



Paper type: Research Article

Hybrid Model Based on Genetic Algorithm, Bayesian Optimization and Machine Learning for Predicting Credit Status of Legal Clients

Pardis Fooladi¹, Mohsen Amini Khozani^{2*}, Zohreh Hajiha³, Shadi Shahverdiani⁴

Received: 2024/12/25

Accepted: 2025/04/08

Abstract

This study presents a novel hybrid model for credit scoring of banking clients, integrating Genetic Algorithm (GA), Bayesian Optimization, and the machine learning technique XGBoost. The primary objective of this model is to enhance accuracy and efficiency in credit risk assessment while minimizing the costs associated with prediction errors. The dataset comprises five-year financial statements and credit profiles of banking clients from 1398 to 1402. After data preprocessing—including normalization and handling of missing values—Genetic Algorithm was employed to select the optimal set of features. Subsequently, Bayesian Optimization was utilized as an advanced tool for fine-tuning the hyperparameters of the XGBoost model. The empirical results demonstrate the superior performance of the proposed hybrid model in comparison to conventional credit scoring methods. Achieving an accuracy of 79.3% and favorable recognition rates for both creditworthy and non-creditworthy clients, the model especially excels in classifying high-risk clients. Statistical analyses and comparative evaluations further confirm the positive impact of feature selection and hyperparameter optimization on the model's performance. This hybrid model offers a practical and effective tool for banks and financial institutions to reduce credit risk and improve customer management.

Keywords: Credit Scoring, Hybrid Model, Genetic Algorithm, Bayesian Optimization, XGBoost Machine Learning Model.



10.71600/jacgr.2024.1194492



How to cite this article: Fooladi, P., Amini Khozani, M., Hajiha, Z., & Shahverdiani, S. (2025). Hybrid Model Based on Genetic Algorithm, Bayesian Optimization and Machine Learning for Predicting Credit Status of Legal Clients. *Accounting and Corporate Governance Researches*, 3(9), 89. (In Persian).

1. PhD Candidate, Department of financial Management, Shahr-e-Qods Branch, Islamic Azad University, Tehran, Iran.
2. Assistant Prof, Department of financial Management, Shahr-e-Qods Branch, Islamic Azad University, Tehran, Iran.
*Corresponding Author: mo.aminikhousani@iau.ac.ir
3. Prof, Department of Accounting, ST,C. Branch, Islamic Azad University, Tehran, Iran.
4. Assistant Prof, Department of financial Management, Shahr-e-Qods Branch, Islamic Azad University, Tehran, Iran.



نوع مقاله: علمی پژوهشی

ارائه‌ی مدل هیبریدی مبتنی بر الگوریتم ژنتیک، بهینه‌سازی بیزین و یادگیری ماشینی جهت پیش‌بینی وضعیت اعتباری مشتریان حقوقی

پردیس فولادی^۱، محسن امینی خوزانی^{۲*}، زهره حاجیه‌ها^۳، شادی شاهوردیانی^۴

تاریخ پذیرش: ۱۴۰۴/۰۱/۱۹

تاریخ دریافت: ۱۴۰۳/۱۰/۰۵

چکیده

این مقاله به ارائه‌ی یک مدل هیبریدی جدید برای اعتبارسنجی مشتریان بانکی می‌پردازد که ترکیبی از الگوریتم ژنتیک، بهینه‌سازی بیزین و مدل یادگیری ماشینی XGBoost است. هدف اصلی این مدل، بهبود دقت و کارایی در ارزیابی ریسک اعتباری مشتریان و کاهش هزینه‌های مرتبط با بروز اشتباه در پیش‌بینی‌هاست. در این تحقیق، داده‌های مربوط به دوره‌های پنج ساله‌ی صورت‌های مالی و وضعیت اعتباری مشتریان بانکی مربوط به سال‌های ۱۳۹۸ الی ۱۴۰۲ مورد استفاده قرار گرفته و پس از پیش‌پردازش داده‌ها شامل نرمال‌سازی و مدیریت داده‌های گم‌شده، از الگوریتم ژنتیک برای انتخاب ویژگی‌های بهینه مورد استفاده قرار گرفته است. سپس، بهینه‌سازی بیزین به عنوان یک ابزار پیشرفته جهت تنظیم دقیق ابرپارامترهای XGBoost به کار گرفته شده است. نتایج حاصل شده از پژوهش، بیان‌گر عملکرد برتر مدل پیشنهادی در مقایسه با روش‌های متداول اعتبارسنجی است. مدل هیبریدی پیشنهادی توانسته است با دقت ۷۹/۳ درصد و نرخ بازشناسی مطلوب برای مشتریان معتبر و غیرمعتبر، به‌ویژه در طبقه‌بندی مشتریان با ریسک بالا، برتری خود را نشان دهد. تحلیل‌های آماری و مقایسه عملکرد مدل با سایر روش‌های موجود، تأثیر مثبت انتخاب ویژگی و تنظیم ارامترها را تأیید می‌کند. این مدل می‌تواند به عنوان یک ابزار عملیاتی برای بانک‌ها و مؤسسات مالی جهت کاهش ریسک اعتباری و بهبود مدیریت مشتریان مورد استفاده قرار گیرد.

واژه‌های کلیدی: اعتبارسنجی، مدل هیبریدی، الگوریتم ژنتیک، بهینه‌سازی بیزین و مدل یادگیری ماشینی XGBoost



10.71600/jacgr.2024.1194492



استناد: فولادی، پردیس؛ امینی خوزانی، محسن؛ حاجیه‌ها، زهره؛ شاهوردیانی، شادی (۱۴۰۳). ارائه‌ی مدل هیبریدی مبتنی بر الگوریتم ژنتیک، بهینه‌سازی بیزین و یادگیری ماشینی جهت پیش‌بینی وضعیت اعتباری مشتریان حقوقی. پژوهش‌های حسابداری و راهبری شرکتی، ۳(۹)، ۸۹-۱۰۵.

۱. دانشجوی دکتری مهندسی مالی، واحد شهرقدس، دانشگاه آزاد اسلامی، تهران، ایران.

۲. استادیار، گروه مهندسی مالی، واحد شهرقدس، دانشگاه آزاد اسلامی، تهران، ایران. *نویسنده مسئول: m.amini@iau.ac.ir

۳. استاد، گروه حسابداری، واحد تهران جنوب، دانشگاه آزاد اسلامی، تهران، ایران.

۴. استادیار، گروه مهندسی مالی، واحد شهرقدس، دانشگاه آزاد اسلامی، تهران، ایران.

مقدمه

اعتبارسنجی به مفهوم ارزیابی و سنجش توان بازپرداخت متقاضیان اعتبار و تسهیلات مالی و احتمال عدم بازپرداخت اعتبارات از سوی آن‌هاست (درواری و همکاران، ۱۴۰۴). اعتبارسنجی دقیق می‌تواند از وقوع ریسک‌های مالی جلوگیری کرده و به اعتباردهندگان اطمینان دهد که وام‌ها و تسهیلات، به مشتریان اعطا می‌شود که توانایی بازپرداخت دارند. این امر نه تنها به کاهش نرخ نکول وام‌ها کمک می‌کند، بلکه به بهبود کیفیت پرتفوی وام‌های بانک نیز منجر خواهد شد (رجبی‌پور میبدی و همکاران، ۱۳۹۲). از یک سو، ریسک مالی و اعتباردهی به مشتریان حقوقی می‌تواند تأثیرات قابل توجهی بر درآمد و ثبات مالی بانک‌ها داشته باشد؛ چنانچه بانک‌ها بدون ارزیابی دقیق اعتبار مشتریان حقوقی اقدام به اعطای وام کنند، احتمال نکول وام‌ها افزایش می‌یابد که این امر ممکن است به زیان‌های مالی جدی منجر گردد. از سوی دیگر، اعتبارسنجی دقیق و مؤثر می‌تواند منجر به افزایش درآمد بانک‌ها از طریق کاهش نرخ نکول و افزایش بازدهی وام‌ها گردد. همچنین، این فرآیند به بانک‌ها امکان می‌دهد تا با اطمینان بیشتری به برنامه‌ریزی مالی و مدیریت ریسک بپردازند، که در نهایت بر ثبات مالی و افزایش اعتماد سرمایه‌گذاران مؤثر خواهد بود (حبیبی و همکاران، ۱۴۰۳). روش‌های سنتی اعتبارسنجی، مانند استفاده از نسبت‌های مالی و مدل‌های آماری ساده، ممکن است در مواجهه با پیچیدگی‌های داده‌های مالی مشتریان حقوقی، عملکرد ضعیفی داشته باشند. این روش‌ها اغلب به داده‌های تاریخی و اطلاعات مالی محدود تکیه می‌کنند که ممکن است نتوانند به‌طور کامل ریسک‌های مالی و اعتباری را ارزیابی کنند (چرن و همکاران، ۲۰۲۱). یکی از مشکلات اصلی روش‌های سنتی اعتبارسنجی این است که معمولاً قادر به پردازش حجم بالای داده‌ها و تنوع آن‌ها نیستند. از آنجایی که داده‌های مالی مشتریان حقوقی می‌تواند شامل اطلاعات پیچیده‌ای مانند جریان‌های نقدی، قراردادهای مالی و تعهدات بلندمدت باشد، استفاده از روش‌های سنتی ممکن است منجر به ارزیابی نادرست ریسک گردد (چن و همکاران، ۲۰۱۶). علاوه بر این، روش‌های سنتی اعتبارسنجی معمولاً انعطاف‌پذیری کمتری در مواجهه با تغییرات سریع در شرایط اقتصادی و مالی دارند. این روش‌ها به دلیل تکیه بر داده‌های تاریخی، ممکن است نتوانند به سرعت به تغییرات در بازار و شرایط اقتصادی واکنش نشان دهند. این امر می‌تواند منجر به ارزیابی نادرست ریسک و تصمیم‌گیری‌های نادرست در اعطای وام شود (لانگ‌وین و همکاران، ۲۰۲۴). در مقابل، استفاده از تکنیک‌های پیشرفته‌تر مانند داده‌کاوی و یادگیری ماشینی می‌تواند به بهبود دقت و کارایی فرآیند اعتبارسنجی کمک کند. این تکنیک‌ها قادر به پردازش حجم بالای داده‌ها و شناسایی الگوهای پیچیده در داده‌های مالی هستند که می‌تواند به ارزیابی دقیق‌تر ریسک‌های مالی و اعتباری منجر شود (مرادی و مخاطب‌رفیعی، ۲۰۱۹). مدل‌های یادگیری ماشینی، به‌طور فزاینده‌ای در شناسایی مشتریان غیرمعتبر و پیش‌بینی ریسک‌های اعتباری استفاده می‌شوند. در این زمینه، استفاده از الگوریتم‌های انتخاب ویژگی برای بهبود دقت مدل‌ها و کاهش پیچیدگی محاسباتی اهمیت ویژه‌ای دارد. روش XGBoost^۱ نیز یک الگوریتم یادگیری گروهی قدرتمند است که طبقه‌بندی مبتنی بر ویژگی را افزایش می‌دهد و برای هر دو وظیفه طبقه‌بندی و رگرسیون استفاده می‌شود. این الگوریتم براساس معیارهای خاصی که عموماً در نگرانی‌های مالی اهمیت بالایی دارند، امکان دستیابی به تصمیم‌گیری‌های اعتباری قابل اعتمادتر و آگاهانه‌تر را فراهم می‌نماید (کندی و گارسیا، ۲۰۲۵). در این میان، مدل‌های ریاضی نیز در مسائل مالی نقش بسیار مهمی ایفا کرده و با استفاده از ابزارهای ریاضی و آماری به تحلیل و پیش‌بینی رفتار بازارهای مالی و ارزیابی ریسک‌های مختلف می‌پردازند. از جمله الگوریتم ژنتیک (GA)^۲ که به‌عنوان یکی از تکنیک‌های محاسبات تکاملی، برای بهینه‌سازی مسائل پیچیده و چندبعدی استفاده می‌شود (مارکز مارزال و همکاران، ۲۰۱۳). به‌علاوه، مدل بهینه‌سازی بیزین^۳ نیز یک روش آماری برای بهینه‌سازی توابع پیچیده و پرهزینه است که از توزیع‌های احتمالی برای مدل‌سازی عدم قطعیت استفاده می‌کند (اسنوک و همکاران، ۲۰۱۲). امروزه، مدل‌های هیبریدی به‌عنوان راهکاری مؤثر برای افزایش دقت و جامعیت اعتبارسنجی شناخته شده‌اند. این مدل‌ها با ترکیب روش‌های سنتی آماری و تکنیک‌های پیشرفته هوش مصنوعی می‌توانند به بهبود عملکرد اعتبارسنجی کمک کنند. مطالعات نشان داده‌اند که مدل‌های هیبریدی می‌توانند با استفاده از داده‌های واقعی، عملکرد بهتری نسبت به مدل‌های تک-گانه داشته باشند (چی و همکاران، ۲۰۱۹). یکی از مزایای اصلی مدل‌های هیبریدی آن است که قادرند با پردازش حجم بالای داده‌ها و

