

بررسی جامع بر روی روش‌های طبقه‌بندی غیرپارامتریک به منظور تفکیک عوارض شهری با استفاده از تلفیق داده‌های لایدار و تصویر هوایی با توان تفکیک مکانی بسیار بالا

زینت قوامی^۱، حسین عارفی^{۲*}، بهناز بیگدلی^۲، میلاد جانعلی پور^۴

- ۱- کارشناس ارشد فتوگرامتری- دانشکده مهندسی نقشه برداری و اطلاعات مکانی- پردیس دانشکده‌های فنی- دانشگاه تهران
- ۲- استادیار دانشکده مهندسی نقشه برداری و اطلاعات مکانی - پردیس دانشکده‌های فنی- دانشگاه تهران
- ۳- استادیار دانشکده مهندسی عمران- دانشگاه صنعتی شاهرود
- ۴- دانشجوی دکتری سنجش از دور- دانشکده مهندسی ژئودزی و ژئوماتیک- دانشگاه خواجه نصیرالدین طوسی

تاریخ دریافت مقاله: ۹۵/۰۴/۱۴ تاریخ پذیرش مقاله: ۹۶/۰۳/۰۹

چکیده

امروزه به دست آوردن اطلاعات پوشش اراضی شهری، یکی از مهم‌ترین ابزارهای مدیریت شهری است و کاربرد زیادی در بررسی تغییرات زمین دارد. طبقه‌بندی تصاویر، یکی از متداول‌ترین روش‌های استخراج اطلاعات از داده‌های سنجش از دور است. وجود نواحی شهری پیچیده و متراکم یکی از مشکلات آنالیزهای سنجش از دور می‌باشد. دقت عملکرد طبقه‌بندی در این مناطق برای محققان مورد توجه بوده و همواره سعی در بهبود این دقت داشته‌اند. با استفاده از تکنیک تلفیق داده‌های مختلف و به کارگیری اطلاعات متنوع از عوارض می‌توان به طبقه‌بندی دقیق‌تر با قابلیت اعتماد بالاتر دست یافت. از جمله روش‌های موفقیت آمیز طبقه‌بندی در سال‌های اخیر، می‌توان به الگوریتم‌های ماشین بردار پشتیبان و الگوریتم‌های یادگیری دسته جمعی مانند بگینگ، بوستینگ و جنگل تصادفی اشاره کرد. در این مقاله در مورد عملکرد این چهار الگوریتم برای شناسایی عوارض شهری با استفاده از نقاط متراکم لایدار و تصویر هوایی با قدرت تفکیک بسیار بالا بحث شده است. در این تحقیق از سه شیوه بر اساس طبقه‌بندی داده لایدار و تصویر هوایی به تنهایی و تلفیق هر دو داده استفاده شده است. نتایج نشان می‌دهد که ترکیب داده لایدار و تصویر هوایی، طبقه‌بندی بهتری را از عوارض شهری به دست می‌دهد. در نهایت طبقه‌بندی عوارض شهری با کمک تلفیق داده‌های لایدار و تصویر هوایی و با استفاده از الگوریتم ماشین‌های بردار پشتیبان با دقت ۹۳/۹۹٪، توانایی بالاتری نسبت به سایر روش‌های طبقه‌بندی مورد استفاده مانند بگینگ، بوستینگ و جنگل تصادفی دارد.

کلیدواژه‌ها: داده لایدار، تصویر هوایی، ماشین‌های بردار پشتیبان، بگینگ، بوستینگ، جنگل تصادفی

۱- مقدمه

با گسترش جمعیت شهرها، طبقه‌بندی پوشش اراضی، تبدیل به یک موضوع مهم در مطالعات شهری شده است [۱ و ۲]. اگرچه مطالعات شهری برای بیشتر محققان جالب می‌باشد، اما مناطق شهری به علت ترکیب انواع عوارض طبیعی و ساخت دست بشر که منجر به یک ارتباط اشتباه بین عوارض و طیفشان می‌شود، با پیچیدگی‌های بسیاری همراه هستند و فرض اینکه استفاده از یک نوع سنجنده می‌تواند تمام اطلاعات برای ویژگی‌های لازم را به دست آورد، بسیار خوش‌بینانه است. بنابراین بیشتر محققان سعی در ترکیب داده‌های به دست آمده از چند منبع مختلف دارند [۳، ۴ و ۵]. در سال ۲۰۱۱، یانگ کا و همکارانش با استفاده از تلفیق داده‌های لایدار و تصویر هوایی با قدرت تفکیک مکانی ۰/۵ متر و با به کارگیری روشی بر پایه‌ی زنجیره‌ی تصادفی مارکوف فازی (FMRF^۱) توانستند طبقه‌بندی پوشش اراضی را با دقت ۸۸/۹٪ انجام دهند، نتایج نشان داد تلفیق این دو داده، اقدامی مناسب به منظور طبقه‌بندی پوشش اراضی می‌باشد [۳]. در سال ۲۰۱۲، سینگ و همکارانش به بررسی میزان تأثیر لایدار در شناسایی کلاس‌های کاربری شهری با استفاده از تلفیق داده‌های لایدار و تصاویر لندست پرداختند. آن‌ها با به کارگیری طبقه‌بندی کننده‌ی نظارت شده‌ی بیشترین شباهت نشان دادند دقت طبقه‌بندی به روش بیشترین شباهت با داده‌ی لایدار ۱ متری و نقشه‌ی موضوعی (TM^۲) سی متری، ۳۲٪ نسبت به استفاده‌ی مجزا از لایدار و ۸٪ نسبت به استفاده‌ی مجزا از TM افزایش یافته است [۵]. در سال ۲۰۱۲، فرانسیسکو خاویر و همکارانش، با تلفیق تصویر هوایی و مدل رقومی نرمال شده‌ی

سطح (nDSM^۳) با قدرت تفکیک مکانی ۰/۵ متر، و استفاده از روش بیشترین شباهت، به دقت ۹۳/۲٪ برای طبقه‌بندی پوشش اراضی دست یافتند [۶]. در سال ۲۰۱۳، مادیوریمما و همکارانش با استفاده از تلفیق ابر نقاط لایدار و تصویر هوایی، پوشش گیاهی در مناطق شهری و ساختمان‌ها را از یکدیگر تفکیک کردند. آن‌ها با حد آستانه‌گذاری جهت قطعه‌بندی ابر نقاط لایدار و استفاده از یک ماسک برای پوشش گیاهی، به دقت بالاتر از ۸۵٪ برای طبقه‌بندی عوارض منطقه که شامل درخت و ساختمان است، دست یافتند [۷].

محمد اواد و همکاران در سال ۲۰۱۷، برای شناسایی عوارض شهری از تلفیق داده‌های اپتیک و لایدار استفاده کردند. در تحقیق انجام شده با استفاده از روشی مبتنی بر تبدیل ویولت به شناسایی عوارض شهری پرداختند. روش پیشنهادی برای شناسایی ساختمان دارای دقت متوسط ۹۶ درصد بود [۸]. در سال ۲۰۱۶، زارع و محمدزاده با استفاده از داده لایدار و اپتیک به شناسایی درختان و مناطق مسکونی پرداختند. برای شناسایی عوارض ذکر شده از روش طبقه‌بندی ماشین بردار پشتیبان و خوشه‌بندی K-Means استفاده شد [۹]. وو و همکاران در سال ۲۰۱۷، از تلفیق داده لایدار و تصویر اپتیک Worldview-2 به شناسایی عوارض پوششی سطح زمین پرداختند. برای تلفیق این دو داده از یک روش کلاسه‌بندی سلسله مراتبی مبتنی بر ماشین بردار پشتیبان استفاده شد. نتیجه حاصل از این تحقیق این بود که تلفیق داده‌های ذکر شده می‌تواند در بهبود دقت اطلاعات تفسیری کمک کند [۱۰]. اصفهانی و محمدزاده در سال ۲۰۱۶، به منظور حل مشکلاتی نظیر شناسایی ساختمان‌ها و درختان کوچک، مرز نامناسب تاج درختان و وجود ساختمان‌ها با پوشش گیاهی از تلفیق داده لایدار و تصویر اپتیک

^۱ Fuzzy Markov Random Field^۲ Thematic Map^۳ Normalized Digital Surface Model

متوسط در مناطق شهری استفاده نموده‌اند. این در حالیکه استفاده از داده‌های با توان تفکیک مکانی بسیار بالا می‌تواند به تولید نقشه‌های دقیق و با جزئیات بالا در مناطق شهری کمک نماید، اما باید در نظر گرفت که پیچیدگی این مناطق از نظر طیفی، هندسی و بافتی باعث ایجاد خطا در اطلاعات استخراجی می‌شود. از این رو هدف اصلی این تحقیق، تلفیق داده‌های لایدار و تصویر هوایی با قدرت تفکیک مکانی بسیار بالا به منظور تولید اطلاعات تفسیری دقیق خواهد بود. به عنوان یک هدف فرعی، در این مطالعه به بررسی کارایی روش‌های طبقه‌بندی ماشین بردار پشتیبان، جنگل تصادفی، بگینگ و بوستینگ در طبقه‌بندی عوارض یک منطقه‌ی شهری پیچیده نیز خواهیم پرداخت.

۲- روش تحقیق

در این مقاله، به منظور طبقه‌بندی عوارض شهری، ابتدا فیلترینگ داده‌های لایدار با استفاده از الگوریتم اکسلسون به منظور تولید مدل رقومی سطح (DSM^۵) انجام خواهد شد [۱۶]. در ادامه، ویژگی‌های بافتی مختلف از داده‌های لایدار و تصویر هوایی استخراج می‌گردد. با استفاده از الگوریتم ژنتیک، ویژگی‌های بهینه و کارآمد در طبقه‌بندی، از میان این ویژگی‌ها انتخاب شده و در نهایت، الگوریتم‌های ماشین بردار پشتیبان، بگینگ، بوستینگ و جنگل تصادفی به منظور طبقه‌بندی عوارض شهری، روی ویژگی‌های بهینه اعمال می‌شود. به منظور دستیابی به عملکرد بهینه برای هر یک از الگوریتم‌های ذکر شده، پارامترهای مؤثر همه‌ی روش‌ها مورد بررسی قرار خواهند گرفت. با کمک روش جست و جوی شبکه‌ای^۶ و با تنظیم پارامترهای بهینه برای الگوریتم ماشین بردار پشتیبان، بهترین عملکرد برای

استفاده کردند. به منظور حل مشکلات ذکر شده از یک استراتژی طبقه‌بندی در سطح شیء استفاده شد. نتایج نشان می‌دهند که تلفیق داده‌های لایدار و اپتیک و طبقه‌بندی در سطح شیء می‌تواند به نتایج با دقت مناسب منجر شود [۱۱].

نا و همکارانش در سال ۲۰۱۰، نشان دادند که الگوریتم جنگل تصادفی^۱ دقت بالاتری نسبت به الگوریتم بیشترین شباهت در طبقه‌بندی پوشش اراضی با استفاده از تصاویر TM دارد [۱۲]. در سال ۲۰۱۰، سزنی و همکارانش، الگوریتم‌های جنگل تصادفی و ماشین بردار پشتیبان^۲ را به منظور طبقه‌بندی پوشش گیاهی در جنگل‌های استوایی با استفاده از تصاویر TM استفاده کردند، نتایج نشان داد که الگوریتم ماشین بردار پشتیبان دقت طبقه‌بندی بالاتری در مقایسه با الگوریتم جنگل تصادفی دارد اما الگوریتم جنگل تصادفی نسبت به ماشین بردار پشتیبان نیاز به تنظیم پارامترهای کمتری برای طبقه‌بندی دارد [۱۳]. در سال ۲۰۰۶، جایلزون از الگوریتم جنگل تصادفی و داده‌های سنجش از دور به منظور طبقه‌بندی انواع پوشش گیاهی در یک منطقه‌ی کوهستانی استفاده کرد و به این نتیجه رسید که جنگل تصادفی برای هدف ذکر شده دقت مشابهی با الگوریتم‌های بگینگ^۳ و بوستینگ^۴ دارد [۱۴]. در سال ۲۰۱۰، پینومهاس و همکارانش دریافتند که جنگل تصادفی و بگینگ نسبت به الگوریتم ماشین بردار پشتیبان دقت بالاتری در طبقه‌بندی پوشش اراضی دارند [۱۵].

