

پردازش تصویر با استفاده از کدگذاری تنک و طبقه‌بندی انطباقی

فریمه شرافتی^۱، دانشجوی کارشناسی ارشد؛ جعفر طهمورث نژاد^۲، استادیار

۱- دانشکده مهندسی فناوری اطلاعات و کامپیوتر- دانشگاه صنعتی ارومیه- ارومیه- ایران- farimah.sherafati@it.uut.ac.ir

۲- دانشکده مهندسی فناوری اطلاعات و کامپیوتر- دانشگاه صنعتی ارومیه- ارومیه- ایران- j.tahmores@it.uut.ac.ir

چکیده: به دلیل افزایش حجم تصاویر تولیدشده توسط دوربین‌ها و دستگاه‌های مختلف، پردازش تصویر در بسیاری از کاربردها از جمله پزشکی، امنیتی و رانندگی اهمیت و جایگاه بالایی یافته است. با این حال بیشتر مدل‌های ایجادشده در حوزه پردازش تصویر کارایی چندانی نداشته و میزان خطای آن‌ها در برخی کاربردها تأثیرگذار است. علت اصلی ناکامی بیشتر مدل‌های ساخته‌شده، اختلاف توزیع بین داده‌های آموزشی (دامنه منبع) و داده‌های تست (دامنه هدف) می‌باشد. درواقع، مدل ساخته‌شده، قابلیت تعمیم‌دهی به داده‌هایی با خصوصیات و توزیع‌های متفاوت از داده‌های آموزشی را ندارد، به همین دلیل در مواجهه با داده‌های جدید دچار افت شدیدی می‌شود. در این مقاله ما یک روش جدید با نام کدگذاری تنک و طبقه‌بندی انطباقی (SADA) پیشنهاد می‌دهیم که یک مدل پردازش تصویری ایجاد می‌کند که در مقابل تغییرات داده‌ای مقاوم می‌باشد. مدل پیشنهادی با ایجاد یک زیر فضای مشترک بین دامنه‌های منبع و هدف اختلاف توزیع آن‌ها را به حداقل رسانده و موجب بهبود کارایی می‌شود. همچنین SADA با انتخاب نمونه‌هایی از دامنه منبع که با دامنه هدف مرتبط می‌باشند اختلاف توزیع بین دامنه‌ها را کاهش می‌دهد. علاوه بر آن، SADA با تطبیق پارامترهای مدل ایجادشده، یک مدل تطبیق‌پذیر برای مواجهه با شیفت داده‌ها ایجاد می‌کند. نتایج به دست آمده از آزمایش‌های متنوع، نشان می‌دهد که روش پیشنهادی ما، برتری قابل ملاحظه‌ای نسبت به تمام روش‌های تطبیق دامنه جدید دارد.

واژه‌های کلیدی: پردازش تصویر، تطبیق دامنه‌های بصری، کدگذاری تنک، وزن‌دهی مجدد نمونه، طبقه‌بندی انطباقی.

Image Processing via Sparse Coding and Adaptive Classification

F. Sherafati¹, MSc Student; J. Tahmoresnezhad², Assistant Professor

1- Faculty of IT & Computer Engineering, Urmia University of Technology, Urmia, Iran, farimah.sherafati@it.uut.ac.ir

2- Faculty of IT & Computer Engineering, Urmia University of Technology, Urmia, Iran, j.tahmores@it.uut.ac.ir

Abstract: Due to the growing increase of generated images via cameras and various instruments, image processing has found an important role in most of practical usages including medical, security and driving. However, most of the available models has no considerable performance and in some usages the amount of error is very effective. The main cause of this failure in most of available models is the distribution mismatch across the source and target domains. In fact, the made model has no generalization to test data with different properties and distribution compared to the source data, and its performance degrades dramatically to face with new data. In this paper, we propose a novel approach entitled Sparse coding and ADaptive classification (SADA) which is robust against data drift across domains. The proposed model reduces the distribution difference across domains via generating a common subspace between the source and target domains and increases the performance of model. Also, SADA reduces the distribution mismatch across domains via the selection of the source samples which are related to target samples. Moreover, SADA adapts the model parameters to build an adaptive model to encounter with data drift. Our variety of experiments demonstrate that the proposed approach outperforms all stat-of-the-art domain adaptation methods.

Keywords: Image processing, visual domains adaptation, sparse coding, sample reweighting, adaptive classification.

تاریخ ارسال مقاله: ۱۳۹۶/۰۷/۰۵

تاریخ اصلاح مقاله: ۱۳۹۶/۰۹/۰۵

تاریخ پذیرش مقاله: ۱۳۹۶/۱۰/۱۲

نام نویسنده مسئول: جعفر طهمورث نژاد

نشانی نویسنده مسئول: ایران - ارومیه - دانشگاه صنعتی ارومیه - دانشکده مهندسی فناوری اطلاعات و کامپیوتر.

۱- مقدمه

یادگیری ماشین، از شاخه‌های پرکاربرد در هوش مصنوعی می‌باشد که با ایجاد الگوریتم و برنامه‌های خودکار به دنبال هوشمندسازی ماشین‌ها و مدل‌ها می‌باشد. در الگوریتم‌های طبقه‌بندی یادگیری ماشین، یک مدل با استفاده از داده‌های برچسب‌دار موجود آموزش داده شده و بر روی داده‌های تست، جهت برچسب‌گذاری نمونه‌های جدید اعمال می‌شود. یک فرض مشترک در تمام الگوریتم‌های طبقه‌بندی وجود دارد، که توزیع و فضای خصیصه‌ای نمونه‌های آموزشی و تست را یکسان در نظر می‌گیرد. اما در کاربردهای دنیای واقعی، این فرضیه همیشه برقرار نیست. به عبارت دیگر، داده‌های آموزشی و تست ممکن است از دامنه‌های کاملاً متفاوتی باشند. در این حالت، مدل ساخته شده روی داده‌های آموزشی، کارایی پایینی روی دامنه تست خواهد داشت. اختلاف توزیع بین نمونه‌های آموزشی و تست از آنجا نشأت می‌گیرد که در برخی سیستم‌ها به تعداد کافی داده آموزشی برچسب‌دار، جهت آموزش مدل وجود ندارد، در این حالت از دامنه‌های برچسب‌دار مرتبط، جهت ایجاد یک مدل، برای برچسب‌گذاری داده‌های تست استفاده می‌شود [۱، ۲]. با این اوصاف داده‌های تست برچسب‌گذاری می‌شوند ولی در بررسی‌های اولیه مشخص می‌شود که اغلب برچسب‌های نسبت داده شده دقیق نبوده و مدل از خطای بالایی رنج می‌برد. به چنین مشکلی، شیفت دامنه‌ها^۱ گفته می‌شود که در آن دامنه‌های آموزشی و تست از توزیع متفاوتی برخوردارند.

یادگیری انتقالی^۲ یا تطبیق دامنه‌ها^۳ یک بهبود یادگیری در یک حوزه جدید از طریق انتقال دانش از یک حوزه مرتبط می‌باشد که قبلاً یاد گرفته شده است. درواقع روش‌های یادگیری انتقالی، یک بهبود بر روی روش‌های یادگیری ماشین هستند، که آن‌ها را جهت مواجهه با مسائلی که در آن‌ها شیفت دامنه‌ها و اختلاف توزیع وجود دارد مقاوم می‌سازند. یادگیری انتقالی بر اساس اطلاعات برچسب دامنه هدف، به دو دسته تقسیم می‌شود: ۱- یادگیری انتقالی نیمه‌نظارت شده که در آن تمام نمونه‌های دامنه منبع دارای برچسب بوده، و تعداد کمی از نمونه‌های دامنه هدف دارای برچسب می‌باشند، ۲- یادگیری انتقالی بدون نظارت که در آن تمام نمونه‌های دامنه هدف بدون برچسب هستند.

وجود اختلاف توزیع بین دامنه‌ها موجب می‌شود مدلی که روی داده‌های آموزشی ایجاد شده است، عملکرد خوبی روی داده‌های تست نداشته باشد. برای سنجش اختلاف توزیع حاشیه‌ای بین دو دامنه از یک معیار غیرپارامتری به نام MMD^۴ استفاده می‌شود. MMD داده‌های منبع و هدف را به یک فضای RKHS^۵ نگاشت می‌کند و با محاسبه میانگین عناصر^۶ هر کدام از دامنه‌ها در فضای مزبور، اختلاف دو دامنه را به دست می‌آورد. اثبات شده است [۳]، میزان فاصله میانگین عناصر دو دامنه در فضای RKHS معادل فاصله دو دامنه در فضای اصلی است. بدین ترتیب در صورتی که فاصله میانگین عناصر داده‌های منبع و هدف

در فضای RKHS کاهش یابد، داده‌ها با یکدیگر تطبیق شده و مدل قوی‌تر و تواناتری ایجاد می‌شود.

کدگذاری تنک^۷ از جمله روش‌های بدون نظارت^۸ است که برای نمایش داده‌ها به صورت کارا مورد استفاده قرار می‌گیرد [۴، ۵]. هدف کدگذاری تنک، یافتن یک مجموعه از بردارهای اصلی است که یک بردار ورودی را بتوان بر اساس یک ترکیب خطی از بردارهای اصلی نمایش داد. مزیت استفاده از بردارهای اصلی در کدگذاری تنک این است که بردارهای اصلی به خوبی توانایی حفظ ساختارها و الگوهای داده ورودی را دارند، بدین ترتیب پردازش بر روی داده کدگذاری شده معادل پردازش بر روی داده‌های ورودی است.

از دیگر روش‌های مقابله با مسئله شیفت دامنه‌ها استفاده از طبقه‌بند انطباقی می‌باشد که با افزایش تطبیق‌پذیری بین دامنه‌های منبع و هدف خطای پیش‌بینی را به حداقل می‌رساند. درواقع طبقه‌بند انطباقی، ابعاد تفکیک‌کننده کلاس‌ها بین دامنه‌های منبع و هدف را با یکدیگر سازگار می‌نماید؛ به عبارت دیگر با بهره‌گیری از توزیع هندسی داده‌ها در ایجاد طبقه‌بند انطباقی، مدل طبقه‌بند ایجاد شده بر روی نمونه‌های دامنه منبع با ساختار هندسی دامنه هدف، سازگاری بیشتری خواهد یافت.

روش پیشنهادی در این مقاله با عنوان SADA^۹، یک روش نیمه‌نظارت شده با بهره‌گیری از کدگذاری تنک، تطبیق خصوصیات و طبقه‌بند انطباقی^{۱۰} است. SADA با استفاده از ماتریس تنک و به کارگیری نرم $l_{2,1}$ نمونه‌هایی از دامنه منبع که مرتبط و نزدیک به دامنه هدف می‌باشند را انتخاب می‌نماید. سپس با استفاده از تطبیق خصوصیات یک زیر فضای جدید به دست می‌آورد که در آن خصوصیات مشترک دو دامنه منبع و هدف نگهداری می‌شود به طوری که اختلاف توزیع حاشیه‌ای دامنه‌ها به حداقل رسانیده می‌شود. در انتها SADA از یک طبقه‌بند انطباقی با هدف حداقل کردن خطای پیش‌بینی مدل روی داده‌های منبع و حداکثر کردن سازگاری بین ساختار منیفلد و مدل پیش‌بینی استفاده می‌نماید.

ادامه این مقاله به صورت زیر سازماندهی شده است. مروری بر کارهای گذشته در بخش دوم مقاله آورده شده است. تعاریف اولیه و مفاهیم مقدماتی در بخش سوم قرار داده شده است. روش پیشنهادی در بخش چهارم معرفی شده است. پایگاه داده‌های مورد آزمایش و نتایج ارزیابی عملکرد روش پیشنهادی در مقایسه با دیگر روش‌های یادگیری ماشین و یادگیری انتقالی در بخش پنجم گزارش شده است. در انتها، نتیجه‌گیری و کارهای پیشنهادی آتی گنجانده شده است.

۲- کارهای پیشین

برای حل مشکل اختلاف توزیع بین دامنه‌ها، روش‌های تطبیق دامنه زیادی پیشنهاد شده است. به طور کلی کارهایی که در حوزه یادگیری انتقالی انجام می‌شوند به سه دسته تقسیم می‌شوند: رویکردهای مبتنی

دارند. ^{۱۹} ARTL [۱۱] یک روش تطبیق طبقه‌بند است که به ایجاد یک مدل انطباقی از طریق ایجاد تطبیق توزیع مشترک، حداقل کردن خطای پیش‌بینی و حداکثر کردن سازگاری هندسی بین دامنه‌های منبع و هدف به‌طور هم‌زمان تطبیق دامنه‌ها را انجام می‌دهد. در این مقاله، مسئله شیفت دامنه‌ها با استفاده از یک راه‌حل نیمه‌نظارت‌شده مورد حل قرار می‌گیرد. درواقع، با ایجاد یک نمایش جدید مبتنی بر خصوصیت به کمک نمایش تنک، فاصله دامنه‌های منبع و هدف کاهش یافته و با افزودن یک طبقه‌بند انطباقی پارامترهای مدل برای مواجهه با اختلاف توزیع‌های بالا قدرتمندتر و مقاوم‌تر می‌شوند.

