

مدلسازی و پیش بینی ترافیک با استفاده از شبکه‌ی عصبی پایه و شبکه‌ی عصبی موجک و به کارگیری سه الگوریتم فراابتکاری ژنتیک، ازدحام ذرات و رقابت استعماری جهت بهینه سازی

زینب قاسم پور^۱، سعید بهزادی^{۲*}

^۱ دانشجوی کارشناسی ارشد سیستم‌های اطلاعات مکانی - دانشکده مهندسی عمران - دانشگاه تربیت دبیر شهید رجایی
zeynabghasempoor8@gmail.com

^۲ استادیار گروه مهندسی نقشه‌برداری - دانشکده مهندسی عمران - دانشگاه تربیت دبیر شهید رجایی
behzadi.saeed@gmail.com

(تاریخ دریافت تیر ۱۳۹۹، تاریخ تصویب آذر ۱۳۹۹)

چکیده

جهت رفتارسنجی ترافیک در کلانشهرها، علاوه بر نیاز به داده‌های ترافیکی به روز و حجیم، نیاز به ارائه و استفاده از روش‌هایی برای تحلیل رفتار ترافیک جهت پیش بینی آن است. موضوع پیش بینی ترافیک، از این جهت مورد توجه است که می‌توان با رسیدن به این هدف، حمل و نقل را رونق بخشید و از بار اقتصادی که هر ساله بر جوامع مختلف وارد می‌شود، کم کرد. به همین دلیل، باید به دنبال راه حل‌هایی برای پیش بینی ترافیک بود. در این میان، استفاده از علمی به نام شبکه‌های عصبی می‌تواند بسیار کاربردی باشد. در این پژوهش، ابتدا سامانه‌ای طراحی شد تا بتواند داده‌های ترافیکی مورد نیاز پژوهش را جمع‌آوری کند. موضوع دستیابی به داده‌های ترافیکی همواره معضلی بوده است که تمامی دغدغه‌مندان این حوزه با آن مواجه بوده‌اند. از این رو در پژوهش حاضر ابتدا با طراحی یک سامانه جهت جمع‌آوری داده‌های ترافیکی به رفع معضل مورد نظر پرداخته شد. در مرحله‌ی بعد، داده‌های ترافیکی جمع‌آوری شده فراخوانی می‌گردد و با استفاده از اعمالی که برای استاندارد سازی و نرمالیزه سازی آن انجام شد، میزان خطا با استفاده از داده‌های تست و آموزشی محاسبه گردید. در این مرحله، میزان خطا با مطلوبیت ۷۲ درصد به دست آمد. در مرحله‌ی بعدی، تحلیل داده‌های ترافیکی با استفاده از ترکیب شبکه عصبی پایه و موجک انجام شد و میزان خطا محاسبه شد. نتایج نشان می‌دهد که محاسبات با استفاده از تبدیلات موجک دقیق تر است ولی به دلیل تعداد ورودی‌های کمتر، مقادیر خطا با استفاده از داده‌های تست و آموزشی، ۲۸ درصد محاسبه گردید. به عبارت دیگر میزان مطلوبیت حدود ۷۲ درصد به دست آمد. در مرحله‌ی آخر هم داده‌های ترافیکی جمع‌آوری شده با استفاده از الگوریتم‌های بهینه‌سازی، بهینه شدند و بهترین نقطه با کمترین خطای ممکن برای هر الگوریتم بهینه‌سازی، محاسبه گردید.

واژگان کلیدی: پیش بینی رفتار ترافیک، شبکه‌ی عصبی، تبدیل موجک، الگوریتم‌های بهینه‌سازی

۱- مقدمه

امروزه، ترافیک در جوامع مختلف، به یک مسئله‌ی بزرگ تبدیل شده است. این موضوع بویژه در کلانشهرها، نمود پر رنگ تری دارد [۱]. تا به امروز، تحقیقات بسیار گسترده‌ای، در حوزه‌ی ترافیک انجام شده است تا بدین وسیله بتوان حجم ترافیک ایجاد را کنترل کرد. اما در پژوهش‌های انجام گرفته، محققان بر روی مسئله‌ی ترافیک و دلایل ایجاد آن، متمرکز شده‌اند و پارامترهای ایجاد ترافیک را مورد بررسی قرار داده‌اند. پارامترهایی از قبیل: آلودگی هوا، میزان جمعیت، دسترسی، میزان آمد و شد و خیلی از پارامترهای دیگر. اما موضوعی که در این میان اهمیت فراوانی دارد، توجه به این نکته است که این پارامترها و عواملی که در ایجاد ترافیک نقش دارند، در جهت پیش‌بینی جریان ترافیک هم می‌توانند مورد استفاده قرار بگیرند. زیرا بدیهی است که اگر بتوان ترافیک را به نحوی پیش‌بینی کرد، می‌توان بسیاری از دغدغه‌های جوامع انسانی بالخصوص کلانشهرها را حل و فصل نمود. اما تا به امروز، تحقیقات مهمی برای رسیدن به این هدف، انجام نگرفته است. محققان، اغلب بر روی علل ایجاد ترافیک، تحقیق نموده‌اند. برای نمونه، شیخ محمدزاده در شهریور ماه ۱۳۸۳، در تحقیقی با موضوع "بررسی تحلیل طراحی شهری با تأکید بر مدل‌سازی ترافیک با استفاده از سیستم‌های اطلاعات مکانی (GIS)" به پیاده‌سازی یک سیستم اطلاعات مکانی، به کمک مدل‌سازی رفتار حرکتی مسافران شهری پرداخته است. نتایج این تحقیق نشان داد که مهمترین معضل در پیاده‌سازی موفق یک سامانه GIS بویژه در کشورهای در حال توسعه، مانند ایران، مسئله‌ی دسترسی به داده‌ها است که تلاش برای اجرای زیر ساختارهای اطلاعات مکانی (SDI)^۱ در این زمینه بسیار راهگشا است [۲]. همانطور که ملاحظه می‌شود، نبود بستری مناسب جهت تولید داده‌های ترافیکی و معضل دسترسی به این داده‌ها، موضوعی بوده است که محققان این حوزه با آن مواجه بوده‌اند که در پژوهش حاضر در ابتدا به رفع این مشکل پرداخته شده است.

یوسفیان در سال ۱۳۹۴، در تحقیقی با موضوع "ارائه یک مدل تخمین حجم جریان ترافیک به تفکیک ناوگان با

استفاده از ماشین‌های بردار پشتیبان"^۲، به مدل‌سازی حجم ترافیک در محدوده‌ی آزاد راهی و به تفکیک ناوگان عبوری پرداخته است. روش استفاده شده در این تحقیق، روش رگرسیون ماشین‌های بردار پشتیبان است. نتایج این تحقیق نشان می‌دهد که روش ماشین بردار پشتیبان، روشی با دقت بالا بوده و حجم ترافیک را با نوسان پایینی پیش‌بینی می‌کند و خطای نسبی با استفاده از این روش، دارای مقادیر کمتری نسبت به روش‌های دیگر می‌باشد [۳].

کاسم^۳ و همکاران در سال ۲۰۱۸ میلادی، پژوهشی تحت عنوان "تجزیه و تحلیل Transcad و تکنیک‌های GIS برای ارزیابی شبکه‌ی حمل و نقل در شهرستان ناصریه"^۴، برای بررسی شبکه‌ی مرکز شهر ناصریه در استان دیرکرد پرداخته است. این تحقیق برای ارزیابی جریان الگوهای شبکه‌ی ترافیکی از طریق چندین برنامه مانند GIS، GPS^۵، Transcad، به منظور جمع‌آوری داده‌های مختلف از قبیل (حجم ترافیک و سرعت جریان آزاد)، انجام گرفته است. نتیجه این پژوهش بیانگر این مطلب است که احداث جاده‌های جدید برای تغییر مسیر سفرهای خارجی، در حین اضافه کردن پل‌های جدید برای خلاص شدن از حمل و نقل که در مرکز شهر ظاهر می‌شوند، پیشنهاد می‌گردد [۴].

اصغری و همکاران در سال ۱۳۸۰، در مقاله‌ای تحت عنوان "ارائه‌ی مدل تخصیص ترافیک به شبکه‌ی حمل و نقل شهری و حل آن با استفاده از الگوریتم ژنتیک" مدلی را که در آن از الگوریتم ژنتیک با یک اندازه‌ی جمعیت و تعداد نسل مشخص برای ارائه‌ی مدل تخصیص ترافیک به شبکه‌ی حمل و نقل است، بررسی کرده‌اند. نتایج این پژوهش نشان می‌دهد که مقدار تابع هدف و وابستگی اعضای نسل‌های جدید به اعضای نسل اولیه، کاهش می‌یابد [۵].

زیانگ^۶ و همکاران در سال ۲۰۱۸ میلادی، تحقیقی را با عنوان "طبقه‌بندی وضعیت ترافیک بزرگراه‌های شهری (یک مطالعه ماشین رویکرد)"، با هدف طبقه‌بندی انجام دادند. در این پژوهش، با استفاده از روش یادگیری ماشین (یعنی استفاده از روش خوشه‌بندی FCM)^۷ برچسب‌های

^۲ Zeynab Qasim^۳ Global Positioning System^۴ Zeyang Cheng^۵ Fuzzy C-Means^۱ Spatial Data Infrastructure

مربوط به خوشه بندی، می توان نتایج طبقه بندی را تعیین نمود. از داده های جریان ترافیک شانگهای، برای هدایت روند مطالعه استفاده شده است که شامل پردازش داده ها، تجزیه و تحلیل خوشه ای و مقایسه های روشی است. نتایج همچنین نشان داد که الگوریتم خوشه بندی بهبودیافته ی FCM، هم در درجه ی عضویت فازی و هم در درجه ی وزنی بهبود پیدا کرده است [۶].

جوادزاده و همکاران در سال ۱۳۹۱، در مقاله ای تحت عنوان "انتقال دانش در محیط سیار با رویکرد پدافند غیرعامل برای پیش بینی ترافیک جاده ای"، به پیش بینی وضعیت ترافیکی گذرگاه ها با راهنمایی رانندگان در انتخاب مسیر برای اتخاذ تمهیدات لازم برای تصمیم گیری های مناسب برای روزهای آتی پرداختند. برای این پژوهش، از اطلاعات به دست آمده از شبکه سیار، شامل چندین دستگاه جی پی اس استفاده شد. نتایج این پژوهش نشان می دهد که هرگاه از یک سامانه، کاربری درخواست خود را مبنی بر پیش بینی میزان ترافیک یک مسیر ارسال نماید، با توجه به تاریخی که کاربر درخواست کرده، دانش حرکت خودروها در گذرگاه ها، دانش زمینه ی محیطی و با استفاده از به کارگیری هوش مصنوعی، اقدام به پاسخگویی مناسب به کاربر می کند [۷].

آقائزادیان در سال ۱۳۹۰، در تحقیقی با عنوان "استخراج الگوهای ترافیکی شهر calgary با استفاده از الگوریتم Fuzzy C-Means"، به بررسی چگونگی جمع آوری داده - های ترافیکی پرداخته است. در واقع در این تحقیق، الگوهای ترافیک که شامل: اطلاعات موقعیت مکانی و سرعت وسایل نقلیه است، با استفاده از روش داده کاوی به دست آمده است. از یک الگوریتم خوشه بندی به نام الگوریتم FCM، جهت یافتن نقاط مشابه و قرار دادن آن ها در خوشه ها، گروه ها و دسته های مشابه استفاده شده است. نتایج این تحقیق نشان می دهد که ترکیب روش داده کاوی و سیستم های اطلاعات مکانی برای تجزیه و تحلیل داده های حاصل از ترافیک و ذخیره و نگهداری و پردازش آن ها بسیار کارآمد می باشد [۸].

