

دریافت مقاله: ۱۴۰۴/۰۵/۰۸

پذیرش مقاله: ۱۴۰۴/۰۸/۱۱

نوع مقاله: پژوهشی

ارائه یک روش ترکیبی برای داده‌افزایی و متوازن‌سازی مجموعه داده‌های نامتوازن

پرستو محقق

کارشناس ارشد مهندسی فناوری اطلاعات، دانشکده مهندسی برق و کامپیوتر، دانشگاه سیستان و بلوچستان، زاهدان، ایران

پست الکترونیکی: P.mohaghegh@pgs.usb.ac.ir

مهری رجائی*

استادیار، دانشکده مهندسی برق و کامپیوتر، دانشگاه سیستان و بلوچستان، زاهدان، ایران

پست الکترونیکی: rajayi@ece.usb.ac.ir

سمیرا نوفرستی

دانشیار، دانشکده مهندسی برق و کامپیوتر، دانشگاه سیستان و بلوچستان، زاهدان، ایران

پست الکترونیکی: snoferesti@ece.usb.ac.ir

چکیده

جنگل تصادفی، K نزدیک‌ترین همسایه و شبکه عصبی پرسپترون چندلایه و براساس معیارهای مختلف شامل فراخوانی، دقت، AUC و F1 عملکرد بهتری در مقایسه با دو روش پایه دارد. برای نمونه، روش HDAM بر اساس معیار F1 در مقایسه با H-SMOTE که عملکرد بهتری نسبت به LoRAS داشته است، کارایی طبقه‌بندی جنگل تصادفی، K نزدیک‌ترین همسایه و شبکه عصبی را به ترتیب ۱/۶۴، ۰/۰۵ و ۶/۱ درصد بهبود داده است. واژه‌های کلیدی: یادگیری ماشین، طبقه‌بندی، مجموعه داده‌های نامتوازن، داده‌افزایی، روش ترکیبی

۱- مقدمه

طبقه‌بندی یکی از تکنیک‌های یادگیری ماشین است که با آموزش بر روی یک مجموعه داده آموزشی از نمونه‌های برچسب‌خورده به ساخت یک مدل پیش‌بینی‌کننده می‌پردازد

یکی از چالش‌های مهم در یادگیری ماشین، داده‌های نامتوازن است. در مجموعه داده‌های نامتوازن، توزیع ناهمگون نمونه‌ها بین کلاس‌های مختلف، باعث می‌شود مدل‌های طبقه‌بندی عملکرد ضعیفی در شناسایی کلاس اقلیت (کلاس با تعداد نمونه کمتر) داشته باشند. این مشکل زمانی برجسته‌تر می‌شود که در یک کاربرد مشخص کلاس اقلیت دارای اهمیت بیشتری است. در این مقاله، یک روش ترکیبی برای حل مسئله داده‌های نامتوازن پیشنهاد می‌شود. این روش که HDAM نامیده می‌شود با هدف تولید نمونه‌های موثرتر برای کلاس اقلیت، دو تکنیک شناخته شده داده‌افزایی به نام‌های H-SMOTE و LoRAS را ترکیب می‌کند. ارزیابی‌های انجام شده نشان می‌دهند به‌صورت میانگین روش پیشنهادی برای هر سه طبقه‌بند

که قادر است برچسب کلاس نمونه‌های جدید را تعیین کند. یکی از مشکلاتی که الگوریتم‌های طبقه‌بندی با آن مواجه هستند، توزیع نامتوازن کلاس‌های مختلف در مجموعه داده آموزشی است. مشکل عدم توازن زمانی رخ می‌دهد که تعداد نمونه‌ها در یک کلاس (که کلاس اقلیت نامیده می‌شود) در مقایسه با سایر کلاس‌ها بسیار کمتر باشد، در حالی که در بسیاری از مسائل واقعی، پیش‌بینی صحیح برچسب نمونه‌های کلاس اقلیت دارای اهمیت بیشتری است. برای نمونه، در مسائلی نظیر شناسایی و طبقه‌بندی آفات کشاورزی، برخی از آفات بسیار نادرند یا در تشخیص تقلب در کارتهای اعتباری، تراکنش‌های مشکوک اندک هستند. در حالی که هدف اصلی این کاربردها شناسایی آفات نادر یا تراکنش‌های مشکوک است. به‌طور کلی مسئله داده‌های نامتوازن یکی از چالش‌های اصلی در یادگیری ماشین محسوب می‌شود، به‌ویژه در حوزه‌هایی مانند تشخیص نفوذ و طبقه‌بندی ناهنجاری‌ها که در آن رخداد‌های ناهنجار بسیار کمتر از رویدادهای عادی رخ می‌دهند.

الگوریتم‌های طبقه‌بندی سنتی دقت لازم را در پیش‌بینی نمونه‌های کلاس اقلیت ندارند، زیرا نمونه‌های آموزشی کافی برای کلاس اقلیت وجود ندارد و این کلاس به درستی آموزش داده نمی‌شود. برای حل مشکل داده‌های نامتوازن، رویکردهای متعددی در سه سطح داده، مدل و ترکیبی ارائه شده‌اند [۱]. در این مقاله بر رویکرد سطح داده متمرکز می‌شویم که در مقایسه با رویکرد سطح مدل رایج‌تر است. هدف رویکرد سطح داده، کاهش نسبت عدم تعادل بین کلاس‌های اکثریت و اقلیت و به تبع آن بهبود دقت طبقه‌بندی است [۱].

به‌طور کلی، تکنیک‌های سطح داده برای حل مسئله مجموعه داده‌های نامتوازن به دو دسته انتخاب نمونه و داده‌افزایی تقسیم می‌شوند. روش‌های انتخاب نمونه سعی در کاهش تعداد نمونه‌های کلاس اکثریت به‌منظور رسیدن به تعادل نسبی در اندازه کلاس‌ها دارد [۲-۴]. در مقابل، روش‌های داده‌افزایی، به‌صورت تصادفی یا در مناطقی

که اهمیت بیشتری برای کاربرد دارد نمونه‌های کلاس اقلیت را کپی‌برداری یا تولید می‌کنند تا زمانی که تعادل بین نمونه‌های کلاس اقلیت و اکثریت حاصل شود [۵-۱۰]. روش‌های انتخاب نمونه موجب حذف نمونه‌هایی می‌شوند که برای جمع‌آوری و برچسب‌گذاری آن‌ها وقت و هزینه بسیاری صرف شده است و گاهی نمونه‌های مهمی هستند که در آموزش طبقه‌بند نقش به‌سزایی دارند. از طرف دیگر، ممکن است مشکل انتخاب نمونه بیش از حد رخ می‌دهد که منجر به کاهش عملکرد طبقه‌بندی می‌شود. به همین دلیل، در پژوهش‌های اخیر، روش‌های داده‌افزایی بیشتر مورد توجه قرار گرفته‌اند.

در این مقاله، یک روش داده‌افزایی ترکیبی به نام HDAM^۱ پیشنهاد می‌شود که با ترکیب دو روش برجسته H-SMOTE [۹] و LoRAS [۱۰] به گونه‌ای که هر داده جدید به‌صورت تصادفی از یکی از این دو روش تولید می‌شود، بر گسترش نمونه‌های کلاس اقلیت می‌پردازد. هدف از ترکیب دو روش نامبرده، استفاده از مزایا و کاهش مشکلات هر کدام از این روش‌ها است. به‌طور مشخص، H-SMOTE [۹] برای تولید نمونه‌های جدید با توزیع پایدار به‌کار می‌رود، در حالی که LoRAS تنوع بیشتری را با ترکیب خطی نمونه‌های محلی ایجاد می‌کند. ادغام این دو رویکرد منجر به ایجاد تعادلی میان کیفیت و تنوع نمونه‌های تولیدی می‌شود.

