

پیش‌بینی ریسک مشتریان در صنعت لیزینگ با استفاده از تکنیک‌های داده‌کاوی در راستای خوشه‌بندی

عبدالحسین کرمپور^۱ عبدالحسین زیارتنی^۲ سارا سهرابی مطلق فاضلی^۳

تاریخ ارسال: ۱۴۰۳/۱۱/۰۹

چکیده:

هدف اصلی هدایت پژوهش، ارائه روشی جامع برای طبقه‌بندی مشتریان بر اساس اعتبار آنها در چارچوب کشورهای در حال توسعه بود. در این تحقیق، جهت بخش‌بندی صنعت لیزینگ بر اساس ارزش اعتباری مشتری از مدل آر. اف. ام. و برخی تکنیک‌های داده‌کاوی و روش‌های مختلف علمی در قالب فرایندی خاص، بهره‌گرفته شد. استفاده از جدول تحلیل واریانس جهت تعیین میزان اثرگذاری هر یک از شاخص‌ها در تفکیک خوشه‌ها، ارائه تحلیل درون خوشه‌ای به منظور بررسی دقیق‌تر ویژگی‌های مشتریان در خوشه دارای بیشترین مشتری و همچنین بخش‌بندی نهایی مشتریان بر اساس ارزش اعتباری مشتریان در قالب جدول ارزش اعتبار مشتریان از موارد خاصی است که در این مطالعه از آنها بهره‌گرفته شده است، درحالی که در مطالعات گذشته به آنها اشاره‌ای نگردیده است. روش پیشنهادی با ارائه یک مطالعه موردی نشان داد که از پایایی بالایی برخوردار است. همچنین که در مقایسه تطبیقی با دیگر روش‌های درختی و خوشه‌بندی از صحت و دقت بیشتری برخوردار بود. استفاده از سیستم RFM هنوز هم در فرآیندهای طبقه‌بندی بصورت قاطع کاربرد دارد و می‌تواند با دیگر روشها بصورت ترکیبی مدلی قویتر را ارائه نماید.

واژه‌های کلیدی: پیش‌بینی ریسک مشتریان، تکنیک‌های داده‌کاوی، خوشه‌بندی، صنعت لیزینگ.
طبقه‌بندی JEL: بازارهای مالی، سرمایه‌گذاری مالی، حکمرانی و تأمین مالی شرکتی (o16).

^۱ استادیار و عضو هیأت علمی گروه مدیریت بازرگانی دانشکده مدیریت دانشگاه آزاد اسلامی واحد تهران شمال (a.karampour@iau-tnb.ac.ir)

^۲ کارشناس ارشد مدیریت بازرگانی_بازاریابی دانشکده فارابی دانشگاه تهران (ah.ziyarati@gmail.com)

^۳ کارشناسی ارشد مدیریت کسب و کار_بازاریابی دانشگاه آزاد اسلامی واحد تهران شمال (sara.sohrab1981@gmail.com)

۱. مقدمه

صورت یک مساله غالباً اساسی تر از حل آن است. حل مساله ممکن است فقط مستلزم مهارت های تجربی و ریاضی باشد، حال آنکه طرح سوالات نو و بررسی مسائل قدیمی از دیدگاهی تازه، نیازمند ذهنی خلاق است و این نکته بیانگر آزمودگی فرد در علم است. هر پژوهش در واقع با قصد پاسخگویی و راه حل یابی برای یک مساله اصلی در قالب یک پرسش آغاز می شود. بدین لحاظ برای اینکه انسجام، هدفمندی و کاربردی بودن پژوهش حفظ شود، بایستی حول مساله اصلی سازماندهی شود (زارعی نژاد، ۱۳۹۲). با در نظر گرفتن مبانی و اهداف خاص تحقیق، سوال اصلی تحقیق بدین صورت مطرح می گردد:

چگونه می توان با استفاده از مدل های داده کاوی یک الگوی بهینه در زمینه طبقه بندی مشتریان بر اساس ریسک اعتباری در صنعت لیزینگ طراحی نمود؟

۲. بدنه اصلی

همانطور که در ابتدا اشاره شد فرآیند شناسایی مشتریان هدف در شرایط رقابتی دارای اهمیت زیادی است. دسته بندی مشتریان بر اساس تکنیک های داده کاوی جهت شناسایی مشتریان هدف به کار می رود. هدف اصلی این تحقیق بررسی روش های مختلف داده کاوی در خوشه بندی مشتریان است. درخت تصمیم، رگرسیون لجستیک، شبکه عصبی، نزدیک ترین همسایگی و ...، از جمله تکنیک های داده کاوی هستند که در این رساله مورد تحقیق و بررسی قرار خواهند گرفت. اهداف اختصاصی تحقیق عبارتند از:

هدف اول پژوهش، بررسی و تجزیه و تحلیل مسائل و چالش های درگیر در طبقه بندی مشتریان در کشورهای در حال توسعه است.

هدف دوم بررسی و تجزیه و تحلیل مناسب و همچنین تلفیق روش های داده کاوی، یعنی، RFM و شبکه عصبی که می تواند برای طبقه بندی مشتریان بر اساس ریسک اعتباری استفاده شود.

هدف سوم تدوین و فرموله کردن روش های داده کاوی برای طبقه بندی مشتریان بالقوه بر اساس ریسک اعتباری در چارچوب مباحث نظری.

هدف چهارم. مقایسه روش های داده کاوی ترکیبی با رگرسیون لجستیک.

۳-۴. سؤالات تحقیق

سوال اصلی تحقیق:

چگونه می توان در یک شرکت یا سازمان، مشتریان را به لحاظ ریسک اقتصادی با استفاده از الگوریتم های داده کاوی طبقه بندی کرد؟

سؤالات فرعی

۱. کدام الگوریتم خوشه بندی صحیح تری را نسبت به سایر الگوریتم های خوشه بندی داده کاوی ارائه

می کند؟

۲. کدام شاخص مهمترین شاخص در خوشه بندی مشتریان می باشد؟

۳. کدام مدل الگوریتم داده کاوی نسبت به سایر الگوریتم ها از صحت بالاتری برخوردار میباشد؟

۴. آیا بهینه سازی خوشه بندی بر مبنای الگوریتم های هوشمند (ژنتیک، ازدحام ذرات، جستجوی فاخته، رقابت استعماری و...) از صحت بالاتری برخوردار نیست؟

۳-۵. روش تحقیق

۱-۵-۳ مدل مفهومی پژوهش

اندازه گیری ارزش دوره عمر مشتری (CLV) یکی از ابزارها برای خوشه بندی مشتریان است که به کمک آن می توان با هدف گیری مشتریان سودآور و حتی حذف برخی از مشتریان و یا کاهش فعالیت های بازاریابی برای بخشی از آنها، سودآوری را تضمین نمود. در این تحقیق از الگوی RFM^۴ بولت (۲۰۱۴)، که شامل سه متغیر اصلی تازگی خرید (R)، فراوانی خرید (F) و ارزش مالی خریدهای مشتری (M) است استفاده شده است و اهمیت نسبی سه متغیر مذکور به عنوان دیگر مقادیر لازم با استفاده از روش فرآیند تحلیل سلسله مراتبی AHP تعیین خواهد شد و اهمیت هر یک از این شاخص ها در مدل WRFM جهت خوشه بندی و رتبه بندی مشتریان به کار برده خواهد شد. همچنین روش شبکه عصبی نیز در این تحقیق بکار برده می شود.

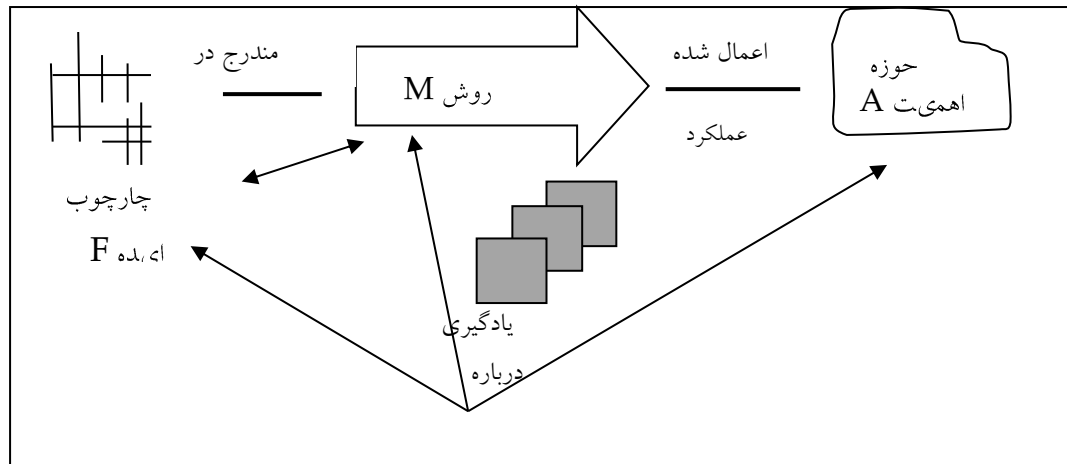
معمولا تحقیقات طبق معیارهای خاص در طبقه بندی های متفاوتی جای می گیرند. تحقیقات علمی بر اساس هدف تحقیق به سه دسته بنیادی، کاربردی و تحقیق و توسعه تقسیم می شود. همچنین تحقیقات علمی را بر اساس چگونگی به دست آوردن داده های مورد نیاز (طرح تحقیق) می توان به دو دسته تحقیق توصیفی و تحقیق اکتشافی تقسیم نمود. به منظور برقراری رابطه علت-معلولی میان دو یا چند متغیر از طرح های آزمایشی استفاده می شود.

