

خوشه‌بندی محورهای برون‌شهری بر اساس الگوهای تردد، سرعت، ترکیب ناوگان و تخلفات رانندگی با استفاده از داده‌های ترددشماری (مطالعه موردی: استان خراسان رضوی، سال ۱۴۰۴)

مقاله علمی-پژوهشی

سحر علیشاهی، دانش آموخته کارشناسی ارشد، گروه مهندسی عمران، دانشکده مهندسی، دانشگاه فردوسی، مشهد، ایران
*سید علی ضیائی (نویسنده مسئول)، دانشیار، گروه مهندسی عمران، دانشکده مهندسی، دانشگاه فردوسی مشهد، مشهد، ایران
حسین محمدی، دانش آموخته کارشناسی ارشد، گروه مهندسی عمران، دانشکده مهندسی، دانشگاه فردوسی، مشهد، ایران
*پست الکترونیکی نویسنده مسئول: sa-ziaee@um.ac.ir

دریافت: ۱۴۰۵/۰۱/۲۲ - پذیرش: ۱۴۰۵/۰۵/۰۵

صفحه ۳۵۴-۳۴۳

چکیده

مدیریت ایمنی و بهره‌برداری راه‌های برون‌شهری نیازمند شناخت تفاوت‌های رفتاری محورهای شبکه است. رتبه‌بندی ساده محورهای پرخطر اگرچه برای اولویت‌بندی مقید است، اما قادر به آشکارسازی تیپ‌های رفتاری مشابه از نظر حجم تردد، سرعت، ترکیب ناوگان و الگوی تخلفات نیست. هدف این پژوهش، خوشه‌بندی محورهای برون‌شهری استان خراسان رضوی بر اساس داده‌های تردد شماری سامانه ۱۴۱ در سال ۱۴۰۴ است. در این مطالعه، ۱۳۶ محور یک جبهه پس از کنترل کیفیت داده و تجمیع ویژگی‌های محوری وارد تحلیل شدند. ویژگی‌های مورد استفاده شامل حجم تردد سالانه، نرخ تخلف فاصله غیرمجاز، نرخ تخلف سرعت، نرخ تخلف سبقت، سرعت متوسط وزنی و سهم وسایل نقلیه تجاری/سنگین بود. برای کاهش اثر مقیاس، داده‌ها استاندارد سازی شدند و برای نمایش ساختار داده از تحلیل مؤلفه‌های اصلی استفاده شد. سپس دو رویکرد K -Means و خوشه‌بندی سلسله‌مراتبی Ward برای تعداد خوشه‌های ۲ تا ۷ ارزیابی شدند. بر اساس شاخص $Silhouette$ ، مقدار $k=4$ با ضریب ۰.۲۶۹، بهترین پیکربندی K -Means را ارائه کرد. نتایج چهار تیپ رفتاری شامل «پرریسک با تخلف فاصله بالا»، «پرتردد با ریسک رفتاری متوسط»، «تجاری/سنگین با سرعت کنترل‌شده» و «کم‌ریسک تر و پایدارتر» را نشان داد. خوشه نخست با ۴۷ محور، نرخ میانگین تخلف فاصله ۲۰۵.۵۵ تخلف به ازای هر ۱۰۰۰ وسیله نقلیه و نرخ وزنی ۲۴۴.۹۲ داشت و از دید مدیریت ایمنی در اولویت نخست پایش و مداخله قرار می‌گیرد. نتایج این پژوهش نشان می‌دهد خوشه‌بندی داده محور می‌تواند مکمل مناسبی برای شاخص‌های ریسک و مدل‌های پیش‌بینی باشد و به راه‌داری کمک کند محورهای مشابه را با سیاست‌های هماهنگ‌تر مدیریت کند.

واژه‌های کلیدی: خوشه‌بندی، K -Means، ترددشماری، تخلف فاصله غیرمجاز، ترکیب ناوگان

۱-مقدمه

[et al,2025]. گزارش‌های بین‌المللی ایمنی راه نیز بر استفاده از داده‌های قابل اعتماد، پایش پیوسته عملکرد شبکه و تصمیم‌گیری مبتنی بر شواهد تأکید می‌کنند [Calinski et al,2024]. در چنین فضایی، داده‌های ترددشماری سامانه ۱۴۱ می‌تواند علاوه بر گزارش حجم تردد، برای شناسایی الگوهای پنهان رفتاری، طبقه‌بندی محورهای مشابه و طراحی

