

شناسایی گروه‌های سخت‌قانون در پیام‌رسان تلگرام با استفاده از شبکه عصبی گرافی گراف‌سیج و داده‌های

شبیه‌سازی شده دارای ویژگی‌های زمانی

مائه عرب بافرانی^۱، محمدعلی زارع چاهوکی^{۲*}

^۱ دانشجوی دکتری، دانشکده مهندسی کامپیوتر، گروه مهندسی نرم‌افزار، دانشگاه یزد، یزد، ایران

^۲ دانشیار، دانشکده مهندسی کامپیوتر، گروه مهندسی نرم‌افزار، دانشگاه یزد، یزد، ایران

مقاله پژوهشی

چکیده

تاریخ دریافت:

۱۴۰۵/۰۱/۱۲

تاریخ پذیرش:

۱۴۰۵/۰۵/۲۷

کلیدواژه‌ها:

گروه‌های سخت‌قانون، پیام‌رسان
تلگرام، شبکه‌های عصبی گرافی،
گراف‌سیج، داده‌های
شبیه‌سازی شده، ویژگی‌های زمانی،
تشخیص رفتار اجباری، یادگیری
عمیق گراف‌محور

نویسنده مسئول:

chahooki@yazd.ac.ir

این پژوهش به شناسایی «گروه‌های سخت‌قانون» در پیام‌رسان تلگرام می‌پردازد، گروه‌هایی که کاربران را برای عضویت یا ادامه فعالیت به انجام اقدامات اجباری مانند دعوت از دیگران یا عضویت در کانال‌های خاص ملزم کرده و روند رشد طبیعی گروه‌ها را دست‌کاری می‌کنند. تشخیص این گروه‌ها برای حفظ یکپارچگی پلتفرم و جلوگیری از سوءاستفاده اهمیت بالایی دارد. کمبود داده‌های برچسب‌دار واقعی، ماهیت پویا و زمانی رفتارها و ترکیب هم‌زمان سیگنال‌های ساختاری و رفتاری از چالش‌های این حوزه هستند. برای غلبه بر این چالش‌ها، مجموعه داده‌ای شبیه‌سازی شده در مقیاس بزرگ شامل ۱۰۰,۰۰۰ کاربر، ۱۰,۰۰۰ گروه و ۱۲ تصویر زمانی تولید شد که ویژگی‌های ایستا و زمانی کاربران و گروه‌ها را در برمی‌گیرد. برچسب‌ها با استفاده از یک امتیاز ریسک مبتنی بر تحلیل مؤلفه‌های اصلی تعیین و یک گراف دوبخشی کاربر-گروه ساخته شد. سپس یک مدل شبکه عصبی گرافی مبتنی بر گراف‌سیج با پروتکل ارزیابی کاملاً استقرایی، تابع زیان آنروپی متقاطع دودویی وزن‌دار و بهینه‌سازی آستانه تصمیم آموزش داده و در کنار شش مدل پایه ارزیابی گردید. مدل پیشنهادی بر روی مجموعه آزمون مستقل به صحت ۹۸/۲ درصد، دقت کلاس مثبت ۹۷/۲ درصد، فراخوانی ۹۳/۷ درصد، امتیاز F1 برابر ۹۵/۴ درصد، AUC برابر ۹۹/۸۷ درصد و میانگین دقت برابر ۹۹/۵۱ درصد دست‌یافت و در تمام بذرها پایدار بود. نتایج نشان داد به دلیل وابستگی ساختاری برچسب‌ها به ویژگی‌های ایستای گروه، مدل‌های جدولی با هزینه محاسباتی کمتر عملکرد اندکی بهتر دارند، درحالی‌که گراف‌سیج در گراف‌های بزرگ و پر درجه مقاوم می‌ماند. این یافته‌ها صرفاً بر روی داده‌های شبیه‌سازی شده تأیید شده‌اند و تعمیم آن‌ها به داده‌های واقعی تلگرام، به دلیل نبود داده‌های برچسب‌دار واقعی، همچنان چالشی باز است.

doi : 10.22034/ABMIR.2026.24432.1231

E-ISSN: [2821-2037](https://doi.org/10.22034/ABMIR.2026.24432.1231)

/The Author 2026. Published by Yazd University This is an open

access article under the CC BY 4.0 License <https://creativecommons.org/licenses/by/4.0/>.



۱- مقدمه

۱-۱ شکاف علمی و چالش‌ها

روش‌های سنتی مبتنی بر ویژگی‌های آماری یا تحلیل ایستا، توانایی محدودی در استخراج الگوهای پیچیده رفتاری دارند. علت اصلی این محدودیت این است که رفتارهای اجباری اغلب در تعامل هم‌زمان میان ساختار گراف و روابط همسایگی کاربران، ویژگی‌ها و الگوهای گروهی و تغییرات زمانی رفتارها رخ می‌دهند [۲]. علاوه بر این، در بسیاری از کاربردهای واقعی، دسترسی به داده‌های برجسب‌دار و قابل اعتماد بسیار محدود است. نبود داده‌های واقعی با برجسب معتبر برای رفتارهای غیرعادی، یکی از موانع اصلی توسعه مدل‌های یادگیری عمیق در این حوزه محسوب می‌شود. این مسئله ضرورت استفاده از داده‌های مصنوعی کنترل‌شده را برجسته می‌سازد، داده‌هایی که بتوانند ویژگی‌های آماری، ساختاری و رفتاری شبکه‌های اجتماعی موجود را شبیه‌سازی کرده و امکان ارزیابی اصولی روش‌های یادگیری را فراهم سازند [۳]، [۵]، [۶]، [۷].

سپس به‌کارگیری آن در داده‌های واقعی است. استفاده از داده مصنوعی، نه تنها محدودیت دسترسی به داده‌های واقعی را برطرف می‌کند، بلکه امکان کنترل دقیق پارامترهای تولید، تحلیل حساسیت و بررسی اصولی رفتار مدل را نیز فراهم می‌آورد. درنهایت، این پژوهش تلاش می‌کند با تحلیل نتایج، پیشنهادهایی برای توسعه مجموعه داده‌های واقعی ارائه دهد و مشخص کند چه نوع ویژگی‌ها و اطلاعاتی باید به داده‌های اجتماعی افزوده شوند تا تشخیص رفتارهای ساختاری و ناهنجار در مقیاس واقعی امکان‌پذیر گردد.

۲- پیشینه پژوهش

مرور مطالعات پیشین نشان می‌دهد که باوجود پیشرفت‌های قابل‌توجه در هر یک از حوزه‌های تولید گراف‌های مصنوعی و شبیه‌سازی شبکه‌های اجتماعی، تشخیص رفتارهای غیرعادی و هماهنگ، کاربرد شبکه‌های عصبی گرافی در تحلیل سامانه‌های اجتماعی و شناسایی ناهنجاری‌ها و الگوهای ساختاری در گراف‌ها، ترکیب شبیه‌سازی اجتماعی با یادگیری مبتنی بر گراف برای تحلیل

شبکه‌های اجتماعی به یکی از پیچیده‌ترین بسترهای تعامل انسانی تبدیل شده‌اند، ساختاری پویا که در آن گروه‌ها، کاربران و الگوهای رفتاری به‌صورت پیوسته شکل می‌گیرند و تغییر می‌کنند. در میان ساختارهای موجود در شبکه‌های اجتماعی، بخشی از گروه‌ها به‌صورت طبیعی رشد می‌کنند و تعداد اعضای آن‌ها به‌مرور زمان افزایش می‌یابد، درحالی‌که برخی دیگر از الگوهای رفتاری غیرعادی، هدایت‌شده و یا اصطلاحاً اجباری پیروی می‌کنند [۱]، [۲]، [۳]، [۴]. شناسایی این نوع ساختارها از اهمیت بالایی در تحلیل دقیق تعاملات اجتماعی، قابل‌اطمینان بودن پیام‌رسان، شناسایی فعالیت‌های هماهنگ شده و غیرواقعی و پیشگیری از سوءاستفاده‌های هدفمند برخوردار است [۱]، [۲]. بااین‌حال، پیچیدگی ساختاری شبکه‌های اجتماعی، وابستگی زمانی رفتارها و وجود نویز و عدم قطعیت، تشخیص چنین الگوهایی را به یک مسئله چالش‌برانگیز در حوزه یادگیری ماشینی و تحلیل گراف تبدیل کرده است [۲]، [۴].

۲-۱ رویکرد پیشنهادی و چارچوب تحقیق

در این پژوهش، چارچوبی برای شبیه‌سازی و یادگیری مبتنی بر گراف برای بررسی رفتارهای اجباری در گروه‌های اجتماعی ارائه می‌شود. ابتدا یک مدل تولید داده مصنوعی طراحی شده که ویژگی‌های آماری و رفتاری کاربران و گروه‌ها را با استفاده از توزیع‌های واقع‌گرایانه و شاخص‌های تعامل اجتماعی شبیه‌سازی می‌کند [۳]، [۵]، [۶]، [۷]. سپس، با ساخت یک گراف دوبخشی شامل کاربران و گروه‌ها، ساختار ارتباطی شبکه اجتماعی بازسازی می‌شود. برای تحلیل این ساختار، یک مدل یادگیری عمیق مبتنی بر شبکه‌های عصبی گرافی از نوع گراف‌سیج با مکانیزم توازن کلاس و آموزش با نظارت توسعه داده‌شده است [۸]. این روش امکان استخراج الگوهای ساختاری، رفتاری و تعاملات پنهان در گراف را فراهم می‌سازد [۳]، [۵]، [۶]، [۷].

۳-۱ هدف تحقیق و کاربردپذیری

هدف اصلی این پژوهش، ارزیابی دقت و کارایی روشی گراف‌محور برای شناسایی رفتارهای اجباری در پیام‌رسان‌ها و

قاعده‌محور مانند R-MAT در حفظ پویایی‌های اجتماعی عملکرد ضعیف‌تری نشان می‌دهند. باین حال، هزینه محاسباتی بالاتر این مدل‌ها و وابستگی آن‌ها به داده‌های آموزشی واقعی، چالش‌هایی برای کاربرد در مقیاس بزرگ ایجاد می‌کند.

در سوی دیگر، پژوهش‌های مبتنی بر مدل‌های عامل‌محور بر تولید جمعیت مصنوعی همراه با شبکه‌های اجتماعی تمرکز دارند [۶]. [۷]. این دسته از مطالعات، جمعیت‌های مصنوعی بزرگ با ساختارهای اجتماعی مشخص تولید می‌کنند تا شبیه‌سازی تعاملات انسانی، مانند انتشار بیماری یا تصمیم‌گیری اجتماعی، امکان‌پذیر شود. نتایج این پژوهش‌ها نشان می‌دهد که در صورت اعتبار آماری مناسب، داده‌های مصنوعی می‌توانند جایگزین قابل‌قبولی برای داده‌های واقعی باشند. مزیت اصلی این رویکرد، توانایی مدل‌سازی رفتار اجتماعی در مقیاس بزرگ است، اما ساده‌سازی بیش‌ازحد رفتار انسانی و وابستگی به قوانین از پیش تعریف‌شده، از محدودیت‌های آن به شمار می‌رود. به‌طور مشابه، چارچوب‌های عامل‌محور جدید نیز نشان داده‌اند که بازسازی پویایی تعاملات اجتماعی می‌تواند برای تحلیل سامانه‌های اجتماعی پیچیده سودمند باشد [۳].

علاوه بر این، مطالعات اخیر تأکید کرده‌اند که کیفیت داده‌های مصنوعی صرفاً به حفظ ویژگی‌های توپولوژیکی محدود نمی‌شود، بلکه میزان شباهت توزیع ویژگی‌های رفتاری و پویایی‌های زمانی نیز در اعتبار این داده‌ها نقش اساسی دارد. درواقع، ارزیابی داده‌های مصنوعی باید علاوه بر معیارهای ساختاری، بر اساس قابلیت استفاده در وظایف پایین‌دستی و میزان انطباق آماری با داده‌های واقعی صورت گیرد [۵]، [۲۸].

۲-۲ تشخیص رفتارهای غیرعادی و اجباری در شبکه‌های اجتماعی

تشخیص رفتارهای هماهنگ و فعالیت‌های از پیش تعیین‌شده، یکی از مسائل مهم در حوزه امنیت و تحلیل شبکه‌های اجتماعی به شمار می‌رود. در پژوهش‌های اولیه نشان داده شد که رفتارهای مصنوعی، از جمله فعالیت بات‌ها، معمولاً در الگوهای زمانی، تراکم تعاملات و هماهنگی رفتاری پنهان می‌شوند. این مطالعات تأکید می‌کنند که اتکا بر ویژگی‌های آماری ساده مانند نرخ فعالیت برای

رفتارهای غیرعادی در سطح گروه‌های اجتماعی، همچنان به‌عنوان یک شکاف پژوهشی مهم مطرح است. بخش عمده پژوهش‌های موجود یا بر داده‌های واقعی و در سطح کاربران تمرکز داشته‌اند، یا صرفاً یکی از ابعاد ساختاری، رفتاری یا زمانی را مدنظر قرار داده‌اند. ازاین‌رو، توسعه چارچوب‌هایی که بتوانند به‌صورت هم‌زمان از داده‌های مصنوعی واقع‌گرایانه، ویژگی‌های رفتاری و زمانی، و یادگیری مبتنی بر گراف برای تحلیل رفتارهای گروهی بهره بگیرند، همچنان نیازمند توجه بیشتر است.

۱-۲ تولید گراف مصنوعی و شبیه‌سازی شبکه‌های اجتماعی

تولید گراف‌های مصنوعی به‌عنوان راهکاری برای غلبه بر محدودیت‌های دسترسی به داده‌های واقعی شبکه‌های اجتماعی، در سال‌های اخیر نقش مهمی در پژوهش‌ها ایفا کرده است. این رویکرد علاوه بر فراهم کردن امکان کنترل دقیق پارامترهای مدل، ریسک‌های اخلاقی مرتبط با حریم خصوصی را نیز کاهش می‌دهد. یکی از کارهای بنیادی در این حوزه، مدل NetGAN [۹] است که با استفاده از مسیرهای تصادفی و مدل‌های مولد تخصصی، توزیع ساختاری گراف را یاد می‌گیرد و گراف‌هایی با ویژگی‌های توپولوژیکی مشابه نمونه‌های واقعی تولید می‌کند. این روش در حفظ ساختارهای مرتبه‌بالا مانند خوشه‌بندی و الگوهای اتصال موفق عمل می‌کند، اما فاقد مدل‌سازی ویژگی‌های گره و تغییرات زمانی است؛ درنتیجه برای تحلیل رفتارهای اجتماعی پیچیده محدودیت دارد.

مدل GraphRNN [۱۰] با استفاده از شبکه‌های بازگشتی، چارچوبی ترتیبی برای تولید گراف معرفی کرد که بازسازی گراف‌های پیچیده‌تر را ممکن می‌سازد. این مدل در بازتولید ساختارهای پویا عملکرد بهتری نشان می‌دهد، اما همچنان به مفروضات ساده‌سازی شده متکی است و قادر به پوشش کامل پیچیدگی رفتار انسانی نیست.

مطالعات جدیدتر نشان داده‌اند که مدل‌های مبتنی بر توجه گرافی می‌توانند ساختارهای اجتماعی واقعی را با دقت بالاتری بازسازی کنند و شباهت آماری بیشتری با شبکه‌های واقعی نظیر فیسبوک و گیت‌هاب داشته باشند [۵]، [۱۶]. درحالی‌که مدل‌های کلاسیک

و شباهت الگوهای تولید محتوا می‌توانند در تشخیص رفتارهای هماهنگ مؤثر باشند.