سیلوا^۱ و همکاران در سال ۲۰۱۸ میلادی، در مقاله ای با عنوان "تکنیک های هوش مصنوعی برای کنترل ترافیک شهری"، به برآورد جریان ترافیک برای جاده های شهری،

بر اساس اصول آموزش ماشین پرداخته اند. در این مقاله، به پیش بینی جریان ترافیک بر اساس google distance و داده های api^۲، داده های جریان ترافیکی و داده های هندسی پرداخته اند. در مرحله ی بعدی، برای تجزیه و تحلیل از روش رگرسیون استفاده شده است. نتایج نشان می دهد که پیش بینی جریان ترافیک برای اکثر کشورها با استفاده از مدل برآورد ترافیکی، بر مبنای k، نزدیک ترین رگرسیون همسایه و زمان سفر، داده های سرعت به دست آمده است. از google distance و داده های هندسی جاده برای به دست آوردن این اطلاعات استفاده شده است [۹].

الفار^۳ و همکاران در سال ۲۰۱۸ میلادی، بر روی موضوعی با عنوان "روش یادگیری ماشین برای پیش بینی ترافیک کوتاه مدت در یک محیط متصل"، به بررسی اختلاف ناشی از وسایل نقلیه و رابطه ی بین این اختلافات و شکل گیری شوک پرداختند. از سه تکنیک یادگیری ماشین، رگرسیون لجستیک و شبکه های عصبی برای تراکم ترافیکی کوتاه مدت برای این پیش بینی استفاده شده است. نتایج نشان دادند که مدل های مختلف ایمنی و ترافیک می توانند با الگوریتم های کنترل ترافیک یکپارچه بشوند و عملکرد آن ها را بهبود بدهند [۱۰].

علوی و همکاران در سال ۱۳۸۹، در مقاله ای با عنوان "مدلسازی مکانی تقاضای سفر مبتنی بر روشی جدید برای پیش بینی و کاهش ترافیک"، از روش علمی ترکیبی سیستم های اطلاعات جغرافیایی (GIS)، سنجش از دور (RS)^۴ و روش تحلیلی رگرسیون خطی چند متغیره برای مدلسازی مکانی ترافیک استفاده کردند. هدف در این پژوهش، پیش بینی تقاضای سفر برای کاهش ترافیک و کنترل آن است. نتایج تحقیق نشان می دهد که متغیر جمعیت به عنوان مهمترین پارامتر با ضریب ۱۴،۲۳ درصد، بیشترین درجه همبستگی را با متغیر وابسته در مدلسازی مکانی دارد. پس از آن به ترتیب، متغیرهای کاربری تجاری و تعداد کارمندان شاغل در منطقه با ضرایب ۱۱.۹- و ۳،۱۰- درصد، به ترتیب دومین و سومین پارامترهای مهم در مدلسازی مکانی تقاضای سفر هستند. متغیر تعداد واحدهای کسی، کمترین تأثیر را در این مدلسازی دارد [۱۱].

^۲ Programming Interface Application

^۳ Amr Elfar

^۴ Remote Sensing

^۱ Silva

ملایری در سال ۱۳۹۴، تحقیقی را بر روی "سامانه برآورد جریان ترافیک و تشخیص خودرو با استفاده از مدل مخلوط گوسی و شبکه عصبی مصنوعی"، انجام داد. در این تحقیق، برای هر یک از مراحل سامانه ی برآورد ترافیک، از سه روش موثر استفاده شده است که عبارت اند از: ۱. روش گوسی ۲. روش جداسازی وسیله نقلیه از سایر اشیای تصویر با استفاده از شبکه ی عصبی ۳. دنبال کردن خودرو با استفاده از جست و جوی درشت به ریز [۱۲].

همانطور که اشاره شد، مسئله ی پیش بینی رفتار ترافیک، در دوره های مختلف، همواره موضوعی بوده است که مورد توجه دغدغه مندان این حوزه بوده است. اما تحقیقات و پژوهش های انجام گرفته در این زمینه، از گذشته های دور تا کنون، فقط بر روی بررسی رفتار ترافیک متمرکز شده اند و در جهت پیش بینی آن، تحقیقات قابل توجهی انجام نگرفته است. به عبارت بهتر، در پژوهش های مطرح شده، معمولاً به بررسی رفتار ترافیک با توجه به پارامترهای تاثیرگذار بر روی آن پرداخته شده است، اما به موضوع پیش بینی ترافیک به عنوان معضلی که کلانشهرها با آن روبه رو هستند، پرداخته نشده است. به علاوه، در اکثر این تحقیقات، مشکل داده های ترافیکی نمود بسیار پررنگی دارد. به عبارت دیگر، همواره نبود بستری برای برداشت و ذخیره ی داده های ترافیکی در شبانه روز معضل پژوهشگران این حوزه بوده است. بنابراین در پژوهش حاضر، ابتدا سامانه ای جهت جمع آوری اطلاعات ترافیکی به صورت متوالی طراحی شد و سپس بر روی داده های ترافیکی جمع آوری شده، تحلیل های گوناگونی انجام شده است. بدین صورت که داده ها ابتدا در شبکه ی عصبی پایه تست شد. جهت پیش بینی ترافیک، در پژوهش حاضر، از تبدیلات موجک^۱ هم استفاده گردید. تبدیلات موجک، کاربردهای فراوانی در شاخه های علوم مهندسی بویژه هوش مصنوعی، بازشناسی الگوها و پیش بینی سری زمانی دارند. به عبارتی دیگر، موجک ها دسته ای از توابع ریاضی هستند که برای تبدیل سیگنال پیوسته به فرکانس به کار می روند. ترکیب شبکه ی عصبی با موجک، تابعی را به دست می دهد که با استفاده از آن می توان رفتار ترافیک را به صورت دقیق تر مورد ارزیابی قرار داد. با توجه به تحلیل های انجام گرفته

در این مرحله، میزان خطاها محاسبه شد. بر طبق محاسبات انجام گرفته در پژوهش حاضر، دقت به دست آمده نشان دهنده ی ۷۲ درصد مطلوبیت بین داده های واقعی ترافیک و داده های حاصل از شبکه ی آموزش دیده می باشد. در مرحله ی آخر، از الگوریتم های بهینه سازی ژنتیک، رقابت استعماری و ازدحام ذرات، جهت بهینه کردن داده های ترافیکی استفاده شد. نتایج حاصل از بهینه سازی در این مرحله نیز گویای این موضوع است که سه الگوریتم مذکور با دقت بسیار بالایی (مطلوبیت ۷۲ درصد)، جواب ها را بهینه می کند.

۲- داده ها و روش تحقیق

۲-۱- محدوده ی مورد مطالعه

برای انجام این پژوهش، نیاز به دسترسی به داده های ترافیکی محدوده ی مورد مطالعه می باشد. در این تحقیق، کلانشهر تهران به عنوان محدوده ی مطالعاتی در ۵۱ درجه و ۶ دقیقه تا ۵۱ درجه و ۳۸ دقیقه طول شرقی و ۳۵ درجه و ۳۴ دقیقه تا ۳۵ درجه و ۵۱ دقیقه عرض شمالی قرار گرفته است و ارتفاع آن از سطح آب های آزاد بین ۱۸۰۰ متر در شمال تا ۱۲۰۰ متر در مرکز و ۱۰۵۰ متر در جنوب متغیر است. تهران بین دو کوه و کویر و در دامنه های جنوبی رشته کوه البرز گسترده شده است. از جنوب به کوه های ری و بی بی شهریانو و دشت های هموار شهریار و ورامین و از شمال توسط کوهستان محصور شده است. ویژگی مهم زمین شناسی تهران، قرار گرفتن آن بین توده عظیم رشته کوه البرز (متعلق به دوران سوم زمین شناسی) و فلات ایران (متعلق به دوران چهارم زمین شناسی) است. تهران در حدود ۷۳۰ کیلومتر مربع مساحت دارد. شهرداری تهران از دیدگاه تأمین نیازمندی ها و اداره بهتر، همانگونه که در شکل ۱ نشان داده شده است، سطح شهر را به ۲۲ منطقه شهرداری و ۱۱۲ ناحیه بخش کرده که شهر ری و تجریش را نیز شامل شده است. از سوی دیگر وزارت کشور هم برای انجام انتخابات های مختلف و کارهای مدیریتی دیگر، تفسیر خاص خود را از شهرستان تهران دارد که بر ۲۲ منطقه شهرداری منطبق نیست و شمیران و ری را نیز شهرستان هایی جداگانه در نظر گرفته است. این شهر دارای نواحی هفت گانه ی مخابراتی است. در تقسیم بندی وزارت کشور شهرستان تهران در مرکز کلانشهر تهران قرار دارد. از

^۱ wavelet

<https://developers.neshan.org/> استفاده شده است. در این پلتفرم، ترافیک با چهار رنگ قرمز پر رنگ، قرمز، نارنجی و زرد نمایش داده شده است که به ترتیب معرف: ترافیک خیلی سنگین، ترافیک سنگین، ترافیک نیمه سنگین و ترافیک سبک می باشند. برای هر کدام از این ۴ رنگ، ۴ کد به اسامی ۵، ۴، ۳ و ۲ در نظر گرفته شد تا در پایگاه داده، وضعیت ترافیکی، با این اعداد ذخیره بشود. روش ذخیره‌ی داده‌ها هم به این صورت است که ابتدا از صفحه‌ی نقشه‌ی ترافیک به صورت اتوماتیک اسکرین شات گرفته شده و سپس با تحلیل‌هایی که بر روی عکس انجام می گیرد، اطلاعات ترافیکی و مختصات عکسی و زمینی نقاط ترافیکی و علاوه بر آن اطلاعات عکسی از جمله: نوع عکس (png, jpeg, tiff)، اندازه‌ی عکس (عرض و ارتفاع)، زمان و ساعت دقیق عکس گرفته شده و آدرس ذخیره‌ی عکس در پایگاه داده ذخیره می شود. در شکل ۲ و ۳، به ترتیب نمونه ای از جدولی که شامل اطلاعات ترافیکی است، و جدولی دیگر که شامل اطلاعات عکسی است، مشاهده می کنید.

