

امتیازدهی اعتباری مشتریان با استفاده از تحلیل پوششی داده ها و تحلیل ممیز در محیط فازی (مطالعه موردی: صنعت لیزینگ)

سعید کاشانی فر^۱ علیرضا علی نژاد^۲

تاریخ ارسال: ۱۴۰۳/۱۱/۱۶

۱- چکیده

در این تحقیق به منظور مدیریت و کنترل ریسک اعتباری مشتریان، از تلفیق دو مدل تحلیل ممیز و تحلیل پوششی داده ها تحت عنوان DEA/DA برای تشخیص وجود و یا عدم وجود یک همپوشانی میان دو گروه به وسیله یک ابرصفحه جداکننده استفاده و با فرض وجود n مشاهده هر یک با k مشخصه مستقل با حضور داده های فازی، مشاهدات به دو گروه G_1 (مشتریان خوش حساب) و G_2 (مشتریان بد حساب) دسته بندی شدند. متغیرهای این تحقیق از روش $6C$ انتخاب و از تعداد ۱۷ شاخص منتخب، با استفاده از روش دلفی تعداد ۸ شاخص تأثیرگذار و مؤثر مورد تأیید و وارد مدل تحقیق گشته است. این شاخص ها برای ۸۳ نفر از مشتریان حقیقی یک شرکت لیزینگ که طی سال های ۱۴۰۰ و ۱۴۰۱ از این شرکت تسهیلات دریافت کرده اند مورد استفاده قرار گرفت. داده های جمع آوری شده در دو گروه G_1 با تعداد ۵۹ مشاهده و گروه G_2 با تعداد ۲۴ مشاهده که هر یک دارای ۸ عامل مستقل $(I_1, \dots, I_8)$ هستند دسته بندی شده اند. نتایج نشان می دهند که هر یک از مشاهدات به طور قطعی در گروه G_1 و G_2 قرار گرفته و با ورود مشاهده جدید وضعیت اعتباری آن پیش بینی می شود. به منظور بررسی عملکرد و ارزیابی اعتبار مدل، با استفاده از یک مجموعه از داده های واقعی به مقایسه عملکرد بسط DEA/DA با شش روش DA پرداخته و مشخص گردید مدل ارایه شده در این تحقیق به طور میانگین ۹۹/۳٪ طبقه بندی درست داده ها را در مقابل شش روش دیگر DA انجام داده است.

کلمات کلیدی: تحلیل پوششی داده ها^۳/تحلیل ممیز^۴ (DEA/DA)، داده های فازی^۵، ریسک اعتباری^۶، امتیازدهی اعتباری^۷

^۱ مدیرعامل و عضو هیأت مدیره شرکت واسپاری صنعت یاران محسنین saeeid.kashanifar@gmail.com

^۲ عضو هیأت علمی دانشگاه آزاد اسلامی واحد قزوین

^۳ Data Envelopment Analysis

^۴ Discriminant Analysis

^۵ Fuzzy Data

^۶ Credit Risk

^۷ Credit Scoring

۱- مقدمه

بررسی، سنجش و اندازه‌گیری اعتبار مشتریان در مؤسسات اعتباری، امروزه یکی از مهمترین تصمیمات مالی به شمار می‌آید. در گذشته تصمیم در خصوص اعطای اعتبار به افراد حقیقی و یا حقوقی، اغلب به عهده فردی خبره یا گروهی از خبرگان بوده است و این امر نیز توسط بخش‌های مرتبط با امور پولی و اعتباری به اجرا در می‌آمده است. از آنجایی که روش‌های قضاوتی مذکور بسیار وقت‌گیر، پرهزینه و ذهنی هستند از اعتبار علمی و پایایی لازم برخوردار نمی‌باشند. اکثر مؤسسات مالی امروزه اقدام به طراحی سیستم‌های امتیازدهی اعتبار عینی و معتبری مبتنی بر الگوها و مدل‌های علمی نموده‌اند.

لیزینگ به عنوان یکی از مهمترین روش‌های تأمین اعتبار در سطح اقتصاد داخلی و بین‌المللی جایگاه خود را تثبیت نموده است و بخش عمده‌ای از عملیات مالی و تسهیلات و سرمایه‌گذاری و تأمین مالی در جهان، با استفاده از روش‌های متنوع لیزینگ صورت می‌گیرد. رشد اقتصادی سال ۲۰۰۵، صنعت لیزینگ را به عنوان یک محصول مالی مکمل محصولات بانکی، معرفی کرده است. در واقع لیزینگ توانسته است به عنوان منبع مالی میان‌مدت که توانایی حمایت از توسعه بخش‌های مختلف اقتصاد یک کشور را دارد، شناخته شود. مقدار سرمایه‌گذاری در کالاهای سرمایه‌ای توسط شرکت‌های لیزینگ موجب افزایش اهمیت لیزینگ در اقتصاد جهان شده است. به‌طوری‌که در سال‌های اخیر، اکثر کشورها، لیزینگ را به عنوان ابزاری جهت بهبود صنایع و تولید بیشتر، شغل‌های جدید و بهبود وضعیت زندگی مردم به کار گرفته‌اند. با توجه به اهمیت صنعت لیزینگ، بحث شناسایی بازار لیزینگ و ایجاد روند بهبود در آن یکی از بحث‌های مهم در این صنعت است.

در گذشته بسیاری از مؤسسات مالی، ریسک اعتباری خود را فقط برای دوره‌ای طولانی و با توجه به ارزش اعتباری مشتری مدیریت می‌کردند. همچنین فرآیند تصمیم‌گیری اعتباری بسیار ابتدایی بود و اغلب اوقات تصمیمات بر اساس اعطا و یا عدم اعطای تسهیلات قرار داشت. حتی اگر مشتریان ورشکسته هم می‌شدند، زیان‌های آنها اغلب به وسیله وثیقه جبران می‌شد ولی بعد از مدتی هزینه‌های اعتباری در مؤسسات مالی به علت افزایش تعداد مشتریان ورشکسته و کاهش ارزش وثیقه‌ها به طور معنی‌داری افزایش یافت. به همین سبب کنترل ریسک اعتباری مهم‌ترین نیازمندی مدیریت مؤثر شناخته می‌شود. لذا توسعه روش‌های اعتبارسنجی و همچنین فرهنگ‌سازی در استفاده از این روش‌ها بخشی از زیرساخت‌های مورد نیاز مؤسسات مالی و اعتباری می‌باشند. از طرف دیگر عدم بازپرداخت به موقع تسهیلات، بیانگر این است که دریافت‌کننده تسهیلات در بهره‌برداری از آن موفقیت چندانی نداشته است. به بیانی دیگر بازده حاصل از به کارگیری تسهیلات به اندازه کافی نبوده و دریافت‌کننده در موعد پرداخت تسهیلات با مشکلاتی مواجه است. بنابراین تسهیلات مذکور به مطالبات معوق تبدیل می‌شوند. برای هر مؤسسه اعطاکننده اعتبار و تسهیلات، توانایی تشخیص مشتریان بدحساب از خوش‌حساب بسیار مهم و ضروری است. (طالبی و شیرزادی، ۱۳۹۰).

