

تعیین عوامل موثر در هزینه و زمان حمل و نقل جاده‌ای و ریلی با قواعد انجمنی

مقاله علمی - پژوهشی

*مهدی یوسفی نژاد عطاری (نویسنده مسئول)، گروه مهندسی صنایع، واحد بناب، دانشگاه آزاد اسلامی، بناب، ایران

انسبه نیشابوری جامی، گروه مهندسی صنایع، دانشگاه بناب، بناب، ایران

*پست الکترونیکی نویسنده مسئول: 1689597631@iaui.ir

دریافت: ۱۴۰۵/۰۳/۲۱ - پذیرش: ۱۴۰۵/۰۵/۲۵

صفحه ۵۳۸-۵۱۳

چکیده

حمل و نقل به جابه جایی انسان، حیوان یا کالا از یک مکان به مکان دیگر گفته می‌شود. به بیان دیگر، حمل و نقل به عنوان حرکتی مشخص از نقطه A به نقطه B تعریف می‌شود. از انواع روش‌های حمل و نقل می‌توان به حمل و نقل هوایی، زمینی (ریلی و جاده‌ای)، دریایی، کابلی، خط لوله و فضایی اشاره کرد. در حمل و نقل زمینی، به دلیل اهمیت این کسب و کار و رقابت میان سازمان‌های فعال در این حوزه، به کارگیری فناوری‌ها و تکنولوژی‌های نوین در مدیریت و اتخاذ تصمیمات بهتر می‌تواند بسیار سودمند باشد. در این مقاله، نخست به معرفی داده‌کاوی، تکنیک‌های آن و نرم‌افزارهای مرتبط با داده‌کاوی می‌پردازیم. سپس، با استفاده از جدولی از حمل‌های صادراتی انجام شده در بازه زمانی یک سال و شش ماه گذشته که از یک شرکت حمل و نقل بین‌المللی تهیه شده است، قوانین انجمنی را با استفاده از نرم‌افزار کلمنتاین در دو مدل استخراج کرده و قوانین به دست آمده را مورد ارزیابی قرار داده شود.

واژه‌های کلیدی: داده‌کاوی، زمان، حمل و نقل جاده‌ای و ریلی، خوشه‌بندی، هزینه مقدمه

۱-مقدمه

موجود، جهت کشف دانش‌های سازمانی و یا پیش بینی وضعیت آینده است. حمل و نقل نقش مهمی را در رشد اقتصادی و جهانی شدن ایفا می‌کند.

در دنیای امروز داده‌ها و اطلاعات بعنوان ثروت سازمانی محسوب گشته و همواره شرکت‌ها و سازمان‌های بزرگ و موفق دنیا به دنبال استفاده مناسب‌تر و تجاری‌تر از این منابع مجازی می‌باشند. از سوی دیگر با پیچیده شدن محیط‌های کسب و کار، ماهیت و حجم داده‌های سازمانی بسیار متفاوت گشته و نگاه یکپارچه و مدیریتی به آنها ضروری می‌گردد. یکی از راه‌حلی که سازمان‌های موفق دنیا در این زمینه اتخاذ می‌نمایند، ایجاد یک سیستم جامع داده‌ای و آماری و بکارگیری صحیح تکنیک‌های داده‌کاوی و کشف دانش می‌باشد (اولسون، ۲۰۰۳). یکی از راه‌حلی که می‌تواند تصمیم‌گیری در چنین فعالیت‌هایی را تسریع بخشیده و دقت تصمیم‌گیری را افزایش دهد، بکارگیری صحیح تکنیک‌های داده‌کاوی بر روی داده‌های

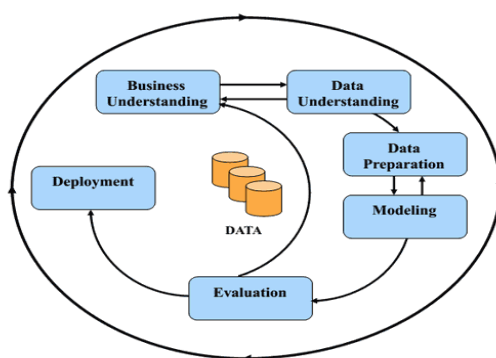
۲-پیشینه تحقیق

تا کنون تعاریف متعددی برای داده‌کاوی ارائه شده است. در برخی از این تعاریف داده‌کاوی در حد ابزار نیست که کاربران را قادر به ارتباط مستقیم با حجم عظیم داده‌ها می‌سازد و در برخی تعاریف دقیق‌تر به کاوش در داده‌ها نیز توجه شده است. در ادامه به برخی از این تعاریف اشاره می‌کنیم: اصطلاح داده‌کاوی به فرایند نیم‌خودکار تجزیه و تحلیل پایگاه داده‌های

۲-۱- فرآیند داده کاوی

بر اساس فرایند استاندارد داده کاوی کریسپ می توان پنج فاز را برای پروژه های داده کاوی در نظر گرفت که عبارتند از: شناخت کسب و کار، شناخت داده، آماده سازی داده ها، مدلسازی، ارزیابی و بکارگیری (جانگ، ۲۰۱۶).

این فرایند استاندارد که شمای آن در شکل ۱ موجود است فرایند بصورت خطی حرکت نمی کند بلکه در هر مرحله امکان بازگشت به عقب وجود دارد و این روند آنقدر تکرار می شود تا در نهایت به بهترین نتیجه برسیم.



شکل ۱. مدل استاندارد غیر خطی داده کاوی (جانگ، ۲۰۱۶)

می شود. با تکرار همین روال می توان در هر تکرار با میانگین گیری از داده ها مراکز جدیدی برای آن ها محاسبه کرد و مجدداً داده ها را به خوشه های جدید نسبت داد. این روند تا زمانی ادامه پیدا می کند که دیگر تغییری در داده ها حاصل نشود (کریمیک و سوبای، ۲۰۱۵).

۲-۳- ارزیابی خوشه بندی (شاخص دیویس بولدین)

برای ارزیابی الگوریتم های خوشه بندی معیارهای متفاوتی وجود دارد که می توانیم آن ها را به دو دسته بدون ناظر و با ناظر تقسیم کنیم. معیارهای ارزیابی بدون ناظر که گاهی در متون علمی از آن با نام معیارهای داخلی نیز یاد می شود آن دسته عملیاتی هستند که با توجه به اطلاعات موجود در مجموعه داده به کار می روند. مهم ترین وظیفه یک الگوریتم خوشه بندی این است که بتواند به بهترین شکل فاصله درون خوشه ای را کمینه و فاصله بین خوشه ها را بیشینه کند. از این رو دو فاکتور مهم که در تمامی معیارهای ارزیابی به کار می رود: تراکم خوشه ای و جدایی خوشه ای است؛ و برآورده شدن اهداف کمینه سازی

بزرگ به منظور یافتن الگوهای مفید اطلاق می شود (هان و کمبار، ۲۰۱۱). به طور خلاصه می توان داده کاوی را جستجو در مجموعه ای از داده ها برای یافتن الگوهایی در بین آن ها نامید و یا داده کاوی استخراج دانش جدید و قابل استناد از مجموعه پایگاه داده های بزرگ نیز معرفی شده است. (دیوید و هاند، ۲۰۰۱). از منظر داده کاوی یعنی تجزیه و تحلیل مجموعه داده های قابل مشاهده برای یافتن روابط مطمئن بین داده ها تعریف می شود (هربرت، ۱۹۹۶).

۲-۲- الگوریتم کا- میانگین

خوشه بندی یکی از شاخه های یادگیری بدون نظارت است و فرآیند خودکاری است که در طی آن، نمونه ها به دسته هایی که اعضای آن مشابه یکدیگر می باشند تقسیم می شوند که به این دسته ها خوشه گفته می شود. بنابراین خوشه مجموعه ای از اشیاء است که در آن اشیاء با یکدیگر مشابه بوده و با اشیاء موجود در خوشه های دیگر غیر مشابه می باشند. برای مشابه بودن می توان معیارهای مختلفی را در نظر گرفت مثلاً می توان معیار فاصله را برای خوشه بندی مورد استفاده قرارداد و اشیائی را که به یکدیگر نزدیک تر هستند را به عنوان یک خوشه در نظر گرفت که به این نوع خوشه بندی، خوشه بندی مبتنی بر فاصله نیز گفته می شود. یکی از این روش ها خوشه بندی k-means است (هان و کمبار، ۲۰۱۱).

در نوع ساده ای از این روش ابتدا به تعداد خوشه های مورد نیاز نقاطی به صورت تصادفی انتخاب می شود. سپس در داده ها با توجه به میزان نزدیکی (شباهت) به یکی از این خوشه ها نسبت داده می شوند و بدین ترتیب خوشه های جدیدی حاصل

خوشه است.

فاکتور تراکم یا Coho به صورت $\frac{1}{m} \sum_{x \in C_i} dis(x, C_i)$ تعریف می‌شود. که در آن x نقطه داخل خوشه و c مرکز همان خوشه است.

فرمول کلی دیویس بولدین به صورت زیر است: مقدار NC برای تمام نقاط است.

$$\frac{1}{nc} \sum_{i=1}^{nc} \left\{ \frac{coh(i) + coh(j)}{Sep(i, j)} \right\}$$

فاصله درون خوشه‌ای و بیشینه سازی فاصله میان خوشه‌ای به ترتیب در گروه بیشینه نمودن تراکم هر خوشه و نیز بیشینه سازی جدایی میان خوشه‌ها است.

یکی از این معیارها دیویس بولدین است که بر اساس این دو فاکتور مهم فرمول آن را به این صورت تعریف می‌شود.

فاکتور جدایی یا Sep در این معیار به صورت $dis(C_i, C_j)$ تعریف می‌شود که در آن dis فاصله اقلیدسی بین مرکز دو

(۱)

۲-۴- الگوریتم FP-Growth

برای تولید قوانین انجمنی الگوریتم‌های زیادی وجود دارند دو تا از پرکاربردترین روش‌ها الگوریتم Apriority و FP-Growth است. که در این تحقیق از روش دوم استفاده می‌کنیم. که در آن بدون تولید مجموعه اقلام کاندید و با کمک یک ساختمان داده‌های درختی الگوهای مکرر را شناسایی می‌کند. این الگوریتم که استراتژی تقسیم و غلبه را دنبال می‌کند، در ابتدا مجموعه داده را به صورت یک درخت موسوم به FP-tree تبدیل می‌کند. پس از آن به صورت مستقیم به استخراج الگوهای مکرر از این درخت می‌پردازد.

به طور خلاصه می‌توان گفت یک FP-tree نمایش فشرده‌ای از داده‌های مجموعه تراکنش است. برای ساخت این درخت الگوریتم، اقلام تراکنش‌ها را یک به یک خوانده و به صورت یک مسیر بر روی درخت نگاشت می‌کند. از آنجا که معمولاً تراکنش‌ها از اقلام مشترکی استفاده می‌کنند، لذا مسیرهای تراکنش‌ها بر روی درخت همپوشانی دارند و به همین دلیل درخت نسبت به داده‌های اولیه فشرده‌تر خواهد بود. در اولین اسکن پایگاه داده مجموعه آیتم‌های یک عضوی و پشتیبانی آن‌ها مشخص می‌شود و نیز مجموعه اقلام پرتکرار به ترتیب نزولی پشتیبانی‌شان مرتب می‌شوند. سپس یک درخت به این صورت ساخته می‌شود که: ابتدا ریشه درخت با پرچسب null ساخته می‌شود. بعد از آن مجموعه داده برای بار دوم اسکن می‌شود. اقلام هر تراکنش به ترتیب L پردازش می‌شوند و یک شاخه برای هر تراکنش ایجاد می‌شود. به منظور تسهیل پیمایش درخت، یک جدول ساخته می‌شود که هر قلم در آن به محل خودش در درخت اشاره می‌کند. درخت پس از اسکن همه تراکنش‌ها کامل می‌شود (کرمیک و سوای، ۲۰۱۵).

۲-۵- الگوریتم دسته بندی

تئوری بیز یکی از روش‌های آماری دسته‌بندی به شمار می‌آید. در این روش رده‌های مختلف، هر کدام به شکل یک فرضیه دارای احتمال در نظر گرفته می‌شوند. هر رکورد آموزشی جدید، احتمال در ست بودن فرضیه‌های پیشین را افزایش و یا کاهش می‌دهد و در نهایت، فرضیاتی که دارای بالاترین احتمال شوند، به عنوان یک کلاس در نظر گرفته شده و برچسبی بر آن‌ها زده می‌شود. این فن با ترکیب تئوری بیز و رابطه سببی بین داده‌ها، به طبقه‌بندی می‌پردازد (هان و کمبار، ۲۰۱۱).

