

توسعه و ارزیابی یک الگوریتم کاهش نوفه به منظور بهبود کارایی و دقت طبقه بندی تصاویر ابرطیفی

احسان لاله زاری^۱، علی اسماعیلی^{۲*}، سعید همایونی^۳

^۱ کارشناس ارشد مهندسی سنجش از دور - دانشکده مهندسی عمران و نقشه برداری - دانشگاه تحصیلات تکمیلی صنعتی و فناوری پیشرفته
ehsanlalehzari@gmail.com

^۲ استادیار دانشکده مهندسی عمران و نقشه برداری - دانشگاه تحصیلات تکمیلی صنعتی و فناوری پیشرفته
aliesmaeily@kgut.ac.ir

^۳ استادیار گروه جغرافیا، محیط زیست و ژئوماتیک - دانشگاه اتاوا - کانادا
saeid.homayouni@uOttawa.ca

(تاریخ دریافت آذر ۱۳۹۴، تاریخ تصویب شهریور ۱۳۹۶)

چکیده

تصویربرداری ابرطیفی، به عنوان یکی از فناوری‌های نوین سنجش از دوری، منبع ارزشمندی برای کاربردهای مختلف علوم زمین، از جمله تهیه نقشه‌های پوششی، شناسایی و اکتشاف معادن، نظارت زیست‌محیطی به شمار می‌رود. با این وجود، به دلایل سخت افزاری و فناوری این داده‌ها دارای مشکلات ذاتی هستند. از آنجایی که بهبود سیستم سخت افزاری سنجنده‌های ابرطیفی بسیار پرهزینه است، روش‌های سنجش از دوری پردازش تصویر مانند کاهش نویز، استخراج ویژگی و غیره به دلیل هزینه کم و موثر بودن مورد توجه قرار گرفته‌اند. یکی از جدیدترین و کارآمدترین این روش‌ها، روش پیش‌بینی فرضیه چندگانه است. نقطه ضعف این روش عدم استفاده از روشی موثر در انتخاب باندهای با شباهت بیشتر است که هدف از این مقاله بررسی روش پیش‌بینی فرضیه چندگانه^۱ و اتخاذ روشی مناسب برای انتخاب باندهای طیفی بر مبنای رگرسیون خطی است. به دلیل انعطاف زیاد روش پیش‌بینی رگرسیونی در تعیین ضرایب شباهت بین باندها، برای انتخاب باندهای طیفی مشابه، این روش انتخاب و پیاده‌سازی شد. داده‌های مورد استفاده در این تحقیق داده‌های رایج جهت کار بر روی تصاویر ابرطیفی است که توسط دانشگاه باسک اسپانیا جمع‌آوری شده‌اند. این داده‌ها شامل تصویر سایت‌های آزمایشی مزارع ایالت ایندیانا از سنجنده AVIRIS و تصویر دانشگاه پاویا از سنجنده ROSIS است. نتایج حاصله از پیاده‌سازی روش پیشنهادی نشان داد که صحت کلی طبقه‌بندی ماشین بردار پشتیبان^۲ و k نزدیکترین همسایگی^۳ برای مجموعه داده‌های ابرطیفی Indian Pines و دانشگاه Pavia به ترتیب برابر با ۹۵/۸۲، ۹۹/۴۳ و ۹۸/۸۸ و ۹۲/۸۹ است که در طبقه‌بندی SVM به ترتیب ۰/۴ و ۰/۳ و در طبقه‌بندی KNN به ترتیب ۸/۲۲ و ۲ درصد افزایش را نشان می‌دهد که نشان دهنده کارآمدی روش پیشنهادی به طور ویژه در مورد طبقه‌بندی KNN است.

واژگان کلیدی: تصاویر ابرطیفی، کاهش نویز، روش پیش‌بینی فرضیه چندگانه، طبقه‌بندی SVM، طبقه‌بندی KNN

* نویسنده رابط

^۱ Multi Hypothesis Prediction

^۲ Support Vector Machine

^۳ K Nearest Neighbor

۱- مقدمه

تصویربرداری ابرطیفی منبع ارزشمندی از اطلاعات مکانی برای گستره‌ی فراوانی از کاربردها مانند، تهیه نقشه‌های پوششی، شناسایی و اکتشاف معادن، نظارت زیست-محیطی است [۱-۴]. با این وجود، تأثیرات نویز روی تصاویر ابرطیفی باعث کاهش روش‌های استخراج اطلاعات و طبقه‌بندی این تصاویر می‌شود [۵و۶]. داده‌های حاصل از سنجنده‌های ابرطیفی سنجنش از دوری به دلیل ماهیت الکترواپتیک سنجنده و همچنین شرایط محیطی چون تأثیرات اتمسفری و تغییرات توپوگرافی دارای خطاهای گوناگونی هستند. این خطاها دارای دو منبع شناخته شده سیستماتیک و ناشناخته تصادفی هستند. نحوه برخورد یا این دو نوع خطا در تصاویر ابرطیفی متفاوت است و هرکدام نیاز به مدلسازی ریاضی و آماری متفاوتی دارد. بنابراین یکی از مراحل مهم پیش پردازش تصاویر ابرطیفی کاهش نویز و افزایش محتوای اطلاعاتی این تصاویر و کمک به بهبود قابلیت‌های تحلیلی تصاویر است [۷]. کاهش خطاهای ناخواسته در تصاویر سنجنش از دوری و بهبود کیفیت این داده‌ها از ابتدای توسعه‌ی روش‌های پردازش تصویر مورد توجه پژوهشگران گوناگونی بوده است. هر چند با پیشرفت دانش و فناوری ساخت سنجنده‌های الکترواپتیک کیفیت داده‌های این سیستم‌ها نیز افزایش یافته است، به دلیل طبیعت این سیستم‌ها و همچنین اثرات محیطی، کاهش اثرات نویز در تصاویر یکی از زمینه‌های مهم کاربردی در سنجنش از دور و پردازش تصویر است. این مهم در مورد تصاویر ابرطیفی دارای اهمیت ویژه‌ای است.

زیرا این سنجنده‌ها دارای ویژگی‌ها و توانایی‌های منحصر به فردی هستند که از جمله می‌توان به میزان اخذ داده‌های طیفی-مکانی با توان تفکیک رادیومتریک بالا آنها اشاره کرد. به همین دلیل این داده‌ها به صورت ذاتی دارای میزان اطلاعات به نویز نسبتاً پایینی نسبت به داده‌های سنجنش از دوری چندطیفی هستند [۸]. از این رو روش پردازش، مدلسازی و کاهش نویز در این تصاویر تا حدودی با تصاویر متداول سنجنش دوری متفاوت بوده است. تا به امروز تحقیقات زیادی روی تأثیر کاهش نویز ابرطیفی صورت گرفته است [۹-۱۳] که نمایانگر اهمیت انجام پیش‌پردازش کاهش نویز قبل از انجام پردازش-هایی نظیر طبقه‌بندی، آشکارسازی هدف، شاخص‌های گیاهی و غیره است [۱۴-۱۶].

Green و همکاران (۱۹۸۸) برای پرداختن به نویز و تعداد باند های تصاویر ابرطیفی روش حداقل کسر نویز (MNF)^۱ را ارائه کردند. تبدیل MNF مولفه‌های جدید کیفیت تصویر را تولید می‌کند و عوارض طیفی بهتری در مولفه‌های اساسی نسبت به تبدیل PCA فراهم می‌آورد، و همچنین توزیع نویز طیفی در این روش اهمیت ندارد [۱۷]. Van Aardt و Wynne (۲۰۰۷) به طور موفقیت آمیزی از داده‌های ابرطیفی سنجنده AVIRIS به منظور تمایز بین گونه های مختلف درختان کاج استفاده کردند، با دقیق‌ترین طبقه‌بندی که در نتیجه‌ی آنالیزهای تمایز اعمال شده روی تبدیل MNF بود. در این تحقیق آن‌ها توانستند با این آنالیزهای اعمال شده روی تبدیل MNF به دقت طبقه‌بندی ۸۵ درصد دست پیدا کنند [۱۸].

Philips و Watson (۲۰۰۸) یک فیلتر میانه تطبیقی (AF) ارائه دادند و این فیلتر میانه‌ی تطبیقی بر مبنای مشتق (AFD) برای تقسیم باندها به زیر باندها و اعمال فیلترهای میانه مکانی بزرگتر به محدوده باندها انجام شد که با کاهش نسبت سیگنال به نویز همراه بود. نتیجه‌ی اعمال این فیلترها دستیابی به دقت طبقه‌بندی بالای ۸۶ درصد بود [۱۹].