تحقیق توصیفی شامل مجموعه روش هایی است که هدف آنها توصیف کردن شرایط یا پدیده های مورد بررسی است. اجرای تحقیقات توصیفی می تواند صرفا برای شناخت بیشتر شرایط موجود یا یاری دادن به فرآیند تصمیم گیری باشد. تحقیق توصیفی شامل مجموعه روش هایی است که هدف آنها توصیف کردن شرایط یا پدیده های مورد بررسی است. تحقیقات توصیفی خود به انواع: پیمایشی، همبستگی، اقدام پژوهشی، بررسی موردی و پس رویدادی (علی-مقایسه ای) تقسیم می شود.

روش تحقیق حاضر از نظر هدف، کاربردی و از نظر نوع جمع آوری اطلاعات و داده ها، تحقیقی اسنادی است. همچنین برای شناسایی متغیرهای تأثیرگذار در رفتار مشتریان از مصاحبه آزاد با خبرگان صنعت مورد مطالعه استفاده شده است. نوع داده های این پژوهش به صورت عددی (پیوسته) و اسمی (گسسته) است.

بر اساس نظر چکلند و هولول (۱۹۹۸)، تمام قسمتهای تحقیق بر طبق مدل ساده ای از اجزاء نشان داده شده در شکل ۱-۳ است. اجزاء ضروری این مدل شامل یک چارچوب هوشمندانه ای از ایده های مرتبط است (F)، که خود یک تئوری است که حاوی اطلاعاتی در مورد مسأله "سؤال" تحقیق است (A). این چهارچوب در شیوه ای

قرار می‌گیرد که استفاده از روشهای متنوع تحقیقی (M) و تکنیکهای مناسب این چهارچوب جهت بازرسی نواحی کاربردی را ازتقاء می‌بخشد. این همان سؤال تحقیق است.



شکل ۱. عناصر مربوط به هر بخش از پژوهش (چکلند و هالول، ۱۹۹۸)

براساس تجزیه و تحلیل علمی در مورد ریسک اعتباری مشتریان نیاز به یک روش داده کاوی برای حل مسأله، حتمی و ضروری بنظر می‌رسید. بنابراین با مشورت کارشناسان حوزه ارتباط با مشتری و ریسک اعتباری، به طور تدریجی چهار چوبی از نظرات (F) شکل گرفت و توسعه یافت. در این پژوهش از روش های داده کاوی موسوم به RFM و SOM برای خوشه بندی مشتریان استفاده شد.

۳-۵-۳. مورد کاوی

همان گونه که قبلا نیز ذکر شد، تمرکز این تحقیق بر روی درک بهتر فرایند های طبقه بندی مشتریان با هدف تقویت کارکرد های کنونی است. مرحله ی بعد شامل آزمایش مدل تلفیقی و پیشنهادی داده کاوی بصورت کلی در یک موقعیت واقعی است، که مستلزم یک روش مورد کاوی است. بر طبق نظر زارعی نژاد (۱۳۹۲) استراتژی های مورد کاوی در بسیاری از موقعیت ها برای کمک به اطلاعات شخصی، گروهی، سازمانی، اجتماعی، سیاسی و مسائل مربوطه استفاده می شود. وی گفته است که مورد کاوی شبیه به یک آزمایش است و بسیاری از شرایط که یک آزمایش را تصدیق می کند، این تحقیق را نیز تصدیق می کند. ایشان استفاده از یک تحقیق مورد کاوی را تنها به پنج دلیل ذکر شده در پایین توجیه و تایید کرد که عبارتند از:

۱- زمانی که این تحقیق یک مورد انتقادی را در آزمایش یک تئوری بدون مشکل نشان دهد.

۲- زمانی که این تحقیق یک مورد افراطی و یا یک مورد منحصر به فرد را نشان دهد.

۳- زمانی که مورد مطالعه ای یک مورد شاخص و یا معرف باشد، و هدفش این باشد که موقعیت ها و شرایط عادی را بدست آورد.

۴- زمانی که مورد تکاملی است، و محقق دارای فرصتی است که پدیده ای را که قبلاً جهت تحقیق علمی در دسترس نبوده مشاهده و تجزیه و تحلیل کند.

۵- زمانی که مورد وابسته به جغرافیا باشد، تا بتوان در دو یا چند نقطه ی مختلف آن را بررسی کرد .

در این مرحله صنعت لیزینگ به عنوان یک تحقیق مورد کاوی انتخاب شد تا اجرای عملی مدل داده کاوی پیشنهادی کلی را برای طبقه بندی مشتریان و پیش بینی ریسک آنها را توصیف کند که در آن فرآیند طبقه بندی بصورت یک فعالیت چالشی نمودار خواهد شد. تحقیق مورد کاوی برای طراحی، ترتیب، تحقیق، تایید واقعیت، مصاحبه اداری و کارشناسان زمان زیادی را صرف کرد. برای تحقیق این مسئله لازم بود با طراحان حوزه مدیریت ارتباط با مشتری انتخاب شده که معمولاً در طرح‌های انتخابی مدیریت ریسک ماهرترند، تعداد زیادی جلسات مذاکره ای برگزار شود در نهایت این امر کارگاهی را می طلبد، که بواسطه ی آن طراحان مدیریت ارتباط با مشتری قادر به مشارکت باشند و ارتباط و تعامل پیوسته و منسجمی را با نویسنده که به عنوان یک تسهیل کننده عمل میکند، بوجود آورد.

۴-۵-۳. مدل RFM

مدل آر.اف.ام. اولین بار توسط هوگس (۱۹۹۴) معرفی گردید . وی برای تحلیل آر.اف.ام. از رفتار گذشته مشتری که به آسانی قابل پیگیری و دسترسی است، استفاده نمود. این مدل از سه بعد مربوط به داده های مبادلاتی مشتریان، برای تحلیل رفتار آنها استفاده مینماید یعنی تاخیر، فرکانس و پول. مدل RFM یک مدل مبتنی بر رفتار است که برای آنالیز رفتار یک مشتری و سپس پیش بینی او بر اساس رفتارش در بانک اطلاعات استفاده می‌شود(یه و همکاران ۲۰۰۹). در بین متغیرهای RFM شاخص تاخر اغلب مهمترین متغیر شناسایی می شود. با این حال بر اساس مطالعات انجام شده در گذشته متغیرهای RFM در واقع Firm-specific هستند و بر اساس طبیعت محصولات شرکت اهمیت آنها متفاوت است (لومسدن ۲۰۰۸). در این الگو r فاصله زمانی آخرین خرید مشتری تا زمان حال، F تعداد خریده‌ها در یک دوره زمانی مشخص و M مبلغ اسمی خریداری شده در دوره مدنظر تعریف می شود (بولت ۲۰۱۴). این پژوهش از مدل RFM برای تبدیل داده های اولیه به فرم دلخواه برای استفاده در الگوریتم خوشه بندی استفاده کرده است در نتیجه تغییراتی در نحوه محاسبه آنها صورت می گیرد. شاخص تاخیر: تازگی خرید (R)، و ارزش مالی خریدهای مشتری، شاخص فراوانی خرید (F): تعداد ماههایی که مشتری تراکنش خرید بیشتری دارد. ارزش مالی خریدهای مشتری: مجموع گردش مشتری را در طول این دوره زمانی در نظر گرفتیم .

۵-۵-۳. روش SOM

برای امتیاز اعتباری و تجزیه و تحلیل امتیازی رفتاری، بسیاری از مطالعات ارائه شده اند که شبکه های عصبی اصلی ترین [نها است که بنظر میرسد بهتر از روش های آماری باشد (هسی ۲۰۰۴). الگوریتم SOM خوشه بندی از نوع شبکه عصبی است که در سال ۱۹۸۱ توسط پژوهشگر فنلاندی کوهنن اختراع شد. این الگوریتم بطور معمول متشکل از دو لایه نورون های ورودی و خروجی است (حسینی، ۲۰۱۳). بطور کلی شبکه عصبی از لایه های نورونی تشکیل شده اند. این نوع نورون ها از طریق ورودی های خود با جهان واقعی در ارتباطند و گروه دیگر از نورون ها نیز از طریق خروجی های خود، جهان بیرون را میسازند (تیسای ۲۰۰۹).

