



مروری بر طبقه‌بندی جامع مدل‌های سنجش اعتبار

پردیس فولادی^۱

محسن امینی‌خوزانی^۲

زهره حاجیها^۳

شادی شاهرودیانی^۴

تاریخ پذیرش مقاله: ۱۴۰۳/۰۵/۰۳ تاریخ دریافت مقاله: ۱۴۰۳/۰۹/۱۰

چکیده

مدل‌های سنجش اعتبار به‌طور قابل‌توجهی در طول زمان تکامل یافته‌اند و تکنیک‌های مختلفی را برای ارزیابی اعتبار افراد و کسب‌وکارها ترکیب می‌کنند. امتیازدهی اعتباری فرآیند با اهمیتی برای پرداخت‌کنندگان اعتبار به شمار می‌رود چراکه احتمال نکول یا عدم‌نکول متقاضیان را تعیین می‌کند و نقش مهمی در تعیین توانایی افراد برای دریافت تسهیلات و اعتبار دارد. مدل‌های سنجش اعتبار یا امتیازدهی اعتباری سنتی، مانند مدل‌هایی که تاریخچه اعتباری و رفتار بازپرداخت وام‌گیرندگان را در برمی‌گیرد، برای دهه‌ها در مدیریت ریسک اعتباری موردتوجه بوده است. این مدل‌ها با ارزیابی اعتبار، تأثیرگذاری بر تأییدیه‌های وام و شکل‌دهی اقدامات پیشگیرانه در صورت وقوع نکول، به فرآیندهای تصمیم‌گیری کمک می‌کنند. با استفاده از چارچوب‌ها و رویکردهای روبه رشد یادگیری ماشینی و منابع داده‌های جایگزین، روش‌های امتیازدهی اعتباری به‌مرورزمان در بهینه‌سازی مدیریت ریسک، بهبود فرآیندهای تأیید وام، کاهش ریسک سرمایه‌گذاری برای وام‌دهندگان و رسیدگی مؤثر به نیازهای مالی وام‌گیرندگان موفق‌تر خواهند بود. در این پژوهش به ارائه‌ی یک طبقه‌بندی جامع از مدل‌های سنجش اعتبار، با استفاده از تحلیل استنادی کتاب‌سنجی پرداخته‌شده و شامل بیش از ۱۰۰ مقاله‌ی داخلی و خارجی دارای استناد است.

کلمات کلیدی

امتیازدهی اعتباری، تحلیل آماری، رگرسیون، یادگیری ماشینی، شبکه‌های عصبی

JEL: E51, C14, C38, C5, C5.

۱- دانشجوی دکتری، گروه مهندسی مالی، واحد شهر قدس، دانشگاه آزاد اسلامی، تهران، ایران. fooladi.pardis@gmail.com

۲- استادیار، گروه مهندسی مالی، واحد شهر قدس، دانشگاه آزاد اسلامی، تهران، ایران. (نویسنده مسئول) amini_k_m@yahoo.com

۳- دانشیار، گروه حسابداری، واحد تهران شرق، دانشگاه آزاد اسلامی، تهران، ایران. drzhajiha@gmail.com

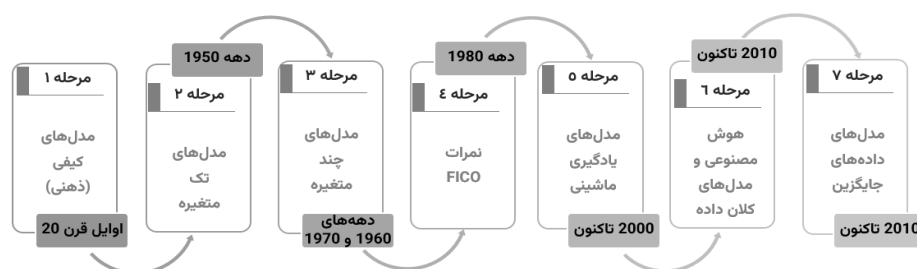
۴- استادیار، گروه مهندسی مالی، واحد شهر قدس، دانشگاه آزاد اسلامی، تهران، ایران. shshahverdiani@gmail.com

صنعت اعتبار در دو دهه‌ی اخیر رشد قابل توجهی را تجربه کرده و از مدل‌های امتیازدهی اعتباری به‌طور گسترده‌ای در صنعت مالی به‌منظور بهبود وجه نقد و کاهش ریسک اعتباری استفاده شده است. از مزایای امتیازدهی اعتباری می‌توان به کاهش هزینه‌ی تجزیه و تحلیل اعتبار، امکان تصمیم‌گیری سریع‌تر در اعطای اعتبار، نظارت دقیق‌تر بر حساب‌های موجود و اولویت‌بندی در اعطای اعتبار اشاره کرد. با رشد خدمات مالی، زیان‌های فزاینده‌ای از وام‌های معوق به اعطاکنندگان اعتبار وارد شده است. از این‌رو، بسیاری از سازمان‌ها در صنعت اعتبار، در واکنش به این زیان‌ها به توسعه‌ی مدل‌های جدیدی از اعتبارسنجی روی آورده‌اند (دیوید وست، ۲۰۰۰). در مجموع می‌توان بیان نمود که هدف از مدل‌های امتیازدهی اعتباری جدید افزایش دقت است، به این معنی که اعطای اعتبار به متقاضیان معتبر صورت گیرد و متقاضیان غیرقابل اعتبار از اعتبار محروم شوند و از این طریق موجب کاهش ضرر و زیان گردد. امتیازدهی اعتباری یکی از موفق‌ترین کاربردها و تکنیک‌های تحقیق در عملیات است که در بانکداری و امور مالی مورد استفاده قرار می‌گیرد و یکی از اولین ابزارهای مدیریت ریسک مالی به شمار می‌رود که توسعه‌یافته است (لانگ و همکاران، ۲۰۰۷). امتیازدهی اعتباری توسط فیر و آیساک در اوایل دهه ۱۹۶۰ توسعه یافت و به عبارت ساده با ایجاد امتیازی مطابقت دارد که می‌تواند برای طبقه‌بندی مشتریان به دو گروه جداگانه استفاده شود: (۱) گروه معتبر یا خوب که احتمالاً منجر به بازپرداخت وام می‌شود؛ و (۲) «گروه غیر معتبر یا بد» که احتمالاً منجر به نکول وام می‌شود (ناسیمنتو و واسکونسلس، ۲۰۱۲). قرار است ارزیابی و تحلیل ریسک در صدور وام انجام شود. قرار است اعتبار فردی که وام به او اعطا می‌شود با صدور یک امتیاز اعتباری ارزیابی و تجزیه و تحلیل شود. با صدور امتیاز اعتباری برای یک فرد حقیقی و حقوقی، تعیین‌کننده‌ی تصمیم‌گیری در خصوص اعطا/عدم اعطای وام به آن است. تکنیک‌های متعددی در حوزه‌ی امتیازدهی اعتباری در طول سال‌ها پیشنهاد و اجرا شده که از تکنیک‌های مبتنی بر آمار تا تکنیک‌های مبتنی بر هوش مصنوعی گسترده است. تکنیک‌های مبتنی بر آمار از تجزیه و تحلیل تشخیصی خطی تا رگرسیون لجستیک را شامل شده و تکنیک‌های مبتنی بر هوش مصنوعی، توانایی یادگیری از داده‌های خاص و حافظه را دارند و مدل‌های متنوعی از جمله شبکه‌های عصبی مصنوعی، الگوریتم ژنتیک و ... را شامل می‌شوند (ام‌پافو و موکاسرا، ۲۰۱۴).

سیر تطور مدل‌های کاربردی سنجش اعتبار

تکامل مدل‌های امتیازدهی اعتباری نشان‌دهنده‌ی تغییر از «ارزیابی‌های ذهنی و کیفی» به

«رویکردهای بسیار پیچیده و مبتنی بر داده» است. هر مرحله بهبودهایی را در دقت، سازگاری و جامعیت به ارمغان آورده است و به وام‌دهندگان کمک می‌کند تا تصمیمات آگاهانه‌تری بگیرند و دسترسی افراد بیشتری را به اعتبار افزایش دهند. در این بخش یک مرور زمانی از تکامل مدل‌های امتیازدهی اعتباری و تفاوت‌های آن‌ها بیان خواهیم کرد.



