



طبقه‌بندی تصاویر با استفاده از نمایش تُنک و تطبیق زیرفضا

فریمه شرافتی و جعفر طهمورث‌نژاد*

دانشکده مهندسی فناوری اطلاعات و کامپیوتر، دانشگاه صنعتی ارومیه، ارومیه، ایران

چکیده

در بیش‌تر الگوریتم‌های یادگیری ماشین و پردازش تصویر، فرض اولیه بر این است که توزیع احتمال داده‌های آموزشی (دامنه منبع) و آزمایش (دامنه هدف) یکسان است؛ اما در کاربردهای دنیای واقعی، برخی معیارها نظیر حالت تصویر، روشنایی یا کیفیت تصویر، موجب ایجاد اختلاف قابل‌توجهی بین دو مجموعه آموزشی و آزمایش می‌شود. به همین دلیل، اغلب مدل‌های ایجادشده بر روی داده‌های آموزشی عملکرد ضعیفی بر روی داده‌های آزمایش خواهند داشت؛ با این حال، روش‌های تطبیق دامنه، راه‌حل بسیار مؤثری برای کاهش اختلاف توزیع بین دامنه‌های آموزشی و آزمایش هستند. در این مقاله یک روش تطبیق دامنه با عنوان نمایش تُنک و تطبیق زیرفضا (SRSA) پیشنهاد شده است، که با وزن‌دهی مجدد نمونه‌های آزمایش و نگاشت داده‌ها به یک زیرفضای جدید مشکل اختلاف توزیع داده‌ها را به‌خوبی مرتفع می‌سازد. SRSA با استفاده از یک نمایش تُنک، بخشی از مجموعه داده‌های هدف را که ارتباط قوی‌تری با داده‌های منبع دارند، انتخاب می‌کند؛ علاوه بر آن، SRSA با نگاشت داده‌های تُنک هدف و داده‌های منبع به زیرفضاهای مستقل، اختلاف توزیع آنها را در فضای به‌دست آمده کاهش می‌دهد؛ در نهایت با برروی هم‌گذاری زیرفضاهای نگاشت‌شده، SRSA اختلاف توزیع بین داده‌های آموزشی و آزمایش را به کمینه می‌رساند. ما روش پیشنهادی خود را با ترتیب‌دادن چهارده آزمایش بر روی پایگاه داده‌های بصری مختلف مورد ارزیابی قرار داده و با مقایسه نتایج به‌دست آمده، نشان داده‌ایم که SRSA عملکرد بهتری در مقایسه با جدیدترین روش‌های یادگیری ماشین و تطبیق دامنه دارد.

واژگان کلیدی: طبقه‌بندی تصویر، تطبیق دامنه‌های بصری، نمایش تُنک، تطبیق زیرفضا.

Image Classification via Sparse Representation and Subspace Alignment

Farimah Sherafati & Jafar Tahmoresnezhad*

Faculty of IT & Computer Engineering, Urmia University of Technology, Urmia, Iran

Abstract

Image representation is a crucial problem in image processing where there exist many low-level representations of image, i.e., SIFT, HOG and so on. But there is a missing link across low-level and high-level semantic representations. In fact, traditional machine learning approaches, e.g., non-negative matrix factorization, sparse representation and principle component analysis are employed to describe the hidden semantic information in images, where they assume that the training and test sets are from same distribution. However, due to the considerable difference across the source and target domains result in environmental or device parameters, the traditional machine learning algorithms may fail.

Transfer learning is a promising solution to deal with above problem, where the source and target data obey from different distributions. For enhancing the performance of model, transfer learning sends the

* Corresponding author

*نویسنده عهده‌دار مکاتبات

knowledge from the source to target domain. Transfer learning benefits from sample reweighting of source data or feature projection of domains to reduce the divergence across domains.

Sparse coding joint with transfer learning has received more attention in many research fields, such as signal processing and machine learning where it makes the representation more concise and easier to manipulate. Moreover, sparse coding facilitates an efficient content-based image indexing and retrieval.

In this paper, we propose image classification via Sparse Representation and Subspace Alignment (SRSA) to deal with distribution mismatch across domains in low-level image representation. Our approach is a novel image optimization algorithm based on the combination of instance-based and feature-based techniques. Under this framework, we reweight the source samples that are relevant to target samples using sparse representation. Then, we map the source and target data into their respective and independent subspaces. Moreover, we align the mapped subspaces to reduce the distribution mismatch across domains. The proposed approach is evaluated on various visual benchmark datasets with 14 experiments. Comprehensive experiments demonstrate that SRSA outperforms other latest machine learning and domain adaptation methods with significant difference.

Keywords: Image classification, Visual domains adaptation, Sparse representation, Subspace alignment.

۱- مقدمه

امروزه با گسترش روزافزون روش‌های مختلف جمع‌آوری اطلاعات دیجیتال مانند پوشش‌گرها^۱ و دوربین‌های دیجیتالی، پردازش تصویر کاربرد فراوانی یافته است. پردازش تصویر بیشتر به موضوع پردازش تصاویر دیجیتال گفته می‌شود که شاخه‌ای از علم رایانه است که با پردازش سیگنال‌های دیجیتال که نماینده تصاویر گرفته‌شده با دوربین دیجیتال یا پوشش‌شده توسط پوشش‌گر هستند، سر و کار دارد. با افزایش تصاویر دیجیتال مسأله طبقه‌بندی تصاویر به یکی از نیازهای اساسی تبدیل شده است. حوزه‌های کار با تصاویر در چهار زمینه بهبود کیفیت ظاهری^۲، بازسازی تصاویر مختل شده^۳، فشرده‌گی و رمزگذاری تصویر^۴ و درک و طبقه‌بندی تصویر توسط ماشین متمرکز می‌شود.

یادگیری ماشین^۵ مجموعه‌ای از الگوریتم‌های رایانه‌ای است که با استفاده از تصاویر آموزشی برچسب‌دار موجود، یک مدل یادگیری ساخته و از آن برای برچسب‌گذاری تصاویر آزمایش استفاده می‌کند. فرض اصلی در اغلب الگوریتم‌های یادگیری ماشین این است که داده‌های آموزشی و آزمایش دارای توزیع یکسانی هستند، اما در کاربردهای دنیای واقعی این فرض همیشه برقرار نبوده و ممکن است داده‌های آموزشی و آزمایش توزیع به‌طورکامل متفاوتی داشته باشند. در این حالت، مدل ساخته‌شده روی داده‌های آموزشی، عملکرد خوبی روی داده‌های آزمایش ندارد [1, 2].

یک راه حل ساده برای حل این مشکل جمع‌آوری مجدد داده‌های برچسب‌دار است؛ اما جمع‌آوری داده‌های برچسب‌دار نیازمند صرف هزینه و وقت بسیار بوده و بهترین راه حل ممکن نیست. راه حل دیگر استفاده از داده‌های برچسب‌دار از منابع دیگر است، که این راه حل هم منجر به مشکل دیگری به نام انتقال دامنه‌ها^۶ می‌شود. انتقال دامنه‌ها درواقع به معنای اختلاف توزیع در داده‌های آموزشی و آزمایش است که باعث کاهش کارایی مدل می‌شود. برای کاهش اثرات انتقال دامنه‌ها بر روی مدل از یک راه حل شناخته‌شده به نام یادگیری انتقالی^۷ یا تطبیق دامنه‌ها^۸ استفاده می‌شود. یادگیری انتقالی با انتقال دانش از یک مجموعه برچسب‌دار به یک مجموعه بدون برچسب باعث افزایش کارایی مدل‌های طبقه‌بندی تصاویر می‌شود. درکل مسائلی که در حوزه یادگیری انتقالی مطرح می‌شوند با توجه به نوع داده مجموعه آزمایش به دو دسته کلی تقسیم می‌شوند: (۱) یادگیری بدون نظارت^۹، (۲) یادگیری نیمه‌نظارت‌شده^{۱۰}. در یادگیری بدون نظارت اگرچه مجموعه داده آموزشی به‌طورکامل برچسب‌دار است ولی مجموعه داده آزمایش به‌طورکامل بدون برچسب بوده و مدل ایجادشده بایستی بتواند این مجموعه را برچسب‌گذاری کند. در یادگیری نیمه‌نظارت‌شده مجموعه آموزشی به‌طورکامل دارای برچسب بوده و علاوه بر آن بخش کوچکی از مجموعه داده آزمایش نیز دارای برچسب است. در این مجموعه، تعداد داده‌های برچسب‌دار مجموعه آزمایش، به‌قدری نیست که بتوان با استفاده از آن یک مدل دقیق و درست ایجاد کرد.