برای تعیین بهترین تعداد خوشه ها از روشی بنام SSE استفاده میشود. در این روش نخست مراکز خوشه در نظر گرفته می شود و سپس فاصله نقطه مورد نظر از مراکز خوشه محاسبه میشود. خوشه ای که SSE پایین تری دارد نشان دهند بهترین خوشه بندی است برای این امر از فرمول زیر استفاده میشود:

$$SSE = \sum_{i=1}^k \sum_{p \in c} d(P, m_i)^2$$

۳-۵-۶. روش امتیازدهی روش پیشنهادی

ارزش هر مشتری را میتوان بر اساس تازگی، تکرار و ارزش پول بصورت زیر محاسبه نمود (غضنفری ۲۰۱۱):

$$V(c_i) = W^R \times R(c_i) + W^F \times F(c_i) + W^M \times M(c_i)$$

که در آن $R(c_i)$ ، $F(c_i)$ و $M(c_i)$ به ترتیب امتیازات مشتری c_i با توجه به معیارهای RFM هستند. W^R ، W^F و W^M اهمیت وزنی برای معیارهای RFM را نشان میدهند.

$$W^F + W^R + W^M = 1$$

سودآوری خوشه O_n با محاسبه میانگین ارزش همه مشتری های خوشه O_n حاصل می شود. از این رو میتوان آن را از طریق معادله زیر تعریف نمود (غضنفری ۲۰۱۱).

$$V(O^n) = W^R \times R(O^n) + W^F \times F(O^n) + W^M \times M(O^n)$$

$$R(O^n) = \frac{\sum_{C_i \in O^n} R(C_i)}{\|O^n\|}$$

$$F(O^n) = \frac{\sum_{C_i \in O^n} F(C_i)}{\|O^n\|}$$

$$M(O^n) = \frac{\sum_{C_i \in O^n} M(C_i)}{\|O^n\|}$$

که در آن $R(O^n)$ ، $F(O^n)$ ، $M(O^n)$ امتیاز O_n خوشه با توجه به معیارهای RFM است.

۳-۵-۷. مدل رگرسیون لجستیک

روش رگرسیون لجستیک زمانی که متغیر وابسته بوده و دارای دو حالت است مورد استفاده قرار میگیرد. این روش همچنین آزمونهای آماری مختلفی را ارائه داده و توانایی های تشخیص گسترده ای را در خود دارد. یکی از منافع رگرسیون لجستیک مقادیر احتمال برای احتمال وقوع رخداد های متغیر وابسته است که نمودی عینی به نتایج مدل میدهد. در صورتی که احتمال وقوع بیشتر از حد آستانه که معمولاً ۰/۵ در نظر گرفته میشود، پیش بینی شود. در این صورت وقوع پدیده مورد نظر محتمل تلقی شده و در غیر این صورت غیر محتمل خواهد بود.

نوعی از مدل رگرسیون است که متغیرهای پیش بین (مستقل) هم در مقیاس کمی و هم در مقیاس مقوله ای می تواند باشد و متغیر وابسته، مقوله ای دو سطحی است. این دو مقوله به گونه ای معمول به عضویت یا عدم عضویت در یک گروه اشاره دارد.

برای توضیح مدل رگرسیون می توان از تابع توزیع تجمعی استفاده نمود. توابع توزیع تجمعی از تغییراتی که مقدار P را در فاصله صفر و یک قرار میدهد به وجود می آید در حالیکه همچنان خواص یکنواختی دارد. فرض کنید که یک توزیع نرمال استاندارد برای بیان احتمال انتخاب شده است (کوسمنت و همکاران، ۲۰۱۴).

$$P(y/x) = \Phi(b'x) = \int_{-\infty}^{b'x} \phi(z) dz$$

که در آن X بردار متغیر توضیحی، P احتمال تجمعی و وقوع پیشآمد، b بردار ضریب که در آن $\Phi(Z)$ تابع چگالی نرمال استاندارد است.

این تابع معروف مدل پروبیت است. اگر تابع توزیع لجستیک برای بیان احتمال تجمعی وقوع پیشآمد P مورد استفاده قرار گیرد به مدل لجستیک منجر می شود (کوسنت و همکاران ۲۰۱۴). آنگاه:

$$p(y/x) = \phi(b'x) = \int_{-\infty}^{b'x} \eta(z) dz = \frac{1}{1 + e^{-b'x}} = \frac{e^{b'x}}{1 + e^{b'x}}$$

$$1 - p_i = \frac{1}{1 + e^{b'x}}$$

که در آن $\eta(Z)$ تابع چگالی لجستیک یا لوجیت است. در مدل‌های لوجیت و پروبیت متغیر وابسته بصورت صفر و یک تعریف می شود. لگاریتم نسبت بخت یا لوجیت از رابطه زیر حاصل می شود.

$$L = \ln\left(\frac{P}{1-P}\right) = B_0 + B_1X_1 + \dots + B_nX_n$$

در رگرسیون لجستیک متغیر وابسته یک متغیر صفر و یک است که دو مقدار را به خود اختصاص می دهد. اگر فرض کنیم Y متغیر تصادفی بیاشد که میتواند مقادیر صفر و یک را بخود اختصاص دهد در این صورت احتمال وقوع Y را می توان بشکل رابطه زیر در نظر گرفت:

$$P(Y=1) = P = \frac{e^{B'X}}{1 + e^{B'X}} P(Y=0) = (1-P) = \frac{1}{1 + e^{B'X}}$$

که در آن B بردار سطری ضرایب و X بردار ستونی متغیرهای مستقل است.

در حالت کلی برآورد ضرایب متغیرهای مستقل یعنی بردار B از طریق حداکثرسازی رابطه زیر بدست می آید که توسط مشتق گیری نسبت به هریک از ضرایب متغیرهای مستقل و مساوی صفر قرار دادن هر یک از مشتق ها محاسبه می شود.

$$\ln L = l = \sum_{i=1}^n y_i \ln\left(\frac{e^{\beta X}}{1 + e^{\beta X}}\right) + \sum_{i=1}^n (1 - y_i) \ln\left(\frac{1}{1 + e^{\beta X}}\right)$$

۸-۵-۳. فرآیند تحلیل سلسله مراتبی

روش AHP را توماس ساتی توسعه داد که او کار پیشروی اش را در اوایل دهه ۷۰ آغاز کرد. این یک شیوه تصمیم گیری چند معیاره است که یک مسئله پیچیده را به یک سلسله مراتب که متشکل از عناصر خاص است، تجزیه می کند. این ابزاری موثر است که هم می تواند ویژگی های قابل لمس و هم ویژگی های نامرئی (غیر قابل لمس) را اداره کند، خصوصاً وقتی در حال بررسی مسائل چند جانبه است. به علت کاربرد وسیع، سهولت استفاده و به کار گیری آن، AHP، در همه جا مورد مطالعه قرار گرفته می شود و در عرصه های گوناگون برای حل مسائل تصمیم گیری از آن استفاده می شود. بعلاوه، استفاده از مقایسات دو به دو که بر اساس قضاوت های کارشناسان خبره، برای مقایسه جایگزینهاست، به فرد تصمیم گیرنده این اجازه را می دهد که بر روی مقایسه فقط دو موضوع پردازد و بررسی و مشاهده را تقریباً آزاد از هر گونه تأثیرات بی ربط انجام دهد. اگر تصمیم گیری توسط گروه صورت پذیرد، می توان AHP را به سادگی تنظیم و طراحی کرد تا اینکه بصورت انفرادی انجام شود. راپر - لو ۵ و شارپ ۶ (۱۹۹۰)، به این موضوع پی بردند که طراحی یک مسئله مانند سلسله مراتب، کمک مفیدی به درک مسائل و پیشبرد بحث درباره آنها می کند. این پروسه می تواند مسائلی را که قبلاً صراحتاً توضیح داده نشده اند را فاش کند. علاوه بر این، درک این پروسه اسان است و افراد تصمیم گیرنده با آن راحت هستند. این روش در ۴ مرحله بکار می رود (ساتی ۱۹۸۰):

(۱) طراحی یک سلسله مراتب که به توصیف مسئله پردازد.

(۲) طرح ریزی ماتریس های برای مقایسات دو به دو میان سطوح (رده های) متوالی

(۳) ارائه اولویت ها یا مقادیر نسبی عناصر در هر سطح از سلسله مراتب که از بردارهای ویژه استفاده می کنند.

(۴) ترکیب مقادیر نسبی سطوح مختلف که از طریق مرحله سوم بدست می آیند تا یک نمره کلی از جایگزین های تعمیم ایجاد کنند.

