



پردازش تصویر بین دامنه‌ای با استفاده از تحلیل تفکیک خطی و تطبیق دامنه مبتنی بر نمونه

مژده زندی‌فر و جعفر طهمورث‌نژاد*

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

چکیده

پردازش تصویر روشی برای اعمال برخی عملیات‌ها بر روی یک تصویر است به طوری که با استفاده از آن، تصاویری با کیفیت بالاتر به دست آمده یا برخی اطلاعات مفید از تصویر استخراج می‌شود. الگوریتم‌های سنتی پردازش تصویر در شرایطی که تصاویر آموزشی (دامنه منبع) که برای یاددهی مدل استفاده می‌شوند، توزیع متفاوتی از تصاویر آزمایش (دامنه هدف) داشته باشند، نمی‌توانند عملکرد خوبی داشته باشند. با این حال، بسیاری از برنامه‌های کاربردی دنیای واقعی به علت کمبود داده‌های برچسب‌دار آموزشی دارای محدودیت هستند؛ از این رو از داده‌های برچسب‌دار دامنه‌های دیگر استفاده می‌کنند. به این ترتیب به خاطر اختلاف توزیع بین دامنه‌های منبع و هدف، طبقه‌بند یادگرفته شده بر اساس مجموعه آموزشی بر روی داده‌های آزمایشی عملکرد ضعیفی خواهد داشت. یادگیری انتقالی و انطباق دامنه، با به کارگیری مجموعه داده‌های موجود دو راه حل برجسته برای مقابله با این چالش هستند، و حتی با وجود اختلاف توزیع قابل ملاحظه بین دامنه‌ها می‌توانند دانش را از دامنه‌های مرتبط به دامنه هدف انتقال دهند. فرض اصلی در مسئله تغییر دامنه این است که توزیع حاشیه‌ای یا توزیع شرطی داده‌های منبع و هدف متفاوت باشد. تطبیق دامنه به طور صریح با استفاده از معیار فاصله از پیش تعیین شده تفاوت در توزیع حاشیه‌ای، توزیع شرطی یا هر دو توزیع را کاهش می‌دهد. در این مقاله، ما به یک سناریوی چالش برانگیز می‌پردازیم که در آن تصاویر دامنه‌های منبع و هدف در توزیع‌های حاشیه‌ای متفاوت بوده و تصاویر هدف دارای برچسب نیستند. بیش تر روش‌های قبلی دو استراتژی یادگیری تطابق ویژگی‌ها و وزن‌دهی مجدد نمونه‌ها را به طور مستقل برای تطبیق دامنه‌ها مورد بررسی قرار داده‌اند. در این مقاله، ما نشان می‌دهیم زمانی که تفاوت دامنه‌ها به طور قابل توجهی بزرگ باشد، هر دو استراتژی مهم و اجتناب‌ناپذیر هستند. روش پیشنهادی ما تحت عنوان تطبیق دامنه مبتنی بر نمونه برای طبقه‌بندی تصاویر (DAIC)، یک فرایند کاهش بُعد بوده که با کاهش اختلاف توزیع تصاویر آموزشی و آزمایشی و به کارگیری هم‌زمان تطابق ویژگی‌ها و وزن‌دهی مجدد کارایی مدل را افزایش می‌دهد. ما با گسترش واگرایی برگمن غیرخطی برای اندازه‌گیری تفاوت توزیع حاشیه‌ای و اعمال آن به الگوریتم کاهش بُعد آنالیز تفکیک خطی فشر، از آن برای ساخت یک نمایش ویژگی مؤثر و قوی برای تفاوت‌های توزیع قابل ملاحظه بین دامنه‌ها استفاده می‌کنیم؛ همچنین، DAIC از مزیت برچسب‌گذاری اولیه برای داده‌های هدف به صورت تکرار شونده برای هم‌گرایی مدل استفاده می‌کند. آزمایش‌های گسترده ما نشان می‌دهد که DAIC به طور قابل توجهی بهتر از الگوریتم‌های یادگیری ماشین پایه و دیگر روش‌های یادگیری انتقالی در نه مجموعه داده بصری تحت سناریوهای مختلف عمل می‌کند.

واژگان کلیدی: پردازش تصویر، یادگیری انتقالی، واگرایی برگمن، کاهش اختلاف توزیع حاشیه‌ای، کاهش ابعاد

Sample-oriented Domain Adaptation for Image Classification

Mojdeh Zandifar & Jafar Tahmoresnezhad*

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

Abstract

Image processing is a method to perform some operations on an image, in order to get an enhanced image or to extract some useful information from it. The conventional image processing algorithms cannot perform

* Corresponding author

* نویسنده عهده‌دار مکاتبات

سال ۱۳۹۸ شماره ۳ پیاپی ۴۱

• تاریخ ارسال مقاله: ۱۳۹۶/۱۲/۲۲ • تاریخ پذیرش: ۱۳۹۸/۴/۱۲ • تاریخ انتشار: ۱۳۹۸/۱۰/۰۷ • نوع مطالعه: پژوهشی

فصلی



well in scenarios where the training images (source domain) that are used to learn the model have a different distribution with test images (target domain). Also, many real world applications suffer from a limited number of training labeled data and therefore benefit from the related available labeled datasets to train the model. In this way, since there is the distribution difference across the source and target domains (domain shift problem), the learned classifier on the training set might perform poorly on the test set. Transfer learning and domain adaptation are two outstanding solutions to tackle this challenge by employing available datasets, even with significant difference in distribution and properties, to transfer the knowledge from a related domain to the target domain. The main assumption in domain shift problem is that the marginal or the conditional distribution of the source and the target data is different. Distribution adaptation explicitly minimizes predefined distance measures to reduce the difference in the marginal distribution, conditional distribution, or both. In this paper, we address a challenging scenario in which the source and target domains are different in marginal distributions, and the target images have no labeled data. Most prior works have explored two following learning strategies independently for adapting domains: feature matching and instance reweighting. In the instance reweighting approach, samples in the source data are weighted individually so that the distribution of the weighted source data is aligned to that of the target data. Then, a classifier is trained on the weighted source data. This approach can effectively eliminate unrelated source samples to the target data, but it would reduce the number of samples in adapted source data, which results in an increase in generalization errors of the trained classifier. Conversely, the feature-transform approach creates a feature map such that distributions of both datasets are aligned while both datasets are well distributed in the transformed feature space. In this paper, we show that both strategies are important and inevitable when the domain difference is substantially large. Our proposed using sample-oriented Domain Adaptation for Image Classification (DAIC) aims to reduce the domain difference by jointly matching the features and reweighting the instances across images in a principled dimensionality reduction procedure, and construct new feature representation that is invariant to both the distribution difference and the irrelevant instances. We extend the nonlinear Bregman divergence to measure the difference in marginal, and integrate it with Fisher's linear discriminant analysis (FLDA) to construct feature representation that is effective and robust for substantial distribution difference. DAIC benefits pseudo labels of target data in an iterative manner to converge the model. We consider three types of cross-domain image classification data, which are widely used to evaluate the visual domain adaptation algorithms: object (Office+Caltech-256), face (PIE) and digit (USPS, MNIST). We use all three datasets prepared by and construct 34 cross-domain problems. The Office-Caltech-256 dataset is a benchmark dataset for cross-domain object recognition tasks, which contains 10 overlapping categories from following four domains: Amazon (A), Webcam (W), DSLR (D) and Caltech256 (C). Therefore $4 \times 3 = 12$ cross domain adaptation tasks are constructed, namely $A \rightarrow W, \dots, C \rightarrow D$. USPS (U) and MNIST (M) datasets are widely used in computer vision and pattern recognition tasks. We conduct two handwriting recognition tasks, i.e., usps-mnist and mnist-usps. PIE is a benchmark dataset for face detection task and has 41,368 face images of size 3232 from 68 individuals. The images were taken by 13 synchronized cameras and 21 flashes, under varying poses, illuminations, and expressions. PIE dataset consists five subsets depending on the different poses as follows: PIE1 (C05, left pose), PIE2 (C07, upward pose), PIE3 (C09, downward pose), PIE4 (C27, frontal pose), PIE5 (C29, right pose). Thus, we can construct 20 cross domain problems, i.e., $P1 \rightarrow P2, P1 \rightarrow P3, \dots, P5 \rightarrow P4$. We compare our proposed DAIC with two baseline machine learning methods, i.e., NN, Fisher linear discriminant analysis (FLDA) and nine state-of-the-art domain adaptation methods for image classification problems (TSL, DAM, TJM, FIDOS and LRSR). Due to these methods are considered as dimensionality reduction approaches, we train a classifier on the labeled training data (e.g., NN classifier), and then apply it on test data to predict the labels of the unlabeled target data. DAIC efficiently preserves and utilizes the specific information among the samples from different domains. The obtained results indicate that DAIC outperforms several state-of-the-art adaptation methods even if the distribution difference is substantially large.

