

کاهش ابعاد داده‌های ابرطیفی به منظور افزایش جدایی پذیری طبقه‌ها و حفظ ساختار داده

مریم ایمانی و حسن قاسمیان*

دانشکده مهندسی برق و کامپیوتر، دانشگاه تربیت مدرس، تهران، ایران



چکیده

امروزه تصویربرداری ابرطیفی به منظور طبقه‌بندی داده‌های سطح زمین با دقت و جزئیات بالا بسیار مورد توجه است. به دلیل کمبود نمونه آموزشی در دسترس، کاهش ابعاد داده ابرطیفی به عنوان یک گام مهم پیش پردازش در تحلیل و طبقه‌بندی تصاویر ابرطیفی به شمار می‌رود. در این مقاله یک روش استخراج ویژگی پیشنهاد شده که سعی می‌کند، علاوه بر افزایش جدایی پذیری طبقه‌ها، ساختار داده را نیز حفظ کند. برای این منظور، دو تابع هدف پیشنهاد شده است. تابع هدف نخست از نمونه‌های آموزشی برچسب‌دار بهره می‌برد و سعی می‌کند نمونه‌های هم طبقه را در فضای کاهش یافته تا جای ممکن به هم نزدیک کند. تابع هدف دوم از نمونه‌های بدون برچسب خوشه‌بندی شده بهره برده و سعی می‌کند نمونه‌های متعلق به یک خوشه را در فضای کاهش یافته، تا جای ممکن به هم نزدیک گرداند. روش پیشنهادی بر روی سه داده ابرطیفی واقعی مورد آزمایش قرار گرفته و برتری آن از نظر دقت طبقه‌بندی نسبت به تعدادی از روش‌های پرکاربرد استخراج ویژگی نشان داده شده است.

واژگان کلیدی: ابعاد بالا، نمونه آموزشی کم، ابرطیفی، طبقه‌بندی، کاهش ویژگی

Feature reduction of hyperspectral data for increasing of class separability and preserving of data structure

Maryam Imani and Hassan Ghassemian*

Faculty of Electrical and Computer Engineering, Tarbiat Modares University, Tehran, Iran

Abstract

Hyperspectral imaging with gathering hundreds spectral bands from the surface of the Earth allows us to separate materials with similar spectrum. Hyperspectral images can be used in many applications such as land chemical and physical parameter estimation, classification, target detection, unmixing, and so on. Among these applications, classification is especially interested. A hyperspectral image is a cube data containing two spatial dimensions and a spectral one. Generally, the Hughes phenomenon is occurred in the supervised classification of hyperspectral images due to the limited available labeled samples and the curse of dimensionality. So, feature reduction is an important preprocessing step for analysis and classification of hyperspectral data. Feature reduction methods are categorized into feature selection approaches and feature extraction ones. Our main focus in this paper is on feature extraction. The feature extraction methods are also divided into three main groups: supervised (with labeled samples), unsupervised (without labeled samples), and semi-supervised (with both labeled and unlabeled samples). The first group of feature extraction methods usually suffers from problems due to limited available training samples. These methods often consider the separability between classes, and so are efficient for classification applications. The second group has no need for training samples, but they often do not consider the separability between

different classes and so, are not appropriate for classification. These methods are usually used for signal representation or preserving the local structure of data. The use of both labeled and unlabeled samples in the third group can increase the abilities of the feature extractor. A feature extraction method is proposed in this paper which belongs to the third group. The proposed method increases the class separability and tries to preserve the structure of data. The proposed feature extraction method uses the ability of unlabeled samples in addition to available limited training samples to improve the classification performance. The experimental results on three real hyperspectral images show the better performance of proposed method compared to some popular feature extraction methods in terms of classification accuracy.

Keywords: high dimension, small training set, hyperspectral, classification, feature reduction.

استاندارد، برای انتخاب ویژگی، به طور معمول از توابع معیار بر مبنای معیار فاصله‌ی آماری که جدایی‌پذیری میان توزیع طبقه‌ها را ارزیابی می‌نمایند، استفاده می‌کنند. از جمله این معیارها می‌توان به واگرایی^۸ و فاصله باتاچاریا^۹ اشاره کرد. بررسی الگوریتم‌های جستجو در انتخاب ویژگی در [9]-[8] انجام شده است. در [10] یک روش انتخاب ویژگی بر اساس الگوریتم جستجوی ژنتیک پیشنهاد شده است. مراجع [12]-[11] از معیار ارزیابی اطلاعات متقابل برای انتخاب زیرمجموعه مناسب ویژگی‌ها استفاده کرده‌اند. نویسندگان در [13] برای انتخاب زیرمجموعه مناسب باندها از ترکیب چندین معیار طبقه‌بندی از جمله جدایی‌پذیری طیفی شامل فاصله اقلیدسی، زاویه طیفی، همبستگی طیفی و هم‌چنین معیار اطلاعات متقابل شامل همبستگی باندها استفاده کرده‌اند. در [14] برای انتخاب ویژگی علاوه بر معیار جدایی‌پذیری طبقه‌ها از معیار دیگری نیز استفاده شده است که بر اساس این عقیده شکل گرفته است: نمونه‌های آموزشی هر طبقه که از جاهای مختلف تصویر انتخاب شده‌اند، در ویژگی‌های مورد نظر ممکن است به هم شبیه یا متفاوت باشند. هر چه تفاوت بین نمونه‌های آموزشی هر طبقه در ویژگی‌های انتخابی کمتر باشد، بهتر است و این باعث افزایش دقت طبقه‌بندی، افزایش پایوری و قدرت تعمیم طبقه‌بندی می‌شود.

دسته دوم روش‌های کاهش ویژگی، روش‌های استخراج ویژگی می‌باشند. در این دسته از روش‌ها، یک تبدیل خطی یا غیر خطی بر روی فضای ویژگی اولیه اعمال می‌شود تا ویژگی‌های مطلوب از آن استخراج شوند. تحلیل مؤلفه‌های اصلی^{۱۰} (PCA) و تحلیل ممیز خطی^{۱۱} (LDA) از جمله معروف‌ترین و پرکاربردترین روش‌های استخراج ویژگی در مسائل شناسایی الگو هستند [15]. اگر یک تصویر N پیکسل

۱- مقدمه

با پیشرفت فناوری سنجنده‌های اخذ تصاویر ابرطیفی^۱، امروزه این تصاویر به منظورهای مختلف شناسایی الگو، آشکارسازی سطح زمین، تحلیل نوع خاک، دیده‌بانی و کنترل مناطق جنگی و کشاورزی، مطالعات محیط زیست و غیره مورد استفاده فراوان قرار گرفته‌اند. تصاویر ابرطیفی که تصویر یک منطقه از سطح زمین را با قدرت تفکیک طیفی و مکانی بالا در صدها باند طیفی مجاور فراهم می‌کنند، قادر به تفکیک دقیق پدیده‌های مختلف سطح زمین و تشخیص جزئیات آن‌ها است؛ اما همه باندهای طیفی موجود در تصاویر ابرطیفی نمی‌توانند نقش مهمی در تشخیص و طبقه‌بندی داده ایفا کنند؛ چون باندهای طیفی مجاور، حاوی اطلاعات زائد و هم‌پوشان هستند. به علاوه، برای تحلیل و طبقه‌بندی حجم زیادی از ویژگی‌های طیفی به تعداد زیادی از نمونه‌های آموزشی نیاز است. از طرف دیگر، جمع‌آوری و اخذ داده آموزشی دشوار و نیازمند صرف زمان و هزینه است. تحلیل و طبقه‌بندی داده با ابعاد بالا و نمونه‌ی آموزشی محدود منجر به پدیده‌ی هیوز^۲ خواهد شد [1]. به این معنی که با افزایش بعد داده و نمونه آموزشی ثابت، تا یک جایی دقت افزایش و از آن جا به بعد دقت کاهش می‌یابد. کاهش ویژگی سبب افزایش دقت طبقه‌بندی، کاهش حجم داده و تفسیر راحت‌تر داده ابرطیفی خواهد شد. کاهش بعد به دو روش کلی انتخاب ویژگی^۳ [2] و استخراج ویژگی^۴ [5]-[3] امکان‌پذیر است.