⁶ Domains Shift

⁷ Transfer Learning

⁸ Domains Adaptation

⁹ Unsupervised Learning

¹⁰ Semi-supervised Learning

¹ Scanners

² Enhancement

³ Restoration

⁴ Compression and Coding

⁵ Machine Learning

منبع و هدف به زیرفضای کم‌بعد مبتنی بر نظریه کاهش بُعد و تطبیق زیرفضاهای جدید انجام می‌شود که توزیع دامنه‌ها را به یکدیگر نزدیک می‌کند.

عملکرد روش پیشنهادی بر روی مجموعه داده‌های تشخیص اشیا (آفیس و کالتک) و اعداد دست‌نویس (MNIST, USPS) مورد بررسی قرار گرفته و با روش‌های مشابه در این حوزه مقایسه شده است. نتایج به‌دست‌آمده از آزمایش‌های جامع انجام‌شده، نشان‌دهنده برتری قابل‌ملاحظه روش پیشنهادی است.

ادامه این مقاله به‌صورت زیر سازماندهی شده است. مروری بر کارهای گذشته در بخش دوم مقاله و تعاریف در بخش سوم آورده شده است. بخش چهارم روش پیشنهادی، آزمایش‌های جامع انجام‌گرفته بر روی پایگاه‌داده‌های مختلف در بخش پنجم و در انتها، نتیجه‌گیری و کارهای پیشنهادی آینده آورده شده است.

۲- مروری بر کارهای گذشته

تطبیق و انتقال دامنه‌ها دو مسأله اساسی در یادگیری ماشین هستند و در سال‌های اخیر توجه بسیاری از پژوهش‌گران را در حوزه‌هایی همچون پردازش زبان‌های طبیعی و بینایی ماشین به‌خود جلب کرده است. روش‌های زیادی برای کاهش اختلاف توزیع بین دامنه‌ها پیشنهاد شده است، ولی با این‌حال سه رویکرد کلی در این حوزه وجود دارد: رویکردهای مبتنی بر نمونه، رویکردهای مبتنی بر خصوصیت و رویکردهای مبتنی بر مدل. در این مقاله، از ترکیب رویکردهای مبتنی بر نمونه و مبتنی بر خصوصیت استفاده شده که در زیر به‌تفصیل توضیح داده شده است.

۲-۱- رویکردهای مبتنی بر نمونه

در رویکردهای مبتنی بر نمونه، هدف، انتخاب نمونه‌هایی از دامنه منبع است که توزیع نزدیک‌تری با نمونه‌های دامنه هدف داشته باشند تا اختلاف توزیع بین دو دامنه کاهش یابد. لندمارک [5] یکی از روش‌های مبتنی بر نمونه است که برای محاسبه میزان اختلاف توزیع بین دامنه‌ها از MMD⁵ استفاده می‌کند و به نمونه‌های منبع که اختلاف توزیع کمتری با نمونه‌های هدف دارند، وزن بیشتری اختصاص می‌دهد.

⁵ Maximum Mean Discrepancy

در چنین شرایطی از یادگیری انتقالی جهت انتقال دانش بین دو مجموعه داده استفاده می‌شود. در این مقاله، بر روی مسأله یادگیری بدون نظارت تمرکز شده است.

در سال‌های اخیر شاهد توسعه روزافزون استفاده از مفهوم نمایش تنگ^۱ در کاربردهای مختلف پردازش تصویر هستیم. استفاده از نمایش تنگ در مسائل طبقه‌بندی تصاویر نتایج موفقیت‌آمیزی به همراه داشته است. ایده اصلی در روش‌های طبقه‌بندی تصاویر مبتنی بر نمایش تنگ، نمایش هر داده به‌صورت ترکیب خطی تنگ از سایر داده‌ها است؛ به‌گونه‌ای که داده‌های مشابه با داده مورد نظر در این ترکیب خطی بیشترین وزن را به خود اختصاص دهند [3].

در این مقاله، یک روش برای تطبیق دامنه بدون نظارت بین داده‌های آموزشی و آزمایش به‌نام طبقه‌بندی تصاویر با استفاده از نمایش تنگ و تطبیق زیرفضا (SRSA²) پیشنهاد شده است. روش پیشنهادی با استفاده از نمایش تنگ، نمونه‌هایی از دامنه منبع که شباهت بیشتری را با نمونه‌های دامنه هدف دارند انتخاب می‌کند، و با استفاده از PCA³ [4] که یک روش شناخته‌شده برای کاهش ابعاد فضای ویژگی داده‌ها می‌باشد، داده‌های انتخاب‌شده از دامنه منبع را به زیرفضای منبع و داده‌های دامنه هدف را به زیرفضای هدف نگاشت می‌کند، به‌گونه‌ای که خطای بازسازی^۴ داده‌ها در زیرفضاهای جدید کمینه باشد. همچنین روش پیشنهادی برای کاهش تغییرات بین دامنه‌ای، زیرفضاهای دامنه‌های منبع و هدف را با یکدیگر تطبیق می‌دهد. این هدف با یافتن یک ماتریس نگاشت امکان‌پذیر می‌شود، به‌طوری که با اعمال این ماتریس نگاشت بر روی داده‌های انتخاب‌شده از دامنه منبع در زیرفضای جدید، اختلاف میان داده‌های آموزشی و آزمایش در زیرفضاهای مجزا کاهش می‌یابد.

به‌طور کلی هر کدام از مراحل روش پیشنهادی شامل فرایندهای زیر جهت ایجاد تطبیق بین دامنه‌های منبع و هدف هستند: (۱) برخلاف روش‌های دیگر که از همه داده‌های دامنه منبع برای ساخت طبقه‌بند استفاده می‌کنند، روش ما با استفاده از نمایش تنگ ابتدا داده‌های دامنه منبع را که شباهت کمی به داده‌های دامنه هدف دارند، حذف می‌کند؛ سپس، داده‌های هدف را بر اساس ترکیب خطی تنگ از داده‌های منبع نمایش می‌دهد، (۲) نگاشت دامنه‌های

¹ Sparse Representation

² image classification via Sparse Representation and Subspace Alignment (SRSA)

³ Principal Component Analysis

⁴ Reconstruction error

طبقه‌بندی‌کننده مبتنی بر نمایش تُنک¹ (SRC) [6, 7] یکی دیگر از روش‌های مبتنی بر نمونه است که در سامانه‌های تشخیص چهره نتایج قابل توجهی را به دست آورده است. ایده اصلی SRC به این صورت است که هر تصویر چهره به صورت ترکیب خطی تُنک از سایر تصاویر چهره نمایش داده می‌شود. انتخاب تصاویر با استفاده از ضرایب آن‌ها صورت می‌گیرد؛ به گونه‌ای که تصاویر با ضرایب صفر حذف شده و تصاویری که ضرایب غیرصفر و وزن بیشتری به آن‌ها تعلق گرفته است، انتخاب می‌شوند.

روش² SSSC [8]، برای کاهش اختلاف توزیع بین دامنه‌های منبع و هدف، با یافتن یک ماتریس انتخاب، نمونه‌های منبع مشابه با نمونه‌های هدف را انتخاب می‌کند و با استفاده از نمایش تُنک هر نمونه هدف را بر اساس یک ترکیب خطی تُنک از نمونه‌های منبع نمایش می‌دهد.