برای ارزیابی روش پیشنهادی، آزمایش‌ها روی چند مجموعه داده مرجع از جمله مجموعه داده تشخیص تقلب در کارتهای اعتباری با ۲۸۴,۸۰۷ رکورد و عدم توازن شدید، مجموعه داده Pima Diabetes با ۷۶۸ نمونه و مجموعه داده نفوذ UNSW-NB15 با بیش از ۲/۵ میلیون رکورد، انجام شده است. نتایج نشان می‌دهد روش پیشنهادی نسبت به H-SMOTE که عملکرد بهتری در مقایسه با LoRAS داشته است، در معیار AUC به‌طور میانگین در طبقه‌بند‌های جنگل تصادفی، K نزدیک‌ترین همسایه و شبکه عصبی به‌ترتیب ۴/۹۵، ۰/۹۶ و ۶/۴ درصد بهبود داشته است. همچنین روش پیشنهادی، معیار F1 طبقه‌بند‌های مذکور را در

1-Hybrid Data Augmentation Method

پایه مانند الگوریتم ژنتیک بر روی کلیه خوشه‌ها عمل پالایش را انجام می‌دهد و یک مجموعه داده کاهش‌یافته ایجاد می‌شود. نمونه‌های نوفه مجموعه کاهش‌یافته کنار گذاشته می‌شود و مجموعه داده کاهش‌یافته نهایی از کلاس اکثریت با نمونه‌های کلاس اقلیت ترکیب شده و به‌عنوان مجموعه آموزشی جدید در نظر گرفته می‌شود.

در [۳] روشی به نام DCIS معرفی شده است که از تکنیک تقسیم و غلبه برای تقسیم مجموعه داده آموزشی اصلی به چندین زیرمجموعه، اعمال الگوریتم انتخاب نمونه بر روی هر یک از زیرمجموعه‌ها و سپس ترکیب زیرمجموعه‌های کاهش‌یافته جهت ساخت مجموعه داده متوازن استفاده می‌کند. برای ترکیب زیرمجموعه‌های کاهش‌یافته دو تکنیک اجتماع و اشتراک پیشنهاد شده است. نتایج آزمایش‌های انجام گرفته نشان می‌دهد اجتماع زیرمجموعه‌ها به نتایج بهتری دست یافته است.

در [۴] روشی مبتنی بر یادگیری تقویتی برای انتخاب نمونه پیشنهاد شده است که ابتدا خوشه‌هایی از نمونه‌ها را بر اساس یک معیار شباهت می‌سازد. سپس عاملی در نظر می‌گیرد که می‌تواند به‌صورت تدریجی داده‌ها را به مجموعه کاهش‌یافته نهایی اضافه کند. در یک حلقه، این عامل خوشه‌ای از نمونه‌ها (عمل) را برای اضافه شدن به نمونه‌های برگزیده (حالت) انتخاب می‌کند و پاداشی متناسب با کاهش خطای مدل پیش‌بینی‌کننده دریافت می‌کند. خروجی الگوریتم (سیاست عامل) ماتریسی است که توازن بین بهبود مدل و اندازه داده‌ها را نشان می‌دهد. هر درایه از ماتریس ارزش افزودن یک خوشه به نمونه‌های موجود را نشان می‌دهد. این ماتریس توانایی متعادل کردن اهداف کاهش نوفه و حجم داده را دارد.

۲-۲- تکنیک‌های داده‌افزایی

در روش‌های داده‌افزایی ایده اصلی، گسترش تعداد نمونه‌های کلاس اقلیت برای متوازن‌سازی توزیع نمونه‌ها بین کلاس‌ها می‌باشد. روش SMOTE در بین روش‌های

مقایسه با H-SOMTE به میزان ۱/۶۴، ۰/۰۵ و ۶/۱ درصد افزایش داده است. این نتایج بیانگر کارایی روش پیشنهادی در حوزه‌های مختلف است.

به‌طورکلی، نوآوری‌های اصلی مقاله به‌صورت زیر ارائه می‌شود:

- برای نخستین بار ترکیب [۹] H-SMOTE و [۱۰] LoRAS به‌منظور مقابله با عدم توازن داده‌ها با هدف ایجاد تعادلی میان کیفیت و تنوع نمونه‌های تولید شده پیشنهاد شده است.
- ارزیابی جامعی روی چندین مجموعه داده از حوزه‌های مختلف (سلامت، مالی و امنیت سایبری) انجام گرفته است.
- نشان داده‌ایم که روش پیشنهادی بر اساس معیارهای AUC، دقت، فراخوانی و F1 در اغلب موارد عملکرد بهتری نسبت به روش‌های متداول دارد، بجز در برخی طبقه‌بندی‌های مبتنی بر فاصله که محدودیت‌های آن نیز به‌صراحت بحث شده است.

ادامه مقاله به‌صورت زیر سازمان‌دهی شده است: در بخش ۲ تحقیقات پیشین در زمینه طبقه‌بندی نامتوازن مرور می‌شوند. در بخش ۳ جزئیات روش پیشنهادی شرح داده می‌شود. در بخش ۴ ارزیابی کارایی روش پیشنهادی ارائه می‌گردد. در پایان، بخش ۵ به نتیجه‌گیری می‌پردازد.

۲- کارهای مرتبط

همان‌طور که پیشتر گفته شد، روش‌های موجود برای متعادل‌سازی مجموعه داده‌های نامتوازن به‌طورکلی به دو دسته انتخاب نمونه و داده‌افزایی تقسیم می‌شوند. در ادامه به معرفی کارهای مرتبط در هر دسته پرداخته می‌شود.

۲-۱- تکنیک‌های انتخاب نمونه

Tsai و همکاران [۲]، روشی با نام CBIS پیشنهاد داده‌اند که در گام اول نمونه‌های کلاس اکثریت را خوشه‌بندی می‌کند و هر خوشه به‌عنوان زیرکلاسی از کلاس اکثریت برچسب‌گذاری می‌شود. در گام بعد، الگوریتم انتخاب نمونه

این است که رابطه بین کلاس اکثریت و اقلیت اهمیت دارد. Bunkhumpornpat و همکاران [۸]، در روشی به نام SL-SMOTE، با محاسبه سطح ایمنی برای هر نمونه اقلیت که براساس فراوانی نمونه‌های اکثریت در K نزدیک‌ترین همسایگی نمونه اقلیت انجام می‌شود، داده‌افزایی را انجام داده‌اند.

Chao و Zhang [۹]، روش H-SMOTE را با هدف تولید نمونه‌های مصنوعی یکنواخت‌تر و نزدیک‌تر به مرکز کلاس اقلیت، معرفی کردند. این روش از دو عمل نمونه‌برداری و پالایش نمونه‌های دارای نوفه تشکیل شده است. ابتدا نمونه مرکزی کلاس اقلیت با میانگین‌گیری از ویژگی‌ها مشخص می‌شود و سپس فاصله منتهن هر نمونه از کلاس اقلیت با مرکز محاسبه می‌شود و نمونه جدید در سطح اول تولید می‌گردد. این عمل برای نمونه‌های جدید و نمونه‌های کلاس اقلیت در سطح دوم نیز تکرار می‌شود.