همچنین، مطالعات میان‌پلتفرمی نیز وجود هماهنگی‌های مشترک میان شبکه‌هایی نظیر X، فیسبوک و تلگرام را تأیید کرده‌اند [۱۵]. برای نمونه، پژوهشگران در مطالعات مبتنی بر مدل‌های عامل‌محور و اتوماتای رفتاری نشان داده‌اند که داده‌های مصنوعی، در صورت طراحی مناسب، می‌توانند ساختار شبکه و الگوهای رفتاری کاربران را با درجه‌ای قابل قبول از واقع‌گرایی بازسازی کنند [۳، ۷]. مزیت اصلی این رویکرد، کنترل کامل پارامترها و امکان بررسی سناریوهای غیرقابل مشاهده در داده‌های واقعی است؛ با این حال، خطر ساده‌سازی بیش‌ازحد رفتار انسانی همچنان وجود دارد.

۲-۳ شبکه‌های عصبی گرافی در تحلیل سامانه‌های اجتماعی

شبکه‌های عصبی گرافی به دلیل توانایی در استخراج هم‌زمان الگوهای ساختاری و رفتاری، پیشرفت‌های چشمگیری در تحلیل شبکه‌های اجتماعی ایجاد کرده‌اند. مدل گراف‌سیج در [۸] با رویکرد یادگیری القایی، امکان پردازش گراف‌های بزرگ و پویا را فراهم می‌کند. این مدل با تجمع اطلاعات همسایگان، ویژگی‌های محلی گره‌ها را به‌طور مؤثر یاد می‌گیرد، اما فاقد سازوکار توجه برای اولویت‌بندی میزان اهمیت همسایگان است.

در مقابل، مدل شبکه توجه بر گراف^۱ در [۱۶] با معرفی مکانیزم توجه، نشان داد که همه همسایگان سهم یکسانی در یادگیری نمایش گره ندارند. این رویکرد دقت مدل را در بسیاری از کاربردها افزایش داد، اما در گراف‌های نویزی یا با تراکم بالا ممکن است با ناپایداری مواجه شود.

پژوهش‌های جدیدتر نیز نشان داده‌اند که استفاده از ساختارهای ناهمگن و بازسازی گراف می‌تواند در شناسایی الگوهای انتشار غیرعادی و شایعات مؤثر باشد [۱۷].

در حوزه کاربردهای اجتماعی، مدل GNN-CARE برای تشخیص تقلب در شبکه‌های اجتماعی توسعه یافت [۴]. این مدل با ترکیب اطلاعات ساختاری گراف و ویژگی‌های رفتاری کاربران، توانایی شناسایی الگوهای پنهان را بهبود می‌بخشد. نتایج گزارش شده نشان

شناسایی چنین رفتارهایی کافی نیست و باید ساختار تعاملات نیز در تحلیل لحاظ شود. با این حال، محدودیت اصلی این رویکردها در ناتوانی آن‌ها در مدل‌سازی روابط غیرخطی و پیچیده رفتاری است.

روش CopyCatch که در [۲] ارائه شد، یکی از نخستین سامانه‌های مبتنی بر گراف برای کشف فعالیت‌های هماهنگ محسوب می‌شود. این روش با شناسایی زیرگراف‌های متراکم از رفتارهای مشابه، الگوهای سازمان‌یافته را آشکار می‌کند. نقطه قوت CopyCatch در تحلیل ساختاری شبکه است، اما فاقد مدل‌سازی پویا بوده و به داده‌های واقعی وابسته است. مطالعات بعدی نشان داده‌اند که رفتارهای هماهنگ معمولاً در تعامل هم‌زمان میان ساختار شبکه، بعد زمانی و ویژگی‌های رفتاری نمایان می‌شوند؛ موضوعی که به کارگیری مدل‌های پیشرفته‌تر را ضروری می‌سازد.

در پژوهش [۱۱] هماهنگی رفتاری بر اساس انحراف الگوهای اشتراک‌گذاری از رفتار متعارف کاربران مورد بررسی قرار گرفت. نتایج این مطالعه نشان داد که رویکردهای مبتنی بر ناهنجاری می‌توانند وابستگی به آستانه‌های ثابت را کاهش داده و کمپین‌های هماهنگ را با انعطاف‌پذیری بیشتری شناسایی کنند. این یافته‌ها بیانگر آن است که تشخیص رفتارهای هماهنگ، بیش از آنکه به یک معیار منفرد متکی باشد، نیازمند ترکیب چندین سیگنال رفتاری و ساختاری است.

مطالعات انجام‌شده در سایر شبکه‌های اجتماعی نیز بر همین موضوع توجه دارند. در پژوهش [۱۲] نشان داده شد که رفتارهای هماهنگ در فیسبوک صرفاً محدود به عملیات نفوذ سازمان‌یافته نیستند و می‌توانند در قالب شبکه‌های رسانه‌ای، تبلیغاتی و سیاسی نیز ظاهر شوند.

همچنین، مطالعه [۱۳] نشان داد که حساب‌های هماهنگ در آبشارهای انتشار اطلاعات، رفتار متفاوتی نسبت به کاربران عادی دارند و نقش پررنگ‌تری در گسترش محتوا ایفا می‌کنند.

افزون بر این، پژوهش [۱۴] بیان کرد که در محیط‌های ویدئو‌محور، نشانه‌هایی نظیر هم‌زمانی انتشار، بازاستفاده از محتوای چندرسانه‌ای

¹ Graph Attention Network

۳- روش پیشنهادی

این بخش روش‌شناسی پیاده‌سازی کامل سیستم پیشنهادی را تشریح می‌کند، از ساخت یک معیار^۱ مصنوعی کنترل‌شده تا طراحی، آموزش و ارزیابی استقرایی شبکه عصبی گرافی. ما ابتدا مسئله یادگیری را صورت‌بندی می‌کنیم (بخش ۳-۱) و نمای کلی پاپ‌لین دومرحله‌ای را ارائه می‌دهیم (بخش ۳-۲). سپس الگوریتم‌های تولید داده‌های مصنوعی و توجیه‌های آماری آن‌ها (بخش ۳-۳)، نمایش گراف و کدگذاری ویژگی‌ها (بخش ۳-۴)، پروتکل استقرایی بدون نشت اطلاعات (بخش ۳-۵)، معماری‌های مدل شامل آشکارساز پیشنهادی گراف‌سیج و تمام رقبا آن (بخش ۳-۶) و در نهایت الگوریتم‌های آموزش، بهینه‌سازی و ارزیابی (بخش‌های ۳-۷ تا ۳-۸) را به تفصیل بیان می‌کنیم. این طراحی از رویه‌های تثبیت‌شده در یادگیری نمایش گراف [۸]، [۲۲]، [۲۳] و تشخیص ناهنجاری و قلب مبتنی بر گراف [۴]، [۲۴]، [۲۵] پیروی می‌کند.

۳-۱ صورت‌بندی مسئله و نمادگذاری

در این پژوهش یک پلتفرم آنلاین به‌عنوان یک گراف ناهمگن و در حال تکامل زمانی مدل‌سازی شده است. فرض بر آن بوده است که $U = \{u_1, \dots, u_n\}$ مجموعه n کاربر و $G = \{g_1, \dots, g_m\}$ مجموعه m گروه (جامعه) باشد. عضویت در گام زمانی گسسته t یک رابطه دوبخشی $E_B^t \subseteq U \times G$ است، که در آن $(u, g) \in E_B^t$ اگر و تنها اگر کاربر u در تصویر زمانی t عضو گروه g باشد. از رابطه عضویت، یک رابطه فرافکنی^۲ گروه-گروه $E_P \subseteq G \times G$ القا می‌شود که یال‌های آن گروه‌هایی را که کسر نسبتاً بزرگی از اعضای مشترک دارند به هم متصل می‌کند. هر گروه g دارای یک بردار ویژگی ایستای $x_g^{stat} \in \mathbb{R}_s^d$ و برای هر تصویر زمانی $t \in \{0, \dots, T-1\}$ دارای یک بردار ویژگی زمانی $x_g^{temp, (t)} \in \mathbb{R}_t^d$ است؛ هر کاربر نیز دارای یک بردار ویژگی $x_u \in \mathbb{R}_u^d$ است. با تعداد کاربر ۱۰۰ هزار، تعداد گروه ۱۰ هزار و تعداد تصویر زمانی برابر با ۱۲، مجموعه گره‌های ترکیبی دارای $|V| = n + m$ برابر با ۱۱۰ هزار عنصر است.

می‌دهد که این مدل در داده‌های واقعی عملکرد مطلوبی دارد، اما وابستگی آن به داده‌های برچسب‌خورده، محدودیت‌هایی برای کاربرد در سناریوهای دارای کمبود داده معتبر ایجاد می‌کند. علاوه بر این، مطالعات اخیر بر توسعه معماری‌های فضایی-زمانی و مقاوم در برابر عدم توازن کلاس متمرکز شده‌اند. پژوهش [۱۸] نشان داد که مدل‌سازی هم‌زمان روابط ساختاری و تغییرات زمانی می‌تواند عملکرد تشخیص گروه‌های متقلب را بهبود بخشد. همچنین، مدل BalancerGNN باهدف ارتقای عملکرد شبکه‌های عصبی گرافی در مجموعه داده‌های نامتوازن ارائه شده است [۱۹]. از سوی دیگر، پژوهش‌هایی نظیر [۲۰] و [۲۱] اهمیت تلفیق بعد زمانی با ساختار گراف را در تشخیص الگوهای پیچیده رفتاری برجسته ساخته‌اند. با این وجود، بخش عمده این مطالعات بر تشخیص قلب در سطح کاربران، تراکنش‌ها یا انتشار محتوا متمرکز بوده و توجه کمتری به تحلیل ناهنجاری در سطح گروه‌های شبکه‌های اجتماعی، به‌ویژه اجتماعات تلگرامی، داشته‌اند.

۲-۴ شکاف پژوهشی

مرور مطالعات پیشین نشان می‌دهد که اگرچه پیشرفت‌های قابل توجهی در زمینه تولید داده‌های مصنوعی، تشخیص رفتارهای هماهنگ و بهره‌گیری از شبکه‌های عصبی گرافی حاصل شده است، اما همچنان شکاف‌هایی در این حوزه وجود دارد. بخش عمده پژوهش‌های موجود یا مبتنی بر داده‌های واقعی و در سطح کاربران انجام شده‌اند، یا تنها یکی از ابعاد ساختاری، رفتاری و زمانی را مورد توجه قرار داده‌اند. همچنین، پژوهش‌های مرتبط با رفتارهای هماهنگ عمدتاً بر پلتفرم‌هایی نظیر فیسبوک، X و تیک‌تاک متمرکز بوده و مطالعات محدودی به تحلیل ناهنجاری در سطح گروه‌های تلگرامی پرداخته‌اند. از این رو، توسعه چارچوبی مبتنی بر داده‌های مصنوعی که بتواند به‌صورت هم‌زمان ویژگی‌های ایستا، رفتاری و زمانی گروه‌ها را مدل‌سازی کرده و از قابلیت‌های یادگیری مبتنی بر گراف برای طبقه‌بندی الگوهای غیرعادی بهره گیرد، می‌تواند بخشی از این شکاف پژوهشی را پوشش دهد.

² Projection

¹ Benchmark

سوگیری استقرایی رابطه‌ای شبکه‌های عصبی مبتنی بر تبادل پیام^۱ همخوانی دارد [۸]، [۲۶]، [۲۳].

۲-۳ نمای کلی سیستم

سیستم پیشنهادی یک پایپ‌لاین دومرحله‌ای است. در مرحله اول، یک مولد^۲ یک محیط مصنوعی خودسازگار از کاربران، گروه‌ها، عضویت‌ها در طی ۱۲ تصویر زمانی، یک فراکنی گراف، آمارهای توصیفی و تقسیم‌بندی‌های لایه‌بندی‌شده^۳ برای آموزش/اعتبارسنجی/آزمون تولید می‌کند. در مرحله دوم، یک موتور آموزش و ارزیابی این مصنوعات را دریافت کرده، یک شیء گراف Geometric-PyTorch [۲۷] می‌سازد، مدل پیشنهادی گراف‌سیج را به همراه شش رقیب آموزش می‌دهد و یک مجموعه ارزیابی کامل (معیارها، منحنی‌های یادگیری، مطالعه حذف، پایداری چندبذری^۴، آزمون معناداری آماری، تحلیل خطا، تشخیص‌های اهمیت ویژگی و اعتبارسنجی توزیع و زمان‌سنجی اجرا) ارائه می‌دهد.

الگوریتم (۱): تولید سرتاسری داده‌های مصنوعی^۵

```
Input: n users, m groups, T snapshots, C communities, target
prevalence  $\pi=0.20$ , seed=42
Output: users, groups, memberships{0..T}, temporal features,
projection edges, splits
1: seed all RNGs(42)
2: community_id  $\leftarrow$  SBM partition of users into C blocks
( $p_{intra}=0.15$ ,  $p_{inter}=0.005$ )
3: for each user u: type_u  $\sim$  Cat(ACTIVE.45, PASSIVE.52, BOT.03)
invite_rate,reply_rate,activity_power  $\sim$  type-specific
Exp/Beta/Uniform laws
4: for each group g: sample size, activity_ratio, reply_count, age, ...
($3.3.3)
derive silent_ratio, engagement_level, interaction_anomaly, flags
5: R1..R5  $\leftarrow$  five raw risk components; standardise; fit 1-component
PCA
 $w \leftarrow |PC1 \text{ loadings}| / \Sigma |loadings|$ ; risk(g)  $\leftarrow$  clip( $\Sigma w_i R_{i,g}$ , 0,1)
6:  $\tau \leftarrow$  Percentile_{100(1-\pi)}(risk); label_g  $\leftarrow$  1[risk(g)  $\geq \tau$ ]
7: for each group g: sample members  $\propto$  homophily weights
(Algorithm 2)
8: for t = 0..T-1: AR(1) update of series + burst injection + dynamic
join/leave (Algorithm 3)

9: build bipartite graph; build group projection via inverted index
(Algorithm 4)
10: drop leakage columns R1..R5, risk_score
11: stratified 70/15/15 split preserving prevalence  $\pi$ 
12: return all artefacts
```

جدول (۱): فهرست نمادها

نماد	تعریف
$U = \{u_1, \dots, u_n\}$	مجموعه n کاربر پلتفرم
$G = \{g_1, \dots, g_m\}$	مجموعه m گروه (جامعه)
t	گام زمانی گسسته (شماره تصویر زمانی)، $t \in \{1, \dots, T\}$
$E_B^t \subseteq U \times G$	رابطه عضویت دوبخشی کاربر-گروه در تصویر زمانی t
$E_P \subseteq G \times G$	رابطه فراکنی گروه-گروه، برخاسته از عضویت مشترک
$x_g^{stat} \in \mathbb{R}^{d_s}$	بردار ویژگی ایستای گروه g
$x_g^{temp,(t)} \in \mathbb{R}^{d_t}$	بردار ویژگی زمانی گروه g در تصویر زمانی t
$x_u \in \mathbb{R}^{d_u}$	بردار ویژگی کاربر u
n, m, T	به ترتیب تعداد کاربران، تعداد گروه‌ها و تعداد تصاویر زمانی
$ V = n + m$	تعداد کل گره‌های گراف ترکیبی
$f_g: G \rightarrow \{0, 1\}$	تابع طبقه‌بند گروه، پارامترمند با θ
$y_g \in \{0, 1\}$	برچسب گروه
π	نرخ شیوع هدف گروه‌های سخت‌قانون
τ	آستانه صدکی امتیاز ریسک برای برچسب‌گذاری
$R_1, \dots, R_5$	پنج مؤلفه خام امتیاز ریسک پیش از استانداردسازی
w	بردار وزن مشتق‌شده از بارهای مؤلفه اصلی اول (PC1)

وظیفه موردنظر، طبقه‌بندی دودویی گره بر روی گره‌های گروه است: یادگیری $\theta: f: G \rightarrow \{0, 1\}$ که هر گروه را به $y_g \in \{0, 1\}$ نگاشت می‌کند. یک گروه Law-Hard یا همان سخت‌قانون گروهی است که پروفایل رفتاری آن سکوت بالا، مشارکت پایین، فعالیت هدایت‌شده توسط ربات‌ها و رشد ناهنجار، آن را به عنوان یک گروه پرخطر و نیازمند نظارت سخت‌گیرانه‌تر نشان می‌دهد. فرضیه مرکزی این است که انتشار اطلاعات در سراسر رابطه دوبخشی کاربر-گروه و رابطه فراکنی گروه-گروه، تخمین ریسک بهتری نسبت به در نظر گرفتن هر گروه به عنوان یک بردار ویژگی مستقل به دست می‌دهد، که با

⁵ End-to-End Synthetic Data Generation

¹ Message-Passing

² Generator

³ Stratified

⁴ Multi-Seed

هم‌راستا است، تا شکل‌گیری عضویت به‌جای یکنواخت بودن، به‌طور واقع‌گرایانه‌ای همسان‌پسندانه باشد.