جنگل تصادفی درخت تصمیم‌های زیادی را تولید می‌کند. برای طبقه‌بندی یک شیء جدید، برداری ورودی را در انتهای هر یک از درختان جنگل تصادفی قرار می‌دهد. هر درختی که طبقه‌بندی را ارائه می‌کند و اصطلاحاً گفته می‌شود این درخت به آن کلاس «رأی» می‌دهد. جنگل طبقه‌بندی را که بیشترین رأی داشته باشد (بین همه درخت‌های جنگل) انتخاب می‌کند. (لی و چن، ۲۰۱۴)

هر درخت به صورت زیر تشکیل می‌شود:

- اگر N تعداد حالت‌ها در مجموعه داده‌های آموزشی (مجموعه‌ی کار) باشد، N حالت را به صورت تصادفی با جایگذاری از داده‌های اصلی، نمونه‌گیری می‌کنیم. این نمونه مجموعه‌ی کار برای این درخت است.

- اگر M متغیر داشته باشیم و m را کوچک‌تر از M در نظر بگیریم به طوری که در هر گره، m متغیر به صورت تصادفی از M انتخاب می‌شوند و بهترین جداسازی روی این m متغیر برای جداسازی گره استفاده می‌شود. مقدار m در طول ساخت جنگل ثابت در نظر گرفته می‌شود.

- هر درخت به اندازه‌ی ممکن بزرگ می‌شود. هیچ هرسی وجود ندارد.

۲-۶- داده کاوی در حمل و نقل

در حوزه مهندسی حمل و نقل مقدار زیادی از داده‌ها در زمینه ترافیک، تحلیل حوادث و تصادفات، شرایط آسفالت، دوام جاده‌ها، مدیریت پل‌ها و غیره استخراج شده است. بر اساس این داده‌ها، تصمیم‌گیرندگان تصمیم می‌گیرند تا یک مشکل مربوط را حل کنند. تصمیم‌گیرندگان همیشه به دنبال روش‌هایی برای دسترسی راحت به داده‌ها و استفاده از داده‌های متفاوتی هستند. توانایی شناسایی داده‌های در دسترس، تعیین ویژگی‌های داده‌ها، استخراج داده‌های مورد نیاز، تبدیل داده‌ها به فرمت‌های لازم جزو الزامات لازم برای برنامه‌ریزی هستند. در دنیای واقعی حمل و نقل اطلاعات مختلف زیادی باید با هم ادغام شده تا به راه حل برسیم. تحقیقات اخیر در زمینه داده‌کاوی در حمل و نقل افق جدیدی را برای تصمیم‌گیرندگان مهندسان حمل و نقل ایجاد کرده است.

در سال ۲۰۱۴ لی و چن با ادغام سه روش داده‌کاوی، خوشه بندی کامپانگین، درخت‌های تصمیم‌گیری، و شبکه‌های عصبی، به پیش‌بینی زمان سفر آزادراه با ازدحام غیر مکرر پرداختند. (لی و چن، ۲۰۱۴).

احمدوند و همکاران با استفاده از الگوریتم اپریوری، قوانین انجمنی موجود در بازه‌های زمانی/مکانی رخداد مشکلات به دست آورده‌اند. همچنین با استفاده از یک تکنیک ۲ مرحله‌ای از این الگوریتم، مشکلاتی که با هم و به صورت همزمان در یک بازه زمانی مشخص اتفاق می‌افتند را شناسایی کرده و در نهایت نشان داده‌اند که از قوانین بدست آمده می‌توان برای شناسایی نیازهای شهروندان و بهبود مدیریت خدمات شهری استفاده نمود (احمدوند، ۱۳۸۸).

پیتر ژانگ، در تحقیق خود در زمینه پیش‌بینی سری‌های زمانی از مدل ترکیبی میانگین متحرک خودهمبسته یکپارچه و شبکه‌های عصبی استفاده کرد و به این نتیجه رسید که استفاده از مدل‌های سری زمانی به تنهایی از دقت و قدرت شبکه‌های عصبی برخوردار نیست. لذا ترکیب مدل‌های شبکه عصبی مصنوعی و ARIMA را به عنوان یک راهکار بهینه در پیش‌بینی سری‌های زمانی مطرح می‌کند و مدعی می‌شود که ترکیبی از مدل‌های شبکه عصبی مصنوعی به عنوان یک مدل

غیرخطی و مدل ARIMA به عنوان یک مدل خطی در پیش‌بینی سری‌های زمانی، عملکرد بهتری نسبت به پیش‌بینی صورت گرفته با استفاده از هر یک از مدل‌ها به تنهایی دارد. او برای اثبات این ادعا با استفاده از چند داده تجربی، ابتدا مجموعه داده را با استفاده از مدل ARIMA مدل‌سازی و خروجی را با خروجی مدل فرضیه خود مقایسه کرده و بر اساس آن، فرضیه خود را به اثبات رسانید (لاتیسا و همکاران، ۲۰۱۰).

لاتیسا، فرولیچ و کاپرا در مقاله‌ای تحت عنوان داده‌کاوی سیستم حمل و نقل عمومی و هوشمند نمودن سیستم‌های حمل و نقل، از ابزارها و تکنیک‌های داده‌کاوی استفاده نموده و مدت زمان سفر مسافران برای رسیدن به مقصد را کاهش داده و حتی الگوهای سفر مسافران را پیش‌بینی نمودند. مطالعه موردی این تحقیق با استفاده از داده‌های موجود در پایگاه داده سیستم حمل و نقل متروی شهر لندن صورت گرفته و مدت زمان سفر هر مسافر را تخمین زده و همچنین ایستگاه‌هایی که بیشتر توسط هر مسافر مورد استفاده قرار می‌گیرد را استخراج نموده است و با ارائه مدلی، مسیری که کمترین مدت زمان سفر را دارد به مسافر پیشنهاد می‌دهد. از اهداف دیگر این مقاله، تعیین ایستگاه‌هایی بود که بیشتر مورد استفاده قرار می‌گیرد زیرا این اطلاعات در جهت بهبود و تغییر برنامه‌های سفر آن ایستگاه موثر است (فیضا، ۲۰۱۸).

آدامز، آکانو و آسموتا در مقاله خود به پیش‌بینی میزان برق مورد نیاز برای تولید در کشور نیجریه پرداختند. هدف آن بود که با در اختیار داشتن میزان برق تولیدی در فاصله سال‌های ۱۹۷۰ تا ۲۰۰۹، میزان برق مورد نیاز تولیدی این کشور را در ۱۰ سال آینده پیش‌بینی نمایند. علت نیاز به این تحقیق آن بود که در سال ۲۰۰۵ به علت عدم تولید برق کافی و همچنین کاهش میزان بارندگی، تنها ۴۰ درصد برق مورد نیاز این کشور تامین گردید. لذا استفاده از مدلی که بتواند این پیش‌گویی را انجام دهد هدف این گروه قرار گرفت. این گروه در این تحقیق از مدل ARIMA برای پیش‌بینی استفاده کرده و دلیل آن را اینگونه بیان نموده‌اند که مدل ARIMA مدلی عمومی است و توانایی مدل‌سازی داده‌ها و سری‌های با دوره و بدون دوره فصلی را دارا می‌باشد و جهت تعیین درجه این مدل از نمودارهای ACF و PACF استفاده نمودند. این گروه پس از مدل‌سازی داده‌های تست با درجه‌های مختلف اتورگرسیون،

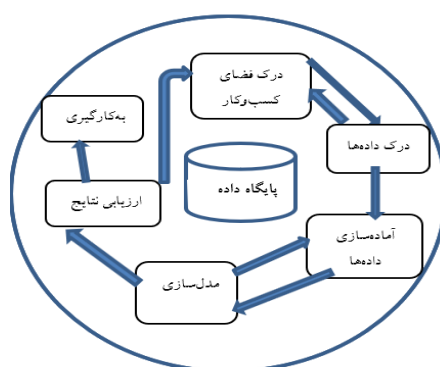
(قصری، ۲۰۱۷)

میانگین متحرک و تفاضل‌گیری و ارزیابی خروجی‌ها بر اساس پارامترهای Stationary R2 و Bayesian Information Criteria (BIC). مدل برتر را انتخاب و بر اساس آن میزان برق مورد نیاز برای تولید در ۱۰ سال بعد را پیش‌بینی نمودند

۳-فرآیند کریسپ

کریسپ داده‌کاوی، به روش کشف دانش که توسط اسامه فیاد و همکارانش ارائه شد، در طول شش مرحله اجرا می‌شود؛ که در ادامه، روش اجرای پژوهش بر اساس این مراحل تشریح خواهد شد (فیاد و همکارانش، ۲۰۰۸)

بزرگ‌ترین شرکت الکترونیکی نرم‌افزار و سخت‌افزار آمریکا، به‌عنوان بخشی از هدفش، برای دادن ارزش بیشتر به مشتریان انبار داده خود، گروه‌هایی از مشاوران و متخصصین فنی داده‌کاوی را برای سرویس‌دهی به نیازمندی‌های مشتریان تأسیس نمود (کیانیان، ۱۳۹۵) فرآیند داده‌کاوی در استاندارد



شکل ۲. مدل فرآیندی کریسپ

۳-۱-درک فضای کسب و کار

داده‌ها را منابع مختلف با پرسشنامه و یا سیستمی جمع‌آوری می‌کنند.

تشریح داده‌ها: تشریح داده‌ها را می‌توان بررسی سطحی و مقدماتی داده‌ها و استخراج مشخصات اولیه مجموعه داده‌ها، تعریف کرد. در این پژوهش در این مرحله فرمت داده، تعداد متغیرها، نوع متغیرها (پیوسته و گسسته) داده و مفهوم متغیرهای موردبررسی قرار می‌گیرد. همچنین در این بخش، با توجه تعداد رکوردها مشخص گردید، آیا نوع و حجم داده‌ها ملزومات پژوهش را برآورده می‌سازد.

اکتشاف داده: پس از تشریح مقدماتی داده‌ها، باید دید دقیق‌تری به داده‌ها پیدا کرد. در این مرحله باید با استفاده از روش‌های آماری مشخصات اصلی داده‌ها را استخراج و به کمک ابزارهای مصورسازی خصوصیات داده‌ها را به تصویر کشیده شد.

بررسی کیفیت داده‌ها: خروجی این مرحله معمولاً به‌صورت گزارشی هست که در آن اشتباهات داده‌ها، مقادیر از دست‌رفته، وجود رکوردهای تکراری، مشخص می‌گردد. منظور از اشتباه‌های داده‌ها، مقادیر بی‌معنی و داده دورافتاده و پرت است که در نتایج تأثیر می‌گذارد.

شناخت اهداف کسب‌وکار: مرحله اول هر پروژه داده‌کاوی این است که مشخص شود چه هدفی از اجرای پروژه مورد انتظار است.

درک اهداف: اهداف تجاری در محدوده کسب‌وکار و یا حوزه سلامت، مطرح هستند و اهداف داده‌کاوی باید در حوزه داده‌کاوی و تکنیک‌های آن بررسی شوند. مهم‌ترین خروجی در این مرحله لیست اهداف داده‌کاوی است. از طرفی ضوابط موفقیت داده‌کاوی نیز باید تعیین گردد.

تهیه برنامه پروژه: در انتهای مرحله نخست کریسپ باید برنامه مانند هر پژوهش دیگری تنظیم گردد. در این برنامه علاوه بر مشخص کردن جزئیات فعالیت‌های پژوهش و زمان اجرای آن، باید روش‌ها و ابزارهای موردنیاز نیز در پروژه انتخاب گردد. خروجی این بخش از فرایند به‌صورت مستندات اجرا و زمان‌بندی پژوهش است. برنامه پژوهش موردنظر در قالب پروپوزال، پیش‌ازاین ارائه شده است.

۳-۲-درک داده‌ها

جمع‌آوری داده‌های اولیه: در این مرحله بر اساس نوع پروژه،

۳-۳- آماده سازی داده ها

انتخاب داده ها: در برخی مسائل به علت حجم بسیار زیاد رکورد های موجود، از روش های نمونه گیری به منظور کاهش حجم داده استفاده می شوند.

پاک سازی داده ها: در این بخش باید با استفاده از گزارش خروجی بررسی کیفیت داده ها در مرحله قبل، ابتدا مقادیر از دست رفته را جایگزین کرد. هم چنین در مرحله دوم باید مقادیر بی معنی موجود در متغیرها با توجه به گزارش مفاهیم متغیرها در مرحله قبل، حذف و یا با مقادیر بامعنی جایگزین شوند.