در فرآیند انتخاب ویژگی، از میان ویژگی‌های اصلی داده، زیر مجموعه مناسبی از آن‌ها انتخاب می‌شود [6]. روش‌های انتخاب ویژگی به سه گروه کلی دسته‌بندی می‌شوند: ۱- مدل فیلتر^۵ ۲- مدل پوشه^۶ ۳- مدل ترکیبی^۷ [7]. روش‌های فیلتر

¹ Hyperspectral

² Hughes phenomenon

³ Feature selection

⁴ Feature extraction

⁵ Filter

⁶ Wrapper

⁷ Hybrid

⁸ Divergence

⁹ Bhattacharyya distance

¹⁰ Principal Component Analysis

¹¹ Linear Discriminant Analysis

به نمونه آموزشی زیاد برای یادگیری، با استفاده از نمونه آموزشی محدود کارا نیست. استخراج توأم ویژگی‌های طیفی و مکانی نیز در [26] پیشنهاد و بررسی شده است.

روش‌های استخراج ویژگی خطی از طریق نگاشت داده از فضای ویژگی اولیه به فضای ویژگی با بعد بالاتر، می‌توانند به روش‌های غیرخطی بسط داده شوند. از توابع هسته^۲ برای محاسبه ضرب داخلی در فضای بعد بالاتر استفاده می‌شود. تحلیل ممیز تعمیم یافته^۳ (GDA) بسط غیرخطی LDA است [27]. روش GDA هم مانند LDA بیشینه‌ساز به استخراج $n_c - 1$ ویژگی است. روش استخراج ویژگی وزن‌دار غیرپارامتریک^۴ (NWFE)، یک روش نظارت‌شده و غیرپارامتریک است و اساس آن اختصاص وزن‌های متفاوت به هر نمونه برای محاسبه میانگین وزن‌دار و تعریف ماتریس‌های پراکندگی درون طبقه‌ای و بین طبقه‌ای غیرپارامتریک جدید برای استخراج بیش از $n_c - 1$ ویژگی می‌باشد [28]. عیب اساسی روش NWFE، زمان محاسباتی بسیار زیاد آن است.

روش پیشنهادی [29] تحلیل تمیز محلی نیمه نظارت شده^۵ (SELD) است که سعی می‌کند هم همسایگی محلی را حفظ و هم تمیز بین طبقه‌ها را بیشینه کند و برای این کار یک روش بدون نظارت مثل جاسازی حفظ همسایگی (NPE^۶) را با یک روش نظارت شده مثل LDA ترکیب می‌کند و هیچ پارامتر آزادی که لزوم به تنظیم آن باشد، نیز ندارد.

روش SELD برخلاف روش‌های نیمه نظارت شده دیگر، ترکیب نظارت شده و بدون نظارت را به صورت خطی انجام نمی‌دهد. روش SELD به جای این که نمونه‌های برچسب‌دار و بدون برچسب را با هم استفاده کند، ابتدا نمونه‌ها را به دو دسته‌ی برچسب‌دار و بدون برچسب تقسیم می‌کند و سپس از نمونه‌های برچسب‌دار فقط در روش نظارت شده و از نمونه‌های بدون برچسب فقط در روش بدون نظارت استفاده می‌کند. رگرسیون خطی روشی برای مدل کردن رابطه خطی بین متغیر وابسته و یک یا چند متغیر مستقل است. در مدل رگرسیون خطی، داده با استفاده از توابع پیشگوی خطی مدل شده و پارامترهای ناشناخته مدل از داده تخمین زده می‌شود. به طور معمول این پارامترها با استفاده از روش کمینه مربعات (LS^۷) به دست می‌آیند. روش استخراج ویژگی بر مبنای رگرسیون ستیغی (FERR^۸) که در [30] پیشنهاد شده است،

d باند دارد، در بیان برداری N بردار d بعدی یا d بردار N تایی داریم که اولی بیان طیفی و دومی بیان مکانی تصویر هستند. در بیان تنسور، به جای استفاده از بردار، برای بیان تصویر ابرطیفی از یک ماتریس دوبعدی استفاده می‌کنیم؛ یعنی بیان طیفی و مکانی کنار هم. استخراج ویژگی‌های طیفی و مکانی با استفاده از بیان تنسور در [16] پیشنهاد شده است. نویسندگان در [17] منحنی طیفی هر پیکسل از داده ابرطیفی را مانند یک سری زمانی در نظر گرفته و دو ویژگی سری زمانی که یکی ویژگی محلی و دیگری ویژگی کلی را بیان می‌کند، جایگزین آن کرده‌اند؛ یعنی ابعاد داده را از d به دو کاهش داده‌اند. برای استفاده از طبیعت غیرخطی داده‌ی ابرطیفی، می‌توان الگوریتم کاهش ویژگی را به جای فضای اقلیدسی در فضای غیراقلیدسی اعمال کرد. استفاده از فضای کروی برای استخراج ویژگی و طبقه‌بندی داده ابرطیفی در [19]-[18] بیان شده است. روش پیشنهادی برای استخراج ویژگی در [20] بر اساس بخش‌بندی منحنی طیفی هر پیکسل است. در این روش، منحنی طیفی هر پیکسل، به چند بخش تقسیم شده و سپس برای هر بخش، دو ویژگی که یکی موقعیت (میانگین) و دیگری شکل (واریانس) آن بخش را تعیین می‌کنند، استخراج شده است. اگر منحنی طیفی هر نمونه به k بخش تقسیم شود، $2k$ ویژگی برای آن نمونه استخراج خواهد شد. تعداد دیگری از روش‌های استخراج ویژگی در [21]-[22] بیان شده‌اند.

استفاده از نمونه‌های آموزشی مجازی، راه دیگری برای مقابله با تعداد محدود نمونه‌های آموزشی است. روش پیشنهادی در [23]، جاسازی نزدیک‌ترین خط ویژگی^۱ (NFLE) نام دارد و از مفاهیم خط ویژگی برای تولید نمونه‌های آموزشی مجازی استفاده می‌کند و از این نمونه‌های مجازی تولید برای تولید ماتریس‌های پراکندگی درون طبقه‌ای و بین طبقه‌ای استفاده می‌کند. نویسندگان در [24] برای استخراج ویژگی نظارت شده از مفهوم اطلاعات متقابل استفاده کرده‌اند. از آنجایی که تخمین اطلاعات متقابل، نسبت به آمارگان مرتبه نخست و دوم، به نمونه آموزشی بسیار بیشتری احتیاج دارد، در نتیجه، این روش با استفاده از اندازه کوچک نمونه آموزشی، قابل اجرا نخواهد بود. روش استخراج ویژگی نظارت شده دیگری با استفاده از شبکه‌های عصبی در [25] پیشنهاد شده است که این روش هم به دلیل نیاز شبکه عصبی

⁵ Semi-supervised local discriminant analysis

⁶ Neighborhood preserving embedding

⁷ Least square

⁸ Feature extraction based on ridge regression

¹ Nearest Feature Line Embedding

² Kernel

³ Generalized Discriminant Analysis

⁴ Nonparametric Weighted Feature Extraction

$$Y = [y_1, y_2, \dots, y_n]_{d \times n} \quad (4)$$

می توان تابع Φ را به شکل ماتریسی بازنویسی کرد.

$$\Phi = \text{tr}(YG_l Y^T) \quad (5)$$

$$\min_A [\Phi = \text{tr}(A_l X G_l X^T A_l^T)] \quad (6)$$

که در روابط بالا، $Y = A_l X$ و $G_l = D_l - W_l$ می باشد. W_l ماتریس شباهت است که عناصر آن را $w_{ij}^l (i = 1:n, j = 1:n)$ تشکیل می دهند و D_l یک ماتریس قطری است که (مان های روی قطر آن، مجموع سطری ماتریس W_l است. برای استخراج m ویژگی از d ویژگی اولیه ($m < d$)، تنها کافی است که m بردار ویژه متناظر با m تا کوچک ترین مقادیر ویژه ماتریس $X G_l X^T$ سطرهای ماتریس A_l را تشکیل دهند.