+ Options														
← T →					id	color	x	y	xasli	yasli	Xzamini	Yzamini	date	screenshot
					1	3	270	77	0.071172	0.0202972	532887.26680688	3960924.6576749	2020-07-28 12:29:15	117
					2	3	271	77	0.0714356	0.0202972	532889.20726905	3960924.6650564	2020-07-28 12:29:15	117
					3	3	272	77	0.0716992	0.0202972	532891.14773121	3960924.6724378	2020-07-28 12:29:15	117
					4	3	273	77	0.0719628	0.0202972	532893.08819338	3960924.6798192	2020-07-28 12:29:15	117
					5	3	274	77	0.0722264	0.0202972	532895.02865554	3960924.6872006	2020-07-28 12:29:15	117
					6	3	275	77	0.07249	0.0202972	532896.96911771	3960924.6945821	2020-07-28 12:29:15	117

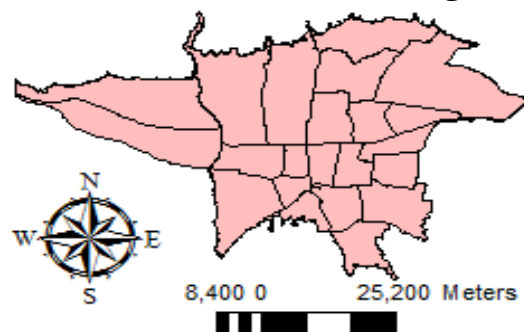
شکل ۲- نمونه ای از جدول شامل ذخیره‌ی اطلاعات ترافیکی و مختصات آن‌ها

+ Options													
				sid	width	height	type	zoom	paths			date	screenshot
☐	 Edit	 Copy	 Delete	1	1342	768	resource	16	http://screenly.com/storage/5ebf7c068a3cf_58ucjW...			2020-05-16 10:07:21	124
☐	 Edit	 Copy	 Delete	2	1342	768	resource	16	http://screenly.com/storage/5ebf7c0689687_PKw3Xi...			2020-05-16 10:07:21	125
☐	 Edit	 Copy	 Delete	3	1342	768	resource	16	http://screenly.com/storage/5ebf7c0688aa3_xFVKcN...			2020-05-16 10:07:21	121
☐	 Edit	 Copy	 Delete	4	1342	768	resource	16	http://screenly.com/storage/5ebf7c06b9b5e_Y0DwZ...			2020-05-16 10:07:21	123
☐	 Edit	 Copy	 Delete	5	1342	768	resource	16	http://screenly.com/storage/5ebf7c07c5224_jKdLAw...			2020-05-16 10:07:28	122
☐	 Edit	 Copy	 Delete	6	1342	768	resource	16	http://screenly.com/storage/5ebf7c4451c4e_rzf5DL...			2020-05-16 10:08:22	146

شکل ۳- نمونه ای از جدول شامل ذخیره‌ی اطلاعات عکسی برداشت شده با زوم مشخص

را به فرمت اکسل از پایگاه داده فراخوانی کرده و ذخیره می کند. از آنجایی که حجم داده‌ها خیلی زیاد است، جهت استخراج داده‌ها، باید برنامه ای نوشته بشود و در آن شرطی تعریف گردد که داده‌ها را در حجم مشخصی برداشت و ذخیره بکند.

شمال به شهرستان‌های کرج و شمیرانات، از مشرق به شهرستان دماوند، از جنوب به شهرستان‌های ورامین و ری و اسلامشهر و از مغرب به شهرستان‌های شهریار و کرج محدود می‌شود. مرکز آن بخش مرکزی است.



شکل ۱- مناطق ۲۲ گانه‌ی کلانشهر تهران مورد مطالعه در پژوهش با مقیاس ۱:۵۰۰۰۰

در پژوهش حاضر، جهت پیش بینی ترافیک، در ابتدا نیاز به داده‌های ترافیکی می‌باشد. برای رسیدن به این منظور، ابتدا سامانه ای طراحی شد که بتواند داده‌های ترافیک را جمع آوری کند. با استفاده از روش‌های برنامه نویسی در محیط‌های مختلف، کدی نوشته و جدولی در پایگاه داده تنظیم شد. لازم به ذکر است که برای این هدف از پلتفرم نقشه نشان به آدرس

۲-۲- استخراج داده‌های ترافیکی جمع‌آوری شده از پایگاه داده

در این مرحله، داده‌های اطلاعات ترافیکی که در پایگاه داده به صورت برخط و متوالی ذخیره شده است، باید استخراج بشود. برای این مهم، کدی نوشته شد که داده‌ها

۳-۲- استاندارد سازی داده های ترافیکی

داده های ترافیکی استخراج شده، برای استفاده و تحلیل در شبکه های عصبی و بهینه سازی با الگوریتم های مختلف، نیاز دارند که تغییر داده بشوند. تغییر دادن به این معنا که فیلد مربوط به زمان ذخیره داده های ترافیکی، باید شکسته شود و به چهار ستون مجزا، تقسیم بندی گردد. فایل اکسل خروجی در این مرحله، شامل ۸ ستون خواهد بود که دو ستون اول شامل مختصات طول و عرض

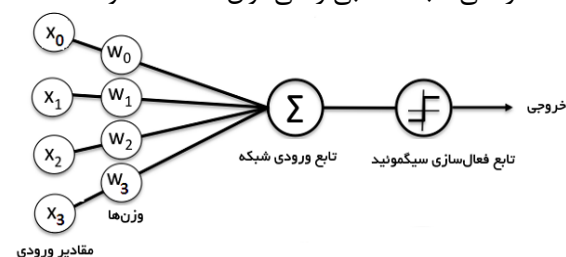
زمینی نقاط برحسب UTM، ستون سوم شامل سال ذخیره، ستون چهارم ماه ذخیره، ستون پنجم روز ذخیره، ستون ششم زمان ذخیره داده های ترافیکی که شامل ساعت و دقیقه و ثانیه، ستون هفتم روزهای هفته، و ستون هشتم رنگ های ترافیکی برداشتی می باشند. در جدول ۱ نمونه ای از این داده ها را مشاهده می کنید.

جدول ۱- نمونه ای از خروجی اکسل داده های ذخیره شده در پایگاه داده پس از اعمال تغییرات لازم

ترافیک برداشتی	روز به هفته	ساعت	روز	ماه	سال	عرض جغرافیایی	طول جغرافیایی
۳	سه شنبه	۹:۳۶:۱۰	۸	۱۲	۱۳۹۸	۳۹۵۲۹۳۷	۵۱۶۹۳۵،۵۳۷۹
۳	سه شنبه	۹:۳۶:۱۰	۸	۱۲	۱۳۹۸	۳۹۵۶۶۲۸	۵۱۷۹۳۶،۹۵۷۸
۴	سه شنبه	۹:۴۶:۱۰	۸	۱۲	۱۳۹۸	۳۹۵۶۶۲۸	۵۱۷۹۶۸،۰۱۸
۳	سه شنبه	۹:۴۶:۱۰	۸	۱۲	۱۳۹۸	۳۹۵۶۶۲۸	۵۱۷۹۹۹،۰۷۸۲

۴-۲- پیاده سازی اطلاعات ترافیکی در شبکه ی عصبی پایه

در پژوهش حاضر، جهت پیش بینی ترافیک، از داده های ترافیکی ذخیره شده در پایگاه داده استفاده می شود. در این مرحله با استفاده از مفاهیم شبکه های عصبی، اقدام به تحلیل داده های ترافیکی می گردد. مغز انسان در خود تعداد بسیار زیادی از نورون ها را جای داده است تا اطلاعات مختلف را پردازش کرده و جهان اطراف را بشناسد. به صورت ساده، نورون ها در مغز انسان اطلاعات را از نورون های دیگر به وسیله ی دندرویدها می گیرند. کار نورون ها این است که ورودی ها را جمع آوری می کند و اگر مقادیر این ورودی های جمع آوری شده، از حد آستانه فراتر برود، نورون فعال شده و از طریق آکسون ها به نورون های دیگر منتقل می شود. در علوم کامپیوتر به این فرآیند شبکه ی عصبی می گویند [۱۳]. در شکل ۴، کارکرد و ساختار کلی شبکه عصبی را می توان مشاهده نمود.



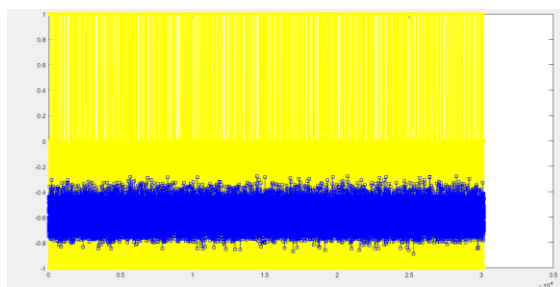
شکل ۴- نمودار مربوط به کارکرد کلی شبکه های عصبی

در شکل ۴، x ها، ورودی ها هستند که از طریق جمع آوری داده ها، حاصل شده اند. w ها، وزن ها هستند که همانطور که در شکل ۴ مشخص است، هر ورودی در یک وزن ضرب می شود. حاصل ضرب ورودی ها در وزنشان، باید با هم جمع بشوند که حاصلشان در سیگما قرار می گیرد. تابع فعال سازی، تابعی است که اعمالی روی سیگما انجام می دهد تا خروجی را تولید بکند. جهت تحلیل داده های ترافیکی موجود، از کارکردهای شبکه های عصبی استفاده شده است. مراحل کاری به این صورت است که ابتدا داده های ورودی که استاندارد سازی شده است، فراخوانی می شود. باید به این نکته توجه داشت که جهت تحلیل های لازم در شبکه های عصبی، نیاز است که اطلاعات نرمالیزه شوند.

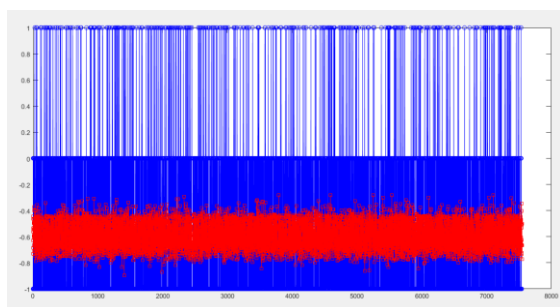
۵-۲- عملیات نرمال سازی داده ها

یکی از اصول علم پایگاه داده ها، از بین بردن افزونگی است. افزونگی به این معناست که یک داده خاص در چند محل مختلف پایگاه داده ذخیره شود. این امر موجب می شود که این خطر بالقوه به وجود آید که داده ها هر لحظه با هم در تضاد قرار گیرند و استخراج واقعیت از آن ها غیرممکن شود. به بیان دیگر، نرمال سازی، فرایندی است که بر اساس آن داده ها و اطلاعات در واحدهای منطقی به نام جدول به شکلی توزیع می شود که علاوه بر حفظ

در رابطه‌ی ۱، جذر مقادیر خطای ms (خطای کم‌ترین مربعات)، را گرفته و حاصل را بر ۲ تقسیم کرده و در نهایت برای محاسبه‌ی خطا بر حسب درصد، حاصل را در ۱۰۰ ضرب کرده تا مقدار خطا به دست بیاید. نمودار مربوط به داده‌های آموزشی و ارزیابی به ترتیب، در شکل ۵ و شکل ۶ نشان داده شده است.



