

الگوریتم یکپارچه بهینه سلسله‌مراتبی برای رتبه‌بندی اعتباری: تلفیق بهینه یادگیری عمیق و متا-کلاسیفایر مبتنی بر جنگل تصادفی

نوع مقاله: پژوهشی

مهدی فرضی^۱

یعقوب پورکریم^۲

سیدعلی پایتخی اسکوئی^۳

مهدی زینالی^۴

رسول برادران حسن‌زاده^۵

تاریخ دریافت: ۱۴۰۴/۱/۳۰

تاریخ پذیرش: ۱۴۰۴/۴/۲۱

چکیده

در حوزه مدیریت ریسک مالی، رتبه‌بندی اعتباری به عنوان مکانیسمی حیاتی برای پیش‌بینی احتمال بازپرداخت تسهیلات توسط متقاضیان شناخته می‌شود. اگرچه مدل‌های سنتی مبتنی بر یادگیری ماشین در این زمینه مورد استفاده گسترده قرار گرفته‌اند، ادغام راهکارهای نوین یادگیری عمیق با پارادایم‌های یادگیری جمعی به عنوان گامی تحول‌آفرین در افزایش دقت پیش‌بینی مطرح شده است. این پژوهش، الگوریتم یکپارچه بهینه سلسله‌مراتبی HUOA را معرفی می‌کند که از سینرژی بین سه لایه پردازشی پیشرفته بهره می‌برد. در لایه پایه، سه کلاسیفایر مبتنی بر یادگیری جمعی شامل Bagging، AdaBoost و شبکه حافظه بلند - کوتاه‌مدت LSTM به صورت موازی جهت استخراج ویژگی‌های سطح اول به کار گرفته می‌شوند. خروجی این لایه وارد لایه متاآموزش می‌شود که در آن یک فراآموزش‌دهنده مبتنی بر جنگل تصادفی با معماری تطبیق‌پذیر، اقدام به ترکیب غیرخطی

mehdi.farzi3357@iau.ac.ir

^۱ گروه حسابداری، واحد تبریز، دانشگاه آزاد اسلامی، تبریز، ایران

^۲ گروه حسابداری، واحد تبریز، دانشگاه آزاد اسلامی، تبریز، ایران (نویسنده مسئول)

pour Karim @ iaut.ac.ir

^۳ گروه اقتصاد، واحد تبریز، دانشگاه آزاد اسلامی، تبریز، ایران

oskoee@iaut.ac.ir

^۴ گروه حسابداری، واحد تبریز، دانشگاه آزاد اسلامی، تبریز، ایران

zeynali@iaut.ac.ir

^۵ گروه حسابداری، واحد تبریز، دانشگاه آزاد اسلامی، تبریز، ایران

baradaran313@iaut.ac.ir

✉ Hierarchical Unified Optimal Algorithm

پیش‌بینی‌ها و تولید امتیاز نهایی ریسک اعتباری می‌نماید. در این پژوهش، یک چارچوب یادگیری عمیق یکپارچه برای رتبه‌بندی اعتباری مشتریان بانک پیشنهاد شده است که بر اساس یادگیری جمعی و بهینه‌سازی یکپارچه پارامترها و انتخاب ویژگی‌ها با استفاده از الگوریتم ژنتیک (GA) طراحی گردیده است. ارزیابی تجربی بر روی داده‌های مشتریان حقوقی یک بانک ایرانی همراه با مجموعه داده بین‌المللی UCI استرالیا و آلمان با معیارهای مختلف و به ویژه دسته‌بندی اشتباه (MC) نشان می‌دهد که HUOA نسبت به روش‌های ترکیبی موجود در بازه زمانی ۲۰۲۳-۲۰۲۵، بهبودی قابل توجه حاصل نموده است. این معماری سلسله‌مراتبی نه تنها قابلیت تفسیرپذیری مدل را از طریق تحلیل اهمیت ویژگی در لایه متا حفظ می‌کند، بلکه با کاهش واریانس پیش‌بینی در سناریوهای نامتوازن کلاس، چارچوبی مقاوم برای تصمیم‌گیری در محیط‌های بانکی پویا ارائه می‌دهد. یافته‌ها حاکی از آن است که تلفیق هوشمندانه LSTM با استراتژی‌های نمونه‌برداری پویا در لایه جمعی، همراه با مکانیزم انتخاب ویژگی تطبیقی در لایه متا، می‌تواند به عنوان پارادایمی جدید در سیستم‌های امتیازدهی اعتباری نسل چهارم مورد توجه قرار گیرد.

واژه‌های کلیدی: رتبه‌بندی اعتباری، یادگیری جمعی، متاکلاسیفایر، بهینه‌سازی، الگوریتم ژنتیک، انتخاب ویژگی

طبقه‌بندی JEL: G17, G21, G32, C45, C63, O33

مقدمه

بانک‌ها به‌عنوان موتور محرکه اقتصاد، نقشی اساسی در تجهیز پس‌اندازها و هدایت آن‌ها به سمت فعالیت‌های مولد اقتصادی ایفا می‌کنند. با این حال، این نقش حیاتی همواره با ریسک اعتباری همراه است؛ ریسکی که ریشه در احتمال عدم بازپرداخت تسهیلات توسط وام‌گیرندگان دارد و در صورت تحقق، نه تنها سودآوری مؤسسات مالی را تحت تأثیر قرار می‌دهد، بلکه ثبات نظام مالی را به مخاطره می‌اندازد. مطالعاتی مانند پژوهش داس^۱ (۲۰۱۹) نشان می‌دهد که بخش عمده دارایی بانک‌ها را مطالبات تشکیل می‌دهد و شکست در مدیریت این دارایی‌ها می‌تواند پیامدهای سنگینی همچون کاهش حقوق صاحبان سهام یا حتی ورشکستگی را در پی داشته باشد. در پاسخ به این چالش، چارچوب‌های نظارتی مانند بازل II با تأکید بر «حساسیت به ریسک»، بانک‌ها را ملزم به توسعه مدل‌های داخلی مبتنی بر داده‌های تاریخی یا استفاده از مؤسسات رتبه‌بندی اعتباری کرد. این تحول، الزامات سرمایه‌ای را از حالت ثابت به متغیر تبدیل نمود و آن را به سطح ریسک پرتفوی اعتباری بانک‌ها گره زد (لاندو^۲، ۲۰۰۹).

در ایران، این چالش‌ها با پیچیدگی‌های خاصی همراه است. نظام مالی کشور به‌طور تاریخی بانک‌محور بوده و بیش از ۹۰ درصد تأمین مالی اقتصاد از طریق بانک‌ها و مؤسسات اعتباری غیربانکی صورت می‌گیرد. این تمرکز بالا، همراه با بحران‌های مکرر مالی، نوسانات ارزی، و تحریم‌های بین‌المللی، احتمال نکول بانک‌ها را به معضلی ساختاری تبدیل کرده است. بررسی‌ها نشان می‌دهد که ناتوانی در ایفای تعهدات توسط یک بانک، نه تنها سپرده‌گذاران و سهامداران را متضرر می‌کند، بلکه می‌تواند به سرایت بحران به کل نظام مالی بینجامد. در چنین بستری، توسعه مدل‌های رتبه‌بندی اعتباری دقیق نه یک انتخاب، که ضرورتی اجتناب‌ناپذیر است. سیستم‌های سنتی امتیازدهی که بر پایه شاخص‌های ایستایی مانند درآمد یا سابقه اعتباری عمل می‌کنند، در مواجهه با داده‌های نامتوازن، وابستگی‌های زمانی پیچیده، و ریسک‌های نوظهور ناشی از تحولات ژئوپلیتیک یا سیاست‌های پولی، اغلب ناتوان می‌مانند (امیری و همکاران، ۱۴۰۰).

سیستم‌های امتیازدهی اعتباری به‌عنوان ابزاری کلیدی در این فرایند، با تحلیل شاخص‌های کیفی و کمی نظیر سابقه مالی، درآمد، و نسبت بدهی به دارایی، احتمال نکول را پیش‌بینی می‌کنند. این سیستم‌ها نه تنها امکان شناسایی مشتریان پرریسک را فراهم می‌سازند، بلکه با تعیین نرخ سود متناسب با سطح ریسک، تا ۳۰ درصد از هزینه‌های ناشی از مطالبات معوقه را کاهش می‌دهند (توماس^۳ و همکاران، ۲۰۱۳). با این حال، جهانی‌شدن فعالیت‌های بانکی، پیچیدگی نوآوری‌های مالی،

^۱ Das

^۲ Lando

^۳ Thomas

و وقوع بحران‌های غیرمنتظره، لزوم بازنگری در رویکردهای سنتی را آشکار ساخته است. بحران مالی ۲۰۰۸ نمادی بارز از این واقعیت بود که حتی پیشرفته‌ترین مدل‌ها نیز در غیاب توجه به ریسک‌های سیستماتیک و عوامل کلان اقتصادی می‌توانند با شکست مواجه شوند. از این رو، بانک‌های امروزی، فراتر از رتبه‌بندی فردی وام‌گیرندگان، به تحلیل روندهای اقتصادی، تحولات ژئوپلیتیک، و سیاست‌های پولی روی آورده‌اند تا ریسک‌های نوظهور را پیش‌بینی کنند (ندری و محرابی، ۱۳۹۷). در این راستا، فناوری‌های تحلیلی پیشرفته مانند یادگیری عمیق و هوش مصنوعی، تحولی بنیادین در حوزه رتبه‌بندی اعتباری ایجاد کرده‌اند. شبکه‌های عصبی عمیق، به ویژه معماری‌هایی مانند LSTM، با قابلیت شناسایی الگوهای پویا در داده‌های زمانی، امکان مدل‌سازی رفتار مالی مشتریان را در بازه‌های طولانی فراهم می‌کنند. با این وجود، چالش‌هایی همچون وابستگی به ساختار داده، حساسیت به عدم تعادل کلاس‌ها، و نیاز به تنظیم دقیق پارامترها، محدودیت‌های قابل توجهی را به همراه دارند. برای نمونه، یک مدل LSTM ممکن است در داده‌های متوازن عملکردی مطلوب داشته باشد، اما در شرایط نامتوازن که نمونه‌های نکول اندک هستند، به سمت کلاس اکثریت سوگیری کند (دو و شو، ۲۰۲۲؛ آلا‌راج و همکاران، ۲۰۲۲).

پژوهش حاضر با هدف غلبه بر این چالش‌ها، چارچوبی یکپارچه مبتنی بر یادگیری جمعی عمیق و بهینه‌سازی هوشمند ارائه می‌دهد. در این روش، ترکیبی از مدل‌های پایه شامل LSTM، AdaBoost، و Bagging به صورت موازی آموزش دیده و خروجی آن‌ها به یک متاکلاسیفایر مبتنی بر جنگل تصادفی منتقل می‌شود. هسته اصلی نوآوری این پژوهش، استفاده از الگوریتم ژنتیک به عنوان موتور بهینه‌سازی یکپارچه است که دو وظیفه کلیدی را همزمان انجام می‌دهد: نخست، تنظیم هایپرپارامترهای هر مدل پایه و دوم، انتخاب ویژگی‌های مرتبط از طریق مکانیزم وزن‌دهی تطبیقی. این رویکرد نه تنها دقت پیش‌بینی را افزایش می‌دهد، بلکه با کاهش ابعاد داده و حذف نویز، امکان استقرار در سیستم‌های بانکی بلادرنگ را فراهم می‌سازد. ارزیابی تجربی بر پایه داده‌های معتبر بین‌المللی و مجموعه داده‌های واقعی یک بانک ایرانی، گواه بهبود معنادار معیارهای عملکردی نظیر دقت، فراخوانی، و کاهش خطای طبقه‌بندی در مقایسه با روش‌های موجود است. این چارچوب، با حفظ تعادل میان دقت و تفسیرپذیری، گامی به سوی نسل چهارم سیستم‌های امتیازدهی اعتباری است که پاسخگوی نیازهای پیچیده بانکداری مدرن است.