Bourenane و همکاران (۲۰۱۰) روشی بر مبنای بهبود الگوریتم فیلتر چندبعدی وینر (MWF)^۲ بر مبنای تنسور Tucker ارائه دادند، که به صورت همزمان روی مولفه‌های مکانی و طیفی تصاویر ابرطیفی عمل می‌کرد و به بهبود همزمان کیفیت تصویر و دقت طبقه‌بندی می‌پرداخت. به طوری که باعث بهبود دقت طبقه‌بندی بوسیله‌ی ماشین بردار پشتیبان بر روی تصاویر ابرطیفی سنجنده AVIRIS شد. با این حال استفاده از یک الگوریتم با هسته‌ی تنسور n حالتی، ممکن است به فشردسازی اطلاعات و از دست دادن جزئیات مکانی منجر شود [۲۰].

Qian و همکاران (۲۰۱۲) روشی سازمان یافته برای مقابله با مشکل طبقه بندی تصاویر ابرطیفی، ارائه کردند. روش آن‌ها شامل دو جزء کلی بود، یک توصیف کننده بافت مکانی-طیفی بر مبنای تبدیل موجک ۳بعدی برای جمع-آوری ویژگی‌های ذاتی و یک مدل رگرسیون منطقی پراکنده برای انتخاب ویژگی و طبقه‌بندی پیکسل‌ها. مزیت این روش ساده‌تر کردن تخمین پارامترهای آن و معتبر بودن آن روی داده‌های ابرطیفی واقعی و کاهش نویز آن‌ها بود [۲۱].

^۱ Minimum Noise Fraction

^۲ Multidimensional Wiener Filtering

۲-۱-۱- داده های Indian pines

این صحنه توسط سنجنده AVIRIS از سایت آزمایشی Indian pines در شمال غربی ایالت ایندیانا جمع آوری شده است و شامل ۱۴۵ سطر و ۱۴۵ ستون در ۲۲۴ باند بازتاب طیفی در محدوده ۰٫۴ تا ۲٫۵ میکرومتر می-باشد. این صحنه از یک صحنه بزرگتر بریده شده است. دو سوم صحنه Indian pines کشاورزی، یک سوم جنگل و بقیه پوشش گیاهی طبیعی دائمی است. واقعیت زمینی موجود در تصویر به شانزده کلاس تقسیم بندی شده (شکل ۱) و تعداد باند ها به علت وجود ناحیه جذبی بخار آب به ۲۰۰ باند کاهش یافته است. باندهای کاهش یافته شامل باندهای ۱۰۴ تا ۱۰۸ و ۱۵۰ تا ۱۶۳ است.



شکل ۱- نقشه واقعیت زمینی داده ی Indian Pines

به علت کافی نبودن تعداد پیکسل های بعضی از کلاس ها، تعدادی از کلاس ها حذف شده و به صورت بدون برچسب تلفیق شدند. در نهایت همانطور که در جدول ۱ این کلاس ها به صورت ضخیم و پر رنگ در آمده اند، ۹ کلاس باقی ماند.

جدول ۱- تعداد پیکسل کلاس های داده ی Indian Pines

#	نام کلاس	تعداد نمونه ها
1	Alfalfa	46
2	Corn-notill	1428
3	Corn-mintill	830
4	Corn	237
5	Grass-pasture	483
6	Grass-trees	730
7	Grass-pasture-mowed	28
8	Hay-windrowed	478
9	Oats	20
10	Soybean-notill	972
11	Soybean-mintill	2455
12	Soybean-clean	593
13	Wheat	205
14	Woods	1265
15	Buildings-Grass-Trees-Drives	386
16	Stone-Steel-Towers	93

Xu و همکاران (۲۰۱۳) نه تنها حذف همبستگی مکانی، بلکه حذف همبستگی طیفی را با استفاده از تبدیل موجک مد نظر داشتند. آن ها روشی بر مبنای رگرسیون چندخطی (MLR)^۱ و تبدیل موجک به منظور تخمین نویز تصاویر ابرطیفی ارائه دادند که روی تعدادی از داده های شبیه سازی شده ی سنجنده AVIRIS و داده های واقعی سنجنده Hyperion اعمال شد. نتایج تجربی آنها نشان داد که این روش انطباق و دقت بیشتری نسبت به سایر روش های طبقه بندی دارد [۲۲].

Chen و همکاران (۲۰۱۴) یک الگوریتم پیش پردازشی مقاوم به نویز تصاویر ابرطیفی طراحی کردند که بر مبنای همسایگی پیکسل های مورد نظر عمل می کرد. این روش می توانست دقت طبقه بندی هر دو طبقه بند بیشترین شباهت و ماشین بردار پشتیبان را بهبود دهد، مخصوصا در شرایطی که اندازه نمونه ها کوچک انتخاب شود و تصاویر نیز نویز داشته باشند [۲۳].

Zhao و Yang (۲۰۱۵) یک الگوریتم کاهش نویز تصاویر ابرطیفی، بوسیله ی افزودگی و همبستگی (RAC)^۲ در هر دو دامنه ی مکانی و طیفی ارائه کردند. نتایج کاهش نویز بوسیله ی روش پیشنهادی آن ها بهتر از سایر نتایج گرفته شده با روش های جدید برای کاهش نویز تصاویر ابرطیفی بود [۲۴].

بنابر تحقیقات بعمل آمده تا کنون روش هایی که الگوریتم های کاهش نویز را به طور همزمان روی مکان و طیف تصاویر ابرطیفی اعمال می کنند، دارای نتیجه ی بهتری روی دقت طبقه بندی این تصاویر هستند. بنابراین در تحقیق پیش رو، هدف ارائه الگوریتمی بهبود یافته به منظور کاهش نویز در تصاویر ابرطیفی و ارزیابی الگوریتم مورد نظر در افزایش دقت طبقه بندی این تصاویر است.

۲- مواد و روش ها

۲-۱- داده های مورد استفاده

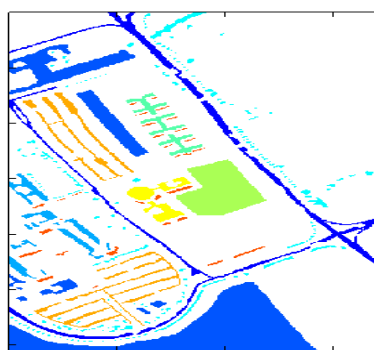
داده های مورد استفاده، داده های رایج جهت کار بر روی تصاویر ابرطیفی است که توسط دانشگاه باسک اسپانیا جمع آوری شده اند. این داده ها شامل تصویر مزارع ایالت ایندیانا از سنجنده AVIRIS و تصویر دانشگاه پلویا از سنجنده ROSIS است که در ادامه توضیح داده خواهد شد [۲۵].

^۱ Multi Linear Regression

^۲ Redundancy And Correlation

۲-۱-۲- داده های دانشگاه Pavia

این صحنه تصویربرداری توسط سنجندهی ROSIS در خلال یک پرواز روی شهر Pavia در شمال ایتالیا برداشت شده است. تعداد باندهای طیفی برای دانشگاه Pavia برابر ۱۰۳ است. ابعاد تصویر دانشگاه Pavia ۳۴۰*۶۱۰ پیکسل است، و تعدادی از ستونهای تصویر در صحنه تصویربرداری که شامل هیچگونه اطلاعاتی نبود قبل از پردازش تصویر حذف شد. قدرت تفکیک هندسی تصویر ۱/۳ متر است. همانطور که در شکل ۲ دیده می شود این داده دارای واقعیت زمینی شامل ۹ کلاس است.



شکل ۲- واقعیت زمینی داده دانشگاه Pavia

تعداد پیکسل هر کلاس از داده دانشگاه Pavia در جدول ۲ آورده شده است.