شکل ۵- میزان تطابق داده‌های آموزشی با مقادیر واقعی داده‌ها



شکل ۶- میزان تطابق داده‌های تست با مقادیر واقعی داده‌ها

در هر دو شکل ۵ و ۶، محور x و y داده‌های ورودی می‌باشند که به صورت یک ماتریس تعریف می‌شود، محور x شامل ورودی‌های شبکه عصبی می‌باشد و محور y هم به عنوان خروجی شبکه عصبی می‌باشد. این خروجی همان رنگ ترافیک می‌باشد که به صورت نرمال محاسبه شده است. همانطور که در این شکل‌ها مشهود است، داده‌های پیش‌بینی شده که شامل رنگ‌های ترافیکی می‌باشد در شکل ۵ به رنگ آبی و در شکل ۶ به رنگ زرد نشان داده شده است. در شکل ۵، داده‌های زرد رنگ به عنوان داده‌های آموزشی و تست به شبکه داده شده‌اند و در شکل ۶، داده‌های آبی رنگ به عنوان داده‌های آموزشی و تست هستند. همان‌طور که مشخص است، شبکه‌ی آموزش داده شده، با مطلوبیت ۷۲ درصد و خطای ۲۸ درصد داده‌ها را برآورد کرده است. نمودار مربوط به داده‌های اجرایی در این مرحله، در شکل ۷ نشان داده شده است. همانطور که در شکل ۷ مشخص است، میزان تطابق داده‌های آموزشی و بهترین برآورد که به

موجودیت داده‌ها از ایجاد پدیده افزونگی جلوگیری به عمل می‌آورد. به این منظور فرم‌های نرمال متعددی تعریف و مورد استفاده قرار می‌گیرد. برای نرمال‌سازی یک جدول، لازم است ابتدا در فرم اول نرمال شده، سپس در فرم فرم‌های بعدی بررسی گردد، در هنگام نرمال‌سازی جداول، تعداد جداول در پایگاه داده افزایش می‌یابد. عملیات نرمال‌سازی قبل از بسیاری از الگوریتم‌های داده‌کاوی مانند شبکه‌های عصبی، SVM^۱، KNN^۲ و KMeans بایستی انجام بگیرد تا ابعاد مختلف به صورت عادلانه توسط الگوریتم بررسی شوند و تأثیر، یکی بیشتر از بقیه نباشد [۱۴]. برای رسیدن به این مقصود، ابتدا با استفاده از دستورات موجود، داده‌های ترافیکی نرمالیزه شدند. نرمالیزه شدن به این مفهوم که همه‌ی داده‌ها در بازه‌ی اعداد بین ۱- و ۱ قرار داده شدند. برای انجام این مرحله، از حدود ۳۸۰۰۰ داده‌ی ترافیکی جمع‌آوری شده استفاده شد. در هنگام آموزش، داده‌ها را به دو دسته‌ی آموزشی و تست، می‌توان تقسیم بندی کرد. حال الگوریتم از روی داده‌های آموزشی یادگیری را انجام می‌دهد و از روی داده‌های ارزیابی یا همان داده‌های تست می‌توان فهمید که الگوریتم و مدل ساخته شده توسط آن، چقدر دقت داشته است. وقتی الگوریتم، عملیات یادگیری را انجام داد و در واقع یک مدل را از روی این داده‌ها ساخت، سپس می‌توان از روی این مدل، عملیات داده‌کاوی را بر روی داده‌های جدید انجام داد. در پژوهش حاضر، ۸۰ درصد داده‌ها، داده‌های آموزشی و ۲۰ درصد داده‌های ترافیکی به عنوان داده‌های تست، مورد استفاده قرار گرفتند. ورودی‌های شبکه‌ی عصبی، ۷ ستون اول داده‌های ترافیکی هستند که همان‌طور که اشاره شد شامل طول و عرض زمینی نقاط، تاریخ و ساعت و روز برداشت داده‌ها و خروجی شامل رنگ ترافیکی می‌باشد.

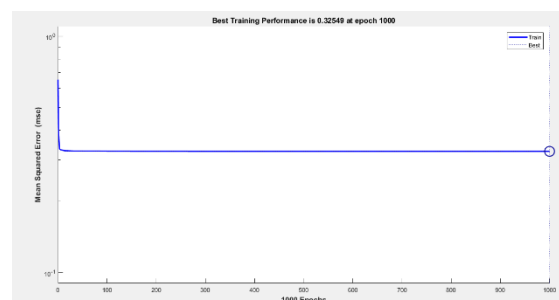
با اجرای این قسمت، خطای $msets=0.3222$ و $msetr=0.3225$ به دست می‌آیند. با استفاده از رابطه‌ی ۱، میزان خطا بر حسب درصد، برای هر کدام از داده‌های آموزشی و ارزیابی ۲۸،۵۲۶۳ درصد و ۲۸،۳۸۱۳ درصد به دست می‌آید.

$$(1) \quad \frac{\sqrt{ms}}{2} \times 100$$

^۱ Support vector machines

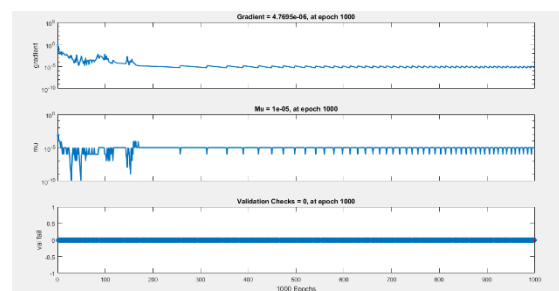
^۲ k-nearest neighbors algorithm

ترتیب با خطوط آبی و خط چین نشان داده شده است، تطابق خیلی نزدیکی با یکدیگر دارند و تقریباً نمودار این دو برآورد، با یکدیگر مماس شده اند.



شکل ۷- نمودار میزان تطابق داده‌های آموزشی و داده‌های برآورد شده

در میانگین و میزان گرادیان محاسبه می‌شود. در شکل شماره ۸، داده‌های اعتبارسنجی شده، مقدار میانگین و گرادیان، نشان داده شده است.



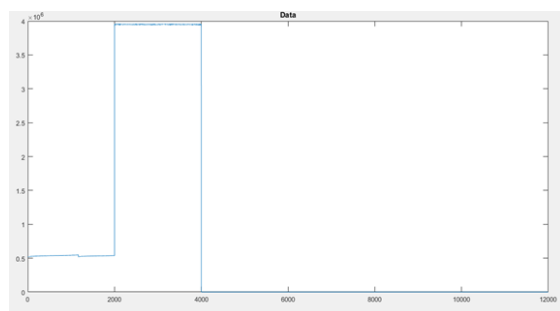
شکل ۸- نمودار نشان دهنده مقادیر داده‌های اعتبارسنجی، میانگین و گرادیان برای ۱۰۰۰ آیتم و دوره

۲-۶- پیاده‌سازی داده‌های ترافیکی در شبکه‌ی عصبی موجک

موجک، دسته‌ای از توابع ریاضی هستند که در آن، قدرت تفکیک هر مؤلفه برابر با مقیاس آن است. تبدیل موجک تجزیه یک تابع، بر مبنای توابع موجک می‌باشد. موجک‌ها، نمونه‌های انتقال یافته و مقیاس شده یک تابع (موجک مادر) با طول متنهای و نوسانی شدیداً میرا هستند. دلیل استفاده از شبکه‌ی عصبی موجک در پژوهش حاضر، قدرت پیش‌بینی و دقت حاصل از این شبکه در موضوع یادگیری است. چرا که شبکه‌ی عصبی موجک، یک مدل ترکیبی از از تئوری فازی و تئوری موجک است.

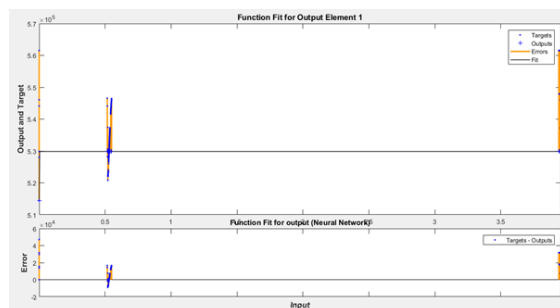
ترکیب شبکه عصبی مصنوعی و موجک، الگوی جدیدی از هوش مصنوعی بنام شبکه عصبی موجکی را پدید آورده است. در بسیاری از موارد، تبدیلات موجک، نتایج بسیار

بهتری از شبکه‌های عصبی مصنوعی را به وجود می‌آورند. معمولاً شبکه‌های عصبی بخصوص شبکه‌های عصبی موجکی را به منظور تقریب توابع غیر خطی ناشناخته و شبیه‌سازی رفتار یک سیستم نامشخص و پیش‌بینی رفتار آن در زمان‌های آینده بکار می‌برند. یک تابع موجک، از نظر زمانی، به طرز مناسبی باید محدود باشد، دارای میانگین صفر باشد و نرم غیر صفر داشته باشد. برای تحلیل داده‌های ترافیکی، با استفاده از تبدیلات موجک، ابتدا داده‌های ترافیکی را که شامل ۲۰۰۰ داده‌ی ورودی است، فراخوانی کرده و ۷۰ درصد داده‌ها به عنوان داده‌های آموزشی، ۱۰ درصد داده‌ها، به عنوان داده‌های اعتبارسنجی و ۲۰ درصد داده‌ها به عنوان داده‌های تست در نظر گرفته شد. با اجرای این برنامه، نمودار داده‌های ترافیکی برای ۲۰۰۰ داده‌ی ورودی به صورت شکل ۹ می‌باشد.



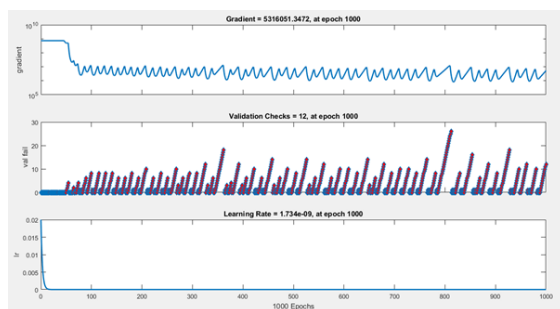
شکل ۹- نمودار داده‌ها بر حسب مقادیر ورودی در تابع موجک

برای اینکه بتوان مقدار تطابق داده‌های واقعی ترافیک را با مقادیر خروجی پیش‌بینی شده، و تفاوت این دو را برآورد کرد، برنامه‌ای نوشته و نمودار مربوط به آن ترسیم شد. میزان تطابق این مقادیر ورودی و پیش‌بینی شده با استفاده از تبدیل موجک، در شکل ۱۰ نشان داده است.



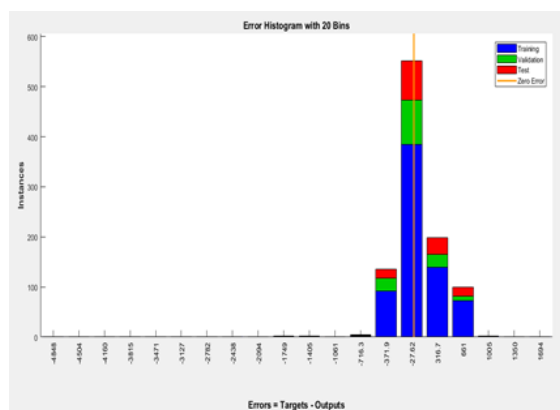
شکل ۱۰- نمودار میزان تطابق و خطا برای داده‌های ترافیکی ورودی و پیش‌بینی شده

همانگونه که در شکل ۱۰ نشان داده شده است، داده‌های آبی رنگ دایره‌ای شکل، تارگت‌ها یا به اصطلاح، داده‌های واقعی رنگ‌های ترافیکی و تارگت‌های آبی رنگی که با علامت



شکل ۱۲- نمودار مربوط به مقادیر اعتبارسنجی و گرادیان در تبدیلات موجک

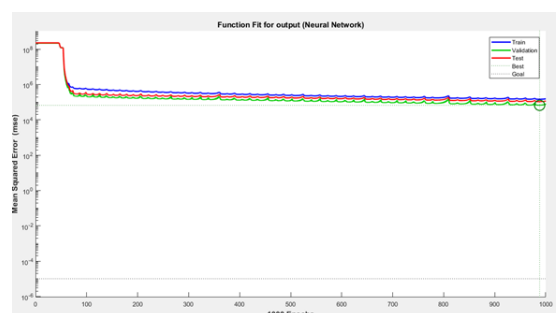
اما برای تحلیل و آموزش با استفاده از تبدیلات موجکی، می‌توان نمودار هیستوگرام را هم ترسیم کرد تا به صورت شهودی، مقادیر هر یک از داده‌ها، و میزان خطای حاصل را برای برنامه‌ی نوشته شده محاسبه و ترسیم کند (شکل ۱۳).