1 Du & Shu

2 Ala'raj

۱. پیشینه پژوهش

بانک‌ها، به‌عنوان نهادهای محوری در گردش مالی اقتصاد، مسئولیت خطیر تبدیل پس‌اندازها به سرمایه‌گذاری‌های مولد را بر عهده دارند. این نقش حیاتی، همواره در معرض تهدید ریسک اعتباری قرار دارد؛ ریسکی که نه تنها سودآوری مؤسسات مالی را تحت تأثیر قرار می‌دهد، بلکه مانند دومینویی، ثبات کل نظام اقتصادی را به خطر می‌اندازد. بر اساس گزارش اخیر صندوق بین‌المللی پول، نزدیک به ۴۰ درصد از زبان‌های بانک‌های در حال توسعه، ناشی از شکست در پیش‌بینی دقیق نکول مشتریان است. این چالش در اقتصادهایی مانند ایران، که بیش از ۹۰ درصد تأمین مالی از طریق نظام بانکی صورت می‌گیرد، ابعاد پیچیده‌تری به خود می‌گیرد. بحران‌های مکرر مالی، نوسانات نرخ ارز، و تحریم‌های بین‌المللی، لایه‌های جدیدی از ریسک‌های سیستماتیک را به مدل‌های سنتی رتبه‌بندی تحمیل کرده‌اند که پاسخگویی به آن‌ها نیازمند تحولی بنیادین در معماری سیستم‌های امتیازدهی است.

تکامل تاریخی روش‌های رتبه‌بندی اعتباری، از مدل‌های تجربی مبتنی بر قضاوت کارشناسان تا الگوریتم‌های پیچیده یادگیری عمیق، گواهی بر این تلاش مستمر است. در دهه ۱۹۹۰، ظهور شبکه‌های عصبی مصنوعی امیدی تازه برای مدل‌سازی غیرخطی روابط بین متغیرها ایجاد کرد، اما محدودیت‌هایی مانند سوگیری در داده‌های نامتوازن و ناتوانی در پردازش الگوهای زمانی، راه را برای نوآوری‌های بعدی هموار ساخت. در این راستا، معماری‌های پیشرفته‌ای مانند LSTM (دو و شو، ۲۰۲۲) و ترنسفورمرها (ژیانگ و همکاران، ۲۰۲۴) با قابلیت استخراج خودکار ویژگی‌های پیچیده از داده‌های دنباله‌ای، دقت پیش‌بینی را به سطح بی‌سابقه‌ای رساندند. با این حال، مطالعات اخیر نشان می‌دهند که حتی این مدل‌ها نیز در مواجهه با داده‌های پراکنده و نامتوازن نمونه‌های نکول (که گاهی کمتر از ۲ درصد مجموعه داده را تشکیل می‌دهند)، به سمت کلاس اکثریت سوگیری می‌کنند و هزینه‌های سنگینی را به سیستم بانکی تحمیل می‌نمایند (هارتومو و همکاران، ۲۰۲۵).

در سال‌های ۲۰۲۴ تا ۲۰۲۵، پژوهشگران با تمرکز بر دو محور تفسیرپذیری مدل و ادغام ریسک‌های سیستماتیک، گام‌های بلندی برداشته‌اند. برای مثال، شتاکیس^۳ و همکاران (۲۰۲۴) با ترکیب محاسبات کوانتومی و شبکه‌های عصبی کانولوشنی، موفق به کاهش ۳۰ درصدی خطای نوع II در شناسایی مشتریان پرریسک شدند، اما ناتوانی این مدل در تحلیل همزمان داده‌های ساختاریافته و غیرساختاریافته (مانند اخبار اقتصادی)، محدودیت جدیدی را آشکار ساخت. از سوی دیگر، هارتومو و همکاران (۲۰۲۵) با معرفی LSTM مبتنی بر مکانیسم توجه، امکان تفسیر تصمیم‌های مدل را با

^۱ Xiong

^۲ Hartomo

^۳ Schetakis

شناسایی وزن ویژگی‌های کلیدی (مانند سابقه پرداخت‌های تأخیری) فراهم کردند، اما نادیده گرفتن اثرات ژئوپلیتیک بر ریسک اعتباری، همچنان به عنوان نقطه ضعفی جدی باقی ماند. خلاصه پیشینه پژوهش در جدول ۱ نشان داده شده است.

جدول ۱. خلاصه پیشینه پژوهش

سال انتشار	نویسنده (ها)	روش	نتایج
۱۴۰۰	ابتیاع و همکاران	یادگیری گروهی بر پایه ماشین بردار پشتیبان	نتایج بدست آمده بر روی داده های بانک پاسارگاد، برتری روش ماشین بردار پشتیبان گروهی بر روش معمولی ماشین بردار پشتیبان و روش جنگل تصادفی را تأیید می‌کند.
۱۴۰۲	خانجانی و همکاران	یادگیری عمیق	استفاده از هوش مصنوعی در فرآیند رتبه بندی اعتباری، باعث بهبود دقت و کیفیت تصمیم گیری مدل‌های رتبه بندی می‌شود. همچنین، استفاده از روش‌های یادگیری عمیق در مقایسه با یادگیری ماشینی، در بهبود دقت و کیفیت تصمیم گیری موثرتر بوده است.
۲۰۲۳	ژانگ و همکاران	یک مدل امتیازدهی اعتباری گروهی بر اساس رگرسیون لجستیک با اثرات تعادلی و وزن دهی ناهمگن	نتایج نشان می‌دهد که در مقایسه با ده مدل امتیازدهی اعتباری نماینده در شش مجموعه داده عمومی، مدل لجستیک-BWE قوی‌ترین توانایی را برای تشخیص نمونه‌های پیش‌فرض دارد و بهترین توانایی تعمیم را در اکثر مجموعه‌های داده با حفظ قابلیت تفسیر دارد.

^۱ Zhang

سال انتشار	نویسنده (ها)	روش	نتایج
۲۰۲۳	خلیلی و رستگار	مدل بهینه حساس به هزینه ترکیبی با مدل‌های KNN، DT، SVM، AdaBoost شبکه عصبی احتمالاتی PNN	نتایج این تحقیق نشان داد که روش پیشنهادی دارای دقت بالایی برای هر دو داده متوازن و نامتوازن است.
۲۰۲۴	ژیانگ و همکاران	مدل ارزیابی ریسک اعتباری جدید مبتنی بر یادگیری عمیق با معماری Deep Q-Network	نتایج نشان دهنده سازگاری قابل توجه مدل پیشنهادی در سناریوهای پویای شبیه سازی شده است و میانگین AUC-ROC 0.89 را در شوک‌های مختلف اقتصادی و تغییرات بازار حفظ می‌کند.
۲۰۲۵	هارمونو و همکاران	مدل ترنسفورمر XAI با تکنیک‌های کاهش وزن	این روش برای مجموعه داده‌های BISAI و اعتبار آلمانی، پیشرفت‌های قابل توجهی را در دقت، از ۸۶،۳۵ درصد به ۸۹،۲۷ درصد و ۹۳ درصد به ۹۵ درصد، همراه با بهبود AUC کلاس اقلیت و معیارهای فراخوان دقیق نشان داد.

۲. روش‌شناسی پژوهش

این پژوهش از نظر هدف، کاربردی بوده و از لحاظ روش‌شناسی در دسته تحقیقات توصیفی قرار می‌گیرد. داده‌های مورد استفاده شامل اطلاعات مالی و ویژگی‌های دموگرافیک اشخاص حقیقی و حقوقی دریافت‌کننده وام از بانک‌ها هستند. برای پیاده‌سازی مدل نیز از نرم‌افزار Python بهره گرفته می‌شود.

۱.۲. داده‌ها

این پژوهش از داده‌های مشتریان حقوقی یک بانک ایرانی به همراه دو مجموعه داده مرجع از کشورهای آلمان و استرالیا استفاده می‌کند. بررسی الگوهای توزیع مشتریان خوش حساب و بدحساب در این سه کشور نشان می‌دهد که ساختار اعتباری هر کشور دارای تفاوت‌های قابل توجهی است. به عنوان مثال، در ایران از میان ۲,۹۸۷ مشتری حقوقی، ۱,۷۲۱ مشتری (معادل ۵۷,۶۲ درصد) خوش حساب و ۱,۲۶۶ مشتری (معادل ۴۲,۳۸ درصد) بدحساب محسوب می‌شوند؛ بدین ترتیب نسبت خوش حساب به بدحساب برابر با ۱,۳۶ به ۱ است. در آلمان، بررسی داده‌های یک مجموعه ۱,۰۰۰ نفری نشان می‌دهد که ۷۰۰ مشتری (۷۰ درصد) خوش حساب و ۳۰۰ مشتری (۳۰ درصد) بدحساب هستند؛ به طوری که نسبت خوش حساب به بدحساب در این کشور برابر با ۲,۳۳ به ۱ بوده و نشان‌دهنده سطح بالاتر تعهد مالی مشتریان می‌باشد. در مقابل، در استرالیا از میان ۶۹۰ مشتری، تنها ۳۰۷ نفر (۴۴,۴۹ درصد) خوش حساب و ۳۸۳ نفر (۵۵,۵۱ درصد) بدحساب تشخیص داده شده‌اند؛ بنابراین نسبت خوش حساب به بدحساب در استرالیا برابر با ۰,۸۰ به ۱ است که بیانگر ریسک اعتباری بالاتر در این کشور می‌باشد. جدول ۲ متغیرهای اعتباری مشتریان حقوقی بانک ایرانی و جدول ۳ متغیرهای کیفی به همراه مقادیر مربوطه را ارائه می‌دهند که در تحلیل رتبه‌بندی اعتباری این پژوهش به کار گرفته شده‌اند.

جدول ۲. متغیرهای اعتباری مشتریان حقوقی بانک ایرانی

ردیف	متغیر	ردیف	متغیر
۱	وضعیت وام	۱۹	نسبت حساب دریافتی به فروش خالص
۲	میزان وام (ریال)	۲۰	نسبت حساب دریافتی به بدهی‌ها
۳	موجودی نقد در بانک (ریال)	۲۱	نسبت حساب پرداختی به فروش خالص
۴	نسبت بدهی به دارایی	۲۲	نسبت فروش به دارایی
۵	نسبت حقوق صاحبان به دارایی	۲۳	نسبت فروش به دارایی ثابت
۶	نسبت بدهی بلند مدت به دارایی	۲۴	نسبت سود خالص به دارایی‌ها
۷	نسبت سود خالص به هزینه مالی	۲۵	نسبت سود خالص به فروش خالص
۸	نسبت دارایی جاری به بدهی جاری	۲۶	نسبت سود خالص به دارایی ثابت

نسبت سود خالص به حقوق صاحبان	۲۷	نسبت حساب دریافتی و موجودی نقد به بدهی جاری	۹
نسبت هزینه فروش به فروش خالص	۲۸	نسبت سرمایه در گردش به بدهی‌ها	۱۰
اندازه بنگاه	۲۹	نسبت دارایی جاری به بدهی‌ها	۱۱
وضعیت مالکیت ملک (اجاره، سرقفلی، دارای سند مالکیت غیر قابل ترهین)	۳۰	نسبت بدهی جاری به دارایی‌ها	۱۲
ارتباط بین فعالیت و مدرک (۰ یا ۱)	۳۱	نسبت موجودی نقد به دارایی‌ها	۱۳
حسن شهرت (عددی بین ۰ تا ۱)	۳۲	نسبت موجودی نقد به فروش خالص	۱۴
نوع وثیقه (عددی بین ۰ تا ۱ برحسب نوع)	۳۳	نسبت سرمایه در گردش به فروش خالص	۱۵
مدت بازپرداخت (سال)	۳۴	نسبت دارایی جاری به فروش خالص	۱۶
نرخ سود	۳۵	نسبت موجودی نقد به بدهی جاری	۱۷
		نسبت سرمایه در گردش به بدهی‌های جاری	۱۸