۳-۳-۲ مدل‌سازی رفتاری کاربر

هر کاربر یکی از سه گونه رفتاری است که نرخ‌های آن‌ها از توزیع‌های مختص به هرگونه نمونه‌برداری می‌شود و ناهمگنی رفتاری واقعی را تولید می‌کند. نکته مهم این است که ربات‌ها به‌شدت دعوت می‌کنند اما تقریباً هرگز پاسخ نمی‌دهند، که این امضای خودکار و هرزگونه^۳ مستندشده در ادبیات ربات‌های اجتماعی است [۱]. به‌طور رسمی، با $t = \text{type}(u)$ روابط زیر تعریف می‌شوند:

$$\begin{aligned} \text{invite}_{\text{rate}} &\sim \text{Exp}(\lambda_t), \\ \text{reply}_{\text{rate}} &\sim \text{Beta}(\alpha_t, \beta_t), \\ \text{activity}_{\text{power}} &\sim \text{Beta/Uniform}_t, \end{aligned} \quad (2)$$

جدول (۳): گونه‌های رفتاری و توزیع‌های مولد آن‌ها

گونه کاربر	احتمال	$\text{invite}_{\text{rate}}$	$\text{reply}_{\text{rate}}$	$\text{activity}_{\text{power}}$
ACTIVE	۰/۴۵	$\text{Exp}(۰/۵)$	$\text{Beta}(۵/۳) \text{ — high}$	$\text{Beta}(۴/۲) \text{ — high}$
PASSIVE	۰/۵۲	$\text{Exp}(۰/۱)$	$\text{Beta}(۱/۵) \text{ — low}$	$\text{Beta}(۱/۵) \text{ — low}$
BOT	۰/۰۳	$\text{Exp}(۵/۰) \text{ — very high}$	$\text{Beta}(۱/۲۰) \text{ — } \approx ۰$	$\text{Uniform}(۰/۸, ۱/۰)$

نرخ‌های نمایی، رفتار دعوت دارای دنباله سنگین^۴ را به تصویر می‌کشند، درحالی‌که قوانین بتا در بازه $[۰, ۱]$ انتخاب طبیعی برای نسبت‌های محدود مانند تمایل به پاسخ و فعالیت هستند.

۳-۳-۳ مهندسی ویژگی‌های گروه

هر گروه یک بردار ویژگی ایستا دریافت می‌کند که از متغیرهای اولیه قابل تفسیر ساخته شده است. اندازه گروه از یک قانون لاگ‌نرمال برش‌یافته^۵ پیروی می‌کند:

$$\text{participants} \sim \text{clip}(\text{Lognormal}(\mu = 6, \sigma = 1.2), 30, 2 \times 10^5) \quad (3)$$

که منعکس‌کننده توزیع‌های دم‌پهن اندازه گروه مشاهده‌شده در پلتفرم‌های آنلاین است؛ علاوه بر این:

$$\begin{aligned} \text{activity}_{\text{ratio}} &\sim \text{Beta}(2, 4), \\ \text{reply}_{\text{count}} &\sim \text{Poisson}(2), \\ \text{group}_{\text{age}} &\sim \text{Exp}(900) \end{aligned} \quad (4)$$

استفاده از یک معیار کاملاً مصنوعی یک انتخاب آگاهانه است: هیچ مجموعه داده واقعی برچسب‌داری از گروه‌های تعدیل‌شده در دسترس نبود و تولید مصنوعی کنترل دقیقی بر فرآیند تولید، دسترسی به واقعیت مبنا و تکرارپذیری کامل را فراهم می‌کند [۲۸]. پایپ‌لاین به‌صورت سراسری با بذر برابر با ۴۲ مقداردهی اولیه شده است. الگوریتم ۱ رویه تولید سرتاسری را که در بخش ۳-۳ بسط داده‌شده است، خلاصه می‌کند.

جدول (۲): مقیاس مجموعه داده‌های مصنوعی

کلاس مثبت	پنجره‌های زمانی	یال‌های گراف (میلیون)	گروه‌ها	کاربران
۲۰٪	۱۲	۷/۶۴	۱۰۰۰۰	۱۰۰۰۰۰

۳-۳-۳ پایپ‌لاین تولید داده‌های مصنوعی

مولد، چهار ویژگی از اکوسیستم‌های آنلاین واقعی را که محرک یادگیری رابطه‌ای هستند بازتولید می‌کند: ساختار جامعه پنهان، رفتار ناهمگن کاربران شامل حساب‌های خودکار، شکل‌گیری عضویت همسان‌پسندانه^۱ و پویایی‌های زمانی پایدار اما دارای فوران^۲. هر مؤلفه با منطق آماری، فرمول‌ها و الگوریتم مربوطه در ادامه شرح داده‌شده است.

۳-۳-۱ ساختار جامعه پنهان (مدل بلوک تصادفی)

کاربران به $C = 8$ جامعه پنهان با اندازه تقریباً مساوی افزاز می‌شوند. ستون فقرات جامعه از یک مدل بلوک تصادفی [۲۹] پیروی می‌کند، که در آن احتمال وجود پیوند بین کاربران i و j تنها به عضویت آن‌ها در بلوک‌های $b(i)$ و $b(j)$ بستگی دارد:

$$\begin{aligned} (i, j) \in E &= B_{b(i), b(j)}, \\ B_{c,c} &= p_{\text{intra}} = 0.15, \\ B_{c,c'} &= p_{\text{inter}} = 0.005 \quad (c \neq c') \end{aligned} \quad (1)$$

از آنجاکه $p_{\text{inter}} \ll p_{\text{intra}}$ است، پیوندها در درون جوامع متمرکز می‌شوند و توپولوژی همسو و ماژولار مشخصه سیستم‌های اجتماعی را بازتولید می‌کنند [۲۹]. همچنین به هر کاربر یک ترجیح موضوعی اختصاص داده می‌شود که تا حدودی با جامعه او

³ Spam

⁴ Heavy-Tailed

⁵ Clipped Log-Normal

¹ Homophilous

² Bursty

فرض کنید $R = [R_1, \dots, R_5]$ به میانگین صفر و واریانس یک استاندارد شود، $\tilde{R} = (R - \mu R) / \sigma R$ یک PCA تک‌مؤلفه‌ای روی $\tilde{R}$ برازش می‌شود، و بارهای مطلق بردار ویژه پیشرو ϕ از ماتریس کوواریانس آن به وزن‌های زیر نرمال می‌شوند:

$$w_i = \frac{|\phi_i|}{\sum_{j=1}^5 |\phi_j|}, \quad \sum_i w_i = 1 \quad (12)$$

از آنجاکه مؤلفه اول (PC1) بیشترین میزان از واریانس مشترک را در میان همه مؤلفه‌ها به خود اختصاص می‌دهد، روش مذکور بیشترین ضریب وزنی را برای آن محوری در نظر می‌گیرد که نماینده اصلی تغییرات مشترک است [۳۰]. امتیاز ریسک و برجسب عبارت‌اند از:

$$\text{risk}(g) = \text{clip}\left(\sum_{i=1}^5 w_i \cdot R_i(g), 0, 1\right) \quad (13)$$

$$y_g = 1[\text{risk}(g) \geq \tau], \quad (14)$$

$$\tau = \text{Percentile}_{80}(\{\text{risk}(g) : g \in G\})$$

که شیوع دقیق حدود ۲۰٪ گروه‌های مثبت را بدون تنظیم دستی اعمال می‌کند و به نگرانی‌های مربوط به تعادل کلاس که معیارهای ناهنجاری را تحریف می‌کنند، می‌پردازد [۳۱]. الگوریتم ۲ نحوه برجسب‌گذاری را ترسیم می‌کند.

الگوریتم (۲): امتیاز ریسک مبتنی بر PCA و عضویت همسان‌پسندانه

Part A — Labelling
1: build $R = [R_1..R_5]$ for all groups # Table 3
2: $\tilde{R} \leftarrow \text{StandardScaler.fit_transform}(R)$
3: $\phi \leftarrow \text{PCA}(n=1).\text{fit}(\tilde{R}).\text{components_}[0]$
4: $w \leftarrow |\phi| / \text{sum}(|\phi|)$
5: $\text{risk} \leftarrow \text{clip}(R @ w, 0, 1)$
6: $\tau \leftarrow \text{percentile}(\text{risk}, 80)$; label $\leftarrow (\text{risk} \geq \tau)$
Part B — Initial membership ($t=0$), per group g
7: $w(u,g) \leftarrow 0.5 + 2 \cdot 1[\text{com}(u)=\text{com}(g)] + 1.5 \cdot 1[\text{topic}(u)=\text{topic}(g)] + 1 \cdot 1[\text{bot}(u)]$
8: $w \leftarrow w / \text{sum}(w)$
9: $M_g \leftarrow \text{sample_without_replacement}(U, \text{size}=\text{participants}_g, p=w)$
10: add edges $\{(u,g) : u \in M_g\}$ to bipartite graph

از آنجاکه برجسب یک تابع قطعی از نرخ سکوت^۳، نرخ مشارکت^۴ و ناهنجاری رفتاری در تعامل^۵ است و این توصیف‌گرها در معرض مدل‌ها قرار می‌گیرند، این وظیفه ذاتاً برای هر مدلی که آن‌ها را

و یک متغیر رشد گاوسی کوچک:

$$\text{growth}_{\text{proxy}} \sim \mathcal{N}(0.012, 0.015^2) \quad (5)$$

توصیف‌گرهای مشتق‌شده اصلی عبارت‌اند از:

$$\text{active}_{\text{users}} = \text{activity}_{\text{ratio}} \times \text{participants} \quad (6)$$

$$\text{silent}_{\text{ratio}} = 1 - \text{activity}_{\text{ratio}} \quad (7)$$

$$\text{engagement}_{\text{level}} = \frac{\text{reply}_{\text{count}} + \text{active}_{\text{users}}}{\text{participants} + 1} \quad (8)$$

$$\text{interaction}_{\text{score}} = \log(1 + \text{reply}_{\text{count}}) \times \text{activity}_{\text{ratio}} \quad (9)$$

$$\text{interaction}_{\text{anomaly}} = (\text{participants} / 2 \times 10^5) \times [1 - \text{engagement}_{\text{level}}] \quad (10)$$

$$\text{active}_{\text{density}} = \frac{\text{active}_{\text{users}}}{\log(1 + \text{participants}) + 10^{-5}} \quad (11)$$

پرچم‌های خطر دودویی^۱ و کدهای وضعیت/نوع دسته‌بندی‌شده، بردار را کامل می‌کنند و شمارش‌های خام را به سیگنال‌های نرمال‌شده عدم فعالیت، عدم مشارکت و ناهنجاری اندازه تبدیل می‌کنند.

۳-۳-۴ ساخت برجسب از طریق امتیاز ریسک مبتنی بر تحلیل مؤلفه‌های اصلی^۲

برجسب‌ها به صورت دستی اختصاص داده نمی‌شوند؛ آن‌ها از یک امتیاز ریسک ترکیبی مشتق می‌شوند که وزن‌های آن به جای انتخاب دلخواه، با استفاده از تحلیل مؤلفه‌های اصلی [۳۰] از داده‌ها آموخته می‌شود. پنج مؤلفه ریسک خام استاندارد شده محاسبه می‌شود.

جدول (۴): پنج مؤلفه ریسک خام

مؤلفه	تعریف (intuition)
R_1 (silent)	silent_ratio — fraction of inactive members
R_2 (anomaly)	$\text{clip}(\text{interaction_anomaly}, *, *)$ — large size $\times$ low engagement
R_3 (growth)	$1 - \text{minmax}(\text{growth_proxy})$ — stagnation risk (inverted)
R_4 (engagement)	$1 - \text{engagement_level} / \text{max}(\text{engagement_level})$ — low-engagement risk
R_5 (bot)	$\cdot \sqrt{\text{bot_activity_flag}} + \cdot \sqrt{\text{low_activity_large_group_flag}}$

³ Silent_ratio

⁴ Engagement_level

⁵ Interaction_anomaly

¹ Bot_activity_flag, Low_activity_large_group_flag, High_silent_ratio_flag, Zero_reply_flag

² PCA

خودکار هستند [۱]. عضویت پویا: پیوستن‌ها با دعوت‌ها و مشارکت مقیاس می‌شوند و خروج‌ها با عدم فعالیت:

$$\text{join}_t \sim \text{Poisson}(\text{clip}(2 \cdot \text{invite}_t + 0.5 \cdot \text{engagement} \times 50, 1, 300)), \quad (17)$$

$$\text{leave}_t \sim \text{Poisson}(\text{clip}((1 - \text{activity}_{\text{ratio}}) \cdot 10 + 2, 1, 50)), \quad (18)$$

که در آن، کاربران پیوسته (مانده در شبکه) با وزن‌دهی مبتنی بر جامعه، موضوع و ربات نمونه‌برداری می‌شوند و کاربران خارج‌شونده (ترک‌کننده) نیز به‌صورت جامعه‌محور انتخاب می‌گردند، به‌گونه‌ای که حذف کاربران نامتناسب و ربات‌ها در اولویت قرار دارد. پس از هر گام، تعداد اعضا دقیقاً از مجموعه عضویت زنده خوانده می‌شود تا شمارش‌های زمانی و گراف عضویت باهم سازگار بمانند. الگوریتم ۳ جزئیات یک تصویر زمانی را نشان می‌دهد.

الگوریتم (۳): یک تصویر زمانی (AR(1) + فوران + عضویت پویا)

```
Input: prev series  $x_{t-1}$ , member sets  $\{M_g\}$ , burst schedule
1: for each series  $s \in \{\text{msg}, \text{inv}, \text{reply}, \text{silent}, \text{mem}\}$ :
 $x_t[s] \leftarrow \text{clip}(\mu_s + \phi(x_{t-1}[s] - \mu_s) + \varepsilon, \text{bounds}_s)$ 
# AR(1),  $\phi=0.85$ 
2: for  $g$  with burst at  $t$ :  $\text{msg}_t, \text{inv}_t \leftarrow \times \text{BURST\_MULT}(3..8)$ ;
 $\text{reply}_t \leftarrow \times 0.3$ 
3: for each group  $g$ :
 $n_{\text{join}} \leftarrow \text{Poisson}(\text{clip}(2 \cdot \text{inv}_t + 0.5 \cdot \text{engagement} \times 50, 1, 300))$ 
 $n_{\text{leave}} \leftarrow \text{Poisson}(\text{clip}((1 - \text{activity\_ratio}) \cdot 10 + 2, 1, 50))$ 
 $\text{leave\_score}(u) \leftarrow 2 \cdot I[\text{com} \neq] + 1.5 \cdot I[\text{topic} \neq] + 1 \cdot \text{bot}$ ;
remove  $n_{\text{leave}} \propto \text{leave\_score}$ 
sample  $n_{\text{join}}$  new users  $\propto$  community/topic/bot weights,
exclude current members
 $M_g \leftarrow (M_g - \text{leavers}) \cup \text{joiners}$ ;  $\text{member\_count}_g \leftarrow |M_g|$ 
4: persist 7 temporal features per group:
msg, inv, reply, silent, mem, net_change, is_burst
```

۷-۳-۳ ساخت گراف و فرافکنی گروه

یک گراف دوبخشی کاربر-گروه از عضویت‌های $t = 0$ ساخته می‌شود. فرافکنی گروه-گروه از ضریب ژاکارد [۳۴] استفاده می‌کند که برای گروه‌های g_i و g_j با مجموعه‌های اعضای M_i و M_j برابر است با:

$$J(g_i, g_j) = |M_i \cap M_j| / |M_i \cup M_j|, \text{ edge created iff } J \geq 0.05. \quad (19)$$

برای جلوگیری از هزینه $O(m^2)$ مقایسه همه باهم، اشتراک‌ها از طریق یک نمایه معکوس (کاربر $\leftarrow$ گروه‌ها) تجمع می‌شوند، بنابراین تنها جفت‌هایی که واقعاً دارای یک عضو مشترک هستند

مشاهده می‌کند آسان است. شش ستون ریسک داخلی (R_1-R_5) و امتیاز خام) قبل از ایجاد تقسیم‌بندی‌ها حذف می‌شوند، اما مواد سازنده آن‌ها باقی می‌مانند. این ویژگی، کلید تفسیر تمام نتایج پایین‌دستی (بخش ۴-۱۰) است و به‌صورت شفاف و سازگار با ارزیابی آگاه از نشت گزارش می‌شود.

