

آموزش منیفلد با استفاده از تشکیل گراف منیفلد مبتنی بر بازنمایی

تنک

رسول حاجی زاده^۱ علی آقاگل زاده^۲ مهدی ازوجی^۳

۱- دانشکده مهندسی برق و کامپیوتر - دانشگاه صنعتی نوشیروانی بابل - بابل

r.hajizadeh@stu.nit.ac.ir

۲- دانشکده مهندسی برق و کامپیوتر - دانشگاه صنعتی نوشیروانی بابل - بابل

aghagol@nit.ac.ir

۳- دانشکده مهندسی برق و کامپیوتر - دانشگاه صنعتی نوشیروانی بابل - بابل

m.ezoji@nit.ac.ir

چکیده: در این مقاله، یک روش آموزش منیفلد مبتنی بر بازنمایی تنک معرفی می‌شود. تشکیل گراف منیفلد در فضای با ابعاد بالا، مهمترین مرحله در روش‌های آموزش منیفلد، جهت استخراج داده‌ها در فضاها با ابعاد پایین است که عموماً به دو دسته محلی و سراسری تقسیم می‌گردند. گراف منیفلد پیشنهادی، به استخراج هم‌زمان ویژگی‌های محلی و سراسری می‌پردازد. پس از تشکیل گراف منیفلد مبتنی بر بازنمایی تنک، دو روش خطی و غیرخطی جهت استخراج داده‌های تعبیه شده در منیفلد، معرفی می‌شوند. روش پیشنهادی، با روش‌های متداول آموزش منیفلد، مانند LLE، LEM، LPP و PCA، مقایسه و ارزیابی شده است. ارزیابی‌های انجام شده بر روی دو پایگاه داده‌های حروف و ارقام دست‌نویس فارسی HODA و IFHCDB، بیان‌گر کارایی بهتر روش پیشنهادی، مبتنی بر معیار نرخ تشخیص درست بوده و نرخ تشخیص درست ۹۱/۸۹ و ۹۳/۸۹، به ترتیب برای پایگاه داده‌های HODA و IFHCDB به دست آمده است. در ادامه، جهت کاهش پیچیدگی محاسباتی روش پیشنهادی، شکل تغییر یافته آن نیز معرفی گردیده است، که نتایج آن بر روی پایگاه داده HODA، نشان‌دهنده کارایی مناسب آن بوده و پیچیدگی محاسباتی تا حدود ۶ برابر کاهش داده است.

کلمات کلیدی: آموزش منیفلد، بازنمایی تنک، کاهش ابعاد، بازشناسی حروف و ارقام دست‌نویس فارسی

تاریخ ارسال مقاله: ۱۳۹۶/۳/۲۸

تاریخ پذیرش مشروط مقاله: ۱۳۹۶/۵/۲۷

تاریخ پذیرش مقاله: ۱۳۹۶/۷/۵

نام نویسنده مسئول: دکتر علی آقاگل زاده

نشانی نویسنده مسئول: ایران - مازندران - بابل - دانشگاه صنعتی نوشیروانی بابل - دانشکده‌ی برق و کامپیوتر

۱- مقدمه

امروزه، پردازش تصویر و شناسایی الگو، در کاربردهای فراوانی، مانند بازشناسایی افراد مبتنی بر تصویر چهره، بازشناسایی حروف و ارقام دستنویس، ادغام تصاویر پزشکی، خوشه‌بندی، مورد توجه بوده و تاثیر به سزایی داشته‌اند [۱-۳]. بیشتر داده‌ها در کاربردهای پردازش تصویر و یا دیگر حوزه‌ها، در فضایی با ابعاد بالا اندازه‌گیری و نمایش داده می‌شوند. اما، بیشتر این داده‌ها، از ابعاد ذاتی پایین‌تری برخوردار بوده و می‌توان آنها را در فضایی با ابعاد بسیار پایین‌تر، به گونه‌ای که کمترین مقدار اطلاعات از دست برود، بازنمایی نمود.

آموزش منیفلد، ابزاری است که در بازنمایی داده‌های با ابعاد بالا، در فضایی با ابعاد پایین در بسیاری از کاربردها مورد استفاده قرار می‌گیرد. در واقع، آموزش منیفلد، نگاشتی از فضای اندازه‌گیری با ابعاد بالا، به فضای بازنمایی با ابعاد پایین است که بیشترین اطلاعات ساختاری منیفلد داده‌ها را از فضای با ابعاد بالا به فضای با ابعاد پایین انتقال می‌دهد. از جمله مزایای به کارگیری آموزش منیفلد و یا دیگر روش‌های کاهش ابعاد، می‌توان به نیاز به حافظه کمتر جهت ذخیره‌سازی داده‌ها، کاهش پیچیدگی محاسباتی و افزایش سرعت پردازش اشاره نمود.

روش‌های آموزش منیفلد، در کاربردهای فراوانی مانند نمایش داده‌ها، دسته‌بندی داده‌ها و حذف نویز به کار گرفته شده‌اند. از آموزش منیفلد، در بازنمایی، بهینه‌سازی و تحلیل مساله‌های مهندسی مکانیک استفاده شده است [۴]. با تعریف منیفلد مبتنی بر شکل ظاهری اجسام، از آموزش منیفلد، به بازنمایی ساختار ظاهری اشیای پیچیده، در فضایی با ابعاد پایین پرداخته شده است [۴]. یکی دیگر از کاربردهای آموزش منیفلد، بازشناسی و طبقه‌بندی افراد، بر اساس تصویر چهره آنها است. روش‌های مبتنی بر آموزش منیفلد^۱ LPP [۷] و OLPP^۲، که روش‌های آموزش منیفلد خطی هستند، جهت بازشناسی و طبقه‌بندی افراد بر حسب تصویر چهره آنها معرفی گردیده‌اند [۵ و ۶]. همچنین، روش توسعه یافته‌ای مبتنی بر LPP و تعریف ماتریس فاصله، جهت تشخیص و بازشناسی افراد با استفاده از تصویر چهره، معرفی شده است [۸]. لوکاره^۳ و همکاران با استفاده از آموزش منیفلد، به تشخیص انحنای ردپاها و استخراج آن پرداخته و در شناسایی و تفکیک حیوانات گوناگون از آموزش منیفلد بهره برده است [۹]. یک روش آموزش نیمه نظارتی^۴، جهت دسته‌بندی تصاویر ابرطیفی^۵، با استفاده از روش‌های آموزش منیفلد محلی معرفی شده است [۱۰]. جهت تشکیل گراف ارتباطی میان داده‌ها، از دو روش آموزش منیفلد محلی^۶ LLE [۱۱] و LTSA^۷ [۱۲] استفاده می‌نماید، که گاهی مهم در روش‌های آموزش نیمه نظارتی است [۱۰]. عموماً، روش‌های آموزش منیفلد، با استخراج ویژگی‌های ساختاری منیفلد و تشکیل گرافی از منیفلد داده‌ها در فضای نمایش با ابعاد بالا، به دنبال بازنمایی داده‌ها در فضایی با ابعاد پایین (و یا ابعاد ذاتی داده‌ها) هستند.

گراف منیفلد در فضای با ابعاد بالا می‌تواند به صورت محلی و یا سراسری تشکیل گردد. در گراف منیفلد محلی، هر داده، تنها به داده‌های همسایگی خود مرتبط است و ضرایب میزان وابستگی میان هر داده و همسایه‌های آن، در روش‌های گوناگون، با استفاده از توابع مبتنی بر فاصله متفاوتی، محاسبه می‌شوند. در تشکیل گراف منیفلد محلی، هر داده با داده‌های خارج از همسایگی خود در ارتباط مستقیم نبوده و ضریب مابین آنها برابر با صفر است. اما با توجه به این‌که هر داده به همسایه‌های خود متصل است، ساختار منیفلد یکپارچه‌ای به دست می‌آید. در گراف منیفلد محلی، ویژگی‌های محلی ساختار منیفلد، متناسب با ضرایب و داده‌های متصل به یکدیگر استخراج می‌شوند. به روش‌های آموزش منیفلد مبتنی بر گراف منیفلد محلی، روش‌های آموزش منیفلد محلی گفته شده و ماتریس گراف محلی حاصل را ماتریس گراف همسایگی می‌نامند. روش‌های آموزش منیفلد محلی، بر اساس حفظ ویژگی‌های ساختاری محلی حاصل از داده‌های منیفلد در فضای با ابعاد بالا، به استخراج داده‌ها در فضای با ابعاد پایین می‌پردازند. روش‌های آموزش منیفلد محلی، تنها به ساختار محلی توجه داشته و در بازنمایی داده‌هایی که ساختار غیرخطی دارند، از عمل‌کرد بهتری برخوردارند. از جمله روش‌های محلی می‌توان به روش‌های LLE^۸، LEM^۹ [۱۳]، LTSA، LPP، SNE^۹ [۱۴] و HLL^{۱۰} [۱۵] اشاره نمود.

روش‌های آموزش منیفلد مبتنی بر گراف منیفلد سراسری را نیز روش‌های آموزش منیفلد سراسری می‌نامند. در تشکیل گراف منیفلد سراسری، بر خلاف گراف منیفلد محلی، هر داده با تمامی داده‌ها در ارتباط بوده و ضرایب میان هر داده با دیگر داده‌ها، متناسب با تابع به کار گرفته شده، تعیین می‌شوند. گراف منیفلد حاصل به استخراج ویژگی‌های ساختاری منیفلد پرداخته و روش‌های آموزش منیفلد سراسری، مبتنی بر حفظ این ساختار، به استخراج داده‌های تعبیه شده در فضای با ابعاد پایین می‌پردازند. روش‌های آموزش منیفلد سراسری، در استخراج منیفلد داده‌هایی که ساختاری خطی دارند، از عمل‌کرد بهتری برخوردار هستند. Isomap^{۱۱} [۱۶]، MVU^{۱۲} [۱۷]، PCA^{۱۳} [۱۹]، از جمله روش‌های آموزش منیفلد سراسری هستند. لازم به ذکر است که روش‌های آموزش منیفلد را می‌توان از نظر خطی و یا غیرخطی بودن، نظارتی و یا غیر نظارتی بودن نیز به گروه‌های گوناگونی دسته‌بندی نمود.

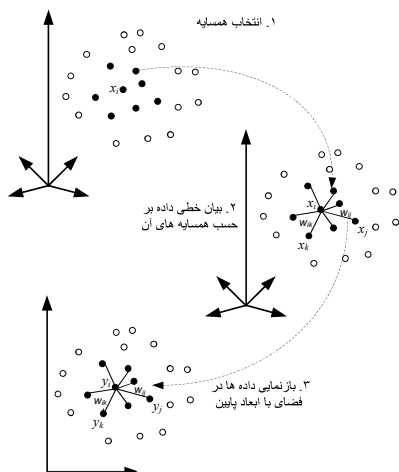
افزون بر اهمیت چگونگی تشکیل گراف منیفلد در استخراج ویژگی‌های ساختاری منیفلد، بیان هر چه تنک‌تر هر داده در گراف نیز، می‌تواند مفید باشد [۱۰]. بازنمایی تنک در کاربردهای فراوانی مورد توجه قرار گرفته است. به عنوان مثال، از مدل‌سازی تنک، جهت کاهش پیچیدگی محاسباتی و ایجاد فضایی برای جهت‌یابی گوینده‌ها استفاده می‌شود [۲۰]. روش‌های آموزش منیفلد محلی، از ساختار گراف منیفلد تنکی برخوردار هستند، اما توجهی به ویژگی‌های سراسری منیفلد ندارند. همچنین، برخی از داده‌های موجود در

از X می‌پردازد (منیفولد آموزش $X \in R^{D \times N}$ $Y \in R^{d \times N}$) به گونه‌ای که منیفولد حاصل از Y در فضای با ابعاد پایین، بیشترین تشابه را با منیفولد حاصل از X در فضای با ابعاد بالا، دارا است.