Bej و همکاران [۱۰]، یک روش داده‌افزایی با نام LoRAS را معرفی کردند. در این روش برای هر نمونه P از کلاس اقلیت، ناحیه‌ای چندضلعی در نظر گرفته شده و سپس K نزدیک‌ترین همسایه‌های نمونه P در آن ناحیه انتخاب می‌شوند. برای هر نمونه که در همسایگی نمونه P قرار دارد، به اندازه معینی، نمونه سایه ایجاد می‌شود. نمونه‌های سایه با ایجاد نوفه در محدوده تابع توزیع نرمال با انحراف معیار استاندارد ایجاد می‌شوند. سپس به‌صورت تصادفی به تعداد ویژگی‌ها از لیست نمونه‌های سایه، نمونه انتخاب می‌شود و در وزن‌های بردار آفین ضرب می‌شوند و خروجی بردار به‌عنوان نمونه جدید در نظر گرفته می‌شود. Temraz و Keane [۱۱]، روشی مبتنی بر جفت‌های متضاد برای داده‌افزایی پیشنهاد دادند. این روش که CFA نام دارد، ابتدا هر دو نمونه از کلاس اکثریت (نمونه x و اقلیت (نمونه p) را که حداکثر در دو ویژگی تفاوت دارند، به‌صورت جفت در مجموعه داده آموزشی در نظر می‌گیرد که این جفت‌ها را جفت‌های متضاد می‌نامند. سپس برای هر نمونه از کلاس اکثریت که با هیچ نمونه‌ای از کلاس اقلیت

داده‌افزایی از محبوبیت بیشتری برخوردار است [۵] زیرا نسخه‌های توسعه‌یافته زیادی برای آن ارائه شده است. ایده اصلی SMOTE پایه این است که برای هر نمونه کلاس اقلیت با توجه به میزان شباهتی که با دیگر نمونه‌های کلاس اقلیت دارد، K نزدیک‌ترین همسایگانش تعیین می‌شوند. سپس بین هر دو نمونه از همسایه‌ها، نمونه‌های جدید به‌صورت تصادفی تولید می‌شوند. مشکلی اصلی SMOTE این است که به دلیل محدود بودن ناحیه تولید نمونه‌ها نمی‌تواند به‌خوبی بر چالش توزیع نامتوازن داده‌ها غلبه کند. برای حل این مشکل، انواع مختلفی از SMOTE پیشنهاد شده است که به سه دسته کلی روش‌های مبتنی بر شناسایی مناطق امن در طبقه اقلیت، روش‌های متکی بر اهمیت ناحیه مرزی بین کلاس اقلیت و اکثریت و روش‌های مبتنی بر رابطه بین کلاس اقلیت و اکثریت تقسیم می‌شوند. در ادامه به معرفی برخی از روش‌های پیشنهاد شده برای هر دسته می‌پردازیم.

روش پایه SMOTE به این اهم رسیده است که همه مناطق در کلاس اقلیت برابر نیستند و برخی مناطق ممکن است مهم‌تر یا ایمن‌تر از سایر مناطق باشد. Douzas و Bacao [۶]، روشی به نام G-SMOTE پیشنهاد داده‌اند که در آن با تعیین یک شعاع امن در اطراف هر نمونه از کلاس اقلیت یک منطقه هندسی به نام ابر کروی برای تولید نمونه‌های مصنوعی جدید ایجاد می‌شود.

روش‌هایی که بر اهمیت ناحیه مرزی بین کلاس اقلیت و اکثریت تمرکز دارند، مبتنی بر این دیدگاه هستند که نمونه‌های کلاس اقلیت واقع در نزدیکی مرز تصمیم‌گیری طبقه‌بند از اهمیت بیشتری برخوردار هستند، بنابراین تولید نمونه‌های مصنوعی در این ناحیه مرزی می‌تواند به بهبود عملکرد طبقه‌بند کمک کند. Saez و همکاران [۷]، روش SMOTE-IPF را ارائه دادند که از یک پالایه بخش‌بندی، برای اصلاح نمونه‌های مرزی و دارای نوفه تولید شده توسط SMOTE استفاده می‌کند.

سومین بینش در راستای بهبود عملکرد SMOTE

جفت نشده است (نمونه x')، با استفاده از فاصله اقلیدسی نزدیک‌ترین همسایه آن از جفت‌های متضاد، شناسایی شده و یک نمونه مصنوعی (نمونه p') با استفاده از ویژگی‌های این جفت متضاد تولید می‌شود. نمونه p' ویژگی‌های مشابه را از p و ویژگی‌های متفاوت خود را از x' می‌گیرد. اگرچه CFA عملکرد خوبی داشته است اما همچنان با محدودیت‌هایی همراه است. از جمله این محدودیت‌ها نیاز به شرکت حداقل ۵ درصد از نمونه‌های کلاس اکثریت در جفت‌های متضاد است. بنابراین این روش بدون جفت‌های متضاد کافی در تولید نمونه‌های اقلیت مصنوعی به شدت با مشکل روبه‌رو خواهد شد.

۳- روش پیشنهادی

در روش پیشنهادی با نام HDAM، برای متوازن‌سازی مجموعه داده، از ترکیب دو الگوریتم H-SMOTE و LoRAS برای تولید نمونه‌های مصنوعی در کلاس اقلیت استفاده می‌شود. دلیل انتخاب دو الگوریتم H-SMOTE و LoRAS، در انتهای بخش ۲ عنوان شده است.

در روش پیشنهادی تا زمانی که تعداد نمونه‌های کلاس اقلیت (C_{min}) برابر با تعداد نمونه‌های کلاس اکثریت (C_{maj}) نشده است، به صورت تصادفی یکی از دو الگوریتم H-SMOTE و LoRAS انتخاب می‌شود و با الهام گرفتن از ایده آن الگوریتم یک نمونه اقلیت جدید ایجاد می‌گردد. پس از بررسی تکراری نبودن، این نمونه به کلاس اقلیت گسترش یافته که با C'_{min} نشان داده می‌شود و در ابتدا شامل نمونه‌های C_{min} است، اضافه می‌شود. در پایان، کلاس اقلیت گسترش یافته تعداد نمونه‌هایی برابر با کلاس اکثریت خواهد داشت.

شکل (۱) شبه کد روش پیشنهادی را نشان می‌دهد. مطابق شکل (۱)، نمونه‌های دو کلاس اقلیت و اکثریت به‌عنوان ورودی دریافت و پارامترهای الگوریتم LoRAS مقداردهی می‌شوند. سپس تا زمانی که تعداد نمونه‌های کلاس اقلیت برابر با تعداد نمونه‌های کلاس اکثریت نشده است (خط ۱)، به تصادف یکی از دو الگوریتم H-SMOTE و LoRAS انتخاب می‌شود (خط ۲) و یک نمونه اقلیت بر اساس نمونه‌های موجود در C_{min} تولید شده و پس از

در سال‌های اخیر، به‌کارگیری مدل‌های زبانی بزرگ برای داده‌افزایی مورد توجه محققان قرار گرفته است. در این رویکرد، با طراحی پرسش‌های مناسب، از مدل زبانی بزرگ خواسته می‌شود نمونه‌های مصنوعی مرتبط با داده‌های اولیه تولید کند. این نمونه‌ها به همراه داده‌های اصلی برای افزایش حجم و تنوع داده‌ها در فرآیند آموزش مدل‌های یادگیری ماشین به‌کار می‌روند [۱۲، ۱۳].