کاهش متغیرها: در این مرحله پس از پاک سازی داده ها، باید از روش های موجود انتخاب ویژگی جهت یافتن متغیرهای تأثیرگذار و حذف متغیرهای کم تأثیر و بی تأثیر استفاده کرد. این مرحله بسیار کاربردی است و در مواقعی که تعداد متغیرها بالا باشد، زمان اجرا و کارایی مدل سازی را پایین می آورد، از این رو مرحله بسیار مهمی است.

۴- مدل سازی

انتخاب روش مدل سازی: در این مرحله ابتدا چندین تکنیک دسته بندی یا خوشه بندی بر اساس نوع داده ها انتخاب خواهد شد. با توجه به مرحله پیش پردازش ممکن است در این مرحله

بر اساس پارامترهای ارزیابی نیاز به بازگشت به عقب داریم و تکنیک های مختلف دسته بندی و پارامترها را بررسی کنیم.

۴-۱- ارزیابی و به کارگیری

ارزیابی نتایج: در این گام می توان با توجه حوزه پژوهش و مطالعه مقالات معتبر و پژوهش های کار شده روی داده ها، شاخص های ارزیابی مدل را انتخاب نمود؛ و بهترین روش و دسته بندی در هر مرحله بر اساس این ویژگی انتخاب شد. در مواردی که هزینه کم باشد و بودجه کافی در اختیار باشد، می توان مدل ساخته شده را بر روی مجموعه داده های واقعی دیگر آزمود و صحت و دقت مدل ساخته شده را بیشتر بررسی نمود.

۴-۲- جمع آوری داده های مدل

مجموعه داده جمع آوری شده شامل ۱۳۸۹ اطلاعات مرتبط به حمل نقل انواع کالا یک شرکت حمل و نقل در انواع ریلی و جاده ای است. که کالاهای شامل شمش روی، و شمش آلومینیومی را به کشور ترکیه صادر می کند. داده ها شامل ۱۴ ستون عددی و گسسته است. ویژگی ها به شرح جدول ۱ تعریف شده اند.

جدول ۱. داده ها و نوع آنها

نوع متغیر	عنوان
تاریخ شمسی	تاریخ تقاضا
عددی	تعداد تقاضا
عددی	تناژ تقاضی (به کیلوگرم)
گسسته	مبدا
گسسته	مقصد
گسسته	نوع بار
عددی	تعداد عرضه
عددی	تناژ بارگیری (به کیلوگرم)
عددی	زمان حمل تخمینی (به روز)
عددی	زمان حمل از روز تقاضا تا روز تخلیه
عددی	هزینه حمل (به ریال)
گسسته دو مقداری حمل ریلی: ۰ حمل کامیونی: ۱	نوع حمل
گسسته دو مقداری معتدل: ۰ نامعتدل: ۱	وضعیت آب و هوایی
گسسته دو مقداری دارد: ۱ ندارد: ۰	نیاز به مجوز

۴-۳-پیش پردازش: پاکسازی و آماده سازی داده ها

در این گام پس از تشریح داده با نمودار، در صورت وجود داده از دست رفته، بر مبنای تعداد داده های میسینگ رویکرد حذف و یا جایگزینی داده از دست رفته انتخاب می شود. هم چنین وجود سطرهای تکراری و متغیرهای اضافی را در نرم افزار بررسی می کنیم.

۴-۴-خوشه بندی

جهت تحلیل داده های حمل نقل، گام اول خوشه بندی و

استخراج دانش قوانین انجمنی

فرآیند کشف قوانین وابستگی، یکی از رویکردهای مهم در علم نوین داده کاوی برای یافتن قواعد و الگوها در پایگاه داده است که بسیار مورد توجه تحقیقگران قرار گرفته است. در گام استخراج دانش، هدف تولید قواعد انجمنی از مجموعه داده است تا روابط بین ویژگی ها در حوزه حمل و نقل استخراج شوند.

فرضیات ابتدایی از این قرارند: مجموعه $I = I_1, I_2, I_3, \dots$ را به عنوان مجموعه اقلام در نظر بگیرید. مجموعه D را به عنوان مجموعه داده ها یعنی تراکنش های پایگاه داده ها در نظر بگیرید بطوریکه هر تراکنش شامل مجموعه ای از اقلام باشد، یعنی هر تراکنش T زیرمجموعه ای از I است ($T \subseteq I$). هر تراکنش

گروه بندی اطلاعات بار است. جهت مدلسازی تحلیل از خوشه بندی کامینز k-means استفاده کردیم. ابتدا تعداد پارامترها را به طور دستی وارد کرده و با پارامتر ارزیابی دیویس بولدین ابتدا تعداد بهینه خوشه ها را مشخص می کنیم. در مرحله تحلیل، با بهینه ترین تعداد خوشه ها را با خروجی جدول مرکزی ارزیابی می کنیم.

شناسه ای بنام TID دارد. اگر A یک مجموعه از اقلام باشد، می گوییم. تراکنش T شامل A است اگر و فقط اگر $A \subseteq T$ باشد. یک قانون/انجمنی گزاره ای است به صورت $A \Rightarrow B$ در حالی که $A \subseteq I, B \subseteq I, A \cap B = \emptyset$. معیار مقبولیت قوانین انجمنی دو پارامتر مهم یعنی پشتیبانی و قابلیت اطمینان قوانین را معرفی می کنیم.

قانون $A \Rightarrow B$ در مجموعه تراکنش های D دارای پشتیبانی برابر با s است اگر s درصد از تراکنش های D شامل $A \cup B$ باشند و این قانون دارای قابلیت اطمینان برابر c است اگر c درصد از تراکنش هایی که شامل A هستند شامل B نیز باشند. یا به قول دیگر:

$$\text{Support}(A \Rightarrow B) = P(A \cup B)$$

$$\text{Confidence}(A \Rightarrow B) = P(B|A)$$

قوانین). مثلاً قانون زیر بیانگر این است که ۷۰ درصد کسانی که تلویزیون خریداری می کنند دستگاه پخش ویدئو نیز می خرند و ۸۰ درصد از کل تراکنشات شامل هر دو قلم تلویزیون و دستگاه پخش ویدئو است.

$$\text{Television} \Rightarrow \text{VCR}$$

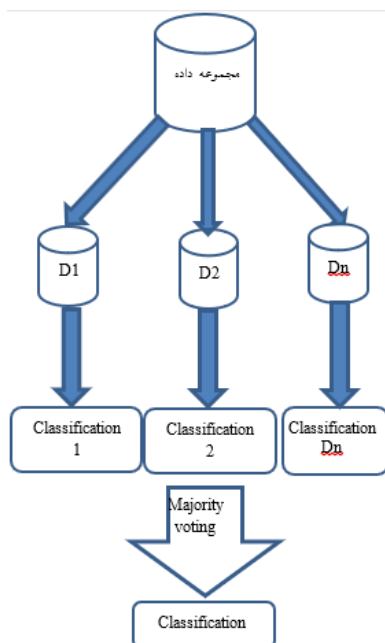
$$[\text{Support} = 80\% \text{ Confidence} = 70\%]$$

تصادفی و روش احتمالی بیز ساده استفاده کنیم و نتایج و قوانین خروجی را با هم مقایسه کنیم.

قوانینی که حد پایین Support یا min-sup و حد پایین Confidence یا min-conf را دارا باشند قوانین انجمنی قوی نامیده می شوند و هدف تمام الگوریتم ها چنین قوانینی است (یا چنانکه خواهیم دید در بعضی موارد زیرمجموعه ای از این

استخراج دانش (الگوریتم های دسته بندی)

در گام دسته بندی پس از گروه بندی داده های بار و تحلیل خوشه ها، از الگوریتم های دسته بندی درختی جمعی مانند جنگل



شکل ۲. مدل الگوریتم‌های تجمعی

می‌شود و همه داده‌ها دقیقاً یک‌بار برای آموزش و یک‌بار برای اعتبارسنجی بکار می‌روند. در نهایت میانگین نتیجه این K بار اعتبارسنجی، به عنوان یک تخمین نهایی برگزیده می‌شود.

معیارهای ارزیابی الگوریتم‌های طبقه‌بندی

این ماتریس چگونگی عملکرد الگوریتم طبقه‌بندی را با توجه به مجموعه داده ورودی به تفکیک انواع رده‌های مسئله طبقه‌بندی نشان می‌دهد. با توجه به جدول ۲، مفاهیم مثبت درست، مثبت غلط، منفی درست و منفی غلط به شرح ذیل است.

جهت اجرای دسته‌بندی باید داده‌ها را به دو دسته آموزشی و آزمون تقسیم کنیم (شکل ۲). با استفاده از روش اعتبارسنجی ضربدری مجموعه داده‌ها را به دودسته آموزشی و آزمون تقسیم می‌کنیم. در روش اعتبارسنجی ساده در تمام مراحل اجرا می‌توان مثلاً ۷۰ درصد مجموعه داده آموزشی و مابقی را برای آزمایش و محاسبه شاخص‌ها بکار گرفت؛ اما در روش هوشمندانه‌تر برای اینکه کل مجموعه داده یک‌بار هم در آموزش و هم در آزمایش باشد، از روش اعتبارسنجی ضربدری استفاده می‌کنیم که در این پژوهش نیز استفاده شد. در این نوع اعتبارسنجی داده‌ها به K زیرمجموعه افراز می‌شوند؛ از این K زیرمجموعه، هر بار یکی برای اعتبارسنجی و $K-1$ تای دیگر برای آموزش بکار می‌روند. این روال K بار تکرار

جدول ۲. ماتریس درهم‌ریختگی

		رکوردهای تخمینی	
رکوردهای واقعی		رده مثبت (+)	رده منفی (-)
	رده منفی (-)	TN	FP
	رده مثبت (+)	FN	TP

-مثبت درست: این مقدار بیانگر تعداد رکوردهایی است که رده واقعی آنها مثبت بوده و دسته‌بند نیز رده آنها را به درستی مثبت تشخیص داده است.

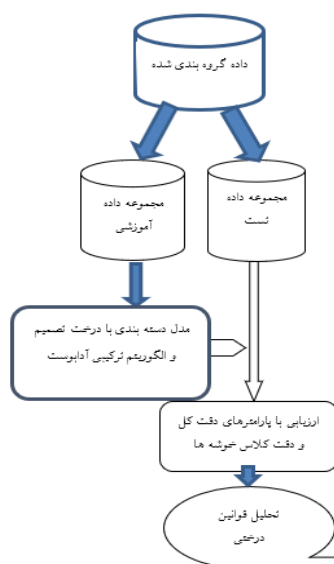
-مثبت غلط: این مقدار بیانگر تعداد رکوردهایی است که رده واقعی آنها منفی بوده و دسته‌بند رده آنها را به نادرستی مثبت تشخیص داده است.

-منفی درست: این مقدار بیانگر تعداد رکوردهایی است که رده واقعی آنها منفی بوده و دسته‌بند نیز رده آنها را به درستی منفی تشخیص داده است.

-منفی غلط: این مقدار بیانگر تعداد رکوردهایی است که رده واقعی آنها مثبت بوده و دسته‌بند رده آنها را به نادرستی منفی تشخیص داده است.

اکنون به تشریح انواع مهم معیارهای ارزیابی الگوریتم‌های طبقه‌بندی می‌پردازیم. مهم‌ترین معیار، دقت یا نرخ طبقه‌بندی است و به این مفهوم است که طبقه‌بند، چند درصد رکوردهای آزمایشی را به درستی طبقه‌بندی کرده است. این دقت بر اساس مفاهیم ماتریس با توجه به رابطه قابل محاسبه می‌شود.

$$Accuracy = \frac{Tp+TN}{Tp+TN+FP+FN} \quad (2)$$



شکل ۳. مدل دسته‌بندی (استخراج دانش)

ارزیابی

مصورسازی داده، آماده سازی و پاکسازی داده را با اجرا در نرم افزار تشریح می‌کنیم. سپس خوشه بندی و دسته بندی نتایج را در قالب جدول گزارش می‌دهیم. نتایج تحلیل خوشه‌های بهینه، جدول مرکزی خوشه ها را در انتها تشریح می‌کنیم. سپس در بخش دسته بندی دو روش جنگل تصادفی، و بیز ساده روی مجموعه داده خوشه بندی شده اعمال کرده و دقت مدل ها را گزارش داده و سپس ارزیابی می‌نماییم.

در ادامه بر اساس تعریف ماتریس درهم‌ریختگی و پارامترهای ارزیابی دقت کل، نتایج را از خروجی نرم‌افزار گزارش داده و سپس تحلیل می‌کنیم. در نهایت با نتیجه‌گیری و تحلیل خوشه‌بندی، استراتژی‌های حوزه حمل‌ونقل را گزارش خواهیم داد.