ماتریس تبدیل A_u در روش پیشنهادی از حل مسئله بهینه سازی زیر به دست می آید:

$$\min (\Phi = \sum_{i=1}^n \sum_{j=1}^n \|y_i - y_j\|^2 w_{ij}^u) \quad (7)$$

که

$$w_{ij}^u = \begin{cases} 1 & x_i^u, x_j^u \text{ belong to a cluster} \\ 0 & x_i^u, x_j^u \text{ do not belong to a cluster} \end{cases} \quad (8)$$

در رابطه بالا، x_i^u و x_j^u دو نمونه بدون برچسب هستند. در روش پیشنهادی، ما به همان تعداد نمونه آموزشی (n)، نمونه بدون برچسب در نظر می گیریم. هم چنین برای انجام خوشه بندی از روش پر کاربرد و معمول k-means (به دلیل سادگی و کارایی خوب آن) با $k = n_c$ بهره برده ایم. با توجه به اینکه، نتیجه خوشه بندی k-means به انتخاب مراکز خوشه ی اولیه وابسته است و این مراکز در هر بار اجرای الگوریتم k-means به طور تصادفی از کل داده انتخاب می شود، ما ده مرتبه الگوریتم k-means را تکرار کرده و از نتایج خوشه بندی که بالاترین دقت طبقه بندی را به دست آورده اند، بهره برده ایم.

توجه کنید که در رابطه (۲)، در صورتی که دو نمونه آموزشی، متعلق به یک طبقه باشند، یک لبه با وزن یک بین آنها قرار می گیرد و در رابطه (۸)، زمانی که دو نمونه بدون برچسب متعلق به یک خوشه باشند، یک لبه با وزن یک بین

به ازای هر باند طیفی یک بردار ویژگی تعریف کرده و رابطه بین بردارهای ویژگی را با استفاده از رگرسیون، مدل می کند. روش پیشنهادی استخراج ویژگی در این مقاله سبب افزایش جدایی پذیری طبقه ها و حفظ ساختار داده در فضای کاهش یافته می شود. روش پیشنهادی، به علت تعداد محدود نمونه های آموزشی در دسترس، علاوه بر نمونه های آموزشی، از قدرت نمونه های بدون برچسب نیز به منظور افزایش دقت طبقه بندی بهره می برد. در ادامه این مقاله در بخش ۲، روش پیشنهادی با جزئیات بیشتر شرح داده می شود و در بخش ۳، نتایج آزمایش ها بیان خواهد شد. در نهایت، بخش ۴ به جمع بندی مقاله خواهد پرداخت.

۲- روش پیشنهادی

در روش پیشنهادی، دو دسته داده خواهیم داشت:

۱- مجموعه نمونه آموزشی که ماتریس تبدیل A_l را با استفاده از آن می سازیم.

۲- مجموعه نمونه های بدون برچسب که ماتریس تبدیل A_u را با استفاده از آن می سازیم.

در ابتدا زیر مجموعه ای از نمونه های بدون برچسب را به دلخواه و به صورت تصادفی انتخاب می کنیم و یک خوشه بندی بر روی آنها انجام می دهیم. به این ترتیب، مجموعه نمونه های بدون برچسب را در چندین خوشه جای می دهیم. داده در فضای ورودی را با x_i (یک بردار d بعدی) نشان می دهیم که d تعداد باندهای طیفی (ویژگی های اولیه) است و $y_i = A x_i$ را تبدیل یافته x_i با بعد کاهش یافته در نظر می گیریم. ماتریس تبدیل A_l در روش پیشنهادی از حل مسئله بهینه سازی زیر به دست می آید:

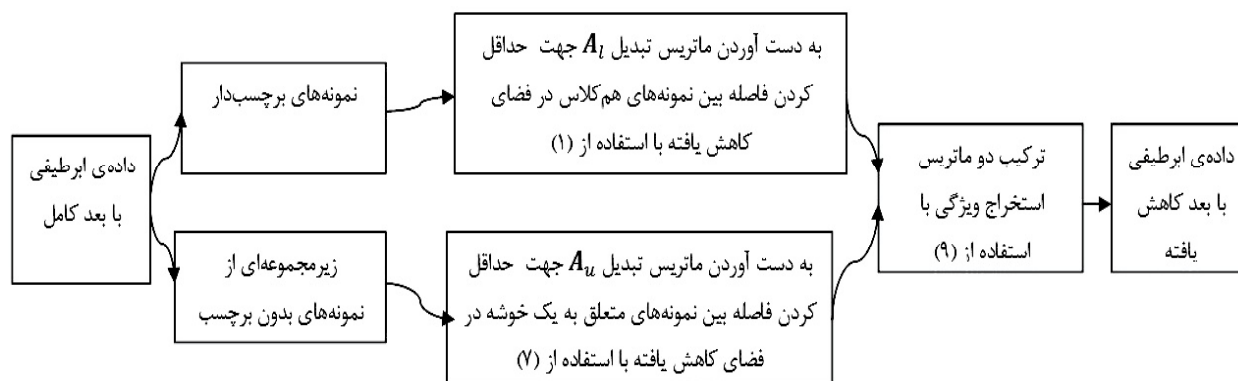
$$\min (\Phi = \frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n \|y_i - y_j\|^2 w_{ij}^l) \quad (9)$$

که

$$w_{ij}^l = \begin{cases} 1 & x_i, x_j \text{ belong to a class} \\ 0 & x_i, x_j \text{ do not belong to a class} \end{cases} \quad (2)$$

در رابطه بالا، $n = \sum_{c=1}^{n_c} n_{tc}$ تعداد کل نمونه های آموزشی، n_c تعداد طبقه ها و n_{tc} تعداد نمونه های آموزشی کلاس c است. اگر برای داده با بعد بالا در فضای اولیه و داده با بعد کاهش یافته در فضای ثانویه، به ترتیب فرم های ماتریسی X و Y را به شکل زیر در نظر بگیریم:

$$X = [x_1, x_2, \dots, x_n]_{d \times n} \quad (3)$$



(شکل-۱): نمودار جعبه‌ای روش پیشنهادی
(Figure-1): Block diagram of the proposed method

ابرطیفی دانشگاه پاویا از محوطه دانشگاه پاویا در ایتالیا با استفاده از سنجنده نوری ROSIS اخذ شده است. این تصویر، دارای تفکیک مکانی ۱/۳ متر بر پیکسل و دارای ۱۱۵ باند طیفی در بازه فرکانسی ۰/۴۳ تا ۰/۸۶ میکرومتر بوده و شامل ۹ طبقه است. ابعاد تصویر دانشگاه پاویا ۶۱۰×۳۴۰ است. بعد از حذف کانال‌های نوفه‌ای، آزمایش‌ها بر روی ۱۰۳ باند باقیمانده انجام شده است. داده ابرطیفی سالیانس جنوبی سنجنده AVIRIS از دهکده سالیانس در کالیفرنیا جنوبی اخذ شده است. این تصویر ۲۱۷×۵۱۲ بوده، شامل شانزده طبقه، و سائز پیکسل ۳٫۷ متر است. داده سالیانس حاوی ۲۲۴ باند طیفی در بازه فرکانسی ۰٫۴ تا ۲٫۵ میکرومتر، با قدرت تفکیک طیفی ۱۰ نانومتر است که بعد از حذف بیست باند جذب آب، ۲۰۴ باند از این تصویر در آزمایش‌ها مورد استفاده قرار می‌گیرند.