۵-۳-۳ شکل‌گیری عضویت همسان‌پسندانه ($t = 0$)

عضویت‌های اولیه همسان‌پسندی را نمونه‌سازی می‌کنند، تمایل بازیگران مشابه به ارتباط با یکدیگر [۳۲]. برای هر گروه، کاربران نامزد بدون جایگذاری با احتمالی متناسب با وزن شباهت نمونه‌برداری می‌شوند (الگوریتم ۲، بخش B):

$$w(u|g) = 0.5 + 2 \times I[\text{com}(u) = \text{com}(g)] + 1.5 \times I[\text{topic}(u) = \text{topic}(g)] + I[\text{bot}(u)] \quad (15)$$

که روی تمام کاربران نرمال شده است. تعداد نمونه‌برداری‌شده برابر با تعداد شرکت‌کنندگان واقعی گروه (بدون سقف مصنوعی) است، که یک گراف عضویت دوبخشی تولید می‌کند که آمیختگی آن منعکس‌کننده SBM زیربنایی و ساختار همسان‌پسندانه است [۲۹]. [۳۲]

۶-۳-۳ پویایی‌های زمانی: AR(1)، فوران‌ها و عضویت

پویا

دوازده سری زمانی تولید می‌شود. هر سری برای هر یک از گروه‌ها و به‌ازای متغیرهای تعداد پیام، تعداد دعوت، نسبت پاسخ، نرخ سکوت و تعداد اعضا، به‌عنوان یک فرآیند خودهمبسته مرتبه اول تکامل می‌یابد [۳۳]:

$$x_t = \mu + \phi(x_{t-1} - \mu) + \varepsilon_t, \quad \phi = 0.85, \varepsilon_t \sim \mathcal{N}(0, \sigma^2), \quad (16)$$

که در آن $\phi = 0.85$ ماندگاری قوی اما بازگشت به میانگین را تحمیل می‌کند و مدل متعارف برای سیگنال‌های زمانی با حافظه کوتاه [۳۳] دو مکانیزم واقع‌گرایی را می‌افزاید. رویدادهای فوران: با احتمال ۰/۱۵، یک گروه جهش ناگهانی ۳ تا ۸ برابری را در پیام‌ها و دعوت‌ها در یک تصویر تصادفی دریافت می‌کند، درحالی‌که نسبت پاسخ آن سرکوب می‌شود (۰/۳ برابر)، که موج‌های اسپم/حمله را مدل‌سازی کرده و از طریق پرچم is_burst ثبت می‌شود؛ چنین موج‌های کم‌پاسخی منعکس‌کننده تقویت

۳-۳-۹ کنترل نشت و تخصیص لایه‌بندی شده

شش ستون ریسک داخلی حذف می‌شوند، پس از آن یک تخصیص لایه‌بندی شده به صورت ۷۰ درصد آموزش ۱۵ درصد اعتبارسنجی و ۱۵ آزمون (با بذر ۴۲) منجر به ۷۰۰۰ گروه آموزشی، ۱۵۰۰ گروه اعتبارسنجی و ۱۵۰۰ گروه آزمون می‌شود، که هر کدام نرخ تقریباً برابر با ۲۰٪ مثبت را حفظ می‌کنند. لایه‌بندی از جابجایی شیوع ناشی از تخصیص جلوگیری می‌کند، که یک منبع شناخته شده از خوش‌بینی در ارزیابی نامتوازن است [۳۱].

۳-۴ نمایش گراف و کدگذاری ویژگی‌ها

مولد یک شیء واحد همگن داده از [۲۷] Geometric-PyTorch می‌سازد. ویژگی‌های طبقه‌بندی کاربر به صورت hot-one کدگذاری شده و با ویژگی‌های عددی کاربر الحاق می‌شوند. برای گروه‌ها، مجموعه ایستا شامل تمام ستون‌های عددی به جز شناسه، برچسب، موضوع و ستون‌های نشت حذف شده و مجموعه زمانی شامل هفت سیگنال زمانی در یک تصویر زمانی انتخاب شده است. از آنجا که فضاهای ویژگی کاربر و گروه در ابعاد متفاوت هستند ($\mathbf{g} \neq \mathbf{u}$)، هر دو با صفر به یک عرض مشترک $\mathbf{g} \mathbf{d} + \mathbf{u} \mathbf{d} = \mathbf{d}$ حاشیه‌گذاری شده و به صورت عمودی روی هم چیده می‌شوند:

$$\mathbf{X} = [[\mathbf{X}_U | \mathbf{0}]; [\mathbf{0} | \mathbf{X}_G]] \in \mathbb{R}^{(n+m) \times d}, \quad (22)$$

بنابراین یک گره کاربر، ویژگی‌های خود را در بلوک اول و صفرها را در بلوک گروه حمل می‌کند و بالعکس. این کار امکان تبادل پیام همگن را فراهم می‌سازد [۸]، [۲۶]. مقیاس‌بندی ویژگی‌ها بدون نشت است: یک مقیاس‌ساز استاندارد تنها روی گروه‌های آموزشی (در تمام تصاویر زمانی) برازش شده و بر روی گروه‌های اعتبارسنجی/آزمون اعمال می‌شود تا از تأثیر آمارهای آزمون بر آموزش جلوگیری کند.

۳-۵ پروتکل ارزیابی استقرایی

یک پروتکل استقرایی دقیق تضمین می‌کند که گروه‌های آزمون هرگز در طول آموزش مشاهده نمی‌شوند، و از نشت تراگذاری^۵ که بسیاری از نتایج یادگیری گراف را تحت تأثیر قرار می‌دهد،

ارزیابی می‌گردند. الگوریتم ۴ یک واقعیت ساختاری برای تحلیل‌های بعدی مرکزی است: با ۱۰۰۰۰ گروه که اعضای خود را از ۱۰۰۰۰۰ کاربر جذب می‌کنند، جفت‌های بسیار کمی به $J \geq 0.05$ می‌رسند که تنها ۶۹ یال باقی می‌مانند و فرافکنی به ۹۹۶۳ مؤلفه تقسیم می‌شود. بنابراین مدل‌ها از رابطه دوبخشی متراکم یاد می‌گیرند (بخش ۴-۱۰).

الگوریتم (۴): فرافکنی گروه از طریق نمایه معکوس (ژاکارد دقیق)

```
Input: memberships (t=0), threshold  $\theta=0.05$ 
1:  $\text{inv} \leftarrow \{\}$ ; for (u,g) in memberships:  $\text{inv}[u].\text{append}(g)$ 
2:  $\text{inter} \leftarrow \text{defaultdict}(\text{int})$ 
3: for u, groups_u in inv: for each pair (g_i,g_j) in groups_u:
    $\text{inter}[(g_i,g_j)] += 1$ 
4: for (g_i,g_j), c in inter:
    $J \leftarrow c / (|M_{\{g_i\}}| + |M_{\{g_j\}}| - c)$ 
   if  $J \geq \theta$ : add projection edge (g_i,g_j, weight=J)
5: return projection edges
```

۳-۳-۸ اعتبارسنجی توزیع داخلی

از آنجا که هیچ بدنه مرجع واقعی در دسترس نبود، مولد دو کلاس برچسب را با دو معیار مکمل در تقابل قرار می‌دهد: آماره کولموگوروف-اسمیرنوف^۱ دو نمونه‌ای [۳۵]، که بیشترین فاصله بین توابع توزیع تجمعی تجربی است:

$$KS = \sup_x |F_0(x) - F_1(x)|, \quad (20)$$

و دیورژانس جنسن-شانون^۲ [۳۶]، که یک نسخه هموار شده، متقارن و کران‌دار از دیورژانس کولبک-لایبلا^۳ بر روی هیستوگرام‌های ۶۰ بازه‌ای است:

$$JS(P||Q) = \frac{1}{2} KL(P||M) + \frac{1}{2} KL(Q||M), M = \frac{1}{2}(P+Q). \quad (21)$$

این معیارها تأیید می‌کنند که ویژگی‌های تعریف‌کننده برچسب، کلاس‌ها را از هم جدا می‌کنند درحالی‌که ویژگی‌های تصادفی این کار را نمی‌کنند. به‌طور شفاف گزارش می‌شود که این یک مقایسه داخلی بین برچسب ۰ و برچسب ۱ است، نه یک اعتبارسنجی مصنوعی در برابر واقعی.

⁴ Pad

⁵ Transductive

¹ Kolmogorov-Smirnov Test

² Jensen-Shannon Divergence

³ Kullback-Leibler Divergence

$$h_{\mathcal{N}(v)}^{(k)} = \text{AGG}_k \left(\{h_u^{(k-1)} : u \in \mathcal{N}(v)\} \right), \quad (23)$$

$$\text{AGG} = \text{mean}(\text{default})$$

$$h_v^{(k)} = \sigma \left(W_k \cdot \left[h_v^{(k-1)} \parallel h_{\mathcal{N}(v)}^{(k)} \right] \right) \quad (24)$$

که در آن $\parallel$ نشان‌دهنده الحاق و $\sigma = \text{ReLU}$ است. تجمیع‌کننده میانگین^۲ به‌صورت زیر پیاده‌سازی می‌شود:

$$h_v^{(k)} = \text{ReLU} \left(W_k \cdot \left[h_v^{(k-1)} \parallel \text{mean}_{u \in \mathcal{N}(v)} h_u^{(k-1)} \right] \right) \quad (25)$$

این شبکه دولایه SAGEConv و یک سر پرسپترون چندلایه دوبلوکه را با حذف تصادفی [۳۷] پس از هر بلوک در هم می‌آمیزد: $\text{SAGE} \rightarrow \text{ReLU} \rightarrow \text{Dropout} \rightarrow \text{SAGE} \rightarrow \text{ReLU} \rightarrow \text{Dropout} \rightarrow \text{Linear} \rightarrow \text{ReLU} \rightarrow \text{Dropout} \rightarrow \text{Linear} \rightarrow \text{logit}$.

دولایه انتشار، به هر گروه یک میدان پذیرش دوهایی^۳ می‌دهد که اعضای آن و سایر گروه‌هایی که آن اعضا به آن‌ها تعلق دارند را شامل می‌شود و سیگنال رابطه‌ای که فرضیه گراف بر آن تکیه دارد، قابل‌دسترس خواهد بود [۸]، [۲۳]. الگوریتم ۶ عبور روبه‌جلو^۴ را نشان می‌دهد.

الگوریتم (۶): عبور روبه‌جلو گراف سیج (به ازای گره v)

```
1:  $h(v) \leftarrow x_v$ 
2: for  $k = 1, 2$ :
  a  $\leftarrow \text{mean}_{u \in \mathcal{N}(v)} h_{\{k-1\}}(u)$ 
   $h_k(v) \leftarrow \text{ReLU}(W_k \cdot \text{concat}(h_{\{k-1\}}(v), a))$ 
   $h_k(v) \leftarrow \text{dropout}(h_k(v), p=0.5)$ 
3:  $z \leftarrow \text{ReLU}(W_3 \cdot h_2(v))$ ;  $z \leftarrow \text{dropout}(z)$ ;  $\text{logit}(v) \leftarrow W_4 \cdot z$ 
4: return  $\text{logit}(v)$ 
```

۳-۶-۲ شبکه عصبی کانولوشنی گراف

یک خط پایه GCN [۲۲] با عمق و سر یکسان، لایه‌های SAGEConv را با دولایه GCNConv جایگزین می‌کند که انتشار نرمال‌شده متقارن را اعمال می‌نماید:

$$H^{(k)} = \sigma \left(\tilde{D}^{-1/2} \tilde{A} \tilde{D}^{-1/2} H^{(k-1)} W_k \right), \quad (26)$$

$$\tilde{A} = A + I$$

که در آن $\tilde{A}$ ماتریس مجاورت با حلقه‌های خودی و $\tilde{D}$ ماتریس درجه آن است [۲۲]. این میانگین‌گیری ثابت و نرمال‌شده با درجه، در برابر هموارسازی بیش‌ازحد^۵ روی گراف‌های با درجه بالا آسیب‌پذیر است [۳۸]، [۳۹].

جلوگیری می‌نماید [۸]. سه مجموعه یال تودرتو با فیلتر کردن هر دو رابطه به گروه‌های قابل‌مشاهده در هر فاز ساخته می‌شوند (الگوریتم ۵).

جدول (۵): زیرگراف‌های تفکیک‌شده براساس فاز که یک محیط

استقرایی واقعی را اعمال می‌کنند

مرحله	گراف قابل‌مشاهده
آموزش	users + training groups only
اعتبارسنجی	users + training + validation groups
آزمون	full graph (users + all groups)

الگوریتم (۵): ساخت زیرگراف استقرایی

```
Input: full edge_index E, num_users n, allowed group sets per phase
1: train_groups = train_ids; val_groups = train_ids U val_ids;
   test_groups = all
2: for phase in {train, val, test}:
   keep edge (a,b) iff every group endpoint of (a,b) ∈ allowed(phase)
   (bipartite edge: its single group endpoint; projection edge: both endpoints)
3: return  $E_{\text{train}} \subseteq E_{\text{val}} \subseteq E_{\text{test}}$ 
```

از آنجاکه گراف سیج توابع تجمیع را به‌جای جاسازی‌های^۱ خاص هر گره یاد می‌گیرد، می‌تواند گروه‌های دیده‌نشده در زمان آموزش را نمایش دهد [۸].

۳-۶-۳ معماری مدل‌ها

تمام معماری‌ها در تمام آزمایش‌ها ثابت هستند و تنها پیکربندی آموزش و ابرپارامترها تنظیم می‌شوند.

