

طراحی یک مدل یادگیری تمیزدهنده جهت بهبود مسیرهای اتوبوس در شبکه‌ی حمل و نقل بهینه

محسن جهانشاهی^{۱*}، مجید قلی‌پور^۲، روزبه ابرازی^۳، مهرداد الماسی^۴

۱- دانشیار، دانشگاه آزاد اسلامی واحد قزوین، گروه مهندسی کامپیوتر و فناوری اطلاعات، قزوین، ایران

۲- دانشجوی دکتری، دانشگاه آزاد اسلامی واحد قزوین، گروه مهندسی کامپیوتر و فناوری اطلاعات، قزوین، ایران

۳- دانشجوی دکتری، دانشگاه صنعتی امیرکبیر، گروه ریاضی و علوم کامپیوتر، تهران، ایران

۴- دانشجوی دکتری، دانشگاه تربیت مدرس، گروه مهندسی برق و کامپیوتر، تهران، ایران

رسید مقاله: ۲۸ دی ۱۳۹۴

پذیرش مقاله: ۴ خرداد ۱۳۹۵

چکیده

طراحی شبکه‌ی حمل و نقل اتوبوس (BTN) یکی از مهم‌ترین مباحث در مدیریت شهری است. پارامترهای تاثیرگذار زیادی در این طراحی، موثر می‌باشند. پارامترهایی که در مجموع منجر به برآورده شدن مجموعه‌ای از اهداف مدیریت شهری می‌شود. بهبود دسترسی پذیری شهروندان، پوشش مساحت بیش‌تری از سطح شهر، کاهش زمان انتظار و هزینه و همچنین کاستن از تعداد تعویض خطوط اتوبوس برای رسیدن به مقصد یک مسافر، از جمله‌ی این اهداف است. طراحی یک BTN یک مساله‌ی NP-hard می‌باشد و بنابراین رسیدن به یک پاسخ بهینه در طراحی‌های با ابعاد بالا کاری دشوار است. راه‌حل معمول در طراحی یک BTN به این صورت می‌باشد: کاهش فضای جستجوی ممکن در ابتدا و سپس ساخت شبکه بر اساس اولویت‌های مدیریت شهری. در این مقاله یک روش جدید برای ارتقای طراحی یک BTN ارائه می‌شود که مبتنی بر یادگیری آماری می‌باشد. این مدل به کمک روش‌های یادگیری آماری و ترکیب آن‌ها با یکدیگر تولید می‌شود. در این تحقیق دانش متخصصان انسانی از شبکه‌ی BTN فعلی استخراج می‌شود، سپس این دانش برای کوچک کردن فضای جستجوی طراحی یک BTN به محدوده‌ای کوچک از معابر به کار گرفته می‌شود. این معابر ویژگی‌های لازم برای شرکت در BTN را دارند و می‌توانند برای مساله‌ی طراحی شبکه‌ی اتوبوس رانی BTNDP یا توسعه‌ی BTN فعلی به کار گرفته شوند. در این مقاله از Naïve Bayesian و دو روش دیگر رگرسیون پایه و ورژن ترکیبی آن‌ها برای تولید مدل خود بهره گرفته شده است. ارزیابی مدل تولیدی بر اساس دقت، False positive و True positive صورت می‌گیرد. مقادیر به دست آمده از این معیارها قابل اعتماد بودن روش پیشنهادی را تصدیق می‌کند. دیتاست مورد استفاده در این مقاله، شامل اطلاعات شبکه‌ی اتوبوس رانی شهر تهران می‌باشد.

کلمات کلیدی: رگرسیون، Naïve Bayesian، شبکه‌ی حمل و نقل اتوبوس (BTN)، مساله‌ی طراحی شبکه‌ی حمل و نقل اتوبوس (BTNDP).

* عهده‌دار مکاتبات

آدرس الکترونیکی: mjahanshahi@iauctb.ac.ir

۱ مقدمه

بهبود سیستم‌های حمل و نقل عمومی یکی از اهداف اصلی در مدیریت شهری می‌باشد. این بهبود می‌تواند باعث کاستن از چالش‌ها در حوزه‌های انرژی، اقتصاد و محیط شود. طراحی BTN یکی از رهیافت‌ها برای رسیدن به این هدف است. فضای جستجوی بزرگی به اندازه‌ی تمام معابر موجود در یک شهر، در طراحی BTN مطرح است. بزرگ بودن این فضا در آبر شهرها کار طراحی را بغرنج‌تر می‌کند. مساله‌ی طراحی شبکه‌ی حمل و نقل اتوبوس که از این به بعد به اختصار آن را BTNDP می‌نامیم و ایجاد شبکه‌ی ایده آل مساله‌ای NP-hard است. افزون بر این باید در نظر گرفت که یک BTN ضعیف مسافران را برای استفاده از شبکه ترغیب نمی‌کند و سودی را نصیب مجریان BTN نخواهد کرد.

جنبه‌های مختلفی برای ساخت یک شبکه‌ی کارا باید مورد توجه قرار گیرد. مواردی چون پوشش مراکز خرید و مراکز خدمات درمانی، کاهش زمان سفر مسافران، مسائل زیست محیطی، انتخاب معابری با پهنای مناسب، راحتی کاربران و اجتناب از ترافیک از جمله‌ی این جنبه‌ها هستند. در این تحقیق BTN موجود به عنوان یک شبکه‌ی مطلوب - شبکه‌ای که توسط متخصصان ماهر طراحی شده است - در نظر گرفته می‌شود، سپس روش جدیدی برای حل BTNDP ارائه می‌شود. این کار از طریق استخراج دانش متخصصانی که BTN فعلی را طراحی کرده‌اند، صورت می‌گیرد. اندازه‌ی یک BTN بتدریج و با گذشت زمان رشد می‌کند؛ بنابراین یک BTN شامل اطلاعات ارزشمندی از متخصصان مختلف در طول سال‌های مختلف است. استخراج دانش [۱] یک راه جدید برای حل مسائل مدیریت شهری/ترافیک می‌باشد [۲ و ۳]. در این مقاله، یک روش مبتنی بر دانش ارائه می‌شود و از آن به عنوان ابزاری برای کاستن مسیرهای شهری به یک مجموعه‌ی شبه بهینه از معابر مناسب راندن اتوبوس، استفاده می‌شود.

با توجه به اینکه در این مقاله روشی مبتنی بر دانش ارائه می‌شود، معابر انتخاب شده دارای همه‌ی ویژگی‌های مورد نیاز از دیدگاه متخصصان مختلف خواهد بود. هدف اصلی این مقاله، پیشنهاد روش‌های یادگیری است که دانش متخصصان انسانی را از BTN‌های موجود استخراج کند. سازمان دهی باقی این مقاله به این صورت است: بخش ۲ بررسی ادبیات پیشین در حوزه‌ی تحقیقاتی این مقاله و بخش ۳ توضیح چارچوب اصلی فرآیند یادگیری این مقاله می‌باشد. مشخصات داده‌ی مورد استفاده در این مقاله و شرح مدل‌های یادگیری به همراه جزییات آن‌ها به ترتیب در بخش‌های ۴ و ۵ بیان می‌شود. بخش ۶ شامل نتایج و تحلیل می‌باشد. در نهایت در بخش ۷ به نتیجه‌گیری پرداخته شده است.