نتایج تحقیق

در این بخش بر اساس مدل فرآیند کریسپ، مراحل درک داده،

درک داده‌ها

اطلاعات آماری مجموعه داده با توجه به شکل ۴ به شرح زیر می‌باشد.

Name	Type	Missing	Statistics
رابط	Integer	0	1
تاریخ تقاضا	Polynomial	0	Least: 96/04/01 (1), Most: 96/01/08 (46), Values: 96/01/08 (46), 97/12/28 (40), ... (391 more)
تعداد تقاضا	Integer	0	Min: 1, Max: 18, Average: 1.699
تولید سفارشی	Integer	0	Min: 5000, Max: 480000, Average: 50007.940
مبدأ	Polynomial	0	Least: RASHT (1), Most: SAHLAN (742), Values: SAHLAN (742), ZANJAN (425), ... (16 more)
مکان	Polynomial	0	Least: DENZLI (1), Most: VAN (532), Values: VAN (532), ANKARA (181), ... (18 more)
نوع بار	Polynomial	0	Least: WIRE (1), Most: container (444), Values: container (444), ZHG NGOT (333), ... (22 more)
تعداد عرضه	Integer	2	Min: 1, Max: 18, Average: 1.696
تولید بازرگانی شده	Integer	0	Min: 4800, Max: 487070, Average: 47853.705
زمان حمل سفارشی	Integer	0	Min: 3, Max: 9, Average: 4.617
زمان حمل از روز تقاضا تا روز تحویل	Integer	0	Min: 3, Max: 12, Average: 5.413
فرمانه حمل	Integer	0	Min: 2874000, Max: 1587029000, Average: 112224154.488
نوع حمل	Binomial	0	Least: 1 (649), Most: 0 (739), Values: 0 (739), 1 (649)

شکل ۴. اطلاعات آماری داده

تنها در ستون عرضه ۲ سطر داده از دسته رفته وجود دارد. متغیر تاریخ تقاضا از نوع تاریخ شمسی است که در نرم افزار از نوع گسسته در نظر گرفته می‌شود. تعداد عرضه و تقاضا از رنج ۱ تا ۱۸ و میانگین ۱,۶۹۸ مشخص شده است. تناژ بارگیری از ۴۹۰۰ کیلوگرم تا ۴۸۷۰۷۰ کیلو تعریف شده است.

با توجه به شکل ۴ در بخش اطلاعات آماری ستون اول نام داده، ستون دوم نوع داده، ستون سوم تعداد داده‌های از دست رفته هر متغیر، ستون چهارم و پنجم ماکزیمم و مینیمم متغیرهای عددی و کمترین و بیشترین مقدار متغیرهای گسسته و ستون آخر مقدار میانگین و انحراف معیار متغیرهای عددی و مقدار مد متغیرهای گسسته را نشان می‌دهد.

Index	Nominal value	Absolute count	Fraction
1	SAHLAN	742	0.535
2	ZANJAN	425	0.306
3	QAZVIN	104	0.075
4	TABRIZ	69	0.050
5	ISFAHAN	13	0.009
6	YAZD	13	0.009
7	QOM	4	0.003
8	TEHRAN	4	0.003

Index	Nominal value	Absolute count	Fraction
11	SALAFCHEGAN	2	0.001
12	ARAK	1	0.001
13	BANDAR EMAM	1	0.001
14	ESFAHAN	1	0.001
15	ESHTEHARD	1	0.001
16	HAMEDAN	1	0.001
17	KASHAN	1	0.001
18	RASHT	1	0.001

شکل ۵. فراوانی مبدأ شهرها در حمل و نقل

در شهرهای قزوین، اصفهان، یزد، قم، تهران، ایلام، سلفچگان و اراک و ... تعریف شده است.

با توجه به شکل ۵، ۵۳ درصد حمل و نقل از واحد مبدأ از شهر سهران، ۳۰ درصد از زنجان و مابقی با درصد بسیار کمتر

Index	Nominal value	Absolute count	Fraction
1	VAN	532	0.383
2	ANKARA	181	0.130
3	MERSIN	141	0.102
4	GAZIANTEP	94	0.068
5	MURATBEY	93	0.067
6	KONYA	80	0.058
7	ADANA	75	0.054
8	ISKENDERUN	53	0.038

Index	Nominal value	Absolute count	Fraction
9	ERZURUM	49	0.035
10	DERINCE	33	0.024
11	IZMIR	18	0.013
12	KAYSERI	18	0.013
13	BURSA	6	0.004
14	ERENKOY	5	0.004
15	GEBZE	5	0.004
16	HATAY	4	0.003

شکل ۶. فراوانی مقصد شهرها در حمل و نقل

ریلی تا میانگین ۵ روز و نهایت ۸ روز فاصله زمانی داریم. -بیشترین هزینه حمل و نقل مربوط به حمل کامیونی تا ۱,۵۰۰,۰۰۰,۰۰۰ ریال است و تفاوت هزینه در حمل ریلی و کامیونی مشهود است.

پیش پردازش و آماده سازی داده‌ها

متغیرهای اضافی

در بخش تحلیل داده های حمل و نقل، گام اول در بخش خوشه‌بندی نیاز به متغیرهای توصیفی عددی و پیوسته داریم. به دلیل وجود متغیرهای زمان تخمینی و فاصله زمانی تاریخ حمل تا تخلیه، وجود ویژگی تاریخ تقاضا که از نوع شمسی است در مدلسازی و استخراج دانش اضافی و از داده کنار گذاشته می‌شود.

حذف مقدار از دست رفته

همان‌طور که در بخش تشریح داده اشاره کردیم تنها در ستون تعداد عر ضه ۲ سطر داده از دست رفته داریم، یکی از راه‌های حذف سطرهای با داده میسینگ است که ایراد این روش حذف اطلاعات مفید در ویژگی‌های دیگر است که دارای مقدار است. یکی از بهترین روش‌ها جایگزینی مقدار از دست رفته با مقدار میانگین همان ستون است تا انحراف معیار داده حفظ شود. از این رو در نرم‌افزار برای متغیر عددی عرضه از جایگزین داده از دست رفته استفاده می‌کنیم.

با توجه به شکل ۶، ۳۸ درصد مقصد بارها به شهر وان ترکیه، ۱۳ درصد آنکارا، ۱۰ درصد مرسین ترکیه، ۰,۰۶ درصد Gaziantep، ۰,۰۶ درصد Muratbey، ۰,۰۵ درصد Konya و مابقی شهرهای Adana، Izmir و سایر شهرهای مرزی ترکیه است.

-زمان حمل تخمین زده شده، از مینیمم ۳ روز تا ۹ روز و میانگین ۴,۶۱ تعریف شده است.

-زمان حمل از روز تقاضا تا روز تخلیه از مینیمم ۳ روز تا ۱۲ روز تعریف شده است.

-وضعیت آب و هوایی در ۲۳۲ سطر در حالت نامعتدل و ۱۱۵۶ سطر در حالت معتدل تعریف شده است.

-۴۶ درصد حمل و نقل‌ها در حالت کامیونی و ۵۳ درصد در حالت ریلی هستند.

-۹۶ درصد بارها در حمل و نقل نیاز به مجوز ندارند.

-اطلاعات از سال ۹۶ تا ۹۸ تعریف شده است.

-ماکزیمم رنج تناژ متقاضی و تناژ بارگیری شده، برای حمل کامیونی تا میانگین ۴۷۵۰۰۰ تعریف شده و برای حمل ریلی تنها از رنج ۲۵۰۰۰ تا ۵۰۰۰۰ کیلوگرم تعریف شده است.

-برای حمل ریلی تنها شهر سهلان مبدا حمل و نقل بوده و سایر بارها در حمل کامیونی در سایر شهرها تعریف شده‌اند.

-در حمل ریلی تنها شهرهای وان، قاضی انتپ، مرسین و اسکندرون مقصد حمل و نقل و مابقی مربوط به حمل کامیونی هستند.

-بیشترین زمان حمل از روز تقاضا تا روز تخلیه، در فاصله زمانی سه روز تا ۱۲ روز برای حمل کامیونی و برای حمل

✓ نوع بار	Polynomial	0	Least WIRE (1)
✓ تعداد عرضه	Integer	0	Min 1
✓ تعداد بارگیری شده	Integer	0	Max 4900
✓ زمان حمل نخشی	Integer	0	Min 3
✓ زمان حمل از روز تقاضا تا روز تخلیه	Integer	0	Min 3

شکل ۷. جایگزینی داده از دست رفته

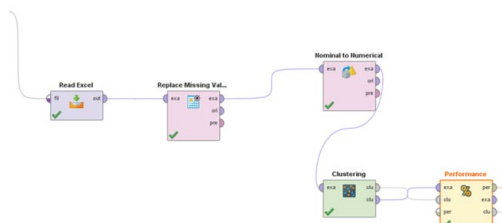
تعداد خوشه را ارزیابی می‌کنیم.

مدلسازی (خوشه بندی)

خوشه بندی کامیانگین

جهت خوشه بندی باید داده ها را به مقدار عددی تبدیل کنیم و مدل به شرح شکل ۸ است.

در این گام الگوریتم خوشه بندی کامیانگین را به داده های پاکسازی شده اعمال کرده و در هر مرحله تعداد k (خوشه ها) شبیه سازی و براساس پارامتر دیویس بولدین ارزیابی می‌نماییم و براساس مقدار ارزیابی و در بخش تحلیل نتایج بهینه ترین



شکل ۸. مدل خوشه بندی

نتایج پارامتر ارزیابی دیویس بولدین را به ازای تعداد خوشه های ۳ و ۴ گزارش می‌دهیم و در بخش مقایسه تحلیل می‌کنیم.

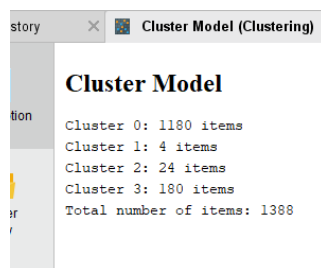
جدول ۳. مقایسه دو روش خوشه بندی

Medel(K)	Davis Bouldin
Kemans(k=3)	0.007
Kmeans(k=4)	0.005

تحلیل خوشه بندی 4-means

در این مرحله با توجه به نتایج نرم افزار به تحلیل خوشه ها، تعداد نمونه ها در هر خوشه می‌پردازیم. با توجه به شکل ۸ تعداد نمونه ها به تفکیک خوشه ها مشخص است. در خوشه صفر ۱۱۸۰ آیتم یعنی بیشتر نمونه ها، در خوشه ۱ تنها ۴ آیتم، در خوشه دو ۲۴ آیتم و در خوشه سه ۱۸۰ آیتم وجود دارد.

ارزیاب دیویس بولدین نسبت فاصله درون خوشه ای به فاصله بین خوشه ها را گزارش می‌دهد و عددی بین صفر و یک را خروجی می‌دهد در این صورت مدلی موفق است که نمونه ها در یک خوشه به یکدیگر نزدیک تر و متراکم تر و از نمونه های خوشه های دیگر دورتر باشند. یعنی هر چه دیویس بولدین به صفر نزدیکتر باشد مدل موفق تر و تعداد خوشه ها بهینه تر است. با توجه به جدول ۴ کمترین مقدار ارزیاب با تعداد ۴ خوشه و الگوریتم کامینز است. از این رو در بخش بعد نتایج خوشه بندی 4-means را برای گروه بندی، گزارش می‌دهیم.



شکل ۸. تفکیک نمونه خوشه

همچنین خروجی زیر (شکل ۹) مشخص می‌کند که هر تراکشن حمل بار در چه خوشه‌ای قرار دارد.

Row No.	ID	cluster	SARJAN	ZANJAN	GAZVIN	ASAQ	RASHT	TABRIZ	YAZD
1	1	cluster_0	1	0	0	0	0	0	0
2	2	cluster_3	0	1	0	0	0	0	0
3	3	cluster_0	0	1	0	0	0	0	0
4	4	cluster_0	0	1	0	0	0	0	0
5	5	cluster_0	1	0	0	0	0	0	0
6	6	cluster_0	1	0	0	0	0	0	0
7	7	cluster_0	1	0	0	0	0	0	0
8	8	cluster_0	1	0	0	0	0	0	0
9	9	cluster_0	1	0	0	0	0	0	0
10	10	cluster_0	1	0	0	0	0	0	0
11	11	cluster_0	0	1	0	0	0	0	0
12	12	cluster_0	0	1	0	0	0	0	0
13	13	cluster_0	0	1	0	0	0	0	0
14	14	cluster_1	0	0	1	0	0	0	0
15	15	cluster_0	0	1	0	0	0	0	0
16	16	cluster_0	0	1	0	0	0	0	0
17	17	cluster_0	0	0	1	0	0	0	0
18	18	cluster_0	0	1	0	0	0	0	0
19	19	cluster_0	0	0	0	1	0	0	0
20	20	cluster_0	1	0	0	0	0	0	0
21	21	cluster_0	1	0	0	0	0	0	0
22	22	cluster_3	0	1	0	0	0	0	0
23	23	cluster_0	1	0	0	0	0	0	0
24	24	cluster_3	0	1	0	0	0	0	0
25	25	cluster_0	0	0	1	0	0	0	0

شکل ۹. خوشه‌بندی نمونه‌ها

این جدول با روش 4-means به صورت جدول ۴ می‌باشد.