در این مقاله از چندین معیار مختلف برای ارزیابی کارایی طبقه‌بندی استفاده می‌کنیم: دقت متوسط، اعتبار متوسط و ضریب کاپا [31]. دقت متوسط، میانگین دقت به‌دست آمده برای تمام طبقه‌ها و اعتبار متوسط برابر میانگین اعتبار به‌دست‌آمده برای تمام طبقه‌ها است. اعتبار برای هر کلاس این‌گونه تعریف شده است: تعداد نمونه‌هایی که درست طبقه‌بندی شده‌اند، تقسیم بر تعداد کل نمونه‌هایی که به آن طبقه تعلق گرفته‌اند.

ضریب کاپا^۴، یکی از پارامترهای دقت است که دقت طبقه‌بندی را نسبت به یک طبقه‌بندی به‌طور کامل تصادفی محاسبه می‌کند. به این معنی که مقدار کاپا، دقت طبقه‌بندی را نسبت به حالتی که یک تصویر به‌طور کامل به‌صورت تصادفی طبقه‌بندی شود، به دست می‌دهد. به این ترتیب، یک کاپا

آن‌ها در نظر گرفته می‌شود. مشابه روند قبلی انجام شده برای نمونه‌های آموزشی، ماتریس تبدیل A_u نیز محاسبه می‌شود. در نهایت ماتریس تبدیل نهایی که برای استخراج ویژگی داده مورد استفاده قرار می‌گیرد، از رابطه زیر محاسبه می‌شود:

$$A = \alpha A_l + (1 - \alpha) A_u \quad (9)$$

که در رابطه بالا پارامتر $0 \leq \alpha \leq 1$ به‌منظور ایجاد مصالحه بین نقش نمونه‌های آموزشی و نمونه‌های بدون برچسب در تهیه ماتریس تبدیل استخراج ویژگی استفاده می‌شود. با استفاده از ماتریس تبدیل A هر نمونه در فضای کاهش یافته، یک بردار m بعدی به شکل $y_{m \times 1} = A_{m \times d} x_{d \times 1}$ خواهد بود. نمودار جعبه‌ای روش پیشنهادی در شکل (۱) مشاهده می‌شود.

۳- ارزیابی و آزمایش‌ها

ما برای انجام آزمایش‌های خود در این قسمت از سه داده ابرطیفی ایندیانا^۱، دانشگاه پاویا^۲ و سالیانس^۳ استفاده کرده‌ایم. داده ابرطیفی ایندیانا مربوط به منطقه‌ای کشاورزی- جنگلی است که توسط سنجنده AVIRIS از شمال شرقی ایالت ایندیانا گرفته شده است. این داده حاوی ۲۲۰ باند طیفی با پهنای ده نانومتر در بازه فرکانسی ۰٫۴ تا ۲٫۵ میکرومتر، ابعاد ۱۴۵×۱۴۵ پیکسل، دقت مکانی بیست متر بر پیکسل و دقت رادیومتریک هشت بیت است. بیست کانال نوفه‌ای این داده حذف شده و آزمایش‌ها بر روی دویست باند باقیمانده انجام می‌گیرد. این داده حاوی شانزده طبقه است. داده

³ Salinas

⁴ Kappa Coefficient (KC)

¹ Indian

² University of Pavia

معادل ۷۵ درصد به این معنی است که نتایج طبقه‌بندی ۷۵ درصد بهتر از موقعی است که پیکسل‌ها به‌طور تصادفی برچسب‌دهی شوند. مقدار صفر برای ضریب کاپا به این معنی است که طبقه‌بندی بدون هیچ ضابطه‌ای و به‌طور کامل تصادفی انجام شده است. مقدار یک به معنی یک طبقه‌بندی به‌طور کامل صحیح بر اساس نمونه‌های گرفته شده است. مقادیر منفی کاپا به معنی ضعف طبقه‌بندی و نتایج بسیار بد تفسیر می‌شود. به‌طور معمول این‌گونه عنوان می‌شود که ضریب کاپا برآوردی بدبینانه بوده و دقت را کمتر از مقدار واقعی بیان می‌کند. ضریب کاپا به‌شکل زیر محاسبه می‌شود:

$$KC = \frac{N \sum_{c=1}^{n_c} t_{cc} - \sum_{c=1}^{n_c} t_{c+} t_{+c}}{N^2 - \sum_{c=1}^{n_c} t_{c+} t_{+c}} \quad (10)$$

در رابطه بالا، N و n_c به‌ترتیب تعداد نمونه‌های آزمایشی و تعداد طبقه‌ها هستند. t_{cc} تعداد نمونه‌هایی است که به‌طور صحیح در طبقه c طبقه‌بندی شده‌اند، t_{c+} تعداد نمونه‌های آزمایشی است که برچسب طبقه c را خورده‌اند. t_{+c} تعداد نمونه‌هایی است که پیش‌گویی می‌شود به طبقه c تعلق دارند. آزمون MCNemars برای بیان معناداری تفاوت بین روش‌های طبقه‌بندی از لحاظ آماری بسیار مفید است [32]. این آزمون برای مقایسه بین یک جفت طبقه‌بند مورد استفاده قرار می‌گیرد. پارامتر Z_{12} به‌صورت زیر تعریف می‌شود:

$$Z_{12} = \frac{f_{12} - f_{21}}{\sqrt{f_{12} + f_{21}}} \quad (11)$$

f_{12} تعداد نمونه‌هایی است که توسط طبقه‌بند ۱ درست طبقه‌بندی و توسط طبقه‌بند ۲ به‌اشتباه طبقه‌بندی شده‌اند. تفاوت بین طبقه‌بندهای ۱ و ۲ از نظر آماری، با اهمیت گفته می‌شود اگر $|Z_{12}| > 1.96$. علامت پارامتر Z_{12} می‌گوید که طبقه‌بند ۱ دقیق‌تر از طبقه‌بند ۲ است ($Z_{12} > 0$) و برعکس ($Z_{12} < 0$).

از آن جایی که به‌طور معمول تعداد نمونه‌های آموزشی در دسترس محدود است، ما در این مقاله برای بررسی کارایی روش‌های استخراج ویژگی با استفاده از مجموعه نمونه‌ی آموزشی کوچک، تنها از شانزده نمونه‌ی آموزشی در هر طبقه بهره برده‌ایم. نمونه‌های آموزشی به‌طور تصادفی از کل صحنه انتخاب شده‌اند. هر آزمایش بیست بار تکرار شده (با بیست مجموعه نمونه‌ی آموزشی متفاوت) و میانگین نتایج حاصل از آزمایش‌ها گزارش شده است.

بعد داده‌های ابرطیفی ایندینا، پاولا و سالیانس را به‌ترتیب از ۲۰۰، ۱۰۳ و ۲۰۴ به ۶ کاهش داده و سپس داده با بعد

کاهش‌یافته را توسط طبقه‌بند بردار پشتیبان^۱ (SVM) طبقه‌بندی و نتایج حاصل از آن را بیان کرده‌ایم. همه آزمایش‌ها با استفاده از نرم‌افزار MATLAB 7.1 انجام شده است.

$\alpha = 1$ به معنای استفاده تنها از نمونه‌های آموزشی و $\alpha = 0$ به معنی استفاده تنها از نمونه‌های آزمایشی است. ما مقدار پارامتر α را در بازه ۰ تا ۱ تغییر داده و دقت متوسط طبقه‌بندی حاصل را برای هر سه داده ابرطیفی مورد آزمایش، به‌ازای شش ویژگی استخراج‌شده، در شکل (۲) نشان داده‌ایم. مقدار مناسب پارامتر α برای سه داده‌ی ایندینا، پاولا و سالیانس به‌ترتیب برابر ۰،۴، ۰،۶ و ۰،۵ به‌دست آمد. همان‌طور که از شکل (۲) پیدا است، کمترین دقت طبقه‌بندی به‌ازای $\alpha = 0$ به‌دست می‌آید، یعنی زمانی که از هیچ داده آموزشی (اطلاعات با ارزش نمونه‌های برچسب‌دار) بهره‌ای نمی‌گیریم. بهترین دقت طبقه‌بندی حول مقدار $\alpha = 0.5$ به‌دست می‌آید؛ یعنی زمانی که علاوه بر اطلاعات با ارزش نمونه‌های برچسب‌دار، از اطلاعات با ارزش موجود در نمونه‌های بدون برچسب خوشه‌بندی‌شده نیز استفاده می‌کنیم. به‌عنوان نمونه، دقت متوسط طبقه‌بندی را در برابر تعداد ویژگی‌های استخراج‌شده (m) با استفاده از روش‌های مختلف استخراج ویژگی برای داده ایندینا به‌دست آورده‌ایم. نتایج حاصل از آزمایش‌ها در شکل (۳) مشاهده می‌شوند.