Keywords: Image processing, Transfer learning, Bregman divergence, Marginal distribution difference reduction, Dimensionality reduction

کم داده‌های برچسب‌دار^۱ موجود، برای آموزش محدود می‌شود. از طرف دیگر هزینه برچسب‌زدن داده‌ها با توجه به دخالت داشتن انسان در این زمینه، مقرون به‌صرفه نبوده و از لحاظ زمانی وقت‌گیر است؛ درحالی‌که جمع‌آوری اطلاعات تصویری بدون برچسب به‌مراتب ساده‌تر بوده و درواقع داده‌های بدون برچسب به‌طور مداوم در حال تولید شدن هستند.

^۱ Label

۱- مقدمه

پردازش تصویر یکی از مسائل مهم در هوش مصنوعی و یادگیری ماشین است، که هدف اصلی آن طراحی و ساخت سامانه‌های خبره مبتنی بر هوش مصنوعی و بینایی ماشین است. الگوریتم‌های پردازش تصویر موجود، منجر به توسعه روش‌های قوی در زمینه‌های یادگیری ماشین و بینایی رایانه‌ای شده‌اند. اگرچه این الگوریتم‌ها تاکنون به‌طور قابل توجهی پیشرفت داشته‌اند، اما عملکرد آن‌ها با توجه به تعداد

توزیع داده‌های آموزشی و آزمایشی با یکدیگر متفاوت است. در انطباق دامنه، به جای فرض یکسان بودن توزیع داده‌ها در دو مجموعه آموزشی و آزمایشی، فرض می‌شود که دارای توزیع متفاوت ولی مرتبط هستند؛ سپس، سعی می‌شود تا با توجه به ارتباط بین داده‌های آموزشی و آزمایشی، طبقه‌بندی آموزش داده شود که بر روی داده‌های آزمایشی از دقت قابل قبولی برخوردار باشد. تطبیق دامنه در انواع مسائل بازشناسی الگو^{۱۱}، پردازش زبان طبیعی^{۱۲}، پردازش گفتار^{۱۳}، مکان‌یابی WI-FI^{۱۴}، زمینه بینایی ماشین در تشخیص چهره، آشکارسازی اشیاء، تشخیص عابر پیاده^{۱۵} و بازشناسی رویداد^{۱۶} کاربرد دارد. مسأله تطبیق دامنه تاکنون در حوزه پردازش زبان طبیعی بسیار موفق بوده، اما تلاش‌هایی که برای این مسأله در حوزه بینایی رایانه‌ای شده، محدود بوده و قابل گسترش است [3].

مهم‌ترین چالش در انتقال مفاهیم یادگرفته‌شده جلوگیری از بروز انتقال منفی^{۱۷} است، به این معنی که استفاده از روش‌های تطبیق دامنه نه تنها مفید نبوده بلکه باعث کاهش کیفیت فرآیند یادگیری شده است. مهم‌ترین کاری که باید انجام شود، این است که از انتقال منفی پیش‌گیری شود. چالش دیگری که در تطبیق دامنه وجود دارد، اندازه‌گیری اختلاف توزیع (تفاوت) بین دامنه‌های منبع و هدف است [4]؛ زیرا با توجه به این تفاوت می‌توان از انتقال منفی جلوگیری کرد. بنابراین معیار اندازه‌گیری اختلاف توزیع دامنه‌ها با یکدیگر بسیار مهم است؛ یعنی با چه معیاری فاصله بین دامنه‌ها و میزان تفاوت آن‌ها با یکدیگر را می‌توان اندازه گرفت. اختلاف توزیع بین دامنه‌های منبع و هدف، شامل اختلاف در توزیع حاشیه‌ای^{۱۸} و شرطی^{۱۹} است. برای مثال می‌توان گفت اگر مسأله یک مسأله دسته‌بندی اسناد باشد، اختلاف توزیع حاشیه‌ای متناظر با حالتی است که دو مجموعه اسناد منبع و هدف روی موضوع‌های^{۲۰} متفاوتی متمرکز شده باشند. همچنین برای اختلاف توزیع شرطی می‌توان حالتی را در نظر گرفت که اسناد دو دامنه منبع و هدف از لحاظ دسته‌های تعریف‌شده توسط کاربران بسیار نامتوازن^{۲۱} باشند. در سامانه‌های یادگیری ماشین، وجود اختلاف توزیع حاشیه‌ای و

الگوریتم‌های یادگیری موجود اغلب با این پیش‌فرض عمل می‌کنند که داده‌های آموزشی^۱ و داده‌های آزمایشی^۲ هر دو از یک توزیع^۳ و فضای ویژگی^۴ یکسانی هستند [1]. با این حال در عمل و در بسیاری از برنامه‌های کاربردی دنیای واقعی، این فرض همواره نقض می‌شود و عملکرد این روش‌ها برای داده‌هایی از دامنه‌های متفاوت کاهش می‌یابد که این امر ناشی از وجود تفاوت در توزیع داده‌های آموزشی استفاده‌شده برای آموزش طبقه‌بند با توزیع داده‌های آزمایشی است. تفاوت در توزیع نمونه‌های آموزشی و آزمایشی می‌تواند ناشی از شرایط تهیه و یا گردآوری آن‌ها باشد. تفاوت در روشنایی، زاویه دید، دقت، تصویر زمینه و شرایط عکس‌برداری می‌تواند منجر به تغییر در توزیع نمونه‌های دو مجموعه داده شود. به عنوان مثال مسأله تشخیص اشیاء^۵ یکی از مسائل پرکاربرد حوزه بینایی ماشین و پردازش تصویر به‌شمار می‌رود. امروزه با رشد روزافزون سامانه‌های بینایی ماشین در کاربردهای روزمره، تشخیص اشیاء بیش از پیش مورد توجه قرار گرفته است. به عبارتی، تفاوت در توزیع نمونه‌های آموزشی و آزمایشی (برای نمونه تفاوت در روشنایی، زاویه دید، دقت، تصویر زمینه و شرایط عکس‌برداری تصاویر) می‌تواند منجر به افت کارایی مدل‌های تشخیص اشیاء در پردازش تصویر شود. برای نمونه اگر یک مدل با استفاده از تصاویر سی تی اسکن^۶ و ایکس-ری^۷ آموزش دیده باشد و از آن جهت تشخیص اعضای بدن در یک تصویر ام آر آی^۸ استفاده شود، به‌طور طبیعی در شرایط عادی مدل کارایی قابل قبولی نخواهد داشت [2]؛ در چنین شرایطی اگر بتوان از دانش به‌دست‌آمده از داده‌های آموزشی (که دارای توزیع یا فضای ویژگی متفاوتی با داده‌های آزمایشی است) به‌نحوی برای داده‌های آزمایشی استفاده کرد، می‌توان عملکرد مدل ایجادشده را در فضای آزمون بهبود بخشید. راه حلی که در سال‌های اخیر مورد توجه قرار گرفته، انطباق دامنه است.