۳-۶-۱ مدل پیشنهادی گراف سیج

در این پژوهش آشکارساز پیشنهادی یک شبکه گراف سیج است [۸]، که کانولوشن گراف کامل را با تجمیع همسایگی مبتنی بر نمونه‌برداری قابل‌یادگیری جایگزین می‌کند و امکان مقیاس‌پذیری روی گراف با ۷/۶۴ میلیون یال و تعمیم استقرایی را فراهم می‌سازد. برای گره v با همسایگی $\mathcal{N}(v)$ هر لایه k مقدار زیر را محاسبه می‌کند:

^۴ Forward-Pass

^۵ Over-Smoothing

^۱ Embeddings

^۲ Mean

^۳ Two-Hop Receptive Field

۷-۳ روند آموزش و بهینه‌سازی

تمام مدل‌های عصبی یک تابع زیان آنتروپی متقاطع دودویی با توازن کلاس را حداقل می‌کنند. با توجه به عدم توازن کلاس‌ها، یک وزن کلاس مثبت از فرکانس‌های آموزشی تنظیم می‌شود:

$$\mathcal{L} = -\frac{1}{N} \sum_{i=1}^N [w_+ y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)], \quad (31)$$

$$w_+ = \frac{\text{freq}(0)}{\text{freq}(1)}$$

که بدون نمونه‌برداری مجدد با عدم توازن شیوع مقابله می‌کند [۳۱]. بهینه‌سازی از Adam [۴۲] استفاده می‌کند، یک زمان‌بند ReduceLROnPlateau (ضریب ۰/۵، صبر^۲ ۱۵ روی زیان اعتبارسنجی) نرخ را کاهش می‌دهد و گرادینان‌ها در حداکثر نرم ۵ برش داده می‌شوند. Dropout سنگین با پارامتر p برابر ۰/۵ [۳۷] و کاهش وزن، مدل را تنظیم می‌کنند. برای گراف پویا، مدل در این پژوهش «چرخشی» آموزش داده می‌شود، به این صورت که در هر دوره، یک تصویر زمانی (از میان ۱۲ تصویر) انتخاب شده و این تصاویر به ترتیب چرخش می‌یابند، درحالی‌که ارزیابی همواره روی آخرین تصویر زمانی انجام می‌گیرد. تمام مدل‌های عصبی از یک بودجه مشترک ۳۰۰ دوره‌ای بهره می‌برند و دارای صبر برابر با ۴۰ بر اساس معیار F1 بر روی داده‌های اعتبارسنجی هستند و در نهایت وزن‌های مربوط به بهترین مقدار F1 بازیابی می‌شوند. آستانه تصمیم‌گیری در بازه [۰/۹۵، ۰/۰۵] (شامل ۹۱ مقدار) به گونه‌ای جستجو می‌شود که بیشترین مقدار F1 را روی داده‌های اعتبارسنجی به دست دهد، سپس همین آستانه، بدون تغییر، روی مجموعه آزمون اعمال می‌گردد. این امر، انتخاب یک نقطه عملیاتی مهم در شرایط عدم توازن داده‌ها محسوب می‌شود [۵۲، ۶۴]. الگوریتم ۷، روند کلی آموزش را خلاصه کرده است.

۱-۷-۳ تنظیم ابرپارامترها

هر مدل توسط جستجوی شبکه‌ای^۳ بر روی F1 اعتبارسنجی تنظیم می‌شود. از آنجاکه گراف فراتر از ۵۰۰۰۰۰ یال است، شبکه‌ها به‌طور خودکار به ۲ تا ۳ پیکربندی برای هر مدل کاهش می‌یابند و حالت

۳-۶-۳ شبکه CARE-GNN

مدل CARE-GNN [۴] یک آشکارساز قلب چندرابطه‌ای است. کانولوشن درون‌رابطه‌ای آن یک تجمیع آگاه از برچسب و فیلترشده با شباهت را انجام می‌دهد: برای هدف i و همسایه j ، پیام را با شباهت کسینوسی یکسوشده^۱ وزن‌دهی می‌کند و در طول آموزش، تنها p درصد برتر از شبیه‌ترین همسایگان را نگه می‌دارد:

$$s_{ij} = \max(0, \cos(h_i, h_j)), \quad (27)$$

$$m_{ij} = (Wh_j) \cdot s_{ij} \cdot 1[s_{ij} \geq Q_{1-p}], \quad p = 0.7 \quad (28)$$

که در آن، Q_{1-p} برابر با چندک $1-p$ ام (یا صدک $1-p$ ام) توزیع شباهت‌ها است. دو کانولوشن مشابه، به‌طور جداگانه بر روی روابط دویخی و فرافکنی اعمال می‌شوند. سپس، یک لایه Softmax با وزن‌های قابل یادگیری، خروجی این دو کانولوشن را به‌صورت وزن‌دار با یکدیگر ادغام می‌کند و حاصل، پیش از ورود به بخش پرسپترون چندلایه، ترکیب می‌شود:

$$h = \text{ReLU}(\alpha_1 h_{\text{bipartite}} + \beta_1 h_{\text{projection}}), \quad (29)$$

$$(\alpha_1, \beta_1) = \text{softmax}(\text{inter_weights})$$

CARE-GNN تنها رقیبی است که به‌طور صریح از رابطه فرافکنی بهره‌برداری می‌کند که با توجه به تقریباً خالی بودن آن (بخش ۳-۷)، در بخش ۴-۱۰ تعیین‌کننده است.

۴-۶-۳ مدل‌های پایه جدولی کلاسیک و عصبی

چهار خط پایه غیرگرافی، مقایسه را کامل می‌کنند. مدل لایه‌های GroupMLP در رابطه ۳۰ درج شده است.

$$\text{Linear}(d) \rightarrow 128 \rightarrow \text{ReLU} \rightarrow \text{Dropout} \rightarrow \text{Linear}(128 \rightarrow 64) \rightarrow \text{ReLU} \rightarrow \text{Linear}(64 \rightarrow 1) \quad (30)$$

مدل را بر روی بردار تخت گروه بدون گراف اعمال می‌کند. خطوط پایه کلاسیک عبارت‌اند از رگرسیون لجستیک، جنگل تصادفی [۴۰] و یک گروه تقویت‌شده با گرادینان تنظیم‌شده به نام XGBoost [۴۱]. یادگیرنده‌های جدولی قوی و پرکاربرد که به‌عنوان مراجع چالش‌برانگیز عمل می‌کنند [۴۰]، [۴۱].

³ Grid Search

¹ Rectified

² Patience

جدول (۶): ابرپارامترهای انتخاب شده برای هر مدل

مدل	ابروپارامترهای تنظیم شده
گراف‌سیج (پیشنهادی)	hidden ۱۲۸, dropout ۰/۵, lr 5×10^{-4} , weight-decay 5×10^{-4} , aggr = mean
GCN	hidden ۱۲۸, dropout ۰/۵, lr 5×10^{-4} , weight-decay 5×10^{-4}
CARE-GNN	hidden ۶۴, dropout ۰/۴, lr 1×10^{-3} , weight-decay 1×10^{-4} , top _p ۰/۷
Random Forest	۲۰۰ trees, unlimited depth, min_split ۲
XGBoost	۲۰۰ trees, max_depth ۴, lr ۰/۰۵, subsample ۰/۸

۴- نتایج و ارزیابی

این بخش یافته‌های تجربی را گزارش می‌دهد. عملکرد مطلق مدل گراف‌سیج پیشنهادی (بخش ۴-۱)، پویایی‌های یادگیری آن (بخش‌های ۴-۲ و ۴-۳)، مقایسه در برابر تمام خطوط پایه از نظر صحت و هزینه (بخش‌های ۴-۴ و ۴-۵) و تحلیل‌ها از طریق مطالعه حذف، پایداری چندبذری و آزمون معناداری، تحلیل خطا، اعتبارسنجی توزیع، اهمیت ویژگی‌ها و تشخیص‌های ساختار گراف (بخش‌های ۴-۶ تا ۴-۱۰) ارائه خواهد یافت.

۴-۱ عملکرد مدل پیشنهادی

مدل گراف‌سیج تنظیم شده برای بودجه کامل آموزش داده شد، و در دوره ۲۵۸ از ۲۹۸ با آستانه تصمیم بهینه شده ۰/۸۹ متوقف گردید. بر روی گراف آزمون کاملاً استقرایی و نگاه داشته شده^۶، این مدل به نتایج جدول ۷ دست می‌یابد.

جدول (۷): عملکرد مجموعه آزمون گراف‌سیج (۱۵۰۰ گروه)

میانگین دقت	AUC	F1	فراخوانی	دقت	صحت
۰/۹۹۵۱	۰/۹۹۸۷	۰/۹۵۴۲	۰/۹۳۶۷	۰/۹۷۲۳	۰/۹۸۲۰

چرخش استفاده می‌شود که جامعیت را فدای زمان اجرای قابل مدیریت می‌کند. پیکربندی‌های انتخاب شده در جدول ۶ آمده است.

الگوریتم (۷): آموزش GNN استقرایی با بهینه‌سازی آستانه

```

Input: model, E_train/E_val/E_test, snapshots {0..11},
budget=300, patience=40
1: criterion ← BCEWithLogits(pos_weight=freq0/freq1); opt ← Adam(5e-4, wd=5e-4)
2: best_f1 ← -1; wait ← 0
3: for epoch = 1..300:
  snap ← rotate: epoch mod 12
  X_G ← scaler.transform(features(snap)); update group node features
  logits ← model(X, E_train); loss ← criterion(logits[train], y[train])
  backprop; clip_grad_norm(5); opt.step()
  evaluate on E_val (final snapshot); scheduler.step(val_loss)
  if val_f1 > best_f1: best_f1 ← val_f1; save weights; wait ← 0
  else wait += 1
  if wait ≥ 40: break
4: restore best weights
5:  $\tau^* \leftarrow \operatorname{argmax}_{\tau \in [0.05, 0.95]} F1\_val(\tau)$ ; report test metrics at  $\tau^*$ 

```

۳-۸ پروتکل ارزیابی

مدل‌ها با مجموعه استاندارد برای طبقه‌بندی دودویی نامتوازن ارزیابی می‌شوند: صحت^۱، دقت^۲، فراخوانی^۳، معیار F1، مساحت زیر منحنی ROC^۴ و میانگین دقت (مساحت زیر منحنی دقت-فراخوانی^۵). از آنجاکه صحت تحت شرایط عدم توازن گمراه‌کننده است، بر F1 و معیارهای مستقل از آستانه AP/AUC تأکید می‌شود؛ نمای دقت-فراخوانی برای مسائل با مثبت‌های نادر بسیار آگاهی‌بخش است [۴۳]، [۴۴]. برای ارزیابی استحکام^۶، مدل پیشنهادی روی پنج بذر تصادفی (۴۲، ۵۲، ۶۲، ۷۲، ۸۲) مجدد آموزش داده شده و میانگین $\pm$ انحراف معیار گزارش می‌گردد. ادعای اینکه مدل گراف خط پایه جدولی را شکست می‌دهد با استفاده از آزمون رتبه علامت‌دار ویلکاکسون^۷ [۴۵] یک‌طرفه بر روی امتیازهای زوجی بذرها آزمایش می‌شود، که روش ناپارامتری توصیه شده برای مقایسه طبقه‌بندهاست [۴۶]. زمان‌های آموزش و استنتاج به صورت زمان واقعی^۸ برای هر مدل ثبت می‌شود.

^۶ Robustness

^۷ WilcoxonS rank-ignedTest

^۸ Wall-clock

^۹ Held-out

^۱ Accuracy

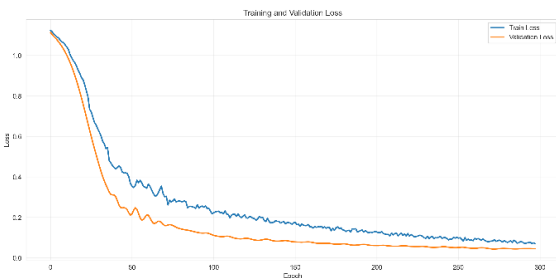
^۲ Precision

^۳ Recall

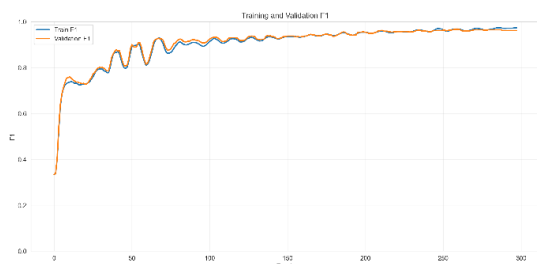
^۴ AUC

^۵ AP

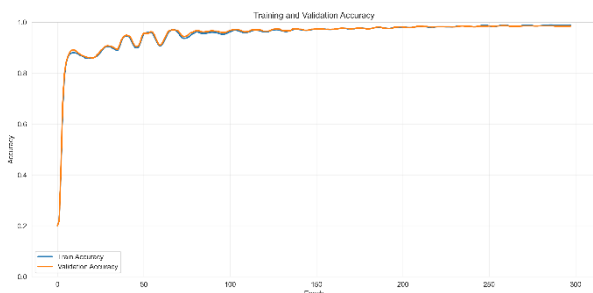
(۶) در طول آموزش روی ~ 0.99 ثابت می‌ماند. هیچ شواهدی از بیش‌برازش وجود ندارد [۳۷].



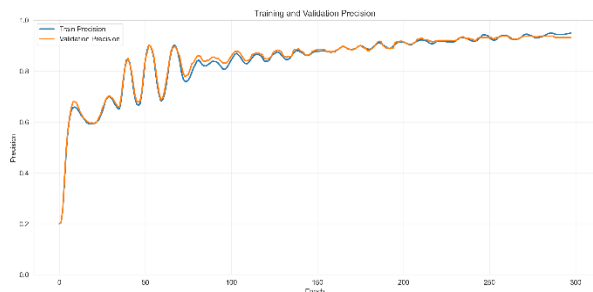
شکل (۲): زیان آموزش در برابر اعتبارسنجی در طول ۲۹۸ دوره



شکل (۳): معیار F1 آموزش در برابر اعتبارسنجی — منحنی‌ها به دقت یکدیگر را دنبال می‌کنند

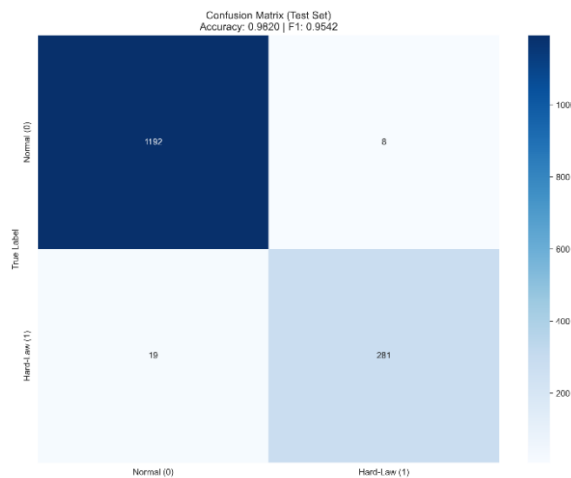


شکل (۴) صحت آموزش در برابر اعتبارسنجی، افزایش سریع سپس یک سطح پایدار نزدیک به $0.99 - 0.98$



شکل (۵): منحنی‌های دقت رسیدن همگرایی به $\sim 0.93 - 0.90$

در شرایط مطلق، عملکرد مدل قوی است: بر روی ۱۵۰۰ گروه آزمون (۱۲۰۰ عادی، ۳۰۰ سخت‌قانون) این مدل تنها ۲۷ خطا (۸ مثبت کاذب و ۱۹ منفی کاذب) مرتکب می‌شود که منجر به ۱۱۹۲ منفی صحیح و ۲۸۱ مثبت صحیح می‌گردد (شکل ۱).



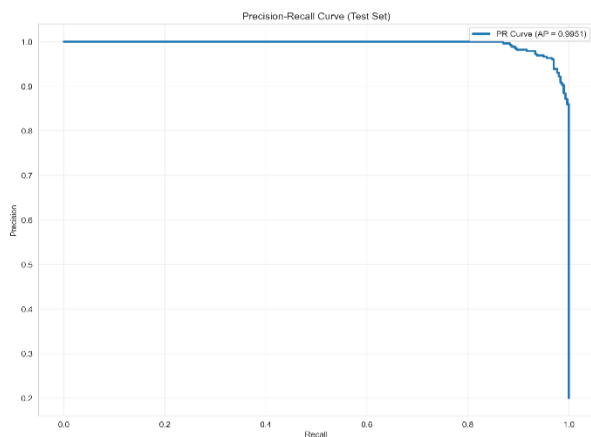
شکل (۱): ماتریس درهم‌ریختگی مجموعه آزمون در مدل گراف‌سیج تعداد بسیار پایین مثبت کاذب (دقت 0.972) به این معنی است که تقریباً هر گروهی که پرچم‌گذاری شده، واقعاً پرخطر است، درحالی‌که ۱۹ منفی کاذب (فراخوانی 0.937) خطاهای پرهزینه‌تری برای تعدیل‌سازی^۱ هستند و در بخش ۴-۸ توصیف می‌شوند.