نتایج طبقه‌بندی داده‌ی ایندینا، پاولا و سالیانس با شش ویژگی استخراجی به‌ترتیب در جدول‌های (۱) تا (۳) مشاهده می‌شوند. جدول (۴)، نتایج آزمون MCNemars را نشان می‌دهد و شکل‌های ۴-۶ نقشه‌ی واقعی زمین^۲ (GTM) و نقشه‌های طبقه به‌دست‌آمده با استفاده از روش‌های مختلف را نمایش می‌دهند. در رابطه با روش‌های استخراج ویژگی مورد مقایسه، توجه به نکات زیر ارزشمند است:

- همان‌طور که در جدول‌های (۱) تا (۳) مشاهده می‌شود، کمترین دقت طبقه‌بندی توسط روش LDA به‌دست آمده است. این روش، با استفاده از تعداد نمونه‌ی آموزشی کم دارای کارایی بسیار ضعیفی است.
- روش‌های LDA، GDA و NWFE روش‌هایی نظارت‌شده هستند که تنها از نمونه‌های آموزشی در دسترس برای استخراج ویژگی استفاده می‌کنند. این روش‌ها با تخمین ماتریس‌های پراکندگی، سعی در کمینه‌کردن فاصله‌های درون گروهی و بیشینه‌کردن فاصله‌های بین گروهی دارند.

² Ground Truth Map

¹ Support Vector Machine

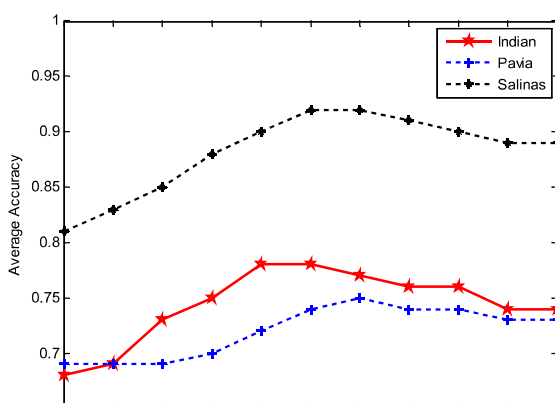
روش پیشنهادی و سایر روش‌ها همواره مقداری مثبت و بزرگتر از ۱/۹۶ را حاصل کرده است.

ما علاوه بر روش‌های استخراج ویژگی بیان شده، روش پیشنهادی خود را با دو روش استخراج ویژگی که در همین اواخر پیشنهاد شده‌اند، نیز مقایسه کرده‌ایم. روش استخراج ویژگی مورد مقایسه نخست روش SELD است که علاوه بر نمونه‌های برچسب‌دار از قدرت نمونه‌های بدون برچسب نیز بهره می‌برد و روش دوم مورد مقایسه، روش FERR است که از رگرسیون برای مدل‌کردن رابطه باند‌های طیفی نسبت به یکدیگر عمل می‌کند. دقت متوسط طبقه‌بندی در برابر تعداد ویژگی‌های استخراج شده برای داده پاویا در شکل (۷) نشان داده شده است.

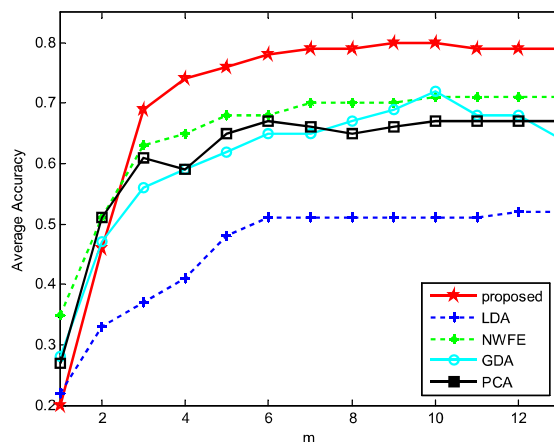
به این ترتیب، این روش‌ها، بر اساس معیار جدایی‌پذیری طبقه‌ها، ویژگی‌ها را استخراج می‌کنند.

- روش PCA یک روش بدون نظارت است که از اطلاعات نمونه‌های آموزشی بهره‌ای نمی‌برد. روش PCA با استخراج مؤلفه‌های اصلی که بیشترین قدرت و واریانس را دارند، سعی در حفظ ساختار داده در فضای بعد کاهش یافته می‌کند.
- روش پیشنهادی، علاوه بر استفاده از نمونه‌های آموزشی برچسب‌دار، از قدرت نمونه‌های آموزشی بدون برچسب نیز بهره می‌برد. روش پیشنهادی، با در نظر گرفتن برچسب طبقه‌ها در تولید گراف داده و محاسبه ماتریس شباهت، علاوه بر افزایش جدایی‌پذیری سعی در حفظ ساختار داده می‌کند. به این ترتیب دقت طبقه‌بندی با استفاده از روش پیشنهادی افزایش می‌یابد. نه تنها در جدول‌های (۱) تا (۳)، بالاترین دقت طبقه‌بندی توسط روش پیشنهادی به دست آمده، بلکه اعداد جدول (۴) نشان می‌دهند که این برتری از نظر آماری نیز معنادار و قابل توجه است. پارامتر Z بین

(شکل-۳): دقت طبقه‌بندی در برابر پارامتر m
(Figure-3): Classification accuracy versus parameter m



(شکل-۲): دقت طبقه‌بندی در برابر پارامتر α
(Figure-2): Classification accuracy versus parameter α



(جدول-۱): نتایج طبقه‌بندی داده‌ی ایندیانا
(Table-1): Classification results for Indian dataset

class			Proposed		LDA		NWFE		GDA		PCA	
No	Name of class	# samples	Acc.	Rel.	Acc.	Rel.	Acc.	Rel.	Acc.	Rel.	Acc.	Rel.
1	یونجه	46	94.0	26.0	70.0	11.0	93.0	14.0	94.0	20.0	81.0	19.0
2	ذرت-بدون شخم	1428	61.0	45.0	32.0	37.0	60.0	56.0	57.0	39.0	45.0	44.0
3	ذرت-کم شخم	830	76.0	57.0	24.0	17.0	40.0	24.0	38.0	30.0	43.0	29.0
4	ذرت	237	79.0	42.0	46.0	14.0	77.0	40.0	58.0	35.0	76.0	43.0
5	سیزه-درختان	483	85.0	78.0	46.0	50.0	58.0	32.0	62.0	48.0	72.0	61.0
6	سیزه-چمنزار	730	91.0	84.0	50.0	51.0	82.0	87.0	65.0	68.0	64.0	74.0
7	سیزه-چمنزار کوتاه	28	00.1	53.0	88.0	13.0	92.0	51.0	92.0	55.0	96.0	58.0
8	کاه و خاشاک	478	79.0	00.1	44.0	77.0	79.0	99.0	79.0	99.0	74.0	96.0