با توجه به توضیحات ذکر شده، در این مقاله به دنبال پاسخ به سؤالات زیر هستیم:

۱. چگونه می‌توان با استفاده از تکنیک تحلیل پوششی داده‌ها/تحلیل ممیز، یک جایگاه اعتباری (درجه عضویت) برای مشتریان تعیین نمود؟
۲. چگونه می‌توان با استفاده از تکنیک تحلیل پوششی داده‌ها/تحلیل ممیز، یک تابع تمایز بین مشتریان خوش‌حساب و مشتریان بدحساب تعیین نمود؟
۳. آیا روش تحلیل پوششی داده‌ها/تحلیل ممیز از کارایی لازم برای امتیازدهی اعتباری مشتریان در صنعت لیزینگ برخوردار است؟

اهمیت این مسأله باعث شده است تا مقاله حاضر به امتیازدهی اعتباری مشتریان یکی از شرکتهای لیزینگ با استفاده از تحلیل پوششی داده‌ها و تحلیل ممیز در محیط فازی با در نظر گرفتن شاخص‌های مؤثر در امتیازدهی مشتریان بپردازد.

۲- پیشینه پژوهش

به رغم بسط و توسعه تکنیک تحلیل پوششی داده‌ها، تاکنون کمتر مطالعه‌ای به تعیین جایگاه اعتباری مشتریان در صنعت لیزینگ با استفاده از این روش پرداخته است. از این رو، مطالعات انجام شده در خصوص رتبه‌بندی اعتباری و بررسی عوامل مؤثر بر ریسک اعتباری مشتریان در بانکداری با استفاده از روش تحلیل پوششی داده‌ها و سایر روش‌های موجود همچون تحلیل ممیزی، رگرسیون ساده و لجستیک و شبکه‌های عصبی مورد بررسی قرار می‌گیرد.

در اواخر سال ۱۹۹۰ روش تحلیل پوششی داده‌ها برای تحلیل رتبه‌بندی اعتباری از سوی ترات^۱ و همکاران، معرفی شد که این روش صرفاً نیازمند اطلاعات ثبتي پیشین یعنی مجموعه مشاهده شده ورودی‌ها و خروجی‌ها برای محاسبه رتبه اعتباری مشتریان بود. بنابراین این امر افق جدیدی برای رتبه‌بندی اعتباری باز نمود (عریانی، ۱۳۸۴).

یولالان و همکاران در سال ۲۰۰۳ در مقاله‌ای تحت عنوان "امتیازدهی اعتباری رهیافتی در بانکداری تجاری" به امتیازدهی اعتباری مشتریان یکی از بانک‌های ترکیه در هفت مرحله پرداختند. آنها در ابتدا تعداد ۱۰۰ شرکت را انتخاب و شرکت‌هایی را که اطلاعات آماری آنها تفاوت معناداری با میانگین صنعت داشت حذف و در نهایت ۸۲ شرکت تولیدی صنعتی را برای بررسی نهایی انتخاب نمودند. سپس با استفاده از معیار 5C، تعداد ۴۶ نسبت مالی اصلی را جهت رتبه‌بندی اعتباری ۸۲ مشتری خود انتخاب و با استفاده از تکنیک تحلیل پوششی داده‌ها به امتیازدهی اعتباری مشتریان بانک پرداختند (Yalaln et al., 2003).

در سال ۲۰۰۸، مین و لی در مقاله‌ای با هدف رتبه‌بندی اعتباری، از تکنیک تحلیلی پوششی داده‌ها استفاده کردند. به این منظور محققان از داده‌های مالی حسابرسی شده تعدادی از شرکت‌های تولیدی استفاده نمودند. در این مقاله ورودی‌ها شامل نسبت‌های هزینه‌های مالی به فروش، بدهی‌های جاری به سرمایه و کل بدهی به کل دارایی و خروجی‌ها شامل نسبت‌های سرمایه به کل دارایی و دارایی جاری به بدهی جاری می‌باشد. به اعتقاد محققین، نتیجه این پژوهش که شامل رتبه اعتباری به دست آمده با استفاده از شیوه تحلیل پوششی داده‌ها بوده دارای نتیجه قابل اعتماد و قابل اتکایی است (Min and Lee, 2008).

مارسین توماس در سال ۲۰۰۹ در مقاله‌ای تحت عنوان "کاربردی از تحلیل پوششی داده‌ها در رتبه‌بندی اعتباری" به رتبه‌بندی اعتباری با استفاده از روش تحلیل پوششی داده‌ها در پنج مرحله پرداخت. در ابتدا هفت نسبت مالی انتخاب و سپس ۱,۴۰۸ شرکت دارای اطلاعات کامل و مناسب از بین تعداد ۲۶۴,۱۰۳ شرکت انتخاب شدند. در مرحله بعد ماتریس همبستگی میان نسبت‌ها محاسبه و با استفاده از مدل CCR کارایی شرکت‌ها محاسبه و مشخص گردید فقط ۲۷ شرکت کاملاً کارا بوده‌اند. در این مقاله به منظور اعتباربخشی رتبه‌های حاصل از تحلیل پوششی داده‌ها، از روش‌های تحلیل رگرسیون و تحلیل ممیزی استفاده گردیده است (Marcin, 2009).

^۱- Trout

در سال ۲۰۱۰، فوکویاما و ماتوسک در مقاله‌ای، با استفاده از تحلیل پوششی داده‌های شبکه‌ای با ساختار دومرحله‌ای، اقدام به تجزیه و تحلیل رفتار اعتباری در بانک‌های ترکیه طی سال‌های ۱۹۹۱ الی ۲۰۰۷ نمودند. در این مقاله، در مرحله اول بانک‌ها با استفاده از چندین ورودی شامل تعداد کارکنان، سرمایه و دارایی‌های ثابت یک خروجی میانی تولید نمودند که به عنوان ورودی مرحله دوم جهت تولید خروجی‌های نهایی وام‌ها و اوراق بهادار محسوب می‌گردند. در این مقاله همچنین به تجزیه و تحلیل تغییرات کارایی در بانک‌ها در ترکیه، قبل و بعد از بحران مالی و همچنین مقایسه بهره‌وری بانک‌های داخلی و خارجی پرداخته است. این مقاله نشان داد که دقت و ثبات نتایج حاصل از استفاده از تحلیل پوششی داده‌های شبکه‌ای بهبود یافته است و بانک‌های ترکیه نسبت به فرایند تثبیت و بازسازی، واکنش مثبت نشان داده‌اند و کارایی آنها به تدریج بهبود یافته است (Fukuyama and Matousek, 2010).

در سال ۲۰۱۳، وانگ و همکاران در مقاله‌ای، با استفاده از ساختار شبکه‌ای ترکیبی، سودآوری و کارایی قابلیت عرضه در بازار در ۶۵ شرکت با فناوری‌های پیشرفته تایوان را مورد مطالعه قرار دادند. بر خلاف مطالعات دو مرحله‌ای متداول، مرحله نخست در این مقاله از دو پردازش شامل تولید اولیه و تأثیرات R&D با عملکرد موازی تشکیل شد. کمینه‌سازی میانگین وزن‌دار پارامترهای فاصله ورودی در مرحله نخست و پارامتر فاصله خروجی در مرحله دوم به شکل منفی اجرا گردید (Wang, et al., 2013).

در سال ۲۰۱۴، وربراکن و همکاران به توسعه مدل تعیین رتبه اعتباری مصرف‌کنندگان با استفاده از پارامترهای مبتنی بر سود پرداختند. نتایج این مدل نشان می‌دهد که مدل ارایه شده در مجموعه مورد آزمون، از نظر دقت اندازه‌گیری نتایج دقیق‌تر از سایر مدل‌های موجود می‌باشد و استقرار آن در سیستم سهل‌تر می‌باشد (Verbraken, et al., 2014).

در سال ۲۰۱۵، نعمتی کوتنایی و همکاران یک مدل داده‌کاوی ترکیبی جهت امتیازدهی اعتباری گروه‌های یادگیری با استفاده از الگوریتم‌های مختلف ارایه نمودند. نتایج این مدل نشان داد که الگوریتم تجزیه و تحلیل مولفه‌های اصلی، بهترین الگوریتم انتخاب ویژگی است. نتایج طبقه‌بندی نیز نشان داد که دقت طبقه‌بندی در شبکه‌های عصبی تطبیقی بالاتر است. در نهایت، مدل ترکیبی به عنوان یک مدل قوی جهت سنجش اعتبار ارائه شده است (Koutanaei, et al., 2015).