نمودار ۱: سیر تطور مدل‌های کاربردی سنجش اعتبار

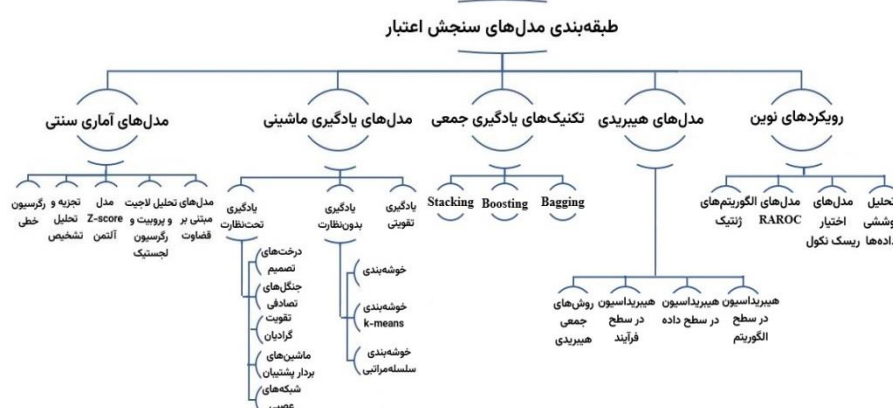
در اوایل قرن بیستم، صنعت مالی و اعتباری شاهد پیدایش مدل‌های کیفی (ذهنی) بود. در ابتدا تصمیمات اعتباری بر اساس قضاوت ذهنی توسط کارشناسان وام گرفته می‌شد که این تصمیمات به شدت بر تعاملات شخصی و ارزیابی‌های کیفی شخصیت و شهرت وام‌گیرندگان متکی بود. ویژگی بارز این مرحله، ذهنی بودن و ناسازگاری آن بود که منجر به سوءگیری‌های ذهنی و خطاهای بالقوه می‌گردید. در دهه‌ی ۱۹۵۰، مدل‌های تک متغیره که از نسبت‌ها یا معیارهای مالی منفرد (مثلاً نسبت بدهی به درآمد) برای ارزیابی اعتبار استفاده می‌کردند، پدید آمدند که مزیت آن‌ها نسبت به مرحله‌ی قبل، پیاده‌سازی ساده و آسان آن بود اما این مدل‌ها فاقد ارزیابی جامع ریسک بودند. پس از آن، مدل‌های چندمتغیره در دهه‌های ۱۹۶۰ و ۱۹۷۰ با استفاده از چندین نسبت و معیار مالی به‌طور هم‌زمان معرفی گردید. در این دوران، تکنیک‌هایی نظیر مدل‌های تجزیه و تحلیل تشخیصی، لاجیت و پروبیت رایج شدند که تفاوت بارز آن‌ها با مدل‌های تک متغیره این بود که دقیق‌تر بوده و دید وسیع‌تری از ریسک اعتباری ارائه دادند. مرحله‌ی بعد در دهه‌ی ۱۹۸۰ بود که نمرات FICO توسط Fair Isaac Corporation توسعه یافت و به استاندارد صنعت تبدیل شد. این امتیازها ترکیبی از تاریخچه‌ی اعتباری، سطوح بدهی فعلی، طول تاریخ اعتبار، اعتبار جدید و انواع اعتبار مورد استفاده بود. مشخصه‌ی این امتیازها استاندارد شدن و پذیرش گسترده‌ی آن توسط صنعت مالی و اعتباری بود که امکان ارزیابی اعتباری سازگار در بین وام‌دهندگان را فراهم نمود. از سال ۲۰۰۰ تاکنون با روی کار آمدن مدل‌های یادگیری ماشینی، صنعت شاهد ترکیب الگوریتم‌های پیشرفته مانند درختان تصمیم،

طبقه‌بندی جامع مدل‌های امتیازدهی اعتباری / فولادی، امینی خوزانی، حاجیها و شاهوردیانی

جنگل‌های تصادفی، ماشین‌های بردار پشتیبان و شبکه‌های عصبی بوده است. این مدل‌ها قادرند حجم وسیعی از داده‌ها، از جمله منابع داده غیرسنتی مانند فعالیت رسانه‌های اجتماعی و رفتار آنلاین را تجزیه و تحلیل کنند. نقطه‌ی تمایز مدل‌های موجود در این مرحله این است که بسیار دقیق و سازگار بوده و قادر به شناسایی الگوها و روندهای پیچیده در داده‌ها می‌باشند. علاوه بر آن، از سال ۲۰۱۰ تاکنون، هوش مصنوعی و مدل‌های کلان داده تحول قابل توجهی در این حوزه به همراه داشته است. پژوهشگران این زمینه، از هوش مصنوعی و تجزیه و تحلیل داده‌های بزرگ جهت امتیازدهی اعتبار استفاده کردند و موفق شدند مجموعه داده‌های عظیمی را در زمان واقعی، پردازش و تجزیه و تحلیل نموده و ارزیابی‌های اعتباری پویاتر و دقیق‌تری ارائه دهند. به بیان ساده‌تر، فراهم آوردن بینش در زمان واقعی و سازگاری با شرایط اقتصادی در حال تغییر و رفتارهای مصرف‌کننده ویژگی‌هایی هستند که می‌توان به مدل‌های این دوره نسبت داد. همچنین در همین دوره، مدل‌های داده‌های جایگزین نیز از منابع داده غیرسنتی مانند پرداخت‌های آب و برق، تاریخچه‌ی پرداخت اجاره و حتی فعالیت در رسانه‌های اجتماعی برای ارزیابی اعتبار-به‌خصوص برای افراد دارای سابقه اعتباری محدود- استفاده نموده که از مزایای آن می‌توان به فراگیرتر بودن آن و دسترسی به اعتبار برای جمعیت‌های محروم اشاره کرد.

طبقه‌بندی مدل‌های سنجش اعتبار

مدل‌های امتیازدهی اعتباری، مدل‌های کمی هستند که از ویژگی‌های مشاهده‌شده‌ی وام‌گیرنده برای محاسبه‌ی امتیازی که احتمال نکول متقاضی را نشان می‌دهد یا برای دسته‌بندی وام‌گیرندگان به طبقات مختلف ریسک نکول استفاده می‌کنند و شامل پنج نوع گسترده به شرح زیر می‌باشند:



نمودار ۲: طبقه‌بندی مدل‌های سنجش اعتبار

مدل‌های آماری سنتی^۱

برجسته‌ترین تکنیک‌های مورد استفاده برای توسعه‌ی امتیازدهی اعتباری، روش‌های تشخیصی آماری و طبقه‌بندی است و از آنجا که سطوح کاربردی و انسانی با درک نتیجه‌گیری یک مدل (به ترتیب از سوی یک متخصص یا یک فرد عادی) مرتبط هستند، سطح عملکرد مربوط به ارزیابی قوانین تصمیم‌گیری از دیدگاه آماری قابل بررسی می‌باشد. مدل‌های آماری سنتی شامل چهار مدل هستند که هر یک به تفصیل در این بخش بررسی خواهند شد.

رگرسیون خطی^۲

تجزیه و تحلیل رگرسیون مزایای بسیاری به‌ویژه در امتیازدهی اعتباری دارد زیرا رویکردی آسان به‌منظور توضیح و پیش‌بینی پارامترهای ریسک، مانند احتمال نکول، می‌باشد. رگرسیون خطی از داده‌های گذشته (تجربه‌ی بازپرداخت وام‌های قدیمی) به‌عنوان ورودی‌های مدل، به‌منظور پیش‌بینی احتمال بازپرداخت وام‌های جدید را استفاده می‌کند. مدل زیر به‌عنوان رگرسیون خطی تخمین زده می‌شود

$$PD_i = \sum_{j=1}^n \beta_j X_{ij} + error(1)$$

که در آن β_j اهمیت تخمینی متغیر z (مانند اهرم)؛ PD_i برای نکول یا عدم نکول وام‌های قدیمی (i) به ترتیب ۱ یا ۰ هستند. سپس احتمال نکول برای وام‌گیرنده به‌صورت $E(PD_i) = (1 - P_i)$ تفسیر شده و برابر با احتمال نکول مورد انتظار (P_i : احتمال بازپرداخت وام) است. (ساندرز و همکاران، ۲۰۲۱).

تجزیه و تحلیل تشخیصی^۳

درحالی‌که مدل‌های خطی، مقداری را برای احتمال نکول مورد انتظار در صورت اعطای وام پیش‌بینی می‌کنند، مدل‌های تشخیصی، وام‌گیرندگان را بسته به ویژگی‌های مشاهده‌شده (X_j)، به کلاس‌های ریسک نکول بالا یا پایین تقسیم می‌کنند. این مدل‌ها از داده‌های گذشته به‌عنوان ورودی‌های مدل برای توضیح تجربه‌ی بازپرداخت وام‌های قدیمی استفاده می‌کنند. آنالیز تشخیصی خطی توسط فیشر (۱۹۳۶) معرفی شد که به دنبال بهترین راه برای جداسازی دو گروه با استفاده از ترکیب خطی متغیرها بود. آیزنیس (۱۹۷۷) این روش را با بیان این‌که این قانون فقط برای دسته‌ی کوچکی از توزیع‌ها بهینه است، مورد انتقاد قرار داد. با این حال، هند و هنلی (۱۹۹۷) ادعا کردند که اگر

طبقه‌بندی جامع مدل‌های امتیازدهی اعتباری / فولادی، امینی خوزانی، حاجیها و شاهوردیانی

متغیرها از یک توزیع بیضوی چندمتغیره پیروی کنند (که توزیع نرمال یک مورد خاص از آن است)، آنگاه قاعده‌ی آنالیز تشخیصی خطی بهینه است.