1 eXtreme Gradient Boosting

2 Genetic Algorithm

3 Bayesian Optimization

شناسایی الگوهای پیچیده، ریسک‌های اعتباری را با دقت بیشتری ارزیابی کند و بدین وسیله می‌تواند به بانک‌ها و مؤسسات مالی در مدیریت بهتر ریسک‌های اعتباری کمک کند (لی و همکاران، ۲۰۲۴).

پژوهش حاضر، به ارائه‌ی یک مدل هیبریدی خواهد پرداخت که از ترکیب الگوریتم ژنتیک، بهینه‌سازی بیزین و الگوریتم XGBoost برای ارزیابی و سنجش اعتبار مشتریان حقوقی استفاده می‌کند. هدف این است که با بهره‌گیری از توانایی الگوریتم ژنتیک در انتخاب بهترین ویژگی، کاربرد بهینه‌سازی بیزین در بهبود جستجوی پارامترهای مدل و قدرت پیش‌بینی XGBoost به یک مدل با دقت بالا دست یابیم. یکی از جنبه‌های مهم و نوآورانه این پژوهش، استفاده از الگوریتم‌های بهینه‌سازی پیشرفته مانند الگوریتم ژنتیک و بهینه‌سازی بیزین برای تنظیم پارامترهای مدل است؛ این روش‌ها به‌طور ویژه برای تعیین پارامترهای بهینه در مدل‌های پیچیده مانند XGBoost به‌کار رفته‌اند. با استفاده از این روش‌های بهینه‌سازی، دقت و کارایی مدل در پیش‌بینی ریسک اعتباری افزایش یافته و قابلیت پیش‌بینی مدل در شرایط پیچیده و متغیرهای نامشهود تقویت می‌گردد. علاوه‌برآن، مدل پیشنهادی قادر است به‌طور هم‌زمان معیارهای مختلفی مانند سودآوری، ریسک و کیفیت خدمات را در نظر بگیرد و با تحلیل داده‌های تاریخی و شناسایی الگوهای مهم، به اعتباردهندگان کمک کند تا تصمیمات بهتری در خصوص اعتبارسنجی مشتریان بگیرند. افزون برآن، مدل هیبریدی ارائه شده قادر به شناسایی نقاط ضعف و قوت هر مشتری بوده و راهکارهایی برای بهبود عملکرد ارائه آنان می‌دهد.

مبانی نظری

در این بخش به مرور مدل‌های موجود در اعتبارسنجی مشتریان می‌پردازد و مطالعات مشابهی را که از مدل‌های الگوریتم ژنتیک، بهینه‌سازی بیزین، XGBoost، روش‌های ترکیبی در تحلیل‌های مالی و اعتبارسنجی استفاده شده، مورد بررسی قرار خواهد داد و پس از آن به مقایسه مدل‌های هیبریدی با مدل‌های سنتی نیز اشاره خواهد نمود.

مدل‌های سنتی اعتبارسنجی در بانکداری: بسیاری از بانک‌ها به‌منظور ارزیابی مشتریان از مدل‌های آماری نظیر رگرسیون لجستیک، تحلیل تشخیصی و مدل‌های امتیازدهی سنتی استفاده کرده‌اند. این مدل‌ها تلاش می‌کنند با استفاده از متغیرهای مالی و اقتصادی، رفتار اعتباری مشتریان را پیش‌بینی کنند. روش‌های آماری و کلاسیک برای ارزیابی مشتریان حقوقی در بانک‌ها و مؤسسات مالی به‌طور گسترده‌ای مورد استفاده قرار گرفته‌اند (جودیچی، ۲۰۰۵). در سال‌های اخیر، با پیشرفت‌های تکنولوژی و افزایش حجم داده‌ها، نیاز به روش‌های پیشرفته‌تر و دقیق‌تر برای ارزیابی مشتریان حقوقی احساس شده است. مدل‌های آماری کلاسیک همچنان به‌عنوان ابزارهای پایه‌ای مورد استفاده قرار می‌گیرند، اما ترکیب آن‌ها با روش‌های مدرن‌تر مانند الگوریتم‌های یادگیری ماشینی می‌تواند دقت و کارایی ارزیابی‌ها را بهبود بخشد. مدل‌های کلاسیک در تحلیل مشتریان حقوقی به دلیل پیچیدگی‌های بالای داده‌ها و تغییرات اقتصادی، ممکن است عملکرد کافی نداشته باشند. این روش‌ها بیشتر به فرضیات خطی متکی‌اند و نمی‌توانند به‌طور دقیق رفتار مشتریان پیچیده و بزرگ را پیش‌بینی کنند. این محدودیت‌ها باعث می‌شود که دقت پیش‌بینی‌ها کاهش یابد و تصمیمات نهایی ممکن است بهینه نباشند (روزنزوایگ، ۲۰۱۴). علاوه بر این، مدل‌های کلاسیک معمولاً به داده‌های تاریخی متکی هستند و نمی‌توانند به‌خوبی با تغییرات سریع اقتصادی و بازار سازگار شوند. در نهایت، برای مواجهه با چالش‌های جدید و پیچیدگی‌های بیشتر، ترکیب این مدل‌ها با روش‌های مدرن‌تر و پیشرفته‌تر ضروری است (بوک، ۲۰۱۵).

کاربرد الگوریتم ژنتیک در اعتبارسنجی: الگوریتم ژنتیک به‌عنوان یکی از روش‌های بهینه‌سازی مبتنی بر اصول تکامل زیستی، نقش برجسته‌ای در اعتبارسنجی ایفا می‌کند. این الگوریتم برای انتخاب ویژگی‌ها^۱ و بهینه‌سازی پارامترهای مدل‌های اعتبارسنجی مورد استفاده قرار می‌گیرد. به دلیل قابلیت بالای الگوریتم ژنتیک در جستجوی فضای گسترده و غیرخطی، این روش برای کاربردهایی که شامل داده‌های بزرگ یا متغیرهای متعدد است، بسیار مفید بوده است (اورسکی، ۲۰۱۴). یکی دیگر از کاربردهای کلیدی الگوریتم ژنتیک، تنظیم بهینه وزن‌ها یا ضرایب در مدل‌های اعتباری است. این رویکرد به‌ویژه در سیستم‌های ترکیبی مانند مدل‌های شبکه عصبی یا سیستم‌های مبتنی بر تصمیم‌گیری کاربرد دارد. (ژو و همکاران، ۲۰۲۰). الگوریتم ژنتیک

همچنین در ارزیابی ریسک اعتباری برای شناسایی بهترین سیاست‌های تخصیص منابع و کاهش ریسک کلی بانک یا مؤسسات مالی استفاده می‌شود. در چنین کاربردهایی، الگوریتم ژنتیک به‌خاطر سرعت و انعطاف‌پذیری بالا، از روش‌های متعارف مانند تحلیل‌های ریاضی پیشی گرفته است (لئو و همکاران، ۲۰۱۹). تحقیقات اخیر نشان داده‌اند که الگوریتم ژنتیک در ترکیب با مدل‌های دیگر، قابلیت رقابت بالایی با جدیدترین روش‌های یادگیری ماشین دارد. الگوریتم ژنتیک به جای استفاده از تمام ویژگی‌ها، زیرمجموعه‌ای از آن‌ها را انتخاب می‌کند که به‌طور بهینه در پیش‌بینی متغیر خروجی (Y) مؤثر هستند. این کار می‌تواند منجر به کاهش پیچیدگی مدل، کاهش نویز و افزایش تعمیم‌پذیری شود. در مدل پیشنهادی این پژوهش، از الگوریتم ژنتیک کلاسیک^۱ (CGA) استفاده شده است که شامل مراحل استاندارد انتخاب، تولید مثل، تقاطع^۲ و جهش^۳ است. این نوع الگوریتم ژنتیک به‌طور خاص با ترکیب مدل XGBoost و معیارهای ارزیابی دقیق (مانند AUC و F1-Score)، بهینه‌سازی قابل‌قبولی را برای مدل اعتبارسنجی پژوهش ارائه کرده است.