انطباق دامنه^۹ یکی از زمینه‌های یادگیری انتقالی^{۱۰} است و یک مسأله مهم در یادگیری ماشین محسوب می‌شود که هدف آن یادگیری طبقه‌بند مناسبی است که بتواند دقت قابل قبولی بر روی داده‌های آزمایشی داشته باشد، در حالی که

¹² Natural-language processing(NLP)

¹³ Speech processing

¹⁴ WI-FI localization

¹⁵ Pedestrian detection

¹⁶ Event recognition

¹⁷ Negative transfer

¹⁸ Marginal distribution

¹⁹ Conditional distribution

²⁰ Topic

²¹ Unbalanced

¹ Training data

² Testing data

³ Distribution

⁴ Feature space

⁵ Object detection

⁶ CT scan

⁷ X-Ray

⁸ MRI

⁹ Domain adaptation

¹⁰ Transfer learning

¹¹ Pattern recognition

شرطی بین دامنه‌های منبع و هدف موجب می‌شود، مدل طبقه‌بند ایجادشده در دامنه منبع، دقت پایینی در پیش‌بینی برچسب‌های نمونه‌های دامنه هدف داشته باشد [4].

بیش‌تر روش‌های یادگیری انتقالی موجود، از فاصله اقلیدسی^۱ یا روش بیشینه اختلاف میانگین‌ها^۲ (MMD) برای محاسبه اختلاف توزیع بین نمونه‌های دامنه منبع و هدف استفاده می‌کنند. این دو متریک دارای محدودیت‌هایی هستند که ممکن است، روی انتقال دانش تأثیر منفی بگذارند. به‌عنوان مثال توابع هدفی که با استفاده از فاصله اقلیدسی تعریف می‌شوند، ممکن است به‌خوبی بهینه نشوند و دچار مشکل زیربهینه^۳ شوند. درواقع فاصله اقلیدسی در فضای داده‌های ورودی، لزوماً نمی‌تواند به‌خوبی عدم شباهت داده‌ها را نشان دهد، زیرا ممکن است برخی از ویژگی‌ها نسبت به بقیه اهمیت بیشتری داشته باشند. علاوه‌براین روش بیشینه اختلاف میانگین‌ها نمی‌تواند برای مسائل رگرسیون مناسب باشد؛ چون ممکن است، داده‌های برچسب‌دار از یک جامعه آماری بزرگ برآورد شده باشند. همچنین این دو متریک نمی‌توانند اطلاعات تفکیک‌کننده اولیه (به‌عنوان مثال اطلاعات هندسی و اطلاعات آماری داده‌ها) به‌دست‌آمده از دامنه منبع را به دامنه هدف انتقال دهند. در این پژوهش برای اجتناب از این‌گونه مسائل از واگرایی برگمن^۴ برای سنجش و کاهش اختلاف توزیع بین دامنه‌های منبع و هدف استفاده می‌شود. واگرایی برگمن [5] یک واگرایی جامع است که تعداد زیادی از فواصل معروف و پرکاربرد از جمله تابع هزینه جمع مربعات^۵، واگرایی کولبک - لایبلر^۶، تابع هزینه منطقی^۷، فاصله ماهالانوبیس^۸ را شامل می‌شود. واگرایی برگمن با کاهش تفاوت بین توزیع نمونه‌های آموزشی و آزمایشی می‌تواند دانش تفکیک‌کننده به‌دست‌آمده از نمونه‌های آموزشی را به نمونه‌های آزمایش انتقال دهد. همچنین واگرایی برگمن، اطلاعات تفکیک‌کننده به‌دست‌آمده از نمونه‌های آموزشی را برای تفکیک‌پذیری هرچه بهتر رده‌های دامنه هدف، حفظ می‌کند.

در مسائل یادگیری انتقالی می‌توان از یک یا چندین دامنه منبع دانش موردنیاز را به دامنه هدف انتقال داد که در این پژوهش فرض شده است برای هر مسئله تنها یک دامنه منبع D_s و یک دامنه هدف D_t وجود دارد. مسائل تطبیق دامنه می‌توانند در دو حالت کلی زیر مورد بررسی قرار بگیرند:

(۱) تطبیق دامنه نیمه‌نظارت شده^۹: که در آن، دامنه هدف شامل تعداد کمی داده برچسب‌دار است، یعنی اگر تعداد داده‌های دامنه منبع n_s باشد، دامنه منبع به‌صورت $D_s = D_t = T_u \cup T_l$ و دامنه هدف به‌صورت $D_t = T_u \cup T_l$ تعریف می‌شوند. اگر تعداد داده‌های برچسب‌دار و بدون برچسب دامنه هدف به‌ترتیب n_{tl} و n_{tu} فرض شود، $T_l = \{(x_i, y_i)\}_{i=1}^{n_{tl}}$ مجموعه داده‌های برچسب‌دار دامنه هدف و $T_u = \{(x_i)\}_{i=1}^{n_{tu}}$ مجموعه داده‌های بدون برچسب دامنه هدف هستند. همچنین در این مسائل فرض می‌شود تعداد زیادی داده بدون برچسب در دامنه هدف موجود است؛ به‌طوری‌که $n_{tl} \ll n_{tu}$ و علاوه‌براین فرض می‌شود داده‌های برچسب‌دار دامنه هدف به‌تنهایی برای ساخت یک طبقه‌بند مناسب، مفید نیستند.

(۲) تطبیق دامنه بدون نظارت^{۱۰}: که در آن، هیچ یک از داده‌های دامنه هدف، برچسب ندارند، یعنی اگر تعداد داده‌های دامنه‌های منبع و هدف به‌ترتیب n_s و n_t باشد، دامنه‌های منبع و هدف به‌ترتیب به‌صورت $D_s = \{(x_i, y_i)\}_{i=1}^{n_s}$ و $D_t = \{(x_i)\}_{i=1}^{n_t}$ تعریف می‌شوند [6]. در این پژوهش، مسأله تطبیق دامنه بدون نظارت مورد بررسی قرار گرفته است، به‌عبارتی فرض شده است، فضای ویژگی دامنه‌های منبع و هدف یکسان هستند ($X_s = X_t$) و فقط توزیع آن‌ها با یکدیگر متفاوت است. هدف یادگیری یک تابع پیش‌بینی f_t برای دامنه هدف، با استفاده از داده‌های برچسب‌دار دامنه منبع است؛ به‌طوری‌که تابع پیش‌بینی حاصل، عملکرد مطلوبی روی دامنه هدف داشته باشد. در روش پیشنهادی برای ایجاد تطبیق بین دامنه‌ها از روش‌های تطبیق خصوصیات استفاده می‌شود. روش‌های تطبیق خصوصیات، الگوریتم‌هایی هستند که می‌توانند نمایش خصیصه‌ای نمونه‌ها در فضای اصلی یا در فضای پنهان را تغییر دهند؛ این نوع الگوریتم‌ها در بخش‌های بعدی به‌تفصیل توضیح داده می‌شوند.

