

طبقه بندی تصاویر ابرطیفی مبتنی بر تلفیق ویژگی‌های مستخرج از روش‌های کدگذاری تنک، تبدیلات خطی و غیر خطی

سمیرا الهیاری بک^{۱*}، علیرضا صفدری نژاد^۲، روح‌اله کریمی^۳

^۱ دانشجوی کارشناسی ارشد مهندسی نقشه‌برداری، فتوگرامتری - گروه ژئودزی و مهندسی نقشه‌برداری - دانشگاه تفرش
geo96.allahyari@tafreshu.ac.ir

^۲ استادیار گروه ژئودزی و مهندسی نقشه‌برداری - دانشگاه تفرش
safdarinezhad@tafreshu.ac.ir

^۳ استادیار گروه ژئودزی و مهندسی نقشه‌برداری - دانشگاه تفرش
karimi@tafreshu.ac.ir

(تاریخ دریافت آذر ۱۳۹۸، تاریخ تصویب اسفند ۱۳۹۸)

چکیده

طبقه‌بندی یکی از مهم‌ترین روش‌های استخراج اطلاعات از تصاویر ابرطیفی است. در این مقاله، راهکاری نوین با هدف تولید ویژگی‌های بمنظور طبقه‌بندی این تصاویر پیشنهاد شده است. این راهکار تلفیقی از تبدیلات خطی، غیرخطی و نمایش تنک بمنظور تولید ویژگی‌های موثر در فرایند طبقه‌بندی تصاویر ابرطیفی است. در روند پیشنهادی، ابتدا با رویکردی جدید و نظارت شده از تبدیل غیرخطی تحلیل مؤلفه‌های اصلی (NLPCA) بمنظور انتقال داده‌های طیفی به فضایی با ابعاد بیشتر استفاده شده است. در مرحله دوم، بکمک تبدیل تحلیل تفکیک‌پذیری خطی (LDA) فرامکعب حاصل از مرحله قبل به فضایی با بعد کمتر انتقال می‌یابد. در ادامه با هدف هم مقیاس کردن ویژگی‌های تولیدی و بهره‌گیری از پتانسیل تمامی داده‌های آموزشی، داده‌ها از طریق روش‌های تخمین تنک سیگنال به فضای ویژگی جدیدی با بعدی متنظر با تعداد کلاس‌های طبقه‌بندی منتقل می‌شوند. در این تحقیق از طبقه‌بندی کننده‌ی k نزدیکترین همسایه‌ی وزندار برای طبقه‌بندی فضای ویژگی استفاده شده است. این راهکار در دو داده‌ی ابرطیفی پیاده‌سازی شده و به طور متوسط بهبود دقت ۶ درصدی را نسبت به باندهای طیفی و سایر زیر مجموعه‌های تلفیق ویژگی از روش پیشنهادی نشان داده است. کسب دقت کلی تا ۹۹ درصد و همچنین تفکیک پذیری بیشتر کلاس‌های با داده‌های آموزشی اندک از ویژگی‌های این روش محسوب می‌شود.

واژگان کلیدی: تصویر ابرطیفی، طبقه‌بندی، نمایش تنک، استخراج ویژگی

* نویسنده رابط

۱- مقدمه

بهره‌گیری از فن‌آوری سنجش از دور و انواع تصاویر ماهواره‌ای در طی سال‌های اخیر به عنوان یکی از مهم‌ترین منابع جمع‌آوری اطلاعات به منظور مطالعه و بررسی منابع زمینی و بهره‌برداری بهینه از آن‌ها شناخته شده است. این فن‌آوری، توجه بسیاری از کارشناسان و متخصصان علوم مختلف از جمله زمین‌شناسی، معدن، محیط‌زیست، هواشناسی، کشاورزی، هیدرولوژی و مواردی از این دست را به خود جلب نموده است. از سوی دیگر، امروزه امکان طراحی و ساخت سنجنده‌هایی با توان تفکیک مکانی و طیفی مطلوب به منظور اخذ داده‌هایی دقیق‌تر و با ارزش اطلاعاتی بیشتری فراهم شده است. یکی از انواع این داده‌ها که از قابلیت تفکیک طیفی بسیار زیادی برخوردار است، تصاویر ابرطیفی می‌باشند.

در تصویربرداری ابرطیفی باندهای طیفی باریک و نسبتاً پیوسته‌ای در محدوده‌ی ۴۰۰ تا ۲۵۰۰ نانومتر ثبت می‌گردد. از آنجا که مواد و پدیده‌ها پاسخ طیفی منحصر به فردی را در طیف الکترومغناطیسی برخوردارند؛ پردازش تصاویر ابرطیفی امکان شناسایی مواد و پدیده‌ها را با اتکا به تمایز پاسخ طیفی آن‌ها فراهم می‌آورد.

از جمله کاربردهای تصویربرداری ابرطیفی می‌توان به استخراج اطلاعات از طریق روش‌هایی مانند آشکارسازی ناهنجاری^۱، آشکارسازی هدف^۲، طبقه‌بندی^۳ و تجزیه طیفی به اجزا خالص^۴ اشاره داشت [۱]. طبقه‌بندی یکی از پرکاربردترین و مهم‌ترین روش‌های استخراج اطلاعات از تصاویر است [۲]. روش‌های طبقه‌بندی به دنبال کشف کلاسی هستند که پیکسل موردنظر را با حداکثر اطمینان به آن نسبت داده و یا درصد تعلق یک پیکسل به یک کلاس را تشخیص دهند.

روش‌های طبقه‌بندی به‌طور مرسوم به دو دسته طبقه‌بندی‌های نظارت شده^۵ و نظارت نشده^۶ تقسیم می‌شوند. روش‌های نظارت شده به اطلاعات اولیه‌ای نظیر تعداد کلاس‌ها، خصوصیات آن‌ها و همچنین تعدادی نمونه

ی معلوم از هر کلاس نیاز دارند. در مقابل، روش‌های نظارت نشده به نمونه‌های معلوم نیاز نداشته و بر اساس مقادیر خود پیکسل‌ها در مورد خوشه‌بندی آن‌ها تصمیم‌گیری می‌کنند.

طبقه‌بندی موفق داده‌های سنجش از دوری یک چالش بزرگ است. چراکه عواملی همچون انتخاب مناسب داده‌های ورودی و روش‌های پیش‌پردازش و پردازش تصاویر ممکن است منجر به تولید نتایج متفاوتی گردد. علیرغم پتانسیل زیاد داده‌های ابرطیفی در تولید نقشه‌های طبقه‌بندی، مسائلی مانند: اول، بعد زیاد فضای ویژگی و افزونگی بیشتر در داده [۳]، دوم، تمایز بعد واقعی داده‌ها با بعد فضای ویژگی [۴]، سوم، وجود باندهای طیفی نویز آلود [۵]، چهارم، وجود باندهای طیفی که توان تفکیک مناسب برای کلاس‌های طبقه‌بندی ندارند [۶]، پنجم، ناهم‌آهنگی بین ابعاد زیاد تصویر و تعداد محدود داده‌های آموزشی در طبقه‌بندی [۷] و ششم، تنوع بیشتر در روش‌های طبقه‌بندی، از جمله چالش‌های موجود در پردازش داده‌های ابرطیفی محسوب می‌شود.

تاکنون راهکارهای توسعه یافته در مورد طبقه‌بندی دو رویکرد اصلی: اول، افزایش تفکیک پذیری فضای ویژگی از طریق انتخاب یا تولید ویژگی و دوم، توسعه و بکارگیری تکنیک‌های مؤثر طبقه‌بندی را در دستور کار خود داشته‌اند. انتخاب و یا استخراج ویژگی از تکنیک‌هایی است که برای بهبود تفکیک‌پذیری ویژگی‌های ورودی به الگوریتم طبقه‌بندی به کار می‌رود. در حالت نخست، طی فرآیندی زیرمجموعه‌ی مناسبی از ویژگی‌های اولیه (عمدتاً باندهای طیفی) انتخاب و مابقی ویژگی‌ها کنار گذاشته می‌شوند. مزیت این روش حفظ تعبیر فیزیکی ویژگی انتخاب شده است. هدف فرآیند انتخاب ویژگی، انتخاب مؤثرترین زیرمجموعه از ویژگی‌های موجود به منظور استفاده در طبقه‌بندی بر طبق یک تابع معیار و الگوریتم جستجو است [۸]. در رویکرد استخراج ویژگی، مجموعه‌ای از تبدیلات خطی و غیرخطی برای محاسبه‌ی ویژگی‌های جدید استفاده می‌شود. تبدیلاتی خطی همچون LDA^۷ و PCA^۸ [۹] و یا غیرخطی مانند NLPCA^۹ و KPCA^{۱۰} [۱۰] نمونه‌ایی از این

^۱ Anomaly Detection

^۲ Target detection

^۳ Classification

^۴ Unmixing

^۵ Supervised

^۶ Unsupervised

^۷ Linear Discriminant Analysis (LDA)

^۸ Principal Component Analysis (PCA)

^۹ Nonlinear Principal Component Analysis (NLPCA)

^{۱۰} Kernel Principal Component Analysis (KPCA)

در روش SRC معرفی کردند. در این روش، اطلاعات مکانی به صورت مجموعه‌ای از پیکسل‌های همسایه برای پیکسل مرکزی در یک پنجره مربعی با اندازه ثابت تعریف می‌شود. فرض بر این است که پیکسل‌های مربوط به یک پنجره‌ی کوچک به احتمال زیاد متعلق به اجزا خالص مشابهی بوده و می‌توان با مجموعه یکسانی از اتم‌ها در یک دیکشنری، اما با مجموعه ضرایب مختلف آن‌ها را بازسازی نمود [۱۴، ۱۵]. از محدودیت‌های این روش، تصمیم‌گیری درباره اندازه‌ی پنجره بهینه برای صحنه‌های مختلف تصویربرداری است. به عنوان مثال، یک پنجره با اندازه‌ی کوچک برای پیکسل‌های نزدیک لبه مناسب بوده، درحالی که در نواحی یکنواخت، پنجره‌های بزرگ برای طبقه‌بندی دقیق ترجیح پذیر خواهد بود.

