

شناسایی ساختمان در مناطق شهری در تصاویر با قدرت تفکیک بالای سنجش از دوری با استفاده از روش آدابوست توسعه یافته و ویژگی های سطح بالا (شبه عمیق)

مینا حمیدی^{۱*}، حمید عبادی^۲، عباس کیانی^۳

^۱ دانشجوی کارشناسی ارشد فتوگرامتری و سنجش از دور - دانشکده مهندسی نقشه برداری - دانشگاه صنعتی خواجه

نصیرالدین طوسی

hamidi.mina94@gmail.com

^۲ استاد دانشکده مهندسی نقشه برداری - دانشگاه صنعتی خواجه نصیرالدین طوسی

ebadi@kntu.ac.ir

^۳ دانشجوی دکتری فتوگرامتری و سنجش از دور - دانشکده مهندسی نقشه برداری - دانشگاه صنعتی خواجه

نصیرالدین طوسی

abbasekiani@yahoo.com

(تاریخ دریافت آذر ۱۳۹۶، تاریخ تصویب فروردین ۱۳۹۷)

چکیده

شناسایی ساختمان از تصاویر سنجش از دور در بروزرسانی نقشه ها، نظارت شهری و طیف وسیعی از کاربردها اهمیت زیادی دارد. تصاویر با قدرت تفکیک مکانی بالا یک منبع داده مهم، برای استخراج اطلاعات مکانی است. این تصاویر امکانات فوق العاده ای برای استخراج عوارض از جمله ساختمان و تجزیه و تحلیل های مکانی در مناطق شهری فراهم کرده اند؛ اما این کار معمولاً به دلیل پیچیدگی ها و ناهمگونی های این داده ها مانند تغییرات درون کلاسی زیاد و تغییرات بین کلاسی کم، با دشواری هایی همراه است. با وجود تلاش های زیادی که برای توسعه روش های اتوماتیک شناسایی ساختمان از این تصاویر طی دهه های گذشته انجام شده است؛ روش های با کارایی بالا به دلیل عدم قطعیت هایی چون انتخاب ویژگی های بهینه هنوز در دسترس نیستند و از سویی به دلیل افزایش قدرت تفکیک داده های مورد استفاده، زمان پردازش نیز بالا می باشد. از این رو، بهبود صحت شناسایی اتوماتیک ساختمان از داده های سنجش از دور و در عین حال زمان پردازش کمتر انگیزه اصلی تحقیق حاضر است. روش پیشنهادی این مقاله، به این صورت است که ابتدا با بکارگیری ساختارهای بافتی شبه عمیق، ویژگی های سطح بالایی را جهت آشکارسازی بهینه ساختمان استخراج می نماید. سپس بر اساس ادغام الگوریتم آدابوست توسعه یافته با روش ماشین بردار پشتیبان بهینه سازی شده با ازدحام ذرات، ویژگی های بهینه را انتخاب کرده و به طبقه بندی باینری عارضه ساختمان و زمینه می پردازد. روش پیشنهادی بر روی مجموعه داده استاندارد واهینگن اجرا و سپس نتایج حاصل از آن با روش های کارآمد یادگیری ماشین مقایسه شده است. همچنین مقایسه ای بین روش مجموعه ویژگی شبه عمیق با روش متداول ویژگی های بافت GLCM صورت گرفته است. نتایج تجربی نشان دادند که به طور میانگین بیشترین صحت کلی و ضریب کاپا حاصل از روش پیشنهادی به ترتیب، ۹۳،۲۵ و ۸۳،۰۶ درصد می باشد و نسبت به روش های مرسوم افزایش دقت ۷،۲۷ درصد در ضریب کاپا دارا می باشد که نشان از اعتبار و توانمندی روش پیشنهادی بوده بعلاوه اینکه زمان محاسبات را حدوداً به نصف کاهش می دهد.

واژگان کلیدی: شناسایی ساختمان، سنجش از دور، انتخاب ویژگی، طبقه بندی

۱- مقدمه

ساختمان‌ها به دلیل فراوانی، تنوع و پیچیدگی، بنیادی‌ترین سازه در مناطق شهری هستند. از این‌رو، شناسایی اتوماتیک ساختمان‌ها در تصاویر هوایی و ماهواره‌ای به یک موضوع مهم برای ایجاد و به‌روزرسانی نقشه‌ها و پایگاه داده سیستم اطلاعات مکانی، به‌منظور عملکرد بهتر در تشخیص تغییرات^۱، تجزیه و تحلیل کاربری زمین^۲ و برنامه‌های نظارت شهری^۳، تبدیل شده است [۱]. تلاش‌های بسیاری به‌منظور سرعت بخشی به فرایند شناسایی و استخراج ساختمان‌ها از تصاویر رقومی توسط الگوریتم‌های نیمه اتوماتیک و اتوماتیک صورت گرفته است، اما این امر به دلایل هندسی ناشی از ساختارهای پیچیده‌ی ساختمان‌ها، دلایل رادیومتریکی، وجود سایه عوارض و ... همچنان یک مسئله چالش‌برانگیز است [۲، ۳].

تصاویر سنجش‌ازدور با قدرت تفکیک مکانی بالا، به‌عنوان یک راه سریع و اقتصادی، امکانات فوق‌العاده‌ای برای استخراج عوارض و تجزیه و تحلیل‌های مکانی در مناطق شهری فراهم کرده‌اند [۴]. اگرچه چنین داده‌هایی برای کاربردهای شهری سنجش‌ازدور بسیار باارزش هستند، اما به دلیل ظهور اشیاء کوچک (همچون دودکش‌ها بر روی بام ساختمان‌ها و ...) که در تصاویر با قدرت تفکیک مکانی پایین پنهان بودند، چالش‌های جدیدی ایجاد کرده‌اند. با توجه به اینکه این تصاویر، دارای واریانس درون کلاسی زیاد و واریانس بین کلاسی کم می‌باشند؛ تفسیر اتوماتیک این تصاویر لزوماً منجر به دقت بالا نمی‌گردد و رسیدن به دقت تفسیر بالا، نیازمند طراحی الگوریتم‌های اتوماتیک به نحوی است که توانایی مقابله با مشکلات ناشی از پیچیدگی صحنه تصویر را دارا باشند [۵].

نخستین رویکردها برای تشخیص اتوماتیک ساختمان‌ها عمدتاً بر استفاده از یک تصویر هوایی یا ماهواره‌ای تکیه داشتند. به‌عنوان مثال، McKeown در [۶]، یک روش تفسیر مبتنی بر دانش^۴ برای شناسایی ساختمان‌ها از تصاویر هوایی ارائه کرد. پس‌از آن، Matsuyama [۷]، یک سیستم خبره^۵ برای استخراج

ساختمان توسعه داد. Peng و همکارانش [۸]، از مدل‌های مارک^۶ برای تشخیص ساختمان از تصاویر هوایی استفاده نمودند. احمدی و همکارانش [۹]، مدل منحنی‌های فعال^۷ را به‌منظور استخراج مرزهای ساختمان از تصاویر هوایی توسعه دادند. با توجه به نحوه ظهور ساختمان‌ها در این داده‌ها می‌توان از ویژگی‌های مفید آن جهت شناسایی و تفکیک عارضه ساختمان از سایر عوارض استفاده کرد؛ اما این رویکردها در معرض مشکلاتی چون ساختمان‌های پیچیده و وجود پوشش گیاهی قرار دارند که عمدتاً به این دلیل است که استفاده از تک تصویر، اطلاعات کافی برای الگوریتم‌ها فراهم نمی‌کند [۱۰].

داده‌های ارتفاعی در قالب مدل رقومی سطح^۸ که از طریق اطلاعات تصاویر پوشش‌دار، تولید شده یا به‌طور مستقیم توسط اسکنر لیزری اخذ می‌شوند نیز در بسیاری از روش‌های تشخیص ساختمان استفاده شده‌اند. Weidner و Förstner در [۱۱]، روشی ساده بر اساس اعمال حد آستانه به DSM نرمال شده برای شناسایی ساختمان‌ها معرفی کردند. Forlani و همکارانش [۱۲]، یک چارچوب مبتنی بر قاعده^۹ برای طبقه‌بندی اتوماتیک داده‌های خام لیدار به ساختمان‌ها، زمین و پوشش گیاهی توسعه دادند. باین‌حال، ساختمان‌ها و برخی دیگر از عوارض روی زمین ممکن است تقریباً هم‌ارتفاع باشند. در این حالت، لازم است برخی از ویژگی‌های تصویری مانند ویژگی‌های طیفی و بافت تصویر برای تفکیک آن‌ها از هم معرفی گردد [۱۳]. اخیراً با دسترسی گسترده به داده‌های ارتفاعی و تصاویر اپتیکی با چند باند طیفی، بکارگیری روش‌های تلفیق داده‌ها برای شناسایی ساختمان، توجه زیادی را به‌سوی خود جلب کرده است. در [۱۴]، Chen و همکارانش جهت شناسایی ساختمان، ابتدا صفحات سه‌بعدی از ابر نقاط استخراج کرده و سپس لبه‌های اولیه ساختمان را از داده‌های لیدار توسط الگوریتم تشخیص لبه‌ی Canny شناسایی نمودند. در ادامه بر اساس لبه‌های تقریبی، به استخراج لبه‌های دقیق ساختمان در فضای تصویر از طریق تبدیل هاف^{۱۰} پرداختند. Khoshelham و همکارانش [۱۵]، به‌منظور استخراج سقف ساختمان‌ها،

^۶ Snake models

^۷ Active contour model

^۸ Digital Surface Model (DSM)

^۹ Rule-based framework

^{۱۰} Hough transform

^۱ Change detection

^۲ Land use analysis

^۳ Urban monitoring

^۴ Knowledge-based

^۵ Expert system

روشی برای انطباق سطوح مسطح به داده‌ی ارتفاعی در داخل نواحی تصویر هوایی قطعه‌بندی شده توسعه دادند. Rottensteiner و همکارانش [۱۶]، برای شناسایی ساختمان، الگوریتم دمپستر-شیفر^۱ را بر پایه‌ی تلفیق تصاویر هوایی و داده‌های لیدار ارزیابی کردند. استفاده از این دو منبع داده به‌عنوان مکمل، ایده مناسبی جهت رفع نواقص هر یک از آن‌ها و استفاده از قابلیت‌های هر دو به‌صورت هم‌زمان می‌باشد.

علاوه بر این، مسئله‌ی شناسایی و استخراج ساختمان می‌تواند به تکنیک‌های سطح پایین^۲ و سطح بالا^۳ تقسیم‌بندی شود [۳]. تکنیک‌های سطح پایین عمدتاً بر اساس تشخیص لبه و استخراج از تصاویر بوده که توسط فرایندهای تعریف قوانین و فرضیه‌ها برای شناسایی ساختمان‌ها دنبال می‌شوند. این روش‌ها مزیت ارائه یک طراحی نسبتاً ساده و هزینه محاسباتی کمی دارند، اما به دلیل محدودیت‌های روش‌شناختی ذاتی آن‌ها، از قابلیت اطمینان برخوردار نیستند. در مقابل، تکنیک‌های سطح بالا تلاش می‌کنند تا فرایند شناخت انسان و مهارت‌های تصمیم‌گیری را که بر اساس تجزیه و تحلیل اطلاعات است، تقلید کنند. اکثر روش‌های سطح بالای تشخیص ساختمان بر اساس طبقه‌بندی تصویر می‌باشد. مسئله طبقه‌بندی، معمولاً با نوع و تعداد ویژگی‌های بکار رفته در آن سروکار دارد [۱۷]. در مطالعات، انواع مختلفی از ویژگی‌ها مانند شاخص تفاضلی پوشش گیاهی نرمال‌شده^۴ به‌منظور تشخیص ساختمان‌ها از درختان و ویژگی‌های ارتفاعی به‌منظور جداسازی عوارض مرتفع از غیرمرتفع تعریف و بررسی شده‌اند و تحقیقات عمدتاً در ایجاد فضاهای ویژگی متمرکز شده‌اند [۱۸-۲۰]. در [۱۸] Matikainen و همکارانش به‌منظور شناسایی ساختمان از تصاویر هوایی و داده لیدار، درخت طبقه‌بندی با معیار گینی را بکار بردند. در این روش، ابتدا DSM حاصل از آخرین بازگشت داده لیدار، به دو بخش زمین و عوارض مرتفع تقسیم‌بندی شد. سپس ساختمان‌ها و درختان، با استفاده از ترکیبات مختلف ویژگی‌های استخراجی از DSM حاصل از اولین و آخرین بازگشت لیدار و تصویر هوایی، از یکدیگر تفکیک

شدند. در ادامه تحقیق فوق، Salah و همکارانش [۱۹]، عملکرد درختان طبقه‌بندی مختلف را برای شناسایی ساختمان از تلفیق تصاویر هوایی و داده لیدار، مورد بررسی قرار دادند. آن‌ها در تحقیق خود، به تولید ویژگی‌های قابل استخراج از سطح داده‌های موردنظر شامل ویژگی‌های طیفی، بافتی به‌دست‌آمده از داده‌های تصویری، شدت داده‌های لیدار و داده ارتفاعی نرمال شده، اقدام نمودند. درروش ارائه‌شده توسط Haiyan Guan و همکارانش [۲۰]، جهت جداسازی نقاط ساختمانی از گیاهان مرتفع، به روش طبقه‌بندی بیشترین شباهت^۵ و بهره‌گیری از داده‌های آموزشی تهیه‌شده از سطح داده‌های طبقه‌بندی‌شده لیدار اقدام صورت گرفت. نهایتاً به‌منظور بهبود نتایج حاصل از طبقه‌بندی نهایی، از قوانین قابل استخراج از سطح داده‌های لیدار بهره گرفته و به‌صورت پس پردازش، نتایج طبقه‌بندی بهبود داده شدند.