۳- تعاریف اولیه و مفاهیم مقدماتی

در این بخش، ابتدا تعریف مسئله ارائه شده و سپس مفاهیم مقدماتی برای ارائه روش پیشنهادی گنجانده شده است.

۳-۱- هدف تحقیق

در حوزه یادگیری ماشین، تطبیق دامنه‌ها به یک ضرورت اجتناب‌ناپذیر تبدیل شده است، زیرا دیگر، الگوریتم‌های کلاسیک یادگیری ماشین با سخگویی مسائل موجود در دنیای واقعی نیستند. در این مقاله، ما به دنبال پیشنهاد روشی هستیم که بتواند بر محدودیت‌های الگوریتم‌های استاندارد یادگیری ماشین غلبه کرده و با انتقال دانش از یک دامنه به دامنه دیگر، بازدهی روش‌های موجود را افزایش دهد. به‌طور کلی ما در این مقاله به دنبال برآورده کردن اهداف ذیل هستیم: (۱) انتقال دانش از یک دامنه به دامنه دیگر، (۲) کاهش اختلاف توزیع داده‌های آموزشی و تست، و (۳) مقاوم کردن مدل‌ها در برابر تغییرات احتمالی داده‌ها. در ادامه تعریف مسئله به‌صورت ریاضی بیان شده است.

۳-۲- تعریف مسئله

دامنه ^{۲۰} - هر دامنه D دارای دو عنصر فضای خصیصه‌ای X و توزیع احتمال حاشیه‌ای $P(X)$ می‌باشد؛ بنابراین دامنه D به‌صورت زیر نشان داده می‌شود: $D = \{X, P(X)\}$. بدین ترتیب، دامنه‌های منبع و هدف زمانی با یکدیگر تفاوت دارند که هر کدام یا از فضای خصیصه‌ای مختلفی برخوردار بوده و یا دارای توزیع احتمال حاشیه‌ای متفاوتی باشند.

وظیفه ^{۲۱} - برای هر دامنه D یک وظیفه T وجود دارد که از دو مجموعه Y برچسب و تابع پیش‌بینی $f(x)$ تشکیل می‌شود؛ بنابراین وظیفه T به‌صورت زیر نشان داده می‌شود: $T = \{Y, f(x)\}$. مجموعه برچسب Y برای نمونه‌های ورودی x توسط تابع $f(x)$ پیش‌بینی می‌شود و احتمال شرطی آن به‌صورت $P(Y|x)$ تعریف می‌شود. بدین ترتیب، زمانی دو وظیفه متفاوت هستند که یا مجموعه برچسب و یا تابع توزیع احتمال شرطی متفاوتی داشته باشند.

بر نمونه [۴]، رویکردهای مبتنی بر خصوصیت [۶] و رویکردهای مبتنی بر تطبیق طبقه‌بند [۷].

در رویکردهای مبتنی بر نمونه، نمونه‌های منبع به‌گونه‌ای وزن‌دهی مجدد می‌شوند تا توزیع نزدیک‌تری با نمونه‌های دامنه هدف داشته باشند. یکی از جدیدترین روش‌ها در این رویکرد، روش لندمارک [۸] می‌باشد. در روش لندمارک، آن نمونه‌هایی از دامنه منبع انتخاب می‌شوند که از لحاظ توزیع دارای بیشترین شباهت با نمونه‌های دامنه هدف باشند. برای به‌دست‌آوردن اختلاف توزیع بین دو دامنه از MMD استفاده می‌شود و به نمونه‌هایی از دامنه منبع که اختلاف توزیع کمتری با نمونه‌های دامنه هدف دارند، وزن بیشتری تعلق می‌گیرد. روش ^{۱۲} TJM [۹] یک روش ترکیبی از روش‌های مبتنی بر نمونه و مبتنی بر خصوصیت است که برای مسائل با اختلاف توزیع بالا به‌کار گرفته می‌شود. TJM ابتدا اختلاف توزیع حاشیه‌ای بین دامنه‌ها را کاهش داده، سپس نمونه‌ها را وزن‌دهی مجدد می‌کند و نمونه‌های منبع که توزیع نزدیک‌تری با نمونه‌های هدف دارند را انتخاب می‌کند. TJM، دارای یک کرنل نسبتاً پیچیده می‌باشد و حل تابع هدف آن دشوار می‌باشد. همچنین TJM به دلیل حذف بخشی از داده ورودی، ساختار داده‌ها را حفظ نمی‌کند. روش ^{۱۳} SSSC [۴]، ابتدا نمونه‌های غیرمرتبط دامنه منبع را حذف کرده، سپس با یافتن یک ماتریس انتخاب، اختلاف توزیع بین داده‌های منبع و هدف را کاهش می‌دهد. SSSC از کدگذاری تنک برای انتخاب نمونه‌های دامنه منبع استفاده کرده و هر نمونه دامنه هدف را با حداقل تعداد نمونه دامنه منبع نمایش می‌دهد.

در رویکردهای مبتنی بر خصوصیت، هدف یادگیری یک نمایش جدید از ویژگی‌ها است. ^{۱۴} GFK [۶] یک روش مبتنی بر خصوصیت است که داده‌های منبع و هدف را بر روی منیفولد گرسمن ^{۱۵} نگاشت می‌کند تا اختلاف توزیع حاشیه‌ای آن‌ها کاهش یابد. GFK بین زیرفضاهای منبع و هدف یک مسیر کوتاه ^{۱۶} بر روی منیفولد ایجاد می‌کند تا تغییرات دامنه از منبع تا هدف را مورد ارزیابی قرار دهد. از آنجایی که حرکت بر روی مسیر کوتاه بایستی به‌صورت روان و پیوسته باشد GFK اندازه زیرفضا را کوچک در نظر می‌گیرد که موجب از بین رفتن بخشی از داده‌های ورودی می‌شود. ^{۱۷} VDA [۱۰] برای کاهش اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌ها، یک نمایش کم بعد از دامنه‌های منبع و هدف ایجاد می‌کند. همچنین، VDA از خوشه‌بندی مستقل از دامنه ^{۱۸} جهت حفظ ساختار هندسی و آماری بین دامنه‌ها در فضای جدید استفاده می‌کند.

رویکردهای مبتنی بر تطبیق طبقه‌بند، با ایجاد طبقه‌بند‌های انطباقی، مدل‌های مقاوم‌تری در مقابل اختلاف توزیع بین دامنه‌های منبع و هدف ایجاد می‌کنند. درواقع طبقه‌بند انطباقی بدون ایجاد تغییر در فضای ویژگی، پارامترهای مدل طبقه‌بند ایجاد شده بر روی داده‌های دامنه منبع را با دامنه هدف سازگار می‌سازد. با این حال، عمده روش‌های مبتنی بر مدل بر روی حل مسائل تطبیق دامنه نیمه‌نظارت شده تمرکز

۳-۴- کاهش بعد

در روش‌های کاهش بعد، داده‌ها از فضای اصلی به یک فضای نگاشت‌شده منتقل می‌شوند به‌شرطی که هزینه بازگرداندن آن‌ها به فضای اصلی^{۲۴} حداقل باشد. PCA^{۲۵} یکی از مرسوم‌ترین روش‌ها برای انتقال داده‌ها به یک فضای کم بعد می‌باشد، به‌طوری‌که ساختار اصلی داده‌ها را در فضای جدید حفظ می‌کند. درواقع PCA داده‌ها را بر روی اجزای اصلی خود نگاشت کرده به‌طوری‌که داده‌ها بر روی این اجزا دارای حداکثر واریانس باشند. اجزای اصلی به‌دست‌آمده، جهت نگاشت داده‌ها به فضای کم‌بعد مورد استفاده قرار می‌گیرند. اگر p ابعاد زیرفضای جدید و $U_s \in R^{d_s \times p}$ و $U_t \in R^{d_t \times p}$ ماتریس‌های نگاشت به فضای کم‌بعد باشند، رابطه ۲ به‌صورت زیر بازنویسی می‌شود:

$$\min_{A, U_t, U_s} \|U_t^T X_t - U_s^T X_s A\|_F^2 + \rho \|A\|_{2,1} \quad (3)$$

با استفاده از رابطه ۳ یک ماتریس انتخاب A و دو ماتریس نگاشت U_t و U_s به‌دست می‌آید که فاصله بین داده‌های منبع و هدف را در زیرفضای جدید کاهش می‌دهد. همچنین PCA برای حفظ واریانس داده‌های منبع و هدف از شرط زیر در کنار رابطه ۳ استفاده می‌کند که در آن V ماتریس هم‌مرکز نمودن است که برای هر دو دامنه با استفاده از رابطه $V = I - \frac{1}{n_s} 11^T$ به‌دست می‌آید و از پراکندگی داده‌ها در زیرفضای تعبیه شده جلوگیری می‌کند که در آن I ماتریس همانی و $\vec{1}$ بردار یک‌ها است.

$$\max_U \text{tr}(U^T X V X^T U) \quad (4)$$

که $U = [U_s, U_t]$ ماتریس نگاشت منبع و هدف و $X = [X_s, X_t]$ مجموعه داده‌های منبع و هدف می‌باشد.

۳-۵- حفظ توپولوژی

به اثبات رسیده است که حفظ ساختار و توپولوژی داده‌ها در فضای نگاشت شده می‌تواند کارایی مدل را به صورت قابل توجهی افزایش دهد [۱۰]. به‌منظور حفظ ساختار توپولوژی هر دامنه، از فرضیات منیفلد استفاده می‌شود. به‌همین دلیل در نگاشت داده‌ها با استفاده از ماتریس U به فضای جدید توپولوژی و ساختار داده‌های نگاشت شده با استفاده از فرضیه منیفلد حفظ می‌شود. مطابق با فرضیه منیفلد، اگر دو داده X_i و X_j در فضای اصلی در همسایگی هم قرار داشته باشند، بایستی در زیرفضای جدید نیز همسایگی آن‌ها حفظ شود. با استفاده از رابطه $S(X_i, X_j) = e^{-\frac{\|X_i - X_j\|_F^2}{\delta^2}}$ که در آن δ میانگین فاصله‌های داده‌ها است، می‌توان میزان همسایگی هر دو نمونه را به‌دست آورد. از روی ماتریس همسایگی ماتریس لاپلاسی به‌دست می‌آید که فاصله هر نمونه از تمام نمونه‌ها به‌جز خودش را محاسبه می‌کند. بدین ترتیب با حداقل کردن رابطه زیر، می‌توان نگاشت U را به‌گونه‌ای به‌دست آورد

اگر $P_s(x)$ و $P_t(x)$ به‌ترتیب نشان‌دهنده توزیع احتمال حاشیه‌ای داده‌های منبع و هدف باشند، در صورت بروز شیفت دامنه‌ها خواهیم داشت: $P_s(x) \neq P_t(x)$. بدین ترتیب، این مقاله به‌دنبال یک نمایش جدید و کارا برای داده‌های منبع و هدف است که پس از نگاشت نمونه‌ها به زیر فضای جدید از توزیع مشابهی برخوردار باشند، یعنی: $P_s(x) \approx P_t(x)$. در ادامه متدولوژی‌های مورد نیاز جهت ایجاد زیرفضای جدید به تفکیک معرفی شده‌اند.