منبع: نتایج تحقیق

جدول ۳. متغیرهای کیفی و مقادیر آنها

مقدار	متغیر
(۰/۴، ۰/۶، ۰/۸، ۱)	X1: وضعیت وام (تسویه شده، معوق، سررسید گذشته، تمدید مدت، استمهال شده)
(۰/۲)	
(۰، ۰/۵، ۱)	X30: وضعیت مالکیت ملک (سرقفلی، دارای سند مالکیت غیر قابل ترهین، اجاره)
(۰، ۱)	X31: ارتباط بین فعالیت و مدرک (مرتبط، غیرمرتبط)

X32: نوع وثیقه (ملک (غیرمنقول)، سپرده بلندمدت، اوراق مشارکت، چک و سایر)	(۱، ۰/۷۵، ۰/۵، ۰/۲۵)
---	----------------------

منبع: نتایج تحقیق

۲.۲. روش پیشنهادی

در این پژوهش، به منظور تحلیل داده‌ها از مدلی مبتنی بر یادگیری عمیق حساس به هزینه استفاده شده است که با بهره‌گیری از الگوریتم فراابتکاری ژنتیک (GA) بهینه‌سازی شده است. در این روش، ترکیبی از مدل‌های پایه شامل LSTM، AdaBoost، و Bagging به‌صورت موازی آموزش دیده و خروجی آن‌ها به یک متاکلاسیفایر مبتنی بر جنگل تصادفی منتقل می‌شود و با در نظر گرفتن خطای دسته‌بندی (MC) به عنوان تابع هدف، بهینه‌سازی صورت می‌گیرد. شماتیک روش پیشنهادی که در شکل ۱ ارائه شده است، شامل چهار مرحله می‌باشد؛ به‌طوری که سه مرحله نخست به‌صورت یکپارچه اجرا می‌شوند: نرمال‌سازی داده‌ها، انتخاب ویژگی و کلاسه‌بندی، و در نهایت مرحله بهینه‌سازی با استفاده از الگوریتم GA.

نرمال‌سازی داده

در الگوریتم‌های خوشه‌بندی و دسته‌بندی، مقیاس‌بندی متغیرها به‌عنوان نخستین و مهم‌ترین گام در پردازش داده‌ها محسوب می‌شود. استفاده از داده‌های خام بدون اعمال مقیاس مناسب می‌تواند منجر به تأثیر نامتعادل متغیرهایی با مقادیر بزرگ‌تر نسبت به سایر ویژگی‌ها شود که این امر عملکرد الگوریتم‌ها را تحت تأثیر قرار می‌دهد. یکی از روش‌های رایج برای مقیاس‌بندی داده‌ها، نرمال‌سازی است؛ فرآیندی که در آن تمامی مقادیر متغیرها به بازه‌ای مشخص مانند ۰ تا ۱ تبدیل می‌شوند. در این پژوهش، نرمال‌سازی با استفاده از روش min-max انجام می‌شود که فرمول آن به صورت زیر تعریف می‌گردد (سوکسامای و همکاران، ۲۰۲۲):

$$X = \frac{X - \min(X)}{\max(X) - \min(X)} \quad (1)$$

که در آن، X بردار مربوط به داده‌های اعتباری است.

انتخاب ویژگی بهینه

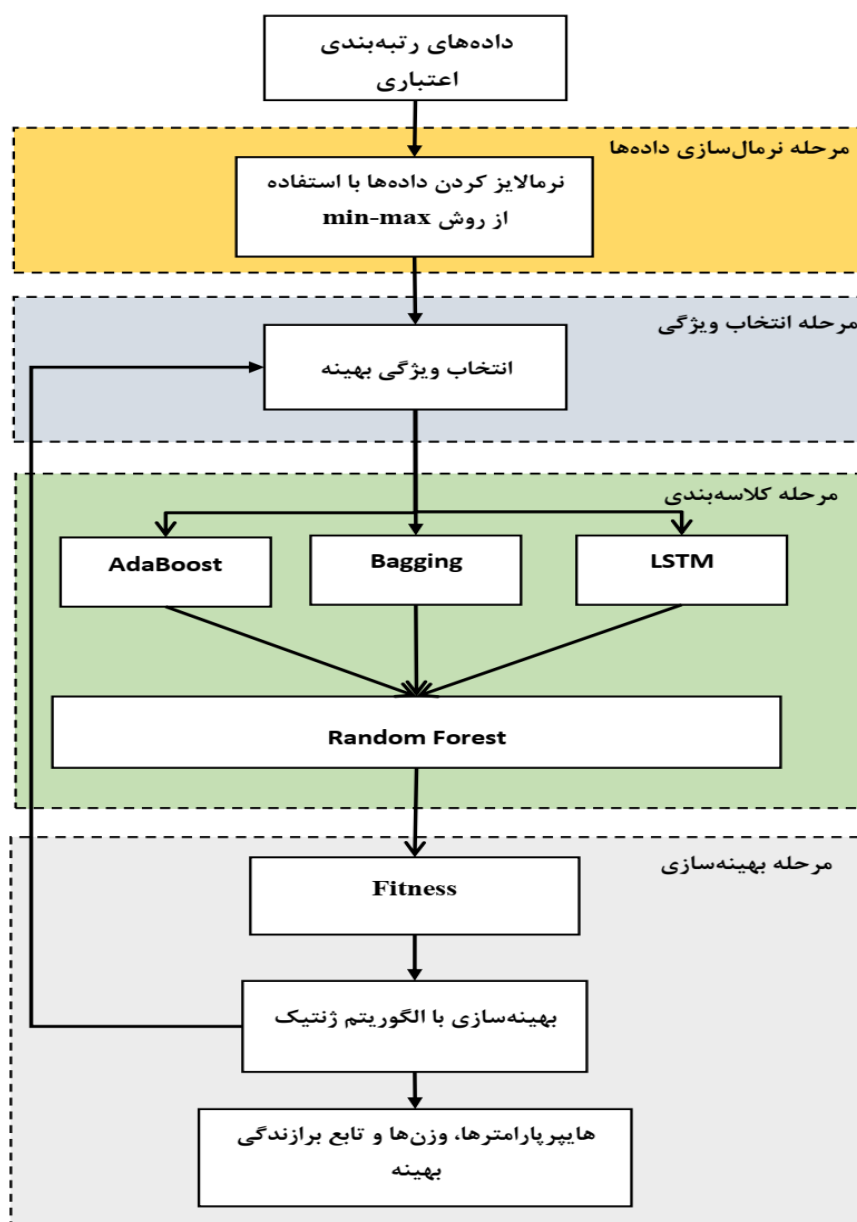
در مسائل کلاسه‌بندی، ممکن است تمام متغیرهای اعتباری به‌طور یکسان در فرایند یادگیری تأثیرگذار نباشند و حتی می‌توانند عملکرد الگوریتم را مختل کنند. برای رفع این مشکل، اجرای یک فرآیند انتخاب ویژگی ضروری است که می‌توان با استفاده از الگوریتم‌های بهینه‌سازی مانند الگوریتم ژنتیک (GA) به آن دست یافت. در این رویکرد، هر ویژگی به صورت دودویی کدگذاری می‌شود؛ به این معنا که ویژگی مورد نظر یا انتخاب شده (۱) یا انتخاب نمی‌شود (۰). در نهایت، ویژگی‌هایی که برای آن‌ها مقدار ۱ تعیین شده باشد، به عنوان ویژگی‌های منتخب در نظر گرفته می‌شوند و سایر ویژگی‌ها حذف می‌گردند.

کلاسه‌بندی بهینه

در روش پیشنهادی، از چارچوب متاکلاسیفایر به‌عنوان استراتژی اصلی طبقه‌بندی بهره گرفته می‌شود. در این رویکرد، سه الگوریتم پایه LSTM، Boosting و Bagging به‌عنوان طبقه‌بندی‌کننده‌های اولیه به کار گرفته می‌شوند. پیش‌بینی‌های حاصل از این الگوریتم‌ها سپس توسط الگوریتم متاکلاسیفایر جنگل تصادفی تجمیع می‌شود. این الگوریتم متا با تخصیص وزن‌های بهینه بر اساس دقت هر یک از الگوریتم‌های پایه، هم‌زمان مسئله داده‌های نامتوازن را حل کرده و اطمینان حاصل می‌کند که کلاس‌های دارای تعداد نمونه کمتر نیز به درستی شناسایی شوند؛ بدین ترتیب خطای طبقه‌بندی (MC) برای کلاس‌های کمتر نمونه به حداقل می‌رسد.

علاوه بر این، برای ارزیابی عملکرد متاکلاسیفایر، از روش اعتبارسنجی متقابل^۱ k-fold استفاده می‌شود. طبق این روش، داده‌ها به ۵ بخش مساوی ($k=5$) به‌صورت تصادفی تقسیم شده و در هر دور، یک بخش به‌عنوان مجموعه تست و چهار بخش به‌عنوان مجموعه آموزش مورد استفاده قرار می‌گیرند. در نهایت، میانگین خطای دسته‌بندی به‌دست آمده از هر یک از این تقسیم‌بندی‌ها به عنوان شاخص نهایی (MC) در نظر گرفته می‌شود.

^۱ Cross-validation



شکل ۱. شماتیک روش پیشنهادی

منبع: محقق ساخته

بهینه‌سازی یکپارچه با الگوریتم GA

در این پژوهش، به‌طور یکپارچه، هایپرپارامترهای هر الگوریتم به همراه وزن‌های بهینه تخصیصی و فرآیند انتخاب ویژگی بهینه با استفاده از الگوریتم GA تنظیم می‌شوند. تابع برازندگی به عنوان شاخص عملکرد، میزان MC را محاسبه می‌کند؛ به‌طوری که این مقدار برابر است با نسبت مجموع خطاهای طبقه‌بندی کلاس منفی (مشتری بدحساب) و کلاس مثبت (مشتری خوش حساب) به کل تعداد نمونه‌های مجموعه تست:

$$MC = \frac{FN + FP}{TP + TN + FP + FN} \quad (2)$$

که تعریف هر یک از متغیرهای فرمول فوق مطابق با جدول ۴ است.