در ادامه، ابتدا روش‌های آموزش منیفولد محلی LLE، LEM، LPP و آموزش منیفولد سراسری PCA به صورت مختصر معرفی می‌گردند. سپس به معرفی الگوریتم OMP و مراحل آن می‌پردازیم.

۲-۱- آموزش منیفولد LLE

روش‌های آموزش منیفولد محلی، بر اساس حفظ ویژگی‌های ساختاری محلی، که با استفاده از رابطه میان هر داده و همسایه‌های آن به دست می‌آید، گرافی از داده‌ها ایجاد نموده و با استفاده از ماتریس گراف همسایگی حاصل، به استخراج داده‌های با ابعاد پایین می‌پردازند. LLE، یک روش آموزش منیفولد محلی پرکاربرد است که در سه گام به استخراج داده‌های تعبیه شده با ابعاد پایین می‌پردازد [۱۱]. شکل (۱) نشان‌دهنده مراحل الگوریتم LLE است.



شکل (۱): مراحل الگوریتم LLE [11].

در گام نخست، همسایه‌های هر داده تعیین خواهد شد. عموماً، از فاصله اقلیدسی، به عنوان پارامتری جهت تعیین همسایه‌ها و فاصله داده‌ها از یکدیگر استفاده می‌شود. دو روش متداول برای تعیین داده‌های همسایه، معیار k نزدیکترین همسایه^{۱۵} و توپ کوچک^{۱۶} هستند. در گام دوم، هر داده بر حسب همسایه‌های خود، به صورت خطی بیان گردیده و ماتریس گراف همسایگی (W_{LLE}) برحسب رابطه کمینه‌سازی (۱) به دست می‌آید:

$$W_{LLE} = \arg \min_W \left(\sum_{i=1}^N \left\| x_i - \sum_{j \in N_i} w_j x_j \right\|^2 \right), \quad s.t. \quad \sum_{j \in N_i} w_j = 1, \quad \forall i = 1, \dots, N. \quad (1)$$

که N_i شامل اندیس همسایه‌های داده i ام و w_j بیانگر ضریب میان همسایه j ام و داده i ام، در بازنمایی خطی داده i ام است. همچنین ضریب مابین داده i ام و داده‌هایی که در همسایگی آن قرار ندارند برابر با صفر خواهد بود. ضرایب بازنمایی به دست آمده برای هر داده، نسبت

همسایگی هر داده، از نظر اطلاعات ساختاری تکراری بوده و چندان در استخراج داده‌های تعبیه شده، مفید نخواهند بود. گراف منیفولد سراسری، در استخراج ویژگی‌های محلی کارایی ضعیف‌تری دارد و همچنین، از ساختاری تنک برخوردار نیست.

در این مقاله، روش آموزش منیفلدی مبتنی بر بازنمایی تنک و الگوریتم ردگیری تطبیقی متعامد^{۱۴} (OMP)، جهت بازشناسی حروف و ارقام دست‌نویس فارسی معرفی می‌شود. در روش پیشنهادی، ویژگی‌های ساختاری محلی و سراسری منیفلد، به طور همزمان مورد توجه قرار گرفته و گرافی تنک از منیفلد داده‌ها مبتنی بر الگوریتم ردگیری تطبیقی متعامد، تشکیل می‌گردد. سپس با بهره‌گیری از گراف منیفلد به دست آمده، طی دو تابع کمینه‌سازی پیشنهادی متفاوت، به صورت خطی و غیرخطی، داده‌های تعبیه شده، در فضای با ابعاد پایین محاسبه و بازنمایی می‌شوند. همچنین یکی از ضعف‌های روش‌های آموزش منیفلد غیرخطی، نگاشت داده‌های جدید یا آزمون به فضای با ابعاد پایین است، که عموماً در کاربردهای بازشناسی با آن روبه‌رو می‌شویم. در این مقاله، الگوریتمی جهت نگاشت داده‌های جدید یا آزمون از فضای با ابعاد بالا به فضای با ابعاد پایین، برای روش آموزش منیفلد پیشنهادی، معرفی خواهد شد. نتایج به دست آمده، بیانگر کارایی مناسب روش پیشنهادی است. بنابراین، نوآوری‌های موجود در مقاله، به صورت زیر قابل بیان است:

- استخراج توام ویژگی‌های محلی و سراسری گراف منیفلد،
- ترکیب مفاهیم تنک بودن و آموزش منیفلد،
- ارائه روشی جهت محاسبه داده‌های آزمون.

در ادامه، در بخش دو مبهم معرفی روش‌های آموزش منیفلد متداول و پرکاربرد LLE، LEM، LPP، PCA، که ساختاری شبیه به روش پیشنهادی در این مقاله دارند، و الگوریتم OMP می‌پردازیم. سپس در بخش سوم، روش آموزش منیفلد پیشنهادی و الگوریتم نحوه نگاشت داده جدید و یا داده آزمون، معرفی خواهند شد. نتایج شبیه‌سازی‌ها و مقایسه‌های صورت گرفته، در بخش چهارم مقاله ارائه می‌شوند. در پایان نیز، به کنکاش، نتیجه‌گیری و جمع‌بندی پرداخته می‌شود.

۲- روش‌های آموزش منیفلد LLE، LEM، LPP، PCA و الگوریتم OMP

با فرض x_i ، داده‌ای در فضای نمایشی با ابعاد بالا، دسته داده-های $X = \{x_1, x_2, \dots, x_N\} \in R^{D \times N}$ بیانگر N داده در فضای با ابعاد D است. روش‌های آموزش منیفلد به دنبال بازنمایی این دسته داده‌ها در فضایی با ابعاد پایین‌تر (d) است که بسیار کوچکتر از ابعاد داده‌ها در فضای نمایش اولیه است؛ یعنی $d \ll D$. با فرض y_i به عنوان داده‌ای در فضای با ابعاد پایین d ، دسته داده‌های متناظر در فضای با ابعاد پایین به صورت $Y = \{y_1, y_2, \dots, y_N\} \in R^{d \times N}$ قابل بیان است. به این ترتیب، آموزش منیفلد، فرآیندی است که به محاسبه Y با استفاده

به چرخش و تغییر مقیاس مستقل بوده و شرط $\sum_{j \in N_i} w_{ij} = 1$ نیز موجب استقلال ضرایب به دست آمده نسبت به جابه جایی می شوند. در واقع، ضرایب ماتریس همسایگی، حامل ویژگی های ساختار محلی منیفلد است.

در مرحله سوم، با استفاده از ماتریس ضرایب به دست آمده و رابطه (۲)، که مشابه رابطه محاسبه ضرایب بازسازی هر داده بر حسب همسایه های آن در فضای با ابعاد بالا است، داده های تعبیه شده محاسبه می شوند [۱۱].

$$Y_{LE} = \arg \min_Y \left(\sum_{i=1}^N \left\| y_i - \sum_{j \in N_i} w_{ij} y_j \right\|^2 \right), \quad s.t. \quad \sum_i y_i = 0, \quad (2)$$

$$\frac{1}{N} (Y^T Y) = I$$

که y_i ، بازنمایی داده λ_m ، در فضای با ابعاد پایین بوده و I ماتریس همانی است. رابطه (۲)، به صورت یک مساله مقدار ویژه قابل حل بوده و دسته داده های $Y_{LE} \in R^{d \times N}$ برابر با d بردار ویژه متناظر با d کوچکترین مقادیر ویژه مثبت ماتریس $M = (I - W_{LE})^T (I - W_{LE})$ است [۱۱].

۲-۲- آموزش منیفلد LEM

الگوریتم LEM، از دیگر روش های محلی مطرح در آموزش منیفلد است. این الگوریتم نیز مانند بیشتر روش های محلی شامل سه مرحله است. گام نخست، تعیین همسایه های هر داده که دقیقاً مانند الگوریتم LLE است. گام دوم، تشکیل گراف ماتریس همسایگی (W_{LEM}) با استفاده از داده ها و همسایه های آنها، که در LEM، وزن w_{ij} میان داده x_i و داده ی همسایه آن یعنی x_j ، با استفاده از رابطه زیر به دست می آید [۱۳]:

$$w_{ij} = \exp \left(- \frac{\|x_i - x_j\|^2}{2t^2} \right) \quad (3)$$

که در آن t ، پارامتر رابطه کرنل لاپلاسی است. همچنین، ضریب میان هر داده و داده هایی که در همسایگی آن قرار ندارند، مانند الگوریتم LLE، برابر با صفر است. در ویرایش هایی از گراف منیفلد LEM، وزن میان هر داده و همسایه های آن، برابر با یک قرار داده می شود. در گام سوم، با استفاده از رابطه (۴)، داده های تعبیه شده در فضای با ابعاد پایین، به دست می آید:

$$Y_{LEM} = \arg \min_Y \sum_i \sum_j \|y_i - y_j\|^2 w_{ij}, \quad s.t. \quad Y^T D Y = I \quad (4)$$

که D یک ماتریس قطری با $D_{ii} = \sum_j w_{ij}$ است. در رابطه بالا به ازای w_{ij} های بزرگتر، که نشان دهنده نزدیک بودن داده λ_m و همسایه آن در فضای با ابعاد بالا است، y_i و همسایه متناظر آن نیز باید به یکدیگر نزدیکتر بوده تا شرایط کمینه سازی تابع فراهم گردد. برعکس، هنگامی که ضرایب w_{ij} کوچک باشند نشان دهنده دورتر بودن داده ها

از یکدیگر بوده که در تابع کمینه سازی تعریف شده نیز می توانند از یکدیگر دورتر قرار گیرند.

مقادیر Y_{LEM} با استفاده از رابطه (۴)، مانند روش LLE، می تواند به صورت یک مساله مقدار ویژه محاسبه گردد. در LEM دسته داده های $Y_{LEM} \in R^{d \times N}$ برابر با d بردار ویژه متناظر با d کوچکترین مقادیر ویژه مثبت ماتریس $D^{-1}L$ است، که در آن $L = D - W_{LEM}$ است [۱۳].

۳-۲- آموزش منیفلد LPP

آموزش منیفلد LPP یک روش محلی است، که مانند دیگر روش های محلی شامل سه گام اصلی یافتن همسایه، تشکیل گراف و استخراج داده های تعبیه شده است. تعیین همسایه و چگونگی تشکیل گراف منیفلد LPP، کاملاً با روش آموزش منیفلد LEM یکسان بوده و تنها در گام سوم، یعنی استخراج داده های تعبیه شده، با روش LEM متفاوت است. در واقع آموزش منیفلد LPP، یک روش آموزش منیفلد خطی است که در آن، ماتریس نگاشت داده ها از فضای با ابعاد بالا به فضای با ابعاد پایین، از رابطه (۵) به دست می آید [۷]:

$$U_{LPP} = \arg \min_U U^T X L X^T U, \quad s.t. \quad U^T X D X^T U = I, \quad (5)$$

که U_{LPP} ماتریس نگاشت، $L = D - W$ ، گراف منیفلد محلی مشابه روش LEM و D ماتریس قطری با ضرایب روی قطر اصلی برابر با $\sum_j w_{ij}$ است. در این روش نیز، ماتریس نگاشت به صورت یک مساله مقدار ویژه قابل محاسبه است [۷]. پس از محاسبه ماتریس نگاشت، بازنمایی داده ها در فضای با ابعاد پایین، به صورت $Y = U_{LPP}^T X$ خواهد بود. از جمله ویژگی های این روش، نگاشت داده های آزمون به فضای با ابعاد پایین، با استفاده از ماتریس نگاشت حاصل از داده های آموزش است.