با وجود این تحقیقات، برخی از ویژگی‌های متنی^۶ موجود در داده‌ها توسط آن‌ها مدل‌سازی نشده‌اند درحالی‌که لازم است الگوریتم‌های طبقه‌بندی تا حد امکان از اطلاعات سطح بالایی استفاده نمایند. از طرفی، تجزیه و تحلیل یک مجموعه داده با تعداد زیادی از ویژگی‌ها، از نظر محاسباتی گران است و می‌تواند عملکرد الگوریتم طبقه‌بندی را کاهش دهد. در این مورد، کاهش ابعاد، یک گام اساسی است که می‌تواند با استراتژی انتخاب ویژگی حاصل شود [۲۱].

الگوریتم‌ها و مدل‌های بسیاری در تحقیقات یادگیری ماشین به‌منظور طبقه‌بندی در دسترس است. باین‌حال، تعداد اندکی از آن‌ها برای طبقه‌بندی داده‌های با حجم بالا مناسب هستند. روش آدا بوست^۷ [۲۲]، یکی از روش‌های مطرح یادگیری ماشین است که علاوه بر قابلیت انتخاب ویژگی حین فرایند یادگیری، می‌تواند در برخورد با داده‌های با حجم بالا کارآمد واقع شود. باین‌وجود، این روش ممکن است در معرض مشکل بیش‌برازش^۸ قرار گیرد. این مشکل در صورتی رخ می‌دهد که مدل یادگیرنده به‌صورت سخت‌گیرانه به طبقه‌بندی درست تمام نمونه‌های آموزشی بپردازد و حال‌آنکه ممکن است

^۵ Maximum Likelihood

^۶ Contextual features

^۷ Adaptive Boosting (AdaBoost)

^۸ Overfitting

^۱ Dempster-Shafer

^۲ Low-level vision techniques

^۳ High-level vision techniques

^۴ Normalized difference vegetation index (NDVI)

داده‌های آموزشی حاوی داده‌های نویزی و پرت^۱ باشند. در نتیجه مدل حاصل از این نوع آموزش، منجر به افزایش خطای تعمیم و کاهش دقت طبقه‌بندی می‌گردد [۲۳].

الگوریتم‌های تکاملی از جمله الگوریتم ژنتیک و بهینه‌سازی ازدحام ذرات^۲ از روش‌های مطلوب در تحقیقات انتخاب زیرمجموعه ویژگی بهینه می‌باشند. در برخی از مطالعات این روش‌ها به صورت توأمان با طبقه‌بندی کننده‌ی قوی و مطرح ماشین بردار پشتیبان^۳ بکار رفته‌اند [۲۴] و در عین حال که دقت خوبی حاصل می‌کنند اما از نظر محاسباتی پرهزینه می‌باشند. از این رو، عملاً برای تصاویر بزرگ مقیاس نمی‌توان ویژگی‌های مختلف و زیاد را با استفاده از آن‌ها اجرا نمود. با توجه به اینکه انگیزه اصلی این پژوهش، بهبود صحت شناسایی اتوماتیک ساختمان از داده‌های سنجش از دور و در عین حال زمان پردازش کمتر می‌باشد، بکارگیری یک رویکرد ترکیبی^۴، برای بهره‌گیری از قابلیت‌های هر الگوریتم، ایده مناسبی به نظر می‌رسد.

بنابراین در پژوهش حاضر، با هدف بهبود شناسایی اتوماتیک ساختمان از داده‌های سنجش از دور، یک رویکرد ترکیبی جدید برای انتخاب ویژگی‌های بهینه از مجموعه داده بزرگ^۵ در یک زمان معقول ارائه شده است. روش پیشنهادی، ابتدا با بکارگیری ساختارهای بافتی شبه عمیق که روشی میانی بین بانک‌های فیلتر متداول [۲۵، ۲۶] و یادگیری عمیق^۶ است، ویژگی‌های سطح بالایی را از ترکیب تصویر چندطیفی و داده ارتفاعی استخراج می‌نماید. سپس بر اساس ادغام الگوریتم آدابوست توسعه یافته با روش ماشین بردار پشتیبان بهینه‌سازی شده با ازدحام ذرات (CB-SVMps0)، ویژگی‌های بهینه را انتخاب کرده و به طبقه‌بندی باینری تصویر در دو کلاس ساختمان و زمینه می‌پردازد.

همچنین در این تحقیق، در رویکردی دیگر به مقایسه و تحلیل الگوریتم‌های بکار رفته در روش پیشنهادی با روش‌های متداول همچون جنگل‌های تصادفی^۷ [۲۷]، پرداخته شده است. علاوه بر این، مقایسه‌ای بین روش

مجموعه ویژگی شبه عمیق با روش متداول ویژگی‌های بافت ماتریس رخداد توأمان سطح خاکستری^۸ صورت گرفته است. نتایج بیانگر این است که ضمن افزایش دقت از حیث سرعت نیز روند پیشنهادی، نتایج مطلوب‌تری را نسبت به روش‌های دیگر کسب نموده است.

ساختار کلی مقاله بدین شکل است که در ادامه در بخش مبانی نظری تحقیق، به معرفی و بیان تئوری روش‌های بکار گرفته شده، پرداخته شده است. سپس در بخش ۳، جزئیات مراحل پیاده‌سازی و نتایج حاصل از آن‌ها ارائه و در ادامه‌ی آن به مقایسه و تحلیل نتایج پرداخته شده است. در نهایت، نتیجه‌گیری و پیشنهادات، ارائه شده است.

۲- مبانی نظری تحقیق

یکی از بهترین روش‌ها برای شناسایی عوارض شهری مانند عارضه ساختمان از تصاویر سنجش از دور، طبقه‌بندی تصویر می‌باشد. طبقه‌بندی، به جداسازی ویژگی‌های مشابه و تقسیم‌بندی طبقاتی آن‌ها که دارای رفتار یکسانی باشند، اطلاق می‌گردد. روش‌های مبتنی بر طبقه‌بندی معمولاً از ویژگی‌های طیفی، بافت و غیره، به منظور شناسایی عارضه ساختمان، بهره می‌گیرند. دقت طبقه‌بندی ساختمان، به دلیل طبقه‌بندی اشتباه بین پیکسل‌های ساختمان و سایر عوارض با طیف مشابه آن در تصویر، ممکن است رضایت‌بخش نباشد؛ بنابراین، دقت روش‌های طبقه‌بندی، به شدت به نوع ویژگی‌های بکار رفته در آن متکی است. بر این اساس، مقاله حاضر بر روی فرایندهای تولید ویژگی، انتخاب ویژگی و طبقه‌بندی متمرکز شده است که در این بخش، به معرفی و بیان جزئیات روش‌های بکار رفته در هر مرحله، پرداخته شده است.

۲-۱- تولید ویژگی

منظور از تولید ویژگی، فرایند ساخت ویژگی‌های جدید از یک یا چند ویژگی خام اولیه است. تولید ویژگی با هدف بهبود دقت طبقه‌بندی، ابعاد فضای ویژگی اولیه را افزایش می‌دهد. سپس ویژگی‌های بهینه از این فضا در مرحله انتخاب ویژگی، انتخاب می‌گردند. ویژگی‌های

^۱ Outlier

^۲ Particle Swarm Optimization (PSO)

^۳ Support Vector Machine (SVM)

^۴ Hybrid approach

^۵ Big Data

^۶ Deep Learning

^۷ Random Forests (RF)

^۸ Gray Level Co-occurrence Matrix (GLCM)

متعددی از جمله بافت تصویر، شاخص‌های طیفی و غیره می‌توان از تصاویر با بکارگیری روش‌های مختلف استخراج کرد. ماتریس GLCM [۲۸]، از جمله روش‌های کلاسیک تولید ویژگی‌های بافت می‌باشد که ارتباط بین دو پیکسل همسایه را در یک فاصله و جهت مشخص در نظر می‌گیرد. این ماتریس، برای یک فاصله معین، معمولاً در ۴ جهت (۰، ۴۵، ۹۰ و ۱۳۵ درجه) تولید می‌گردد. سپس ویژگی‌های آماری از جمله میانگین، واریانس، کنتراست و غیره را می‌توان از آن استخراج کرد.

بانک‌های فیلتر همچون فیلترهای لاپلاسین گوسین^۱ نیز از جمله روش‌های تولید ویژگی بافت هستند که نسبتاً ساده بوده و می‌توانند نتایج مناسبی را حاصل کنند. با این حال، هر بانک فیلتر برای نوع خاصی از بافت طراحی شده است. بنابراین، یک رویکرد متفاوت نیاز است که قادر باشد به‌طور خودکار با انواع مختلف بافت سازگار گردد. یادگیری عمیق روشی برای حل این مشکل است که در آن استخراج‌گرهای ویژگی به‌طور مستقیم آموزش داده می‌شوند. اما این روش نسبت به تنظیمات شکننده است و اغلب نتایج آن در مقایسه با بانک‌های فیلتر ساده بهبود نمی‌یابد. از این رو، روشی تحت عنوان شبه عمیق ارائه گردید [۲۹]؛ که مشابه با روش‌های یادگیری عمیق، ویژگی‌ها را در مقیاس‌های مختلف بر اساس پچ^۲ استخراج می‌کند و نیز مزیت سادگی بانک‌های فیلتر متداول را دارا می‌باشد.

مجموعه ویژگی‌های شبه عمیق، شامل ویژگی‌های وسیعی هستند که از طریق یک پنجره متحرک^۳ به‌عنوان واحد پایه که پیکسل به پیکسل بر روی تصویر جابجا می‌شود تولید می‌گردند. ابعاد پنجره متحرک را می‌توان از 3×3 تا 25×25 در نظر گرفت که به قدرت تفکیک تصویر و اندازه اشیاء تصویری بستگی دارد. این مجموعه ویژگی قادر است با تصاویر با شرایط نوری مختلف، انواع کلاس‌های شیء و ساختارهای مختلف صحنه هماهنگ گردد. مجموعه ویژگی شبه عمیق، به‌طور طبیعی قادر است طیف وسیعی از مقیاس‌ها و فرکانس‌های بافت را پوشش دهد. همچنین این روش، از مزایای اطلاعات رنگ از طریق اختلاف یا نسبت بین باندهای طیفی بهره می‌برد. به‌طور کلی، این روش قادر است اطلاعات زیادی از محتوای

تصاویر با قدرت تفکیک مکانی بالا را در اختیار قرار دهد؛ از این رو، ویژگی‌های استخراجی توسط آن‌ها، ویژگی‌های سطح بالا نامیده می‌شوند. بکارگیری چنین ویژگی‌هایی به سیستم‌های رایانه‌ای قابلیت رقابت با درک بصری انسان را می‌دهد. در ادامه به بیان جزئیات نحوه استخراج ویژگی‌های شبه عمیق مورد استفاده در این تحقیق، پرداخته شده است.

حوزه‌ی نخست (الگوی تصادفی متقارن)، در طی یک روند تکراری، پچ‌های تصادفی داخل پنجره متحرک به‌عنوان الگوی بافت در تصویر تعیین می‌گردند، به این نحو که در هر تکرار یک پچ با اندازه و موقعیت تصادفی داخل پنجره متحرک انتخاب‌شده و سپس قرینه آن نسبت به پیکسل مرکزی پنجره تعیین می‌گردد؛ درنهایت، اختلاف میانگین جفت پچ‌ها محاسبه گردیده و به پیکسل مرکزی اختصاص داده می‌شود. قابل ذکر است که نمونه‌برداری تصادفی پچ در هر تکرار، یک‌بار صورت می‌گیرد یعنی با تغییر موقعیت پنجره متحرک در تصویر، جفت پچ‌های انتخابی در آن تکرار تغییر نمی‌کنند. علاوه بر این، پچ‌ها می‌توانند با یکدیگر هم در همان مرحله و هم در مراحل متوالی همپوشانی داشته باشند. ایجاد پچ‌های تصادفی این امکان را می‌دهد که بتوان طیف وسیعی از بافت‌ها را با کاهش زمان اجرا و حافظه مورد نیاز استخراج کرد.