۲-۲- رویکردهای مبتنی بر خصوصیت

در رویکردهای مبتنی بر خصوصیت، فضای ویژگی^۳ برای نمایش بهتر ویژگی‌های مشترک دو دامنه تغییر می‌کند. GFK⁴ [9] و TCA⁵ [10] از جمله این روش‌ها هستند. داده‌های منبع و هدف را بر روی منیفلد گرسمن^۶ نگاشت می‌کند و بین زیرفضاهای منبع و هدف یک مسیر کوتاه^۷ بر روی منیفلد ایجاد می‌کند تا تغییرات دامنه از منبع تا هدف را مورد ارزیابی قرار دهد. از آنجایی که حرکت بر روی مسیر کوتاه بایستی به صورت روان و پیوسته باشد، GFK اندازه زیرفضا را کوچک در نظر می‌گیرد که موجب از بین رفتن بخشی از داده‌های ورودی می‌شود. TCA اختلاف توزیع حاشیه‌ای بین دامنه‌ها را با استفاده از MMD کم می‌کند و با بیشینه‌کردن میزان واریانس داده‌ها، ویژگی‌های اصلی آن‌ها را حفظ می‌کند. SA⁸ [11] داده‌های منبع را به زیرفضای منبع و داده‌های هدف را به زیرفضای هدف نگاشت می‌کند، و با استفاده از یک ماتریس نگاشت زیرفضاهای جدید را تطبیق می‌دهد و بدین صورت اختلاف توزیع بین دامنه‌ها را کاهش می‌دهد. SCA⁹ [12] با بیشینه‌سازی اختلاف بین نمونه‌ها در رده‌های مختلف و کمینه‌سازی اختلاف بین

نمونه‌های هر رده، دامنه‌ها را به زیرفضای جدید نگاشت می‌دهد. SDA¹⁰ [13] هم‌زمان با ایجاد تطبیق بین زیرفضاهای دامنه منبع و هدف، اختلاف توزیع بین دامنه‌ها را از طریق تطبیق واریانس مؤلفه‌های اساسی دامنه‌ها به کمینه می‌رساند. VDA¹¹ [14] برای کاهش اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌ها، یک نمایش کم‌بعد از دامنه‌های منبع و هدف ایجاد می‌کند. همچنین، VDA از خوشه‌بندی مستقل از دامنه^{۱۲} جهت حفظ ساختار هندسی و آماری بین دامنه‌ها در فضای جدید استفاده می‌کند.

روش پیشنهادی در این مقاله، یک راه حل بدون نظارت با بهره‌گیری از روش‌های مبتنی بر نمونه و مبتنی بر خصوصیت است. در این روش یک مدل طبقه‌بند توسط کاهش اختلاف توزیع بین نمونه‌های دامنه منبع و هدف با استفاده از نمایش تُنک و همچنین با استفاده از تطبیق زیرفضاها ایجاد می‌شود که دارای کمینه خطای پیش‌بینی است.

۳- تعاریف و مفاهیم اولیه

۳-۱- هدف پژوهش

تطبیق دامنه یکی از مسائل مطرح در حوزه یادگیری ماشین است. این مقاله، به دنبال پیشنهاد روشی است که بتواند بر محدودیت‌های الگوریتم‌های یادگیری ماشین کلاسیک غلبه کند. در سال‌های اخیر استفاده از نمایش تُنک در کاربردهای مختلف پردازش تصویر رواج زیادی پیدا کرده است؛ بنابراین، نمایش تُنک به عنوان مدلی مناسب جهت مدل‌سازی داده‌ها در فضای جدید استفاده می‌شود.

۳-۲- تعریف دامنه^{۱۳} و وظیفه^{۱۴}

در یادگیری انتقالی هر دامنه از دو مفهوم زیر تشکیل شده است: فضای ویژگی^{۱۳} X و توزیع احتمال حاشیه‌ای $P(X)$. بنابراین دامنه D به صورت $D = \{X, P(X)\}$ نمایش داده می‌شود. بدین ترتیب، زمانی دو دامنه متفاوت هستند که یا فضای ویژگی و یا توزیع حاشیه‌ای متفاوتی داشته باشند.

برای هر دامنه D یک وظیفه T وجود دارد که شامل مجموعه برچسب Y و تابع پیش‌بینی $f(x)$ است، یعنی

¹⁰ Subspace Distribution Alignment

¹¹ Visual Domain Adaptation via transfer feature learning

¹² Domain invariant clustering

¹³ Domain

¹⁴ Task

¹ Sparse Representation-based Classifier

² Sample Selection Sparse Coding

³ Feature Space

⁴ Geodesic Flow Kernel

⁵ Transfer Component Analysis

⁶ Grassman Manifold

⁷ Geodesic

⁸ Subspace Alignment

⁹ Scatter Component Analysis

دامنه منبع با $D_S = \{X_S, Y_S\}$ و دامنه هدف با $D_T = \{X_T, Y_T\}$ نمایش داده می‌شوند که در آن $X_T \in R^{d_T \times n_T}$ و $X_S \in R^{d_S \times n_S}$ به ترتیب داده‌های منبع و هدف هستند و $Y_S \in R^{n_S \times 1}$ و $Y_T \in R^{n_T \times 1}$ برچسب‌های منبع و هدف، d_S ابعاد نمونه‌های منبع، d_T ابعاد نمونه‌های هدف، n_S و n_T به ترتیب نشان‌دهنده تعداد نمونه‌های منبع و هدف است. فرض شده $d_S = d_T$ است و فضای ویژگی دامنه‌های منبع و هدف یکسان است.

۳-۳-۳- کاهش بُعد

برای به‌دست‌آوردن زیرفضاهای مشترک یا مستقل می‌توان از روش‌های کاهش بعد برای قرار دادن داده‌ها در یک زیرفضای تعبیه‌شده جدید استفاده کرد. تجزیه مؤلفه‌های اصلی (PCA) یک روش کاربردی برای استخراج ویژگی از داده‌ها (سیگنال یا تصویر) است. به تعبیر دیگر می‌توان با PCA کاهش بعد انجام داده و ویژگی‌های جدید استخراج کرد. PCA داده‌ها را به‌گونه‌ای به یک زیرفضای جدید نگاشت می‌کند که ساختار اصلی داده‌ها حفظ شود. بدین ترتیب، داده‌ها بر روی اجزای اصلی خود نگاشت می‌شوند که شامل اجزایی است که دارای واریانس بالایی هستند. از بین اجزای به‌دست‌آمده، اجزایی که دارای حداکثر واریانس باشند به‌عنوان اجزای اصلی جهت نگاشت داده‌ها به فضای کم‌بعد استفاده می‌شوند.

تابع بهینه‌سازی PCA با استفاده از رابطه $H = I_{n_S+n_T} - \max_{A^T A = I} \text{tr}(A^T X H X^T A)$ به‌دست می‌آید، که $\frac{1}{n_S+n_T} \vec{1} \vec{1}^T$ تابع مرکزیت بوده، I ماتریس همانی و $\vec{1}$ بردار ستونی از یک‌ها است [4]. tr نشان‌دهنده حاصل جمع عناصر قطر اصلی ماتریس است.

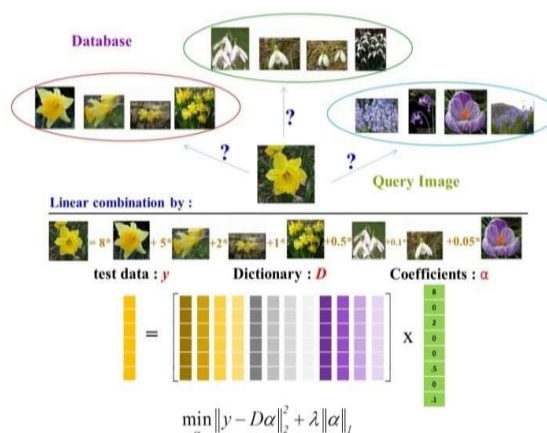
۴- روش پیشنهادی

در این مقاله یک روش نوین با عنوان SRSA پیشنهاد شده است که با استفاده از نمایش تنک نمونه‌های منبع مرتبط با نمونه‌های هدف را انتخاب کند. SRSA یک نمایش کم‌بعد برای نمونه‌های منبع انتخاب‌شده و نمونه‌های هدف، ایجاد می‌کند و برای کاهش بعد داده‌ها از PCA استفاده می‌کند. سپس، SRSA زیرفضاهای ایجادشده منبع و هدف را با یکدیگر تطبیق داده تا توزیع آنها بر روی یکدیگر منطبق شود.