9	جو دوسر	20	00.1	63.0	90.0	06.0	00.1	44.0	00.1	16.0	00.1	30.0
10	سویا-بدون شخم	972	60.0	61.0	31.0	22.0	54.0	63.0	74.0	45.0	67.0	54.0
11	سویا-کم شخم	2455	43.0	80.0	18.0	52.0	33.0	63.0	20.0	64.0	40.0	70.0
12	سویا-شخم کامل	593	55.0	57.0	18.0	21.0	55.0	37.0	17.0	24.0	23.0	23.0
13	گندم	205	98.0	95.0	78.0	50.0	89.0	68.0	94.0	72.0	75.0	79.0
14	بیشه	1265	86.0	95.0	58.0	90.0	56.0	83.0	88.0	91.0	88.0	91.0
15	ساختمان-سبزه	386	54.0	54.0	42.0	25.0	27.0	58.0	21.0	56.0	31.0	32.0
16	سنگ-برج	93	91.0	46.0	87.0	26.0	98.0	48.0	87.0	48.0	89.0	40.0
Average Acc. and Average Rel.			0.78	0.65	0.51	0.35	0.68	0.54	0.65	0.51	0.67	0.53
Kappa coefficient			0.62		0.29		0.48		0.47		0.50	

(جدول-۲): نتایج طبقه‌بندی داده دانشگاه پاولیا
(Table-2): Classification results for University of Pavia dataset

class			Proposed		LDA		NWFE		GDA		PCA	
No	Name of class	# samples	Acc.	Rel.	Acc.	Rel.	Acc.	Rel.	Acc.	Rel.	Acc.	Rel.
1	آسفالت	6631	0.51	0.89	0.31	0.55	0.64	0.77	0.52	0.84	0.15	0.54
2	چمنزار	18649	0.62	0.89	0.53	0.84	0.54	0.82	0.70	0.81	0.75	0.86
3	شن و سنگ ریزه	2099	0.60	0.46	0.49	0.31	0.56	0.49	0.47	0.42	0.71	0.50
4	درختان	3064	0.63	0.81	0.60	0.67	0.65	0.75	0.56	0.74	0.68	0.72
5	صفحات فلزی	1345	1.00	0.98	0.89	1.00	0.99	0.99	0.99	0.86	0.99	0.96
6	زمین دست نخورده	5029	0.89	0.41	0.60	0.22	0.75	0.32	0.54	0.34	0.72	0.47
7	قبر	1330	0.86	0.30	0.28	0.14	0.87	0.40	0.86	0.30	0.86	0.18
8	آجر	3682	0.64	0.64	0.26	0.42	0.59	0.73	0.69	0.68	0.57	0.79
9	سایه	947	0.99	0.98	0.79	0.60	0.99	1.00	1.00	1.00	1.00	1.00
Average Acc. and Average Rel.			0.75	0.71	0.53	0.53	0.73	0.70	0.70	0.67	0.71	0.67
Kappa coefficient			0.58		0.38		0.54		0.55		0.55	

(جدول-۳): نتایج طبقه‌بندی داده سالیناس
(Table-3): Classification results for Salinas dataset

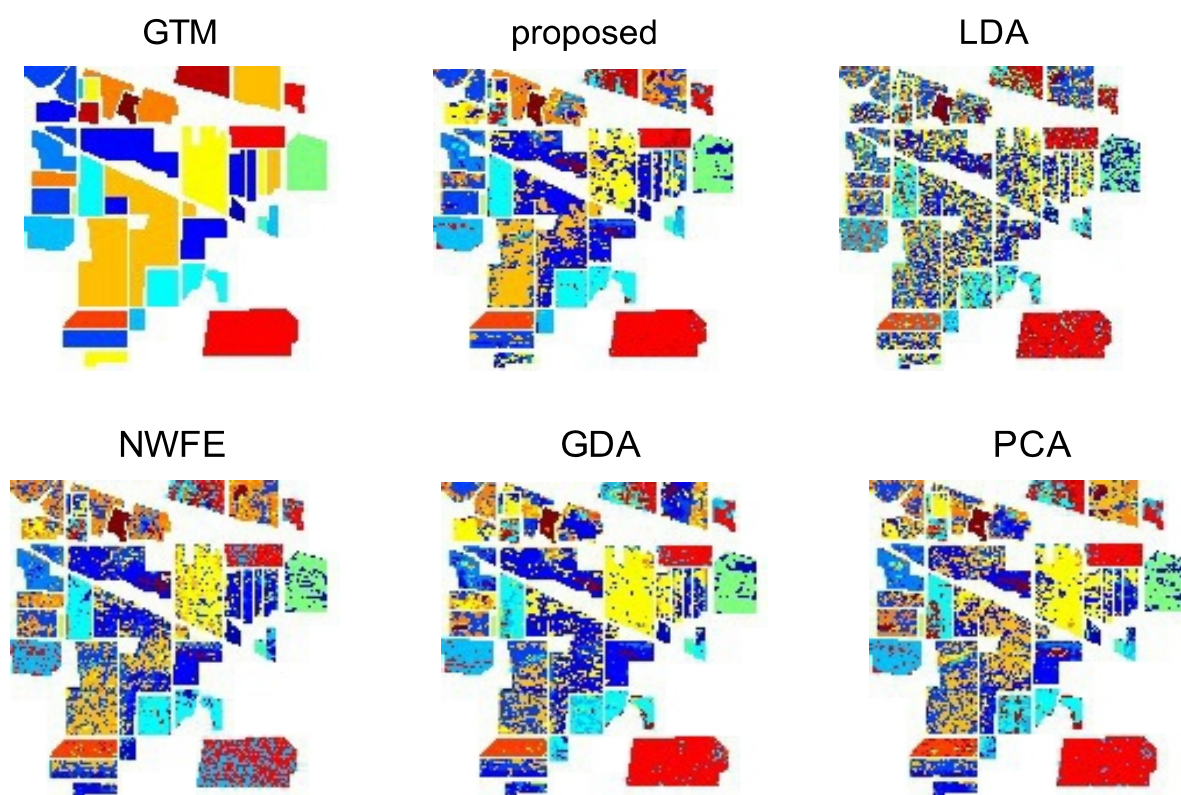
class			Proposed		LDA		NWFE		GDA		PCA	
No	Name of class	# samples	Acc.	Rel.	Acc.	Rel.	Acc.	Rel.	Acc.	Rel.	Acc.	Rel.
1	بروکلی- سبزه-علف هرز ۱	2009	0.96	1.00	0.96	0.99	0.98	0.91	0.95	1.00	0.96	0.98
2	بروکلی-سبزه-علف هرز ۲	3726	0.99	0.98	0.96	0.98	0.93	0.98	0.99	0.96	0.98	0.97
3	زمین شخم شده	1976	0.95	0.91	0.67	0.55	0.97	0.81	0.84	0.79	0.97	0.93
4	زمین شخم شده-زیر-گاواهن	1394	0.99	0.97	0.67	0.81	0.98	0.92	0.99	0.94	0.99	0.99
5	زمین شخم شده-هموار	2678	0.95	0.96	0.71	0.72	0.95	0.97	0.92	0.89	0.97	0.98
6	کاه بن	3959	1.00	1.00	0.93	1.00	0.99	1.00	0.99	0.98	0.98	0.99
7	کرفس	3579	0.99	0.93	0.97	1.00	0.98	0.94	0.96	0.93	0.98	0.93
8	انگور	11271	0.74	0.74	0.53	0.70	0.51	0.73	0.77	0.70	0.57	0.67
9	خاک-تاکستان	6203	0.96	0.99	0.69	0.87	0.98	0.98	0.92	0.99	0.92	0.99
10	ذرت- سبزه پیر-علف هرز	3278	0.89	0.87	0.84	0.43	0.72	0.93	0.84	0.83	0.77	0.81
11	کاهو-۴ هفته	1068	0.98	0.75	0.85	0.94	0.75	0.87	0.93	0.73	0.98	0.71
12	کاهو-۵ هفته	1927	0.98	0.95	0.50	0.39	0.97	0.95	0.90	0.95	1.00	0.89
13	کاهو-۶ هفته	916	0.96	0.97	0.76	0.59	0.97	0.80	0.93	0.96	0.97	0.78
14	کاهو-۷ هفته	1070	0.94	0.94	0.73	0.73	0.94	0.89	0.94	0.93	0.94	0.80
15	تاکستان	7268	0.59	0.63	0.55	0.50	0.70	0.47	0.50	0.61	0.58	0.48
16	تاکستان-شاغولی-حاریست	1807	0.86	0.97	0.90	0.99	0.82	0.92	0.83	0.93	0.71	0.98
Average Acc. and Average Rel.			0.92	0.91	0.76	0.76	0.88	0.88	0.89	0.88	0.89	0.87
Kappa coefficient			0.85		0.69		0.79		0.82		0.79	