عیسی‌زاده و عریانی نیز در سال ۱۳۸۹، با هدف رتبه‌بندی مشتریان حقوقی بانک کشاورزی بر حسب ریسک اعتباری، به بررسی کاربرد تحلیل پوشش داده‌ها در این حوزه پرداختند. در این تحقیق با استفاده از روش نمونه‌گیری خوشه‌ای از مناطق شرق و غرب بانک کشاورزی استان تهران، ۲۸۶ شرکت تسهیلات‌گیرنده مورد بررسی قرار گرفته و پس از خارج کردن داده‌های نامناسب تنها ۷۵ شرکتی که با استفاده از عقد فروش اقساطی ۲۴ ماهه با سررسید پایان اردیبهشت‌ماه سال ۱۳۸۴ تسهیلات دریافت کرده بودند، برای تجزیه و تحلیل نهایی مورد استفاده قرار گرفتند. در نهایت با استفاده از روش تحلیل پوششی داده‌ها، کارایی‌های فنی محاسبه و شرکت‌ها رتبه‌بندی شدند. نتایج به دست آمده نشان داد که ۱۵ شرکت (معادل ۲۰ درصد شرکت‌های مورد بررسی) روی مرز کارایی قرار داشته و کاملاً کارا قلمداد شده‌اند. همچنین میانگین کارایی فنی معادل ۷۸ درصد بوده است، به این معنا که در مجموع شرکت‌های مورد بررسی، ۲۲ درصد بیش از میزان مورد نیاز، ورودی‌ها و عوامل تولید را مورد استفاده قرار داده و دارای سودآوری پایینی بوده‌اند (عیسی‌زاده و عریانی، ۱۳۸۹).

در سال ۱۳۹۱ شورورزی و همکاران به ارزیابی ریسک اعتباری مشتریان حقوقی بانک صادرات بر مبنای دو مدل شبکه‌های عصبی مصنوعی و لجیت و مقایسه آنها با هم پرداختند. در این پژوهش با استفاده از دو مدل

رگرسیون لجستیک و شبکه‌های عصبی پرسپترون چند لایه يك نمونه ۱۲۰ تایی (۷۱ مشتری خوش حساب و ۴۹ مشتری بد حساب) که از بانک صادرات خراسان رضوی تسهیلات اعتباری دریافت نموده‌اند، بررسی گردید. ابتدا ۲۶ متغیر مستقل شامل متغیرهای کیفی و مالی شناسایی و در نهایت ۱۳ متغیر را که دارای اثر معناداری بر ریسک اعتباری و تفکیک بین دو گروه از مشتریان بودند، انتخاب و مدل نهایی به وسیله آنها برازش شد. به منظور سنجش کارایی مدل شبکه عصبی در مقایسه با رگرسیون لجستیک، نتایج حاصل از رگرسیون لجستیک و شبکه عصبی با یکدیگر مقایسه و بررسی نتایج نشان داد که هر دو مدل در ارزیابی ریسک اعتباری مشتریان حقوقی بانک از کارایی بالا و قابلیت مشابهی برخوردار می‌باشند (شورورزی و همکاران، ۱۳۹۱).

احمدیان در سال ۱۳۹۲ به بررسی عوامل موثر بر تعیین رتبه اعتباری مشتریان سیستم بانکی در ایران پرداخت. شاخص‌های مالی، شاخص‌های اعتباری، فاکتورهای مدیریتی صنعت، استمرار فعالیت و بازارهای هدف متغیرهایی است که در این پژوهش مورد بررسی قرار گرفت. یافته‌های پژوهش نشان‌دهنده این بود که شاخص‌های مالی، وضعیت اعتباری مشتری نزد سیستم بانکی، فاکتورهای مدیریتی و استمرار فعالیت بر رتبه اعتباری متقاضی اثرگذار است و همچنین بین صنعت و بازارهای هدف اصلی بنگاه و رتبه اعتباری آن رابطه معنی‌داری وجود ندارد (احمدیان، ۱۳۹۲).

در سال ۱۳۹۳، علیزاده به منظور تحلیل بهتر مشتریان و شناسایی نیازهای بالفعل و بالقوه و تعیین اعتبار آنها، با استفاده از تکنیک‌های داده‌کاوی به ارایه یک چارچوب بهینه به منظور افزایش دقت رتبه‌بندی ریسک اعتباری مشتریان بانکی پرداخت. در این تحقیق پس از بررسی برخی تکنیک‌های پرکاربرد داده‌کاوی جهت رتبه‌بندی ریسک اعتباری مشتریان در بانک‌ها و موسسات مالی و ساخت مدل‌های مربوط به هر یک از آنها، با ارائه یک چارچوب بهینه در ترکیب این مدل‌ها میزان دقت صحت دسته‌بندی ارتقا یافت (علیزاده، ۱۳۹۳). در سال ۱۳۹۴، استیری به بررسی و بکارگیری شاخص‌های مالی موثر جهت رتبه‌بندی شرکت‌های صنعت محصولات غذایی بورس اوراق بهادار تهران با استفاده از مدل‌های تصمیم‌گیری چند شاخصه پرداخت (استیری، ۱۳۹۴).

۳- ریسک و انواع آن در صنعت لیزینگ

لیزینگ یکی از پیچیده‌ترین شیوه‌های تأمین مالی در جهان امروزی است و اتخاذ تصمیمات مؤثر و موفق در این صنعت نیازمند تجربه فراوان است. در طول مدت یک اجاره، ریسک و بازده آن و همچنین عوامل مؤثر بر این ریسک و بازده تغییر می‌کنند.

ریسک‌های لیزینگ تحت سه دسته کلی ریسک تجهیزات (ریسک عدم تطابق ارزش تجهیزات با میزان پیش‌بینی شده)، ریسک اعتباری (ریسک عدم تحقق سود مورد انتظار از مجرای درآمد اجاره بها) و ریسک مالیاتی (ریسک عدم تحقق منافع مالیاتی پیش‌بینی شده) و ریسک اعتباری (عدم اطمینان موجود در مورد اینکه آیا مستأجر، پرداخت‌های باقیمانده اجاره بهای خود را انجام خواهد داد یا خیر) تقسیم‌بندی خواهند شد.

امتیازدهی اعتباری^۱ روشی برای ارزیابی ریسک اعتباری متقاضیان تسهیلات است. این روش با استفاده از داده‌های گذشته و روش‌های آماری سعی بر مجزا کردن آثار مشخصه‌های متفاوت متقاضیان بر نکول تسهیلات دارد.

^۱- Credit Scoring

امتیازدهی اعتباری با ایجاد امتیاز^۱ روشی را به وجود می‌آورد که بانک یا مؤسسه مالی با کمک آن، متقاضیان یا مشتریان اعتباری خود را بر اساس ریسک‌شان رتبه‌بندی کنند (امارنات، ۱۹۹۷). در دنیای مدل‌های آماری امتیازدهی اعتباری، مدل‌های گوناگونی وجود دارد که هر کدام مشخصات، مزایا و معایب خاص خود را دارند. این مدل‌ها به مدل‌های قضاوتی و آماری تقسیم‌بندی و خود مدل‌های آماری نیز به چند بخش زیر تقسیم می‌شوند:

روش‌های آماری مبتنی بر نسبت‌های حسابداری^۲، مدل‌های ساختاری^۳، مدل‌های کاهش‌یافته^۴ و مدل‌های ترکیبی^۵. در طبقه‌بندی دیگر، مدل‌های آماری به روش‌های آماری پارامتری و ناپارامتری تقسیم می‌شوند. مدل‌هایی که هدف آنها به دست‌آوردن و تخمین پارامترها به منظور طبقه‌بندی مشتریان است، مدل‌های پارامتری و مدل‌هایی که در آنها پارامتری محاسبه نشده و هدف آنها طبقه‌بندی مشتریان با توجه به روابط بین متغیرهاست، مدل‌های ناپارامتری هستند.