Fang و همکاران بمنظور بهره‌مندی از مزایای اندازه‌های مختلف پنجره‌های مکانی برای طبقه‌بندی، مدل MASR^۷ را پیشنهاد کردند. در مقایسه با روش JSRC، این مدل بهبود آشکاری را از لحاظ صحت طبقه‌بندی به دست آورد. با این حال، این روش به دلیل فرآیند چند مقیاسه بودن پنجره بسیار زمان‌بر بود. [۱۶].

Song و همکاران برای استفاده کامل از اطلاعات مکانی غیرمحمولی پیکسل‌های تصویر و بهبود عملکرد JSRC، روش جدید طبقه‌بندی تصاویر ابرطیفی با استفاده از JSRC مبتنی بر KNN^۸ را پیشنهاد کردند. در این روش با استفاده از تبدیل PCA ابعاد داده‌ی اولیه کاهش یافته و یک فضای ویژگی بر اساس مؤلفه‌های اصلی تصویر و مختصات مکانی پیکسل‌ها تعریف می‌گردد. سپس در این فضای ویژگی، K نمونه از همسایه‌های هر پیکسل با روش جستجوی KNN شناسایی شده و از طریق نمایش تنک مشترک، تمامی K پیکسل برچسب‌دهی می‌شوند. از مزایای این روش عدم نیاز به تصمیم‌گیری درباره اندازه بهینه پنجره است [۱۷]. از محدودیت‌هایی این روش می‌توان به استفاده KNN از فاصله اقلیدسی به عنوان معیار شناسایی همسایگان و یکسان بودن وزن هر ویژگی اشاره داشت. چراکه ممکن است مقیاس و توزیع هر ویژگی متفاوت باشد.

روش‌های استخراج ویژگی بشمار می‌روند. در این روش‌ها عمدتاً تولید ویژگی با هدف کاهش بعد دنبال می‌شود. هر چند که در مواردی با انتقال داده‌ها به فضایی با ابعاد بیشتر، افزایش تفکیک‌پذیری دنبال می‌گردد [۱۱]. در ادامه برخی از روش‌های طبقه‌بندی متکی به استخراج ویژگی مرور شده‌اند. Giorgio و همکاران با بکاربردن ویژگی‌های تولید شده توسط تبدیلات PCA، NLPCA و KPCA بر روی دو داده‌ی ابرطیفی به دست آمده از سنجنده‌های Chris و HyMap و بکارگیری روش‌های طبقه‌بندی ANN^۱ و SVM^۲ به این نتیجه رسیدند که ویژگی‌های استخراج شده از روش NLPCA به طریق نظارت نشده منجر به دقت بیشتر طبقه‌بندی در مقایسه با ویژگی‌های استخراج شده از تبدیلات PCA، KPCA و همچنین داده خام اولیه شده است [۱۰]. مزیت روش NLPCA نسبت به PCA و KPCA در این مقاله حفظ محتوای اطلاعاتی بیشتر و همچنین توزیع یکنواخت اطلاعات بین ویژگی‌های مستخرج، گزارش شده است [۱۲].

یکی دیگر از روش‌های تولید ویژگی، روش‌های مبتنی بر نمایش تنک (SR)^۳ است. این تکنیک بر این فرض استوار بوده که می‌توان پیکسل‌های مورد آزمون را با تعداد معدودی اتم (جزء خالص) از یک دیکشنری از اتم‌ها بازنمایی نمود. بر این اساس، Bhuvaneswari و همکاران با مطالعه‌ی جامع در مورد روش طبقه‌بندی بکمک بازیابی تنک سیگنال (SRC)^۴، بهبود دقت این روش را نسبت به روش‌های سنتی مانند SVM نشان دادند. سپس برای بیشتر کردن دقت روش SRC زمانیکه تعداد محدود نمونه‌ی آموزشی موجود است؛ روش KS^۵ را پیشنهاد نموده‌اند. در این روش، داده‌ها با استفاده از تابع نگاشت غیرخطی به فضایی با ابعاد بیشتر با جدایی پذیری زیاد نگاشت می‌شوند. سپس در این فضای جدید با استفاده از تکنیک نمایش تنک ویژگی‌های جدیدی برای ورود به طبقه‌بندی استخراج می‌گردد [۱۳].

Chen و همکاران مدل نمایش تنک مشترک^۶ (JSRC) در طبقه‌بندی تصاویر ابرطیفی را با تلفیق اطلاعات مکانی

^۱ Artificial Neural Network (ANN)

^۲ Support Vector Machine (SVM)

^۳ Sparse Representation (SR)

^۴ Sparse Representation Classification (SRC)

^۵ Kernel Sparse Representation (KSR)

^۶ Joint Sparse Representation Classification (JSRC)

^۷ Multiscale Adaptive Sparse Representation (MASR)

^۸ K Nearest Neighbors (KNN)

با توجه به محدودیت‌ها و مزایای روش‌های پیشین، در این تحقیق رویکردی جدید بمنظور تولید ویژگی‌های تفکیک‌پذیر برای ورود به الگوریتم طبقه‌بندی‌کننده‌ی KNN با هدف افزایش دقت طبقه‌بندی پیشنهاد شده است. عبارت بهتر، ایده‌ی به‌کارگیری ویژگی‌های تولید شده توسط تخمین تنک، تبدیلات خطی و تبدیلات غیرخطی راهکار پیشنهاد شده در این پژوهش می‌باشد. ساختار مقاله پیش رو مشتمل بر پنج بخش می‌باشد. بعد از بخش نخست بعنوان مقدمه، بخش دوم به بیان داده‌های مورد استفاده و پیش پردازش‌ها اختصاص دارد. در بخش سوم مبانی نظری ارائه شده و در بخش چهارم به تشریح روش پیشنهادی پرداخته شده است. در بخش پنجم نتایج کسب شده از روش پیشنهادی ارائه شده و مورد بحث قرار گرفته و آخرین بخش از مقاله نیز به نتیجه‌گیری و ارائه پیشنهادات برای کارهای آتی اختصاص دارد.

۲- داده‌های مورد استفاده و پیش پردازش

در این مقاله از دو تصویر ابرطیفی مشهور به‌همراه نقشه‌ی واقعیت زمینی آنها استفاده شده است. این تصاویر در تحقیقات گسترده‌ای بمنظور ارزیابی روش‌های طبقه‌بندی استفاده شده‌اند. در ادامه هر داده معرفی و روند پیش پردازش‌های مرتبط با هر کدام ارائه شده است.

۲-۱- تصویر منطقه Indian Pines

تصویر Indian pines توسط سنجنده AVIRIS در سال ۱۹۹۲ از یک سایت آزمایشی در شمال غربی ایندیانا جمع‌آوری شده است (شکل ۱). این تصویر از ۲۲۴ باند طیفی در محدوده‌ی طول‌موج ۰/۴ تا ۲/۵ میکرومتر با ابعاد 145×145 پیکسل و حد تفکیک مکانی ۲۰ متر تشکیل شده است. پوشش این منطقه شامل دو سوم محصولات کشاورزی و یک سوم جنگل یا سایر پوشش‌های گیاهی است. از آنجا که این عکس در ماه ژوئن گرفته شده، برخی از محصولات موجود مانند ذرت و سویا در مراحل اولیه‌ی رشد با پوشش کم‌تر از ۵٪ قرار داشته‌اند. با توجه به مقدار خاشاکی که از محصولات کاشته شده از سال‌های قبل بر روی زمین باقی مانده و به دلیل متفاوت بودن خاک موجود در سطح مناطقی که گیاهان یکسان دارند، نواحی منطقه موردبررسی به سه دسته ۱- ناحیه بدون شخم، ۲- ناحیه کم

شخم و ۳- ناحیه با شخم کامل تقسیم شده‌اند. با این شرایط، در این تصویر ۱۶ کلاس از انواع محصولات کشاورزی وجود دارد (جدول ۱). برای پردازش این تصویر با حذف باندهایی که در منطقه‌ی جذب آب قرار دارند؛ تعداد باندها به ۲۰۰ باند کاهش یافته است [۱۸].

در این پژوهش مقادیر پیکسل‌های تصویر ایندیانا در بازه‌ی بین صفر و یک نرمال شده و از هر کلاس ۲۵٪ نمونه‌ها به عنوان داده‌ی آموزشی استفاده شده است.



شکل ۱- تصویر نمونه باند داده‌ی ایندیانا

جدول ۱- نام و تعداد نمونه‌های آزمایشی و تست هر کلاس داده ایندیانا

کلاس	نام	آزمون	تست
۱	یونجه	۱۲	۳۴
۲	ذرت-بدون شخم	۳۵۷	۱۰۷۱
۳	ذرت-کم شخم	۲۰۸	۶۲۲
۴	ذرت	۵۹	۱۷۸
۵	سبزه-علفزار	۱۲۱	۳۶۲
۶	سبزه-درختان	۱۸۳	۵۴۷
۷	سبزه-علفزار-درو شده	۷	۲۱
۸	گاه و خاشاک	۱۲۰	۳۵۸
۹	جودوسر	۵	۱۵
۱۰	سویا-بدون شخم	۲۴۳	۷۲۹
۱۱	سویا-کم شخم	۶۱۴	۱۸۴۱
۱۲	سویا-شخم کامل	۱۴۸	۴۴۵
۱۳	گندم	۵۱	۱۵۴
۱۴	جنگل	۳۱۶	۹۴۹
۱۵	ساختمان-سبزه-درخت	۹۷	۲۸۹
۱۶	سنگ-فولاد-برج	۲۳	۷۰

۲-۲- تصویر دانشگاه Pavia

تصویر دانشگاه پاویا در سال ۲۰۰۱ توسط سنجنده نوری ROSIS از محوطه‌ی دانشگاه پاویا در ایتالیا اخذ شد (شکل ۲). این تصویر حاوی ۱۱۵ باند طیفی در محدوده‌ی طول‌موج ۰/۴۳ تا ۰/۸۶ میکرومتر، ابعاد 610×340 پیکسل، توان تفکیک مکانی ۱/۳ متر و شامل ۹ کلاس است (جدول ۲).