۴-۲- آموزش منیفلد PCA

آنالیز اجزای اصلی (PCA)، یکی از متداولترین روش های سراسری و خطی آموزش منیفلد و کاهش ابعاد است. ایده اصلی PCA، پیدا کردن زیرفضای خطی در فضای با ابعاد پایین است، که بیشترین تناسبات را با میزان پراکندگی داده ها در فضای با ابعاد بالا دارا است [۱۹]. با تعریف ماتریس کوواریانس داده ها در فضای با ابعاد بالا، به صورت $\bar{X} = \frac{1}{N} \sum_{i=1}^N x_i$ که $Cov(X) = \sum_{i=1}^N (x_i - \bar{X})(x_i - \bar{X})^T$ صورت دسته داده های X است و با توجه به نامنفی و متقارن بودن ماتریس کوواریانس داریم:

$$Cov(X) = U D U^T \quad (6)$$

که $U \in R^{d \times d}$ ، یک ماتریس همانی متعامد و یک ($U^T U = I$) شامل بردارهای ویژه و D ماتریس قطری شامل مقادیر ویژه است. با فرض $U = [u_1, u_2, \dots, u_D]$ بردارهای ویژه ی متناظر با مقادیر ویژه $\lambda_1, \lambda_2, \dots, \lambda_D$ ، اثبات می شود که $0 \leq \lambda_D \leq \lambda_{D-1} \leq \dots \leq \lambda_1$ بیانگر میزان

دیگری نیز، جهت اتمام حلقه تکرار این الگوریتم، معرفی شده است [۲۱]. در زیر به بیان مراحل الگوریتم OMP متداول، می‌پردازیم:

الگوریتم (۱): الگوریتم ردگیری تطبیقی متعامد (OMP)

ورودی‌ها: دیکشنری X و داده مورد بررسی y .
خروجی‌ها: α بردار ضرایب بازنمایی تنک و D ماتریس اتم‌های انتخاب شده از دیکشنری.

۱- مقداردهی اولیه: $r=y$, $\alpha=0$, $D=\emptyset$, $\Lambda=\emptyset$ مجموعه اندیس اتم‌های انتخاب شده از دیکشنری، که $\emptyset$ نماد تهی بودن است.

۲- محاسبه نرم r ، یعنی $\|r\|_2$ ، چنانچه $\|r\|_2 > \varepsilon$ و یا تعداد اتم انتخابی کوچکتر مساوی k است برو به گام ۳، در غیراینصورت اتمام الگوریتم و رفتن به انتهای آن. (ε یک مقدار مثبت بسیار کوچک است)

۳- $j^* = \arg \max_{j \in \Lambda} \langle r, x_j \rangle$ که $\langle \cdot, \cdot \rangle$ به ترتیب نماد ضرب داخلی و قدرمطلق است.

۴- به روز نمودن Λ و D : $\Lambda = \Lambda \cup \{j^*\}$ و $D = [D, x_{j^*}]$

۵- به روز نمودن بردار ضرایب تنک: $\alpha^* = \arg \min_{\alpha} \|y - D\alpha\|_2^2$

۶- محاسبه بردار باقی‌مانده $r = y - D\alpha^*$ و رفتن به مرحله ۲.

با اجرای الگوریتم ردگیری تطبیقی متعامد، ضرایب تنک و اتم‌های انتخاب شده از دیکشنری به دست می‌آیند که وابسته به مقدار ε و یا حداکثر تعداد اتم‌های انتخابی مجاز (k) است.

روش‌های آموزش منیفلد محلی، مانند LLE، LEM و LPP، تنها به ویژگی‌های ساختاری محلی منیفلد توجه داشته و در استخراج ویژگی‌ها، تنها به داده‌های موجود در همسایگی هر داده توجه می‌نماید و از داده‌هایی که در همسایگی آنها قرار ندارند، هیچ استفاده‌ای نمی‌کنند. در روش‌های آموزش منیفلد سراسری نیز، هر داده با تمامی داده‌های دیگر در ارتباط است و توجهی به ویژگی‌های ساختاری محلی ندارد. همچنین، بیان هر چه تنک‌تر گراف منیفلد، می‌تواند در استخراج ویژگی‌های منیفلد بهتر عمل نماید [۱۰]. از این‌رو، نیاز به روش آموزش منیفلدی، که بتواند به طور همزمان به استخراج ویژگی‌های محلی و سراسری پرداخته، و همچنین از ساختاری تنک برخوردار باشد، احساس می‌گردد.

۳- روش پیشنهادی: آموزش منیفلد مبتنی بر بازنمایی تنک^۹ (SRML)

همان‌طور که اشاره شد، توانایی اصلی بیشتر روش‌های آموزش منیفلد، در چگونگی تشکیل گراف منیفلد داده‌ها در فضای با ابعاد بالا است. در روش پیشنهادی، گراف منیفلدی در فضای با ابعاد بالا، مبتنی بر مفهوم بازنمایی تنک تشکیل گردیده و پس از آن به استخراج داده‌های تعبیه

پراکندگی داده‌های نگاشت یافته توسط ماتریس U ، یعنی $Y = U^T X$ است؛ یا به عبارت دیگر $Cov(Y) = D$. بنابراین، به راحتی می‌توان نشان داد که داده‌های نگاشت یافته به فضای با ابعاد پایین، توسط ماتریس نگاشت $W_{PCA} = [u_1, u_2, \dots, u_d]$ ، بیشترین پراکندگی داده‌ها را حفظ خواهد نمود. در نتیجه، محاسبه داده‌ها در فضای با ابعاد پایین به صورت زیر است:

$$Y = W_{PCA}^T X \quad (7)$$

۲-۵ الگوریتم ردگیری تطبیقی متعامد (OMP)

بیشترین کاربرد الگوریتم OMP در حل مساله بهینه‌سازی تنک زیر است [21]:

$$\alpha^* = \arg \min_{\alpha} \|\alpha\|_0, \quad s.t. \quad \|y - X\alpha\|_2^2 \leq \varepsilon \quad (8)$$

و یا،

$$\alpha^* = \arg \min_{\alpha} \|y - X\alpha\|_2^2, \quad s.t. \quad \|\alpha\|_0 \leq k \quad (9)$$

که X دیکشنری پایه و یا ماتریس شامل کلیه داده‌های اندازه‌گیری شده است، α ضرایب اتم‌ها و یا داده‌های اندازه‌گیری شده‌ی موجود در دیکشنری بوده، و y داده‌ای است که به دنبال بازنمایی تنک آن هستیم و در دیکشنری وجود ندارد. $\|\cdot\|_p$ نیز بیانگر نرم p است.

الگوریتم OMP، عموماً در حل مسائل بهینه‌سازی بازنمایی تنک با محدودیت نرم صفر به کار گرفته می‌شود. الگوریتم OMP، یک الگوریتم اصلاح شده ردگیری تطبیقی^{۱۷} (MP) است که سرعت همگرایی آن نسبت به الگوریتم متداول MP بهتر است. الگوریتم OMP، به بازنمایی خطی هر داده برحسب اتم‌های انتخاب شده از یک دیکشنری بیش از کامل^{۱۸} پرداخته و از ساختاری تکرار شونده برخوردار است. در OMP، ابتدا اتمی از دیکشنری (X)، که بیشترین شباهت را با داده مورد بررسی (y) دارد، انتخاب شده و اتم انتخابی به مجموعه اتم‌های انتخاب شده (D) از دیکشنری، که از ابتدا به صورت یک مجموعه تهی ($D = \emptyset$) بارگذاری شده است، اضافه می‌شود. سپس y بر حسب مجموعه اتم یا اتم‌های انتخاب شده D و با استفاده از رابطه کمینه‌سازی (۱۰) بازنمایی گردیده و ضریب یا ضرایب بردار ضرایب بازنمایی تنک (α) محاسبه می‌شود:

$$\alpha^* = \arg \min_{\alpha} \|y - D\alpha\|_2^2 \quad (10)$$

پس از محاسبه ضرایب α ، اختلاف میان داده y و بازنمایی آن با استفاده از ضرایب α و مجموعه اتم‌های انتخاب شده D ، به صورت $r = y - D\alpha$ محاسبه می‌گردد. در ادامه، r به عنوان داده مورد بررسی جایگزین گردیده و فرآیند فوق تکرار خواهد شد، با این تفاوت که اتم‌های انتخابی از دیکشنری X نمی‌تواند تکراری باشد. همچنین، معیار توقف این الگوریتم، بر اساس کمتر شدن نرم r از یک مقدار مشخص (عموماً مقدار مثبت کوچک) و یا انتخاب حداکثر تعداد اتم‌های انتخابی مجاز (k) از دیکشنری X است. معیارهای گوناگون

شده، خواهیم پرداخت. روش پیشنهادی، شامل سه گام اساسی زیر است: ۱- تشکیل گراف منیفلد داده‌ها در فضای با ابعاد بالا، مبتنی بر مفهوم بازنمایی تنک، ۲- استخراج داده‌های تعبیه شده و بازنمایی آن در فضای با ابعاد پایین، ۳- بازنمایی داده آزمون. در ادامه به شرح هر یک از این گام‌ها پرداخته می‌شود:

۳-۱- تشکیل گراف منیفلد داده‌ها در فضای با ابعاد بالا با استفاده از مفهوم بازنمایی تنک

در روش پیشنهادی، جهت تشکیل گراف منیفلد، به بازنمایی k -تنک هر داده $(x_i, \forall i = 1, \dots, N)$ بر حسب اتم‌های دیکشنری حاصل از بقیه داده‌ها $(\hat{X}_i = X - \{x_i\})$ و با استفاده از رابطه (۱۱) می‌پردازیم:

$$w_i^* = \arg \min_{w_i} \|x_i - \hat{X}_i w_i\|_2, \text{ s.t. } \|w_i\|_0 = k \quad (11)$$

که در آن، x_i داده نام X در فضای با ابعاد بالا است، مجموعه داده‌های در فضای با ابعاد بالا به جز داده نام است و w_i بردار ضرایب تنک است که تنها k مقدار غیر صفر دارد و یا به عبارت دیگر هر داده را تنها به k داده‌ی دیگر در گراف منیفلد، اتصال خواهد داد. الگوریتم پیشنهادی جهت محاسبه بردار ضرایب w_i ، دارای ساختاری مشابه الگوریتم ردگیری تطبیقی متعامد است. در واقع، x_i همان داده مورد بررسی و $\hat{X}_i$ شامل اتم‌های دیکشنری خواهد بود. در ادامه، مانند الگوریتم ردگیری تطبیقی متعامد، بردار ضرایب w_i محاسبه خواهد شد، با این تفاوت که معیار توقف الگوریتم، بر اساس انتخاب k اتم از دیکشنری $\hat{X}_i$ است. این اتم‌های انتخاب شده، داده‌هایی هستند که در گراف منیفلد با داده نام مرتبط بوده و بردار ضرایب w_i وزن ارتباطی مابین آنها را بیان می‌کند. در ادامه، با محاسبه ضرایب تنک برای تمام داده‌های موجود در فضای با ابعاد بالا، ماتریس گراف منیفلد مبتنی بر بازنمایی تنک، W_{SR} ، با استفاده از رابطه (۱۲) به دست می‌آید:

$$W_{SR} = \sum_{i=1}^N \arg \min_{w_i} \|x_i - \hat{X}_i w_i\|_2 = \arg \min_W \sum_{i=1}^N \|x_i - \hat{X}_i w_i\|_2^2 \quad (12)$$

$$\text{s.t. } \|w_i\|_0 = k, \forall i = 1, \dots, N.$$

از نکات بارز روش پیشنهادی، متفاوت بودن دیکشنری به کار گرفته شده، جهت محاسبه ضرایب گراف منیفلد به ازای هر داده است، که با کنار گذاشتن داده مورد بررسی از مجموعه داده‌های در فضای با ابعاد بالا، به دست می‌آید. از دیگر ویژگی‌های روش پیشنهادی، تنک بودن گراف منیفلد پیشنهادی، مانند روش‌های آموزش منیفلد محلیاست؛ در حالی که در تشکیل گراف، تنها به داده‌های در همسایگی هر داده محدود نبوده و بدین ترتیب گراف پیشنهادی می‌تواند در استخراج ویژگی‌های سراسری منیفلد نیز موثر باشد. الگوریتم تشکیل گراف منیفلد پیشنهادی، به طور مختصر در زیر بیان شده است:

الگوریتم (۲): الگوریتم تشکیل گراف منیفلد مبتنی بر بازنمایی تنک

ورودی: منیفلد داده‌ها در فضای با ابعاد بالا (X) ، تعیین مقادیر k خروجی: گراف منیفلد مبتنی بر بازنمایی تنک (W_{SR}) .