گفتن این مطلب که کدامیک از این مدل‌ها توانایی بهتری در پیش‌بینی نکول دارند دشوار است. هر کدام معایب و مزایای خاص خود را دارند و انتخاب هر کدام از این مدل‌ها به شرایط تجاری متقاضیان و همچنین پورتنوی خاص لیزینگ بستگی دارد. حتی بعضی وقت‌ها ممکن است از دو یا چند مدل برای اندازه‌گیری ریسک اعتباری لیزینگ استفاده شود.

اما در این میان یکی از معروف‌ترین روش‌هایی که در سال‌های اخیر مورد توجه بسیاری از صاحب‌نظران و تحلیلگران قرار گرفته، روش تحلیل پوششی داده‌هاست. این روش یک روش ناپارامتری و بر اساس روش‌های برنامه‌ریزی خطی است که برای تعیین کارایی بنگاه‌های اقتصادی مورد استفاده قرار می‌گیرد.

۴- متدولوژی تحقیق

با توجه به اینکه در تحلیل ممیز (DA) ابرصفحه جداکننده توسط یک تابع خطی بیان می‌شود، این ابرصفحه ممکن است دو گروه را از هم جدا کند یا اینکه بعضی از مشاهدات مربوط به یک گروه را در گروه دیگر طبقه‌بندی کند و یا تعدادی از مشاهدات را به طور ناصریح و اشتباه طبقه‌بندی کند.

۴-۱ توسعه مدل DEA/DA

توسعه مدل DEA/DA یک روش غیرپارامتری دو مرحله‌ای است که تابع جداسازی خطی (به جای ابرصفحه جداساز) را ارائه می‌دهد و با هر یک از مقادیر وزنی منفی و مثبت مرتبط با هر عامل مشاهده سر و کار دارد.

۴-۲ مجموعه‌های فازی

در دنیایی که ما در آن زندگی می‌کنیم، بسیاری از تعاریف و مفاهیم غیرقطعی و نسبی هستند. در این حالت شخص تصمیم‌گیرنده با نوعی عدم قطعیت روبه‌رو است که به عدم مرزبندی دقیق و محکم مفاهیم مربوط می‌شود. نظریه مجموعه‌های فازی قادر است به بسیاری از مفاهیم، متغیرها و سیستم‌هایی که نادقیق و مبهم هستند، صورت‌بندی ریاضی ببخشد و زمینه را برای استنتاج و تصمیم‌گیری در شرایط عدم قطعیت فراهم کند.

^۱- Score

^۲- Accounting – Based Statistical Methods

^۳- Structural Models

^۴- Reduced form Models

^۵- Hybrid Models

هنگامی که برخی مشاهدات فازی است، تابع هدف و محدودیت‌های مدل تصمیم‌گیری نیز به ماهیت فازی تبدیل می‌شود. در واقع، یک عدد فازی دلخواه توسط یک جفت مرتب شده توابع $(\underline{u}(r), \bar{u}(r))$ نمایش داده می‌شود که به آن فرم پارامتری اعداد فازی گفته می‌شود.

۴-۳ رتبه‌بندی فازی اعداد مثلثی

یکی از کارآمدترین رویکردهای مرتب‌سازی مؤلفه‌های $F(\mathfrak{R})$ ، تعریف تابع رتبه‌بندی $\tau: F(\mathfrak{R}) \rightarrow \mathfrak{R}$ است که هر عدد فازی را به صورت یک عدد قطعی ترسیم می‌کند. رتبه‌ها در $F(\mathfrak{R})$ به صورت ذیل تعریف می‌شود:

$$\begin{aligned} \tau(a) \geq \tau(b) & \quad \tilde{a} \succeq \tilde{b} \text{ اگر و فقط اگر} \\ \tau(a) > \tau(b) & \quad \tilde{a} \succ b \text{ اگر و فقط اگر} \\ \tau(a) = \tau(b) & \quad \tilde{a} \approx \tilde{b} \text{ اگر و فقط اگر} \end{aligned} \quad (۱)$$

در اینجا، $\tilde{a}\tilde{b}$ در $F(\mathfrak{R})$ وجود دارد.

قضیه: τ یک تابع رتبه‌بندی خطی در نظر گرفته می‌شود. آنگاه:

$$\begin{aligned} 1. \quad \tilde{a} \succeq \tilde{b} & \quad \text{iff} \quad \tilde{a} - \tilde{b} \succeq 0 \quad \text{iff} \quad -\tilde{b} \succeq -\tilde{a} \\ 2. \quad \tilde{a} \succeq \tilde{b} \text{ and } \tilde{c} - \tilde{d} \succeq 0 & \quad \text{iff} \quad \tilde{a} + \tilde{c} \succeq \tilde{b} + \tilde{d} \end{aligned} \quad (۲)$$

در اینجا، توجه به تابع رتبه‌بندی خطی محدود می‌شود که تابع رتبه‌بندی τ با شرایط زیر است:

$$\begin{aligned} \tau(k\tilde{a} + \tilde{b}) &= k\tau(\tilde{a}) + \tau(\tilde{b}) \text{ به ازای هر } \tilde{a} \text{ و } \tilde{b} \text{ متعلق به } \tau(\mathfrak{R}), \text{ و هر } k \in \mathfrak{R} \text{ برقرار است. به ازای عدد فازی} \\ \tilde{a} &= (\underline{a}(r), \bar{a}(r)) \text{ تابع رتبه‌بندی } \tau(\tilde{a}) = \frac{1}{2} \int_0^1 (\underline{a}(r), \bar{a}(r)) dr \text{ بکار گرفته می‌شود. این تابع به} \\ \tau(\tilde{a}) &= \frac{1}{2} \left(a^m + \frac{1}{2}(a^l + a^u) \right) \text{ کاهش می‌یابد. برای اعداد فازی مثلثی } \tilde{a} = (a^m, a^l, a^u) \text{ و } \tilde{b} = (b^m, b^l, b^u) \\ \text{خواهیم داشت:} \end{aligned}$$

$$(\tilde{a} \succeq \tilde{b}) \Leftrightarrow \left(a^m + \frac{1}{2}(a^l + a^u) \geq b^m + \frac{1}{2}(b^l + b^u) \right) \quad (۳)$$

۴-۴ مدل تحقیق

فرض کنید n مشاهده فازی وجود دارد که توسط $\tilde{z}_j$ ($j=1, \dots, n$) نمایش داده می‌شود و هر یک با k عامل مستقل به صورت $\tilde{z}_{ij}$ ($i=1, \dots, k$) برای شاخص z ام نشان داده می‌شود. فرض بر آن است که مشاهدات در دو گروه G_1 و G_2 ، به ترتیب هر کدام با n_1 و n_2 مشاهده طبقه‌بندی می‌شوند. $G_1 \cup G_2 = G$ و $n_1 + n_2 = n$ برقرار است. عضویت هر یک از مشاهدات G معین است. هدف، یافتن قاعده‌ای برای طبقه‌بندی صحیح‌تر مشاهدات نادرست می‌باشد. قاعده مورد نظر در شناسایی عضویت هر یک از مشاهدات نادرست جدید مؤثر و سودمند است. فرآیند یافتن این قاعده در دو مرحله صورت می‌گیرد. با تلاش برای کمینه‌سازی انحراف کل در مشاهداتی که به