شکل ۱۳- نمودار هیستوگرام تبدیلات موجک

همانطور که در شکل ۱۳ مشخص است، بیشترین داده‌ها، مربوط به داده‌های آموزشی هستند که در نمودار، به رنگ آبی نشان داده شده است. رنگ‌های سبز و قرمز نیز به ترتیب، نشان دهنده‌ی داده‌های اعتبارسنجی و داده‌های تست هستند. خط نارنجی رنگ که در شکل ۱۳ پدیدار است، نمایانگر میزان خطا است. به عبارتی دیگر، مقدار خطا در این مرحله از پژوهش، در نمودار هیستوگرام به صورت اختلاف بین تارگت‌ها (داده‌های ورودی) و خروجی‌ها می‌باشد. در شکل ۱۴، چهار نمودار مجزا برای هر یک از داده‌های آموزشی، تست، اعتبارسنجی و حاصل جمع این سه داده نشان داده شده است. همانطور که در شکل ۱۴ مشاهده می‌شود، مقادیر داده‌های ترافیکی با دقت بسیار بالایی بر روی بهترین مقادیر (مقادیر فیت) برآورد شده است. مقادیر داده‌های خروجی نیز همانطور که

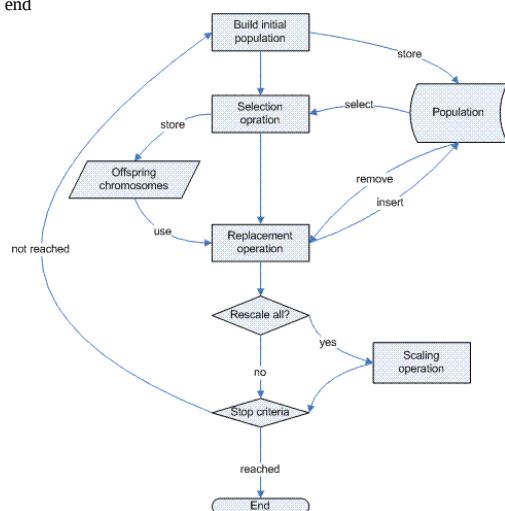
(+) نشان داده شده اند، مقادیر خروجی یا همان داده‌های ترافیکی پیش بینی شده هستند. همان‌گونه که در شکل ۱۰ هم مشهود است، مقادیر داده‌های واقعی و داده‌های محاسبه شده تقریباً روی هم افتاده اند و قسمت‌های نارنجی رنگ در نمودار شکل ۱۰، میزان اختلاف و خطا بین این دو داده را نشان می‌دهد که نتایج قابل توجه و خوبی را کسب کرده است. بدیهی است که هر مقدار، تعداد داده‌های ترافیکی ورودی برای آموزش شبکه بیشتر باشد، طبیعتاً نتایج بهتری را هم در پی خواهد داشت. ولی یکی از نکات قابل ذکر در رابطه با تبدیلات موجک این است که این تبدیلات برای داده‌های حجیم بسیار محاسبات کند و طولانی مدت دارند. به همین منظور، برای انجام این پژوهش، از ۲۰۰۰ داده استفاده شده است. یکی از نمودارهای بسیار مهم دیگر که در این مرحله از پژوهش ترسیم و مورد تحلیل قرار گرفت، در شکل ۱۱ نشان داده شده است. در شکل ۱۱، نمودار مربوط به داده‌های آموزشی، داده‌های تست و داده‌های اعتبارسنجی ترسیم گردیده است و در این میان بهترین حالت هم روی نمودار مشخص شده است. همانطور که در شکل ۱۱ مشاهده می‌شود، خطوط آبی مربوط به داده‌های آموزشی، خطوط سبز مربوط به داده‌های اعتبارسنجی، خطوط قرمز مربوط به داده‌های تست و خط چین‌ها مربوط به بهترین داده‌ها می‌باشند. قابل ملاحظه است که این خطوط، اختلاف بسیار اندکی با یکدیگر دارند، به طوری که میزان خطای پژوهش را به طرز قابل ملاحظه‌ای بسیار پایین آورده اند.



شکل ۱۱- نمودار مربوط به مقادیر خروجی(رنگ‌های ترافیکی پیش بینی شده) با تبدیلات موجک

آنچه که در این پژوهش، اهمیت دارد، عملکرد شبکه-ی آموزش داده شده و میزان یادگیری آن در یک دوره‌ی ۱۰۰۰ تایی است. میزان این یادگیری و گرادیان مربوطه، در شکل ۱۲ نشان داده شده است.

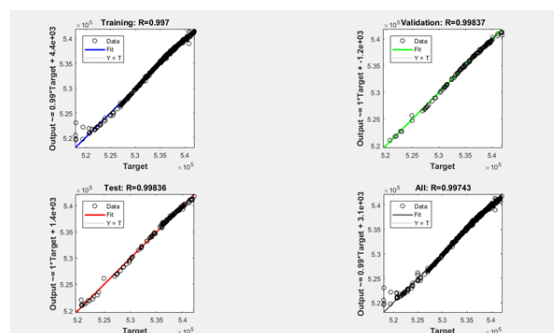
- 3 Choose initial population
- 4 repeat
- 5 Evaluate the individual fitness of a certain proportion of the population
- 6 Select pairs of best-ranking individuals to reproduce
- 7 Apply crossover operator
- 8 Apply mutation operator
- 9 until terminating condition
- 10 end



شکل ۱۵- شمای کلی الگوریتم ژنتیک

در الگوریتم‌های بهینه‌سازی مانند الگوریتم ژنتیک، همواره یک تابعی به عنوان تابع هدف یا تابع هزینه تعریف می‌شود. کار تابع هدف این است که یک سری محاسبات بر روی سطرها و ستون‌های ماتریس ورودی انجام می‌دهد و به عنوان مثال برای هر سطر ماتریس، بهترین را انتخاب می‌کند. دلیل استفاده از الگوریتم ژنتیک در پژوهش حاضر، این است که تمامی پاسخ‌ها در این الگوریتم دقیق هستند و هیچ مقادیر تقریبی در آن وجود ندارد. علاوه بر این، الگوریتم ژنتیک از لحاظ آموزش بسیار ساده است و سریع و موثر عمل می‌کند. در این مرحله، فایل اکسل داده‌های ترافیکی با حدود ۳۸۰۰۰ رکورد فراخوانی شده است. در واقع ماتریسی، با ۳۸۰۰۰ سطر و تعداد هشت ستون. تابع هدف بر روی ۳۸۰۰۰ سطر ماتریس ورودی، محاسبات را انجام داده و در نهایت یک ماتریسی که حاصل تابع هزینه است، با تعداد ۳۸۰۰۰ سطر و یک ستون تولید می‌کند. با انجام این کار، جمعیت جدیدی تولید خواهد شد که در پژوهش حاضر حدود ۱۰ درصد از داده‌های ورودی به طور مستقیم از جمعیت اولیه به صورت کاملاً رندم و تصادفی به نسل بعد منتقل شدند. ۷۰ درصد داده‌ها حاصل از عمل crossover (عمل تقاطع) بر روی جمعیت اولیه، به نسل بعد منتقل شدند و مابقی جمعیت (یعنی حدود ۲۰ درصد) از طریق عمل mutation (عمل جهش)، به نسل بعد منتقل می‌شوند.

در شکل ۱۴ مشخص است، بر طبق رابطه ای ریاضی برای هر کدام از چهار نمودار محاسبه و بیان می‌گردد.



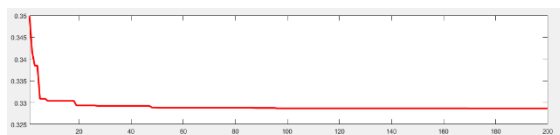
شکل ۱۴- نمودارهای مربوط به داده‌ها و میزان تطابق آن‌ها با بهترین حالت ممکن

۲-۷- بهینه‌سازی داده‌های ترافیکی با الگوریتم بهینه‌سازی ژنتیک

الگوریتم ژنتیک^۱، تکنیک جستجو در علم رایانه برای یافتن راه‌حل تقریبی برای بهینه‌سازی مدل ریاضی و مسائل جستجو است. الگوریتم ژنتیک، نوع خاصی از الگوریتم‌های تکاملی است که از تکنیک‌های زیست‌شناسی برگشتی مانند وراثت، جهش زیست‌شناسی و اصول انتخابی داروین برای یافتن فرمول بهینه جهت پیش‌بینی یا تطبیق الگو استفاده می‌شود. الگوریتم‌های ژنتیک، اغلب گزینه خوبی برای تکنیک‌های پیش‌بینی بر مبنای رگرسیون هستند. در مدل‌سازی، الگوریتم ژنتیک، یک تکنیک برنامه‌نویسی است که از تکامل ژنتیکی به عنوان یک الگوی حل مسئله استفاده می‌کند. مسئله‌ای که باید حل شود دارای ورودی‌هایی می‌باشد که طی یک فرایند الگو برداری شده از تکامل ژنتیکی به راه‌حل‌ها تبدیل می‌شود. سپس راه‌حل‌ها به عنوان کاندیداها توسط تابع ارزیابی^۲ مورد ارزیابی قرار می‌گیرند و چنانچه شرط خروج مسئله فراهم شده باشد الگوریتم خاتمه می‌یابد. به‌طور کلی یک الگوریتم، مبتنی بر تکرار است که اغلب بخش‌های آن به صورت فرایندهای تصادفی انتخاب می‌شوند که این الگوریتم‌ها از بخش‌های تابع برازش، نمایش، انتخاب و تغییر، تشکیل می‌شوند [۱۵]. شمای کلی شبه کد الگوریتم ژنتیک به صورت شکل ۱۵ است.

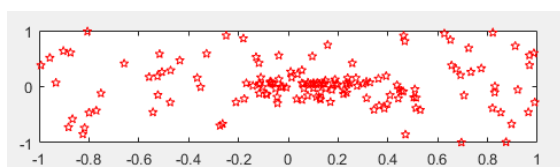
¹ Genetic Algorithm
² begin

^۱ Genetic algorithm
^۲ Fitness Function



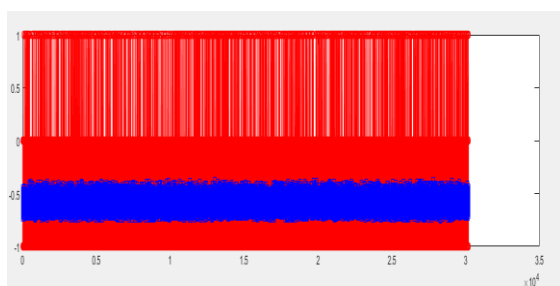
شکل ۱۶- نمودار اجرای تابع هزینه برای الگوریتم ژنتیک در شبکه عصبی پایه برای نمونه ی جمعیت ۲۰۰ تایی

در شکل ۱۶، محور طول‌ها تعداد تکرار الگوریتم می باشد. در این پژوهش، جمعیت ۲۰۰ تایی در نظر گرفته شده است و محور عرض‌ها بیانگر مقادیری است که تابع هدف به خود اختصاص داده است. همان طور که در شکل ۱۶ مشهود است مقادیر تابع هزینه از ۰.۳۵ شروع شده و در انتهای اجرای الگوریتم به مقدار تقریباً ۰.۳۲۷ رسیده است. می توان نحوه ی اجرای الگوریتم ژنتیک را به صورت دو بعدی در شکل ۱۷ مشاهده نمود.