در جدول (۱)، خلاصه روش‌های داده‌افزایی معرفی شده به همراه نقاط ضعف و قوت هر کدام و کارایی آنها گزارش داده شده است. علی‌رغم تحقیقات گسترده، مجموعه داده‌های نامتوازن همچنان یک مسئله چالش‌برانگیز در حوزه یادگیری ماشین به‌شمار می‌آید. دلیل این امر آن است که بسیاری از روش‌های موجود یا صرفاً برای یک کاربرد خاص طراحی شده‌اند، یا از مقیاس‌پذیری کافی برخوردار نیستند، یا توانایی تولید نمونه‌های کافی برای رفع عدم توازن شدید را ندارند و یا در شرایط خاص، پیچیدگی محاسباتی بالایی ایجاد می‌کنند. روش H-SMOTE با وجود تولید نمونه‌های جدید با توزیع پایدار و معقول و عملکرد بهتر نسبت به الگوریتم‌های پایه مانند SMOTE، همچنان در سطوح بالای عدم‌توازن با مشکل کمبود نمونه‌های مصنوعی مطابق با الگوی واقعی داده‌ها مواجه است. از سوی دیگر، روش LoRAS می‌تواند مجموعه داده‌های با نرخ

جدول ۱. مقایسه روش‌های پیشین داده‌افزایی

مرجع	ایده اصلی	مزایا	معایب	نتایج	
				روش ارزیابی	طبقه‌بند
SMOTE [۵]	درون‌یابی خطی بین نمونه‌های کلاس اقلیت	پایاده‌سازی ساده، تنوع داده‌های تولید شده	تولید داده‌های نامعتبر در صورت وجود نمونه‌های نویز یا مرزی در همسایگی، تولید نمونه‌های تکراری	میانگین نتایج حاصل از روش اعتبارسنجی متقابل با ۵ حلقه بر روی چند مجموعه داده مختلف از مخزن UCI	DT
					KNN
G-SMOTE [۶]	درون‌یابی خطی بین نمونه‌های موجود در مناطق امن تعیین شده برای کلاس اقلیت	کاهش نمونه‌های همپوشان با کلاس اکثریت، تنوع بیشتر نمونه‌های تولید شده در مقایسه با SMOTE	محاسبات پیچیده، نیاز به تنظیم دقیق پارامترهای متعدد، تولید نمونه‌های تکراری	میانگین نتایج حاصل از روش اعتبارسنجی متقابل با ۵ حلقه بر روی چند مجموعه داده دودویی از مخازن UCI و KEEL	DT
					KNN
SMOTE-IPF [۷]	به‌کارگیری فیلتر پارتیشن‌بندی برای اصلاح نمونه‌های مرزی و دارای نویز	کاهش نمونه‌های نوفه و ایجاد مرز تصمیم یکنواخت‌تر نسبت به SMOTE، عملکرد بهتر روی مجموعه داده‌های نویزدار	هزینه محاسباتی و زمان اجرای بالا برای مجموعه‌داده‌های بزرگ، پایاده‌سازی پیچیده، حذف نمونه‌های مرزی مفید توسط پالایه‌ها در برخی مجموعه‌ها	میانگین نتایج حاصل از روش اعتبارسنجی متقابل با ۵ حلقه بر روی چند مجموعه داده واقعی	C۴,۵
SL-SMOTE [۸]	تولید نمونه‌های اقلیت بر اساس فراوانی نمونه‌های اکثریت در همسایگی آن	کاهش تولید نمونه‌های نامعتبر در مرز تصمیم	نیاز به برآورد دقیق سطح امن، پیچیدگی بالاتر نسبت به SMOTE	میانگین نتایج حاصل از روش اعتبارسنجی متقابل با ۱۰ حلقه بر روی مجموعه داده Satimage	C۴,۵
H-SMOTE [۹]	نمونه‌برداری و فیلتر نمونه‌های دارای نوفه	توزیع یکنواخت‌تر و پایدارتر نمونه‌ها نسبت به SMOTE	پیچیدگی محاسباتی بیشتر در مقایسه با SMOTE پایه، عدم تطابق نمونه‌های تولید شده با الگوی واقعی داده‌ها در مجموعه داده‌های بسیار نامتوازن	میانگین نتایج بر روی ۱۲ مجموعه داده مخزن UCI که با روش تقسیم تصادفی به نسبت ۳۰ به ۷۰ به دو مجموعه آزمون و آموزش تقسیم شده‌اند	SVM
LoRAS [۱۰]	تولید نمونه‌های اقلیت با میانگین‌گیری از زیرمجموعه‌های همسایه به جای خط مستقیم	عملکرد بالا در مجموعه‌داده‌های بسیار نامتوازن	احتمال تولید نمونه‌های غیرقابل اعتماد، پیچیدگی محاسباتی بالاتر نسبت به H-SMOTE، نیاز به تنظیم پارامترهای متعدد	میانگین نتایج حاصل از روش اعتبارسنجی متقابل با ۳ و ۱۰ حلقه بر روی ۱۴ مجموعه داده از مخازن UCI و KEEL	KNN, LR, RF
CFA [۱۱]	تولید نمونه‌های اقلیت با نمونه‌های کلاس مقابل به جای درون‌یابی	قابل توضیح، سازگار با قیود دامنه	نیاز به تعریف سازوکار معتبر برای جفت‌های متضاد در هر کاربرد	میانگین نتایج حاصل از روش اعتبارسنجی متقابل با ۵ حلقه بر روی ۲۵ مجموعه داده از مخازن UCI و KEEL	KNN
					MLP
IoT-IDS [۱۲]	ترکیب داده‌سازی مبتنی بر مدل زبانی (توصیف/تولید) ردیف‌های عددی با روش‌های درون‌یابی برای طبقه‌بندی حملات	بهبود پوشش الگوهای نادر، مناسب برای طبقه‌بندی چندکلاسه	خاص کاربرد، ریسک عدم تطابق با توزیع واقعی داده‌ها	میانگین نتایج بر روی ۸ مجموعه داده نفوذ در IIoT که با روش تقسیم تصادفی به نسبت ۲۰ به ۸۰ به دو مجموعه آزمون و آموزش تقسیم شده‌اند	RF
					MLP

HDAM Algorithm	
Inputs:	
C_{min} :	Minority class data points
C_{maj} :	Majority class data points
k :	Number of nearest neighbors to be considered per parent data point (default value : 30 if $ C_{min} > 100$, 5 otherwise)
$ S_p $:	Number of generated shadow samples per parent data point (default value : $\max\left(\left\lceil \frac{2 F }{k} \right\rceil, 40\right)$)
N_{aff} :	Number of shadow points to be chosen for a random affine combination (default value : $ F $)
Outputs:	
C'_{min} :	Minority data points and non_duplicate generated samples
1	WHILE $ C'_{min} < C_{maj} $ do
2	Method $\leftarrow$ Randomly select an algorithm (H-SMOTE or LoRAS)
3	IF Method = LoRAS then
4	Execute LoRAS Algorithm to generate a new minority sample (p)
5	ELSE
6	Execute H-SMOTE Algorithm to generate a new minority sample (p)
7	IF p is not a duplicate sample then
8	Add p to C'_{min}
9	End WHILE
10	Return C'_{min}

شکل ۱. شبه کد روش پیشنهادی HDAM

به صورت $\alpha_1 + \alpha_2 + \alpha_3 = 1$ تعریف می‌شود و نمونه جدید از طریق رابطه (۲) به دست می‌آید [۱۰]:

(۲)

$$L = \alpha_1 s_1 + \alpha_2 s_2 + \alpha_3 s_3$$

اگر الگوریتم انتخابی H-SMOTE باشد (خط ۵ شبه کد)، به تصادف، تولید نمونه سطح اول یا دوم انتخاب می‌شود. اگر تولید نمونه سطح اول انتخاب شود، برای نمونه P_i از کلاس اقلیت که به تصادف انتخاب شده است، مطابق رابطه (۳)، نمونه جدید x_i میان نمونه مرکزی Z (که با میانگین گرفتن از نمونه‌های کلاس اقلیت محاسبه می‌شود) و نمونه P_i تولید می‌شود و در کلاس اقلیت قرار می‌گیرد [۹].

$$x_i = p_i + 0.5 \times (Z - p_i) \quad (3)$$

اگر تولید نمونه سطح دوم انتخاب شود، دو نمونه سطح اول x_i و x_j مطابق رابطه (۳) تولید می‌شوند. سپس مطابق رابطه (۴) نمونه جدید x_{ii} در سطح دوم ایجاد می‌شود [۹].

$$x_{ii} = x_i + 0.5 \times (x_j - x_i) \quad (4)$$

بررسی تکراری نبودن، این نمونه به مجموعه C'_{min} اضافه می‌گردد (خطوط ۳ تا ۸).