جدول ۲- تعداد پیکسل کلاسهای دادهی دانشگاه Pavia

#	نام کلاس	تعداد نمونه ها
1	Asphalt	6631
2	Meadows	18649
3	Gravel	2099
4	Trees	3064
5	Painted metal sheets	1345
6	Bare Soil	5029
7	Bitumen	1330
8	Self-Blocking Bricks	3682
9	Shadows	947

۲-۲- روش ها

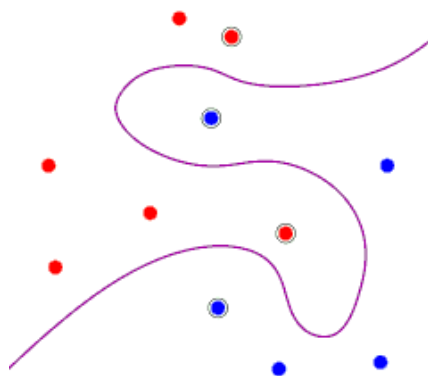
در این تحقیق تمامی پیش پردازش های لازم برای تصاویر ابرطیفی در نظر گرفته شده است. از جمله تصحیح رادیومتریکی (کالیبراسیون سنجنده)، کاهش نویز، استخراج ویژگی و اعتبارسنجی متقابل. برای رفع مشکلات ناشی از نویز روشی بر مبنای پیش بینی فرضیه چندگانه ارائه شده است که در ادامه به تشریح آن پرداخته خواهد شد. به منظور طبقه بندی تصاویر از الگوریتم های طبقه بندی ماشین

بردار پشتیبان و نزدیک ترین k همسایگی بنا به مزیت هایی که توضیح داده می شود، مورد استفاده قرار گرفتند. پیاده سازی تمام الگوریتم ها در محیط برنامه نویسی Matlab انجام شد. برای طبقه بندی SVM، از ابزار LibSVM [۲۶] که در محیط Matlab نصب گردید، کمک گرفته شد.

۲-۲-۱- ماشین بردار پشتیبان (SVM)

ماشین های بردار پشتیبان یک گروه از الگوریتم های طبقه بندی نظارت شده هستند که پیش بینی می کنند یک نمونه در کدام کلاس یا گروه قرار می گیرد. این الگوریتم برای تفکیک دو کلاس از یکدیگر، از یک صفحه استفاده می کند. به طوری که این صفحه از هر طرف بیشترین فاصله را تا هر دو کلاس داشته باشد. نزدیکترین نمونه های آموزشی به این صفحه، بردارهای پشتیبان نام دارند.

این الگوریتم حساسیت کمتری نسبت به پدیده های با ابعاد بالا دارد، از این رو در طبقه بندی داده های ابرطیفی روش مناسبی به شمار می رود. به طور کلی ماشین بردار پشتیبان، یک طبقه بندی کننده ی باینری و خطی است که با توسعه ی آن و استفاده از توابع کرنل، می توان به عنوان طبقه بندی کننده ی چندکلاسی و غیر خطی نیز از آن استفاده کرد (شکل ۳).



شکل ۳- مسئله فضای غیرخطی در طبقه بندی SVM

برای برقراری شرایط یکسان در مقایسه بین روش ها، در اینجا از کرنل گوسی با پارامتر تنظیم C و عرض کرنل γ استفاده شده است. برای تخمین پارامترهای بهینه ی C و γ از اعتبارسنجی متقابل ۵ لایه استفاده شد.

۲-۲-۲- K نزدیکترین همسایه (KNN)

k نزدیکترین همسایه (KNN) یک الگوریتم یادگیری است که در روش بازشناسی الگو طی چندین دهه مطالعه

شده توسط اعتبارسنجی متقابل استفاده کرد. در اینجا از اعتبارسنجی متقابل ۵ لایه (تقسیم مشاهدات به ۵ قسمت) استفاده شد. این کار برای برقراری شرایط یکسان در روش‌های مورد مقایسه انجام می‌شود.

۲-۲-۴- پیش‌بینی فرضیه چندگانه (MHPrediction)

اگر تصویر ابرطیفی ورودی را مجموعه‌ای از M بردار پیکسلی در نظر بگیریم: $X = \{x_m \in R^N, m = 1, 2, \dots, M\}$ به طوری که N تعداد باند های طیفی باشد. فرض کنید که x یک بردار پیکسلی (امضای طیفی) از تصویر باشد. هدف، پیدا کردن یک ترکیب خطی بهینه برای تمام پیش‌بینی‌های ممکن یا فرض‌هایی برای نمایش مجدد x است. نمایش مجدد بهینه می‌تواند به صورت زیر فرموله شود:

$$\hat{W} = \operatorname{argmin} \|X - Hw\|_2^2 \quad (2)$$

به طوری که H یک ماتریس با ابعاد $N \times L$ است که ستون‌هایی با L فرض ممکن دارد و $\hat{W} = [\hat{w}_1, \dots, \hat{w}_L]^T$ یک بردار $L \times 1$ از ضرایب مربوط به تمام ستون‌های H است. به دلیل اینکه برای پنجره‌های با ابعاد بزرگ، $N < L$ است، ماهیت بد مطرح شده مسئله نیازمند بعضی از انواع منظم‌سازی است که برای تفکیک در میان تعداد بی‌شماری از ترکیبات خطی ممکن، تکیه بر فضای جواب معادله (۲) دارد.

رایج‌ترین رویکرد برای منظم‌سازی مسئله کمترین مربعات، روش منظم‌سازی تیخونوف است [۲۸] که یک ضریب جریمه L_2 روی نرم $\hat{W}$ اعمال می‌کند.

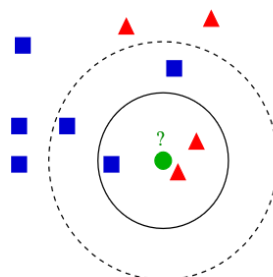
$$\hat{W} = \operatorname{argmin} \|X - Hw\|_2^2 + \lambda \|F_w\|_2^2 \quad (3)$$

به طوری که F ماتریس تیخونوف و λ پارامتر منظم‌سازی است. قسمت F اجازه تحمیل معلومات قبلی را روی جواب مسئله می‌دهد.

در این جا رویکردهایی پیشنهاد شده است که در آن فرض‌هایی که بیشترین عدم شباهت از بردار پیکسلی اصلی را دارند، نسبت به فرض‌هایی که بیشترین شباهت را دارند باید وزن کمتری به آن‌ها داده شود. به طور مشخص، یک F قطری به صورت زیر اتخاذ می‌شود:

$$F = \begin{bmatrix} \|x - h_1\|_2 & \dots & \emptyset \\ \vdots & \ddots & \vdots \\ \emptyset & \dots & \|x - h_L\|_2 \end{bmatrix} \quad (4)$$

شده است. KNN به عنوان یکی از کاراترین و ساده‌ترین روش‌ها شناخته شده است. این الگوریتم، k همسایه نزدیک در میان پیکسل‌های آموزشی بر اساس یک معیار شباهت پیدا کرده و کلاس‌های این k همسایه را بر اساس فراوانی آن‌ها انتخاب می‌کند (شکل ۴).



شکل ۴- نحوه شناسایی کلاس‌ها در همسایگی‌های مختلف

یک مشکل در روش KNN تعیین مقدار k (تعداد همسایه‌ها) است و برای تعیین آن باید یک سری آزمایشات با مقادیر مختلف k انجام شود تا بهترین مقدار برای k تعیین گردد. در اینجا از روش اعتبارسنجی متقابل ۵ لایه برای پیدا کردن مقدار بهینه‌ی k استفاده شده است.

۲-۲-۳- اعتبارسنجی متقابل (Cross Validation)

در این تحقیق برای ارزیابی مدل‌های طبقه‌بندی از روش اعتبارسنجی متقابل [۲۷] استفاده شده است. در روش اعتبارسنجی متقابل یک مشاهده از مجموعه‌ی مشاهدات حذف شده و مقدار آن بر اساس $N-1$ مشاهده‌ی باقی‌مانده پیش‌بینی می‌شود. این عمل را برای تمام N موقعیت تکرار کرده و ملاک ریشه‌ی دوم میانگین توان دوم خطاهای اعتبارسنجی متقابل به صورت:

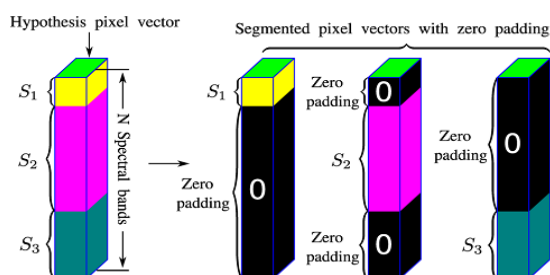
$$RMSE = \left[\frac{1}{N} \sum_{i=1}^N e_{(-i)}^2 \right]^{\frac{1}{2}} \quad (1)$$

محاسبه می‌شود، که در آن $e_{(-i)}$ خطای پیش‌بینی محاسبه شده بر اساس مشاهدات به جز مشاهده‌ی i ام است. سپس دقت کلی اعتبارسنجی به صورت $1-RMSE$ بدست می‌آید. بدیهی است که درصد این مقدار هرچه به ۰۰۱ نزدیکتر باشد، پیش‌بینی از دقت بیشتری برخوردار است.