روش پیشنهادی این مقاله با عنوان DAIC^{۱۱}، یک روش بدون نظارت با بهره‌گیری از تطبیق خصوصیات و وزن‌دهی مجدد نمونه‌ها است. DAIC با استفاده از الگوریتم FLDA^{۱۲} و واگرایی برگمن یک نمایش جدید از داده‌ها ایجاد می‌کند تا اختلاف توزیع حاشیه‌ای دامنه‌ها را کاهش داده و به‌صورت هم‌زمان از روش وزن‌دهی مجدد برای انتخاب

⁸ Mahalanobis distance

⁹ Semi-supervised

¹⁰ Unsupervised

¹¹ Sample-oriented Domain Adaptation for Image Classification (DAIC)

¹² Fisher's linear discriminant analysis

¹ Euclidean distance

² Maximum mean discrepancy (MMD)

³ Suboptimal

⁴ Bregman divergence

⁵ Squared loss

⁶ Kullback-Leibler (KL)

⁷ Logistic loss

۲- کارهای پیشین

رویکردهای یادگیری انتقالی را می‌توان به سه دسته کلی تقسیم کرد:

- رویکردهای مبتنی بر انتقال پارامتر
- رویکردهای مبتنی بر نمونه
- رویکردهای مبتنی بر ویژگی

رویکردهای مبتنی بر انتقال پارامتر فرض می‌کنند که دو وظیفه^۲ در دامنه منبع و هدف دارای تعدادی پارامتر یکسان یا توزیع پیشین^۳ با ابر پارامترهای^۴ یکسان است. بدین ترتیب با پیدا کردن این پارامترها یا توزیع‌های پیشین مشترک می‌توان یادگیری را به دامنه جدید انتقال داد. از کارهای انجام‌شده در این دسته می‌توان به ARTL^۵ [7] و JACRL [31] اشاره کرد. ARTL یک روش تنظیم تطبیقی است که برای یادگیری یک طبقه‌بند انطباقی با کاهش خطای طبقه‌بند در دامنه منبع، افزایش انطباق هندسی دامنه‌ها در فضای جدید و ایجاد تطبیق در توزیع مشترک بین دامنه‌ها یک مدل مقاوم در برابر تغییرات ایجاد می‌کند. در این روش هدف ساخت یک رده‌بند انطباقی است؛ به‌طوری‌که این معیارها محقق شوند: (۱) اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌های منبع و هدف کاهش یابد، (۲) ریسک ساختاری مدل در داده‌های برچسب‌دار دامنه منبع کمینه شود، و (۳) نرخ سازگاری مدل و ساختار هندسی داده‌ها پیشنهاد شود.

JACRL یک رویکرد دومرحله‌ای مبتنی بر مدل برای حل مشکل اختلاف توزیع بین دامنه‌ها است. در مرحله نخست، JACRL با بهره‌گیری از تطبیق خصوصیات، داده‌های منبع و هدف براساس تجزیه و تحلیل اجزای اصلی (PCA) به یک زیرفضای مشترک با بُعد کم نگاشت می‌شوند؛ سپس در این فضا با استفاده از معیار بیشینه اختلاف میانگین (MMD)، اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌ها کاهش می‌یابد. در مرحله دوم، یک رده‌بند انطباقی بر روی داده‌های منبع و هدف به‌منظور کمینه‌کردن ریسک تجربی تابع پیش‌بینی در داده‌های برچسب‌دار دامنه منبع و همچنین بیشینه‌کردن نرخ سازگاری تابع پیش‌بینی و ساختار هندسی داده‌ها، ایجاد می‌شود.

رویکردهای مبتنی بر نمونه براساس وزن‌دهی مجدد نمونه‌ها و یا انتخاب نمونه‌هایی که تفاوت بین توزیع دامنه‌های منبع و هدف را کمینه می‌سازند، عمل می‌کنند. روش‌های مبتنی بر نمونه با تخصیص وزن به نمونه‌ها، تابع ضرر^۶ را روی

نمونه‌هایی از دامنه منبع که بیشترین شباهت را به نمونه‌های دامنه هدف دارند (برای ایجاد تطبیق‌پذیری بیشتر بین دامنه‌ها) استفاده کند.

اهداف اصلی رویکرد پیشنهادی بدین صورت خلاصه شده است: (۱) حفظ ساختار اصلی داده‌ها: DAIC به کمک فرآیند FLDA، داده‌های دامنه منبع و هدف را به یک زیرفضای ویژگی کم‌بعد نگاشت می‌کند به‌طوری‌که ساختار اصلی داده‌ها حفظ شود، (۲) تطبیق توزیع حاشیه‌ای: DAIC با استفاده از واگرایی برگمن اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف را در زیرفضای نگاشت‌شده کاهش می‌دهد، و (۳) انتخاب نمونه: با کاهش اختلاف توزیع حاشیه‌ای بین نمونه‌ها، همچنان ممکن است، نمونه‌هایی در دامنه منبع وجود داشته باشند که مرتبط با نمونه‌های هدف نباشند. برای حل این مشکل DAIC با کمینه‌کردن نرم $l_{2,1}$ بر روی داده‌های دامنه منبع، نمونه‌هایی از دامنه منبع را که مرتبط با دامنه هدف است، انتخاب می‌کند. درنهایت، تمامی این اهداف با یکدیگر ادغام می‌شوند و یک نمایش ویژگی جدید به دست می‌آید؛ به‌طوری‌که در مقابل اختلاف توزیع بین دامنه‌ها مقاوم باشد.

در این مقاله یک چارچوب جدید برای تطبیق دامنه‌های بصری پیشنهاد شده است که در آن تطبیق خصوصیات و وزن‌دهی به نمونه‌ها به‌طور هم‌زمان بهینه می‌شوند. DAIC با استفاده از تطبیق خصوصیات و واگرایی برگمن یک زیرفضای جدید به‌دست می‌آورد که در آن خصوصیات مشترک دو دامنه منبع و هدف حفظ می‌شوند و سپس با استفاده از ماتریس وزن‌دهی به نمونه‌ها و به‌کارگیری نرم $l_{2,1}$ نمونه‌هایی را از دامنه منبع که مرتبط و نزدیک به دامنه هدف هستند، انتخاب می‌کند؛ به‌طوری‌که اختلاف توزیع حاشیه‌ای بین دامنه‌ها به کمینه می‌رسد و زمینه برای اجرای الگوریتم‌های کلاسیک یادگیری ماشین فراهم می‌شود.

ساختار ادامه مقاله بدین صورت است: در بخش ۲ پژوهش‌های مرتبط مورد بررسی قرار می‌گیرند. در بخش ۳ روش پیشنهادی آورده شده و بخش ۴ شامل پایگاه‌داده‌های مورد ارزیابی و فرضیات آزمایش‌ها است. در بخش ۵ نتایج حاصل از عملکرد الگوریتم پیشنهادی و مقایسه آن با دیگر روش‌های موجود قرار دارد. و درنهایت در بخش ۶ نتیجه‌گیری و کارهای پیشنهادی آینده ارائه شده است.

⁴ Hyperparameters

⁵ Adaptation regularization for transfer learning

⁶ Loss function

¹ $l_{2,1}$ norm

² Task

³ Prior Distribution

توزیع داده‌ها کمینه می‌کنند. وزن‌دهی مجدد نمونه^۱ و نمونه‌برداری اهمیت^۲ دو روش مهم مبتنی بر این رویکرد هستند. از جمله این روش‌ها می‌توان به SSSC^۳ [8] و LSSA [30] اشاره کرد.

روش SSSC ابتدا نمونه‌های غیر مرتبط دامنه منبع را حذف کرده، سپس با یافتن یک ماتریس انتخاب، اختلاف توزیع بین داده‌های منبع و هدف را کاهش می‌دهد. SSSC از کدگذاری تُنک^۴ برای انتخاب نمونه‌های دامنه منبع استفاده کرده و هر نمونه دامنه هدف را با کمینه تعداد نمونه دامنه منبع نمایش می‌دهد.

LSSA یک رویکرد جدید دومرحله‌ای برای تطبیق دامنه بدون نظارت است که از مفهوم لندمارک استفاده می‌کند. در مرحله نخست لندمارک‌ها شناسایی می‌شوند؛ برای این کار هر نمونه چه از دامنه منبع باشد و چه از دامنه هدف به‌عنوان یک نامزد لندمارک در نظر گرفته می‌شود؛ سپس یک معیار کیفی برای هر نامزد محاسبه می‌شود که اگر مقدار این معیار بزرگ‌تر از مقدار آستانه تعریف شده باشد، می‌توان نتیجه گرفت که این نامزد لندمارک است. در مرحله دوم با استفاده از یک کرنل گوسی و با توجه به لندمارک‌های انتخاب‌شده در مرحله قبل، تمام نقاط دامنه منبع و دامنه هدف به یک فضای مشترک تصویر می‌شوند.