در حوزه‌ی دوم (سطوح مقیاس)، ویژگی‌ها از طریق محاسبه‌ی مقادیر میانگین پچ‌های مربع هم‌مرکز با پیکسل مرکزی پنجره متحرک، با ابعاد مختلف (از 3×3 تا ابعاد پنجره متحرک) استخراج می‌گردد. در حوزه‌ی سوم (الگوی معین بین‌باندی)، ویژگی‌ها بر اساس اختلاف مقادیر میانگین پچ‌های مربع هم‌اندازه و هم‌مرکز با مرکز پنجره متحرک بین باندهای طیفی مختلف استخراج‌شده و در حوزه‌ی بافتی چهارم (الگوی معین مقیاسی)، بر اساس اختلاف بین مقادیر میانگین پچ‌های مربع با اندازه‌های مختلف در باند طیفی مشابه، تولید می‌گردند. ابعاد فضای ویژگی در این سه حوزه، به تعداد باندهای تصویر و تعداد پچ‌های مربع بستگی دارد.

موارد فوق شامل ویژگی‌های خطی هستند. درحالی‌که در حوزه‌ی آخر، ویژگی‌های غیرخطی تولید می‌شود. بدین نحو که اختلاف مقادیر میانگین پچ‌های مربع بین دو باند طیفی مختلف، تقسیم بر جمع آن‌ها می‌گردد. این پچ‌های مربع (از 3×3 تا ابعاد پنجره متحرک) با مرکز پنجره

^۱ Laplacian of Gaussian (LoG)

^۲ Patch-Based

^۳ Sliding window

متحرک هم‌مرکز و هردو جفت پیچ (در دو باند طیفی متفاوت) هم‌اندازه هستند. این ویژگی‌ها با در نظر گرفتن نسبت بین باندهای طیفی شبیه به شاخص‌های طیفی می‌باشند اما از آنجایی که این کار را هم‌زمان با بکارگیری اطلاعات مجاورت در هر باند طیفی انجام می‌دهند، از بهره‌وری بالاتری برخوردارند.

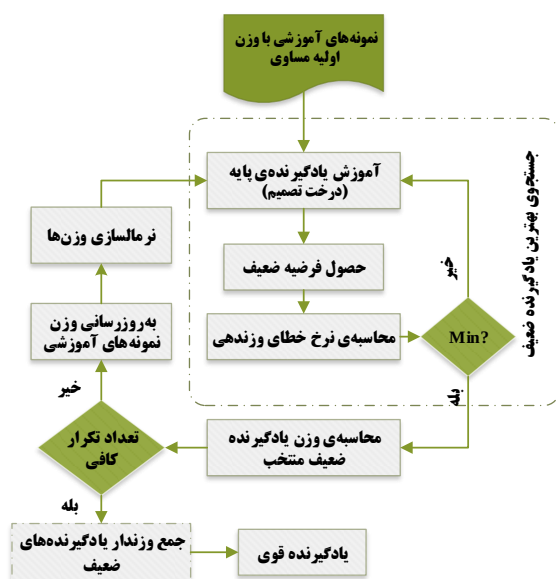
۲-۲- انتخاب ویژگی و طبقه‌بندی

یکی از مراحل مهم در پیاده‌سازی موفق طبقه‌بندی تصویر، انتخاب ویژگی‌های مناسب می‌باشد. در بحث طبقه‌بندی، می‌توان از ویژگی‌هایی نظیر اثر طیفی، اطلاعات بافت و غیره استفاده کرد. بکارگیری ویژگی‌های متنوع و متعدد در امر طبقه‌بندی، ممکن است به دلیل همبستگی بین آن‌ها، منجر به کاهش دقت طبقه‌بندی گردد؛ بنابراین، انتخاب ویژگی‌های مفید و مناسب برای تفکیک کلاس‌ها از یکدیگر، از اهمیت بالایی برخوردار است. روش آدابوست، از جمله روش‌های یادگیری ماشین نظارت‌شده است که هم‌زمان با طبقه‌بندی قادر به انتخاب ویژگی‌های مناسب می‌باشد. کارآمدی این روش طی سالیان در تحقیقات متنوع و حوزه‌های متفاوت به اثبات رسیده است [۳۰-۳۳]. همچنین نسخه‌های متفاوتی از این روش به منظور افزایش عملکرد آن ارائه شده است [۳۴-۳۶]. جدیدترین نسخه‌ی توسعه‌یافته‌ی آن، آدابوست مبتنی بر اطمینان^۱ [۳۷]، می‌باشد که در ادامه به معرفی آن پرداخته شده است.

۲-۲-۱- الگوریتم آدابوست توسعه‌یافته

ایده اصلی بوستینگ^۲ [۳۸]، این است که یک طبقه‌بندی کننده ساده و کارآمد به الگوریتم یادگیری معرفی گردد، به طوری که دقت آن بر روی نمونه‌های آموزشی، کمی بیشتر از حالت تصادفی باشد. چنین طبقه‌بندی کننده‌هایی، طبقه‌بندی کننده ضعیف نامیده می‌شوند که حتی می‌توانند به سادگی یک حد آستانه تصمیم‌گیری یک‌بعدی باشند. بوستینگ، با هدف بهبود دقت طبقه‌بندی، گروهی از طبقه‌بندی کننده‌های ضعیف را به یک طبقه‌بندی کننده قوی (طبقه‌بندی کننده با دقت

بالا) تبدیل می‌کند. از میان انواع روش بوستینگ، روش آدابوست (بوستینگ تطبیقی) محبوب‌ترین آن‌ها است. آدابوست، یک روش طبقه‌بندی باینری نظارت‌شده با مجموعه آموزشی $\{(x_i, y_i), i=1, \dots, N\}$ است که نمونه‌ها با کلاس مثبت یعنی کلاس هدف ($y_i = +1$) را از نمونه‌ها با کلاس منفی یعنی کلاس زمینه ($y_i = -1$) تفکیک می‌کند. مراحل کلی الگوریتم آدابوست در شکل (۱) نشان داده شده است.



شکل ۱- روند نمای کلی الگوریتم آدابوست

در ابتدای آموزش، احتمال انتخاب هر یک از نمونه‌ها یکسان است و تمام نمونه‌ها وزن اولیه‌ی برابر $1/N$ می‌گیرند. سپس در یک روند تکراری، طبقه‌بندی کننده‌های ضعیف انتخاب می‌شوند؛ در هر تکرار الگوریتم t یادگیرنده پایه، با توجه به توزیع وزن نمونه‌ها w_t آموزش می‌بیند. هدف الگوریتم آموزشی، یافتن بهترین یادگیرنده ضعیف h_t با کمترین نرخ خطای وزندهی بر روی نمونه‌های آموزشی در هر تکرار می‌باشد. نرخ خطای وزندهی ϵ_t جمع وزن نمونه‌هایی است که توسط فرضیه ضعیف h_t اشتباه طبقه‌بندی شده‌اند. به منظور تعیین درجه اهمیت و اثربخشی هر کدام از یادگیرنده‌های ضعیف در طبقه‌بندی کننده نهایی، در هر تکرار الگوریتم t پس از محاسبه نرخ خطای وزندهی ϵ_t ، وزن یادگیرنده ضعیف α_t طبق رابطه (۱) محاسبه می‌گردد که به آن پارامتر تطبیقی نیز می‌گویند.

^۱ Confidence Based AdaBoost (CB-AdaBoost)

^۲ Boosting

^۳ Weak Hypothesis

این اساس، نسخه جدیدی از این الگوریتم با نام CB-AdaBoost ارائه شده است که این موضوع در آن، در نظر گرفته شده است. در الگوریتم CB-AdaBoost، تابع ضرر^۱ الگوریتم آدابوست بر مبنای مقدار اطمینان از برچسب نمونه‌های آموزشی، اصلاح شده است. تابع ضرر اصلاح شده، سبب شده است که وزن نمونه‌ها بر اساس یک قاعده جدید به‌روزرسانی گردند.

الگوریتم CB-AdaBoost، فرض می‌کند که برچسب واقعی نمونه‌های آموزشی z مجهول است و آنچه به‌عنوان برچسب نمونه‌های آموزشی در مجموعه آموزشی به آن‌ها اختصاص داده شده است را برچسب مشاهده شده y و برای هر نمونه آموزشی، احتمال درستی برچسب مشاهده شده y را به‌عنوان یک پارامتر معلوم در نظر می‌گیرد. احتمال درستی برچسب مشاهده شده به‌صورت رابطه (۴) تعریف می‌گردد که مقدار آن عددی بین صفر و یک است.

$$\gamma = p(z = y|x) \quad (4)$$

$$1 - \gamma = p(z = -y|x) \quad (5)$$

که در رابطه (۵)، $1-\gamma$ بیانگر احتمال نادرستی برچسب y و احتمال درستی برچسب $-y$ می‌باشد. مقدار $|\gamma - (1-\gamma)|$ بیانگر قابلیت اطمینان و علامت آن بیانگر اطمینان در جهت درستی یا نادرستی می‌باشد؛ بنابراین از $\text{sign}(2\gamma - 1)y$ به‌عنوان درستی برچسب با سطح اطمینان $|2\gamma - 1|$ می‌توان استفاده کرد. این موضوع از قانون بیز الهام گرفته شده است. اگر مقدار γ یک باشد به معنی اطمینان ۱۰۰ درصد از درستی برچسب y و مقدار صفر به معنی اطمینان ۱۰۰ درصد از نادرستی y می‌باشد. مقدار ۰٫۵ حالت فازی با اطمینان صفر است. با توجه به فرض مذکور، الگوریتم CB-AdaBoost تابع ضرر الگوریتم آدابوست را به‌صورت رابطه (۶) اصلاح می‌کند.

$$\hat{L} = \frac{1}{N} \sum_{i=1}^N [\gamma_i e^{-y_i f(x_i)} + (1 - \gamma_i) e^{y_i f(x_i)}] \quad (6)$$

که در آن، $f(x_i)$ ترکیب خطی یادگیرنده‌های ضعیف و N تعداد کل نمونه‌ها می‌باشد. بر اساس تابع ضرر اصلاح شده، رابطه (۱) نیز به‌فرم رابطه (۷) اصلاح خواهد شد.

$$\alpha_t = \frac{1}{2} \ln \left(\frac{1 - \varepsilon_t}{\varepsilon_t} \right) \quad (1)$$

هر چه یادگیرنده ضعیف بهتر باشد، وزن بیشتری در طبقه‌بندی کننده‌ی نهایی خواهد داشت. پس از محاسبه پارامتر تطبیقی، وزن هر نمونه آموزشی مطابق با رابطه (۲) به‌روزرسانی می‌گردد که مقدار آن برای نمونه‌های آموزشی که اشتباه طبقه‌بندی شده‌اند بزرگ‌تر از یک و برای نمونه‌های آموزشی که به‌درستی طبقه‌بندی شده‌اند کوچک‌تر از یک می‌باشد. این مکانیسم، طبقه‌بندی درست نمونه‌های سخت را در تکرارهای بعدی تضمین می‌کند. قبل از شروع تکرار بعدی وزن نمونه‌ها بین صفر و یک نرمالایز می‌گردند. سپس در تکرار بعدی، یادگیرنده ضعیفی که بر روی توزیع وزن به‌روزرسانی شده، بهتر عمل کند، انتخاب می‌گردد. این روند تا زمانی که تعداد تکرار دلخواه حاصل شود، ادامه می‌یابد.

$$w_i^{t+1} = w_i^t \cdot e^{-\alpha_t y_i h_t(x_i)} \quad (2)$$

درنهایت، طبقه‌بندی کننده‌ی نهایی یا به عبارتی یادگیرنده‌ی قوی H ، از ترکیب خطی یادگیرنده‌های ضعیف، طبق رابطه (۳) ایجاد می‌شود که دقت را تا سطح دلخواه افزایش می‌دهد. نتیجه طبقه‌بندی، علامت H است و مقدار H بیانگر اطمینان از نتیجه طبقه‌بندی می‌باشد.

$$H(x) = \sum_{t=1}^T \alpha_t h_t(x) \quad (3)$$

آدابوست، هم‌زمان با طبقه‌بندی، قابلیت انتخاب ویژگی‌های مناسب را از بین ویژگی‌های متعدد و اضافی دارد که این کار را بر اساس یادگیرنده‌های ضعیف انجام می‌دهد؛ بنابراین یادگیرنده‌ی پایه‌ی مورد استفاده در آدابوست، علاوه بر سادگی، باید قابلیت انتخاب ویژگی داشته باشد. از آنجایی که الگوریتم درخت تصمیم‌گیری کم‌عمق، هر دو خصوصیت سادگی و قابلیت انتخاب ویژگی را دارا است، متداول‌ترین انتخاب برای یادگیرنده پایه در آدابوست می‌باشد.