۳-۳- کدگذاری تنک

در تطبیق دامنه نیمه‌نظارت‌شده، تمام نمونه‌های دامنه منبع دارای برچسب بوده اما از بین نمونه‌های دامنه هدف فقط تعداد کمی از آن‌ها برچسب داشته که به‌تنهایی قادر به ایجاد یک مدل جهت برچسب‌گذاری نمونه‌های هدف نمی‌باشند. اگر $X_s \in R^{d_s \times n_s}$ و $X_t \in R^{d_t \times n_t}$ که n_t و n_s به‌ترتیب تعداد نمونه‌های منبع و هدف و $d_s = d_t$ ابعاد فضای خصوصیات باشند، هدف، یافتن نمونه‌های منبع مرتبط با نمونه‌های هدف است تا بتوان طبقه‌بندی ساخت که به‌کمک آن به نمونه‌های بدون برچسب دامنه هدف برچسب مناسب نسبت داد. برای انتخاب نمونه‌ها از دامنه منبع از کدگذاری تنک استفاده می‌شود. با استفاده از نرم ۱ می‌توان هر نمونه هدف را با حداقل تعداد نمونه منبع نمایش داد. رابطه زیر ماتریس انتخاب $A \in R^{n_s \times n_t}$ را به‌دست می‌آورد که اختلاف توزیع بین داده‌های دامنه منبع و هدف را حداقل می‌سازد.

$$\min_A \|X_t - X_s A\|_F^2 + \rho \|A\|_1 \quad (1)$$

که $\|.\|_F$ نرم فروبنیوس^{۲۲} می‌باشد و ماتریس A یک ماتریس تنک است، که با استفاده از مسئله بهینه‌سازی فوق به‌دست می‌آید و با وزندهی مجدد به نمونه‌های منبع، فاصله بین دامنه‌های منبع و هدف را کاهش می‌دهد. ρ پارامتر تنظیم پراکندگی^{۲۳} است. با به‌کارگیری نرم ۱ تعداد ضرایب غیر صفر حداقل می‌شود ولی نرم ۱ منجر به انتخاب نمونه نمی‌شود زیرا دقیقاً مشخص نیست هر نمونه منبع دقیقاً با کدام یک از نمونه‌های هدف مرتبط است. با استفاده از رابطه ۲ از نرم $l_{2,1}$ به‌جای نرم ۱ استفاده می‌شود که می‌تواند هر سطر ماتریس را هم‌زمان صفر یا غیر صفر کند. سطرهای غیر صفر هر نمونه هدف را با حداقل تعداد نمونه منبع نشان می‌دهند، بنابراین سطرهای غیر صفر به عنوان نمونه‌های منبع مرتبط انتخاب می‌شوند.

$$\min_A \|X_t - X_s A\|_F^2 + \rho \|A\|_{2,1} \quad (2)$$

بدین ترتیب با استفاده از نرم $l_{2,1}$ یک مجموعه محدود از نمونه‌های منبع به‌دست می‌آید که می‌تواند همه نمونه‌های هدف را به‌خوبی توصیف کند.

قسمت اول رابطه (۸) یک انتخاب نمونه انجام می‌دهد. قسمت دوم رابطه (۸) یک محدودیت روی کدگذاری تنک است که بتواند هر نمونه هدف را با حداقل تعداد نمونه منبع توصیف کند. قسمت سوم رابطه (۸) به منظور حفظ هندسه اصلی نمونه‌ها در فضای نگاشت شده است و قسمت چهارم رابطه (۸) برای جلوگیری از جواب‌های بدیهی است و محدودیت رابطه (۸) برای ماکزیمم کردن واریانس داده‌ها است تا کلاس‌ها به خوبی از یکدیگر متمایز شوند.

روش پیشنهادی SADA با استفاده از رابطه ۸ دو ماتریس نگاشت U_s و U_t را با استفاده از PCA به دست می‌آورد و ماتریس نگاشت A را با حل معادله جهت نگاشت داده‌های منبع و هدف به زیرفضای جدید ایجاد می‌کند. باین حال، در بیشتر مسائل دنیای واقعی اختلاف توزیع دامنه‌های منبع و هدف به اندازه‌ای است که صرفاً با نگاشت داده‌ها به فضای جدید، توزیع داده‌ها یکسان نشده و کارایی مدل بهبود داده نمی‌شود. از این رو، در مرحله دوم الگوریتم از یک طبقه‌بند انطباقی جهت مقاوم کردن مدل در مقابل تغییرات توزیع داده‌ها بهره برده می‌شود.

۴-۲- طبقه‌بند انطباقی

هدف از ایجاد طبقه‌بند انطباقی رسیدن به دو هدف زیر است: (۱) حداقل کردن خطای پیش‌بینی مدل روی داده‌های منبع و (۲) حداکثر کردن سازگاری بین ساختار منیفلد و مدل پیش‌بینی. درواقع، طبقه‌بند انطباقی پارامترهای مدل را به گونه‌ای به دست می‌آورد که مقاومت مدل در مقابل تغییرات داده بین دو دامنه افزایش یابد و کارایی آن با شیف داده‌ها متأثر نشود. از همین رو، تابع خطای مدل به صورت زیر تعریف می‌شود که در آن f طبقه‌بند انطباقی و $g(x)$ داده‌های نگاشت شده توسط ماتریس نگاشت A می‌باشند.

$$\min_{f \in H_k} \sum_{i=1}^{n_s} (f(g(x_i)), y_i) + \sigma f_k^2 + \gamma M_f(P_s(g(x)), P_t(g(x))) \quad (9)$$

که در آن H_k مجموعه طبقه‌بندها و σ و γ پارامترهای تنظیم هستند. بخش اول رابطه (۹) پارامترهای مدل را به گونه‌ای تعیین می‌کند که خطای پیش‌بینی مدل حداقل شود. بخش دوم رابطه (۹) شامل مربع طبقه‌بند انطباقی جهت جلوگیری از جواب‌های بدیهی است. بخش سوم رابطه (۹) جهت حفظ ساختار و توپولوژی داده‌ها در قالب طبقه‌بند از لاپلاسیان داده‌ها استفاده می‌کند.

جهت به دست آوردن مقادیر بهینه پارامترهای مدل f از رابطه (۹) مشتق گرفته و برابر صفر قرار داده شده است تا پارامترهای انطباقی طبقه‌بند به صورت زیر به دست آیند:

$$\alpha = (\sigma I + (R + \gamma L) X)^{-1} R Y^T \quad (10)$$

که در آن I ماتریس همانی است. R یک ماتریس قطری است که نشان‌دهنده نمونه‌های منبع است و Y برچسب‌های نمونه‌های منبع است.

که همسایگی نمونه‌های منبع و هدف در فضای نگاشت شده حفظ شده باشد:

$$\min_U \text{tr}(U^T X L X^T U) \quad (5)$$

که در آن L ماتریس لاپلاسیان نمونه‌های منبع و هدف است.

۴- روش پیشنهادی

این مقاله به دنبال پیشنهاد روشی است که بتواند یک نمایش با بعد کم برای داده‌های منبع و هدف ایجاد کند که در آن فاصله دو دامنه به حداقل مقدار ممکن برسد. همچنین نمونه‌هایی از دامنه منبع که به دامنه هدف مرتبط نمی‌باشند، حذف شوند. برای این اهداف، از نمایش تنک برای انتخاب نمونه‌ها و PCA جهت کاهش بعد داده‌ها استفاده شده است. به علاوه، برای کاهش فاصله بین دامنه‌های منبع و هدف با استفاده از یک روش نوین، توزیع حاشیه‌ای بین دو دامنه را به حداقل می‌رساند.

اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف با استفاده از یک روش غیر پارامتری به نام MMD کاهش می‌یابد. در این روش داده‌ها به فضای RKHS برده شده و در آنجا اختلاف بین میانگین نمونه‌های دامنه‌های منبع و هدف محاسبه می‌شود. اگر A تابع نگاشت به فضای جدید در نظر گرفته شود، اختلاف توزیع حاشیه‌ای دامنه‌های منبع و هدف به صورت رابطه زیر تعریف می‌شود [۱۱، ۱۲]:

$$\text{Mrg}(X_s, X_t) = \left\| \frac{1}{n_s} \sum_{i=1}^{n_s} A X_i - \frac{1}{n_t} \sum_{j=n_s+1}^{n_s+n_t} A X_j \right\|_H^2 \quad (6)$$

برای حل ساده‌تر رابطه (۶) می‌توان آن را به فرم بسته زیر نوشت:

$$\text{Mrg}(X_s, X_t) = \text{tr}(A^T X M_0 X^T A) \quad (7)$$

که ماتریس $M_0 \in R^{(n_s+n_t) \times (n_s+n_t)}$ همان ماتریس MMD است و اگر

$$(M_0)_{ij} = \frac{1}{n_s n_t} \quad \text{اگر } x_i, x_j \in X_s \text{ باشد}$$

$$(M_0)_{ij} = \frac{-1}{n_s n_t} \quad \text{در غیر این صورت}$$

آنگاه $(M_0)_{ij} = \frac{1}{n_t n_t}$ است.

۴-۱- فرموله‌سازی روش SADA

روش پیشنهادی SADA ترکیبی از رویکردهای مبتنی بر نمونه و رویکردهای مبتنی بر خصوصیت است. به همین دلیل از روابط ۳ و ۴ جهت انتخاب نمونه و یافتن فضای خصوصیات استفاده می‌شود. همچنین با توجه به این که در نگاشت داده‌ها به فضای جدید بایستی توپولوژی داده‌ها حذف شود و اختلاف دو دامنه نیز کاهش یابد، ماتریس نگاشت A به گونه‌ای حاصل می‌شود که اختلاف حاشیه‌ای دامنه‌ها با استفاده از روابط ۵ و ۶ حداقل شود. بدین ترتیب تابع هدف SADA به صورت زیر به دست می‌آید:

$$\min_{A, U_t, U_s} \|U_t^T X_t - U_s^T X_s A\|_F^2 + \rho \|A\|_{2,1} + \mu \text{tr}(U^T X L X^T U) + \|U\|_2^2 \quad (8)$$

$$\text{s.t. } U^T X V X^T U = I$$

که μ و $\mathbf{M} + \mathbf{Z} = \begin{bmatrix} \mathbf{A}\mathbf{A}^T + \mu\mathbf{L}_s + \tau\mathbf{Z}_{ss} & -\mathbf{A} + \tau\mathbf{Z}_{st} \\ -\mathbf{A}^T + \tau\mathbf{Z}_{ts} & \eta\mathbf{I} + \mu\mathbf{L}_t + \tau\mathbf{Z}_{tt} \end{bmatrix}$ است و μ پارامتر تنظیم گراف است. روش SADA با جزئیات کامل، در شکل ۲ ارائه شده است.

Algorithm 2. Image processing via sparse coding and adaptive classification (SADA)

1: Input: Source data $\mathbf{X}_s$, Target data $\mathbf{X}_t$, Projection matrix $\mathbf{U}_s^T, \mathbf{U}_t^T$, subspace dimension P , $\mathbf{Y}_s$ source data label.
 2: Parameters: Regularization parameters δ, γ
 3: Output: Target domain labels $\mathbf{Y}_t$
 4: Randomly initialize $\mathbf{U} \in \mathbf{R}^{(d_s+d_t)*p}$.
 6: Compute selection matrix $\mathbf{A}$ via Algorithm 1.
 7: Compute projection matrix $\mathbf{U}$ via solving Eq. (13).
 8: Repeat 6 and 7 until the maximum number of iterations is reached or optimum values of Eq. (8) is found.

// Data projection

9: $\mathbf{X}_s' = \mathbf{U}_s^T \mathbf{X}_s \mathbf{A}$

10: $\mathbf{Y}_s' = \mathbf{Y}_s \mathbf{A}$

11: $\mathbf{X}_t' = \mathbf{U}_t^T \mathbf{X}_t$

12: $\mathbf{X}' = [\mathbf{X}_s', \mathbf{X}_t']$

// Adaptive classification

13: Compute weight matrix $\mathbf{W}$ by $\mathbf{w}_{ij} = e^{-\frac{(x_i - x_j)^2}{\delta}}$ where $(x_i, x_j) \in \mathbf{X}'$

14: Compute diagonal matrix $\mathbf{D}$ by $\mathbf{D}_{ii} = \sum_{j=1}^{n_s+n_t} \mathbf{w}_{ij}$.

15: $\mathbf{L} \leftarrow \mathbf{I} - \mathbf{D}^{-\frac{1}{2}} \mathbf{W} \mathbf{D}^{-\frac{1}{2}}$

16: Compute diagonal matrix $\mathbf{R}$ where $R_{ii} = \begin{cases} 1 & x_i \in \mathbf{X}_s' \\ 0 & \text{otherwise} \end{cases}$

17: $\alpha \leftarrow (\sigma \mathbf{I} + (\mathbf{R} + \gamma \mathbf{L}) \mathbf{X})^{-1} \mathbf{R} \mathbf{Y}'^T$

18: Learn an adaptive classifier f .

19: Return target domain labels $\mathbf{Y}_t$ using classifier f .