بهینه‌سازی مدل‌ها با استفاده از الگوریتم‌های بهینه‌سازی و اهمیت انتخاب ویژگی و بهینه‌سازی پارامترها:

در بهینه‌سازی مدل‌ها، استفاده از بهینه‌سازی بیزین به‌عنوان روشی برای بهبود جستجوی پارامترهای مدل شناخته شده است. این روش به‌ویژه در شرایطی که فضای جستجو بسیار بزرگ است، می‌تواند به‌طور کارآمدی بهترین ترکیب پارامترها را پیدا کند (اسنوک و همکاران، ۲۰۱۲).^۴ در این پژوهش، با استفاده از بهینه‌سازی بیزین، پارامترهایی مانند `max_depth`، `learning_rate`، `n_estimators` و `gamma` بهینه می‌شوند تا دقت مدل برای پیش‌بینی ریسک‌های اعتباری به حداکثر برسد. افزون بر آن، انتخاب ویژگی‌ها و بهینه‌سازی پارامترهای مدل از اهمیت ویژه‌ای برخوردار است، چرا که می‌تواند به کاهش پیچیدگی و افزایش دقت مدل کمک کند. برخلاف روش‌های جستجوی تصادفی یا جستجوی شبکه‌ای^۵، روش بیزین با تعداد تکرارهای کمتر، نتایج بهتری ارائه می‌دهد. این مسأله باعث می‌شود منابع محاسباتی بهینه‌تر مصرف شوند و زمان اجرای مدل کاهش یابد. از آنجایی که روش XGBoost یک الگوریتم پیچیده با تعداد زیادی از هایپرپارامترهاست، روش بهینه‌سازی بیزین می‌تواند به‌طور خودکار تأثیرات متقابل بین پارامترها را مدل‌سازی کند و بهترین ترکیب را ارائه دهد. از این‌رو می‌توان اذعان داشت روش بیزین نقشی کلیدی در بهینه‌سازی بخش XGBoost ایفا می‌کند. روش بهینه‌سازی بیزین مورد استفاده در مدل ما، با توجه به ضرورت پیاده‌سازی مدل در کتابخانه `skopt` از `Gaussian Process` استفاده می‌کند. این مدل به دلیل توانایی‌اش در مدل‌سازی ارتباطات پیچیده میان پارامترها و ارائه تخمین‌های دقیق، انتخابی مناسب برای بهینه‌سازی مدل XGBoost در این پروژه بوده است. به بیان دیگر، استفاده از روش بهینه‌سازی بیزین با تمرکز بر `Gaussian Process` به ما کمک کرده است تا با تعداد تکرارهای کمتر، تنظیمات بهینه برای XGBoost را پیدا کنیم. الگوریتم بهینه‌سازی بیزین در یک حلقه انجام می‌شود که در هر مرحله پارامترهای بهینه را جستجو می‌کند و از یک مدل احتمالاتی برای پیش‌بینی عملکرد تابع هدف استفاده می‌کند و از فرآیند زیر پیروی می‌کند:

(۱) **تعریف تابع هدف:** ابتدا تابعی که باید بهینه شود (مانند دقت مدل یا معیار خطا) تعریف می‌شود. این تابع به

ابزار پارامترهای مدل وابسته است. فرض کنید در این مدل، $f(x)$ تابع هدف است که باید بهینه شود.

$$x^* = \arg \max_{x \in X} f(x) \quad (۱)$$

(۲) **تعریف فضای جستجو:** فضای ابزار پارامترهایی که باید بررسی شوند مشخص می‌شود. در این مدل، X مجموعه

ابزار پارامترهایی است که در فضای جستجو X قرار دارد.

(۳) **مدل‌سازی اولیه:** یک مدل احتمال مانند فرآیند گاوسی^۶ یا رگرسیون تصادفی^۷ برای تقریب تابع هدف ساخته می‌شود.

این مدل تخمینی از عملکرد تابع هدف در فضای جستجو ارائه می‌دهد. در واقع، تابع هدف به‌عنوان یک فرآیند

گاوسی $g(x)$ مدل می‌شود که در آن میانگین تخمینی تابع و $k(x, x')$ تابع همبستگی یا کرنل می‌باشد:

1 Canonical Genetic Algorithm

2 Crossover

3 Mutation

4 Snoek et al

5 Grid Search

6 Gaussian Process

7 Random Forest

$$g(x) \sim \mathcal{GP}(\mu(x), k(x, x')) \quad (2)$$

۴) **اکتساب داده‌های جدید:** با استفاده از یک تابع اکتساب^۱ بهترین نقطه بعدی در فضای جستجو برای ارزیابی انتخاب می‌شود. این تابع بین بهره‌برداری (استفاده از نواحی با عملکرد خوب) و اکتشاف (بررسی نواحی ناشناخته) تعادل ایجاد می‌کند. در این مدل، نقطه بعدی برای ارزیابی x_{next} از طریق بهینه‌سازی تابع اکتساب $\alpha(x)$ انتخاب می‌شود:

$$x_{next} = \arg \max_{x \in X} \alpha(x) \quad (3)$$

۵) **به‌روزرسانی مدل احتمال:** مدل احتمال بر اساس داده‌های جدید به‌روزرسانی می‌شود و این فرآیند تکرار می‌شود تا زمانی که یک معیار توقف (مانند تعداد تکرار یا عدم بهبود قابل توجه) برآورده شود.

مدل‌های یادگیری ماشینی و کاربرد آن‌ها در اعتبارسنجی مشتری: مدل‌های یادگیری ماشینی به‌ویژه در حوزه شبیه‌سازی ریسک و ارزیابی اعتبار توانسته‌اند تحولات چشمگیری ایجاد کنند. به‌طور خاص، مدل‌هایی همچون XGBoost به دلیل دقت بالا، توانایی پردازش داده‌های بزرگ و قابلیت مقابله با ویژگی‌های پیچیده به یکی از محبوب‌ترین الگوریتم‌ها در این زمینه تبدیل شده است (چن و گاسترین، ۲۰۱۶^۲). این مدل به‌طور خاص برای طبقه‌بندی و پیش‌بینی استفاده می‌شود و می‌تواند با ترکیب ویژگی‌های مختلف و وزن‌دهی مناسب، پیش‌بینی دقیقی از وضعیت مالی مشتریان به‌دست دهد. در مدل‌های اعتبارسنجی مالی، مهم‌ترین چالش‌ها شامل شناسایی دقیق مشتریان غیرمعتبر (مشتریانی که احتمال بازپرداخت وام را ندارند) و مدیریت داده‌های ناقص و نابرابر است (فردمن، ۲۰۰۱^۳). از دلایل انتخاب XGBoost برای مدل هیبریدی پیشنهادی در این پژوهش می‌توان به مواردی همچون دقت بالا، مقابله با overfitting (کاهش بروز اضافه‌آموزی)، سریع و کارا بودن، قابلیت توضیح‌پذیری^۴ و قابلیت گنجاندن ویژگی‌های متنوع اشاره نمود. در مدل هیبریدی، XGBoost معمولاً به‌عنوان یک مدل پایه برای انجام پیش‌بینی‌ها و اعتبارسنجی‌ها استفاده می‌شود. این مدل می‌تواند با سایر مدل‌ها ترکیب شود تا نتایج بهینه‌تری حاصل شود. از آنجایی که در پژوهش حاضر، هدف پیش‌بینی یک متغیر طبقه‌بندی (ورود یا عدم ورشکستگی شرکت) می‌باشد، از روش XGBClassifier استفاده شده تا پیش‌بینی‌های دقیق‌تری برای طبقه‌بندی ورشکستگی انجام دهد. روش XGBoost در واقع یک الگوریتم است که از مدل‌های ریاضی برای بهینه‌سازی و پیش‌بینی استفاده می‌کند. در مدل پیشنهادی، XGBoost به‌عنوان یک الگوریتم با پایه ریاضی عمل می‌کند که در پس‌زمینه از توابع ریاضی برای تقویت پیش‌بینی‌ها و کاهش خطا استفاده نموده و ورودی‌ها عبارتند از:

$$D = \{(x_i, y)\}_{i=1}^n \leftarrow \text{داده‌های آموزشی شامل ویژگی‌ها } (x_i) \text{ و برچسب } (y)$$

$$K: \text{تعداد تکرارها یا درخت‌های تصمیم} \leftarrow$$

$$f_k(x): \text{مدل پیش‌بینی درختی در تکرار } k \leftarrow$$

$$Obj(\theta): \text{تابع هدف شامل خطای پیش‌بینی و منظم‌کننده} \leftarrow$$

خروجی این روش، پیش‌بینی نهایی با ترکیب درخت‌های تصمیم خواهد بود و نمایش ریاضی معادلات کلیدی الگوریتم به‌صورت زیر می‌باشد:

$$\text{تابع هدف: } Obj = \sum [g_i * f(x_i) + 1/2 * h_i * f(x_i)^2] + \lambda \Omega(f) \leftarrow$$

$$\Omega(f) = \gamma T + 1/2 \alpha \sum (w) \quad \text{که در آن } \gamma \text{ و } \alpha \text{ پارامترهای تنظیم‌شده هستند}$$

$$F(x) = F_{-1}(x) + \eta * f(x) \quad \text{به‌روزرسانی مدل: } \leftarrow$$

پیشینه پژوهش

شن و وو (۲۰۲۵) در مطالعه‌ای یک مدل هیبریدی برای ارزیابی ریسک اعتباری شرکت‌های بورسی ارائه داده‌اند که ترکیبی از شبکه عصبی کانولوشنی زمانی (TCN) و مدل DilateFormer است. نتایج تجربی نشان داده است که این مدل ترکیبی عملکرد بهتری نسبت به روش‌های سنتی و برخی مدل‌های یادگیری عمیق در پیش‌بینی ریسک اعتباری دارد و می‌تواند ابزار مفیدی برای