جدول ۴. جدول مرکزی خوشه‌ها بر اساس شاخص‌های عددی و کیفی حمل‌ونقل

شاخص های حمل و نقل	امتیاز cluster0	امتیاز cluster1	امتیاز cluster2	امتیاز cluster3
هزینه حمل	۷۲۰۴۸۴۶۳٫۲	۱۴۷۱۴۴۹۷۵۰	۶۶۷۱۳۱۸۴۷٫۵	۲۷۲۲۷۷۱۳۱٫۵
زمان حمل از روز تقاضا تا روز تخلیه	۵٫۲۳۱	۶٫۷۵۰	۶٫۳۵۰	۶٫۶۶۱
زمان حمل تخمینی	۴٫۴۴۵	۶٫۲۵۰	۵٫۶۰۰	۵٫۲۷۱
تناژ بارگیری	۳۷۷۵۱	۱۹۵۶۲۷	۱۴۸۲۷۸٫۱	۸۲۳۳۸٫۲
تعداد عرضه	۱٫۲۵۴	۸٫۲۵	۶٫۱۵۰	۳٫۳۲۷
تناژ متقاضی	۴۰۹۹۷٫۵	۱۹۸۰۰۰	۱۴۷۰۷۵	۷۹۸۳۳٫۳
تعداد تقاضا	۱٫۲۵۶	۸٫۲۵۰	۶٫۱۵۰	۳٫۳۲۱
نیاز به مجوز: ندارد	۰٫۹۷۰	۱	۰٫۹۵۰	۰٫۹۳۹
نیاز به مجوز: دارد	۰٫۰۳۰	۰	۰٫۰۵۰	۰٫۰۶۱
وضعیت آب و هوایی: معتدل	۰٫۸۴۴	۰٫۷۵۰	۰٫۶۰۰	۰٫۷۷۰
وضعیت آب و هوایی: نامعتدل	۰٫۱۵۶	۰٫۲۵۰	۰٫۴۰۰	۰٫۲۳۰
نوع حمل: ریلی	۰٫۶۲۵	۰	۰	۰٫۰۰۶
نوع حمل: کامیونی	۰٫۳۷۵	۱	۱	۰٫۹۹۴

با توجه به جدول ۴ تحلیل هر خوشه به شرح زیر است. خوشه ۱ (۲۸،۰ درصد): بارهای حمل و نقل شده در این گروه از نظر ویژگی حمل پرهزینه ترین بارها نسبت به خوشه های دیگر (تا مبلغ ۱ میلیارد بیشتر)، است و میانگین تناژ بارگیری بیشترین مقدار، هم چنین بیشترین تعداد عرضه و تقاضا در این گروه است. ۱۰۰ درصد بارهای حمل و نقلی این گروه نیاز به مجوز ندارد و نوع حمل تمام تراکنش ها ۱۰۰ درصد از نوع کامیونی است و حمل ریلی را شامل نمیشود. مبدا تمامی بارها از شهر قزوین و ۷۵ درصد به شهر آدانا و ۲۵ درصد به شهر MURATBAY است. نوع بارها در ۷۵ درصد از نوع میلگرد و ۲۵ درصد از نوع شمش روی است. ماکزیمم تایم زمانی حمل از روزتقاضا تا تخلیه در بارهای این خوشه است. تنها ۴ آیتم بار با ویژگی های خیلی متفاوت و ویژه در این گروه قرار دارند.

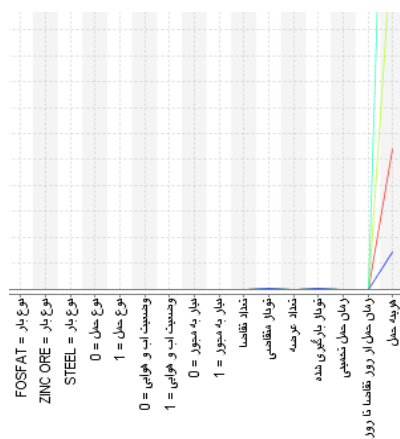
خوشه صفر (فراوانی ۸۵،۰۱ درصد): بیشترین بارها (۱۱۸۰ آیتم) در این گروه هستند که به طور میانگین کم هزینه ترین بارها نسبت به خوشه های دیگر، کمترین زمان حمل از روزتقاضا تا روز تخلیه در مورد بارهای این گروه است. ۹۹ درصد بارهای ریلی در این خوشه هستند و در گروه های دیگر به جز درصد بسیار کم در خوشه ۳، بار از نوع ریلی نداریم. اما ۳۷ درصد نیز بار کامیونی داریم. با تفاوت زیادی کمترین مقدار عرضه و تقاضا در این خوشه است. مبدا شهرها ۶۲ درصد از سهلان و ۲۶ درصد از زنجان و مابقی با درصد کم از قزوین، تبریز و یزد است و تقریباً به تمامی شهرها مقصد داریم اما بیشترین مکان خالی شده بار در شهر وان ترکیه (۴۲ درصد) است. ۳۷ درصد نوع بارها کانتینر و ۲۲ درصد از نوع شمش روی است.

خوشه دو (۱،۷۲ درصد): بعد از خوشه ۱ هزینه های نسبتاً بالای حمل، مربوط به بارهای این گروه است. بیشترین تناژ بارگیری بعد از خوشه ۱، و بیشترین تعداد عرضه و تقاضا با میانگین ۶،۱ در مورد آیتم های این گروه است. ۰،۰۵ درصد بارها برای حمل و نقل در این گروه نیاز به مجوز دارند و ۴۰ درصد حمل و نقل ها در وضعیت آب و هوایی نامعتدل است. در صورتیکه در خوشه های دیگر بیشتر حمل و نقل ها در وضعیت معتدل بوده است. ۴۱ درصد بارها از شهر رشت، ۲۹ درصد قزوین و ۲۰ درصد قزوین ارسال شده و مقصد با درصد ۳۳ درصد به وان و مابقی آدانا و آنکارا بوده است. ۱۰۰ درصد بارها از نوع کامیونی است.

خوشه ۳ (۱۳،۰۷ درصد): هزینه های حمل و نقل در این گروه متوسط، از خوشه ۱ و ۲ کمتر از خوشه صفر بیشتر است. تفاوت شاخص این گروه نسبت به گروه های دیگر در تفاوت فاصله زمانی نسبت زمان حمل تخمینی تا زمان حمل از روزتقاضا و روز تخلیه است که نسبت به گروه های دیگر بیشتر است. تناژ بارگیری، تعداد عرضه و تقاضا در این خوشه متوسط و ۰،۰۶ درصد بارها نیاز به مجوز دارد. ۲۳ درصد در وضعیت آب و هوایی نامعتدل بوده است. در ۰،۰۶ بارها از نوع ریلی است. ۵۶ درصد بارها از زنجان و به طور میانگین به تمامی شهرهای ترکیه ارسال شده است.

متغیرهای تاثیرگذار در تحلیل بارها

در بحث تحلیل و گروه بندی خوشه ها، با نمودار deviation (شکل ۱۰) فاکتورهای تاثیرگذار در طبقه بندی بارهای در سیستم حمل و نقل را بررسی می کنیم.



شکل ۱۰. نمودار deviation خوشه ها

این است که این سه فاکتور نسبت به فاکتورهای وضعیت هوا، نوع بار و مبدا و مقصد تاثیرگذاری بیشتری در جداسازی خوشه‌ها دارند.

مدلسازی (قواعد انجمنی)

جهت تولید قواعد انجمنی و اقلام مکرر، با الگوریتم fp-growth در ریپیدماینر، فیچرهای گسسته و باینری را برای تولید قوانین مکرر و آیتم‌هایی که با هم تکرار می‌شوند را انتخاب و نتایج به شرح زیر است.

Size	Support	Item 1	Item 2	Item 3
1	0.535	SHALAN = ۱۰۰		
1	0.468	فرع جل		
1	0.383	VAQ = ۱۰۰		
1	0.320	container = ۱۰۰		
1	0.306	ZANJAN = ۱۰۰		
1	0.240	ZHIG INQOT = ۱۰۰		
1	0.187	مقصود از بار و نقل		
2	0.343	SHALAN = ۱۰۰	VAQ = ۱۰۰	
2	0.320	SHALAN = ۱۰۰	container = ۱۰۰	
2	0.306	فرع جل	ZANJAN = ۱۰۰	
2	0.240	فرع جل	ZHIG INQOT = ۱۰۰	
2	0.320	VAQ = ۱۰۰	container = ۱۰۰	
2	0.183	ZANJAN = ۱۰۰	ZHIG INQOT = ۱۰۰	
3	0.320	SHALAN = ۱۰۰	VAQ = ۱۰۰	container = ۱۰۰
3	0.183	فرع جل	ZANJAN = ۱۰۰	ZHIG INQOT = ۱۰۰

Parameters

FP.Growth

☒ find min number of itemsets

min number of itemsets
100

max number of retries
15

positive value

min support
0.7

max items
-1

must contain

شکل ۱۱. اقلام مکرر حمل و نقل

شکل ۱۱ بخشی از آیتم‌های پرتکرار را مشخص می‌کند که براساس مقدار $\text{support} = 70\%$ پرتکرارترین آیتم‌ها، به شرح زیر است.

۱. ۵۳ درصد سیستم حمل و نقل مبدا از شهر سهران و ۳۸ درصد ا مقصد بارها شهر وان بوده است.

۲. ۴۶٫۵ درصد بارها از نوع حمل کامیونی است.

۳. ۳۲ درصد سیستم حمل نوع بار کانتینر و ۲۴ درصد شمش روی بوده است.

۴. در ۳۴ درصد حمل و نقل، مبدا از سهران و مقصد به شهر وان بوده است.

۵. در ۳۲ درصد حمل و نقل با مبدا از شهر سهران، نوع بار کانتینر است.

۶. در ۳۰ درصد نوع حمل کامیونی، بار از نوع شمش روی بوده است.

۷. در ۳۰ درصد نوع حمل کامیونی، مبدا از شهر زنجان است.

۸. در ۳۲ مقصد به شهر وان، بار از نوع کانتینر است.

۹. در ۳۲ درصد مبدا سهران، مقصد شهر وان و نوع بار کانتینر است.

با $\text{confidence}=70\%$ قوانین انجمنی؛ اگر آنگاه؛ به شرح زیر است.

No.	Premises	Conclusion	Support	Confidence
4	VAN = مقصد	container = نوع بار	0.320	0.835
5	VAN = مقصد	container = نوع بار SAHLAN = مبدا	0.320	0.835
6	VAN = مقصد	SAHLAN = مبدا	0.343	0.895
7	VAN = مقصد SAHLAN = مبدا	container = نوع بار	0.320	0.933
8	container = نوع بار	SAHLAN = مبدا	0.320	1
9	ZANJAN = مبدا	نوع حمل	0.306	1
10	ZING INGOT = نوع بار	نوع حمل	0.240	1
11	container = نوع بار	VAN = مقصد	0.320	1
12	container = نوع بار	VAN = مقصد SAHLAN = مبدا	0.320	1
13	container = نوع بار SAHLAN = مبدا	VAN = مقصد	0.320	1
14	container = نوع بار VAN = مقصد	SAHLAN = مبدا	0.320	1
15	ZING INGOT = نوع بار ZANJAN = مبدا	نوع حمل	0.193	1

شکل ۱۲. قواعد انجمنی

۱. با اطمینان ۸۳,۵ درصد اگر مقصد بار شهر وان است، آنگاه نوع بار از نوع کانتینر است.
۲. با اطمینان ۸۳,۵ درصد اگر مقصد بار شهر وان است، آنگاه نوع بار از نوع کانتینر و مبدا از شهر سهران است.
۳. با اطمینان ۸۹,۵ درصد اگر مقصد بار شهر وان است، آنگاه مبدا از شهر سهران است.
۴. با اطمینان ۹۳,۳ درصد، اگر مقصد بار شهر وان و مبدا از شهر، نوع بار کانتینر است.
۵. با اطمینان ۱۰۰ درصد، اگر مبدا شهر زنجان است، نوع حمل کامیونی است.
۶. با اطمینان ۱۰۰ درصد، اگر نوع بار شمش روی باشد، نوع حمل کامیونی است.
۷. با اطمینان ۱۰۰ درصد، اگر نوع بار کانتینر باشد، مقصد شهر وان است.
۸. با اطمینان ۱۰۰ درصد، اگر نوع بار کانتینر باشد، مقصد شهر وان و مبدا شهر سهران است.
۹. با اطمینان ۱۰۰ درصد، اگر نوع بار شمش روی باشد، مبدا زنجان و نوع حمل کامیونی است.