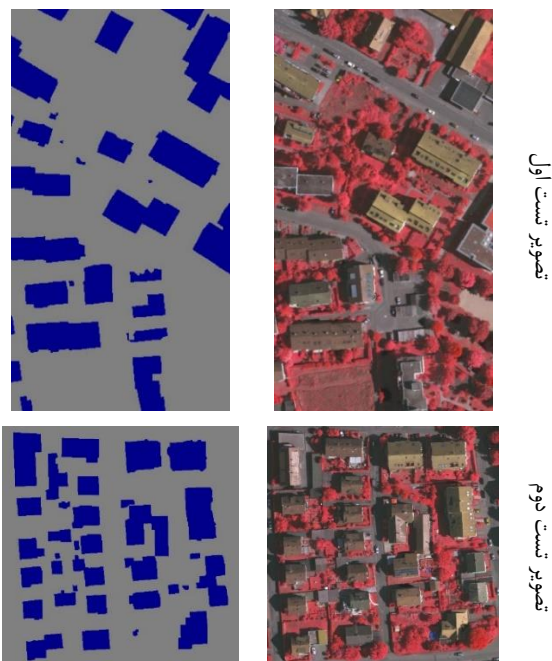
همان‌طور که ذکر شد، آدابوست، وزن اولیه نمونه‌ها را در شروع آموزش، برابر در نظر می‌گیرد؛ اما در حقیقت، به دلیل همپوشانی توزیع کلاس‌ها و نیز وجود نویز در داده‌ها، نمونه‌های آموزشی، وزن برابر با یکدیگر ندارند. بر

^۱ Loss Function

۳- پیاده‌سازی و ارزیابی نتایج

۳-۱- داده‌های مورد استفاده و منطقه مطالعاتی

داده‌های مورد استفاده در این تحقیق، مربوط به منطقه واهینگن از کشور آلمان می‌باشد که توسط کمیسیون III گروه کاری ۴ از جامعه بین‌المللی فتوگرامتری و سنجش از دور به منظور انجام تحقیقات در زمینه استخراج عارضه، آماده‌سازی گردیده است [۴۰]. تصاویر هوایی موجود از منطقه، توسط دوربین دیجیتالی Intergraph/ZI DMC اخذ شده‌اند که دارای قدرت تفکیک مکانی ۸ سانتی‌متر و دارای سه باند طیفی مادون قرمز IR، قرمز R و سبز G می‌باشند. همچنین، از روی تصاویر هوایی اصلی، مدل رقومی سطح و اورتوفتوموزائیک واقعی نیز تهیه شده است. در تهیه این دو نوع داده، از یک شبکه با قدرت تفکیک زمینی ۹ سانتی‌متر، استفاده شده است. همچنین، این دو مجموعه داده، به منظور استفاده یکجا از آن‌ها، هم مرجع شده‌اند. داده‌ی ارتفاعی بکار رفته در این تحقیق، مدل رقومی سطح نرمالایز شده (nDSM) است که توسط همین گروه، تهیه شده است. داده‌ی واقعیت زمینی که توسط همین گروه به صورت دستی برچسب‌گذاری معنایی شده، برای ارزیابی عملکرد استفاده می‌شود. در شکل (۲) تصاویر هوایی محدوده مطالعاتی و تصویر واقعیت زمینی آن‌ها، نشان داده شده است. به منظور تولید ویژگی‌های کارآمد، nDSM منطقه به عنوان یک باند تصویر به تصویر هوایی اورتو با سه باند طیفی، اضافه شده است.



شکل ۲- منطقه مطالعاتی، واهینگن آلمان
الف) تصویر رنگی کاذب
ب) داده واقعیت زمینی

$$\alpha_t = \frac{1}{2} \ln \left(\frac{\sum_{i:h_t(x_i)=y_i} w_{1i}^t + \sum_{i:h_t(x_i) \neq y_i} w_{2i}^t}{\sum_{i:h_t(x_i) \neq y_i} w_{1i}^t + \sum_{i:h_t(x_i)=y_i} w_{2i}^t} \right) \quad (7)$$

که در آن، $w_{1i} = \gamma_i$ و $w_{2i} = 1 - \gamma_i$ می‌باشد. در نهایت، وزن نمونه‌ها مطابق با روابط (۸) و (۹) به روزرسانی شده و بر اساس رابطه (۱۰) توزیع وزن نمونه‌ها D_i ، برای آموزش الگوریتم محاسبه می‌گردد.

$$w_{1i}^{t+1} = w_{1i}^t \cdot e^{-\alpha_t y_i h_t(x_i)} \quad (8)$$

$$w_{2i}^{t+1} = w_{2i}^t \cdot e^{\alpha_t y_i h_t(x_i)} \quad (9)$$

$$D_i^{t+1} = \frac{|w_{1i}^{t+1} - w_{2i}^{t+1}|}{\sum_{i=1}^N |w_{1i}^{t+1} - w_{2i}^{t+1}|} \quad (10)$$

به طور کلی، الگوریتم CB-AdaBoost نسبت به آدا بوست با اطلاعات اولیه‌ی بیشتری آموزش داده می‌شود.

۳-۲- ارزیابی دقت طبقه‌بندی

ارزیابی دقت طبقه‌بندی، معمولاً از طریق ماتریس ابهام که بیانگر رابطه بین داده‌ی مرجع و داده‌های برچسب‌گذاری شده توسط طبقه‌بندی کننده است، انجام گرفته و شاخص‌هایی نظیر صحت کلی^۱، ضریب کاپا^۲ و غیره را می‌توان محاسبه نمود [۳۹]. معیار صحت کلی طبق رابطه (۱۱)، نشان‌دهنده درصد نمونه‌هایی است که طبقه‌بندی آن‌ها به طور صحیح انجام شده است. این معیار یک برآورد خوش‌بینانه بوده و همیشه دقت را بالاتر از مقدار واقعی محاسبه می‌کند. به منظور ارزیابی دقت با استفاده از تمام اطلاعات ماتریس ابهام، می‌توان از ضریب کاپا استفاده کرد. ضریب کاپا که از طریق رابطه (۱۲) محاسبه می‌گردد؛ دقت طبقه‌بندی را نسبت به یک طبقه‌بندی کاملاً تصادفی حساب می‌کند. این معیار، برآوردی بدبینانه بوده و دقت را کمتر از مقدار واقعی بیان می‌کند.

$$O.A = \frac{\sum_{i=1}^n x_{ii}}{N} \times 100 \quad (11)$$

$$K.C = \frac{N \sum_{i=1}^n x_{ii} - \sum_{i=1}^n x_{i+} x_{+i}}{N^2 - \sum_{i=1}^n x_{i+} x_{+i}} \times 100 \quad (12)$$

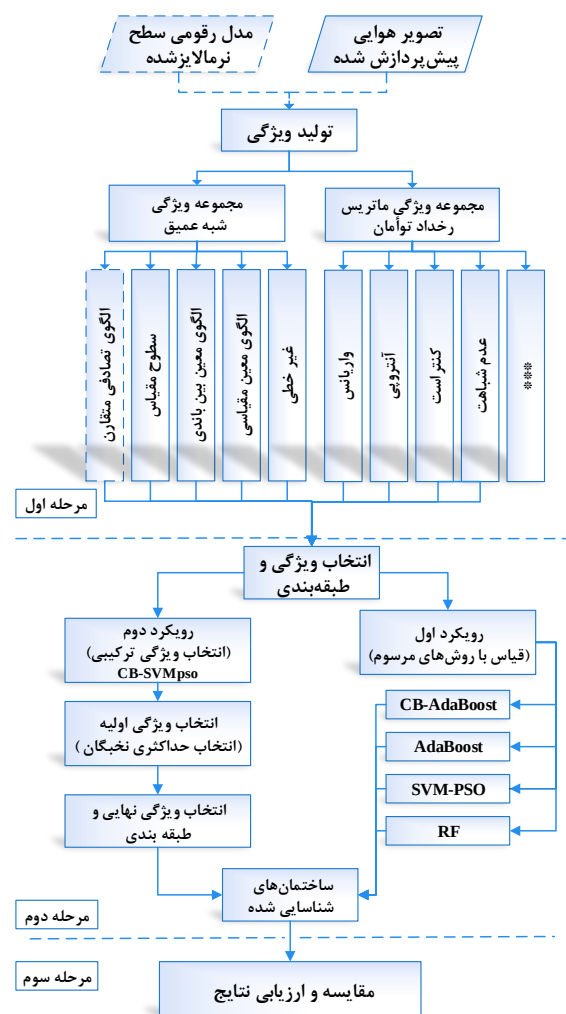
که در این دو رابطه، N تعداد کل نمونه‌ها، n تعداد کلاس‌ها و x_{ii} تعداد نمونه‌های درست طبقه‌بندی شده در هر کلاس، x_{i+} و x_{+i} به ترتیب، مجموع کل سطر i و ستون i است.

^۱ Overall Accuracy (O.A)

^۲ Kappa Coefficient (K.C)

۳-۲- پیاده‌سازی روش مورد استفاده در این تحقیق

طرح کلی روش، به این صورت است که ویژگی‌های متعددی از تصویر هوایی با سه باند طیفی و nDSM منطقه مطالعاتی، با استفاده از روش مجموعه ویژگی شبه عمیق و روش ماتریس GLCM تولید می‌گردد. سپس، نتایج حاصل از تولید ویژگی، وارد مرحله انتخاب ویژگی می‌شود که هدف آن کاهش ابعاد فضای ویژگی و طبقه‌بندی تصویر در دو کلاس ساختمان و غیر ساختمان، به نحوی است که این دو کلاس به بهترین نحو از هم تفکیک گردند. درنهایت، به مقایسه و ارزیابی نتایج حاصل از روش‌های پیاده‌سازی شده پرداخته می‌شود. روند نمای کلی روش پیشنهادی در شکل (۳) نشان داده شده است. جزئیات مراحل پیاده‌سازی در ادامه تشریح شده است.



شکل ۳- روند نمای کلی الگوریتم پیشنهادی

به‌منظور تعیین بهترین ابعاد پنجره متحرک، مقادیر مختلفی از 7×7 تا 25×25 امتحان گردید و ابعاد 15×15 به‌عنوان اندازه مناسب پنجره متحرک، تعیین شد. با توجه به اینکه در زیرگروه الگوی تصادفی متقارن، ویژگی‌ها بر اساس پیچ‌های تصادفی، طی روند تکراری تولید می‌گردند، تعداد ویژگی‌های تولیدشده به تعداد تکرار انتخابی و تعداد باندهای تصویر بستگی دارد. در این تحقیق، برای تصویر تست اول و تصویر تست دوم، به ترتیب تعداد تکرار ۱۵ و ۱۰ در نظر گرفته شد که ۶۰ و ۴۰ ویژگی از این طریق تولید شد. تعداد ویژگی‌های تولیدشده توسط زیرگروه‌های سطوح مقیاس، الگوی معین بین‌باندی و الگوی معین مقیاسی، به تعداد باندهای تصویر و تعداد پیچ‌های مربع (از ابعاد 3×3 تا ابعاد پنجره متحرک)، بستگی دارد. تعداد کل ویژگی‌های تولیدشده با استفاده از این سه زیرمجموعه، برای هر دو تصویر تست اول و تصویر تست دوم، ۲۸۰ ویژگی است. تعداد ویژگی‌های تولیدشده برای زیرگروه آخر نیز که شامل ویژگی‌های غیرخطی است به تعداد باندهای تصویر و تعداد پیچ‌های مربع، بستگی دارد که ۸۴ ویژگی از این طریق تولید شد. درمجموع با توجه به ابعاد 15×15 پنجره متحرک، برای تصویر تست اول و تصویر تست دوم، به ترتیب، تعداد ۴۲۴ و ۴۰۴ ویژگی، از طریق روش شبه عمیق تولید شده است که با تلفیق آن‌ها با باندهای طیفی و داده ارتفاعی به ترتیب ۴۲۸ و ۴۰۸ ویژگی حاصل می‌شود. به‌منظور استخراج ویژگی‌های بافتی از ماتریس GLCM، ۸ ویژگی میانگین، واریانس، همگنی، کنتراست، عدم شباهت، آنتروپی، همبستگی و گشتاور دوم استخراج گردید. در این تحقیق، ابعاد کرنل برابر 3×3 ، 5×5 و 7×7 در ۴ جهت ۰، ۴۵، ۹۰ و ۱۳۵ درجه در نظر گرفته شده است. به‌منظور حذف تأثیر جهات مختلف زوایا، از هر ویژگی در ۴ جهت میانگین‌گیری شده است. درنهایت، تعداد کل ویژگی‌های بافتی تولیدشده، ۹۶ ویژگی می‌شود که با تلفیق آن‌ها با باندهای طیفی و داده ارتفاعی، ۱۰۰ ویژگی، برای ورود به مرحله طبقه‌بندی آماده می‌گردد.