1 Acquisition Function
2 Chen & Guestrin
3 Friedman
4 Interpretability

مؤسسات مالی در ارزیابی دقیق‌تر اعتبار شرکت‌ها باشد. حافظ و همکاران (۲۰۲۵) نیز یک مرور سیستماتیک از تکنیک‌های مبتنی بر هوش مصنوعی برای شناسایی تقلب در کارت‌های اعتباری ارائه داده و به بررسی روش‌های مختلف یادگیری ماشین و یادگیری عمیق، از جمله شبکه‌های عصبی، درخت‌های تصمیم، الگوریتم‌های دسته‌بندی و تکنیک‌های مبتنی بر داده‌های نامتوازن پرداخته‌اند. نتایج حاکی از آن است که ترکیب الگوریتم‌های مختلف و استفاده از مدل‌های ترکیبی می‌تواند دقت تشخیص تقلب را افزایش داده و بهبود عملکرد سیستم‌های مالی را تسهیل کند. منگ (۲۰۲۴) در مقاله خود به بررسی یک مدل ترکیبی جدید برای پیش‌بینی قیمت‌های بازار مالی پرداخته‌اند که از ترکیب روش پرسپترون چندلایه (MLP) و بهینه‌سازی (ALO) استفاده می‌کند. هدف اصلی این مدل، بهبود دقت پیش‌بینی قیمت‌های سهام در بازارهای مالی است که به دلیل نوسانات بالا و پیچیدگی‌های غیرخطی، پیش‌بینی آن‌ها چالش‌برانگیز است. اوزوپک و همکاران (۲۰۲۴) نیز یک مدل ترکیبی جدید به نام EMD-TI-LSTM را معرفی نموده‌اند که از ترکیب EMD^۱، شاخص‌های تکنیکال (TI) و LSTM^۲ برای پیش‌بینی مالی استفاده کرده و هدف آن، بهبود دقت پیش‌بینی‌های مالی با استفاده از تکنیک‌های یادگیری ماشینی است. سروانتس و همکاران (۲۰۲۴) در مقاله خود به کاربرد الگوریتم‌های ژنتیک در اعتبارسنجی اشاره نموده‌اند. مسأله مورد بررسی این مطالعه شامل شناسایی مجموعه‌ای از رئوس با حداقل تعداد برای یک گراف داده شده است. الگوریتم ژنتیک رتبه‌ای استفاده شده در این مطالعه توانسته است از بهینه‌های محلی فرار کند و به بهینه‌سازی راه‌حل‌های موجود بپردازد. طبق یافته‌های این تحقیق، الگوریتم‌های ژنتیک می‌توانند در حل مسائل نظری نیز مفید باشند. ویسبند و همکاران (۲۰۲۳) نیز در پژوهشی به بررسی اعتبارسنجی واریانت‌های ژنتیکی با استفاده از داده‌های NGS و شبکه‌های عصبی عمیق پرداخته‌اند. نتایج نشان می‌دهد که این مدل می‌تواند به دقتی مشابه با محققان آموزش‌دیده دست یابد و بهبود قابل توجهی در فرآیند اعتبارسنجی واریانت‌های ژنتیکی ایجاد کند. علاوه بر آن، مطالعه انجام‌شده توسط کاتوچ و همکاران (۲۰۲۱) مروری بر تاریخچه، وضعیت فعلی و آینده الگوریتم‌های ژنتیک دارد. در این مقاله، الگوریتم‌های معروف و پیاده‌سازی‌های آن‌ها با مزایا و معایب‌شان بررسی شده‌اند و حوزه‌های مختلف تحقیقاتی که در آن‌ها از الگوریتم‌های ژنتیک استفاده می‌شود، پوشش داده شده است. چن و گاسترین (۲۰۱۶) نیز در مقاله‌ای به معرفی سیستم XGBoost که یک سیستم تقویت درختی مقیاس‌پذیر است، پرداخته‌اند که این سیستم به‌طور گسترده‌ای توسط دانشمندان علم داده برای دستیابی به نتایج پیشرفته در بسیاری از چالش‌های یادگیری ماشین استفاده می‌شود. این مقاله نشان می‌دهد که XGBoost به دلیل کارایی بالا و توانایی در مدیریت داده‌های بزرگ و پیچیده، به یکی از ابزارهای اصلی در بسیاری از کاربردهای یادگیری ماشین تبدیل شده است. همچنین، لسمن و همکاران (۲۰۱۵) در پژوهشی چندین الگوریتم طبقه‌بندی جدید را با الگوریتم‌های پیشرفته موجود مقایسه کرده‌اند که شامل روش‌های سنتی مانند رگرسیون لجستیک و تحلیل تشخیصی خطی و همچنین الگوریتم‌های جدیدتر مانند ماشین‌های بردار پشتیبان (SVM) و شبکه‌های عصبی می‌باشد. ژانگ و ما (۲۰۱۲) در کتاب خود روش‌ها و کاربردهای یادگیری جمعی^۳ را مورد بررسی قرار داده‌اند. این کتاب شامل مباحثی نظیر مبانی یادگیری جمعی، مرور روش‌ها و کاربردهای مختلف الگوریتم‌های تقویتی، بررسی یکی از محبوب‌ترین الگوریتم‌های یادگیری جمعی یعنی جنگل تصادفی و مثال‌هایی از کاربرد یادگیری جمعی در تشخیص چهره، ردیابی اشیاء، بیوانفورماتیک و غیره بوده و برای محققان و متخصصان بسیار مفید می‌باشد و به ارائه دانش عمیق و کاربردی در زمینه یادگیری جمعی پرداخته است. یانوفسکی و بیکل (۲۰۱۰) در پژوهشی به بررسی روش‌های مختلف اعتبارسنجی الگوریتم‌های تشخیص بیان ژنی پرداخته‌اند و دو روش اصلی برای ارزیابی خطای پیش‌بینی توسط آنان معرفی شده که عبارتند از: روش اعتبارسنجی متقابل و روش پیش‌بینی پسین. نتایج نشان می‌دهد که روش‌های مبتنی بر مدل‌های سلسله‌مراتبی عملکرد بهتری نسبت به سایر الگوریتم‌ها دارند. ادوارد آلتمن (۱۹۶۸) در یکی از پژوهش‌های خود که به مقاله‌ای پایه‌ای در زمینه پیش‌بینی ورشکستگی شرکت‌ها تبدیل گشته، از تحلیل تشخیصی چندگانه (MDA) برای بررسی مجموعه‌ای از نسبت‌های مالی و اقتصادی استفاده نموده است. در این مطالعه که هدف اصلی آن توسعه‌ی یک مدل پیش‌بینی ورشکستگی بر اساس داده‌های حسابداری بود، با استفاده از داده‌های مالی شرکت‌های ورشکسته و غیرورشکسته، پنج نسبت مالی کلیدی شناسایی شد که توانست با دقت بالایی

1 Empirical Mode Decomposition

2 Empirical Mode Decomposition

3 Ensemble Learning

ارائه‌ی مدل هیبریدی مبتنی بر الگوریتم ژنتیک، بهینه‌سازی بیزین و... شرکت‌های ورشکسته و غیرورشکسته را از هم تفکیک کند و به‌عنوان یک ابزار مهم در تحلیل ریسک مالی مورد استفاده قرار گیرد.

روش‌شناسی پژوهش

مدل هیبریدی طراحی‌شده در این پژوهش بر پایه سه تکنیک اصلی الگوریتم ژنتیک، بهینه‌سازی بیزین و یادگیری ماشینی XGBoost استوار است. در این بخش، ساختار مدل هیبریدی و نحوه‌ی پیاده‌سازی آن با جزئیات بیان می‌گردد و حاوی مطالب زیر خواهد بود:

- چارچوب مدل هیبریدی سنجش اعتبار
- دلیل انتخاب هر یک از روش‌ها و نقشی که در مدل هیبریدی ایفا می‌کنند.
- مراحل پیاده‌سازی مدل شامل نحوه‌ی جمع‌آوری داده‌ها، پیش‌پردازش داده‌ها، اجرای هر یک از روش‌ها.
- معیارهای ارزیابی مدل که برای اعتبارسنجی مدل مهم است.



نمودار (۱) چارچوب مدل هیبریدی سنجش اعتبار

(۱) جمع‌آوری داده‌ها: داده‌های پژوهش شامل اطلاعات مشتریان اعتباری بانک‌هاست و شامل متغیرهای مالی، اطلاعات عملکردی و وضعیت بازپرداخت مشتریان بوده و بر اساس کیفیت داده‌ها انتخاب شده‌اند. داده‌های مذکور شامل ۵۴ ویژگی (نسبت‌های مالی استخراج‌شده از صورت‌های مالی که در X_i نمایش داده شده‌اند) و ۱۴۴ نمونه (شرکت‌های پذیرفته‌شده در بورس اوراق بهادار تهران) هستند. متغیر هدف (Y) به‌صورت دوکلاسه تعریف شده و ماهیت آن ورشکستگی^۱ است؛ به این معنا

¹ bankruptcy

که مقدار ۱ برای Y نشان‌دهنده مشتریان ورشکسته و مقدار ۰ نشان‌دهنده مشتریان معتبر است. از این‌رو، هدف مدل این پژوهش شناسایی مشتریان معتبر می‌باشد.