ها به فضایی نگاشت شده که علاوه بر کاهش ویژگی از تفکیک پذیری بهتری برخوردار می‌شوند. به عبارت بهتر، این تبدیل داده‌ها را از فضای اولیه به فضایی نگاشت می‌دهد که در آن مقدار پراکندگی بین کلاس‌ها افزایش و از پراکندگی داخل کلاس‌ها کاسته می‌شود. در LDA ماتریس‌های پراکندگی درون کلاسی (S_w) و بین کلاسی (S_b) به شکل زیر تعریف می‌شوند.

$$S_b = \sum_{k=1}^{n_c} n_k (\mu_k - \mu)(\mu_k - \mu)^T \quad (1)$$

$$S_w = \sum_{k=1}^{n_c} \sum_{i=1}^{n_k} (x_{i,k} - \mu_k)(x_{i,k} - \mu_k)^T \quad (2)$$

در روابط فوق μ_k ، n_k ، μ و $x_{i,k}$ به ترتیب بردار میانگین کلاس k ام، تعداد نمونه‌های کلاس k ام، میانگین کل نمونه‌های آموزشی و نمونه i ام در کلاس k ام بوده و همچنین تعداد کلاس‌ها برابر n_c می‌باشد.

هدف از تبدیل LDA یافتن راستاهایی است که با تصویر کردن نمونه‌های فضای ویژگی به آنها تابع هدف ارائه شده در رابطه‌ی (۳) بیشینه گردد.

$$J(w) = \frac{W^T S_b W}{W^T S_w W} \quad (3)$$

در این رابطه W راستای مورد نظر بمنظور بیشینه سازی تابع هدف (J) است. اثبات می‌شود که راستای مورد نظر معادل با بردار ویژه‌ی متناظر با بزرگترین مقدار ویژه-ی ماتریس $S_w^{-1} S_b$ می‌باشد. سایر بردارهای ویژه‌ی این ماتریس از نظر تفکیک پذیری در رده‌های بعدی قرار خواهند داشت [۲۰].

۳-۲- تحلیل مؤلفه اصلی غیرخطی (NLPCA)

یکی از روش‌های خطی تولید ویژگی نظارت نشده روش تحلیل مؤلفه‌های اصلی است. در این تکنیک بردار-های ویژه‌ی ماتریس کوواریانس داده‌ها بعنوان مبنای تبدیل استفاده شده و مؤلفه‌های اصلی متناظر با مقادیر ویژه‌ی بزرگتر، پراکندگی بیشتری را در خود جای می‌دهند. کاهش تدریجی محتوای اطلاعاتی در ویژگی‌های تولید شده توسط تبدیل PCA، حذف همبستگی میان داده‌ها با فرضیات خطی، نظارت نشده بودن و عدم لحاظ شدن رفتارهای پیچیده و غیرخطی در توزیع داده‌ها را می‌توان از محدودیت‌های روش PCA قلمداد نمود [۲۱]

برای پردازش این تصویر با حذف باندهای نویزی و باندهای جذب آب، آزمایش‌ها بر روی ۱۰۳ باند باقیمانده انجام شده است [۱۹].

در این تصویر نیز مقادیر پیکسل‌های تصویر پایوا در بازه‌ی بین صفر و یک نرمال شده و از هر کلاس ۱۰٪ نمونه‌ها به عنوان داده‌ی آموزشی استفاده شده است.



شکل ۲- تصویر داده‌ی دانشگاه پایوا

جدول ۲- نام و تعداد نمونه‌های آزمایشی و تست هر کلاس داده‌ی پایوا

کلاس	نام	آزمون	تست
۱	آسفالت	۶۶۳	۵۹۶۸
۲	چمن‌زار	۱۸۶۵	۱۶۷۸۴
۳	شن و سنگ‌ریزه	۲۱۰	۱۸۸۹
۴	درختان	۳۰۶	۲۷۵۸
۵	صفحات فلزی	۱۳۵	۱۲۱۰
۶	خاک ساده	۵۰۳	۴۵۲۶
۷	قیر	۱۳۳	۱۱۹۷
۸	آجر	۳۶۸	۳۳۱۴
۹	سایه	۹۵	۸۵۲

۳- مبانی نظری تحقیق

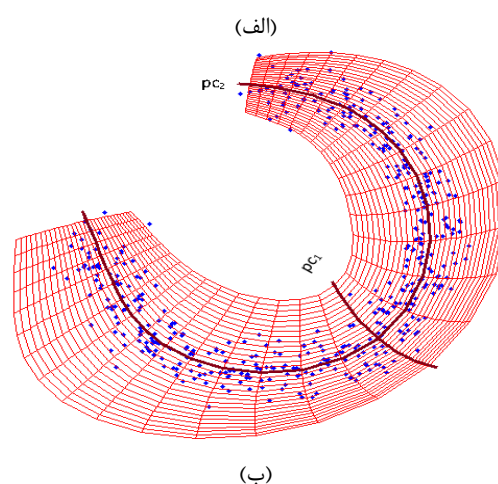
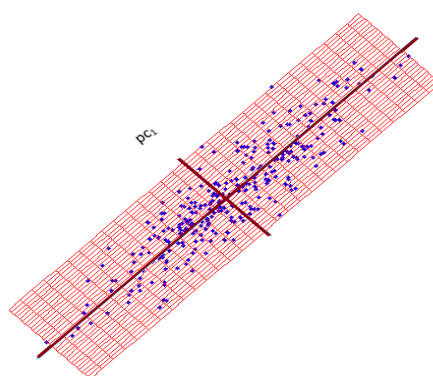
در این بخش مبانی تئوری ابزارهای استفاده شده در این تحقیق تشریح شده است. این موارد شامل: ۱- روش تبدیل مبتنی بر تحلیل تفکیک‌پذیری خطی (LDA)، ۲- تحلیل مؤلفه‌های اصلی غیرخطی (NLPCA) و ۳- تخمین تنک سیگنال می‌باشند.

۳-۱- تبدیل تحلیل تفکیک پذیری خطی (LDA)

یکی از روش‌های خطی تولید ویژگی نظارت شده که از معیار جدایی پذیری کلاس‌ها برای استخراج ویژگی استفاده می‌کند تبدیل LDA می‌باشد. در این تبدیل داده-

(شکل ۳-الف) تاکنون راهکارهای مختلفی بمنظور یافتن مؤلفه‌های اصلی در شرایط توزیع پیچیده داده‌ها در فضای ویژگی توسعه یافته است. یکی از این نسخه‌ها، تبدیل غیرخطی تحلیل مؤلفه‌ای اصلی (NLPCA) است. این تبدیل با تصویر کردن داده‌های ورودی به سطوح غیر مستوی، پیچیده و چندبعدی، تلاش می‌کند تا مؤلفه‌های غیرخطی در توزیع داده‌ها را شناسایی کند. شکل (۳) مقایسه‌ای میان تبدیل PCA و NLPCA و نحوه انطباق یافتن مؤلفه‌های اصلی در این دو تبدیل را نشان می‌دهد.

مطابق با شکل (۳-ب) تبدیل NLPCA یک تعمیم غیر خطی از تحلیل مؤلفه‌های اصلی (PCA) می‌باشد. این تکنیک همبستگی خطی و غیرخطی بین ویژگی‌ها را بدون توجه به ماهیت غیرخطی داده برآورد کرده و داده‌ها را توسط مؤلفه‌های غیرخطی از فضای اولیه به فضایی با ابعاد پایین‌تر نگاشت می‌کند.

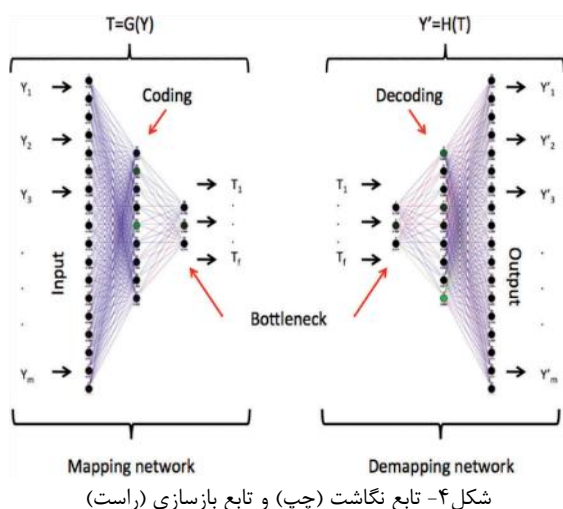


شکل ۳- (الف) تحلیل مؤلفه‌های اصلی خطی (ب) تحلیل مؤلفه‌های اصلی غیرخطی

تبدیل NLPCA توسط کرامر معرفی شده است. در این روش یک شبکه عصبی مصنوعی چند لایه بمنظور تعیین نگاشت غیرخطی داده‌ها آموزش داده می‌شود.

ورودی و خروجی این شبکه‌ی عصبی در روند آموزش یکسان است. عبارت بهتر، این شبکه‌ی عصبی مصنوعی، داده‌های ورودی را به خود آنها نگاشت می‌کند. معماری این شبکه از سه لایه ورودی، خروجی و میانی تشکیل شده، که در لایه میانی از نرون‌های کمتری نسبت به لایه‌های ورودی و خروجی استفاده می‌شود [۲۲].