به صورت دقیق‌تر، اگر الگوریتم LoRAS انتخاب شود (خط ۳)، یک نمونه P_i از کلاس اقلیت C_{min} انتخاب و ناحیه‌ای مکعب مستطیل شکل اطراف آن در نظر گرفته می‌شود. سپس k نزدیک‌ترین همسایه‌های نمونه P_i تعیین می‌شوند. برای هر نمونه q والد که در همسایگی نمونه P_i قرار دارد، مطابق رابطه (۱) به اندازه $|S_p|$ نمونه سایه ایجاد می‌شود [۱۰]:

(۱)

$$|S_p| = \left(\max \left(\left\lceil \frac{2|F|}{k} \right\rceil, 40 \right) \right)$$

که F تعداد ویژگی‌ها است. نمونه‌های سایه به وسیله تولید نوفه با کمک تابع توزیع نرمال با انحراف معیار استاندارد L_σ ایجاد می‌شوند و مقدار پیش‌فرض آن $[0.005, \dots, 0.005]$ می‌باشد. سپس به صورت تصادفی به اندازه N_{aff} که به تعداد ویژگی‌ها است از لیست نمونه‌های سایه، نمونه انتخاب می‌شود و مطابق رابطه (۲) در وزن‌های بردار آفین ضرب می‌شوند. خروجی بردار به عنوان نمونه جدید به مجموعه نمونه‌های مصنوعی افزوده می‌شود. یک بردار آفین با وزن‌های تصادفی

جدول ۲. مشخصات مجموعه داده‌های مورد استفاده

مجموعه داده	شناسه	تعداد ویژگی‌ها	تعداد کلاس‌ها	تعداد نمونه‌ها	تعداد نمونه‌های اقلیت	تعداد نمونه‌های اکثریت	IR	OVO/OVR
Vehicle	D1	۱۹	۴	۸۴۶	۱۹۹	۶۴۷	۳/۲۵	OVR
Yeast3	D2	۹	۱۰	۱۴۸۴	۱۶۳	۱۳۲۱	۸/۱	OVR
Page-blocks0	D3	۱۱	۵	۵۴۷۲	۵۵۹	۴۹۱۳	۸/۷۸	OVR
Ecoli3	D4	۸	۸	۳۳۶	۳۵	۳۰۱	۸/۶	OVR
Yeast4	D5	۹	۱۰	۱۴۸۴	۵۱	۱۴۳۳	۲۸/۰۹	OVR
Yeast6	D6	۹	۱۰	۱۴۸۴	۳۵	۱۴۴۹	۴۱/۴	OVR
Glass3	D7	۱۰	۷	۲۱۴	۱۷	۱۹۷	۱۱/۵۸	OVR
Abalone9-vs-13	D8	۹	۲	۸۹۲	۲۰۳	۶۸۹	۳/۳۹	OVO
Abalone13	D9	۹	۲۸	۴۱۷۷	۲۰۳	۳۹۷۴	۱۹/۵۷	OVR
WineQuality-White3-vs-7	D10	۱۲	۲	۹۰۰	۲۰	۸۸۰	۴۴	OVO

نمونه‌های کلاس اکثریت را نشان می‌دهد. هر چه نسبت عدم توازن (IR) به ۱ نزدیک‌تر باشد آن مجموعه داده متعادل‌تر است.

جهت ارزیابی دقت الگوریتم‌های طبقه‌بندی نیاز به یک مجموعه آزمون است که نمونه‌های آن برای آموزش طبقه‌بند مورد استفاده قرار نگرفته باشند، برای این منظور هر یک از مجموعه داده‌های جدول (۲) با روش اعتبارسنجی متقابل با k حلقه با نمونه‌گیری متوازن^۲ به دو بخش آموزش و آزمون تقسیم شده‌اند. در این مقاله مقدار k برابر ۵ در نظر گرفته شده است.

جهت پیاده‌سازی روش پیشنهادی از زبان برنامه‌نویسی پایتون استفاده شده است. برای ارزیابی روش پیشنهادی، خروجی این روش به‌عنوان مجموعه آموزش الگوریتم‌های طبقه‌بندی K نزدیک‌ترین همسایه (KNN)، جنگل تصادفی (RF) و شبکه عصبی پرسپترون چند لایه (MLP) به‌کار گرفته شده است. مقادیر فرایارامترهای طبقه‌بندهای مذکور با روش جستجوی شبکه‌ای^۳ تنظیم شده است. در جدول (۳) مقادیر بررسی شده برای فرایارامترهای الگوریتم‌های طبقه‌بندی مختلف و مقدار انتخاب شده (مقادیر پررنگ

این روال تا زمان برآورده شدن شرط توقف (برابری اندازه کلاس اقلیت و اکثریت) ادامه خواهد داشت. کلاس اقلیت نهایی تعداد نمونه‌های برابر با تعداد نمونه‌های کلاس اکثریت دارد و بنابراین مشکل عدم توازن مجموعه داده‌ای حل می‌شود. در پایان، مجموعه آموزش متوازن جدید که متشکل از نمونه‌های دو کلاس C'_{min} و C_{maj} است برای ساخت مدل پیش‌بینی کننده استفاده می‌شود.

۴ - نتایج

در ادامه ابتدا مشخصات مجموعه داده‌های مورد استفاده و نیز ابزار و تنظیمات شبیه‌سازی معرفی می‌شود. سپس نتایج ارزیابی کارایی روش پیشنهادی ارائه می‌گردد

۴-۱- مجموعه داده‌ها، ابزار و تنظیمات شبیه‌سازی

جدول (۲) ویژگی‌های مجموعه داده‌های نامتوازن انتخاب شده از دو مخزن [۱۴] UCI و [۱۵] KEEL را نشان می‌دهد. روش‌های موجود داده‌افزایی عمدتاً روی مسائل طبقه‌بندی دودویی تمرکز دارند. بنابراین مجموعه داده‌های چند کلاسه بایستی به مجموعه‌هایی با کلاس‌های دودویی تبدیل شوند. برای انجام این تبدیل از دو روش یک به یک (OVO) [۱۶] و یک به چند (OVR) [۱۷] استفاده می‌شود. روش OVO یک مجموعه داده چندکلاسه را به یک مجموعه داده دودویی براساس هر جفت کلاس تقسیم می‌کند. در مقابل، رویکرد OVR یکی از کلاس‌های چندگانه را انتخاب می‌کند و آن را در برابر همه کلاس‌های دیگر قرار می‌دهد، به‌طوری که یکی از کلاس‌ها به‌عنوان کلاس مثبت (اقلیت) و سایر کلاس‌ها به‌عنوان کلاس منفی (اکثریت) تلقی می‌شود و بدین‌ترتیب نسبت عدم توازن کلاس‌ها تغییر می‌کند. نسبت عدم توازن (IR) مطابق رابطه (۵) محاسبه می‌گردد:

$$IR = \frac{N_{maj}}{N_{min}} \quad (5)$$

که N_{min} تعداد نمونه‌های کلاس اقلیت و N_{maj} تعداد

2-Stratified k-fold cross validation
3-Grid Search

جدول ۴. مقادیر پارامترهای روش LoRAS [۱۰]