تکرار مراحل زیر برای هر داده موجود در منیفلد $(i = 1, \dots, N)$. تعداد کل داده‌های موجود در منیفلد است.

۱- تشکیل دیکشنری $\hat{X}_i = X - \{x_i\}$ ، مقداردهی اولیه: $r = x_i$ ، $w_i = 0$ بردار ضرایب تنک، $D = \emptyset$ اتم‌های انتخاب شده از دیکشنری، $\Lambda = \emptyset$ مجموعه اندیس اتم‌های انتخاب شده از دیکشنری.

۲- محاسبه نرم صفر بردار ضرایب تنک $(\|w_i\|_0)$. چنانچه شرط $\|w_i\|_0 = k$ برقرار است، اتمام محاسبه بردار ضرایب w_i و در غیر این صورت رفتن به مرحله ۳.

۳- یافتن بهترین نمونه از دیکشنری $(\hat{X}_i)$ یا به عبارت دیگر $j^* = \arg \max_{j \in \Lambda} \langle r, \hat{x}_j \rangle$

۴- به روزرسانی Λ و D : $\Lambda = \Lambda \cup \{j^*\}$ و $D = [D, \hat{x}_{j^*}]$

۵- محاسبه و به روز نمودن بردار ضرایب تنک (w_i) : $w_i^* = \arg \min_{w_i} \|x_i - Dw_i\|_2^2$

۶- محاسبه بردار باقی‌مانده $r = x_i - Dw_i^*$ و رفتن به مرحله ۲.

۳-۲- استخراج و بازنمایی داده‌های تعبیه شده در فضای با ابعاد پایین

ماتریس گراف منیفلد تنک تشکیل شده (W_{SR}) ، بیانگر ویژگی‌های ساختاری منیفلد داده‌ها در فضای با ابعاد بالا است. در روش پیشنهادی، دو ساختار متفاوت جهت محاسبه داده‌ها در فضای با ابعاد پایین، مبتنی بر حفظ ساختار گراف W_{SR} معرفی می‌گردد.

۳-۲-۱- استخراج داده‌ها، مبتنی بر حفظ بازنمایی خطی هر داده

در الگوریتم تشکیل گراف منیفلد پیشنهادی، ضرایب تنک برای هر داده در فضای با ابعاد بالا، به صورت ترکیب خطی آن داده، بر حسب داده‌های متصل به آن، بیان می‌شود. بنابراین، در نمایش داده‌های تعبیه شده در فضای با ابعاد پایین نیز، تا حد ممکن به دنبال حفظ و برقراری این ترکیب خطی خواهیم بود. بر همین اساس و با تکیه بر حفظ ساختار گراف منیفلد به دست آمده، تابع کمینه‌سازی زیر جهت استخراج داده‌ها در فضای با ابعاد پایین پیشنهاد می‌شود:

$$Y_{SRML} = \arg \min_Y \left(\sum_{i=1}^N \|y_i - Y w_i\|_2^2 \right) \quad (13)$$

که در آن w_i بردار ضرایب تنک محاسبه شده برای داده نام، در فضای با ابعاد بالا است. رابطه فوق، از ساختاری شبیه به رابطه LLE برخوردار بوده و به صورت رابطه (۱۴) قابل بازنویسی است [11]:

$$Y_{SRML} = \arg \min_Y \left(\sum_{i=1}^N \|y_i - Y w_i\|_2^2 \right) = \arg \min_Y Y^T (I - W_{SR})^T (I - W_{SR}) Y \quad (14)$$

با استفاده از دو شرط $\sum_i y_i = 0$ و $\frac{1}{N}(Y^T Y) = I$ ، رابطه بالا به صورت یک مساله مقدار ویژه قابل حل بوده و ماتریس داده‌ها در فضای با ابعاد پایین (Y)، برابر با d بردار ویژه متناظر با d کوچکترین مقادیر ویژه مثبت ماتریس $M_{SR} = (I - W_{SR})^T (I - W_{SR})$ خواهد بود. این روش استخراج داده‌های تعبیه شده، غیرخطی است و در مواجهه با داده‌ی آزمون، که نقشی در استخراج داده‌های تعبیه شده‌ی آموزش نداشته‌اند، دچار مشکل می‌شود، که در بخش ۳-۳، الگوریتمی جهت محاسبه بازنمایی داده آزمون در فضای با ابعاد پایین، معرفی خواهد شد.

۳-۲-۲- استخراج داده‌ها با استفاده از ماتریس نگاشت

در این بخش به معرفی روشی خطی، جهت استخراج داده‌های تعبیه شده در فضای با ابعاد پایین و با استفاده از ماتریس نگاشت خواهیم پرداخت، به گونه‌ای که ساختار گراف تنک به‌دست آمده در فضای با ابعاد بالا، تا حد ممکن حفظ گردد. در روش پیشنهادی، پس از تشکیل ماتریس گراف منیفلد تنک W_{SR} ، ماتریس نگاشت U_{SRML} ، از رابطه‌ای مشابه رابطه (۵) حاصل می‌گردد که تنها در آن به جای ماتریس ضرایب محلی W_{LPP} ، از ماتریس گراف منیفلد تنک مبتنی بر الگوریتم ردگیری تطبیقی متعامد، یعنی W_{SR} استفاده شده، که در زیر نشان داده شده است:

$$U_{SRML} = \arg \min_U U^T X L_{SR} X^T U, \quad s.t. \quad U^T X D_{SR} X^T U = I \quad (15)$$

که در آن، $L_{SR} = D_{SR} - W_{SR}$ ، و D_{SR} ماتریس قطری با ضرایب روی قطر اصلی برابر با $\sum_j W_{SR_{ij}}$ است. ماتریس نگاشت U_{SRML} به صورت یک مساله مقدار ویژه قابل حل بوده و ماتریس $U_{SRML} \in R^{D \times d}$ ، متناسب با d بردار ویژه متناظر با d کوچکترین مقدار ویژه مثبت ماتریس $(X D_{SR} X^T)^{-1} X L_{SR} X^T$ است. بنابراین، مختصات داده‌ها در فضای با ابعاد پایین، به صورت $Y_{SRML} = U_{SRML}^T X$ به دست خواهد آمد.

۳-۳- بازنمایی داده‌های آزمون در فضای با ابعاد پایین

یکی از مشکلات اساسی بیشتر روش‌های آموزش منیفلد غیرخطی، محاسبه مختصات داده‌ای در فضای با ابعاد پایین است که در محاسبه بازنمایی داده‌های تعبیه شده‌ی آموزش، دخیل نبوده است. همچنین، در کاربردهای طبقه‌بندی، نگاشت داده‌های آزمون، از مشکلات اساسی روش‌های آموزش منیفلد غیرخطی است. یکی از راه‌های غیرمتعارف در برخورد با داده‌های جدید، تشکیل دوباره گراف ساختار منیفلد به همراه داده جدید و محاسبه مجدد داده‌های تعبیه شده آموزش و آزمون به طور همزمان در فضای با ابعاد پایین است، که منطقی نبوده و از نظر محاسباتی پیچیده بوده و زمان اجرای بالایی را در بر خواهد داشت. در [22]، روشی مبتنی بر ساختار مساله ویژه، برای روش‌های آموزش منیفلد LLE، LEM، Isomap، MDS و Spectral Clustering معرفی شده است. در [۲۲]، با ارائه الگوریتمی مبتنی بر مساله ویژه، به

محاسبه داده‌های آزمون در فضای با ابعاد پایین پرداخته و سپس به تطبیق روش‌های آموزش منیفلد، با الگوریتم پیشنهادی مبتنی بر مساله ویژه می‌پردازد.

در آموزش منیفلد پیشنهادی، پس از تشکیل گراف منیفلد تنک، دو روش متفاوت، جهت استخراج داده‌ها در فضای با ابعاد پایین، معرفی گردیده است. روش مبتنی بر ماتریس نگاشت، از ساختاری خطی، جهت محاسبه و بازنمایی داده‌ها در فضای با ابعاد پایین، بهره می‌برد. در این روش، با استفاده از ماتریس نگاشت به دست آمده از رابطه (۱۵)، که بر حسب داده‌های آموزش به دست آمده است، می‌توان به بازنمایی داده‌های آزمون، در فضای با ابعاد پایین به صورت $y_i = U_{SRML}^T x_i$ پرداخت که در آن x_i داده جدید یا آزمون در فضای با ابعاد بالا و y_i داده متناظر آن در فضای با ابعاد پایین است. در شکل ۲-الف، نحوه محاسبه داده آزمون، در روش مبتنی بر ماتریس نگاشت، نشان داده شده است. اما روش استخراج داده‌های مبتنی بر حفظ بازنمایی خطی هر داده، از ساختاری غیرخطی برخوردار است و به راحتی نمی‌توان به بازنمایی داده آزمون پرداخت. در ادامه به معرفی روشی جهت محاسبه‌ی داده‌های آزمون، در فضای با ابعاد پایین و بدون نیاز به تشکیل دوباره گراف می‌پردازیم. شکل ۲-ب چگونگی محاسبه مختصات داده آزمون در روش SRML مبتنی بر حفظ بازنمایی خطی هر داده را نشان می‌دهد. در روش پیشنهادی، ابتدا به بازنمایی k -تنک داده آزمون x_i با استفاده از رابطه (۱۶) و دیکشنری حاصل از تمام داده‌های آموزش (X) پرداخته و ضرایب بازنمایی تنک داده آزمون، w_i ، محاسبه می‌گردد (بلوک اول شکل ۲-ب):

$$w_i^* = \arg \min_{w_i} \|x_i - X w_i\|_2, \quad s.t. \quad \|w_i\|_0 = k \quad (16)$$

ضرایب w_i ، با استفاده از الگوریتم (۲) قابل محاسبه است. سپس، با استفاده از w_i و داده‌های آموزش بازنمایی شده در فضای با ابعاد پایین (Y)، y_i با استفاده از رابطه (۱۷) به دست می‌آید:

$$y_i^* = \arg \min_{y_i} \|y_i - Y w_i\|_2 \quad (17)$$

که یک تابع بهینه‌سازی درجه دوم بوده و دارای جواب فرم بسته‌ای است که با مشتق‌گیری از رابطه فوق بر حسب y_i ، و برابر با صفر قرار دادن مشتق به دست آمده، داریم:

$$\frac{\partial}{\partial y_i} (\|y_i - Y w_i\|_2^2) = \frac{\partial}{\partial y_i} (y_i - Y w_i)^T (y_i - Y w_i) = 2(y_i - Y w_i) = 0. \quad (18)$$

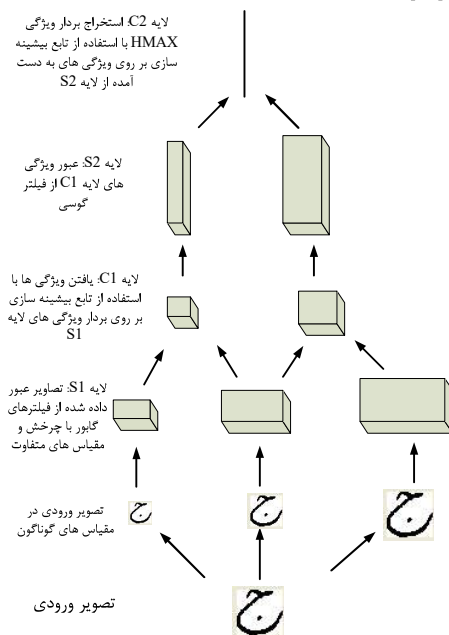
در نتیجه، داده متناظر با داده آزمون در فضای با ابعاد پایین، به صورت $y_i = Y w_i$ قابل محاسبه است که در آن، y_i به صورت ترکیب خطی از داده‌های آموزش بازنمایی شده در فضای با ابعاد پایین، قابل بیان است (بلوک دوم شکل ۲-ب). بنابراین، با استفاده از روش مطرح شده در بالا، می‌توان داده‌های آزمون و یا داده‌هایی را که در آموزش منیفلد آموزش نقشی نداشته‌اند، بدون نیاز به تشکیل گراف منیفلد، به فضای با ابعاد پایین نگاشت نمود.