بیشتر روش‌های موجود در زمینه تطبیق دامنه بدون نظارت روش‌های مبتنی بر ویژگی هستند که ایده کلی این روش‌ها تغییر ویژگی‌های اولیه و به‌دست‌آوردن یک نمایش جدید برای داده‌ها است، با این هدف که ویژگی‌های جدید، بازنمایی^۵ بهتری برای خصوصیات مشترک دامنه‌های منبع و هدف ارائه دهند. درواقع روش‌های مبتنی بر ویژگی بر دو فرض کلی استوارند: (۱) برخی ویژگی‌های مختص یک دامنه^۶ و برخی مستقل از دامنه^۷ هستند. (۲) از فضای ویژگی اصلی به فضای ویژگی نهان^۸ که توسط دامنه منبع و دامنه هدف به اشتراک گذاشته شده، یک نگاشت وجود دارد.

در این شاخه از روش‌های انطباقی، فرض می‌شود، دامنه‌های منبع و هدف در یک فضای ویژگی جدید دارای توزیع حاشیه‌ای و یا توزیع شرطی مشابه هستند. هدف اصلی در این شاخه از الگوریتم‌ها، یافتن تبدیلی است؛ به‌نحوی که نمایش فضای ویژگی را عوض کرده و داده‌ها را به فضایی برود که در آن تشابه توزیع دامنه منبع و هدف وجود داشته باشد.

رویکردهای کاهش بُعد از جمله روش‌های یادگیری بازنمایی جدید از داده‌ها هستند که اختلاف توزیع بین دامنه‌های منبع و هدف را کاهش می‌دهند تا خطای دسته‌بندی برای مدل کاهش یابد؛ به‌طور کلی، رویکردهای کاهش بُعد از دو چارچوب اصلی پیروی می‌کنند: (۱) الگوریتم‌هایی که براساس روش PCA هستند و داده‌ها را به یک فضای کم‌بُعد نگاشت می‌کنند؛ به‌طوری‌که واریانس داده‌های نگاشت‌شده بیشینه شود. از جمله این الگوریتم‌ها می‌توان به VDA [9]، JDA [10] و CLGA [22] اشاره کرد که از PCA برای نگاشت داده‌ها به یک زیرفضای کم‌بُعد استفاده می‌کنند. CLGA یک رویکرد تطبیق دامنه بدون نظارت چندمنبعی است که به بررسی نحوه بهره‌گیری مؤثر و هم‌زمان از اطلاعات رده‌ها و دامنه‌ها می‌پردازد. CLGA به‌منظور افزایش افزایش انطباق دامنه‌ها، ابتدا چند دامنه را به‌عنوان یک دامنه در نظر می‌گیرد، سپس هم‌زمان اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌های منبع و هدف را به کمینه می‌رساند (تطبیق‌پذیری کلی). علاوه‌براین، CLGA به‌منظور افزایش قابلیت تفکیک‌پذیری، ارتباط بین رده‌ها و دامنه‌های مجزا را بررسی و از ساختار هندسی دامنه‌ها و رده‌ها استفاده می‌کند. درواقع، CLGA به‌منظور افزایش قابلیت تفکیک‌پذیری، فاصله بین نمونه‌های متعلق به یک رده اما دامنه متفاوت را کاهش می‌دهد، و فاصله بین نمونه‌های متعلق به یک دامنه اما رده متفاوت را افزایش می‌دهد (تطبیق‌پذیری محلی). (۲) الگوریتم‌هایی که براساس روش FLDA عمل می‌کنند. این الگوریتم‌ها در تلاش هستند داده‌ها را به‌گونه‌ای به فضای کم‌بُعد نگاشت کنند که رده‌های مختلف تفکیک‌پذیری بیشتری نسبت به هم داشته باشند. از جمله این الگوریتم‌ها می‌توان به SCA [11] اشاره کرد که این روش با استفاده از FLDA یک نمایش جدید از نمونه‌های دامنه منبع و هدف ایجاد می‌کند که در آن فاصله بین رده‌های مختلف (تفکیک‌پذیری) بیشینه شده و اختلاف توزیع بین دامنه‌ها کمینه شود. گفتنی است که PCA بر روی خود داده‌ها عمل می‌کند در حالی که FLDA براساس پیدا کردن روابط بین رده‌های متفاوت داده، انجام می‌پذیرد.

دراین مقاله، یک روش تطبیق دامنه بدون نظارت برای حل مسأله یادگیری انتقالی ارائه شده است. درواقع یک چارچوب کلی با ترکیبی از دو روش مبتنی بر خصوصیت و مبتنی بر نمونه پیشنهاد می‌شود. در این روش داده‌های دامنه

⁵ Representation

⁶ Domain-specific

⁷ Domain-invariant

⁸ Latent space

¹ Instance Reweighting

² Importance Sampling

³ Sample selection sparse coding

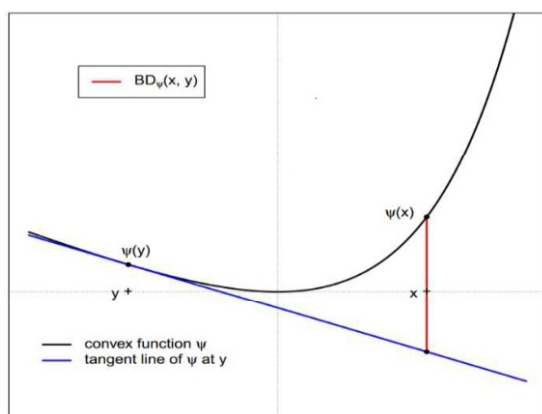
⁴ Sparse coding

[32]، [33]، [23]. بیش‌تر روش‌های یادگیری انتقالی از فاصله اقلیدسی برای اندازه‌گیری تفاوت بین نمونه‌ها در دامنه منبع و هدف استفاده می‌کنند [34]، [35]، [36]، [37]. با این حال، در بسیاری از کاربردهای دنیای واقعی، ممکن است فاصله اقلیدسی نتواند شباهت‌های ذاتی یا ناهم‌انگهی‌های بین نمونه‌ها را انتقال دهد. واگرایی برگمن^۱ یکی از معیارهای اندازه‌گیری فاصله بین دو توزیع احتمال است که دامنه گسترده‌ای از فاصله‌ها از جمله تابع هزینه جمع مربعات، واگرایی کولیک - لایبلر، تابع هزینه منطقی، فاصله مایلانویس را در بر می‌گیرد. واگرایی برگمن با کاهش اختلاف بین توزیع نمونه‌های آموزشی و آزمایشی می‌تواند دانش تفکیک‌کننده به‌دست‌آمده از نمونه‌های آموزشی را به نمونه‌های آزمایش انتقال دهد.

تعریف: فرض کنید $\Omega \rightarrow \mathbb{R}^m$: یک تابع محذب در مجموعه محذب $\Omega \subseteq \mathbb{R}^m$ است که فرض می‌شود ناتهی و تمایزپذیر است. بنابراین برای $x, y \in \mathbb{R}^m$ واگرایی برگمن با توجه به ψ به‌صورت زیر تعریف می‌شود:

$$BD_{\psi}(x, y) = \psi(x) - \psi(y) - \langle x - y, \nabla \psi(y) \rangle \quad (1)$$

که در آن $\nabla \psi(y)$ نشان‌دهنده گرادیان ψ در نقطه y و $\langle \dots \rangle$ ضرب داخلی است. واگرایی برگمن تعریف‌شده در (۱) را می‌توان تفاوت بین مقدار تابع محذب در x و نخستین مرتبه بسط تیلور^۲ آن در نقطه y ، تفسیر کرد، به‌طور معادل، بخش باقی‌مانده از نخستین مرتبه بسط تیلور ψ در y دانست. مفهوم هندسی واگرایی برگمن در شکل (۱) نشان داده شده است. براساس شکل (۱) واضح است که واگرایی برگمن فاصله بین مقدار تابع محذب در x و تانژانت آن در y است.



