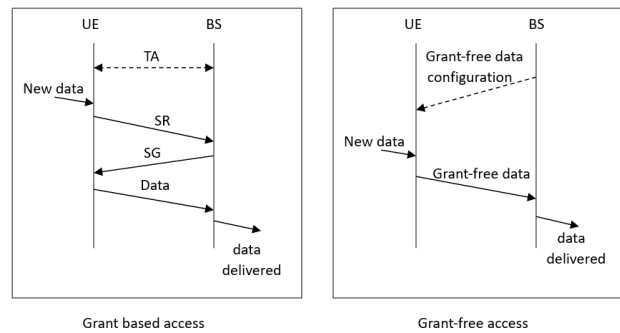
شکل ۲. نمودار زمانی طرح‌های بدون تخصیص توقف و انتظار و ارسال مکرر کور [۲۴]

بدین صورت که برای اطمینان از ارسال بسته‌ی مورد نظر، بسته چندین بار ارسال می‌شود. در این تحقیق بر دسترسی بدون تخصیص مبتنی بر ALOHA با K ارسال مکرر تمرکز شده است. روش مبتنی بر ALOHA به این دلیل انتخاب شده است که نسبت به روش‌های دیگر پرباربردتر بوده و تامین امنیت آن در دسترسی بدون تخصیص قبلاً مطالعه نشده است.

۳- مدل سیستم

روش پیشنهادی برای تشخیص وجود بدرفتاری، با فرض دسترسی بدون تخصیص مبتنی بر ALOHA شکاف‌دار با K ارسال مکرر ارائه شده است. در پروتکل ALOHA شکاف‌دار سنتی، یک ایستگاه به محض دریافت یک بسته جدید، آن را ارسال می‌کند. اگر برخوردی اتفاق بیفتد، بسته در صف ایستگاه ذخیره می‌شود و در شکاف‌های بعدی ارسال می‌شود؛ که در این روش احتمال برخورد میان بسته‌ها بالا است. در این تحقیق به منظور کاهش برخورد میان بسته‌ها، این پروتکل را طوری تغییر می‌دهیم که همه‌ی ایستگاه‌ها همیشه بسته‌های خود را با احتمال p ارسال کنند. همچنین به منظور اطمینان از ارسال موفقیت آمیز بسته، ایستگاه‌ها بسته‌های خود را K بار ارسال می‌کنند.

مدل سیستم با الهام از مدل موجود در مراجع [۲۸، ۷] طراحی می‌شود. در این سیستم، چندین UE یک کانال مشترک را به اشتراک می‌گذارند و بسته‌های داده‌ی با محدودیت تاخیر را از طریق کانال مشترک به یک ایستگاه پایه (BS) ارسال می‌کنند. همه‌ی بسته‌ها اندازه یکسانی دارند و اگر یک بسته طبق محدودیت تاخیرش به گیرنده تحویل داده نشود، منقضی شده و از سیستم حذف می‌شود. مشابه [۷]، کاربران یک باند فرکانسی W هر تری را به اشتراک گذاشته‌اند که این پهنای باند به تعدادی کانال فرکانسی با اندازه یکسان تقسیم شده است. همچنین زمان به شکاف‌هایی با اندازه یکسان تقسیم شده است. مشابه [۲۸]، زمان هر شکاف، شامل مدت زمان ارسال بسته از یک ایستگاه به گیرنده و دریافت بازخورد از گیرنده در مورد تصدیق دریافت بسته است. ایستگاه‌ها می‌توانند در هر شکاف بسته‌ی جدید برای ارسال داشته باشند. هنگامی که یک گره بسته‌ای برای ارسال داشته باشد، یکی از کانال‌های فرکانسی



شکل ۳. روش‌های دسترسی مبتنی بر تخصیص و بدون تخصیص

روش توقف و انتظار (SAW)

در این روش به محض آماده شدن بسته، UE تلاش می‌کند که بسته خود را روی یکی از منابع موجود و بدون اخذ اجازه از BS، ارسال کند. پس از انتقال، UE برای مدت زمانی منتظر دریافت یک پیام تصدیق از سوی BS می‌ماند. انتقال بدون تخصیص، در صورت شناسایی صحیح پیام UE، موفقیت آمیز است. لذا تصدیق دریافت (Ack) از BS صادر می‌شود. دو مدل توقف و انتظار مطرح شده است:

۱. SAW A (ارسال مجدد روی منابع مشترک): BS در صورت عدم انتقال موفقیت آمیز، بازخوردی صادر نمی‌کند. لذا UE پس از گذشت یک بازه‌ی زمانی از ارسال اول، بار دیگر بسته‌ی خود را بر روی منابع اشتراکی ارسال می‌کند. منابع اشتراکی منابعی است که بین گره‌ها به اشتراک گذاشته شده است و گره‌ها برای دسترسی به این منابع با یکدیگر رقابت می‌کنند.
۲. SAW B (ارسال مجدد روی منابع اختصاصی): در این شیوه، BS قادر به صدور SG برای ارسال مجدد بسته ناموفق است. در صورتی این اتفاق می‌افتد که فعالیت UE به صورت صحیح شناسایی شود، اما تشخیص داده‌ی آن موفقیت آمیز نباشد. در این صورت، ارسال مجدد به صورت برنامه‌ریزی شده بر روی منابع اختصاصی بدون تداخل، رخ خواهد داد. در صورت عدم شناسایی UE، هیچ SG با ACK توسط BS صادر نمی‌شود. لذا UE بعد از یک بازه‌ی زمانی، ارسال مجدد روی منابع مشترک را مشابه طرح SAW A انجام خواهد داد.

در هر دو مدل تعداد دفعات ارسال مجدد محدود است [۲۶].

روش ارسال مکرر کور

در این روش UE چند مرتبه بسته‌ی خود را با فواصل زمانی معین روی منابع مشترک ارسال می‌کند؛ و نیازی به دریافت بازخورد از سوی BS نیست. اگر بخواهیم این روش را با روش توقف و انتظار از نظر زمانی مقایسه کنیم، آخرین ارسال کور همزمان با دومین ارسال از روش SAW است (در شکل ۲ نشان داده شده است) [۲۶].

روش‌های مبتنی بر ALOHA

از دیگر روش‌های مورد استفاده برای دسترسی بدون تخصیص، ALOHA و ALOHA شکاف‌دار می‌باشند که به شیوه‌های گوناگون استفاده شده‌اند و در مراجع [۳-۸] بررسی شده‌اند. در بسیاری از مقالات موجود برای این روش، برای مقابله با برخورد میان بسته‌ها و تامین سطح اطمینان مورد نیاز، روش‌های مبتنی بر K ارسال مکرر را در کنار آن‌ها به کار گرفته‌اند.

در این تحقیق از شبکه‌های عصبی عمیق (DNN)^{۱۰} برای تشخیص بدرفتاری استفاده خواهد شد. اگرچه مدل‌های مختلفی از شبکه عصبی عمیق ارائه شده است، لیکن از آنجایی که رفتار دسترسی به رسانه‌ی یک دستگاه، در گذر زمان قابل تحلیل است، شناسایی گره بدرفتار به کمک شبکه‌های عصبی عمیق دارای حافظه، امکان‌پذیر است و شبکه‌های عصبی بازگشتی یا مکرر در این خصوص مطرح هستند. شبکه عصبی بازگشتی نوعی شبکه عصبی است که حافظه‌ی داخلی دارد؛ ولی شبکه عصبی بازگشتی در مسائلی که نیاز به حافظه‌ی بلند مدت داشته باشند، نمی‌تواند خوب عمل کند. در واقع شبکه عصبی بازگشتی حافظه‌ی کوتاه مدت خوبی دارد، اما حافظه‌ی بلند مدت ندارد. در عوض، شبکه عصبی حافظه کوتاه مدت طولانی (LSTM)^{۱۱} نوع خاصی از شبکه‌های عصبی بازگشتی است که نیاز به حافظه‌ی بلند مدت شبکه‌ی بازگشتی را حل می‌کند و قادر به یادگیری وابستگی‌های طولانی مدت است.