است. همچنین مرجع [۳۴]، به معرفی روشی مبتنی بر ویژگی HMAX، جهت بازنمایی ارقام دست‌نویس فارسی پرداخته است. ویژگی HMAX، یک مدل محاسباتی سلسله مراتبی است که از ۴ لایه تشکیل می‌شود: لایه اول (S1)، لایه فیلتر گابور است. در این لایه، تصویر از فیلترهای گابور متعدد با پارامترهای چرخش و اندازه گوناگون عبور داده شده و خروجی فیلترها، بردار ویژگی‌های لایه S1 را تشکیل خواهد داد. لایه دوم (C1)، لایه تغییرناپذیر محلی^{۲۰} است. در این لایه، با استفاده از تابع بیشینه بر روی دسته‌های مختلف از بردارهای ویژگی به دست آمده در لایه S1، به استخراج ویژگی‌های تغییرناپذیر با چرخش، مکان و اندازه می‌پردازد. لایه سوم (S2)، لایه ویژگی میانی^{۲۱} است. مربعی از چهار ویژگی کنار هم و بدون هم‌پوشانی به دست آمده از لایه C1، به عنوان ورودی یک فیلتر گوسی در این لایه مورد استفاده قرار می‌گیرد. همچنین در این لایه تمام حالت‌های ممکن چینش چهار ویژگی کنار هم C1 مورد ارزیابی قرار می‌گیرند و یک دیکشنری از ویژگی‌های HMAX به دست می‌آید. لایه چهارم (C2)، لایه تغییرناپذیر سراسری است. این لایه مانند لایه دوم عمل نموده و با استفاده از تابع بیشینه و خروجی‌های لایه S2، به استخراج بردار ویژگی تغییرناپذیر با مکان و مقیاس خواهد پرداخت. عموماً از خروجی لایه چهارم، یعنی بردار ویژگی C2، به عنوان ویژگی HMAX در کاربردهای گوناگون استفاده گردیده است. شکل ۴، خلاصه‌ای از نحوه استخراج ویژگی HMAX را نشان می‌دهد. در شکل ۴، تصویر هر حرف دست‌نویس در مقیاس‌های متفاوت به عنوان ورودی به لایه S1 وارد شده و پس از عبور از فیلترها و توابع بیشینه‌ساز لایه‌های مختلف، خروجی لایه C2 به عنوان بردار ویژگی استخراج می‌شود. بردار ویژگی C2 به کارگرفته شده در این بررسی، یک بردار ۴۰۰ بعدی به ازای هر تصویر حروف دست‌نویس فارسی است، که تشکیل منیفلد داده‌ای ۴۰۰ بُعدی را خواهد داد. در مرجع [۳۳] توضیحات بیشتر در مورد HMAX بیان شده است. در این مقاله، حروف دست‌نویس فارسی، به صورت جدول (۱) کلاس‌بندی شده‌اند که شامل ۱۸ کلاس متفاوت است.

۴-۱- مقایسه نتایج نرخ تشخیص درست به دست آمده از روش‌های آموزش منیفلد و روش پیشنهادی SRML

جدول ۲، نتایج نرخ تشخیص درست به دست آمده از روش پیشنهادی SRML و روش‌های آموزش منیفلد متداول LEM، LLE، PCA و LPP، بر روی حروف دست‌نویس پایگاه داده‌های HODA را بیان می‌کند. تعداد داده‌های همسایه در روش‌های آموزش منیفلد محلی (LPP، LEM، LLE) و مقدار k در روش آموزش منیفلد پیشنهادی SRML، با یکدیگر برابر فرض می‌شوند. با این فرض، گراف حاصل از روش‌های محلی و روش SRML، دارای تعداد داده‌های

متصل به یکدیگر یکسانی بوده و در نتیجه در شرایط مشابهی می‌توان به مقایسه و ارزیابی آنها پرداخت. در جدول ۲، نتایج به ازای مقادیر مختلف k برابر با ۱۵، ۲۰ و ۳۰، نشان داده شده و به ازای هر k ، مقادیر نرخ تشخیص درست در فضای با ابعاد پایین، به ازای ابعاد گوناگون ارائه گردیده است. در نتایج به دست آمده، زیربخش خطی روش SRML، مربوط به استخراج داده‌ها مبتنی بر ماتریس نگاشت است و زیربخش غیرخطی مربوط به استخراج داده‌ها مبتنی بر حفظ بازنمایی خطی هر داده در گراف، است.



شکل (۴): نحوه محاسبه ویژگی HMAX [۳۳]

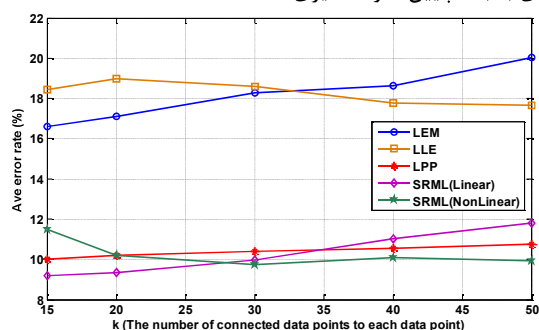
نتایج نشان داده شده در جدول ۲، بیانگر کارایی مناسب آموزش منیفلد SRML، در استخراج داده‌های تعبیه شده، نسبت به دیگر روش‌های آموزش منیفلد است و هر دو ساختار پیشنهادی، جهت استخراج داده‌های تعبیه شده، از کارایی مناسبی برخوردار هستند. همچنین، نتایج به دست آمده بر روی پایگاه داده‌های IFHCDB، در جدول ۳ نشان داده شده است که باز هم بیانگر عمل‌کرد مناسب الگوریتم SRML در استخراج داده‌های تعبیه شده، مبتنی بر میزان نرخ تشخیص درست است.

میزان نرخ تشخیص درست به دست آمده بر روی دو پایگاه داده‌های HODA و IFHCDB در فضای اصلی (با بعد ۴۰۰) به ترتیب برابر با ۹۰/۴۶ و ۸۹/۳۷ است که در مقایسه با نتایج نرخ تشخیص درست در فضاهای با ابعاد پایین نشان داده شده در جدول‌های ۲ و ۳، روش پیشنهادی SRML به نرخ تشخیص درست در فضای اصلی نزدیکتر است. همچنین، به ازای ابعاد بازنمایی بزرگتر در فضاهای با ابعاد پایین، روش پیشنهادی، از نرخ تشخیص درست بالاتری نسبت به نرخ تشخیص درست به دست آمده در فضای اصلی برخوردار است، که نشان‌دهنده قابلیت مناسب روش پیشنهادی، در

انتخاب و اتصال داده‌های مشابه و هم‌کلاس در گراف منیفولد است. البته باید اشاره داشت که روش LPP نیز، از نتایج قابل قبولی برخوردار است. بنابراین، با توجه به نزدیکی مقادیر نرخ تشخیص در فضای اصلی و فضای بازنمایی با ابعاد پایین، می‌توان نتیجه گرفت که روش پیشنهادی در استخراج و انتقال ساختار منیفولد، از فضای با ابعاد بالا به فضای با ابعاد پایین از عمل کرد مناسبی برخوردار است. از دلایل عمل کرد مناسب روش پیشنهادی، محدود نبودن انتخاب داده‌های متصل به هم در گراف به داده‌های در همسایگی هر داده و انتخاب داده‌های با بیشترین میزان تشابه است که احتمال انتخاب داده‌های هم‌کلاس را بالاتر می‌برد. همچنین روش SRML مبتنی بر ماتریس نگاشت، به دلیل برخورداری از ساختاری خطی و نگاشت داده‌های آزمون به فضای با ابعاد پایین با استفاده از ماتریس نگاشت، می‌تواند به راحتی در کاربردهای بازنمایی به کار گرفته شود.

۴-۲- بررسی رفتار الگوریتم پیشنهادی SRML بر حسب ابعاد در فضای با ابعاد پایین و k های مختلف

در این بخش، رفتار الگوریتم SRML و دیگر روش‌های تحت بررسی، نسبت به تغییر مقدار k و تغییر ابعاد فضای نگاشت داده‌ها (d) در فضای با ابعاد پایین، بر روی پایگاه داده‌های HODA، تحت بررسی قرار گرفته است. در جدول ۴، نتایج به دست آمده به ازای مقادیر ۱۵ و ۲۰، ۳۰ برای k و مقدار d از ۴۰ تا ۱۵۰ نشان داده شده است. همچنین جهت بررسی اثر k مقادیر نرخ تشخیص در ستروشه‌های آموزش منیفولد، به ازای k های برابر با ۴۰ و ۵۰ نیز در جدول ۵ نشان داده شده است. مقادیر دو جدول ۴ و ۵ نیز نشان‌دهنده کارایی روش پیشنهادی SRML است. در ادامه، جهت بررسی اثر k بر روی عمل کرد روش‌های مورد بررسی، متوسط نرخ خطای تشخیص، بر حسب k های مختلف بر روی پایگاه داده‌های HODA در شکل ۵ بیان گشته است که در آن به ازای هر k ، خطاهای نرخ تشخیص به ازای d های گوناگون، در فضای با ابعاد پایین متوسط‌گیری شده‌اند.

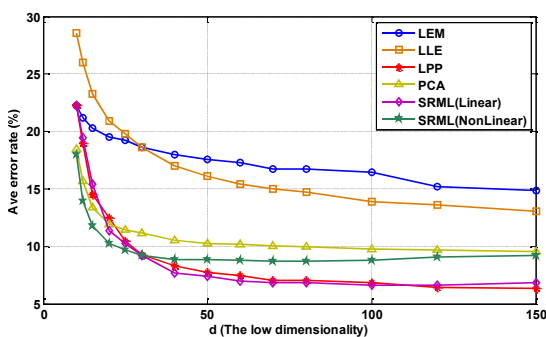


شکل (۵): متوسط نرخ خطای تشخیص روش‌های آموزش منیفولد بر حسب k از ۱۵ تا ۵۰ بر روی پایگاه داده‌های HODA.

شکل ۵ نشان می‌دهد که روش پیشنهادی SRML به همراه روش LPP از عمل کرد مناسبی برخوردار هستند. به خصوص در k های

کوچکتر و یا به عبارت دیگر در گراف‌های تنک‌تر، روش SRML پیشنهادی مبتنی بر ماتریس نگاشت، بهتر بوده و با افزایش k و یا به عبارت دیگر گراف‌های منیفولد پیچیده‌تر، SRML مبتنی بر حفظ بازنمایی خطی هر داده، از عمل کرد بهتری برخوردار است. یکی دیگر از نتایجی که می‌توان از شکل ۵ استخراج نمود، ثبات میزان نرخ خطای تشخیص روش SRML پیشنهادی به ازای k های متفاوت است و در عین حال کمترین میزان متوسط نرخ خطای تشخیص را دارد. همچنین، فاصله زیاد میان متوسط نرخ خطای تشخیص روش‌های LLE و LEM از روش پیشنهادی SRML نیز نشان‌دهنده عمل کرد مناسب روش پیشنهادی SRML است.