مرحله انتخاب ویژگی و طبقه‌بندی تصویر بر اساس دو رویکرد مختلف پیاده‌سازی شده است. در رویکرد اول با استفاده از روش CB-AdaBoost انتخاب ویژگی و طبقه‌بندی به‌صورت هم‌زمان انجام شده است. از آنجایی‌که در الگوریتم CB-AdaBoost، پارامتر ۷ نیز به‌عنوان ورودی به الگوریتم معرفی می‌گردد، باید از قبل تعیین شود. بدین

پیاده‌سازی روش تولید ویژگی شبه عمیق، مطابق توضیحات بخش (۱-۲)، در چند زیرگروه انجام شده است.

منظور، با استفاده از روش KNN، طبق رابطه‌ی (۱۳)، احتمال درستی برچسب هر نمونه از بین K نزدیک‌ترین همسایگی‌اش، محاسبه گردید.

$$\gamma = p(z = y|x) = \frac{1}{K} \sum_{j=1}^K \sum_{x_j \in N(x)} I(y_j = y) \quad (13)$$

که در آن، $N(x)$ بیانگر مجموعه K نزدیک‌ترین همسایگی x است. بهترین مقدار برای K را به روش سعی و خطا می‌توان به دست آورد. در این تحقیق، مقدار $K=5$ در نظر گرفته شده است. پس از تعیین γ ، مراحل پیاده‌سازی، مطابق با الگوریتم CB-AdaBoost که جزئیات آن در بخش (۲-۲-۱) بیان شد، انجام گردیده است. یادگیرنده‌ی پایهی مورد استفاده در این روش، درخت تصمیم با یک گره و دو برگ می‌باشد که به آن استامپ تصمیم^۱ می‌گویند. پارامتر تنظیمی این روش، تعداد تکرار الگوریتم یادگیری است که در اینجا بر روی ۵۰ تکرار تنظیم شده است. به منظور ارزیابی مقایسه‌ای این روش نسبت به سایر روش‌های یادگیری ماشین، روش آدابوست اصلی، RF و SVM-PSO که روش‌های مطرحی در جامعه‌ی یادگیری ماشین و سنجش از دور می‌باشند، نیز پیاده‌سازی شده است. پارامتر تنظیمی روش آدابوست و RF نیز تعداد تکرار الگوریتم یادگیری است که در این تحقیق، به ترتیب، تعداد ۵۰ و ۱۰۰ تکرار در نظر گرفته شده است. در طبقه‌بندی SVM از تابع کرنل پایه شعاعی استفاده شده است [۴۱].

در رویکرد دوم، روش CB-AdaBoost به عنوان الگوریتم انتخاب ویژگی و روش ماشین بردار پشتیبان به عنوان طبقه‌بندی کننده در نظر گرفته شد، همچنین با ترکیب روش CB-AdaBoost با روش SVM-PSO، انتخاب ویژگی در دو مرحله صورت گرفت، بدین نحو که ۵۰ درصد ویژگی‌های انتخابی توسط الگوریتم CB-AdaBoost وارد الگوریتم انتخاب ویژگی و طبقه‌بندی SVM-PSO گردید. به منظور ارزیابی مقایسه‌ای، رویکرد ترکیبی فوق با استفاده از روش آدابوست نیز پیاده‌سازی گردید. همچنین الگوریتم SVM-PSO بدون دخالت دادن الگوریتم‌های آدابوست و CB-AdaBoost نیز پیاده‌سازی شده است. پیاده‌سازی روش جنگل‌های تصادفی در محیط

نرم‌افزار EnMAP-Box و سایر روش‌ها در محیط نرم‌افزار MATLAB R2015a انجام شده است.

۳-۳- ارزیابی و تحلیل نتایج

به منظور ارزیابی نتایج روش‌های پیاده‌سازی شده، تصویر طبقه‌بندی شده با داده‌ی واقعیت زمینی مورد ارزیابی قرار گرفته و از معیارهای صحت کلی، ضریب کاپا، دقت کاربر^۲، دقت تولیدکننده^۳ و درصد کیفیت^۴ کلاس ساختمان که از طریق ماتریس ابهام به دست می‌آیند، استفاده شده است. در این تحقیق دو استراتژی دنبال شده است. در رویکرد اول، انتخاب ویژگی و طبقه‌بندی به طور هم‌زمان توسط روش‌های آدابوست، RF، CB-AdaBoost و SVM-PSO بر روی مجموعه ویژگی‌های شبه عمیق و GLCM که هر کدام با ۴ ویژگی IR، R، G و nDSM تلفیق شده‌اند، انجام شد که نتایج ارزیابی این روش‌ها برای دو تصویر مطالعاتی در جدول (۱) ارائه شده است.

مطابق با این جدول، نتایج ارزیابی‌ها را می‌توان از دو دیدگاه مقایسه کرد. از یکسو با مقایسه روش‌های انتخاب ویژگی و طبقه‌بندی بر اساس معیارهای صحت کلی، ضریب کاپا و کیفیت در هر دو تصویر مطالعاتی، مشاهده می‌گردد که روش SVM-PSO در ۱۰ مورد و CB-AdaBoost در ۲ مورد از ۱۲ حالت پیاده‌سازی شده، عملکرد بهتری نسبت به سایر روش‌ها دارد. همچنین، روش CB-AdaBoost در ۱۱ و ۱۰ مورد از ۱۲ حالت پیاده‌سازی شده، به ترتیب دقت‌های بهتر از RF و آدابوست حاصل کرده است. مقایسه بین دو روش آدابوست و RF نیز نشان می‌دهد که هر کدام، در ۶ مورد از ۱۲ حالت نسبت به دیگری دقت بهتری حاصل کرده است که می‌توان گفت بر اساس شرایط پیاده‌سازی در این تحقیق، این دو روش تقریباً از عملکرد مساوی برخوردارند.

به طور کلی، بر اساس میانگین دقت روش‌های مذکور مطابق با شکل (۴)، مشاهده می‌شود که روش CB-AdaBoost از عملکرد بهتری نسبت به دو روش آدابوست و RF برخوردار است و روش SVM-PSO نسبت به سایر روش‌ها، دقت‌های بیشتری حاصل کرده است.

^۲ User Accuracy (U.A)

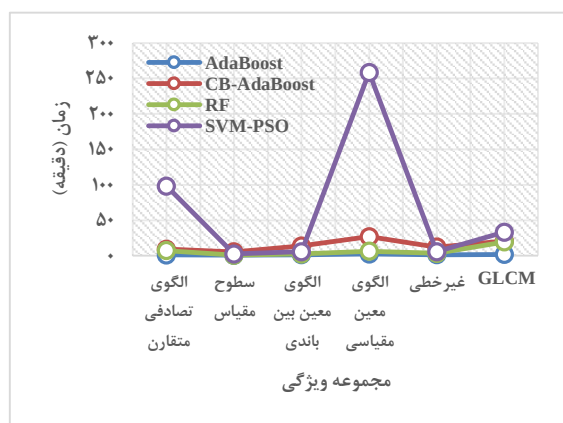
^۳ Producer Accuracy (P.A)

^۴ Quality Percent (Q.P)

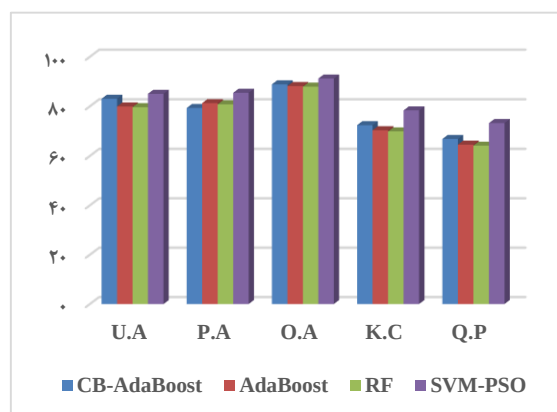
^۱ Decision Stump

جدول ۱- ارزیابی دقت نتایج روش‌های ارائه شده بر اساس رویکرد اول

روش‌های مختلف	میارهای دقت	تصویر تست اول						تصویر تست دوم					
		الگوی تصادفی متقارن ۶۴	سطوح مقیاس ۳۲	الگوی معین بین بانندی ۸۸	الگوی معین مقیاسی ۱۷۲	غیر خطی ۸۸	GLCM ۱۰۰	الگوی تصادفی متقارن ۴۴	سطوح مقیاس ۳۲	الگوی معین بین بانندی ۸۸	الگوی معین مقیاسی ۱۷۲	غیر خطی ۸۸	GLCM ۱۰۰
AdaBoost	U.A	۶۴,۷۶	۶۴,۷۱	۶۱,۷۷	۶۴,۹۱	۷۴,۸۲	۶۳,۹۵	۹۲,۷۳	۹۶,۰۷	۹۲,۷۱	۹۳,۲۱	۹۶,۵۸	۹۲,۹
	P.A	۹۲,۴۵	۹۴,۲۶	۹۳,۲۲	۹۲,۹۳	۹۳,۶۵	۹۳,۹۹	۶۹,۸۸	۶۷,۳۰	۷۰,۸۲	۷۱,۷۱	۶۶,۱۲	۶۸,۹۱
	O.A	۸۹,۶۸	۹۰,۰۱	۸۹,۱۲	۸۹,۸۰	۹۲,۲۹	۸۹,۷۷	۸۶,۴۸	۸۵,۵۲	۸۶,۹۸	۸۷,۵۳	۸۴,۸۶	۸۵,۹۷
	K.C	۶۹,۸۴	۷۰,۶۵	۶۷,۷۶	۷۰,۱۸	۷۸,۲۶	۶۹,۹۰	۶۹,۸۷	۶۸,۵۸	۷۰,۸۳	۷۲	۶۷,۴۴	۶۸,۹۱
	Q.P	۶۱,۵۱	۶۲,۲۵	۵۹,۱۲	۶۱,۸۵	۷۱,۲۱	۶۱,۴۳	۶۶,۲۵	۶۵,۵	۶۷,۰۸	۶۸,۱۵	۶۴,۶۱	۶۵,۴۶
CB-AdaBoost	U.A	۶۶,۹۰	۶۹,۸۱	۷۴,۵۷	۷۹,۱۲	۷۶,۲۶	۶۹,۵۸	۹۲,۲۲	۹۶,۰۹	۹۴,۰۵	۸۹,۶۲	۹۶,۴۱	۹۱,۹۰
	P.A	۸۵,۱۱	۸۸,۰۵	۸۸,۷۴	۸۵,۹۶	۸۶,۶۴	۸۶,۱۰	۷۲,۰۹	۶۸,۹۹	۷۰,۴۱	۷۸,۳۶	۶۶,۹۴	۷۴,۶۸
	O.A	۸۸,۵۹	۸۹,۸۹	۹۱,۰۰	۹۱,۳۹	۹۰,۹۶	۸۹,۳۹	۸۷,۵۶	۸۶,۵۲	۸۶,۹۹	۸۹,۹۵	۸۵,۳۵	۸۸,۷۷
	K.C	۶۷,۶۶	۷۱,۴۴	۷۴,۹۲	۷۶,۷۱	۷۵,۲۱	۷۰,۱۷	۷۱,۸۹	۷۰,۴۸	۷۱,۰۶	۷۶,۴۱	۶۸,۳۱	۷۴,۲۸
	Q.P	۵۹,۸۹	۶۳,۷۷	۶۷,۶۹	۷۰,۰۶	۶۸,۲۴	۶۲,۵۵	۶۷,۹۶	۶۷,۱۲	۶۷,۴۱	۷۱,۸۴	۶۵,۳۱	۷۰,۰۷
RF	U.A	۶۱,۳۸	۶۳,۱۷	۷۵,۴۵	۶۱,۰۱	۷۳,۷۵	۶۴,۲۰	۹۰,۰۱	۹۵,۳۴	۹۵,۴۱	۸۹,۵	۹۷,۰۸	۸۹,۶۱
	P.A	۹۷,۰۹	۹۱,۵۶	۸۵,۹۸	۹۶,۳۳	۸۸,۰۲	۹۳,۶۳	۷۲,۷۵	۶۹,۱۹	۶۲,۳۲	۷۴,۷۱	۶۲,۵۶	۷۵,۵۳
	O.A	۸۹,۶۹	۸۹,۱۴	۹۰,۶۱	۸۹,۴۸	۹۰,۷۵	۸۹,۷۷	۸۷,۴۹	۸۶,۵۲	۸۲,۱۸	۸۸,۳۲	۸۲,۵۴	۸۸,۷۲
	K.C	۶۹,۱۲	۶۸,۱۴	۷۴,۲۳	۶۸,۴۸	۷۴,۲۸	۶۹,۹۴	۷۱,۴۲	۷۰,۳۵	۶۲,۳۶	۷۳,۰۲	۶۳,۳۳	۷۳,۸۵
	Q.P	۶۰,۲۷	۵۹,۷۰	۶۷,۱۸	۵۹,۶۳	۶۷,۰۲	۶۱,۵۱	۶۷,۳۲	۶۶,۹۳	۶۰,۵۱	۶۸,۶۸	۶۱,۴۱	۶۹,۴۵
SVM-PSO	U.A	۸۱,۵۴	۷۹,۷۳	۸۲,۳	۸۱,۲۶	۸۳,۲۴	۷۷,۵۷	۸۳,۷۷	۹۶,۵۹	۸۷,۳۴	۸۴,۱۱	۹۸,۱۱	۸۴,۹۴
	P.A	۹۵,۸۹	۹۴,۶۵	۹۳,۵۹	۹۵,۱۳	۹۰,۱۸	۹۷,۲۶	۷۵,۸۰	۶۸,۸۲	۸۰,۱۱	۸۲,۹۰	۶۵,۰۳	۸۶,۲۰
	O.A	۹۴,۴۱	۹۳,۶۹	۹۴,۰۵	۹۴,۱۷	۹۳,۴۲	۹۳,۷۳	۸۷,۷۰	۸۶,۵۰	۹۰,۱۷	۹۰,۴۹	۸۴,۳۶	۹۱,۸۰
	K.C	۸۴,۵۱	۸۲,۴۷	۸۳,۶۹	۸۳,۸۶	۸۲,۲۳	۸۲,۳۱	۷۰,۸۲	۷۰,۵۲	۷۶,۵۸	۷۶,۸۲	۶۶,۷۹	۷۹,۸۴
	Q.P	۷۸,۷۸	۷۶,۲۹	۷۷,۹۱	۷۸,۰۱	۷۶,۳۳	۷۵,۹۱	۶۶,۰۹	۶۷,۱۹	۷۱,۷۷	۷۱,۶۸	۶۴,۲۳	۷۴,۷۷