۴-۲ پویایی‌های یادگیری و تحلیل بیش‌برازش

هر دو زیان آموزش و اعتبارسنجی به آرامی از $\sim 1/12$ به ~ 0.05 طی ۲۹۸ دوره کاهش می‌یابند (شکل ۲). زیان اعتبارسنجی به‌طور پیوسته در زیر زیان آموزش قرار دارد که پیامد موردانتظار استفاده از Dropout سنگین ($p = 0.5$) و وزن‌دهی متوازن کلاسی است که زیان آموزش اندازه‌گیری‌شده را برجسته می‌کند، همراه با یک زیرگراف اعتبارسنجی که اندکی آسان‌تر است. صحت و F1 (شکل‌های ۳ و ۴) در ۲۵ دوره اول به‌سرعت افزایش می‌یابند، از یک مرحله اکتشافی کوتاه تا حدود دوره ۷۰ می‌گذرند و سپس در یک سطح پایدار نزدیک به $0.98 - 0.99$ مستقر می‌شوند، به‌طوری‌که منحنی‌های آموزش و اعتبارسنجی اساساً روی هم قرار می‌گیرند. دقت (شکل ۵) به $\sim 0.93 - 0.90$ همگرا می‌شود و فراخوانی (شکل

¹ Moderation

هستند که همگی از گراف‌سیج پیشی می‌گیرند و در سطح ۲ مدل‌ها گراف‌سیج و پرسپترون چندلایه ($F1 \approx 954/0 - 956/0$) قرار دارند. مدل گراف‌سیج در صحت با پرسپترون چندلایه برابری می‌کند اما دقت به‌طور چشمگیری بهتری دارد، آن هم با هزینه محاسباتی زیاد. در سطح ۳ GCN و GNN-CARE به $F1 = 333/0$ سقوط می‌کنند (بخش ۴-۱۰). تسلط مدل‌های جدولی پیامد مستقیم این است که برچسب تابعی از ویژگی‌های ایستای گروه است (بخش ۴-۹)، که مدل‌های تقویت گرادیان و خطی قوی آن را تقریباً به‌طور کامل بازیابی می‌کنند [۴۱]، [۴۰].

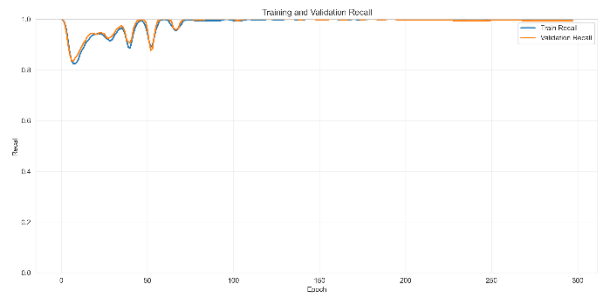


شکل (۸): منحنی دقت-فراخوانی (آزمون)، میانگین دقت

$$(AP) = 0.9951$$

۴-۵ هزینه محاسباتی

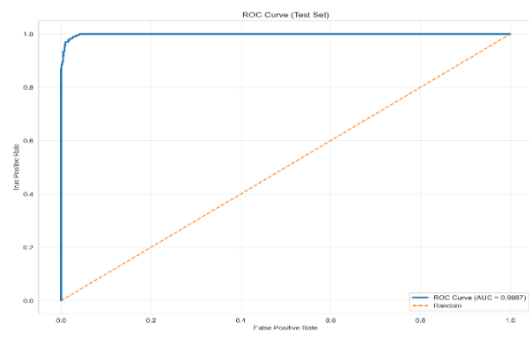
شکاف هزینه-فایده اصلی‌ترین یافته عملی است (شکل ۱۰). مدل گراف‌سیج برای آموزش به حدود ۹۸۲/۱ ثانیه (≈ 33 دقیقه) زمان نیاز دارد، درحالی‌که این زمان برای XGBoost برابر ۱/۴۲ ثانیه و برای رگرسیون لجستیک ۰/۰۳ ثانیه است، اختلاف ۳۵۰ برابری تا ۶۰۰۰ برابری، درحالی‌که هیچ مزیت صحتی بر روی این معیار ارائه نمی‌دهد. برای وظیفه‌ای که سیگنال آن جدولی است، مدل گراف گزینه‌ای گران‌تر و اندکی ضعیف‌تر است، این تبادلی کمی^۲ به‌نوبه خود یک نتیجه قابل گزارش است.



شکل (۶): منحنی‌های فراخوانی که در طول آموزش به $\approx 99/0$ رسیدند و ثابت می‌مانند

۴-۳ تفکیک‌پذیری: ROC و Precision-Recall

هر دو منحنی مستقل از آستانه، تفکیک‌پذیری تقریباً کامل را تأیید می‌کنند: $AUC = 9987/0$ (شکل ۷) و میانگین دقت برابر با 0.9951 (شکل ۸). مدل، دو کلاس را در فضای امتیاز مرتب می‌کند، بنابراین $F1$ برابر با 0.954 تنها منعکس‌کننده تبادلی^۱ در نقطه کار با آستانه 0.89 است، نه رتبه‌بندی ضعیف. پایین آوردن آستانه در صورتی‌که از دست دادن گروه‌های پرخطر پرهزینه‌تر از هشدارهای کاذب باشد، دقت را فدای فراخوانی می‌کند، تصمیمی که منحنی دقت-فراخوانی آن را به‌صراحت نشان می‌دهد [۴۳].



شکل (۷): منحنی ROC (آزمون)، $AUC = 0.9987$ ، مماس بر

گوشه بالا-چپ

۴-۴ نتایج مقایسه‌ای

جدول ۸ و شکل ۹ تمام هفت مدل را بر روی مجموعه آزمون گزارش می‌دهند. سه سطح واضح پدیدار می‌شود. سطح ۱ یادگیرنده‌های جدولی XGBoost ($F1 = 992/0$) و جنگل تصادفی (0.988) و رگرسیون لجستیک (0.988) قوی‌ترین

² Quantified Trade-off

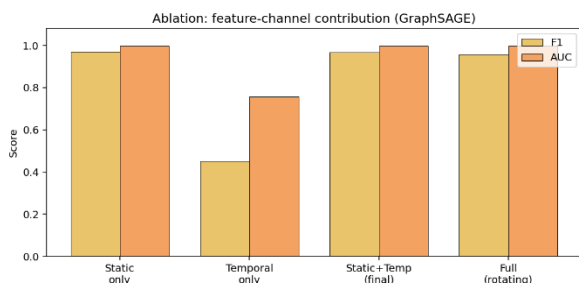
¹ Trade-off

جدول (۸): نتایج مقایسه‌ای مجموعه آزمون برای تمام مدل‌ها

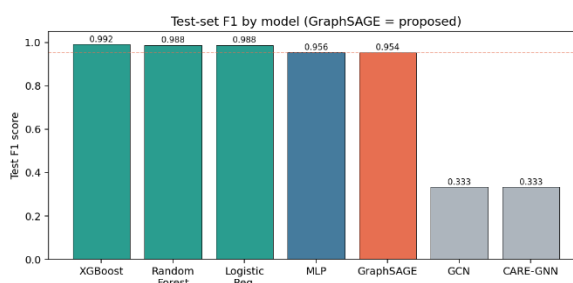
مدل	صحت	دقت	فراخوانی	F1	AUC	آموزش (ثانیه)	استنتاج (ثانیه)
XGBoost (تنظیم شده)	۰/۹۹۶۷	۰/۹۸۳۶	۱/۰۰۰۰	۰/۹۹۱۷	۱/۰۰۰۰	۱/۴۲	۰/۰۰۷
Random Forest	۰/۹۹۵۳	۰/۹۹۰۰	۰/۹۸۶۷	۰/۹۸۸۳	۱/۰۰۰۰	۵/۵۳	۰/۰۶۳
Logistic Regression	۰/۹۹۵۳	۰/۹۹۳۳	۰/۹۸۳۳	۰/۹۸۸۳	۰/۹۹۹۹	۰/۰۳	۰/۰۰۰
MLP	۰/۹۸۲۰	۰/۹۳۳۳	۰/۹۸۰۰	۰/۹۵۶۱	۰/۹۹۸۵	۳/۷۲	۰/۰۰۰
گراف سیج (پیشنهادی)	۰/۹۸۲۰	۰/۹۷۲۳	۰/۹۳۶۷	۰/۹۵۴۲	۰/۹۹۸۷	۱۹۸۲/۳	۴/۵۷
GCN (تنظیم شده)	۰/۲۰۰۰	۰/۲۰۰۰	۱/۰۰۰۰	۰/۳۳۳۳	۰/۵۴۱۲	۵۰۲/۳	۶/۸۷
CARE-GNN (تنظیم شده)	۰/۲۰۰۰	۰/۲۰۰۰	۱/۰۰۰۰	۰/۳۳۳۳	۰/۴۹۵۴	۳۹۷/۳	۴/۴۵

جدول (۹): مطالعه حذف بر روی کانال‌های ویژگی

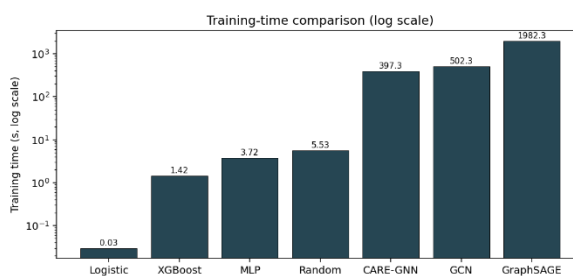
پیکربندی	صحت	دقت	فراخوانی	F1	AUC
فقط ویژگی‌های ایستا	۰/۹۸۸۷	۰/۹۷۳۲	۰/۹۷۰۰	۰/۹۷۱۶	۰/۹۹۹۲
فقط ویژگی‌های زمانی	۰/۶۳۶۰	۰/۳۲۲۸	۰/۷۴۶۷	۰/۴۵۰۷	۰/۷۵۷۲
ایستا + زمانی (final snapshot)	۰/۹۸۷۳	۰/۹۷۳۱	۰/۹۶۳۳	۰/۹۶۸۲	۰/۹۹۸۸
کامل (ایستا + زمانی، چرخشی)	۰/۹۸۳۳	۰/۹۶۶۱	۰/۹۵۰۰	۰/۹۵۸۰	۰/۹۹۸۷



شکل (۱۱): ویژگی‌های ایستا تقریباً تمام سیگنال را حمل می‌کنند، ویژگی‌های زمانی به‌تنهایی ضعیف هستند و هنگام اضافه‌شدن اندکی به عملکرد آسیب می‌زنند. ویژگی‌های ایستا اساساً تمام سیگنال تفکیک‌پذیر را حمل می‌کنند (به‌تنهایی $F1 = ۰/۹۷۲$)، که پیامد مستقیم تعریف برچسب است. ویژگی‌های زمانی به‌تنهایی ضعیف هستند ($F1 = ۰/۴۵۱$, $AUC =$



شکل (۹): معیار F1 مجموعه آزمون بر اساس مدل. یادگیرنده‌های جدولی (سبز) پیشرو هستند؛ گراف‌سیج (نارنجی) مدل پیشنهادی است و مدل‌های GCN/CARE-GNN فرومی‌باشند



شکل (۱۰): مقایسه زمان آموزش (مقیاس لگاریتمی): گراف‌سیج بین ۳۵۰ تا ۶۰۰۰ برابر کندتر از مدل‌های جدولی است

۴-۶ مطالعه حذفی^۱

برای جداسازی مشارکت هر کانال ویژگی، گراف‌سیج تحت چهار پیکربندی ویژگی مجدد آموزش داده شد (جدول ۹ و شکل ۱۱).

^۱ Ablation Study

(۵۱۵ شرکت‌کننده) با فعالیت/مشارکت بالاتر هستند: گروه‌های کوچک با فعالیت مرزی که به اشتباه پرچم‌گذاری شده‌اند. مثبت‌های صحیح و منفی‌های صحیح میانگین ویژگی تقریباً یکسانی دارند، که تأیید می‌کند مدل از یک همبستگی جعلی بهره‌برداری نمی‌کند. جدول (۱۱): میانگین ویژگی‌ها بر اساس نتیجه پیش‌بینی

دسته بندی	تعداد	activity	silent	engage	reply	participants
مثبت صحیح	۲۸۱	۳۳۹/۰	۶۶۱/۰	۳۴۳/۰	۹۷/۱	۸۳۷
منفی صحیح	۱۱۹۲	۳۲۸/۰	۶۷۲/۰	۳۳۳/۰	۹۷/۱	۸۴۵
مثبت کاذب	۸	۳۹۵/۰	۶۰۵/۰	۴۰۵/۰	۳۸/۲	۵۱۵
منفی کاذب	۱۹	۲۸۵/۰	۷۱۵/۰	۲۹۵/۰	۵۳/۲	۱۵۲۱

۴-۹ اعتبارسنجی توزیع و اهمیت ویژگی‌ها

اعتبارسنجی توزیع داخلی (جدول ۱۲) درواقع یک عیب‌یابی برچسب است. سه ویژگی که امتیاز ریسک را تعریف می‌کنند، کلاس‌ها را تقریباً به‌طور کامل از هم جدا می‌کنند ($KS \approx 98/0$ و $JS \approx 81/0$)، درحالی‌که ویژگی‌های غایب در امتیاز از نظر آماری بین کلاس‌ها غیرقابل تشخیص هستند [۳۵]، [۳۶] ($KS \approx 0.2/0$ و $p \ll 0.05$).

تحلیل اهمیت ویژگی‌ها این موضوع را از سمت مدل نیز تأیید می‌کند (شکل ۱۳). هم دیدگاه مبتنی بر ناخالصی (جنگل تصادفی [۴۰]) و هم دیدگاه خطی (رگرسیون لجستیک) توافق دارند که `engagement_level` و `activity_ratio` و `silent_ratio` غالب هستند و در مجموع حدود ۷۹٪ از اهمیت جنگل تصادفی و بزرگ‌ترین ضرایب لجستیک را با علائم موردانتظار تشکیل می‌دهند (جدول ۱۳). طبقه‌بندها به‌درستی تعریف ریسک مولد را بازایی کرده‌اند.

۷۵۷/۰: پویایی‌های AR(۱)/فوران تنها همبستگی ضعیفی با برچسب دارند. جالب‌توجه است که افزودن اطلاعات زمانی اندکی عملکرد را کاهش می‌دهد، بنابراین کانال زمانی برای این برچسب بیش از سیگنال، نویز وارد می‌کند. درنتیجه، مدل‌سازی زمانی در اینجا باید به‌جای منبع دقت، به‌عنوان یک مطالعه عمومیت‌بخشی در نظر گرفته شود [۴۷]، [۴۸].

۴-۷ پایداری و معناداری آماری

در سراسر پنج بذر، گراف‌سیج پایدار است (جدول ۱۰)، اما مقایسه زوجی با بهترین خط پایه ادعای برتری را تأیید نمی‌کند.

جدول (۱۰): پایداری پنج‌بذری برای گراف‌سیج

معیار	گراف‌سیج (mean $\pm$ std)
صحت	$0.976/0 \pm 0.013/0$
دقت	$0.919/0 \pm 0.060/0$
فراخوانی	$0.967/0 \pm 0.011/0$
F1	$0.942/0 \pm 0.028/0$
AUC	$0.998/0 \pm 0.001/0$

آزمون یک‌طرفه رتبه‌ای علامت‌دار ویلکاکسون [۴۵]، [۴۶] برای فرضیه « $XGBoost < GraphSAGE$ » بر روی پنج امتیاز ۱F زوجی، آماره‌ای برابر با ۰/۰ و مقدار p برابر با ۱/۰ را نشان می‌دهد که در سطح معناداری $\alpha = 0.05$ معنادار نیست. در هر بذر، مقدار ۱F مربوط به $XGBoost$ (حدود ۰/۹۸۸ تا ۰/۹۹۲) از گراف‌سیج (۰/۸۹۷ تا ۰/۹۷۰؛ میانگین ۰/۹۴۱۵ در برابر ۰/۹۹۰۱) بالاتر بوده است. بنابراین، در این معیار مصنوعی، مدل گراف‌سیج از نظر آماری بهتر از خط‌پایه جدولی عمل نمی‌کند و درعین حال به‌مراتب گران‌تر از آن است.