بنابر پیشینه تحقیق، دقت روش‌های مختلف طبقه‌بندی در مناطق مختلف دارای نتایج متفاوتی بوده است. به عبارت دیگر دقت طبقه‌بندی کننده‌ها به منطقه مورد مطالعه و داده‌های مورد استفاده وابسته است. از طرف دیگر روش‌های داده‌های با توان تفکیک مکانی

^۱ Random Forest

^۲ Support Vector Machines

^۳ Bagging

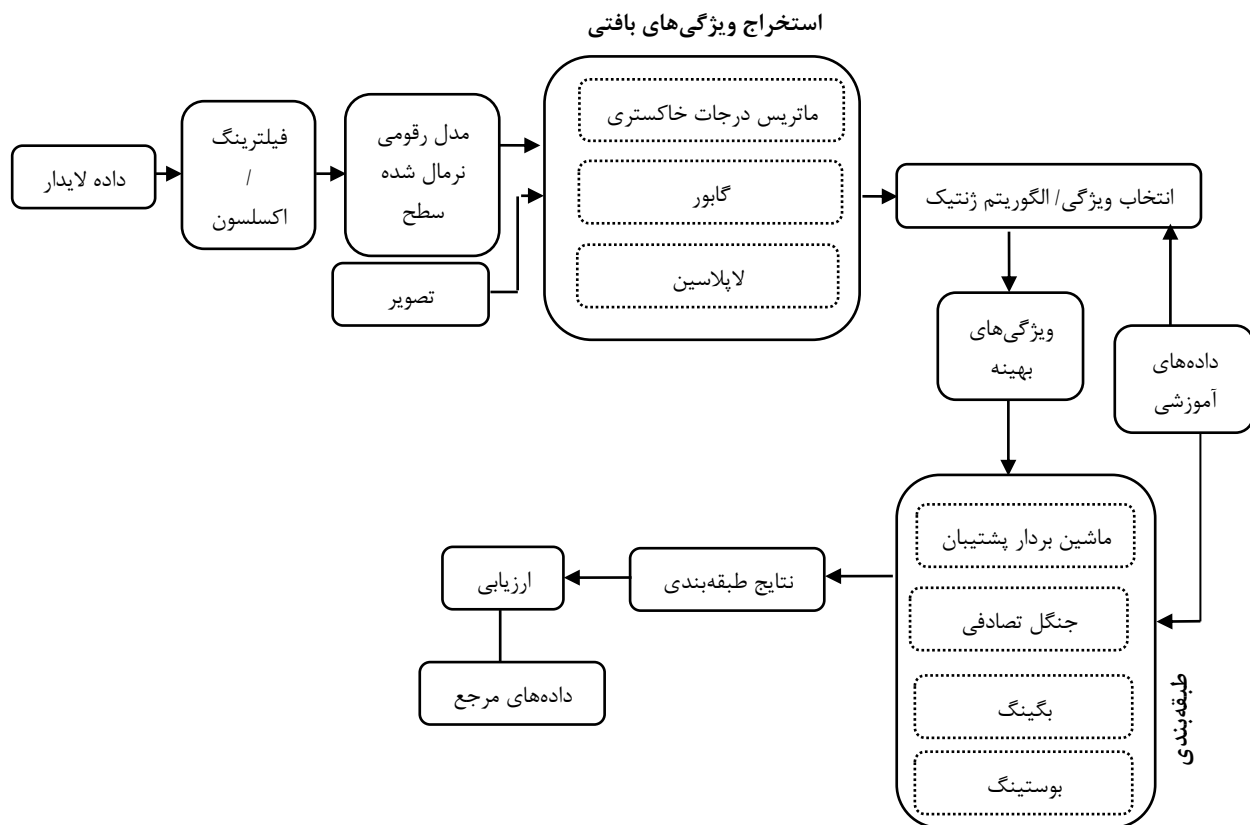
^۴ Boosting

^۵ Digital Surface Model

^۶ Grid Searching

خواهد شد. روند کلی تحقیق در شکل (۱) نشان داده شده است.

این الگوریتم به دست خواهد آمد. علاوه بر این با تعیین تعداد درخت بهینه برای الگوریتم‌های یادگیری دسته‌جمعی، بهترین عملکرد این الگوریتم‌ها نیز حاصل



شکل ۱: روند کلی تحقیق

۲-۱- فیلترینگ داده‌های لایدار

در تمام کاربردهای لایدار، فیلترینگ داده لایدار، یک گام ضروری در تفکیک بازگشت‌های زمینی و غیرزمینی است [۱۵]. روش‌های متعددی برای این منظور با توجه به شرایط محیطی مختلف استفاده شده‌اند. از آنجاییکه منطقه مورد مطالعه این تحقیق، یک منطقه شهری نسبتاً مسطح است، نتایج حاصل از مطالعات پیشین نشان می‌دهد که در محیط شهری، الگوریتم اکسلسون به دقت بهتری نسبت به سایر الگوریتم‌های فیلترینگ دست می‌یابد [۱۵]. بنابراین به منظور تفکیک زمین از غیر زمین در منطقه مورد مطالعه، الگوریتم اکسلسون استفاده شده است. در این روش، یک شبکه مثلثی

نامنظم و پراکنده از نقاط ایجاد می‌شود که اولین تخمین از زمین لخت می‌باشد. این شبکه‌ی مثلثی در یک روند تکراری با اضافه نمودن نقاط از طریق حد آستانه گذاری، متراکم می‌گردد. این روند تا زمانی که تمام نقاط به عنوان زمین یا عارضه کلاس‌بندی شوند ادامه می‌یابد [۱۶]. در این مرحله، ورودی الگوریتم، DSM و خروجی آن، nDSM می‌باشد. این مدل از سطح، در مراحل بعدی برای استخراج ویژگی‌های مورد نظر استفاده می‌شود.

۲-۲- استخراج ویژگی از داده‌های لایدار و

تصویر هوایی

با توجه به محدودیت‌های استفاده از داده‌های خام

خاکستری z قرار می گیرد. با محاسبه ی ماتریس روابط همسایگی پیکسل ها، ویژگی های آماری بافتی نظیر آنروپی^۶، وابستگی^۷، کنتراست^۸، میانگین^۹، واریانس^{۱۰}، همگنی^{۱۱}، عدم شباهت^{۱۲} و ممان دوم^{۱۳} از ماتریس GLCM استخراج می شوند [۲۱].

۲-۲-۲- ویژگی های بافتی گابور

فیلتر گابور یک تکنیک گسترده به منظور تجزیه و تحلیل بافت می باشد. این فیلتر در واقع به صورت یک صفحه ی سینوسی با فرکانس و جهت مشخص می باشد که با یک پوش گوسی مدوله شده است. هر فیلتر گابور را می توان هم در حوزه ی مکان و هم در حوزه ی فرکانس تعریف کرد. فیلتر گابور، محدودیت های خاص خودش را دارد. از جمله این محدودیت ها، می توان به وابستگی این فیلترها به پارامترهای مختلف و تنظیم درست این پارامترها اشاره کرد. این مشکل، گاهی به طراحی بانک فیلتر ارجاع داده می شود و شامل انتخاب تعدادی فیلتر مناسب در جهت ها و فرکانس های مختلف می باشد. با اجرای فیلتر گابور در طول موج ها و جهات مختلف، می توان میانگین و انحراف معیار که دو ویژگی گابور هستند، را محاسبه نمود [۲۲].

۲-۲-۳- ویژگی بافتی لاپلاسی

یکی از روش های مبتنی بر مشتق دوم جهت آنالیز بافت تصویر، فیلتر لاپلاسی می باشد. این فیلتر در پیدا کردن جزئیات در تصویر بسیار موفق عمل می کند. فیلتر لاپلاسی با اعمال کرنلی با ابعاد پنجره مختلف روی تصویر اجرا می شود [۲۳].

به دلیل وجود تعداد پیکسل های زیاد و محاسبات بیش از حد، با استفاده از تحلیل و آنالیز تصویر، ویژگی هایی را تولید و سپس این ویژگی ها به منظور طبقه بندی استفاده می شوند [۱۷]. با توجه به تحقیقات انجام شده در زمینه ی طبقه بندی و کارآیی بالای ویژگی های بافتی در استخراج عوارض شهری و نقش مؤثر این ویژگی ها در تفکیک عوارض [۱۸، ۱۹ و ۲۰]، در این تحقیق ویژگی های بافتی مبتنی بر ماتریس هم اتفاق^۱، ویژگی های بافتی گابور^۲ و نیز ویژگی های بافتی لاپلاسی^۳ از داده ی لایدار و تصویر هوایی استخراج خواهد شد، لذا در ادامه مفاهیم این روش ها بیان خواهند شد.

۲-۲-۱- ویژگی های بافتی مبتنی بر ماتریس هم اتفاق درجات خاکستری (GLCM^۴)

بافت یک منطقه از تصویر، از طریق سطوح خاکستری توزیع شده در سرتاسر پیکسل های منطقه، تعیین می شود. یک راه مناسب برای دستیابی به بافت یک تصویر، اندازه گیری آن در یک پنجره است. برای اندازه گیری بافت در یک پنجره، می توان از توصیفگرهای آماری مختلفی استفاده کرد. اطلاعاتی از این قبیل می تواند از هیستوگرام های^۵ مرتبه ی دوم و با در نظر گرفتن پیکسل ها به صورت جفت، استخراج شوند [۱۷]. هیستوگرام های تغییرات روشنایی، به عنوان تابعی از فاصله و جهت در نظر گرفته می شوند. ماتریس هم اتفاق، تقریبی از احتمال توأم مرتبه دوم است. اگر I تصویر و N تعداد درجات خاکستری باشد، آن گاه ماتریس GLCM، ماتریس مربع G با ابعاد N می باشد که در آن (i, j) امین عنصر، تعداد دفعاتی است که درجه خاکستری i در همسایگی درجه

⁶ Entropy

⁷ Correlation

⁸ Contrast

⁹ Mean

¹⁰ Variance

¹¹ Homogeneity

¹² Dissimilarity

¹³ Second Moment

¹ Co-Occurrence Matrix

² Gabor

³ Laplacian

⁴ Gray Level Co-occurrence Matrix

⁵ Histogram

۲-۳- انتخاب ویژگی با الگوریتم ژنتیک^۱

یکی از راه‌های کاهش ابعاد یک مجموعه داده، انتخاب یک زیر مجموعه از ویژگی‌ها از مجموعه‌ی در دسترس می‌باشد. تکنیک‌های انتخاب ویژگی، هیچ‌گونه تغییری در مشخصات و خصوصیات داده‌ی ورودی ایجاد نمی‌کند و تنها ویژگی‌های حاوی اطلاعات مفید را با در نظر گرفتن معیار، انتخاب می‌کند. اگر تعداد کل ویژگی‌ها برابر N باشد، تعداد کل زیرمجموعه‌های ممکن برابر 2^N می‌شود. این تعداد برای N های متوسط هم خیلی زیاد است. بر اساس نحوه جستجو در میان این تعداد زیر مجموعه، روش‌های ختلف انتخاب ویژگی را می‌توان به سه دسته‌ی جست و جوی کامل، جست و جوی مکاشفه‌ای و جست و جوی تصادفی تقسیم کرد [۲۴]. در روش‌هایی که از جستجوی کامل استفاده می‌کنند، تابع تولید کننده بر اساس تابع ارزیابی استفاده شده، تمام فضای جواب (زیرمجموعه‌های ممکن) را برای یافتن جواب بهینه جستجو می‌کند. در روش‌های با جستجوی مکاشفه‌ای، در هر بار اجرای الگوریتم، یک ویژگی به مجموعه ویژگی انتخاب شده، اضافه و یا از آن حذف می‌شود. به همین دلیل پیچیدگی زمانی آنها محدود می‌باشد. در اینگونه موارد، اجرای الگوریتم خیلی سریع و پیاده سازی آنها نیز بسیار ساده است. روش‌هایی که از جستجوی تصادفی استفاده می‌کنند، محدوده کمتری از فضای کل حالات را جستجو می‌کنند، که اندازه این محدوده به حداکثر تعداد تکرار الگوریتم بستگی دارد. در این روش‌ها پیدا شدن جواب بهینه به اندازه منابع موجود و زمان اجرای الگوریتم بستگی دارد. در هر بار تکرار، تابع تولید کننده تعدادی از زیرمجموعه‌های ممکن از فضای جستجو را به صورت تصادفی انتخاب و در اختیار تابع ارزیابی قرار می‌دهد [۲۴].