۲ مروری بر تحقیقات مشابه

تحقیقات بر روی BTNDP را می‌توان به دودسته‌ی رهیافت‌های دقیق و رهیافت‌های مکاشفه‌ای (تقریبی) تقسیم نمود [۴]. روش‌های دقیق بر جنبه‌های ریاضی تمرکز دارند. این دسته از روش‌ها تلاش دارند تا زمان انتظار، انتقال و تعداد دفعات تعویض خطوط مسافران را حداقل کنند. برای به دست آوردن این اهداف روش‌های دقیق از

ویژگی‌هایی چون طول مسیر بهینه، زمان سرباره و فاصله‌ی دسترسی به عنوان ویژگی‌های تصمیم‌گیری استفاده می‌کند [۵]. رهیافت‌های دقیق با دشواری‌هایی روبرو هستند. دشواری‌هایی چون پیچیدگی محاسباتی بالا [۶]، [۷] نیاز به تخصیص زیر مدل [۸] و ماهیت چندهدفه‌ی BTNDP. مدل Guan و همکاران [۹] مثالی از رهیافت‌های دقیق است. این مدل تلاش می‌کند تا تابعی از طول خطوط اتوبوس و فاصله‌ی پیموده شده توسط مسافران برای رسیدن به ایستگاه‌های اتوبوس را حداقل کند. رهیافت‌های مکاشفه‌ای نیز مشابه با رهیافت‌های دقیق هستند با این تفاوت که در این دسته از روش‌ها هزینه و سود اپراتورهای خطوط اتوبوس و مسافران به عنوان معیارهای اصلی در نظر گرفته می‌شوند. در این دسته از روش‌ها دانش متخصصان برای تعیین تابع ارزیابی روش‌ها مورد استفاده قرار می‌گیرد. به دلیل استفاده از دانش متخصصان در روش‌های دسته‌ی دوم، روش پیشنهادی در این مقاله در این دسته قرار می‌گیرد. روش‌های مکاشفه‌ای عموماً مساله‌ی اصلی را به دو بخش تقسیم می‌کنند: (۱) تعیین یک مجموعه از خطوط اتوبوس‌کاندید (۲) محاسبه‌ی هزینه‌ی اجرا. [۱۰-۱۴] از جمله‌ی پژوهش‌هایی هستند که در دسته‌ی روش‌های مکاشفه‌ای قرار می‌گیرند. در این دسته از پژوهش‌ها رهیافت حل دارای دو گام می‌باشد [۱۵]:

۱- انتخاب یک راه‌حل تقریباً بهینه

۲- بهبود راه‌حل تقریباً بهینه‌ی انتخاب شده

از بین دو گام فوق، گام دوم دارای اهمیت بیش‌تری است. بهبود کامل به راه‌حل تقریباً بهینه برای یک شهر کوچک با معیارهای ساده، کاری آسان است. با این حال در شهرهای بزرگ اجرای این گام بسیار پیچیده خواهد بود. روش‌های مکاشفه‌ای (تقریبی) [۱۰-۱۴] عموماً از شبکه‌های عصبی، الگوریتم ژنتیک و یا دیگر الگوریتم‌های فرامکاشفه‌ای برای بهبود راه‌حل‌هایشان بهره می‌گیرند؛ البته در به کارگیری الگوریتم‌های فرامکاشفه‌ای برای حل BTNDP چالش‌هایی نیز وجود دارد [۶، ۸، ۱۶ و ۱۷]:

۱- به دلیل متفاوت بودن جمعیت‌های اولیه، الگوریتم‌های فرامکاشفه‌ای نتایج متفاوتی در اجراهای متفاوت تولید می‌کنند. به علاوه، فضای جستجوی بزرگی وجود دارد و جواب نهایی وابستگی زیادی به مقدار دهی اولیه‌ی پارامترهای الگوریتم دارد.

۲- عموماً جمعیت اولیه توسط متخصص انسانی تعیین می‌شود؛ بنابراین فرد متخصص باید حتماً از تجربه‌ی کافی برای تعیین خطوط‌کاندید اتوبوس رانی برخوردار باشد.

در ادامه‌ی این بخش تعدادی از رهیافت‌هایی که برای BTNDP به کار گرفته شده‌اند مورد بررسی قرار می‌گیرند. یانگ‌نگ و همکاران [۱۸] از تفکیک‌های تخمینی تسلسلی، برای حل BTNDP بهره گرفته‌اند. دیتاست به کار گرفته شده در پژوهش [۱۸] حاوی اطلاعات شبکه‌ی حمل و نقل بین شهری می‌باشد. محققان این پژوهش شبکه‌ی خود را طوری طراحی کردند که بتواند به تقاضاهای با حجم بالا پاسخگو باشد. تحقیق [۱۹] تلاش خود را بر غلبه بر چالش‌های تقاضای فصلی مسافران متمرکز کرده است. وی و همکاران [۱۹] یک مدل تکاملی جدید را ارائه دادند و به دلیل NP-hard بودن BTNDP، استفاده از رهیافت‌های تکاملی را توصیه پذیر دانستند [۱۹].

پژوهش [۲۰] تلاشی برای بهینه‌سازی زمان سفر، با تأکید بر زمان مورد نیاز برای تغییر خطوط توسط مسافران می‌باشد. رسیدن به این هدف در [۲۰] از طریق الگوریتم ژنتیک صورت می‌گیرد. زمان سفر تنها

پارامتری نیست که نیاز است تا در طراحی شبکه به صورت بهینه درآید. پژوهش [۲۱] تحقیقی برای یافتن شبکه‌ی مطلوب با در نظر گرفتن هزینه‌ی کل به عنوان فاکتور اصلی است. پژوهش [۲۱] نیز به مانند [۲۰] از الگوریتم ژنتیک بهره گرفته است. دویس و همکاران [۴]، BTNDP را از دیدی متفاوت از دیگر محققان حل کردند، آن‌ها تلاش می‌کردند شبکه‌شان را با تقاضای موجود منطبق کنند. در پژوهش [۴] هرگاه معابر شهر یا تقاضای مسافران تغییر می‌کرد، شبکه با شبکه‌ی جدیدی جایگزین می‌شد.

زاهو و همکاران [۲۲] یک پژوهش در دسته‌ی روش‌های دقیق را انجام دادند. آن‌ها یک متدولوژی ریاضی برای BTNDP ارائه دادند و اهداف خود را کمینه کردن تعداد جابه‌جایی مسافران بین خطوط، کمینه کردن فاصله‌ی مسیر و بیشینه کردن پوشش سرویس شبکه قرار دادند. دنگ و همکاران [۲۳] تلاش داشتند تا تعداد مسافران بر روی کوتاه‌ترین مسیر را نیز حداکثر کنند. این ویژگی باعث می‌شود تا زمان سفر کاهش یابد. تحقیق [۲۴]، BTNDP را با هدف حداقل کردن هزینه‌ی کل کاربران و پیمانکاران حل می‌کند. در شبکه‌ی تولیدی توسط [۲۴] مسیریابی انتخاب می‌شود که کم‌ترین هزینه را به همراه دارد. در [۲۴] تابع هزینه‌ی کل یک تابع convex نیست و طبیعتاً مسیرهای با تقاضای بیش‌تر نیاز به سرویس بیش‌تری خود دارند. پژوهش [۲۴] این فرض را که خطوط مستقیم لزوماً مسیرهای بهتری نیستند در نظر گرفته و BTNDP را حل می‌کند.

برخی از مواقع، برای حل BTNDP نیاز به انجام فرآیندهای پیش پردازش و پس پردازش است. از جمله‌ی این فرآیندها انتخاب route-technology است که در پژوهش [۲۵] مورد استفاده قرار گرفته است. رهیافت [۲۵] مشخصه‌های سخت افزاری شبکه را با سیاست‌های اجرایی در نظر گرفته و از طریق یک روش انتخاب route-technology تعیین می‌کند چه hardware-headway ای برای هر جریان در طول مسیر ترانزیت، استفاده شود. مکانیزم حل [۲۵]، یک مدل تکرار شونده است. تکرار تا زمانی ادامه می‌یابد که دیگر نیاز به هیچ تغییری در سطح اتصال سرویس‌ها نباشد. پژوهش [۲۶] یک پژوهش مبتنی بر فرآیندهای پس پردازش است. پدیدآورندگان این پژوهش محل ایستگاه‌ها را به صورت بهینه در سراسر شبکه‌ی خطوط اتوبوس تعیین می‌کنند. اهداف پژوهش [۲۶]، پوشش تقاضا و حداقل کردن زمان سفر است. این کار از طریق تعیین محل ایستگاه‌ها به صورت بهینه صورت می‌گیرد. شیه و همکاران [۷] یک مدل مکاشفه‌ای برای طراحی شبکه‌های اتوبوس ارائه دادند. این مدل دارای چهار روال بود: ۱) روال تولید یک مسیر ۲) روال تحلیل شبکه که شامل یک مدل تخصیص سفر می‌باشد ۳) یک روال انتخاب مرکز حمل و نقل که شامل مراکز انتقال مناسب برای هماهنگی مسیرها است. ۴) یک روال بهبود شبکه را که مسیریابی که توسط پروسیجر ۱ انتخاب شدند بهبود می‌بخشد.