۴-۸ تحلیل خطا

میانگین مقادیر ویژگی به ازای دسته پیش‌بینی (جدول ۱۱) نشان می‌دهد که خطاها به‌جای تکیه بر میان‌برهای تصادفی، در مرز واقعی کلاس متمرکز شده‌اند.

منفی‌های کاذب در ۱۹ گروه پرخطری که نادیده گرفته شده‌اند؛ به‌طور چشمگیری بزرگ‌تر هستند (میانگین ۱۵۲۱ شرکت‌کننده در برابر ≈ 840 در کل) و بالاترین نرخ سکوت (۰/۷۱۵) را دارند: گروه‌های بزرگ و خاموش نزدیک به مرز که با یک آستانه آگاه از اندازه یا ویژگی قابل بازایی هستند. مثبت‌های کاذب (۸) کوچک‌تر

جدول (۱۲): اعتبارسنجی توزیع داخلی برچسب ۰ در برابر برچسب ۱

ویژگی	KS	مقدار آماره p-value	JS	μ (۰ برچسب)	μ (۱ برچسب)
activity_ratio	۹۸۲/۰	۰/۰	۸۱۲/۰	۳۸۷/۰	۱۰۹/۰
silent_ratio	۹۸۲/۰	۰/۰	۸۱۲/۰	۶۱۳/۰	۸۹۱/۰
engagement_level	۹۷۷/۰	۰/۰	۸۱۰/۰	۳۹۳/۰	۱۱۴/۰
reply_count	۰۱۸/۰	۷۰۷/۰	۰۳۱/۰	۰۲/۲	۹/۱
participants	۰۲۴/۰	۳۳۱/۰	۰۳۷/۰	۸۲۸	۸۳۶
growth_proxy	۰۱۹/۰	۵۹۰/۰	۰۷۲/۰	۰۱۲/۰	۰۱۲/۰

فراکنی تنها شامل ۶۹ یال است که میان ۹۹۶۳ مؤلفه مجزا توزیع شده‌اند. در چنین ساختاری، هر دو مدل GCN و CARE-GNN در عمل به یک طبقه‌بند ساده تبدیل شدند که همه نمونه‌ها را در کلاس «سخت‌قانون» قرار می‌داد. همچنین، آستانه‌های بهینه تصمیم‌گیری این دو مدل به کمترین مقدار ممکن رسید؛ موضوعی که نشان می‌دهد هیچ الگوی امتیازدهی معناداری برای تفکیک دو کلاس توسط این مدل‌ها شکل نگرفته است. دو مکانیزم مسئول این پدیده هستند:

الف) هموارسازی بیش‌ازحد در گرافی بزرگ و با گره‌های پردرجه: علت اصلی این رفتار را نمی‌توان صرفاً به عمق شبکه نسبت داد؛ بلکه نحوه ترکیب اطلاعات گره با اطلاعات همسایگان نقش تعیین‌کننده‌ای در آن دارد. در GCN، انتشار اطلاعات از طریق ماتریس نرمال‌سازی متقارن $\tilde{D}^{-1/2} \tilde{A} \tilde{D}^{-1/2}$ انجام می‌شود. در نتیجه، گره‌هایی با درجه بالا (که در این گراف تا ۵۰۰۷ نیز می‌رسد) هنگام به‌روزرسانی نمایش خود، اطلاعاتشان را با میانگین تعداد بسیار زیادی از همسایگان ترکیب می‌کنند [۲۲]، [۳۸]. این فرآیند باعث می‌شود ویژگی‌های متمایزکننده گره به‌تدریج کمرنگ شود و پس از تنها دو مرحله انتشار، نمایش برداری گروه‌ها تا حد زیادی به یکدیگر شبیه شود. در چنین شرایطی، توانایی مدل در تفکیک کلاس‌ها از بین می‌رود؛ پدیده‌ای که با عنوان هموارسازی بیش‌ازحد در GCN های به‌کاررفته روی گراف‌های مترکم شناخته می‌شود [۳۸]، [۳۹]. پیامد نهایی این وضعیت نیز گرایش مدل به پیش‌بینی کلاس غالب برای همه نمونه‌ها است. در مقابل، روش تجمیع قابل‌یادگیری مبتنی بر الحاق در گراف‌سیج، که راه‌حل توصیه‌شده در این پژوهش است [۸]، دچار این فروپاشی نمی‌شود.

جدول (۱۳): برترین ویژگی‌ها بر اساس اهمیت جنگل تصادفی و ضریب رگرسیون لجستیک

ویژگی	اهمیت جنگل تصادفی	ضریب رگرسیون لجستیک
silent_ratio	۳۱۴/۰	۲۵/۷+
activity_ratio	۲۵۸/۰	۲۵/۷-
engagement_level	۲۱۶/۰	۱۷/۸-
interaction_score	۰۸۶/۰	۱۳/۲-
active_density	۰۳۵/۰	۹۶/۱-

۴-۱۰ ساختار گراف و فروپاشی GCN/CARE-GNN

آماره‌های گراف آموزشی (جدول ۱۴) هم ضعیف بودن رابطه فراکنی و هم علت شکست دو مورد از مدل‌های گراف را توضیح می‌دهد.

جدول (۱۴): آماره‌های گراف آموزشی

مقدار	معیار
۱۱۰,۰۰۰	گره‌های ترکیبی
۷,۶۴۱,۳۲۹	یال‌های ترکیبی
۰۰۱۲۶/۰	چگالی ترکیبی
۷,۶۴۱,۲۶۰	یال‌های دوبخشی
۶۹	یال‌های فراکنی
۹/۱۳۸	درجه میانگین
۵,۰۰۷	حداکثر درجه
۱	مؤلفه‌های همبند (ترکیبی)
۹,۹۶۳	مؤلفه‌های همبند (فراکنی)
۰۰۲۴/۰	خوشه‌بندی میانگین (فراکنی)

گراف ترکیبی از یک مؤلفه همبند واحد تشکیل شده است که بیش از ۶۴/۷ میلیون یال دوبخشی را در برمی‌گیرد؛ در مقابل، گراف

ارائه نمی‌دهد، زیرا برچسب، طبق طراحی، یک تابع قطعی از چند ویژگی گره ایستا است که به‌طور مستقل توسط مطالعه حذف، تحلیل‌های اعتبارسنجی توزیع و اهمیت ویژگی‌ها و آزمون غیرمعمول و یلکاکسون تأیید شد. علاوه بر این، GCN و CARE-GNN تحت گراف دوبخشی بزرگ با درجه بالا و یک رابطه فرافکنی خالی فرومی‌باشند (بخش ۴-۱۰). این نتایج مسیر روشنی را به‌سوی سهم قوی‌تر گراف ترسیم می‌کنند: بازتعریف برچسب به‌گونه‌ای که به ویژگی‌های همسایگی/زمانی بستگی داشته باشد (مانند نسبت ربات‌ها در میان اعضای مشترک، یا خطر ناشی از فوران) و غنی‌سازی گراف فرافکنی (یک آستانه ژاکارد پایین‌تر، یا هم‌عضویتِ جمع‌شده در تمام تصاویر زمانی)، تا اطلاعات رابطه‌ای و زمانی به‌جای اینکه مازاد باشند، ضروری گردند [۴، ۲۴، ۴۷، ۴۸]. زیرساخت برای چنین آزمایشی در حال حاضر در پایپ‌لاین وجود دارد.

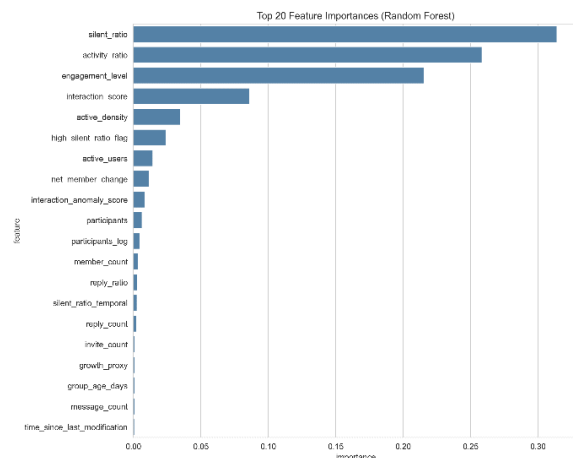
۵- بحث

هدف این بخش، تفسیر نتایج در چارچوب یادگیری گراف و تبیین دامنه کاربرد رویکرد پیشنهادی است. مهم‌ترین یافته پژوهش، برتری مدل‌های جدولی با هزینه محاسباتی حدود ۳۵۰ تا ۶۰۰۰ برابر کمتر نسبت به گراف‌سیج از نظر F1 است، نتیجه‌ای که بیش از آنکه ناشی از ضعف GNN ها باشد، به نحوه تعریف برچسب‌ها و ماهیت داده‌ها بازمی‌گردد.

۵-۱ پارادوکس مدل پیچیده در برابر مدل ساده

برچسب گروه‌های سخت‌قانون، همان‌گونه که در بخش ۳-۳ بیان شد، از امتیاز ریسکی مبتنی بر PCA روی پنج ویژگی ایستای گروه (نرخ سکوت، ناهنجاری تعامل، رکود رشد، کم‌مشارکتی و فعالیت ربات) استخراج شده است. بنابراین، برچسب به توپولوژی گراف، ویژگی‌های همسایگان یا روابط بین‌گروهی وابسته نیست. در نتیجه، مدل‌های جدولی با دسترسی مستقیم به این ویژگی‌ها تمام سیگنال لازم برای طبقه‌بندی را در اختیاردارند و انتشار پیام در گراف اطلاعات جدیدی تولید نمی‌کند. این نتیجه با مطالعه حذفی (بخش ۴-۶)، اعتبارسنجی توزیع (بخش ۴-۹) و آزمون و یلکاکسون (بخش ۴-۷) نیز تأیید شد.

ب) ناکافی بودن سیگنال رابطه‌ای در CARE-GNN: معماری CARE-GNN برای بهره‌گیری هم‌زمان از دو نوع رابطه، یعنی رابطه دوبخشی و رابطه فرافکنی، طراحی شده است. با این حال، از آنجاکه گراف فرافکنی تنها ۶۹ یال دارد، یکی از مهم‌ترین منابع اطلاعاتی این مدل عملاً فاقد محتوای مؤثر است. از سوی دیگر، فرآیند جمع‌بندی مبتنی بر فیلتر شباهت در رابطه دوبخشی نیز تحت تأثیر همان پدیده هموارسازی بیش‌ازحد قرار می‌گیرد و در نتیجه، توان مدل برای استخراج الگوهای متمایزکننده کاهش می‌یابد [۴]. بنابراین، حتی انتخاب زیرمجموعه‌ای از همسایگان نیز نمی‌تواند اطلاعات کافی برای جداسازی دو کلاس فراهم کند. در مجموع، این نتایج را نباید به‌عنوان نشانه‌ای از ضعف ذاتی GCN یا CARE-GNN تفسیر کرد، بلکه بیشتر بیانگر آن است که در گراف‌های دوبخشی بسیار بزرگ و نامتوازن، گراف‌سیج به دلیل نحوه حفظ اطلاعات گره و شیوه جمع‌بندی همسایگان، پایداری و توان تفکیک بالاتری از خود نشان می‌دهد [۴، ۲۲].



شکل (۱۲): ویژگی برتر از نظر اهمیت در جنگل تصادفی که سه ویژگی غالب هستند

۴-۱۱ خلاصه یافته‌ها

آشکارساز گراف‌سیج پیشنهادی تحت یک پروتکل استقرایی دقیق بدون هیچ‌گونه بیش‌برازش به عملکرد مطلق بسیارخوبی دست می‌یابد (صحت ۰/۹۸۲، $F1 = ۰/۹۵۴$ ، $AUC = ۰/۹۹۸۷$) و در بین بذرها پایدار است. با این حال، مهم‌ترین و صادقانه‌ترین یافته دارای ظرافت‌هایی است: بر روی این معیار مصنوعی، ساختار گراف هنوز مزیت قابل‌اندازه‌گیری نسبت به مدل‌های جدولی بسیار ارزان‌تر

[۳۲]. این داده‌ها همچنان فاقد برخی خصوصیات مهم شبکه‌های واقعی از جمله نویزهای توپولوژیک، ساختارهای پیچیده اجتماع‌محور که به‌طور کامل با مدل بلوک تصادفی بازتولید نمی‌شوند و رفتارهای تطبیقی مهاجمان هستند [۲۹]، [۴]. از این رو، تعمیم نتایج به شبکه‌های واقعی تلگرام باید با احتیاط انجام شود و کمبود داده‌های واقعی برچسب‌خورده همچنان یکی از چالش‌های اصلی این حوزه است.

هم‌چنین، تعریف برچسب بر پایه PCA (بخش ۳-۳-۴) مستقیماً بر ویژگی‌های ایستای گروه استوار است. هرچند این انتخاب کنترل آزمایش و تکرارپذیری را افزایش می‌دهد، دامنه اعتبار نتایج را نیز به مسائلی محدود می‌کند که برچسب با ویژگی‌های محلی گره همسو باشد. در کاربردهای واقعی، ممکن است گروه‌های سخت‌قانون تنها از طریق الگوهای رفتاری هماهنگ، وابستگی‌های زمانی یا ارتباط با گروه‌های پرخطر قابل‌شناسایی باشند؛ در چنین شرایطی انتظار می‌رود مزیت مدل‌های جدولی کاهش یافته و GNN های فضایی-زمانی عملکرد بهتری ارائه دهند.

۶- نتیجه‌گیری و محدودیت‌ها

این پژوهش با موفقیت رویکردی مبتنی بر شبکه‌های عصبی گرافی، به‌ویژه گراف‌سیج را برای شناسایی گروه‌های سخت‌قانون که با رفتارهای اجباری و رشد غیرطبیعی مشخص می‌شوند، توسعه داد. به دلیل کمبود داده‌های واقعی و برچسب‌دار، یک پایپ‌لاین تولید داده‌های مصنوعی ایجاد شد که ساختارهای پیچیده شبکه‌های اجتماعی واقعی، الگوهای رفتاری متنوع، پویایی‌های زمانی و روابط دویخشی کاربر-گروه را با دقت مدل‌سازی می‌کند. ارزیابی‌ها نشان داد که مدل پیشنهادی گراف‌سیج با دستیابی به دقت فراخوانی ۹۳/۷٪ و معیار AUC برابر با ۹۹/۸۷٪ در یک محیط کاملاً استقرایی، عملکرد قدرتمندی از خود نشان می‌دهد و نسبت به شبکه‌های گرافی دیگری نظیر GCN و CARE-GNN که در مواجهه با گراف‌های متراکم و دارای درجات بالا دچار مشکل هموارسازی بیش‌ازحد می‌شوند، بسیار مقاوم‌تر است.

با این وجود، تحلیل‌های دقیق مقایسه‌ای و مطالعه حذف نشان داد که در چارچوب فعلی، ویژگی‌های ایستای گروه، غالب‌ترین سیگنال‌های تشخیص را فراهم می‌کنند؛ در نتیجه مدل‌های کلاسیک

باین‌حال، این نتیجه به همه مسائل تشخیص گروه‌های ناهنجار قابل تعمیم نیست. بر اساس این پژوهش و مطالعات پیشین، GNN ها در شرایط زیر برتری دارند [۸]، [۲۲]، [۲۳]، [۲۴]:

الف) روابط توپولوژیک غنی: زمانی که برچسب به الگوی ارتباط میان گروه‌ها، مانند اعضای مشترک یا شبکه‌های هماهنگ، وابسته باشد و این اطلاعات در ویژگی‌های محلی وجود نداشته باشد [۴]، [۲۴].