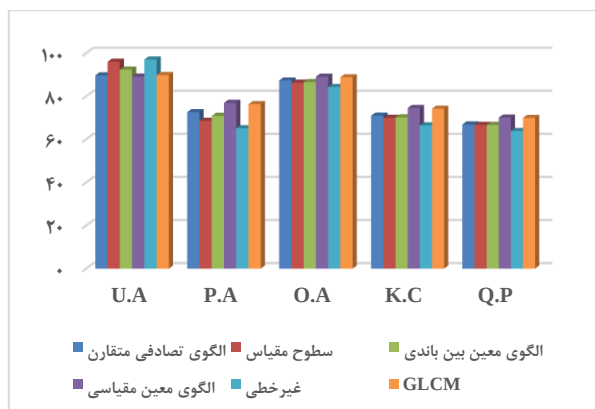
شکل ۵- نمودار زمان محاسبات بر اساس ویژگی‌های مختلف برای روش‌های موردقیاس در رویکرد اول



شکل ۴- نمودار میانگین دقت نتایج روش‌های ارائه شده در رویکرد اول بر اساس معیارهای ارزیابی مختلف

از سوی دیگر، با مقایسه دقت‌ها از دیدگاه مجموعه ویژگی‌های بکار رفته، بر اساس میانگین دقت‌ها طبق نمودارهای شکل (۶) مشاهده می‌گردد که در تصویر تست اول، مجموعه ویژگی شبه عمیق غیرخطی و در تصویر تست دوم، مجموعه ویژگی الگوی معین مقیاسی بیشترین دقت‌ها را حاصل کرده‌اند.

درحالی‌که با در نظر گرفتن زمان محاسبات این روش‌ها، طبق نمودار شکل (۵) مشاهده می‌شود که با افزایش تعداد ویژگی‌ها به ۱۷۲ ویژگی در مجموعه ویژگی الگوی معین مقیاسی، روش‌های آدابوست، CB-AdaBoost و RF به‌طور قابل‌توجهی سریع‌تر از روش SVM-PSO می‌باشند.



ب) تصویر تست دوم



الف) تصویر تست اول

شکل ۶- نمودار میانگین دقت نتایج ارزیابی بر اساس مجموعه ویژگی‌های مختلف

بر روی هر دو تصویر مطالعاتی، با استفاده از الگوریتم SVM-PSO بدون دخالت دادن الگوریتم‌های آدابوست و CB-AdaBoost نیز انجام شده است.

بر اساس نتایج ارزیابی رویکردهای ترکیبی فوق، از مقایسه‌ی دقت نتایج روش‌های ترکیبی Ada-SVMpsو و CB-SVMpsو در دو حالت با و بدون استفاده از CV مشاهده می‌گردد که در هر دو تصویر تست اول و دوم هر یک از این روش‌ها در حالت بدون CV نسبت به حالت با CV خود، دقت‌های بهتری حاصل کرده‌اند. این موضوع می‌تواند به این دلیل باشد که هنگام استفاده از روش CV، تعداد ویژگی‌های بیشتری توسط الگوریتم‌های آدابوست و CB-AdaBoost امتیازدهی و انتخاب می‌گردند که ممکن است فضای مسئله را بدون افزودن اطلاعات مفید افزایش دهند. همچنین مقایسه بین این دو روش نسبت به هم نشان می‌دهد که در تصویر تست اول Ada-SVMpsو دقت بهتری نسبت به CB-SVMpsو حاصل کرده است و در تصویر تست دوم عکس این حالت رخ داده است. علاوه بر این، با مقایسه این دو روش نسبت به روش SVM-PSO مشاهده می‌گردد که در هر دو تصویر تست اول و دوم، دو روش فوق دقت‌های بیشتری نسبت به SVM-PSO حاصل کرده‌اند.

از دیدگاه دیگر که آدابوست و CB-AdaBoost به‌عنوان الگوریتم انتخاب ویژگی، ویژگی‌های انتخابی را در ۳ مقیاس مختلف وارد طبقه‌بندی کننده SVM می‌کنند، تفاوت دقت‌های ناشی از بکارگیری تصاویر مختلف و نیز تعداد و نوع ویژگی‌های انتخابی مشاهده می‌شود. از جمله در تصویر تست اول، روش آدابوست با تعداد ۲۰ و ۵۰ درصد از بهترین ویژگی‌هایش دقت‌های بهتری از CB-AdaBoost با تعداد ۲۰ ویژگی حاصل کرده است و حال آنکه در تصویر تست دوم

در اینجا یک نکته قابل‌توجه این است که ویژگی‌های شبه عمیق، بر اساس نحوه استخراج ویژگی در ۵ مجموعه به‌صورت جدا از هم مورد ارزیابی قرار گرفته‌اند و حال آنکه ویژگی‌های بافت موجود در مجموعه ویژگی GLCM شامل ویژگی‌های متنوعی از جمله واریانس، کنتراست، عدم شباهت و غیره می‌باشد. بر این اساس و با توجه به دقت نتایج هر یک از زیرمجموعه ویژگی‌های شبه عمیق، می‌توان گفت که این مجموعه ویژگی حاوی ویژگی‌های کارآمدی به‌منظور شناسایی ساختمان می‌باشد.

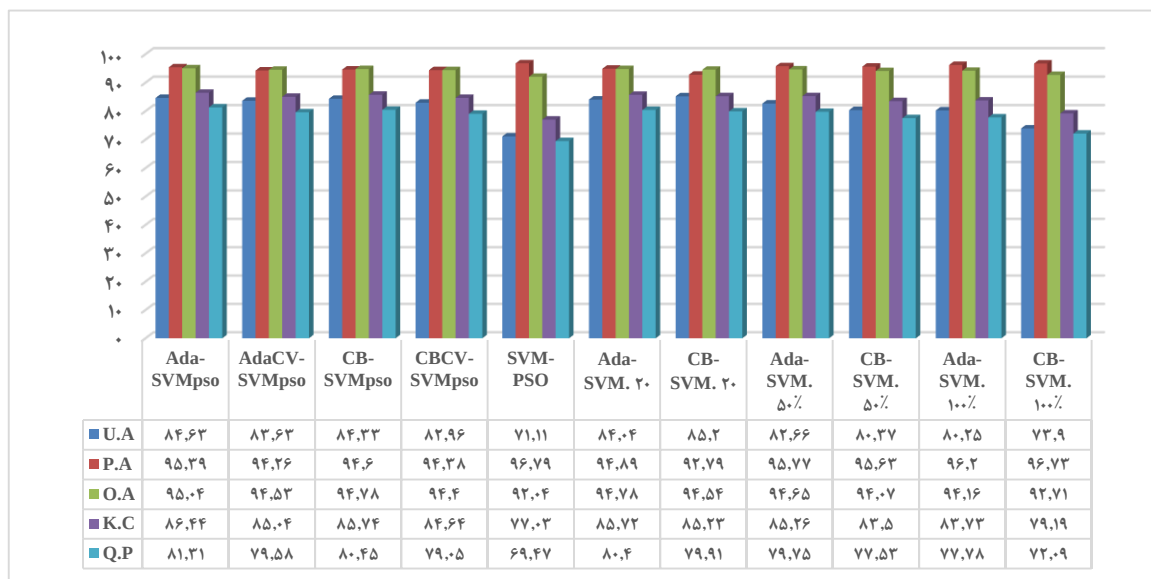
در رویکرد دوم، مطابق با روش‌های ارائه‌شده در اشکال (۷ و ۸)، ترکیب‌بندی روش‌ها، به این صورت است که ابتدا با استفاده از الگوریتم‌های آدابوست و CB-AdaBoost که بر روی ۵۰۰ تکرار تنظیم‌شده‌اند، ویژگی‌های مناسب از مجموعه تمام ویژگی‌های شبه عمیق به همراه ویژگی‌های اصلی IR، R، G و ویژگی ارتفاعی nDSM که برای تصویر تست اول شامل ۴۲۸ ویژگی و برای تصویر تست دوم شامل ۴۰۸ ویژگی می‌باشند، انتخاب می‌گردد. سپس در یک رویکرد، تعداد ۲۰، ۵۰ درصد و ۱۰۰ درصد از بهترین ویژگی‌های انتخابی وارد طبقه‌بندی کننده SVM شده و در رویکرد دیگر، ۵۰ درصد از بهترین ویژگی‌های انتخابی وارد الگوریتم انتخاب ویژگی و طبقه‌بندی SVM-PSO شده است. در نسخه‌ای دیگر از رویکرد ترکیبی آدابوست و CB-AdaBoost با SVM-PSO، از روش اعتبارسنجی متقابل^۱ (CV) با ۵ بخش اجرا یا اصطلاحاً فولد^۲ در مرحله‌ی انتخاب اولیه‌ی ویژگی‌ها توسط الگوریتم‌های انتخاب ویژگی مذکور، استفاده شده است. علاوه بر روش‌های فوق‌الذکر، آزمایشات

^۱ Cross Validation (CV)

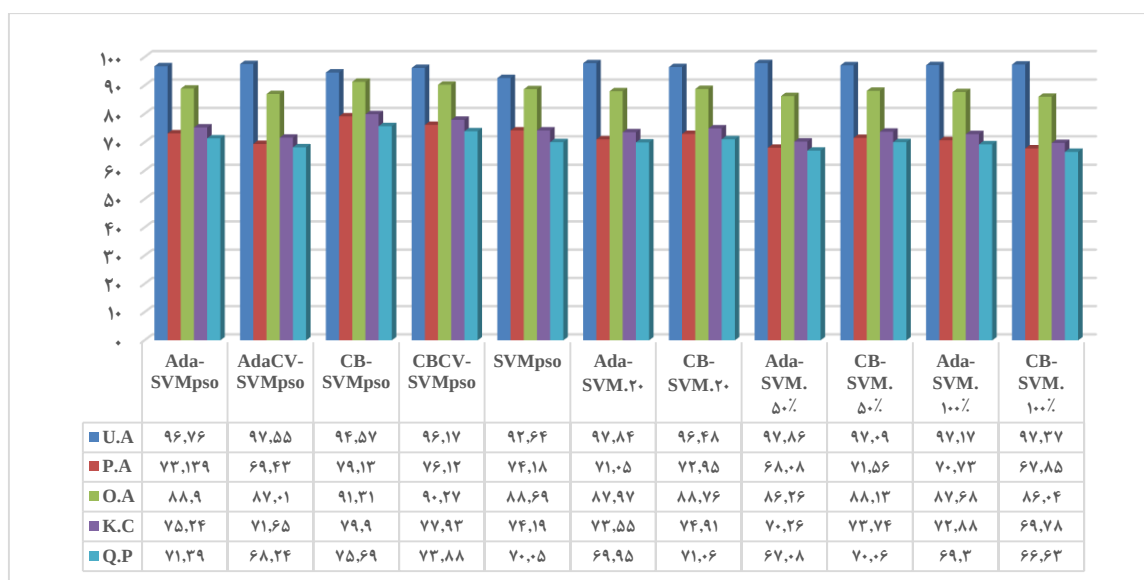
^۲ Fold

ناشی از در نظر نگرفتن برخی از بهترین ویژگی‌ها هنگام استفاده تنها از ۲۰ ویژگی برتر باشد. در مقایسه‌ای دیگر بین روش‌های ترکیبی فوق با روش SVM-PSO برای تصویر تست اول، بیشترین دقت‌های حاصل از روش‌های ترکیبی مذکور، با تفاوت دقت حدود ۱ تا ۱۰ درصد نسبت به SVM-PSO از عملکرد بهتری برخوردارند و برای تصویر تست دوم عملکردی تقریباً برابر حاصل کرده است.