$T = \{Y, f(x)\}$ تابع $f(x)$ برای نمونه‌های ورودی x ، مجموعه برچسب Y را پیش‌بینی می‌کند و می‌توان آن را به‌صورت توزیع احتمال شرطی $P(Y|x)$ بیان کرد. بدین ترتیب، زمانی دو وظیفه متفاوت هستند که با مجموعه برچسب و یا تابع توزیع احتمال شرطی متفاوتی داشته باشند؛ بدین معنی که $Y_S \neq Y_T$ یا $P_S(Y_S|X_S) \neq P_T(Y_T|X_T)$.

۳-۳-۳- نمایش تنک

امروزه نمایش تنک به‌عنوان ابزاری قوی و کارآمد برای نمایش سیگنال‌ها تبدیل شده و در کاربردهایی همچون نمونه‌برداری، بازیابی و طبقه‌بندی تصاویر مورد استفاده قرار گرفته است. هدف این نوع نمایش، تقریب داده‌های آزمایش (تصویر جدید) با استفاده از تعداد کمی (تنک) از داده‌های آموزشی است. طبقه‌بندی با استفاده از نمایش تنک به این صورت است که داده‌های آموزشی به‌عنوان دیکشنری^۱ در نظر گرفته می‌شوند و داده‌های هدف به‌صورت ترکیب خطی از آنها نمایش داده می‌شوند. همان‌طور که در شکل (۱) نشان داده شده، داده آزمایش y ، بر اساس ترکیب خطی از دیکشنری (D) نمایش داده شده است.



(شکل-۱): نحوه عملکرد نمایش تنک. y داده آزمایش، D دیکشنری، α ماتریس ضرایب و λ پارامتر تنظیم می‌باشد. داده آزمایش y با استفاده از دیکشنری و ماتریس ضرایب، به‌صورت ترکیب خطی از دیکشنری نمایش داده می‌شود. (Figure-1): Sparse representation methodology. y is test data, D shows dictionary, α is the coefficient matrix and λ illustrates the tuning parameter. Test data y is represented by a linear combination of the dictionary and coefficient matrix.

¹ Dictionary

۴-۱- انتخاب نمونه

SRSA برای انتخاب نمونه‌ها از یک نمایش تُنک استفاده می‌کند و از نرم $l_{2,1}$ برای محدود کردن ضرایب ماتریس انتخاب A بهره می‌برد [15]. SRSA با استفاده از رابطه زیر، ماتریسی به دست می‌آورد که اختلاف توزیع بین داده‌های منبع و هدف را کمینه می‌سازد و می‌توان هر نمونه هدف را به صورت ترکیب خطی از کمینه نمونه‌های منبع نمایش داد.

$$\min_A \|X_T - X_S A\|_F^2 + \rho \|A\|_{2,1} \quad (۱)$$

که در این رابطه، $\|\cdot\|_F$ نرم فروبنیوس^۱ و ماتریس $A \in R^{n_s \times n_r}$ یک ماتریس تُنک است، که با استفاده از مسأله بهینه‌سازی بالا به دست می‌آید و ρ پارامتر تنظیم پراکندگی است. با استفاده از نرم $l_{2,1}$ هر سطر ماتریس همزمان صفر یا غیرصفر می‌شود که سطرهای غیرصفر به عنوان نمونه‌های منبع مرتبط با نمونه‌های هدف انتخاب و نمونه‌های منبع انتخاب شده با X_{S1} نمایش داده می‌شوند.

۴-۲- تولید زیرفضا

با وجود اینکه داده‌های منبع و هدف در فضای D بعدی یکسان قرار دارند، با این حال توزیع حاشیه‌ای متفاوتی دارند. بنابراین، داده‌های منبع و هدف با استفاده از PCA به زیرفضاهای مربوط به خود نگاشت داده می‌شوند تا اختلاف توزیع آن‌ها کاهش یابد. برای هر دامنه، d بردار ویژه مربوط به بزرگ‌ترین مقادیر ویژه PCA انتخاب می‌شوند. این بردارهای ویژه به عنوان زیرفضاهای منبع و هدف استفاده می‌شوند که به ترتیب با P_S و P_T نمایش داده می‌شوند $(P_S, P_T \in R^{D \times d})$.

۴-۳- تطبیق دامنه با استفاده از تطبیق زیرفضا

استراتژی اصلی در زیرفضای مبتنی بر تطبیق دامنه، این گونه است که داده‌های منبع و هدف به یک زیرفضای مشترک نگاشت می‌شوند، اما در این حالت فقط ویژگی‌های مشترک از دو دامنه استخراج می‌شوند و بعضی از اطلاعات که مختص هر کدام از دامنه‌ها هستند از بین می‌روند. با این حال، SRSA داده‌های منبع و هدف هر کدام از دامنه‌ها را به زیرفضای مربوط به خود نگاشت داده و سپس، زیرفضاهای

^۱ Frobenius norm

به دست آمده از هر دامنه را با یکدیگر تطبیق می‌دهد که برای این هدف از استراتژی تطبیق زیرفضا استفاده می‌کند. با استفاده از PCA، نمونه‌های منبع انتخاب شده (X_{S1}) و نمونه‌های هدف (X_T) به زیرفضاهای مجزا نگاشت داده می‌شوند؛ سپس، زیرفضاهای منبع (P_S) و هدف (P_T) با استفاده از ماتریس نگاشت M با یکدیگر تطبیق داده می‌شوند. ماتریس نگاشت M با حداقل سازی واگرایی ماتریس برگمن^۲ زیر به دست می‌آید:

$$F(M) = \|P_S M - P_T\|_F^2 \quad (۲)$$

فرم بسته رابطه (۲) را می‌توان به صورت زیر نوشت:

$$M^* = \operatorname{argmin}_M (F(M)) \quad (۳)$$

که M^* مقدار بهینه ماتریس تطبیق را نشان می‌دهد. به دلیل این که نرم فروبنیوس برای عملیات متعامد بودن^۳ ثابت^۴ است، فرمول (۲) را می‌توان به صورت زیر بازنویسی کرد:

$$F(M) = \|P_S' P_S M - P_S' P_T\|_F^2 = \|M - P_S' P_T\|_F^2 \quad (۴)$$

با توجه به رابطه (۴)، می‌توان نتیجه گرفت که M^* بهینه با استفاده از رابطه $M^* = P_S' P_T$ به دست می‌آید و $X_a = P_S' P_S P_T$ است که X_a دامنه منبع تطبیق شده با دامنه هدف است.

ماتریس نگاشت M^* سامانه مختصات زیرفضای منبع را با سامانه مختصات زیرفضای هدف، با استفاده از بردارهای پایه منبع تطبیق می‌دهد. عملکرد روش تطبیق زیرفضا در شکل (۲) نشان داده شده است.

۴-۴- الگوریتم روش SRSA

با مشتق گیری از رابطه (۱) و مساوی قرار دادن آن با صفر، مقدار A بهینه به صورت زیر به دست می‌آید:

$$A = (\rho R + X_S^T X_S)^{-1} X_S^T X_T \quad (۵)$$

که در رابطه بالا، $R = \frac{1}{2 \|A^t\|_2}$ یک ماتریس قطری و وابسته به A است و A با یک روش تکراری به دست می‌آید. الگوریتم روش SRSA با جزئیات کامل، در الگوریتم (۱) ارائه شده است.