پژوهش [۲۷]، حل BTNDP را با شکستن مساله به یک سری زیر مساله انجام می‌دهد. این زیر مساله‌ها به صورت ترتیبی حل می‌شود. انتخاب یک زیر مجموعه‌ی متصل از مسیرها از یک مجموعه‌ی معین از مسیرها، چالش اصلی در [۲۷] است. در [۲۷] مجموع زمان سفر کاربران در حالی که ناوگان مورد نیاز در حداقل وضعیت خود قرار دارد، بهینه (حداقل) می‌شود.

[۲۸] پژوهش مکاشفه‌ای دیگری است که شبکه‌ی خود را با تمرکز بر روی سود پیمانکار کردن هزینه‌ی مسافران و کاهش زمان انتظار مسافران طراحی کرده است. پژوهش [۲۸] از Ant system برای حل BTNDP بهره برده است. پژوهش [۲۹] نیز با رویکردی مشابه و با در نظر گرفتن محل مناسب برای ایستگاه‌های اتوبوس مساله‌ی BTNDP را حل کرده است. هدف اصلی در [۲۹] حداقل کردن مجموع هزینه‌ی مسافران و پیمانکاران شبکه‌ی اتوبوس رانی می‌باشد.

پژوهش [۳۰] یک روش کاربردی برای شبکه‌های بزرگ می‌باشد. این پژوهش با تقاضاهای پویا در BTNDP مواجه است. این پژوهش که در دسته‌ی روش‌های مکاشفه‌ای جای می‌گیرد، از الگوریتم ژنتیک برای یافتن زیر مجموعه‌های بهینه از مسیرها بهره می‌گیرد. جهت گیری معیارهای طراحی در [۳۰] بر تولید شبکه‌ای متراکم به جای طراحی یک شبکه‌ی گسترده، متمرکز است. شبکه‌ی متراکم همراه با ویژگی‌هایی چون بهبود کارایی، ادغام شدن مسیرهای مستقیم با نقاط نقل و انتقال موثر می‌باشد. این ویژگی‌ها نهایتاً باعث بهبود کیفیت سرویس BTN خواهد شد. پژوهش‌های [۳۱] و [۸] دو پژوهش مشابه از دسته‌ی پژوهش‌های مکاشفه‌ای است. پژوهش [۳۱] یک نسخه‌ی تغییر یافته از پژوهش [۸] می‌باشد. پژوهش [۳۱] نسبت به پژوهش [۸] دارای دو تغییر است. ۱) در نظر گرفتن سرعت اقتصادی اتوبوس‌ها (۲) بروز کردن دفعات تکرار یک مسیر در حین فرآیند تولید مسیرها.

در این مقاله، ما یک روش یادگیری جدید برای حل BTNDP ارائه می‌دهیم. در رهیافت این مقاله طراحی شبکه‌ی فعلی خطوط اتوبوس رانی (خطوط تندرو) شهر تهران به عنوان شبکه‌ی قابل اعتماد و مطلوب در نظر گرفته می‌شود و از دانش استخراج شده از این شبکه برای پیش‌بینی مسیرهای جدید و افزودن به شبکه استفاده می‌شود. روش‌های به کار گرفته شده در این مقاله از جمله روش‌های غیر robust است، نسخه‌ی robust این روش‌ها در پژوهش‌های [۳۲]، [۳۳] بیان شده است. استفاده از آن‌ها را به عنوان ادامه‌ی این کار پژوهشی پیشنهاد می‌کنیم.

داشتن دانش پیشین از شبکه‌ی فعلی اتوبوس رانی، می‌تواند فرآیند تصمیم‌گیری پیمانکاران و مدیران شهری را آسان‌تر کند. از آنجایی که استخراج دانش مزیت اصلی در روش پیشنهادی ما می‌باشد، می‌تواند در فرآیند تصمیم‌گیری به مدیران شهری کمک قابل توجهی کند. در ادامه به بیان فرآیند یادگیری می‌پردازیم.

۳ فرآیند یادگیری

در این بخش به بررسی رهیافت حل این مقاله برای BTNDP می‌پردازیم. با توجه به آنچه در بخش قبل به آن پرداختیم می‌توان فقدان توجه کافی به دانش متخصصان انسانی را استنباط نمود. رهیافت پیشنهادی این مقاله به دنبال غلبه بر این چالش است. دیتاست به کار گرفته شده در این مقاله یک دیتاست دو کلاسه می‌باشد. هر رکورد از دیتاست، بیانگر مشخصات یک معبر می‌باشد و برچسب همراه با آن رکورد، بیان‌کننده‌ی حضور آن معبر در BTN می‌باشد (۱ نشانه‌ی حضور و ۰ عدم حضور). فرآیند یادگیری روش پیشنهادی این مقاله دارای شش گام می‌باشد. در بخش‌های بعدی هر گام با جزئیات مورد بررسی قرار می‌گیرد.

- ۱- جمع آوری داده‌ها: این بخش داده‌ی مورد نیاز را در فرمت مطلوب تهیه می‌کند. این داده‌ها در بخش یادگیری و تست مورد استفاده قرار می‌گیرد (در بخش ۴ به این گام پرداخته شده است).
- ۲- تعیین ویژگی‌ها: در این گام نحوه‌ی انتخاب ویژگی‌ها بیان می‌شود، سپس به تحلیل ویژگی‌ها از طریق الگوریتم Information Gain می‌پردازیم. این تحلیل موثر بودن ویژگی‌های انتخاب شده را نشان می‌دهد.
- ۳- تعیین نوع طبقه‌بند: طبقه‌بندها را می‌توان در دو دسته‌ی آماری و قطعی تقسیم بندی نمود. در این مقاله از روش‌های طبقه‌بندی آماری بهره گرفته می‌شود. به دلیل ماهیت دیتاست مورد استفاده‌مان، طبقه‌بندهای قطعی نتایج ضعیف‌تری نسبت به طبقه‌بندهای آماری تولید کردند.
- ۴- پارامترهای طبقه‌بند: تعیین پارامترهای مناسب به کمک دیتاست یادگیری، در این گام انجام می‌گیرد.
- ۵- بهبود کارایی طبقه‌بند: این گام از طریق داده‌های بخش یادگیری انجام می‌گیرد. در این مرحله از ترکیب طبقه‌بندها برای تولید نتایجی بهتر استفاده می‌شود.
- ۶- کارایی طبقه‌بندها: محاسبه‌ی کارایی طبقه‌بندها بر اساس دقت و پارامترهای دیگری چون precision. در بخش ۶ به این گام پرداخته می‌شود.

۴ خاص سازی داده

در روش پیشنهادی این مقاله، داده‌های خام خود را به صورت shape فایل‌هایی از موسسه‌ی حمل و نقل و سیستم‌های هوشمند وابسته به دانشگاه امیرکبیر (ITSRI) و سازمان کنترل ترافیک شهر تهران تهیه کردیم shape- فایل‌ها فرمتی متداول از فایل‌ها در نرم‌افزار ArcGIS می‌باشد، سپس shape فایل‌ها را به هم متصل کردیم و فایل حاصل را به فرمت یک فایل excel در آورديم. در نهایت با توجه به پژوهش‌های پیشین [۵، ۸، ۳۴-۳۷]، مشورت با متخصصان انسانی و استفاده از روش‌های رتبه‌بندی ویژگی (در این مقاله از Information Gain استفاده نمودیم)، به ویژگی‌های ذکر شده در جدول ۱ دست یافتیم. این بخش معادل با گام‌های ۱ و ۲ ذکر شده در فرآیند یادگیری (بخش سوم) می‌باشد. همه‌ی ویژگی‌های جدول ۱ در رهیافت پیشنهادی این مقاله به کار گرفته است (الحاقیه‌ی ۲).