از طرف دیگر برای پیش‌بینی تعداد همسایگی‌های بهینه در طبقه‌بندی KNN و پارامترهای بهینه‌ی C و γ در طبقه‌بندی SVM نیز می‌توان از پارامترهای پیش‌بینی

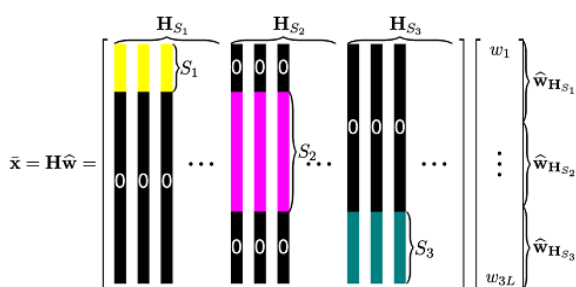
همبستگی بین باندها تقسیم شوند، به طوری که در هر گروه با یکدیگر همبستگی بالایی داشته باشند.

به طور کلی فرض می‌شود یک بردار پیکسلی N بعدی بر طبق ماتریس ضرایب همبستگی متقابل باندها به J بخش تقسیم شده است. با نگه داشتن فقط یکی از بخش‌های J و جا به جا کردن بقیه بخش‌ها با مقدار صفر، شکل یک بردار N بعدی ساخته می‌شود (Zero Padding). این پردازش برای تولید فرضیه‌هایی بر مبنای بخش‌بندی باندهای طیفی است که در شکل ۶ برای $J=3$ نشان داده شده است.



شکل ۶- بخش‌بندی باندهای طیفی

اگر L فرضیه از پنجره جستجو بدست آید و J بخش برای تقسیم بندی باندهای طیفی استفاده شود، بنابراین تعداد کل فرضیه‌های H برابر با $J \times L$ خواهد بود. انگیزه این تولید فرضیه از باندهای طیفی بخش‌بندی شده به گونه‌ای است که وزن‌های محاسبه شده برای فرضیه‌ها قابل متعادل شدن برای بخش‌های طیفی مختلف می‌شود. جزئیات در شکل ۷ با $J=3$ نشان داده شده است.



شکل ۷

۲-۵- آنالیز تفکیک‌پذیری فضای ویژگی (FSDA)^۱

ایمانی و قاسمیان (۲۰۱۵) الگوریتم آنالیز تفکیک‌پذیری فضای ویژگی (FSDA) را ارائه دادند [۲۹]. روش

به طوری که $h_1, h_2, \dots, h_L$ ستون‌های ماتریس H هستند. با این ساختار، Γ وزن‌های دامنه‌ی بزرگ اختصاص یافته به آن فرضیه‌هایی که دارای یک فاصله قابل توجه از x باشند را جریمه می‌کند. سپس برای هر بردار پیکسلی، $\hat{W}$ می‌تواند مستقیماً بوسیله‌ی راه حل معمولی تیخونوف محاسبه شود:

$$\hat{W} = (H^T H + \lambda \Gamma^T \Gamma)^{-1} H^T x \quad (5)$$

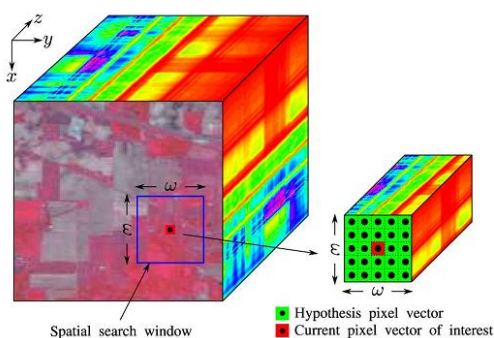
بنابراین، بردار پیکسلی پیش‌بینی شده که x را تقریب می‌زند به صورت زیر محاسبه می‌شود:

$$\bar{x} = H \hat{W} \quad (6)$$

و مجموعه داده‌ی پیش‌بینی شده $\bar{X} = \{\bar{x}_m \in R^N, m = 1, 2, \dots, M\}$ بوسیله جایگذاری هر بردار پیکسلی در X با بردار پیش‌بینی شده مربوط به آن تولید می‌شود.

تولید مجموعه‌ی فرضیه‌ها:

یک مجموعه داده‌ی ابرطیفی، معمولاً نشان دهنده بعضی درجات از پیوستگی به صورت تکه‌ای مکانی است. برای هر نمونه، بردارهای پیکسلی همسایگی مکانی آن، خصوصیات طیفی مشابهی نسبت به آن خواهند داشت. برای ارائه رویکرد فرضیه چندگانه، برای یک بردار پیکسلی مورد نظر با در نظر گرفتن تمام همسایگی‌های بردار پیکسلی داخل یک پنجره، جستجوی مکانی با اندازه $w \times w$ تولید می‌شود (شکل ۵).



شکل ۵- پنجره جستجوی مکانی

سپس این بردارهای پیکسلی همسایگی، به عنوان ستون‌های ماتریس فرضیه H قرار داده می‌شوند. به دلیل اینکه باندهای طیفی تصویر ابرطیفی دارای همبستگی هستند، می‌توانند به چندین گروه بر مبنای ضرایب

^۱ Feature Space Discriminant Analysis

مراحل پیاده سازی الگوریتم FSDA

۱- بدست آوردن ماتریس پراکندگی بین طیفی از روی مجموعه داده ی اولیه:

$$X_j = \begin{bmatrix} x_{11j} & x_{12j} & \dots & x_{1cj} \\ x_{21j} & x_{22j} & \dots & x_{2cj} \\ \vdots & \vdots & \ddots & \vdots \\ x_{d1j} & x_{d2j} & \dots & x_{dcj} \end{bmatrix}$$

$$j = 1, \dots, n_t$$

$$S_f = \sum_{j=1}^{nt} \sum_{i=1}^d (h_{ij} - \bar{h}_j)(h_{ij} - \bar{h}_j)^T$$

$$\bar{h}_j = \frac{1}{d} \sum_{i=1}^d h_{ij}$$

۲- بیشینه سازی $tr(S_f)$ و بدست آوردن ماتریس تصویر کننده W به طوری که نگاشت آن برابر است با:

$$(g_{ij})_{c \times c} = W_{c \times c} (h_{ij})_{c \times 1}$$

$$j = 1, \dots, n_t$$

$$i = 1, 2, \dots, d$$

و فضای ویژگی جدید برابر است با:

$$R_j = \begin{bmatrix} r_{11j} & r_{12j} & \dots & r_{1cj} \\ r_{21j} & r_{22j} & \dots & r_{2cj} \\ \vdots & \vdots & \ddots & \vdots \\ r_{d1j} & r_{d2j} & \dots & r_{dcj} \end{bmatrix}$$

$$j = 1, \dots, n_t$$

$$R_{kj} = [r_{1kj} \quad r_{2kj} \quad \dots \quad r_{dkj}]$$

$$k = 1, 2, \dots, c$$

۳- محاسبه ماتریس های پراکندگی بین کلاسی و درون کلاسی:

$$S_b = \sum_{j=1}^{nt} \sum_{k=1}^c (R_{kj} - \bar{R})(R_{kj} - \bar{R})^T$$

$$S_w = \sum_{k=1}^c \sum_{j=1}^{nt} \sum_{i=1}^d (R_{ki} - R_{Kj})(R_{ki} - R_{Kj})^T$$

برای مقابله با کاهش مرتبه S_w :

$$S_w = 0.5S_w + 0.5diag(S_w) \quad (11)$$

۴- بدست آوردن ماتریس تصویر کننده ثانویه

$$y_{p \times 1} = A_{p \times d} \times X_{d \times 1}$$

$$A = \max tr(S_w^{-1} S_b) \quad (12)$$

FSDA از دو ماتریس تصویر کننده برای استخراج ویژگی استفاده می کند. ماتریس تصویر کننده اولیه برای بیشینه کردن پراکندگی بین طیفی استفاده شده و ماتریس تصویر کننده ثانویه که روی فضای ویژگی که در مرحله اول بدست آمده است اعمال می شود، تفکیک پذیری بین کلاس ها را بیشینه می کند.

روش های استخراج ویژگی مانند LDA^۱، GDA^۲، NWFE^۳ و MMLDA^۴ تنها از دو اندازه گیری استفاده می کنند: پراکندگی درون کلاسی و پراکندگی بین کلاسی. روش FSDA علاوه بر پراکندگی های درون کلاسی و بین کلاسی، از یک اندازه گیری سومی با عنوان پراکندگی بین طیفی نیز استفاده می کند. در تعداد نمونه های آموزشی کم روش FSDA دارای عملکرد بهتری نسبت به سایر روش ها است.