مدل Z-score آلتمن و توسعه‌های آن

در سال ۱۹۶۸، ادوارد آلتمن موضوعی را منتشر کرد که به معروف‌ترین مدل پیش‌بینی‌کننده‌ی ورشکستگی تبدیل شد. این پیش‌بینی‌کننده، یک مدل آماری به نام Z-score است که پنج نسبت مالی را برای تولید یک محصول ترکیب می‌کند و ثابت شده که این مدل ابزار قابل‌اعتمادی در پیش‌بینی ورشکستگی برای ترکیب متنوعی از واحدهای تجاری است (آنجم، ۲۰۱۲). مدل اصلی آلتمن

$$Z = 0.012 (X_1) + 0.014 (X_2) + 0.033 (X_3) + 0.006 (X_4) + 0.999 (X_5) \quad (۲)$$

است که در آن X_1 نسبت سرمایه در گردش به کل دارایی‌ها؛ X_2 نسبت سود انباشته به کل دارایی‌ها؛ X_3 نسبت سود قبل از بهره و مالیات (EBIT) به کل دارایی‌ها؛ X_4 نسبت ارزش بازار حقوق صاحبان سهام به ارزش دفتری کل بدهی و X_5 نسبت فروش به کل دارایی می‌باشد (آلتمن، ۲۰۰۰ و ۲۰۰۲، چوواخین و جرمانیا، ۲۰۰۳). در سال ۱۹۸۳، آلتمن مدل Z-score اصلاح‌شده را برای شرکت‌های خصوصی توسعه داد. این مدل، ارزش دفتری حقوق صاحبان سهام را جایگزین ارزش بازار در X_4 کرد و به شرح زیر است:

$$Z = 0.717 (X_1) + 0.847 (X_2) + 3.107 (X_3) + 0.42 (X_4) + 0.998 (X_5) \quad (۳)$$

آلتمن مدل اصلی و مدل زتای اصلاح‌شده‌ی خود را کامل ندانسته و به چند موضوع از جمله وجود ذهنیت در وزن دهی‌ها، ابهام در مدل، رویکرد تک متغیره و برخی نسبت‌های گمراه‌کننده اشاره کرد. علاوه بر آن، از دیدگاه او، نسبت پنجم (فروش به کل دارایی‌ها) تفاوتی بین شرکت‌های شکست‌خورده و شکست‌نخورده و همچنین از صنعتی به صنعت دیگر نشان نمی‌دهد و مدل قادر به پیش‌بینی دقیق مشکلات مالی برای شرکت‌های غیر تولیدی و فرم‌های غیرعمومی نیست (شفر، ۲۰۰۰). مدل Z-score تجدیدنظر شده از نسبت ارزش دفتری به کل بدهی‌ها برای X_4 استفاده می‌کند تا قابلیت کاربرد خود را برای شرکت‌های خصوصی حفظ کند. سه متغیر اول بدون تغییر هستند؛ با این حال، فاکتور وزن دوباره محاسبه می‌شود. از این‌رو مدل Z-score اصلاح‌شده در سال ۱۹۹۳ به صورت

$$Z = 6.56 (X_1) + 3.26 (X_2) + 6.72 (X_3) + 1.05 (X_4) \quad (۴)$$

معرفی شد که در آن، امتیاز Z برای شرکت‌های ورشکسته کمتر از ۱/۱۰، برای شرکت‌های غیر ورشکسته بیشتر از ۲/۶۰ و منطقه‌ی خاکستری نیز بین ۱/۱۰ و ۲/۶۰ می‌باشد که نتایج مدل

نهایی، از دقت بالایی در پیش‌بینی ورشکستگی برخوردار است.

تحلیل لاجیت و پروبیت و رگرسیون لجستیک

مدل‌های تحلیل لاجیت^۴

از آنجاکه توزیع داده‌های اعتباری معمولاً غیر نرمال بوده و این واقعیت ممکن است از نظر تئوری مشکلی را در هنگام انجام LDA ایجاد کند، یکی از راه‌های غلبه بر مشکلات مربوط به غیر نرمال بودن داده‌ها، استفاده از یک مدل خطی تعمیم‌یافته از LDA است که امکان توزیع پارامتری را فراهم می‌کند و به مدل لاجیت معروف است. با توجه به بردار ویژگی‌های x ، احتمال نکول p با رابطه

$$\log\left(\frac{p}{1-p}\right) = w_0 + \sum w_i \log x_i \quad (5)$$

به بردار x مربوط می‌شود. از مزایای آن نسبت به آنالیز تشخیصی خطی، می‌توان به استفاده از روش حداکثر درست‌نمایی برای تخمین پارامترهای w_i ، فراهم کردن احتمالات حضور در یک کلاس مشخص و امکان سروکار داشتن با داده‌های طبقه‌بندی‌شده اشاره کرد. مطالعاتی که در این خصوص انجام شده، نشان داده‌اند که مدل لاجیت از روش آنالیز تشخیصی خطی بهتر عمل می‌کند (ویگینتون، ۱۹۸۰). اخیراً، تحلیل لاجیت به یکی از رویکردهای اصلی طبقه‌بندی در امتیازدهی اعتباری در عملکرد بانک‌ها تبدیل شده است. ضرایب به‌دست‌آمده دارای همان مقادیری هستند که در مطالعات از قانون تصمیم‌گیری LDA استفاده می‌کردند. از معایب این روش این است که به همبستگی بالا بین متغیرهای توضیحی حساس بوده و باید اطمینان حاصل کرد که هیچ متغیری در مجموعه آموزشی باقی نمانده است. همچنین این روش به مقادیر ازدست‌رفته حساس بوده و تمام مشاهدات دارای مقادیر گم‌شده باید حذف شوند. لارنس و ارشدی (۱۹۹۵) از مدل لاجیت برای تحلیل مدیریت وام‌های مشکل‌دار استفاده کردند. در حوزه‌ی وام مسکن، کمپبل و دیتریش (۱۹۸۳) از یک مدل لاجیت استفاده کردند تا نشان دهند که سن وام، نسبت وام به ارزش، نرخ بهره و نرخ بیکاری در توضیح پیش‌پرداخت‌ها، معوقات و نکول وام مسکن مهم هستند. گاردنر و میلز (۱۹۸۹) با تشخیص این موضوع که وام‌های معوق لزوماً به نکول ختم نمی‌شوند، از مدل رگرسیون لاجیت برای تخمین احتمال نکول برای وام‌های معوق فعلی استفاده کردند. توصیه‌ی آنان به بانکداران، استفاده از این روش برای شناسایی شدت مشکل و در نتیجه فرموله کردن پاسخ مناسب به بزهکاری است. اخیراً کاریتو و همکاران (۲۰۰۴) نیز دریافتند که روش لاجیت در پیش‌بینی نکول‌ها نسبت به سایر روش‌ها برتری دارد.

رگرسیون لجستیک^۵

به‌طور کلی، رگرسیون لجستیک احتمال یک نتیجه باینری را تخمین می‌زند که معمولاً برای پیش‌بینی نکول استفاده می‌شود. گنگ دونگ و همکاران (۲۰۱۰) در مطالعه‌ی خود رویکرد جدیدی از امتیازدهی اعتباری را، باهدف بهبود دقت پیش‌بینی مدل‌های رگرسیون لجستیک مورد استفاده در امتیازدهی اعتباری، بدون به خطر انداختن ویژگی‌های مطلوب آن‌ها مانند قوی بودن مدل و شفافیت، ارائه داده‌اند. نویسندگان این مقاله، به پیشنهاد یک مدل رگرسیون لجستیک با ضرایب تصادفی (LRR) پرداخته‌اند که فرض می‌کند ضرایب، به‌طور مستقل و نرمال در جامعه توزیع شده‌اند. محققان از مقایسه‌ی یک مدل رگرسیون لجستیک با ضرایب ثابت (LRF) و مدل LRR دریافتند که مدل LRR دارای دقت پیش‌بینی بهبود یافته‌ای نسبت به مدل LRF می‌باشد.

مدل‌های مبتنی بر قضاوت

روش‌های متعددی برای استخراج مدل‌های مبتنی بر قضاوت متخصصان وجود دارد که یکی از آن‌ها، فرآیند تحلیل سلسله‌مراتبی (AHP) نام دارد که فرآیندی ساختاریافته برای سازمان‌دهی و تجزیه و تحلیل تصمیمات پیچیده است. مدل AHP بر این اصل استوار است که هنگامی که تصمیمی در مورد یک موضوع خاص مدنظر باشد، تصمیم‌گیرندگان مسأله، تصمیم خود را به سلسله‌مراتبی از زیر مسائل ساده‌تر تجزیه می‌کنند که هر یک می‌تواند به‌طور مستقل تجزیه و تحلیل شود. نکته‌ی کلیدی AHP این است که قضاوت‌های انسانی، نه تنها برای اطلاعات زیربنایی، بلکه برای ارزیابی استثنائات و مواردی که اولویت ندارند یا به‌طور قابل توجهی در داده‌ها نشان داده نشده‌اند، نیز مورد استفاده قرار می‌گیرد و از این حیث دارای اهمیت بسیاری است. کاستا و همکاران (۲۰۰۲) در مطالعه‌ی خود یک مدل امتیازدهی کیفی اعتباری برای وام‌های تجاری بر اساس مفاهیم AHP ارائه داده‌اند (ناتسون، ۲۰۲۰).

مدل‌های یادگیری ماشینی

چالش مدل‌های پیش‌بینی ساده‌ی سنتی این است که اگرچه به دلیل شفافیت‌شان هنوز در امتیازدهی اعتباری رایج هستند، از قابلیت‌های پیش‌بینی قدرتمندتر الگوریتم‌های یادگیری ماشین مدرن برخوردار نبوده و منجر به زیان‌های احتمالی یا افزایش نکول تعهدات می‌گردند. یکی از مزایای اصلی آن، توانایی در کار با مجموعه داده‌های بدون ساختار می‌باشد که بسیاری از وظایف را بهبود بخشیده و پیشرفت‌هایی را در پردازش متن، گفتار، تصویر، ویدئو و صدا به ارمغان آورده است. دنیای کسب‌وکار نیز این پیشرفت‌ها را مورد توجه قرار داده و یادگیری عمیق به‌طور فزاینده‌ای برای بهبود

عملکردهای تحلیل تجاری علی‌الخصوص مدیریت ریسک اعتباری مورد استفاده قرار گرفته است. مدل یادگیری ماشینی را می‌توان به سه دسته‌ی اصلی تقسیم کرد: (۱) یادگیری تحت نظارت، (۲) یادگیری بدون نظارت و (۳) یادگیری تقویتی (نیمه نظارتی که اساساً ترکیبی از یادگیری تحت نظارت و بدون نظارت است). مارک اسمیت (۲۰۲۲) در مطالعه‌ای دو مدل برجسته‌ی یادگیری ماشینی یعنی یادگیری عمیق (DL) و ماشین‌های تقویت گرادیان (GBM) را در زمینه‌ی امتیازدهی اعتبار مقایسه کرده است. یافته‌های این مقاله نشان می‌دهد که GBM عموماً از نظر قدرت و سرعت محاسباتی از DL بهتر عمل می‌کند و آن را به انتخاب ارجح برای برنامه‌های امتیازدهی اعتباری تبدیل می‌کند. در پژوهش ارائه‌شده توسط اسمیرنو و استوپنیکو (۲۰۲۳)، کاربرد مدل‌های یادگیری عمیق، به‌ویژه شبکه‌های عصبی مکرر (RNNs) برای بهبود امتیازدهی اعتباری با استفاده از داده‌های سوابق اعتباری مشتریان مورد بررسی قرار گرفته و با مقایسه‌ی روش‌های سنتی امتیازدهی اعتباری با مدل‌های RNN پیشنهادی برتری عملکرد مدل LSTM را اثبات کرده‌اند.