ایمنی راه و بهره‌برداری شبکه حمل‌ونقل، به‌ویژه در محورهای برون‌شهری، به شدت وابسته به شناخت تفاوت‌های رفتاری محورهای مختلف است. گزارش جهانی ایمنی راه نشان می‌دهد که تصادفات جاده‌ای همچنان یکی از چالش‌های جدی سلامت عمومی در جهان است و سالانه حدود ۱.۱۹ میلیون نفر در اثر تصادفات جان خود را از دست می‌دهند [Abduljabar

۲- پیشینه تحقیق

۲-۱- خوشه‌بندی در تحلیل ترافیک

K-Means یکی از پرکاربردترین روش‌های خوشه‌بندی است که با کمینه‌سازی مجموع مربعات درون‌خوشه‌ای، داده‌ها را به k گروه تقسیم می‌کند [Davies et al, 1979 and Etemad, 2023]. خوشه‌بندی سلسله‌مراتبی Ward نیز با کمینه‌سازی افزایش واریانس درون‌خوشه‌ای، ساختاری درختی از گروه‌بندی داده‌ها ایجاد می‌کند [Federal Highway Administration Road Safety Fundamentals, 2024]. در کاربردهای ترافیکی، این روش‌ها برای شناسایی الگوهای مشابه جریان، سرعت، تقاضا و وضعیت شبکه استفاده شده‌اند. مطالعات جدید نیز نشان می‌دهند که ترکیب روش‌های خوشه‌بندی با داده‌های پرتعداد ترافیکی می‌تواند برای مدیریت هوشمند ترافیک، تشخیص الگوهای زمانی و برنامه‌ریزی عملیاتی مفید باشد [MacQueen, 1967 and Rousseeuw 1987].

۲-۲- اعتبارسنجی خوشه‌ها

از آنجا که خوشه‌بندی برچسب مرجع ندارد، ارزیابی درونی خوشه‌ها اهمیت زیادی دارد. شاخص Silhouette انسجام درون‌خوشه‌ای و جدایی بین خوشه‌ای را با دامنه نظری -1 تا $+1$ می‌سنجد؛ مقادیر نزدیک به صفر معمولاً نشان‌دهنده همپوشانی خوشه‌هاست و در داده‌های واقعی و چندبعدی حمل‌ونقل انتظار مقادیر بسیار بالا همیشه واقع‌بینانه نیست [Fredriksso et al, 2025 and Federal Highway Administration, 2008]. شاخص Calinski-Harabasz نسبت جدایی بین خوشه‌ای به پراکندگی درون‌خوشه‌ای را نشان می‌دهد و مقدار بالاتر آن مطلوب‌تر است [Ikotun et al, 2023]. شاخص Davies-Bouldin نیز میانگین شباهت نسبی خوشه‌ها را می‌سنجد و مقدار کمتر آن مطلوب‌تر است [Road Safety Annual Report, 2024].

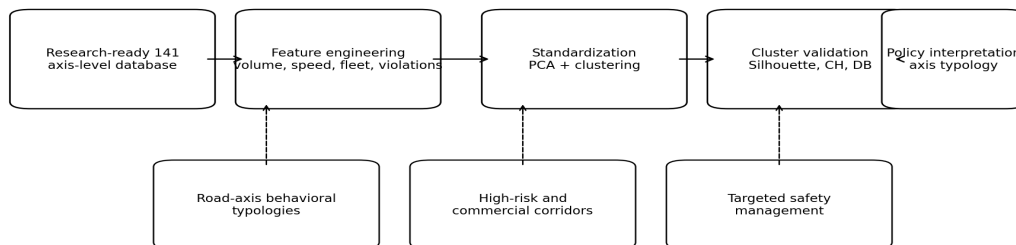
سیاست‌های مداخله‌ای هدفمند استفاده شوند. در مطالعات حمل‌ونقل، خوشه‌بندی یکی از ابزارهای اصلی یادگیری بدون ناظر برای کشف ساختارهای پنهان در داده‌هاست. خوشه‌بندی به پژوهشگر اجازه می‌دهد بدون تعریف قبلی برچسب ریسک، محورهایی را که از نظر حجم، سرعت، ترکیب ناوگان و تخلف رفتاری مشابه هستند در یک گروه قرار دهد. این نوع تحلیل برای شبکه‌هایی مانند خراسان رضوی اهمیت ویژه دارد، زیرا محورهای زیارتی، تجاری، حومه‌ای و بین‌شهری ممکن است از نظر میانگین تردد یا سرعت به تنهایی قابل تفکیک نباشند، اما در ترکیب چندمتغیره رفتارهای متفاوتی نشان دهند. نوآوری مقاله در تبدیل دیتابیس ساعتی/روزانه سامانه ۱۴۱ به یک پایگاه داده محور-محور برای خوشه‌بندی رفتاری است. سؤال اصلی پژوهش این است: آیا می‌توان محورهای برون‌شهری استان خراسان رضوی را بر اساس الگوهای چندبعدی تردد، سرعت، ترکیب ناوگان و تخلفات رانندگی به تیپ‌های رفتاری معنادار تقسیم کرد؟ سؤال فرعی دوم این است که هر خوشه چه دلالتی برای مدیریت ایمنی و اولویت‌بندی مداخلات دارد؟