یکی از روش‌های جست و جوی تصادفی شناخته شده، الگوریتم ژنتیک می‌باشد. این الگوریتم توسط جان هنری هالند در سال ۱۹۶۷ ابداع شده‌است. بعدها این روش با تلاش‌های گلدبرگ در سال ۱۹۸۹، جایگاه خویش را یافته و امروزه نیز به واسطه توانایی‌های خویش، از اهمیت بالایی در میان دیگر روش‌ها برخوردار است [۲۴]. استخراج تعداد زیادی ویژگی به منظور تفکیک عوارض شهری و انتخاب تعدادی از این ویژگی‌ها به عنوان ویژگی‌های کارآمد و استفاده از آن‌ها در مراحل بعد، در واقع یک مسئله‌ی بهینه‌سازی است. بنابر این، الگوریتم ژنتیک گزینه‌ی مناسبی به منظور انتخاب ویژگی‌های بهینه و کارآمد می‌باشد. الگوریتم ژنتیک، از روند ژنتیک ارگانیسم‌های بیولوژیکی، الهام گرفته شده‌است. این الگوریتم، شامل چندین راه حل بالقوه است که کروموزوم یا فرد می‌باشند. یک مجموعه از کروموزوم‌ها برای تشکیل جمعیت به کار گرفته می‌شود. برای هدف انتخاب ویژگی بر اساس ژنتیک، طول هر کروموزوم باید برابر با ویژگی‌های ورودی باشد. در این حالت، مقدار هر ژن، ۱ و ۰، حضور یا عدم حضور ویژگی مربوطه را نمایش می‌دهد. شایستگی هر کروموزوم با استفاده از یک تابع شایستگی^۲، ارزیابی می‌شود. کروموزوم‌های مناسب‌تر، از طریق گام انتخاب، به عنوان والد برای تولید کروموزوم‌های جدید انتخاب می‌شوند. در این گام، دو کروموزوم مناسب انتخاب و از طریق گام هم‌گذاری^۳، هم‌گذاری^۴، با یکدیگر تلفیق می‌شوند. سپس جهش^۴ روی فرزندان برای افزایش میزان تصادفی بودن افراد، که به منظور کاهش احتمال به دام افتادن در بهینه‌ی محلی صورت می‌گیرد، انجام می‌شود [۲۴]. الگوریتم ژنتیک یک روند تکراری است، بنابراین،

² Fitness Function³ Crossover⁴ Mutation¹ Genetic Algorithm

شناخته شده‌ای در حل مسائل محدودیت‌دار هستند صورت می‌گیرد [۲۵].

یک مشخصه‌ی خاص الگوریتم ماشین بردار پشتیبان اینست که به طور همزمان خطای تجربی طبقه‌بندی را به حداقل و حاشیه‌ی هندسی را به حداکثر می‌رساند. بنابراین، ماشین‌های بردار پشتیبان، طبقه‌بندی کننده‌ی بیشترین حاشیه^۱ نیز نامیده می‌شود [۲۵]. ماشین‌های بردار پشتیبان، بردار ورودی را به فضایی با ابعاد بالاتر که یک ابر صفحه‌ی جداکننده‌ی بیشینه در آن ساخته شده‌است، نگاشت می‌کند. دو ابر صفحه‌ی موازی در هر طرف ابر صفحه‌ای که داده‌ها را از هم جدا می‌کند، ساخته شده‌است. صفحه‌ی جداکننده، صفحه‌ای است که فاصله‌ی بین دو ابر صفحه‌ی موازی را بیشینه می‌کند [۲۶]. فرض بر اینست که حاشیه‌ی بیشتر یا فاصله‌ی بیشتر بین ابر صفحه‌های موازی منجر به خطای کمتری برای طبقه‌بندی کننده خواهد شد. در مواردی که با یک ابر صفحه‌ی خطی نمی‌توان داده‌ها را از یکدیگر جدا نمود، بردارهای آموزشی از طریق یک تابع به فضایی با ابعاد بالاتر نگاشت می‌شوند. نقاط داده به‌صورت زیر در نظر گرفته می‌شوند:

$$\{(x_1, y_1), \dots, (x_n, y_n)\} \quad \text{رابطه (۱)}$$

y_n در واقع ۱ یا -۱ است که نشان دهنده‌ی کلاسی است که x_n به آن تعلق دارد. n نیز در واقع تعداد نمونه‌هاست. می‌توان این نمونه‌های آموزشی را با استفاده از ابر صفحه‌ی جداکننده مشاهده کرد.

معادله‌ی این ابر صفحه به صورت رابطه (۲) است:

$$Wx + b = 0 \quad \text{رابطه (۲)}$$

در این معادله x نقطه‌ای از ابر صفحه، بردار W عمود بر ابر صفحه‌ی جداکننده می‌باشد و b نیز فاصله‌ی نزدیکترین نقطه‌ی ابر صفحه تا مبدأ می‌باشد. آموزش SVM شامل پیدا کردن موقعیت ابر صفحه‌ی جدا کننده، تخمین مقدار بردار W و اسکالر b با

این روند تا زمانی که معیار توقف برآورد نشود، ادامه می‌یابد. در این روش، معیارهای مختلفی را می‌توان به‌عنوان معیار توقف در نظر گرفت. به‌عنوان مثال، اگر اختلاف بین بهترین مقدار ارزیابی و میانگین مقادیر تمام ارزیابی‌ها در یک تکرار، کمتر از یک حد آستانه‌ی از قبل تعیین شده باشد، الگوریتم می‌تواند متوقف شود. علاوه بر این، تعداد تکرارها می‌تواند از قبل توسط کاربر به‌منظور توقف روند الگوریتم، تعیین شوند [۲۴].

۲-۴- رویکردهای مختلف طبقه‌بندی

در سال‌های اخیر، رویکردهای مختلفی به منظور طبقه‌بندی پوشش اراضی بر مبنای معیارهای مختلف ارائه شده است. یکی از رویکردهای طبقه‌بندی پوشش اراضی کاربری، طبقه‌بندی پارامتریک و غیر پارامتریک می‌باشد. در روش‌های پارامتریک فرضی از توزیع داده‌ها مطرح است این در حالی است که در روش‌های غیر پارامتریک به هیچ فرضی از توزیع داده‌ها نیاز نیست. از آنجایی که در داده‌های لایدار اطلاع دقیقی از توزیع داده‌ها موجود نیست بنابراین به‌منظور طبقه‌بندی عوارض با استفاده از این داده‌ها، روش‌های غیر پارامتریک از کارایی بالاتری برخوردار هستند. از جمله این روش‌ها می‌توان به ماشین بردار پشتیبان و الگوریتم‌های یادگیری دسته جمعی مانند جنگل تصادفی و بگینگ و بوستینگ اشاره نمود. در ادامه روش کار کلی این طبقه‌بندی کننده‌ها ارائه می‌گردد.

۲-۴-۱- طبقه‌بندی داده‌ها با استفاده از الگوریتم ماشین بردار پشتیبان

ماشین‌های بردار پشتیبان، اولین بار توسط وپنیک در سال ۱۹۶۳ ارائه شد. مبنای کاری طبقه‌بندی کننده‌ی ماشین‌های بردار پشتیبان دسته‌بندی خطی داده‌ها می‌باشد و در تقسیم خطی داده‌ها، سعی بر انتخاب خطی است که حاشیه اطمینان بیشتری داشته باشد. حل معادله‌ی پیدا کردن خط بهینه برای داده‌ها به‌وسیله روش‌های کوادراتیک که روش‌های

¹ Maximal Margin

برای اهداف کلی، بعضی توابع کرنل معروف عبارتند از کرنل خطی^۱، کرنل چند جمله‌ای^۲، کرنل گوسین^۳ و کرنل سیگموئید^۴. در میان این کرنل‌ها، کرنل گوسین، عملکرد بهتری نسبت به سایر کرنل‌ها داشته است، این کرنل بر خلاف کرنل خطی، به صورت غیر خطی نمونه‌ها را به فضایی با ابعاد بالاتر نگاشت می‌کند و فرایامترهای کمتری نسبت به کرنل چند جمله‌ای دارد. کرنل گوسین مشکلات عددی کمتری نسبت به سایر کرنل‌ها دارد [۲۷]. با وجود اینکه برای بعضی از داده‌ها، عملکرد ماشین‌های بردار پشتیبان نسبت به پارامترهای کرنل بسیار حساس است، اما این الگوریتم هم‌چنان از قدرت بالایی برخوردار است. بنابراین، کاربران معمولاً نیاز به انجام ارزیابی، به منظور سنجیدن تنظیمات بهینه‌ی پارامترها دارند [۲۶].

برای طبقه‌بندی با استفاده از الگوریتم ماشین بردار پشتیبان، ابتدا باید پارامترهای هسته و پارامتر تنظیم، تعیین شود. مناسب‌ترین مقادیر برای این پارامترها مشخص نیستند. پارامترهایی که باید تعریف شوند، در واقع پهنای گوسین در کرنل گوسین و ترم پناستی می‌باشند. با استفاده از روش جست و جوی شبکه‌ای و ارزیابی عملکرد الگوریتم با رشد نمایی پارامترها، می‌توان مقادیر آن‌ها را برآورد کرد [۲۸].

۲-۴-۲- طبقه‌بندی داده‌ها با استفاده از

الگوریتم‌های یادگیری دسته‌جمعی^۵

یک الگوریتم یادگیری دسته‌جمعی، یک سیستم طبقه‌بندی چند کلاسه است که با هدف به‌دست آوردن طبقه‌بندی بهتر از طریق ترکیب

توجه به جواب یک مسئله‌ی بهینه‌سازی می‌باشد [۲۶] ابر صفحه‌ی کانونیکال بهینه، ابر صفحه‌ای است که بیشترین حاشیه را دارا می‌باشد. برای تعیین فراصفحه‌ی بهینه برای داده‌های تفکیک پذیر خطی، باید مسئله‌ی بهینه‌سازی زیر حل شود (رابطه (۳)):

$$\min \frac{1}{2} \|W\|^2 \quad (3)$$

$$y_i [W^T \cdot x_i + b] \geq 1 \quad i = 1, 2, \dots, l$$

از آنجاییکه بررسی قیود نابرابری مشکل است، معمولاً با استفاده از ضرایب لاگرانژ $\alpha_{i=1}^N$ مسئله‌ی بهینه‌سازی بالا به یک فضای دوگانه تصویر می‌شود. این مسئله‌ی بهینه‌سازی، توسط نقاط زینی تابع لاگرانژ حل می‌شود. (رابطه (۴))

$$\max [\sum_{i=1}^l \alpha_i - \frac{1}{2} \sum_{i=1}^l \alpha_i \alpha_j x_i y_j x_i^T y_j] \quad (4)$$

$$\alpha_i \geq 0 \quad i = 1, 2, \dots, l$$

$$\sum_{i=1}^l \alpha_i y_i = 0$$

به‌علت وجود نویز و ترکیب کلاس‌ها در هنگام انتخاب نمونه‌های آموزشی، از متغیر $\xi_i > 0$ جهت در نظر گرفتن اثر نمونه‌های آموزشی بدون طبقه‌بندی استفاده می‌شود. بنابر این، معادله‌ی فراصفحه در این حالت برای دو کلاس با رابطه (۵) مشخص می‌شود.

$$y_i (W x_i - b) \geq 1 - \xi_i \quad (5)$$

این فراصفحه با حل مسئله‌ی بهینه‌سازی قید دار زیر حل می‌شود:

$$\min \left\{ \frac{1}{2} \|W\|^2 + C \sum_{i=1}^k \xi_i \right\} \quad (6)$$

$$y_i (W x_i - b) \geq 1 - \xi_i \quad i = 1, 2, \dots, k$$

$C > 0$ پارامتر جریمه‌ی ترم خطا می‌باشد. این پارامتر، اجازه‌ی کنترل جریمه‌ی مربوط به خطاها را می‌دهد. بنابراین، این پارامتر، تعادل بین تعداد نمونه‌های آموزشی برچسب‌گذاری شده و پهنای حاشیه را کنترل می‌کند [۲۶].