یادگیری تحت نظارت

رایج‌ترین روشی که در امور مالی و سایر صنایع به کار گرفته می‌شود، یادگیری تحت نظارت است. این روش، یک قانون کلی طبقه‌بندی را یاد می‌گیرد که از آن برای پیش‌بینی برچسب‌ها برای مشاهدات دیگر در مجموعه داده استفاده می‌کند. در تحقیقی که توسط مایکل بوکر و همکاران (۲۰۲۱) انجام شده، نیاز به مدل‌های امتیازدهی اعتباری را نه تنها برای ارائه‌ی پیش‌بینی‌های دقیق ریسک، بلکه برای برآورده کردن خواسته‌های نظارتی برای شفافیت و قابلیت حسابرسی، مورد توجه قرار داده و مروری بر تکنیک‌های قابل درک و کاربرد این مدل‌ها در امتیازدهی اعتباری ارائه می‌کند.

درخت‌های تصمیم

درخت‌های تصمیم روش‌های ناپارامتریکی هستند که به بریمن، فریدمن، اولشن و استون (۱۹۸۴) منتسب بوده و در قالب نمودارهای شماتیک و درختی شکل برای نشان دادن احتمال آماری استفاده می‌شوند. این روش یک تکنیک انعطاف‌پذیر و قوی است؛ هرچند، عمدتاً در عملیات بانکی تنها به‌عنوان یک ابزار پشتیبان در کنار روش‌های تخمین پارامتری، استفاده می‌شود (ترور، تیشیرانی، فریدمن، ۲۰۰۱). اولین کسانی که از این روش در حوزه‌ی امتیازدهی اعتباری استفاده کردند، فریدمن، آلتمن و کائو (۱۹۸۵) بودند که آن را بهتر از LDA یافتند. جدیدترین روشی که در رابطه با درخت‌های تصمیم به متون تخصصی اضافه شد توسط فلدمن و گراس (۲۰۰۵) بود که از این روش برای

طبقه‌بندی جامع مدل‌های امتیازدهی اعتباری / فولادی، امینی خوزانی، حاجیپا و شاهوردیانی

داده‌های نکول وام مسکن استفاده کرده و مزایا و معایب آن را در رابطه با روش‌های سنتی مورد بحث قرار داده‌اند. (دنيسون، مالیک و اسمیت، ۱۹۸۸).

جنگل‌های تصادفی

جنگل‌های تصادفی، ترکیبی از پیش‌بینی کننده‌های درختی هستند، به‌طوری‌که هر درخت به یک نمونه (زیرمجموعه) از داده‌های توسعه‌ی مدل (داده‌های آموزشی) که به‌طور تصادفی انتخاب شده است، بستگی دارد (بريمن، ۲۰۰۱). به بیان دیگر، جنگل تصادفی فرآیند تولید درختان نامرتبط بر روی مجموعه‌ای از نمونه‌های بوت‌استرپ و میانگین‌گیری آن‌هاست که برای هر درخت، ویژگی‌ها به‌طور تصادفی به‌عنوان کاندیدای تقسیم انتخاب می‌شوند. الگوریتم جنگل تصادفی برای رگرسیون و طبقه‌بندی در پژوهشی توسط ترور، تبشیرانی و فردمن (۲۰۱۷) بیان شده است.

تقویت گرادیان^۶

تقویت گرادیان یک روش مجموعه‌ای برای مسائل رگرسیون و طبقه‌بندی؛ و ترکیبی خطی از یک سری مدل‌های ضعیف می‌باشد که به‌صورت تناوبی برای ایجاد یک مدل نهایی قوی ساخته شده است. این روش به خانواده الگوریتم‌های یادگیری جمعی تعلق دارد و عملکرد آن همواره از الگوریتم‌های اساسی یا ضعیف یا تکنیک‌های Bagging بهتر است. تقویت گرادیان، از درخت‌های رگرسیون برای اهداف پیش‌بینی استفاده کرده و با برازش یک مدل بر روی باقی‌مانده‌ها، مدل را به‌طور مکرر تولید می‌کند و با بهینه‌سازی یک تابع هدف تعمیم می‌یابد. (هستی، تبشیرانی و فریدمن ۲۰۱۷). ژو و گو (۲۰۲۲) در مطالعه‌ای نسخه‌ی پیشرفته‌ای از درخت تصمیم تقویت گرادیان (GBDT) را مورد بحث قرار داده‌اند که مزایای استراتژی تکنیک Bagging و تقویت الگوی بهینه‌سازی آن را در برمی‌گیرد. این رویکرد بهبود قابل توجهی در امتیازدهی اعتباری در مجموعه داده‌های عمومی نشان داده است. ژیا و همکاران (۲۰۲۱) به‌علاوه، لو و همکاران (۲۰۲۲) با انجام تحقیقاتی بر روی درخت‌های تصمیم‌گیری تقویت‌کننده گرادیان نشان داده‌اند که این روش به دلیل ترکیب ویژگی‌ها و قابلیت‌های انتخابی که برای داده‌های اعتباری با ابعاد بالا و همبستگی پیچیده مناسب است، عملکرد رضایت‌بخشی در امتیازدهی اعتبار دارد.

ماشین‌های بردار پشتیبان (SVM)^۷:

یک الگوریتم یادگیری ماشینی است که برای طبقه‌بندی، رگرسیون و تشخیص نقاط پرت استفاده می‌شود و با یافتن بهترین راه، به تفکیک داده‌ها به کلاس‌های مختلف می‌پردازد. گو و لی (۲۰۱۹) به

بررسی جامع SVM و رویکردهای فرا ابتکاری در مدل‌های امتیازدهی اعتباری از سال ۱۹۹۷ تا ۲۰۱۸، پرداخته و اثربخشی این تکنیک‌های هوش مصنوعی را نشان داده‌اند. این بررسی انواع مدل، روش‌های ارزیابی و رویکرد مدل‌سازی ترکیبی پیشرفته را مورد بحث قرار می‌دهد. ژانگ‌وهو (۲۰۰۹) در مطالعه‌ی دیگری به بررسی کاربرد SVM در امتیازدهی اعتباری پرداخته و آن را با روش‌های شبکه عصبی پس انتشار (BNN) مقایسه کرده‌اند. این مطالعه نشان داد که SVM و BNN هر دو به دقت پیش‌بینی ۸۰ درصدی در مجموعه داده‌های اعتباری استرالیا و آلمان از دست‌یافته‌اند. استکینگ و اسکچ (۲۰۰۳) نیز تحقیقاتی در مورد مقایسه‌ی عملکرد SVM با روش‌های سنتی مانند تحلیل تشخیصی خطی و رگرسیون لجستیک انجام داده و نشان داده‌اند که SVM می‌تواند جایگزینی قوی برای امتیازدهی اعتبار باشد.

شبکه‌های عصبی

مدل شبکه‌های عصبی (NNW) یک نمایش ریاضی است که در آن از مغز انسان و توانایی آن برای انطباق بر اساس جریان اطلاعات جدید، الهام گرفته شده است. از نظر ریاضی، NNW یک ابزار بهینه‌سازی غیرخطی است و تاکنون انواع مختلفی از آن در متون تخصصی مشخص شده است (بیشاپ، ۱۹۹۵). مشکل عمده‌ی NNW ها فقدان قابلیت توضیح آن‌هاست. اگرچه آن‌ها می‌توانند به میزان دقت پیش‌بینی بالایی دست یابند، دلیل این‌که چرا و چگونه تصمیم گرفته شده است قابل توضیح نیست. به عنوان مثال، در مورد وام رد شده، نمی‌توان تعیین کرد که دقیقاً کدام مشخصه (ها)ی کلیدی برای رد درخواست مدنظر است. در نتیجه، توضیح نتایج تصمیم‌گیری برای مدیران بسیار دشوار است (بیزنس و همکاران، ۲۰۰۳). پژوهش میرتا بنسیچ و همکارانش (۲۰۰۵) با استفاده از یک مجموعه داده نسبتاً کوچک، بر شناسایی ویژگی‌های کلیدی برای امتیازدهی اعتباری در وام‌دهی مشاغل کوچک، به‌ویژه در اقتصادهایی که در حال گذار هستند، تمرکز دارد. این مطالعه به مقایسه‌ی عملکرد مدل‌های مختلف رگرسیون لجستیک، الگوریتم‌های مختلف شبکه عصبی (NN) و درخت‌های تصمیم‌گیری طبقه‌بندی و رگرسیون درختی (CART) می‌پردازد. علاوه بر آن، شبکه‌های عصبی به‌طور چشم‌گیری در زمینه‌ی امتیازدهی اعتباری اهمیت یافته‌اند. هایاشی (۲۰۲۲) به بررسی سیستماتیک روندهای نوظهور در یادگیری عمیق برای امتیازدهی اعتباری پرداخته و بر دقت برتر الگوریتم‌های شبکه‌های عصبی نسبت به رویکردهای یادگیری ماشین سنتی تأکید کرده است. به‌ویژه، شبکه‌های اعتقادی عمیق ۸ (DBN)، به دلیل قابلیت‌های بالقوه‌ی طبقه‌بندی، دقت بالاتری از خود نشان داده‌اند. مطالعه‌ی دیگری توسط چانگرن و ژیا (۲۰۲۲) یک مدل مبتنی بر ترانسفورماتور برای امتیازدهی