مدل شبکه عصبی عمیق استفاده شده در این پژوهش، یک شبکه LSTM با دو لایه مخفی است که از تابع فعال‌سازی $\tanh^{12}$ استفاده می‌کند. این تابع به صورت زیر تعریف می‌شود:

$$\tanh(x) = (e^x - e^{-x}) / (e^x + e^{-x}) \quad (1)$$

رفتار هر کاربر به صورت دنباله‌ای از اعداد (با مقادیر تعریف شده‌ی ۰، ۱ و ۲) می‌باشد، که هر عدد از دنباله، بیانگر وضعیت ارسال کاربر در یکی از شکاف‌های زمانی است. با فرض وجود Ack، تشخیص وضعیت تلاش هر کاربر ممکن است. مقدار صفر به معنی تلاش مواجهه شده با شکست در یک شکاف، مقدار ۱ به معنی ارسال بسته موفقیت‌آمیز در یک شکاف، و مقدار ۲ به معنی عدم تلاش برای ارسال در آن شکاف است. هر کدام از این دنباله‌های اعداد که رفتار کاربر را نشان می‌دهند، دارای یک برچسب می‌باشد که بدرفتاری یا خوشرفتاری آن کاربر را در داده آموزشی مشخص می‌کند. به منظور تامین ورودی شبکه LSTM مد نظر، مراحل زیر انجام می‌شود:

(۱) رفتارها که اعداد ۰، ۱ و ۲ هستند، به صورت زیر به مقادیر دودویی تبدیل می‌شوند.

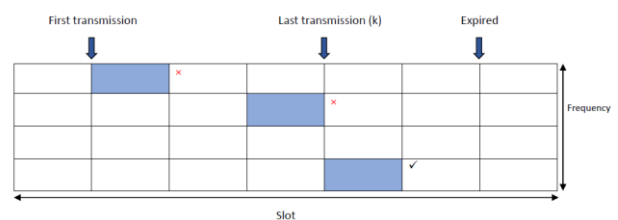
0 → [0, 1, 0]
1 → [1, 0, 0]
2 → [0, 0, 1]

(۲) دنباله ورودی پنجره‌بندی می‌شود. اندازه پنجره‌ی مذکور W در نظر گرفته شده و هر بار این پنجره‌ی کشویی S تا جابجا می‌شود و به شبکه LSTM تزریق می‌شود. هر پنجره از داده استفاده شده برای آموزش، دارای برچسب تعیین‌کننده بدرفتاری یا خوشرفتاری کاربر مربوطه می‌باشد که خروجی مورد انتظار شبکه LSTM را معین می‌کند. لذا لایه ورودی LSTM دارای ابعاد $W * I$ می‌باشد که I تعداد ویژگی‌های هر عمل دسترسی به رسانه در مجموعه داده است. در مدل تهدید اول و دوم یعنی ارسال بسته با احتمال $P=1$ و $P=0.9$ تعداد ویژگی‌های مذکور در مجموعه داده، ۳ ویژگی می‌باشد. ولی در مدل تهدید سوم یعنی جعل در تعداد ارسال مکرر K ، تعداد ویژگی‌های مجموعه داده ۴ ویژگی در نظر گرفته شده است. در واقع، به منظور تشخیص بدرفتاری جعل در مقدار K ، شماره ترتیب^{۱۳} بسته نیز به مجموعه ویژگی‌های هر تلاش برای دسترسی به رسانه، اضافه شد.

لایه خروجی شبکه عصبی طراحی شده، دارای یک نورون است که تعیین‌کننده بدرفتاری یا خوشرفتاری کاربر است. میزان دقت عملکرد راهکار پیشنهادی بر اساس میزان تشخیص‌های درست و به کمک شبه کد الگوریتم روی مجموعه داده آزمایشی محاسبه می‌شود. به منظور بررسی عملکرد روش

را به صورت تصادفی برای ارسال انتخاب می‌کند. ایستگاه‌ها بسته‌های خود را در ابتدای شکاف با احتمال p ارسال می‌کنند و ایستگاه در پایان شکاف می‌داند که بسته با موفقیت به گیرنده تحویل داده شده است یا خیر (بر اساس تصدیق دریافتی). اگر بیش از یک ایستگاه در یک شکاف اقدام به ارسال بسته خود کنند، براساس کانال‌های فرکانسی که به صورت تصادفی انتخاب کرده‌اند ممکن است برخورد پیش بیاید؛ یعنی اگر دو یا چند ایستگاه، کانال فرکانسی مشترکی را انتخاب کرده باشند، برخورد رخ می‌دهد و بسته‌های این ایستگاه‌ها از بین می‌روند. در صورت رخ دادن برخورد، ایستگاه‌ها می‌توانند بسته‌های خود را در شکاف‌های بعدی مجدداً ارسال کنند. تعداد ارسال‌های مجدد برای هر بسته حداکثر K می‌باشد که به صورتی تعیین می‌شود که از مدت زمان محدودیت تاخیر بسته‌ها گذر نکند.

در شکل ۳ یک UE تلاش می‌کند بسته‌ی خود را از طریق کانال مشترک به ایستگاه پایه ارسال کند. هر بسته دارای محدودیت تاخیر ۵ شکاف زمانی می‌باشد و پس از گذشت ۵ شکاف از زمان ورود بسته به سیستم، بسته منقضی شده و از سیستم حذف می‌شود. همچنین UE هر بسته‌ی خود را حداکثر $K=3$ بار می‌تواند ارسال کند. در این شکل به محض ورود بسته، UE تلاش می‌کند که بسته‌ی خود را ارسال کند (با احتمال p). در تلاش اول با شکست مواجه می‌شود. در شکاف‌های بعدی نیز تلاش می‌کند که بسته‌ی خود را با احتمال p ارسال کند. سرانجام در تلاش سوم که آخرین تلاش او می‌باشد موفق می‌شود بسته را با موفقیت برای گیرنده ارسال کند.



شکل ۳. اشتراک ۴ کانال بین N کاربر با محدودیت تاخیر ۵ شکاف زمانی و حداکثر تکرار $K=3$

مدل بدرفتاری که در این پژوهش در نظر گرفته شده شامل ۳ نوع بدرفتاری است. بدین صورت که گره بدرفتار: (۱) به جای ارسال با احتمال p ، هر زمان که بسته داشته باشد، بسته‌ی خود را ارسال کند؛ و در واقع p را معادل ۱ در نظر بگیرد؛ (۲) به جای ارسال با احتمال p ، بسته‌ی خود را با احتمال ثابت نزدیک به یک (یعنی ۰.۹) ارسال کند؛ و (۳) تعداد ارسال مکرر K را بیشتر از مقدار توافق شده قرار دهد.

۴- روش پیشنهادی برای تشخیص بدرفتاری

روش‌های تشخیص بدرفتاری در MAC به استفاده از تکنیک‌های مختلف برای شناسایی و جلوگیری از بدرفتاری در پروتکل‌های لایه MAC اشاره دارند. نظریه بازی‌ها، روش‌های آماری و احتمالات، سه نمونه از روش‌های شناسایی و مقابله با بدرفتاری در MAC هستند که می‌توانند برای شناسایی و جلوگیری از بدرفتاری در شبکه‌های سلولی هم استفاده شوند. نمونه‌ای دیگر از روش‌های شناسایی بدرفتاری، روش‌های یادگیری عمیق می‌باشند. در پژوهش‌های متعددی استفاده از روش‌های هوشمند و یادگیر برای کشف بدرفتاری مورد توجه بوده است [۲۹، ۳۰].