کرنل‌های زیادی در الگوریتم ماشین بردار پشتیبان وجود دارد. بنابراین انتخاب یک تابع کرنل مناسب یک موضوع مهم در این الگوریتم است [۲۶]. با این وجود،

¹ Linear kernel

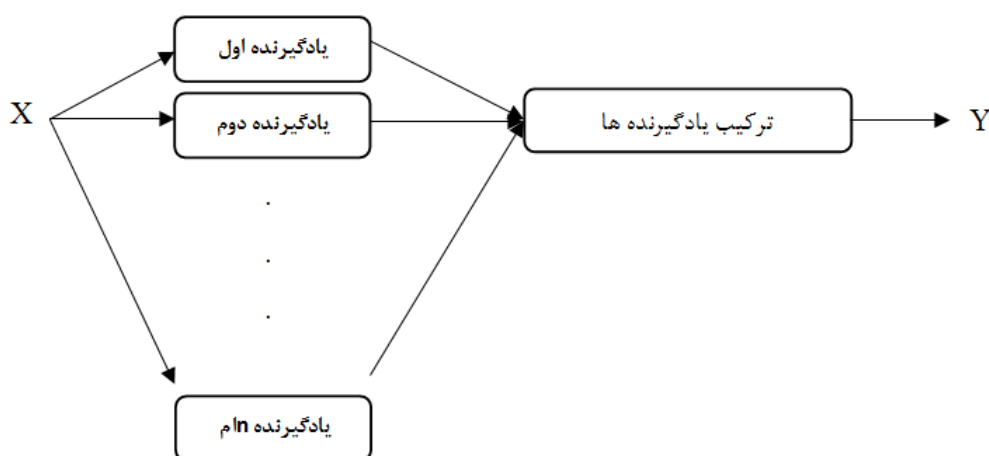
² Polynomial kernel

³ Gaussian kernel

⁴ Sigmoid kernel

⁵ Ensemble learning algorithms

برچسب نهایی کلاس برای داده‌ی ورودی توسط رأی اکثریت طبقه‌بندی کننده‌ها تعیین می‌شود [۲۹]. ماهیت همگرایی و تکراری بودن الگوریتم‌های یادگیری دسته جمعی، پتانسیل خوبی را برای بهبود طبقه‌بندی مناطق پیچیده تأمین می‌کند [۳۰]. شکل (۲)، طراحی کلی یک الگوریتم یادگیری دسته‌جمعی معمول را نمایش می‌دهد [۳۱]



شکل ۲: طرح کلی سیستم‌های یادگیری دسته جمعی [۳۱]

به منظور اختصاص دهی هر پیکسل به یک کلاس براساس بیشترین تعداد رأی که از گروه طبقه‌بندی کننده‌ها دریافت می‌کنند، استفاده می‌شود. هر درخت با استفاده از درصد مشخصی از نمونه‌هایی که به صورت تصادفی انتخاب شده‌اند و درصد باقی مانده از نمونه‌ها، آموزش داده می‌شوند و به‌منظور برآورد دقت طبقه‌بندی مورد استفاده قرار می‌گیرند. علاوه بر این، بهترین انشعاب در هر گره، می‌تواند توسط تعداد ویژگی‌ها و انشعابات تعریف شده توسط کاربر به‌دست آید [۳۰]. نکته‌ی کلیدی الگوریتم جنگل تصادفی، ساده‌سازی محاسبات پیچیده از طریق کاهش تعداد ویژگی‌های ورودی در هر گره می‌باشد. در الگوریتم جنگل تصادفی، نیاز به تعیین دو پارامتر می‌باشد: تعداد درخت‌های تصمیم‌گیری که باید ایجاد شوند (K) و تعداد ویژگی‌هایی که به‌طور تصادفی

خروجی‌های چندین طبقه‌بندی کننده، ایجاد می‌شود. این الگوریتم‌ها در دو گام اجرا می‌شود، این گام‌ها شامل تولید یادگیرنده‌های پایه و سپس ترکیب آن‌ها به‌منظور دستیابی به یک الگوریتم یادگیری خوب می‌باشند. استراتژی‌های مختلفی به‌منظور ترکیب خروجی‌های طبقه‌بندی کننده‌ها وجود دارد، یکی از این استراتژی‌ها، استراتژی رأی اکثریت می‌باشد که اجرای آسانی دارد و معمولاً استفاده می‌شود.

۲-۴-۳- الگوریتم جنگل تصادفی

الگوریتم جنگل تصادفی در سال ۲۰۰۱ توسط برایمن ارائه شد. این الگوریتم، یک طبقه‌بندی کننده‌ی دسته‌جمعی بر مبنای درخت می‌باشد. الگوریتم جنگل تصادفی، ترکیبی از درخت‌های تصمیم‌گیری است که هر درخت تصمیم، یک رأی برای اختصاص دادن بیشترین کلاس به یک بردار ورودی ارائه می‌دهد. این الگوریتم، تنوع درخت‌های تصمیم‌گیری را به‌منظور ارتقای آن‌ها از طریق تغییر مجموعه آموزشی افزایش می‌دهد. درخت تصمیم، نمونه‌های آموزشی را در هر گره با استفاده از قوانین تصمیم‌گیری، به زیر مجموعه‌های کوچکتر تقسیم می‌کند. در هر گره، به منظور پیدا کردن بهترین متغیرها برای تقسیم بندی، بررسی‌هایی انجام می‌شود [۲۹]. قانون بیشترین رأی

در نظر گرفته می‌شود (m تعداد کل ویژگی‌هاست) [۲۹]. شکل (۳) روند کلی الگوریتم جنگل تصادفی را نشان می‌دهد.

برای انشعاب در هر گره در یک درخت در نظر گرفته می‌شود (m). این الگوریتم، نسبت به مقدار m حساسیتی ندارد و اغلب به صورت $\sqrt{m}$

ورودی: مجموعه داده‌ی $D = \{(x_1, y_1), (x_2, y_2), \dots, (x_m, y_m)\}$	
اندازه‌ی زیر مجموعه ویژگی K	
پروسه:	
i.	یک گره درختی بر روی داده D ایجاد کن و در N قرار بده.
ii.	اگر تمام نمونه‌ها در یک کلاس باشند، N را برگردان.
iii.	مجموعه‌ای از ویژگی‌ها که می‌توان باز هم تقسیم کرد در f قرار بده.
iv.	اگر f خالی بود، N را برگردان.
v.	K ویژگی را به طور تصادفی از f انتخاب کن و در f قرار بده.
vi.	ویژگی‌ای که بهترین گره درختی را در f دارد، در $N.f$ قرار بده.
vii.	بهترین گره درختی را در روی $N.f$ در $N.P$ قرار بده.
viii.	زیر مجموعه D با مقادیر $N.f$ کوچکتر از $N.P$ را در D_l قرار بده.
ix.	پروسه را با پارامترهای (D_l, K) فراخوانی کن و N_l را ایجاد کن.
x.	پروسه را با پارامترهای (D_r, K) فراخوانی کن و N_r را ایجاد کن.
xi.	N را برگردان.
xii.	خروجی: یک درخت تصمیم تصادفی

شکل ۳: روند کلی الگوریتم جنگل تصادفی [۲۹]

طبقه‌بندی بهینه، نیاز به تنظیم تنها یک پارامتر دارد [۲۹]. در شکل (۴)، روند کلی الگوریتم بگینگ نشان داده شده است.

ورودی: مجموعه داده‌ی	
$D = \{(x_1, y_1), (x_2, y_2), \dots, (x_m, y_m)\}$	
الگوریتم یادگیری پایه $\mathcal{S}$	
تعداد یادگیرنده‌های پایه T	
پروسه:	
i.	برای $T = 1, \dots, t$
ii.	$h_t = \mathcal{S}(D, D_{bs})$ ، توزیع D_{bs} bootstrap می‌باشد.
iii.	پایان

شکل ۴: روند کلی الگوریتم بگینگ [۲۹]

۲-۴-۵- الگوریتم بوستینگ

الگوریتم بوستینگ، یک تکنیک یادگیری دسته‌جمعی

۲-۴-۴- الگوریتم بگینگ

الگوریتم بگینگ در سال ۱۹۹۶ توسط برایمن به منظور بهبود دقت طبقه‌بندی و تعمیم الگوهای داده معرفی شد. این الگوریتم، مانند الگوریتم جنگل تصادفی، از گروهی از طبقه‌بندی کننده‌های بر مبنای درخت تشکیل شده است. اختصاص دادن یک کلاس به هر پیکسل بر اساس بیشترین تعداد رأی که از گروه طبقه‌بندی کننده‌ها برای هر پیکسل دریافت شده است، انجام می‌شود. تفاوت بین این دو الگوریتم اینست که در الگوریتم بگینگ، تمام ویژگی‌ها برای انشعاب در هر گره، کاندید هستند. در حالیکه در جنگل تصادفی، از یک تعداد ویژگی ثابت که توسط کاربر به صورت تصادفی تعریف می‌شود، استفاده می‌گردد. بنابراین، الگوریتم بگینگ، به منظور دستیابی به نتیجه‌ی

مناطق متراکم مسکونی، پوشش گیاهی، جاده و اتومبیل می‌باشد. ساختمان‌ها با ارتفاعات و رنگ‌های متنوع که به علت موادهای مختلف تشکیل دهنده‌ی ساختمان‌ها مثل سیمان، موزاییک و فلز ایجاد شده‌اند، استخراج عوارض را بسیار مشکل کرده‌است. شکل (۶)، ارتوفتوی رنگی منطقه را نمایش می‌دهد.

مجموعه داده‌ی مورد استفاده در این تحقیق، در سال ۲۰۱۱ با استفاده از لیزر اسکنر هوایی در ارتفاع ۳۰۰ متری زمین بر فراز منطقه‌ی شهری زیرگ در کشور بلژیک جمع‌آوری شده‌است. ارتوفتوی رنگی به دست آمده در سه باند قرمز، سبز و آبی با قدرت تفکیک مکانی برابر ۵ سانتی‌متر، می‌باشد. فواصل نقاط لایدار جمع‌آوری شده نیز ۱۰ سانتی‌متر و بسیار متراکم می‌باشد. چگالی داده‌های لایدار، ۶۵ نقطه در هر مترمربع است. جمع‌آوری داده‌ها به صورت همزمان انجام شده و مجموعه‌ی داده‌ها به سیستم ژئودتیک جهانی ۱۹۸۴ (WGS84^۱) هم مرجع شده‌اند. داده‌های آموزشی و تست به منظور انجام طبقه‌بندی نیز توسط یک عامل خبره از طریق تفسیر تصویر هوایی با قدرت تفکیک مکانی بالا در شش کلاس شهری چمن، درخت، جاده، ساختمان با سقف شیروانی، ساختمان با سقف مسطح و اتومبیل به دست آمد. این داده‌ها در شکل (۷) نمایش داده شده‌اند.