طبقه‌بندی جامع مدل‌های امتیازدهی اعتباری / فولادی، امینی خوزانی، حاجیها و شاهوردیانی

اعتباری را معرفی کرد که از داده‌های رفتاری آنلاین استفاده می‌کند که از روش‌های سنتی در پیش‌بینی ریسک نکول کاربر بهتر عمل کرده است. لی و همکاران (۲۰۰۶) در پژوهشی به فراآموزی ۹ با شبکه‌های عصبی به‌عنوان راهی برای بهبود دقت مدل با استفاده از داده‌های محدود اشاره کرده‌اند که این رویکرد از شبکه‌های عصبی نظارت‌شده برای طراحی سیستم‌های امتیازدهی اعتباری استفاده می‌کند و می‌تواند دقت پیش‌بینی را افزایش دهد. سوخارو و همکاران (۲۰۲۱) با همکاری با یک بانک بزرگ اروپایی، یک شبکه‌ی عصبی ایجاد کرده‌اند که از داده‌های بانکی تراکنش برای امتیازدهی اعتباری مشتری استفاده می‌کند. این مدل در عملکرد از راه‌حل‌های موجود پیشی گرفته است. تسای و هونگ (۲۰۱۴) نیز به مقایسه‌ی مجموعه‌های شبکه‌ی عصبی و شبکه‌های عصبی ترکیبی بر روی مجموعه داده‌های امتیازدهی اعتباری معیار پرداخته و دریافتند که این مجموعه‌ها و مدل‌های ترکیبی بهتر از مدل‌های شبکه‌ی عصبی منفرد عمل می‌کنند.

یادگیری بدون نظارت

یادگیری بدون نظارت به روش‌هایی اشاره دارد که در آن داده‌های ارائه‌شده به الگوریتم، حاوی برچسب (رویدادها) نیستند. به عبارت دیگر، این الگوریتم‌ها به جای پیش‌بینی داده‌های جدید یا ناشناخته، در خدمت کاوش ویژگی‌های داده‌های بررسی‌شده هستند و شامل تکنیک‌های زیر می‌باشد.

خوشه‌بندی

خوشه‌بندی فرآیند به دست آوردن گروه‌های طبیعی از داده‌ها است و برخلاف طبقه‌بندی، به جای تجزیه و تحلیل داده‌های برچسب‌گذاری شده، خوشه‌بندی داده‌ها را برای تولید این برچسب تجزیه و تحلیل می‌کند. گروه‌ها به گونه‌ای تشکیل می‌شوند که اشیاء یک گروه بسیار شبیه به یکدیگر هستند و در عین حال با اشیاء گروه دیگر تفاوت زیادی دارند.

k-means خوشه‌بندی

خوشه‌بندی K-means روشی برای کمی سازی بردارهاست، برگرفته از پردازش سیگنال بوده و روش مطلوبی برای آنالیز خوشه‌بندی در داده‌کاوی محسوب می‌شود. این روش، باهدف تجزیه‌ی n مشاهده به k خوشه استفاده می‌شود که در آن هر مشاهده به خوشه‌ای با نزدیک‌ترین میانگین تعلق دارد. تعریف دقیق ریاضی آن به شکل زیر است که در آن μ_i میانگین نقاط در S_i (خوشه i ام) می‌باشد:

$$\arg \min_S \sum_{i=1}^k \sum_{x \in S_i} \|x - \mu_i\|^2 = \arg \min_S \sum_{i=1}^k |S_i| \text{Var}(S_i) \quad (6)$$

و معادل است با به حداقل رساندن دوجه‌دوی مربع انحراف از نقاط در همان خوشه یعنی

$$\sum_{Cluster C_i} \sum_{Dimension d} \sum_{x,y \in C_i} (x_d - y_d)^2 \quad (7)$$

شی‌نا (۲۰۱۰) در پژوهش خود به بهبود الگوریتم خوشه‌بندی k-means پرداخته و روشی را برای تعیین کارآمد میانگین‌های اولیه‌ی بهینه پیشنهاد کرده‌اند که گامی مهم در الگوریتم k-means به شمار می‌آید. سیناگا و یانگ (۲۰۲۰) نیز به پیشنهاد یک نسخه‌ی بدون نظارت از الگوریتم خوشه‌بندی k-means پرداخته‌اند که می‌تواند به‌طور خودکار تعداد بهینه‌ی خوشه‌ها را بدون هیچ مقدار اولیه یا انتخاب پارامتر قبلی پیدا کند.

خوشه‌بندی سلسله‌مراتبی

الگوریتم‌های خوشه‌بندی سلسله‌مراتبی هر نقطه‌ی داده را در ابتدا به‌عنوان یک خوشه‌ی واحد در نظر می‌گیرند و سپس به‌طور پی‌درپی جفت خوشه‌ها را ادغام می‌کنند تا زمانی که تمامی خوشه‌ها در یک خوشه‌ی واحد ادغام شوند که شامل تمام نقاط داده است. سایمون‌لیز و همکاران (۲۰۲۳) در مطالعه‌ی خود به بررسی روش‌های خوشه‌بندی سلسله‌مراتبی، برای طبقه‌بندی مسائل مربوط به مدل‌های ریسک اعتباری شناسایی‌شده از طریق گزارش‌های اعتبارسنجی پرداخته‌اند و اثربخشی خوشه‌بندی سلسله‌مراتبی را در بهبود مدل‌های امتیازدهی اعتباری با امکان بخش‌بندی بهتر مشتری که برای ارزیابی دقیق ریسک اعتباری حیاتی هستند، موردبررسی قرار داده‌اند. جدوال و همکاران (۲۰۲۲) نیز در پژوهشی اهمیت یادگیری بدون نظارت را در تقسیم مشتریان به گروه‌های مشابه بر اساس پارامترهای مختلف بیان کرده و معتقدند با خوشه‌بندی بهینه‌ی مشتریان، مدل‌های یادگیری ماشینی می‌توانند به‌دقت بالاتری دست یابند. این مقاله به معرفی یک رویکرد جدید پرداخته که خوشه‌بندی k-means و خوشه‌بندی سلسله‌مراتبی را ترکیب می‌کند.

یادگیری تقویتی

یادگیری تقویتی یک روش نوظهور است که بین یادگیری تحت نظارت و بدون نظارت قرار می‌گیرد و در فرآیندهای اعتبارسنجی کاربرد دارد (ناتسون، ۲۰۲۰). هراسیموچ (۲۰۱۸) در پژوهشی به بهینه‌سازی آستانه‌ی پذیرش در امتیازدهی اعتباری با استفاده از یادگیری تقویتی پرداخته و دریافته‌اند که الگوریتم یادگیری تقویتی از روش‌های استاتیک سنتی مبتنی بر بهینه‌سازی حساس به هزینه بهتر عمل می‌کند و منجر به سود بیشتر در محیط‌های شبیه‌سازی‌شده و داده‌های دنیای واقعی می‌گردد. الپو و همکاران (۲۰۲۲) نیز در مقاله‌ای پتانسیل یادگیری تقویتی را در افزایش فرآیندهای

طبقه‌بندی جامع مدل‌های امتیازدهی اعتباری / فولادی، امینی خوزانی، حاجیها و شاهوردیانی

تصمیم‌گیری در امتیازدهی اعتباری و ارائه‌ی راه‌حل‌های پویا و سودآورتر در مقایسه با روش‌های سنتی نشان داده و به جنبه‌های اساسی یادگیری تقویتی مانند تخصیص اعتبار که برای اثربخشی الگوریتم در برنامه‌های کاربردی دنیای واقعی ضروری است، پرداخته است.

تکنیک‌های یادگیری جمعی

تکنیک‌های یادگیری جمعی یک تکنیک یادگیری ماشینی است که با ادغام پیش‌بینی‌های چند مدل، دقت و انعطاف‌پذیری را در پیش‌بینی افزایش می‌دهد و هدف آن کاهش خطاها یا سوگیری‌هایی است که ممکن است در مدل‌های فردی وجود داشته باشد. تکنیک‌های یادگیری جمعی عبارت‌اند از: Stacking، Boosting، Bagging. دومیترسکو و همکاران (۲۰۱۴) در پژوهش خود روشی را برای افزایش مدل‌های امتیازدهی اعتباری موردبحث قرار داده و یک روش رگرسیون درختی لجستیک مقیدشده با تابع جریمه (PLTR) را معرفی می‌کند که با ترکیب قوانین درخت‌های تصمیم رگرسیون لجستیک به بهبود قدرت پیش‌بینی رگرسیون لجستیک برای امتیازدهی اعتباری با ادغام اثرات درخت تصمیم غیرخطی می‌پردازد.