(جدول-۴): نتایج آزمون MCNemars
(Table-4): MCNemars test results

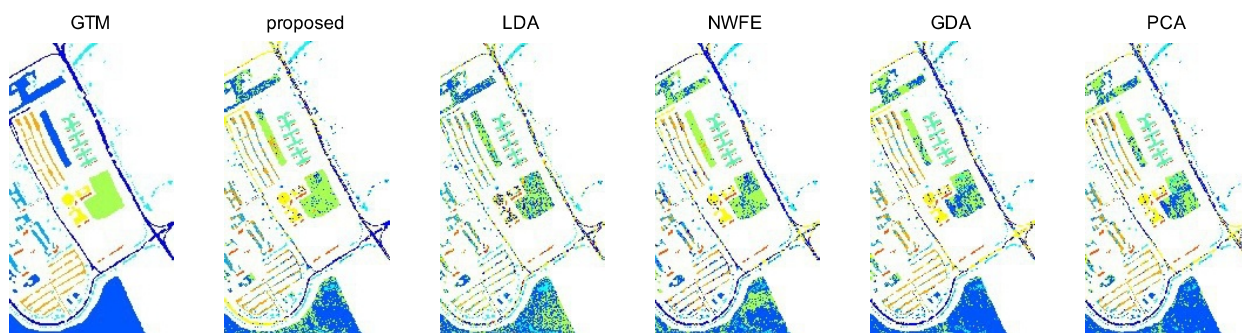
نتایج داده‌ی دانشگاه پاویا					
	Proposed	LDA	NWFE	GDA	PCA
Proposed	0	54.59	11.88	4.66	6.70
LDA	-54.59	0	-41.46	-48.38	-49.56
NWFE	-11.88	41.46	0	-7.90	-5.54
GDA	-4.66	48.38	7.90	0	1.86
PCA	-6.70	49.56	5.54	-1.86	0

نتایج داده‌ی ایندیانا					
	Proposed	LDA	NWFE	GDA	PCA
Proposed	0	85.46	40.23	62.26	35.19
LDA	85.46-	0	02.28-	62.26-	24.31-
NWFE	40.23-	02.28	0	86.2	85.3-
GDA	62.26-	62.26	86.2-	0	31.7-
PCA	35.19-	24.31	85.3	31.7	0

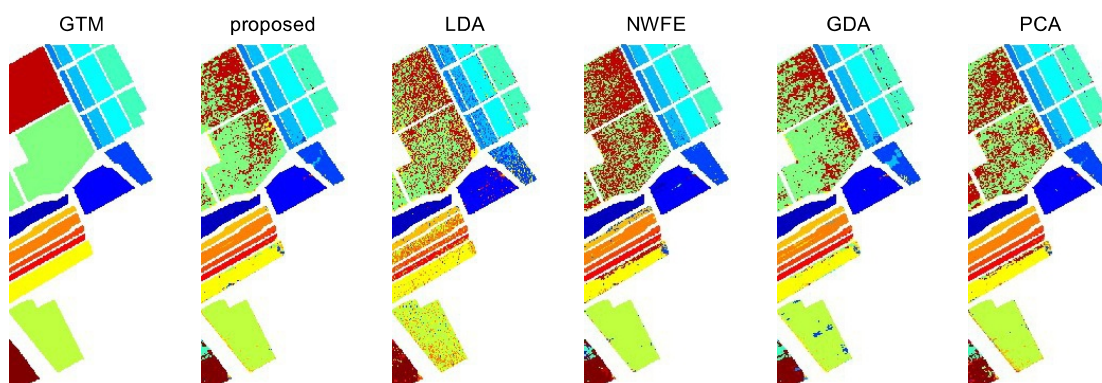
نتایج داده سالیناس						
Proposed	LDA	NWFE			GDA	PCA
Proposed	0	64.24	29.53	17.83	28.66	
LDA	-64.24	0	-40.95	-54.92	-41.39	
NWFE	-29.53	40.95	0	-12.52	1.78	
GDA	-17.83	54.92	12.52	0	14.64	
PCA	-28.66	41.39	-1.78	-14.64	0	



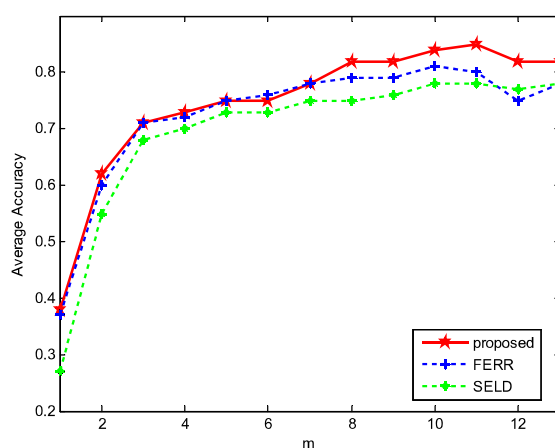
(شکل-۴): نقشه‌های طبقه‌بندی برای داده‌ی ایندیانا
(Figure-4): Classification maps for Indian dataset



(شکل-۵): نقشه‌های طبقه‌بندی برای داده‌ی دانشگاه پاویا
(Figure-5): Classification maps for University of Pavia dataset.



(شکل-۶): نقشه‌های طبقه‌بندی برای داده‌ی سالیناس
(Figure-6): Classification maps for Salinas dataset.



(شکل-۷): مقایسه روش پیشنهادی با روش‌های FERR و SELD در تعداد ویژگی‌های استخراج شده مختلف
(Figure-7): Comparison of the proposed method with FERR and SELD in different number of extracted features

برچسب خوشه‌بندی‌شده، باعث افزایش دقت طبقه‌بندی می‌شود. نتایج آزمایش‌ها بر روی سه داده ابرطیفی ایندیانا، پاویا و سالیناس، برتری روش پیشنهادی را نسبت به سایر روش‌های استخراج ویژگی از نظر معیار دقت، اعتبار، ضریب کاپا و آزمون MCNemars نشان می‌دهند.

۴- جمع‌بندی

در این مقاله، یک روش استخراج ویژگی برای طبقه‌بندی داده‌های ابرطیفی پیشنهاد شد که علاوه بر افزایش قدرت تفکیک بین طبقه‌ها، ساختار داده را در فضای با بعد کوچک‌تر نیز حفظ می‌کند و با استفاده از اطلاعات موجود در داده بدون