۱-۸-۵-۳. مقایسات دو به دو (زوجی)

پس از طرح ریزی سلسله مراتبی که مسئله را شرح می دهد، مرحله دوم مرحله اندازه گیری و جمع آوری داده هاست که شامل اداره مقایسات دو به دو است. دو نوع مقایسه وجود دارد که انسان ها بوجود می آورند: مقایسه مطلق و نسبی. اندازه گیری مطلق (معروف به امتیاز بندی) برای دسته بندی گزینه ها بر حسب معیارها یا درجه بندی معیار بکار برده می شود؛ برای مثال: عالی، خیلی خوب، خوب، متوسط، زیر متوسط، بد، خیلی بد. در مقایسات مطلق، گزینه ها با یک نمونه استاندارد یا یک خط مبنا که در حافظه فرد وجود دارد و از طریق تجزیه رشد یافته است تطبیق می شوند. اندازه گیری نسبی در مقایسات دو به دو به دوایی دو عنصر نسبت به یک ویژگی

⁵- Roper – Lowe

⁶- Sharp

مشترک بکار برده می شود. در مقایسات نسبی، گزینه ها بصورت دو تا دو تا براساس یک ویژگی مشترک مقایسه می شوند. AHP با هر دو نوع از مقایسه ها برای گرفتن مقیاس نسبی اندازه گیری استفاده شده است (ساعتی، ۱۹۸۰).

مقایسه دو به دو، فرایند جمع اوری داده ای ورودی عناصر تصمیم گیری است، طوریکه بتوان اولویت های مقیاس نسبی را استنتاج کرد. این موضوع در مشکلات (مسائل) تعمیم گیری حائز اهمیت است، تا جائیکه انتخاب یک گزینه خاص بطور فوری، معمولاً غیر ممکن است. به ندرت گزینه ای وجود دارد که از لحاظ تمام معیارها برگزیده دارای ارجحیت باشد. عملاً، این عادت که گزینه های خاص از لحاظ برخی معیارها نسبت به بقیه بهتر هستند در حالیکه بقیه گزینه ها بر حسب معیارهای باقی مانده ارجح تر در نظر گرفته می شوند. این مشکل از طریق مقایسات دو به دویی انتخاب گزینه ها حل می شود.

اصطلاحات کلامی مقیاس اصولی ساتی که دارای شدت ۹ امتیازی مطلق است و در جدول ۳-۲ نشان داده شده است. برای برارود کردن نسبت ها همانند اعداد استفاده می شود، یعنی اهمیت نسبی عناصر را در هر سطحی از سلسله مراتب با عدد نشان می دهد.

جدول ۱. اصطلاحات کلامی مقایسات زوجی

ردیف	قضاوت	مقدار عددی
۱	کاملاً مهم تر و مطلوب تر	۹
۲	ترجیح یا اهمیت خیلی قوی	۷
۳	ترجیح یا اهمیت قوی	۵
۴	کمی مهم تر یا کمی مطلوب تر	۳
۵	اهمیت یکسان و مطلوبیت برابر	۱
۶	ترجیحات بین فواصل	۲ و ۴ و ۶ و ۸

مقدار معکوس نیز به مقایسه معکوس نسبت داده می شود، یعنی $a_{ij} = 1/a_{ji}$.

ماتریس مقایسات زوجی بصورت معادله (۱) نگاره می شود. که در آن a_{ij} نشان دهنده قضاوت های کلامی ست.

$$A = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{bmatrix} \quad (۱)$$

در صورتی که در معادله (۱) رابطه ی $a_{jk} = a_{ik} / a_{ij}$ $i, j, k = 1, \dots, n$ برای تمام مقایسات برقرار باشد، ماتریس A سازگار است (ساتی، ۱۹۸۰).

هنگامی که تمام مقایسات دو به دو در هر سطح انجام شدند و ماتریس های مقایسه بوجود آمدند، مقیاس اولویت های نسبی از طرف مقایسات زوجی کسب می شود. از این رو، بجای تعیین دو عدد w_i و w_j که نسبت w_i / w_j را شکل می دهد، یک عدد مطلق مفرد از مقیاس بنیادی (که بیان می کند چند بار بزرگتر بر کوچکتر

تسلط می یابد) برای تقریب نسبت $(w_i / w_j) / 1$ معین شده است. نزدیک ترین تقریب عدد صحیح نسبت به w_i / w_j است. بنابراین، فرایند مقایسه دو به دو از اندازه واقعی برای عنصرهایی که تحت مقایسه قرار گرفتند استفاده می کند. یعنی، وزن های $w = (w_1, w_2, \dots, w_n)$ که قبلاً بدست آمده اند به سمت ماتریس معکوس A که در معادله (۲) نشان داده شده است، برده می شود (ساتی، ۱۹۸۰):

$$A = \begin{bmatrix} w_1 / w_1 & w_1 / w_2 & \dots & w_1 / w_n \\ w_2 / w_1 & w_2 / w_2 & \dots & w_2 / w_n \\ \dots & \dots & \dots & \dots \\ w_n / w_1 & w_n / w_2 & \dots & w_n / w_n \end{bmatrix} \quad (۳-۲)$$

از آنجا که $a_{ij} = w_i / w_j$ و همچنین $a_{ij} \cdot a_{jk} = \frac{w_i}{w_j} \cdot \frac{w_j}{w_k} = \frac{w_i}{w_k} = a_{ik}$ ، بنابراین بمنظور اثبات مقادیر معکوس خواهیم داشت:

$$a_{ji} = \frac{w_j}{w_i} = \frac{1}{w_i / w_j} = \frac{1}{a_{ij}}$$

تصور کلی ساتی که در بالا بیان شد، منجر به تقریب حکم ماتریس A شد که در معادله (۳-۱) با ماتریس نسبت های داده شده در معادله (۲) نشان داده می شود، که عناصرشان نسبت های مقادیر w_i / w_j هر عنصر (n)

نسبت به بقیه عناصر هستند. ورودی‌های درون ماتریس، سلطه برتری عنصر مسیر سطری را نسبت به عنصر مسیر ستونی بیان می‌کند.

۳. جمع‌بندی و نتیجه‌گیری

در این بخش به توصیف روش پژوهش، اهداف، سوالات و روند تکاملی تحقیق پرداختیم. یک تحقیق زمانی پایایی مناسب را خواهد داشت که از نظر روش شناسی منطق مناسبی داشته باشد. در ادامه روش‌های داده‌کاوی استفاده شده در تحقیق مورد بحث قرار گرفتند. روش RFM یکی از متداول‌ترین روش‌های داده‌کاوی است که برای پیش‌بینی بر اساس سه متغیر تاخر، تناوب و ارزش پول، بکار می‌رود. روش SOM شبکه عصب نیز ارزش روش‌های داده‌کاوی بر پایه تشکیل نوروهای عصبی با روابط پیچیده است که دنیای خارج مسایل ارتباط برقرار می‌کند. از طرفی رگرسیون لجستیک یکی دیگر از روش‌های داده‌کاوی است که بر پایه آمار و احتمالات و متغیرهای وابسته عمل مینماید. در این تحقیق این روشها مورد استفاده قرار گرفتند و مورد مقایسه واقع شدند. برای اتفاقات سازمانی، نمی‌توان توجیه‌های ساده یافت و برای آنها به سادگی علت‌هایی پیدا کرد. بعلاوه غالباً یافتن سرچشمه‌ها مقدور نیست زیرا منشأها معمولاً در فاصله‌ای بسیار دورتر از نتایج و علایم قرار دارند و نتایج و علایم بواسطه حلقه‌های تشدید کننده انحراف از معیار با اصل خود تفاوتی فاحش و غیر قابل تصور می‌یابند. پژوهشگر پس از اینکه روش تحقیق خود را مشخص کرد و با استفاده از ابزارهای مناسب، داده‌های مورد نیاز را برای آزمون فرضیه‌های خود جمع‌آوری کرد، نوبت آن فرا می‌رسد که با بهره‌گیری از تکنیک‌های داده‌پردازی مناسبی که با روش تحقیق، نوع متغیرها، ... سازگاری دارد، داده‌های جمع‌آوری شده را دسته‌بندی و تجزیه و تحلیل نماید و در نهایت فرضیه‌های سوالاتی را که تا این مرحله او را در تحقیق هدایت کرده‌اند در بوته آزمون قرار دهد و تکلیف آنها را روشن کند و سرانجام بتواند پاسخی (راه‌حلی) برای پرسشی که تحقیق (تلاشی سیستماتیک برای بدست آوردن آن بود) بیابد. (خاکی، ص ۳۰۴-۳۰۳ : ۱۳۸۲)

۲-۴. آماده‌سازی و پیش‌پردازش داده‌ها

در این پژوهش از داده‌های تراکنشی مشتریان مربوط به صنعت لیزینگ استفاده شد. در مجموع تعداد ۱۲۳۲ مشتری در فاصله زمانی ۱۴۰۳-۱۴۰۲ جمع‌آوری شد. با توجه به محدودیت‌هایی که در تحویل داده‌های دموگرافیک وجود داشت، تنها داده‌های تراکنشی در اختیار قرار گرفتند. مرحله آماده‌سازی از مهمترین و پیچیده‌ترین مراحل داده‌کاوی است. مراحل بکار رفته شامل فرآیند پاکسازی و کاهش بعد می‌باشد. در نهایت تعداد ۱۶۰ مشتری برای انجام عملیات اعتبار سنجی باقی ماندند. با مدل RFM داده‌های آماده برای استفاده در عملیات اعتبار سنجی شدند.