¹² hyperbolic tangent

¹³ sequence number

¹⁰ Deep Neural Network

¹¹ Long Short-Term Memory

$$TNR = \frac{TN}{N} = \frac{TN}{TN + FP} \quad (4)$$

برای پیاده‌سازی الگوریتم تشخیص بدرفتاری مبتنی بر LSTM از زبان پایتون و کتابخانه‌های pandas و torch استفاده شده است. مدل LSTM استفاده شده در ارزیابی‌ها دو لایه مخفی دارد که هر یک از این دو لایه مخفی، ۱۰۰ نورون دارند. ارزیابی‌ها برای هر یک از ۳ مدل تهدید ارایه شده در بخش ۴ جداگانه انجام شده است. در مدل تهدید اول و دوم، یعنی ارسال با احتمال ۱ و ۰.۹، دنباله ورودی با پنجره‌ای به اندازه $W = 10$ پنجره‌بندی می‌شود و هر بار این پنجره کشویی $S = 5$ تا جابجا می‌شود و لذا لایه ورودی 10×3 نورون دارد. در مدل تهدید سوم، یعنی افزایش تعداد ارسال مکرر (K) به مقداری بیشتر از مقدار توافق شده، دنباله ورودی با پنجره‌ای به اندازه $W = 50$ پنجره‌بندی می‌شود. لذا لایه ورودی نیز 50×4 نورون دارد. مجموعه داده به دو بخش آموزش و آزمون تقسیم شده است که فرآیند یادگیری مدل روی قسمت آموزش آن انجام می‌شود. مدل برای ۱۵۰ تکرار با استفاده از تابع میانگین مربعات خطا (MSELoss) و بهینه‌ساز آدام با نرخ آموزش ۰.۰۰۱، آموزش داده شده است.

همچنین روش پیشنهادی با یک طرح تشخیص بدرفتاری کلاسیک مبتنی بر مشخصه‌های آماری ترافیک مقایسه شده است. در روش تشخیص بدرفتاری مبتنی بر مشخصه‌های آماری، میانگین، میانه و واریانس نمونه ترافیکهای آموزشی خوش‌رفتار و بدرفتار محاسبه می‌شود و بر اساس آن، آستانه‌هایی برای هر معیار آماری تنظیم می‌شود. این آستانه‌ها برای طبقه‌بندی گره‌ها (به‌عنوان خوش‌رفتار یا بدرفتار) بر اساس ترافیک آنها استفاده می‌شود. به این صورت که ویژگی‌های آماری مذکور (میانگین، میانه و واریانس) از ترافیک واقعی استخراج می‌شوند و با آستانه‌های تعیین شده در مرحله قبل مقایسه می‌شوند و براساس این مقایسه، یک گره به‌عنوان خوش‌رفتار یا بدرفتار طبقه‌بندی می‌شود.

۵-۱- نتایج ارزیابی با فرض مدل تهدید اول (ارسال بسته با احتمال $P=1$)

در بخش اول از ارزیابی‌ها، مدل تهدید اول که در آن گره‌های بدخواه با احتمال ۱ بسته ارسال می‌کنند، در نظر گرفته شده است. در ابتدا نحوه عملکرد روش پیشنهادی برحسب تعداد گره‌های بدرفتار بررسی می‌شود. در نمودار شکل ۴ تمامی شرایط محیط اعم از تعداد شکاف‌ها، تعداد کانال‌های فرکانسی، زمان منقضی شدن بسته‌ها و تعداد مجاز ارسال مکرر هر بسته ثابت بوده (مطابق جدول ۱) و صرفاً تعداد گره‌ها از ۱۰ تا ۱۰۰ افزایش یافته است. در همه این سناریوها، تعداد گره‌های بدرفتار در نظر گرفته شده، ۲۰ درصد از تعداد کل گره‌ها است.

جدول ۱. مقادیر ثابت لحاظ شده برای پارامترها در شبیه‌سازی‌ها

تعداد شکاف‌ها	۱۰۰
تعداد مجاز ارسال مکرر بسته	۱۰
زمان منقضی شدن بسته	۱۵ شکاف
تعداد کانال‌های فرکانسی	۴۰

همان‌طور که از نمودارها مشخص است وقتی شرایط محیط ثابت باشد، با افزایش تعداد گره‌ها دقت مدل اندکی افزایش می‌یابد و سپس مجدداً کاهش می‌شود. افزایش اولیه به این علت می‌باشد که با افزایش گره‌ها، داده‌های

پیشنهادی، از معیار دقت^{۱۴} استفاده می‌شود، چرا که عملکرد یک سیستم تشخیص‌دهنده بدرفتاری را در برابر نمونه‌های مثبت و منفی به‌خوبی نشان می‌دهند.

الگوریتم ۱: ارزیابی عملکرد روش پیشنهادی

ورودی: مدل آموزش دیده شده با LSTM (M) و مجموعه داده آزمایشی (D)

خروجی: ماتریس درهم‌ریختگی

مراحل زیر را N_{epochs} بار تکرار کن:

- یک لیست جدید به نام $predictions$ ایجاد کن
- مراحل زیر را $N_{batches}$ بار تکرار کن:
 - o با استفاده از مدل آموزش دیده‌ی M ، خروجی را برای یک دسته منتخب (b) از داده‌های D به دست آور
 - o خروجی دسته b را به لیست $predictions$ اضافه کن
- $n_{correct}$ را برابر با مجموع تعداد خروجی‌هایی که به‌درستی پیش‌بینی شده‌اند قرار بده
- دقت را به صورت $((n_{correct} / len(predictions)) * 100)$ محاسبه کن

پایان

۵- ارزیابی

برای تعلیم مدل LSTM طراحی شده در این پژوهش، ابتدا مدل سیستم پیشنهادی در یک شبیه‌ساز توسعه داده شده توسط نویسندگان، شبیه‌سازی شده و مجموعه داده‌ای از رفتار کاربران درست‌کار و بدخواه از آن استخراج شده است. در این شبیه‌سازی، فرض شده است که ورود بسته‌های کاربران از توزیع پواسن تبعیت می‌کند. همان‌طور که گفته شد، رفتار هر کاربر به‌صورت دنباله‌ای از اعداد نشان داده می‌شود که هر رقم از آن دنباله، وضعیت کاربر در یک شکاف را با یکی از مقادیر ۰، ۱ و ۲ برای تلاش مواجه شده با شکست، ارسال موفقیت-آمیز، و عدم تلاش برای ارسال، نشان می‌دهد. هر دنباله بسته به کاربر تولید کننده آن، یک برچسب بدخواه بودن یا نبودن را هم می‌گیرد.

به منظور بررسی عملکرد سیستم‌های طبقه‌بند، از ماتریس درهم‌ریختگی استفاده می‌شود. در این ماتریس ۴ معیار به شرح زیر محاسبه داده می‌شود:

- مثبت واقعی (TP): گره‌های بدرفتار با موفقیت توسط مدل شناسایی شوند.
- منفی واقعی (TN): گره‌های خوش‌رفتار به‌درستی توسط مدل شناسایی شوند.
- مثبت کاذب (FP): گره‌های خوش‌رفتار به اشتباه به‌عنوان گره بدرفتار شناسایی شوند.
- منفی کاذب (FN): گره‌های بدرفتار به اشتباه به‌عنوان گره خوش-رفتار شناسایی شوند.