شکل ۲: الگوریتم SADA

در الگوریتم SADA، هدف به دست آوردن ماتریس انتخاب $\mathbf{A}$ و ماتریس نگاشت $\mathbf{U}$ و در نهایت پیش‌بینی برچسب‌های دامنه هدف با استفاده از طبقه‌بند انطباقی است. SADA در ابتدا با استفاده از الگوریتم ۱ ماتریس انتخاب $\mathbf{A}$ را به دست آورده و سپس با استفاده از الگوریتم ۲ یک مقدار اولیه تصادفی به ماتریس نگاشت $\mathbf{U}$ می‌دهد و با استفاده از رابطه ۱۳ مقادیر دقیق ماتریس نگاشت $\mathbf{U}$ را به دست می‌آورد. با به دست آمدن مقادیر $\mathbf{A}$ و $\mathbf{U}$ در یک فرایند تکرارشونده، داده‌های ورودی به صورت $\mathbf{X}_s' = \mathbf{U}_s^T \mathbf{X}_s \mathbf{A}$ و $\mathbf{X}_t' = \mathbf{U}_t^T \mathbf{X}_t$ به زیرفضای جدید نگاشت می‌شوند. در زیرفضای به دست آمده $\mathbf{X}_s'$ و $\mathbf{X}_t'$ دارای

با به دست آوردن پارامترهای طبقه‌بند انطباقی با استفاده از رابطه (۱۰) می‌توان مدل را با استفاده از یک طبقه‌بند استاندارد (نظیر SVM) ایجاد کرد.

۴-۳ الگوریتم روش SADA

برای بهینه‌سازی پارامترهای مدل، به دلیل اینکه نمی‌توان هم‌زمان چند پارامتر را با هم بهینه کرد به صورت یک‌به‌یک بهینه‌سازی پارامترها انجام می‌شود. با استفاده از یک الگوریتم بهینه‌سازی متناوب در دو مرحله می‌توان مسئله را حل کرد. در مرحله اول ماتریس تنک $\mathbf{A}$ را با ثابت نگه داشتن $\mathbf{U}_s$ و $\mathbf{U}_t$ می‌توان محاسبه کرد. سپس در مرحله دوم ماتریس‌های نگاشت $\mathbf{U}_s$ و $\mathbf{U}_t$ با ثابت نگه داشتن $\mathbf{A}$ محاسبه می‌شوند. با مشتق‌گیری از رابطه (۸) نسبت به $\mathbf{A}$ و مساوی قرار دادن آن با صفر، مقدار $\mathbf{A}$ بهینه با استفاده از رابطه زیر به دست می‌آید:

$$\mathbf{A} = (\rho \mathbf{N} + \mathbf{X}_s^T \mathbf{U}_s \mathbf{U}_s^T \mathbf{X}_s)^{-1} \mathbf{X}_s^T \mathbf{U}_s \mathbf{U}_t^T \mathbf{X}_t \quad (11)$$

که در رابطه فوق، $\mathbf{N} = \frac{1}{2 \|\mathbf{A}'\|_2}$ یک ماتریس قطری و وابسته به $\mathbf{A}$ است و $\mathbf{A}$ با یک روش تکراری به دست می‌آید. نحوه به دست آوردن $\mathbf{A}$ در شکل ۱ ارائه شده است.

Algorithm 1. An algorithm for solving selection matrix $\mathbf{A}$

1: Input: Source data $\mathbf{X}_s$, Target data $\mathbf{X}_t$, Projection matrix $\mathbf{U}_s^T, \mathbf{U}_t^T$

2: Parameter: Regularization parameters ρ

3: Output: Selection matrix $\mathbf{A}$.

5: Compute selection matrix $\mathbf{A}$ via solving Eq. (11).

6: Compute diagonal matrix $\mathbf{N} = \frac{1}{2 \|\mathbf{A}'\|_2}$ based on $\mathbf{A}$.

7: Repeat 5 and 6 until the maximum number of iterations is reached or optimum values of Eq. (8) is found.

شکل ۱: الگوریتم به دست آوردن ماتریس $\mathbf{A}$

همچنین، با فرض اینکه متغیر $\mathbf{A}$ در رابطه (۸) ثابت باشد، بایستی نسبت به متغیر $\mathbf{U}$ مشتق گرفته شود. با توجه به اینکه محدودیت مسئله خود بر حسب متغیر $\mathbf{U}$ می‌باشد به کمک لاگرانژ مقدار بهینه متغیر $\mathbf{U}$ به دست می‌آید. مقدار تابع لاگرانژ به صورت زیر محاسبه می‌شود:

$$\min_{\mathbf{U}} tr(\mathbf{U}^T \mathbf{X} (\mathbf{M} + \mathbf{Z}) \mathbf{X}^T \mathbf{U}) - tr((\mathbf{U}^T \mathbf{X} \mathbf{V} \mathbf{X}^T \mathbf{U} - \mathbf{I}) \mathbf{Q}) \quad (12)$$

که $\mathbf{Q} = \text{diag}(q_1, \dots, q_p) \in \mathbf{R}^{p \times p}$ ضرایب لاگرانژ بوده و یک ماتریس قطری است. با مشتق‌گیری از رابطه (۱۲) بر حسب $\mathbf{U}$ و برابر صفر قرار دادن آن، رابطه زیر به دست می‌آید:

$$\mathbf{X} (\mathbf{M} + \mathbf{Z}) \mathbf{X}^T \mathbf{U} = \mathbf{X} \mathbf{V} \mathbf{X}^T \mathbf{U} \mathbf{Q} \quad (13)$$

نمونه از داده‌های دامنه MNIST به‌عنوان داده‌های تست به‌کارگرفته می‌شود. با جای نمونه‌های آموزشی و تست دامنه MNIST_vs_USPS(M_U) برای یک آزمایش دیگر ایجاد شده است.

پایگاه داده پای [۱۷]، پایگاه داده معیار برای تشخیص چهره است که شامل ۴۱۳۶۸ تصویر چهره از ۶۸ فرد با اندازه تصاویر ۳۲×۳۲ است. این پایگاه داده شامل ۵ دامنه مختلف است که عبارتند از: پای ۱ (حالت چپ)، پای ۲ (حالت بالا)، پای ۳ (حالت پایین)، پای ۴ (حالت روبرو)، پای ۵ (حالت راست). در مجموع ۲۰ آزمایش بین دامنه‌ای بر روی پایگاه داده پای قابل طراحی است که از میان ۵ دامنه، دو دامنه متفاوت به‌عنوان دامنه‌های منبع و هدف انتخاب می‌شوند و کارایی الگوریتم پیشنهادی در شرایط متفاوت مورد بررسی قرار می‌گیرد. اطلاعات جزئی‌تر از پایگاه داده‌ها در جدول ۱ نشان داده شده است.

جدول ۱: معرفی پایگاه داده‌ها

پایگاه داده	تعداد نمونه‌ها	تعداد ابعاد	تعداد کلاس‌ها	اختصار
USPS	۱۸۰۰	۲۵۶	۱۰	U
MNIST	۲۰۰۰	۲۵۶	۱۰	M
Amazon	۹۵۸	۸۰۰	۱۰	A
Webcam	۲۹۵	۸۰۰	۱۰	W
Dslr	۱۵۷	۸۰۰	۱۰	D
Caltech256	۱۱۲۳	۸۰۰	۱۰	C
PIE1	۳۳۳۲	۱۰۲۴	۶۸	P1
PIE2	۱۶۲۹	۱۰۲۴	۶۸	P2
PIE3	۱۶۳۲	۱۰۲۴	۶۸	P3
PIE4	۳۳۲۹	۱۰۲۴	۶۸	P4
PIE5	۱۶۳۲	۱۰۲۴	۶۸	P5

۵-۱-۲- تنظیمات آزمایش

در تمام آزمایش‌ها، از طبقه‌بند NN به‌عنوان طبقه‌بند پایه استفاده شده است. ضریب MMD می‌باشد که به صورت تجربی برای پایگاه داده آفیس و کالتک ۱، اعداد ۵، پای ۰/۰۰۵ و پارامتر η ، ۰/۰۰۱ در نظر گرفته شده است. پارامتر تنظیم تنگی ρ ، پارامتر تأثیرگذار در انتخاب نمونه‌های دامنه منبع می‌باشد که برای همه پایگاه داده‌ها با ۰/۰۰۱ تنظیم شده است. به علت اینکه در محدوده وسیعی از مقادیر پارامتر تنظیم گراف μ ثابت می‌باشد، بنابراین برای همه پایگاه داده‌ها، ۰/۰۰۱ فرض شده است. همچنین، پیاده‌سازی روش پیشنهادی SADA توسط نرم‌افزار متلب^{۳۲} انجام گرفته است.

۵-۲- تأثیر پارامترها

برای به‌دست‌آوردن تنظیمات مدل، روش SADA با مقادیر مختلف پارامترها مورد ارزیابی قرار گرفته است. سه پارامتر مختلف در این روش

حداقل اختلاف توزیع می‌باشند. در مرحله بعد طبقه‌بند انطباقی بر روی داده‌های نگاشت‌شده X'_s و X'_t اعمال می‌شود تا مدل مقاوم‌تری در برابر تغییرات داده‌ای ایجاد شود.

۵-۳- آزمایش‌ها

در این بخش، روش پیشنهادی از جنبه‌های مختلف بررسی شده و تأثیر پارامترهای مختلف مدل بر روی کارایی آن بیان شده است. سپس، ضرورت هر پارامتر در الگوریتم نشان داده شده و عملکرد SADA با دیگر الگوریتم‌ها، مورد مقایسه و ارزیابی قرار گرفته است و نتایج به دست آمده به تفصیل مورد بحث قرار گرفته است.

۵-۱- تنظیمات آزمایش و مجموعه داده‌ها

۵-۱-۱- معرفی مجموعه داده‌ها

پایگاه داده‌های مورد استفاده در این مقاله عبارتند از: (۱) آفیس^{۲۶} و کالتک^{۲۷}، (۲) اعداد (USPS و MNIST)، (۳) چهره (پای^{۲۸}). مجموعه داده آفیس [۱۳] و کالتک [۱۴]، از مجموعه پایگاه داده‌های شناخته شده برای مسئله تطبیق دامنه‌های بصری است که شامل چهار دامنه آمازون^{۲۹}، وبکم^{۳۰}، DSLR^{۳۱} و کالتک است. تصاویر موجود در دامنه آمازون از وبسایت‌های تجاری آنلاین دانلود شده است، تصاویر دامنه وبکم توسط دوربین‌های کم کیفیت وب گرفته شده است و تصاویر دامنه DSLR توسط دوربین‌های با کیفیت SLR دیجیتال گردآوری شده‌اند. تصاویر موجود در دامنه کالتک از وبسایت گوگل جمع‌آوری شده‌اند و تنها از ۱۰ کلاس مشترک بین دامنه‌ها استفاده شده است. تمام داده‌های منبع در الگوریتم استفاده شده‌اند و ۱۰ درصد از داده‌های هدف به صورت تصادفی به‌عنوان داده برچسب‌دار آموزشی استفاده شده است و بقیه داده‌های هدف به‌عنوان داده تست مورد استفاده قرار گرفته‌اند. بر روی این مجموعه داده، ۱۲ آزمایش طراحی شده است که در هر یک از این آزمایش‌ها یکی از مجموعه داده‌ها (به‌عنوان نمونه آمازون)، به‌عنوان دامنه منبع و یکی از سه مجموعه داده دیگر به‌عنوان دامنه هدف انتخاب می‌شوند.

مجموعه داده‌های USPS [۱۵] و MNIST [۱۶]، شامل اعداد دست‌نویس ۰ تا ۹ می‌باشند که USPS شامل تقریباً ۹۰۰۰ تصویر با اندازه ۱۶×۱۶ و MNIST شامل در حدود ۷۰۰۰۰ تصویر با اندازه ۲۸×۲۸ می‌باشد. تمام تصاویر به‌اندازه ۱۶×۱۶ تبدیل شده و دارای ۲۵۶ بعد می‌باشند. همه داده‌های منبع در الگوریتم استفاده شده‌اند و تنها ۱۰ درصد از داده‌های هدف به صورت تصادفی به‌عنوان داده برچسب‌دار آموزشی استفاده شده و بقیه داده‌های هدف به‌عنوان داده تست در نظر گرفته شده‌اند. این پایگاه داده شامل ۱۰ کلاس مختلف می‌باشد و به منظور آزمایش هر دو دامنه در شرایط یکسان، دامنه MNIST_vs_USPS(U_M) ایجاد شده است که به‌طور تصادفی ۱۸۰۰ نمونه از داده‌های دامنه USPS به‌عنوان داده‌های آموزشی و ۲۰۰۰

وجود دارد که مقادیر بهینه آن‌ها برای پایگاه‌داده‌های مختلف در جدول ۲ نشان داده شده است.