به‌طور کلی این شبکه‌ی عصبی از دو نیمه تشکیل شده است. نیمه اول با استفاده از لایه‌های نگاشت (Mapping) ویژگی‌هایی با ابعاد کم‌تر استخراج می‌کند (شکل ۴).



$$T = G(Y) \quad (۴)$$

به طریقی مشابه، نیمه دوم شبکه‌ی عصبی با استفاده از لایه‌های بازیابی (Demapping) داده‌ها را از فضای ویژگی با ابعاد کمتر در لایه میانی بازسازی کرده و به ابعاد اولیه بازمی‌گرداند (شکل ۴).

$$Y' = H(T) \quad (۵)$$

در روابط (۴) و (۵)، Y و Y' به ترتیب داده‌های ورودی به شبکه عصبی و داده‌ی تخمین‌زده شده توسط آن می‌باشند. بردار مؤلفه‌های اصلی متناظر با هر مشاهده در فضای Y است. توابع G و H نیز به ترتیب تبدیلات غیرخطی نگاشت و بازیابی می‌باشند.

توپولوژی شبکه‌ی عصبی در NLPCA بنحوی انتخاب شده که در لایه‌ی خروجی، کمترین خطا در بازسازی داده‌ی اولیه تولید گردد. از مزایای روش NLPCA علاوه بر کاهش بعد، فشرده‌سازی اطلاعات و همچنین توزیع اطلاعات در مؤلفه‌های غیرخطی استخراج شده است. در این

روش، تعداد مؤلفه‌های تولیدی برابر با بعد فضای ورودی نیست. به همین دلیل در مقایسه با تبدیل PCA میزان از دست رفت اطلاعات متناسب با کیفیت برازش شبکه عصبی اندازه‌گیری می‌شود. از مشکلات این روش، مسائل مربوط به آموزش شبکه‌های عصبی و همچنین جستجو بمنظور شناسایی توپولوژی بهینه برای آن می‌باشد [۲۳].

۳-۳- روش کدگذاری تنک

کدگذاری تنک شاخه‌ای از علوم مرتبط با حسگری فشرده بوده که در آن تلاش می‌شود یک سیگنال بکمک ترکیب خطی تعداد معدودی جزء خالص از یک کتابخانه از عناصر بازنمایی گردد. فرو معین بودن دستگاه معادلات مربوط به تخمین تنک و عدم مشارکت بخش بسیار زیادی از عناصر کتابخانه در روند کدگذاری یک سیگنال از ویژگی‌های این تکنیک محسوب می‌شود [۲۴، ۲۵]. مصطلح است که بردار عناصر کتابخانه، اتم و کتابخانه، دیکشنری نامیده شوند. دیکشنری متشکل از تعداد زیادی اتم با طول واحد می‌باشد [۲۶] (رابطه ۶).

$$[D]_{b \times n} = [d_1, d_2, \dots, d_n] \quad \text{s.t. } b \ll n \quad (6)$$

در رابطه ۶، $[d_i]_{b \times 1}$ اتم‌های دیکشنری، b تعداد باندهای طیفی و n تعداد اتم‌های دیکشنری می‌باشند.

هدف از فرآیند تخمین تنک یک سیگنال، یافتن بردار تنک $[\alpha]_{n \times 1}$ از طریق حل دستگاه معادلات فرو معین ارائه شده در رابطه (۷) می‌باشد [۲۷].

$$s = D \times \alpha \quad (7) \\ \text{s.t. } \|s - D\alpha\| \leq \varepsilon \quad \text{or} \quad \|\alpha\|_0 \leq LOS$$

در رابطه ۷، s سیگنال مشاهداتی، $\|\cdot\|$ نرم دوم بردار و $\|\cdot\|_0$ نرم صفر یک بردار (شمارنده‌ی تعداد عناصر غیر صفر یک بردار) می‌باشند. LOS اصطلاحاً حداکثر میزان تنکی یک بردار و ε حد مجاز خطا در بازسازی یک سیگنال در روش تخمین تنک است.

از نقطه نظر تئوری، راهکاری صریح بمنظور حل چنین دستگاه معادلاتی (رابطه ۷) وجود نداشته و بکارگیری الگوریتم‌های ابتکاری و مبتنی بر جستجو، یکی از راهکارهای یافتن پاسخ برای این دستگاه‌های معادلات محسوب می‌شود [۲۷]. یکی از راهکارهای توسعه یافته

بمنظور تخمین تنک یک سیگنال الگوریتم OMP است. در این رویکرد، با روندی تکراری، اتم‌های سازنده‌ی یک سیگنال همزمان با تأمین قیود رابطه‌ی (۷) شناسایی و ضریب (فراوانی) هر اتم برآورد می‌گردد [۲۸].

بعبارت بهتر، در تکنیک OMP، در روندی تکراری و در هر مرحله، اتمی از دیکشنری که دارای کم‌ترین زاویه‌ی طیفی با بردار خطای تخمین سیگنال (برآورد شده توسط اتم‌های یافته شده در تکرارهای قبلی) داشته باشد، به عنوان اتم جدید به مجموعه اتم‌های انتخاب شده‌ی قبلی افزوده می‌شود. بعبارت دیگر، با در نظر گرفتن r بعنوان بردار خطای تخمین توسط اتم‌های انتخاب شده‌ی قبلی برای سیگنال s (رابطه ۸)، در هر تکرار، اتمی که $\| [d_i]^T \times [r] \|_2$ را بیشینه سازد، بعنوان اتم جدید به مجموعه اتم‌های انتخاب شده‌ی قبلی افزوده می‌شود.

$$[r] = [s] - D \times [A] \quad (8)$$

شایان ذکر است در تکرار اول و بمنظور انتخاب اولین اتم در این تکنیک، $r = s$ در نظر گرفته می‌شود [۲۹].

۴- روش پیشنهادی

در این مقاله راهکاری سه مرحله‌ای بمنظور تولید ویژگی به تعداد کلاس‌های طبقه‌بندی (n_c) پیشنهاد شده است. در ادامه، ویژگی‌های تولید شده بکمک طبقه‌بندی کننده‌ی KNN وزندار برجسبدهی می‌شوند.

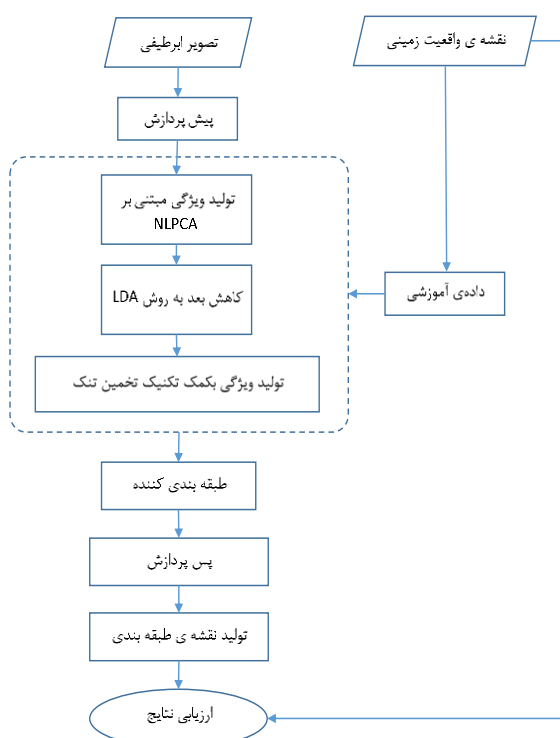
در این روند سه مرحله‌ای، ابتدا از تبدیل NLPCA بصورت نظارت شده بمنظور انتقال باندهای طیفی به فضایی با ابعاد بیشتر استفاده شده است. این مرحله برای اولین بار در مورد این تبدیل بکارگرفته شده و تحلیل مؤلفه‌ی اصلی غیرخطی نظارت شده (SNLPCA)^۱ نامگذاری شده است. در مرحله‌ی دوم، ویژگی‌های تولیدی بکمک یک تبدیل ساده‌ی LDA به فضایی با بعد $n_c - 1$ انتقال می‌یابد. در مرحله‌ی سوم، یک دیکشنری بکمک تمامی داده‌های آموزشی تولید و هر نمونه از ویژگی‌های بدست آمده از مرحله‌ی دوم بصورت تنک تخمین زده می‌شود. در ادامه، از بردار تخمین تنک هر نمونه، به تعداد کلاس‌های طبقه‌بندی ویژگی تولید می‌شود. سازوکار روش پیشنهادی در این مقاله در فلوچارت شکل ۵ ارائه شده است.

^۱ Supervised Non-Linear Principal Component Analysis

در ادامه، هر یک از بخش‌های مربوط به فلوچارت شکل ۵ به تفصیل تشریح شده است. بعد از اعمال پیش پردازش های ذکر شده در بخش ۲، اولین اقدام بمنظور تولید ویژگی، SNLPCA است. در این روند، تعداد n_c شبکه‌ی عصبی جهت آموزش یافتن به داده‌های آموزشی هر کلاس از طبقه‌بندی بکار گرفته می‌شود. برای این اقدام از توبلاکس NLPCA در محیط نرم‌افزار متلب استفاده شده است. معماری شبکه‌ی عصبی به نحوی طراحی شده که در مورد داده‌های آموزشی بیشترین انطباق سراسری تأمین گردد. لازم به ذکر است که آموزش شبکه‌ی عصبی از طریق روش پس‌انتشارخطا و مبتنی بر تکنیک گرادیان نزولی انجام می‌شود. بعد از آموزش یافتن شبکه‌های عصبی، پاسخ تمامی پیکسل‌های تصویر ابرطیفی مربوط به لایه‌ی مؤلفه های اصلی (T در شکل ۴) از تمامی شبکه‌های عصبی آموزش یافته، به عنوان ویژگی‌های این مرحله تولید می گردد. در این تحقیق ابعاد مؤلفه‌های اصلی مستخرج از شبکه‌های عصبی ثابت و معادل ۲۰ مؤلفه‌ی اصلی انتخاب شده است. بر این اساس، انتظار می‌رود که ویژگی‌های تولید شده در این مرحله معادل $20 \times n_c$ باشد. محتمل است که در این روند با افزایش تعداد کلاس‌های طبقه‌بندی بعد فضای ویژگی تولید شده بیشتر از بعد داده‌های خام ورودی باشد. هرچند که در صورت وجود کلاس‌های محدود در طبقه‌بندی و یا بکارگیری تعداد کمتری از مؤلفه‌های اصلی در معماری شبکه‌ی عصبی، ممکن است این انتقال غیرخطی منجر به تولید فضایی با بعدی بیشتر از تعداد باندهای طیفی نگردد. ایده‌ی اصلی چنین اقدامی، احتمال پاسخ متفاوت نمونه‌های مربوط به یک کلاس در هنگام استفاده از شبکه‌ی عصبی آموزش یافته برای کلاسی دیگر می‌باشد. ضمن اینکه با چنین اقدامی، بکمک یک تبدیل غیرخطی (مشابه با کرنل‌های غیرخطی) فضای جدیدی برای بازنمایی داده‌های خام ایجاد می‌گردد. بدلیل وجود مشکلات فضاهای ویژگی با ابعاد بالا^۱، دومین اقدام از روش پیشنهادی، کاهش بعد نظارت شده بکمک روش خطی LDA انتخاب شده است. این تبدیل تلاش می‌کند تا بکمک آماره‌هایی همچون ماتریس کوواریانس داخل کلاسی (S_W) و بین کلاسی (S_B)، راستاهایی را بمنظور تصویر کردن نمونه های فضای ورودی به فضایی با بیشترین تفکیک‌پذیری