روش‌های مختلفی برای رتبه‌بندی ویژگی‌ها وجود دارد: Gini Index [۳۸]، Information Gain [۳۹]، Likelihood-Ration [۴۰] از آن جمله هستند. ما روش Information Gain را بر روی ویژگی‌های اولیه اعمال نمودیم و در نهایت به ویژگی‌های جدول ۱ رسیدیم. از میان ویژگی‌های اولیه، ویژگی‌هایی چون وجود پل عابر پیاده، مدرسه یا فرودگاه کنار گذاشته شد. در میان ویژگی‌های جدول ۱، سه ویژگی تاثیرگذاری بیش‌تری دارند: تعداد معابر متصل به یک مسیر خاص، پهنای معبر، نوع معبر (بزرگراه، بلوار، کوچه و ...). با این که رتبه‌بندی ویژگی‌های جدول ۱ تنها به عنوان یک کار آماری صورت گرفت و همه‌ی آن‌ها در فرآیند یادگیری شرکت

دارند؛ اما توجه پذیر بودن برتری این سه ویژگی می تواند تصدیق دیگری باشد بر کارایی ویژگی های انتخاب شده. در زیر بخش بعدی به معرفی معیار Information Gain پرداخته شده است.

۴-۱ معیار Information Gain

این معیار، در اصل معیاری مبتنی بر نظریه ی اطلاعات است [۳۹] و اطلاعات مورد نیاز برای طبقه بندی یک رکورد معین را محاسبه می کند. اطلاعات مورد نیاز برای کد کردن برچسب های ممکن از کلاس (Y) توسط رابطه ی (۱) محاسبه می شود. در رابطه ی (۱) مقدار $P(Y = C_i)$ یک مقدار احتمال با مقداری غیر از صفر است و احتمال آن که یک رکورد متعلق را به کلاس C_i باشد بیان می کند. در معیار Information Gain یک تابع لگاریتمی باینری به کار گرفته می شود و اطلاعات در قالب بیت هایی ذخیره می شود.

$$H(Y) = - \sum_{i=1}^{i=k} P(Y = C_i) \log_2(P(Y = C_i)) \quad (1)$$

در رابطه ی (۱) $H(Y)$ تحت عنوان آنتروپی داده، شناخته می شود. در صورت برقرار بودن توزیع یکنواخت بر روی داده های یادگیری مقدار $H(Y)$ بیش تر خواهد بود. مقدار اطلاعات مورد نیاز برای طبقه بندی یک رکورد بر اساس ویژگی های دانسته شده ی x_a از طریق رابطه ی (۲) قابل محاسبه می باشد.

$$H(Y | x_a) = - \sum_{u \in \text{val}(x_a)} P(x_a = u) H(Y | x_a = u) \quad (2)$$

مقدار معیار Information Gain برابر با میزان تفاوت بین دو رابطه ی (۱) و (۲) و مطابق با رابطه ی (۳) می باشد.

$$I(x_a) = H(Y) - H(Y | x_a) \quad (3)$$

جدول ۱. ویژگی های دیتاست

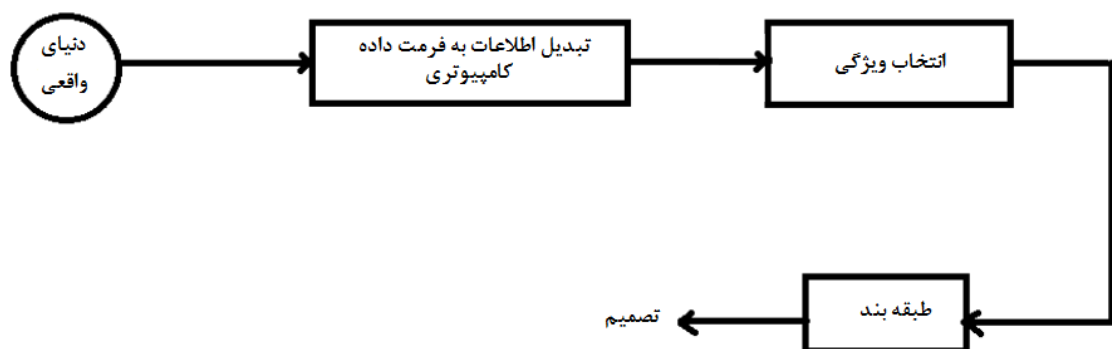
شماره سطر	نام ویژگی
۱	طول معبر
۲	تعداد معابر متصل به یک معبر
۳	پوشش منطقه ای توسط خطوط اتوبوس
۴	تعداد بیمارستان های در دسترس یک معبر. ناحیه ی همسایگی ۲۰۰ متر مربع در نظر گرفته شده است.
۵	تعداد مراکز خرید در دسترس یک معبر (ناحیه ی همسایگی ۲۰۰ متر مربع در نظر گرفته شده است)
۶	عرض معبر
۷	یک طرفه یا دو طرفه بودن معبر
۸	مقدار جریان در یک معبر
۹	چگالی جمعیت در همسایگی یک معبر
۱۰	نوع مسیر (کوچه، خیابان یا بزرگراه)
۱۱	جمعیت کل پوشش داده شده توسط معبر

۵ مدل یادگیری

طبقه‌بند، عموماً بخش اصلی در رهیافت‌های یادگیری ماشین است [۴۱]. تولید یک مدل، هدف اصلی از یک مساله طبقه‌بندی می‌باشد. این مدل تولیدی که طبقه‌بند نامیده می‌شود، بردار مقادیر ویژگی‌ها را دریافت می‌کند و آن‌ها را به یک برچسب از کلاس هدف نگاشت می‌کند (رابطه‌ی (۴)). در رابطه‌ی (۴) بردار $X = (x_1, x_2, x_3, \dots, x_p)$ بردار مقادیر ویژگی‌ها و بردار $Y \in (C_1, C_2, C_3, \dots, C_p)$ بردار برچسب کلاس هدف می‌باشد.

$$(X, Y) = ((x_1, x_2, x_3, \dots, x_p), Y) \quad (4)$$

نحوه‌ی تولید مدل رهیافت پیشنهادی این مقاله در شکل (۱) بیان شده است. در ادامه و در زیر بخش‌های مدل یادگیری، در ابتدا هر روش را از لحاظ نظری معرفی می‌کنیم؛ سپس نحوه‌ی به کارگیری آن روش و پارمترهای به کار گرفته شده در آن روش را بیان می‌داریم.



شکل ۱. نحوه‌ی تصمیم‌گیری

۵-۱ روش PLS regression

هدف اصلی در PLS، تولید یک مجموعه از متغیرهای ناهمبسته است. این متغیرها که مولفه‌های PLS هم نامیده می‌شود [۴۲]، [۴۳] ارتباط میان بردار ویژگی‌های ورودی (X) و برچسب کلاس‌ها (Y) را می‌یابد. PLS را می‌توان جانشینی برای PCR (principal component regression) دانست. PCR که یک روش مبتنی بر ایده‌ی کاهش ابعاد است [۴۲]، یک مدل خطی $Y = (\beta X) + \varepsilon$ را تخمین می‌زند (β و ε به ترتیب یک بردار ثابت یک بردار تصادفی از متغیرهایی با توزیع نرمال می‌باشد). PCR با معیار Least-square یک مساله با پاسخ یکتا نمی‌باشد [۴۴]، PLS از طریق جایگزین کردن معیار Least-square با ماگزیمم کوواریانس بین X و Y توسط این چالش را حل می‌کند.

تولید مولفه‌های PLS بخش اصلی روش PLS را شکل می‌دهد. مولفه‌های PLS بیان ارتباطات بین X و Y هستند. رابطه‌ی اصلی در PLS به صورت رابطه‌ی (۵) می‌باشد. در این رابطه W^x و W^y بردارهای وزن می‌باشند [۴۵]. این بردارهای ویژه مولفه‌های اصلی X و Y هستند.

$$W^X Z = \int_0^T E(X_t Z) X_t dt, \quad W^Y Z = \sum_{i=1}^p E(Y_i Z) Y_i, \quad \forall Z \in \mathcal{H}. \quad (5)$$

اگر $X_1, \dots, X_n$ یک نمونه از X باشند. تخمین زننده W^X یک ماتریس $\hat{W}^X$ با سائز $n \times n$ خواهد بود. درایه‌های این ماتریس از طریق $\hat{w}_{i,j} = (X_i, X_j)_{L^2([0,T])}$ قابل محاسبه هستند [۴۶].