داده‌ها و آماده‌سازی ویژگی‌ها

داده‌های این پژوهش از دیتابیس نهایی پژوهشی سامانه ۱۴۱ برای سال ۱۴۰۴ استخراج شد. واحد تحلیل، محور یک‌جهته است؛ بنابراین هر کد شش‌رقمی محور به عنوان یک جهت حرکت مستقل در نظر گرفته شد. این انتخاب با تعریف رسمی ترددشماری سازگار است و از ادغام نادرست رفت و برگشت جلوگیری می‌کند. پس از کنترل کیفیت داده‌های ساعتی، شاخص‌های محوری در سطح سالانه محاسبه شدند. تخلف فاصله غیرمجاز بر اساس تعریف سامانه، عدم رعایت حداقل فاصله زمانی دو ثانیه‌ای با وسیله نقلیه جلویی است. از آنجا که تعداد خام تخلف تحت تأثیر حجم تردد قرار دارد، در این پژوهش نرخ تخلف فاصله، سرعت و سبقت به ازای هر ۱۰۰۰ وسیله نقلیه محاسبه شد. همچنین برای کنترل اثر مقیاس محور، لگاریتم حجم تردد سالانه در کنار متغیرهای نرخ و سهم ناوگان استفاده شد.

نقش در مقاله	متغیر/روش	گروه ویژگی
نمایش اندازه و شدت تقاضای محور	log_total_vehicles, total_vehicles	حجم تردد
تمایز محورهای آزادجریان و کنترل شده	avg_speed_weighted	رفتار سرعتی
بازنمایی رفتار پرریسک رانندگان	unsafe_distance_rate_per_۱۰۰۰، speeding_rate_per_۱۰۰۰، overtaking_rate_per_۱۰۰۰	تخلفات رفتاری
شناخت محورهای تجاری/سنگین	commercial_vehicle_share_class_۲_۵	ترکیب ناوگان
حذف اثر مقیاس و نمایش دوبعدی	StandardScaler, PCA	آماده سازی آماری
انتخاب تعداد خوشه و مقایسه الگوریتم ها	Silhouette, Calinski-Harabasz، Davies-Bouldin	اعتبارسنجی خوشه

جدول ۱. ویژگی های مورد استفاده در خوشه بندی محورهای برون شهری



شکل ۱. چارچوب روش شناسی مقاله خوشه بندی

۳-روش شناسی

برای نمایش تصویری ساختار کلی داده مناسب است. در انتخاب تعداد خوشه، تمرکز اصلی بر Silhouette بود، اما Calinski-Harabasz و Davies-Bouldin نیز برای کنترل پایداری تفسیر گزارش شدند. با توجه به ماهیت واقعی داده های راه، همپوشانی رفتاری محورهای طبیعی است؛ بنابراین انتخاب k صرفاً بر اساس بیشینه عددی شاخص ها انجام نشد، بلکه تفسیرپذیری مهندسی خوشه ها نیز معیار نهایی بود.