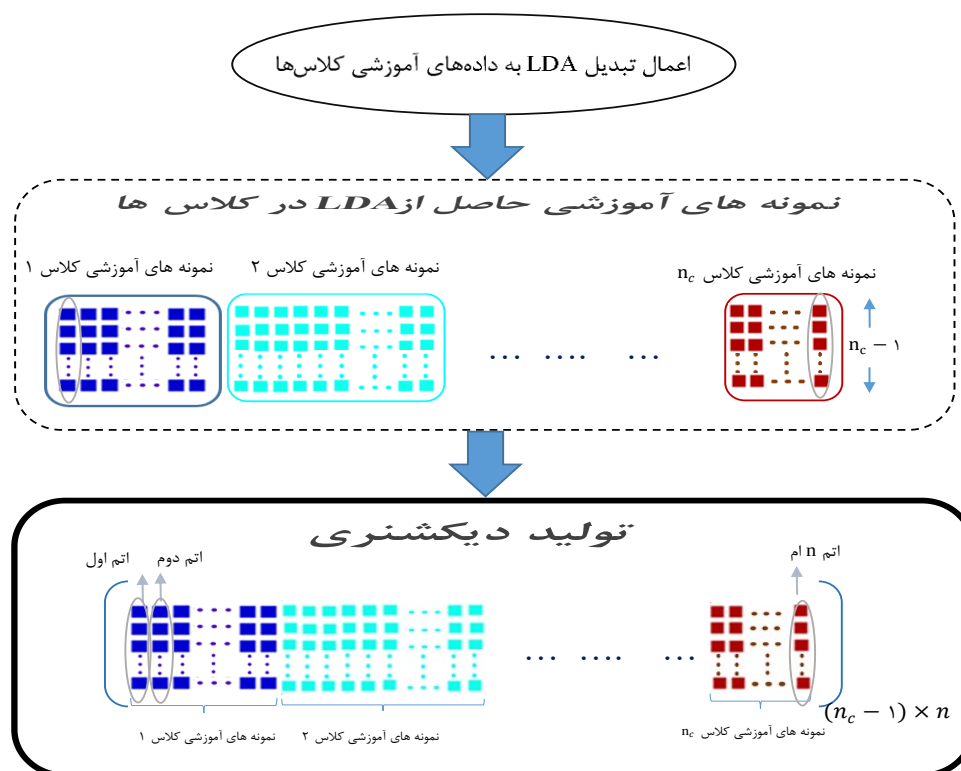
خطی بیابد. بدیهی است که داده‌های آموزشی مورد نیاز در تبدیل LDA نیز از خلال عبور نمونه‌های آموزشی در فضای باندهای طیفی از لایه‌ی نگاشت (Mapping) شبکه‌های عصبی مصنوعی (به ازای همه‌ی کلاس‌ها) ایجاد می‌گردند (شکل ۴).

در گام سوم، بکمک موقعیت نمونه‌های آموزشی و ویژگی‌های تولید شده از تبدیل LDA، یک دیکشنری از رفتار کلاس‌های طبقه‌بندی تولید می‌گردد. اتم‌های دیکشنری شامل تمامی نمونه‌های آموزشی بوده که بصورت منظم و به تفکیک کلاس، ستون‌های ماتریس دیکشنری را تولید می‌کنند (شکل ۶).



شکل ۵- الگوریتم روش پیشنهادی

^۱ Curse of Dimensionality



شکل ۶- تولید ماتریس دیکشنری

لحاظ کردن قرابت مکانی برچسب پیکسل‌ها در تصاویر ابرطیفی از جمله مواردی است که در تحقیقات گذشته به عنوان راهکار بهبود دقت نتایج طبقه‌بندی معرفی شده است. در مرحله سوم از روش پیشنهادی، با اعمال یک فیلتر میانگین در فضای مکان به ویژگی‌های مستخرج از مرحله دوم (تبدیل LDA) تلاش شده تا اثر پیکسل‌های همسایه در روند برچسب‌دهی هر پیکسل لحاظ شده و به نحوی نویزهای ناخواسته در داده‌ها کاهش یابد.

با توجه به تولید ویژگی‌های مبتنی بر کدگذاری تنک، انتظار می‌رود که در تصویر ویژگی‌های مرتبط با هر کلاس، یک کلاس با پاسخ بزرگتری از سایر کلاس‌ها ظاهر گردد. از سوی دیگر، بردار تخمین تنک مربوط به داده‌های آموزشی هر کلاس نیز صرفاً یک درایه‌ی غیرصفر متناظر با محل قرارگیری اتم متناظرش در دیکشنری خواهد داشت. در چنین شرایطی، مسأله‌ی برچسب‌دهی برای بخش زیادی از پیکسل‌ها محدود به یافتن نزدیکترین همسایه از داده‌های آموزشی خواهد بود. با اینحال بخشی از پیکسل‌ها بدلیل وجود نویز، سهم حضور اندک از کلاس‌های طبقه‌بندی و شباهت احتمالی زیاد با سایر کلاس‌ها، برچسب صحیحی از طریق نزدیک‌ترین همسایه نمی‌یابند. به همین دلیل، الگوریتم طبقه‌بندی KNN وزندار بعنوان طبقه‌بندی کننده انتخاب شده است. در این الگوریتم، تعداد K

با در نظر گرفتن n بعنوان مجموع تمامی داده‌های آموزشی و $b = n_c - 1$ بعنوان بعد اتم‌های دیکشنری (بدست آمده از تبدیل LDA)، ابعاد ماتریس دیکشنری $b \times n$ خواهد بود. بدیهی است که در چنین شرایطی $b \gg n$ است.

هر درایه از بردار α متناظر با یک اتم از دیکشنری است (رابطه‌ی ۷). متناسب با نحوه‌ی چیدمان داده‌های آموزشی در ماتریس دیکشنری (شکل ۶)، می‌توان درایه‌های متناظر از بردار α مربوط به داده‌های آموزشی هر کلاس را شناسایی کرد. بعد از بکارگیری تکنیک OMP بمنظور تخمین تنک هر پیکسل، در صورتیکه درایه‌های متناظر با هر کلاس از بردار α مقداردهی شوند؛ ماکزیمم مقادیر درایه‌های مربوط به هر کلاس بعنوان ویژگی مربوط به آن کلاس تولید می‌گردد.

بعبارت بهتر، تعداد ویژگی‌های مستخرج از روش تخمین تنک برابر با تعداد کلاس‌های طبقه‌بندی است. هر ویژگی متناظر با یک کلاس از طبقه‌بندی بوده که مقادیر ثبت شده در آن، ماکزیمم درایه‌های مرتبط با آن کلاس از بردار تنک α می‌باشد. بدیهی است که در صورت عدم انتخاب هیچکدام از داده‌های آموزشی یک کلاس در در روند تخمین تنک، مقدار صفر در ویژگی متناظر با آن کلاس برای آن پیکسل ثبت خواهد شد.

نزدیکترین همسایه از داده‌های آموزشی برای هر پیکسل در فضای ویژگی شناسایی می‌شود. معکوس فاصله میان آن پیکسل و همسایگان یافت شده بعنوان وزن در نظر گرفته می‌شود. وزن کلاس‌های مشابه تجمیع شده و نهایتاً برچسب کلاسی که بیشترین وزن تجمعی را کسب کرده باشد به آن پیکسل اختصاص می‌یابد.