در روش FSDA ابتدا با استفاده از یک تصویر کننده اولیه فضای ویژگی به یک فضا که دارای تفکیک پذیری بیشتری است تبدیل می شود و سپس با استفاده از یک تصویر کننده ثانویه فضای ویژگی بدست آمده از مرحله اول به یک فضایی که کلاس های مختلف دارای بیشترین تفکیک پذیری در آن هستند، تبدیل می شود. تفاوت اصلی روش FSDA با سایر روش ها همین تصویر کننده اولیه است. زمانی که تعداد نمونه های آموزشی کم باشد، این تصویر کننده اولیه عملکرد طبقه بندی را به طور قابل توجهی افزایش می دهد.

در این روش بر خلاف سایر روش ها، به جای استفاده از میانگین کلاس ها از تمام نمونه های آموزشی برای تشکیل ماتریس های پراکندگی بین طیفی، بین کلاسی و درون کلاسی استفاده می شود که باعث افزایش کارایی طبقه بندی با تعداد نمونه های آموزشی کم می شود.

در اغلب روش ها به دلیل محدودیت مرتبه ماتریس پراکندگی بین کلاسی می توان حداکثر C-1 (C تعداد کلاس ها است) ویژگی از داده ها استخراج کرد ولی در روش FSDA به دلیل استفاده از نمونه های آموزشی به جای میانگین کلاس ها، این اجازه را می دهد تا بیش از C-1 ویژگی از داده ها استخراج کرد.

^۱ Linear Discriminant Analysis

^۲ Global Discriminant Analysis

^۳ Nonparametric weighted feature extraction

^۴ Median-Mean Line based Discriminant Analysis

به طوری که A ماتریس تصویرکننده ثانویه و y مجموعه ویژگی‌های استخراج شده و X مجموعه داده اولیه (خام) است.

۳- نتایج و بحث و بررسی

۳-۱- پیش پردازش داده‌ها

ابتدا تصحیح رادیومتریکی بر روی داده‌ها انجام شد و سپس الگوریتم پیش‌بینی فرضیه چندگانه با استفاده از انتخاب رگرسیونی باندهای طیفی، برای کاهش نویز روی داده‌ها اعمال گردید. با استفاده از الگوریتم FSDA به تعداد باندهای داده‌ها ویژگی استخراج شد و در نهایت پیش از اعمال طبقه‌بندی، اعتبارسنجی متقابل ۵ لایه جهت جلوگیری از خطای Overfitting و بدست آوردن پارامترهای بهینه در طبقه‌بندی‌های SVM و KNN اعمال شد.

در روش پیش‌بینی فرضیه چندگانه روش‌های مختلفی جهت تقسیم‌بندی باندهای طیفی اعمال شد و در نهایت نتیجه گرفته شد که می‌توان از برازش خطی رگرسیونی بین باندهای طیفی به نتایج بهتری، هم در افزایش نسبت سیگنال به نویز و هم در افزایش دقت طبقه‌بندی دست یافت.

در این تحقیق روش پیش‌بینی فرضیه چندگانه برای کارایی بهتر مورد بازبینی قرار گرفت. با توجه به آزمایش‌هایی که در این روش انجام گرفت به این نتیجه رسیدیم که به منظور انتخاب محدوده باندهای با شباهت بیشتر به جای استفاده از ضرایب همبستگی بین باندها، برای محدوده‌بندی آن‌ها از رگرسیون خطی بدین منظور استفاده کنیم. این اقدام باعث بهبودهای چشم‌گیری در نسبت سیگنال به نویز و بهبود دقت طبقه‌بندی تصاویر ابرطیفی شد که در ادامه نتایج ارائه خواهد شد.

ایده‌ی این روش بر مبنای تحقیق کیان‌دو و هی‌یانگ (۲۰۰۸) در زمینه‌ی انتخاب باندهای ابرطیفی بر مبنای شاخص‌های شباهت، ارائه شده است. در این مقاله سعی شده است که از رگرسیون خطی برای پیدا کردن باندهای مشابه استفاده شود [۳۰].

به منظور ارزیابی بین انتخاب محدوده‌های باندی داده‌ها و بهبود نسبت سیگنال به نویز (رابطه ۱۳) جدول ۳ ارائه می‌شود. این نسبت به صورت

$$SNR(x_m, \hat{x}_m) = \log_{10} \frac{Var(x_m)}{MSE(x_m, \hat{x}_m)} \quad (13)$$

بدست می‌آید به طوریکه x_m بردار پیکسلی اصلی، $\hat{x}_m$ بردار پیکسلی حاوی نویز و $Var(x_m)$ واریانس اجزای بردار x_m است. همچنین میانگین مربع خطا (MSE) به صورت

$$MSE(x_m, \hat{x}_m) = \frac{1}{M} \|x_m - \hat{x}_m\|_2^2 \quad (14)$$

بدست می‌آید [۲۵].

جدول ۳- مقایسه نسبت سیگنال به نویز داده‌های پیش‌پردازش شده به روش‌های مختلف

روش	مجموعه داده	محدوده باندها	SNR
MH	Indian Pines	{۱۰۶:۲۰۰ و ۷۶:۱۰۵ و ۳۶:۷۵ و ۱:۳۵}	۱۰/۸۰
Reg-MH	Indian Pines	{۱۴۶:۲۰۰ و ۱:۱۴۵}	۲۴/۰۷
MH	Pavia	{۷۶:۱۰۳ و ۱:۷۵}	۱۱/۷۴
Reg-MH	Univesity Pavia Univesity	{۹۱:۱۰۳ و ۱:۹۰}	۱۳/۲۹

همانطور که در جدول ۳ دیده می‌شود نسبت سیگنال به نویز در مورد داده Indian Pines بعد از اعمال الگوریتم پیشنهادی بسیار بهتر شده است. در مورد داده دانشگاه Pavia نیز این نسبت در حد قابل قبولی افزایش نشان می‌دهد.

۳-۲- نتایج طبقه‌بندی SVM

۳-۲-۱- داده‌ی Indian Pines

پیش از انجام طبقه‌بندی، پارامترهای طبقه‌بندی SVM به روش کرنل 'RBF' گوسی، شامل پارامترهای C و γ توسط روش اعتبارسنجی متقابل ۵ لایه برای مجموعه داده‌های تولید شده به روش‌های مختلف، بدست آمد. این پارامترها برای مجموعه داده‌ی Indian Pines در جدول ۴ قابل مشاهده است.

جدول ۴- پارامترهای بدست آمده از اعتبارسنجی متقابل ۵ لایه

روش	C	γ	دقت اعتبارسنجی
MH	۲۶۲۱۴۴	۰/۰۰۳۹	۹۹/۸۸
FSDA	۱۶	۱	۹۹/۰۷
MH-FSDA	۱۶	۰/۲۵	۹۹/۸۹
Reg-MH-FSDA	۲۶۲۱۴۴	۰/۰۰۳۹	۹۹/۹۱

یافته است. این افزایش حاصل بهبود ۱۳/۲۷ دسیبلی (جدول ۳) در نسبت سیگنال به نویز بوده است. این روش تلفیقی و بهبود یافته، دقت بالای طبقه‌بندی با تعداد نمونه‌های آموزشی کم را از الگوریتم استخراج ویژگی FSDA و دقت بالای پایدار در طبقه‌بندی را از الگوریتم کاهش نویز پیش‌بینی فرضیه چندگانه به ارث برده است. همچنین افزایش چشمگیر نسبت سیگنال به نویز خود را مرهون انتخاب بهینه‌ی باندهای طیفی از روش رگرسیون است.

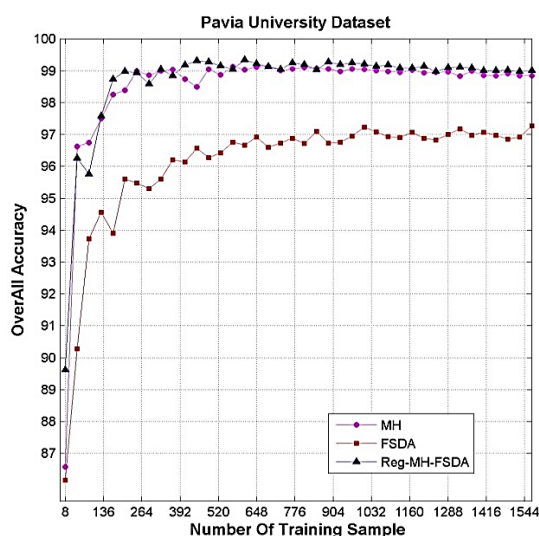
۳-۲-۲- داده‌های دانشگاه Pavia

پیش از انجام طبقه‌بندی، پارامترهای طبقه‌بندی SVM به روش کرنل RBF گوسی، شامل پارامترهای C و γ توسط روش اعتبارسنجی متقابل ۵ لایه^۱ برای مجموعه داده‌های تولید شده به روش‌های مختلف، بدست آمد. این پارامترها برای مجموعه داده‌ی دانشگاه Pavia در جدول ۵ قابل مشاهده است.