است که در سال ۱۹۹۶ توسط فرویند و شاپیرو با هدف بهبود دقت طبقه‌بندی و با تبدیل یک گروه ضعیف از طبقه‌بندی کننده‌ها به یک طبقه‌بندی کننده‌ی قوی انجام شد [۲۹]. این الگوریتم، توسط آموزش یک مجموعه از یادگیرنده‌ها به طور متوالی و ترکیب آن‌ها به منظور پیش بینی، کار می‌کند. بر خلاف الگوریتم‌های جنگل تصادفی و بگینگ، الگوریتم بوستینگ به منظور طبقه‌بندی، از تمام مجموعه داده، به جای نمونه‌برداری از یک بخش خاصی از مجموعه داده استفاده می‌کند. ابتدا در اولین تکرار، وزن‌های برابری به هر نمونه اختصاص داده می‌شود و بعد از آن، در تکرارهای بعدی، وزن‌ها بر اساس برآزش به داده‌ها، دوباره تنظیم می‌شود. در هر تکرار نمونه‌هایی که در تکرارهای قبلی به درستی طبقه‌بندی نشده‌اند، وزن‌های بیشتری نسبت به نمونه‌هایی که در تکرارهای قبلی به درستی طبقه‌بندی شده‌اند، به خود می‌گیرند و خطاهای طبقه‌بندی که در تکرارهای قبلی اتفاق افتاده است با تکرارهای بعدی، تصحیح می‌شود. شکل (۵) روند کلی الگوریتم بوستینگ را نمایش می‌دهد.

ورودی: توزیع نمونه D

الگوریتم یادگیری پایه $\mathcal{S}$

تعداد تکرارهای یادگیری T

پروسه:

i. $D_1 = D$ % توزیع اولیه

ii. برای $t = 1, \dots, T$

iii. $h_t = \mathcal{S}(D_t)$ % آموزش یک یادگیرنده‌ی ضعیف از

توزیع D_t

iv. $\epsilon_t = P_{x \sim D_t}(h_t(x) \neq f(x))$ % ارزیابی خطای h_t

v. $D_{t+1} = \text{Adjust_Distribution}(D_t, \epsilon_t)$

vi. پایان

شکل ۵: روند کلی الگوریتم بوستینگ [۲۹]

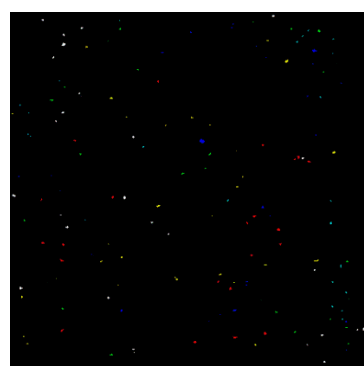
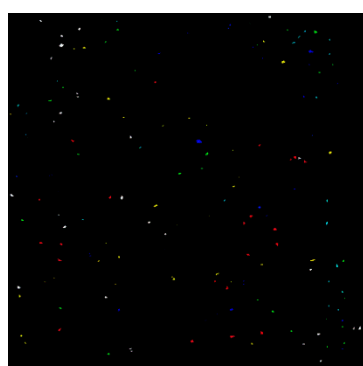
۳- منطقه‌ی مورد مطالعه و داده‌ها

منطقه‌ی مورد مطالعه در این تحقیق، در شهر زیربرگ در کشور بلژیک واقع شده‌است. این منطقه شامل

¹ World Geodetic System 1984



شکل ۶: ارتوفتوی رنگی با قدرت تفکیک مکانی ۵ سانتیمتر

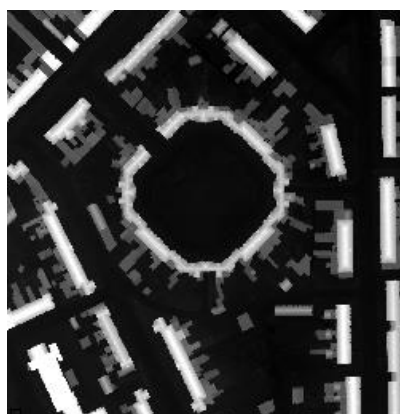


چمن	سبز
درخت	آبی
جاده	زرد
ساختمان با سقف شیروانی	سفید
ساختمان با سقف مسطح	قرمز
اتومبیل	آبی روشن

ب. داده‌های تست

الف. داده‌های آموزشی

شکل ۷: داده‌های آموزشی و تست منطقه



شکل ۸: مدل رقومی نرمال شده‌ی سطح منطقه

در مرحله‌ی بعد، از تصویر هوایی و nDSM منطقه، ویژگی‌های بافتی استخراج شدند. با محاسبه‌ی تغییرات بافت در یک پنجره‌ی 3×3 و میانگین‌گیری در چهار جهت ۰، ۴۵، ۹۰، ۱۳۵ در تصویر، ویژگی‌های آماری مرتبه دوم از ماتریس GLCM استخراج شدند.

۴- پیاده‌سازی و بحث

در مرحله‌ی اول تحقیق، به منظور فیلترینگ داده‌های لایدار، از الگوریتم اکسلسون استفاده شد. در تنظیمات این الگوریتم، باید نوع منطقه‌ی مورد مطالعه و گرانولیت‌ی سطح را مشخص نمود. در این مطالعه، تمام حالت‌های موجود را بررسی نموده و از مقایسه‌ی نتایج به این نتیجه رسیدیم که حالت شهری و گرانولیت‌ی نرم، بهترین نتیجه را در مقایسه با سایر حالت‌ها به دست می‌دهد. در این مرحله، ورودی الگوریتم، DSM و خروجی آن nDSM می‌باشد. در شکل (۸)، مدل رقومی نرمال شده‌ی سطح، نمایش داده شده‌است.

کتر شود، الگوریتم متوقف خواهد شد. علاوه بر این اگر به ۱۰۰ تکرار برسیم، در روند اجرایی الگوریتم توقف خواهیم داشت. پارامترهای الگوریتم ژنتیک استفاده شده برای این منظور در جدول (۱) آورده شده است. با تنظیم این پارامترها، از میان ۴۱ ویژگی استخراج شده، ۱۴ ویژگی بهینه از ویژگی‌های اپتیک و ۷ ویژگی بهینه از ویژگی‌های لایدار به منظور استفاده در طبقه‌بندی، انتخاب شدند.

شکل (۹) نمودار همگرایی الگوریتم ژنتیک را با پارامترهای تنظیم شده نمایش می‌دهد. نقاط آبی رنگ، نشان دهنده بهترین مقادیر تابع ارزیابی در هر نسل و نقاط قرمز رنگ نمایشگر میانگین مقادیر تابع ارزیابی در هر نسل می‌باشد. پس از هر بار تکرار، جواب‌های بد حذف و جواب‌های نسل جدید که مقدار تابع ارزیابی بالایی دارند جایگزین آن‌ها می‌شوند و به مرور فاصله بین مقدار تابع ارزیابی بین حداقل و حداکثر کم و میانگین جواب‌ها به بهترین جواب میل می‌کند این روند در کل نشان دهنده نزدیک شدن به جواب بهینه می‌باشد. با اجرای الگوریتم ژنتیک از میان ۴۱ ویژگی استخراج شده، ۱۴ ویژگی بهینه از ویژگی‌های تصویر هوایی و ۷ ویژگی بهینه از ویژگی‌های لایدار به منظور استفاده در طبقه‌بندی، انتخاب شدند.

جدول ۱: تنظیم پارامترهای الگوریتم ژنتیک

پارامتر	مقدار
جمعیت اولیه	۵۰
تعداد نسل‌ها	۱۰۰
نرخ همگرایی	۰/۸
نرخ جهش	۰/۲
شرط همگرایی	$1e^{-4}$
نوع همگرایی	Scattered
نوع جهش	Gaussian
نوع انتخاب	Tournament Selection

در شکل‌های (۱۰، ۱۱، ۱۲ و ۱۳) به ترتیب ویژگی بافتی کنتراست استخراج شده از nDSM منطقه،

این ویژگی‌ها عبارتند از: آنتروپی، وابستگی، کنتراست، میانگین، واریانس، همگنی، عدم شباهت و گشتاور دوم. از داده‌ی لایدار نرمال شده، هشت ویژگی و از ارتوفتوی رنگی، ۲۴ ویژگی (هشت ویژگی برای هر باند) استخراج گردید.

علاوه بر این، فیلتر گابور در یک پنجره 3×3 و در هشت جهت اصلی $\frac{\pi}{8} 2(n-1)$ با $n = 1, 2, \dots, 8$ برای طول موج‌های ۸، ۱۲، ۱۶ و ۲۴ اجرا گردید. بنابراین، ۳۲ ویژگی از داده‌ی لایدار و ۹۶ ویژگی از داده‌ی اپتیک (۳۲ ویژگی برای هر باند) استخراج شد. با استفاده از تمام خروجی‌های به‌دست آمده از طول موج‌ها و جهات مختلف، می‌توان میانگین و انحراف معیار که ویژگی‌های گابور هستند، را محاسبه نمود. در واقع دو ویژگی گابور برای داده‌های لایدار و دو ویژگی گابور برای داده‌های اپتیک به‌دست آمد.

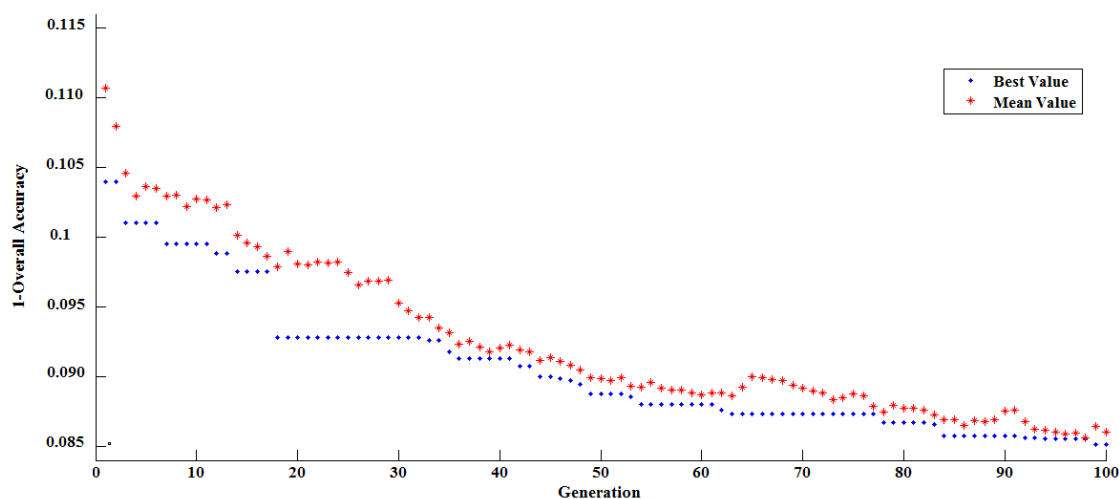
علاوه بر این به‌منظور استخراج ویژگی بافتی لاپلاسی، با اعمال فیلتر لاپلاسی در یک پنجره 3×3 ، یک ویژگی لاپلاسی از داده‌ی لایدار و سه ویژگی از تصویر هوایی (یک ویژگی برای هر باند) استخراج گردید.

با توجه به اینکه تمامی ویژگی‌های استخراج شده در بالا به بررسی تغییرات ارتفاعی در حیطه‌ی مکان پرداخته‌اند، از nDSM منطقه نیز به‌عنوان یک ویژگی به‌منظور بررسی ارتفاع مطلق در تصویر، استفاده کردیم.

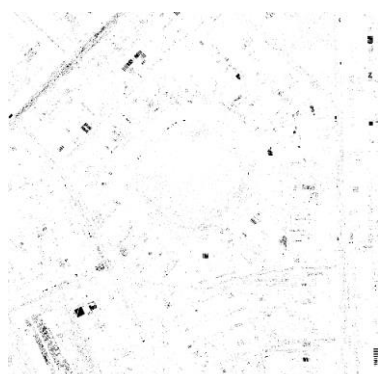
با استخراج ۲۹ ویژگی از داده‌های اپتیک (دو ویژگی گابور، سه ویژگی لاپلاسی، بیست و چهار ویژگی GLCM) و ۱۲ ویژگی از داده‌های لایدار (دو ویژگی گابور، یک ویژگی لاپلاسی، یک ویژگی nDSM و هشت ویژگی GLCM)، در مجموع ۴۱ ویژگی داریم. به‌منظور انتخاب ویژگی‌های بهینه، از الگوریتم ژنتیک با شرط همگرایی $1e^{-4}$ استفاده نمودیم. در واقع اگر تغییرات دقت کلی الگوریتم طبقه‌بندی (تابع شایستگی الگوریتم ژنتیک) از این مقدار

استخراج شده از nDSM به عنوان نمونه‌ی ویژگی‌های استخراج شده‌ی بهینه آورده شده است.