مدلسازی (دسته بندی: استخراج دانش)

در این بخش دو دسته بند روش احتمالی بیزساده و روش جمعی درختی جنگل تصادفی به مجموعه داده با برچسب خوشه ها اعمال می کنیم. و نتایج هر دسته بند را گزارش می دهیم.

نتایج بیزساده

مجموعه داده (خروجی خوشه بند) را با روش اعتبار سنجی ضربداری در ده مرحله به داده های آموزشی و تست تقسیم می کنیم و متغیر برچسب (هدف) را cluster در نظر می گیریم که یک متغیر گسسته چند مقداری است. (Nominal).

بعد از تقسیم نمونه ها از بیزساده برای دسته بندی استفاده می کنیم.

در سمت راست نرم افزار پارامترهای دقت که در خروجی مشخص می شود را مشخص می کنیم. پارامترهای دقت کل (accuracy), Recall پس از اجرا، نتایج به شرح زیر است.

Accuracy: 89.77%

نتیجه ماتریس درهم ریختگی به صورت زیر است.

accuracy: 88.77% +/- 3.00% (mknn: 88.77%)

	true cluster_0	true cluster_3	true cluster_1	true cluster_2	class precision
pred cluster_0	1078	6	0	0	99.45%
pred cluster_3	94	147	0	2	61.49%
pred cluster_1	0	0	2	1	66.67%
pred cluster_2	10	27	2	21	35.00%
class recall	91.19%	81.67%	50.00%	87.50%	

شکل ۱۳. نتایج دقت بیز ساده و ماتریس درهم‌ریختگی

-دقت کلاس خوشه صفر: ۹۱,۱۹ درصد

-دقت کلاس خوشه ۳: ۸۱,۶۷ درصد

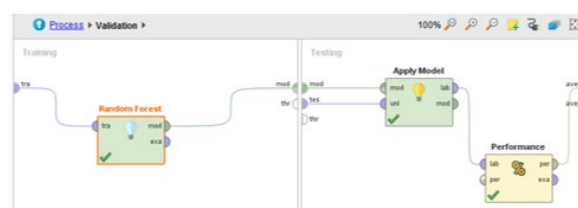
-دقت کلاس خوشه ۲: ۵۰ درصد

-دقت کلاس خوشه ۱: ۸۷,۵۰ درصد.

نتایج جنگل تصادفی

یک طبقه‌بندی می‌دهد و می‌گوییم این درخت به آن کلاس «رای» می‌دهد. جنگل طبقه‌بندی ای که بیشترین رای را داشته باشد (بین همه درخت‌های جنگل) انتخاب می‌کند.

جنگل تصادفی درخت تصمیم‌های زیادی را تولید می‌کند. برای طبقه‌بندی یک شیء جدید از برداری ورودی را در انتهای هر یک از درختان جنگل تصادفی قرار می‌دهد. هر درخت به ما



Parameters

Random Forest

number of trees: 10

criterion: gain_ratio

maximal depth: 10

☒ apply pruning

confidence: 0.25

☒ apply prepruning

[Hide advanced parameters](#)

[Change compatibility \(7.1.001\)](#)

شکل ۱۴. عملگر جنگل تصادفی

الف) دقت خوشه ۰: ۹۹,۷۵ درصد

-از این رو تعداد نمونه‌هایی که واقعاً در خوشه صفر هستند و به درستی پیش‌بینی شده‌اند: ۱۱۷۷ نمونه

-ازاینرو تعداد نمونه‌هایی که واقعاً در خوشه ۰ هستند و به‌اشتباه در خوشه ۱ پیش‌بینی شده‌اند: ۰ نمونه
 -ازاینرو تعداد نمونه‌هایی که واقعاً در خوشه ۰ هستند و به‌اشتباه در خوشه ۲ پیش‌بینی شده‌اند: ۰ نمونه
 -ازاینرو تعداد نمونه‌هایی که واقعاً در خوشه ۰ هستند و به‌اشتباه در خوشه ۳ پیش‌بینی شده‌اند: ۳ نمونه

ب) دقت خوشه ۳: ۹۱,۶۷ درصد

-ازاینرو تعداد نمونه‌هایی که واقعاً در خوشه ۳ هستند و به‌درستی پیش‌بینی شده‌اند: ۱۶۵ نمونه
 -ازاینرو تعداد نمونه‌هایی که واقعاً در خوشه ۳ هستند و به‌اشتباه در خوشه ۰ پیش‌بینی شده‌اند: ۱۵ نمونه
 -ازاینرو تعداد نمونه‌هایی که واقعاً در خوشه ۳ هستند و به‌اشتباه در خوشه ۲ پیش‌بینی شده‌اند: ۰ نمونه
 -ازاینرو تعداد نمونه‌هایی که واقعاً در خوشه ۳ هستند و به‌اشتباه در خوشه ۱ پیش‌بینی شده‌اند: ۰ نمونه

پ) دقت خوشه ۱: ۱۰۰ درصد

-ازاینرو تعداد نمونه‌هایی که واقعاً در خوشه ۱ هستند و به‌درستی پیش‌بینی شده‌اند: ۴ نمونه
 -ازاینرو تعداد نمونه‌هایی که واقعاً در خوشه ۰ هستند و به‌اشتباه در خوشه ۱ پیش‌بینی شده‌اند: ۰ نمونه
 -ازاینرو تعداد نمونه‌هایی که واقعاً در خوشه ۰ هستند و به‌اشتباه در خوشه ۲ پیش‌بینی شده‌اند: ۰ نمونه
 -ازاینرو تعداد نمونه‌هایی که واقعاً در خوشه ۰ هستند و به‌اشتباه در خوشه ۳ پیش‌بینی شده‌اند: ۰ نمونه

ج) دقت خوشه ۲: ۵۴,۱۷ درصد

-ازاینرو تعداد نمونه‌هایی که واقعاً در خوشه ۲ هستند و به‌درستی پیش‌بینی شده‌اند: ۱۳ نمونه
 -ازاینرو تعداد نمونه‌هایی که واقعاً در خوشه ۲ هستند و به‌اشتباه در خوشه ۱ پیش‌بینی شده‌اند: ۰ نمونه
 -ازاینرو تعداد نمونه‌هایی که واقعاً در خوشه ۲ هستند و به‌اشتباه در خوشه ۳ پیش‌بینی شده‌اند: ۸ نمونه
 -ازاینرو تعداد نمونه‌هایی که واقعاً در خوشه ۲ هستند و به‌اشتباه در خوشه ۰ پیش‌بینی شده‌اند: ۳ نمونه

تحلیل نتایج

تفاوت است و عدم توازن در خوشه‌ها مشاهده می‌شود. دقت مدل‌ها در خوشه‌ها تفاوت زیادی دارد. در روش بیز ساده روی خوشه ۱ که تنها ۴ آیتم دارد قطعیت پایین ۵۰ درصد دارد. اما در خوشه‌ها دیگر نیز میانگین ۸۱ درصد داریم که خطای ۱۵ درصد بیشتر از جنگل تصادفی است. به طور کلی در روشهای جمعی، به دلیل تقسیم مجموعه داده، روی داده‌های نامتوازن دقت بهتری دارند و نتایج را بهبود می‌دهند.

نتایج دو مدل درخت تصمیم و جنگل تصادفی در خروجی خوشه‌بندی را به‌طور خلاصه جهت تحلیل در جدول ۵ گزارش شده است.

با دقت در نتایج مشاهده می‌کنیم که دقت روش جمعی جنگل تصادفی که از ۱۰ درخت تصمیم برای پیش‌بینی استفاده می‌کند، با دقت ۹۱,۹۷٪ دقت بالاتری از روش احتمالی بیز ساده دارد. از آنجا که نسبت نمونه‌ها در خوشه‌های مختلف بسیار

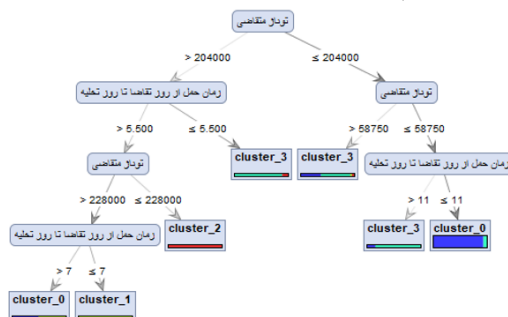
جدول ۵. نتایج دسته‌بندی

Model	Recall (cluster0)	Recall (cluster3)	Recall (cluster2)	Recall (cluster 1)	Accuracy
NB	91.19%	81.67%	87.50%	50%	89.77%
Random forest	99.75%	91.67%	54.17%	100%	97.91%

تحلیل دانش درختان جنگل تصادفی

سیاه مانند شبکه عصبی، نتایج آن را می توان به صورت قوانین «اگر آنگاه» بررسی کرد و ارتباط بین فاکتورها را در حوزه حمل و نقل تحلیل کرد. چند مدل درختان جنگل تصادفی که در بخش دسته بندی اشاره کردیم، به شرح زیر است.

درخت تصمیم از جمله روش هایی است که خروجی آن به صورت قوانین قابل تحلیل بررسی می شود، از این رو جزو روش های جعبه سفید است و در مقایسه با الگوریتم های جعبه



شکل ۱۵. درخت ۱

Tree 1

تناژ متقاضی < 204000

| زمان حمل از روز تقاضا تا روز تخلیه $< 5,500$

| | تناژ متقاضی < 228000

| | | زمان حمل از روز تقاضا تا روز تخلیه < 7 : $\{cluster_0=1, cluster_1=0, cluster_2=0, cluster_3=0\}$

| | | $\{cluster_3=0, cluster_1=1, cluster_2=0\}$

| | | زمان حمل از روز تقاضا تا روز تخلیه ≥ 7 : $\{cluster_0=0, cluster_3=0, cluster_1=3, cluster_2=0\}$

| | | تناژ متقاضی ≥ 228000 : $\{cluster_0=0, cluster_3=0, cluster_1=0, cluster_2=2\}$

| | | زمان حمل از روز تقاضا تا روز تخلیه $\geq 5,500$: $\{cluster_0=0, cluster_3=17, cluster_1=0, cluster_2=2\}$

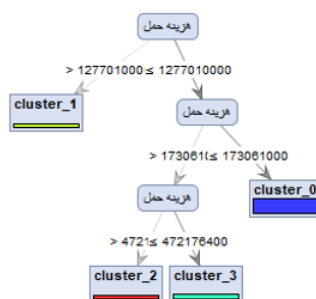
تناژ متقاضی ≥ 204000

| | تناژ متقاضی < 58750 : $\{cluster_0=47, cluster_3=66, cluster_1=4, cluster_2=5\}$

| | تناژ متقاضی ≥ 58750

| | | زمان حمل از روز تقاضا تا روز تخلیه < 11 : $\{cluster_0=1, cluster_3=5, cluster_1=0, cluster_2=0\}$

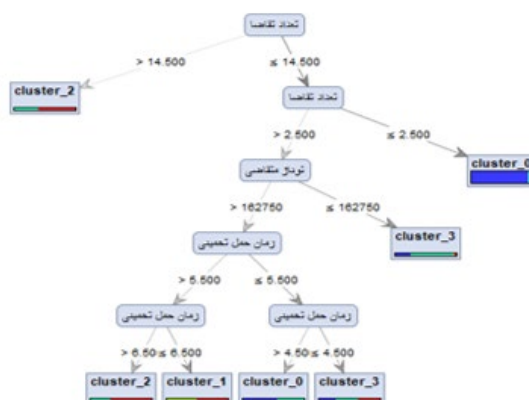
| | | زمان حمل از روز تقاضا تا روز تخلیه ≥ 11 : $\{cluster_0=1157, cluster_3=77, cluster_1=0, cluster_2=0\}$