جدول ۵- پارامترهای بدست آمده از اعتبارسنجی متقابل ۵ لایه

روش	C	γ	دقت اعتبارسنجی
MH	۸۱۹۲	۰/۱۲۵	۹۹/۹۹
FSDA	۸	۱۲۸	۹۹/۹۶
MH-FSDA	۳۲	۸	۹۹/۹۹
Reg-MH-FSDA	۸	۳۲	۹۹/۹۹

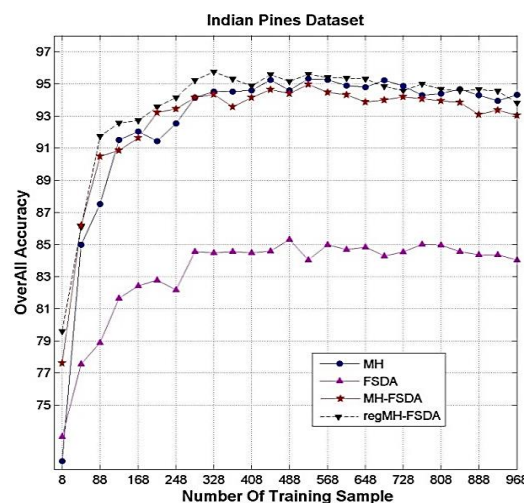
طبقه‌بندی SVM برای تعداد نمونه‌های آموزشی متفاوت جهت بررسی کارکرد پیش‌پردازش‌ها انجام شد. نتایج در شکل ۹ قابل مشاهده است.



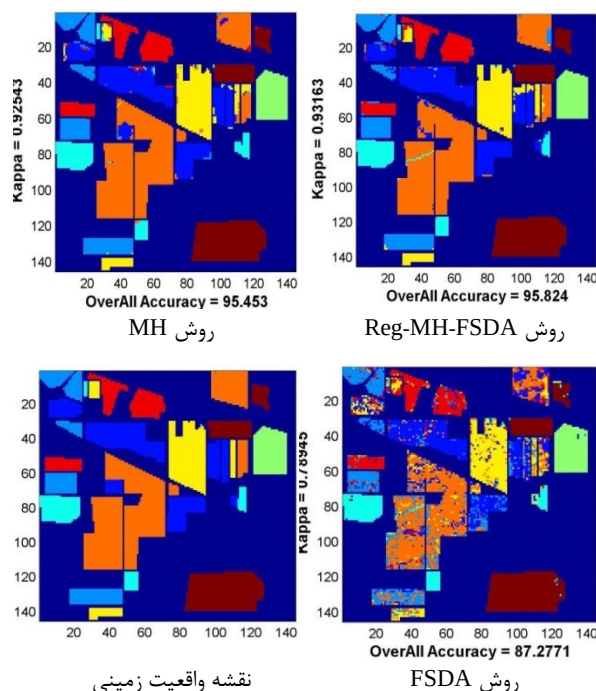
شکل ۹- نتایج طبقه‌بندی SVM برای نمونه‌های آموزشی متفاوت

^۱ 5 Fold Cross Validation

طبقه‌بندی SVM برای تعداد نمونه‌های آموزشی متفاوت جهت بررسی کارکرد پیش‌پردازش‌ها انجام شد. نتایج در شکل ۸ قابل مشاهده است.



شکل ۸- نتایج طبقه‌بندی SVM برای نمونه‌های آموزشی متفاوت

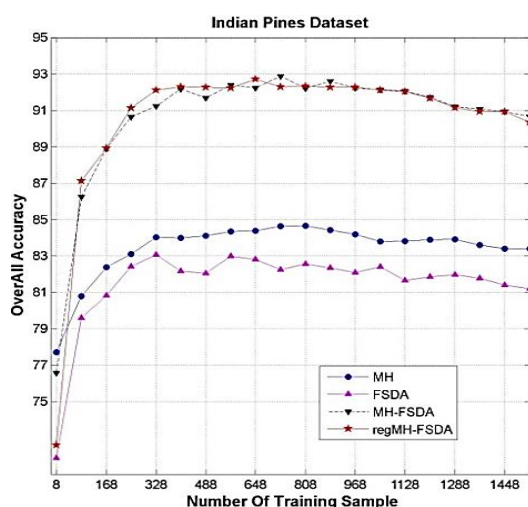


شکل ۹- کلاس‌های پیش‌بینی شده‌ی داده‌ی Indian Pines در مقابل واقعیت زمینی برای طبقه‌بندی SVM پیش‌پردازش‌های مختلف

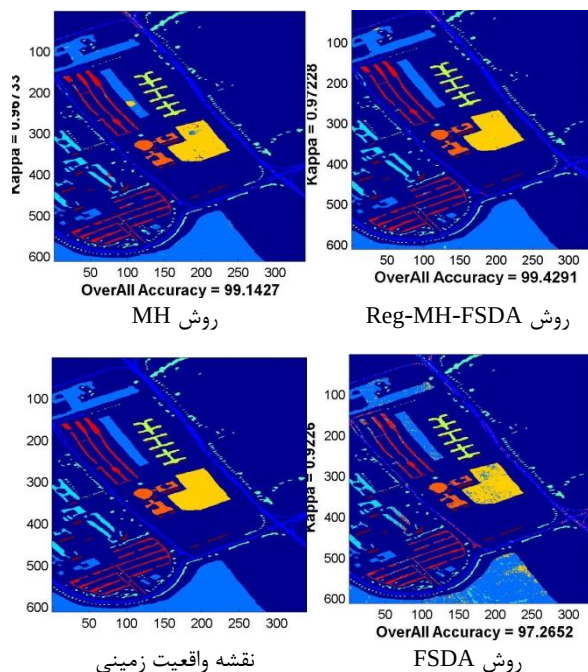
نتایج طبقه‌بندی SVM در مورد داده‌های پیش‌پردازش شده با روش‌های مختلف همراه با نقشه‌ی واقعیت زمینی، با ۵۰ بار تکرار و میانگین‌گیری از آن‌ها جهت حذف خطای انتخاب تصادفی نمونه‌های آموزشی در شکل ۹ آورده شده است. نتایج نشان می‌دهد که دقت طبقه‌بندی SVM داده‌های پیش‌پردازش شده با روش پیش‌بینی فرضیه چندگانه با انتخاب باند رگرسیونی تا حدود ۰/۴ درصد افزایش

نتایج مربوط به دقت طبقه‌بندی KNN برای تعداد نمونه‌های آموزشی متفاوت در شکل ۱۱ آمده است. همانطور که در شکل مشخص است، الگوریتم پیش‌پردازش فرضیه چندگانه با انتخاب باند رگرسیونی در مورد طبقه‌بندی KNN بسیار عالی عمل کرده است. نتیجه‌ی این عملکرد افزایش ۷ الی ۱۰ درصدی دقت طبقه‌بندی به روش KNN برای داده‌های پیش‌پردازش شده با روش پیش‌بینی فرضیه چندگانه با انتخاب باند رگرسیونی در مقابل روش‌های FSDA و پیش‌بینی MH بوده است.

نتایج طبقه‌بندی KNN داده‌ی Indian Pines برای روش‌های پیش‌پردازش با استخراج ویژگی FSDA، پیش‌بینی فرضیه چندگانه و پیش‌بینی فرضیه چندگانه با انتخاب باند رگرسیونی در برابر نقشه‌ی واقعیت زمینی در شکل ۱۲ نشان داده شده است که برای تعداد نمونه آموزشی برابر روش پیشنهادی دارای افزایش ۸/۲ درصدی در دقت طبقه‌بندی بوده است.



شکل ۱۱- نتایج طبقه‌بندی KNN برای نمونه‌های آموزشی متفاوت



شکل ۱۰- کلاس‌های پیش‌بینی شده‌ی داده‌ی دانشگاه Pavia در مقابل واقعیت زمینی برای طبقه‌بندی SVM برای پیش‌پردازش‌های مختلف

نتایج طبقه‌بندی SVM در مورد داده‌های پیش‌پردازش شده دانشگاه Pavia با روش‌های مختلف همراه با نقشه‌ی واقعیت زمینی، با ۵۰ بار تکرار و میانگین‌گیری از آن‌ها جهت حذف خطای انتخاب تصادفی نمونه‌های آموزشی در شکل ۱۰ آورده شده است.