شکل ۶ نیز، نشان‌دهنده چگونگی رفتار روش‌های آموزش منیفولد، به ازای d های متفاوت در فضای با ابعاد پایین است که در آن، نرخ خطای تشخیص، به ازای k های مختلف، متوسط‌گیری شده است. نرخ خطای تشخیص تمام روش‌های آموزش منیفولد، با افزایش d کاهش یافته و نتایج نیز، بیانگر کارایی مناسب روش پیشنهادی SRML و روش LPP است. همچنین به ازای d های بیشتر از ۶۰، نرخ تشخیص درست روش‌های SRML، LPP و PCA بهبود چندانی نداشته و مقدار d برابر با ۵۰ را می‌توان به عنوان بُعد ذاتی داده‌های منیفولد، از نظر اطلاعات طبقه‌بندی نمودن کلاس‌های گوناگون، مد نظر قرار داد. از طرفی به ازای d های کوچکتر از ۳۰، روش SRML مبتنی بر حفظ بازنمایی خطی هر داده، از بهترین نرخ تشخیص درست برخوردار است که نشان‌دهنده کارایی بهتر این روش، در استخراج و انتقال اطلاعات تمایزی، از فضای اصلی به فضای با ابعاد پایین، خصوصاً در ابعاد بسیار پایین است. نتیجه دیگری که می‌توان از شکل ۶ درک نمود، رفتار کاهشی میزان متوسط نرخ خطا است که در روش‌های مختلف، کمی با یکدیگر متفاوت هستند. اما ذات کلی تغییرات متوسط نرخ خطای تشخیص بر حسب افزایش بُعد بازنمایی، به صورت کاهش است که قابل انتظار بوده است. این نوع رفتار می‌تواند بُعد ذاتی قابل استخراج توسط روش‌های مختلف را تعیین نماید که البته در روش‌های مختلف چندان متفاوت نیست. در اینجا نیز، بُعد ذاتی داده‌ها در بازه ۴۰ تا ۶۰ قرار گرفته است.



شکل (۶): متوسط نرخ خطای تشخیص روش‌های آموزش منیفولد بر حسب d از ۱۰ تا ۱۵۰ بر روی پایگاه داده‌های HODA.

بنابراین، با توجه به نمودارهای شکل‌های ۵ و ۶، d و k مناسب جهت به کارگیری روش SRML بر روی پایگاه داده‌های HODA، به ترتیب برابر با ۶۰ و ۲۰ خواهد بود و افزایش مقدار آنها تنها موجب افزایش پیچیدگی محاسباتی و پیچیده‌تر شدن گراف منیفلد می‌شود. همچنین، شکل ۷ نشان‌دهنده‌ی انحراف معیار تغییرات روش‌های تحت بررسی، بر حسب تغییرات k و d است. نمودارهای به‌دست آمده در شکل ۷-ب به همراه شکل ۵، نشان‌دهنده‌ی عمل‌کرد مناسب روش SRML نسبت به تغییرات مقدار k است. شکل ۵ نشان می‌دهد که روش SRML، از کمترین نرخ خطای تشخیص برخوردار بوده و روش LPP، دارای نتایجی نزدیک به آن است. در شکل ۷-ب، به ازای هر k ، انحراف معیار مقدار نرخ تشخیص به ازای d های گوناگون محاسبه گشته و مقدار کم آن، نشان می‌دهد که می‌توان به ازای ابعاد پایین‌تر نیز، به نرخ تشخیص قابل قبولی دست یافت. با توجه به شکل ۷-ب، میزان انحراف معیار مقادیر نرخ خطای تشخیص روش SRML، از مقدار پایینی برخوردار بوده که نشان‌دهنده‌ی حساسیت پایین روش پیشنهادی، به بُعد فضای بازنمایی در فضا با ابعاد پایین است.

همچنین، نمودارهای شکل ۷-الف در کنار شکل ۶ نشان می‌دهد که روش‌های LPP و SRML مبتنی بر حفظ بازنمایی خطی هر داده، از کمترین متوسط نرخ خطای تشخیص بر حسب d های گوناگون برخوردار هستند و در عین حال، دارای کمترین میزان تغییرات نرخ تشخیص به ازای تغییر مقدار k هستند. متوسط و انحراف معیار نرخ خطای تشخیص روش SRML پیشنهادی، به ازای d های بزرگتر از ۵۰ تقریباً ثابت است.

بنابراین، با توجه به شکل‌های ۵ تا ۷، با استفاده از روش SRML، می‌توان از کمترین مقادیر k و d ، جهت بازنمایی داده‌ها در فضا با ابعاد پایین، بهره برد که به ترتیب منجر به ساختار گراف منیفلدی تنک‌تر و کاهش حجم ذخیره‌سازی و پیچیدگی محاسباتی می‌گردند.

۴-۳- نتایج الگوریتم M-SRML پیشنهادی

در این بخش، به مقایسه نتایج حاصل از روش پیشنهادی SRML و روش اصلاح شده‌ی آن، یعنی M-SRML، می‌پردازیم. نتایج شبیه‌سازی بر روی پایگاه داده‌های HODA، به ازای k برابر با ۵ و k_i برابر با ۳، ۴ و ۶ و همچنین k برابر با ۱۰ و k_i برابر با ۲ و ۳، در جدول ۶ نشان داده شده است. نتایج به‌دست آمده، نشان‌دهنده‌ی عمل‌کرد تقریباً یکسان و حتی بهتر روش پیشنهادی M-SRML نسبت به روش پیشنهادی SRML در برخی از ابعاد بازنمایی است. برای نمونه، نرخ تشخیص درست در فضا با ابعاد پایین ۳۰، و به ازای $k=30$ در روش SRML غیرخطی، برابر با ۹۱٫۲۶ و به ازای $k=5$ و $k_i=6$ و $k_i=3$ در روش M-SRML غیرخطی، که دارای گراف منیفلدی با تعداد عناصر متصل به هم یکسان با روش SRML به ازای $k=30$ است، به ترتیب برابر با ۹۱٫۹۷ و ۹۱٫۲۶ است که برابر و بهتر از نرخ تشخیص درست روش SRML است؛ چراکه به نظر می‌رسد در روش M-

SRML، با انتخاب چند داده مشابه، به‌خصوص در مراحل اولیه اجرای الگوریتم، داده‌های هم‌کلاسیب‌تری در گراف منیفلد، به یکدیگر متصل خواهند بود. همچنین روش M-SRML به ازای $k=5$ و $k_i=6$ و $k_i=3$ ، به ترتیب، در حدود ۶ و ۳ برابر پیچیدگی محاسباتی را کاهش می‌دهد.

۵- نتیجه‌گیری

در این مقاله، روش آموزش منیفلدی مبتنی بر بازنمایی تنک معرفی گردیده است. اساسی‌ترین و مهم‌ترین بخش در روش‌های آموزش منیفلد، تشکیل گراف منیفلد در فضای با ابعاد بالا و سعی در حفظ آن در فضای با ابعاد پایین است. در روش پیشنهادی، گراف منیفلدی با استفاده از الگوریتم SRML، که مبتنی بر الگوریتم تعقیب تطبیقی متعامد پیشنهاد شده است، تشکیل می‌گردد. این گراف تنک بوده و از نظر ساختاری به صورت ترکیبی از داده‌های محلی و سراسری است که مبتنی بر نتایج به دست آمده، در استخراج همزمان ویژگی‌های محلی و سراسری منیفلد، از کارایی مناسبی برخوردار است. در واقع، در تشکیل گراف منیفلد، این‌که هر داده با چه داده‌هایی در ارتباط بوده و میزان وابستگی داده با آن‌ها، که توسط ضریب مابین آنها بیان می‌شود، چگونه است، اصل و اساس استخراج ویژگی‌های ساختاری منیفلد است. در روش پیشنهادی، پس از تشکیل گراف منیفلد، روش خطی مبتنی بر ماتریس نگاشت و روش غیرخطی مبتنی بر حفظ بازنمایی خطی هر داده، جهت استخراج داده‌های تعبیه شده معرفی شده است. همچنین روشی جهت برخورد با داده‌های آزمون و بازنمایی آنها در فضای با ابعاد پایین، معرفی گردیده است. نتایج شبیه‌سازی و مقایسه آن با دیگر روش‌های متداول آموزش منیفلد، بر روی دو پایگاه داده‌های حروف و ارقام دست‌نویس فارسی HODA و IFHCDB، بیانگر کارایی بهتر روش پیشنهادی SRML مبتنی بر معیار نرخ تشخیص درست است. شکل‌های ۵ تا ۷ نیز، نشان‌دهنده‌ی نمودار متوسط و انحراف معیار نرخ خطای تشخیص به ازای k ها (تعداد داده‌های متصل به هر داده در گراف) و d های (بُعد فضای بازنمایی با ابعاد پایین) گوناگون است. بر اساس این نمودارها، به ازای k برابر با ۲۰ و d برابر با ۶۰، الگوریتم SRML بر روی پایگاه داده‌های HODA از عمل‌کرد مناسبی برخوردار بوده و افزایش d و k کمکی به بهبود نرخ تشخیص درست نکرده و تنها موجب افزایش پیچیدگی گراف منیفلد خواهد شد.

همچنین روش M-SRML، جهت کاهش میزان پیچیدگی محاسباتی الگوریتم پیشنهادی معرفی گردیده است که نتایج آن بر روی پایگاه داده‌های HODA نشان داده شده است. نتایج به دست آمده، نشان می‌دهد که با کاهش در حدود ۶ برابری میزان پیچیدگی محاسباتی، نتایجی مشابه روش SRML قابل حصول است.

بنابر نتایج به‌دست آمده، چگونگی تعریف و تشکیل گراف منیفلد، نقشی اساسی در درک و استخراج داده‌های تعبیه شده داشته و بیان

هرچه بهتر گراف منیفلد حاصل از داده‌ها در فضای با ابعاد بالا، موضوع می‌تواند مورد توجه محققان قرار گیرد. می‌تواند به بهبود کارایی روش‌های آموزش منیفلد کمک نماید، که این

جدول (۱): کلاس‌بندی حروف دست‌نویس فارسی.

شماره کلاس	۱	۲	۳	۴	۵	۶	۷	۸	۹	۱۰	۱۱	۱۲	۱۳	۱۴	۱۵	۱۶	۱۷	۱۸
اعضای کلاس	الف	ب، پ، ت، ث	ج، چ، ح، خ	د، ذ	ر، ز، ژ	س، ش	ص، ض	ط، ظ	ع، غ	ف	ق	ک، گ	ل	م	ن	و	ه	ی

جدول (۲): نرخ تشخیص درست روش‌های PCA، LPP، LEM، LLE و روش پیشنهادی SRML بر روی پایگاه داده‌های HODA.

PCA	k=۳۰						k=۲۰						k=۱۵						d
	LPP	LEM	LLE	SRML		LPP	LEM	LLE	SRML		LPP	LEM	LLE	SRML					
				خطی	غیر خطی				خطی	غیر خطی				خطی	غیر خطی				
۸۱/۵۲	۷۷/۵۷	۷۷/۸۱	۷۲/۲۱	۷۷/۵۸	۸۲/۴۲	۷۸/۲۲	۷۸/۶۷	۷۳/۵۷	۷۹/۵۹	۸۳/۴۷	۷۸/۸۵	۷۹/۵۰	۷۵/۷۲	۸۰/۲۱	۸۱/۴۷	۱۰			
۸۴/۳۰	۸۰/۸۸	۷۸/۸۴	۷۵/۱۱	۸۱/۳۹	۸۶/۳۲	۸۱/۳۱	۸۰/۰۷	۷۴/۵۷	۸۲/۶۶	۶/۹۵	۸۱/۴۷	۸۰/۵۰	۷۵/۳۹	۸۲/۵۹	۸۵/۲۹	۱۲			
۸۶/۶۴	۸۵/۵۱	۷۹/۷۹	۷۷/۴۲	۸۴/۸۴	۸۸/۸۸	۸۵/۷۳	۸۰/۹۲	۷۶/۰۹	۸۶/۳۹	۸۸/۹۷	۸۶/۱۷	۸۱/۳۱	۷۸/۰۳	۸۶/۸۰	۸۷/۱۵	۱۵			
۸۸/۰۷	۸۷/۴۹	۸۰/۵۶	۷۹/۷۲	۸۹/۳۶	۹۰/۹۸	۸۷/۷۱	۸۱/۵۴	۷۷/۶۶	۸۹/۸۱	۹۰/۳۲	۸۸/۰۶	۸۱/۸۶	۷۹/۴۸	۸۹/۸۳	۸۷/۹۴	۲۰			
۸۸/۵۸	۸۹/۶۴	۸۰/۸۲	۸۰/۰۱	۸۹/۹۳	۹۱	۸۹/۸۲	۸۱/۷۰	۷۸/۷۶	۹۰/۷۴	۹۰/۳۸	۹۰/۰۲	۸۲/۴۵	۷۹/۹۴	۹۱/۱۶	۸۸/۵۷	۲۵			
۸۸/۸۴	۹۰/۶۳	۸۱/۵۶	۸۱/۰۱	۹۱	۹۱/۲۶	۹۱/۰۸	۸۲/۳۲	۸۰/۶۱	۹۱/۸۲	۹۰/۸۵	۹۱/۰۲	۸۲/۹۸	۸۰/۳۹	۹۱/۸۹	۸۹/۳۹	۳۰			

جدول (۳): نرخ تشخیص درست روش‌های PCA، LPP، LEM، LLE و روش پیشنهادی SRML بر روی پایگاه داده‌های JFHCDB.