N_{off}	K	تعداد نمونه‌های اقلیت	مجموعه داده
۱۸	۳۰	۱۹۹	Vehicle
۸	۳۰	۱۶۳	Yeast3
۱۰	۳۰	۵۵۹	Page-blocks0
۷	۵	۳۵	Ecoli3
۸	۵	۵۱	Yeast4
۸	۵	۳۵	Yeast6
۹	۵	۱۷	Glass3
۸	۳۰	۲۰۳	Abalone9-vs-13
۸	۳۰	۲۰۳	Abalone13
۱۱	۵	۲۰	WineQuality-White3-vs-7

جدول ۵. مقایسه کارایی روش پیشنهادی بر اساس معیار AUC برای طبقه‌بند RF

مجموعه داده	بدون داده‌افزایی	LoRAS	H-SMOTE	HDAM
D1	۰/۹۳۷۴	۰/۹۰۲۶	۰/۹۵۳۴	۰/۹۷۷۳
D2	۰/۷۸۳۳	۰/۷۹۷۶	۰/۸۲۲۹	۰/۸۶۰۳
D3	۰/۹۳۰۳	۰/۸۷۳۱	۰/۹۲۳۵	۰/۹۴۳۱
D4	۰/۶۸۹۰	۰/۵۹۹۲	۰/۷۰۶۹	۰/۷۸۳۳
D5	۰/۵۲۷۰	۰/۵۰۰۰	۰/۵۵۷۳	۰/۶۱۶۱
D6	۰/۵۸۳۳	۰/۶۱۲۱	۰/۷۲۵۷	۰/۷۶۹۳
D7	۰/۵۰۰۰	۰/۴۹۷۵	۰/۵۰۰۰	۰/۶۱۶۶
D8	۰/۶۷۶۴	۰/۶۱۶۳	۰/۶۵۴۶	۰/۷۴۲۵
D9	۰/۴۹۹۵	۰/۴۹۸۴	۰/۴۹۷۸	۰/۵۰۲۲
D10	۰/۵۰۰۰	۰/۴۹۷۱	۰/۴۹۸۲	۰/۵۲۴۴
میانگین	۰/۶۶۲۶	۰/۶۳۹۳	۰/۶۸۴۰	۰/۷۳۳۵

جداول (۵ تا ۷) مقادیر معیار AUC را برای سه طبقه‌بند KNN، RF و MLP نشان می‌دهند، زمانی که این طبقه‌بندها روی مجموعه داده پایه (بدون داده‌افزایی) و مجموعه داده‌های تولید شده توسط روش‌های H-SMOTE، HDAM و LoRAS آموزش دیده‌اند. در این جداول، بهترین عملکرد به صورت پررنگ و رتبه دوم به صورت زیرخطدار نشان داده شده است.

جدول (۵) نشان می‌دهد که برای طبقه‌بند RF، روش HDAM در تمامی مجموعه داده‌ها بهتر از سایر روش‌های موجود عمل کرده است. به طور میانگین، روش پیشنهادی HDAM با امتیاز ۰/۷۳۳۵ از AUC بالاتری برای طبقه‌بند RF برخوردار است.

مطابق جدول (۶)، برای طبقه‌بند KNN، روش پیشنهادی

جدول ۳. مقادیر مورد بررسی برای فرآیندهای طبقه‌بندهای MLP و KNN، RF

طبقه‌بند	مقادیر فرآیندهای
جنگل تصادفی (RF)	ntree = {50, 100, 200, 400, 600} max_depth = {None, 4, 6, 10, 20, 30, 50, 80, 100}
-K نزدیک‌ترین همسایه (KNN)	n_neighbors = {3, 5, 7, 10, 20, 30, 50} metric = {'Minkowski', 'Manhattan', 'Euclidean'}
پرسپترون چند لایه (MLP)	hidden_layer_sizes = (100, 50) activation = {'identity', 'logistic', 'tanh', 'relu'} solver = {'sgd', 'lbfgs', 'adam'} alpha = 0/0001

در جدول) نشان داده شده است. مقادیر فرآیندها بر اساس مقادیر پیش‌فرض برای هر طبقه‌بند در سایت scikit-learn.org و مقادیر رایج در تحقیقات مرتبط، انتخاب شده‌اند [۱۱].

مقادیر پارامترهای روش LoRAS برای هر مجموعه داده مطابق جدول (۴) و بر اساس مرجع [۱۰] انتخاب شده است. در روش LoRAS، به k در صورتی که $|C_{min}| \geq 100$ باشد، به صورت پیش‌فرض مقدار ۳۰ داده می‌شود و در غیر این صورت مقدار ۵ می‌گیرد. اندازه N_{off} برابر تعداد ویژگی‌ها است [۱۱].

۴-۲- ارزیابی روش پیشنهادی

برای ارزیابی کارایی روش پیشنهادی از معیارهای رایج دقت، فراخوانی، F1 و AUC استفاده شده است. همچنین کارایی روش پیشنهادی با دو روش H-SMOTE و LoRAS مقایسه شده است. تمامی روش‌های مذکور بر روی ۱۰ مجموعه داده که مشخصات آنها در جدول (۲) آورده شده است و با به کارگیری ۳ طبقه‌بند یادگیری ماشین جنگل تصادفی K (RF)، نزدیک‌ترین همسایه (KNN) و شبکه عصبی پرسپترون چند لایه (MLP) مورد ارزیابی قرار گرفته‌اند. نتایج ارائه شده در این بخش میانگین ۱۰ بار اجرای روش پیشنهادی است.

جدول ۶. مقایسه کارایی روش پیشنهادی بر اساس معیار AUC برای

طبقه‌بند KNN

مجموعه داده	بدون داده‌افزایی	LoRAS	H-SMOTE	HDAM
D1	۰/۹۰۸۴	۰/۸۷۴۸	۰/۹۳۲۶	۰/۹۳۲۲
D2	۰/۷۹۸۵	۰/۸۶۴۳	۰/۸۶۲۹	۰/۸۷۴۱
D3	۰/۸۲۴۷	۰/۸۶۶۳	۰/۸۷۱۱	۰/۸۹۰۸
D4	۰/۷۸۳۵	۰/۸۰۷۹	۰/۸۷۷۷	۰/۸۶۶۷
D5	۰/۵۲۹۴	۰/۷۶۰۴	۰/۷۶۷۱	۰/۷۶۰۸
D6	۰/۵۸۱۰	۰/۸۳۹۱	۰/۸۵۸۳	۰/۸۵۷۹
D7	۰/۵۰۰۰	۰/۶۲۹۳	۰/۷۶۱۵	۰/۷۷۰۰
D8	۰/۶۳۹۴	۰/۶۲۱۷	۰/۶۶۱۱	۰/۷۱۳۱
D9	۰/۴۹۹۵	۰/۵۴۹۳	۰/۵۰۰۷	۰/۵۵۳۹
D10	۰/۵۰۰۰	۰/۶۱۳۵	۰/۵۸۶۸	۰/۵۵۶۲
میانگین	۰/۶۵۶۴	۰/۷۴۷۴	۰/۷۶۷۹	۰/۷۷۷۵