فرایند روش شناسی شامل پنج مرحله بود: (۱) استخراج شاخص های محوری از دیتابیس ۲ (research-ready)، استانداردسازی متغیرها برای حذف اثر مقیاس، (۳) اجرای K-Means و Agglomerative-Ward برای $k=2$ تا ۷، (۴) انتخاب تعداد خوشه با استفاده از شاخص های درونی، و (۵) تفسیر خوشه ها بر اساس پروفایل رفتاری آن ها. برای نمایش دوبعدی ساختار داده، تحلیل مؤلفه های اصلی نیز به کار رفت. دو مؤلفه نخست PCA در مجموع ۶۴،۶۷ درصد واریانس ویژگی ها را توضیح دادند که

جدول ۲. ارزیابی الگوریتم‌ها و تعداد خوشه‌های مختلف

الگوریتم	تعداد خوشه	Silhouette	Calinski-Harabasz	Davies-Bouldin
KMeans	۲	۰,۲۶۶	۴۲,۸۱۵	۱,۶۱۰
KMeans	۳	۰,۲۴۲	۴۳,۰۳۰	۱,۳۸۰
KMeans	۴	۰,۲۶۹	۴۴,۳۹۷	۱,۲۱۵
KMeans	۵	۰,۲۵۱	۴۱,۹۹۰	۱,۲۵۱
KMeans	۶	۰,۲۵۵	۴۲,۴۸۳	۱,۱۷۶
KMeans	۷	۰,۲۵۵	۴۱,۰۷۶	۱,۰۸۹
Agglomerative-Ward	۲	۰,۲۶۴	۴۲,۴۸۷	۱,۶۱۷
Agglomerative-Ward	۳	۰,۲۰۲	۳۹,۰۹۲	۱,۴۳۰
Agglomerative-Ward	۴	۰,۲۲۴	۳۶,۴۰۳	۱,۵۵۰
Agglomerative-Ward	۵	۰,۲۳۵	۳۶,۱۲۲	۱,۲۸۳
Agglomerative-Ward	۶	۰,۲۲۷	۳۵,۷۹۴	۱,۲۰۷
Agglomerative-Ward	۷	۰,۲۱۴	۳۶,۰۳۳	۱,۲۲۳

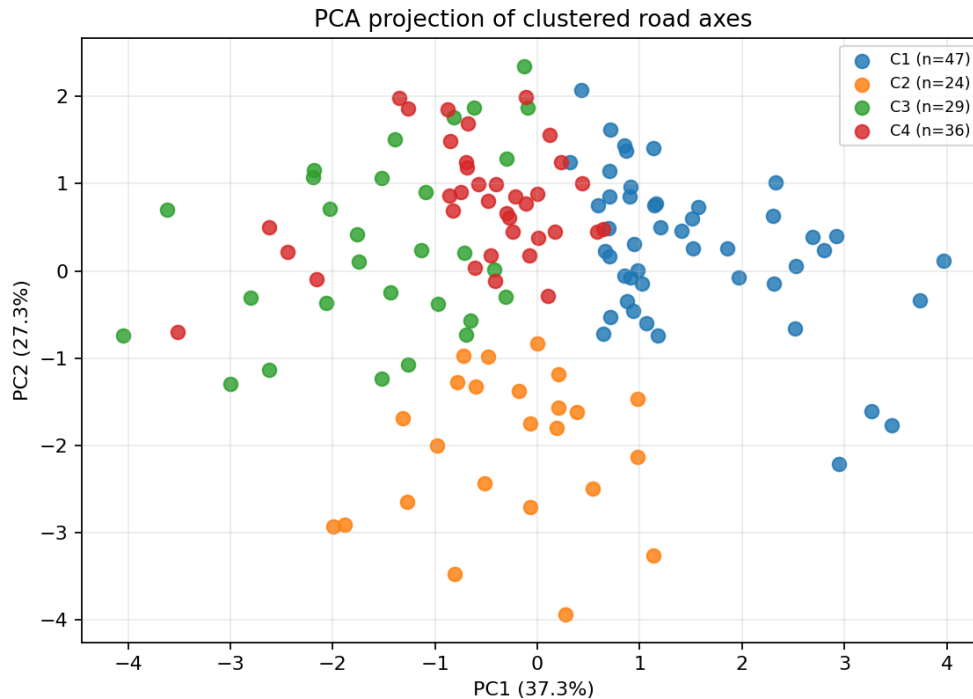


شکل ۲. ارزیابی تعداد خوشه بر اساس شاخص‌های درونی

۴- تعداد خوشه و ساختار دوبعدی

برچسب خوشه‌ها نشان می‌دهد. پراکندگی نقاط تأیید می‌کند که خوشه‌ها کاملاً جدا و مصنوعی نیستند، بلکه تیپ‌هایی از رفتارهای ترافیکی با مرزهای همپوشان در شبکه واقعی هستند؛ این موضوع در تحلیل ترافیکی طبیعی و قابل انتظار است.