^۲ Bregman matrix divergence

^۳ Orthonormal

^۴ Invariant

۱-۵- مجموعه داده‌های دنیای واقعی

کارایی روش پیشنهادی در این مقاله، بر روی دو نوع پایگاه داده زیر مورد ارزیابی قرار گرفته است: (۱) آفیس و کالتک^۱، (۲) اعداد (USPS و MNIST).

مجموعه داده آفیس [16] و کالتک [17]، از مجموعه داده‌های بسیار مشهور برای تطبیق دامنه هستند. مجموعه داده آفیس از سه دامنه آمازون^۲، وبکم^۳ و DSLR^۴ تشکیل شده است که شامل تصاویری با کیفیت، رنگ و روشنایی متفاوت هستند. تصاویر هر مجموعه داده تفاوت‌های زیادی با یکدیگر دارند. دامنه آمازون شامل تصاویری است که از سایت‌های تجاری بازرگاری شده است. دامنه وبکم شامل تصاویری با وضوح پایین است که توسط دوربین‌های کم کیفیت وب گرفته شده‌اند و دامنه DSLR شامل تصاویری با وضوح بالا است که توسط دوربین‌های دیجیتالی حرفه‌ای گرفته شده‌اند. تصاویر موجود در دامنه کالتک، از وبسایت گوگل جمع‌آوری شده‌اند. با وجود این که دامنه‌های مختلف توزیع‌های متفاوتی دارند اما در آزمایش‌های ما از ده رده معنایی مشترک استفاده شده است. بر روی پایگاه داده آفیس و کالتک، دوازده آزمایش طراحی شده است که در هر یک از این آزمایش‌ها یکی از مجموعه داده‌ها (به عنوان نمونه وبکم)، به عنوان دامنه منبع و یکی از سه مجموعه داده دیگر به عنوان دامنه هدف انتخاب می‌شوند.

پایگاه داده اعداد شامل دو دامنه USPS [18] و MNIST [19] است. دامنه USPS، شامل اعداد دست‌نویس پویش شده از نامه‌های اداره پست آمریکا در اندازه‌های ۱۶*۱۶ پیکسلی است که ۷۲۶۱ داده آموزشی و ۲۰۰۷ داده آزمایش دارد. دامنه MNIST، شامل اعداد دست‌نویس جمع‌آوری شده از دانش‌آموزان دبیرستانی آمریکا و کارمندان سازمان‌های مالیاتی آمریکا در اندازه ۲۸*۲۸ پیکسلی است که شامل شش هزار داده آموزشی و هزار داده آزمایش است. این پایگاه داده شامل ده رده مختلف می‌باشد و به منظور آزمایش هر دو دامنه در شرایط یکسان، دامنه USPS_vs_MNIST(U_M) ایجاد شده است که به طور تصادفی ۱۸۰۰ نمونه از داده‌های دامنه USPS به عنوان داده‌های آموزشی و دوهزار نمونه از داده‌های دامنه MNIST به عنوان داده‌های آزمایش به کار گرفته می‌شود. با جابه‌جایی نمونه‌های آموزشی و آزمایش

¹ Office and Caltech

² Amazon

³ Webcam

⁴ Digital single-lens reflex

Algorithm 1. Algorithm of image classification via sparse representation and subspace alignment (SRSA)

Input: Source data X_S , Target data X_T , Source labels L_S , Subspace dimension d

Parameter: Regularization parameters ρ

Output: Predicted target labels L_T

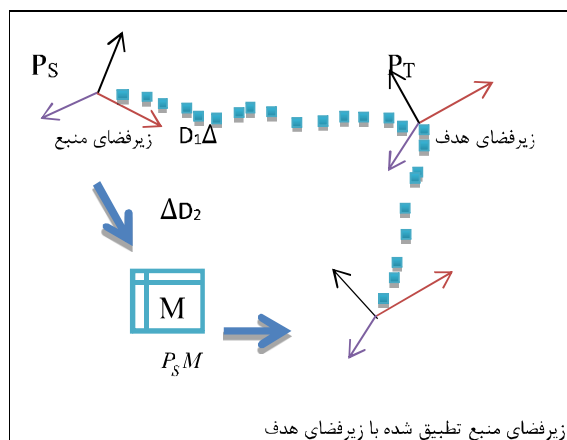
1: Repeat until the maximum number of iterations is reached

2: Compute selection matrix A via solving Eq. (5).

3: Compute diagonal matrix $N = \frac{1}{2 \|A^i\|_2}$ based on A .

4: END Repeat

5: Return $L_T = kNN(X_S A P_S P_S^T P_T, X_T P_T, L_S)$



(شکل-۲): عملکرد روش تطبیق زیرفضا. زیرفضای منبع با P_S و زیرفضای هدف با P_T نشان داده شده است. زیرفضای منبع با زیرفضای هدف تطبیق داده شده و $X_a = P_S M$ است.

(Figure-2): The performance of the subspace matching method. The source subspace is represented by P_S and the target subspace is showed by P_T . The source subspace is aligned with target subspace via $X_a = P_S M$.

در الگوریتم SRSA، هدف پیش‌بینی برچسب‌های نمونه‌های دامنه هدف با استفاده از نمایش تنگ و تطبیق زیرفضا است. SRSA در ابتدا نمونه‌های منبع را با استفاده از ماتریس انتخاب A به دست آورده، سپس، زیرفضاهای هر کدام از دامنه‌های منبع و هدف را به طور مجزا محاسبه کرده و با یکدیگر تطبیق می‌دهد و در نهایت با استفاده از طبقه‌بند NN برچسب نمونه‌های هدف را پیش‌بینی می‌کند.

۵- تنظیمات اولیه محیط آزمایش

در این بخش مجموعه داده‌ها، نتایج آزمایش‌ها و عملکرد SRSA با دیگر الگوریتم‌ها مورد مقایسه، مورد ارزیابی قرار گرفته است.

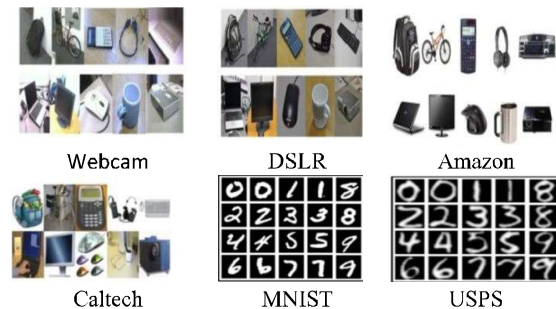
MNIST_vs_USPS(M_U) برای یک آزمایش دیگر ایجاد شده است.

به طور کلی کارایی الگوریتم پیشنهادی بر روی چهارده مجموعه تصاویر بین دامنه‌ای مختلف مورد ارزیابی قرار گرفته است. اطلاعات جزئی‌تر از پایگاه داده‌ها در جدول (۱) و تصاویر نمونه از هر دامنه در شکل (۳) نشان داده شده است.

(جدول-۱): معرفی پایگاه داده‌ها

(Table-1): The datasets

پایگاه داده	تعداد نمونه‌ها	تعداد ابعاد	تعداد کلاس‌ها	اختصار
USPS	1800	256	10	U
MNIST	2000	256	10	M
Amazon	958	800	10	A
Webcam	295	800	10	W
DSLR	157	800	10	D
Caltech256	1123	800	10	C



(شکل-۳): تصاویر نمونه از پایگاه داده‌های آفیس و

کالتک و اعداد

(Figure-3): Sample images from Office & Caltech and digits datasets

۵-۲- مفروضات پیاده‌سازی

برای به دست آوردن تنظیمات مدل، روش SRSA با مقادیر مختلف پارامترها مورد ارزیابی قرار گرفته است. طبقه‌بند پایه در تمام آزمایش‌ها، طبقه‌بند NN است. به علاوه، پیاده‌سازی روش پیشنهادی SRSA، توسط نرم‌افزار Matlab انجام گرفته است. دقت طبقه‌بندی بر روی نمونه‌های دامنه هدف برای مقایسه الگوریتم‌ها توسط رابطه زیر محاسبه می‌شود:

$$\text{Accuracy} = \frac{|\{x: x \in D_T \wedge f(x) = y(x)\}|}{n_T}$$

که D_T ، $f(x)$ و $y(x)$ به ترتیب نشان‌دهنده دامنه هدف، تابع پیش‌بینی، برچسب واقعی داده و تعداد نمونه‌های دامنه هدف هستند.