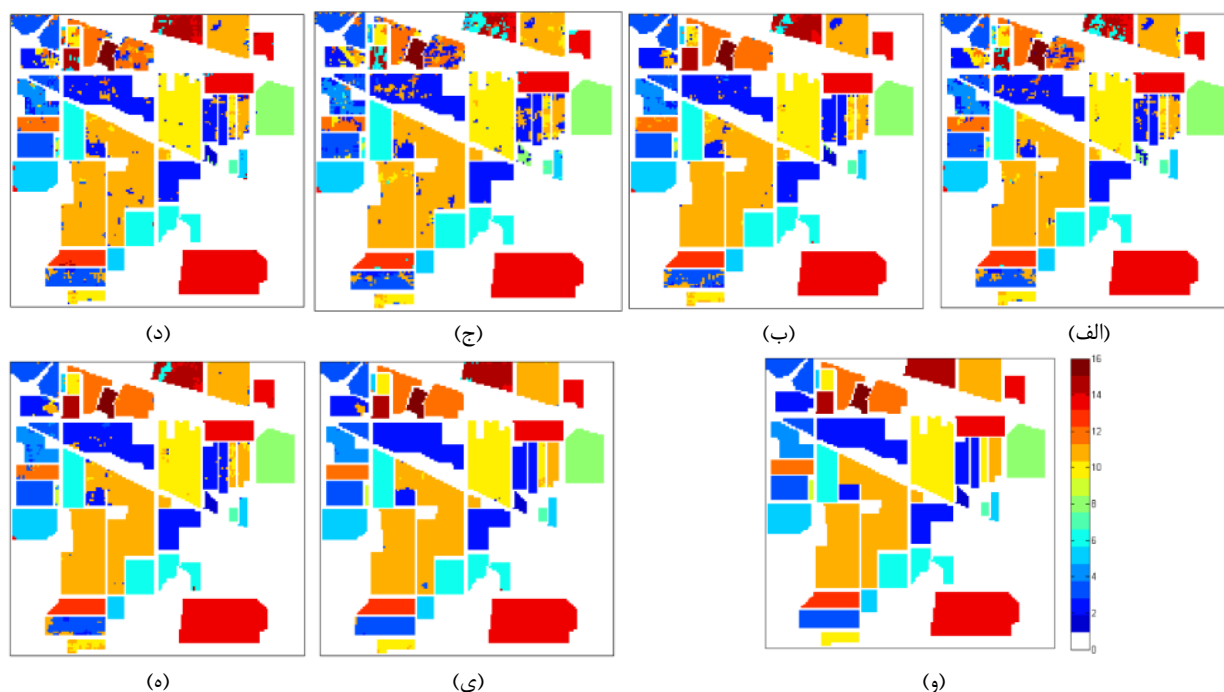
آخرین گام همانند تمام روش‌های طبقه‌بندی، استفاده از پس پردازش‌های رایجی همچون فیلتر اکثریت بمنظور کاهش نویز در نتایج خواهد بود.

بمنظور ارزیابی و مقایسه‌ی نتایج در این تحقیق، از ۵ ترکیب مختلف (به غیر از ترکیب پیشنهاد شده در این مقاله) از ویژگی‌های تولید شده در این مقاله بعنوان ورودی فرایند طبقه‌بندی استفاده شده است. عبارت بهتر، به غیر از ترکیب پیشنهاد شده، از: ۱- باندهای طیفی بطور مستقیم، ۲- ویژگی‌های مستخرج از اعمال تبدیل LDA به باندهای طیفی، ۳- ویژگی‌های مستخرج از روند نظارت شده‌ی

NLPCA، ۴- ویژگی‌های مستخرج از اعمال تبدیل LDA به مؤلفه‌های اصلی بدست آمده از تبدیل SNLPCA و ۵- بکارگیری مستقیم روش استخراج ویژگی بکمک کدگذاری تنک، بعنوان ورودی‌های طبقه‌بندی استفاده شده که نتایج هریک در بخش نتایج، ارائه و مورد بحث قرار گرفته است.

۴- ارزیابی نتایج

برای ارزیابی کارایی الگوریتم پیشنهادی در این تحقیق، از دو داده‌ی واقعی ایندیانا و پاولیا (معرفی شده در بخش ۲) استفاده شده است. برای این هدف، طبقه‌بندی KNN وزن‌دار علاوه بر داده خام در پنج سری ویژگی استخراج شده از داده‌های مذکور پیاده‌سازی شده است. شکل ۷ و ۸ نقشه‌ی طبقه‌بندی بدست آمده از باندهای طیفی و ویژگی‌های استخراج شده از آن به همراه نقشه‌ی واقعیت زمینی را در دو مجموعه داده نشان می‌دهد.

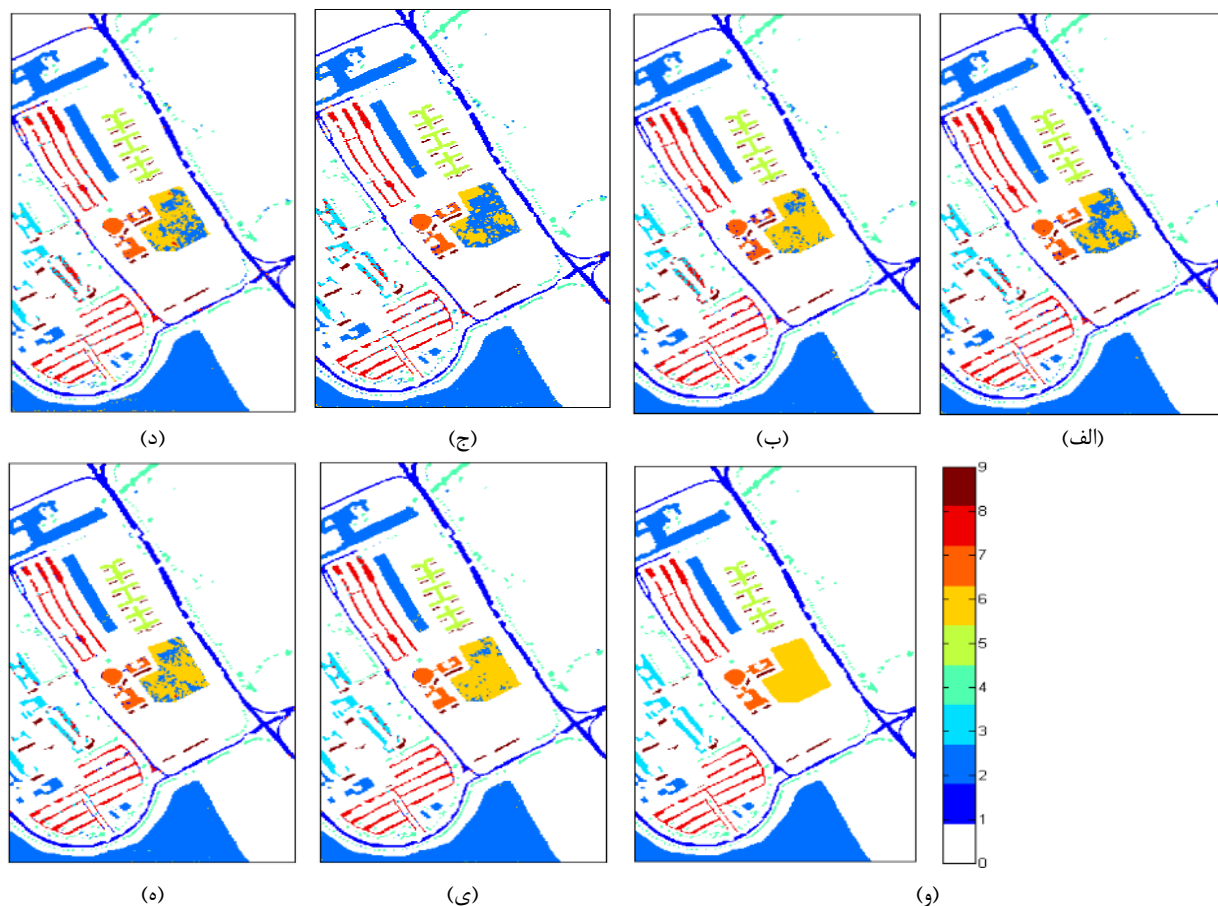


شکل ۷- نتایج طبقه‌بندی داده‌های Indiana Pine؛ (الف) داده خام، (ب) LDA، (ج) SNLPCA، (د) SRC، (ه) SNLPCA+LDA، (ی) الگوریتم پیشنهادی، (و) نقشه واقعیت زمینی

دقت کاربر^۱، دقت کلی^۲ و ضریب کاپا^۳ استفاده شده است [۳۰]. جداول ۳ و ۴ بیانگر این پارامترها در تصاویر ایندیانا و پاولیا می‌باشند.

مقایسه‌ی بصری نتایج در شکل ۷ و ۸ نشان دهنده‌ی بهبود نتیجه‌ی طبقه‌بندی روش پیشنهادی در مقایسه با روش‌های دیگر استخراج ویژگی می‌باشد. در این بین استخراج ویژگی به صورت نظارت شده در روش NLPCA، طبقه‌بندی ضعیف‌تری را تولید کرده است. در ادامه، به منظور ارزیابی و مقایسه کمی نتایج، از پارامترهای آماری

^۱ Producers Accuracy
^۲ Overall Accuracy
^۳ Kappa Coefficient



شکل ۸- نتایج طبقه‌بندی داده‌های دانشگاه Pavia؛ (الف) داده خام، (ب) LDA، (ج) SNLPCA، (د) SRC، (ه) SNLPCA+LDA، (ی) الگوریتم پیشنهادی، (و) نقشه واقعیت زمینی