نادرستی طبقه‌بندی شده است، مرحله اول به تابع غیرپارامتری دست می‌یابد که مشاهدات را مجدداً طبقه‌بندی می‌کند. مرحله دوم نحوه تعیین عضویت مشاهداتی را مورد بحث و بررسی قرار می‌دهد که در مرحله اول با موفقیت طبقه‌بندی نشده بوده‌اند. از منظر ریاضی، فرمول‌بندی این دو مرحله به‌صورت ذیل بیان می‌شود:

$$\begin{aligned}
 \text{Min} \quad & \sum_{j \in G_1} s_{1j}^+ + \sum_{j \in G_2} s_{2j}^- \\
 \text{s.t.} \quad & \sum_{i=1}^k (\lambda_i^+ - \lambda_i^-) \tilde{z}_{ij} + s_{1j}^+ - s_{1j}^- \approx d+1 \quad j \in G_1 \\
 & \sum_{i=1}^k (\lambda_i^+ - \lambda_i^-) \tilde{z}_{ij} + s_{2j}^+ - s_{2j}^- \approx d \quad j \in G_2 \\
 & \sum_{i=1}^k (\lambda_i^+ + \lambda_i^-) = 1 \\
 & s_{1j}^+, s_{1j}^-, s_{2j}^+, s_{2j}^- \geq 0 \quad j = 1, \dots, n \\
 & \lambda_i^+, \lambda_i^- \geq 0 \quad i = 1, \dots, k
 \end{aligned} \tag{۴}$$

مدل فوق‌فازی است. به منظور استخراج مدل قطعی متناظر، تابع رتبه‌بندی بکار گرفته می‌شود. با توجه به ویژگی $A \approx B$ ، اگر و فقط اگر $\tau(A) = \tau(B)$ برقرار باشد، می‌توان مدل تحقیق را به‌صورت ذیل تغییر داد:

$$\begin{aligned}
 \text{Min} \quad & \sum_{j \in G_1} s_{1j}^+ + \sum_{j \in G_2} s_{2j}^- \\
 \text{s.t.} \quad & \tau \left(\sum_{i=1}^k (\lambda_i^+ - \lambda_i^-) \tilde{z}_{ij} + s_{1j}^+ - s_{1j}^- \right) = \tau(d+1) \quad j \in G_1 \\
 & \tau \left(\sum_{i=1}^k (\lambda_i^+ - \lambda_i^-) \tilde{z}_{ij} + s_{2j}^+ - s_{2j}^- \right) = \tau(d) \quad j \in G_2 \\
 & \sum_{i=1}^k (\lambda_i^+ + \lambda_i^-) = 1 \\
 & s_{1j}^+, s_{1j}^-, s_{2j}^+, s_{2j}^- \geq 0 \quad j = 1, \dots, n \\
 & \lambda_i^+, \lambda_i^- \geq 0 \quad i = 1, \dots, k
 \end{aligned} \tag{۵}$$

با توجه به ویژگی‌های تابع رتبه‌بندی (τ)، مقادیر متناظر ذیل به جای دو محدودیت نخست به‌دست می‌آید:

$$\begin{aligned}
 \sum_{i=1}^k (\lambda_i^+ - \lambda_i^-) \tau(\tilde{z}_{ij}) + s_{1j}^+ - s_{1j}^- &= d+1 \quad j \in G_1 \\
 \sum_{i=1}^k (\lambda_i^+ - \lambda_i^-) \tau(\tilde{z}_{ij}) + s_{2j}^+ - s_{2j}^- &= d \quad j \in G_2
 \end{aligned} \tag{۶}$$

بنابراین، با در نظر گرفتن متناظر محاسبه شده در تابع (τ)، فرمول‌بندی مدل نهایی به‌دست آمده در مرحله یک به‌صورت ذیل بیان می‌شود:

$$\begin{aligned}
 & \text{Min} \quad \sum_{j \in G_1} s_{1j}^+ + \sum_{j \in G_2} s_{2j}^- \\
 & \text{s.t.} \quad \frac{1}{2} \sum_{i=1}^k (\lambda_i^+ - \lambda_i^-) \left[z_{ij}^m + \frac{z_{ij}^l + z_{ij}^u}{2} \right] + s_{1j}^+ - s_{1j}^- = d + 1 \quad j \in G_1 \\
 & \quad \frac{1}{2} \sum_{i=1}^k (\lambda_i^+ - \lambda_i^-) \left[z_{ij}^m + \frac{z_{ij}^l + z_{ij}^u}{2} \right] + s_{2j}^+ - s_{2j}^- = d \quad j \in G_2 \\
 & \quad \sum_{i=1}^k (\lambda_i^+ + \lambda_i^-) = 1 \\
 & \quad s_{1j}^+, s_{1j}^-, s_{2j}^+, s_{2j}^- \geq 0 \quad j = 1, \dots, n \\
 & \quad \lambda_i^+, \lambda_i^- \geq 0 \quad i = 1, \dots, k
 \end{aligned} \tag{۷}$$

در اینجا، λ_i^* و d^* راه حل های بهینه مدل فوق به شمار می روند:

$$\begin{aligned}
 1. \quad & \frac{1}{2} \sum_{i=1}^k (\lambda_i^{*+} - \lambda_i^{*-}) \left[z_{ij}^m + \frac{z_{ij}^l + z_{ij}^u}{2} \right] \leq d^* \Rightarrow z_m \in G_2 \\
 2. \quad & \frac{1}{2} \sum_{i=1}^k (\lambda_i^{*+} - \lambda_i^{*-}) \left[z_{ij}^m + \frac{z_{ij}^l + z_{ij}^u}{2} \right] \geq d^* + 1 \Rightarrow z_m \in G_1
 \end{aligned} \tag{۸}$$

در غیر این صورت، z_m به منطقه همپوشانی تعلق دارد. به منظور طبقه بندی z_m ، مرحله دوم آغاز می شود. پیش از آغاز مرحله دوم، مجموعه ذیل تعریف می شود:

$$\begin{aligned}
 R_1 &= \left\{ j \in G : \frac{1}{2} \sum_{i=1}^k \lambda_i^* \left(z_{ij}^m + \frac{z_{ij}^l + z_{ij}^u}{2} \right) \geq d^* + 1 \right\} \\
 R_2 &= \left\{ j \in G : \frac{1}{2} \sum_{i=1}^k \lambda_i^* \left(z_{ij}^m + \frac{z_{ij}^l + z_{ij}^u}{2} \right) \leq d^* \right\} \\
 R_0 &= \left\{ j \in G : d^* < \frac{1}{2} \sum_{i=1}^k \lambda_i^* \left(z_{ij}^m + \frac{z_{ij}^l + z_{ij}^u}{2} \right) < d^* + 1 \right\} \\
 C_1 &= \{ j \in R_1 : j \in G_1 \}, C_2 = \{ j \in R_2 : j \in G_2 \}, G'_1 = G_1 - C_1, G'_2 = G_2 - C_2
 \end{aligned} \tag{۹}$$