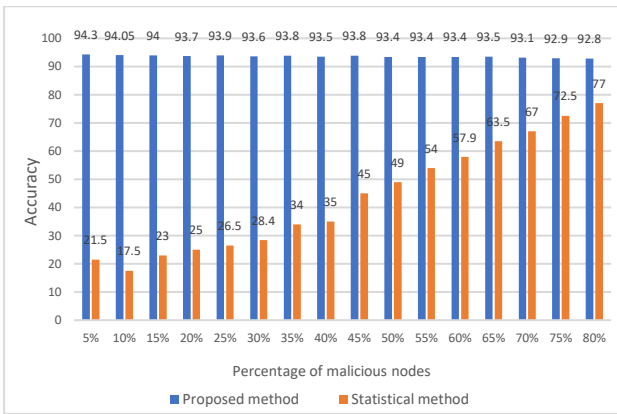
معیار دقت که در ارزیابی مدل پیشنهادی مورد توجه قرار گرفته است، نسبت تمام نمونه‌های صحیح پیش‌بینی شده، چه مثبت و چه منفی، به تعداد کل نمونه‌ها، دقت نامیده می‌شود. ارتباط دقت با معیارهای چهارگانه فوق در رابطه زیر آمده است:

$$accuracy = \frac{TP + TN}{P + N} = \frac{TP + TN}{TP + TN + FP + FN} \quad (2)$$

نرخ مثبت واقعی و منفی واقعی نیز به‌صورت زیر قابل محاسبه خواهند بود:

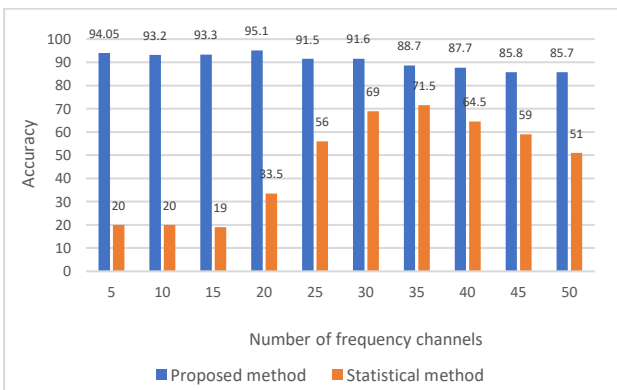
$$TPR = \frac{TP}{P} = \frac{TP}{TP + FN} \quad (3)$$

¹⁴ Accuracy



شکل ۵. دقت روش پیشنهادی بر حسب درصد گره‌های بدرفتار در مدل تهدید اول

همچنین تعداد کل گره‌ها و تعداد گره‌های بدرفتار ثابت فرض شده است و در عوض تعداد کانال‌های فرکانسی در دسترس از ۵ تا ۵۰ افزایش یافته است. همان‌طور که مشاهده می‌شود، با افزایش تعداد کانال‌های فرکانسی، دقت روش پیشنهادی کاهش می‌یابد. به این علت که با افزایش تعداد کانال‌های فرکانسی، احتمال انتخاب کانال فرکانسی مشترک بین گره‌ها کاهش می‌یابد و در نتیجه احتمال برخورد میان بسته‌ها کمتر می‌شود. بنابراین گره‌های بدرفتار اختلال کمتری در ارسال گره‌های خوش‌رفتار ایجاد می‌کنند. به همین دلیل تشخیص بدرفتاری یا خوش‌رفتاری کاربر دشوارتر می‌شود و دقت کاهش می‌یابد. نتایج مقایسه روش پیشنهادی با روش پایه آماری نیز نشان می‌دهد که دقت روش پیشنهادی از دقت روش آماری بیشتر است. اگرچه دقت روش آماری ابتدا رو به رشد است، ولی با افزایش منابع رادیویی، مجدداً روند کاهشی به خود می‌گیرد.

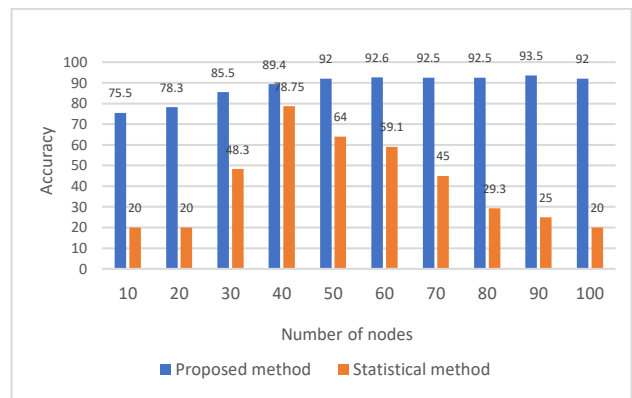


شکل ۶. دقت روش پیشنهادی بر حسب تعداد کانال‌های فرکانسی در مدل تهدید اول

در نهایت در نمودار شکل ۷ نیز منحنی مشخصه عملکرد (ROC)^{۱۵} برای روش پیشنهادی (به ازای آستانه‌های مختلف اعمال شده روی مدل یادگیرنده برای تعیین خوش‌رفتاری/بدرفتاری)، ارائه شده است. مقادیر آستانه‌ها روی نمودار مشخص شده است. این نمودار از بالا بودن نرخ مثبت واقعی روش پیشنهادی، به نسبت نرخ مثبت کاذب آن حکایت دارد.

آموزشی مدل زیاد شده و دقت آن بهتر می‌شود. اما زیاد شدن بیش از حد گره‌ها مجدداً باعث کاهش دقت خواهد شد. چرا که افزایش برخورد های ناشی از افزایش بار شبکه، دقت آن در تشخیص برخورد ناشی از بدرفتاری تغییر احتمال را کاهش می‌دهد. میزان دقت مدل به طور متوسط ۸۸.۳ درصد بوده و قابل قبول می‌باشد.

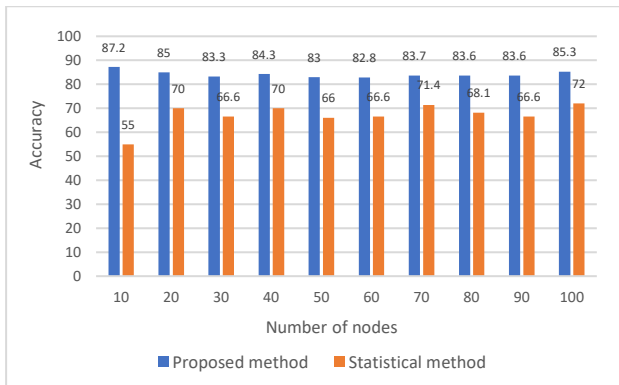
نتایج مقایسه روش پیشنهادی با روش پایه آماری نیز نشان می‌دهد که دقت روش پیشنهادی با اختلاف قابل توجه، بیشتر از دقت روش تشخیص بدرفتاری مبتنی بر مشخصه‌های آماری می‌باشد. همچنین مشاهده می‌شود که با افزایش تعداد گره‌ها، دقت تشخیص بدرفتاری براساس مشخصه‌های آماری کاهش می‌یابد، در حالی که این دقت در روش پیشنهادی تغییر چندانی نکرده است.