عکس این حالت مشاهده می‌شود. همچنین مقایسه دقت بین تعداد ویژگی‌های ورودی در سه مقیاس مختلف برای هر الگوریتم نشان می‌دهد که در ۳ مورد از ۴ حالت، با استفاده از تعداد ۲۰ ویژگی برتر نسبت به استفاده از ۵۰ و ۱۰۰ درصد از بهترین ویژگی‌ها نتایج بهتری حاصل شده است. در صورتی که تعداد ۵۰ و ۱۰۰ درصد بهترین ویژگی‌ها، دقت‌های بیشتری از ۲۰ ویژگی برتر حاصل کنند؛ می‌تواند



شکل ۷- نتایج ارزیابی رویکرد ترکیبی با استفاده از مجموعه ویژگی شبه عمیق بر روی تصویر تست اول

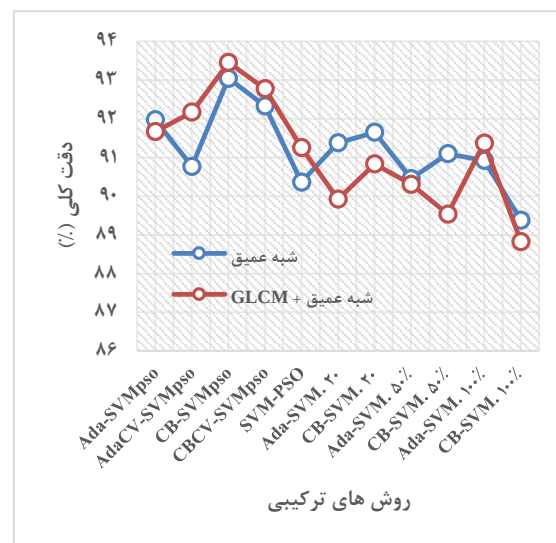


شکل ۸- نتایج ارزیابی رویکرد ترکیبی با استفاده از مجموعه ویژگی شبه عمیق بر روی تصویر تست دوم

مشاهده می‌گردد که در برخی از روش‌ها بهبود دقت حاصل شده است اما در برخی دیگر، نتیجه نسبت به استفاده تنها از مجموعه ویژگی‌های شبه عمیق بهتر نشده است. این امر

مطابق با نمودار شکل (۹)، با افزودن ویژگی‌های بافت حاصل از ماتریس GLCM به مجموعه ویژگی‌های شبه عمیق و تکرار آزمایشات با استفاده از این مجموعه ویژگی

ممکن است ناشی از جایگزینی ویژگی‌های بافت GLCM با برخی از بهترین ویژگی‌های شبه عمیق در انتخاب اولیه‌ی ویژگی‌ها باشد که در این حالت نادیده گرفته شده‌اند.

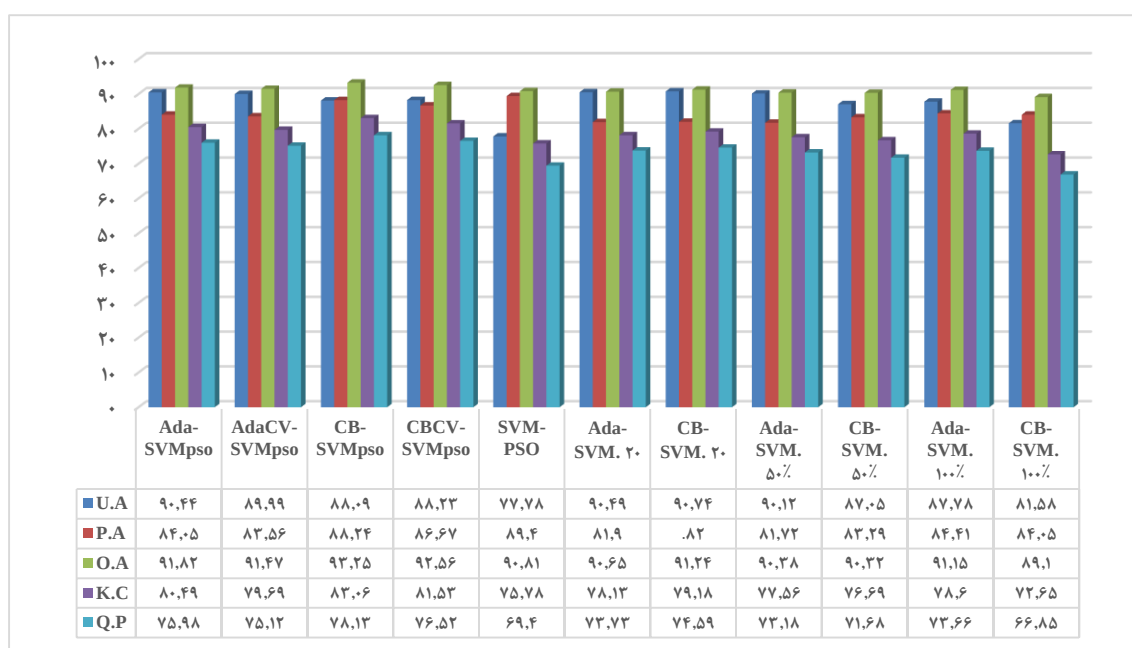


شکل ۹- نمودار میانگین صحت کلی نتایج ارزیابی رویکرد ترکیبی با استفاده از دو مجموعه ویژگی مختلف

درنهایت، با در نظر گرفتن تمام روش‌های ارائه شده در رویکرد دوم، بر اساس میانگین دقت نتایج ارزیابی‌ها در هر دو تصویر مطالعاتی و هر دو مجموعه ویژگی، مطابق با نمودار شکل (۱۰) مشاهده می‌گردد که روش CB-SVMpso با صحت کلی ۹۳,۲۵ درصد، ضریب کاپا ۸۳,۰۶ درصد و کیفیت کلاس ساختمان ۷۸,۱۳ درصد نسبت به سایر روش‌ها از عملکرد بهتری برخوردار است و در مقایسه

با روش SVM-PSO به ترتیب، ۲,۴۴، ۷,۲۷ و ۸,۷۳ درصد بهبود یافته است.

علاوه بر این، با مقایسه‌ی این روش‌ها از نقطه نظر سرعت محاسبات، طبق نمودار شکل (۱۱) مشاهده می‌شود که روش CB-SVMpso به طور متوسط ۲ برابر سریع تر از روش SVM-PSO می‌باشد. زمان محاسبات با ابعاد تصویر و تعداد باندها به صورت تصاعدی در ارتباط است که این موضوع در نمودار شکل (۵) نیز مشاهده شده است. حال با افزایش ابعاد تصویر و به تبع آن افزایش ابعاد باندهای ویژگی تولید شده، حجم محاسبات تا حد زیادی ممکن است بالا رود که عملاً روش‌های فرا ابتکاری مانند PSO، GA و ... که ماهیت جمعیت مینا و تکراری دارند قابلیت اجرا نداشته باشند ولی در عوض روش پیشنهادی به دلیل انتخاب اولیه‌ی ویژگی‌های نخبه توسط الگوریتم CB-AdaBoost، حجم محاسبات، زمان اجرا و حافظه مورد نیاز را به طور قابل توجهی کاهش می‌دهد. از این رو، رویکرد پیشنهادی، برای طبقه بندی داده‌های با حجم بالا مناسب می‌باشد. لازم به ذکر است که در حالت استفاده از روش CV، به دلیل افزایش حجم محاسبات، زمان اجرا نیز افزایش یافته است. به طور کلی، این نتایج، کارآمدی روش پیشنهادی را چه از نظر دقت و چه از نظر زمان محاسبات، اثبات می‌کند.



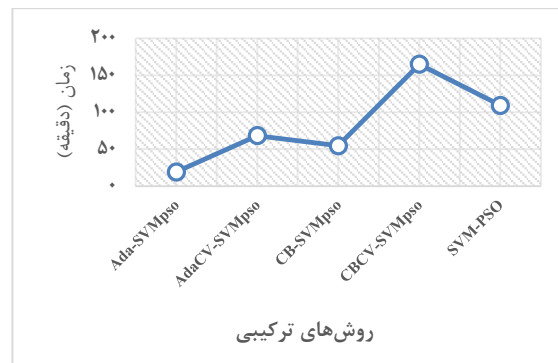
شکل ۱۰- نمودار میانگین دقت نتایج ارزیابی رویکرد ترکیبی

دو مجموعه ویژگی به خوبی شناسایی شده‌اند، روش SVM-PSO در شناسایی آن‌ها دچار خطا شده است. همچنین روش Ada-SVMpsو در تفکیک برخی درختان از ساختمان‌ها، خطا ایجاد کرده است. این موضوع به ویژگی‌های انتخاب‌شده توسط این روش‌ها برمی‌گردد. بر این اساس می‌توان گفت، روش پیشنهادی در امر انتخاب ویژگی بهتر عمل کرده است. به‌طور کلی، کیفیت تصاویر ساختمان‌های آشکارسازی شده توسط روش پیشنهادی نسبت به روش SVM-PSO و Ada-SVMpsو بهبود یافته است. همچنین مشاهده می‌شود که با افزودن ویژگی‌های بافت GLCM به ویژگی‌های شبه‌عمیق، برخی ساختمان‌های شناسایی‌شده توسط روش‌های مذکور، نسبت به حالتی که این ویژگی‌ها در مجموعه ویژگی شبه عمیق دخالت داده نمی‌شوند، کمی بهبود یافته‌اند، اما در کل بهبود چندانی حاصل نشده است. از این‌رو، نیاز به بهبود کیفیت تصاویر با استفاده از عملیات پس‌پردازش مانند استفاده از عملگرهای مورفولوژی احساس می‌شود.

۴- نتیجه‌گیری

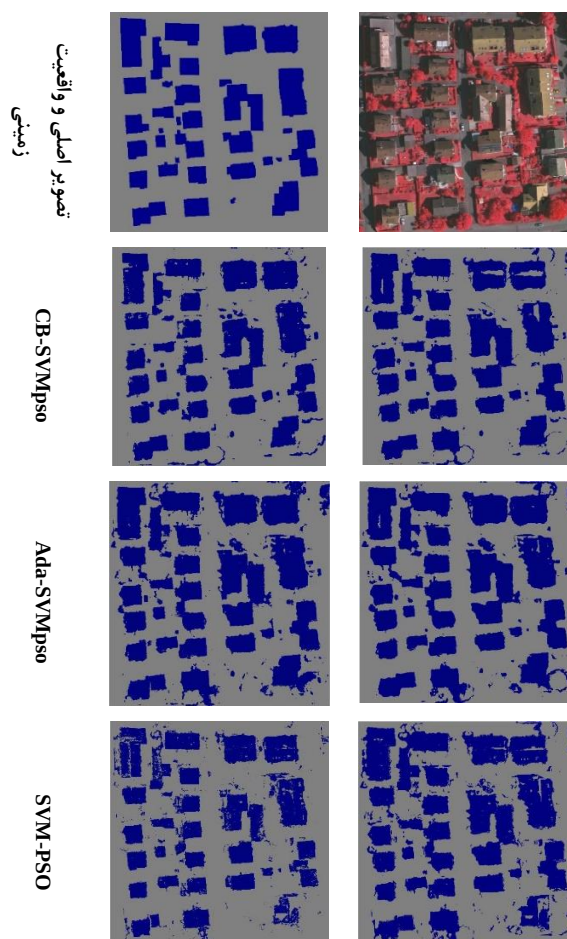
شناسایی اتوماتیک ساختمان با استفاده از داده‌های سنجش‌ازدور، از برخی مشکلات بخصوص در مناطق شهری ناشی از پیچیدگی طیفی در صحنه رنج می‌برد. تصاویر با قدرت تفکیک مکانی بالا حاوی جزئیات زیادی از صحنه تصویر بوده و سقف ساختمان‌ها که غیر همگن، شیب‌دار، صاف و غیره می‌باشد، می‌تواند خواص طیفی مختلفی را در میان مسائل دیگر ایجاد کند. همچنین با توجه به استفاده از مواد مشابه، برخی از ساختمان‌ها نمی‌توانند از محوطه‌ها و پارکینگ‌ها کاملاً تفکیک گردند. برای غلبه بر این مسائل، استفاده از اطلاعات مجاورت و ارتفاع ضروری است. بر این اساس، بخش عمده این تحقیق، استفاده از ویژگی‌های مکانی پیکسل‌های مجاور تصویر چندطیفی و داده ارتفاعی، برای افزایش دقت طبقه‌بندی ساختمان می‌باشد.