جدول ۷. مقایسه کارایی روش‌های پیشنهادی بر اساس معیار AUC

برای طبقه‌بند MLP

مجموعه داده	بدون داده‌افزایی	LoRAS	H-SMOTE	HDAM
D1	۰/۸۵۹۲	۰/۸۰۱۶	۰/۸۳۶۹	۰/۹۶۲۸
D2	۰/۸۱۲۶	۰/۸۴۶۸	۰/۸۸۵۲	۰/۸۹۷۷
D3	۰/۸۵۶۷	۰/۸۳۶۸	۰/۸۶۷۷	۰/۸۸۷۸
D4	۰/۷۸۹۰	۰/۸۳۳۲	۰/۸۶۰۱	۰/۸۹۷۰
D5	۰/۵۲۸۲	۰/۶۷۵۵	۰/۶۸۶۰	۰/۷۵۸۰
D6	۰/۵۰۰۰	۰/۷۳۵۹	۰/۷۶۲۰	۰/۷۹۳۳
D7	۰/۵۰۰۰	۰/۵۱۱۱	۰/۵۹۳۷	۰/۷۳۱۶
D8	۰/۷۳۷۰	۰/۶۹۹۳	۰/۷۱۴۳	۰/۸۰۱۴
D9	۰/۵۰۰۰	۰/۵۵۷۴	۰/۵۰۳۹	۰/۵۳۶۹
D10	۰/۵۰۰۰	۰/۵۹۴۲	۰/۵۸۵۷	۰/۶۶۸۱
میانگین	۰/۶۵۸۲	۰/۷۰۹۱	۰/۷۲۹۶	۰/۷۹۳۶

جدول ۸. کارایی روش پیشنهادی بر اساس معیارهای دقت و فراخوانی

دقت			
طبقه‌بند	HDAM	H-SMOTE	LoRAS
RF	۱۰	۰	۰
KNN	۵	۴	۳
MLP	۱۰	۰	۰
مجموع	۲۵	۴	۳
فراخوانی			
طبقه‌بند	HDAM	H-SMOTE	LoRAS
RF	۱۰	۰	۰
KNN	۳	۵	۲
MLP	۹	۱	۰
مجموع	۲۲	۶	۲

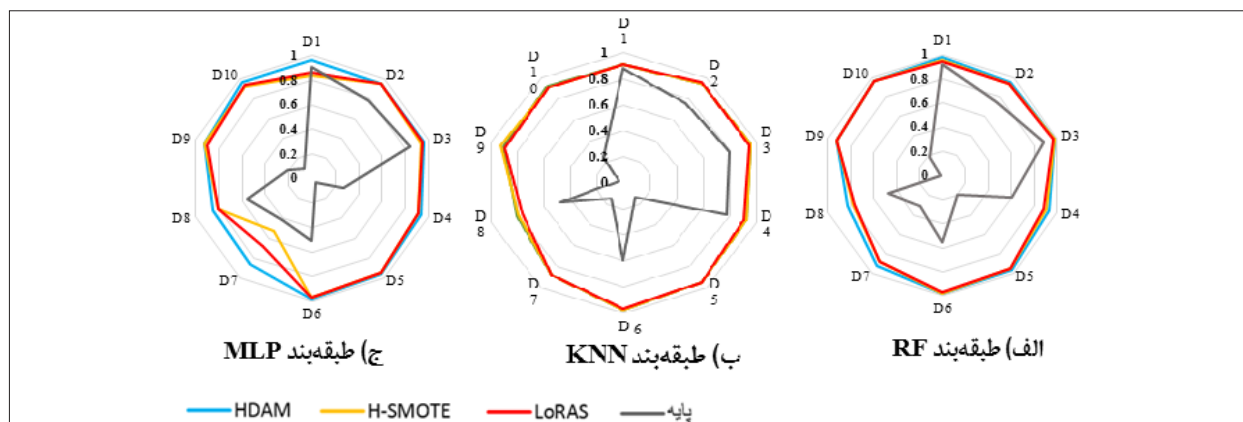
در چهار مجموعه داده به AUC بیشتری دست یافته است. در مجموعه داده‌های D1، D4، D5 و D6 روش H-SMOTE و در مجموعه داده‌های D9 و D10 روش LoRAS بهتر عمل کرده است. به‌طور میانگین، روش پیشنهادی HDAM با مقدار ۰/۷۷۷۵ برای AUC عملکرد بهتری نسبت به سایر روش‌ها دارد.

جدول (۷) نیز نشان می‌دهد که برای طبقه‌بند MLP، روش پیشنهادی HDAM با میانگین امتیاز ۰/۷۹۳۶ برای AUC به نتایج بهتری در مقایسه با سایر روش‌ها دست یافته است. برای این طبقه‌بند، روش پیشنهادی در ۹ مورد از ۱۰ مجموعه داده مورد بررسی، عملکرد بهتری در مقایسه با دیگر روش‌ها داشته است.

جدول (۸) نشان می‌دهد هر روش در چند مجموعه داده بیشترین کارایی برحسب معیارهای دقت/فراخوانی را در مقایسه با سایر روش‌ها به‌دست آورده است. مطابق جدول (۸)، روش HDAM در ۲۵ مورد از ۳۰ آزمایش صورت گرفته دارای دقت بالاتری است و روش H-SMOTE با ۴ مورد در جایگاه دوم قرار دارد. همچنین، روش HDAM در ۲۲ مورد از مجموع ۳۰ آزمایش انجام شده فراخوانی بالاتری در مقایسه با دو روش دیگر به‌دست آورده است. به دلیل این که برخی روش‌ها بر روی بعضی از مجموعه داده‌ها عملکرد یکسانی دارند، در برخی سطرهای جدول

(۸) حاصل جمع مقادیر از ۱۰ بیشتر است. به‌عنوان مثال از مقایسه دقت طبقه‌بند KNN این نتیجه حاصل شد که برای مجموعه داده D5، دو روش LoRAS و HDAM مقدار یکسان و برتر ۰/۹۶۰۶ را کسب کرده‌اند.

شکل (۲) معیار F1 حاصل از اعمال روش‌های متوازن‌سازی مورد بررسی را روی ۱۰ مجموعه داده با استفاده از سه طبقه‌بند KNN، RF و MLP مقایسه می‌کند. همچنین جدول (۹)، میانگین معیار F1 بر روی مجموعه داده‌ها را برای هر طبقه‌بند نشان می‌دهد. مطابق شکل (۲-الف) نتایج حاصل از طبقه‌بند RF نشان‌دهنده آن است که



شکل ۲. مقادیر F1 حاصل از سه طبقه‌بند (الف) جنگل تصادفی، (ب) K نزدیک‌ترین همسایه و (ج) پرسپترون چندلایه

مورد آزمایش قرار گرفت. نتایج نشان داد که روش پیشنهادی در تمام موارد با طبقه‌بندهای RF و MLP موجب بهبود قابل توجه معیارهای AUC، دقت، فراخوانی و F1 نسبت به روش‌های موجود شده است. با این حال، در طبقه‌بند نزدیک‌ترین همسایه (KNN) با وجود این که به‌طور میانگین بهبود اندکی داشته است، در برخی مجموعه‌داده‌ها عملکرد اندکی ضعیف‌تر از سایر روش‌ها بود. دلیل این امر آن است که KNN به‌شدت به توزیع داده‌ها در فضای ویژگی حساس است و افزایش نمونه‌های مصنوعی در مناطق مرزی گاهی می‌تواند منجر به طبقه‌بندی اشتباه شود. بنابراین، هر چند روش پیشنهادی در اکثر طبقه‌بندها به‌ویژه مدل‌های مبتنی بر شبکه عصبی و درخت تصمیم، نتایج بهتری را ارائه داده است، اما باید محدودیت‌های آن در مواجهه با طبقه‌بندهای مبتنی بر فاصله نیز مورد توجه قرار گیرد.