از این رو، فرمول بندی مرحله دوم به صورت مدل (۱۰) بیان می شود:

$$\begin{aligned}
 & \text{Min} \quad \sum_{j \in G'_1} s_{1j}^+ + \sum_{j \in G'_2} s_{2j}^- \\
 & \text{s.t.} \quad \frac{1}{2} \sum_{i=1}^k (\lambda_i^+ - \lambda_i^-) \left[z_{ij}^m + \frac{z_{ij}^l + z_{ij}^u}{2} \right] \geq d+1 \quad j \in C_1 \\
 & \quad \frac{1}{2} \sum_{i=1}^k (\lambda_i^+ - \lambda_i^-) \left[z_{ij}^m + \frac{z_{ij}^l + z_{ij}^u}{2} \right] + s_{1j}^+ - s_{1j}^- = c \quad j \in G'_1 \\
 & \quad \frac{1}{2} \sum_{i=1}^k (\lambda_i^+ - \lambda_i^-) \left[z_{ij}^m + \frac{z_{ij}^l + z_{ij}^u}{2} \right] + s_{2j}^+ - s_{2j}^- = c \quad j \in G'_2 \\
 & \quad \frac{1}{2} \sum_{i=1}^k (\lambda_i^+ - \lambda_i^-) \left[z_{ij}^m + \frac{z_{ij}^l + z_{ij}^u}{2} \right] \leq d \quad j \in C_2 \\
 & \quad \sum_{i=1}^k (\lambda_i^+ + \lambda_i^-) = 1 \\
 & \quad d \leq c \leq d+1 \\
 & \quad s_{1j}^+, s_{1j}^-, s_{2j}^+, s_{2j}^- \geq 0 \quad j = 1, \dots, n \\
 & \quad \lambda_i^+, \lambda_i^- \geq 0 \quad i = 1, \dots, k
 \end{aligned} \tag{10}$$

در اینجا، c^* و $\lambda_i^* (i=1, \dots, k)$ راه حل های بهینه به دست آمده از مرحله دوم در نظر گرفته می شوند. آنگاه:

$$\begin{aligned}
 1. \quad & \frac{1}{2} \sum_{i=1}^k (\lambda_i^{+*} - \lambda_i^{-*}) \left[z_{im}^m + \frac{z_{im}^l + z_{im}^u}{2} \right] \geq c^* \Rightarrow z_m \in G_1 \\
 2. \quad & \frac{1}{2} \sum_{i=1}^k (\lambda_i^{+*} - \lambda_i^{-*}) \left[z_{im}^m + \frac{z_{im}^l + z_{im}^u}{2} \right] \leq c^* \Rightarrow z_m \in G_2
 \end{aligned} \tag{11}$$

۴-۵ متغیرهای تحقیق

با توجه به اینکه در این تحقیق به امتیازدهی اعتباری مشتریان در شرکت های لیزینگ می پردازیم و از آنجایی که بخش عمده ای از درخواست کنندگان دریافت تسهیلات از شرکت های لیزینگ، اشخاص حقیقی می باشند، تعداد ۱۷ متغیر اصلی و مؤثر بر ریسک اعتباری با استفاده از روش 6C انتخاب و مورد ارزیابی قرار گرفته اند. در نهایت با توجه به معیارهای بررسی شده و قضاوت خبرگان از روش نظرسنجی دلفی، تعداد ۸ شاخص $(I_1, \dots, I_8)$ به ترتیب سن متقاضی، میزان مبلغ تسهیلات، نرخ سود تسهیلات، دوره بازپرداخت تسهیلات، میزان مبلغ قسط ماهیانه، میزان مطالبات معوق و زمان دیربازپرداخت اقساط انتخاب و تعیین گردید.

۴-۶ مقادیر داده‌ها

در این تحقیق، جامعه آماری مشتریان حقیقی یک شرکت لیزینگ می‌باشند. به این منظور از میان ۳۶۰ مشتری حقیقی متقاضی اخذ تسهیلات، با استفاده از نمونه‌گیری تصادفی ساده، تعداد ۸۳ مشتری که طی سال‌های ۱۴۰۰ و ۱۴۰۱ از این شرکت تسهیلات دریافت کرده‌اند انتخاب گردید.

۴-۷ تبدیل داده‌های نامطلوب فازی به داده‌های مطلوب فازی

با توجه به حضور داده‌های نامطلوب در این تحقیق، برای مطلوب نمودن آنها می‌توان از روش‌های متفاوتی استفاده نمود که در روش انتخابی، معکوس کلیه مقادیر عامل نامطلوب را با بیشترین مقدار (L, M, U) در بین واحدهای تصمیم‌گیرنده برای آن عامل جمع و سپس مقادیر نهایی را با مقداری ثابت جمع می‌کنیم.

۴-۸ اجرای مدل تحقیق

با توجه به داده‌های این تحقیق، براساس میانگین زمان دیر بازپرداخت اقساط هر مشتری، مشتریان به دو گروه ذیل دسته‌بندی می‌شوند:

- مشتریان خوش حساب (G_1) : آن دسته از مشتریانی که مطالبات آنها در سرفصل مطالبات جاری قرار داشته و میانگین زمان دیر بازپرداخت اقساط آنها از ۳۰ روز بیشتر نباشد.

- مشتریان بدحساب (G_2) : آن دسته از مشتریانی که مطالبات آنها در سرفصل غیرجاری بوده و یا میانگین زمان دیر بازپرداخت اقساط آنها بیش از ۳۰ روز باشد.

اکنون داده‌های جمع‌آوری شده در دو گروه G_1 با تعداد ۵۹ مشاهده و گروه G_2 با تعداد ۲۴ مشاهده که هر یک دارای ۸ عامل مستقل $(I_1, \dots, I_8)$ هستند دسته‌بندی شده‌اند. همچنین مقادیر z_{ij}^l ، z_{ij}^m و z_{ij}^u به ازای $i = 1, \dots, 8$ و $j = 1, \dots, 83$ با هر مشاهده و بالطبع با هر عامل مشاهده مطابقت دارد.

با حل مدل، مقادیر بهینه به دست آمده با استفاده از مدل مرحله یک (رابطه ۷) عبارت است از:

جدول ۱- مقادیر λ_i^* محاسبه شده در مرحله اول

Unit	1	2	3	4	5	6	7	8
λ^{+*}	0.0000	0.4860	0.0000	0.0043	0.0000	0.0000	0.0000	0.0000
λ^{-*}	0.0080	0.4860	0.0000	0.0000	0.0107	0.0000	0.0000	0.0049

همچنین مقدار $d^* = 4.85$ می‌باشد. اکنون با استفاده از راه‌حل‌های بهینه مدل فوق (رابطه ۸) و در اختیار داشتن مقادیر z_m ، مقادیر λ_i^* و d^* که مقادیر به دست آمده به شرح جدول (۲) می‌باشد:

جدول ۲ - مقادیر z_m محاسبه شده در مرحله اول

Unit	z_m	Unit	z_m	Unit	z_m	Unit	z_m	Unit	z_m
1	5.6533	18	5.8353	35	6.1931	52	5.8785	69	4.4269
2	5.9773	19	5.6401	36	5.5635	53	5.9421	70	4.4807
3	5.9577	20	6.0290	37	6.0320	54	5.8515	71	4.1289
4	5.7665	21	5.9879	38	5.7874	55	5.8466	72	4.6691
5	5.9434	22	5.6876	39	5.9312	56	5.9174	73	4.4624
6	5.6975	23	5.8626	40	5.7588	57	6.0722	74	4.6621
7	5.8144	24	5.7930	41	5.7782	58	5.4854	75	4.5840
8	5.7449	25	5.8571	42	5.9149	59	5.9011	76	4.6705
9	5.8099	26	5.7628	43	5.8332	60	4.2218	77	4.4097
10	5.4172	27	6.0009	44	5.9517	61	4.5362	78	4.6587
11	5.4172	28	5.7870	45	5.8743	62	4.4574	79	4.7349
12	5.4172	29	5.5641	46	5.8693	63	4.1703	80	4.7553
13	5.4172	30	5.8266	47	5.7782	64	4.5869	81	4.2447
14	5.4172	31	6.0071	48	5.8428	65	4.4321	82	4.6405
15	5.4189	32	5.9532	49	5.7921	66	4.4805	83	4.8880
16	5.9386	33	5.7760	50	5.8106	67	4.5393	-	-
17	5.8958	34	5.8388	51	6.0947	68	4.6554	-	-

مطابق راه‌حل‌های بهینه ارائه شده، در صورتیکه z_m به هر یک از گروه‌های G_1 و G_2 تعلق نداشته باشد به منطقه همپوشانی تعلق دارد. مانند Unit10 که z_m آن معادل 5.4172 است و از آنجایی که این عدد مابین مقدار عددی d^* و $d^* + 1$ است، به منظور طبقه‌بندی باید وارد مرحله دوم شود.