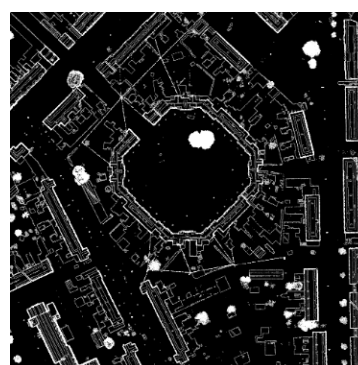
ویژگی بافتی واریانس استخراج شده از باند R تصویر هوایی، ویژگی بافتی آنتروپی استخراج شده از باند B تصویر هوایی و ویژگی بافتی میانگین گابور



شکل ۹: نمودار همگرایی الگوریتم ژنتیک



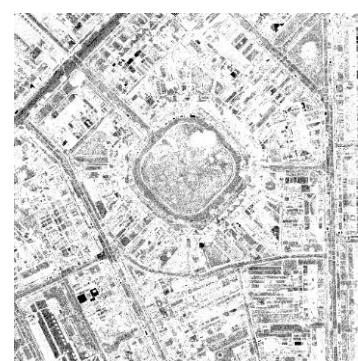
شکل ۱۲: ویژگی بافتی آنتروپی استخراج شده از باند B تصویر هوایی



شکل ۱۰: ویژگی بافتی کنتراست استخراج شده از nDSM با استفاده از ماتریس GLCM



شکل ۱۳: ویژگی بافتی میانگین گابور استخراج شده از DSM



شکل ۱۱: ویژگی بافتی واریانس استخراج شده از باند R تصویر هوایی

توسط کاربر حساس هستند (به عنوان مثال درخت‌ها). انتخاب تعداد کم درخت می‌تواند دقت طبقه‌بندی را به علت افزایش خطای تعمیم بخشی که به خاطر ناتوانی الگوریتم طبقه‌بندی در تعمیم الگوهای داده، ایجاد شده‌است، کاهش دهد. از طرف دیگر، انتخاب تعداد زیاد درخت نیز دقت طبقه‌بندی را افزایش نمی‌دهد زیرا خطای تعمیم بخشی به خارج از یک تعداد خاصی از درخت‌ها محدود می‌شود [۲۹]. از آنجایی که الگوریتم‌های جنگل تصادفی، بگینگ و بوستینگ از درخت‌های طبقه‌بندی کننده‌ی پایه استفاده می‌کنند، رشد تعداد درخت‌های طبقه‌بندی به منظور اجرای الگوریتم، امری رایج است و مقایسه‌ی عملکرد این سه الگوریتم بر مبنای تأثیر تعداد درخت‌ها بر دقت کلی انجام می‌گیرد. بنابراین به منظور ارزیابی تأثیر تعداد درختان طبقه‌بندی روی دقت کلی طبقه‌بندی، تعداد درختان این الگوریتم‌ها را از ۱ تا ۱۰۰۰ تغییر می‌دهیم.

دقت کلی طبقه‌بندی عوارض منطقه در شش کلاس شهری مختلف توسط الگوریتم ماشین بردار پشتیبان، برای داده‌ی اپتیک ۸۳/۲۶٪، برای داده‌ی لایدار ۷۱/۰۵٪ و برای تلفیق این دو داده ۹۳/۹۹٪ حاصل شد. نتایج نشان دهنده‌ی افزایش دقت طبقه‌بندی با استفاده‌ی هم‌زمان از داده‌ی لایدار و تصویر هوایی می‌باشد. شکل (۱۳)، تصویر طبقه‌بندی شده از تلفیق داده‌ی لایدار و تصویر هوایی با استفاده از الگوریتم ماشین بردار پشتیبان و تنظیم بهترین پارامترها برای آن، نشان داده شده است. نتایج طبقه‌بندی عوارض شهری با استفاده از الگوریتم ماشین بردار پشتیبان و تلفیق داده‌ی لایدار و اپتیک و دقت تولید کننده و دقت کاربر مربوط به هر کدام از کلاس‌ها در جدول (۲) آورده شده‌است. دقت تولید کننده مربوط به یک کلاس از نسبت تعداد پیکسل‌های واقعیت زمینی که به درستی به کلاس مورد نظر نسبت داده شده‌اند، به تعداد کل پیکسل‌های واقعیت زمینی در آن کلاس محاسبه می‌گردد.

در نهایت، با اعمال الگوریتم‌های ماشین بردار پشتیبان، جنگل تصادفی، بگینگ و بوستینگ روی داده‌های لایدار، اپتیک و تلفیق لایدار و اپتیک، طبقه‌بندی عوارض شهری در شش کلاس مختلف انجام شد. با به کارگیری الگوریتم ماشین بردار پشتیبان روی این داده‌ها و با استفاده از روش جست و جوی شبکه‌ای، طبقه‌بندی عوارض شهری صورت گرفت. داده‌های آموزشی به منظور انجام طبقه‌بندی نیز توسط یک عامل خبره از طریق تفسیر تصویر هوایی با قدرت تفکیک مکانی بالا به دست آمد. با توجه به نتایج به دست آمده از مطالعات قبلی، بهترین کرنل برای الگوریتم ماشین بردار پشتیبان، کرنل گوسین می‌باشد که در این تحقیق نیز از این کرنل استفاده شده است [۲۱]. این کرنل دارای دو پارامتر γ و C می‌باشد که باید قبل از طبقه‌بندی تنظیم شوند. با بررسی رشد پارامترهای این کرنل به صورت نمایی، بهترین مقدار برای پارامتر جریمه، با استفاده از روش جست و جوی شبکه‌ای برای داده‌های لایدار، اپتیک و تلفیق این دو داده، $C = 1000$ به دست آمد. تغییرات پارامتر پهنای باند گوسین (γ) نیز تأثیری در دقت طبقه‌بندی وارد نمی‌کند. بنابراین، با تنظیم پارامترهای بهینه برای الگوریتم ماشین بردار پشتیبان و اعمال این پارامترها به داده‌های آموزشی، طبقه‌بندی کننده اجرا شد.

در این مطالعه از ماتریس ابهام جهت ارزیابی دقت استفاده می‌کنیم و ضریب کاپا برای این منظور محاسبه می‌شود. بنابراین از طبقه‌بندی کننده‌ی ماشین بردار پشتیبان، برای طبقه‌بندی داده‌های تست برای دستیابی به دقت استفاده می‌شود و دقت هر یک از کلاس‌ها با استفاده از ماتریس ابهام به دست می‌آید.

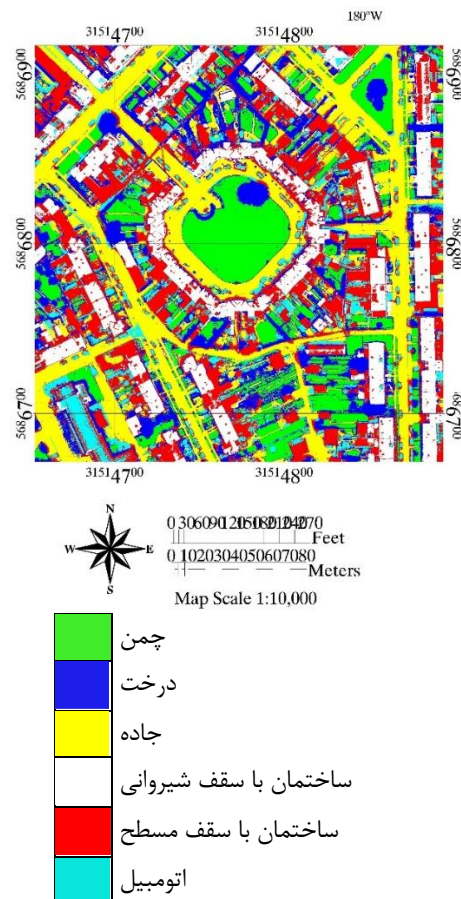
علاوه بر این، الگوریتم‌های یادگیری دسته جمعی که شامل الگوریتم‌های جنگل تصادفی، بگینگ و بوستینگ می‌باشند، روی داده‌های اپتیک، لایدار و تلفیق اپتیک و لایدار به منظور طبقه‌بندی عوارض شهری اجرا شدند. این الگوریتم‌ها نسبت به مقادیر پارامترهای تعریف شده

کلاس شده‌اند. در طبقه‌بندی ساختمان‌های با سقف شیروانی و جاده، پیکسل‌های سایر کلاس‌ها کمتر به این کلاس‌ها اختصاص داده شده‌اند و به همین دلیل این کلاس‌ها از دقت کاربر بالایی برخوردار هستند. علاوه بر این، با تنظیم بهترین تعداد درخت برای الگوریتم جنگل تصادفی، دقت $82/9241\%$ برای داده‌ی اپتیک، $74/9487\%$ برای داده‌ی لایدار و $90/4472\%$ برای تلفیق داده‌های لایدار و اپتیک حاصل شد.

با اعمال الگوریتم بگینگ و بررسی اثر تغییر تعداد درختان روی این الگوریتم به‌منظور دستیابی به تعداد درخت بهینه، دقت $82/1330\%$ برای داده‌ی اپتیک و دقت $71/2687\%$ برای داده‌ی لایدار و دقت $90/6065\%$ برای تلفیق این دو داده حاصل شد. با تنظیم بهترین تعداد درخت برای الگوریتم بوستینگ، دقت $65/5318\%$ برای داده‌ی اپتیک و دقت $60/4805\%$ برای داده‌ی لایدار و دقت $60/4805\%$ برای تلفیق این دو داده حاصل شد.

ملاحظه می‌شود که اعمال تمامی الگوریتم‌های بالا بر روی تلفیق داده‌های اپتیک و لایدار، منجر به دقت کلی بالاتری نسبت به سایر داده‌ها می‌شود. جدول (۲) به ترتیب نتایج طبقه‌بندی عوارض شهری در شش کلاس مختلف با استفاده از الگوریتم‌های ماشین بردار پشتیبان با $C=1000$ و $\gamma=1$ ، جنگل تصادفی با ۹۰۰ درخت، بگینگ با ۹۰۰ درخت و بوستینگ با ۱۰۰۰ درخت و تلفیق داده‌ی لایدار و اپتیک و دقت تولید کننده و دقت کاربر مربوط به هر کدام از کلاس‌ها را نشان می‌دهد. به‌منظور دقت الگوریتم‌های طبقه‌بندی از معیار توافق^۱ ارائه شده توسط [۳۲] استفاده شد. الگوریتم‌های طبقه‌بندی ماشین بردار پشتیبان، جنگل تصادفی، بگینگ و بوستینگ به ترتیب دارای مقادیر 94% ، 90% ، 91% و 47% برای معیار توافق بودند که نشان می‌دهد

دقت کاربر مربوط به یک کلاس نیز از نسبت تعداد پیکسل‌های واقعیت زمینی که به درستی به کلاس مورد نظر نسبت داده شده‌اند، به تعداد کل پیکسل‌های واقعیت زمینی از تمامی کلاس‌ها که به کلاس مورد نظر نسبت داده شده‌اند محاسبه می‌گردد.