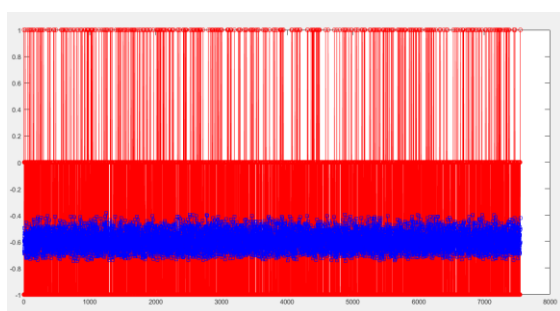


شکل ۱۷- نحوه ی اجرای الگوریتم ژنتیک به صورت دو بعدی

همانگونه که در شکل ۱۷ مشاهده می شود، الگوریتم به تدریج به سمت بهتر شدن و در واقع تولید جمعیت نسل بعد حرکت می کند. نمودار مربوط به میزان تطابق داده‌های واقعی با مقادیر آموزشی و تست در شکل‌های ۱۸ و ۱۹ نمایش داده شده است.



شکل ۱۸- نمودار تطابق مقادیر واقعی داده‌های ترافیکی با داده‌های آموزشی با بهینه ساز الگوریتم ژنتیک



شکل ۱۹- نمودار تطابق مقادیر واقعی داده‌های ترافیکی با داده‌های تست با بهینه ساز الگوریتم ژنتیک

عمل تقاطع در این پژوهش، خود شامل دو بخش اصلی است. بخش انتخاب والدین و بخش تولید فرزندان از روی والدین. برای انجام این کار، روش های مختلفی وجود دارد. به عنوان نمونه می توان این دو بخش را به صورت کاملاً تصادفی از روی جمعیت اولیه انتخاب و به نسل بعد منتقل کرد. اما روش دیگری که در این مرحله وجود دارد، انتخاب جمعیت والدین و تولید فرزندان از طریق انتخاب آن‌ها به صورت تصادفی وزن دار است. بدین مفهوم که آن‌هایی که دارای موقعیت بهتری هستند، با احتمال بیشتری انتخاب بشوند. در تحقیق حاضر، ۷۰ درصد جمعیت به روش تصادفی وزن دار انتخاب شدند. این روش خود شامل چندین روش متفاوت است که از میان آن‌ها، بهترین روش یعنی روش تقاطع وزن دار به کار گرفته شده است. یعنی از جمعیت اولیه، آن‌هایی انتخاب شدند که تابع هزینه ی آن‌ها مقدار بیشتری دارد. بعد از انتخاب والدین، نوبت تولید فرزند از والدین می باشد. بنابراین جمعیت جدید تولید می شود و نیاز است که دوباره برای این جمعیت جدید، تابع هزینه ی جدید تعریف و دوباره بهترین‌ها به نسل بعد منتقل بشود. این عمل تا زمانی که به شرط مورد نیاز مسئله ی ما برسد ادامه پیدا می کند تا در مرحله آخر، الگوریتم متوقف و جمعیت نهایی که شامل بهترین‌های نسل، همراه با تابع هزینه ی آن‌ها، تولید خواهد شد. در گام بعدی، برای انجام این عمل، برنامه ی شبکه عصبی پایه استفاده می شود. تابع هزینه مربوط به الگوریتم بهینه سازی ژنتیک به این قسمت برنامه، اضافه می گردد، سپس برنامه اجرا می شد تا نقطه ای به عنوان بهترین و بهینه ترین نقطه برای این داده‌های ترافیکی به دست آید. تابع هزینه در پژوهش حاضر به صورت رابطه ی ۲ تعریف شده است.

$$\text{Cost_ANN_EA}(XX, X_{tr}, Y_{tr}, \text{Network}) \quad (2)$$

در رابطه ی ۲، xx تعداد کل جمعیت اولیه، x_{tr}, y_{tr} به ترتیب تعداد ستون‌ها و سطرها ی ماتریس آموزشی و $network$ هم شبکه ی عصبی آموزش داده شده به صورت ساده است. کاری که باید انجام شود این است که تابع هدف تعریف شده برای این پژوهش را در برنامه ی الگوریتم ژنتیک قرار داده تا جمعیت جدید را برای ما تولید بکند. نمودار مربوط به اجرای این الگوریتم در شکل ۱۶ قابل مشاهده است.

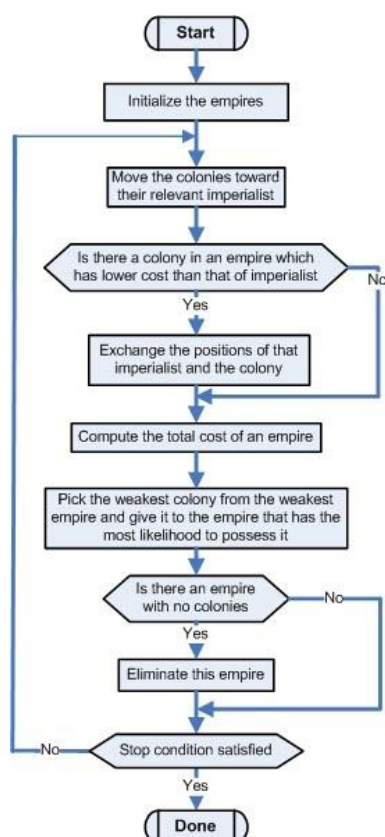
مانند مرحله‌ی قبل، در هر دو شکل ۱۸ و ۱۹ محور x و y داده‌های ورودی می‌باشند که به صورت یک ماتریس تعریف می‌شود، محور x شامل ورودی‌های شبکه عصبی می‌باشد و محور y هم به عنوان خروجی شبکه عصبی می‌باشد. با اجرای این برنامه، مقادیر $MSE_{tr}=0,3287$ و $MSE_{ts}=0,3422$ به دست می‌آیند. برای محاسبه‌ی مقادیر خطا و تطابق داده‌های واقعی ترافیکی و داده‌های پیش‌بینی شده، مقادیر خطاهای کمترین مربعات را در رابطه‌ی ۱، قرار داده تا مقدار خطا بر حسب درصد محاسبه گردد.

اگر مقادیر MSE_{tr} و MSE_{ts} در رابطه‌ی ۱ قرار داده شود، آنگاه مقادیر خطا برای آن‌ها به ترتیب $28,6662$ و $29,2489$ درصد محاسبه می‌شود. یعنی با مطلوبیت تقریباً بیشتر از هفتاد درصد بین داده‌های ترافیکی موجود و داده‌های ترافیکی پیش‌بینی شده توسط شبکه‌ی آموزش دیده با الگوریتم زنتیک.

۲-۸- بهینه‌سازی داده‌های ترافیکی با الگوریتم بهینه‌سازی رقابت استعماری

الگوریتم رقابت استعماری^۱، روشی در حوزه محاسبات تکاملی است که به یافتن پاسخ بهینه مسائل مختلف بهینه‌سازی می‌پردازد. این الگوریتم با مدل‌سازی ریاضی فرایند تکامل اجتماعی-سیاسی، الگوریتمی برای حل مسائل ریاضی بهینه‌سازی ارائه می‌دهد. از لحاظ کاربرد، این الگوریتم، در دسته الگوریتم‌های بهینه‌سازی تکاملی همچون الگوریتم‌های ژنتیک (Genetic Algorithms)، روش بهینه‌سازی ازدحام ذرات، الگوریتم کلونی مورچگان (Ant Colony Optimization)، الگوریتم تبرید شبیه‌سازی شده (Simulated Annealing)، الگوریتم تکامل تفاضلی (Differential Evolution)، الگوریتم فرهنگی (Cultural Algorithm)، الگوریتم ممیتیک (Memetic Algorithm)، الگوریتم زنبورها (Bees Algorithm)، الگوریتم بهینه‌سازی کاوش مبتنی بر باکتری (Bacterial Foraging Optimization Algorithm) و غیره قرار می‌گیرد. همانند همه الگوریتم‌های قرار گرفته در این دسته، الگوریتم رقابت استعماری نیز مجموعه اولیه‌ای از جواب‌های احتمالی را تشکیل می‌دهد. این جواب‌های اولیه در الگوریتم ژنتیک با عنوان «کروموزوم»، در الگوریتم ازدحام

ذرات با عنوان «ذره» و در الگوریتم رقابت استعماری نیز با عنوان «کشور» شناخته می‌شوند. الگوریتم رقابت استعماری با روند خاصی که در ادامه می‌آید، این جواب‌های اولیه (کشورها) را به تدریج بهبود داده و در نهایت جواب مناسب مسئله بهینه‌سازی (کشور مطلوب) را در اختیار می‌گذارد. پایه‌های اصلی این الگوریتم را سیاست همسان‌سازی^۲، رقابت استعماری^۳ و انقلاب^۴ تشکیل می‌دهند. این الگوریتم با تقلید از روند تکامل اجتماعی، اقتصادی و سیاسی کشورها و با مدل‌سازی ریاضی بخش‌هایی از این فرایند، عملگرهایی را در قالب منظم به صورت الگوریتم ارائه می‌دهد که می‌توانند به حل مسائل پیچیده بهینه‌سازی کمک کنند. در واقع این الگوریتم جواب‌های مسئله بهینه‌سازی را در قالب کشورها نگریسته و سعی می‌کند در طی فرایندی تکرار شونده این جواب‌ها را رفته رفته بهبود داده و در نهایت به جواب بهینه مسئله برساند [۱۶]. شمای کلی الگوریتم رقابت استعماری در شکل ۲۰ نمایش داده شده است.

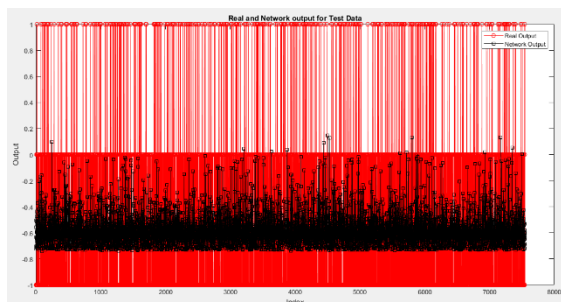


شکل ۲۰- فلوچارت الگوریتم رقابت استعماری

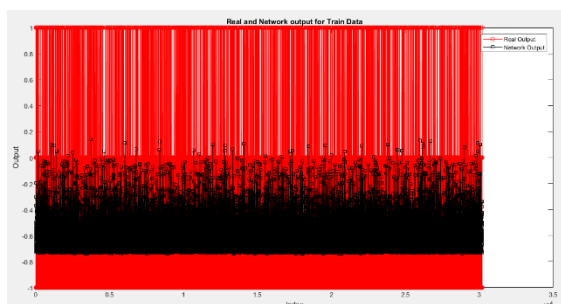
^۲ Assimilation
^۳ Imperialistic Competition
^۴ Revolution

^۱ Imperialist Competitive Algorithm – ICA

۲۱ و ۲۲ به رنگ مشکی نشان داده شده است. در شکل ۲۱ و ۲۲ داده‌های قرمز رنگ به عنوان داده‌های آموزشی و تست (داده‌های واقعی) به شبکه داده شده اند. همان طور که مشخص است، داده‌های پیش بینی شده توسط شبکه ی عصبی پایه، با مطلوبیت حدود ۷۰ درصد و خطای کمتر از ۲۸ درصد، داده‌ها را برآورد کرده است.