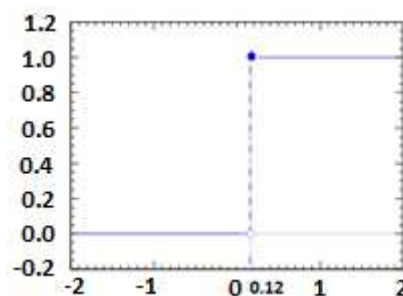
گام‌های زیر چگونگی به کارگیری PLS در پژوهش ما را بیان می‌کند:

۱- اجرای PLS بر روی داده‌های بخش یادگیری.

۲- محاسبه‌ی یک مقدار برای Y (برای هر رکورد از دیتاست یک مقدار محاسبه می‌شود).

۳- اجرای یک تابع پله‌ای بر روی Y برای پیش‌بینی برچسب کلاس یک معبر.

بهترین مقدار تمیز دهنده برای نگاشت یک مقدار به یک برچسب در PLS برابر با 0.12 می‌باشد (شکل ۲). این مقدار از طریق جستجو در مقادیر تولیدی در PLS و انتخاب مقداری که به بهترین شکل تفکیک را بر روی معابر انجام دهد به دست آمده است. پارامترهای مورد استفاده برای PLS در جدول ۲ بیان شده. الگوریتم Nipals به کار گرفته شده در PLS برای محاسبه‌ی مولفه‌ی اصلی دیتاست می‌باشد. استفاده از این الگوریتم زمان اجرا را افزایش می‌دهد، در عین حال باعث افزایش دقت در محاسبه‌ی ماتریس کوواریانس هم خواهد شد.



شکل ۲. تابع پله‌ای PLS

جدول ۲. مقادیر پارامترها در PLS

پارامترهای PLS	مقادیر
الگوریتم مورد استفاده برای تخمین زدن وزن‌ها	الگوریتم Nipals
تعداد مولفه‌ها	۲
حداکثر تعداد تکرارها در الگوریتم Nipals	۴۰۰

۲-۵ روش Lasso regression

روش Lasso یکی از انواع least-square regression می‌باشد. Lasso را می‌توان یک تخمین زننده‌ی مدل‌های خطی دانست. Lasso باقی‌مانده‌ی مجموع مربعات را با در نظر گرفتن مجموع مقادیر قدر مطلق ضرایب مینیمایز می‌کند. مقدار مطلق ضرایب باید از مقدار ثابتی کم‌تر باشد [۴۷]. به دلیل توانایی Lasso برای ساختن یک مدل آماری کارا برای حل BTNDP از آن بهره گرفتیم. اگر (X_i, Y_i) ، $i = 1, 2, \dots, N$ بیان‌کننده‌ی رکوردهای یک

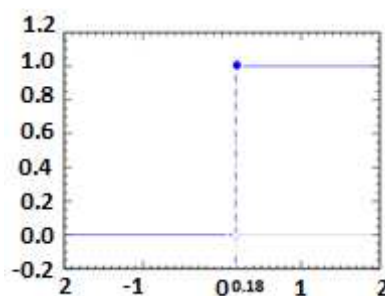
دیتاست باشند $X_i = (x_{i1}, \dots, x_{ip})^T$ بردار ویژگی ها و Y_i برچسب کلاس رکورد را نشان دهد). اگر Lasso به عنوان یک مدل رگرسیون معمول در نظر گرفته شود، آن وقت مشاهدات مستقل از یکدیگر بوده و x_{ij} استاندارد سازی می شود. نتیجتاً $\sum_i \frac{x_{ij}^2}{N} = 1$ و $\sum_i \frac{x_{ij}}{N} = 0$ برقرار خواهد بود. رابطه‌ی (۶) بیان کننده‌ی مقداری است که در Lasso محاسبه می شود.

$$(\hat{\alpha}, \hat{\beta}) = \arg \min \left\{ \sum_{i=1}^N \left(y_i - \alpha - \sum_j \beta_j x_{ij} \right)^2 \right\} \quad \text{subject to } \sum_j |\beta_j| \leq t \quad (6)$$

رابطه‌ی (۶) یک مسالهی بهینه‌سازی از درجه‌ی ۲ می باشد. پارامتر t یک پارامتر میزان سازی (تنظیم) است. مقدار این پارامتر باید مقداری بیش از صفر باشد. t بیان کننده‌ی میزان انقباض در رابطه‌ی (۶) است. در رهیافت این مقاله از رابطه‌ی (۷) به عنوان تابعی که قرار است در طی فرآیند پیش‌بینی مینماید شود، بهره گرفته ایم.

$$\min \left\{ \frac{1}{N} \|Xw - Y\|_2^2 + \frac{\lambda}{p} \|w\|_1 \right\} \quad (7)$$

در رابطه‌ی بالا، X (داده‌ی یادگیری) یک ماتریس با ابعاد $N \times p$ می باشد. λ پارامتر تنظیمی است که اگر دارای مقدار صفر باشد رابطه‌ی (۶) را تبدیل به یک برآورد least-square می کند. Y بردار برچسب‌های کلاس و w ضرایب رگرسیون را نشان می دهد. تابع پله‌ای (مقدار تفکیک کننده در فرآیند پیش‌بینی) و مقادیر پارامترهای مورد استفاده توسط Lasso به ترتیب در شکل ۳ و جدول ۳ بیان شده است.



شکل ۳. تابع پله‌ای Lasso regression

جدول ۳. مقادیر پارامترها در Lasso regression

مقادیر	پارامترهای Lasso
۰/۱۸	پارامتر تنظیم سازی
۱۵۰۰۰	حداکثر تعداد تکرارهای متد بهینه سازی
۰/۰۰۰۰۱	مقداری که با رسیدن به مقداری کم‌تر از آن بهینه سازی متوقف می شود

۳-۵ روش Naive Bayesian

روش Naive Bayesian یک طبقه‌بند آماری است. در این طبقه‌بند، مقدار احتمال برای همه‌ی برچسب‌های کلاس با فرض استقلال ویژگی‌ها محاسبه می‌شود و کلاس دارای بیش‌ترین احتمال، برچسب کلاس رکورد را تعیین می‌کند. استقلال ویژگی‌ها از یکدیگر باعث می‌شود تا توزیع احتمال یک ویژگی بر ویژگی دیگر اثر نداشته باشد؛ بنابراین نیازی به محاسبه‌ی توزیع احتمال توأم نیست و هر توزیع احتمال به صورت جداگانه محاسبه می‌شود (رابطه‌ی ۸).

$$P(x_i | C_m, x_j) = P(x_i | C_m), i \neq j \quad (8)$$

با در نظر گرفتن تئوری Bayesian، مقدار $P(Y = C_m | X)$ توسط رابطه‌ی (۹) محاسبه می‌شود. در این رابطه X بردار داده و Y حاوی یکی از برچسب‌های کلاس می‌باشد.

$$P(Y = C_m | X) = \frac{P(Y = C_m) P(X | Y = C_m)}{P(X)} \quad (9)$$

صورت کسر رابطه‌ی (۹) را می‌توان به توزیع احتمال توأم تبدیل نمود؛ بنابراین رابطه‌ی (۹) به صورت رابطه‌ی (۱۰) در خواهد آمد.

$$P(C_m) \times P(x_1 | C_m) \times P(x_2 | C_m, x_1) \times P(x_3 | C_m, x_1, x_2) \times \dots \times P(x_n | C_m, x_1, x_2, \dots, x_{n-1}) \quad (10)$$

با در نظر گرفتن رابطه‌ی (۸) معادله‌ی (۱۰) به صورت معادله‌ی (۱۱) قابل تبدیل است.

$$P(C_m) \times P(x_1 | C_m) \times P(x_2 | C_m) \times P(x_3 | C_m) \times \dots \times P(x_n | C_m) = P(C_m) \times \prod_{i=1}^n P(x_i | C_m) \quad (11)$$

طبقه‌بند Naive Bayesian کلاس یک رکورد تست را از طریق رابطه‌ی (۱۲) تعیین می‌کند.

$$\hat{Y}(X) = \operatorname{Argmax}_{C_m} \left(P(C_m) \times \prod_{i=1}^n P(x_i | C_m) \right) \quad (12)$$

در مساله‌ی این مقاله برچسب کلاس با احتمال بزرگ‌تر تعیین کننده، برچسب پیش‌بینی برای رکورد تست خواهد بود. در تعیین احتمال‌های پیشین از تخمین زننده‌ی لاپلاس استفاده نمودیم [۴۸].