۵-۳- تأثیر پارامترها

برای به دست آوردن تنظیمات مدل، SRSA را با مقادیر مختلف پارامترها بررسی کرده‌ایم. دو پارامتر در این روش وجود دارد که مقادیر بهینه آن‌ها در جدول (۲) آورده شده است. پارامتر ρ ، پارامتر تنظیم تنگی می‌باشد که پارامتر تأثیرگذار در انتخاب نمونه‌های منبع است و تعداد نمونه‌های انتخاب شده را از دامنه منبع مشخص می‌کند. پارامتر K ، ابعاد زیرفضای جدید را مشخص می‌کند. پارامتر ρ در محدوده $[0.05, 0.1]$ و پارامتر K در محدوده $[10, 220]$ در نظر گرفته شده‌اند. روش SRSA با پارامترهای بیان شده بر روی چهارده مجموعه داده مورد آزمایش قرار می‌گیرد.

دقت طبقه‌بندی و تعداد نمونه‌های انتخاب شده با توجه به پارامتر ρ در جدول (۳) نشان داده شده است. هنگامی که $\rho = 0.05$ است ماتریس A تنگی کمتری دارد و در نتیجه نمونه‌های منبعی که توزیع نزدیک‌تری با نمونه‌های هدف دارند انتخاب می‌شوند.

نتایج به دست آمده بر روی پایگاه داده‌های آفیس و کالتک در شکل (۴) نشان‌دهنده این است که پایگاه داده آفیس و کالتک دارای رفتار متفاوتی نسبت به مقادیر مختلف پارامتر K هستند. برای به دست آوردن مقدار محدوده بهینه پارامتر، محدوده‌ای انتخاب می‌شود که دارای بهینه‌ترین دقت بر روی تعداد بیشتری از پایگاه داده‌ها باشد. برای بیشتر پایگاه داده‌ها محدوده بهینه پارامتر K در پایگاه داده آفیس و کالتک، $[40, 60]$ می‌باشد. نتایج به دست آمده برای مقادیر مختلف K ، در پایگاه داده اعداد نشان‌دهنده این است که پایگاه داده اعداد دارای حساسیت زیادی نسبت به مقادیر مختلف K و بهترین محدوده برای پارامتر K در پایگاه داده اعداد، نیز $[40, 60]$ است.

(جدول-۲): مقادیر بهینه پارامترها برای دو پایگاه داده بصری

(آفیس و کالتک، اعداد)، ρ پارامتر تنظیم تنگی، K ابعاد

زیرفضای جدید

(Table-2): Optimal parameters for two visual datasets (office & Caltech, digits), ρ is the sparsity parameter and K is the new subspace dimension

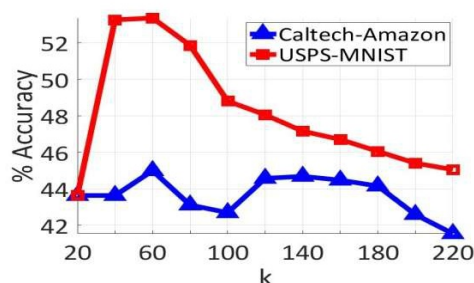
SRSA	l_1 norm	Orj	S-T
54.35	53.95	54.40	U-M
68.56	66.11	66.44	M-U
61.46	60.03	60.42	Avg

(جدول-۳): تأثیر پارامتر تنظیم تنگی ρ بر تعداد نمونه‌های

انتخاب‌شده از دامنه منبع در پایگاه داده اعداد

(Table-3): The effect of sparsity regularization on the number of selected samples from the source domain on the digit dataset

ρ	0.001	0.005	0.01	0.05	0.1	0.5
متوسط دقت U-M	54.40	54.40	54.40	54.40	54.40	54.35
U-M #	1800	1800	1800	1800	1800	1780
متوسط دقت M-U	66.44	66.44	66.44	66.44	66.44	68.56
M-U #	2000	2000	2000	2000	2000	1366



(شکل-۴): ارزیابی دقت پایگاه داده‌های USPS-MNIST

Caltech-Amazon با مقادیر مختلف پارامتر k

(Figure-4): USPS-MNIST, Caltech-Amazon datasets accuracy assessment with different values of k parameter

۴-۵- ضرورت هر پارامتر در الگوریتم

در این بخش، ضرورت دو پارامتر بر روی پایگاه داده اعداد بررسی شده و نتایج آن در جدول (۴) گزارش شده است.

(جدول-۴): ضرورت هر پارامتر در الگوریتم

S-T دامنه منبع و هدف، Orj: همه نمونه‌های منبع

(Table-4): The necessity of each parameter in the algorithm

پارامترها	K	ρ
آفیس و کالتک	50	0.5
اعداد	50	0.5

۱. استفاده از همه نمونه‌های منبع (Orj): انتخاب همه نمونه‌های منبع برای آموزش مدل به جای انتخاب برخی از آن‌ها.

۲. استفاده از نرم l_1 به جای نرم $l_{2,1}$: با استفاده از نرم l_1 به جای نرم $l_{2,1}$ تعداد ضرایب غیرصفر در ماتریس A افزایش یافته و به‌طور دقیق نمی‌توان هر نمونه هدف را با تعداد کمی از نمونه‌های منبع نمایش داد.

نتایج آزمایش‌ها حاکی از این است که استفاده از نمایش تنگ برای انتخاب نمونه‌ها تأثیر به‌سزایی بر روی

دقت مدل دارد. استفاده از همه نمونه‌های منبع کارایی مدل را کاهش می‌دهد، به‌طوری‌که انتخاب برخی از نمونه‌ها با استفاده از نمایش تنگ برای ساخت مدل پیش‌بینی، متوسط بهبود دقت را ۱/۰۴ افزایش می‌دهد.

همچنین، همان‌طور که در جدول نشان داده شده است، استفاده از نرم l_1 برای تعیین ضرایب ماتریس تنگ کارایی را به‌طور ملموس پایین می‌آورد. زمانی که از نرم $l_{2,1}$ استفاده شود، به دلیل این که نمونه‌های منبع مرتبط با هدف دقیقاً مشخص می‌شود، متوسط دقت مدل ۱/۴۳ افزایش می‌یابد.

۵-۵- مقایسه با الگوریتم‌های موجود

الگوریتم SRSA با هشت الگوریتم زیر مورد مقایسه قرار گرفته است: طبقه‌بندی نزدیک‌ترین همسایه (NN) [9] SA، PCA [3] DAM، [20] SDA، [11] GFK، [7] TCA [8] و SCA [10]. به دلیل اینکه تمامی روش‌های نام‌برده شده (به جز NN)، روش‌های کاهش بُعد هستند، از طبقه‌بندی استاندارد نزدیک‌ترین همسایه برای آموزش یک طبقه‌بند بر روی داده‌های منبع جهت پیش‌بینی داده‌های بدون برچسب دامنه هدف استفاده می‌شود.

دقت طبقه‌بندی SRSA و هشت روش دیگر بر روی پایگاه داده‌های آفیس و کالتک و اعداد در جداول (۵ و ۶) نشان داده شده است. همان‌طور که مشاهده می‌شود SRSA کارایی بهتری نسبت به هشت روش دیگر از خود نشان داده و در بیشتر حالات عملکرد مطلوبی ارائه می‌دهد. نتایج گزارش شده شامل میانگین برای همه روش‌ها است.

در پایگاه داده اعداد، SRSA دارای ۴/۹ متوسط بهبود دقت نسبت به بهترین الگوریتم مورد مقایسه SCA و ۶/۱۴ نسبت به الگوریتم استاندارد NN است. در پایگاه داده آفیس و کالتک، SRSA دارای ۱/۴۳ متوسط بهبود دقت نسبت به بهترین الگوریتم مورد مقایسه SCA و ۱۵/۲۲ نسبت به الگوریتم استاندارد NN است. در ادامه SRSA را با تمام روش‌ها به صورت جزئی مقایسه می‌کنیم.