جدول ۳- نتایج طبقه‌بندی تصویر ایندیانا حاصل از روش‌های تولید ویژگی

روش پیشنهادی	SNLPCA+LDA	SPARSE	SNLPCA	LDA	داده خام	کلاس
۱۰۰	۱۰۰	۸۶/۹۶	۳۰/۴۳	۸۹/۱۳	۴۷/۸۳	۱
۹۸/۴۶	۹۳/۹۱	۸۹/۳۶	۸۳/۲۶	۹۳/۹۸	۸۷/۱۱	۲
۹۹/۰۴	۹۰/۱۲	۸۷/۷۱	۸۰/۳۶	۸۷/۸۳	۸۵/۰۶	۳
۹۸/۳۱	۹۰/۳	۷۴/۲۶	۶۷/۰۹	۸۴/۸۱	۷۴/۲۶	۴
۹۹/۷۹	۹۸/۵۵	۹۴/۸۲	۹۵/۲۴	۹۸/۱۴	۹۶/۰۷	۵
۹۹/۸۶	۹۹/۷۳	۹۹/۷۳	۱۰۰	۹۹/۸۶	۱۰۰	۶
۱۰۰	۱۰۰	۱۰۰	۹۶/۴۳	۹۶/۴۳	۹۶/۴۳	۷
۱۰۰	۱۰۰	۹۹/۷۹	۱۰۰	۱۰۰	۱۰۰	۸
۱۰۰	۸۰/۰۰	۵۰/۰۰	۲۵/۰۰	۴۵/۰۰	۳۵/۰۰	۹
۹۸/۳۵	۹۰/۹۵	۹۲/۷۰	۹۱/۵۶	۹۱/۹۸	۹۲/۸۰	۱۰
۹۹/۲۳	۹۸/۱۷	۹۴/۶۲	۹۲/۷۵	۹۵/۸۹	۹۵/۰۷	۱۱
۹۹/۸۳	۹۶/۶۳	۸۱/۱۱	۶۲/۳۹	۹۶/۲۹	۷۶/۰۵	۱۲
۱۰۰	۱۰۰	۹۵/۶۱	۹۷/۰۷	۱۰۰	۹۹/۰۲	۱۳
۱۰۰	۹۹/۹۲	۹۹/۱۳	۹۹/۱۳	۹۹/۹۲	۹۹/۷۶	۱۴
۹۲/۲۳	۷۹/۰۲	۷۷/۹۸	۳۸/۶۰	۷۵/۳۹	۴۷/۶۷	۱۵
۹۸/۹۲	۹۸/۹۲	۹۸/۹۲	۱۰۰	۹۸/۹۲	۹۷/۸۵	۱۶
۹۹/۰۰	۹۵/۷	۹۲/۳۹	۸۷/۴۳	۹۴/۶۵	۹۰/۵۶	OA
۹۸/۸۵	۹۵/۰۸	۹۱/۳۰	۸۵/۵۹	۹۳/۸۹	۸۹/۱۸	Kappa

جدول ۴- نتایج طبقه بندی تصویر پاپویا حاصل از روش های تولید ویژگی

کلاس	داده خام	LDA	SNLPCA	SPARSE	SNLPCA+LDA	روش پیشنهادی
۱	۹۶/۵۰	۹۸/۸۱	۹۶/۷۹	۹۲/۹۴	۹۶/۹۸	۹۸/۷۶
۲	۹۹/۶۸	۹۹/۷۳	۹۹/۵۴	۹۹/۵۹	۹۹/۷۶	۹۹/۷۴
۳	۸۶/۴۷	۷۹/۷۵	۷۹/۲۸	۷۸/۴۷	۸۹/۷۶	۹۴/۹۵
۴	۸۷/۵۰	۹۲/۴۹	۸۳/۸۱	۸۹/۵۹	۹۶/۵۴	۹۶/۹۰
۵	۱۰۰	۱۰۰	۱۰۰	۱۰۰	۱۰۰	۱۰۰
۶	۵۶/۹۷	۸۳/۸۳	۴۱/۳۸	۶۳/۷۳	۷۳/۲۶	۸۹/۷۴
۷	۹۲/۹۳	۸۳/۱۶	۸۹/۶۲	۹۱/۲۰	۹۲/۶۳	۹۷/۴۴
۸	۹۲/۷۸	۹۳/۸۶	۹۰/۳۶	۹۲/۳۱	۹۰/۸۷	۹۶/۴۴
۹	۹۹/۸۹	۹۹/۵۸	۹۹/۸۹	۹۹/۸۹	۹۹/۸۹	۹۹/۸۹
OA	۹۱/۸۶	۹۵/۲۱	۸۹/۰۸	۹۱/۷۲	۹۴/۵۲	۹۷/۶۳
Kappa	۸۸/۹۳	۹۳/۵۷	۸۵/۰۲	۸۸/۸۱	۹۲/۶۳	۹۶/۸۴

مطابق با موارد ذکر شده در بخش دوم، در تصویر ایندیانا به دلیل باقی ماندن خاشاک گیاهان کاشته شده در سال های قبل، ناهنجاری های طیفی اندکی در مزارع با کشت مشابه ثبت شده است. این موضوع در برخی از زمین های کشاورزی باعث کاهش دقت طبقه بندی شده است. با اینحال بکارگیری ویژگی های پیشنهاد شده در این مقاله حتی در این موارد نیز توانسته در مقایسه با سایر ترکیبات ویژگی، دقت های بهتری را در نتایج طبقه بندی کسب نماید. از سوی دیگر، روش پیشنهادی توانسته دقت های حداکثری برای کلاس های با داده ی آموزشی کم (کلاس های ۱ و ۹ از داده ی ایندیانا) تولید نماید. این در حالیکه همین موارد، دقت های بسیار پایین تری را هنگام بکارگیری باندهای طیفی و سایر ترکیبات مربوط به ویژگی های تولیدی، کسب کرده اند.

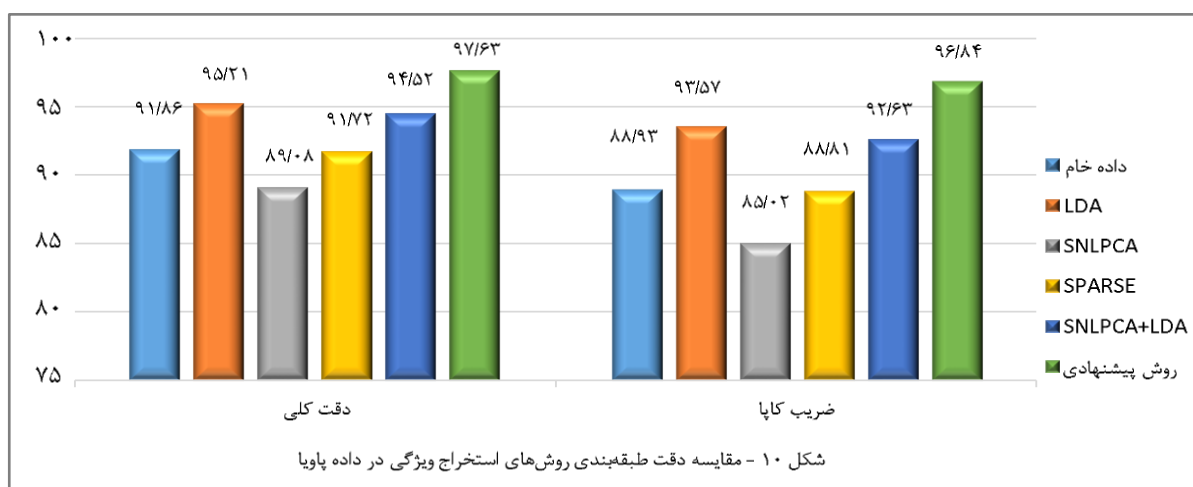
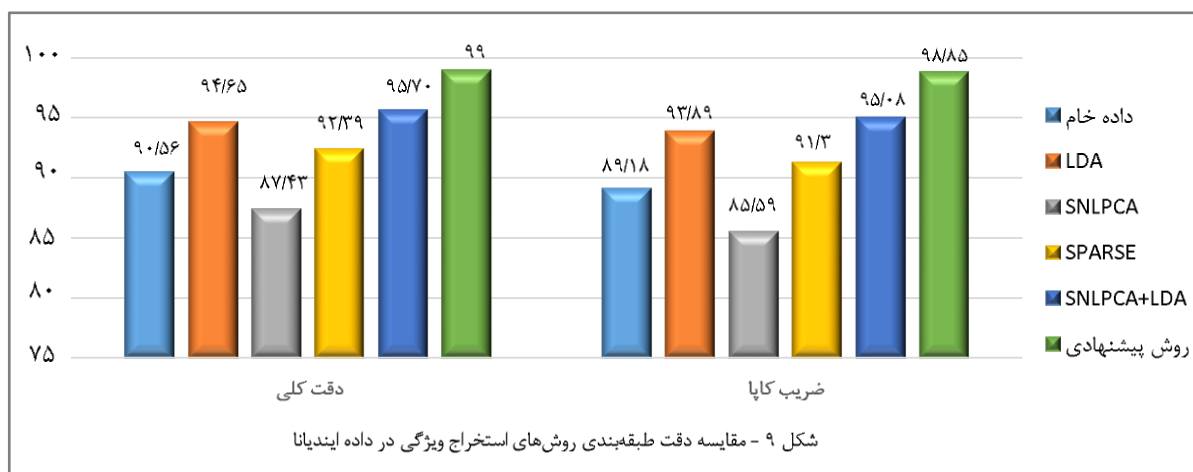
از سوی دیگر، علیرغم بکارگیری فیلتر میانگین با ابعاد پنجره ی ۳×۳ در زمان تولید ویژگی در روش تخمین تنک، کماکان دقت زیادی در برچسب دهی نمونه های واقع در لبه ی کلاس ها در این روش پیشنهادی مشاهده می شود. عبارت بهتر، اعمال فیلتر میانگین با هدف مشارکت مکانی پیکسل ها در برچسب دهی نتوانسته در مناطق مرزی کلاس ها خللی ایجاد کند و ویژگی های تولید شده کماکان برای برچسب دهی در مناطق مرزی مناسب بودند.

از دیگر نتایج قابل استنباط از جدول (۳) و (۴)، دقت اندک استفاده از ویژگی های SNLPCA در نتایج طبقه بندی است. به نظر می رسد بعد زیاد فضای ویژگی تولید شده توسط این تبدیل و احتمال هم مقیاس نبودن تمامی ویژگی های تولیدی دلایل عدم کفایت استفاده

مستقیم از تبدیل SNLPCA باشد. این ویژگی ها در بیشتر موارد دقت های پایین تری را نسبت به باندهای طیفی داشتند. با اینحال زمانیکه ویژگی های SNLPCA بکمک تبدیلات خطی، کاهش بعد می یابند، تفکیک پذیری بهتری را نسبت به داده های خام (باندهای طیفی) از خود نشان می دهند. بطور متوسط دقت کلی در زمان استفاده از روش SNLPCA+LDA در مقایسه با LDA بهبود ۴ درصدی را در مورد دو داده ی این تحقیق نشان می دهد.