نتایج نشان داد بهترین مقدار K-Means بر اساس Silhouette برابر با $k=4$ و مقدار شاخص برابر با ۰,۲۶۹ است. شاخص Davies-Bouldin برای $k=4$ برابر با ۱,۲۱۵ بود که نسبت به بسیاری از حالات دیگر جدایی مناسب‌تری را نشان می‌دهد. شکل ۳ نمایش PCA محورها را همراه با



شکل ۳. نمایش دوبعدی خوشه‌ها با PCA

پروفایل و تفسیر خوشه‌ها

جدول ۳ و شکل‌های ۴ و ۵ پروفایل چهار خوشه نهایی را نشان می‌دهند. خوشه ۱ با ۴۷ محور و نرخ وزنی تخلف فاصله ۲۴۴,۹۲، پرریسک‌ترین خوشه از نظر فاصله‌گیری رانندگان است. خوشه ۲ شامل ۲۴ محور با ریسک رفتاری متوسط و حجم پایین‌تر است. خوشه ۳ با ۲۹ محور، نرخ

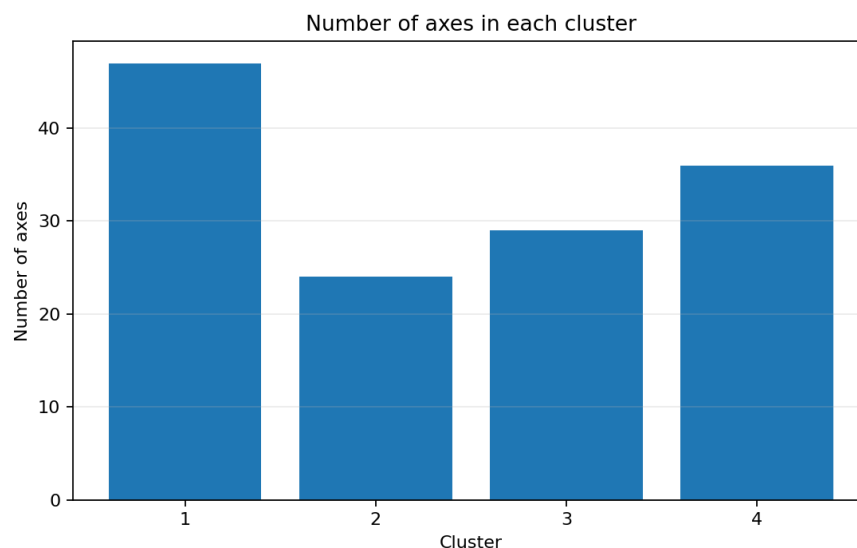
سرعت بسیار بالا و نرخ فاصله کمتر دارد و می‌تواند نماینده محورهای آزادجریان‌تر یا دارای الگوی سرعتی خاص باشد. خوشه ۴ با ۳۶ محور، سهم تجاری بالاتر و نرخ فاصله پایین‌تر دارد و از منظر تخلف فاصله کم‌ریسک‌تر و پایدارتر تفسیر می‌شود.

جدول ۳. خلاصه و تفسیر چهار خوشه نهایی محورهای برون‌شهری

خوشه	برچسب تفسیری	تعداد محور	تردد کل	میانگین نرخ فاصله/۱۰۰۰	نرخ وزنی فاصله/۱۰۰۰	میانگین نرخ سرعت/۱۰۰۰	میانگین سرعت	سهم تجاری
۱	پرریسک با تخلف فاصله بالا	۴۷	۲۹۹,۶۸۳,۰۶۲	۲۰۵,۵۵۰	۲۴۴,۹۱۸	۱۰۸,۷۹۸	۸۶,۳۵۱	۰,۱۱۸
۲	پرتردد با ریسک رفتاری متوسط	۲۴	۲۹,۳۸۹,۴۴۸	۱۵۹,۷۰۹	۱۷۶,۶۲۱	۸۳,۷۰۹	۶۷,۲۲۸	۰,۱۱۵
۳	تجاری/سنگین با سرعت کنترل‌شده	۲۹	۳۲,۶۰۰,۸۰۶	۸۷,۱۰۰	۱۱۵,۳۸۷	۳۷۰,۹۶۶	۸۷,۲۱۴	۰,۱۱۴
۴	کم‌ریسک‌تر و پایدارتر	۳۶	۶۵,۸۶۷,۲۴۴	۶۲,۹۶۳	۷۳,۲۸۸	۱۳۰,۶۴۹	۸۸,۵۲۳	۰,۲۶۵