جدول ۴. ماتریس درهم ریختگی

کلاس واقعی	کلاس پیش‌بینی شده	
	مشتریان بد	مشتریان خوب
مشتریان بد	منفی درست ^۲ (TN)	مثبت کاذب ^۱ (FP)
مشتریان خوب	منفی کاذب ^۴ (FN)	مثبت درست ^۳ (TP)

منبع: نتایج تحقیق

فلوچارت الگوریتم GA در شکل ۲ نمایش داده شده است؛ بنابراین مسئله بهینه‌سازی مورد بررسی در پژوهش حاضر به صورت زیر تعریف می‌شود:

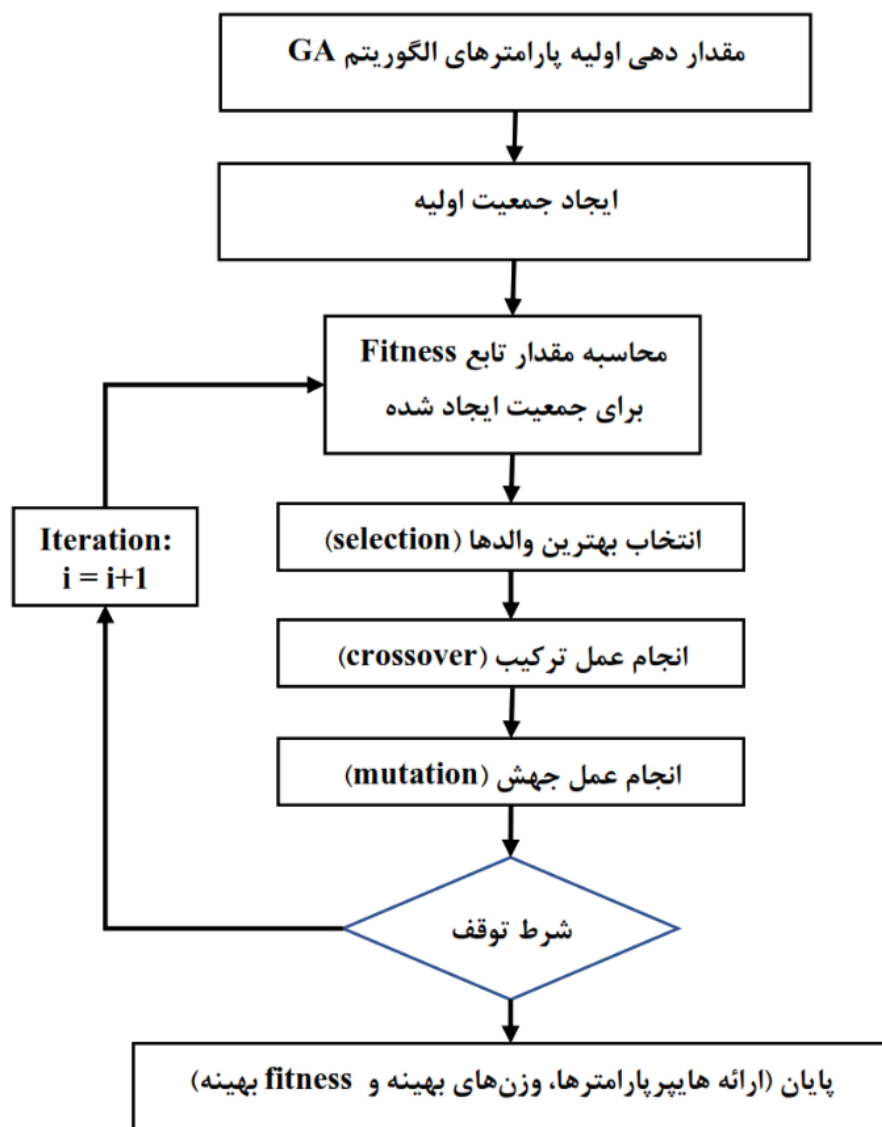
$$Fitenss = \min MC \quad (3)$$

1 False positive

2 True negative

3 True positive

4 False negative



شکل ۲. فلوچارت الگوریتم GA

منبع: نتایج تحقیق

ارزیابی عملکرد روش‌های کلاسه‌بندی از طریق معیارهای متنوعی انجام می‌شود. در این پژوهش، هفت معیار متفاوت مورد استفاده قرار گرفته‌اند که عبارتند از:

$$\text{Precision} = \frac{TP}{TP+FP} \quad (۴)$$

$$\text{ACC} = \frac{TP}{TP+TN+FP+FN} \quad (۵)$$

$$\text{Recall or Sensitivity} = \frac{TP}{TP+FN} \quad (۶)$$

$$\text{F-Score} = \frac{2 \times \text{Recall} \times \text{Precision}}{(\text{Recall} + \text{Precision})} \quad (۷)$$

$$\text{Error}_{\text{Type I}} = \frac{FP}{TN+FP} \quad (۸)$$

$$\text{Error}_{\text{Type II}} = \frac{FN}{FN+TP} \quad (۹)$$

$$\text{MC} = \frac{FN+FP}{TP+TN+FP+FN} \quad (۱۰)$$

۳. یافته‌های پژوهش

در این بخش نتایج به‌دست‌آمده از روش پیشنهادی بررسی می‌شود. همانطور که قبلاً اشاره شد، رویکرد ارائه‌شده یک روش بهینه‌سازی یکپارچه است که با استفاده از تکنیک‌های متعدد و انتخاب هوشمند ویژگی‌ها، به ارزیابی اعتبار مشتریان می‌پردازد. در این پژوهش، چهار الگوریتم اصلی شامل AbaBoost، Bagging، RF و LSTM مورد استفاده قرار گرفته‌اند و بهینه‌سازی هایپرپارامترهای مربوط به هر یک از آن‌ها از طریق الگوریتم ژنتیک انجام می‌شود. جدول ۵، هایپرپارامترهای مدل به همراه فضای جستجو (کران پایین و کران بالا) را نمایش می‌دهد و همچنین جدول ۶، پارامترهای مربوط به الگوریتم ژنتیک را نشان می‌دهد. در فرآیند انتخاب ویژگی، مقدار ۰ نشان‌دهنده عدم انتخاب و مقدار ۱ نمایانگر انتخاب ویژگی می‌باشد. علاوه بر این، سه هایپرپارامتر وزنی نیز برای تخصیص وزن به الگوریتم‌ها در نظر گرفته شده است؛ به‌گونه‌ای که هر کدام مقداری بین ۰ تا ۱ دارند و مجموع این سه پارامتر برابر با ۱ می‌شود. از این رو، با احتساب ۱۱ هایپرپارامتر ذکرشده در جدول ۵، سه پارامتر وزنی و تعداد ویژگی‌های موجود در دیتاست (n)، تعداد کل هایپرپارامترهایی که باید از طریق بهینه‌سازی با الگوریتم GA تنظیم شوند، برابر با ۱۴ + n خواهد بود.

نتایج تجربی حاکی از آن است که استفاده از الگوریتم ژنتیک برای بهینه‌سازی همزمان هایپرپارامترهای چندین الگوریتم طبقه‌بندی از جمله Random، Bagging، AdaBoost

Forest و LSTM، تاثیر بسزایی در بهبود عملکرد نهایی سیستم داشته است. همانطور که در شکل ۳ نشان داده شده، فرآیند همگرایی در الگوریتم ژنتیک بسیار سریع اتفاق می افتد؛ به طوری که در نسل ۲۵ به نقطه همگرایی مطلوب دست یافته ایم. این امر، علاوه بر تأیید سازگاری بالای پارامترهای تنظیم شده، نشان دهنده فضای جستجوی بهینه تر در فضای هایپرپارامتر است که می تواند منجر به عملکرد دقیق تر و کارآمدتر سیستم طبقه بندی گردد. از جنبه عملکرد، نتایج تجربی بر روی دیتاست ایران که در جدول ۷ ارائه شده اند، بسیار چشمگیر هستند. میانگین دقت (ACC) به مقدار ۰,۹۴۸، همراه با Precision برابر با ۰,۹۶۳، Recall معادل ۰,۹۷۱ و F1-Score به میزان ۰,۹۶۷ نشانگر توانمندی بالای سیستم در تشخیص الگوهای پنهان موجود در داده های اعتباری است. این ارقام، علاوه بر نشان دادن صحت عملکرد مدل، بیانگر تعادل مطلوب بین دسته بندی صحیح نمونه های مثبت (مشتریان خوش حساب) و منفی (مشتریان بد حساب) هستند. کاهش خطاهای نوع I و II به ترتیب به ۰,۰۴۸ و ۰,۰۳۴، و همچنین رسیدن میزان کلی خطای طبقه بندی (MC) به ۰,۰۵۲، شواهدی بر بهبود قابل توجه عملکرد مدل در مواجهه با داده های نامتوازن ارائه می دهد.

یکی از جنبه های کلیدی رویکرد پیشنهادی، بهینه سازی دقیق هایپرپارامترهای الگوریتم های استفاده شده است. تنظیمات بهینه در الگوریتم AdaBoost نشان می دهد که نرخ یادگیری به ۰,۱۹۳ و تعداد تخمین گر ها به ۲۳۳ تعیین شده اند؛ در حالی که در مدل LSTM، تنظیمات به گونه ای بهینه سازی شده اند که تعداد لایه های پنهان برابر با ۲۵، نرخ dropout برابر با ۰,۲۴۷، نرخ یادگیری ۰,۱۴۳ و تعداد epoch برابر با ۸۶ در نظر گرفته شده است. این تنظیمات به صورت دقیق و با استفاده از الگوریتم ژنتیک به دست آمده اند، که جزئیات آن ها در جدول ۵ و جدول مربوط به تنظیمات الگوریتم ژنتیک (جدول ۶) به تفصیل آمده است. در الگوریتم های Bagging و مدل متای RF نیز پارامترهایی مانند تعداد تخمین گر ها و عمق درخت ها به گونه ای بهینه شده اند که در نهایت عملکرد سیستم را بهبود بخشند. ارزیابی عملکرد مدل از منظر شاخص های آماری و معیارهای دقیق نیز از اهمیت ویژه ای برخوردار است. به عنوان نمونه، مقدار شاخص AUC برابر ۰,۹۸، نشان می دهد که مدل متا در استخراج الگوهای پیچیده و تفکیک دقیق میان کلاس های مختلف، عملکرد برتری دارد (شکل ۴). این مقدار قابل ملاحظه در AUC، که به عنوان یکی از شاخص های مهم در ارزیابی عملکرد سیستم های طبقه بندی محسوب می شود، تأکید بر توانایی سیستم در تشخیص صحیح نمونه های مثبت و منفی در محیط های واقعی دارد. همچنین، استفاده از تکنیک اعتبارسنجی متقابل k-fold (که در این پژوهش به صورت ۵-fold پیاده سازی شده است) به عنوان یک روش استاندارد برای تقسیم بندی داده ها، امکان اطمینان از تعمیم پذیری مدل را فراهم آورده است. نتایج حاصل از این اعتبارسنجی متقابل نشان می دهد که عملکرد مدل در چندین تقسیم بندی داده ای مختلف ثابت

و پایدار است. به عبارت دیگر، میانگین عملکرد در پنج دسته‌بندی مختلف که در جدول ۷ آمده است، از ثبات و صحت مدل در شرایط مختلف سخن می‌گوید.