نتایج مشابه کسب شده از پیاده سازی روش پیشنهادی در داده های دانشگاه پاپویا نیز موید تعمیم پذیر بودن این راهکار می باشد. عبارت بهتر، راهکار پیشنهادی نسبت به تمایز در محتوای تصویری، تعداد باندهای طیفی و همچنین بازه ی طیفی سیستم تصویربرداری، رفتار پایداری را از خود نشان داده است.

با مقایسه نتایج کمی حاصل از داده پاپویا و ایندیانا مشاهده می شود که در روش پیشنهادی بطور میانگین بهبود ۶/۰۱ درصدی نسبت به سایر روش های استفاده از ویژگی ها و باندهای طیفی حاصل شده است. شکل ۹ و ۱۰ نمودار مقایسه دقت روش ها را در هر یک از دو داده نشان می دهد.



۵- نتیجه‌گیری و پیشنهادات

در این مقاله، روشی تلفیقی بمنظور تولید ویژگی از تصاویر ابرطیفی پیشنهاد شد. در این روش نظارت شده، از تبدیلات خطی (LDA)، غیرخطی (SNLPCA) و تخمین تنک استفاده گردید. ویژگی‌های نهایی در این روش به تعداد کلاس‌های طبقه‌بندی بود. آزمون‌های مختلفی بمنظور ارزیابی کفایت این ترکیب پیشنهادی صورت گرفت. با در نظر گرفتن باندهای طیفی و همچنین سه روش SNLPCA، LDA و تخمین تنک، ۵ ترکیب از ویژگی به غیر از روش پیشنهادی بمنظور طبقه‌بندی استفاده شد. نتایج ارزیابی نشان داد که روش پیشنهادی از توان تأمین دقت‌های بیشتری در نتایج طبقه‌بندی برخوردار است. از دیگر نتایج قابل توجه، کسب دقت‌های بسیار خوب در مورد کلاس‌هایی است که تعداد داده‌های آموزشی اندکی از آنها در روند تولید ویژگی حضور دارند. این درحالی است که باندهای طیفی و همچنین اکثر

ویژگی‌های مقایسه‌ای، توان تولید نتایج دقیق را در چنین کلاس‌هایی برخوردار نبودند. پیشنهاداتی همچون: اول، بررسی اثر تغییر تعداد مؤلفه‌های اصلی تبدیل SNLPCA در نتایج، دوم، بکارگیری قرابت مکانی پیکسل‌های مجهول به مکان داده‌های آموزشی در روند تخمین تنک سیگنال‌ها، سوم، بهبود راستای اتم‌های موجود در ماتریس دیکشنری بکمک تکنیک‌های آموزش دیکشنری مثل KSVD، چهارم، بکارگیری روش‌های پایدار مبتنی بر تولید نمونه‌های تقویتی در محاسبه‌ی ماتریس کوواریانس داخل کلاسی و بین‌کلاسی تبدیل LDA، در زمان وجود بعد زیاد فضای ویژگی بمنظور بهبود عملکرد این تبدیل و پنجم، خوشه‌بندی فضای ویژگی و تولید ویژگی‌های تفکیک‌پذیر هر کلاس در هر خوشه از مواردی بشمار رفته که در ادامه‌ی این تحقیق در دست‌ورکار نویسندگان این مقاله قرار دارد.

مراجع

- [1] D.-W. Sun, *Hyperspectral imaging for food quality analysis and control*. Elsevier, 2010
- [2] J. Sauvola and M. Pietikainen, "Page segmentation and classification using fast feature extraction and connectivity analysis," in *Proceedings of 7rd International Conference on Document Analysis and Recognition*, 1995, vol. 2: IEEE, pp. 1127-1131.
- [3] R. Liu and D. F. Gillies, "Overfitting in linear feature extraction for classification of high-dimensional image data," *Pattern Recognition*, vol. 53, pp. 73-86, 2016.
- [4] B. Ghaddar and J. Naoum-Sawaya, "High dimensional data classification and feature selection using support vector machines," *European Journal of Operational Research*, vol. 265, no. 3, pp. 993-1004, 2016
- [5] A. Mahmood, A. Robin, and M. Sears, "Modified residual method for the estimation of noise in hyperspectral images," *IEEE Transactions on Geoscience and Remote Sensing*, vol. 55, no. 3, pp. 1451-1460, 2016.
- [6] H. Li, Z. Ye, and G. Xiao, "Hyperspectral image classification using spectral-spatial composite kernels discriminant analysis," *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing*, vol. 8, no. 6, pp. 2341-2350, 2014.
- [7] S. A. Medjahed, T. A. Saadi, A. Benyettou, and M. Ouali, "Gray wolf optimizer for hyperspectral band selection," *Applied Soft Computing*, vol. 40, pp. 178-186, 2016.
- [8] S. Li, H. Wu, D. Wan, and J. Zhu, "An effective feature selection method for hyperspectral image classification based on genetic algorithm and support vector machine," *Knowledge-Based Systems*, vol. 24, no. 1, pp. 40-48, 2011.
- [9] M. Kamandar and H. Ghassemian, "Linear feature extraction for hyperspectral images based on information theoretic learning," *IEEE Geoscience and Remote Sensing Letters*, vol. 10, no. 4, pp. 702-706, 2012.
- [10] G. Licciardi, P. R. Marpu, J. Chanussot, and J. A. Benediktsson, "Linear versus nonlinear PCA for the classification of hyperspectral data based on the extended morphological profiles," *IEEE Geoscience and Remote Sensing Letters*, vol. 9, no. 3, pp. 447-451, 2011.
- [11] H. Chen, H. Zhao, J. Shen, R. Zhou, and Q. Zhou, "Supervised machine learning model for high dimensional gene data in colon cancer detection," in *2015 IEEE International Congress on Big Data*, 2015: IEEE, pp. 134-141.
- [12] G. Licciardi, F. Del Frate, and R. Duca, "Feature reduction of hyperspectral data using autoassociative neural networks algorithms," in *2009 IEEE International Geoscience and Remote Sensing Symposium*, 2009, vol. 1: IEEE, pp. I-176-I-179.
- [13] N. Bhuvaneshwari and V. Sivakumar, "A comprehensive review on sparse representation for image classification in remote sensing," in *2016 International Conference on Communication and Electronics Systems (ICCES)*, 2016: IEEE, pp. 1-4.
- [14] Y. Chen, N. M. Nasrabadi, and T. D. Tran, "Hyperspectral image classification using dictionary-based sparse representation," *IEEE transactions on geoscience and remote sensing*, vol. 49, no. 10, pp. 3973-3985, 2011.
- [15] B. Tu, X. Zhang, X. Kang, G. Zhang, J. Wang, and J. Wu, "Hyperspectral image classification via fusing correlation coefficient and joint sparse representation," *IEEE Geoscience and Remote Sensing Letters*, vol. 15, no. 3, pp. 340-344, 2018.
- [16] L. Fang, S. Li, X. Kang, and J. A. Benediktsson, "Spectral-spatial hyperspectral image classification via multiscale adaptive sparse representation," *IEEE Transactions on Geoscience and Remote Sensing*, vol. 52, no. 12, pp. 7738-7749, 2014.
- [17] W. Song, S. Li, X. Kang, and K. Huang, "Hyperspectral image classification based on KNN sparse representation," in *2016 IEEE international geoscience and remote sensing symposium (IGARSS)*, 2016: IEEE, pp. 2411-2414.
- [18] W. Zhang, X. Li, and L. Zhao, "Band priority index: A feature selection framework for hyperspectral imagery," *Remote Sensing*, vol. 10, no. 7, p. 1095, 2018.
- [19] C. Chen, N. Chen, and J. Peng, "Nearest regularized joint sparse representation for hyperspectral image classification," *IEEE Geoscience and Remote Sensing Letters*, vol. 13, no. 3, pp. 424-428, 2016.
- [20] Q. Wang, Z. Meng, and X. Li, "Locality adaptive discriminant analysis for spectral-spatial classification of hyperspectral images," *IEEE Geoscience and Remote Sensing Letters*, vol. 14, no. 11, pp. 2077-2081, 2017.

- [21] F. P. Shah and V. Patel, "A review on feature selection and feature extraction for text classification," in 2016 International Conference on Wireless Communications, Signal Processing and Networking (WiSPNET), 2016: IEEE, pp. 2264-2268.
- [22] M. A. Kramer, "Nonlinear principal component analysis using autoassociative neural networks," AIChE journal, vol. 37, no. 2, pp. 233-243, 1991.
- [23] G. Licciardi and J. Chanussot, "Spectral transformation based on nonlinear principal component analysis for dimensionality reduction of hyperspectral images," European Journal of Remote Sensing, vol. 51, no. 1, pp. 375-390, 2018.
- [24] R. Roscher and B. Waske, "Shapelet-based sparse image representation for landcover classification of hyperspectral data," in 2014th IAPR Workshop on Pattern Recognition in Remote Sensing, 2014: IEEE, pp. 1-6.
- [25] M. Elad, Sparse and redundant representations: from theory to applications in signal and image processing. Springer Science & Business Media, 2010.
- [26] S. Huang, H. Zhang, and A. Pižurica, "A robust sparse representation model for hyperspectral image classification," Sensors, vol. 17, no. 9, p. 2087, 2017.
- [27] J. A. Tropp and S. J. Wright, "Computational methods for sparse solution of linear inverse problems," Proceedings of the IEEE, vol. 98, no. 6, pp. 948-958, 2010.
- [28] N. Aravind, K. Abhinandan, V. V. Acharya, and S. S. David, "Comparison of OMP and SOMP in the reconstruction of compressively sensed hyperspectral images," in 2011 International Conference on Communications and Signal Processing, 2011: IEEE, pp. 188-192.
- [29] S. Soofbaf, M. Sahebi, and B. Mojaradi, "A sliding window-based joint sparse representation (swjsr) method for hyperspectral anomaly detection," Remote Sensing, vol. 10, no. 3, p. 434, 2018.
- [30] S. Jog and M. Dixit, "Supervised classification of satellite images," in 2016 Conference on Advances in Signal Processing (CASP), 2016: IEEE, pp. 93-98.