۳-۳- نتایج طبقه‌بندی KNN

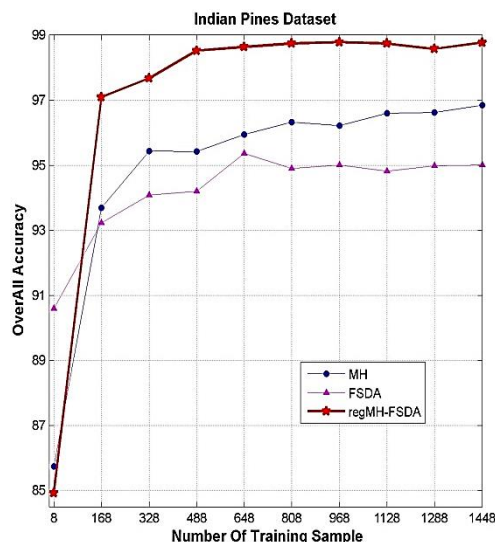
۳-۳-۱- داده‌های Indian Pines

پیش از طبقه‌بندی داده‌ی Indian Pines، از اعتبارسنجی متقابل ۵ لایه جهت جلوگیری از مسئله‌ی Overfitting و همچنین بدست آوردن پارامتر بهینه‌ی طبقه‌بندی KNN، یعنی تعداد بهینه‌ی همسایگی‌ها جهت طبقه‌بندی استفاده شد.

جدول ۶ نتایج این اعتبارسنجی را نشان می‌دهد.

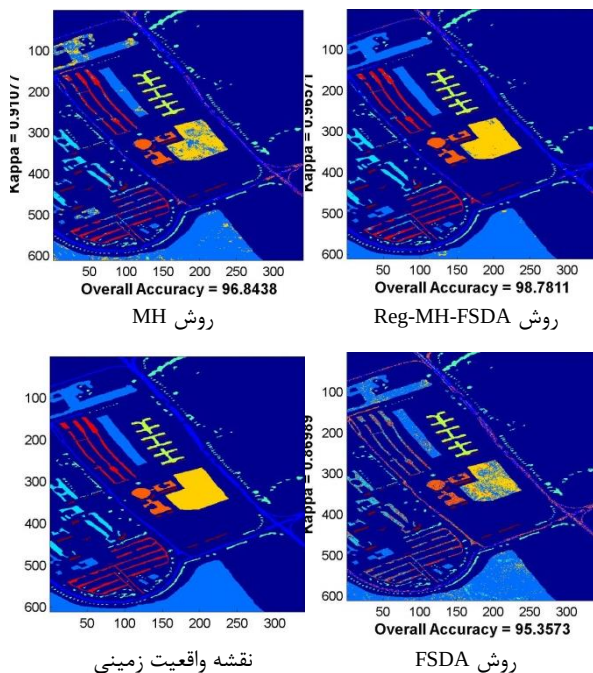
جدول ۶- پارامترهای بدست آمده از اعتبارسنجی متقابل ۵ لایه

روش	همسایگی بهینه	دقت اعتبارسنجی
MH	۷	۹۹/۲۹
FSDA	۱	۹۹/۳۴
MH-FSDA	۱	۹۹/۸۵
Reg-MH-FSDA	۱	۹۹/۸۵



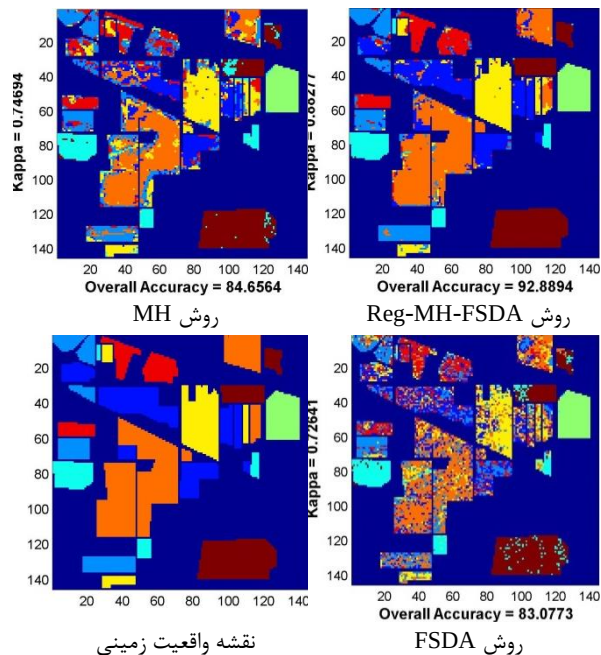
شکل ۱۳- نتایج طبقه‌بندی KNN برای نمونه‌های آموزشی متفاوت

همانطور که در شکل ۱۳ مشاهده می‌شود، طبقه‌بندی KNN برای داده‌ی دانشگاه Pavia در مورد پیش‌پردازش با روش پیشنهادی نتایج قابل توجهی نسبت به سایر روش‌های پیش‌پردازش گرفته است، که این امر نشان دهنده‌ی عملکرد بهتر الگوریتم پیش‌پردازش پیشنهادی است.



شکل ۱۴- کلاس‌های پیش‌بینی شده‌ی داده‌ی دانشگاه Pavia در مقابل واقعیت زمینی برای طبقه‌بندی KNN برای پیش‌پردازش‌های مختلف

همانطور که در شکل ۱۴ نشان داده شده است، دقت طبقه‌بندی KNN بر روی داده‌های دانشگاه Pavia پیش‌پردازش شده با روش پیشنهادی حدود ۲ درصد بهبود یافته است.



شکل ۱۲- کلاس‌های پیش‌بینی شده‌ی داده‌ی Indian Pines در مقابل واقعیت زمینی برای طبقه‌بندی KNN برای پیش‌پردازش‌های مختلف

۳-۲-۳- داده‌های دانشگاه Pavia

پیش از طبقه‌بندی داده‌ی دانشگاه Pavia، از اعتبارسنجی متقابل ۵ لایه جهت جلوگیری از مسئله‌ی Overfitting و همچنین بدست آوردن پارامتر بهینه‌ی طبقه‌بندی KNN، یعنی تعداد بهینه‌ی همسایگی‌ها جهت طبقه‌بندی استفاده شد. جدول ۷ نتایج این اعتبارسنجی را نشان می‌دهد.

جدول ۷- پارامترهای بدست آمده از اعتبارسنجی متقابل ۵ لایه

روش	همسایگی بهینه	دقت اعتبارسنجی
MH	۴	۹۹/۹۷
FSDA	۴	۹۹/۹۶
Reg-MH-FSDA	۴	۹۹/۹۹

نتایج طبقه‌بندی KNN برای داده‌ی دانشگاه Pavia در همسایگی‌های بهینه و تعداد نمونه‌های آموزشی متفاوت طبق شکل ۱۳ بدست آمد.

۴- نتیجه گیری

همانطور که در مقدمه گفته شد، تصویربرداری ابرطیفی منبع ارزشمندی برای کاربردهای مختلف به شمار می‌رود. لذا بهبود کیفیت این تصاویر، رویکردی اقتصادی و عاقلانه به نظر می‌رسد. از زمان پیدایش تصاویر رقومی همواره سعی در بهبود کیفیت آن‌ها شده است و تصاویر ابرطیفی نیز از این قاعده مستثنی نیستند. به مرور زمان الگوریتم‌ها و روش‌های مختلفی برای غلبه بر مشکلات تصاویر ابرطیفی ارائه شده است. در سال‌های اخیر اقبال زیادی در مورد الگوریتم‌هایی که در پردازش خود مولفه‌های مکانی و طیفی را به طور توأم مد نظر قرار می‌دهند، وجود داشته است. یکی از جدیدترین و کارآمدترین این روش‌ها، روش پیش‌بینی فرضیه چندگانه است. نقطه ضعف این روش عدم استفاده از روشی موثر در انتخاب باندهای با شباهت بیشتر است. در این مقاله سعی شده است که روش پیش‌بینی فرضیه چندگانه مورد بررسی قرار گرفته و روشی مناسب برای انتخاب باندهای طیفی اتخاذ گردد.