شکل ۱۴: تصویر طبقه‌بندی شده از تلفیق داده‌ی لایدار و تصویر هوایی با استفاده از الگوریتم ماشین بردار پشتیبان برای $C=1000, \gamma=1$

همان‌طور که ملاحظه می‌شود، با اجرای این الگوریتم بر روی تلفیق این دو داده، پیکسل‌هایی از کلاس جاده در کلاس‌های دیگر طبقه‌بندی شده و منجر به کاهش دقت تولید کننده در کلاس جاده شده‌اند. علاوه بر این، در طبقه‌بندی کلاس اتومبیل، پیکسل‌هایی از کلاس‌های دیگر به اشتباه به این کلاس اختصاص داده شده و باعث کاهش دقت کاربر در این

¹ Index of agreement

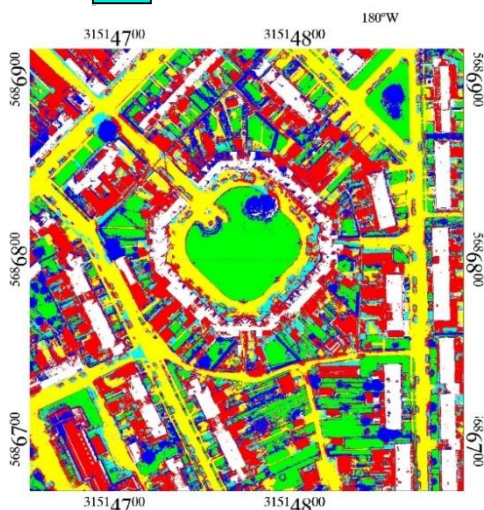
دقت روش ماشین بردار پشتیبان از سایر روش‌ها بهتر بوده است.

جدول ۲: ارزیابی دقت کلاس‌های طبقه‌بندی شده با استفاده از الگوریتم‌های ماشین بردار پشتیبان، جنگل تصادفی، بگینگ و بوستینگ

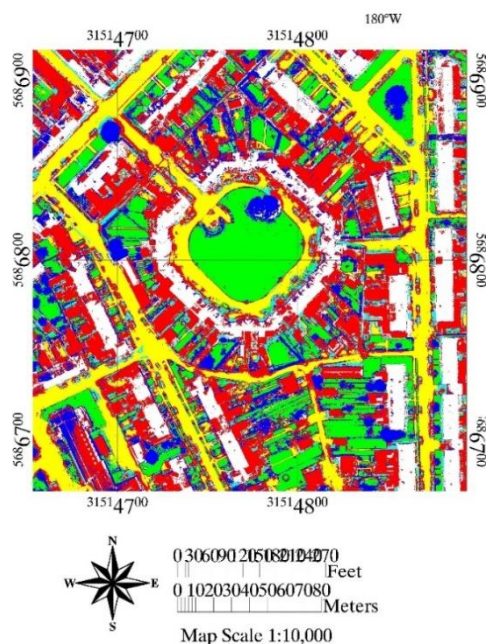
نام الگوریتم	دقت تولید کننده و کاربر (درصد)	چمن	درخت	جاده	ساختمان با سقف شیروانی	ساختمان با سقف مسطح	اتومبیل
ماشین بردار پشتیبان	دقت تولید کننده	۹۸٫۸۵	۹۱٫۷۲	۸۰٫۸۸	۹۰٫۴۹	۹۴٫۸۹	۹۸٫۹۴
	دقت کاربر	۹۶٫۲۲	۹۳٫۴۶	۹۹٫۳۹	۹۹٫۷۷	۸۸٫۹۴	۷۰٫۵۲
جنگل تصادفی	دقت تولید کننده	۹۸٫۴۷	۸۳٫۸۳	۸۳٫۱۰	۸۸٫۳۲	۹۷٫۶۳	۸۸٫۹۵
	دقت کاربر	۹۶٫۰۹	۹۴٫۸۵	۸۳٫۱۰	۹۳٫۳۴	۸۴٫۴۸	۸۸٫۹۵
بگینگ	دقت تولید کننده	۹۸٫۹۸	۸۴٫۶۴	۸۱٫۳۸	۸۷٫۰۳	۹۷٫۳۳	۹۶٫۶۹
	دقت کاربر	۹۵٫۹۱	۹۴٫۷۴	۹۸٫۹۷	۹۴٫۴۵	۸۶٫۵۶	۷۲٫۰۴
بوستینگ	دقت تولید کننده	۰٫۰۰	۵۱٫۷۱	۹۷٫۰۴	۸۲٫۰۹	۷۷٫۹۲	۰٫۰۰
	دقت کاربر	۰٫۰۰	۸۲٫۰۰	۳۳٫۶۸	۸۳٫۳۰	۸۸٫۵۰	۰٫۰۰



در شکل‌های (۱۵، ۱۶ و ۱۷) به ترتیب تصویر طبقه‌بندی شده‌ی حاصل از تلفیق داده‌های لایدار و اپتیک توسط الگوریتم جنگل تصادفی، بگینگ و بوستینگ با تعداد درخت بهینه نمایش داده شده است.



شکل ۱۶: تصویر طبقه‌بندی شده از تلفیق داده‌ی لایدار و تصویر هوایی با استفاده از الگوریتم بگینگ با تعداد درخت ۹۰۰

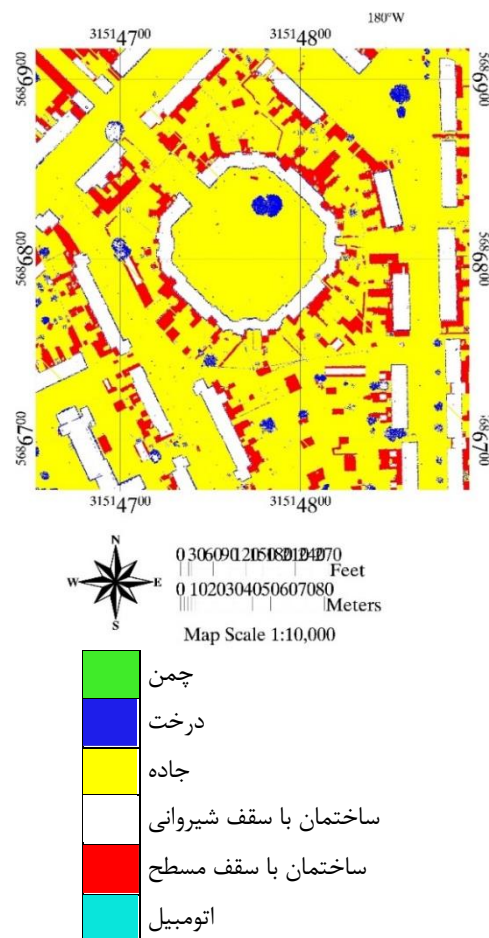


شکل ۱۵: تصویر طبقه‌بندی شده از تلفیق داده‌ی لایدار و تصویر هوایی با استفاده از الگوریتم جنگل تصادفی با تعداد درخت ۹۰۰

جنگل تصادفی و بگینگ دو عارضه‌ی اتومبیل و چمن را با دقت بالایی طبقه‌بندی می‌کنند.

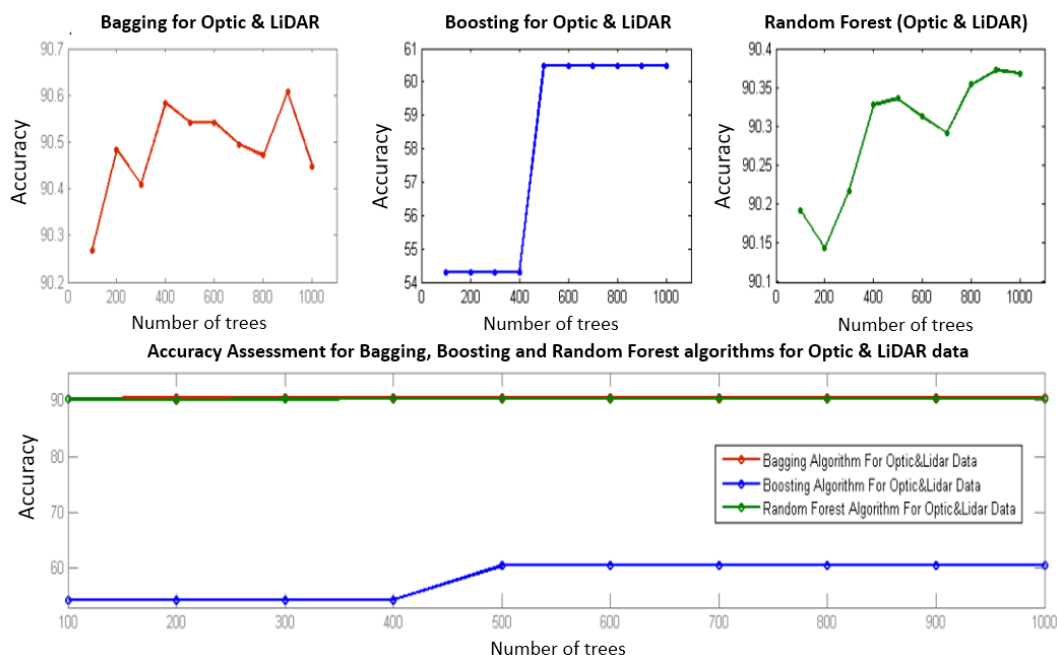
۵- نتیجه گیری

استفاده از داده‌های با توان تفکیک مکانی بالا برای استخراج عوارض شهری به علت پیچیدگی مناطق و جزئیات بالا در داده‌ها به یکی از دغدغه‌های متخصصین تبدیل شده است. داده‌های لایدار و تصاویر هوایی برای استخراج این عوارض شهری با دقت و صحت بالا به‌کار برده شده‌اند. در این تحقیق، نتایج اجرای الگوریتم ماشین بردار پشتیبان و الگوریتم‌های یادگیری دسته جمعی بر روی داده‌های لایدار و تصویر هوایی با قدرت تفکیک مکانی بسیار بالا، به‌منظور طبقه‌بندی عوارض شهری، ارائه شد. اعمال الگوریتم‌های ماشین بردار پشتیبان، جنگل تصادفی، بگینگ و بوستینگ بر روی تلفیق داده‌های اپتیک و لایدار، دقت بالاتری نسبت به اعمال آن بر روی ارتوفتوی رنگی یا داده‌ی لایدار به صورت مجزا به‌دست می‌دهد. دقت کلی طبقه‌بندی کننده‌ی ماشین بردار پشتیبان، $93/99\%$ ، دقت طبقه‌بندی کننده‌ی جنگل تصادفی، $90/4472\%$ ، دقت طبقه‌بندی کننده‌ی بگینگ، $90/6065\%$ و دقت طبقه‌بندی کننده‌ی بوستینگ، $60/4805\%$ به‌دست آمد. الگوریتم ماشین بردار پشتیبان توانایی بالاتری در طبقه‌بندی تلفیق داده‌های لایدار و تصویر هوایی از خود نشان داد. با استفاده از روش جست و جوی شبکه‌ای و ارزیابی عملکرد الگوریتم با رشد نمایی پارامترها، مقادیر بهینه‌ی پارامترها برآورد شد و بهترین عملکرد این الگوریتم حاصل گردید. علاوه بر این به‌منظور ارزیابی تأثیر تعداد درختان طبقه‌بندی روی دقت کلی الگوریتم‌های یادگیری دسته جمعی، تغییرات تعداد درختان این الگوریتم‌ها از ۱ تا ۱۰۰۰ مورد بررسی قرار گرفت و مقدار بهینه‌ی این پارامتر برآورد شد.



شکل ۱۷: تصویر طبقه‌بندی شده از تلفیق داده‌ی لایدار و تصویر هوایی با استفاده از الگوریتم بوستینگ ۱۰۰۰

در شکل (۱۸)، تأثیر تعداد درخت‌ها بر دقت کلی الگوریتم‌های جنگل تصادفی، بگینگ و بوستینگ اعمال شده روی تلفیق داده‌های اپتیک و لایدار نشان داده شده است. ملاحظه می‌شود الگوریتم جنگل تصادفی با تعداد ۹۰۰ درخت، بگینگ با تعداد ۹۰۰ درخت و بوستینگ با تعداد ۱۰۰۰ درخت به بهترین نتایج خود دست می‌یابند. الگوریتم‌های جنگل تصادفی و بگینگ عملکرد نسبتاً مشابهی دارند و دقت کلی بالاتری نسبت به الگوریتم بوستینگ دارند. الگوریتم بوستینگ در طبقه‌بندی کلاس اتومبیل و چمن به مشکل بر می‌خورد و توانایی طبقه‌بندی این دو عارضه را ندارد. این در حالی است که الگوریتم‌های



شکل ۱۸: تأثیر تعداد درخت‌ها روی دقت کلی اجرای الگوریتم‌های جنگل تصادفی، بگینگ و بوستینگ روی تلفیق داده‌های اپتیک و لایدار

در این تحقیق، اغلب ویژگی‌های استخراج شده از داده‌ها در گروه ویژگی‌های بافتی قرار دارند، به‌منظور بررسی هر چه بیشتر تأثیر ویژگی‌ها روی دقت طبقه‌بندی، می‌توان از ویژگی‌های سطح و ارتفاع نیز استفاده نمود و فضای ویژگی‌ها را گسترش داد. به‌منظور بررسی دقیق‌تر عملکرد الگوریتم‌های یادگیری دسته جمعی در تحقیقات بعدی، می‌توان سایر پارامترهای مؤثر بر آن‌ها مانند اندازه‌ی داده‌ها و استحکام داده‌ها را نیز در نظر گرفت و تأثیر این پارامترها را روی دقت نتایج طبقه‌بندی مورد بررسی قرار داد.