۳-۴. وزن شاخص‌ها

سه شاخص تولیدشده برای استفاده در الگوریتم خوشه‌بندی بر اساس روش AHP تعیین وزن می‌شوند. این هدف می‌تواند با مد نظر گرفتن معیارهای تازگی خرید (R)، و ارزش مالی خریدهای مشتری، شاخص فراوانی خرید (F): در ارتباط با خوش‌بندی مشتریان کسب شوند. یکی از مراحل اصلی روش تحلیل سلسله‌مراتبی، مقایسات زوجی است. این مرحله یکی از نقاط قوت اصلی مدل AHP است. این مدل نسبت مقیاس اولویت‌ها را به درستی استخراج می‌کند، برعکس رویکردهای سنتی ارزیابی اوزان که توجیه آنها نیز سخت است. از

متخصصین خواسته شد تا با استفاده از اعداد مقیاس اصلی قضاوت خود را بیان کنند: ۱ = برابر، ۳ = نسبتاً مهمتر، ۵ = به شدت مهمتر، ۷ = خیلی به شدت مهمتر است و ۹ = فوق العاده مهمتر است (ساتی، ۱۹۸۰). نتایج جمع اوری شده، سپس، در یک ماتریس متقارن، جهت تشکیل ماتریس‌های قضاوت مقایسه زوجی مربوطه، وارد شدند. ذکر این نکته ضروری است که AHP در مقایسه با ANP به تعداد مقایسه‌های زوجی کمتری نیاز دارد.

پس از بدست آوردن ماتریس مقایسه زوجیا استفاده از بردارهای ویژه نرمال‌سازی شده‌ی حاصل که دربرگیرنده‌ی اولویت‌های نسبی هستند، ارائه می‌شوند، یعنی اهمیت نسبی عناصر سطوح مشابه سلسله‌مراتب نسبت به عنصر سطح بالاتر در یک ویژگی مشترک مورد مقایسه قرار می‌گیرند. در این فصل، روش بردار ویژه، که بردار ویژه‌ی یک ماتریس دوسویه را تخمین می‌زند، به شکلی که در زیر آمده، جهت تخمین بردارهای اولویت نسبی مورد استفاده قرار می‌گیرد. با توجه به این مساله وزن شاخص‌ها طبق ماتریس جدول ۲ محاسبه شد.

جدول ۲. ماتریس مقایسه زوجی

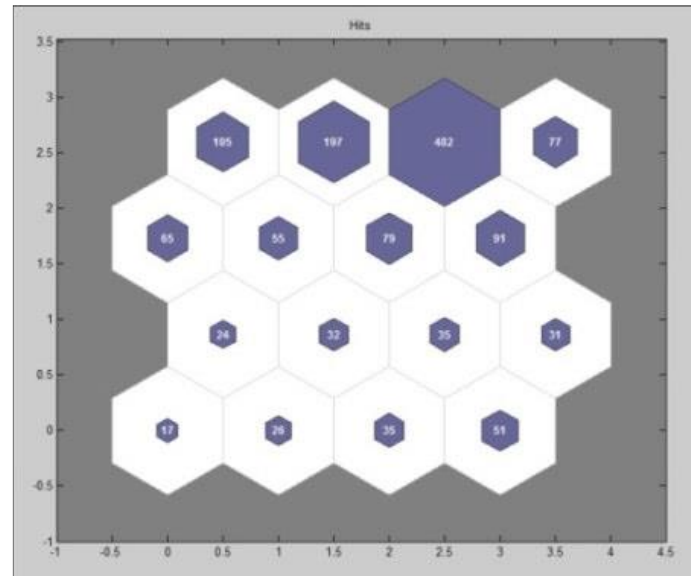
	R	M	F	وزن
R تازگی خرید	1	6	4	0.69
M ارزش مالی خریدهای مشتری	1/6	1	1/3	0.09
F فراوانی خرید	1/4	3	1	0.22

بنابراین وزن شاخص تازگی خرید ۶۹٪، ارزش مالی مشتری ۹٪ و فراوانی خرید ۲۲٪ از مجموع عوامل RFM را دارا هستند.

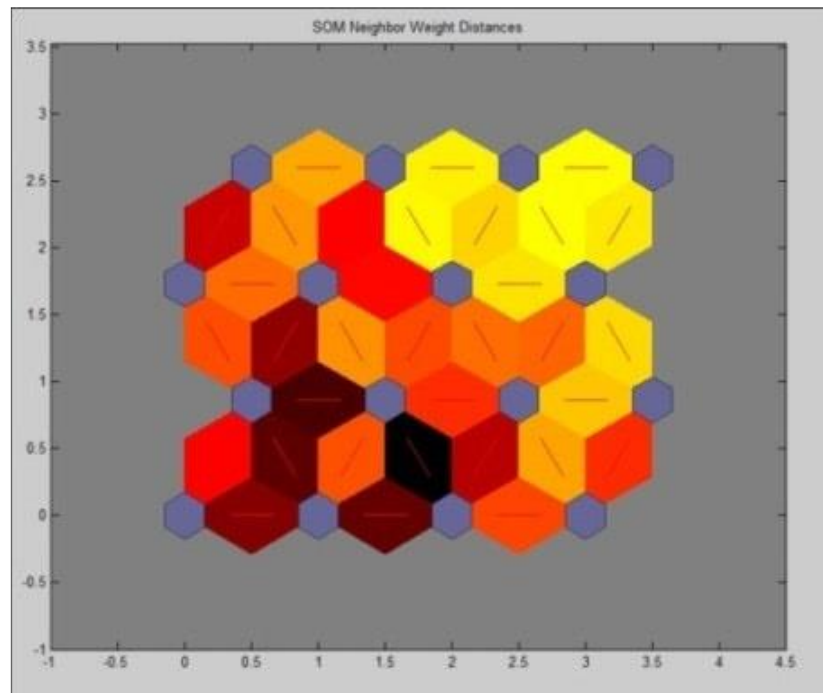
۴-۴. خوشه بندی SOM بر اطلاعات RFM

در این مطالعه در نخست با بکارگیری الگوریتم SOM و شبکه ای در ابعاد چهار در چهار و نورون‌های شش ضلعی تولید شد. هریک از این سلول‌های عصبی از راه وزن‌های سیناپسی که در طول یادگیری به بردار ورودی متصل است، تنظیم شدند. فاز اول برآورد ناهموار است که برای تولید الگوهای داده‌های ناخالص استفاده شد. فاز دوم فاز تنظیم بود که برای تنظیم نقشه شبکه به مدل ویژگی‌های خوب از داده مورد استفاده قرار گرفت (۱۵).

انجام الگوریتم خوشه بندی به این علت است که بتوان گروه‌هایی از مشتریان را برای رتبه بندی شناسایی نمود. نتایج اجرای الگوریتم SOM در نرم افزار متلب در شکل ۲ و ۳ قابل مشاهده است. در شکل ۲ تعداد مشتریانی را که در هر نورون تقسیم شده اند نمایش میدهد. شکل ۳ نشان دهنده میزان فاصله بین نورون‌های همسایه از یکدیگر است که هرچقدر میزان فاصله بین نورون‌های همسایه از یکدیگر بیشتر باشد، با رنگ تیره تر نشان داده می‌شود و هرچقدر فاصله کمتر باشد با رنگ‌های روشن تر نمایش داده می‌شود.



شکل ۲. تعداد اعضای نوروں ها



شکل ۳. میزان فاصله نوروں ها

در جدول ۳ نتایج محاسبات SSE بر اساس تعداد خوشه های مختلف را مشاهده می کنید. در نتیجه می توان گفت که ۸ خوشه اساسی در این شبکه می توان یافت.

جدول ۳. نرخ پارامتر SSE به ازای تعداد خوشه در SOM

تعداد خوشه	SSE
۸	۱۴/۶۲
۹	۱۸/۶۱
۱۱	۱۵/۷۶
۱۰	۱۹/۵۰

بر اساس تعداد خوشه تعیین شده، مشتریان بر اساس شاخص‌های آر. اف. ام. خوشه بندی شدند که نتایج این خوشه بندی در جدول ۴ آورده شده است.