شکل ۴. دقت راهکار پیشنهادی بر حسب تعداد گره‌ها در مدل تهدید اول

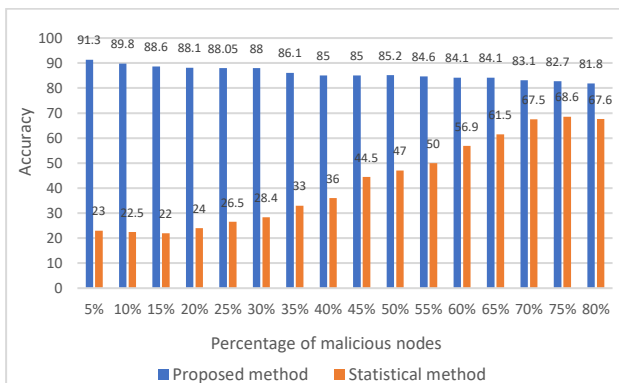
در نمودار شکل ۵ نیز تمامی شرایط محیط اعم از تعداد شکاف‌ها، تعداد کانال‌های فرکانسی، زمان منقضی شدن بسته‌ها و تعداد مجاز ارسال مکرر هر بسته ثابت فرض شده است. همچنین تعداد گره‌ها ثابت فرض شده و از بین گره‌های موجود، درصد گره‌های بدرفتار از ۵ درصد تا ۸۰ درصد افزایش یافته است. همان‌طور که مشاهده می‌شود، با افزایش تعداد گره‌های بدرفتار، دقت مدل کاهش می‌یابد. به این علت که با افزایش تعداد گره‌های بدرفتار، برخورد میان گره‌ها افزایش می‌یابد و تشخیص برخورد ناشی از بدرفتاری نسبت به برخورد ناشی از ازدحام، دشوار می‌شود. فلذا دقت مدل در تشخیص شکست‌های ناشی از بدرفتاری تغییر احتمال ارسال، کاهش می‌یابد. اما این افت دقت قابل توجه نیست، زیرا مدل یادگیرنده براساس شواهد رفتاری گره‌ها اقدام به شناسایی بدرفتاری می‌کند و افزایش درصد گره‌های بدرفتار پس از آموزش مدل، تأثیر چندانی در دقت آن ندارد. نتایج مقایسه روش پیشنهادی با روش پایه آماری نیز نشان می‌دهد که دقت روش پیشنهادی بیشتر از دقت روش تشخیص بدرفتاری پایه می‌باشد. البته افزایش تعداد گره‌های بدرفتار، این اختلاف را کاهش داده است. چون افزایش تعداد گره‌های بدرفتار در نمونه‌های آموزشی، تعیین آستانه‌های روش پایه برای تشخیص بدرفتاری را دقیق‌تر کرده است. نمودار شکل ۶ نیز با فرض ثابت بودن تمامی شرایط محیط اعم از تعداد شکاف‌ها، تعداد کانال‌های فرکانسی، زمان منقضی شدن بسته‌ها و تعداد مجاز ارسال مکرر هر بسته رسم شده است.

¹⁵ Receiver Operating Characteristic



شکل ۸. دقت روش پیشنهادی بر حسب تعداد گره‌ها در مدل تهدید دوم

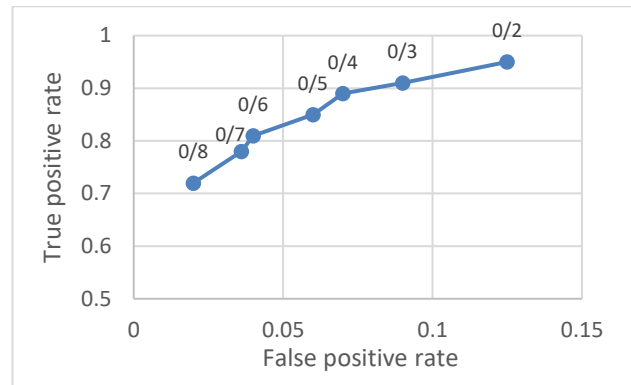
بنابراین، دقت مدل در تشخیص شکست‌های ناشی از بدرفتاری (تغییر احتمال ارسال)، کاهش می‌یابد. اما این افت دقت قابل توجه نیست، زیرا مدل یادگیرنده براساس شواهد رفتاری گره‌ها اقدام به شناسایی بدرفتاری می‌کند و افزایش درصد گره‌های بدرفتار پس از آموزش مدل، تاثیر چندانی در دقت آن ندارد. نتایج مقایسه روش پیشنهادی با روش پایه آماری نیز نشان می‌دهد که در این سناریو نیز دقت روش پیشنهادی بیشتر از دقت روش پایه می‌باشد.



شکل ۹. دقت روش پیشنهادی بر حسب درصد گره‌های بدرفتار در مدل تهدید دوم

نمودار شکل ۱۰ نیز با فرض ثابت بودن تعداد شکاف‌ها، تعداد کانال‌های فرکانسی، زمان منقضی شدن بسته‌ها و تعداد مجاز ارسال مکرر هر بسته رسم شده است. همچنین تعداد گره‌ها و تعداد گره‌های بدرفتار ثابت فرض شده است و در عوض تعداد کانال‌های فرکانسی از ۵ تا ۵۰ افزایش یافته است. همان‌طور که مشخص است، در این مورد نیز مانند مدل تهدید اول (شکل ۶)، با افزایش تعداد کانال‌های فرکانسی دقت مدل کاهش می‌یابد. به این دلیل که با افزایش تعداد کانال‌های فرکانسی، احتمال انتخاب کانال فرکانسی مشترک بین گره‌ها کاهش می‌یابد و در نتیجه احتمال برخورد میان بسته‌ها کمتر می‌شود. بنابراین گره‌های بدرفتار اختلال کمتری در ارسال گره‌های خوش‌رفتار ایجاد می‌کنند و به همین دلیل تشخیص بدرفتاری یا خوش‌رفتاری کاربر دشوارتر می‌شود و دقت کاهش می‌یابد.

نتایج مقایسه روش پیشنهادی با روش پایه آماری نیز نشان می‌دهد که دقت روش پیشنهادی از دقت روش آماری بیشتر است. اگرچه دقت روش مبتنی بر مشخصه‌های آماری ابتدا افزایش می‌یابد، ولی دوباره روند کاهشی دارد و بالا تر رفتن تعداد کانال‌ها و کاهش ازدحام برای استفاده از کانال‌ها، این روند افزایش متوقف می‌شود.



شکل ۷. منحنی مشخصه عملکرد در مدل تهدید اول

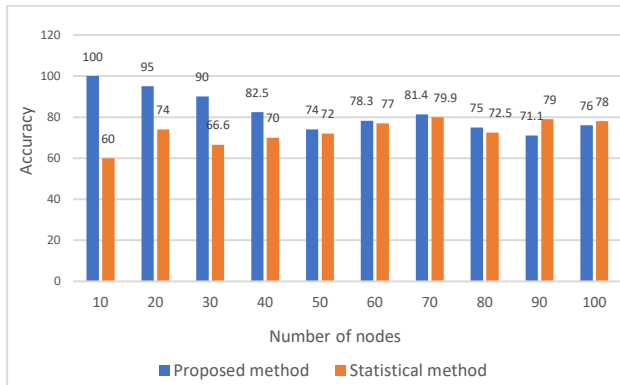
۵-۲- نتایج ارزیابی با فرض مدل تهدید دوم (ارسال بسته با احتمال $P=0.9$)

در بخش دوم از ارزیابی‌ها، مدل تهدید دوم یعنی ارسال با احتمال 0.9 توسط گره‌های بدخواه را در نظر می‌گیریم. در این بخش نیز، ابتدا نحوه عملکرد روش پیشنهادی بر حسب تعداد گره‌های بدرفتار بررسی می‌شود. در نمودار شکل ۸ تمامی شرایط محیط اعم از تعداد شکاف‌ها، تعداد کانال‌های فرکانسی، زمان منقضی شدن بسته‌ها و تعداد مجاز ارسال مکرر هر بسته ثابت بوده و صرفاً تعداد گره‌ها از ۱۰ تا ۱۰۰ افزایش یافته است. در همه این سناریوها، تعداد گره‌های بدرفتار در نظر گرفته شده، ۲۰ درصد تعداد کل گره‌ها است.