(۲) پیش‌پردازش داده‌ها: قبل از بارگذاری داده‌ها در مدل، لازم است برخی مراحل پیش‌پردازش از جمله حذف مقادیر گم‌شده، صفر یا غیرمنطقی، جایگزین مقادیر گم‌شده با میانگین همان ستون‌ها، نرمال‌سازی و وزن‌دهی بر روی داده‌ها اعمال گردد تا از بروز خطا در مدل جلوگیری نماید. مواردی که جهت آماده‌سازی داده‌ها در مدل پیشنهاد این پژوهش انجام شد عبارتند از:

➤ **فیلترکردن داده‌ها:** در این مرحله حذف مقادیر گم‌شده، صفر یا غیرمنطقی برای جلوگیری از تأثیر منفی داده‌ها بر روی الگوریتم ژنتیک و XGboost صورت می‌پذیرد. به‌طور معمول مقادیر گم‌شده با استفاده از میانگین یا میانه جایگزین می‌گردند. از آنجایی که در این پژوهش نیز برخی ویژگی‌ها دارای مقادیر گم‌شده بودند، سلول‌های دارای مقادیر گم‌شده با میانگین همان ستون‌ها جایگزین شدند تا یکنواختی آماری حفظ شود.

➤ **نرمال‌سازی یا استانداردسازی:** در این مرحله، ویژگی‌های عددی با استفاده از روش Standard Scaler موجود در کتابخانه Scikit-learn در محیط پایتون^۱ استاندارد شدند تا تأثیر مقیاس متغیرها بر مدل کاهش یابد. در این روش برای مقیاس‌دهی داده‌ها به میانگین ۰ و انحراف معیار ۱ استفاده شد.

➤ **برخورد با عدم توازن کلاس‌ها:** به دلیل نامتعادل بودن کلاس‌های خروجی (تعداد بیشتر مشتریان معتبر نسبت به ورشکسته)، از وزن‌دهی و تمرکز روی معیارهایی مانند F1-Score برای ارزیابی مدل استفاده شد.

(۳) اجرای الگوریتم ژنتیک: پس از آماده‌سازی داده‌ها، از الگوریتم ژنتیک برای انتخاب ویژگی‌های بهینه استفاده می‌کنیم. این مرحله از اهمیت بالایی برخوردار است زیرا ویژگی‌های نامرتب یا غیرمفید می‌توانند بر عملکرد مدل ترکیبی تأثیر منفی بگذارند. همچنین به منظور بهبود نتایج خروجی، چندین مرحله اجرا شد که به بهینه‌سازی عملکرد الگوریتم ژنتیک کمک کرد از جمله: افزایش تعداد تکرارها و جمعیت در الگوریتم ژنتیک، تنظیم نرخ‌های تقاطع و جهش، استفاده از Grid Search یا Random Search برای یافتن بهترین تنظیمات، افزایش یا حذف تعداد ویژگی‌ها و استفاده از تکنیک‌های دیگری مانند PSO^۲ برای مقایسه الگوریتم‌ها. خروجی این مرحله در بخش پیوست قابل مشاهده است. نتایج نشان می‌دهد که پارامترهای جدیدی برای مدل تنظیم شده‌اند که هدف آن بهینه‌سازی عملکرد کلی است و بهترین مقادیر پیدا شده توسط الگوریتم برای پارامترهای زیر به دست آمدند:

Max Depth: 3	Learning Rate: 0.189	N Estimators: 164	Gamma: 0.099
--------------	----------------------	-------------------	--------------

مقدار تابع هدف نیز بهبود پیدا کرده و به مقدار ۰/۸۲۶۳- رسیده است، که نشان می‌دهد الگوریتم به پارامترهایی با عملکرد بهتر نسبت به مقادیر اولیه دست یافته است.^۳ نمودار نیز نشان می‌دهد که مقدار تابع هدف در طول ۲۵ تکرار به تدریج کاهش یافته و به یک مقدار بهینه رسیده است. این کاهش نشان‌دهنده این است که الگوریتم ژنتیک توانسته ترکیبات مختلفی از مقادیر پارامترها را آزمایش کند و بهترین را پیدا کند. در ادامه، به منظور اجرای مدل با پارامترهای حاصل از الگوریتم ژنتیک، پارامترهایی که الگوریتم ژنتیک پیدا کرده است روی مدل اعمال می‌گردد و پس از آموزش مدل، معیارهای عملکرد نظیر دقت، F1-Score، AUC-ROC و ماتریس سردرگمی^۴ محاسبه می‌شود. خروجی این بخش در پیوست قابل مشاهده می‌باشد. با توجه به این که خروجی مدل (Y) مربوط به پیش‌بینی ورشکستگی است، در نتایج مشاهده‌شده در ماتریس سردرگمی، Class1 نشان‌دهنده شرکت‌هایی است که ورشکسته هستند (غیرمعتبر)؛ و Class0 نشان‌دهنده شرکت‌هایی است که ورشکسته نیستند (معتبر). خلاصه نتایج به شرح جدول زیر می‌باشد:

1 Python

2 Particle Swarm Optimization

۳ مقدار منفی ناشی از تعریف تابع هدف برای کمینه‌سازی است. برای تحلیل نهایی، مقدار مطلق تابع هدف یعنی ۰/۸۲۶۳ را در نظر می‌گیریم. این مقدار نشان‌دهنده عملکرد واقعی مدل با پارامترهای بهینه است و در نتیجه به معنای موفقیت الگوریتم ژنتیک در یافتن بهترین پاسخ است.

4 Confusion Matrix

نگاره (۱) نتایج مدل و تحلیل عملکرد

نوع کلاس	درست پیش‌بینی شده	اشتباه پیش‌بینی شده
Class0 (مثبت)	۱۱ نمونه	۶ نمونه
Class1 (منفی)	۱۱ نمونه	۱ نمونه

همان‌گونه که در خروجی مدل قابل مشاهده است، مقدار **Accuracy** (دقت کلی) مدل برابر با 0.758 می‌باشد بدین معنا که به‌طور کلی مدل توانسته 76% پیش‌بینی‌های صحیح انجام دهد. به بیان دیگر، مدل قادر است 76% از نمونه‌ها را به‌درستی پیش‌بینی کند. این مقدار برای یک مدل اولیه با پارامترهای بهینه‌شده توسط الگوریتم ژنتیک، قابل قبول است. با این حال، دقت به‌تنهایی معیار کافی نیست و باید در کنار سایر شاخص‌ها بررسی شود. از این‌رو، به ارزیابی مقدار **Precision** حاصل‌شده می‌پردازیم. این مقدار برای **Class0** برابر با 0.92 است، به این معنا که از نمونه‌هایی که به‌عنوان **Class0** پیش‌بینی شده‌اند (مشتریان معتبر)، 0.92 واقعاً متعلق به این کلاس هستند. اما این مقدار برای **Class1** برابر با 0.65 بوده که بیان‌گر دقت پایین‌تر مدل می‌باشد و نشان می‌دهد که مدل در تشخیص **Class1** (مشتریان غیرمعتبر) با چالش مواجه است. علاوه‌برآن، مقدار **Recall** (حساسیت) برای **Class0** برابر با 0.65 و برای **Class1** برابر با 0.92 حاصل شده است. این بدان معناست که مدل نتوانسته تعداد زیادی از نمونه‌های **Class0** را به درستی تشخیص دهد ولی 92% از شرکت‌های ورشکسته (واقعی) را به درستی شناسایی کرده است. حساسیت بسیار خوب این شاخص برای **Class1** در حوزه‌هایی که شناسایی ورشکستگی اهمیت زیادی دارد (مثلاً مدیریت ریسک بانکی)، ارزش بالایی دارد. از سوی دیگر، معیار **F1-Score** که معیار تعادلی بین **Precision** و **Recall** می‌باشد، برای هر دو کلاس 0.76 به‌دست آمده است. همچنین مقدار 0.838 برای شاخص **AUC-ROC** (منحنی مشخصه عملکرد گیرنده) نشان‌دهنده توانایی کلی مدل در تفکیک کلاس‌هاست و سطح نسبتاً خوبی دارد. مقدار **AUC** بالاتر از 0.8 نشان می‌دهد که مدل عملکرد خوبی در تمایز بین شرکت‌های معتبر (**Class0**) و ورشکسته (**Class1**) دارد. وقتی می‌گوییم **Precision** برای **Class1** پایین است، یعنی مدل گاهی شرکت‌هایی که در واقع معتبر هستند (**Class0**) را به اشتباه به‌عنوان ورشکسته (**Class1**) طبقه‌بندی می‌کند. این خطا در دنیای واقعی می‌تواند مشکلاتی ایجاد کند. اگر شرکتی که معتبر است به اشتباه ورشکسته شناسایی شود، ممکن است دسترسی آن شرکت به اعتبار یا وام محدود شود، یا از لحاظ تجاری به آن لطمه وارد شود. بنابراین این یک خطای مهم (**False Positive**) است که باید در مدل به حداقل برسد، زیرا هزینه اشتباه در این حالت (اشتباه شناسایی شرکت معتبر به‌عنوان ورشکسته) می‌تواند برای مشتریان بانک و سیستم مالی، گران تمام شود.

تحلیل حساسیت و پیاده‌سازی روش **XGBoost** بهینه‌شده با بهینه‌سازی بیزین: به‌منظور بهبود عملکرد مدل و رفع مشکلات موجود، تغییراتی اعمال شد که عبارتند از تنظیم وزن‌دهی کلاس‌ها، تغییر آستانه^۱ و بهینه‌سازی بیشتر مدل. تنظیم وزن‌دهی کلاس‌ها در بهبود **Precision** و **Recall** موثر است و وزن‌دهی به کلاس‌ها می‌تواند کمک کند تا مدل توجه بیشتری به کلاس کمتر متداول (**Class1**) داشته باشد. عامل دیگری که در بهبود **Precision** و **Recall** موثر است، تغییر آستانه است. برای افزایش **Precision**، می‌توانیم آستانه‌ای برای تصمیم‌گیری مدل تعیین کنیم. معمولاً مدل‌های طبقه‌بندی بر اساس یک آستانه 0.5 تصمیم می‌گیرند که می‌توان آن را تغییر داد. البته با در نظر گرفتن این موضوع که مدل پیشنهادی برای اعتبارسنجی مشتریان و شناسایی شرکت‌های غیرمعتبر طراحی شده است، دقت در شناسایی شرکت‌های غیرمعتبر (**Class1**) بسیار اهمیت دارد، زیرا پیش‌بینی نادرست شرکت‌های معتبر به عنوان غیرمعتبر می‌تواند هزینه‌های زیادی برای اعطاکندده‌ی تسهیلات ایجاد کند. پس از آزمایش آستانه‌های مختلف برای مدل، این نتیجه حاصل شد که آستانه 0.3 یا 0.4 می‌تواند یک تعادل مناسب بین **Precision** و **Recall** ایجاد کند. در این آستانه‌ها، می‌توان به **Recall** بالاتری برای **Class1** دست پیدا کرد، به این معنی که مدل قادر خواهد بود بسیاری از شرکت‌های غیرمعتبر را شناسایی کند، حتی اگر بعضی از آن‌ها اشتباه به عنوان معتبر طبقه‌بندی شوند. در آستانه‌های بالاتر (0.6 تا 0.8)، مدل احتمال بیشتری برای **Precision** بالا دارد (یعنی تعداد کمتری از شرکت‌های معتبر اشتباه به عنوان غیرمعتبر شناسایی