جدول ۴. نتایج خوشه بندی مشتریان

خوشه	تعداد مشتریان
۱	۱
۲	۱
۳	۲
۴	۸
۵	۲
۶	۱۴۰
۷	۲
۸	۴
جمع	۱۶۰

۴-۵. تحلیل واریانس معیارهای RFM

بدلیل آنکه بدانیم کدام فاکتور تاثیر بیشتری بر جداسازی خوشه‌ها داشته است باید از تحلیل واریانس استفاده می‌کردیم. با توجه به جدول تحلیل واریانس جدول ۵ مشخص می‌گردد که شاخص ارزش پولی خریدهای مشتری (M') با بزرگترین مقدار آماره F یعنی ۸۷۷/۲۵ در سطح خطای کمتر از ۱ درصد بیشترین نقش را در جداسازی خوشه‌ها از یکدیگر داشته است. همچنین شاخص تازگی خرید (R') با کوچکترین مقدار F یعنی ۲/۲۷ در سطح خطای کمتر از ۰/۰۲۹ کمترین نقش را در جداسازی خوشه‌ها داشت.

جدول ۵. تحلیل واریانس خوشه ها

Sig.	آماره F	خوشه ها		شاخص ها
		درجه آزادی	میانگین مربعات	
۰/۰۲۹	۲/۲۷	۷	۰/۰۰۱	تازگی خرید
۰	۲۸۱/۸۱	۷	۰/۰۵	شاخص فراوانی خرید
۰	۸۷۷/۲۵	۷	۰/۱۳۹	ارزش پولی خریدهای مشتری (M')

۴-۶. طبقه بندی مشتریان

در ادامه، متوسط ارزش هر یک از شاخص های تازگی خرید و ارزش مالی خریدهای مشتری، شاخص فراوانی خرید در هر خوشه با تقسیم مجموع ارزش شاخص ها در هر خوشه به تعداد مشتریان آن خوشه تعیین گردید که در جدول ۶ آورده شده است.

جدول ۶. متوسط ارزش شاخص های RMF در خوشه ها

خوشه	متوسط ارزش تازگی خرید	متوسط ارزش مالی خرید مشتری	متوسط ارزش فراوانی خرید
۱	۰/۰۷۱	۰/۰۷۳۵	۰/۵۹۹
۲	۰/۰۷۱۲	۰/۱۳۵۶	۰/۳۷۸۳
۳	۰/۰۷۱	۰/۲۶۵۳	۰/۳۲۰۴
۴	۰/۰۶۴۳	۰/۱۰۶	۰/۰۲۵۴
۵	۰/۰۷۱۵	۰/۰۸۷۹	۰/۲۵۵
۶	۰/۰۵۲	۰/۰۰۸۵	۰/۰۰۶۳
۷	۰/۰۷۱	۰/۲۳۶	۰/۱۵۱۳
۸	۰/۰۷۲	۰/۰۹۹	۰/۱۲۱۶

با توجه به خروجیهای مراحل قبل، مقایسه متوسط ارزش شاخصهای آ.اف.ام. در هر خوشه با متوسط ارزش این شاخصها در کل داده ها و همچنین رتبه بندی خوشه ها با توجه به مقدار متوسط ارزش شاخص های آ.اف.ام.، تحلیل خوشه ها در قالب جدول ۷ صورت گرفت.

جدول ۷. تحلیل خوشه ها

تحلیل	خوشه ها						
	۱	۲	۳	۴	۵	۶	۷
وضعیت متوسط ارزش شاخص ها	↑	↑	↑	↑	↑	↓	↑
رتبه متوسط ارزش شاخص MM	۱	۲	۳	۷	۴	۸	۵
رتبه متوسط ارزش شاخص MF	۷	۳	۱	۲	۶	۸	۵
رتبه متوسط ارزش شاخص MR	۱	۵	۱	۵	۵	۸	۱

در این جدول، ابتدا در قسمت وضعیت شاخص ها، به مقایسه متوسط ارزش هر یک از شاخص های RFM. در هر خوشه با متوسط ارزش این شاخصها در کل داده ها پرداخته شده است که این مقایسه مشخص می نماید که متوسط ارزش هر یک از شاخصهای آر.اف.ام. در هر خوشه نسبت به متوسط ارزش این شاخصها در کل داده ها در چه وضعیتی قرار دارد. برای هر یک از شاخصها، در وضعیتی که متوسط ارزش شاخص در یک خوشه بیشتر از متوسط ارزش آن شاخص در کل داده ها باشد، این وضعیت با علامت (↑) (وضعیت مطلوب) و در صورتی که متوسط ارزش شاخصی در یک خوشه کمتر از متوسط ارزش آن در کل داده ها باشد، این (وضعیت با علامت (↓) نشان داده است (وضعیت نامطلوب) در قسمت بعدی جدول ۷ جهت بررسی دقیق تر خوشه ها، رتبه متوسط ارزش شاخصها در هر خوشه در مقایسه با خوشه های دیگر آورده شده است. برای مثال رتبه مربوطه در خوشه ۱ و سطر تازگی خرید مشخص میکند که این خوشه از نظر مقدار متوسط ارزش شاخص های تازگی خرید در رتبه X قرار دارد. پس از انجام تحلیل خوشه ای، ریسک اعتباری مشتری در هر خوشه بر اساس رابطه زیر محاسبه گردید و در نهایت نیز مشتریان بر اساس ارزش اعتباری در خوشه های ایجاد شده در قالب ارزش اعتبار مشتری جدول ۸ نشان داده شده است، بخش بندی شدند.

$$RV = M_{R''} + M_{F''} + M_{M''}$$

جدول ۸. ارزش رتبه بندی شده مشتریان

رتبه بندی	ارزش اعتباری RV	تعداد	خوشه
۱	۰/۷۹۶	۱	۱
۳	۰/۵۹۷	۱	۲
۲	۰/۶۷۸	۲	۳
۷	۰/۱۹۹	۸	۴
۵	۰/۴۳۶	۲	۵
۸	۰/۱۶۵	۱۴۰	۶
۴	۰/۴۶۰	۲	۷
۶	۰/۲۹۲	۴	۸

برای سنجش کارایی مدل ارائه شده صحت نتایج این مدل با الگوریتم های طبقه بندی دیگری چون شبکه عصبی، درخت تصمیم CHAID و رگرسیون لجستیک مورد مقایسه قرار گرفت که نتایج در جدول ۹ قابل مشاهده است.

جدول ۹. مقایسه تطبیقی دقت الگوریتم ها

الگوریتم	RFM پیشنهادی	شبکه عصبی	CHAID	رگرسیون لجستیک
دقت	۷۴%/۹۴	۷۴%/۳۳	۷۳%/۹۹	۶۸%/۳۹

همان طور که در جدول ۹ مشاهده می شود، آر.اف.ام پیشنهادی به میزان ۰/۶۱ درصد دارای دقت بالاتری نسبت به بقیه روش ها می باشد. این الگوریتم از پیچیدگی کمتری برخوردار است که بدلیل گستردگی آن، نیازی به فرم های درختی زیاد نیست و این یک مزیت بشمار می رود. الگوریتم استخراج شده قابلیت طبقه بندی رکوردهای جدید مشتریان را داشته و می توان به میزان صحت آن در تشخیص وضعیت رفتار مشتریان جدید اطمینان کرد. هنگامی که اطلاعات مشتری جدید وارد سیستم می شود، آن را از ابتدای خوشه بندی بر اساس ویژگی های بیوگرافیک هدایت می کنیم تا زمانی که به یک گره برگ برسیم، در آنجا رده مربوطه را به مشتری نسبت داده و می توانیم پیش بینی کنیم مشتری در کدام طبقه قرار می گیرد.

۴-۷. نتیجه گیری

در بازارهای رقابتی امروزی، با گرایش شرکتها به سمت مشتری مداری، مدیریت ارتباط با مشتری نیز به سمت پیچیدگی های خاصی گرایش پیدا کرده است. طبق مطالعات گذشته تخمین زده شده است که هزینه های جذب مشتریان جدید، پنج برابر هزینه های حفظ مشتریان موجود خواهد بود (کوتلر ۱۹۹۴). از سویی دیگر، تمرکز شرکتهای امروزی تنها بر فروش کالاهایشان نیست، آنها در پی خلق و حفظ مشتریان سودآور هستند. در این مبحث پر پیچ و خم طبقه بندی مشتریان اهمیت بسزایی دارد. در این فصل داده ها بر اساس روش تحقیق

پیشنهادی RFM و SOM تحلیل شدند و خوشه بندی مشتریان و در نهایت رتبه بندی آنها بر اساس اعتبار مشتریان انجام شد. این تحلیل بر روی ۱۶۰ مشتری صنعت لیزینگ مورد مطالعه انجام شد و نتایج بسیار ارزنده ای را در بر داشت. در آخر نیز مدل پیشنهادی با دیگر روش های داده کاوی مقایسه شد که مشخص شد دقت این روش پیشنهادی به اندازه ۰/۶۱ درصد بیش از سایر روش ها بوده است.