PCA یک روش کاهش بُعد است که یک فضای مشترک کم بعد بین دامنه‌ها ایجاد می‌کند. PCA اختلاف توزیع بین دامنه‌ها را چندان کاهش نمی‌دهد و عملکرد ضعیفی نسبت به دیگر الگوریتم‌های تطبیق دامنه نشان می‌دهد و دارای دقت پایینی است. در پایگاه داده اعداد، SRSA دارای ۵/۸۷ متوسط بهبود دقت نسبت به PCA و در

¹ Nearest neighbor

پایگاه داده آفیس و کالتک، دارای ۶/۹۴ متوسط بهبود دقت و عمده ترین دلیل برتری SRSA نسبت به PCA استفاده از تطبیق زیرفضا برای کاهش اختلاف بین دامنه ها است. TCA یک روش تطبیق خصوصیات است که با استفاده از اجزای انتقال، یک نمایش مشترک بین دامنه های منبع و هدف ایجاد می کند که در نمایش جدید، اختلاف توزیع حاشیه ای بین دامنه های منبع و هدف کاهش می یابد. در این روش به اختلاف توزیع شرطی بین دامنه های منبع و هدف توجهی نشده و همچنین از داده های برچسب دار دامنه منبع در ایجاد فضای جدید استفاده نمی شود؛ درحالی که در روش SRSA، علاوه بر کاهش اختلاف توزیع حاشیه ای بین دامنه ها، از برچسب های نمونه های دامنه منبع برای ایجاد تفکیک پذیری بهتر بین رده ها استفاده شده است. به همین دلیل، عملکرد روش SRSA نسبت به روش TCA بهتر است. در پایگاه داده اعداد، SRSA دارای ۷/۷۹ متوسط بهبود دقت نسبت به TCA و در پایگاه داده آفیس و کالتک، دارای ۳/۵۶ متوسط بهبود دقت است.

DAM یک روش تطبیق دامنه چندمنبعی است که یک طبقه بند مقاوم برای پیش بینی برچسب نمونه های آزمایش با به کارگیری یک مجموعه از طبقه بندهای پایه از پیش یادگرفته شده ایجاد می کند. در پایگاه داده اعداد، SRSA دارای ۱۳/۷ متوسط بهبود دقت نسبت به DAM و در پایگاه داده آفیس و کالتک، دارای ۴/۰۴ متوسط بهبود دقت است.

در روش GFK، زیرفضاهایی از دامنه های منبع و هدف ایجاد می شود که در این زیرفضاها اختلاف توزیع

حاشیه ای بین دامنه های منبع و هدف به کمینه می رسد، اما به دلیل اینکه زیرفضاهای نگاشت شده دارای ابعاد بسیار کمی است، باعث از بین رفتن بخشی از داده شده و نمایش خوبی از داده ها ایجاد نمی شود. در پایگاه داده اعداد، SRSA دارای ۴/۶۲ متوسط بهبود دقت نسبت به GFK و در پایگاه داده آفیس و کالتک، دارای ۳/۶۴ متوسط بهبود دقت است. از جمله دلایل بهبود این است که در روش SRSA یک زیرفضا برای هر کدام از دامنه ها ایجاد می شود که ساختار داده ها را نیز حفظ می کند و همچنین با استفاده از نمایش تُنک فقط از برخی از نمونه های دامنه منبع برای ساخت طبقه بند استفاده می کند.

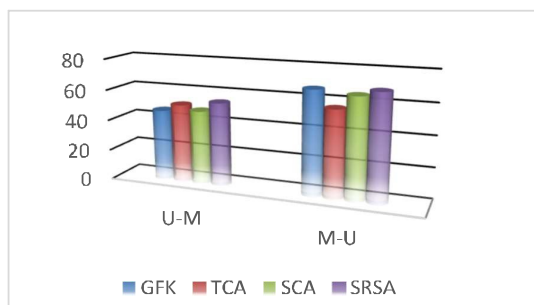
SA نیز یکی از روش های کاهش بعد در تطبیق بین دامنه ها است که زیرفضاهای به دست آمده را با یکدیگر تطبیق می دهد. در پایگاه داده اعداد، SRSA دارای ۸/۷۳ متوسط بهبود دقت نسبت به SA و در پایگاه داده آفیس و کالتک، دارای ۱/۸۷ متوسط بهبود دقت است. دلیل بهبود قابل ملاحظه SRSA نسبت به SA این است که SRSA از نمایش تُنک برای انتخاب نمونه های منبع استفاده می کند.

SDA یک روش تطبیق دامنه بدون نظارت است که شامل تطبیق توزیع به تطبیق زیرفضا است که علاوه بر کاهش اختلاف بین دو زیرفضا، اختلاف داده را نیز کم می کند. در پایگاه داده اعداد، SRSA دارای ۱۱/۲ متوسط بهبود دقت نسبت به SDA و در پایگاه داده آفیس و کالتک، دارای ۲/۰۷ متوسط بهبود دقت و عمده ترین دلیل برتری SRSA نسبت به SDA، علاوه بر استفاده از تطبیق زیرفضا، استفاده از نمایش تُنک برای انتخاب نمونه ها است.

(جدول ۵-): دقت (%) طبقه بند در پایگاه داده آفیس و کالتک

(Table-5): Classification accuracy on Office & Caltech dataset

SRSA	SCA (2016)	SDA (2015)	SA (2013)	GFK (2012)	DAM (2012)	TCA (2011)	PCA (2002)	NN	Dataset
44.57	45.62	49.69	49.27	41.02	42.69	45.82	36.95	23.7	C_A
39.32	40	38.98	40	40.68	34.58	30.51	32.54	25.76	C_W
52.87	47.13	40.13	39.49	38.85	33.12	35.67	38.22	25.48	C_D
38.91	39.72	39.54	39.98	40.25	35.35	40.07	34.73	26	A_C
40.68	34.92	30.85	32.22	38.98	31.86	35.25	35.59	29.83	A_W
37.58	39.49	33.76	33.76	36.31	36.31	34.39	27.39	25.48	A_D
30.37	31.08	34.73	35.17	30.72	33.84	29.92	26.36	19.86	W_C
32.15	29.96	39.25	39.25	29.75	37.58	28.81	29.35	22.96	W_A
90.45	87.26	75.80	75.16	80.89	80.89	85.99	77.07	59.24	W_D
32.15	30.72	35.89	34.55	30.28	32.41	32.06	29.65	26.27	D_C
31.21	31.63	38.73	39.87	32.05	34.34	31.42	32.05	28.5	D_A
88.81	84.41	76.95	76.95	75.59	77.63	86.44	75.93	63.39	D_W
46.59	45.16	44.52	44.72	42.95	42.55	43.03	39.65	31.37	متوسط دقت

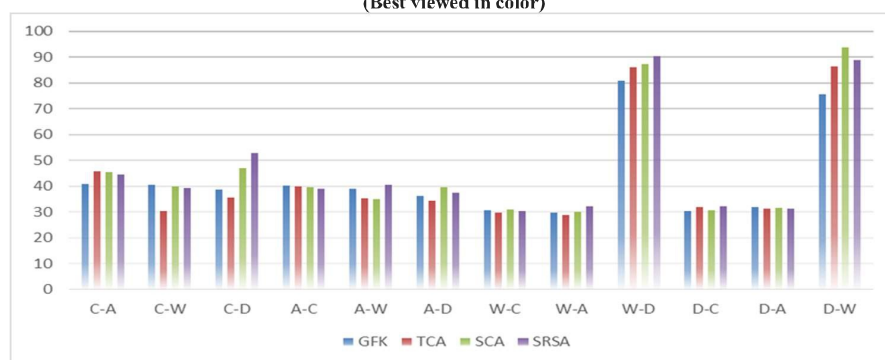


(شکل-۵): مقایسه دقت طبقه‌بندی بر روی پایگاه داده اعداد با استفاده از روش‌های GFK، TCA، SCA و SRSA (نمایش بهتر به صورت رنگی)