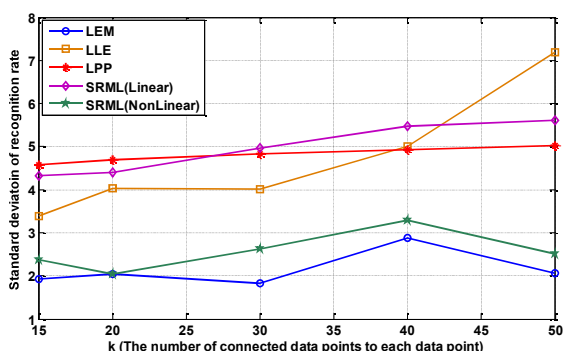
PCA	k=۳۰					k=۲۰					k=۱۵					d
	LPP	LEM	LLE	SRML		LPP	LEM	LLE	SRML		LPP	LEM	LLE	SRML		
				خطی	غیر خطی				خطی	غیر خطی				خطی	غیر خطی	
۸۴/۷۸	۸۳/۱۶	۷۸/۸۱	۷۸/۸۴	۸۴/۳۳	۸۶/۳۱	۸۳/۷۰	۷۹/۳۵	۷۹/۶۱	۸۵/۱۷	۷۸/۶۳	۸۴/۰۲	۷۹/۷۷	۷۸/۳۹	۸۵/۲۴	۷۷/۳۱	۱۰
۸۶/۰۴	۸۶/۲۸	۷۹/۲۱	۷۹/۸۹	۸۶/۹۷	۸۶/۹۱	۸۶/۶۷	۷۹/۵۸	۸۰/۴۴	۷۸/۶۸	۷۹/۶۱	۸۶/۸۵	۸۰/۱۵	۸۰/۳۴	۸۸/۱۰	۸۰/۳۶	۱۲
۸۷/۲۹	۸۷/۹۴	۷۹/۶۲	۸۱/۱۱	۸۹/۲۰	۸۷/۸۵	۸۸/۳۱	۸۰/۳۲	۸۱/۲۶	۹۰/۳۴	۸۰/۸۲	۸۸/۶۷	۸۰/۷۳	۸۰/۸۰	۹۰/۳۰	۸۱/۶۶	۱۵
۸۸/۰۹	۸۹/۷۳	۸۰/۹۵	۸۲/۴۹	۹۰/۴۲	۸۹/۴۳	۹۰/۱۱	۸۱/۳۶	۸۲/۰۳	۹۱/۳۸	۸۲/۴۳	۹۰/۳۵	۸۱/۷۶	۸۲/۲۶	۹۱/۵۶	۸۳	۲۰
۸۸/۵۶	۹۰/۵۴	۸۱/۱۹	۸۳/۵۴	۹۲/۰۱	۸۹/۸۴	۹۱/۰۷	۸۱/۶۹	۸۲/۲۲	۹۲/۴۸	۸۳/۴۱	۹۱/۲۲	۸۲/۵۰	۸۲/۹۸	۹۲/۶۷	۸۳/۵۰	۲۵
۸۸/۸۸	۹۱/۴۱	۸۱/۶۷	۸۳/۸۵	۹۲/۶۸	۹۰/۲۴	۹۱/۶۱	۸۲/۳۰	۸۲/۶۱	۹۳/۲۱	۸۵/۵۰	۹۱/۸۵	۸۲/۷۶	۸۳/۲۶	۹۳/۳۱	۸۵/۸۶	۳۰

جدول (۴): نرخ تشخیص درست روش‌های PCA، LPP، LEM، LLE و روش پیشنهادی SRML بر روی پایگاه داده‌های HODA برای d از ۴۰ تا ۱۵۰.

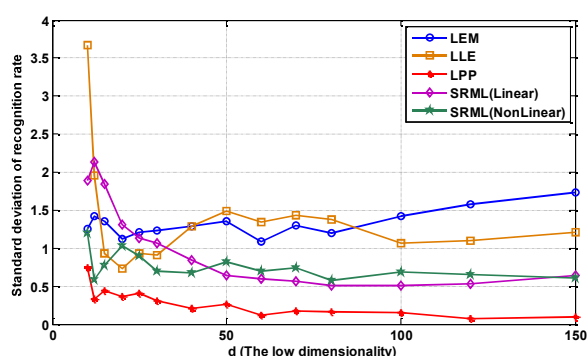
PCA	k=۳۰						k=۲۰						k=۱۵						d
	LPP	LEM	LLE	SRML		LPP	LEM	LLE	SRML		LPP	LEM	LLE	SRML					
				خطی	غیر خطی				خطی	غیر خطی				خطی	غیر خطی				
۸۹/۴۴	۹۱/۶۹	۸۱/۹۱	۸۲/۰۸	۹۲/۷۱	۹۱/۴۶	۹۱/۸۰	۸۲/۹۳	۸۱/۷۳	۹۲/۹۹	۹۰/۹۵	۹۱/۹۴	۸۳/۸۲	۸۲/۲۰	۹۳/۱۶	۸۹/۸۸	۴۰			
۸۹/۷۲	۹۲/۳۶	۸۲/۶۶	۸۲/۶۶	۹۲/۹۶	۹۱/۴۲	۹۲/۴۵	۸۳/۵۶	۸۲/۹۱	۹۲/۹۹	۹۰/۹۹	۹۲/۵۴	۸۴/۰۵	۸۲/۶۷	۹۳/۱۷	۸۹/۶۵	۵۰			
۸۹/۸۱	۹۲/۶۵	۸۲/۵۴	۸۳/۶۷	۹۳/۲۳	۹۱/۶۷	۹۲/۵۰	۸۳/۷۱	۸۳/۷۴	۹۳/۶۰	۹۰/۷۹	۹۲/۶۹	۸۴/۲۱	۸۳/۲۴	۹۳/۶۱	۸۹/۹۸	۶۰			
۸۹/۹۴	۹۲/۹۷	۸۳/۳۹	۸۳/۷۹	۹۳/۵۸	۹۱/۸۲	۹۲/۹۸	۸۴/۲۰	۸۴/۰۴	۹۳/۶۱	۹۰/۶۸	۹۳/۱۸	۸۵/۰۴	۸۴/۰۵	۹۳/۶۴	۹۰/۰۷	۷۰			
۹۰/۰۳	۹۲/۸۲	۸۳/۲۴	۸۴/۳۸	۹۳/۴۳	۹۱/۸۱	۹۳/۰۲	۸۴/۳۳	۸۴/۲۲	۹۳/۶۳	۹۰/۸۹	۹۳/۱۹	۸۵	۸۴/۲۵	۹۳/۵۹	۹۰/۳۰	۸۰			
۹۰/۲۲	۹۳/۰۶	۸۳/۴۲	۸۵/۶۴	۹۳/۶۶	۹۱/۹۴	۹۳/۲۱	۸۴/۶۰	۸۵/۲۸	۹۳/۸۵	۹۰/۹۰	۹۳/۴۴	۸۵/۰۶	۸۵/۱۳	۹۳/۸۱	۹۰	۱۰۰			
۹۰/۳۰	۹۳/۵۴	۸۳/۶۱	۸۵/۶۶	۹۳/۵۹	۹۱/۵۱	۹۳/۶۱	۸۴/۹۱	۸۵/۴۱	۹۳/۷۱	۹۰/۶۵	۹۳/۶۱	۸۵/۶۸	۸۵/۶۳	۹۳/۸۸	۸۹/۷۹	۱۲۰			
۹۰/۴۶	۹۳/۴۹	۸۳/۷۶	۸۶/۲۷	۹۳/۴۲	۹۱/۲۸	۹۳/۶۶	۸۵/۲۹	۸۵/۷۶	۹۳/۶۳	۹۰/۶۰	۹۳/۷۶	۸۵/۸۹	۸۶/۰۴	۹۳/۸۹	۸۹/۶۷	۱۵۰			

جدول (۵): نرخ تشخیص درست روش‌های LLE, LEM, LPP و روش پیشنهادی SRML بر روی پایگاه داده‌های HODA به ازای k برابر با ۴۰ و ۵۰ و ابعاد d از ۱۰ تا ۱۵۰.

d	k=۵۰					k=۴۰				
	LPP	LEM	LLE	SRML		LPP	LEM	LLE	SRML	
				خطی	غیر خطی				خطی	غیر خطی
۱۰	۷۶/۷۵	۷۵/۹۹	۶۴/۸۵	۷۵/۲۷	۸۲/۸۴	۷۷/۲۰	۷۶/۸۲	۷۰/۸۰	۷۶/۲۷	۸۲/۰۲
۱۲	۸۰/۵۹	۷۶/۵۶	۷۰/۱۱	۷۷/۴۵	۸۶/۲۷	۸۰/۸۱	۷۸/۰۹	۷۴/۷۳	۷۸/۵۹	۸۵/۶۰
۱۵	۸۴/۹۴	۷۷/۶۳	۷۵/۴۳	۸۲/۲۵	۸۸/۶۳	۸۵/۱۴	۷۸/۸۳	۷۶/۵۳	۸۲/۷۵	۸۷/۳۸
۲۰	۸۶/۹۷	۷۸/۸۳	۷۹/۴۰	۸۶/۵۲	۸۹/۴۲	۸۷/۳۹	۷۹/۷۷	۷۹/۰۲	۸۷/۷۹	۸۹/۹۶
۲۵	۸۸/۸۸	۷۹/۱۷	۸۱/۴۷	۸۸/۳۱	۹۰/۴۹	۸۹/۲۹	۷۹/۷۳	۸۰/۹۰	۸۸/۷۶	۹۰/۹۷
۳۰	۹۰/۳۱	۷۹/۵۵	۸۲/۶۳	۸۹/۱۱	۹۱/۱۸	۹۰/۴۵	۸۰/۴۸	۸۲/۳۲	۹۰/۰۹	۹۱/۱۱
۴۰	۹۱/۳۳	۸۰/۱۹	۸۵/۱۸	۹۰/۸۴	۹۱/۴۵	۹۱/۷۵	۸۱/۰۹	۸۳/۶۶	۹۲/۰۵	۹۱/۸۴
۵۰	۹۱/۷۸	۸۰/۲۸	۸۶/۳۸	۹۱/۴۵	۹۱/۷۸	۹۲/۱۵	۸۱/۶۷	۸۴/۸۰	۹۲/۲۸	۹۱/۹۵
۶۰	۹۲/۳۹	۸۱/۲۹	۸۶/۷۱	۹۲/۰۵	۹۱/۶۶	۹۲/۴۱	۸۱/۹۰	۸۵/۶۴	۹۲/۶۸	۹۱/۷۸
۷۰	۹۲/۶۷	۸۱/۶۵	۸۷/۵۳	۹۲/۳۸	۹۱/۷۷	۹۲/۸۵	۸۱/۹۲	۸۵/۷۰	۹۲/۷۱	۹۱/۹۳
۸۰	۹۲/۷۵	۸۱/۷۱	۸۷/۷۳	۹۲/۳۰	۹۱/۵۷	۹۲/۸۸	۸۲/۳۳	۸۵/۹۷	۹۲/۸۶	۹۱/۷۰
۱۰۰	۹۳/۱۳	۸۱/۷۷	۸۷/۹۳	۹۲/۵۵	۹۱/۵۶	۹۳/۰۱	۸۲/۳۵	۸۶/۷۹	۹۳/۰۲	۹۱/۶۱
۱۲۰	۹۳/۴۲	۸۲/۴۱	۸۸/۳۴	۹۲/۴۲	۹۱/۳۸	۹۳/۵۱	۸۲/۸۸	۸۶/۸۲	۹۳/۰۵	۹۱/۳۹
۱۵۰	۹۳/۶۴	۸۲/۷۷	۸۸/۹۴	۹۲/۱۸	۹۰/۹۴	۹۳/۵۵	۸۳/۱۷	۸۷/۷۱	۹۲/۶۶	۹۱/۳۴



(ب)



(الف)

شکل (۷): نمودار انحراف معیار نرخ تشخیص درست (الف) بر حسب d و (ب) بر حسب k ، بر روی پایگاه داده‌های HODA.