شکل ۴. پروفایل استانداردشده خوشه‌ها بر اساس ویژگی‌های اصلی



شکل ۵. تعداد محورهای عضو هر خوشه

محورهای شاخص در هر خوشه

یک سیاست یکسان برای همه محورها، می‌توان برای هر خوشه بسته مداخله‌ای متفاوت طراحی کرد: در خوشه پریسک، تابلوهای هشدار فاصله، کنترل هوشمند فاصله و اطلاع‌رسانی ساعات پرتخلف؛ در خوشه سرعتی، کنترل سرعت و بازنگری حد سرعت مجاز؛ در خوشه تجاری، پایش حمل‌ونقل سنگین و مدیریت سبقت؛ و در خوشه کم‌ریسک، پایش نگهدارنده و استفاده به عنوان الگوی مقایسه.

برای کمک به تفسیر میدانی خوشه‌ها، سه محور دارای بیشترین نرخ تخلف فاصله در هر خوشه در جدول ۴ ارائه شده است. در خوشه پریسک، محورهای مرتبط با مشهد-چناران، فریمان-مشهد و چناران-مشهد در صدر قرار دارند. این الگو می‌تواند ناشی از حجم تقاضا، جریان‌های حومه‌ای-زیارتی، نزدیکی به مراکز جذب سفر و فشردگی جریان باشد. البته برای نتیجه‌گیری قطعی، اتصال داده‌های هندسی، حد سرعت مجاز و داده تصادفات ضروری است. نتایج خوشه‌بندی نشان می‌دهد که محورهای شبکه خراسان رضوی را نمی‌توان صرفاً بر اساس حجم تردد یا سرعت متوسط طبقه‌بندی کرد.

برخی محورهای پرتردد دارای نرخ بالای تخلف فاصله هستند، در حالی که برخی محورهای دارای سهم بالاتر وسایل تجاری، نرخ فاصله پایین‌تری نشان می‌دهند. این یافته می‌تواند بیانگر تفاوت در ماهیت جریان، هدف سفر، ترکیب تقاضا و رفتار رانندگان باشد. از دیدگاه مهندسی ترافیک، خوشه ۱ باید در اولویت کنترل و پایش فاصله‌گیری قرار گیرد؛ خوشه ۳ نیازمند بررسی رفتار سرعتی و کنترل سرعت است؛ و خوشه ۴ می‌تواند به عنوان گروه مرجع برای مقایسه محورها استفاده شود. کاربرد مدیریتی این مقاله در طراحی بسته‌های سیاستی مبتنی بر تیپ محور است. به جای اعمال

جدول ۴. سه محور شاخص در هر خوشه بر اساس نرخ تخلف فاصله غیرمجاز

خوشه	برچسب خوشه	کد محور	نام محور	نرخ فاصله/۱۰۰۰	نرخ سرعت/۱۰۰۰	تردد کل
۱	پرریسک با تخلف فاصله بالا	۳۱۳۱۰۳	مشهد - چناران (مشهد - تقاطع آزادراه)	۴۷۱,۸۵	۱۴۲,۵۵	۲۰,۱۸۰,۸۲۱
۱	پرریسک با تخلف فاصله بالا	۳۱۳۳۵۲	فریمان - مشهد (عوارضی فریمان)	۴۴۶,۴۹	۱۸,۹۵	۳,۵۲۸,۳۸۱
۱	پرریسک با تخلف فاصله بالا	۳۱۳۱۵۳	چناران - مشهد (تقاطع آزادراه - مشهد)	۴۰۴,۹۴	۳۶,۲۰	۱۷,۰۹۴,۴۲۲
۲	پرتدد با ریسک رفتاری متوسط	۳۱۴۳۵۱	بردسکن - سبزوار (سبزوار)	۲۵۶,۲۶	۱۱۶,۷۵	۱,۶۹۱,۹۲۳
۲	پرتدد با ریسک رفتاری متوسط	۳۱۳۷۵۱	تایباد - تربت جام (تربت جام)	۲۴۹,۳۱	۱۴,۰۳	۲,۱۷۶,۴۱۴
۲	پرتدد با ریسک رفتاری متوسط	۳۱۴۳۵۴	اسفراین - سبزوار (سبزوار)	۲۴۸,۸۶	۲۱,۲۸	۱,۶۳۲,۵۰۶
۳	تجاری/سنگین با سرعت کنترل شده	۳۱۵۵۰۲	گناباد - قائن (گناباد)	۱۹۳,۹۶	۳۵۹,۳۵	۱,۴۱۶,۸۶۰
۳	تجاری/سنگین با سرعت کنترل شده	۳۱۳۸۵۲	تربت حیدریه - مشهد (رابط سنگ - باغچه)	۱۵۴,۹۰	۵۸۵,۹۰	۵,۷۳۴,۳۷۵
۳	تجاری/سنگین با سرعت کنترل شده	۳۱۳۲۷۴	مشهد - تربت حیدریه (کمربندی تربت حیدریه)	۱۴۸,۳۶	۴۳۶,۲۸	۳,۶۸۴,۳۷۷
۴	کم ریسک تر و پایدارتر	۳۱۵۵۵۲	قائن - گناباد (گناباد)	۲۴۴,۰۷	۱۹۶,۴۲	۱,۴۴۴,۰۹۰
۴	کم ریسک تر و پایدارتر	۳۱۴۳۷۰	نیشابور - سبزوار (زعفرانیه)	۱۱۵,۸۲	۹۲,۰۷	۳,۴۷۴,۴۵۰
۴	کم ریسک تر و پایدارتر	۳۱۵۷۵۱	سبزوار - نیشابور (نیشابور)	۱۱۳,۴۴	۲۱,۷۱	۳,۵۲۲,۱۴۱