جدول ۵. هایپر پارامترهای الگوریتم‌ها

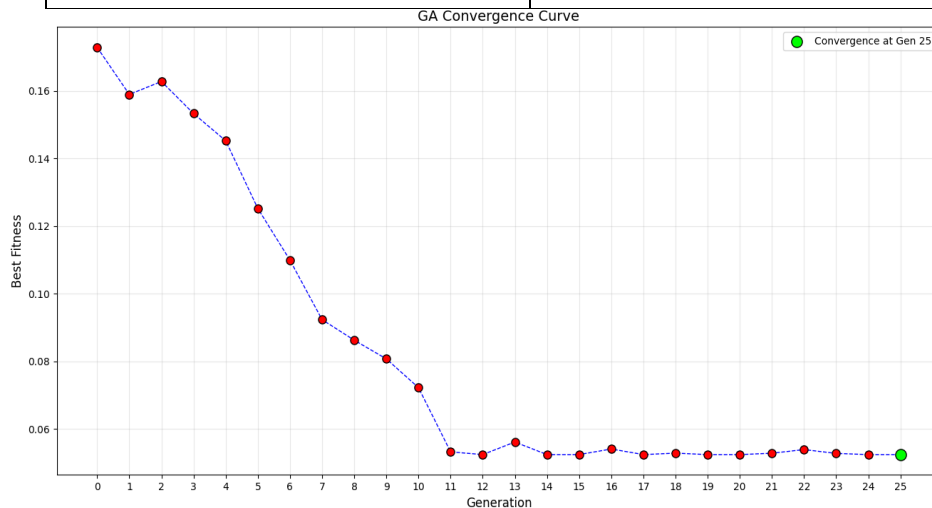
الگوریتم	پارامتر	LB	UB	مقدار بهینه
AdaBoost	نرخ یادگیری	۰,۰۱	۱,۰	۰,۱۹۳
	تعداد تخمین‌گرها	۵۰	۵۰۰	۲۳۳
	ماکزیمم عمق	۱	۱۰	۶
Bagging	تعداد تخمین‌گرها	۵۰	۵۰۰	۳۰۹
	ماکزیمم عمق	۱	۱۰	۵
	نرخ ماکزیمم تعداد سمپل	۰,۰۱	۱,۰	۰,۵۳۱
	نرخ ماکزیمم تعداد ویژگی	۰,۰۱	۱,۰	۰,۴۲۰
RF	تعداد تخمین‌گرها	۵۰	۵۰۰	۴۱۴
	ماکزیمم عمق	۱	۱۰	۵
	نرخ ماکزیمم تعداد سمپل	۰,۰۱	۱,۰	۰,۶۹۴
	نرخ ماکزیمم تعداد ویژگی	۰,۰۱	۱,۰	۰,۲۳۷
LSTM	تعداد لایه‌های پنهان	۱۰	۵۰	۲۵
	نرخ dropout	۰,۰۱	۰,۵	۰,۲۴۷
	نرخ یادگیری	۰,۰۱	۱,۰	۰,۱۴۳
	تعداد epoch	۱۰	۱۰۰	۸۶

منبع: نتایج تحقیق

جدول ۶. پارامترهای الگوریتم ژنتیک

پارامتر	مقدار
اندازه جمعیت	۳۰
نرخ جهش	۰,۲
نرخ ترکیب	۰,۸
عملگر انتخاب	Tournament selection
حداکثر تعداد نسل	۳۰

تلورانس	1e-5
معیار توقف	رسیدن به تلورانس مدنظر یا حداکثر تعداد نسل



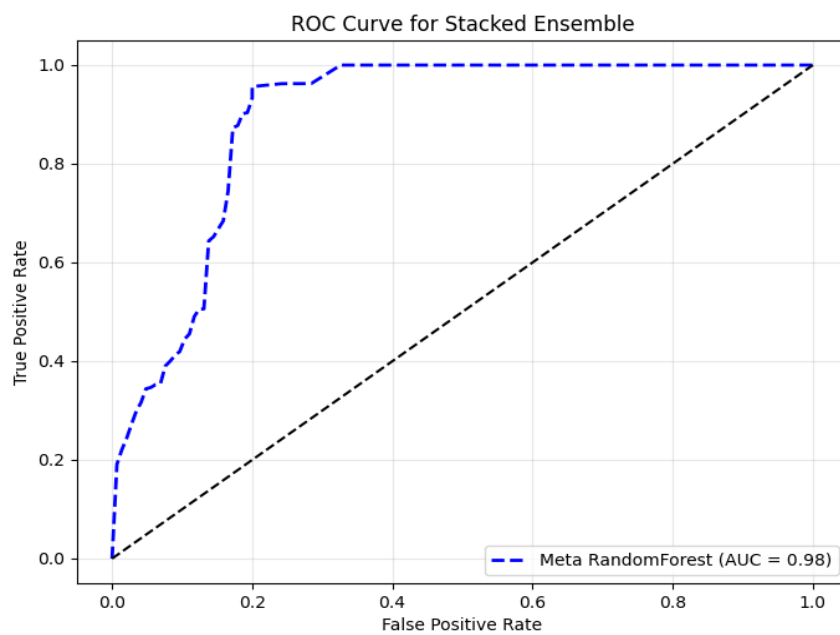
شکل ۳. نمودار همگرایی مدل پیشنهادی برای دیتای ایران

منبع: خروجی نرم افزار

جدول ۷. عملکرد مدل ترکیبی برای دیتای ایران

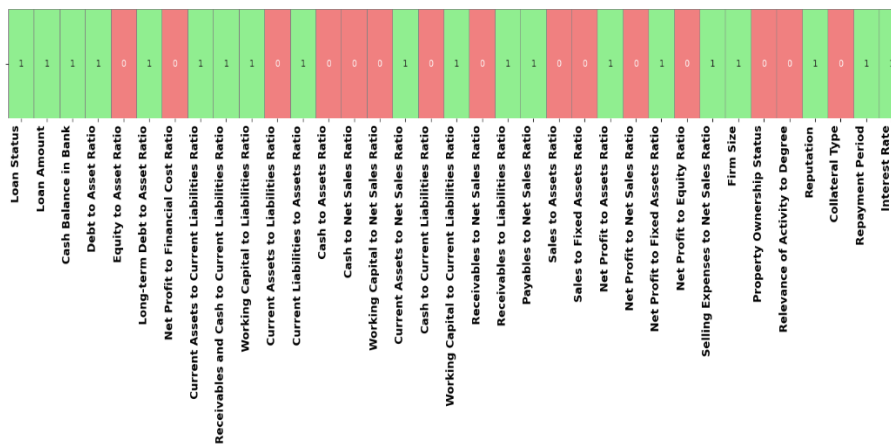
MC	E-Type II	E-Type I	F1-Score	Recall	Precision	ACC	Fold
0.055	0.016	0.017	0.978	0.977	0.978	0.945	1
0.030	0.017	0.015	0.973	0.979	0.968	0.970	2
0.052	0.009	0.105	0.958	0.981	0.936	0.948	3
0.067	0.060	0.086	0.961	0.964	0.958	0.933	4
0.055	0.065	0.017	0.964	0.954	0.974	0.945	5
0.052	0.034	0.048	0.967	0.971	0.963	0.948	Mean

منبع: نتایج تحقیق



شکل ۴. نمودار ROC مدل پیشنهادی با رویکرد دوم برای دیتای ایران

منبع: خروجی نرم افزار



شکل ۵. ویژگی‌های بهینه انتخاب شده مدل پیشنهادی برای دیتای ایران

منبع: نتایج تحقیق

در انتخاب ویژگی‌های بهینه برای رتبه‌بندی اعتباری، همانطور که در شکل ۵ نمایش داده شده است، مدل بهینه پیشنهادی تمرکز ویژه‌ای بر روی استخراج شاخص‌هایی دارد که توانایی تمایز الگوهای پیچیده اعتباری را در ابعاد مختلف اقتصادی نشان می‌دهند. در این مدل، بر مبنای تحلیل‌های آماری و ارزیابی‌های تجربی، مجموعه‌ای از ویژگی‌های کلیدی به‌عنوان عوامل تعیین‌کننده در نظر گرفته شده است. این ویژگی‌ها عبارتند از:

این انتخاب ویژگی‌های بهینه نشان می‌دهد که مدل پیشنهادی به عوامل اقتصادی مرتبط با نقدینگی، سودآوری و شرایط بازپرداخت توجه ویژه‌ای دارد. از آنجایی که اطلاعات مربوط به موجودی نقد، نسبت‌های مرتبط با دارایی و بدهی و همچنین شاخص‌های مرتبط با سودآوری و بازپرداخت، در تمایز الگوهای اعتباری و تعیین ریسک اعتباری نقش کلیدی دارند، این مدل قادر است به شکل دقیق‌تری الگوهای پیچیده را در داده‌های اعتباری شناسایی کند. به عبارت دیگر، انتخاب این ویژگی‌ها نه تنها به بهبود دقت طبقه‌بندی کمک می‌کند، بلکه نشان‌دهنده حساسیت بیشتر مدل پیشنهادی به عواملی است که در واقع منعکس‌کننده سلامت مالی و پتانسیل بازپرداخت مشتریان می‌باشند. در نتیجه، استفاده از شاخص‌هایی مانند موجودی نقد در بانک، نسبت دارایی جاری به بدهی جاری، نسبت سرمایه در گردش به بدهی‌ها، نسبت سود خالص به دارایی‌ها و نرخ سود، ابعاد گسترده‌ای از نقدینگی و سودآوری را پوشش می‌دهد که در نهایت به بهبود عملکرد مدل در رتبه‌بندی اعتباری کمک می‌کند.

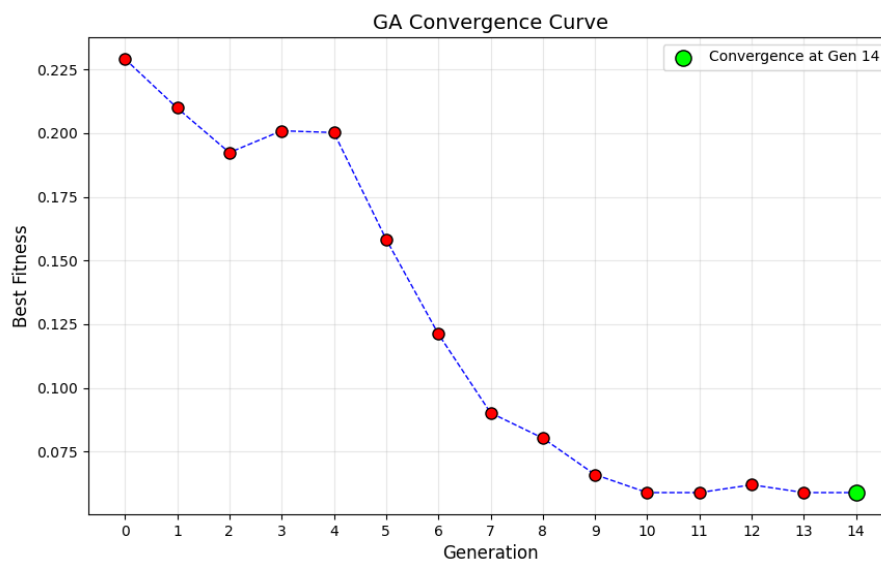
علاوه بر این، ترکیب ویژگی‌های انتخاب‌شده در این مدل، به مدل امکان می‌دهد تا با تمرکز بر جنبه‌های کلیدی مانند وضعیت وام، میزان وام، نسبت‌های مالی مهم و شاخص‌های سودآوری، به شیوه‌ای جامع و چندبعدی به تحلیل ریسک اعتباری بپردازد. این امر نه تنها منجر به افزایش دقت پیش‌بینی مدل می‌شود، بلکه امکان شناسایی بهتر نقاط ضعف و قوت مشتریان را فراهم می‌آورد. بطور کلی، مدل پیشنهادی با انتخاب هوشمندانه ویژگی‌هایی که به طور مستقیم با پتانسیل بازپرداخت و سلامت مالی مشتریان مرتبط هستند، ابزاری قوی برای مدیریت ریسک اعتباری فراهم می‌آورد. این انتخاب به مدل اجازه می‌دهد تا بتواند به شیوه‌ای دقیق و موثر، الگوهای پیچیده موجود در داده‌های اعتباری را استخراج کرده و بهبود قابل‌توجهی در عملکرد طبقه‌بندی ایجاد کند. چنین رویکردی، می‌تواند به عنوان یک چارچوب نوین در توسعه سیستم‌های تصمیم‌گیری هوشمند در حوزه بانکداری و مدیریت مالی مورد استفاده قرار گیرد.

به منظور اعتبارسنجی نتایج به‌دست‌آمده از تحلیل دیتای ایران و اطمینان از تعمیم‌پذیری مدل، دو دیتای بنچمارک از استرالیا و آلمان به کار گرفته شده‌اند. استفاده از این دو مجموعه داده به پژوهشگران امکان مقایسه و ارزیابی عملکرد مدل در شرایط متفاوت و با ویژگی‌های آماری گوناگون

را می‌دهد، که این امر می‌تواند ارزش افزوده رویکرد بهینه‌سازی یکپارچه و تنظیم دقیق هایپرپارامترها را در بهبود عملکرد سیستم رتبه‌بندی اعتباری برجسته نماید.