- [11] H. Peng, F. Long, C. and Ding, "Feature Selection Based on Mutual Information: Criteria of Max Dependency, Max-Relevance, and Min-Redundancy", IEEE Transactions on Pattern Analysis and Machine Intelligence, vol. 27, no. 8, pp. 1226-1238, 2005.
- [12] J.M. Sotoca, and F. Pla, "Band selection using mutual information matrix for hyperspectral data", marmota.dlsi.uji.es/WebBIB/papers, 2006.
- [13] J. Yin, C. Gao, and X. Jia, "Using Hurst and Lyapunov Exponent For Hyperspectral Image Feature Extraction", IEEE Geoscience and Remote Sensing Letters, vol. 9, no. 4, pp. 705 – 709, 2012.
- [14] L. Bruzzone, and C. Persello, "A novel approach to the selection of spatially invariant features for the classification of hyperspectral images with improved generalization capability", IEEE Transactions on Geoscience and Remote Sensing, vol. 47, no. 9, pp. 3180–3191, 2009.
- [15] K. Fukunaga, Introduction to Statistical Pattern Recognition, Academic Press Inc, San Diego, 1990.
- [16] L. Zhang, L. Zhang, D. Tao, and X. Huang, "Tensor Discriminative Locality Alignment for Hyperspectral Image Spectral–Spatial Feature Extraction", IEEE Transactions on Geoscience and Remote Sensing, vol. 51, no. 1, pp. 242–256, 2013.
- [17] J. Yin, Y. Wang, and J. Hu, "A new dimensionality reduction algorithm for hyperspectral image data using evolutionary strategy", IEEE Transactions on Industrial Informatics, vol. 8, no. 4, pp. 935–943, 2012.
- [18] D. Lunga and O. Ersoy, "Spherical Nearest Neighbor Classification: Application to hyperspectral Data", ECE Technical Reports, Purdue University, 2010.
- [19] D. Lunga and O. Ersoy, "Spherical Stochastic Neighbor Embedding of Hyperspectral Data", IEEE Transactions on Geoscience and Remote Sensing, vol. 51, no. 2, pp. 857–871, 2013.
- [20] M. Zortea, V. Haertel, and R. Clarke, "Feature Extraction in Remote Sensing High-Dimensional Image Data", IEEE Geoscience and Remote Sensing Letters, vol. 4, no. 1, pp. 107 – 111, 2007.
- [21] M. Imani and H. Ghassemian, "Assessment of Performance Improvement in Hyperspectral Image Classification Based on Adaptive Expansion of Training Samples", Journal of Information Systems and Telecommunication, vol. 2, no. 2, pp. 63-70, 2014.

5-References

۵-مراجع

- [1] G. F. Hughes, "On the mean accuracy of statistical pattern recognition," IEEE Transactions on Information Theory, vol. IT-14, no. 1, pp. 55–63, 1968.
- [2] Q. Zhang, Y. Tian, Y. Yang, and C. Pan, "Automatic Spatial–Spectral Feature Selection for Hyperspectral Image via Discriminative Sparse Multimodal Learning", IEEE Transactions on Geoscience and Remote Sensing, vol. 53, no. 1, pp. 261-279, 2015.
- [3] J. Xia, J. Chanussot, P. Du, X. He, "Spectral–Spatial Classification for Hyperspectral Data Using Rotation Forests with Local Feature Extraction and Markov Random Fields", IEEE Transactions on Geoscience and Remote Sensing, vol. 53, no. 5, pp. 2532–2546, 2015.
- [4] M. Imani and H. Ghassemian, "Feature Extraction Using Weighted Training Samples", IEEE Geoscience and Remote Sensing Letters, vol. 12, no. 7, pp. 1387-1391, 2015.
- [5] X. Kang, S. Li, L. Fang, J. A. Benediktsson, "Intrinsic Image Decomposition for Feature Extraction of Hyperspectral Images", IEEE Transactions on Geoscience and Remote Sensing, vol. 53, no. 4, pp. 2241-2253, 2015.
- [6] L. Ladha, and T. Deepa, "Feature Selection Methods and Algorithms", International Journal on Computer Science and Engineering, vol. 3, no. 5, pp. 1787–1797, 2011.
- [7] C. Persello, and L. Bruzzone, "Advanced Techniques for the Classification of Very High Resolution and Hyperspectral Remote Sensing Images", PhD Dissertation, university of Trento, 2010.
- [8] D. Korycinski, M. Crawford, J.W. Barnes, J. Ghosh, "Adaptive feature selection for hyperspectral data analysis using a binary hierarchical classifier and tabu search", IEEE Symposium on Geoscience and Remote Sensing, vol. 1, pp. 297- 299, 2003.
- [9] S.B. Serpico and L. Bruzzone, "A New Search Algorithm for Feature Selection in Hyperspectral Remote Sensing Images", IEEE Transactions on Geoscience and Remote Sensing, vol. 39, no. 7, pp. 1360–1367, 2001.
- [10] S. Li, H. Wu, D. Wan, J. Zhu, "An effective feature selection method for hyperspectral image classification based on genetic algorithm and support vector machine", Knowledge-Based Systems, vol. 24, pp. 40-48, 2011.



مریم ایمانی کارشناسی و کارشناسی ارشد خود را در رشته مهندسی برق - گرایش مخابرات به ترتیب در سال های ۱۳۸۸ و ۱۳۹۰ از دانشگاه شاهد دریافت کرد. ایشان در سال ۱۳۹۱ و

۱۳۹۴ به ترتیب در دوره دکتری و پسادکتری در دانشگاه تربیت مدرس به ادامه تحصیل مشغول شد. زمینه های پژوهشی مورد علاقه ایشان، شناسایی آماری الگو، پردازش سیگنال و اطلاعات و مهندسی سنجش از دور است.

نشانی رایانامه ایشان عبارت است از :

maryam.imani@modares.ac.ir



محمد حسن قاسمیان یزدی

تحصیلات کارشناسی خود را در رشته مهندسی مخابرات در سال ۱۳۵۸ در دانشکده مخابرات تهران به پایان رساند. ایشان مدرک کارشناسی ارشد و

دکترای خود را در رشته مهندسی مخابرات در دانشگاه پردو آمریکا به ترتیب در سال ۱۳۶۳ و ۱۳۶۷ اخذ کرد. ایشان هم اکنون استاد دانشکده مهندسی برق و کامپیوتر هستند. آنالیز و پردازش تصاویر از منابع چندگانه، پردازش اطلاعات و شناسایی الگو، مهندسی سنجش از دور، پردازش تصاویر و سیگنال های مهندسی پزشکی از جمله علاقه های پژوهشی ایشان است.

نشانی رایانامه ایشان عبارت است از :

ghassemi@modares.ac.ir

[22] M. Imani and H. Ghassemian, "Band Clustering-Based Feature Extraction for Classification of Hyperspectral Images Using Limited Training Samples", IEEE Geoscience and Remote Sensing Letters, vol. 11, no. 8, pp. 1325 – 1329, 2014.

[23] Y.-L. Chang, J.-N. Liu, C.-C. Han, and Y.-N. Chen, "Hyperspectral Image Classification Using Nearest Feature Line Embedding Approach", IEEE Transactions on Geoscience and Remote Sensing, vol. 52, no. 1, pp. 278–287, 2014.

[24] 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, 2013.

[25] M. J. Mendenhall and E. Merényi, "Relevance-Based Feature Extraction for Hyperspectral Images", IEEE Transactions on Neural Network, vol. 19, no. 4, pp. 658–672, 2008.

[26] F. Tsai and J.-S. Lai, "Feature Extraction of Hyperspectral Image Cubes Using Three-Dimensional Gray-Level Cooccurrence", IEEE Transactions on Geoscience and Remote Sensing, vol. 51, no. 6, pp. 3504–3513, 2013.

[27] G. Baudat, and F. Anouar, "Generalized discriminant analysis using a kernel approach", Neural Comput., vol. 12, pp. 2385–2404, 2000.

[28] B. C. Kuo, and D. A. Landgrebe, "Nonparametric weighted feature extraction for classification", IEEE Transactions on Geoscience and Remote Sensing, vol. 42, no. 5, pp. 1096–1105, 2004.

[29] W. Liao, A. Pižurica, P. Scheunders, W. Philips, Y. Pi, "Semisupervised Local Discriminant Analysis for Feature Extraction in Hyperspectral Images," IEEE Transactions on Geoscience and Remote Sensing, vol. 51, no. 1, pp. 184–198, 2013.

[30] M. Imani and H. Ghassemian, "Ridge regression-based feature extraction for hyperspectral data", International Journal of Remote Sensing, vol. 36, no. 6, pp. 1728–1742, 2015.

[31] J. Cohen, "A coefficient of agreement from nominal scales", Educational and Psychological Measurement, vol. 20, no. 1, pp. 37–46, 1960.

[32] G. M. Foody, "Thematic map comparison: Evaluating the statistical significance of differences in classification accuracy", Photogrammetric Engineering and Remote Sensing, vol. 70, no. 5, pp. 627–633, 2004.