۵- نتیجه گیری

این پژوهش یک چارچوب داده محور برای خوشه بندی محورهای برون شهری استان خراسان رضوی با استفاده از داده های سامانه ۱۴۱ ارائه کرد. نتایج نشان داد چهار تیپ رفتاری قابل تفسیر در شبکه وجود دارد و خوشه بندی می تواند مکمل مناسبی برای مدل های پیش بینی تخلف و شاخص های ریسک باشد. مهم ترین نتیجه کاربردی پژوهش آن است که محورهای یک استان از نظر رفتار ترافیکی و تخلفات همگن نیستند و برنامه ریزی ایمنی باید بر اساس تیپ رفتاری محور انجام شود. بر اساس یافته ها، خوشه پرریسک با تخلف فاصله بالا، اولویت نخست برای بررسی میدانی و مداخلات مدیریتی است.

پیشنهاد های اجرایی

پیشنهاد های اجرایی عبارت اند از: ۱) به روز رسانی ماهانه خوشه بندی محور ها برای پایش تغییر رفتار شبکه؛ ۲) طراحی داشبورد مدیریتی برای نمایش خوشه هر محور؛ ۳) تعریف برنامه کنترل فاصله برای خوشه پرریسک؛ ۴) اتصال نتایج خوشه بندی به داده های تصادفات برای اعتبار سنجی ایمنی؛ ۵) استفاده از خوشه ها برای تخصیص بهینه منابع راهداری و پلیس راه؛ و ۶) تعریف سناریوهای مداخله متفاوت برای هر خوشه.

۶- مراجع

- Fredriksson, H, (2025). Exploring spatio-temporal traffic performance variation in road networks. *Transportation Research\Preprint Report*.
- Ikotun, A. M., Ezugwu, A. E., Abualigah, L., Abuhaija, B., & Heming, J,K (2023). means clustering algorithms: A comprehensive review, variants analysis, and advances in the era of big data. *Information Sciences* 622, 178-220.
- International Transport Forum / OECD, Road Safety Annual Report (2024).
- Jain, A. K, :Data clustering 50 years beyond K-means. *Pattern Recognition Letters*, 31(18), 651-666.
- Jolliffe, I. T. .Principal Component Analysis (2002). *Springer*;
- Kwon, Y., Lee, M., Lee, M. J., & Son, S.-W, Cluster formations of free and congested flows in urban road networks. *arXiv*, 240808122.
- Lloyd, S,Least squares quantization in PCM. ,*IEEE Transactions on Information Theory*28(2), 129-137.
- MacQueen, J, (1967). Some methods for classification and analysis of multivariate observations. *Proceedings of the Fifth Berkeley Symposium*,
- Abduljabbar, R., et al, (2025). Machine learning traffic flow prediction models for smart cities: Review and applications. *Infrastructures*, 7-10.
- Caliński, T., & Harabasz, J (2024). A dendrite method for cluster analysis. *Communications in Statistics*3(1), 1-27.
- Chai, A. B. Z., et al, (2024). Enhancing road safety with machine learning: Current state and future directions. *Engineering Applications of Artificial Intelligence*.
- Clustering Algorithms to Analyse Smart City Traffic Data, (2024). *International Journal of Advanced Computer Science and Applications*, (15)8.
- Davies, D. L., & Bouldin, D. W. A. (1979). cluster separation measure. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, PAMI 2(1), 224-227.
- Etemad, S., et al, (2023).Clustering of urban traffic patterns by K-Means and Dynamic Time Warping: Case study. *arXiv 230909830*.
- Federal Highway Administration,Road Safety Fundamentals: *Measuring Safety*(2024).
- Federal Highway Administration (2008). Surrogate Safety Assessment Model and Validation Final Report.