(شکل-۱): تصویر هندسی واگرایی برگمن
(Figure-1): Geometrical illustration of Bregman divergence

¹ Bregman divergence

² Taylor expansion

منبع و هدف توسط ترکیبی از روش FLDA و متریک واگرایی برگمن، به یک زیرفضای جدید که در آن اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف کمینه‌شده، نگاشت می‌شوند. همچنین به‌طور هم‌زمان نمونه‌هایی از دامنه منبع که شباهت بیشتری با نمونه‌های دامنه هدف دارند، انتخاب و مدل براساس این نمونه‌های انتخاب‌شده ایجاد می‌شود.

۳- روش پیشنهادی

در این بخش، جزئیات الگوریتم پیشنهادی DAIC برای حل مسائل تطبیق دامنه بدون نظارت شرح داده می‌شود.

۳-۱- تعریف مسأله

مسأله شیفت در دامنه‌ها زمانی رخ می‌دهد که دامنه‌های آموزشی و آزمایشی توزیع متفاوتی داشته باشند؛ بنابراین طبقه‌بند ساخته‌شده روی داده‌های دامنه آموزشی صحت خوبی روی داده‌های دامنه آزمایش نخواهد داشت.

با توجه به دامنه منبع برچسب‌دار $D_S = \{(x_{s1}, y_{s1}), \dots, (x_{sn}, y_{sn})\}$ و یک دامنه هدف بدون برچسب $D_t = \{x_{t1}, \dots, x_{tm}\}$ تحت شرایط $X_s = X_t$ ، $Y_s = Y_t$ و $P_s(x_s) \neq P_t(x_t)$ دو دامنه از توزیع حاشیه‌ای متفاوتی برخوردار هستند، هدف ما در این مقاله این است که مدلی طراحی کنیم که توزیع حاشیه‌ای داده‌های آموزشی و آزمایشی را تقریباً برابر کند؛ یعنی $P_s(x_s) \sim P_t(x_t)$ DAIC اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف را با یادگیری یک نمایش جدید کاهش می‌دهد. درواقع DAIC با استفاده توأم از روش کاهش بعد FLDA و وزن‌دهی به نمونه‌های دامنه منبع دقت مدل ایجادشده را برای نمونه‌های دامنه هدف به‌طور قابل ملاحظه‌ای افزایش می‌دهد.

در نمایش جدید، به‌صورت هم‌زمان اختلاف توزیع حاشیه‌ای بین دامنه‌ها کمینه شده و از روش وزن‌دهی به نمونه‌های دامنه منبع، برای بهبود عملکرد مدل طبقه‌بند در پیش‌بینی برچسب‌ها استفاده شده است.

۳-۲- واگرایی برگمن

بیشتر مطالعات قبلی در حوزه یادگیری انتقالی یا براساس ارزش‌گذاری مجدد نمونه‌های دامنه منبع مطابق با توزیع دامنه هدف [24]، [25]، [26] و یا براساس یافتن یک نمایه جدید برای نمونه‌های دامنه منبع و هدف، به منظور کاهش اختلاف توزیع بین دامنه‌های منبع و هدف عمل می‌کنند

درواقع واگرایی برگمن با توجه به نوع تابع محدب Ψ انتخاب شده، تابع ضرر را کاهش می دهد. به عنوان مثال اگر $\Psi(x) = x^T W x$ و یک ماتریس قطعی مثبت^۱ باشد، واگرایی برگمن به صورت رابطه (۲) تعریف می شود:

$$BD_{\Psi}(x, y) = x^T W x - y^T W y - \langle x - y, 2W y \rangle = (x - y)^T W (x - y) \quad (2)$$

که W معکوس ماتریس کوواریانس است. بنابراین باتوجه به رابطه (۲) به این نوع واگرایی برگمن فاصله مالهالانوبیس بین x و y گویند. در صورتی که W یک ماتریس همانی باشد، واگرایی برگمن به مربع فاصله اقلیدسی بین x و y تبدیل، که به صورت زیر نوشته می شود:

$$BD_{\Psi}(x, y) = \|x - y\|^2 \quad (3)$$

به طوری که $\|x\| = \sqrt{x^T x}$ تعریف می شود. بنابراین با توجه به رابطه (۳)، واگرایی برگمن که فاصله بین P_t و P_s را بر طبق فاصله مربع اقلیدسی محاسبه می کند، به صورت زیر تعریف می شود:

$$D_W(P_s \| P_t) = \frac{1}{n_s^2} \sum_{s=1}^{n_s} \sum_{t=1}^{n_s} G_{\Sigma_{11}}(\bar{y}_s - \bar{y}_t) + \frac{1}{n_t^2} \sum_{s=n_s+1}^{n_s+n_t} \sum_{t=n_t+1}^{n_s+n_t} G_{\Sigma_{22}}(\bar{y}_s - \bar{y}_t) - \frac{1}{n_s n_t} \sum_{s=1}^{n_s} \sum_{t=n_t+1}^{n_s+n_t} G_{\Sigma_{12}}(\bar{y}_s - \bar{y}_t) \quad (4)$$

برای تخمین مقادیر P_t و P_s از روش برآورد تراکم هسته^۲ استفاده شده است که تراکم را در $\mathbf{y} \in \mathbf{R}^d$ به عنوان مجموع کرنل های بین $\mathbf{y}$ و هر نمونه $\bar{\mathbf{y}}_i$ تخمین می زند؛ به عنوان مثال $p(\bar{\mathbf{y}}) = \left(\frac{1}{n}\right) G_{\Sigma}(\bar{\mathbf{y}} - \bar{\mathbf{y}}_i)$ که در آن n تعداد نمونه ها و $G_{\Sigma}(\bar{\mathbf{y}})$ کرنل گوسی $\mathbf{d}$ بُعدی با ماتریس کوواریانس Σ است. که در این رابطه $\Sigma = \Sigma_1 + \Sigma_2$ ، $\Sigma_{11} = \Sigma_1 + \Sigma_2$ و $\Sigma_{12} = \Sigma_1 + \Sigma_2$ است.

DAIC از واگرایی برگمن جهت اندازه گیری و کاهش اختلاف توزیع بین دامنه های منبع و هدف استفاده می کند.

۳-۳- کاهش بُعد براساس الگوریتم جداکننده

فیشر خطی

استفاده از تمام ویژگی های داده در بیشتر موارد بسیار پرهزینه و وقت گیر است، از طرفی همه این ویژگی ها ارزش

اطلاعاتی خاصی ندارند؛ لذا نیاز به روش هایی جهت بهینه سازی و کاهش ابعاد داده ها قبل از ورود به سامانه های طبقه بند احساس می شود. این مسأله در بسیاری از کاربردها (مانند طبقه بندی) اهمیت به سزایی دارد؛ زیرا در این کاربردها تعداد زیادی ویژگی وجود دارد، که بسیاری از آنها یا بلااستفاده هستند و یا اینکه بار اطلاعاتی چندانی ندارند. حذف نکردن این ویژگی ها مشکلی از لحاظ اطلاعاتی ایجاد نمی کند؛ ولی بار محاسباتی را برای کاربرد مورد نظر بالا می برد؛ و علاوه بر این باعث می شود که اطلاعات غیر مفید زیادی را به همراه داده های مفید ذخیره کنیم [38]. الگوریتم تحلیل جداکننده فیشر خطی (FLDA) [12] یکی از روش های معروف و پرکاربرد کاهش بُعد خطی است. روش FLDA داده های برداری را به فضایی با ابعاد کم تر نگاشت می کند؛ به طوری که داده های نگاشت شده دارای بیشترین پراکندگی بین رده ای و کم ترین پراکندگی درون رده ای هستند. درواقع FLDA به دنبال پیدا کردن ماتریس نگاشت (W) است که نسبت دترمینان S_b (ماتریس پراکندگی بین رده ای) به S_w (ماتریس پراکندگی درون رده ای) بیشینه شود، در اصطلاح به این ماتریس، مقیاس فیشر می گویند که به صورت زیر تعریف می شود:

$$W = \underset{W}{\operatorname{argmax}} \frac{|W^T S_b W|}{|W^T S_w W|} \quad (5)$$