مدل‌های هیبریدی

مدل‌های هیبریدی در امتیازدهی اعتباری به سیستمی اشاره دارد که روش‌ها یا تکنیک‌های متعددی را برای بهبود دقت پیش‌بینی اعتبار وام‌گیرنده ترکیب می‌کند (روی و شاو، ۲۰۲۱). از انواع روش‌های هیبریدی در امتیازدهی اعتباری می‌توان به هیبریداسیون در سطح داده، هیبریداسیون در سطح الگوریتم، هیبریداسیون در سطح فرآیند و روش‌های جمعی هیبریدی اشاره کرد. روی و شاو (۲۰۲۱) با استفاده از ترکیب مدل‌های BWM و TOPSIS، بر توسعه‌ی یک سیستم امتیازدهی اعتباری برای شرکت‌های کوچک و متوسط (SMEs) تمرکز کرده‌اند که بر محدودیت‌های مدیریت داده‌های مالی غلبه کرده و پیش‌بینی ریسک اعتباری را بهبود بخشیده است. یوسفی و همکاران (۲۰۲۱) نیز در مقاله‌ای یک رویکرد جدید برای امتیازدهی اعتباری را موردبحث قرار داده و به پیشنهاد یک مدل که ترکیبی از سیستم‌های استنتاج عصبی-فازی تطبیقی (ANFIS) و شبکه‌های عصبی مکرر (RNN) برای شناسایی و پیش‌بینی شوک‌های بازار است پرداخته‌اند. علاوه‌برآن، چن و همکاران (۲۰۲۰) در تحقیقی که پیرامون یک رویکرد جامع برای بهبود دقت مدل‌های امتیازدهی اعتباری به‌منظور تشخیص تقلب اعتباری ارائه‌شده، با ادغام رگرسیون لجستیک با شواهد وزنی، یک مدل هیبریدی امتیاز اعتباری جدید ایجاد کرده‌اند و از شاخص‌های اقتصادی برای ایجاد یک سیستم امتیازدهی اعتباری قابل‌اعتمادتر استفاده نموده و به نتایج قابل‌توجهی ازجمله رسیدگی به خطاهای

قابل توجه در رگرسیون لجستیک به دلیل ضعف در سوابق، افزایش نرخ پیش‌بینی امتیازات اعتباری و کاهش وقوع تقلب اعتباری دست‌یافته است.

سایر رویکردهای نوین

الگوریتم‌های ژنتیک

رویکرد الگوریتم ژنتیک (GA)^{۱۰}، شکلی از محاسبات تکاملی است که برای بهینه‌سازی مدل‌های پیش‌بینی، از جمله مدل‌هایی که برای ارزیابی اعتباری استفاده می‌شوند، کاربرد دارد. چند دیدگاه در مورد کاربرد الگوریتم ژنتیک در امتیازدهی اعتباری از جمله بهینه‌سازی مدل‌های پیش‌بینی کننده، انتخاب ویژگی (دیوید هورن، ۲۰۱۶)، مدیریت روابط غیرخطی (آدیسو و همکاران، ۲۰۲۲)، سازگاری (فینلای، ۲۰۰۶)، یادگیری گروهی و قابلیت جستجوی سراسری (رامپون و همکاران، ۲۰۱۳) وجود دارد. بحیرایی و ارشدی (۲۰۱۲) تکنیک جدیدی به نام برنامه‌ریزی ژنتیکی هندسی پویا (DGGP) را برای پیش‌بینی ورشکستگی شرکت‌ها با رفع نواقص روش‌های قبلی و با استفاده از نسبت‌های مالی معرفی کرده‌اند و نشان می‌دهد که چگونه تغییرات در مقادیر شاخص DRS با تغییر در الگوهای ریسک و جهت‌ها مرتبط است و می‌تواند تفسیر اقتصادی سلامت مالی یک شرکت را تغییر دهد. در پژوهش دیگری نیز که توسط چی و سو (۲۰۱۸) انجام شد، ترکیبی از الگوریتم‌های ژنتیک و مدل‌های امتیازدهی دوگانه به‌منظور بهبود دقت امتیازدهی اعتباری موردبررسی قرار گرفته است. رویکرد پیشنهادی، نقاط قوت الگوریتم‌های ژنتیک را با مدل امتیازدهی افزایش داده و با انتخاب مؤثر متغیرهای مهم و بهینه‌سازی توانایی پیش‌بینی مدل، موجب عملکرد بهبودیافته‌ای از امتیازدهی اعتباری می‌گردد.

مدل‌های RAROC

یک مدل بسیار محبوب است که برای ارزیابی و قیمت‌گذاری ریسک اعتباری بر اساس داده‌های بازار استفاده می‌شود. RAROC (بازده سرمایه تعدیل‌شده بر اساس ریسک)^{۱۱} ابتدا در سال ۱۹۹۸ توسط بانک‌داران دویچه در آلمان مورد استفاده قرار گرفت و اکنون تقریباً توسط تمام بانک‌های بزرگ در ایالات متحده و اروپا پذیرفته شده است. RAROC هم به‌عنوان یک معیار ریسک اعتباری و هم به‌عنوان یک ابزار قیمت‌گذاری وام برای مدیر مالی عمل می‌کند. ساندرز و الن (۲۰۱۰) در فصلی از کتاب «سنجش ریسک اعتباری در داخل و خارج از بحران مالی»، مدل‌های RAROC و کاربرد آن‌ها در وام‌دهی توسط بانک‌ها و مؤسسات مالی بزرگ را مورد بحث قرار داده‌اند. دامپس و همکاران (۲۰۱۹)

طبقه‌بندی جامع مدل‌های امتیازدهی اعتباری / فولادی، امینی خوزانی، حاجیها و شاهوردیانی

نیز فصلی از کتاب «مقدمه‌ای بر مدل‌سازی و ارزیابی ریسک اعتباری» را به تشریح عناصر اصلی مدل‌سازی ریسک اعتباری و تخمین احتمال نکول، زیان ناشی از نکول و قرار گرفتن در معرض نکول اختصاص داده و به بررسی اقدامات مالی برای ارزیابی سودآوری وام از روش RAROC پرداخته‌اند.

مدل‌های اختیار ریسک نکول

هنگامی که شرکتی با انتشار اوراق قرضه یا افزایش وام‌های بانکی خود، به جمع‌آوری سرمایه می‌پردازد، اختیار نکول یا بازپرداخت بسیار ارزشمندی دارد (مرتون، ۱۹۷۴ و بلک و شولز، ۱۹۷۳). چنانچه پروژه‌های سرمایه‌گذاری وام‌گیرنده با شکست مواجه شود به‌طوری‌که نتواند به دارنده‌ی اوراق قرضه یا بانک بازپرداخت کند، این اختیار را دارد که در بازپرداخت بدهی خود کوتاهی کند و دارایی‌های باقی‌مانده را به دارنده‌ی بدهی واگذار کند. شرکت KMV این ایده‌ی نسبتاً ساده را به یک مدل نظارت بر اعتبار تبدیل کرد. در حال حاضر، بسیاری از بزرگ‌ترین موسسه‌های مالی ایالات‌متحده از این مدل برای تعیین فرکانس ریسک نکول (EDF) شرکت‌های بزرگ استفاده می‌کنند (بلک و شولز، ۱۹۷۳). مدل KMV درواقع با استفاده از رویکرد مدل قیمت‌گذاری اختیار (OPM) برای استخراج ارزش بازار ضمنی دارایی‌ها و نوسانات دارایی‌های یک شرکت استفاده می‌کند.

تحلیل پوششی داده‌ها ۱۲ (DEA)

اگرچه تحلیل نسبت‌های مالی برای ارزیابی مالی شرکت‌ها، قدمتی دیرینه دارد، اما به علت محدودیت‌های موجود، نمی‌تواند راهنمای مناسبی برای سرمایه‌گذاران، اعتباردهندگان و مدیران واحدهای تجاری باشد. تکنیک تحلیل پوششی داده‌ها می‌تواند این مشکل را برطرف کند چراکه یک تکنیک ریاضی مبتنی بر برنامه‌ریزی خطی است که در آن با استفاده از یک مجموعه‌ی چندتایی از متغیرهای ورودی و خروجی، کارایی یک گروه از واحدهای موردبررسی تعیین می‌شود. در این روش، یک مرز کارا به‌صورت تجربی بر اساس ورودی‌ها مشخص می‌شود، سپس واحدهایی که بر روی مرز کارا قرار می‌گیرند به‌عنوان واحدهای کارا شناخته می‌شوند. چن و همکاران (۲۰۰۹) رویکرد جدیدی را برای دسته‌بندی مشتریان بانک بر اساس سطوح مشارکت‌شان پیشنهاد کرده‌اند. هدف مطالعه‌ی آن‌ها ایجاد یک مدل امتیازدهی رفتاری است که از DEA برای طبقه‌بندی مشتریان به گروه‌های دارای مشارکت بالا و پایین استفاده می‌کند. علی‌نژاد و کاشانی‌فر (۲۰۱۸) نیز در پژوهشی بر مدیریت و کنترل ریسک اعتباری مشتریان با ترکیب دو مدل تجزیه و تحلیل تشخیصی و تحلیل پوششی داده‌ها تمرکز کرده و به طبقه‌بندی مشاهدات به دو گروه مشتریان خوب و بد با تشخیص وجود یا عدم وجود همپوشانی بین گروه‌ها با استفاده از یک ابر صفحه‌ی جداکننده در حضور داده‌های فازی پرداخته‌اند.