۵- نتیجه‌گیری

در این مقاله یک روش ترکیبی برای داده‌افزایی معرفی شد که در هر گام آن، به‌صورت تصادفی یکی از دو الگوریتم H-SMOTE یا LoRAS انتخاب شده و یک نمونه از کلاس اقلیت با روش انتخاب شده تولید می‌شود. این فرآیند تا برآورده شدن شرط برابری تعداد نمونه‌های کلاس اقلیت و اکثریت و متوازن شدن مجموعه داده ادامه می‌یابد. نتایج ارزیابی‌های متعدد انجام شده بر روی ۱۰ مجموعه داده از

جدول ۹. میانگین مقادیر F1 حاصل از ۱۰ مجموعه داده برای سه

طبقه‌بند RF، KNN و MLP

طبقه‌بند	بدون داده‌افزایی	LoRAS	H-SMOTE	HDAM
RF	۰/۴۹۰۴	۰/۹۱۸۰	۰/۹۲۶۰	۰/۹۴۲۴
KNN	۰/۴۸۴۲	۰/۹۰۶۴	۰/۹۱۵۶	۰/۹۱۶۱
MLP	۰/۴۶۳۹	۰/۸۹۲۲	۰/۸۷۵۳	۰/۹۳۶۲

روش پیشنهادی HDAM در کلیه ۱۰ مجموعه داده، از F1 بالاتری نسبت به روش‌های پیشین و حالت پایه برخوردار هستند. جدول (۹) نیز نشان می‌دهد، در مقایسه با حالت پایه (بدون داده‌افزایی)، روش HDAM به‌طور میانگین به بهبود ۴۵/۲ درصد در F1 دست یافته است.

بر اساس جدول (۹) برای طبقه‌بند KNN، روش پیشنهادی HDAM بیشترین میزان بهبود (۴۳/۱۹ درصد) را نسبت به حالت پایه دارد. همچنین بر طبق شکل (۲-ب)، با به‌کارگیری طبقه‌بند HDAM، KNN در ۴ مجموعه داده D10 و D4، D5، D8 مقادیر F1 بالاتری را در مقایسه با روش‌های پیشین کسب نموده است.

در نهایت مطابق شکل (۲-ج)، روش HDAM در ۹ مجموعه داده (بجز D9) بالاترین میزان F1 را در مقایسه با روش‌های پیشین کسب کرده است. جدول (۹) نیز نشان می‌دهد در صورت استفاده از طبقه‌بند MLP بهبود ۴۷/۲۳ درصدی F1 روش پیشنهادی در مقایسه با حالت پایه بدون متوازن‌سازی حاصل گردیده است.

در ارزیابی‌های انجام شده، روش پیشنهادی بر روی مجموعه‌داده‌های مختلف و با استفاده از چندین طبقه‌بند

- [9] Chao, X. & Zhang, L. 2023. Few-shot imbalanced classification based on data augmentation. *Multimedia Systems*, Vol. 29, No. 5, pp.2843-2851.
- [10] Bej, S., Davtyan, N., Wolfien, M., Nassar, M. & Wolkenhauer, O. 2021. "LoRas: an oversampling approach for imbalanced datasets," *Machine Learning*, Vol. 110, No. 2, pp. 279-301.
- [11] Temraz, M. & Keane, M. T. 2022. "Solving the class imbalance problem using a counterfactual method for data augmentation," *Machine Learning with Applications*, vol. 9, pp.100375.
- [12] Melícias, F. S., Ribeiro, T. F., Rabadão, C., Santos, L., & Costa, R. L. D. C. 2024. "GPT and interpolation-based data augmentation for multiclass intrusion detection in IIoT," *IEEE Access*, Vol. 12, pp. 17945-17965.
- [13] Dai, H., Liu, Z., Liao, W., Huang, X., Cao, Y. Wu, Z. & Li, X. 2025. "Auggpt: Leveraging chatgpt for text data augmentation," *IEEE Transactions on Big Data*.
- [14] Blake, C. & Merz, C. 1978. "UCI Machine Learning Repository," Department of Information and Computer Science, [Online] Available: <https://archive.ics.uci.edu/datasets>. [Accessed 2025].
- [15] Derrac, J., Garcia, S., Sanchez, L. & Herrera, F. 2015. "Keel data-mining software tool: Data set repository, integration of algorithms and experimental analysis framework," *J. Mult. Valued Logic Soft Comput*, 17, pp. 255-287.
- [16] Chih-Wei, H. & Chih-Jen, L. 2002. "A comparison of methods for multiclass support vector machines," *IEEE Transactions on Neural Networks*, Vol. 13, No. 2, pp. 415-425.
- [17] Bishop, C. M. & Nasrabadi, N. M. 2006. "Pattern recognition and machine learning," New York: springer, Vol. 4, No. 4, pp. 738.

مخازن UCI و KEEL نشان می‌دهد که طبقه‌بندهای یادگیری ماشین هنگامی که بر روی مجموعه داده‌های متوازن شده با روش پیشنهادی آموزش دیده‌اند، در مقایسه با دو روش HSMOTE و LoRAS عملکرد بهتری بر اساس معیارهای رایج طبقه‌بندی یعنی دقت، فراخوانی، F1 و AUC نشان داده‌اند. این نتایج نشان می‌دهند که ترکیب دو یا چند روش داده‌افزایی با خصوصیات متفاوت و مکمل می‌تواند از مزایای همه روش‌ها بهره‌مند شود و مشکلات استفاده تکی از هر روش را کاهش دهد. به‌عنوان کار آتی، برآینم تا با ترکیب روش‌های داده‌افزایی مناسب برای طبقه‌بندی چندکلاسه، کارایی روش پیشنهادی را برای طبقه‌بندی چندکلاسه مورد بررسی قرار دهیم.

۶. مراجع

- [1] Khan, M. H., Iqbal, M. S., Muhammad Shah, F., Al Mahmud, J., Popel, M. H., Showrov, M. I. H., Ahmed, S. & Rahman, O. 2020. "A Survey of Methods for Managing the Classification and Solution of Data Imbalance Problem," *Journal of Computer Science*, Vol. 16, No. 11, pp. 1546-1557.
- [2] Tsai, C. F., Lin, W. C., Hu, Y. H. & Yao, G. T. 2019. "Under-Sampling Class Imbalanced Datasets by Combining Clustering Analysis and Instance Selection," *Information Sciences*, vol. 477, pp. 47-54, 2019.
- [3] Huang, M. W., Tsai, Ch. F. & Lin, W. Ch. 2021. "Instance selection in medical datasets: Adivide-and-conquer framework," Vol. 90, pp. 106957.
- [4] Moran, M., Cohen, T., Ben-Zion, Y. & Gordon, G. 2022. "Curious instance selection," *Information Sciences*, Vol. 608, pp. 794-808.
- [5] Chawla, N. V., Bowyer, W., Hall, L. O. & Kegelmeyer, W. P. 2002. "SMOTE: Synthetic Minority Over-sampling Technique," *Journal of Artificial Intelligence Research*, Vol. 16, pp. 321-357.
- [6] Douzas, G. & Bacao, F. 2019. "Geometric SMOTE a geometrically enhanced drop-in replacement for SMOTE," *Information Sciences*, Vol. 501, pp. 118-135.
- [7] Sáez, J. A., Luengo, J., Stefanowski, J. & Herrera, F. 2015. "SMOTE-IPF: Addressing the noisy and borderline examples problem in imbalanced classification by a re-sampling method with filtering," *Information Sciences*, Vol. 291, pp. 184-203, 2015.
- [8] Bunkhumpornpat, C., Sinapiromsaran, K. & Lursinsap, C. 2009. "Safe-Level-SMOTE: Safe-Level-Synthetic Minority Over-Sampling Technique for Handling the Class Imbalanced Problem," in *Advances in Knowledge Discovery and Data Mining*, New York.