ب) ضعف ویژگی‌های ایستا: در گروه‌های تازه‌تأسیس یا گروه‌هایی که ویژگی‌های آماری آن‌ها به‌صورت عمدی عادی‌سازی شده است، ساختار همسایگی می‌تواند تنها منبع اطلاعات تفکیک‌کننده باشد.

ج) پویایی‌های زمانی: در صورت دسترسی به مسیرهای زمانی انتشار محتوا یا فعالیت‌های هماهنگ، GNN های فضایی-زمانی قادر به مدل‌سازی هم‌زمان وابستگی‌های ساختاری و زمانی هستند [۱۸]، [۲۰]، [۲۱]، [۴۷]، [۴۸].

د) گراف‌های ناهمگن و چندرابطه‌ای: این مدل‌ها می‌توانند اهمیت انواع مختلف روابط را به‌صورت خودکار بیاموزند، درحالی‌که مدل‌های جدولی به مهندسی دستی ویژگی‌ها وابسته‌اند [۱۷]، [۲۳].

ه) تعمیم استقرایی گراف‌سیج: به دلیل یادگیری تابع تجمیع، توانایی نمایش گره‌های دیده‌نشده را نیز دارد و برای سامانه‌های پویایی مانند تلگرام مناسب‌تر است [۸].

بنابراین، یافته این پژوهش به معنای برتری مطلق مدل‌های جدولی یا ناکارآمدی GNN ها نیست، بلکه نشان می‌دهد زمانی که برچسب مستقیماً از ویژگی‌های ایستای گره مشتق شود و گراف اطلاعات ساختاری محدودی داشته باشد، مدل‌های جدولی انتخاب مناسب‌تری هستند؛ در مقابل، در مسائل واقعی که برچسب از روابط ساختاری، هم‌عضویی و پویایی‌های زمانی ناشی می‌شود، GNN ها همچنان مزیت روش‌شناختی خود را حفظ می‌کنند.

۵-۲ محدودیت‌های روش‌شناختی و دامنه

اعتبارسنجی

تمامی آزمایش‌ها بر روی مجموعه داده‌ای شبیه‌سازی شده انجام شده‌اند. اگرچه توزیع‌های مصنوعی بر اساس ویژگی‌های گزارش شده برای شبکه‌های اجتماعی واقعی طراحی شده‌اند [۲۹]،

فراافکنی متمرکز شوند که وابستگی بیشتری به تعاملات ساختاری و پویایی‌های زمانی دارند، تا ارزش‌افزوده شبکه‌های عصبی گرافی در شناسایی رفتارهای هماهنگ پنهان و ناهنجار به‌وضوح برجسته‌تر گردد.

جدولی مانند XGBoost با صرف هزینه محاسباتی بسیار پایین‌تر توانستند نتایجی رقابتی و حتی اندکی بهتر ارائه دهند. اگرچه استفاده از داده‌های مصنوعی امکان کنترل دقیق، تحلیل مکانیزم‌ها و تکرارپذیری کامل را فراهم کرد، اما از محدودیت‌های اصلی این رویکرد عدم تأیید مستقیم با داده‌های پلتفرم‌های پیام‌رسان واقعی است. تحقیقات آینده باید بر توسعه برچسب‌ها و گراف‌های

References

- [1] D. Javed, N. Z. Jhanjhi, N. A. Khan, S. K. Ray, A. Al Mazroa, F. Ashfaq, and S. R. Das, "Towards the future of bot detection: A comprehensive taxonomical review and challenges on Twitter/X," *Computer Networks*, vol. 254, Art. no. 110808, 2024.
- [2] A. Beutel, W. Xu, V. Guruswami, C. Palow, and C. Faloutsos, "CopyCatch: Stopping group attacks by spotting lockstep behavior in social networks," in *Proc. 22nd Int. World Wide Web Conf. (WWW)*, Rio de Janeiro, Brazil, 2013, pp. 119–130.
- [3] A. Buitrago López et al., "Agent-based simulation of online social networks and disinformation," *arXiv preprint arXiv:2512.22082*, 2025.
- [4] Y. Dou, Z. Liu, L. Sun, Y. Deng, H. Peng, and P. S. Yu, "Enhancing graph neural network-based fraud detectors against camouflaged fraudsters," in *Proc. 29th ACM Int. Conf. on Information and Knowledge Management (CIKM)*, 2020, pp. 315–324.
- [5] A. Davies and N. Ajmeri, "Realistic synthetic social networks with graph neural networks," *arXiv preprint arXiv:2212.07843*, 2022.
- [6] N. Jiang, F. Yin, B. Wang, and A. T. Crooks, "A large-scale geographically explicit synthetic population with social networks for the United States," *Scientific Data*, vol. 11, Art. no. 1204, 2024, doi: 10.1038/s41597-024-03970-1.
- [7] N. Jiang, A. T. Crooks, H. Kavak, A. Burger, and W. G. Kennedy, "A method to create a synthetic population with social networks for geographically-explicit agent-based models," *Computational Urban Science*, vol. 2, Art. no. 7, 2022, doi: 10.1007/s43762-022-00034-1.
- [8] W. L. Hamilton, R. Ying, and J. Leskovec, "Inductive representation learning on large graphs," in *Advances in Neural Information Processing Systems (NeurIPS)*, 2017, pp. 1024–1034.
- [9] A. Bojchevski, O. Shchur, D. Zügner, and S. Günnemann, "NetGAN: Generating graphs via random walks," in *Proc. 35th Int. Conf. on Machine Learning (ICML)*, PMLR, vol. 80, 2018, pp. 610–619.
- [10] J. You, R. Ying, X. Ren, W. L. Hamilton, and J. Leskovec, "GraphRNN: Generating realistic graphs with deep auto-regressive models," in *Proc. 35th Int. Conf. on Machine Learning (ICML)*, PMLR, 2018, pp. 5708–5717.
- [11] A. Zareie et al., "Identifying coordination in online social networks through anomalous sharing behaviour," *Online Social Networks and Media*, vol. 50, Art. no. 100341, 2025.
- [12] R. Rogers and N. Righetti, "Coordinated inauthentic behaviour on Facebook? A typology of manufactured attention," *Platforms & Society*, vol. 2, Art. no. 29768624251369784, 2025.
- [13] M. Cinelli et al., "Coordinated inauthentic behavior and information spreading on Twitter," *Decision Support Systems*, vol. 160, Art. no. 113819, 2022.
- [14] L. Luceri et al., "Coordinated inauthentic behavior on TikTok: Challenges and opportunities for detection in a video-first ecosystem," in *Proc. Int. AAAI Conf. on Web and Social Media*, vol. 20, no. 1, 2026.
- [15] F. Cinus et al., "Exposing cross-platform coordinated inauthentic activity in the run-up to the 2024 US election," in *Proc. ACM Web Conf.*, 2025, 2025.
- [16] P. Veličković, G. Cucurull, A. Casanova, A. Romero, P. Liò, and Y. Bengio, "Graph attention networks," *arXiv preprint arXiv:1710.10903*, 2017.
- [17] Y. Liu et al., "Heterogeneous graph convolutional network for rumor detection with multi-level interactive fusion and graph reconstruction," *Scientific Reports*, vol. 15, no. 1, Art. no. 31639, 2025.
- [18] S. Shehnepoor et al., "Spatio-temporal graph representation learning for fraudster group detection," *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 35, no. 5, pp. 6628–6642, 2022.

- [19] M. Boyapati and R. Aygun, "BalancerGNN: Balancer graph neural networks for imbalanced datasets: A case study on fraud detection," *Neural Networks*, vol. 182, Art. no. 106926, 2025.
- [20] D. Saldaña-Ulloa, G. De Ita Luna, and J. R. Marcial-Romero, "A temporal graph network algorithm for detecting fraudulent transactions on online payment platforms," *Algorithms*, vol. 17, no. 12, Art. no. 552, 2024.
- [21] Z. Qu et al., "Temporal enhanced multimodal graph neural networks for fake news detection," *IEEE Trans. Comput. Soc. Syst.*, vol. 11, no. 6, pp. 7286–7298, 2024.
- [22] T. N. Kipf and M. Welling, "Semi-supervised classification with graph convolutional networks," in *Proc. Int. Conf. on Learning Representations (ICLR)*, 2017.
- [23] Z. Wu, S. Pan, F. Chen, G. Long, C. Zhang, and P. S. Yu, "A comprehensive survey on graph neural networks," *IEEE Transactions on Neural Networks and Learning Systems*, vol. 32, no. 1, pp. 4–24, 2021.
- [24] X. Ma, J. Wu, S. Xue, J. Yang, C. Zhou, Q. Z. Sheng, H. Xiong, and L. Akoglu, "A comprehensive survey on graph anomaly detection with deep learning," *IEEE Transactions on Knowledge and Data Engineering*, vol. 35, no. 12, pp. 12012–12038, 2023.
- [25] H. Qiao, H. Tong, B. An, I. King, C. Aggarwal, and G. Pang, "Deep graph anomaly detection: A survey and new perspectives," *arXiv preprint arXiv:2409.09957*, 2024.
- [26] B. Khemani, S. Patil, K. Kotecha, and S. Tanwar, "A review of graph neural networks: concepts, architectures, techniques, challenges, datasets, applications, and future directions," *Journal of Big Data*, vol. 11, Art. no. 18, 2024.
- [27] M. Fey and J. E. Lenssen, "Fast graph representation learning with PyTorch Geometric," in *ICLR Workshop on Representation Learning on Graphs and Manifolds*, 2019.
- [28] N. Patki, R. Wedge, and K. Veeramachaneni, "The synthetic data vault," in *Proc. IEEE Int. Conf. on Data Science and Advanced Analytics (DSAA)*, 2016, pp. 399–410.
- [29] P. W. Holland, K. B. Laskey, and S. Leinhardt, "Stochastic blockmodels: First steps," *Social Networks*, vol. 5, no. 2, pp. 109–137, 1983.
- [30] I. T. Jolliffe, *Principal Component Analysis*, 2nd ed. New York, NY: Springer, 2002.
- [31] W. Chen, K. Yang, Z. Yu, Y. Shi, and C. L. P. Chen, "A survey on imbalanced learning: latest research, applications and future directions," *Artificial Intelligence Review*, vol. 57, no. 6, Art. no. 137, 2024.
- [32] M. McPherson, L. Smith-Lovin, and J. M. Cook, "Birds of a feather: Homophily in social networks," *Annual Review of Sociology*, vol. 27, pp. 415–444, 2001.
- [33] G. E. P. Box, G. M. Jenkins, G. C. Reinsel, and G. M. Ljung, *Time Series Analysis: Forecasting and Control*, 5th ed. Hoboken, NJ: Wiley, 2015.
- [34] S. Sahu, "A Fast Parallel Approach for Neighborhood-based Link Prediction by Disregarding Large Hubs," *arXiv preprint arXiv:2401.11415*, 2024.
- [35] H. H. Rashidi, S. Albahra, B. Hu, and B. P. Rubin, "A novel and fully automated platform for synthetic tabular data generation and validation," *Scientific Reports*, vol. 14, 2024.
- [36] P. A. Apellániz, A. Jiménez, B. Arroyo Galende, J. Parras, and S. Zazo, "Synthetic Tabular Data Validation: A Divergence-Based Approach," *IEEE Access*, vol. 12, 2024.
- [37] N. Srivastava, G. Hinton, A. Krizhevsky, I. Sutskever, and R. Salakhutdinov, "Dropout: A simple way to prevent neural networks from overfitting," *Journal of Machine Learning Research*, vol. 15, pp. 1929–1958, 2014.
- [38] Q. Li, Z. Han, and X.-M. Wu, "Deeper insights into graph convolutional networks for semi-supervised learning," in *Proc. 32nd AAAI Conf. on Artificial Intelligence*, 2018, pp. 3538–3545.
- [39] K. Oono and T. Suzuki, "Graph neural networks exponentially lose expressive power for node classification," in *Proc. Int. Conf. on Learning Representations (ICLR)*, 2020.
- [40] L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [41] T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *Proc. 22nd ACM SIGKDD Int. Conf. on Knowledge Discovery and Data Mining (KDD)*, 2016, pp. 785–794.
- [42] D. P. Kingma and J. Ba, "Adam: A method for stochastic optimization," in *Proc. Int. Conf. on Learning Representations (ICLR)*, 2015.
- [43] T. Saito and M. Rehmsmeier, "The precision-recall plot is more informative than the ROC plot when evaluating binary classifiers on imbalanced datasets," *PLoS ONE*, vol. 10, no. 3, Art. no. e0118432, 2015.
- [44] J. Davis and M. Goadrich, "The relationship between precision-recall and ROC curves," in



- Proc. 23rd Int. Conf. on Machine Learning (ICML), 2006, pp. 233–240.
- [45] A. Foucart, A. Elskens, and C. Decaestecker, "Ranking the scores of algorithms with confidence," in Proc. ESANN 2025.
- [46] J. Demšar, "Statistical comparisons of classifiers over multiple data sets," *Journal of Machine Learning Research*, vol. 7, pp. 1–30, 2006.
- [47] Y. Zheng, L. Yi, and Z. Wei, "A survey of dynamic graph neural networks," *Frontiers of Computer Science*, vol. 19, no. 6, Art. no. 196323, 2025.
- [48] E. Rossi, B. Chamberlain, F. Frasca, D. Eynard, F. Monti, and M. Bronstein, "Temporal graph networks for deep learning on dynamic graphs," in *ICML Workshop on Graph Representation Learning*, 2020.

Detection of Hard-Law Groups in Telegram Messenger Using GraphSAGE Graph Neural Networks and Simulated Data with Temporal Features

Maedeh Arab Bafrani¹, Mohammad Ali Zare Chahooki^{2*}

¹Ph.D. Candidate, Department of Software Engineering, Faculty of Computer Engineering, Yazd University, Yazd, Iran

²Associate Professor, Department of Software Engineering, Faculty of Computer Engineering, Yazd University, Yazd, Iran

Article Information

Original Research Paper

Received:

2026 April 01

Accepted:

2026 August 18

Keywords:

Hard-law groups, Telegram messenger, Graph neural networks, GraphSAGE, Simulated dataset, Temporal features, Coercive behavior detection, Graph-based deep learning

Corresponding Author*:

chahooki@yazd.ac.ir

Abstract

This study addresses the detection of ‘Hard-law’ groups in the Telegram messaging platform, defined as groups that manipulate their organic growth by requiring users to perform mandatory actions, such as inviting additional users or joining specific channels, as a prerequisite for joining or continuing participation. Detecting such groups is essential for preserving platform integrity and mitigating abusive behaviors. The scarcity of labeled real-world data, the dynamic temporal nature of user behavior, and the need to jointly exploit structural and behavioral signals constitute the primary challenges in this domain. To address these challenges, a large-scale simulated dataset comprising 100,000 users, 10,000 groups, and 12 temporal snapshots was generated, capturing both static and temporal features of users and groups. Group labels were assigned using a principal component analysis (PCA)-based risk score, and a user-group bipartite graph was constructed. A GraphSAGE-based graph neural network was then trained using a fully inductive evaluation protocol, a weighted binary cross-entropy loss function, and decision threshold optimization, and its performance was compared with six baseline models. On an independent test set, the proposed model achieved an accuracy of 98.2%, a positive-class precision of 97.2%, a recall of 93.7%, an F1-score of 95.4%, an area under the ROC curve (AUC) of 99.87%, and an average precision (AP) of 99.51%, while demonstrating stable performance across all random seeds. The results indicate that, because the labels are structurally dependent on static group attributes, tabular machine learning models achieve slightly better performance with lower computational cost, whereas GraphSAGE remains robust on large-scale, high-degree graphs. These findings have been validated only on simulated data, and their generalization to real-world Telegram data remains an open challenge due to the lack of labeled datasets.

 : 10.22034/ABMIR.2026.24432.1231

E-ISSN: [2821-2037](#) /The Author 2026. Published by Yazd University This is an open access article under the CC BY 4.0 License <https://creativecommons.org/licenses/by/4.0/>.