- Sohail, A. M. (2024). Unravelling traffic dynamics with K-means clustering and data preprocessing. *Information Systems and Smart Cities*.
- Ward, J. H., (2022). Hierarchical grouping to optimize an objective function. *Journal of the American Statistical Association*.
- World Health Organization, Global status report on road safety (2023).
- Xu, D., & Tian, Y (2015). A comprehensive survey of clustering algorithms. *Annals of Data Science* (2), 165-193.
- Yumak, A. (2025). A machine learning approach to identify high-risk road segments. *Applied Sciences* (15) 25, 1224-1225.
- Morissette, L., & Chartier, S. (2013). The k-means clustering technique: General considerations and implementation in Mathematica. *Tutorials in Quantitative Methods for Psychology*, (9) 1, 15-24.
- Pavlyshyn, V., et al, (2025). An adaptive machine learning approach to traffic pattern recognition. *Smart Cities*, (5) 4, 151-152.
- Pedregosa, F., et al, (2011). Scikit-learn: Machine learning in Python. *Journal of Machine Learning Research*, (12), 2825-2830.
- Rousseeuw, P. J., Silhouettes, (1987). A graphical aid to the interpretation and validation of cluster analysis. *Journal of Computational and Applied Mathematics*, 20-53-65.

Clustering of Intercity Axes Based on Traffic Patterns, Speed, Fleet Composition, and Driving Violations Using Traffic Count Data (Case Study: Razavi Khorasan Province, 2025)

*Sahar Alishahi, M.Sc., Grad., Department of Civil Engineering,
Faculty of Engineering, Ferdowsi University of Mashhad, Iran.*

*Seyed Ali Ziaee, Associate Professor, Department of Civil Engineering,
Faculty of Engineering, Ferdowsi University of Mashhad, Iran.*

*Hosein Mohammadi, M.Sc., Grad., Department of Civil Engineering,
Faculty of Engineering, Ferdowsi University of Mashhad, Iran.*

E-mail: sa-ziaee@um.ac.ir

Received: March 2026- Accepted: August 2026

ABSTRACT

Safety management and operation of intercity roads require understanding the behavioral differences among network corridors. Simple ranking of high-risk corridors, while useful for prioritization, fails to reveal similar behavioral patterns in terms of traffic volume, speed, fleet composition, and violation profiles. This study aims to cluster the intercity corridors of Razavi Khorasan Province based on traffic count data from the 141 system in the year 1404 (2025–2026). In this study, 136 unidirectional corridors were included in the analysis after data quality control and aggregation of corridor-level features. The features used included annual traffic volume, following-distance violation rate, speed violation rate, overtaking violation rate, weighted average speed, and the share of commercial/heavy vehicles. To reduce scaling effects, the data were standardized, and Principal Component Analysis was employed to visualize the data structure. Subsequently, two approaches—K-Means and Ward's hierarchical clustering—were evaluated for cluster counts ranging from 2 to 7. Based on the Silhouette index, $k=4$ with a coefficient of 0.269 provided the best K-Means configuration. The results revealed four behavioral types: "high-risk with high following-distance violations," "high-traffic with moderate behavioral risk," "commercial/heavy with controlled speed," and "lower-risk and more stable." The first cluster, comprising 47 corridors, had an average following-distance violation rate of 205.55 violations per 1,000 vehicles and a weighted rate of 244.92, placing it as the top priority for safety monitoring and intervention. The findings of this study indicate that data-driven clustering can serve as a suitable complement to risk indices and predictive models, assisting road authorities in managing similar corridors through more coordinated policies.

Keywords: Traffic Clustering, K-Means, Traffic Counting, Unsafe Following Distance Violation, Fleet Composition