جدول ۲: مقادیر بهینه پارامترها برای سه پایگاه داده بصری (آفیس و کالتک، اعداد و پای)، P: ابعاد فضای جدید، δ : پارامتر تنظیم در طبقه‌بند انطباقی، γ : پارامتر تنظیم در طبقه‌بند انطباقی.

پارامترهای بهینه	P	δ	γ
آفیس و کالتک	۶۰	۰/۱	۰/۰۰۱
اعداد	۲۲۰	۰/۰۱	۰/۱
پای	۱۶۰	۰/۰۰۵	۰/۰۰۵

پارامترهای γ و δ در محدوده $[۰/۰۰۱, ۱۰/۰]$ و پارامتر P در محدوده $[۲۰, ۲۲۰]$ در نظر گرفته شده‌اند. نتایج به دست آمده از روش SADA با مقادیر مختلف پارامتر P با روش‌های JDA، TJM و VDA در پایگاه داده‌های آفیس و کالتک و اعداد مقایسه شده است.

نتایج به دست آمده برای مقادیر مختلف پارامتر δ بر روی پایگاه‌داده‌های آفیس و کالتک، اعداد و پای در شکل ۳ نشان می‌دهد که برای پایگاه‌داده آفیس و کالتک، مقادیر بالای δ دقت مدل را کاهش می‌یابد. مقادیر بالای پارامتر δ پیچیدگی مدل را افزایش داده و باعث می‌شود تأثیر سایر عوامل در ایجاد طبقه‌بند انطباقی نادیده گرفته شود. در چنین حالتی از ساختار اصلی داده‌ها، در ایجاد تطبیق بین دامنه‌های منبع و هدف استفاده نشده است. محدوده بهینه در پایگاه داده آفیس و کالتک برای پارامتر δ ، $[۰/۰۵, ۰/۱]$ است. پایگاه‌داده اعداد حساسیت چندانی به مقدار پارامتر δ ندارد و در هر دو آزمایش تغییرات زیادی با مقادیر مختلف متغیر δ بر روی نتایج ایجاد نمی‌شود. با این حال برای مقادیر کمتر از ۰/۱ نتایج مطلوب‌تری به دست می‌آید. پایگاه‌داده پای نسبت به پارامتر δ حساسیت زیادی دارد و در محدوده بالا مدل صحت کاهش قابل‌ملاحظه‌ای پیدا می‌کند، محدوده بهینه پایگاه‌داده پای برای پارامتر δ ، $[۰/۰۰۵, ۰/۰۱]$ است.

نتایج به دست آمده برای مقادیر مختلف پارامتر γ بر روی پایگاه‌داده‌های آفیس و کالتک، اعداد و پای در شکل ۴ نشان می‌دهد که برای پایگاه‌داده آفیس و کالتک، مقادیر بالای γ دقت مدل را کاهش می‌دهد. درواقع مقادیر بالای پارامتر γ ، باعث نادیده گرفتن اطلاعات داده‌های برجسب‌دار دامنه منبع در ایجاد طبقه‌بند انطباقی می‌شود. محدوده بهینه در پایگاه داده آفیس و کالتک برای پارامتر γ ، $[۰/۰۰۱, ۰/۰۱]$ است. پایگاه‌داده اعداد نسبت به مقادیر مختلف پارامتر γ حساسیت کمی دارد و در مقادیر بالا صحت مدل افزایش می‌یابد. پایگاه‌داده پای نسبت به پارامتر γ نیز حساسیت زیادی دارد و در محدوده بالا صحت کاهش قابل‌ملاحظه‌ای پیدا می‌کند. بدین ترتیب محدوده بهینه پایگاه‌داده پای برای پارامتر γ ، $[۰/۰۱, ۰/۰۱]$ است.

نتایج به دست آمده بر روی پایگاه‌داده آفیس و کالتک در شکل ۵ حاکی از آن است که پایگاه‌داده‌های مختلف در آفیس و کالتک دارای رفتار متفاوتی نسبت به مقادیر مختلف پارامتر P می‌باشند. دقت به دست آمده برای پارامتر P در هر روش، نشان‌دهنده صحت روش برای بازسازی داده‌ها در فضای جدید می‌باشد. به عبارتی، نشان‌دهنده صحت انتقال دانش از دامنه منبع به دامنه هدف در فضای جدید می‌باشد. برای به دست آوردن مقدار محدوده بهینه پارامتر، محدوده‌ای انتخاب می‌شود که دارای بهینه‌ترین صحت بر روی تعداد بیشتری از پایگاه داده‌ها باشد. برای بیشتر پایگاه داده‌ها محدوده بهینه پارامتر P در پایگاه داده آفیس و کالتک، $[۰/۰۰۱, ۰/۰۱]$ می‌باشد.

نتایج به دست آمده برای مقادیر مختلف P در پایگاه داده اعداد در شکل ۶ نشان می‌دهد که پایگاه‌داده اعداد دارای حساسیت کمی نسبت به مقادیر مختلف P می‌باشد. بهترین محدوده برای پارامتر P در پایگاه داده اعداد، $[۲۲۰, ۱۴۰]$ می‌باشد.

ρ پارامتر تنظیم تنکی است که یکی از پارامترهای تأثیرگذار در SADA می‌باشد. ρ تنکی ضرایب ماتریس A را کنترل می‌کند و بنابراین تعداد نمونه‌های منبع مرتبط را انتخاب می‌کند. صحت طبقه‌بندی و تعداد نمونه‌های انتخاب شده با توجه به ρ در جدول ۳ گزارش شده است. به عنوان نمونه در پایگاه‌داده U-M هنگامی که $\rho = ۰/۵$ است هیچ نمونه‌ای از دامنه منبع انتخاب نشده است؛ اما زمانی که ρ کوچک است، ماتریس A تنکی کمتری دارد و در نتیجه نمونه‌های منبعی که توزیع نزدیک‌تری با نمونه‌های هدف دارند انتخاب می‌شوند. به عنوان نمونه در پایگاه‌داده U-M هنگامی که $\rho = ۰/۰۰۱$ است، تمام نمونه‌های منبع انتخاب شده‌اند. در تمام آزمایش‌ها مقدار پارامتر ρ برابر با ۰/۰۱ تنظیم شده است.

در شکل ۷ تعدادی از نمونه‌های انتخاب شده و انتخاب نشده از دامنه کالتک را نشان می‌دهد زمانی که دامنه آمازون به عنوان دامنه هدف در نظر گرفته شده است. به عنوان مثال، کوله‌پشتی از دامنه کالتک که شباهت بیشتری به کوله‌پشتی از دامنه آمازون دارند انتخاب شده و آن‌هایی که هیچ‌گونه شباهتی ندارند انتخاب نشده‌اند.

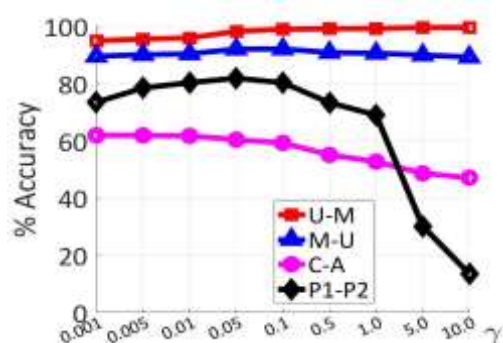
ρ	۰/۰۰۱	۰/۰۰۵	۰/۰۱	۰/۰۵	۰/۱	۰/۵
U-M	۹۸/۹۲	۹۸/۷۰	۹۸/۳۹	۹۷/۲۲	۹۷/۱۴	۹۷/۷۰
U-M #	۱۸۰۰	۱۸۰۰	۱۸۰۰	۷۸۲	۴۰۹	۰
C-A	۶۱/۸۹	۶۱/۹۹	۶۱/۳۳	۵۵/۶۹	۵۸/۰۷	۵۸/۷۴
C-A #	۱۰۶۸	۶۴۶	۴۹۴	۷۳	۲	۰

جدول ۳: تأثیر پارامتر تنظیم تنکی ρ بر تعداد نمونه‌های انتخاب شده از دامنه منبع

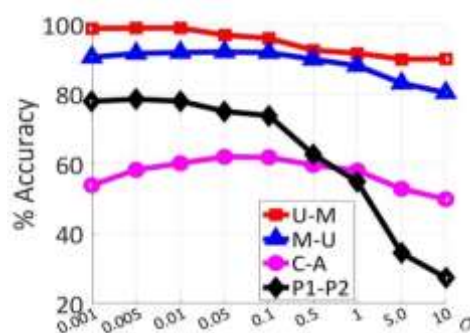
۵-۳- ضرورت هر پارامتر در الگوریتم

در این بخش، اهمیت وجودی برخی از پارامترها با طراحی دو آزمایش بر روی پایگاه‌داده آفیس و کالتک بررسی شده و نتایج آن در جدول ۴ نشان داده شده است.

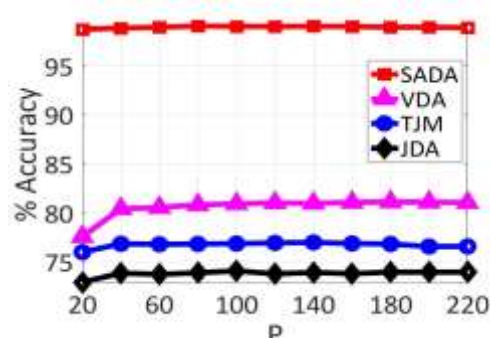
با بررسی تأثیر استفاده از نرم‌های مختلف می‌توان نتیجه گرفت که



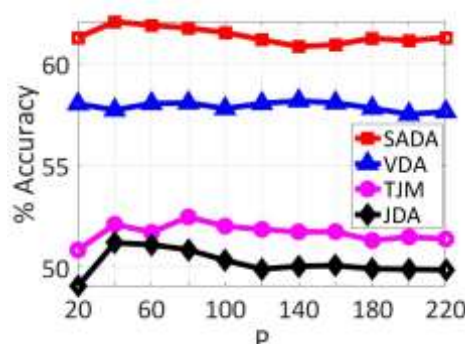
شکل ۴: ارزیابی صحت پایگاه داده‌های U-M، M-U، C-A، P1-P2 با مقادیر مختلف پارامتر σ



شکل ۳: ارزیابی صحت پایگاه داده‌های U-M، M-U، C-A، P1-P2 با مقادیر مختلف پارامتر σ



شکل ۶: ارزیابی دقت روش‌های TJM، JDA، VDA و SADA در پایگاه داده U-M با مقادیر مختلف پارامتر P



شکل ۵: ارزیابی دقت روش‌های TJM، JDA، VDA و SADA در پایگاه داده C-A با مقادیر مختلف پارامتر P

نرم I_1 کارایی قابل توجهی نداشته و استفاده از نرم $I_{2,1}$ باعث می‌شود که هر نمونه هدف با حداقل تعداد نمونه منبع به خوبی توصیف شود.

بررسی تأثیر پارامتر همسایگی نشان‌دهنده این است که، حفظ توپولوژی و ساختار داده‌ها در نگاشت آن‌ها به زیرفضای جدید تأثیر بسیار تعیین‌کننده‌ای در نتایج به دست آمده دارد. درواقع، در صورتی که داده‌ها بدون حفظ توپولوژی به فضای جدید نگاشت شوند، ساختار اصلی خود را از دست داده و مدل به دست آمده قابل اعتماد نمی‌باشد. در نهایت تأثیر طبقه‌بند انطباقی بر روی مدل بررسی شده است. نتایج به دست آمده، نشان می‌دهد که SADA با تطبیق پارامترهای مدل کارایی الگوریتم را به صورت قابل ملاحظه‌ای افزایش می‌دهد. درواقع، طبقه‌بند انطباقی مقاوم‌پذیری مدل را در راستای تغییرات داده به اندازه قابل توجهی افزایش می‌دهد.

جدول ۴: ضرورت هر پارامتر در الگوریتم

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

$SADA_{\mu}$: الگوریتم با حذف پارامتر تنظیم گراف

SADA	SVM	$SADA_{\mu}$	I_1 norm	Orj	S-T
۶۱/۹۲	۵۷/۶	۵۵/۶۲	۶۱/۸۸	۶۱/۶۹	C-A
۱۰۶۹	۱۰۶۹	۲۷۳	۱۰۸۷	۱۱۲۳	C-A #
۶۳/۴۷	۴۸/۶	۴۴/۳۴	۶۲/۳۴	۶۲/۳۸	C-W
۱۰۶۹	۸۸۰	۱۸۸	۹۷۹	۱۱۲۳	C-W#

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

۲. استفاده از نرم I_1 به جای نرم $I_{2,1}$: با استفاده از نرم I_1 به جای نرم $I_{2,1}$ تعداد ضرایب غیر صفر در ماتریس A افزایش یافته و دقیقاً مشخص نمی‌شود که نمونه منبع دقیقاً با کدام نمونه هدف مرتبط است.