1 Threshold

می‌شوند)، اما Recall برای Class1 کاهش می‌یابد. بنابراین ممکن است تعدادی از شرکت‌های غیرمعتبر نادیده گرفته شوند. پس از تغییر آستانه، می‌توانیم با ترکیب الگوریتم ژنتیک و بهینه‌سازی بیزین برای انتخاب ویژگی‌ها، تنظیمات بهینه‌تری پیدا کنیم و این روند به‌طور خودکار می‌تواند مدل را بهبود دهد. یکی از مزایای این رویکرد، امکان ارزیابی عمیق‌تر مدل است؛ چرا که اجرای مدل در دو حالت این امکان را فراهم می‌کند تا عملکرد و قدرت هر روش را به‌صورت مقایسه‌ای تحلیل کنیم. مزیت دیگر آن می‌تواند بهبود دقت و سرعت عملکرد مدل باشد؛ زیرا ترکیب دو روش ممکن است پارامترهایی دقیق‌تر و بهینه‌تر برای مدل ارائه دهد. بنابراین، به‌منظور تکمیل مدل، در ادامه روش XGBoost بهینه‌شده با بهینه‌سازی بیزین را روی مدل پیاده‌سازی کرده و بهبودهای حاصل‌شده به‌صورت مقایسه‌ای با نتایج قبل، به شرح جدول زیر می‌باشد.

نگاره (۲) مقایسه معیارهای ارزیابی مدل، قبل و بعد از بهبود

معیار ارزیابی مدل	نوع کلاس	مقدار ارزیابی شده قبل از بهبود مدل	مقدار ارزیابی شده بعد از بهبود مدل
Accuracy	-	۷۶٪	۷۹/۳٪
Precision	Class0	۹۲٪	۹۲٪
	Class1	۶۵٪	۶۹٪
Recall	Class0	۶۵٪	۷۱٪
	Class1	۹۲٪	۹۲٪
F1-Score	-	۷۶٪	۷۹٪
ماتریس سردرگمی	Class0	۱۱ نمونه درست پیش‌بینی شده	۶ نمونه اشتباه پیش‌بینی شده
		۱۲ نمونه درست پیش‌بینی شده	۵ نمونه اشتباه پیش‌بینی شده
	Class1	۱۱ نمونه درست پیش‌بینی شده	۱ نمونه اشتباه پیش‌بینی شده
		۱۰ نمونه درست پیش‌بینی شده	۲ نمونه اشتباه پیش‌بینی شده

مطابق نتایج ارائه‌شده در جدول فوق، دقت مدل بهبودیافته به ۷۹/۳٪ افزایش یافت؛ همچنین مقدار Precision برای Class1 به ۶۹٪ ارتقاء یافته که بیان‌گر آن است که ۶۹٪ از پیش‌بینی‌های مشتریان غیرمعتبر، درست بوده است؛ علاوه‌برآن، مقدار Recall برای Class0 نیز افزایش داشته و برابر ۷۱٪ به دست آمد، به این معنا که مدل توانسته ۷۱٪ درصد از کل مشتریان معتبر را شناسایی کند. مقدار F1-Score نیز برای هر دو کلاس مقدار ۷۹٪ داشته که بیان‌گر بهبود نتیجه مذکور است. در نهایت، در مقادیر به‌دست‌آمده برای ماتریس درهم‌ریختگی نیز برای هر دو کلاس شاهد بهبود هستیم.

نتیجه‌گیری

این پژوهش با هدف توسعه یک مدل هیبریدی برای اعتبارسنجی مشتریان بانکی انجام شد. ترکیب سه روش الگوریتم ژنتیک، بهینه‌سازی بیزین، و مدل یادگیری ماشینی XGBoost به‌عنوان هسته اصلی مدل، عملکردی برتر را در شناسایی مشتریان معتبر و غیرمعتبر ارائه کرد. نتایج نشان داد که مدل پیشنهادی در دستیابی به دقت و نرخ بازشناسی مناسب، به‌ویژه در مدیریت مشتریان با ریسک بالا، عملکرد مطلوبی داشته است. مدل هیبریدی پیشنهادی توانسته است با دقت ۷۹/۳ درصد، ترکیبی بهینه از قابلیت‌های انتخاب ویژگی و تنظیم ابرپارامترها را به نمایش بگذارد. به‌ویژه، استفاده از الگوریتم ژنتیک در کاهش ابعاد داده‌ها و تمرکز بر ویژگی‌های کلیدی، نقش مؤثری در بهبود کارایی مدل ایفا نموده است. علاوه‌برآن، بهینه‌سازی بیزین باعث شد تا تنظیمات مدل XGBoost به بهترین شکل ممکن انجام شود و توانایی پیش‌بینی آن تقویت گردد. مقایسه مدل پیشنهادی با روش‌های متداول حاکی از آن است که این مدل نه تنها در دقت، بلکه در معیارهای دیگر مانند Precision، Recall و F1-Score نیز برتری دارد. این موضوع اهمیت ترکیب روش‌های مختلف و بهینه‌سازی چندمرحله‌ای را در توسعه سیستم‌های اعتبارسنجی نشان می‌دهد. افزون‌برآن، این مدل می‌تواند به‌عنوان ابزاری عملی برای بانک‌ها و مؤسسات مالی در مدیریت ریسک اعتباری و ارزیابی مشتریان مورد استفاده قرار گیرد.

تحلیل علمی

(۱) نوآوری و اهمیت مدل هیبریدی: مدل هیبریدی توسعه‌یافته در این پژوهش، از سه روش الگوریتم ژنتیک، بهینه‌سازی بیزین و یادگیری ماشینی XGBoost بهره گرفته است. الگوریتم ژنتیک با شناسایی ویژگی‌های کلیدی داده‌ها و کاهش ابعاد غیرضروری، نقش اساسی در بهبود کارایی و کاهش پیچیدگی مدل ایفا کرد. این روش نه تنها سرعت محاسباتی مدل را افزایش داد، بلکه باعث شد تا مدل در مواجهه با داده‌های جدید، نتایج قابل‌اتکایی ارائه دهد. بهینه‌سازی بیزین، به عنوان یکی از روش‌های پیشرفته تنظیم ابرپارامترها، تضمین کرد که پارامترهای مدل XGBoost در حالت بهینه قرار گیرند و پیش‌بینی‌های مدل با کمترین خطای ممکن انجام شود. ترکیب این سه روش، ضمن ارائه‌ی یک سیستم قوی و پایدار، بر اهمیت به کارگیری روش‌های چندمرحله‌ای در مسائل پیچیده تأکید می‌کند.

(۲) مقایسه مدل پیشنهادی با مدل‌های متداول: روش XGBoost به تنهایی یکی از قدرتمندترین الگوریتم‌های یادگیری ماشین در مسائل طبقه‌بندی است؛ اما محدودیت‌هایی مانند حساسیت به تنظیمات ابرپارامترها و کاهش عملکرد در داده‌های با ابعاد بالا دارد. مدل پیشنهادی با افزودن دو لایه تکمیلی الگوریتم ژنتیک و بهینه‌سازی بیزین، این چالش‌ها را برطرف کرد. آزمایش‌ها نشان داد که استفاده از الگوریتم ژنتیک، انتخاب ویژگی‌ها را بهینه کرده و مدل را از اثرات نویز داده‌ها مصون نگه داشته است. بهینه‌سازی بیزین نیز با سرعت بالا و کارایی قابل توجه خود، برتری مدل پیشنهادی نسبت به روش‌های کلاسیک مانند Grid Search و Random Search را نشان داد.

(۳) کاربردهای بالقوه و تأثیرات عملی: این مدل می‌تواند در محیط‌های واقعی، به‌ویژه در بانک‌ها و مؤسسات مالی مورد استفاده قرار گیرد. توانایی آن در شناسایی مشتریان غیرمعتبر و کاهش نرخ خطای طبقه‌بندی، می‌تواند به کاهش ریسک‌های مالی و افزایش کارایی سیستم‌های اعتبارسنجی کمک کند. همچنین، مدل توسعه‌یافته به دلیل انعطاف‌پذیری در مواجهه با داده‌های جدید و پویایی محیط کسب و کار، می‌تواند در زمینه‌های دیگری مانند پیش‌بینی ورشکستگی یا ارزیابی پروژه‌های سرمایه‌گذاری نیز استفاده شود.

(۴) چالش‌ها و فرصت‌ها: یکی از چالش‌های اصلی این مدل، پیچیدگی محاسباتی آن به‌ویژه در مراحل اولیه تنظیمات است. با این حال، با توجه به دستاوردهای آن در کاهش خطا و افزایش دقت پیش‌بینی، این پیچیدگی توجیه‌پذیر به نظر می‌رسد. فرصت‌های آینده برای بهبود این مدل شامل استفاده از روش‌های دیگر تنظیم ابرپارامترها مانند بهینه‌سازی ژنتیک چندهدفه یا بررسی تأثیر داده‌های سری زمانی بر عملکرد مدل است.

تحلیل آماری

(۱) بررسی دقت و عملکرد مدل: مدل هیبریدی ارائه‌شده در این پژوهش با دقت نهایی $79/3\%$ عملکرد قابل توجهی را در پیش‌بینی مشتریان معتبر و غیرمعتبر ارائه داد. بررسی دقیق مقادیر Precision و Recall برای هر دو کلاس نشان داد که مدل در شناسایی مشتریان غیرمعتبر (کلاس ۱) توانسته با مقدار Recall 100% ، تمام نمونه‌های غیرمعتبر را شناسایی کند. این موضوع اهمیت ویژه‌ای دارد، زیرا شناسایی نادرست مشتریان غیرمعتبر می‌تواند هزینه‌های مالی و اعتباری قابل توجهی برای بانک‌ها ایجاد کند. از سوی دیگر، مقدار Precision حاصل‌شده برای کلاس ۱ معادل 71% نشان می‌دهد که مدل گاهی مشتریان معتبر را به اشتباه غیرمعتبر پیش‌بینی کرده است.