اکنون پیش از آغاز مرحله دوم می‌بایست مطابق مجموعه‌های تعریف شده در رابطه (۹)، مشاهده‌های مربوط به هر یک از مجموعه‌های $R_1, R_2, R_0, C_1, C_2, G'_1$ و G'_2 را مشخص نماییم. سپس با استفاده از رابطه (۱۰) مقادیر بهینه را محاسبه می‌نمائیم که طبق جدول (۳) به دست می‌آیند:

جدول ۳ - مقادیر λ_i^* محاسبه شده در مرحله دوم

Unit	1	2	3	4	5	6	7	8
λ^+	0.0000	0.4654	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000
λ^-	0.0118	0.4654	0.0000	0.0168	0.0339	0.0000	0.0000	0.0067

و مقدار $c^* = 0$ و $d^* = -0.222$ می‌باشد. اکنون با استفاده از راه‌حل‌های بهینه مدل فوق (۱۱) و در اختیار داشتن $\lambda_i^* (i=1, \dots, k)$ و c^* ، مقادیر z_m مرحله دوم را محاسبه می‌کنیم.

جدول ۴ - مقادیر z_m محاسبه شده در مرحله دوم

Unit	Z_m	Unit	Z_m	Unit	Z_m	Unit	Z_m	Unit	Z_m
1	0.4451	18	0.7020	35	1.1979	52	0.7977	69	-0.7683
2	0.8728	19	0.4481	36	0.0568	53	0.8703	70	-0.7045
3	0.8373	20	0.9849	37	0.9666	54	0.4912	71	-1.1036
4	0.6243	21	0.9210	38	0.5678	55	0.4833	72	-0.3598
5	0.8048	22	0.4493	39	0.8072	56	0.7825	73	-0.6232
6	0.4811	23	0.7380	40	0.6157	57	1.0220	74	-0.4659
7	0.6375	24	0.6206	41	0.6712	58	0.0430	75	-0.2233
8	0.5475	25	0.7200	42	0.8377	59	0.4671	76	-0.4526
9	0.6506	26	0.4982	43	0.7462	60	-1.0522	77	-0.8035
10	0.4028	27	0.9075	44	0.8874	61	-0.6651	78	-0.7303
11	0.4028	28	0.6277	45	0.8163	62	-0.7338	79	-0.6156
12	0.4028	29	0.3019	46	0.7598	63	-1.1579	80	-0.5797
13	0.4028	30	0.6908	47	0.6712	64	-0.5841	81	-1.1929
14	0.4028	31	0.7265	48	0.7728	65	-0.7821	82	-0.7524
15	0.4154	32	0.8504	49	0.6761	66	-0.1978	83	-0.3816
16	0.8452	33	0.6510	50	0.6668	67	-0.6548	-	-
17	0.7898	34	0.7357	51	1.1020	68	-0.4876	-	-

همانگونه که مشهود است، هر یک از مشاهدات به طور قطعی در گروه G_1 و G_2 قرار گرفته اند و اکنون می توان با ورود مشاهده جدید و محاسبه z_m مرحله اول و در صورت لزوم z_m مرحله دوم، وضعیت اعتباری آن را پیش بینی نمود.

❖ مشتری جدید اول با شاخص های واقعی زیر را در نظر بگیریم:

جدول ۵ - شاخص های مربوط به مشتری جدید اول

Unit	I1	I2	I3	I4	I5	I6	I7	I8
L	32	455750000	2000000000	22	36	74398148	76509833	37
M	32	455750000	2000000000	22	36	74398148	75910571	27
U	32	455750000	2000000000	22	36	74398148	54512293	2

با محاسبه مقدار $z_m = 6.21$ مشاهده می شود که مقدار z_m از مقدار $d^* + 1 = 5.85$ بزرگتر می باشد. لذا این مشتری در گروه G_1 یا همان مشتریان خوش حساب قرار می گیرد.

❖ مشتری جدید دوم با شاخص های واقعی زیر را در نظر بگیریم:

جدول ۶- شاخص‌های مربوط به مشتری جدید دوم

Unit	I1	I2	I3	I4	I5	I6	I7	I8
L	57	2160437500	6700000000	10	60	14004861	15238824	201
M	57	2160437500	6700000000	10	60	14004861	14391625	63
U	57	2160437500	6700000000	10	60	14004861	14023278	3

با محاسبه مقدار $z_m = 4.82$ مشاهده می‌شود که مقدار z_m از مقدار $d^* = 4.85$ کوچکتر می‌باشد. لذا این مشتری در گروه G_2 یا همان مشتریان بدحساب قرار می‌گیرد.

❖ مشتری جدید سوم با شاخص‌های واقعی زیر را در نظر بگیریم:

جدول ۷- شاخص‌های مربوط به مشتری جدید سوم

Unit	I1	I2	I3	I4	I5	I6	I7	I8
L	63	6407613500	27800000000	19	48	80383507	82859622	34
M	63	6407613500	27800000000	19	48	80383507	79396052	9
U	63	6407613500	27800000000	19	48	80383507	62080824	-8

با محاسبه مقدار $z_m = 5.77$ مشاهده می‌شود که مقدار z_m بین مقادیر $d^* = 4.85$ و $d^* + 1 = 5.85$ قرار دارد. لذا این مشتری در منطقه همپوشانی قرار دارد. با اجرای مرحله دوم مدل و محاسبه $c = 0$ و $z_m = 0.074$ ، از آنجایی که مقدار z_m از مقدار c بزرگتر است، نتیجه می‌گیریم که این مشتری در گروه G_1 و جزء مشتریان خوش حساب طبقه‌بندی می‌شود.

۵- نتیجه‌گیری

در این تحقیق با استفاده از یک تکنیک جدید تلفیقی، با روش DEA/DA و با معیار مینیمم کردن مجموع انحرافات و همچنین با حضور داده‌های فازی یک فرآیند محاسباتی دو مرحله‌ای در پیش گرفته شد. هر چند روش‌های تلفیقی DEA/DA توسط سوئیچی-کریه‌ارا (۱۹۹۸) ارایه گردیده ولی مدل به کار رفته شده در این تحقیق مدل قدرتمندتری نسبت به مدل‌های پیشین است و علاوه بر رفع ایرادات، نتایج با مینیمم کردن انحراف کل، با موفقیت بیشتری طبقه‌بندی گردید. نتایج به دست آمده در مرحله اول قادر به طبقه‌بندی مشاهدات در هر یک از گروه‌های خوش حساب و بدحساب بوده و در صورت عدم تشخیص گروه مناسب برای هر کدام از مشاهدات (همپوشانی)، مرحله دوم قادر به طبقه‌بندی مشاهده به صورت دقیق است. با استفاده از نتایج این مدل می‌توان با ورود یک متقاضی جدید جایگاه آن را در هر یک از گروه‌های اعتباری تشخیص داده و در خصوص رفتار اعتباری مشتری جدید پیش‌بینی‌های لازم را اتخاذ نمود. بر اساس نتایج به دست آمده مذکور، می‌توان گفت که بسط DEA/DA بسیار دقیق و سریع عمل می‌کند. نتیجه مهم دیگری که می‌توان گرفت این است که اغلب روش‌های پارامتری خطی و غیرخطی، وضعیت نشدنی زیادی ایجاد می‌کنند. ضمن اینکه دو روش غیرپارامتری DEA/DA و بسط DEA/DA برای هر مجموعه داده شده، شدنی هستند و این نتیجه یک مزیت روش غیرپارامتری را نشان



می‌دهد. بسط DEA/DA از خطای مشخصه نیز جلوگیری می‌کند، بنابراین می‌تواند هر یک از مجموعه داده‌ها را پوشش دهد.