۳. حذف پارامتر تنظیم گراف ($SADA_{\mu}$): با حذف این پارامتر اهمیت نگهداری ساختار هندسی هر دامنه مشخص می‌شود.

۴. استفاده از طبقه‌بند SVM به جای طبقه‌بند انطباقی: با استفاده از یک طبقه‌بند استاندارد SVM، اهمیت استفاده از طبقه‌بند انطباقی مشخص می‌شود.

با بررسی نتایج حاصل‌شده از آزمایش‌های متعدد در جدول ۴ می‌توان نتیجه گرفت که استفاده از تمام نمونه‌های منبع در ساخت مدل باعث کاهش ملموس کارایی مدل می‌شود. از این رو SADA با استفاده از کدگذاری تنک و نرم $I_{2,1}$ اقدام به حذف نمونه‌های غیرمرتبط بین دو دامنه می‌نماید.

۵-۴- اهمیت استفاده از برچسب‌های دامنه هدف (یادگیری

نیمه‌نظارت‌شده)

دقت الگوریتم پیشنهادی بر روی پایگاه‌داده آفیس و کالتک در دو حالت نیمه‌نظارت‌شده (SADA) و بدون نظارت ($SADA^-$)، در شکل ۸ نشان داده شده است. نتایج به‌دست‌آمده به‌وضوح نشان‌دهنده تأثیر قابل‌ملاحظه برچسب‌های دامنه هدف در افزایش دقت مدل دارد. درواقع، در یادگیری نیمه‌نظارت‌شده، داده‌های برچسب‌دار دامنه هدف به‌عنوان یک پل و رابط بین دو دامنه عمل می‌کنند. بدین ترتیب، در یافتن زیر فضای مشترک بین دامنه‌های منبع و هدف اطلاعات تفکیک‌کننده دو منبع نیز مورد استفاده قرار می‌گیرند.

۵-۵- مقایسه با الگوریتم‌های موجود

الگوریتم SADA با هفت الگوریتم زیر مورد مقایسه قرار گرفته است: طبقه‌بند نزدیک‌ترین همسایه (NN^{33})، PCA [۱۱]، TCA^{34} [۱]، GFK [۶]، JDA^{35} [۱۸]، TJM [۹] و VDA [۱۲]. به‌دلیل این که تمامی روش‌های نام‌برده‌شده (به‌جز NN)، روش‌های کاهش بعد می‌باشند، از طبقه‌بند استاندارد نزدیک‌ترین همسایه [۱۹] برای آموزش یک طبقه‌بند بر روی داده‌های دامنه منبع جهت پیش‌بینی برچسب داده‌های دامنه هدف استفاده می‌شود.

صحت طبقه‌بندی SADA و هفت روش دیگر بر روی پایگاه‌داده‌های آفیس و کالتک، اعداد و پای در جدول ۵ نشان داده شده است. همان‌طور که مشاهده می‌شود SADA کارایی بهتری نسبت به هفت روش دیگر از خود نشان داده و در بیشتر حالات عملکرد مطلوبی ارائه می‌دهد. نتایج گزارش‌شده شامل میانگین و انحراف از استاندارد برای همه روش‌ها می‌باشد.

نتایج به‌دست‌آمده حاکی از میانگین صحت $SADA^{34}$ روی ۷۵/۷۵٪، پایگاه داده و بهبود صحت ۴/۵۶٪ نسبت به بهترین الگوریتم مورد مقایسه VDA و ۲۳/۲۹٪ نسبت به الگوریتم استاندارد NN می‌باشد و در ۲۶ آزمایش از ۳۴ آزمایش SADA عملکرد بهتری نسبت به روش‌های دیگر ارائه نموده است. در ادامه عملکرد SADA در مقایسه با هر یک از روش‌های موجود مورد بحث قرار گرفته است.

PCA یکی از روش‌های کاهش بعد است که برای کاهش ابعاد داده‌ها مورد استفاده قرار می‌گیرد. PCA اختلاف توزیع بین دامنه‌ها را چندان کاهش نمی‌دهد و عملکرد ضعیفی نسبت به دیگر الگوریتم‌های تطبیق دامنه نشان می‌دهد. نتایج حاصل از آزمایش‌ها حاکی از این است که SADA به تفکیک نسبت به PCA دارای متوسط بهبود صحت ۹/۳۹٪ در پایگاه‌داده آفیس و کالتک، ۹/۴٪ در پایگاه‌داده اعداد و ۳۹/۳۹٪ در پایگاه‌داده پای است و به‌طور کلی دارای ۲۸/۰۴٪ بهبود صحت است. عمده‌ترین دلیل برتری SADA نسبت به PCA استفاده از تطبیق خصوصیات برای کاهش اختلاف توزیع حاشیه‌ای بین دامنه‌ها می‌باشد.

TCA یک روش تطبیق دامنه بدون نظارت است و از برچسب‌های دامنه منبع برای تطبیق دامنه‌ها استفاده نمی‌کند و تنها اختلاف توزیع حاشیه‌ای بین دامنه‌ها را کاهش می‌دهد و برای حفظ ویژگی‌های اصلی داده‌های ورودی، واریانس داده‌ها را حداکثر می‌سازد. نتایج حاصل از آزمایش‌ها حاکی از این است که SADA به تفکیک نسبت به TCA دارای متوسط بهبود صحت ۴/۶۴٪ در پایگاه‌داده آفیس و کالتک، ۱۲/۱۲٪ در پایگاه‌داده اعداد و ۶۰/۹۳٪ در پایگاه‌داده پای است و به‌طور کلی دارای ۳۹/۱۹٪ بهبود است. دلیل بهبود قابل‌ملاحظه SADA نسبت به TCA این است که SADA علاوه بر استفاده از برچسب‌های دامنه منبع، از بخش کوچکی از برچسب‌های دامنه هدف نیز برای تطبیق دامنه‌ها استفاده می‌کند و با استفاده از یک طبقه‌بند انطباقی خطای پیش‌بینی مدل روی داده‌های منبع را حداقل می‌سازد.

در روش GFK به‌منظور کاهش اختلاف توزیع حاشیه‌ای، داده‌های منبع و هدف به یک زیرفضای کوچک‌تری در منیفولد گر سمن نگاشت می‌شوند، اما به‌دلیل این که زیر فضاهای نگاشت شده دارای ابعاد بسیار کمی می‌باشند باعث از بین رفتن بخشی از داده‌ها شده و نمایش خوبی از داده‌ها ایجاد نمی‌شود. نتایج حاصل از آزمایش‌ها حاکی از این است که متوسط بهبود صحت روش SADA به تفکیک نسبت به روش GFK در پایگاه‌داده آفیس و کالتک ۶/۲۳٪، در پایگاه‌داده اعداد ۸/۸٪ و در پایگاه‌داده پای ۳۷/۵۱٪ است و در مجموع دارای ۲۵/۷۸٪ بهبود است. دلیل برتری SADA نسبت به GFK این است که در روش SADA یک زیر فضای مشترک با حداکثر واریانس ایجاد می‌شود که ساختار اصلی داده‌ها را حفظ می‌کند و با استفاده از تطبیق خصوصیات، اختلاف توزیع حاشیه‌ای کاهش می‌یابد.

روش‌های TJM و JDA از جمله جدیدترین روش‌های تطبیق دامنه می‌باشند. روش TJM یک روش ترکیبی از روش‌های مبتنی بر نمونه و مبتنی بر خصوصیت است که تنها اختلاف توزیع حاشیه‌ای را کاهش می‌دهد. باین حال TJM به‌دلیل اینکه بخشی از داده ورودی را حذف می‌کند، ساختار اصلی داده‌ها را حفظ نمی‌کند. درحالی که روش JDA، یک روش بدون نظارت می‌باشد که اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌های منبع و هدف را به‌طور هم‌زمان کاهش می‌دهد و دارای عملکرد بهتری نسبت به TJM می‌باشد.

متوسط بهبود صحت SADA به تفکیک نسبت به روش TJM در پایگاه‌داده آفیس و کالتک ۵/۲۳٪، در پایگاه‌داده اعداد ۱۶/۵۴٪ و در پایگاه‌داده پای ۵۹/۸۵٪ است. همچنین متوسط بهبود صحت روش SADA به تفکیک نسبت به روش JDA در پایگاه‌داده آفیس و کالتک ۴/۱۱٪، در پایگاه‌داده اعداد ۱۲/۱۰٪ و در پایگاه‌داده پای ۶/۰۶٪ است. دلیل برتری SADA نسبت به روش‌های TJM و JDA این است که SADA ساختار داده‌ها را حفظ کرده و از یک روش نیمه‌نظارت‌شده جهت انطباق دامنه‌ها استفاده می‌کند. همچنین SADA از یک طبقه‌بند انطباقی برای کاهش خطای پیش‌بینی مدل بر روی داده‌ها استفاده می‌کند.