در دیتای استرالیا، نتایج نشان می‌دهد که الگوریتم ژنتیک در نسل ۱۴ به نقطه همگرایی رسیده است (شکل ۶)، که این همگرایی سریع می‌تواند به عنوان نشانه‌ای از فضای جستجوی بهینه و سازگاری بالای پارامترهای تنظیم‌شده تلقی شود. از نظر عملکرد، مدل ترکیبی در این مجموعه داده عملکرد بسیار مطلوبی از خود نشان داده است؛ میانگین دقت (ACC) برابر با ۰,۹۴۱، Precision برابر با ۰,۹۴۶، Recall معادل ۰,۹۵۶ و F1-Score برابر با ۰,۹۵۱ گزارش شده است. همچنین، خطاهای نوع I و II به ترتیب برابر ۰,۰۴۵ و ۰,۰۵۵ است و میزان کلی خطای طبقه‌بندی (MC) به ۰,۰۵۹ می‌رسد (جدول ۸). به علاوه، تنظیمات بهینه برای هر یک از الگوریتم‌های مورد استفاده تغییر یافته است؛ برای مثال، در الگوریتم AdaBoost نرخ یادگیری به ۰,۰۹۲ و تعداد تخمین‌گرها به ۴۶۹ تنظیم شده‌اند، در حالی که در مدل LSTM تعداد لایه‌های پنهان به ۳۸، نرخ dropout به ۰,۲۰۳، نرخ یادگیری به ۰,۱۱۷ و تعداد epoch به ۱۰۱ تعیین شده‌اند (جدول ۹). از سوی دیگر، بهبود قابل‌ملاحظه شاخص AUC به مقدار ۰,۰۹۸، که در نمودار ROC مشاهده می‌شود، بیانگر توانایی مدل متا در استخراج الگوهای پیچیده و بهبود تشخیص نهایی است (شکل ۷).

در دیتای آلمان، الگوریتم ژنتیک در نسل ۳۷ به همگرایی رسیده است (شکل ۸)؛ اگرچه این همگرایی نسبت به دیتای استرالیا کندتر اتفاق می‌افتد، اما همچنان بیانگر فضای جستجوی بهینه و تطبیق دقیق‌تر پارامترها می‌باشد. نتایج عملکرد مدل ترکیبی در این مجموعه داده نشان‌دهنده میانگین دقت (ACC) برابر با ۰,۸۵۵، Precision برابر با ۰,۹۰۹، Recall معادل ۰,۹۶۰ و F1-Score برابر با ۰,۹۳۴ می‌باشد. همچنین، نرخ خطای نوع I به ۰,۱۰۵ و خطای نوع II به ۰,۰۴۵ رسیده است و هزینه اشتباه کلی (MC) به ۰,۱۴۵ می‌رسد (۱۰). تنظیمات بهینه که در جدول ۱۱ نشان داده شده است نیز به صورت دقیق انجام شده است؛ به عنوان نمونه، در الگوریتم AdaBoost نرخ یادگیری به مقدار ۰,۰۱۷ و تعداد تخمین‌گرها به ۸۰ تنظیم شده‌اند؛ در الگوریتم Bagging تعداد تخمین‌گرها ۴۱۲، ماکزیمم عمق ۹، نرخ ماکزیمم تعداد سمپل ۰,۰۸۸ و نرخ ماکزیمم تعداد ویژگی ۰,۱۷۶ تعیین شده‌اند؛ همچنین در الگوریتم RF تعداد تخمین‌گرها ۴۴۴، ماکزیمم عمق ۶، نرخ ماکزیمم تعداد سمپل ۰,۷۲۳ و نرخ ماکزیمم تعداد ویژگی ۰,۰۵۶ تنظیم شده است؛ در نهایت، برای مدل LSTM تعداد لایه‌های پنهان به ۱۴، نرخ dropout به ۰,۱۱۰، نرخ یادگیری به ۰,۱۴۳ و تعداد epoch برابر با ۲۱۹ انتخاب شده‌اند. نمودار ROC در این رویکرد نشان می‌دهد که عملکرد مدل ترکیبی با AUC نزدیک به ۰,۹۱ بهبود یافته است؛ این افزایش در AUC بیانگر توانایی مدل متا در استخراج روابط پیچیده و بهبود تشخیص نهایی است (شکل ۹).



شکل ۶. نمودار همگرایی مدل پیشنهادی برای دیتای استرالیا

منبع: خروجی نرم افزار

جدول ۸. عملکرد مدل ترکیبی در رویکرد دوم برای دیتای استرالیا

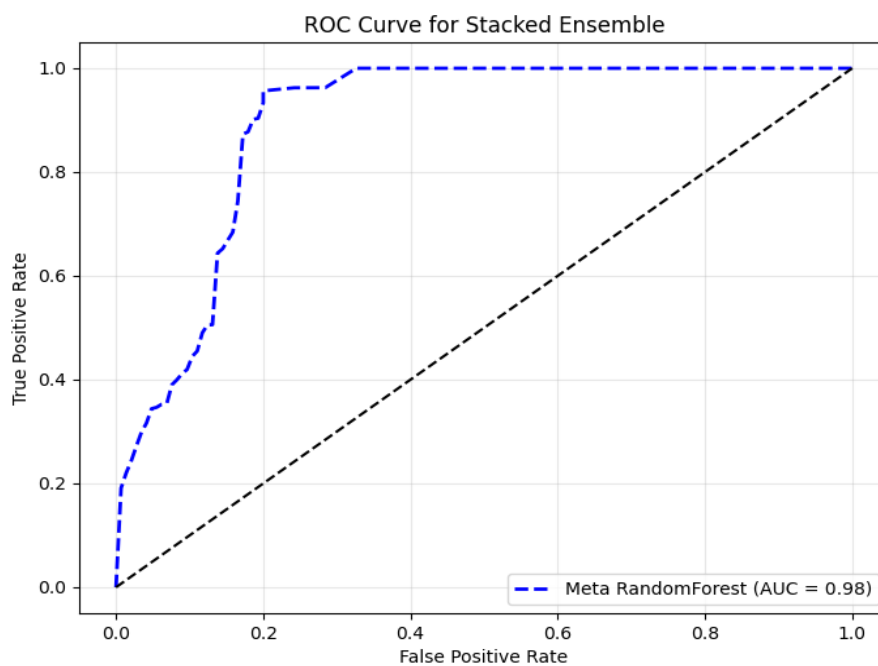
MC	E-Type II	E-Type I	F1-Score	Recall I	Precision	ACC	Fold
0.056	0.015	0.065	0.957	0.987	0.928	0.944	1
0.063	0.018	0.056	0.957	0.972	0.943	0.937	2
0.073	0.092	0.008	0.962	0.938	0.988	0.927	3
0.084	0.130	0.086	0.898	0.900	0.897	0.916	4
0.018	0.018	0.011	0.979	0.985	0.972	0.982	5
0.059	0.055	0.045	0.951	0.956	0.946	0.941	Mean

منبع: نتایج تحقیق

جدول ۹. مقادیر بهینه های پیرامترهای مدل پیشنهادی برای دیتاست استراليا

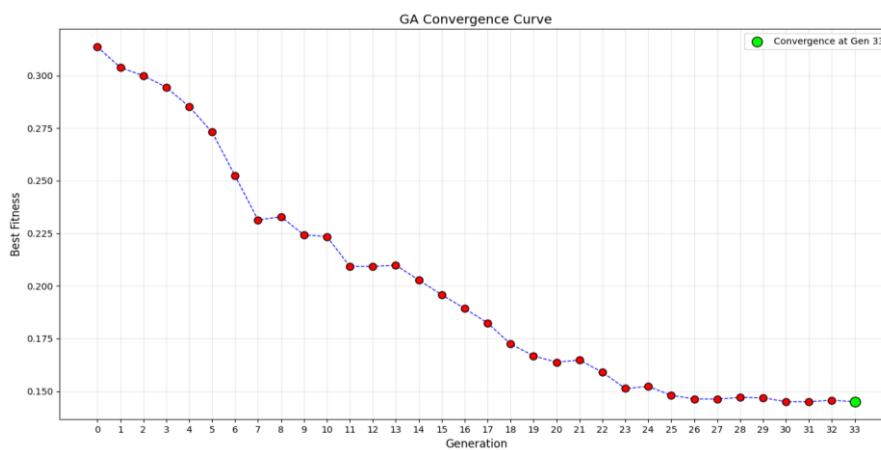
الگوریتم	پارامتر	LB	UB	مقدار بهینه
AdaBoost	نرخ یادگیری	۰,۰۱	۱,۰	۰,۰۹۲
	تعداد تخمین گرها	۵۰	۵۰۰	۴۶۹
	ماکزیمم عمق	۱	۱۰	۸
Bagging	تعداد تخمین گرها	۵۰	۵۰۰	۲۳۹
	ماکزیمم عمق	۱	۱۰	۴
	نرخ ماکزیمم تعداد سمپل	۰,۰۱	۱,۰	۰,۴۱۳
	نرخ ماکزیمم تعداد ویژگی	۰,۰۱	۱,۰	۰,۷۵۷
RF	تعداد تخمین گرها	۵۰	۵۰۰	۳۷۲
	ماکزیمم عمق	۱	۱۰	۵
	نرخ ماکزیمم تعداد سمپل	۰,۰۱	۱,۰	۰,۵۷۹
	نرخ ماکزیمم تعداد ویژگی	۰,۰۱	۱,۰	۰,۳۷۹
LSTM	تعداد لایه های پنهان	۱۰	۵۰	۳۸
	نرخ dropout	۰,۰۱	۰,۵	۰,۲۰۳
	نرخ یادگیری	۰,۰۱	۱,۰	۰,۱۱۷
	تعداد epoch	۱۰	۱۰۰	۱۰۱

منبع: نتایج تحقیق



شکل ۷. نمودار ROC مدل پیشنهادی با رویکرد دوم برای دیتای استرالیا

منبع: خروجی نرم افزار



شکل ۸. نمودار همگرایی مدل پیشنهادی با رویکرد دوم برای دیتای آلمان

منبع: خروجی نرم افزار

جدول ۱۰. عملکرد مدل ترکیبی در رویکرد دوم برای دیتای آلمان

MC	E-Type II	E-Type I	F1-Score	Recall	Precision	ACC	Fold
0.139	0.036	0.069	0.955	0.974	0.937	0.861	1
0.164	0.071	0.114	0.914	0.929	0.899	0.836	2
0.143	0.049	0.124	0.925	0.954	0.898	0.857	3
0.135	0.042	0.086	0.944	0.961	0.926	0.865	4
0.143	0.026	0.131	0.932	0.984	0.886	0.857	5
0.145	0.045	0.105	0.934	0.960	0.909	0.855	Mean