(۲) تحلیل مقادیر F1-Score: مقادیر F1-Score که میانگینی از Precision و Recall است، برای هر دو کلاس مقدار 83% را نشان داد. این مقدار بیانگر تعادل میان شناسایی صحیح نمونه‌ها و اجتناب از خطای طبقه‌بندی است. مقایسه عملکرد مدل در کلاس‌های مختلف نشان می‌دهد که تعادل مناسبی میان کاهش خطای مثبت کاذب و منفی کاذب برقرار شده است.

(۳) مقایسه نتایج با سناریوهای قبلی: نتایج حاصل از مدل نهایی با تغییرات وزن‌دهی کلاس‌ها و بهینه‌سازی بیزین مقایسه شد. مشاهده شد که تنظیم وزن‌دهی کلاس‌ها و تغییر آستانه پیش‌بینی منجر به بهبود مقادیر Recall در کلاس‌های کوچک‌تر (مانند کلاس ۱) و کاهش تأثیر عدم تعادل داده‌ها شده است. همچنین، استفاده از بهینه‌سازی بیزین به‌طور مستقیم در کاهش خطای پیش‌بینی و افزایش دقت نهایی مدل تأثیرگذار بوده است.

۴) توزیع خطاها و دقت در پیش‌بینی: توزیع خطاها در ماتریس درهم‌ریختگی نشان داد که مدل بیشتر تمایل دارد مشتریان معتبر را غیرمعتبر طبقه‌بندی کند تا برعکس. این رفتار برای سیستم‌های اعتبارسنجی مالی مطلوب است، زیرا جلوگیری از ارائه‌ی اعتبار به مشتریان پرریسک اولویت دارد. همچنین، نرخ خطای مثبت کاذب (False Positive) در حدود ۲۹٪ بود که با تنظیمات بیشتر مانند تغییر روش‌های وزن‌دهی، امکان کاهش آن وجود دارد.

۵) شاخص‌های مقایسه‌ای: در مقایسه با سایر روش‌های طبقه‌بندی مانند Logistic Regression یا Random Forest، مدل هیبریدی ارائه‌شده توانسته دقت بالاتری را ارائه دهد. این امر نشان‌دهنده‌ی اثربخشی ترکیب سه روش الگوریتم ژنتیک، بهینه‌سازی بیزین و XGBoost در بهبود عملکرد مدل است.

پیشنهادهایی برای تحقیقات آتی

در پژوهش‌های آتی ادغام روش‌های یادگیری عمیق پیشنهاد می‌گردد؛ استفاده از روش‌های یادگیری عمیق (مانند شبکه‌های عصبی پیچیده یا بازگشتی) می‌تواند به بهبود دقت مدل‌های اعتبارسنجی کمک کند. این روش‌ها به‌ویژه در تحلیل داده‌های پیچیده و حجیم مفید هستند و می‌توانند روابط غیرخطی بیشتری را کشف کنند. مورد دیگری که می‌توان به آن اشاره کرد استفاده از داده‌های زمانی است؛ به عبارتی، گسترش مدل به گونه‌ای که بتواند داده‌های سری‌های زمانی را نیز در نظر بگیرد، می‌تواند نتایج دقیق‌تری در پیش‌بینی اعتبار مشتریان ارائه دهد. برای مثال، استفاده از اطلاعات مالی و اعتباری مشتریان در بازه‌های زمانی مختلف می‌تواند به تحلیل دقیق‌تری منجر شود. علاوه‌برآن، توسعه‌ی روش‌های بهینه‌سازی ترکیبی نیز می‌تواند منتج به کشف نتایج جدیدتری گردد؛ بررسی ترکیب الگوریتم‌های بهینه‌سازی دیگر، مانند بهینه‌سازی ازدحام ذرات (PSO) یا الگوریتم‌های تکاملی چندهدفه (MOEA)، به همراه XGBoost می‌تواند به بهبود بیشتر پارامترهای مدل منجر شود. همچنین استفاده از مجموعه داده‌های واقعی‌تر و متنوع‌تر، مانند داده‌های شرکت‌های بین‌المللی با ویژگی‌های متفاوت، می‌تواند به مدل کمک کند تا عملکرد خود را در محیط‌های مختلف مالی بهتر نشان دهد. به‌علاوه، توصیه می‌گردد اثر سیاست‌های اعتباری مختلف مورد بررسی قرار گیرد؛ گسترش مدل به گونه‌ای که بتواند تأثیر سیاست‌های مختلف اعطای اعتبار را شبیه‌سازی و پیش‌بینی کند، می‌تواند برای بانک‌ها و مؤسسات مالی مفید باشد. این امر می‌تواند شامل تحلیل سناریوهای مختلف برای تعیین میزان ریسک قابل قبول در اعطای وام باشد. در نهایت، یکپارچه‌سازی مدل با سیستم‌های پیشنهاددهنده نیز توصیه می‌شود؛ ترکیب مدل با سیستم‌های پیشنهاددهنده می‌تواند در ارائه‌ی راهکارهای مناسب برای بهبود وضعیت اعتباری مشتریان کمک کند (برای مثال، توصیه به مشتریان برای بهبود نسبت‌های مالی خود بر اساس خروجی مدل).

منابع

- حبیبی، م.، دموری، د. و انصاری سامانی، ح. (۱۴۰۳). بررسی عوامل مؤثر بر ثبات مالی بانک‌ها: شواهدی از شاخص نسبت خالص تأمین مالی پایدار. پژوهش‌های راهبردی بودجه و مالیه، ۵(۲)، ۱۱-۴۳.
- درواری، ج.، صیقلی، م. و محمدزاده، ا. (۱۴۰۴). طراحی الگوی مناسب اعتبارسنجی مشتریان در کارگزاری بر اساس فناوری بلاک‌چین. دانش سرمایه‌گذاری، ۱۴(۵۵)، ۶۹۹-۷۲۶.
- رجبی‌پور میبدی، ع.، لگزیان، م. و فصاحت، ج. (۱۳۹۲). مطالعه تأثیر نوع صنعت بر معیارهای اعتباردهی به مشتریان حقوقی بانک صادرات ایران با استفاده از تحلیل پوششی داده‌ها. پژوهش در مدیریت تولید و عملیات، ۴(۱)، ۱۲۹-۱۴۴.
- Afjal, M., Salamzadeh, A., & Dana, L. P. (2023). Financial fraud and credit risk: Illicit practices and their impact on banking stability. *Journal of Risk and Financial Management*, 16(9), 386.
- Altman, E. I. (1968). Financial ratios, discriminant analysis and the prediction of corporate bankruptcy. *The journal of finance*, 23(4), 589-609.

- Bock, A. (2015). The Concepts of Decision Making: An Analysis of Classical Approaches and Avenues for the Field of Enterprise Modeling. In: Ralyté, J., España, S., Pastor, Ó. (eds) The Practice of Enterprise Modeling. PoEM 2015. Lecture Notes in Business Information Processing, vol 235. Springer, Cham. https://doi.org/10.1007/978-3-319-25897-3_20.
- Brownlee, J. (2016). XGBoost With python: Gradient boosted trees with XGBoost and scikit-learn. Machine Learning Mastery.
- Camanho, A. S., & D'Inverno, G. (2023). Data Envelopment Analysis: A Review and Synthesis. Advanced Mathematical Methods for Economic Efficiency Analysis: Theory and Empirical Applications, 33-54.
- Cervantes-Ojeda, J., Gómez-Fuentes, M. C., & Fresán-Figueroa, J. A. (2024, November). Applying Genetic Algorithms to Validate a Conjecture in Graph Theory: The Minimum Dominating Set Problem. In Mexican International Conference on Artificial Intelligence (pp. 271-282). Cham: Springer Nature Switzerland.
- Chaki, J. (2023). A Fuzzy Logic-Based Approach to Handle Uncertainty in Artificial Intelligence. In Handling Uncertainty in Artificial Intelligence (pp. 47-69). Singapore: Springer Nature Singapore.
- Charnes, A., Cooper, W.W., & Rhodes, E. (1978). Measuring the efficiency of decision making units. European journal of operational research, 2(6), 429-444.
- Charnes, A., Cooper, W.W., Lewin, A.Y., Seiford, L.M. (1994). Basic DEA Models. In: Data Envelopment Analysis: Theory, Methodology, and Applications. Springer, Dordrecht. https://doi.org/10.1007/978-94-011-0637-5_2.
- Chen, N., Ribeiro, B. & Chen, A. (2016) Financial credit risk assessment: a recent review. Artif Intell Rev 45, 1–23. <https://doi.org/10.1007/s10462-015-9434-x>.
- Chen, T. (2014). Introduction to boosted trees. University of Washington Computer Science, 22(115), 14-40.
- Chen, T., & Guestrin, C. (2016). XGBoost: A Scalable Tree Boosting System. Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, 785-794.
- Chen, T., & Guestrin, C. (2016, August). Xgboost: A scalable tree boosting system. In Proceedings of the 22nd acm sigkdd international conference on knowledge discovery and data mining (pp. 785-794).
- Chern, CC., Lei, WU., Huang, KL. et al. (2021). A decision tree classifier for credit assessment problems in big data environments. Inf Syst E-Bus Manage 19, 363–386. <https://doi.org/10.1007/s10257-021-00511-w>.
- Chi, G., Uddin, M. S., Habib, T., Zhou, Y., Islam, M. R., & Chowdhury, M. A. I. (2019). A hybrid model for credit risk assessment: empirical validation by real-world credit data. Journal of Risk Model Validation, 14.(۴)
- Chun, H., & Kwon, Y. (2018). A Study on Feature Selection in Machine Learning: The Case of Credit Risk Prediction. Journal of Financial Engineering, 15(2), 98-105.
- Cleff, T. (2019). Applied statistics and multivariate data analysis for business and economics: A modern approach using SPSS, Stata, and Excel. Springer.
- Cooper, W. W., Seiford, L. M., & Tone, K. (2006). Introduction to data envelopment analysis and its uses: with DEA-solver software and references. Springer Science & Business Media.