همان‌طور که در شکل ۸ مشاهده می‌شود، مقادیر دقت نسبت به مدل تهدید قبلی (شکل ۴) کمتر است. به این علت که در این مدل، گره بدرفتار همیشه بدرفتار نیست و در واقع به طور متوسط می‌توان گفت که در ده درصد مواقع خوش رفتار بوده است. به همین علت تشخیص این نوع بدرفتاری برای مدل یادگیرنده سخت‌تر بوده و در نتیجه دقت کمتر است. همچنین افزایش تعداد گره‌ها اثر منظمی بر دقت نداشته است. میزان متوسط دقت ۸۴.۱ درصد است.

نتایج مقایسه روش پیشنهادی با روش پایه آماری نیز نشان می‌دهد که دقت روش پیشنهادی بیشتر از دقت روش آماری می‌باشد. همچنین با بررسی دقیق‌تر مقادیر مثبت واقعی (TP) و منفی واقعی (TN) مشاهده شد که درصد مثبت واقعی روش پایه در مقایسه با روش پیشنهادی کم می‌باشد، یعنی روش آماری توانایی کمی برای تشخیص گره‌های بدرفتار دارد و تعداد زیادی از گره‌ها را خوش رفتار تشخیص می‌دهد. به همین دلیل است که دقت تا حدودی بالا می‌باشد؛ زیرا اکثر گره‌ها خوش رفتار هستند و درست تشخیص داده شده‌اند، اما گره‌های بدرفتار به درستی تشخیص داده نشده‌اند و از آنجایی که در همه حالات، گره‌های بدرفتار تنها ۲۰ درصد از کل گره‌ها را شامل می‌شوند، دقت افت زیادی ندارد، اگرچه از روش پیشنهادی به مقدار قابل توجهی پایین‌تر است. در شکل ۹ نیز تمامی شرایط محیط اعم از تعداد شکاف‌ها، تعداد کانال‌های فرکانسی، زمان منقضی شدن بسته‌ها، تعداد مجاز ارسال مکرر هر بسته و تعداد گره‌ها ثابت فرض شده است، اما از بین گره‌های موجود، درصد گره‌های بدرفتار از ۵ درصد تا ۸۰ درصد افزایش یافته است. همان‌طور که مشخص است، در این مورد نیز مانند مدل تهدید قبلی (شکل ۵)، با افزایش تعداد گره‌های بدرفتار، دقت مدل کاهش می‌یابد. زیرا با افزایش تعداد گره‌های بدرفتار، برخورد میان گره‌ها افزایش می‌یابد و تشخیص برخورد ناشی از بدرفتاری نسبت به برخورد ناشی از ازدحام کمی دشوار می‌شود.

(TP) و منفی واقعی (TN) مشاهده می‌شود که این روش قادر به تشخیص گره-های بدرفتار (دستکاری در تنظیمات تعداد ارسال مکرر) نیست و درصد مثبت واقعی در حدود صفر است. یعنی این روش تعداد زیادی از گره‌ها را به عنوان گره‌ی خوش‌رفتار طبقه‌بندی می‌کند و از آنجایی که تنها ۲۰ درصد از تعداد کل گره‌ها، بدرفتار هستند، این افت دقت مشهود نیست.



شکل ۱۲. دقت روش پیشنهادی بر حسب تعداد گره‌ها در مدل تهدید سوم

نمودار شکل ۱۳ نیز با فرض ثابت بودن تعداد شکاف‌ها، تعداد کانال‌های فرکانسی و زمان منقضی شدن بسته‌ها رسم شده است. اما تعداد گره‌ها ثابت فرض شده است و از بین گره‌های موجود، درصد گره‌های بدرفتار از ۵ درصد تا ۷۰ درصد افزایش داده شده است. همان‌طور که مشاهده می‌شود با افزایش تعداد گره‌های بدرفتار، دقت مدل کاهش می‌یابد. می‌دانیم که با افزایش تعداد گره‌های بدرفتار، بار شبکه افزایش می‌یابد، چرا که گره‌های بدرفتار بسته‌های خود را به دفعات بیشتری ارسال می‌کنند. در نتیجه، برخورد میان گره‌ها افزایش می‌یابد و تشخیص برخورد ناشی از بدرفتاری نسبت به برخورد ناشی از ازدحام کمی دشوار می‌شود. بنابراین، دقت مدل در تشخیص شکست‌های ناشی از دستکاری بدخواهانه در تنظیمات تعداد ارسال مکرر، کاهش می‌یابد. اما این افت دقت قابل توجه نیست، زیرا مدل یادگیرنده براساس سابقه شواهد رفتاری تک تک گره‌ها اقدام به شناسایی بدرفتاری می‌کند و افزایش درصد گره‌های بدرفتار پس از آموزش مدل، تاثیر چندانی در دقت آن ندارد.

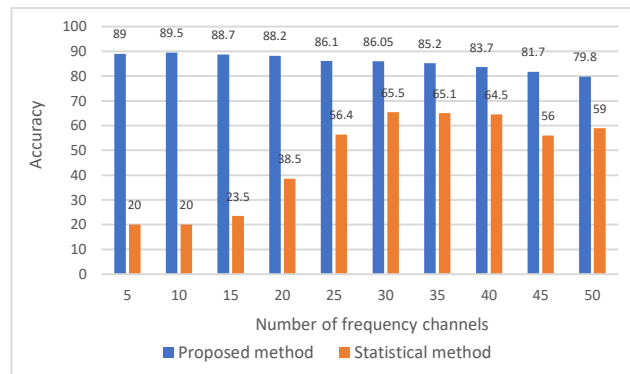
نتایج مقایسه روش پیشنهادی با روش پایه آماری نیز نشان می‌دهد که در ابتدا دقت روش تشخیص بدرفتاری مبتنی بر مشخصه‌های آماری بیشتر از دقت روش پیشنهادی می‌باشد. ولی با افزایش درصد گره‌های بدرفتار، دقت کاهش می‌یابد. البته، در این حالت نیز مانند شکل ۱۰، وقتی مقادیر مثبت واقعی (TP) و منفی واقعی (TN) بررسی شدند، مشاهده شد که روش پایه به خوبی قادر به تشخیص گره‌های بدرفتار (دستکاری کننده در تنظیمات تعداد ارسال مکرر) نمی‌باشد و درصد مثبت واقعی بسیار کم است و روش پایه تعداد زیادی از گره‌ها را به عنوان گره‌ی خوش‌رفتار طبقه‌بندی می‌کند و به همین دلیل، با بالا رفتن درصد گره‌های بدرفتار، دقت کاهش می‌یابد.

در پایان در نمودار شکل ۱۴ نیز منحنی مشخصه عملکرد (RoC) برای روش پیشنهادی (به ازای آستانه‌های مختلف تعیین کننده کلاس خروجی پیش‌بینی شده توسط مدل یادگیرنده)، ارایه شده است. این نمودار از بالا بودن نرخ مثبت واقعی روش پیشنهادی، به نسبت نرخ مثبت کاذب آن حکایت دارد.