شکل ۲۱- نمودار میزان تطابق داده‌های واقعی ترافیک و داده‌های تست با بهینه ساز الگوریتم رقابت استعماری



شکل ۲۲- نمودار تطابق داده‌های واقعی ترافیک با داده‌های آموزشی با بهینه ساز الگوریتم رقابت استعماری

۲-۹- بهینه‌سازی داده‌های ترافیکی با الگوریتم بهینه‌سازی ازدحام ذرات

روش بهینه‌سازی ازدحام ذرات یا به اختصار روش PSO^۱، یک روش سراسری کمینه سازی است که با استفاده از آن می‌توان با مسائلی که جواب آن‌ها یک نقطه یا سطح در فضای n بعدی می‌باشد، برخورد نمود. در اینچنین فضایی، فرضیاتی مطرح می‌شود و یک سرعت ابتدایی به آن‌ها اختصاص داده می‌شود، همچنین کانال‌های ارتباطی بین ذرات در نظر گرفته می‌شود. سپس این ذرات در فضای پاسخ حرکت می‌کنند، و نتایج حاصله بر مبنای یک «ملاک شایستگی» پس از هر بازه‌ی زمانی محاسبه می‌شود. با گذشت زمان، ذرات به سمت ذراتی که دارای ملاک

یکی از دلایل استفاده از این الگوریتم در پژوهش حاضر، تعداد بسیار حجیم داده‌های ترافیکی حاصل از سامانه‌ی تحت‌وب طراحی شده است. زیرا الگوریتم رقابت استعماری توانایی بهینه سازی توابع با مقادیر حجیم داده‌ها را دارد و جواب‌ها رو با سرعت بالایی همگرا می‌کند. جهت پیاده سازی با الگوریتم رقابت استعماری، ابتدا داده‌های ترافیکی با استفاده از شبکه‌ی عصبی پایه اجرا گردید. مانند پیاده‌سازی با الگوریتم ژنتیک، در ابتدای کار تعداد داده‌های تست و آموزشی به صورت کاملاً تصادفی انتخاب گردید.

در الگوریتم‌های بهینه‌سازی مانند الگوریتم رقابت استعماری، همواره یک تابعی به عنوان تابع هدف یا تابع هزینه تعریف می‌شود. کار تابع هدف این است که یک سری محاسبات بر روی سطرها و ستون‌های ماتریس ورودی انجام می‌دهد و به عنوان مثال برای هر سطر ماتریس، بهترین را انتخاب می‌کند. تابع هزینه در این پژوهش، به صورت رابطه ی ۳ تعریف شد.

$$\text{Cost_ANN_EA}(XX, X_{tr}, Y_{tr}, \text{Network}) \quad (3)$$

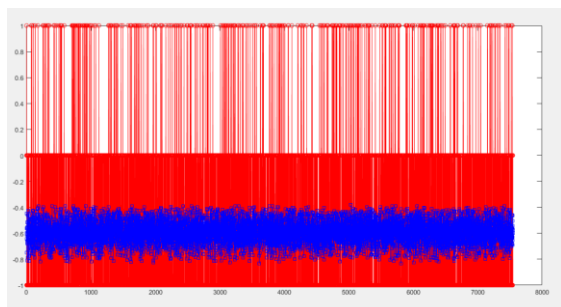
در رابطه ی ۳، xx جمعیت اولیه، y_{tr}, x_{tr} به ترتیب مقادیر داده‌های آموزشی می باشند. با اجرای برنامه، مقادیر $\text{BestCost}=0,3390$ و $\text{MSE}_{tr}=0,3380$ و $\text{MSE}_{ts}=0,3390$ به‌دست آمد. برای محاسبه ی مقادیر خطا و تطابق داده‌های واقعی ترافیکی و داده‌های پیش بینی شده، مقادیر خطاهای کمترین مربعات را در در رابطه ی ۱ قرار داده تا مقدار خطا بر حسب درصد محاسبه گردد.

اگر مقادیر MSE_{tr} و MSE_{ts} در رابطه ی ۱ قرار داده بشود، آنگاه مقادیر خطا برای آن‌ها به ترتیب ۲۹,۱۱۱۹ و ۲۹,۰۶۸۹ درصد محاسبه می‌شود. یعنی با مطلوبیت تقریباً بیشتر از هفتاد درصد بین داده‌های ترافیکی موجود و داده‌های ترافیکی پیش بینی شده توسط شبکه‌ی آموزش دیده با الگوریتم رقابت استعماری. نمودار مربوط به مقادیر واقعی و داده‌های تست و داده‌های آموزشی در شکل‌های ۲۱ و ۲۲ نمایش داده شده است.

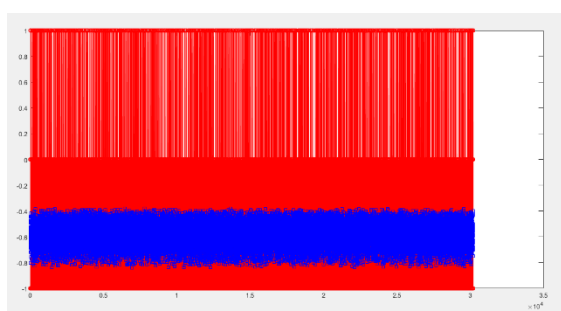
در هر دو شکل ۲۱ و ۲۲، محور x و y داده‌های ورودی می باشند که به صورت یک ماتریس تعریف می‌شود، محور x شامل ورودی های شبکه عصبی می باشد و محور y هم به عنوان خروجی شبکه عصبی می باشد. همانطور که در این شکل‌ها مشهود است، داده‌های پیش بینی شده که شامل رنگ‌های ترافیکی می باشد در شکل

^۱ Particle swarm optimization

های واقعی، داده‌های آموزشی و داده‌های تست به ترتیب در شکل ۲۳ و ۲۴ نمایش داده شده است.



شکل ۲۳- نمودار میزان تطابق داده‌های واقعی ترافیکی با داده‌های train با بهینه ساز الگوریتم ازدحام ذرات



شکل ۲۴- نمودار تطابق داده‌های واقعی ترافیکی با داده‌های test با بهینه ساز الگوریتم ازدحام ذرات

در هر دو شکل ۲۳ و ۲۴، محور x و y داده‌های ورودی می باشند که به صورت یک ماتریس تعریف می شود، محور x شامل ورودی های شبکه عصبی می باشد و محور y هم به عنوان خروجی شبکه عصبی می باشد. همانطور که در این شکل‌ها مشهود است، داده‌های پیش بینی شده که شامل رنگ‌های ترافیکی می باشد در شکل ۲۳ و ۲۴ به رنگ آبی، نشان داده شده است. در شکل ۲۳ و ۲۴ داده‌های قرمز رنگ به عنوان داده‌های آموزشی و تست (داده‌های واقعی) به شبکه داده شده اند. همانطور که مشخص است، شبکه‌ی آموزش دیده، با مطلوبیت حدود ۷۲ درصد و خطای ۲۸ درصد، داده‌ها را برآورد کرده است.

۳- نتایج تحقیق

در تحقیق حاضر، به بررسی مسئله‌ی ترافیک از بعد مکانی و زمانی جهت پیش بینی آن پرداخته شد. بدین مفهوم که ابتدا سامانه‌ای طراحی شد تا بتواند داده‌های ترافیکی را در یک محدوده و مکان مشخص و در یک بازه-ی زمانی مشخص به صورت کاملاً خودکار و لحظه‌ای

شایستگی بالاتری هستند و در گروه ارتباطی یکسانی قرار دارند، شتاب می‌گیرند. علی‌رغم اینکه هر روش در محدوده‌ای از مسائل به خوبی کار می‌کند، این روش در حل مسائل بهینه سازی پیوسته موفقیت بسیاری از خود نشان داده است [۱۷-۱۹]. جهت بهینه‌سازی داده‌های ترافیکی و رسیدن به مطلوب‌ترین نقطه، می‌توان از الگوریتم بهینه‌سازی ازدحام ذرات استفاده نمود. الگوریتم ازدحام ذرات، این توانایی را دارد که در آن هر ذره، از اطلاعات گذشته‌ی خود سود ببرد. به این مفهوم که این الگوریتم، دارای حافظه است و برای تولید نسل جدید، هر ذره از نسل قبل خود استفاده می‌کند. این در حالی است که سایر الگوریتم‌های تکاملی چنین ویژگی ای ندارند. به همین دلیل، در این پژوهش، چون قرار است محقق از داده‌های ترافیکی برداشتی استفاده کند تا ترافیک را در آینده پیش بینی کند، از این الگوریتم استفاده شده است. در این پژوهش، ابتدا برنامه‌ی شبکه عصبی پایه نوشته شد و در مرحله‌ی بعدی، تابعی که معرف الگوریتم بهینه‌سازی ازدحام ذرات می باشد به برنامه اضافه گردید. در الگوریتم‌های بهینه‌سازی مانند الگوریتم ازدحام ذرات، همواره یک تابعی به عنوان تابع هدف یا تابع هزینه تعریف می‌شود. کار تابع هدف این است که یک سری محاسبات بر روی سطرها و ستون‌های ماتریس ورودی انجام می‌دهد و به عنوان مثال برای هر سطر ماتریس، بهترین را انتخاب می‌کند. تابع هزینه در این پژوهش، به صورت رابطه‌ی ۴ تعریف شد.

$$\text{Cost_ANN_EA}(XX, Xtr, Ytr, \text{Network}) \quad (4)$$

در رابطه‌ی ۴، xx جمعیت اولیه، ytr, xtr به ترتیب مقادیر داده‌های آموزشی می باشند. با اجرای برنامه‌ی نوشته شده، مقادیر $MSEtr=0.3305$ و $MSEts=0.3252$ به دست آمد که مقادیر بهینه هستند. برای محاسبه‌ی مقادیر خطا و تطابق داده‌های واقعی ترافیکی و داده‌های پیش بینی شده، مقادیر خطاهای کمترین مربعات را در رابطه‌ی ۱، قرار داده تا مقدار خطا بر حسب درصد محاسبه گردد. اگر مقادیر $MSEtr$ و $MSEts$ در رابطه‌ی ۱ قرار داده بشود، آنگاه مقادیر خطا برای آن‌ها به ترتیب ۲۸,۷۴۴۶ و ۲۸,۵۱۳۲ درصد محاسبه می‌شود. یعنی با مطلوبیت تقریباً هفتاد و دو درصد بین داده‌های ترافیکی موجود و داده‌های ترافیکی پیش بینی شده توسط شبکه‌ی آموزش دیده با الگوریتم رقابت ازدحام ذرات. نمودار مربوط به داده-

۴- نتیجه گیری

آنچه که در پژوهش حاضر مورد ارزیابی و تحلیل قرار گرفت، میزان هماهنگی و تطابق داده‌های ترافیکی تولیدی به صورت اتوماتیک، با داده‌های تولیدی از طریق یادگیری شبکه‌های عصبی بود. همانطور که در پژوهش حاضر، بیان شد، ابتدا سامانه ای تحت وب طراحی شد تا با استفاده از این سامانه، اقدام به جمع‌آوری داده‌های ترافیک به صورت کاملاً خودکار و لحظه ای گردد. از داده‌های ترافیکی جمع‌آوری شده در این سامانه، به عنوان داده‌های اولیه‌ی مورد نیاز برای پژوهش حاضر، استفاده گردید. یدین صورت که ابتدا با استفاده از آموزش ساده، شبکه آموزش داده شد و سپس داده‌هایی به عنوان داده‌های ورودی در اختیار شبکه‌ی آموزش دیده قرار داده شده تا ترافیک را بر طبق شاخصه‌های زمان (که شامل سال، ماه، روز، ساعت، دقیقه، ثانیه) و شاخصه‌های موقعیت جغرافیایی، برای نقاط دیگر ارزیابی کند. میزان خطا برای حدود ۳۸۰۰۰ داده‌ی ورودی در این مرحله حدود ۲۸ درصد، به دست آمد. به عبارتی با حدود ۷۲ درصد تطابق داده‌های واقعی و داده‌های پیش‌بینی شده توسط شبکه عصبی آموزش دیده حاصل گردید. در مرحله‌ی بعدی، همین مراحل با تبدیلات موجک انجام شد. با این تفاوت که میزان داده‌های ورودی ۲۰۰۰ در نظر گرفته و شبکه آموزش داده شد. در این مرحله، میزان تطابق داده‌های ترافیکی با مقادیر پیش‌بینی شده توسط شبکه آموزش دیده، با مقداری خطایی حدود ۲۸ درصد به دست آمد. به عبارت دیگر مطلوبیت بدست آمده به میزان ۷۲ درصد حاصل گردید. این نتایج نشان می‌دهد که در صورت آموزش شبکه موجکی با تعداد ورودی‌های بیشتر، مطمئناً نتیجه‌ی بهتری نسبت به آموزش از طریق شبکه‌ی عصبی پایه به دست خواهد داد. در آخرین مرحله از این پژوهش، اقدام به بهینه‌سازی داده‌ها با استفاده از سه الگوریتم بهینه‌سازی رقابت استعماری، الگوریتم ژنتیک و الگوریتم ازدحام ذرات گردید. نتایج حاصل از این مرحله به خوبی نشان داد که استفاده از الگوریتم‌های بهینه‌ساز در مسائل پیش‌بینی ترافیک، در یافتن بهترین نقطه، کمک شایانی می‌کند. الگوریتم‌های فراابتکاری مذکور، به دلایل شرح داده شده در این پژوهش، جهت پیش‌بینی مقادیر حجیم داده‌های ترافیکی نسبت به الگوریتم‌های دیگر در اولویت