نتیجه‌گیری

مدل‌های امتیازدهی اعتباری به یکی ابزارهای ضروری و بااهمیت برای مؤسسات مالی به‌منظور ارزیابی ریسک‌های مرتبط با یک متقاضی اعتباری تبدیل شده‌اند. این مدل‌ها تحول تاریخی قابل‌توجهی را تجربه کرده و از سیستم‌های مبتنی بر قوانین سنتی به رویکردهای پیچیده‌تر، به‌ویژه با ادغام هوش مصنوعی در فرآیندهای ارزیابی اعتبار، گذار کرده‌اند. (آمیونتولاس و همکاران، ۲۰۲۱). با تکامل بازارهای مالی، محدودیت‌های سیستم‌های امتیازدهی اعتباری مبتنی بر قوانین سنتی آشکار شد. پیدایش هوش مصنوعی یک تغییر پارادایم در امتیازدهی اعتباری را نشان داد. الگوریتم‌های یادگیری ماشین، زیرمجموعه‌ای از هوش مصنوعی، یک رویکرد مبتنی بر داده را معرفی کردند که از محدودیت‌های سیستم‌های مبتنی بر قوانین سنتی فراتر رفت. این الگوریتم‌ها می‌توانند مجموعه داده‌های وسیعی را تجزیه و تحلیل کنند، الگوها را شناسایی کنند و پیش‌بینی‌هایی را با سطحی از دقت و کارایی که قبلاً غیرقابل دستیابی بود انجام دهند (کامیاب و همکاران، ۲۰۲۳). استفاده از هوش مصنوعی در امتیازدهی اعتباری امکان ارزیابی دقیق‌تری از ریسک اعتباری را با در نظر گرفتن طیف وسیع‌تری از عوامل فراتر از عوامل معمولی فراهم می‌کند. مدل‌های هوش مصنوعی، از جمله مدل‌های مبتنی بر شبکه‌های عصبی، درخت‌های تصمیم‌گیری و روش‌های جمعی، می‌توانند به‌طور پویا، با شرایط در حال تغییر سازگار شوند، داده‌های بلادرنگ و یادگیری از الگوهای جدید را ترکیب کنند. این سازگاری به برطرف نقاط ضعف مدل‌های سنتی می‌پردازد و دقت ارزیابی‌های اعتباری را افزایش و خطاها یا سوگیری‌هایی که ممکن است در مدل‌های فردی وجود داشته باشد را کاهش می‌دهد. ظهور هوش مصنوعی با به‌کارگیری الگوریتم‌های یادگیری ماشینی و تجزیه و تحلیل‌های پیش‌بینی‌کننده‌ی پیشرفته، موجب تحولات چشم‌گیری در صنعت اعتباری گشته است. این مدل‌ها بر اساس قابلیت تفسیر، دقت و مقیاس‌پذیری‌شان ارزیابی می‌شوند. انتقال به رویکردهای مبتنی بر هوش مصنوعی امکان ادغام منابع داده‌ی جایگزین، مانند رسانه‌های اجتماعی و شاخص‌های مالی نامتعارف را فراهم کرده و دامنه‌ی ارزیابی اعتباری را گسترش می‌دهد (آدی و همکاران، ۲۰۲۴). رویکرد یادگیری عمیق، به دلیل توانایی‌اش در الگوگیری پیچیده از داده‌های خام که به فرمت‌های موردقبول برای مدل‌های یادگیری ماشینی تبدیل نشده، همچنین زمان‌بر بودن و نیازمند تخصص در حوزه‌ی موردبحث، حائز اهمیت می‌باشد. از مزایای روش درخت‌های تصمیم این است که بسیار شهودی است، به راحتی قابل بیان و توضیح برای مدیریت بوده و قادر به مواجهه با مشاهدات از دست‌رفته می‌باشد، اما نقطه ضعف اصلی آن، بار محاسباتی در مورد مجموعه داده‌های بزرگ و ناپایداری اغلب درخت‌ها است؛ اما در

طبقه‌بندی جامع مدل‌های امتیازدهی اعتباری / فولادی، امینی خوزانی، حاجیها و شاهوردیانی

مقابل، روش تقویت گرادیان که ترکیبی خطی از یک سری مدل‌های ضعیف می‌باشد، به ایجاد یک مدل نهایی قوی کمک می‌کند. در رویکرد خوشه‌بندی K-means روند محاسبه بسیار سریع بوده و تنها محاسبه‌ی موردنیاز، فاصله بین نقاط و مراکز گروه است اما از معایب آن این است که اولاً باید تعداد خوشه‌ها را تعیین کرد، در صورتی که گاهی نیاز داریم الگوریتم این کار را برای ما انجام دهد زیرا می‌خواهیم به دید کلی از داده‌ها دست‌یابیم؛ دوماً این روش با انتخاب تصادفی مراکز خوشه شروع می‌شود و ممکن است نتایج خوشه‌بندی متفاوتی را در اجراهای مختلف الگوریتم به دست آورد. درحالی‌که در الگوریتم‌های خوشه‌بندی سلسله‌مراتبی نیازی به تعیین تعداد خوشه نیست و حتی می‌توانیم بعد از ساختن درخت، تعیین کنیم تعداد خوشه‌ها چه قدر باشد و این الگوریتم به انتخاب اندازه‌ی فاصله نیز حساس نیست. رویکرد یادگیری تقویتی به ماشین اجازه می‌دهد تا بر اساس بازخورد از محیط بیاموزد و یادگیری‌های مجدد خود را با گذشت زمان تطبیق دهد. یادگیری تقویتی در رباتیک و تئوری بازی‌ها رایج است و نیاز کمی به مداخله انسانی دارد. تکنیک‌های یادگیری جمعی پیشرفت قابل‌توجهی در تکنیک‌های امتیازدهی اعتباری محسوب می‌شود که تعادلی بین عملکرد و قابلیت تفسیر که در بخش مالی ضروری است، ارائه می‌دهد. اگرچه این تکنیک‌ها معمولاً پیش‌بینی‌های صحیح‌تری انجام می‌دهند، اما میزان پیچیدگی‌شان در سطح بالایی است. مدل‌های هیبریدی اغلب روش‌های آماری، یادگیری ماشینی و پردازش داده‌های مختلف را برای ایجاد یک سیستم امتیازدهی اعتباری قوی‌تر و قابل‌اعتمادتر ادغام می‌کنند. الگوریتم‌های ژنتیک می‌توانند عملکرد مدل‌های امتیازدهی اعتباری را با بهینه‌سازی پارامترها، انتخاب مرتبط‌ترین ویژگی‌ها و تطبیق با داده‌های جدید به‌طور قابل‌توجهی افزایش دهند و آن‌ها را به ابزاری ارزشمند در زمینه‌ی ارزیابی اعتبار تبدیل کنند. در مدل‌های RAROC وام‌دهنده به‌جای ارزیابی ROA سالانه واقعی یا تعهد شده طبق قرارداد در مورد وام، سود مورد انتظار و درآمد کارمزد را با کسر هزینه وجوه در برابر ریسک مورد انتظار وام محاسبه می‌کند. از طرف دیگر، رویکرد تحلیل پوششی داده‌ها امکان ارزیابی دقیق‌تری از اعتبار مشتری را در محیط‌هایی که داده‌ها دقیق نیستند، فراهم می‌کند و استفاده از داده‌های فازی به محاسبه‌ی عدم قطعیت‌ها و تغییرات در رفتار مشتری و پروفایل‌های مالی کمک می‌کند.

پژوهش حاضر، سعی بر مرور و جمع‌آوری مدل‌های امتیازدهی اعتباری در پنج طبقه‌ی کلی مدل‌های آماری سنتی، مدل‌های یادگیری ماشینی، تکنیک‌های یادگیری جمعی، مدل‌های هیبریدی و سایر رویکردهای نوین داشته است. بدین منظور، در این پژوهش با مطالعه‌ی پژوهش‌های انجام‌گرفته در این حوزه، سیر تحول مدل‌ها موردبررسی قرار گرفته و پژوهشگرانی که دست‌آوردهایی در هر یک از

زمینه‌های مرتبط داشته‌اند، با ذکر منبع معرفی شده‌اند. مبنای این مطالعه، گردآوری حداکثری مدل‌های موجود در صنعت ارزیابی ریسک اعتباری، ارزیابی ورشکستگی و امتیازدهی اعتباری تا زمان حال بوده است، هرچند به دلیل گستردگی و تعدد مدل‌های ریاضی، قابلیت‌های هوش مصنوعی و الگوریتم‌های پیش‌بینی کننده و همچنین به جهت عدم دسترسی به برخی منابع و کتب علمی در سراسر دنیا، ممکن است مدل‌هایی توسط محققان ارائه شده باشد که در این پژوهش آورده نشده است. علاوه بر آن، به علاقه‌مندان به این حوزه پیشنهاد می‌گردد، با توجه به روبه رشد بودن صنعت اعتباری، به‌طور پیوسته دانش خود را از ظهور مدل‌های جدید به‌روز نمایند.