۴-۵ استفاده از روش‌های ترکیبی

برای حل BTNDP تا کنون از سه روش طبقه‌بند آماری استفاده کردیم. هر کدام از آن‌ها در جنبه‌هایی برتر بودند.

- PLS روابط بین متغیرهای پاسخ و متغیرهای پیش‌بینی را پیدا می‌کند.
- Lasso مبتنی بر least-square ها می‌باشد.
- طبقه‌بند Naive Bayesian با بودن استقلال بین ویژگی‌ها کارایی خوبی از خود نشان می‌دهد.

با توجه به سه جنبه‌ی ذکر شده در بالا تصمیم گرفتیم تا Naïve Bayesian را با Lasso و PLS به صورت جداگانه ترکیب کنیم و متدهای ترکیبی زیر را تولید کنیم. نتیجه‌ی این عمل همراه با نرخ خطای کم‌تری بود (در بخش ۶ به صورت دقیق‌تر به تحلیل نتایج این دو روش خواهیم پرداخت).

- ترکیب PLS و Naïve Bayesian که آن را PN نامیده‌ایم.
- ترکیب Lasso و Naïve Bayesian که آن را LN نامیده‌ایم.

در دو روش ترکیبی PN و LN مقادیر تولید شده توسط PLS و Lasso نرمالایز می‌شود به نحوی که مقدار جدید در محدوده‌ی ۱- تا ۱+ قرار گیرد، سپس احتمال برچسب ۱ (شرکت معبر در شبکه) با کمک طبقه‌بند Naïve Bayesian محاسبه می‌شود. این مقدار با مقدار تولید شده توسط PLS و Lasso جمع شده، اگر مقداری بیش از ۰/۶۲ بود برچسب پیش‌بینی برابر با ۱ خواهد شد.

۶ تحلیل و نتایج

پیاده‌سازی روش‌های پیشنهاد شده در این مقاله در این بخش مورد بررسی و تحلیل قرار می‌گیرد. داده‌های استفاده شده در این مقاله، داده‌های شبکه‌ی اتوبوس رانی تهران می‌باشد (الحاقیه‌ی ۱). این داده‌ها شامل ۶۰۶۶ رکورد می‌باشد. ما این داده‌ها را به دو گروه تقسیم نمودیم:

- ۱- داده‌های بخش یادگیری که شامل ۴۵۰۰ رکورد می‌باشد.
- ۱- داده‌های بخش تست، که شامل ۱۵۶۶ رکورد هستند.

پیاده‌سازی PLS، Lasso و Naïve Bayesian توسط زبان برنامه‌نویسی پایتون و با استفاده از Library Orange [۴۹] در Komodo IDE 8.0 صورت گرفته است. پارامترهای استفاده شده در پیاده‌سازی روش‌های PLS و Lasso به ترتیب در جداول ۲ و ۳ بیان شده. فرآیند پیش‌بینی PLS، Lasso و Naïve Bayesian در چارچوبی بدین شکل عمل می‌کند: هر کدام از این روش‌ها، مقداری را به عنوان پاسخ، برای هر رکورد تست تولید می‌کنند، سپس بر حسب مقدار تولید شده توسط آن‌ها، رکورد ورودی به یکی از برچسب‌های کلاس‌ها نگاشت می‌شود. این نگاشت در PLS و Lasso از طریق توابع پله‌ای (شکل ۲ و ۳) و در Naïve Bayesian از طریق میزان احتمال پیش‌بینی شده، صورت می‌گیرد.

جداول ۵، ۶ و ۷ شامل نتایج پیش‌بینی PLS، Lasso و Naïve Bayesian بر روی داده تست می‌باشند. ارزیابی روش‌ها توسط معیارهای دقت، True-Positive و False-Positive صورت گرفته است (روابط ۱۳، ۱۴ و ۱۵). این روابط با در نظر گرفتن جدول ۴ بیان شده‌است.

$$\text{دقت} = \frac{A + D}{A + B + C + D} \quad (۱۳)$$

$$\text{True positive} = \frac{D}{C + D} \quad (۱۴)$$

$$False\ positive = \frac{B}{A+B} \quad (15)$$

جدول ۴. ماتریس درهم ریختگی				جدول ۵. ماتریس درهم ریختگی برای روش PLS			
پیش‌بینی		منفی	مثبت	پیش‌بینی		منفی	مثبت
واقعی	منفی			واقعی	منفی		
	مثبت	A	B		مثبت	۱۱۶۱	۱۴۷
	مثبت	C	D		مثبت	۱۳	۲۴۵

جدول ۶. ماتریس درهم ریختگی برای روش Lasso				جدول ۷. ماتریس درهم ریختگی برای روش Naïve Bayesian			
پیش‌بینی		منفی	مثبت	پیش‌بینی		منفی	مثبت
واقعی	منفی			واقعی	منفی		
	مثبت	۱۲۲۴	۹۵		مثبت	۱۲۳۲	۸۰
	مثبت	۲۹	۲۱۸		مثبت	۳۸	۲۱۶

در دنیای واقعی بین دو نقطه‌ی شهری بیش از یک مسیر وجود دارد و بسته به نوع هدف در BTNDP می‌توان شبکه‌های مختلفی را تولید نموده؛ بنابراین اگر یک مسیر پیش‌بینی شده در شبکه‌ی اتوبوس رانی حضور نداشته باشد، این نشان دهنده‌ی عدم کیفیت آن مسیر نمی‌باشد. با این حال معیار دقت به تنهایی معیار مناسبی نیست (دقت PLS، Lasso و Naïve Bayesian به ترتیب ۰/۸۹۷، ۰/۹۲، ۰/۹۲۴ می‌باشد (جدول ۸)). برای غلبه بر ضعف معیار دقت، از معیارهای True-Positive و False-positive هم استفاده نمودیم.

جدول ۸، نشان می‌دهد، نتایج به دست آمده به میزان قابل توجهی مناسب است. با دقت در مقادیر معیارهای True-Positive و False-Positive در جدول ۸، می‌توان دانست مقادیر True-Positive در PLS و Lasso نسبت به Naïve Bayesian بهتر بوده، و وضعیتی مشابه برای Naïve Bayesian نسبت به PLS و Lasso در معیار False-Positive وجود دارد. این نتایج به دلیل ماهیت این سه روش و جنبه‌هایی که هر کدام از آن‌ها در آن بهتر هستند، می‌باشد (بخش ۴-۵). همان‌طور که در بخش ۴-۵ بیان شد، روش‌های PLS و Lasso به صورت جداگانه با Naïve Bayesian ترکیب می‌شود تا با ترکیب جنبه‌های برتری آن‌ها با یکدیگر بتوانیم نرخ خطا در شرکت مسیرها در BTN را کاهش دهیم. این امر به این دلیل است که کاهش حضور مسیرهای نامناسب در BTN اهمیت بیش‌تری در BTNDP دارد و اندکی کاهش در میزان True-Positive را توجیه پذیر، می‌کند.