با اعمال الگوریتم‌های به‌کار گرفته شده بر روی مناطق دیگر، می‌توان به ارزیابی دقیق‌تری از این الگوریتم‌ها دست یافت. از آنجایی که الگوریتم بوستینگ، نتایج دقیقی از طبقه‌بندی منطقه ارائه نداد، نیاز به بررسی بیشتر و جامع‌تر در مناطق مختلف دارد. باید توجه داشت که امکان دستیابی به دقت بالاتر برای تمامی الگوریتم‌های به‌کار گرفته شده با تنظیم پارامترهای مؤثر بر آن‌ها وجود دارد.

باید توجه داشت که دقت روش‌های مختلف طبقه‌بندی در مناطق مختلف دارای نتایج متفاوتی است. به عبارت دیگر دقت طبقه‌بندی کننده‌ها به منطقه مورد مطالعه و داده‌ی مورد استفاده وابسته است. بنابراین، نتایج ضعیف الگوریتم بوستینگ در این تحقیق دلیلی بر انکار توانایی‌های این الگوریتم نمی‌باشد و امکان این امر که با اجرای این الگوریتم بر روی داده‌های دیگر و در سایر مناطق، طبقه‌بندی عوارض با دقت زیاد انجام شود، وجود دارد.

با در نظر گرفتن مطالعات انجام شده در زمینه طبقه‌بندی پوشش اراضی شهری و نتایج این تحقیق، می‌توان گفت روش‌های پیشنهاد شده در این تحقیق، همواره نیازمند اصلاحات و مطالعات بالاتر به‌منظور تکمیل و گسترش بیشتر هستند. همچنین از آنجایی که عملکرد هر یک از این طبقه‌بندی کننده‌های مورد استفاده در این تحقیق، تحت تأثیر پارامترها و فضای ویژگی‌های ورودی می‌باشند، حل همزمان این دو منجر به بهبود چشمگیری در بهبود دقت طبقه‌بندی می‌گردد.

مراجع

- [1] B. Sirmacek and P. Reinartz, "Automatic crowd density and motion analysis in airborne image sequences based on a probabilistic framework," in Computer Vision Workshops (ICCV Workshops), 2011 IEEE International Conference on, 2011, pp. 898-905.
- [2] B. Sirmacek and P. Reinartz, "Kalman filter based feature analysis for tracking people from airborne images," in ISPRS workshop high-resolution earth imaging for geospatial information, Hannover, Germany, 2011.
- [3] Y. Cao, H. Zhao, N. Li, and H. Wei, "Land-cover classification by airborne LIDAR data fused with aerial optical images," in Multi-Platform/Multi-Sensor Remote Sensing and Mapping (M2RSM), 2011 International Workshop on, 2011, pp. 1-6.
- [4] Y. Chen, W. Su, J. Li, and Z. Sun, "Hierarchical object oriented classification using very high resolution imagery and LIDAR data over urban areas," *Advances in Space Research*, vol. 43, pp. 1101-1110, 2009.
- [5] K. K. Singh, J. B. Vogler, D. A. Shoemaker, and R. K. Meentemeyer, "LiDAR-Landsat data fusion for large-area assessment of urban land cover: Balancing spatial resolution, data volume and mapping accuracy," *ISPRS Journal of Photogrammetry and Remote Sensing*, vol. 74, pp. 110-121, 2012.
- [6] F. J. Mestas-Carrascosa, I. L. Castillejo-González, M. S. de la Orden, and A. G.-F. Porras, "Combining LiDAR intensity with aerial camera data to discriminate agricultural land uses," *Computers and electronics in agriculture*, vol. 84, pp. 36-46, 2012.
- [7] M. Bandyopadhyay, J. A. van Aardt, and K. Cawse-Nicholson, "Classification and extraction of trees and buildings from urban scenes using discrete return LiDAR and aerial color imagery," in *SPIE Defense, Security, and Sensing*, 2013, pp. 873105-873105-9.
- [8] M. M. Awad, "Toward Robust Segmentation Results Based on Fusion Methods for Very High Resolution Optical Image and LiDAR Data," *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing*, 2017.
- [9] A. Zarea and A. Mohammadzadeh, "A novel building and tree detection method from LiDAR data and aerial images," *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing*, vol. 9, pp. 1864-1875, 2016.
- [10] M. Wu, Z. Sun, B. Yang, and S. Yu, "Synergistic Use of WorldView-2 Imagery and Airborne LiDAR Data for Urban Land Cover Classification," in *IOP Conference Series: Earth and Environmental Science*, 2017, p. 012035.
- [11] M. Esfahani and A. Mohammadzade, "Fusion of Pixel-Based and Object-Based Analysis in Order to Buildings and Trees Detection in Urban Area From LiDAR and Optic Data," *Journal of Geomatics Science and Technology*, vol. 6, pp. 27-42, 2016.
- [12] B. Ghimire, J. Rogan, V. R. Galiano, P. Panday, and N. Neeti, "An evaluation of bagging, boosting, and random forests for land-cover classification in Cape Cod, Massachusetts, USA," *GIScience & Remote Sensing*, vol. 49, pp. 623-643, 2012.
- [13] C. Zhang and Z. Xie, "Data fusion and classifier ensemble techniques for vegetation mapping in the coastal Everglades," *Geocarto International*, vol. 29, pp. 228-243, 2014.
- [14] P. O. Gislason, J. A. Benediktsson, and J. R. Sveinsson, "Random forests for land cover classification," *Pattern Recognition Letters*, vol. 27, pp. 294-300, 2006.
- [15] X. Meng, N. Currit, and K. Zhao, "Ground filtering algorithms for airborne LiDAR data: A review of critical issues," *Remote Sensing*, vol. 2, pp. 833-860, 2010.

- [16] P. Axelsson, "DEM generation from laser scanner data using adaptive TIN models," *International Archives of Photogrammetry and Remote Sensing*, vol. 33, pp. 111-118, 2000.
- [17] S. Theodoridis and K. Koutroubas, "Pattern Recognition, Academic Press," New York, 1999.
- [18] M. Janalipour and A. Mohammadzadeh, "Building damage detection using object-based image analysis and ANFIS from high-resolution image (Case study: BAM earthquake, Iran)," *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing*, vol. 9, pp. 1937-1945, 2016.
- [19] M. Janalipour and M. Taleai, "Building change detection after earthquake using multi-criteria decision analysis based on extracted information from high spatial resolution satellite images," *International Journal of Remote Sensing*, vol. 38, pp. 82-99, 2017.
- [20] M. JANALIPOUR, A. MOHAMMADZADEH, Z. M. J. VALADAN, and S. AMIRKHANI, "BUILDINGS'DAMAGE DETERMINATION AFTER THE EARTHQUAKE BY USING ANFIS MODEL AND REMOTE SENSING IMAGERY," 2015.
- [21] F. Albrechtsen, "Statistical texture measures computed from gray level cooccurrence matrices," *Image processing laboratory, department of informatics, university of oslo*, vol. 5, 2008.
- [22] F. Bianconi and A. Fernández, "Evaluation of the effects of Gabor filter parameters on texture classification," *Pattern Recognition*, vol. 40, pp. 3325-3335, 2007.
- [23] D. Zhao, Z. Lin, R. Xiao, and X. Tang, "Linear Laplacian discrimination for feature extraction," in *Computer Vision and Pattern Recognition*, 2007. CVPR'07. IEEE Conference on, 2007, pp. 1-7.
- [24] A. M. Milad Janalipour, Mohammad Javad Valadan zouj., "Application of the Weighted Indexes Using Training Data and Genetic Algorithm on High Resolution Images for Vegetation Detection in Urban Areas," *Journal of Natural Environment*, vol. 67, pp. 135-243, 2014.
- [25] C. Persello, "Advanced techniques for the classification of very high resolution and hyperspectral remote sensing images," *University of Trento*, 2010.
- [26] K. S. Durgesh and B. Lekha, "Data classification using support vector machine," *Journal of Theoretical and Applied Information Technology*, vol. 12, pp. 1-7, 2010.
- [27] V. Vapnik, *The nature of statistical learning theory*: Springer Science & Business Media, 2013.
- [28] M. Chi and L. Bruzzone, "Semisupervised classification of hyperspectral images by SVMs optimized in the primal," *IEEE Transactions on Geoscience and Remote Sensing*, vol. 45, pp. 1870-1880, 2007.
- [29] R. Pino-Mejías, M. D. Cubiles-de-la-Vega, M. Anaya-Romero, A. Pascual-Acosta, A. Jordán-López, and N. Bellinfante-Crocci, "Predicting the potential habitat of oaks with data mining models and the R system," *Environmental Modelling & Software*, vol. 25, pp. 826-836, 2010.
- [30] S. E. Sesnie, B. Finegan, P. E. Gessler, S. Thessler, Z. Ramos Bendana, and A. M. Smith, "The multispectral separability of Costa Rican rainforest types with support vector machines and Random Forest decision trees," *International Journal of Remote Sensing*, vol. 31, pp. 2885-2909, 2010.
- [31] Z.-H. Zhou, *Ensemble methods: foundations and algorithms*: CRC Press, 2012.
- [32] R. G. Pontius Jr and M. Millones, "Death to Kappa: birth of quantity disagreement and allocation disagreement for accuracy assessment," *International Journal of Remote Sensing*, vol. 32, pp. 4407-4429, 2011.



Comprehensive investigation on non-parametric classification methods in order to separate urban objects using the integration of very high spatial resolution LiDAR and aerial data

Zinat Ghavami¹, Hossein Arefi^{*2}, Behnaz Bigdeli³, Milad Janalipour⁴

1- MSc student of photogrammetry in School of Surveying and Geospatial Engineering, College of Engineering, University of Tehran

2- Assistant professor in School of Surveying and Geospatial Engineering, College of Engineering, University of Tehran

3- Assistant professor in College of Civil Engineering, Shahrood University of Technology

4- PhD student of Remote Sensing in Faculty of Geodesy & Geomatics Engineering, K. N. Toosi University of Technology

Abstract

Nowadays, to obtain information covering urban land, the city is one of the most important and widely used management tools in the study of Earth changes. Classification of images is one of the most common methods of extracting information from remote sensing data. Complex and dense urban areas are one of the problems in the analysis of remote sensing. The accuracy of classification performance in these areas is under the attention of the researchers and always tries to improve the accuracy. Using different data and application integration techniques to classify a variety of effects can be more accurate-achieved with higher reliability. Among the successful classification methods in recent years, support vector machines algorithm and ensemble learning algorithms such as Bagging, Boosting and Random Forest noted can be mentioned. In this paper, the performance of the four algorithms to identify the effects of the dense city with a very high resolution aerial LiDAR and image are discussed. The results show that the combination of LiDAR data and aerial image, gives out a better classification of the degree of urban features. The classification of urban features with the help of integrated LiDAR and aerial image information with the use of support vector machine algorithm-precision 93.99% performs higher ability than other classification methods such as Bagging, Boosting and Random Forest.

Key words: LiDAR Data, Areal Image, Support Vector Machines, Bagging, Boosting, Random Forest