جدول (۶): میزان نرخ تشخیص درست با استفاده از الگوریتم M-SRML بر روی پایگاه داده‌های HODA.

d	k=۱۰				k=۵					
	$k_r = 3$		$k_r = 2$		$k_r = 6$		$k_r = 4$		$k_r = 3$	
	خطی	غیر خطی	خطی	غیر خطی	خطی	غیر خطی	خطی	غیر خطی	خطی	غیر خطی
۱۰	۷۲/۹۳	۸۳/۰۴	۷۷/۲۶	۸۲/۲۴	۷۴/۲۲	۸۳/۲۴	۷۶/۳۱	۸۱/۷۸	۷۶/۷۹	۷۷/۷۸
۱۲	۷۹/۶۵	۸۷/۰۹	۸۱/۱۷	۸۴/۱۸	۸۱/۱۶	۸۶/۵۷	۸۲/۲۵	۸۴/۹۱	۸۲/۲۴	۸۲/۵۱
۱۵	۸۴/۶۴	۸۹/۲۰	۸۶/۲۶	۸۷/۲۴	۸۵/۶۳	۸۹/۴۰	۸۵/۸۸	۸۷/۳۴	۸۵/۷۹	۸۶/۰۲
۲۰	۸۸/۶۷	۹۰/۷۹	۸۹/۸۴	۸۸/۵۰	۸۹	۹۱/۴۰	۸۹/۷۷	۸۹/۶۰	۸۹/۸۶	۸۷/۹۳
۲۵	۸۹/۹۲	۹۰/۹۳	۹۱/۴۴	۸۹/۰۳	۹۰/۶۴	۹۱/۷۸	۹۱/۳۱	۹۱	۹۱/۳۹	۸۸/۹۶
۳۰	۹۱/۴۷	۹۱/۲۶	۹۲/۱۴	۸۹/۸۰	۹۱/۴۶	۹۱/۹۷	۹۱/۶۴	۹۰/۹۸	۹۱/۸۹	۸۹/۵۷

مهندسی برق و الکترونیک ایران، جلد ۱۳، شماره ۲، ص. ۹۳-۱۰۲، ۱۳۹۵.

[۲] پارسا پویا، صفابخش رضا، "روش جدید تقطیع تصویر بر مبنای خوشه‌بندی فازی مبتنی بر تکامل تفاضلی چندهدفه"، مجله مهندسی برق و الکترونیک ایران، جلد ۱۳، شماره ۲، ص. ۱۰۳-۱۱۴، ۱۳۹۵.

مراجع

[۱] رحمتی مریم، اصغری بجستانی محمد رضا، "ارائه روشی برای حذف خطای نوارشدگی در تصاویر سنجنده های آرایه خطی"، مجله

[۲۰] آتشبار محمود، کهانی محمدحسین، "جهت یابی چند گوینده با استفاده از روش WCSSDOA"، مجله مهندسی برق و الکترونیک ایران، جلد ۱۳، شماره ۲، ص. ۶۱-۷۴، ۱۳۹۵.

- [21] Cai, T. T., Wang, L., "Orthogonal matching pursuit for sparse signal recovery with noise", IEEE Transactions on Information Theory, Vol. 57, No. 7, pp. 4680-4688, 2011.
- [22] Bengio, Y., Paiement, J. F., Vincent, P., Delalleau, O., Le Roux, N., Ouimet, M., "Out-of-sample extensions for LLE, Isomap, MDS, eigenmaps, and spectral clustering", Advances in Neural Information Processing Systems, Cambridge, MA, MIT Press, Vol. 1, pp. 177-184, 2003.
- [23] Hossein Khosravi, Ehsanollah Kabir., "Introducing a Very Large Dataset of Handwritten Farsi Digits and a Study on their Varieties", Pattern Recognition Letters, Vol. 28, No. 10, pp. 1133-1141, 2007.
- [24] Mozaffari, S., Faez, K., Faradji, F., Ziaratban, M., Golzan, S. M., "A comprehensive isolated Farsi/Arabic character database for handwritten OCR research", In 10th International Workshop on Frontiers in Handwriting Recognition. Suvisoft, pp. 385-389, 2006.
- [25] Salimi, H., Giveki, D., "Farsi/Arabic handwritten digit recognition based on ensemble of SVD classifiers and reliable multi-phase PSO combination rule", International Journal on Document Analysis and Recognition (IJDAR), Vol. 16, No. 4, pp. 371-386, 2013.
- [26] Sajedi, H., "Handwriting recognition of digits, signs, and numerical strings in Persian", Computers & Electrical Engineering, Vol. 49, No. 1, pp. 52-65, 2016.
- [27] Parseh, M. J., Meftahi, M., "A New Combined Feature Extraction Method for Persian Handwritten Digit Recognition", International Journal of Image and Graphics, Vol. 17, No. 2, 2017.
- [28] Parvin, H., Alizadeh, H., "Classifier ensemble based class weighting", American Journal of Scientific Research, Vol 19, pp. 84-90, 2011.
- [29] Niya A. M., Sajed H., "Recognition of individual handwritten letters of the Farsi language using a decision tree", International Journal of Computer Application, Vol. 55, No. 5, pp. 7-11, 2012.
- [30] Arbain, N. A., Azmi, M. S., Melhem, L. B., Kamilah, A., "Enhancement of Triangle Coordinates For Triangle Features For Better Classification", Jordanian Journal of Computers and Information Technology, Vol. 2, No. 2, pp. 107-118, 2016.
- [31] Liu, C. L., Suen, C. Y., "A new benchmark on the recognition of handwritten Bangla and Farsi numeral characters", Pattern Recognition, Vol. 42, No. 12, pp. 3287-3295, 2009.
- [32] Moradi, M., Eshghi, M., Shahbazi, K., "A trainless recognition of handwritten Persian/Arabic letters using primitive elements", International Journal of Computer Applications, Vol. 142, No. 11, 2016.
- [33] Serre, T., Riesenhuber, M., "Realistic modeling of simple and complex cell tuning in the HMAX model, and implications for invariant object recognition in cortex", Technical Report CBCL Paper 239 / AI Memo 2004-017, Massachusetts Institute of Technology, Cambridge, MA, July 2004.
- [34] Borji A, Hamidi M, Mahmoudi F, "Robust handwritten character recognition with features inspired by visual ventral stream", Neural Process Letter, Vol. 28, pp. 97-111, 2008.

[۳] کامران رحیم، نظام آبادی پور حسین، سریزدی سعید، "ترمیم تصاویر رنگی با نواحی مخدوش بزرگ براساس تجزیه تصویر به مولفه های بافت و ساختار"، مجله مهندسی برق و الکترونیک ایران، جلد ۸ شماره ۲، ص. ۱۳-۲۴، ۱۳۹۰.

- [4] Meng, L., Breilkopf, P., Le Quilliec, G., Raghavan, B., Villon, P., "Nonlinear shape-manifold learning approach: concepts, tools and applications", Archives of Computational Methods in Engineering, pp. 1-21, 2016.
- [5] He, X., Yan, S., Hu, Y., Niyogi, P., Zhang, H. J., "Face recognition using laplacianfaces", IEEE Transactions on Pattern Analysis and Machine Intelligence, Vol. 27, No. 3, pp. 328-340, 2005.
- [6] Cai, D., He, X., Han, J., Zhang, H. J., "Orthogonal laplacianfaces for face recognition", IEEE Transactions on Image Processing, Vol. 15, No. 11, pp. 3608-3614, 2006.
- [7] He, X., Niyogi, P., "Locality preserving projections", Advances in Neural Information Processing Systems, Vol. 16, 2003.
- [8] Jing, D., Li, B., "Distance-weighted manifold learning in facial expression recognition", In Industrial Electronics and Applications (ICIEA), 2016 IEEE 11th Conference on, pp. 1771-1775, 2016.
- [9] Lokare, N., Qian G., Wesley S., Zoe J., Sky A., Edgar L., "Manifold learning approach to curve identification with applications to footprint segmentation", In Computational Intelligence for Multimedia, Signal and Vision Processing (CIMSIVP), 2014 IEEE Symposium on, pp. 1-8, 2014.
- [10] Ma, L., Melba M. C., Xiaoquan Y., Yan G., "Local-manifold-learning-based graph construction for semisupervised hyperspectral image classification", IEEE Transactions on Geoscience and Remote Sensing, Vol 53, No. 5, pp. 2832-2844, 2015.
- [11] Saul, L. K., Roweis, S. T., "Think globally, fit locally: unsupervised learning of low dimensional manifolds", Journal of Machine Learning Research, Vol 4, pp. 119-155, 2003.
- [12] Zhang, Z., Zha, H., "Principal manifolds and nonlinear dimensionality reduction via tangent space alignment", SIAM Journal on Scientific Computing, Vol. 26, No. 1, pp. 313-338, 2004.
- [13] Belkin, M., Niyogi, P., "Laplacian eigenmaps for dimensionality reduction and data representation", Neural Computation, Vol. 15, No. 6, pp. 1373-1396, 2003.
- [14] Hinton, G., Roweis, S., "Stochastic neighbor embedding", Advances in Neural Information Processing Systems, Vol. 15, pp. 833-840, 2002.
- [15] Donoho, D. L., Grimes, C., "Hessian eigenmaps: locally linear embedding techniques for high-dimensional data", Proceedings of the National Academy of Sciences, Vol. 100, No. 10, pp. 5591-5596, 2003.
- [16] Borg, I., Groenen, P. J., "Modern multidimensional scaling: theory and applications", Journal of Educational Measurement, Vol. 40, No. 3, pp. 277-280, 2003.
- [17] Weinberger, K. Q., Saul, L. K., "An introduction to nonlinear dimensionality reduction by maximum variance unfolding", AAAI Conference on Artificial Intelligence, Vol. 6, pp. 1683-1686, 2006.
- [18] Tenenbaum, J. B., De Silva, V., Langford, J. C., "A global geometric framework for nonlinear dimensionality reduction", Science, Vol. 290, No. 5500, pp. 2319-2323, 2000.
- [19] Abdi, H., Williams, L. J., "Principal component analysis", Wiley Interdisciplinary Reviews:



- ¹ Locality Preserving Projection
- ² Orthogonal LPP
- ³ Lokare
- ⁴ Semi Supervised
- ⁵ Hyper Spectral
- ⁶ Locally Linear Embedding
- ⁷ Local Tangent Space Alignment
- ⁸ Laplacian Eigenmap
- ⁹ Stochastic Neighborhood Embedding
- ¹⁰ Hessian LLE
- ¹¹ Multidimensional Scaling
- ¹² Maximum Variance Unfolding
- ¹³ Principal Component Analysis
- ¹⁴ Orthogonal Matching Pursuit
- ¹⁵ K-nearest neighbors
- ¹⁶ ϵ -ball
- ¹⁷ Matching Pursuit
- ¹⁸ Over Complete
- ¹⁹ Sparse Representation based Manifold Learning
- ²⁰ Local invariance layer
- ²¹ Intermediate feature layer