جدول ۸. دقت، True-Positive و False-Positive برای روش PLS، Lasso و Naïve Bayesian			
	دقت	True-Positive	False-Positive
PLS	۰/۸۹۷	۰/۹۴۹	۰/۱۱۲۳
Lasso	۰/۹۲۰	۰/۸۸۲	۰/۰۷۲۰
Naïve Bayesian	۰/۹۲۴	۰/۸۵۰	۰/۰۶۰۹

۷ نتیجه‌گیری و جمع‌بندی

BTNDP یک مساله NP-hard است. در نتیجه رسیدن به یک شبکه‌ی کاملاً ایده‌آل در ابرشهرها کاری شدنی نیست. به علاوه حل این مساله فقط با تاکید بر جنبه‌های ریاضی، منجر به تولید یک شبکه‌ی عملیاتی یا توسعه‌ی مناسب شبکه‌های موجود نخواهد شد. در این مقاله یک رهیافت مبتنی بر دانش ارایه شده که بر خلاف دیگر رهیافت‌ها از شبکه‌ی موجود یاد می‌گیرد. به بیان دیگر ما به جای آن که یک مساله‌ی NP-hard را حل کنیم، عوامل موثر در طراحی شبکه‌های جاری را یاد می‌گیریم؛ بنابراین روش ما علاوه بر توانایی برای حل BTNDP توانایی به کارگیری برای توسعه‌ی شبکه را نیز دارد. در این مقاله ما از سه روش یادگیری (PLS، Lasso و Naïve Bayesian) و نسخه‌ی ترکیبی آن‌ها بهره گرفتیم و نتایج را به کمک دقت، True-Positive و False-Positive ارزیابی نمودیم. در نتایج به این دید رسیدیم که PLS و Lasso در معیار True-Positive نسبت به Naïve-Bayesian برتر می‌باشند و شرایطی برعکس در معیار False-Positive بین Naïve-Bayesian با PLS و Lasso برقرار است؛ بنابراین PLS و Lasso را با Naïve Bayesian به صورت جداگانه ترکیب کردیم و توانستیم به نرخ خطای کم‌تری در مسیرهایی که در BTN شرکت دارند، برسیم. به کارگیری دانش استخراج شده از شبکه‌های حمل و نقل دیگر و به کارگیری ورژن Robust طبقه‌بندهای آماری را به عنوان زمینه‌هایی برای پژوهش‌های آینده پیشنهاد می‌کنیم.

سپاسگزاری

داده‌های استفاده شده در این مقاله از طریق موسسه‌ی حمل و نقل و سیستم‌های هوشمند وابسته به دانشگاه امیرکبیر (ITSRI) و سازمان کنترل ترافیک شهر تهران تهیه شده است. نویسندگان تشکر خود را از این دو موسسه به دلیل حمایت‌هایشان اعلام می‌دارند. همچنین لازم است تا از دانشگاه آزاد اسلامی واحد قزوین به دلیل حمایت مالی از انجام این پژوهش تقدیر و تشکر نماییم.

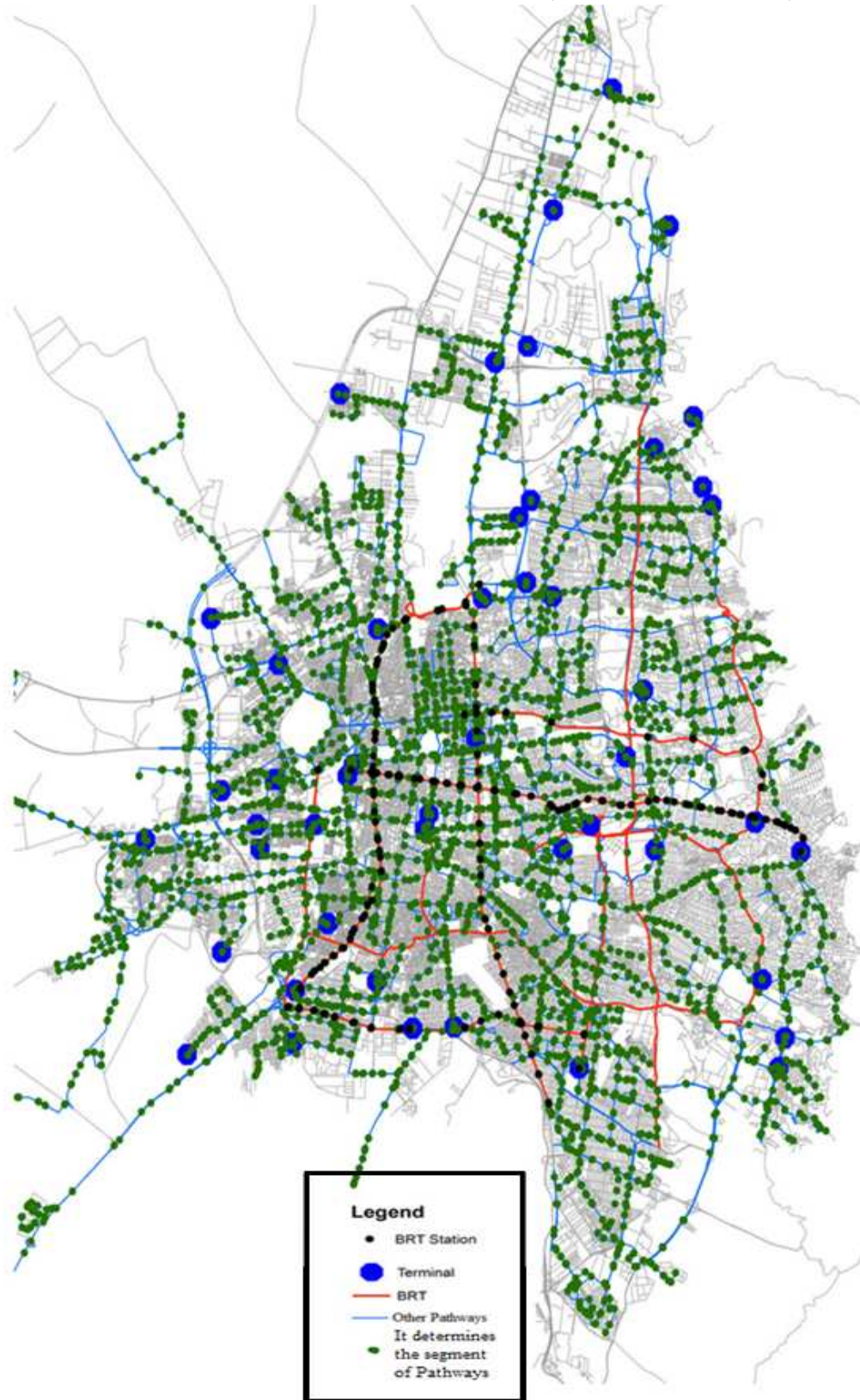
منابع

- [۱] آذر، ع.، مهدوی راد، ع.، موسی خانی، م.، (۱۳۹۴). طراحی مدل ترکیبی داده کاوی و تصمیم‌گیری چند معیاره (مورد مطالعه: بانک اطلاعات یارانه‌های مرکز آمار ایران). مجله تحقیق در عملیات در کاربردهای آن، ۴۴(۱)، ۹۵-۱۱۱.
- [2] Hashemi, S. M., Almasi M., Ebrazi R., Jahanshahi M., (2012). Predicting the Next State of Traffic by Data Mining Classification Techniques. International Journal of Smart Electrical Engineering (IJSEE), 1(3).
- [3] Haiyun, Lu., (2014). Short-Term Traffic Prediction Using Rainfall. International Journal of Signal Processing Systems, 2(1), 70-74.
- [4] Dubois, D., Bel, G., Llibre M., Libre, M., (1979). A Set of Methods in Transportation Network Synthesis and Analysis. Operational Research, 30, 797-808.
- [5] Baaj, M. H., Mahmassani, H. S., (1990). TRUST: A Lisp program for the analysis of transit route configurations. Transportation Research Record, 1283, 125-135.
- [6] Zhao, F., Zeng, X., (2006). Simulated annealing-genetic algorithm for transit network optimization. Journal of Computing in Civil Engineering, 20(1), 57-68.