شکل ۱۶. درخت ۲

Tree 2

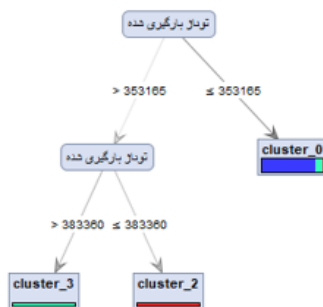
هزینه حمل < 1141070000 : $\{cluster_0=0, cluster_3=0, cluster_1=6, cluster_2=0\}$ cluster_1
 هزینه حمل ≥ 1141070000
 هزینه حمل < 173113500 |
 هزینه حمل < 519008850 : $\{cluster_0=0, cluster_3=0, cluster_1=0, cluster_2=13\}$ cluster_2 |
 هزینه حمل ≥ 519008850 : $\{cluster_0=0, cluster_3=186, cluster_1=0, cluster_2=0\}$ cluster_3 |
 هزینه حمل ≥ 173113500 : $\{cluster_0=1183, cluster_3=0, cluster_1=0, cluster_2=0\}$ cluster_0 |



شکل ۱۷. درخت ۳

Tree 3

تناژ بارگیری شده < 353165
 تناژ بارگیری شده < 383360 : $\{cluster_0=0, cluster_3=9, cluster_1=0, cluster_2=0\}$ cluster_3 |
 تناژ بارگیری شده ≥ 383360 : $\{cluster_0=0, cluster_3=0, cluster_1=0, cluster_2=7\}$ cluster_2 |
 تناژ بارگیری شده ≥ 353165 : $\{cluster_0=1193, cluster_3=157, cluster_1=1, cluster_2=21\}$ cluster_0



شکل ۱۸. درخت ۴

Tree 4

تعداد تقاضا < 14500 : $\{cluster_0=0, cluster_3=2, cluster_1=0, cluster_2=3\}$ cluster_2
 تعداد تقاضا ≥ 14500
 تعداد تقاضا < 2500 |
 تناژ متقاضی < 162750 |
 زمان حمل تخمینی < 6500 |

cluster_2 {cluster_0=0, cluster_3=1, cluster_1=0, cluster_2=2}	زمان حمل تخمینی $< ۶,۵۰۰$:				
cluster_1 {cluster_0=0, cluster_3=0, cluster_1=3, cluster_2=3}	زمان حمل تخمینی $\geq ۶,۵۰۰$:				
	زمان حمل تخمینی $\geq ۵,۵۰۰$:				
cluster_0 {cluster_0=4, cluster_3=3, cluster_1=0, cluster_2=0}	زمان حمل تخمینی $< ۴,۵۰۰$:				
cluster_3 {cluster_0=4, cluster_3=5, cluster_1=0, cluster_2=5}	زمان حمل تخمینی $\geq ۴,۵۰۰$:				
cluster_3 {cluster_0=30, cluster_3=77, cluster_1=2, cluster_2=4}	تناژ متقاضی ≥ ۱۶۲۷۵۰ :				
cluster_0 {cluster_0=1158, cluster_3=79, cluster_1=0, cluster_2=3}	تعداد تقاضا $\geq ۲,۵۰۰$:				

تجزیه و تحلیل نتایج

-در حمل ریلی تنها شهرهای وان، قاضی انتپ، مرسین و اسکندرون مقصد حمل و نقل و مابقی مربوط به حمل کامیونی هستند.

-بیشترین زمان حمل از روز تقاضا تا روز تخلیه، در فاصله زمانی سه روز تا ۱۲ روز برای حمل کامیونی و برای حمل ریلی تا میانگین ۵ روز و نهایت ۸ روز فاصله زمانی داریم.

-بیشترین هزینه حمل و نقل مربوط به حمل کامیونی تا ۱,۵۰۰,۰۰۰,۰۰۰ ریال است و تفاوت هزینه در حمل ریلی و کامیونی مشهود است.

در گام پاکسازی داده تنها در ویژگی تعداد عرضه دو سطر داده از دست رفته داریم، که به جای حذف سطرها که اطلاعات مهم ستون های دیگر حذف شود، از جانشینی داده با مقدار میانگین ستون استفاده کردیم. متغیر تاریخ تقاضا از نوع شمس است و به دلیل وجود فاصله زمانی تقاضا تا روز حمل، نیازی به این ویژگی نداشته و آن را حذف کردیم.

استخراج دانش (خوشه بندی)

جهت بهینه سازی تعداد خوشه ها و بهترین مقدار خوشه بندی از ارزیاب دیوس بولدین و مفاهیم خوشه بندی به دست آمد؛ این ارزیاب هر چه کمتر باشد خوشه ها بهینه تر یعنی فاصله درون خوشه ای متراکم تر و بین خوشه ای بیشتر است. که بر اساس نمودار ۵-۱ با مقدار $k=4$ بهترین خوشه بندی با کامینز را داشتیم.

در ادامه نتایج گروه بندی مشتریان خوشه بندی کامینز را از نظر تعداد نمونه ها در هر خوشه، جدول میانگین مرکزی خوشه ها بررسی کردیم؛ خوشه ها را از نظر مفاهیم حمل و نقل تحلیل کردیم؛ شکل ۲۰.

مجموعه داده جمع آوری شده شامل ۱۳۸۸ اطلاعات مرتبط به حمل نقل انواع کالا یک شرکت حمل و نقل در انواع ریلی و جاده ای است. که کالاهای شامل شمش روی، و شمش آلومینیومی را به کشور ترکیه صادر می کند. داده ها شامل ۱۴ ستون عددی و گسسته است. نوع ویژگی ها مبدأ، مقصد، نوع بار، هزینه حمل، زمان حمل تخمینی و زمان انجام شده، نوع حمل، تعداد عرضه، تناژ متقاضی و تناژ بارگیری است. با تشریح اطلاعات آماری و نمودارهای فراوانی اطلاعات زیر استخراج شد:

-تعداد عرضه و تقاضا از رنج ۱ تا ۱۸ و میانگین ۱,۶۹۸ مشخص شده است.

-تناژ بارگیری از ۴۹۰۰ کیلوگرم تا ۴۸۷۰۷۰ کیلو تعریف شده است.

-۵۳ درصد حمل و نقل از واحد مبدأ از شهر سهران، ۳۰ درصد از زنجان و مابقی با درصد بسیار کمتر در شهرهای قزوین، اصفهان، یزد، قم، تهران، ایلام، سلفچگان و اراک و ... تعریف شده است.

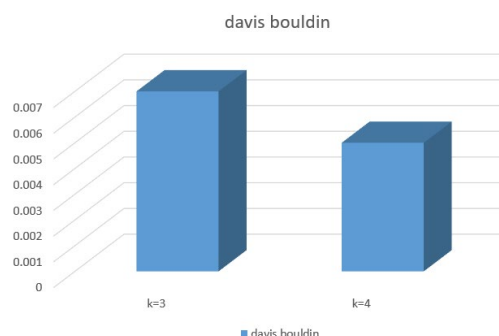
-زمان حمل تخمین زده شده، از مینیمم ۳ روز تا ۹ روز و میانگین ۴,۶۱ تعریف شده است. زمان حمل از روز تقاضا تا روز تخلیه از مینیمم ۳ روز تا ۱۲ روز تعریف شده است.

-۴۶ درصد حمل و نقل ها در حالت کامیونی و ۵۳ درصد در حالت ریلی هستند.

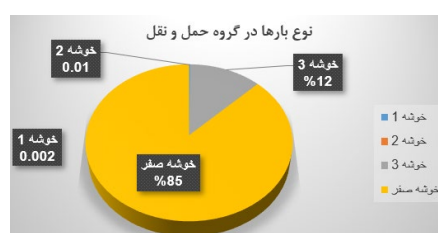
-۹۶ درصد بارها در حمل و نقل نیاز به مجوز ندارند و اطلاعات از سال ۹۶ تا ۹۸ تعریف شده است.

-ماکزیمم رنج تناژ متقاضی و تناژ بارگیری شده، برای حمل کامیونی تا میانگین ۴۷۵۰۰۰ تعریف شده و برای حمل ریلی تنها از رنج ۲۵۰۰۰ تا ۵۰۰۰۰ کیلوگرم تعریف شده است.

-برای حمل ریلی تنها شهر سهران مبدأ حمل و نقل بوده و سایر بارها در حمل کامیونی در سایر شهرها تعریف شده اند.



شکل ۱۹. مقایسه ارزیاب خوشه بندی با دیویس بولدین



شکل ۲۰. نمودار فراوانی گروه بارهای حمل و نقل

میانگین ۶٫۱ در مورد آیتم‌های این گروه با ۲۴ نمونه است. وجه تمایز این خوشه با خوشه ۱ در این است که ۴۰ درصد بارها در وضعیت هوایی نامعادل است که عمدتاً به وان و آدانا و آنکارا ارسال شده است.

خورشه ۳ (شاخص‌های متوسط از نظر حمل و نقل با بیشترین فاصله زمانی حمل): هزینه‌های حمل و نقل در این گروه متوسط، از خوشه ۱ و ۲ کمتر از خوشه صفر بیشتر است. تفاوت شاخص این گروه نسبت به گروه‌های دیگر در تفاوت فاصله زمانی نسبت زمان حمل تخمینی تا زمان حمل از روز تقاضا و روز تخلیه است که نسبت به گروه‌های دیگر بیشتر است. در اصل ۱۸۰ بار عمدتاً از نوع شمش روی، به طور متوسط از زنجان به تمامی شهرهای ترکیه ارسال شده است. در بخش مصورسازی داده‌ها از نمودار deviation استفاده کردیم. با تحلیل این نمودار مشخص شد که زمان حمل، هزینه حمل و تا حدودی تناژ بار و وضعیت هوا تأثیرگذاری بیشتری نسبت به سایر فاکتورها در جداسازی نوع حمل و نقل دارد.

خورشه ۱ (بارهای پرهزینه کامیونی) که با تنها ۴ آیتم که از نظر هزینه حمل، تناژ بارگیری به کیلو، تعداد عرضه و تقاضا، ماکزیمم تایم زمانی حمل از روز تقاضا تا تخلیه بیشترین مقدار را داشتند و متفاوت از خوشه‌های دیگر در یک گروه مجزا قرار داده شدند. همه بارهای کامیونی از شهر قزوین به دو شهر آدانای ترکیه و شهر MURATBAY ارسال شده‌اند.

خورشه صفر (کم هزینه ترین بار با بیشترین بار ریلی): بیشترین تعداد بارها (۸۵ درصد) که به طور میانگین از نظر هزینه حمل، تناژ بارگیری به کیلو، تعداد عرضه و تقاضا، ماکزیمم تایم زمانی حمل از روز تقاضا تا تخلیه کمترین مقدار را داشتند. مبدا بیشتر این بارها که عمدتاً از نوع ریلی و ۳۷ درصد کانتینر و مابقی شمش روی هستند از شهر سهران است و به طور میانگین به تمامی شهرهای ترکیه بار ارسال شده است.

خورشه ۲ (بارهای نسبتاً پرهزینه کامیونی در وضعیت هوای نامعادل): بعد از خوشه ۱ بیشترین هزینه حمل، بیشترین تناژ بارگیری بعد از خوشه ۱، و بیشترین تعداد عرضه و تقاضا با

استخراج دانش (قواعد انجمنی)

ابتدا با روش fp-growth و مینیمم اطمینان ۷۰ درصد الگوهای پرتکرار را ایجاد کردیم و براساس الگوهای پرتکرار ۹ قانون انجمنی را روی شاخص‌های کیفی استخراج کردیم که به شرح

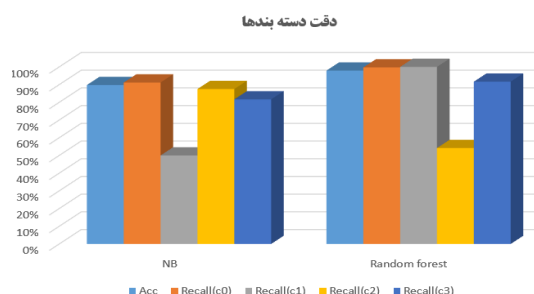
زیر است.

۱. با اطمینان ۸۳٫۵ درصد اگر مقصد بار شهر وان است، آنگاه نوع بار از نوع کانتینر است.