منابع

- ۱) البرزی، م.، پورزندی، م.؛ و خان‌بابایی، م. (۱۳۸۹). به‌کارگیری الگوریتم ژنتیک در بهینه‌سازی درختان تصمیم‌گیری برای اعتبارسنجی مشتریان بانک‌ها. نشریه مدیریت فناوری اطلاعات، دوره ۲، ۲۳-۳۸.
- ۲) راعی، ر.؛ و سروش، ا. (۱۳۹۱). اعتبارسنجی مشتریان حقوقی کوچک و متوسط بانک‌ها با استفاده از مدل‌های لجیت و پروبیت. پژوهشنامه اقتصاد، دوره ۱۲، شماره ۴۴، پیاپی ۱، ۱۴۵-۱۳۱.
- ۳) رجبی پورمبیدی، ع.، لگزیان، م.؛ و فصاحت، ج. (۱۳۹۲). مطالعه تأثیر نوع صنعت بر معیارهای اعتبار دهی به مشتریان حقوقی بانک صادرات با استفاده از DEA. مدیریت تولید و عملیات، دوره ۴، ۱۴۴-۱۲۹.
- 4) Alinezhad, A., & Kashanifar, S. (2018). Customer credit scoring using data envelopment analysis and discriminant analysis in a fuzzy environment (case study: a leasing company affiliated with a private bank). *Iranian Journal of Insurance Research*, 8(1), 27-40.
- 5) Aren, S., & Nayman Hamamcı, H. (2022). The impact of financial defence mechanisms and phantasy on risky investment intention. *Kybernetes*, 51(1), 141-164.
- 6) Bahiraie, A., & Arshadi, A. (2012). Bankruptcy Prediction: Dynamic Geometric Genetic Programming (DGGP) Approach. *Money and Economy*, 101.
- 7) Black, F. and Scholes, M. (1973). The Pricing of Options and Corporate Liabilities, *Journal of Political Economy* 81: 637-659.
- 8) Bücker, M., Szepannek, G., Gosiewska, A., & Biecek, P. (2022). Transparency, auditability, and explainability of machine learning models in credit scoring. *Journal of the Operational Research Society*, 73(1), 70-90.
- 9) Chai, N., Wu, B., Yang, W., & Shi, B. (2019). A multicriteria approach for modeling small enterprise credit rating: evidence from China. *Emerging Markets Finance and Trade*, 55(11), 2523-2543.
- 10) Chang, C. C., Wong, W. K., Lo, S. T., & Liao, Y. H. (2023). The relationship between sovereign credit rating changes and firm risk. *Heliyon*, 9(10).
- 11) Chen, K., Yadav, A., Khan, A., & Zhu, K. (2020). Credit fraud detection based on hybrid credit scoring model. *Procedia Computer Science*, 167, 2-8.
- 12) Chi, G., & Zhang, Z. (2017). Multi criteria credit rating model for small enterprise using a nonparametric method. *Sustainability*, 9(10), 1834.

- 13) Dastile, X., Celik, T., & Potsane, M. (2020). Statistical and machine learning models in credit scoring: A systematic literature survey. *Applied Soft Computing*, 91, 106263.
- 14) Dong, G., Lai, K. K., & Yen, J. (2010). Credit scorecard based on logistic regression with random coefficients. *Procedia Computer Science*, 1(1), 2463-2468.
- 15) Dumitrescu, E., Hué, S., Hurlin, C., & Tokpavi, S. (2022). Machine learning for credit scoring: Improving logistic regression with non-linear decision-tree effects. *European Journal of Operational Research*, 297(3), 1178-1192.
- 16) Otto, G., & Ukpere, W. (2011). Credit and thrift co-operatives in Nigeria: A potential source of capital formation and employment. *African Journal of Business Management*, 5(14), 5675-5680.
- 17) Goh, R. Y., Lee, L. S., Seow, H. V., & Gopal, K. (2020). Hybrid harmony search-artificial intelligence models in credit scoring. *Entropy*, 22(9), 989.
- 18) Hayashi, Y. (2022). Emerging trends in deep learning for credit scoring: A review. *Electronics*, 11(19), 3181.
- 19) Herasymovych, M. (2018). Optimizing Acceptance Threshold in Credit Scoring using Reinforcement Learning (Doctoral dissertation, Master thesis, University of Tartu). Retrieved from <https://core.ac.uk/download/pdf/159135256.pdf>.
- 20) Jadwal, P. K., Pathak, S., & Jain, S. (2022). Analysis of clustering algorithms for credit risk evaluation using multiple correspondence analysis. *Microsystem Technologies*, 28(12), 2715-2721.
- 21) Rozo, B. J. G., Crook, J., & Andreeva, G. (2023). The role of web browsing in credit risk prediction. *Decision Support Systems*, 164, 113879.
- 22) Trivedi, S. K. (2020). A study on credit scoring modeling with different feature selection and machine learning approaches. *Technology in Society*, 63, 101413.
- 23) Li, Y., & Chen, W. (2020). A comparative performance assessment of ensemble learning for credit scoring. *Mathematics*, 8(10), 1756.
- 24) Lis, S., Kubkowski, M., Borkowska, O., Serwa, D., & Kurpanik, J. (2023). Analyzing Credit Risk Model Problems through NLP-Based Clustering and Machine Learning: Insights from Validation Reports. *arXiv preprint arXiv:2306.01618*.
- 25) Rahman, M. J., & Zhu, H. (2024). Predicting financial distress using machine learning approaches: Evidence China. *Journal of Contemporary Accounting & Economics*, 20(1), 100403.
- 26) Moscato, V., Picariello, A., & Sperlí, G. (2021). A benchmark of machine learning approaches for credit score prediction. *Expert Systems with Applications*, 165, 113986.

- 27) Mpofu, T. P., & Mukosera, M. (2014). Credit scoring techniques: a survey. *International Journal of Science and Research (IJSR)*, 2319-7064.
- 28) Roy, P. K., & Shaw, K. (2021). A multicriteria credit scoring model for SMEs using hybrid BWM and TOPSIS. *Financial Innovation*, 7(1), 77.
- 29) Smirnov, V. S., & Stupnikov, S. A. (2023). A Deep Learning Approach to Credit Scoring Using Credit History Data. *Lobachevskii Journal of Mathematics*, 44(1), 198-204.
- 30) Sukharev, I., Shumovskaia, V., Fedyanin, K., Panov, M., & Berestnev, D. (2020, November). Ews-gcn: Edge weight-shared graph convolutional network for transactional banking data. In *2020 IEEE International Conference on Data Mining (ICDM)* (pp. 1268-1273). IEEE.
- 31) Saunders, A., Cornett, M. M., & Erhemjamts, O. (2021). *Financial institutions management: A risk management approach*. McGraw-Hill.
- 32) Wang, C., & Xiao, Z. (2022). A deep learning approach for credit scoring using feature embedded Transformer. *Applied Sciences*, 12(21), 10995.
- 33) Xia, Y., He, L., Li, Y., Fu, Y., & Xu, Y. (2021). A dynamic credit scoring model based on survival gradient boosting decision tree approach. *Technological and Economic Development of Economy*, 27(1), 96-119.
- 34) Zhang, E. X. (2020). The impact of cash flow management versus accruals management on credit rating performance and usage. *Review of Quantitative Finance and Accounting*, 54(4), 1163-1193.
- 35) Xia, Y., Xu, T., Wei, M. X., Wei, Z. K., & Tang, L. J. (2023). Predicting chain's manufacturing SME credit risk in supply chain finance based on machine learning methods. *Sustainability*, 15(2), 1087.
- 36) Xu, J., & Liu, F. (2020). The impact of intellectual capital on firm performance: A modified and extended VAIC model. *Journal of Competitiveness*, (1).
- 37) Xu, Z., Meng, L., He, D., Shi, X., & Chen, K. (2022). Government Support's signaling effect on credit financing for new-energy enterprises. *Energy Policy*, 164, 112921.
- 38) Tezerjan, M. Y., Samghabadi, A. S., & Memariani, A. (2021). ARF: A hybrid model for credit scoring in complex systems. *Expert Systems with Applications*, 185, 115634.
- 39) Zou, Y., & Gao, C. (2022). Extreme learning machine enhanced gradient boosting for credit scoring. *Algorithms*, 15(5), 149.

یادداشت‌ها:

1. Traditional statistical models
2. linear regression
3. discriminant analysis
4. Logit Analysis
5. logistic regression
6. Gradient boosting
7. Support Vector Machine
8. Deep Belief Networks
9. Metalearning
10. genetic algorithm
11. risk-adjusted return on capital
12. Data Envelopment Analysis

Comprehensive Classification of Credit Scoring Models

Pardis Fooladi¹

Mohsen Amini Khozani²

Zohreh Hajiha³

Shadi Shahverdiani⁴

Receipt: 24/07/2024

Acceptance: 30/11/2024

Abstract

Credit scoring models have evolved significantly over time, incorporating various techniques to assess the creditworthiness of individuals and businesses. Credit scoring is an important process for credit providers because it determines the probability of default or non-default of applicants and plays an important role in determining people's ability to receive facilities and credit. Traditional credit scoring models, such as those that incorporate the credit history and repayment behavior of borrowers, have been of interest in credit risk management for decades. These models help decision-making processes by evaluating credit, influencing loan approvals, and shaping preventive measures in case of default. Using emerging machine learning frameworks and approaches and alternative data sources, credit scoring methods will become more successful over time in optimizing risk management, improving loan approval processes, reducing investment risk for lenders, and effectively addressing the financial needs of borrowers. In this research, a comprehensive classification of credit scoring models has been presented using bibliometric citation analysis and includes more than 100 cited articles.

key words

Credit scoring, statistical analysis, regression, machine learning, neural networks. JEL: E51, C14, C38, C5, C5.

1-PhD Student, Department of Financial Engineering, Shahr-e-Qods Branch, Islamic Azad University, Tehran, Iran. fooladi.pardis@gmail.com

2-Assistant Professor, Department of Financial Engineering, Shahr-e-Qods Branch, Islamic Azad University, Tehran, Iran. (Corresponding Author) amini_k_m@yahoo.com

3-Associate Professor, Department of Accounting, East Tehran Branch, Islamic Azad University, Tehran, Iran. drzhajiha@gmail.com

4-Assistant Professor, Department of Financial Engineering, Shahr-e-Qods Branch, Islamic Azad University, Tehran, Iran. shshahverdiani@gmail.com