- [7] Shih, M. C., Mahmassani, H. S., Baaj M. H., (1998). Planning and design model for transit route networks with coordinated operations. *Transportation Research Record: Journal of the Transportation Research Board*, 1623(1), 16-23.
- [8] Baaj, M. H., Mahmassani, H. S., (1995). Hybrid route generation heuristic algorithm for the design of transit networks. *Transportation Research C*, 3(1), 31-50.
- [9] Guan, J.F., Yang, H., Wirasinghe, S.C., (2003). Simultaneous optimization of transit line configuration and passenger line assignment. *Transport. Res, Part B* 40 (10), 885-902.
- [10] Lines, A. H., Lampkin, W., Saalman, P. D., (1966). The design of Routes and Service Frequency for a Municipal Bus Company. *Business Operation Research LTD*, London, U.K.
- [11] Silman, L. A., Barzily, Z., Passy U., (1974). Planning the Route System for Urban Buses. *Computers and Operations Research*, 1(2), 201-211.
- [12] Bansal, A. N., (1981). Optimization of Bus Route Network for a Fixed Spatial Distribution. *Scientific Management of Transportation Systems*; Amsterdam: North Holland Publishing Company, 346-355.
- [13] Ceder, A., Wilson, H. M., (1986). Bus network Design. *Transportation Research*, 20(4), 331-334.
- [14] Dufourd, H., Gendreau, M., Laporte, G., (1996). Locating a Transit Line Using Tabu Search. *Transportation Science*, 4, 1-19.
- [15] Fusco G., Gori S., Petrelli M., (2002). A heuristic transit network design algorithm for medium size towns. *Proceedings of the 9th Euro Working Group on Transportation*, 652-656.
- [16] Bielli, M., Caramia, M., Carotenuto, P., (2002). Genetic Algorithms in Bus Network Optimization., *Transportation Research Part C: Emerging Technologies*, 10(1), 19-34.
- [17] Fan W., Machemehl R. B., (2004). A Tabu Search Based Heuristic Method for the Transit Route Network Design Problem. *Proceedings of the 9th International Conference on Computer-Aided Scheduling of Public Transport (CASPT)*, San Diego, CA.
- [18] Yanfeng O., Nourbakhsh S. M., Cassidy M., (2013). Continuum Approximation Approach to Bus Network Design under Spatially Heterogeneous Demand. *Transportation Research, Part B: Methodological*, 68, 333-344.
- [19] Wei W., Xinmiao Y., Xuewu C., (2002). *Urban Public Transportation System Planning and Management*. Science Publishing Company, Beijing.
- [20] Somnuk N., Lovell D. J., (2003). Optimal time transfer in bus transit route network design using a genetic algorithm. *Journal of Transportation Engineering*, 129(5), 510-521.
- [21] Pattnaik S. B., Mohan S., Tom V. M., (1998). Urban Bus Transit Network Design Using Genetic Algorithm. *Journal of Transportation Engineering*, 124(4), 368-375.
- [22] Zhao F., Gan A., (2003). Optimization of transit network to minimize transfers. prepared for Research Center Florida Department Transportation.
- [23] Deng, Y., Mo, J., Wang J., (2008). An Improvement Route Generation Algorithm for Bus Network Design. *Workshop on Power Electronics and Intelligent Transportation System*, 199-202.
- [24] Newell, G. F., (1979). Some Issues Relating to the Optimal Design of Bus Routes. *Transportation Science*, 13(1), 20-35.
- [25] Rea, J. C., (1972). Designing urban transit systems: An approach to the route-technology selection problem. No. ISBN 0-309-02089-1, 48-59.
- [26] Jahani, M., Hashemi, S.M., Ghatee, M., Jahanshahi M., (2014). A novel model for bus stop location appropriate for Public Transit Network Design: The case of Central Business Districts (CBD) of Tehran. *International Journal of Smart Electrical Engineering*, 2(3).
- [27] Poorzahedy H., Safari F., (2011). An Ant System application to the Bus Network Design Problem: an algorithm and a case study. *Public Transport*, 3(2), 165-187.
- [28] Spasovic, L.N., Boilé M. P., Bladikas A. K., (1994). Bus Transit Service Coverage for Maximum Profit and Social Welfare. *Transportation Research Record*, 1451, 12-22.
- [29] Spasovic, L. N., Schonfeld P., (1993). A Method for Optimizing Transit Service Coverage. *Transportation Research Record*, 1402, 28-39.
- [30] Ernesto, C., Gori, S., Petrelli, M., (2012). A bus network design procedure with elastic demand for large urban areas. *Public Transport*, 4(1), 57-76.
- [31] Foletta, N., Estrada-Romeu, M., Roca-Riu, M., Martí, P., (2010). New Modifications to Bus Network Design Methodology. *Transportation Research Record: Journal of the Transportation Research Board*, 2197(1), 43-53.
- [32] Constantine, C., Mannor, S., Xu, H., (2011). *Robust Optimization in Machine Learning*. Optimization for machine learning MIT Press, 369.

- [33] Ben-Tal, A., El Ghaoui, L., Nemirovski, A. (2009). Robust Optimization. Princeton Series in Applied Mathematics. Princeton University Press.
- [34] Fan, W. D., Randy, B. M., (2008). Some Computational Insights on the Optimal Bus Transit Route Network Design Problem. Journal of the Transportation Research Forum, 47(3), 61-76.
- [35] Lee, Y. J., Vuchic, V.R., (2000). Transit Network Design with Variable Demand. Proc the 79th Annual Meeting of TRB, Washington D.C.
- [36] Zhao, F., Ubaka, I., (2004). Transit network optimization – minimizing transfers and optimizing route directness. Journal of Public Transportation, 7(1), 67-82.
- [37] Sadrsadat, H., Poorzahedi, H., Haghani, A., Sharifi, E., (2012). Bus network design using Genetic algorithm. TRF Foundation The 53rd Annual Forum was held March in Tampa, Florida, 15-17.
- [38] Breiman, L., Friedman, J. H., Olshen R. A., Stone, C. J., (1984). Classification and Regression Trees. Publisher: Chapman & Hall.
- [39] Quinlan, J., (1986). Induction of decision trees. Machine Learning, 1(1), 81-106.
- [40] Attneave, F., (1959). Applications of information theory to psychology: a summary of basic concepts methods and results. University of Oregon.
- [41] Fayyad, U., Gregory, P.S., Smyth, P., (1996). From Data Mining to Knowledge Discovery in Databases. American Association for Artificial Intelligence, 37-54.
- [42] Aguilera, A. M., Ocana, F., Valderrama, M. J., (1997). An approximated principal component prediction model for continuous-time stochastic process. Applied Stochastic Models and Data Analysis, 13(2), 61-72.
- [43] Saporta, G., (1981), Méthodes exploratoires d'analyse de données temporelles., PhD diss. Université Pierre et Marie Curie-Paris.
- [44] Ferraty, C. H. F., Sarda, P., (1999). Functional linear model. Statistics & Probability Letters, 45(1), 11-22.
- [45] Escoufier, Y., (1970), Echantillonnage dans une population de variables aléatoires réelles. Department de math.; Univ. des sciences et techniques du Languedoc.
- [46] Preda, C., Saporta, G., Lévêder, C., (2007). PLS classification of functional data. Computational Statistics, 22(2), 223-235.
- [47] Tibshirani, R., (1996). Regression Shrinkage and Selection via the Lasso. Journal of the Royal Statistical Society, Series B (methodological) 58(1), 267-288.
- [48] Cestnik, B., (1990). Estimating probabilities: A crucial task in machine learning. the Ninth European Conference on Artificial Intelligence, Stokholm, 147-149.
- [49] <http://docs.orange.biolab.si/>

الحاقیه ۱. در این الحاقیه دیتاست استفاده شده در این مقاله نمایش داده شده است. ما داده‌های خود را در فرمت Shape فایل از سازمان کنترل ترافیک شهر تهران دریافت نمودیم. این داده‌ها در نقشه‌ی زیر نشان داده شده است. خطوط قرمز و آبی به ترتیب نمایش مسیرهایی با و بدون خطوط اتوبوس رانی تندرو هستند. داده‌های Shape فایل اولیه را به فرمت excel تبدیل کردیم و در فرآیند یادگیری از آنها استفاده نمودیم. تصویر خطی از دیتاست مورد استفاده در الحاقیه ۲ نشان داده شده است.



الحاقیه ۲. تصویر خطی از دیتاست مورد استفاده در این مقاله