۲-۵. چگونگی پردازش اهداف پژوهش

هدف اصلی هدایت پژوهش، ارائه روشی جامع برای طبقه بندی مشتریان بر اساس اعتبار آنها در چارچوب کشورهای در حال توسعه بود. هدف اول پژوهش، بررسی و تجزیه و تحلیل مسائل و چالش های درگیر در طبقه بندی مشتریان شرکت ها در کشورهای در حال توسعه است. دو روش پژوهش در دستیابی به این هدف مورد استفاده قرار گرفت - بررسی ادبیات سنتی و گذشته و تعامل با کارشناسان عرصه مدیریت ارتباط با مشتری، دو روش پژوهشی هستند که در دستیابی به این هدف مورد استفاده قرار گرفت. به دنبال بحث ارائه شده در ادامه بررسی های صورت گرفته در زمینه طبقه بندی مشتریان مشخص کرد که انواع پژوهش های صورت گرفته طبق مدل های کیفی و کمی هستند. این تحقیق در ردیف مدل های کیفی قرار می گیرد. مدل های کیفی، مدل های توصیفی هستند که بر پایه قوانین ریاضیاتی و فرمول ها نیستند. این مدل ها سیستم های موجود را با توجه به روابط اتفافی، پایه ای، ترکیبی و یا تنظیم شده میان موضوعات (اشیا) و وقایع به تصویر می کشند (کلانسی، ۱۹۸۹). این تحقیقات از حوزه مدیریت ارتباط با مشتری آغاز شده و در نهایت خوشه بندی مشتریان بر اساس اعتبار آنها ارائه شد. با این حال، رتبه بندی مشتریان در کشورهای در حال توسعه بعنوان یک مسئله پیچیده در مسائل اقتصادی و مدیریتی در نظر گرفته می شود. علاوه بر این، طبقه بندی مشتریان در زیرساخت های اعتباری خواستار تحقیق و تفحص از تلفیق عوامل فنی و غیر فنی مؤثر بر خوشه بندی گروه ها به عنوان یک هدف برای اولین بار راه اندازی شد. تحقیقات نشان می دهد که اکثر پژوهش ها در مدل های کیفی مربوط به استفاده از مدل های تصمیم گیری و روش های ژنتیک و شبکه های عصبی در جهت طبقه بندی مشتریان و استراتژیک راهبردی بوده است.

بر اساس مطالعات صورت گرفته هیچ روش طبقه بندی جامعی برای ارزیابی اعتبار مشتریان وجود ندارد که از نظر تحلیلی بتواند به ماهیت پیچیده سیستم ارزیابی اعتبار مشتریان صنعت لیزینگ رسیدگی کند. تحقیقات اولیه در این مطالعه نشان داد که روش های داده کاوی نقش هدایت کلیدی را در توسعه مدل های خوشه بندی مشتریان ایفا می کند. این نتیجه گیری در نهایت منجر به نیاز یک رویکرد چند مرحله برای طبقه بندی مشتریان، در صنعت مورد مطالعه شد. بنابراین لازم بود که روش ها اکتشافی مبنا قرار گیرند تا برای مقابله با مسائل پیچیده مرتبط با ارزیابی و طبقه بندی مشتریان مناسب باشد.

هدف دوم بررسی و تجزیه و تحلیل مناسب و همچنین تلفیق روش های داده کاوی، یعنی RFM و شبکه عصبی، که می تواند برای طبقه بندی مشتریان بر اساس ریسک اعتباری استفاده شود. تمرکز بر تحقیق و تجزیه و تحلیل روش های مناسب و تکنیک های زمینه ای داده کاوی، می تواند برای بهبود روند خوشه بندی جاری مشتریان مورد استفاده قرار گیرد. به منظور دستیابی به این هدف، لازم بود که برای انجام این پژوهش، روش های داده کاوی بازنگری شود. در واقع تا جایی که امکان داشت تکنیک های اتخاذ شده مناسب برای طبقه بندی مشتریان اعتباری در این تحقیق، استفاده شد.

هدف سوم تدوین و فرموله کردن روش‌های داده‌کاوی برای طبقه‌بندی مشتریان بالقوه بر اساس ریسک اعتباری در چارچوب مباحث نظری است. تکنیک‌های شرح داده شده در این تحقیق از طریق ماهیت تحلیلی خود با ارائه درک و بینش عمیق‌تری نسبت به مسائل مرتبط با طبقه‌بندی مشتریان شرکت داشته‌اند. در درجه اول، مدل کلی RFM تدوین شد که در آن کارشناسان مدیریت ارتباط با مشتری دعوت شده بودند تا قضاوت‌های خود را ارائه دهند. در این پژوهش از داده‌های تراکنشی مشتریان مربوط به صنعت لیزینگ استفاده شد. در مجموع تعداد ۱۲۳۲ مشتری در فاصله زمانی ۱۴۰۳-۱۴۰۲ جمع‌آوری شد. با توجه به محدودیت‌هایی که در تحویل داده‌های دموگرافیک وجود داشت، تنها داده‌های تراکنشی در اختیار قرار گرفتند. مرحله آماده‌سازی از مهمترین و پیچیده‌ترین مراحل داده‌کاوی است. مراحل بکار رفته شامل فرآیند پاکسازی و کاهش بعد می‌باشد. در نهایت تعداد ۱۶۰ مشتری برای انجام عملیات اعتبار سنجی باقی ماندند. با مدل RFM داده‌های آماده برای استفاده در عملیات اعتبار سنجی شدند. در مرحله بعد، در مرحله تحلیل، مدل مذکور با استفاده از روش AHP، وزن نهایی هر یک از گزینه‌ها تعیین شد. با این حال، سادگی ساختار سلسله‌مراتبی و روابط یک سو به خطی در میان معیارها و زیر معیارها در روش AHP و پنهان کردن مسائل مهمی، مانند وابستگی متقابل میان عوامل کیفی و تعامل میان سطوح تصمیم‌گیری و پردازش بیش از حد آسان در این مدل سبب پیدایش مشکلاتی در مسأله می‌شود. در مرحله بعد روش SOM برای خوشه‌بندی استفاده شد. آزمون اعتبار سنجی مدل فوق را تایید نمود. بنابراین می‌توان آن را به عنوان مدل عمومی و مرجع در سایر تکنیک‌های خوشه‌بندی مشتریان به کار برد.

هدف چهارم، مقایسه روش‌های داده‌کاوی ترکیبی با رگرسیون لجستیک بود. بر این اساس مدل پیشنهادی با دیگر روش‌های داده‌کاوی مورد مقایسه قرار گرفت. نتایج نشان داد که آ.ف.ام پیشنهادی به میزان ۶۱٪ درصد دارای دقت بالاتری نسبت به بقیه روش‌ها می‌باشد. این الگوریتم از پیچیدگی کمتری برخوردار است که بدلیل گستردگی آن، نیازی به فرم‌های درختی زیاد نیست و این یک مزیت بشمار می‌رود. الگوریتم استخراج شده قابلیت طبقه‌بندی رکوردهای جدید مشتریان را داشته و می‌توان به میزان صحت آن در تشخیص وضعیت رفتار مشتریان جدید اطمینان کرد. هنگامی که اطلاعات مشتری جدید وارد سیستم می‌شود، آن را از ابتدای خوشه‌بندی بر اساس ویژگی‌های بیوگرافیک هدایت می‌کنیم تا زمانی که به یک گره برگ برسیم، در آنجا رده مربوطه را به مشتری نسبت داده و می‌توانیم پیش‌بینی کنیم مشتری در کدام طبقه قرار می‌گیرد.

۳-۵. پردازش سوالات تحقیق

سوال اول: کدام الگوریتم خوشه‌بندی صحیح‌تری را نسبت به سایر الگوریتم‌های خوشه‌بندی داده‌کاوی ارائه می‌کند؟

در این تحقیق با ارائه یک مطالعه موردی اثبات شد که روش پیشنهادی RFM و شبکه مصنوعی SOM بهترین الگوریتم برای اجرای خوشه‌بندی مشتریان صنعت لیزینگ بود.

سوال دوم: کدام شاخص مهمترین شاخص در خوشه‌بندی مشتریان می‌باشد؟ پاسخ: بر اساس روش AHP استفاده شده در این پژوهش مشخص شد که وزن شاخص تازگی خرید ۶۹٪، ارزش مالی مشتری ۹٪ و فراوانی خرید ۲۲٪ از مجموع عوامل RFM را دارا هستند. بنابراین می‌توان گفت که شاخص تازگی خرید با ۶۹٪ مهمترین شاخص در خوشه‌بندی مشتریان بود.