(Figure-5): The comparison of the classification accuracy on Digit dataset using GFK, TCA, SCA and SRSA methods (Best viewed in color)

(شکل-۶): مقایسه دقت طبقه‌بندی بر روی پایگاه داده آفیس و کالتک با استفاده از روش‌های GFK، TCA، SCA و SRSA (نمایش بهتر به صورت رنگی)

(Figure-6): The comparison of the classification accuracy on Office & Caltech dataset using GFK, TCA, SCA, and SRSA (Best viewed in color)



(جدول-۶): دقت (%) طبقه‌بندی در پایگاه داده اعداد

(Table-6): The classification accuracy on Digit dataset

SRSA	SCA (2016)	SDA (2015)	SA (2013)	GFK (2012)	DAM (2012)	TCA (2011)	PCA (2002)	NN	Dataset
54.35	48	35.70	41.50	46.45	42.69	51.05	44.95	44.7	U_M
68.56	65.11	65	63.95	67.22	52.83	56.28	66.22	65.94	M_U
61.46	56.56	50.35	52.73	56.84	47.76	53.67	55.59	55.32	متوسط دقت

نقش نمونه‌هایی از داده‌های منبع را که عامل ایجاد اختلاف بین دو دامنه هستند، کم‌رنگ‌تر می‌کند؛ علاوه بر این، SRSA با تطبیق زیرفضاهای ایجادشده اختلاف توزیع بین زیرفضاها را به کمینه می‌رساند. برای ارزیابی عملکرد SRSA، آزمایش‌های مختلفی بر روی انواع پایگاه داده‌های واقعی ترتیب داده شده و نتایج ارزیابی‌ها نشان از برتری قابل‌ملاحظه SRSA در مقایسه با روش‌های تطبیق دامنه بدون نظارت در این حوزه است. برای ادامه کار از نظریه منیفلدها^۱ جهت برجسته‌سازی بهتر نمونه‌های هدف بهره خواهیم برد.

¹ Manifolds

۶- نتیجه‌گیری

در این مقاله یک روش بدون نظارت جهت حل مسأله شیفیت در دامنه‌ها، با عنوان طبقه‌بندی تصاویر با استفاده از نمایش تنگ و تطبیق زیرفضا برای تطبیق دامنه‌های بصری ارائه شد. SRSA یک روش نوین است که در دو مرحله تلاش می‌کند اختلاف توزیع بین داده‌های آموزشی و آزمایش را کاهش دهد. درواقع SRSA مدلی ایجاد می‌کند که با به‌کمینه‌رساندن ناسازگاری داده‌ها، مقاومت و تحمل‌پذیری بیشتری در برابر تغییرات توزیع داده‌ها بین دامنه‌های منبع و هدف داشته باشد. SRSA در گام نخست با استفاده از نرم‌افزار

- [15] X. Li, M. Fang, J. J. Zhang and J. Wu, "Sample selection for visual domain adaptation via sparse coding", *Signal Processing: Image Communication*, vol. 44, pp. 92-100, 2016.
- [16] K. Saenko, B. Kulis, M. Fritz and T. Darrell, "Adapting visual category models to new domains", *Proceedings of the European Conference on Computer Vision*, pp. 213-226, 2010.
- [17] G. Griffin, A. Holub and P. Perona, "Caltech-256 object category dataset", Technical Report 7694, 2007.
- [18] J. J. Hull, "A", *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 16, no. 5, pp. 550-554, 1994.
- [19] Y. LeCun, L. Bottou, Y. Bengio, P. Haffner, "Gradient-based learning applied to document recognition", *Proc. IEEE*, vol. 86, no. 11, pp. 2278-2324, 1998.
- [20] L. Duan, D. Xu, I.W. Tsang, "Domain adaptation from multiple sources: a domain-dependent regularization approach", *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 23, no. 3, pp. 504-518, 2012.
- [1] Y. Zhang, T. Liu, M. Long, and M. I. Jordan, "Bridging Theory and Algorithm for Domain Adaptation", arXiv preprint arXiv:1904.05801, 2019.
- [2] S. Rezaei and J. Tahmoresnezhad, "Discriminative and domain invariant subspace alignment for visual tasks", *Iran Journal of Computer Science*, pp. 1-12, 2019.
- [3] W. Kumagai and T. Kanamori, "Risk bound of transfer learning using parametric feature mapping and its application to sparse coding", *Machine Learning*, pp. 1-34, 2019.
- [4] I. Jolliffe "Principal component analysis", Wiley, vol. 2, pp. 433-459, 2002.
- [5] M. Jing, J. Li, J. Zhao, and K. Lu, "Learning distribution-matched landmarks for unsupervised domain adaptation", *In International conference on database systems for advanced applications*, pp. 491-508, 2018.
- [6] A. Li, D. Chen, Z. Wu, G. Sun, and K. Lin, "Self-supervised sparse coding scheme for image classification based on low rank representation", *PloS one*, Vol. 13(6), e0199141, 2018.
- [7] A. Mirjalili, V. Abootalebi, M. T. Sadeghi, "improving the performance of sparse representation-based classifier for EEG classification," *JSDP*, pp. 43-55, 2015.
- [8] X. Li, M. Fang, J. J. Zhang and J. Wu, "Sample selection for visual domain adaptation via sparse coding", *Signal Processing: Image Communication*, vol. 44, pp. 92-100, 2016.
- [9] B. Gong, Y. Shi, F. Sha and K. Grauman, "Geodesic flow kernel for unsupervised domain adaptation", *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 2066-2073, 2012.
- [10] S. J. Pan, I. W. Tsang, J. T. Kwok and Q. Yang, "Domain adaptation via transfer component analysis", *IEEE Trans. Neural Netw.*, vol. 22, no. 2, pp. 199-210, 2011.
- [11] B. Fernando, A. Habrard, M. Sebban, and T. Tuytelaars, "Unsupervised visual domain adaptation using subspace alignment", in *Proc. IEEE International Conference on Computer vision*, pp. 2960-2967, 2013.
- [12] M. Ghifary, D. Balduzzi, W. B. Kleijn, and M. Zhang, "Scatter component analysis: A unified framework for domain adaptation and domain generalization", *IEEE Transactions on Pattern Analysis and Machine Intelligence*, pp. 1-1, 2016.
- [13] B. Sun and K. Saenko, "Subspace distribution alignment for unsupervised domain adaptation", in *Proc. British Machine Vision Conference*, 2015.
- [14] J. Tahmoresnezhad and S. Hashemi, "Visual domain adaptation via transfer feature learning", *Knowledge and Information Systems*, vol. 50, no. 2, pp. 585-605, 2017.

فریماه شرافتی مدرک کارشناسی خود



را در رشته مهندسی فناوری اطلاعات از دانشگاه ارسنجان دریافت کردند. ایشان با ادامه تحصیل موفق به اخذ مدرک کارشناسی ارشد در رشته مهندسی فناوری اطلاعات از دانشگاه صنعتی ارومیه شده‌اند. علایق پژوهشی ایشان شامل حوزه‌های یادگیری ماشین، یادگیری انتقالی است.

نشانی رایانامه ایشان عبارت است از:

farimah.sherafati@it.uut.ac.ir

جعفر طهمورث‌نژاد مدرک دکترای



مهندسی کامپیوتر-هوش مصنوعی خود را در سال ۱۳۹۴ از دانشگاه شیراز دریافت کردند. ایشان در راستای فعالیتهای علمی خود، در حال حاضر به‌عنوان استادیار دانشگاه صنعتی ارومیه در حال فعالیت هستند. علایق پژوهشی ایشان شامل حوزه‌های یادگیری ماشین، یادگیری انتقالی، داده‌کاوی و سیستم‌های امنیتی است.

نشانی رایانامه ایشان عبارت است از:

j.tahmores@it.uut.ac.ir