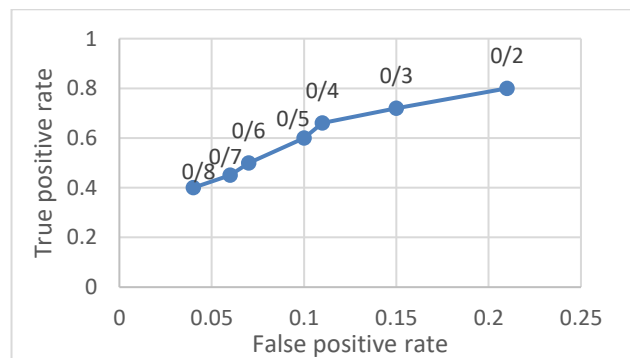
۴- نتیجه‌گیری

در این تحقیق از رویکرد یادگیری رفتار به کمک شبکه‌های عصبی عمیق

در نهایت در نمودار شکل ۱۱ نیز منحنی مشخصه عملکرد (RoC) برای روش پیشنهادی (به ازای آستانه‌های مختلف تعیین کننده کلاس خروجی پیش‌بینی شده توسط مدل یادگیرنده)، ارایه شده است. این نمودار از بالا بودن نرخ مثبت واقعی روش پیشنهادی، به نسبت نرخ مثبت کاذب آن حکایت دارد.



شکل ۱۰. دقت روش پیشنهادی بر حسب تعداد کانال‌های فرکانسی در مدل تهدید دوم



شکل ۱۱. منحنی مشخصه عملکرد در مدل تهدید دوم

۳-۵- نتایج ارزیابی با فرض مدل تهدید سوم (جعل در تعداد ارسال مکرر)

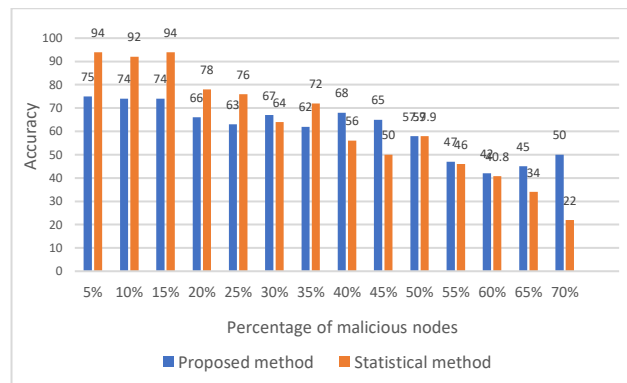
در بخش سوم از ارزیابی‌ها، مدل تهدید سوم، یعنی دستکاری در تنظیمات تعداد ارسال مکرر (K) به مقداری بیشتر از (دو برابر) مقدار توافق شده را در نظر می‌گیریم. همانند بخش‌های دیگر ابتدا نحوه عملکرد روش پیشنهادی بر حسب تعداد گره‌های بدرفتار بررسی می‌شود. در نمودار شکل ۱۰ تمامی شرایط محیط اعم از تعداد شکاف‌ها، تعداد کانال‌های فرکانسی و زمان منقضی شدن بسته‌ها ثابت بوده و صرفاً تعداد گره‌ها از ۱۰ تا ۱۰۰ افزایش یافته است. در همه‌ی این سناریوها، تعداد گره‌های بدرفتار در نظر گرفته شده، ۲۰ درصد تعداد کل گره‌ها است.

همان‌طور که از شکل ۱۲ مشخص است، وقتی شرایط محیط ثابت باشد، با افزایش تعداد گره‌ها دقت مدل روند کاهشی دارد. زیرا با زیاد شدن تعداد گره‌ها و در حضور گره‌های بدخواهی که تعداد ارسال مکرر را افزایش می‌دهند، برخوردهای ناشی از افزایش بار شبکه بیشتر می‌شود و به دنبال آن دقت مدل در تشخیص بدرفتاری ناشی از افزایش تعداد ارسال مکرر، کاهش می‌یابد.

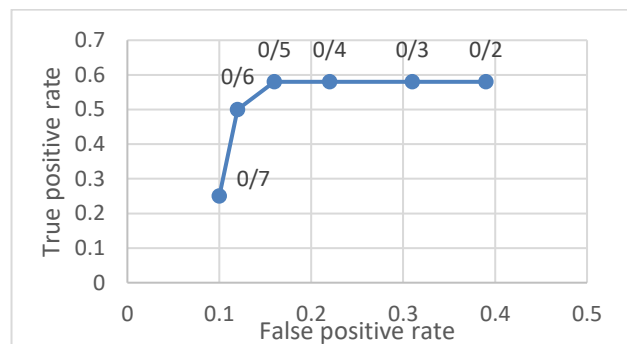
نتایج مقایسه روش پیشنهادی با روش پایه آماری نیز حاکی از آن است که دقت روش پیشنهادی در اغلب موارد بیشتر از دقت روش تشخیص بدرفتاری مبتنی بر مشخصه‌های آماری می‌باشد. ولی همان‌طور که از نمودار مشخص است، دقت روش تشخیص بدرفتاری مبتنی بر مشخصه‌های آماری نیز در این سناریو بهتر از سناریوهای قبل می‌باشد. اما پس از بررسی مقادیر مثبت واقعی

مراجع

- [1] J. A. d. Peral-Rosado, R. Raulefs, J. A. López-Salcedo and G. Seco-Granados, "Survey of Cellular Mobile Radio Localization Methods: From 1G to 5G," *IEEE Communications Surveys & Tutorials*, vol. 20, no. 2, pp. 1124-1148, 2018.
- [2] N. H. Mahmood, R. Abreu, R. Böhnke, M. Schubert, G. Berardinelli and T. H. Jacobsen, "Uplink Grant-Free Access Solutions for URLLC services in 5G New Radio," *16th International Symposium on Wireless Communication Systems (ISWCS)*, pp. 607-612, 2019.
- [3] R. Abbas, T. Huang, B. Shahab, M. Shirvanimoghaddam, Y. Li and B. Vucetic, "Grant-Free Non-Orthogonal Multiple Access: A Key Enabler for 6G-IoT," 2020.
- [4] N. Ye, J. Yu, A. Wang and R. Zhang, "Help from space: grant-free massive access for satellite-based IoT in the 6G era", *Digit. Commun. Netw.*, pp. 215-224, 2022.
- [5] A. Azari, P. Popovski, G. Miao and C. Stefanovic, "Grant-Free Radio Access for Short-Packet Communications over 5G Networks", *GLOBECOM 2017 - 2017 IEEE Global Communications Conference*, Singapore, pp. 1-7, 2017.
- [6] H. M. Gürsu, W. Kellerer and C. Stefanović, "On Throughput Maximization of Grant-Free Access with Reliability-Latency Constraints," *ICC 2019 - 2019 IEEE International Conference on Communications (ICC)*, Shanghai, China., pp. 1-7, 2019.
- [7] H. YU, Z. FEI, C. CAO, M. XIAO, D. JIA and N. YE, "Analysis of irregular repetition spatially-coupled slotted ALOHA," *SCIENCE CHINA Information Sciences*, vol. 62, 2019.
- [8] F. Formaggio, A. Munari and F. Clazzer, "On receiver diversity for grant-free based machine type communications," *Ad Hoc Networks Journal*, vol. 107, 2020.
- [9] S. Moon and J.-W. Lee, "Performance Study of Repetition-Based Grant-Free Schemes in the mMTC Scenario," *2019 34th International Technical Conference on Circuits/Systems, Computers and Communications (ITC-CSCC)*, JeJu, Korea (South), pp. 1-2, 2019.
- [10] B. Singh, O. Tirkkonen, Z. Li and M. A. Uusitalo, "Contention-Based Access for Ultra-Reliable Low Latency Uplink Transmissions," *IEEE Wireless Communications Letters*, vol. 7, pp. 182-185, 2018.
- [11] P. Minero and M. Franceschetti, "Throughput of Slotted ALOHA with Encoding Rate Optimization and Multipacket Reception," *IEEE INFOCOM 2009*, pp. 2971-2975, 2009.
- [12] A. B. Mackenzie and S. B. Wicker, "Selfish users in Aloha: A game-theoretic approach," *IEEE 54th Vehicular Technology Conference*, vol. 3, pp. 1354 - 1357, 2001.
- [13] R. T. B. Ma, V. Misra and D. Rubenstein, "An Analysis of Generalized Slotted-Aloha Protocols," *IEEE/ACM Transactions on Networking*, vol. 17, pp. 936-949, 2009.
- [14] C.-H. Chin, J. G. Kim and D. Lee, "Stability of Slotted Aloha with Selfish Users under Delay Constraint," *KSI Transactions on Internet and Information Systems*, vol. 5, pp. 542-559, 2011.
- [15] P. Antoniadis, S. Fdida, C. Griffin, Y. Jin and G. Kesidis, "Distributed medium access control with conditionally altruistic users," *EURASIP Journal on Wireless Communications and Networking*, vol. 202, 2013.
- [16] Y. Lin, "Jamming-aware Randomized Spectrum Sharing in Cognitive Radio Communication Networks," *Applied Mechanics and Materials*, Vols. 513-517, pp. 834-840, 2014.
- [17] F.-T. Hsu and H.-J. Su, "Throughput improvement in ALOHA networks with power capture and a cheat-proof access control," in *21st Annual IEEE International Symposium on Personal, Indoor and Mobile Radio Communications*, Istanbul, Turkey, 2010.
- [18] D. Wang, C. COMANICIU and U. TURELI, "Cooperation and Fairness for Slotted Aloha," *Wireless Personal Communications*, p. 13-27, 2007.
- [19] Y. Jin, G. Kesidis and J. W. Jang, "A Channel Aware MAC Protocol in an ALOHA Network with Selfish Users," *IEEE Journal on Selected Areas in Communications*, vol. 30, pp. 128-137, 2012.
- [20] W. F. Fihri, H. E. Ghazi, B. A. E. Majd and Bouanani, "A Machine Learning Approach for Backoff Manipulation Attack Detection in Cognitive Radio," *IEEE Access*, vol. 8, pp. 227349-227359, 2020.