جدول ۵. صحت (%) طبقه‌بند بر روی ۳۴ پایگاه داده بصری

SADA	VDA	TJM	JDA	GFK	TCA	PCA	NN	Dataset
۹۸/۹۲±۰/۳۰	۷۳/۵۷±۱/۰۶	۷۵/۹۴±۱/۲۱	۶۷/۸۶±۱/۶۱	۸۴/۰۶±۰/۸۶	۸۰/۳۶±۱/۰۱	۸۳/۷۷±۱/۳۵	۸۲/۶۸±۰/۸۸	U_M
۹۲/۰۰±۱/۷۳	۸۰/۰۰±۱/۲۲	۸۱/۹۰±۲/۰۶	۷۴/۹۹±۴/۷۸	۸۹/۲۷±۰/۳۱	۸۶/۳۱±۰/۷۹	۸۸/۲۸±۰/۸۹	۸۹/۳۶±۰/۴۲	M_U
۷۸/۵۱±۴/۰۱	۷۳/۹۹±۲/۵۵	۱۳/۶۱±۰/۳۳	۷۴/۶۴±۱/۱۲	۳۴/۵۶±۱/۴۲	۱۲/۷۸±۰/۴۶	۳۰/۸۲±۱/۰۶	۴۷/۴۹±۱/۳۶	P1_P2
۷۹/۲۱±۴/۷۶	۷۲/۲۷±۲/۱۸	۱۴/۱۴±۰/۱۵	۷۳/۴۷±۱/۰۶	۳۶/۵۳±۰/۷۵	۱۴/۱۹±۰/۳۲	۳۳/۰۹±۰/۶۶	۴۹/۷۱±۲/۷۴	P1_P3
۹۱/۱۹±۱/۲۹	۸۵/۳۷±۰/۴۰	۲۲/۴۴±۰/۱۰	۸۷/۴۴±۱/۰۷	۴۴/۶۱±۱/۱۵	۱۹/۲۳±۰/۳۴	۴۲/۷۰±۰/۹۹	۵۹/۸۴±۱/۴۷	P1_P4
۷۷/۰۹±۷/۸۰	۶۸/۴۰±۱/۴۱	۱۳/۱۰±۰/۲۱	۶۹/۹۳±۱/۷۶	۳۴/۲۸±۱/۰۷	۱۴/۹۲±۰/۲۳	۳۲/۰۲±۱/۲۱	۴۹/۰۷±۱/۴۲	P1_P5
۹۰/۲۷±۱/۱۶	۸۰/۱۳±۰/۶۱	۲۳/۵۸±۰/۵۳	۷۴/۶۱±۵/۵۰	۴۵/۴۷±۰/۸۸	۱۸/۶۲±۰/۳۸	۴۴/۹۰±۱/۱۶	۶۲/۸۴±۱/۵۰	P2_P1
۷۲/۵۲±۸/۱۹	۸۲/۳۵±۰/۹۳	۲۵/۶۰±۰/۴۰	۷۵/۰۶±۱/۱۱	۴۴/۹۲±۰/۸۵	۲۴/۴۷±۰/۷۷	۴۲/۳۷±۱/۹۲	۵۹/۵۱±۲/۹۷	P2_P3
۸۹/۴۲±۲/۲۶	۸۶/۵۹±۰/۹۷	۳۱/۵۰±۰/۲۹	۸۳/۷۲±۱/۸۱	۵۸/۵۹±۰/۷۶	۳۲/۴۴±۰/۴۴	۵۸/۰۰±۰/۷۰	۷۲/۵۹±۰/۳۱	P2_P4
۶۸/۲۱±۰/۶۵	۷۰/۵۴±۱/۰۰	۱۶/۸۵±۰/۲۰	۶۶/۰۹±۲/۰۷	۳۵/۲۰±۱/۴۲	۱۴/۷۳±۰/۶۴	۳۴/۷۵±۱/۳۲	۵۱/۵۵±۱/۱۸	P2_P5
۹۰/۱۸±۱/۲۰	۷۷/۴۶±۱/۶۰	۲۲/۸۱±۰/۷۶	۷۳/۶۶±۷/۰۰	۴۴/۶۵±۱/۲۷	۱۷/۹۹±۰/۵۱	۴۴/۵۳±۱/۰۴	۶۲/۲۹±۱/۳۹	P3_P1
۷۲/۸۸±۹/۹۰	۷۷/۳۰±۱/۳۳	۲۹/۱۲±۰/۱۰	۷۱/۱۲±۱/۲۰	۴۲/۷۰±۰/۵۳	۲۲/۰۱±۰/۶۷	۳۹/۳۳±۰/۴۷	۵۵/۲۹±۱/۲۰	P3_P2
۸۹/۶۸±۱/۵۳	۸۴/۸۵±۱/۱۷	۳۳/۳۹±۰/۳۷	۸۲/۵۲±۱/۸۵	۵۷/۳۰±۰/۷۹	۳۱/۷۷±۰/۴۳	۵۷/۲۱±۰/۴۶	۷۱/۰۴±۰/۶۹	P3_P4
۷۱/۹۷±۲/۷۴	۶۹/۷۹±۲/۸۳	۱۵/۱۵±۰/۳۹	۶۵/۲۹±۲/۳۱	۳۶/۱۷±۲/۴۷	۱۴/۲۸±۰/۴۰	۳۴/۲۶±۱/۸۶	۵۱/۲۵±۳/۰۴	P3_P5
۹۲/۸۹±۱/۲۱	۸۵/۵۳±۰/۹۵	۳۰/۶۸±۰/۷۹	۸۶/۱۲±۱/۲۹	۴۶/۰۹±۰/۸۲	۱۸/۶۴±۰/۶۶	۴۴/۰۸±۱/۳۲	۶۳/۹۹±۱/۰۶	P4_P1
۸۲/۶۸±۷/۱۲	۸۶/۹۰±۰/۷۲	۲۹/۰۳±۰/۲۳	۸۶/۵۳±۰/۶۹	۵۹/۹۱±۰/۹۳	۳۶/۴۰±۰/۴۷	۵۶/۲۲±۰/۷۲	۷۳/۳۸±۰/۶۲	P4_P2
۸۵/۹۰±۴/۲۳	۸۸/۵۴±۱/۰۱	۳۹/۵۰±۱/۲۵	۹۰/۱۵±۰/۳۳	۶۷/۲۵±۰/۹۵	۴۲/۶۸±۰/۴۳	۶۴/۷۱±۱/۱۶	۸۱/۵۱±۱/۲۴	P4_P3
۷۸/۶۸±۳/۲۰	۷۳/۷۳±۱/۷۲	۱۵/۰۷±۰/۳۰	۷۱/۶۵±۳/۵۷	۳۹/۵۲±۱/۵۶	۱۶/۰۹±۰/۲۵	۳۶/۰۳±۱/۸۶	۵۵/۹۴±۲/۱۷	P4_P5
۸۸/۱۳±۱/۵۳	۷۸/۵۱±۱/۶۷	۱۹/۱۰±۰/۴۳	۷۵/۱۹±۹/۵۲	۴۴/۶۴±۱/۲۹	۱۸/۵۵±۰/۵۴	۴۳/۹۶±۱/۵۵	۶۱/۹۵±۱/۳۷	P5_P1
۷۲/۹۴±۵/۱۶	۷۰/۷۲±۱/۷۲	۱۲/۷۰±۰/۲۵	۶۴/۹۸±۶/۳۰	۳۲/۹۲±۰/۶۵	۱۲/۷۴±۰/۱۵	۳۰/۹۴±۰/۴۳	۴۷/۶۳±۱/۷۲	P5_P2
۷۶/۲۰±۲/۱۷	۷۱/۸۰±۳/۹۹	۱۳/۵۱±۰/۸۱	۶۷/۸۶±۱/۶۱	۳۶/۱۰±۰/۲۶	۱۴/۶۵±۰/۲۰	۳۴/۲۳±۰/۸۷	۵۰/۴۸±۴/۴۱	P5_P3
۸۷/۶۶±۱/۹۵	۷۸/۹۵±۰/۹۶	۱۸/۴۰±۰/۳۰	۷۴/۹۹±۴/۷۸	۴۴/۵۲±۰/۵۸	۲۰/۵۱±۰/۴۰	۴۴/۱۵±۰/۶۰	۶۰/۱۷±۰/۶۸	P5_P4
۶۱/۹۲±۶/۵۷	۵۳/۳۴±۳/۱۱	۵۰/۸۵±۱/۷۳	۵۰/۳۴±۱/۳۳	۴۷/۷۴±۱/۰۰	۵۰/۷۷±۱/۳۲	۴۷/۹۰±۲/۷۰	۲۸/۰۷±۱/۱۲	C_A
۶۳/۴۷±۱۲/۹۴	۵۷/۹۲±۷/۶۶	۵۵/۸۱±۷/۲۰	۵۳/۰۰±۱/۰۶۰	۵۴/۷۹±۷/۹۴	۵۱/۸۵±۴/۰۹	۴۶/۵۳±۹/۹۱	۲۹/۰۲±۱/۱۵	C_W
۵۵/۱۱±۵/۰۰	۵۱/۳۹±۱۰/۹۲	۴۹/۹۳±۴/۵۲	۵۰/۳۶±۴/۵۰	۴۹/۸۵±۲/۳۷	۵۱/۸۲±۲/۹۶	۴۲/۶۳±۶/۵۴	۲۷/۳۷±۰/۸۶	C_D
۴۹/۴۸±۱/۶۷	۴۴/۴۳±۰/۹۱	۴۲/۰۰±۰/۹۱	۴۲/۷۹±۱/۲۲	۴۲/۴۶±۰/۸۲	۴۲/۷۰±۱/۷۰	۳۹/۲۰±۰/۵۲	۲۸/۰۰±۱/۰۶	A_C
۵۹/۲۸±۸/۲۴	۵۹/۴۳±۶/۴۹	۵۳/۴۷±۴/۲۴	۵۲/۲۶±۶/۳۷	۵۲/۹۴±۹/۱۲	۵۰/۸۷±۷/۵۳	۴۶/۱۱±۹/۲۳	۳۴/۳۸±۴/۹۲	A_W
۵۳/۲۱±۱/۰۰	۵۱/۴۶±۵/۵۹	۴۹/۴۹±۲/۳۴	۴۷/۰۸±۴/۰۶	۴۴/۲۳±۴/۵۲	۴۲/۱۲±۵/۴۵	۳۷/۱۵±۶/۸۰	۲۷/۲±۵۲/۳۷	A_D
۴۶/۲۳±۱/۳۱	۴۱/۵۵±۲/۷	۳۸/۵۹±۳/۲۵	۴۰/۶۲±۱/۶۴	۳۸/۷۲±۱/۰۹	۴۰/۰۴±۲/۸۵	۳۷/۵۹±۲/۴۴	۲۵/۳±۳۵/۴۱	W_C
۵۷/۷۰±۴/۸۵	۵۶/۹۵±۲/۷۳	۴۸/۶۹±۳/۲۲	۵۴/۳۹±۶/۳۴	۴۸/۵۸±۴/۸۱	۵۲/۵۴±۳/۵۴	۴۸/۱۴±۴/۱۴	۳۶/۵۵±۴/۱۰	W_A
۸۰/۰۰±۸/۵۵	۹۴/۸۲±۲/۸۹	۸۸/۳۲±۳/۵۵	۹۰/۳۶±۳/۱۷	۸۸/۶۹±۷/۴۹	۹۲/۷۷±۱/۰۰	۸۶/۵۰±۱/۶۹	۶۰/۵۱±۲/۶۶	W_D
۴۵/۳۲±۳/۷۸	۴۰/۸۴±۴/۶۲	۳۷/۸۹±۱/۱۳	۳۹/۶۵±۳/۱۳	۳۷/۵۲±۱/۳۷	۳۹/۹۷±۱/۵۶	۳۷/۸۱±۱/۰۵	۲۵/۵۰±۵/۳۷	D_C
۵۸/۰۵±۸/۱۳	۵۷/۲۰±۱۰/۱۱	۴۸/۵۱±۳/۵۹	۵۲/۵۹±۸/۳۰	۴۶/۸۹±۷/۳۶	۵۲/۵۵±۷/۱۳	۴۷/۹۷±۶/۹۲	۳۷/۰۷±۷/۱۲	D_A
۸۴/۵۷±۲/۲۳	۹۳/۹۶±۰/۵۱	۸۸/۰۴±۵/۰۳	۹۱/۴۳±۰/۹۸	۸۷/۱۳±۷/۹۳	۹۰/۶۸±۰/۶۲	۸۴/۱۵±۰/۹۲	۶۴/۷۲±۹/۶۹	D_W
۷۵/۷۵±۴/۰۷	۷۱/۱۹±۲/۵۵	۳۶/۷۳±۱/۵۳	۶۸/۹۳±۳/۲۰	۴۹/۹۷±۲/۲۸	۳۶/۵۶±۱/۴۸	۴۷/۷۱±۲/۲۵	۵۲/۴۶±۲/۲۸	متوسط صحت

خطای پیش‌بینی مدل روی داده‌های منبع را حداقل ساخته و سازگاری بین ساختار منیفلد و مدل پیش‌بینی را حداکثر می‌سازد. نمودارهای مربوط به مقایسه صحت طبقه‌بند بر روی ۳۴ پایگاه داده با روش‌های SADA، TJM، JDA و VDA در شکل‌های ۹ و ۱۰ نشان داده شده است. به دلیل اینکه الگوریتم SADA، یک الگوریتم تکرار شونده برای تخمین دقیق برج سبب داده‌های دامنه هدف است از ۱۰ تکرار برای به دست آوردن نتایج این الگوریتم استفاده شده است.

روش VDA علاوه بر تطبیق توزیع‌های حاشیه‌ای و شرطی از یک خوشه‌بندی مستقل از دامنه برای بهبود صحت طبقه‌بند استفاده می‌کند. با این حال، به دلیل ویژگی‌های متفاوت داده‌های آموزشی و تست در دامنه‌های منبع و هدف، صحت بالایی در پیش‌بینی داده‌های دامنه هدف را دارا نمی‌باشد. متوسط بهبود صحت روش SADA به تفکیک نسبت به روش VDA در پایگاه داده اعداد ۱۸/۶۷٪ و در پایگاه داده پای ۳/۶۳٪ است. در واقع SADA با استفاده از یک طبقه‌بند انطباقی