جمع‌آوری کند. در این پژوهش، محدوده‌ی مورد نظر، مناطق ۲۲ گانه‌ی کلانشهر تهران در نظر گرفته شد (بعد مکانی). سپس با توجه به موقعیت منطقه، با استفاده از روش‌های ارائه شده، اقدام به جمع‌آوری داده‌های ترافیکی مورد نیاز پژوهش به صورت شبانه‌روزی گردید (بعد زمانی). داده‌های ترافیکی جمع‌آوری شده در این پژوهش شامل طول جغرافیایی و عرض جغرافیایی نقاط ترافیکی در محدوده‌ی مورد مطالعه، طول و عرض عکسی این نقاط، زمان برداشت این داده‌های ترافیکی که شامل سال، ماه، روز، ساعت، دقیقه و ثانیه است، می‌باشد. علاوه بر موارد مذکور، اطلاعات عکس گرفته شده از محدوده‌های ترافیکی نیز که شامل نوع عکس ذخیره شده، حجم عکس، طول و عرض عکس مذکور و مسیر ذخیره‌ی آن می‌باشد، برداشت و در پایگاه داده ذخیره گردیده است. این داده‌ها برای آموزش شبکه‌ی عصبی پایه و موجک مورد استفاده قرار گرفت تا داده‌های ترافیک در آینده توسط شبکه پیش‌بینی شود. جهت آموزش شبکه عصبی پایه، ۸۰ درصد داده‌های ترافیکی به عنوان داده‌های آموزشی و ۲۰ درصد این داده‌ها به عنوان داده‌های تست مورد استفاده قرار گرفت. نتایج حاصل از این پژوهش نشان داد که شبکه‌ی آموزش داده شده با مطلوبیت بیش از ۷۲ درصد ترافیک را پیش‌بینی کرده است. برای پیش‌بینی ترافیک با استفاده از شبکه‌ی عصبی موجک، ۷۰ درصد داده‌ها به عنوان داده‌های آموزشی، ۲۰ درصد داده‌ها جهت داده‌های تست و ۱۰ درصد داده‌ها نیز به عنوان داده‌های اعتبارسنجی مورد ارزیابی قرار گرفت. نتایج حاصل از پژوهش نشان داد که شبکه‌ی عصبی موجک با مطلوبیت بیش از ۷۲ درصد ترافیک را پیش‌بینی کرده است. همچنین مطالعات و تحلیل‌های انجام گرفته بیانگر این مطلب است که هرچه تعداد داده‌های ورودی برای آموزش شبکه‌ی عصبی موجک بیشتر باشد، آنگاه شبکه‌ی عصبی موجک با دقت بیشتری ترافیک را پیش‌بینی خواهد کرد. در ادامه‌ی پژوهش، اقدام به بهینه‌سازی داده‌های ترافیکی توسط سه الگوریتم فراابتکاری گردید. نتایج حاصل از این مراحل به خوبی نشان داد که داده‌های ترافیکی حاصل از شبکه‌ی عصبی موجک با داده‌های واقعی حاصل از سامانه‌ی طراحی شده، اختلاف بسیار کمی دارند که این جواب با به کارگیری الگوریتم‌های بهینه‌سازی به خوبی و با دقت بسیار زیاد بهینه شدند.

قرار دارند. به علاوه، نتایج حاصل از سه الگوریتم بهینه ساز تقریباً به یکدیگر نزدیک و با دقت بالایی جواب را بهینه می‌کنند. سه الگوریتم بهینه ساز مذکور، تقریباً با دقت ۷۲ درصد جواب را بهینه کرده‌اند. این بدین مفهوم است که بین داده‌های واقعی ترافیک برداشت شده توسط سامانه‌ی تحت وب طراحی شده و داده‌هایی که از طریق یادگیری شبکه حاصل شده‌اند، تقریباً ۲۸ درصد خطا وجود دارد.

۵- پیشنهادات

دستاوردهای حاصل در این پژوهش را، در حوزه‌های مدیریتی، کنترل و پیش‌بینی ترافیک، که نقش مهمی در اقتصاد کلانشهرها دارد و بسیاری از حوزه‌های دیگر که نیازمند داده‌های ترافیکی جهت پیش‌بینی هستند، می‌توان مورد استفاده قرار داد. به عنوان نمونه، از سامانه‌ی تحت وب طراحی شده برای تولید داده‌های ترافیکی، می‌توان برای بررسی میزان تاثیر ترافیک تولیدی بر روی موضوع آلودگی

که امروزه معضل بسیاری از کلانشهرها است، استفاده نمود. این مسئله را هم از بعد زمانی و هم از بعد مکانی می‌توان مورد ارزیابی قرار داد و به نتایج قابل توجهی دست یافت. به صورت کلی، همانطور که اشاره شد، از آنجایی که موضوع دستیابی به داده‌های ترافیکی همواره مشکلی بوده است که طی دوره‌های مختلف لاینحل باقی مانده و بسیاری از پژوهش‌ها و تحقیقات انجام گرفته در زمینه‌ی مسئله‌ی ترافیک را همواره با مشکل مواجه کرده است، از این رو، با طراحی این سامانه‌ی تحت وب جهت تولید داده‌های ترافیکی، به جرات می‌توان بیان کرد که مشکل دستیابی به داده‌های اولیه برای انجام تحقیقات مختلف در حوزه‌ی ترافیک، مرتفع گردیده و محقق سعی کرده است که نمونه‌ای از به کارگیری این داده‌های ترافیکی حاصل از سامانه‌ی طراحی شده را با انجام پژوهش حاضر که پیش‌بینی ترافیک توسط آموزش شبکه‌های عصبی و بهینه‌سازی توسط الگوریتم‌های فراابتکاری می‌باشد، به پژوهشگران حوزه‌ی ترافیک معرفی کند.

مراجع

- [1] S. Behzadi, An intelligent location and state reorganization of traffic signal, *Geodesy and Cartography*, 46(3) (2020) 145-150.
- [2] A.S.M. Zadeh, M.A. Rajabi, Analyzing the effect of the street network configuration on the efficiency of an urban transportation system, *Cities*, 31 (2013) 285-297.
- [3] Y. Zhi-Sheng, S. Chun-Fu, X. Zhi-Hua, Y. Hao, Short-term traffic volumes forecasting of road network based on principal component analysis and support vector machine [J], *Journal of Jilin University (Engineering and Technology Edition)*, 1 (2008) 10.
- [4] Z. Qasim, A.-R. Ziboon, K. Fali, TransCad analysis and GIS techniques to evaluate transportation network in Nasiriyah city, in: *MATEC Web of Conferences*, EDP Sciences, 2018, pp. 03029.
- [5] M. Asgharpour, S.Ebrahimnezhad (2001). "Presenting a traffic allocation model to the urban transportation network and solving it using genetic algorithm", *Journal of the College of Engineering, University of Tehran*, pp.587-602, (in Persian)
- [6] Z. Cheng, W. Wang, J. Lu, X. Xing, Classifying the traffic state of urban expressways: A machine-learning approach, *Transportation Research Part A: Policy and Practice*, (2018).
- [7] M. Javadzadeh, M. Kangavari, Knowledge Transfer in the Mobile Environment with a Passive Defence Approach for Road Traffic Forecast, (2013)
- [8] M.A. Silgu, H.B. Celikoglu, Clustering traffic flow patterns by fuzzy C-means method: some preliminary findings, in: *International Conference on Computer Aided Systems Theory*, Springer, 2015, pp. 756-764.
- [9] S.P. Kumarage, R. Rajapaksha, D. De Silva, J. Bandara, Traffic Flow Estimation for Urban Roads Based on Crowdsourced Data and Machine Learning Principles, in: *First International Conference on Intelligent Transport Systems*, Springer, 2017, pp. 263-273.
- [10] A. Elfar, A. Talebpour, H.S. Mahmassani, Machine learning approach to short-term traffic congestion prediction in a connected environment, *Transportation Research Record*, 2672(45) (2018) 185-195.
- [11] A. Alavi, A. Parhizgar, A. Roknoddin eftekhari, M.GHalibaf, M.Poormoosavi (2010). "Spatial modeling of travel demand based on a new method for forecasting and reducing traffic (District 6 of Tehran)", *The Journal of Spatial Planning*, 15(4), pp.43-61, (in Persian)

- [12] M.Malayeri, (2015)," Traffic flow estimation and vehicle detection system using Gaussian mixed model and artificial neural network", Kerman University of Industrial Graduate Studies, Faculty of Electrical and Computer Engineering, (in Persian)
- [13] S. Behzadi, A. Jalilzadeh, Introducing a Novel Digital Train Model Using Artificial Neural Network Algorithm, Civil Engineering Dimension, 22(2) (2020) 47-51.
- [14] H. Jafarian, S. Behzadi, Evaluation of PM2. 5 emissions in Tehran by means of remote sensing and regression models, Pollution, 6(3) (2020) 521-529.
- [15] S. Behzadi, M. Kolbadinejad, INTRODUCING A NOVEL METHOD TO SOLVE SHORTEST PATH PROBLEMS BASED ON STRUCTURE OF NETWORK USING GENETIC ALGORITHM, International Archives of the Photogrammetry, Remote Sensing & Spatial Information Sciences, (2019).
- [16] J. Behnamian, A parallel competitive colonial algorithm for JIT flowshop scheduling, Journal of Computational Science, 5(5) (2014) 777-783.
- [17] K.-L. Du, M. Swamy, Particle swarm optimization, in: Search and optimization by metaheuristics, Springer, 2016, pp. 153-173.
- [18] Z. Chatrsimab, A. Alesheikh, B. Vosoghi, S. Behzadi, M. Modiri, Land Subsidence Modelling Using Particle Swarm Optimization Algorithm and Differential Interferometry Synthetic Aperture Radar, ECOPERSIA, 8(2) (2020) 77-87.
- [19] Z. Chatrsimab, A.A. Alesheikh, B. Voosoghi, S. Behzadi, M. Modiri, Development of a Land Subsidence Forecasting Model Using Small Baseline Subset—Differential Synthetic Aperture Radar Interferometry and Particle Swarm Optimization—Random Forest (Case Study: Tehran-Karaj-Shahriyar Aquifer, Iran), in: Doklady Earth Sciences, Springer, 2020, pp. 718-725.