شکل ۱۳. دقت روش پیشنهادی بر حسب درصد گره‌های بدرفتار در مدل تهدید سوم



شکل ۱۴. منحنی مشخصه عملکرد در مدل تهدید سوم

برای تشخیص بدرفتاری دستگاه‌ها در دسترسی بدون تخصیص به رسانه بی-سیم در شبکه سلولی نسل ۵ بهره گرفته شده است. مدل شبکه عصبی استفاده شده در این پژوهش، یک حافظه کوتاه مدت طولانی (LSTM) با دو لایه مخفی است که رفتار کاربران در دسترسی به رسانه به صورت دنباله‌ای پنجره‌بندی شده، به این شبکه عصبی تزریق و برچسب بدرفتاری/خوشرفتاری کاربر به عنوان خروجی از آن استخراج می‌شود. نتایج شبیه‌سازی دقت مطلوب این روش در تشخیص بدرفتاری را نشان می‌دهد. البته افزایش تعداد گره‌ها، افزایش تعداد کانال‌های فرکانسی در دسترس، افزایش تعداد گره‌های بدرفتار، و حملات خاموش-روشن، دقت مدل را کاهش می‌دهند. همچنین، مقایسه روش پیشنهادی با یک روش آماری پایه، برتری قابل توجه روش پیشنهادی را در تشخیص بدرفتاری در عموم اوقات نشان می‌دهد. البته این بهبود به قیمت تحمیل بار محاسباتی برای تعلیم مدل یادگیرنده حاصل شده است که البته تعلیم مدل می‌تواند به صورت برون-خط انجام شود تا سربار روش پیشنهادی را در هنگام استفاده بالا نبرد. الگوریتم به پارامترهای مختلفی مثل تعداد گره‌های موجود در محیط، تعداد گره‌های بدرفتار و تعداد منابع فرکانسی وابسته است. اما تغییر در تعداد منابع فرکانسی بیشترین تاثیر را در دقت مدل داشته است. لذا پیشنهاد می‌شود در کارهای آینده، معیارهای زمینه‌ای دیگر نظیر منابع فرکانسی موجود شبکه نیز در طراحی مدل یادگیرنده مورد توجه قرار گیرند. همچنین، در ادامه این تحقیق می‌توان عملکرد سایر روش‌های یادگیری عمیق را برای تشخیص بدرفتاری بررسی کرد. همچنین می‌توان به مدل‌سازی ریاضی رفتار دسترسی بدون تخصیص مبتنی بر K بار ارسال مکرر و ارائه راهکارهای مبتنی بر مدل تحلیلی برای تشخیص آن پرداخت و بر اساس آن مدل‌های یادگیرنده توانمندتری برای تشخیص بدرفتاری‌های دسترسی به رسانه بدون تخصیص ایجاد نمود.

- [21] R. T. B. Ma, V. Misra and D. Rubenstein, "Modeling and Analysis of Generalized Slotted-Aloha MAC Protocols in Cooperative, Competitive and Adversarial Environments," *IEEE International Conference on Distributed Computing Systems*, pp. 62-62, 2006.
- [22] S. Gyawali and Y. Qian, "Misbehavior Detection using Machine Learning in Vehicular Communication Networks," *IEEE International Conference on Communications (ICC)*, pp. 1-6, 2019.
- [23] Y. Imamverdiyev and L. Sukhostat, "Anomaly detection in network traffic using extreme learning machine," *IEEE 10th International Conference on Application of Information and Communication Technologies (AICT)*, pp. 1-4, 2016.
- [24] X. Ling, Y. Le, J. Wang and Z. Ding, "Hash Access: Trustworthy Grant-Free IoT Access Enabled by Blockchain Radio Access Networks," *IEEE Network*, vol. 34, pp. 54-61, 2020.
- [25] N. Wang, W. Li, A. Alipour-Fanid, L. Jiao, M. Dabaghchian and K. Zeng, "Pilot Contamination Attack Detection for 5G MmWave Grant-Free IoT Networks," *IEEE Transactions on Information Forensics and Security*, vol. 16, pp. 658-670, 2021.
- [26] G. Berardinelli, N. H. Mahmood, R. Abreu, T. Jacobsen, K. Pedersen, I. Z. Kovács and P. Mogensen, "Reliability Analysis of Uplink Grant-Free Transmission Over Shared Resources," *IEEE Access*, vol. 6, pp. 23602-23611, 2018.
- [27] J. Zhang, L. Lu, Y. Sun, Y. Chen, J. Liang, J. Liu, H. Yang, S. Xing, Y. Wu, J. Ma, I. B. F. Murias and F. J. L. Hernando, "PoC of SCMA-Based Uplink Grant-Free Transmission in UCNC for 5G," *IEEE Journal on Selected Areas in Communications*, vol. 35, no. 6, pp. 1353-1362, 2017.
- [28] J. D. P. -N. C. a. Y. S. H. L. Deng, "On the Asymptotic Performance of Delay-Constrained Slotted ALOHA," in *27th International Conference on Computer Communication and Networks (ICCCN)*, Hangzhou, China, 2018.
- [29] M. VASOUJOUYBARI, E. Ataie, M. Bastam, "An MLP-based Deep Learning Approach for Detecting DDoS Attacks", *Journal of Electrical Engineering of Tabriz University*, vol. 52, no. 3, pp. 195-204, 1401.
- [30] M. Zendedel, J. Hamidzadeh, "Improving intrusion detection in the Internet of Things network using deep learning and chaotic shrimp optimization algorithm", *Journal of Electrical Engineering, University of Tabriz*, vol. 53, no. 2, pp. 127-138, 2023.