منبع: نتایج تحقیق

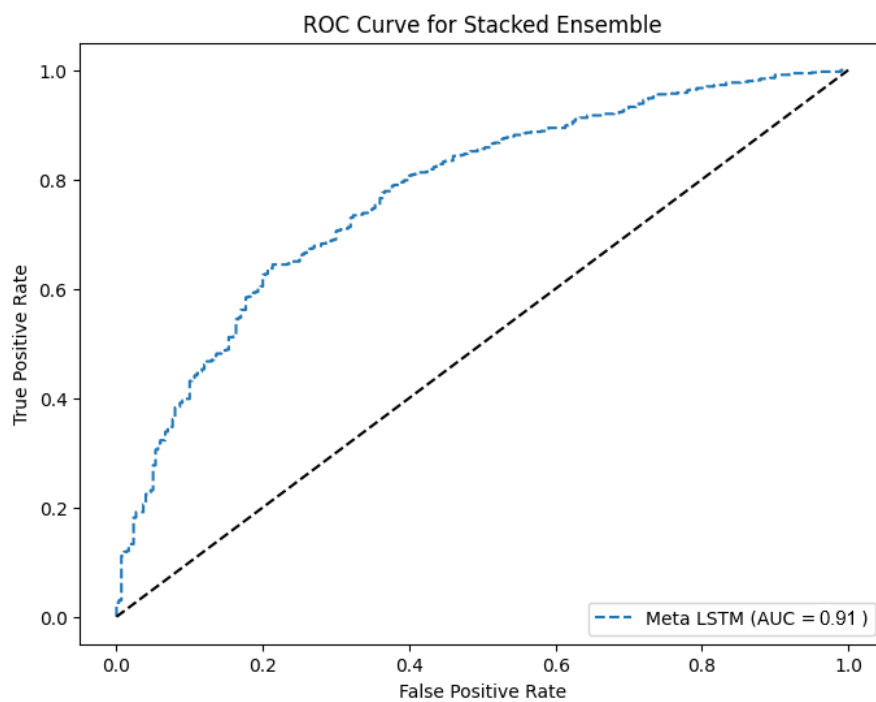
در مجموع، نتایج حاصل از تحلیل‌های انجام‌شده بر روی دیتای استرالیا و آلمان به عنوان دو دیتای بنچمارک، عملکرد مطلوب و بهبود قابل توجه سیستم رتبه‌بندی اعتباری را در شرایط مختلف تأیید می‌کنند. این دستاوردها نشان‌دهنده اثربخشی رویکرد بهینه‌سازی یکپارچه با تنظیم دقیق هایپرپارامترها و استفاده از الگوریتم‌های بهینه‌سازی فراابتکاری می‌باشند و می‌توانند مبنایی قوی برای توسعه مدل‌های پیشرفته در زمینه مدیریت ریسک اعتباری و سیستم‌های تصمیم‌گیری هوشمند در حوزه بانکداری محسوب شوند.

جدول ۱۱. مقادیر بهینه هایپر پارامترهای مدل رویکرد دوم برای دیتاست آلمان

مقدار بهینه	UB	LB	پارامتر	الگوریتم
۰,۰۱۷	۱,۰	۰,۰۱	نرخ یادگیری	AdaBoost
۸۰	۵۰۰	۵۰	تعداد تخمین‌گرها	
۳	۱۰	۱	ماکزیمم عمق	
۴۱۲	۵۰۰	۵۰	تعداد تخمین‌گرها	Bagging
۹	۱۰	۱	ماکزیمم عمق	
۰,۰۸۸	۱,۰	۰,۰۱	نرخ ماکزیمم تعداد سمپل	

۰,۱۷۶	۱,۰	۰,۰۱	نرخ ماکزیمم تعداد ویژگی	
۴۴۴	۵۰۰	۵۰	تعداد تخمین‌گرها	RF
۶	۱۰	۱	ماکزیمم عمق	
۰,۷۲۳	۱,۰	۰,۰۱	نرخ ماکزیمم تعداد سمپل	
۰,۰۵۶	۱,۰	۰,۰۱	نرخ ماکزیمم تعداد ویژگی	
۱۴	۵۰	۱۰	تعداد لایه‌های پنهان	
۰,۱۱۰	۰,۵	۰,۰۱	نرخ dropout	LSTM
۰,۱۴۳	۱,۰	۰,۰۱	نرخ یادگیری	
۲۱۹	۱۰۰	۱۰	تعداد epoch	

منبع: نتایج تحقیق



شکل ۹. نمودار ROC مدل پیشنهادی با رویکرد دوم برای دیتای آلمان

منبع: خروجی نرم افزار

به منظور ارزیابی بیشتر عملکرد مدل پیشنهادی تحقیق حاضر در مقایسه با روش‌های مطرح در ادبیات اخیر، نتایج به‌دست‌آمده برای دو دیتاست استرالیا و آلمان مورد بررسی قرار گرفت. در دیتاست استرالیا، روش پیشنهادی با استفاده از ترکیب خروجی‌های سه الگوریتم Bagging، AdaBoost و LSTM، که در مرحله اول وظیفه طبقه‌بندی را بر عهده داشته‌اند، به کمک یک مدل متا مبتنی بر Random Forest (بدون استفاده از تکنیک‌های نمونه‌برداری) به دقت (ACC) برابر با ۰,۹۴۱ دست یافته است؛ در حالی که در دیتاست آلمان، عملکرد مشابه با ACC برابر با ۰,۸۵۵ گزارش شده است. با مقایسه نتایج، روش‌های مبتنی بر مدل‌های رای‌گیری متعدد مانند NB-DT-QDA و TDNN-PNN در شرایطی که از نمونه‌برداری استفاده نمی‌شود، دقت‌های پایین‌تری (برای استرالیا بین ۰,۹۰۸ تا ۰,۹۳۰ و برای آلمان بین ۰,۷۹۴ تا ۰,۸۴۲) را نشان می‌دهند. همچنین، استفاده از تکنیک‌های نمونه‌برداری در مدل‌هایی نظیر XGBoost + SMOTE-ENN (ACC=0.910) برای استرالیا و ۰,۸۹۵ برای آلمان و Stacked + SMOTE (که تنها برای دیتاست آلمان گزارش شده با ACC=0.832) عملکرد بهتری نسبت به روش‌های بدون نمونه‌برداری ارائه می‌دهد؛ اما همچنان در مقایسه با مدل پیشنهادی، روش‌های مذکور عملکرد کمتری دارند. علاوه بر این، مطالعه‌ای از خلیلی و رستگار (۲۰۲۳) نیز مدل ترکیبی مبتنی بر stacking را بدون استفاده از نمونه‌برداری ارائه نموده و دقت ۰,۹۳۸ برای دیتاست استرالیا و ۰,۸۱۶ برای دیتاست آلمان گزارش کرده است؛ که این عملکرد نسبت به روش پیشنهادی کمی پایین‌تر است. از سویی دیگر، مطالعات اخیر مانند آرولبا و سان (۲۰۲۴) با استفاده از ترکیب مدل XGBoost و تکنیک‌های نمونه‌برداری، عملکرد قابل رقابتی را ارائه داده‌اند که در برخی موارد نزدیک به نتایج تحقیق حاضر هستند.

در مجموع، نتایج نشان می‌دهد که استفاده از مدل متا مبتنی بر RF در مدل پیشنهادی، با وجود عدم استفاده از تکنیک‌های نمونه‌برداری، توانسته است عملکرد نهایی سیستم طبقه‌بندی را در دو دیتاست استرالیا و آلمان بهبود بخشد. این نتایج از نظر دقت و کاهش هزینه‌های اشتباه (MC) به خوبی بیانگر توانایی مدل در استخراج روابط پیچیده و بهبود تصمیم‌گیری‌های اعتباری می‌باشد. از طرفی، تفاوت‌های مشاهده‌شده بین نتایج دیتاست آلمان و استرالیا ممکن است ناشی از تفاوت‌های ساختاری و ویژگی‌های داده‌ای موجود در این دو مجموعه باشد؛ به همین دلیل، بهبودهای آتی در تنظیم دقیق‌تر پارامترها و به‌کارگیری تکنیک‌های نمونه‌برداری متناسب با هر دیتاست می‌تواند زمینه ارتقای عملکرد را فراهم آورد.

جدول **Error! No text of specified style in document.** ۱: مقایسه نتایج مدل

پیشنهاد یا سایر منابع

دیتاست استرالیا	دیتاست آلمان	sampling	روش	رفرنس
0.909	0.794	×	NB-DT-QDA-TDNN-PNN (Majority Voting)	تریپاتی ^۱ و همکاران (۲۰۲۸)
0.909	0.794	×	NB-DT-QDA-TDNN-PNN (Weighted Voting)	
0.908	0.805	×	NB-DT-QDA-TDNN-PNN (Layered Majority Voting)	
0.93	0.842	×	NB-DT-QDA-TDNN-PNN (Layered Weighted Voting)	
0.938	0.816	×	Stacking	خلیلی و رستگار (۲۰۲۳)
0.910	0.895	✓	XGBoost + SMOTE-ENN	آرولبا و سان (۲۰۲۴)
-	0.832	✓	Stacking + SMOTE	روفیک ^۳ و همکاران (۲۰۲۴)
0.941	0.855	×	Stacking	روش پیشنهادی

منبع: نتایج تحقیق

۴. بحث و نتیجه‌گیری

در این پژوهش، مدل پیشنهادی مبتنی بر یادگیری جمعی و استفاده از مدل متا به‌عنوان رویکرد نوینی برای رتبه‌بندی اعتباری ارائه شده است. در شرایطی که رتبه‌بندی اعتباری به‌عنوان یکی از ارکان اساسی در مدیریت ریسک و تصمیم‌گیری‌های مالی، نقش حیاتی در کاهش خسارات ناشی از اعطای تسهیلات و بهبود تخصیص منابع دارد، استفاده از روش‌های سنتی مانند تحلیل عاملی یا مدل‌های بیزی به تنهایی قادر به استخراج الگوهای پیچیده موجود در داده‌های بانکی نبوده است. بنابراین، با توجه به پیچیدگی‌های روابط میان متغیرهای اقتصادی و اعتباری، رویکرد پیشنهادی با ترکیب الگوریتم‌های AdaBoost، Bagging و LSTM در لایه اول و به‌کارگیری یک مدل متا مبتنی بر RF در لایه دوم جهت مدل‌سازی روابط غیرخطی میان خروجی‌های الگوریتم‌های پایه طراحی شده است. این مدل با ترکیب خروجی‌های چندین الگوریتم پایه و به‌کارگیری الگوریتم

¹ Tripathi

² Aruleba & Sun

³ Rofik

ژنتیک جهت بهینه‌سازی همزمان هایپرپارامترها و انتخاب ویژگی‌های بهینه، توانسته است الگوهای پیچیده موجود در داده‌های اعتباری را به‌طور دقیق استخراج کند. اعتبارسنجی نتایج مدل با استفاده از دیتاست‌های متفاوت از جمله داده‌های ایرانی و دو مجموعه بنچمارک از استرالیا و آلمان نشان می‌دهد که مدل پیشنهادی دارای تعمیم‌پذیری و پایداری بالایی است. یافته‌های به‌دست‌آمده، علاوه بر تأکید بر اهمیت استفاده از فناوری‌های نوین در مدیریت ریسک اعتباری، زمینه را برای توسعه مدل‌های پیشرفته‌تر و کاربرد استراتژی‌های بهینه‌سازی جایگزین فراهم می‌آورد. از منظر سیاستی، این مدل می‌تواند به نهادهای مالی در کاهش ریسک‌های مرتبط با اعطای تسهیلات و بهبود فرآیندهای تصمیم‌گیری کمک کند.

از منظر سیاستی، یافته‌های این پژوهش نشان می‌دهد که نهادهای مالی و بانکی می‌توانند با سرمایه‌گذاری در فناوری‌های هوش مصنوعی و بهره‌گیری از سیستم‌های رتبه‌بندی اعتباری پیشرفته، ریسک‌های مرتبط با اعطای تسهیلات را به شکل قابل توجهی کاهش دهند و از این طریق به بهبود فرآیندهای تصمیم‌گیری و تخصیص منابع کمک نمایند. از سوی دیگر، از منظر پژوهشی، پیشنهاد می‌شود تحقیقات آتی با استفاده از الگوریتم‌های بهینه‌سازی جایگزین مانند بهینه‌سازی مبتنی بر یادگیری تقویتی و استراتژی‌های نمونه‌برداری پیشرفته مانند SMOTE و SMOTE-ENN به منظور بهبود عملکرد مدل‌ها در مواجهه با داده‌های نامتوازن و پیچیده، ادامه یابد. همچنین توسعه مدل‌های چندلایه و ترکیب آن‌ها با سایر تکنیک‌های یادگیری عمیق مانند CNN جهت استخراج ویژگی‌های بیشتر از داده‌های اعتباری می‌تواند به تعمیم‌پذیری مدل‌های رتبه‌بندی اعتباری در سطح بین‌المللی کمک نماید.