پیشنهادهای کاربردی این تحقیق، عبارتند از:

۱. به کارگیری و تجزیه و تحلیل اطلاعات به دست آمده جهت تدوین برنامه‌ها و تصمیم‌گیری‌های مربوطه برای مشتریان جدید.
۲. تغییر کاهشی مدت معیار جداسازی گروه‌ها و مقایسه نتایج به دست آمده با نتایج فعلی جهت بررسی ریسک اعتباری مشتری جدید.
۳. تعریف شاخص‌های جدید (مالی و غیرمالی) متناسب با صنعت مورد مطالعه.
۴. تعیین و تثبیت شاخص‌ها و متغیرهای مؤثر بر رتبه‌بندی متقاضیان حقوقی دریافت تسهیلات.

۶- منابع

احمدیان علی اشرف (۱۳۹۲)، عوامل موثر بر تعیین رتبه اعتباری مشتریان سیستم بانکی در ایران، پنجمین کنفرانس نظام تأمین مالی در ایران، تهران، ایران.

استیری امیر (۱۳۹۴)، بررسی و بکارگیری شاخص‌های مالی موثر جهت رتبه‌بندی شرکت‌های صنعت محصولات غذایی بورس اوراق بهادار تهران با استفاده از مدل‌های تصمیم‌گیری چند شاخصه، پایان‌نامه کارشناسی ارشد، دانشگاه آزاد اسلامی، تهران، ایران.

شورورزی محمدرضا، مسیح‌آبادی ابوالقاسم و غیائی شهرکی مهدی (۱۳۹۱)، ارزیابی ریسک اعتباری مشتریان حقوقی بانک صادرات بر مبنای دو مدل شبکه‌های عصبی مصنوعی و لوجیت و مقایسه آنها، اولین کنفرانس ملی حسابداری مدیریت مالی و سرمایه‌گذاری، ۲۶ بهمن ماه، گرگان، گلستان، ایران.

طالبی محمد و شیرزادی نازنین، (۱۳۹۰)، ریسک اعتباری: اندازه‌گیری و مدیریت، سازمان چاپ و انتشارات اوقاف، چاپ اول، ایران.

عریانی بهاره (۱۳۸۴)، رتبه‌بندی ریسک اعتباری مشتریان حقوقی بانک‌ها به روش تحلیل پوششی داده‌ها (مطالعه موردی: شعب مختلف بانک کشاورزی استان تهران سال ۱۳۸۰)، پایان‌نامه کارشناسی ارشد، دانشگاه بوعلی سینا، همدان، ایران.

علیزاده فاطمه (۱۳۹۳)، ارائه یک چارچوب بهینه به منظور افزایش دقت رتبه‌بندی ریسک اعتباری مشتریان بانکی با استفاده از تکنیک‌های داده‌کاوی، پایان‌نامه کارشناسی ارشد، دانشگاه پیام نور، تهران، ایران.

عیسی‌زاده سعید و عریانی بهاره (۱۳۸۹)، رتبه‌بندی مشتریان حقوقی بانک‌ها برحسب ریسک اعتباری به روش تحلیل پوششی داده‌ها (مطالعه موردی: شعب بانک کشاورزی)، فصل‌نامه سیاست‌ها و پژوهش‌های اقتصادی، سال هجدهم، شماره ۵۵، صفحه ۸۶-۵۹.

Jahanshahloo G., Hosseinzadeh Lotfi F., Rezai Balf F., Zhiani Reazai F., (2007), **Discriminant Analysis of Interval Data using Monte Carlo Method in Assessment of Overlap**, Applied Mathematics and Computation, 191, 521-532.

Lin T-U., Chiu Sh-H., (2013), **Using independent component analysis and network DEA to improve bank performance evaluation**, Economic Modelling, 32, 608-616.

Marcin T., (2009), **Application of Data Envelopment Analysis in Credit Scoring**, Master's Thesis in Financial Mathematics, Technical Report, Krzysztofik School.

Min JH., Lee YC., (2007), **A Practical Approach to Credit Rating**, Journal of Expert Systems with Applications.

Nemati Koutanaei F., Sajedi H., Khanbabaei M., (2015), **A hybrid data mining model of feature selection algorithms and ensemble learning classifiers for credit scoring**, Journal of

Retailing and Consumer Services, 27, 11-23.

Sueyoshi T., (1999), **DEA-Discriminant Analysis in the View of Goal Programming**, European Journal Operational Research, 115, 564-582.

Sueyoshi T., (2001), **Extended DEA-Discriminant Analysis**, European Journal Operational Research, 131, 324-351.

Sueyoshi T., (2004), **Mixed Integer Programming Approach of Extended DEA-Discriminant Analysis**, European Journal Operational Research, 152, 45-55.

Sueyoshi T., (2006), **DEA-Discriminant Analysis: Methodological Comprison Among Eight Discriminant Analysis Approaches**, European Journal Operational Research, 169, 247-272.

Verberaken T., Bravo C., Weber R., Baesence B., (2014), **Development and application of consumer credit scoring models using profit-based classification measures**, European Journal of Operational Research, 238, 505-513.

Yalaln R., Emel A.B., Oral M., Reisman A., (2003), **A Credit Scoring Methodology for the Commercial Banking Sector**, Socio Economic Planning Sciences, 37, 103-123.

Customer Credit Scoring Using Data Envelopment Analysis & Discriminant Analysis in Fuzzy Environments (Case Study: Leasing Industry)

Saeed Kashanifar¹

Alireza Alinezhad²

Abstract

In order to manage and control the credit risk of customers, this paper provided a combined model of data envelopment analysis (DEA) and discriminant analysis (DEA/DA) – to distinguish the presence (or absence) of an overlap between two different classes using a separating hyperplane. Under the assumption of n observations for k independent in a set of fuzzy data, the observations were categorized into two groups; namely, G_1 (Customers with a good credit standing) and G_2 (Customers with a poor credit standing). The research variables were selected by the 6C technique from 17 selected criteria. A total of 8 effective criteria were confirmed using the Delphi method and included in the model. The indicators were applied for 83 real customers of Eghtesad Novin Leasing Company, who received loans in 2021 and 2022. The data collected were divided into two groups: G_1 ($n=59$), and G_2 ($N=24$); 8 independent factors were included in each group ($I_1, \dots, I_8$). The results revealed that each observation was definitely placed in the two groups (G_1 and G_2) and its credit scoring was predicted by the arrival of new observation. In order to evaluate and validate the model performance, a comparative analysis was performed between the extended DEA/DA approach and six different discriminant analysis (DA) techniques for a set of real data. The model proposed here obtained a correct classification rate of almost 99.3% against the other methods.

Keywords: Data Envelopment Analysis / Discriminant Analysis (DEA/DA); Fuzzy Data; Credit Risk; Credit Scoring.

¹ Managing Director in Sanat Yaran Mohsenin Leasing Co. saeed.kashanifar@gmail.com

² Associate Professor, Islamic Azad University, Qazvin, Iran