سوال سوم: کدام مدل الگوریتم داده کاوی نسبت به سایر الگوریتم‌ها از صحت بالاتری برخوردار می‌باشد؟ پاسخ: بر اساس مقایسه تطبیقی مشخص شد که مدل پیشنهادی RFM و SOM از صحت بالاتری برخوردار هستند. طوریکه آر.اف.ام پیشنهادی به میزان ۰/۶۱ درصد دارای دقت بالاتری نسبت به بقیه روش‌ها می‌باشد. این الگوریتم از پیچیدگی کمتری برخوردار است که بدلیل گستردگی آن، نیازی به فرم‌های درختی زیاد نیست و این یک مزیت بشمار می‌رود.

سوال چهارم: آیا بهینه‌سازی خوشه‌بندی بر مبنای الگوریتم‌های هوشمند (ژنتیک، ازدحام ذرات، جستجوی فاخته، رقابت استعماری و...) از صحت بالاتری برخوردار نیست؟ پاسخ: بنظر می‌رسد در پیرو سوال قبل روش‌های ذکر شده نتوانند صحت بالاتری نسبت به روش پیشنهادی داشته باشند چراکه نتیجه مقایسات تطبیقی این موضوع را نشان داده است.

۴. پیشنهادهای کاربردی

با توجه به اینکه مشتریان تمامی خوشه‌ها به غیر از خوشه ششم، از نظر شاخص‌های آر.اف.ام. در وضعیت مطلوبی قرار دارند، لذا پیشنهاد می‌گردد که به منظور حفظ این مشتریان، شرکت با برقراری ارتباطات و تعاملات بیشتر با آنها، سعی در تبدیل وفاداری رفتاری این مشتریان به وفاداری نگرشی نماید. اگرچه مشتریان خوشه‌های ۳ و ۷، ارزش اعتباری بالایی ندارند، اما فراوانی خرید آنها در مقایسه با خوشه‌های دیگر، بسیار بالا است. شرکت می‌تواند با در نظر گرفتن تخفیف‌های حجمی خاص برای مشتریان این خوشه‌ها، بر حجم خرید این مشتریان افزوده و از این طریق، تعداد دفعات خرید آنها و متعاقب آن، قسمتی از هزینه‌های تحمیل شده را کاهش دهد. با توجه به اینکه خوشه ششم در پایین‌ترین سطح اعتباری قرار دارند اما احتمال ارتقای این مشتریان به سطوح بالای هرم ارزش مشتری، بسیار بالا می‌باشد و شرکت می‌تواند با در نظر گرفتن مزایایی خاص برای خرید آنها، وفاداری رفتاری را در بین آنها ترویج دهد.

پیشنهاد می‌گردد که در مطالعات آتی، از مجموعه داده وسیع‌تری به لحاظ قلمروی زمانی استفاده گردد که به طور قطع منتج به ایجاد نتایج قویتر و گستره وسیع‌تری از دانش کاربردی پیرامون ویژگیهای رفتاری مشتریان خواهد شد. از طرفی، برشهای زمانی متناوب از پایگاه داده شرکت، امکان اجرای پویای داده کاوی بر مبنای مدل آر.اف.ام. جهت تعیین ارزش اعتبار مشتریان را فراهم خواهد آورد که نتایج آن از این لحاظ که روند تغییرات در رفتار مشتریان را منعکس میکند، میتواند نقش شایانی در اقدامات بازاریابی و بهبود مدیریت ارتباط با مشتری ایفا نماید. همچنین استفاده از روش‌های تلفیقی دیگر همانند تکنیک‌های تصمیم‌گیری چند معیاره مانند ANP میتواند کمک شایانی برای بهبود این تحقیق باشد. لزوم فازی بودن امتیازات در RFM میتواند مسیری برای توسعه مدل مذکور در شرایط عدم قطعیت باشد.

در این تحقیق، جهت بخش‌بندی صنعت لیزینگ بر اساس ارزش اعتباری مشتری از مدل آر.اف.ام. و برخی تکنیک‌های داده کاوی و روش‌های مختلف علمی در قالب فرایندی خاص، بهره گرفته شد. استفاده از جدول تحلیل واریانس جهت تعیین میزان اثرگذاری هر یک از شاخص‌ها در تفکیک خوشه‌ها، ارائه تحلیل درون خوشه‌ای به منظور بررسی دقیق‌تر ویژگی‌های مشتریان در خوشه دارای بیشترین مشتری و همچنین بخش‌بندی نهایی مشتریان بر اساس ارزش اعتباری مشتریان در قالب جدول ارزش اعتبار مشتریان از موارد خاصی است که در این



مطالعه از آنها بهره گرفته شده است، درحالی که در مطالعات گذشته به آنها اشاره ای نگردیده است. روش پیشنهادی با ارائه یک مطالعه موردی نشان داد که از پایایی بالایی برخوردار است. همچنین که در مقایسه تطبیقی با دیگر روش های درختی و خوشه بندی از صحت و دقت بیشتری برخوردار بود. استفاده از سیستم RFM هنوز هم در فرآیندهای طبقه بندی بصورت قاطع کاربرد دارد و می تواند با دیگر روشها بصورت ترکیبی مدلی قویتر را ارائه نماید.

۵. فهرست منابع

- حافظنیا، م، ر(۱۳۸۲). مقدمه‌ای بر روش تحقیق در علوم انسانی، انتشارات سمت، تهران، چاپ هشتم
- زارعی نژاد، محسن، ۱۳۹۲. ارزیابی و انتخاب تأمین‌کننده ثالث لجستیک معکوس با رویکرد تلفیقی ANP و IFG-MCDM. پایان نامه کارشناسی ارشد. دانشگاه آزاد واحد تهران جنوب.
- خاکی، غلامرضا " روش تحقیق با رویکرد پایان نامه نویسی " (تهران ، انتشارات بازتاب ۱۳۸۲)
- Hughes, Arthur M. (2023), Strategic database marketing, Chicago: Probus publishing.
- Coussement, K., Van den Bossche, F. A., & De Bock, K. W. (2024). Data accuracy's impact on segmentation performance: Benchmarking RFM analysis, logistic regression, and decision trees. *Journal of Business Research*, 67(1), 2751-2758.
- Saaty, T. L. (2022). The analytic hierarchy process: planning, priority setting, resources allocation. *New York: McGraw*.
- Kotler, P. (2023), Marketing management: Analysis, planning, implementation, and control, New Jersey: Prentice-Hall.
- Yeh, C., Yang, K. & Ting, T.; "Knowledge Discovery on RFM Model Using Bernoulli Sequence", *Expert Systems with Applications*, Vol. 36, pp. 5866–5871, 2019.
- Lumsden SA, Beldona S, Morison AM. Customer value in an all-inclusive travel vacation club: An application of the RFM framework. *J. Hosp. Leisure Mark.*, 16(3): 2, pp. 70-285, 2024
- N. C. Hsieh, "An integrated data mining and behavioral scoring model for analyzing bank customers", *Expert Systems with Applications*, 27, pp. 623–633, 2024
- S. Y. Hussein, & et al., "Segmenting and Profiling Green Consumers with Use of Self Organizing Maps", *Journal of Management Research in Iran*, Volume 17, Number 2, 2023.
- Tsai, Lu. , "Customer churn prediction by hybrid neural networks", *Expert Systems with Applications*, Vol. 36, pp. 12547-12553, 2019.
- Ghazanfari, M. et al, "Customer segmentation export edible fruits", *Quarterly Journal of Commerce*, No. 55, 151 – 181, 2022

predicting customers ' risk in the leasing industry using data mining techniques to clustering

Authors⁷

Abdolhossien karampour

Abdolhossien ziyarati

Sara sohrabi motlagh fazeli

Abstract

The main objective of this research is to present a comprehensive method for classifying customers based on their validity in developing countries. In this research, for classifying the leasing industry based on customer's credit worthiness RFM. and some data mining techniques and different scientific methods have been used in the form of a specific process. the use of variance analysis table to determine the effect of each indicator in cluster segregation, presenting in-cluster analysis in order to examine the more accurate characteristics of customers in cluster with customer loyalty and customer's final segmentation based on customer credit worthiness in the form of customer credit worthiness is one of the special cases that have been used in this study. the proposed method by presenting a case study showed that it has a high frequency. also, in comparison with other tree methods and clustering, it was more accurate. the use of RFM is still widely used in classification processes and can be combined with other methods in a hybrid model.

Keywords: predicting customers ' risk,
leasing industry, data mining techniques, clustering.

JEL Classification O16 Financial Markets • Saving and Capital Investment • Corporate Finance and Governance

⁷assistant professor and member of the faculty of department of business administration of islamic azad university, north Tehran branch.