دترمینان ماتریس کواریانس میزان پراکندگی یک رده را بیان می کند. دترمینان تنها ضرب مقادیر قطر اصلی بوده که مقادیر مستقل واریانس ها هستند. دترمینان تحت هر نگاشت متعامد نرمالی مقدار برابر دارد؛ بنابراین مقیاس فیشر، سعی می کند تا واریانس میانگین رده ها بیشینه و واریانس داده های هر رده کمینه شوند. S_b ماتریس پراکندگی بین رده ای و S_w ماتریس پراکندگی درون رده ای داده های اصلی بوده و بدین صورت تعریف می شوند:

$$S_b = \sum_{i=1}^C n_i (\bar{\mathbf{m}}^{(i)} - \bar{\mathbf{m}})(\bar{\mathbf{m}}^{(i)} - \bar{\mathbf{m}})^T \quad (6)$$

$$S_w = \sum_{i=1}^C \sum_{j=1}^{n_i} (x_j^{(i)} - \bar{\mathbf{m}}^{(i)})(x_j^{(i)} - \bar{\mathbf{m}}^{(i)})^T \quad (7)$$

در اینجا $\bar{\mathbf{m}}$ و $\bar{\mathbf{m}}^{(i)}$ به ترتیب نشان دهنده میانگین کل داده ها و میانگین داده های رده i ام هستند. روش جداکننده خطی به دنبال یافتن ماتریس $W \in \mathbf{R}^{D \times d}$ است که رابطه (۱) را بیشینه کند. در این صورت هر داده $\mathbf{x} \in \mathbf{R}^D$ به صورت $\mathbf{y} = W^T \mathbf{x}$ به یک فضای کم بُعد d به صورت خطی منتقل می شود.

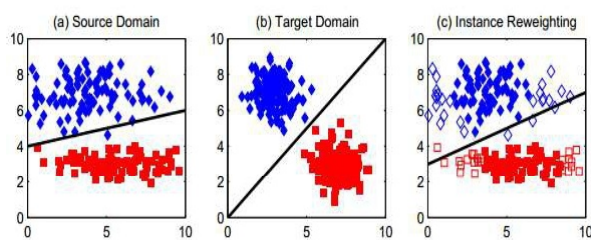
¹ Positive definite matrix

² Kernel density estimation (KDE)

فضای نگاشت شده W محاسبه می‌کند و λ پارامتر تنظیم است که بین $F(W)$ و $D_W(P_S \parallel P_T)$ تنظیم ایجاد می‌کند.

۳-۵- وزن دهی به نمونه‌های دامنه منبع

با این حال، تطبیق توزیع ویژگی‌های داده براساس کمینه‌سازی واگرایی برگمن در معادله (۱۰) به‌تنهایی برای مسائل شیف‌ت دامنه‌ها مناسب نیست و به‌طور کامل نمی‌تواند اختلاف توزیع بین دامنه‌های منبع و هدف را کاهش دهد. درواقع، اگر اختلاف توزیع بین دامنه‌های منبع و هدف زیاد باشد، حتی در زیرفضای نگاشت‌شده نیز همواره نمونه‌هایی از دامنه منبع وجود دارند که با نمونه‌های دامنه هدف مرتبط نیستند و این باعث کاهش تطبیق‌پذیری مدل می‌شود. شکل (۲) به‌خوبی این مشکل را نشان می‌دهد. در این شرایط، تطبیق ویژگی و وزن دهی مجدد نمونه‌ها به‌طور مشترک برای ایجاد سازگاری و تطبیق دامنه بسیار مهم و اجتناب ناپذیر است.



(شکل-۲): (a) دامنه منبع پس از تطابق ویژگی؛ (b) دامنه هدف پس از تطابق ویژگی با توجه به نمونه‌های منبع نامرتبط (نشان داده شده در c، اختلاف دامنه پس از تطابق ویژگی‌ها هنوز بزرگ است. (c) دامنه منبع پس از تطبیق ویژگی و وزن دهی مجدد به نمونه‌ها. در حال حاضر نمونه‌های نامناسب دامنه منبع دارای وزن کم‌تری هستند تا اختلاف دامنه کاهش یابد (b نسبت به c).

(Figure-2): (a) source domain after feature matching; (b) target domain after feature matching. Due to the irrelevant source instances (shown as unfilled markers in c), the domain difference is still large after feature matching. (c) source domain after joint feature matching and instance reweighting. The irrelevant source instances are now down-weighted to further reduce domain difference (b vs c).

بنابراین، با ترکیب یک روش کاهش بعد و یک روش وزن دهی مجدد، نمونه‌هایی از دامنه منبع را در زیرفضای مشترک انتخاب می‌کنیم که تطبیق‌پذیری مدل را افزایش دهد. درواقع با استفاده از دو ماتریس انتخاب A_S (ماتریس تبدیل متناظر با نمونه‌های دامنه منبع) و A_T (ماتریس تبدیل متناظر با نمونه‌های دامنه هدف) و به‌کارگیری نرم $\ell_{2,1}$ و نرم فروبنیوس $\|\cdot\|_F^2$ ، نمونه‌هایی از دامنه منبع که مرتبط و نزدیک

³ Frobenius norm

۳-۴- الگوریتم جداکننده فیشر خطی با

استفاده از متریک واگرایی برگمن

چنان‌که در بخش قبلی دیده شد در روش جداکننده فیشر خطی نیاز به حل مسأله بیشینه‌کردن تریس^۱ به‌صورت (۵) به‌دفعات زیاد است:

$$\arg \max_{W^T W} \frac{\text{tr}(W^T S_b W)}{\text{tr}(W^T S_w W)} \quad (۸)$$

روش متداول برای بیشینه‌کردن نسبت تریس یادشده با استفاده از بردارهای ویژه حاصل از مسأله مقادیر ویژه تعمیم یافته است؛ پس درواقع تابع هدف الگوریتم FLDA برابر است با

$$F(W) = \text{tr}^{-1}(W^T S_b W) \text{tr}(W^T S_w W) \quad (۹)$$

که $\text{tr}^{-1}(x)$ معکوس $\text{tr}(x)$ است.

نکته مهم در مورد روش FLDA این است که در مسائلی مانند تشخیص الگو و طبقه‌بندی تصویر که با ابعاد بزرگ داده مواجه هستیم، ممکن است ماتریس نگاشت W روی داده‌های آموزشی دچار زیادآموزی شود. همچنین FLDA کلاسیک به‌دلیل اینکه اطلاعات اولیه را در مورد تفاوت بین توزیع نمونه‌های آموزشی و آزمایشی درنظر نمی‌گیرد، به‌تنهایی نمی‌تواند پاسخ‌گوی مسأله شیف‌ت دامنه‌ها باشد؛ با این حال برای اجتناب و یا کاهش مشکل زیادآموزی، ما اطلاعات اولیه نمونه‌های آموزشی را با اضافه‌کردن یک بخش تنظیم^۲ منطقی به رابطه (۹) درنظر می‌گیریم. بدین ترتیب، روش پیشنهادی DAIC یک زیرفضای مطلوب پیدا می‌کند که در آن اختلاف توزیع حاشیه‌ای کمینه‌شده و رده داده‌های آموزشی و آزمایشی به‌صورت مستقل از هم جدا می‌شوند. DAIC علاوه بر این که اطلاعات تفکیک‌کننده موجود در نمونه‌های آموزشی را درنظر می‌گیرد، توزیع بایاس بین نمونه‌های آموزشی و آزمایشی را نیز بررسی می‌کند. از این رو واگرایی برگمن و تابع هدف الگوریتم FLDA را ترکیب می‌کنیم؛ پس با توجه به رابطه‌های (۴) و (۸) خواهیم داشت:

$$W = \underset{W \in \mathbb{R}^{D \times d}}{\text{argmin}} F(W) + \lambda D_W(P_S \parallel P_T) \quad (۱۰)$$

$F(W)$ تابع هدف الگوریتم در زیرفضای نگاشت‌شده، $D_W(P_S \parallel P_T)$ واگرایی برگمن که فاصله بین P_S و P_T را در