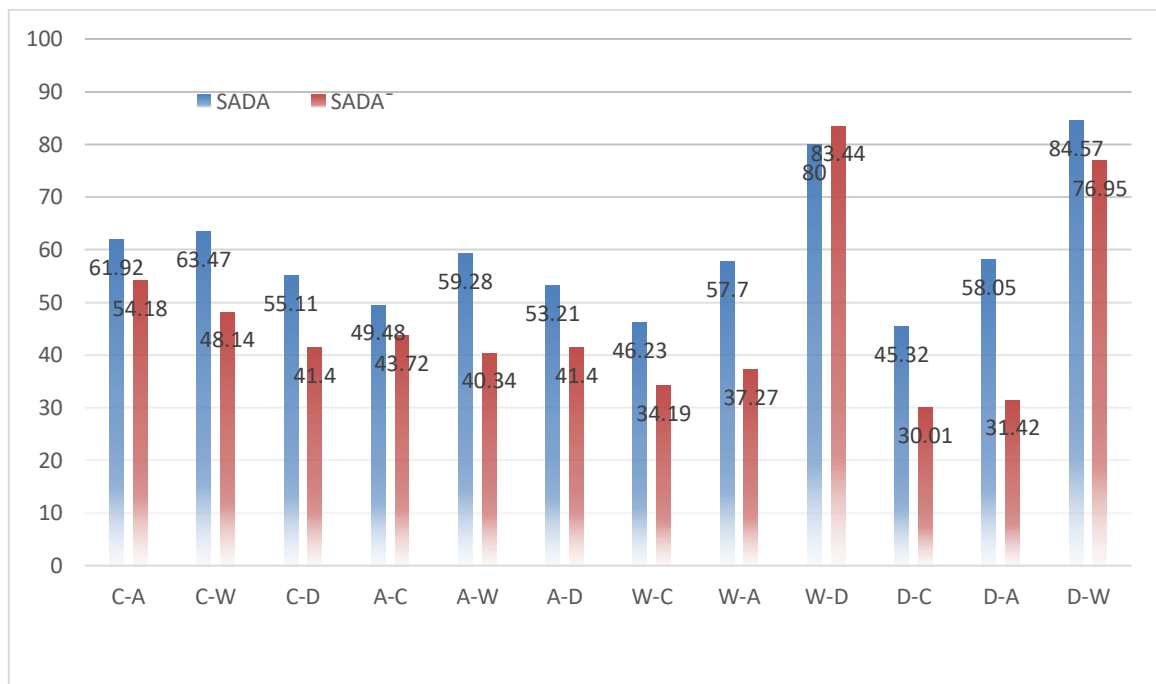


نمونه‌های انتخاب‌نشده از دامنه منبع کالتک

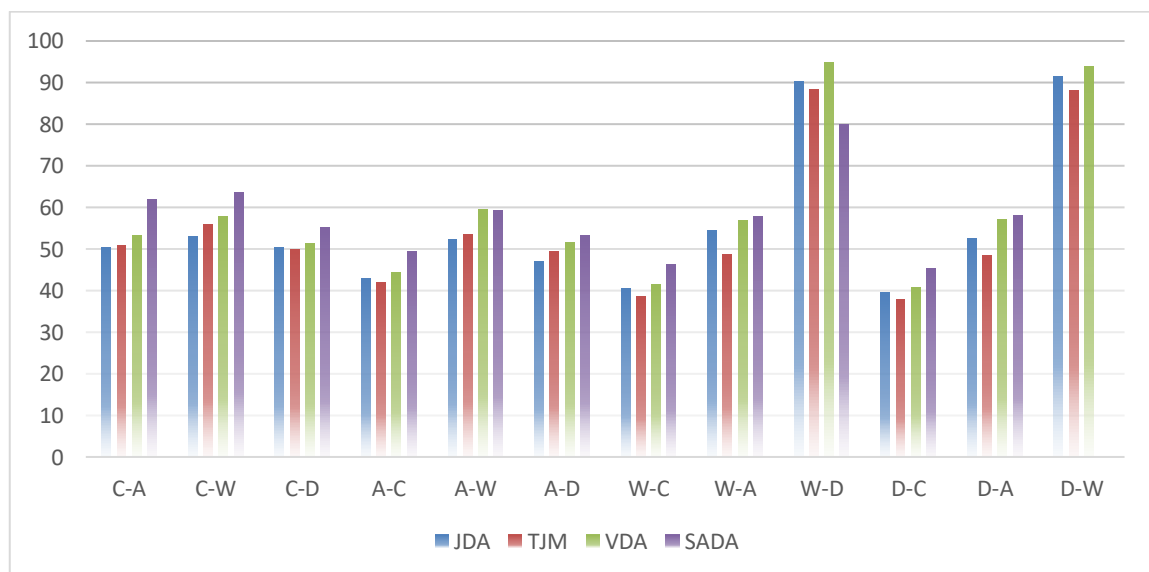
نمونه‌های انتخاب‌شده از دامنه منبع کالتک

نمونه‌ها از دامنه هدف آمازون

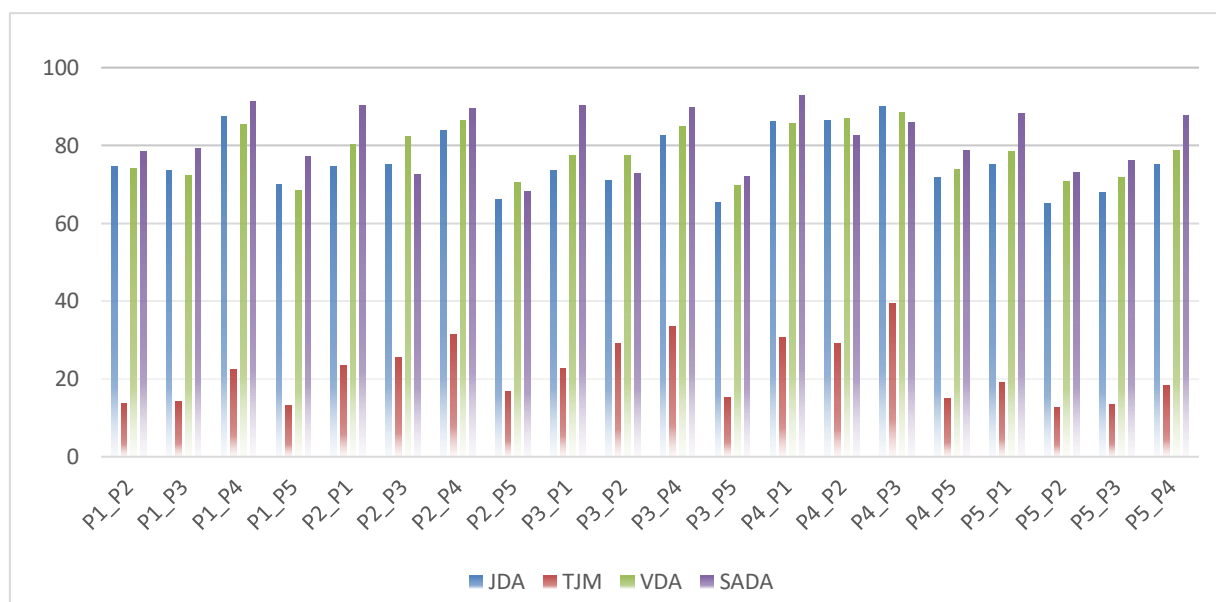
شکل ۷: نمونه‌های انتخاب‌شده از دامنه منبع کالتک برای دامنه هدف آمازون



شکل ۸: مقایسه پایگاه داده آفیس و کالتک در دو حالت نیمه نظارت‌شده (SADA) و بدون نظارت (SADA⁻)



شکل ۹: صحت طبقه‌بندی در پایگاه داده‌های آفیس و کالتک و اعداد با استفاده از روش‌های JDA، TJM، VDA و FMM (نمایش بهتر به صورت رنگی)



شکل ۱۰: صحت طبقه‌بندی در پایگاه داده پای با استفاده از روش‌های JDA، TJM، VDA و FMM (نمایش بهتر به صورت رنگی)

۶- نتیجه‌گیری

در این مقاله، یک روش نوین جهت حل مسئله شیفیت در دامنه‌ها، با عنوان کدگذاری تنک و طبقه‌بندی انطباقی برای تطبیق دامنه‌های بصری (SADA) پیشنهاد شد. SADA یک روش نیمه‌نظارت‌شده می‌باشد که در دو مرحله تلاش می‌کند اختلاف توزیع بین داده‌های آموزشی و تست را کاهش دهد. درواقع SADA مدلی ایجاد می‌کند که با به حداقل رساندن ناسازگاری داده‌ها، مقاومت و تحمل‌پذیری بیشتری در برابر تغییرات توزیع داده‌ها بین دامنه‌های منبع و هدف داشته باشد. SADA در گام اول با استفاده از نرم $l_{2,1}$ نقش نمونه‌هایی از داده‌های منبع را که عامل ایجاد اختلاف بین دو دامنه هستند را کم‌رنگ‌تر می‌کند و با استفاده از MMD اختلاف توزیع حاشیه‌ای دو دامنه را به حداقل می‌رساند. علاوه بر این، SADA با تطبیق پارامترهای مدل به کمک یک طبقه‌بند انطباقی مقاومت و تحمل‌پذیری آن در مقابل تغییرات داده‌ها را به حداکثر می‌رساند.

برای ارزیابی عملکرد SADA، آزمایش‌های مختلفی بر روی انواع پایگاه‌داده‌های واقعی ترتیب داده شده است و نتایج ارزیابی‌ها نشان از برتری قابل‌ملاحظه SADA نسبت به دیگر الگوریتم‌های تطبیق دامنه دارد. برای ادامه کار، برنامه‌ریزی جهت گسترش SADA برای کاهش هم‌زمان توزیع اختلاف‌های حاشیه‌ای و شرطی انجام شده است.

مراجع

- [3] B. Okutmustur, "Reproducing kernel Hilbert spaces", 2005.
- [4] 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.
- [5] طاهره زارع بیدکی و محمدتقی صادقی، «بهینه‌سازی وزن‌ها در کرنل مرکب برای طبقه‌بند مبتنی بر نمایش تنک کرنلی»، مجله مهندسی برق دانشگاه تبریز، جلد ۴۷، شماره ۳، صفحات ۱۰۵۹-۱۰۷۲، ۱۳۹۶.
- [6] 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.
- [7] L. Bruzzone and M. Marconcini, "Domain adaptation problems: a DASVM classification technique and a circular validation strategy", IEEE Trans Pattern Anal Mach Intell, vol. 32, no. 5, pp. 770-787, 2010.
- [8] B. Gong, K. Grauman and F. Sha, "Connecting the dots with landmarks: Discriminatively learning domain-invariant features for unsupervised domain adaptation", Proceedings of the International Conference on Machine Learning, vol. 28, no. 1, pp. 222-230, 2013.
- [9] M. Long, J. Wang, G. Ding, J. Sun and P. S. Yu, "Transfer joint matching for unsupervised domain adaptation", IEEE conference on computer vision and pattern recognition, pp. 1410-1417, 2014.
- [10] J. Tahmoresnezhad and S. Hashemi, "Visual domain adaptation via transfer feature learning", KnowlInf Syst, vol. 50, no. 2, pp. 585-605, 2016.
- [11] M. Long, J. Wang, G. Ding, S. J. Pan and P. Yu, "Adaptation regularization: a general framework for transfer learning", IEEE Trans. Knowl. Data Eng, vol. 26, pp. 1076-1089, 2013.
- [12] Jolliffe I, Principal component analysis, Wiley, vol. 2, pp. 433-459, 2002.
- [13] 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.
- [14] G. Griffin, A. Holub and P. Perona, "Caltech-256 object category dataset", Technical Report 7694, 2007.

- [1] 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.
- [2] J. Tahmoresnezhad and S. Hashemi, "A generalized kernel-based random k-sample sets method for transfer learning", Iran J Sci Technol Trans Electrical Eng, vol. 39, pp. 193-207, 2015.

IEEE international conference on computer vision, pp. 2200-2207, 2013.

[19] مهرداد حیدری ارجلو، سید قدرت اله سیف السادات و مرتضی رزاز، «یک روش هوشمند تشخیص جزیره در شبکه توزیع دارای تولیدات پراکنده مبتنی بر تبدیل موجک و نزدیک‌ترین k -همسایگی (kNN)»، مجله مهندسی برق دانشگاه تبریز، جلد ۴۳، شماره ۱، صفحات ۱۵-۲۶، ۱۳۹۲.

[15] J. J. Hull, "A database for handwritten text recognition research", IEEE Trans. Pattern Anal. Mach. Intell, vol. 16, no. 5, pp. 550-554, 1994.

[16] Y. LeCun, L. Bottou, Y. Bengio and P. Haffner, "Gradient-based learning applied to document recognition", Proc. IEEE, vol. 86, no. 11, pp. 2278-2324, 1998.

[17] T. Sim, S. Baker and M. Bsat, "The CMU pose, illumination, and expression (PIE) database", Proceedings of Fifth IEEE International Conference on Automatic Face Gesture Recognition, pp. 53-58, 2002.

[18] M. Long, J. Wang, G. Ding, J. Sun and S. YuPhilip, "Transfer feature learning with joint distribution adaptation",

زیرنویس‌ها

¹⁹ Adaptation regularization based transfer learning

²⁰ Domain

²¹ Task

²² Frobenius norm

²³ Sparsity regularization parameter

²⁴ Reconstruction error

²⁵ Principal component analysis

²⁶ Office

²⁷ Caltech

²⁸ PIE

²⁹ Amazon

³⁰ Webcam

³¹ Digital single-lens reflex

³² Matlab

³³ Nearest neighbor

³⁴ Transfer component analysis

³⁵ Transfer feature learning with Joint Distribution Adaptation

¹ Domains shift

² Transfer learning

³ Domains adaptation

⁴ Maximum mean discrepancy (MMD)

⁵ Reproducing kernel Hilbert space

⁶ Element mean

⁷ Sparse coding

⁸ Unsupervised

⁹ image processing via Sparse coding and ADaptive classification

¹⁰ Adaptive classifier

¹¹ $l_{2,1}$ norm

¹² Transfer Joint Matching

¹³ Sample selection sparse coding

¹⁴ Geodesic flow kernel

¹⁵ Grassman manifold

¹⁶ Geodesic

¹⁷ Visual Domain Adaptation via transfer feature learning

¹⁸ Domain invariant clustering