در این راستا، از یکسو با گسترش فضای ویژگی به روش شبه عمیق، هدف این بوده است که الگوریتم طبقه‌بندی با اطلاعات سطح بالا و جامع‌تری آموزش ببیند؛ اما تمام ویژگی‌های موجود برای تفکیک عارضه ساختمان از غیر ساختمان مفید و مناسب نیستند. بعلاوه اینکه به دلیل حجم بالای داده‌های ورودی و افزایش زمان



شکل ۱۱- نمودار میانگین زمان محاسبات برای رویکرد ترکیبی

در شکل (۱۲)، تصاویر ساختمان‌های آشکارسازی شده توسط روش پیشنهادی در مقایسه با روش‌های Ada-SVMpsو و SVM-PSO، نشان داده شده است.



الف) ویژگی‌های شبه عمیق و ب) ویژگی‌های شبه عمیق و بافت GLCM

شکل ۱۲- تصاویر ساختمان‌های آشکارسازی شده در تصویر تست دوم: ستون‌ها بیانگر مجموعه ویژگی مختلف، سطرها بیانگر روش‌های متفاوت

مطابق با این شکل، مشاهده می‌شود که برخی از ساختمان‌هایی که توسط روش پیشنهادی با استفاده از هر

محاسبات و حافظه موردنیاز، به منظور کاهش هزینه پردازش، لازم است تا عملیات انتخاب ویژگی صورت گیرد؛ بنابراین در پژوهش حاضر، با هدف بهبود شناسایی اتوماتیک ساختمان از داده‌های سنجش از دور، یک رویکرد ترکیبی جدید برای انتخاب ویژگی‌های بهینه از مجموعه داده بزرگ در یک زمان معقول ارائه شده است. روش پیشنهادی بر اساس ادغام الگوریتم آداپوست توسعه یافته با روش ماشین بردار پشتیبان بهینه سازی شده با ازدحام ذرات، ویژگی‌های بهینه را انتخاب کرده و به طبقه بندی باینری تصویر در دو کلاس ساختمان و زمینه می پردازد. همچنین مقایسه‌ای بین روش‌های کارآمد از جمله روش SVM-PSO انجام شد. روش پیشنهادی بر روی دو مجموعه داده، تست و ارزیابی گردید. در آزمایشات، به منظور خلوص نتایج نهایی، هیچ مرحله‌ی پیش پردازش و

پس پردازشی دخالت داده نشده است. مقادیر صحت کلی و ضریب کاپا نتایج روش پیشنهادی که به طور میانگین به ترتیب برابر با ۹۳٫۲۵ و ۸۳٫۰۶ درصد می باشد؛ در مقایسه با روش SVM-PSO به ترتیب ۲٫۴۴ و ۷٫۲۷ درصد بهبود دقت حاصل کرده است؛ بعلاوه اینکه زمان محاسبات را حدوداً به نصف کاهش می دهد. این موضوع نشان از توانمندی این رویکرد در شناسایی اکثریت ساختمان‌ها و عوارض غیر ساختمان در یک زمان معقول می باشد. با این حال، هنوز هم برخی از ساختمان‌ها به علت طیف‌های مشابه با عوارض مجاورشان شناسایی نشده است. برای این منظور، در تحقیقات آتی می توان علاوه بر ویژگی‌های بکار رفته در این تحقیق از ویژگی‌های تکمیلی مانند ویژگی‌های هندسی بهره برد.

مراجع

- [1] D. J. Kim and P. L. Manjusha, "Building Detection in High Resolution Remotely Sensed Images based on Automatic Histogram-Based Fuzzy C-Means Algorithm," 2017.
- [2] M. Ghanea, P. Moallem, and M. Momeni, "Automatic building extraction in dense urban areas through GeoEye multispectral imagery," *International journal of remote sensing*, vol. 35, pp. 5094-5119, 2014.
- [3] T. Hermosilla, L. A. Ruiz, J. A. Recio, and J. Estornell, "Evaluation of automatic building detection approaches combining high resolution images and LiDAR data," *Remote Sensing*, vol. 3, pp. 1188-1210, 2011.
- [4] N. Shrivastava and P. K. Rai, "Remote-sensing the urban area: Automatic building extraction based on multiresolution segmentation and classification," *Geografia-Malaysian Journal of Society and Space*, vol. 11, 2017.
- [5] X. Huang, W. Yuan, J. Li, and L. Zhang, "A new building extraction postprocessing framework for high-spatial-resolution remote-sensing imagery," *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing*, vol. 10, pp. 654-668, 2017.
- [6] D. M. McKeown and W. A. Harvey, "Automating knowledge acquisition for aerial image interpretation," in *Image Understanding and the Man-Machine Interface*, 1987, pp. 144-165.
- [7] T. Matsuyama, "Knowledge-based aerial image understanding systems and expert systems for image processing," *IEEE Transactions on Geoscience and Remote Sensing*, pp. 305-316, 1987.
- [8] J. Peng, D. Zhang, and Y. Liu, "An improved snake model for building detection from urban aerial images," *Pattern Recognition Letters*, vol. 26, pp. 587-595, 2005.
- [9] S. Ahmadi, M. V. Zoj, H. Ebadi, H. A. Moghaddam, and A. Mohammadzadeh, "Automatic urban building boundary extraction from high resolution aerial images using an innovative model of active contours," *International Journal of Applied Earth Observation and Geoinformation*, vol. 12, pp. 150-157, 2010.
- [10] K. Khoshelham, C. Nardinocchi, E. Frontoni, A. Mancini, and P. Zingaretti, "Performance evaluation of automated approaches to building detection in multi-source aerial data," *ISPRS Journal of Photogrammetry and Remote Sensing*, vol. 65, pp. 123-133, 2010.
- [11] U. Weidner and W. Förstner, "Towards automatic building extraction from high-resolution digital elevation models," *ISPRS journal of Photogrammetry and Remote Sensing*, vol. 50, pp. 38-49, 1995.
- [12] G. Forlani, C. Nardinocchi, M. Scaioni, and P. Zingaretti, "Complete classification of raw LIDAR data and 3D reconstruction of buildings," *Pattern analysis and applications*, vol. 8, pp. 357-374, 2006.
- [13] D. Chai, "A Probabilistic Framework for Building Extraction From Airborne Color Image and DSM," *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing*, vol. 10, pp. 948-959, 2017.

- [14] L. C. Chen, T.-A. Teo, Y.-C. Shao, Y.-C. Lai, and J.-Y. Rau, "Fusion of LIDAR data and optical imagery for building modeling," *International Archives of Photogrammetry and Remote Sensing*, vol. 35, pp. 732-737, 2004.
- [15] K. Khoshelham, Z. Li, and B. King, "A split-and-merge technique for automated reconstruction of roof planes," *Photogrammetric Engineering & Remote Sensing*, vol. 71, pp. 855-862, 2005.
- [16] F. Rottensteiner, J. Trinder, S. Clode, and K. Kubik, "Building detection by fusion of airborne laser scanner data and multi-spectral images: Performance evaluation and sensitivity analysis," *ISPRS Journal of Photogrammetry and Remote Sensing*, vol. 62, pp. 135-149, 2007.
- [17] M. Aladeemy, S. Tutun, and M. T. Khasawneh, "A new hybrid approach for feature selection and support vector machine model selection based on self-adaptive cohort intelligence," *Expert Systems with Applications*, vol. 88, pp. 118-131, 2017.
- [18] L. Matikainen, H. Kaartinen, and J. Hyyppä, "Classification tree based building detection from laser scanner and aerial image data," *International Archives of Photogrammetry, Remote Sensing and Spatial Information Sciences*, vol. 36, p. W52, 2007.
- [19] M. Salah, J. C. Trinder, and A. Shaker, "Performance evaluation of classification trees for building detection from aerial images and LiDAR data: a comparison of classification trees models," *International Journal of Remote Sensing*, vol. 32, pp. 5757-5783, 2011/10/20 2011.
- [20] H. Guan, Z. Ji, L. Zhong, J. Li, and Q. Ren, "Partially supervised hierarchical classification for urban features from lidar data with aerial imagery," *International journal of remote sensing*, vol. 34, pp. 190-210, 2013.
- [21] D. A. A. Gnana, S. Appavu, and E. J. Leavline, "Literature Review on Feature Selection Methods for High-Dimensional Data," *methods*, vol. 136, 2016.
- [22] Y. Freund and R. E. Schapire, "A decision-theoretic generalization of on-line learning and an application to boosting," *Journal of computer and system sciences*, vol. 55, pp. 119-139, 1997.
- [23] A. J. Wyner, M. Olson, J. Bleich, and D. Mease, "Explaining the success of adaboost and random forests as interpolating classifiers," *Journal of Machine Learning Research*, vol. 18, pp. 1-33, 2017.
- [24] C.-L. Huang and C.-J. Wang, "A GA-based feature selection and parameters optimization for support vector machines," *Expert Systems with applications*, vol. 31, pp. 231-240, 2006.
- [25] J. Shao and W. Förstner, "Gabor wavelets for texture edge extraction," in *ISPRS Commission III Symposium: Spatial Information from Digital Photogrammetry and Computer Vision*, 1994, pp. 745-753.
- [26] M. Varma and A. Zisserman, "Texture classification: are filter banks necessary?" in *2003 IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, 2003. *Proceedings*, 2003, pp. II-691-8 vol.2.
- [27] L. Breiman, "Random forests," *Machine learning*, vol. 45, pp. 5-32, 2001.
- [28] R. M. Haralick and K. Shanmugam, "Textural features for image classification," *IEEE Transactions on systems, man, and cybernetics*, pp. 610-621, 1973.
- [29] P. Tokarczyk, J. D. Wegner, S. Walk, and K. Schindler, "Features, color spaces, and boosting: New insights on semantic classification of remote sensing images," *IEEE Transactions on Geoscience and Remote Sensing*, vol. 53, pp. 280-295, 2015.
- [30] P. Viola and M. J. Jones, "Robust real-time face detection," *International journal of computer vision*, vol. 57, pp. 137-154, 2004.
- [31] R. Caruana and A. Niculescu-Mizil, "An empirical comparison of supervised learning algorithms," in *Proceedings of the 23rd international conference on Machine learning*, 2006, pp. 161-168.
- [32] J. Heo and J. Y. Yang, "AdaBoost based bankruptcy forecasting of Korean construction companies," *Applied soft computing*, vol. 24, pp. 494-499, 2014.
- [33] D. S. Prasvita and A. M. Arymurthy, "Classification of LiDAR Images Fused With Aerial Optical Images Using Ensemble Classifier AdaBoost. MH and Post-processing BFS," *International Journal of Technology And Business*, vol. 1, pp. 10-16, 2017.
- [34] S. Wu and H. Nagahashi, "Parameterized adaboost: introducing a parameter to speed up the training of real adaboost," *IEEE Signal Processing Letters*, vol. 21, pp. 687-691, 2014.
- [35] C. Gao, N. Sang, and Q. Tang, "On selection and combination of weak learners in AdaBoost," *Pattern Recognition Letters*, vol. 31, pp. 991-1001, 2010.
- [36] Z. Fu, D. Zhang, X. Zhao, and X. Li, "Adaboost algorithm with floating threshold," 2012.

- [37] Z. Luo, X. Dang, and Y. Chen, "Label confidence based AdaBoost algorithm," in Neural Networks (IJCNN), 2017 International Joint Conference on, 2017, pp. 3617-3624.
- [38] R. E. Schapire, "The strength of weak learnability," Machine learning, vol. 5, pp. 197-227, 1990.
- [39] J. A. Richards and J. Richards, Remote sensing digital image analysis vol. 3: Springer, 1999.
- [40] ISPRS. (2013). Web site of the ISPRS test project on urban classification and 3D building reconstruction. Available: <http://www2.isprs.org/commissions/comm3/wg4/tests.html>
- [41] E. Tamimi, H. Ebadi, and A. Kiani, "Evaluation of different metaheuristic optimization algorithms in feature selection and parameter determination in SVM classification," Arabian Journal of Geosciences, vol. 10, p. 478, 2017.