در اینجا به دلیل انعطاف زیاد روش پیش‌بینی رگرسیونی در تعیین ضرایب شباهت بین باندی، برای

انتخاب باندهای طیفی مشابه، این روش را اتخاذ شد. داده های به کار گرفته شده در این تحقیق از داده های مختلف و متفاوت چه از لحاظ قدرت تفکیک مکانی و چه از لحاظ قدرت تفکیک طیفی استفاده شد. نتایج اخذ شده در این تحقیق بیانگر عملکرد بسیار خوب این روش است. به طوری که صحت کلی طبقه‌بندی ماشین بردار پشتیبان و k نزدیکترین همسایگی برای مجموعه داده‌های ابرطیفی Indian Pines و دانشگاه Pavia حاصل از روش پیشنهادی به ترتیب برابر با $95/82$ ، $99/43$ و $92/89$ و $98/88$ بدست آمد که در طبقه‌بندی SVM به ترتیب $0/4$ ، $0/3$ و در طبقه‌بندی KNN به ترتیب $8/22$ و 2 درصد افزایش را نشان می‌دهد که نشان دهنده‌ی کارآمدی روش پیشنهادی است. روش طبقه‌بندی KNN یکی از روش‌های ساده و کم هزینه از لحاظ زمان است، بنابراین بهبود دقت چنین روشی یک موفقیت بسیار خوب است.

سپاسگزاری

در اینجا بر خود لازم می‌دانم از جناب آقای دکتر چن از دانشگاه تگزاس آمریکا به خاطر راهنمایی هایشان در مورد الگوریتم پیش‌بینی فرضیه چندگانه تقدیر و تشکر بعمل آورم.

مراجع

- [1] Plaza, A., Benediktsson, J.A. Boardman, J.W. and Blackwell, W.J. (2009). "Recent advances in techniques for hyperspectral image processing." Remote Sens. Environ., Vol. 113, No. S1, pp. 110–122.
- [2] Zhao, Y., Zhang, L. and Kong, S.G. (2011). "Band-subset-based clustering and fusion for hyperspectral imagery classification." IEEE Trans. Geosci. Remote Sens., Vol. 49, No. 2, pp. 747–756.
- [3] Bishop, C.A., Liu, J.G. and Mason, P.J. (2011). "Hyperspectral remote sensing for mineral exploration in Pulong, Yunnan Province, China." Int. J. Remote Sens., Vol. 32, No. 9, pp. 2409–2426.
- [4] Yuen P.W. and Richardson, M. (2010). "An introduction to hyperspectral imaging and its application for security, surveillance and target acquisition." Imaging Sci. J., Vol. 58, No. 5, pp. 241–253.
- [5] Skauli, L. (2011). "Sensor noise informed representation of hyperspectral data with benefits for image storage and processing." Opt. Exp., Vol. 19, No. 14, pp. 13 031–13 046.
- [6] Acito, N., Diani, M. and Corsini, G. (2011). "Signal-dependent noise modeling and model parameter estimation in hyperspectral images." IEEE Trans. Geosci. Remote Sens., Vol. 49, No. 8, pp. 2957–2971.
- [7] Zhang, H. (2012). "Hyperspectral image denoising with cubic total variation model, ISPRS Annals of Photogrammetry." Remote Sensing and Spatial Information Sciences No. 7, pp. 95-98, 2012.
- [8] Hisham, O. and Qian, S. (2006). "Noise reduction of hyperspectral imagery using hybrid spatial-spectral derivative-domain wavelet shrinkage." Geoscience and Remote Sensing, IEEE Transactions on, Vol. 44, No. 2, pp. 397-408.
- [9] Guangyi, CH. and Qian, SH. (2008). "Simultaneous dimensionality reduction and denoising of hyperspectral imagery using bivariate wavelet shrinking and principal component analysis." Canadian Journal of Remote Sensing, Vol. 34, No. 5, pp. 447-454.
- [10] Guangyi, CH. and Qian, SH. (2009). "Denoising and dimensionality reduction of hyperspectral imagery using wavelet packets, neighbour shrinking and principal component analysis." International Journal of Remote Sensing, Vol. 30, No. 18, pp. 4889-4895.

- [11] Damien, L. and Bourennane, S. (2008). "Noise removal from hyperspectral images by multidimensional filtering." *Geoscience and Remote Sensing, IEEE Transactions on*, Vol. 46, No. 7, pp. 2061-2069.
- [12] Yi, W., Niu, R. and Yu, X. (2010). "Anisotropic diffusion for hyperspectral imagery enhancement." *Sensors Journal, IEEE*, Vol. 10, No. 3, pp. 469-477.
- [13] Guangyi, CH. and Qian, SH. (2011). "Denoising of hyperspectral imagery using principal component analysis and wavelet shrinkage." *Geoscience and Remote Sensing, IEEE Transactions on*, Vol. 49, No. 3, pp. 973-980.
- [14] Phillips, Rhonda, PH., Christine, D., Layne, E.L., Watson, T. and Randolph H. (2009). "Wynne, An adaptive noise-filtering algorithm for AVIRIS data with implications for classification accuracy." *Geoscience and Remote Sensing, IEEE Transactions on*, Vol. 47, No. 9, pp. 3168-3179.
- [15] Lei, J., Zhang, L., Rover, J., Wylie, B.K. and Chen, X. (2014). "Geostatistical estimation of signal-to-noise ratios for spectral vegetation indices." *ISPRS Journal of Photogrammetry and Remote Sensing*, No. 96, pp. 20-27.
- [16] Nicola, A. and Diani, M. (2014). "Defense & Security Mitigating the impact of signal-dependent noise on hyperspectral target detection." *SPIE Newsroom*, September 2014.
- [17] Green, A.A., Berman, M., Switzer, P., and Craig, M.D. (1988). "A transformation for ordering multispectral data in terms of image quality with implications for noise removal." *Geoscience and Remote Sensing, IEEE Transactions on*, Vol. 26, No. 1, pp. 65-74.
- [18] Aardt, V. and Wynne, R.H. (2007). "Examining pine spectral separability using hyperspectral data from an airborne sensor: An extension of field-based results." *International Journal of Remote Sensing*, Vol. 28, No. 2, pp. 431-436.
- [19] Phillips, R.D., Watson, L.T., Blinn, C.E., and Wynne, R.H. (2008). "An adaptive noise reduction technique for improving the utility of hyperspectral data." In *Proceedings of the 17th William T. Pecora Memorial Remote Sensing Symposium*, pp. 16-20.
- [20] Bourennane, S., Fossati, C., and Cailly, A. (2010). "Improvement of classification for hyperspectral images based on tensor modeling." *Geoscience and Remote Sensing Letters, IEEE*, Vol. 7, No. 4, pp. 801-805.
- [21] Qian, Y., Ye, M., and Zhou, J. (2013). "Hyperspectral image classification based on structured sparse logistic regression and three-dimensional wavelet texture features." *Geoscience and Remote Sensing, IEEE Transactions on*, Vol. 51, No. 4, pp 2276-2291.
- [22] Xu, D., Sun, L., and Luo, J. (2013). "Noise estimation of hyperspectral remote sensing image based on multiple linear regressions and wavelet transform." *Boletim de Ciências Geodésicas*, Vol. 19, No. 4, pp. 639-652.
- [23] Chen, C., Li, W., Tramel, E.W., Cui, M., Prasad, S., and Fowler, J.E. (2014). "Spectral-spatial preprocessing using multihypothesis prediction for noise-robust hyperspectral image classification." *Selected Topics in Applied Earth Observations and Remote Sensing, IEEE Journal of*, Vol. 7, No. 4, pp. 1047-1059.
- [24] Zhao, Y.Q., and Yang, J. (2015). "Hyperspectral image denoising via sparse representation and low-rank constraint." *Geoscience and Remote Sensing, IEEE Transactions on*, Vol. 53, No. 1, pp. 296-308.
- [25] http://www.ehu.eus/ccwintco/index.php?title=Hyperspectral_Remote_Sensing_Scenes
- [26] Chang, CH.CH. and Lin, CH.J. (2011). "LIBSVM: a library for support vector machines." *ACM Transactions on Intelligent Systems and Technology*, 2:27:1--27:27, 2011. Software available at <http://www.csie.ntu.edu.tw/~cjlin/libsvm>
- [27] Hansen P.C. and O'Leary, D.P. (1993). "The use of the L-curve in the regularization of discrete ill-posed problems," *SIAM J. Sci. Comput.*, Vol. 14, No. 6, pp. 1487-1503.
- [28] Tikhonov, A.N., and Arsenin, V.I. (1977). "Solutions of ill-posed problems." Vh Winston.
- [29] Imani, M., AND Ghassemian, H. (2015). "Feature space discriminant analysis for hyperspectral data feature reduction." *ISPRS Journal of Photogrammetry and Remote Sensing*, No. 102, pp. 1-13.
- [30] Du, Q., and Yang, H. (2008). "Similarity-based unsupervised band selection for hyperspectral image analysis." *Geoscience and Remote Sensing Letters, IEEE*, Vol. 5, No. 4, pp. 564-568.