۲. با اطمینان ۸۳,۵ درصد اگر مقصد بار شهر وان است، آنگاه نوع بار از نوع کانتینر و مبدا از شهر سهلان است.
۳. با اطمینان ۸۹,۵ درصد اگر مقصد بار شهر وان است، آنگاه مبدا از شهر سهلان است.
۴. با اطمینان ۹۳,۳ درصد، اگر مقصد بار شهر وان و مبدا از شهر، نوع بار کانتینر است.
۵. با اطمینان ۱۰۰ درصد، اگر مبدا شهر زنجان است، نوع حمل کامیونی است.
۶. با اطمینان ۱۰۰ درصد، اگر نوع بار شمش روی باشد، نوع حمل کامیونی است.
۷. با اطمینان ۱۰۰ درصد، اگر نوع بار کانتینر باشد، مقصد شهر وان است.
۸. با اطمینان ۱۰۰ درصد، اگر نوع بار کانتینر باشد، مقصد شهر وان و مبدا شهر سهلان است.
۹. با اطمینان ۱۰۰ درصد، اگر نوع بار شمش روی باشد، مبدا زنجان و نوع حمل کامیونی است.

استخراج دانش (دسته‌بندی)

نتایج زیر از تحلیل و مقایسه دقت دسته‌بندها در شکل ۲۱ گزارش می‌شود.



شکل ۲۱. دقت دسته‌بندها

۱. اگر هزینه حمل بیشتر از یک میلیارد باشد، بارها در خوشه ۱ (بارهای پرهزینه پرتقاضای کامیونی) هستند.
۲. اگر هزینه حمل کمتر از ۱۷۳ میلیون باشد بار در خوشه بارهای کم هزینه و کم تقاضای ریلی است.
۳. اگر هزینه حمل بیشتر از ۵۱۹ میلیون باشد در خوشه ۲ و کمتر از ۵۱۹ میلیون و بیشتر از ۲۰۰ میلیون در خوشه ۳ است. (تاثیر مستقیم ویژگی هزینه حمل)
۴. گر تناژ بارگیری کمتر از ۳۶۳۱۶۵ کیلو گرم باشد در خوشه صفر و اگر بیشتر از ۳۸۳۳۶۰ کیلو باشد در خوشه ۳ است.
۵. اگر تناژ متقاضی کمتر از ۵۸۷۵۰ کیلو و زمان حمل از روز تقاضا تا روز تخلیه کمتر از ۱۱ روز باشد در خوشه صفر است.
۶. اگر تناژ متقاضی بیشتر از ۲۰۳۰۰۰ کیلو و زمان حمل بیشتر از ۵ روز در خوشه ۲ است.
۷. ولی اگر تناژ بیشتر از ۲۲۸۰۰۰ و زمان حمل بیشتر از ۷ روز باشد در خوشه ۱ است.
۸. اگر تعداد تقاضا کمتر از ۲,۵ باشد در خوشه صفر است.

با دقت در نتایج مشاهده می‌کنیم که دقت روش جمعی جنگل تصادفی که از ۱۰ درخت تصمیم برای پیش بینی استفاده می‌کند، با دقت ۹۱,۹۷ درصد دقت بالاتری از روش احتمالی بیز ساده دارد. از آنجا که نسبت نمونه‌ها در خوشه‌های مختلف بسیار متفاوت است و عدم توازن در خوشه‌ها مشاهده می‌شود. دقت مدل‌ها در خوشه‌ها تفاوت زیادی دارد. در روش بیز ساده روی خوشه ۱ که تنها ۴ آیتم دارد قطعیت پایین ۵۰ درصد دارد. اما در خوشه‌ها دیگر نیز میانگین ۸۱ درصد داریم که خطای ۱۵ درصد بیشتر از جنگل تصادفی است. به دلیل متوازن نبودن داده‌ها از روش تولید قوانین، جعبه سفید و ترکیبی جنگل تصادفی استفاده شد. این روش به دلیل تقسیم مجموعه داده‌ها بین ۱۰ درخت تصمیم موجب توازن داده‌ها شده و دقت را روی تمام کلاس‌ها با مقدار بسیار بالا گزارش داد. ریشه و والد بودن متغیر هزینه حمل در اکثر درخت‌ها نشان دهنده تاثیرگذاری این متغیر در خوشه بندی است. تعدادی از قوانین تولیدی و مفید درختان تصمیم و جنگل تصادفی به شرح ادامه است.

۵- نتیجه گیری

روش جمعی جنگل تصادفی) روی خروجی خوشه بندی استفاده شد. نتایج نشان داد جنگل تصادفی با دقت کل ۹۷,۹۱ درصد نتایج بسیار دقیق تری از بیزساده برای دسته بندی خوشه ها دارد. این روش هم چنین اهمیت الگوریتم های جمعی (رای گیری) را جهت بهبود دقت روی داده های نامتوزان اثبات کرد. با درختان جنگل تصادفی، ۸ قواعد مفید با استخراج شد. نتایج قوانین اثبات کرد که در والد بیشتر درختان هزینه حمل و تناژ بارگیری معیار مقایسه خوشه ها و ساخت درخت است که اهمیت این فاکتورها را در گروه بندی سیستم حمل و نقل جاده ای وریلی اثبات کرد. به طور کلی اگر محققان در آینده بخواهند توسعه ای بر نتایج حاصل این پژوهش پیشنهاد دهند، می توانند موارد زیر را مدنظر قرار دهند: (۱) استفاده از روش های دیگر خوشه بندی مانند som و خوشه بندی فازی. (۲) استفاده از روش های بهینه سازی تعداد خوشه ها مانند الگوریتم فرابتنکاری ژنتیک و ازدحام ذرات.

در این مقاله جهت بررسی فاکتورهای تاثیرگذار در هزینه و زمان حمل جاده و ریلی و استخراج دانش از سه تکنیک داده کاوی: خوشه بندی، قواعد انجمنی و دسته بندی استفاده کردیم. در گام اول با روش خوشه بندی مبتنی بر فاصله کامیانگین ۱۳۸۸ تراکنش حمل و نقل جاده و ریلی را به ۴ خوشه بهینه تقسیم کردیم. نتایج نشان داد با اختلاف، ویژگی های هزینه حمل، تناژ و تعداد عرضه و تقاضا تاثیرگذارترین ویژگی ها در جداسازی خوشه ها هستند. درصد بسیار کمی آب و هوا و نوع حمل و نوع بار تاثیر در گروه بندی بارهای حمل و نقل داشت. در خوشه صفر و یک هزینه حمل و مقدار عرضه و تقاضا تفاوت شاخص با گروه های دیگر داشت. در خوشه ۲ وضعیت هوای نامتعادل و در خوشه ۳ فاصله زمانی از روز تقاضا تا روز تخلیه موجب تفاوت خوشه ای با دیگر خوشه ها شده بود. در فاز قواعد انجمنی، ۹ قواعد با قطعیت بالای ۸۵ درصد گزارش روی متغیرهای کیفی نوع بار و مبدا و مقصد گزارش شد. در گام دسته بندی از دو روش بیز ساده (روش احتمالی و

۶- مراجع

- اسماعیلی، م، (۱۳۹۱). مفاهیم و تکنیک های داده کاوی، انتشارات نیاز دانش.
- ع. م. احمدوند. ب. مینایی بیدگلی. ا. اخوندزاده. (۱۳۸۸). تحلیل رضایتمندی شهروندان با استفاده از تکنیک های داده کاوی: مورد کاوی شهرداری تهران. سومین کنفرانس داده کاوی.
- رستگار، ن.، علیزاده، س.، غضنفری، م.، وزیر، ف. (۱۳۸۷). کاربرد روش خوشه بندی در انتخاب استراتژی های صادرات کالاها در ایران. دومین کنفرانس داده کاوی ایران.
- محمد بیگی، زهرا (۱۳۹۳). کاربرد RFID به منظور کنترل و کاهش تصادفات درون شهری با رویکرد داده کاوی.
- کیانیان، علیرضا (۱۳۹۵). ارزیابی و بهبود کیفیت خدمات حمل و نقل عمومی با استفاده از رویکرد تلفیقی داده کاوی و توسعه عملکرد کیفیت (مطالعه موردی: اتوبوسرانی مشهد).
- Olson, D. (2003). Managerial Issues of Enterprise Resource Planning. McGraw-Hill/Irwin, 1 edition, 45.
- Adams S.O, Akano R.O, Asemota O, J. (2011). Forecasting Electricity Generation in Nigeria using Univariate Time Series Models. *European journal of*
- Scientific Reaserch*. Vol.58. No.1. 30-37.
- Albert Y. Chen, Ting-Yi Yu. (2016). Network based temporary facility location for the Emergency Medical Services considering the disaster induced demand and the transportation infrastructure in disaster response. *Transportation Research Part B. Approaches and Big Data Analytics in Smart Transportation System*.
- C. S. Li and M.-C. Chen (2014). A data mining-based approach for travel time prediction in freeway with non-recurrent congestion. *Neurocomputing*, Vol. 133, 74-83.
- David J. HAND, (2001). Data Mining: Statistics and More? *The American Statistician*, Vol. 52, No. 2. 112-113
- Fayyad.U, Piatetsky-Shapiro.G, Padhraic.S, (2008). From Data Mining to Knowledge Discovery in Databases.
- Feyza Gürbüz, Fatma Turnab (2018). Rule extraction for tram faults via data mining for safe. *Transportation Research Part A*.
- Han.J, Kamber.M, (2011). Data Mining: Concepts and Techniques. San Diego Academic Press, 25-26.
- Hana Ševčíková, Adrian E. Raftery, Paul A. Waddell (2011). Uncertain benefits: Application of

- Soyoung Jung, Xiao Qin, Cheol Oh. (2016). Jung, S., Qin, X., & Oh, C. (2016). Improving strategic policies for pedestrian safety enhancement using classification tree modeling. *Transportation Research Part A: Policy and Practice*, *85*, 53–64.
- Venkata Siva Reddy and R. Vasanth Kumar Mehta. (2019). Study on Computational Intelligence
- Wirth, R. and J. Hipp. CRISP-DM, (2000). Towards a standard process model for data mining. in Proceedings of the 4th international conference on the practical applications of knowledge discovery and data mining. *Citeseer*.
- Yuxuan Ji, Nikolas Geroliminis (2012). On the spatial partitioning of urban transportation networks. *Transportation Research Part B*.
- Zhang G Peter, Min Qi. (2005). Neural network forecasting for seasonal and trend time series. *European Journal of Operational Research*. Vol. 160. 501-515
- Zhang G Peter. (2003). Time series forecasting using a hybrid ARIMA and neural network model. *Neurocomputing* 50. 159-175.
- Bayesian melding to the Alaskan. *Transportation Research Part A*.
- Herbert A. E. (1996). Introduction to Data Mining and Knowledge Discovery. *Two Crows Corporation*, Third Edition. 5-10.
- Hu R. (2010). Medical data mining based on association rules. *Computer and Information Science*. 3(4): 104-8.
- Jithin Raj, Hareesh Bahuleyan, Lelitha Devi Vanajakshi. (2016). Application of data mining techniques for traffic density estimation and prediction.
- Kremic, E. and A. Subasi (2015). Performance of Random Forest and SVM in Face Recognition. *The International Arab Journal of Information Technology*.
- Lathia.N, Froehlich.J, Capra.L, (2010). Mining Public Transport Usage for Personalised Intelligent Transport Systems. *IEEE ICDM*.
- Mehrdad Kargari, Mohammad Mehdi Sepehri. (2012). Stores clustering using a data mining approach for distributing automotive spare-parts to reduce transportation costs.
- Milad Ghasri (2017). Taha Hossein Rashidi, S. Travis Waller. Developing a disaggregate travel demand system of models using. *Transportation Research Part A*.
- Ramageri M. (2011). Data mining techniques and applications. *Indian Journal of Computer Science and Engineering*. 4(1): 301-5.

Determining Factors Affecting the Cost and Time of Road and Rail Transportation Using Association Rules

*Mahdi Yousefi Nejad Attari, Department of Industrial Engineering, Bon. C.,
Islamic Azad University, Bonab, Iran.*

*Ensiyeh Neishabouri Jami, Department of industrial Engineering, University of Bonab,
Bonab, Iran.*

E-mail: 1689597631@iaau.ir

Received: March 2026- Accepted: August 2026

ABSTRACT

Transportation or movement refers to the movement of people, animals, or goods from one place to another. In other words, the act of transportation is defined as a specific movement from point A to point B. Types of transportation methods include air, land (rail and road), sea, cable, pipeline, and space. In land transportation, due to the importance of this business and the competition between organizations active in this field, the use of new technology and technologies in management and making better decisions can be very useful. In this article, we will first introduce data mining and its techniques and data mining software. Then, considering a table of export shipments made in the past one year and six months for an international transportation company, we prepared and extracted association rules in two models using Rapid Miner software and evaluated the relevant rules. The results presented show the factors affecting time and cost and finally, future suggestions are expressed.

Keywords: Data Mining, Time, Road and Rail Transportation, Clustering, Cost