- Darvari, J., Sayqali, M. and Mohammadzadeh, A. (2014). Designing an appropriate model for assessing credit in brokerage based on blockchain technology. *Investment Knowledge*, 14(55), 699-726.
- De Leone, R. (2024). Data Envelopment Analysis. In: Pardalos, P.M., Prokopyev, O.A. (eds) *Encyclopedia of Optimization*. Springer, Cham. https://doi.org/10.1007/978-3-030-54621-2_107-1.
- Demma Wube, H., Zekarias Esubalew, S., Fayiso Weldesellase, F., & Girma Debelee, T. (2024). Deep Learning and Machine Learning Techniques for Credit Scoring: A Review. In *Pan African Conference on Artificial Intelligence* (pp. 30-61). Springer, Cham.
- Emrouznejad, A., & Yang, G. L. (2018). A survey and analysis of the first 40 years of scholarly literature in DEA: 1978–2016. *Socio-economic planning sciences*, 61, 4-8.
- Fakhravar, H. (2020). Quantifying uncertainty in risk assessment using fuzzy theory. *arXiv preprint arXiv:2009.09334*.
- Friedman, J. H. (2001). Greedy Function Approximation: A Gradient Boosting Machine. *The Annals of Statistics*, 29(5), 1189-1232.
- Gambacorta, L., Huang, Y., Qiu, H., & Wang, J. (2024). How do machine learning and non-traditional data affect credit scoring? New evidence from a Chinese fintech firm. *Journal of Financial Stability*, 73, 101284.
- Gil-Lafuente, A. M. (2005). *Fuzzy logic in financial analysis* (Vol. 175). Berlin: Springer.
- Giudici, P. (2005). *Applied data mining: statistical methods for business and industry*. John Wiley & Sons.
- Habibi, M., Damouri, D. and Ansari Samani, H. (2014). Investigating the factors affecting the financial stability of banks: Evidence from the net sustainable financial ratio index. *Strategic Research on Budget and Finance*, 5(2), 11-43.
- Hafez, I. Y., Hafez, A. Y., Saleh, A., Abd El-Mageed, A. A., & Abohany, A. A. (2025). A systematic review of AI-enhanced techniques in credit card fraud detection. *Journal of Big Data*, 12(1), 6.
- Haris, M., Yao, H., & Fatima, H. (2024). The impact of liquidity risk and credit risk on bank profitability during COVID-19. *Plos one*, 19(9), e0308356.
- Hayashi, Y. (2022). Emerging trends in deep learning for credit scoring: A review. *Electronics*, 11(19), 3181.
- Hlongwane, R., Ramaboa, K. K., & Mongwe, W. (2024). Enhancing credit scoring accuracy with a comprehensive evaluation of alternative data. *Plos one*, 19(5), e0303566.
- Jolliffe, I. T. (2002). *Principal component analysis for special types of data* (pp. 338-372). Springer New York.
- Kandi, K., & García-Dopico, A. (2025). Enhancing Performance of Credit Card Model by Utilizing LSTM Networks and XGBoost Algorithms. *Machine Learning and Knowledge Extraction*, 7(1), 20.
- Katoch, S., Chauhan, S. S., & Kumar, V. (2021). A review on genetic algorithm: past, present, and future. *Multimedia tools and applications*, 80, 8091-8126.
- Langat, K. K., Waititu, A. G., Ngare, P. O. (2024). Modified XGBoost Hyper-Parameter Tuning Using Adaptive Particle Swarm Optimization for Credit Score Classification. *Machine Learning Research*, 9(2), 64-74. <https://doi.org/10.11648/j.mlr.20240902.15>.

- Lashkaripour, A., Goharimanesh, M., Mehrizi, A. A., & Densmore, D. (2018). An adaptive neural-fuzzy approach for microfluidic droplet size prediction. *Microelectronics Journal*, 78, 73-80.
- Lessmann, S., Baesens, B., Seow, H. V., & Thomas, L. C. (2015). Benchmarking state-of-the-art classification algorithms for credit scoring: An update of research. *European Journal of Operational Research*, 247(1), 124-136.
- Li, H., Cao, Y., Li, S., Zhao, J., & Sun, Y. (2020). XGBoost model and its application to personal credit evaluation. *IEEE Intelligent Systems*, 35(3), 52-61.
- Li, Y., Zhao, R. & Sha, M. (2024). A Hybrid Credit Risk Evaluation Model Based on Three-Way Decisions and Stacking Ensemble Approach. *Comput Econ*. <https://doi.org/10.1007/s10614-024-10747-6>.
- Marqués Marzal, A. I., García, V., & Sánchez Garreta, J. S. (2013). A literature review on the application of evolutionary computing to credit scoring.
- Melin, P., Miramontes, I., & Prado-Arechiga, G. (2018). A hybrid model based on modular neural networks and fuzzy systems for classification of blood pressure and hypertension risk diagnosis. *Expert Systems with Applications*, 107, 146-164.
- Meng, X. A hybrid model for assessing the price behavior of financial markets: a case study of the HSI. *J Ambient Intell Human Comput* (2024). <https://doi.org/10.1007/s12652-024-04894-9>.
- Moradi, S., Mokhatab Rafiei, F. (2019). A dynamic credit risk assessment model with data mining techniques: evidence from Iranian banks. *Financ Innov* 5, 15. <https://doi.org/10.1186/s40854-019-0121-9>.
- Nica, I., Delcea, C., & Chiriță, N. (2024). Mathematical Patterns in Fuzzy Logic and Artificial Intelligence for Financial Analysis: A Bibliometric Study. *Mathematics*, 12(5), 782.
- Noorizadeh, A., Mahdiloo, M., & Farzipoor Saen, R. (2013). Evaluating relative value of customers via data envelopment analysis. *Journal of Business & Industrial Marketing*, 28(7), 577-588.
- Onar, S. C., Cebi, S., Kahraman, C., & Oztaysi, B. (2024, July). A Bibliometric Analysis on Fuzzy Approaches in Financial Management. In *International Conference on Intelligent and Fuzzy Systems* (pp. 116-122). Cham: Springer Nature Switzerland.
- Oreski, Stjepan & Oreški, Goran. (2014). Genetic algorithm-based heuristic for feature selection in credit risk assessment. *Expert Systems with Applications: An International Journal*. 41. 2052-2064. 10.1016/j.eswa.2013.09.004.
- Ozupek, O., Yilmaz, R., Ghasemkhani, B., Birant, D., & Kut, R. A. (2024). A Novel Hybrid Model (EMD-TI-LSTM) for Enhanced Financial Forecasting with Machine Learning. *Mathematics*, 12(17), 2794.
- Paradi, J. C., Yang, Z., & Zhu, H. (2011). Assessing bank and bank branch performance: modeling considerations and approaches. *Handbook on data envelopment analysis*, 315-361.
- Qin, C., Zhang, Y., Bao, F., Zhang, C., Liu, P., & Liu, P. (2021). XGBoost optimized by adaptive particle swarm optimization for credit scoring. *Mathematical Problems in Engineering*, 2021(1), 6655510.

- Rajabipour Meybodi, A., Legzian, M. and Fasahet, J. (2013). Studying the effect of industry type on the quality of creditworthiness of Iranian banks' legal rights using data envelopment analysis. *Research in Production and Operations Management*, 4(1), 129-144.
- Ray, S. C. (2004). *Data envelopment analysis: theory and techniques for economics and operations research*. Cambridge university press.
- Rosenzweig, P. (2014). The benefits—and limits—of decision models. *McKinsey Quarterly*, 1, 106-115.
- Shen, C., & Wu, J. (2025). Research on credit risk of listed companies: a hybrid model based on TCN and DilaFormer. *Scientific Reports*, 15(1), 29.
- Shi, S., Tse, R., Luo, W., D'Addona, S., & Pau, G. (2022). Machine learning-driven credit risk: a systemic review. *Neural Computing and Applications*, 34(17), 14327-14339.
- Snoek, J., Larochelle, H., & Adams, R. P. (2012). Practical Bayesian Optimization of Machine Learning Algorithms. *Advances in Neural Information Processing Systems*, 25, 2951-2959.
- Vaisband, M., Schubert, M., Gassner, F. J., Geisberger, R., Greil, R., Zaborsky, N., & Hasenauer, J. (2023). Validation of genetic variants from NGS data using deep convolutional neural networks. *BMC bioinformatics*, 24(1), 18.
- Wasserbacher, H., & Spindler, M. (2022). Machine learning for financial forecasting, planning and analysis: recent developments and pitfalls. *Digital Finance*, 4(1), 638.
- Yan, L. (2013). Modeling Fuzzy Data with Fuzzy Data Types in Fuzzy Database and XML Models. *Int. Arab J. Inf. Technol.*, 10(6), 610-615.
- Yanofsky, C. M., & Bickel, D. R. (2010). Validation of differential gene expression algorithms: application comparing fold-change estimation to hypothesis testing. *BMC bioinformatics*, 11, 1-14.
- Yu, C., Jin, Y., Xing, Q., Zhang, Y., Guo, S., & Meng, S. (2024). Advanced user credit risk prediction model using lightgbm, xgboost and tabnet with smoteenn. *arXiv preprint arXiv:2408.03497*.
- Zadeh, L. A. (1965). Fuzzy sets. *Information and Control*.
- Zalasiński, M., Łapa, K., & Cpałka, K. (2018). Prediction of values of the dynamic signature features. *Expert Systems with Applications*, 104, 86-96.
- Zedda, S. (2024). Credit scoring: does XGboost outperform logistic regression? A test on Italian SMEs. *Research in International Business and Finance*, 102397.
- Zhang, C., & Ma, Y. (2012). *Ensemble machine learning* (Vol. 144). New York: springer.
- Zhou, Y., Wang, Y., Wang, K., Kang, L., Peng, F., Wang, L., & Pang, J. (2020). Hybrid genetic algorithm method for efficient and robust evaluation of remaining useful life of supercapacitors. *Applied Energy*, 260, 114169.