چارچوب ارائه‌شده در این پژوهش با ارائه یک مدل ترکیبی مبتنی بر یادگیری جمعی و استفاده از مدل متا، نقش مهمی در بهبود دقت سیستم‌های رتبه‌بندی اعتباری و کاهش خطاهای طبقه‌بندی دارد و می‌تواند به عنوان یک راهکار عملی در مدیریت ریسک اعتباری و توسعه سیستم‌های تصمیم‌گیری هوشمند در صنعت مالی مورد استفاده قرار گیرد. این دستاوردها زمینه را برای تحقیقات آتی در جهت ارتقای تکنیک‌های یادگیری عمیق، مدل‌های انسامبل و استراتژی‌های نمونه‌برداری فراهم آورده و نشان‌دهنده اهمیت بالای فناوری‌های نوین در تحول سیستم‌های اعتبارسنجی است.

منابع

۱. ابتیاع مجید، حسینی سید محمد، خوجانی رامین، (۱۴۰۰). رتبه بندی اعتباری مشتریان بانک به کمک روش جدید گروهی بر پایه ماشین بردار پشتیبان: مطالعه موردی بانک پاسارگاد، *محاسبات نرم*، ۱۰(۲)، ۱۵-۲.
۲. اخباری مهدیه، مخاطب رفیعی فریمه، (۱۳۸۹). کاربرد سیستم های استدلال عصبی - فازی در رتبه بندی اعتباری مشتریان حقوقی بانک ها، *مجله تحقیقات اقتصادی*، ۳(۹۲): ۲۱-۱.
۳. اسراری مصطفی. (۱۳۹۲). مقایسه کارایی مدل های شبکه عصبی، الگوریتم ژنتیک و لاجیت در ارزیابی ریسک اعتباری مشتریان حقیقی: مطالعه موردی در بانک کشاورزی، پایان نامه کارشناسی ارشد، رشته MBA، گرایش مالی، دانشگاه علوم اقتصادی.
۴. امیری هیوا، دهدار فرهاد، عبدلی محمدرضا، (۱۴۰۰). آرایه مدلی بهینه برای تعیین رتبه اعتباری احتمال نکول بانک ها و مؤسسات اعتباری غیربانکی بر مبنای تحلیل تمایزی (تشخیصی) و احتمال شرطی غیرخطی لاجیت و پروبیت، *نشریه اقتصاد و بانکداری اسلامی*، ۱(۳۶): ۱۷۵-۱۴۹.
۵. پاک بین، پرستو؛ پویانفر، احمد، (۱۳۹۷). رتبه بندی اعتباری مشتریان بانک با تکنیک طبقه بندی ترکیبی پویا و احتمال نرم (مطالعه موردی: بانک پاسارگاد)، پایان نامه کارشناسی ارشد، رشته مهندسی مالی و مدیریت ریسک، دانشگاه خاتم.
۶. جیحونی پور مهرداد، اعظمی سمیه، دل انگیزان سهراب، (۱۴۰۳). مدل سازی و شناسایی روابط علی بین عوامل اصلی ریسک اعتباری در سیستم بانکی با استفاده از تکنیک تصمیم گیری دیمتل، *نشریه سیاست گذاری اقتصادی*، ۱۷(۳۳): ۱۸۰-۲۱۱.
۷. حسامی، علی، (۱۳۹۸). امتیازدهی اعتباری مشتریان حقیقی بانکی با استفاده از روش یادگیری عمیق مورد مطالعه: مشتریان اعتباری یکی از بانک های ایرانی، پایان نامه کارشناسی ارشد، مهندسی صنایع - سیستم های مالی، دانشگاه اراک.
۸. خانجانی، محسن و همکاران، (۱۴۰۲). رتبه بندی اعتباری مشتریان حقوقی با استفاده از روش های یادگیری ماشین و یادگیری عمیق، اولین کنفرانس بین المللی توانمندی مدیریت، مهندسی صنایع، حسابداری و اقتصاد.
۹. دادمحمدی دانیال، احمدی عباس. (۱۳۹۳). رتبه بندی اعتباری مشتریان بانک با استفاده از شبکه عصبی با اتصالات جانبی، *فصلنامه توسعه مدیریت پولی و بانکی*، ۲(۳): ۲۸-۱.
۱۰. کوچه باغی، سمیرا؛ رجبی بهجت، امیر. (۱۳۹۶). انتخاب زیرمجموعه ای از ویژگی ها مبتنی بر الگوریتم جستجوی گرانشی باینری برای رتبه بندی اعتباری مشتریان بانکی با استفاده از ترکیب الگوریتم های طبقه بندی، پایان نامه کارشناسی ارشد، دانشگاه آزاد اسلامی، واحد بندرعباس.

۱۱. گل محمدی، نسرین. (۱۳۹۵). رتبه‌بندی اعتباری مشتریان بانک با استفاده از الگوریتم فراابتکاری (ژنتیک) (مورد مطالعه: مشتریان حقیقی بانک مسکن - شعب منطقه غرب تهران)، پایان‌نامه کارشناسی ارشد، مرکز آموزش عالی رجاء، دانشکده مهندسی صنایع.
۱۲. مهدوی، کاوه؛ حری، محمدصادق. (۱۳۹۴). طراحی مدلی جهت پیش‌بینی رتبه اعتباری مشتریان بانکها با استفاده از الگوریتم فراابتکاری و هیبریدی چند معیاره شبکه عصبی فازی - کلونی مورچگان (مطالعه موردی شعب پست بانک استان تهران، پژوهش‌های مدیریت در ایران، ۱۹(۱): ۹۱-۱۱۶.
۱۳. ندری، کامران؛ محرابی، لیلا. (۱۳۹۷). بررسی انواع ریسک و مدیریت ریسک در نظام بانکداری اسلامی، توسعه راهبرد، ۵۴(۱۴): ۱۶۰-۱۸۱.
۱۴. نوری کاوه، امینی صفیاء، محمودی خوشرو امید(۱۴۰۴). ارزیابی تأثیر نوآوری مالی بر عملکرد نهادها و مؤسسات مالی مبتنی بر رویکردهای فازی، نشریه اقتصاد و بانکداری اسلامی، ۱۴(۵۰): ۵۰۲-۴۷۱.
۱۵. نظرآقایی، مهدی. (۱۳۹۸). دسته‌بندی ریسک اعتباری مشتریان حقیقی با استفاده از یادگیری جمعی (مطالعه موردی بانک سپه)، پژوهش‌های پولی - بانکی، ۱۲(۳۹): ۱۶۶-۱۲۹.
16. Abdou, H. A., Mitra, S., Fry, J., & Elamer, A. A. (2019). Would two-stage scoring models alleviate bank exposure to bad debt? *Expert Systems with Applications*, 128, 1-13.
17. Abellán, J., & Mantas, C. J. (2014). Improving experimental studies about ensembles of classifiers for bankruptcy
18. Ala'raj, M., Abbod, M. F., Majdalawieh, M., & Jum'a, L. (2022). A deep learning model for behavioural credit scoring in banks. *Neural Computing and Applications*, 1-28.
19. Aruleba, I., & Sun, Y. (2024). Effective credit risk prediction using ensemble classifiers with model explanation. *IEEE Access*.
20. Attigeri, G. V., Pai, M. M., & Pai, R. M. (2017). Credit risk assessment using machine learning algorithms. *Advanced Science Letters*, 23(4), 3649-3653.
21. Das, S. R. (2019). Credit risk derivatives. In *World Scientific Reference on Contingent Claims Analysis in Corporate Finance: Volume 3: Empirical Testing and Applications of CCA* (pp. 373-400).
22. Du, P., & Shu, H. (2022). Exploration of financial market credit scoring and risk management and prediction using deep learning and bionic algorithm. *Journal of Global Information Management (JGIM)*, 30(9), 1-29.
23. Dumitrescu, E., Hue, S., Hurlin, C., & Tokpavi, S. (2018). *Machine Learning for Credit Scoring: Improving Logistic Regression with Non Linear Decision Tree Effects* (Doctoral dissertation, PhD thesis. Paris Nanterre University, University of Orleans).

24. Eletter, S. F., & Yaseen, S. G. (2017). Loan decision models for the Jordanian commercial banks. *Global Business and Economics Review*, 19(3), 323-338.
25. Hartomo, K. D., Arthur, C., & Nataliani, Y. (2025). A Novel Weighted Loss TabTransformer Integrating Explainable AI for Imbalanced Credit Risk Datasets. *IEEE Access*.
26. Khalili, N., & Rastegar, M. A. (2023). Optimal cost-sensitive credit scoring using a new hybrid performance metric. *Expert Systems with Applications*, 213, 119232.
27. Lando, D. (2009). Credit risk modeling. In *Credit Risk Modeling*. Princeton University Press.
28. Maleki, M. S., Motevallian, S. N., Hosseini, F., Sabokrou, M., & Maleki, H. S. (2021, October). Improvement of credit scoring by lstm autoencoder model. In *2021 11th International Conference on Computer Engineering and Knowledge (ICCKE)* (pp. 182-187). IEEE.
29. Rofik, R., Aulia, R., Musaadah, K., Ardyani, S. S. F., & Hakim, A. A. (2024). The optimization of credit scoring model using stacking ensemble learning and oversampling techniques. *Journal of Information System Exploration and Research*, 2(1).
30. Schetakis, N., Aghamalyan, D., Boguslavsky, M., Rees, A., Rakotomalala, M., & Griffin, P. R. (2024). Quantum machine learning for credit scoring. *Mathematics*, 12(9), 1391.
31. Suksamai, C., Hengpraprom, K., & Silachan, K. (2022, July). A Comparison of Normalization Data Transformation Efficiency Affecting with Bank Customer Credit Data Classification using Data Mining Techniques. The 14th NPRU National Academic Conference Nakhon Pathom Rajabhat University.
32. Sutrisno, H., & Halim, S. (2017, September). Credit Scoring Refinement Using Optimized Logistic Regression. In *2017 International Conference on Soft Computing, Intelligent System and Information Technology (ICSIT)* (pp. 26-31). IEEE.
33. Thomas, L., Crook, J., & Edelman, D. (2017). *Credit scoring and its applications*. Society for industrial and Applied Mathematics.
34. Tripathi, D., Edla, D. R., & Cheruku, R. (2018). Hybrid credit scoring model using neighborhood rough set and multi-layer ensemble classification. *Journal of Intelligent & Fuzzy Systems*, 34(3), 1543-1549.
35. Wang, Y., Jia, Y., Fan, S., & Xiao, J. (2023). Deep Reinforcement Learning Based on Balanced Stratified Prioritized Experience Replay for Customer Credit Scoring in Peer-to-Peer Lending.
36. Xiong, K., Li, L., Wang, Z., & Cao, G. (2024). A Dynamic Credit Risk Assessment Model Based on Deep Reinforcement Learning. *Academic Journal of Natural Science*, 1(1), 20-31.

-
37. Zhang, R., Ligu, X., & Qin, W. An Ensemble Credit Scoring Model Based on Logistic Regression with Heterogeneous Balancing and Weighting Effects. *Available at SSRN 4167821*.
 38. Zhuang, Y., Xu, Z., & Tang, Y. (2015, September). A credit scoring model based on bayesian network and mutual information. In *2015 12th Web Information System and Application Conference (WISA)* (pp. 281-286). IEEE.