^۱ Trace

^۲ Regularization

$$W = \operatorname{argmin} \operatorname{tr}^{-1}(W^T S_b W) \operatorname{tr}(W^T S_w W) + \lambda D_W(P_s \parallel P_t) + \tau (\|A_s\|_{2,1} + \|A_t\|_F^2) \quad (12)$$

τ پارامتر تنظیم بین تطابق ویژگی و وزن دهی به نمونه‌ها است، علاوه بر این شرط $W^T W = I$ نیز باید برآورده شود. این شرط نشان دهنده این است که سطرها و یا ستون‌های ماتریس W مستقل از یکدیگر بوده و برهم عمودند و مجموع آن‌ها برابر یک است. با وجود این شرط در رابطه بهینه‌سازی، جواب‌های بدیهی در مجموعه جواب قرار نمی‌گیرند. برای به دست آوردن زیرفضای خطی بهینه W با توجه به رابطه (۱۲) می‌توان از الگوریتم گرادیان نزولی^۳ استفاده کرد:

$$W_{k+1} = W_k - \eta(k) \left(\frac{\partial F(W)}{\partial W} + \lambda \sum_{i=1}^{n_s+n_t} \frac{\partial D_W(P_s \parallel P_t)}{\partial \bar{y}_i} \frac{\partial \bar{y}_i}{\partial W} \right) \quad (13)$$

∂W گرادیان نسبت به W و η نرخ یادگیری در k امین تکرار است، بدین ترتیب با در نظر گرفتن رابطه (۸) مشتق $F(W)$ بر حسب W برابر است با:

$$\frac{\partial F(W)}{\partial W} = 2 \operatorname{tr}^{-1}(W^T S_b W) S_w W - 2 \operatorname{tr}^{-2}(W^T S_b W) \operatorname{tr}(W^T S_w W) S_b W \quad (14)$$

با توجه به رابطه (۴) مشتق $D_W(P_s \parallel P_t)$ را نسبت به W می‌توان به صورت زیر نوشت:

$$\sum_{i=1}^{n_s+n_t} \frac{\partial D_W(P_s \parallel P_t)}{\partial \bar{y}_i} \frac{\partial \bar{y}_i}{\partial W} = \sum_{i=1}^{n_s} \frac{D_W(P_s \parallel P_t)}{\partial \bar{y}_i} \bar{x}_i^T + \sum_{i=n_s+1}^{n_s+n_t} \frac{D_W(P_s \parallel P_t)}{\partial \bar{y}_i} \bar{x}_i^T \quad (15)$$

در این مقاله، هدف ما کاهش اختلاف توزیع بین دامنه‌ها با استفاده هم‌زمان از تطبیق ویژگی‌ها و وزن‌دهی مجدد نمونه‌های دامنه منبع (برای ایجاد تطبیق پذیری بیشتر) است، تا زمینه‌ای مناسب برای اجرای الگوریتم‌های کلاسیک یادگیری ماشین فراهم شود. مزیت اصلی DAIC، بهینه‌سازی هم‌زمان تطبیق ویژگی و وزن‌دهی مجدد نمونه‌ها است؛ به طوری که اختلاف توزیع بین دامنه‌ها کمینه شده و اطلاعات تفکیک‌کننده نمونه‌های دامنه منبع و هدف به خوبی حفظ شود. در ادامه شکل (۳) روش پیشنهادی را با جزئیات بیشتر نشان می‌دهد.

³ Gradient descent

به دامنه هدف هستند انتخاب می‌شوند. درواقع نرم $\ell_{2,1}$ می‌تواند تنظیم‌کننده تنگی^۱ را روی ماتریس A اعمال کرده و ماتریس A را با سطرهای تنگ نشان دهد. از آنجا که هر سطر ماتریس A مربوط به یک نمونه است، سطر تنگی^۲ می‌تواند وزن‌دهی نمونه‌ها را تسهیل کند. رابطه زیر ماتریس انتخاب A را به دست می‌آورد که اختلاف توزیع بین داده‌های دامنه منبع و هدف را در زیرفضای نگاشت شده کمینه می‌سازد.

$$\|A_s\|_{2,1} + \|A_t\|_F^2 \quad (11)$$

با استفاده از نرم $\ell_{2,1}$ می‌توان هر نمونه هدف را با حداقل تعداد نمونه منبع نمایش داد. نرم $\ell_{2,1}$ را فقط روی نمونه‌های دامنه منبع اعمال می‌کنیم؛ زیرا هدف ما اضافه کردن وزن نمونه‌هایی از دامنه منبع است که شباهت و ارتباط بیشتری با نمونه‌های دامنه هدف دارند. با به کمینه‌رساندن رابطه (۱۱)، رابطه (۱۰) نیز کمینه شده و در این صورت در نمایش جدید نمونه‌هایی از دامنه منبع که مرتبط با نمونه‌های دامنه هدف هستند اهمیت بیشتری دارند. با استفاده از این تنظیم‌کننده (رابطه (۱۱))، DAIC نسبت به اختلاف توزیع ناشی از نمونه‌های نامرتبط دامنه‌های منبع و هدف مقاوم می‌شود.

۳-۶- تابع هدف نهایی و بهینه‌سازی

روش DAIC، یک روش ترکیبی از روش‌های مبتنی بر نمونه و مبتنی بر خصوصیت است که برای مسائلی که دارای اختلاف توزیع زیادی هستند، پیشنهاد شده است. در این روش، ابتدا یک نمایش کم‌بعد از دامنه‌های منبع و هدف ایجاد می‌شود که در این نمایش جدید با بهره‌گیری از واگرایی برگمن اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف کاهش می‌یابد؛ سپس به طور هم‌زمان، به نمونه‌هایی از دامنه منبع که دارای کم‌ترین شباهت از نظر توزیع با نمونه‌های دامنه هدف هستند، وزن کم‌تری در ایجاد مدل طبقه‌بند اختصاص داده می‌شود. درواقع مدل طبقه‌بند براساس نمونه‌هایی از دامنه منبع که بیشترین شباهت را به نمونه‌های دامنه هدف دارند، آموزش می‌بیند؛ در این صورت تطبیق‌پذیری داده‌ها با یکدیگر افزایش می‌یابد و مدل مقاوم‌تری در مقابل اختلاف توزیع بین دامنه‌های منبع و هدف ایجاد می‌شود. بدین ترتیب با ترکیب رابطه‌های (۱۰) و (۱۱) به معادله (۱۲) می‌رسیم که تابع هدف نهایی DAIC است:

¹ Sparsity regularizer

² Row-sparsity

۳-۳- پیچیدگی زمانی

در این بخش، پیچیدگی زمانی روش پیشنهادی نسبت به روش‌های مورد مقایسه، مورد ارزیابی قرار می‌گیرد. پیچیدگی محاسباتی الگوریتم DAIC با علامت O بزرگ تحلیل می‌شود. تحدد تابع هدف که در رابطه (۱۰) تعریف می‌شود، به $F(W)$ و $D(W)$ وابسته است. تحدد $F(W)$ وابسته به یک الگوریتم خاص یادگیری زیرفضا مانند FLDA است. درحالی‌که تحدد $D(W)$ به مدل توزیع مجموعه نمونه‌های آموزشی و آزمایشی بستگی دارد. اگر N تعداد کل نمونه‌ها و D واگرایی برگمن باشد رابطه (۱۲) دارای پیچیدگی زمانی $O((D+N)^2)$ است. پس به‌طور کلی پیچیدگی زمانی DAIC برابر با $O((D+N)^2)$ است. اگر d و C به ترتیب اندازه ابعاد فضای اصلی، اندازه ابعاد فضای جدید و تعداد رده‌ها باشند، روش NN (با در نظر گرفتن تعداد همسایگی برابر با یک برای نمونه‌ها)، دارای پیچیدگی زمانی $O(mN)$ است. روش TJM یک روش ترکیبی از روش‌های مبتنی بر نمونه و مبتنی بر خصوصیت است، دارای پیچیدگی $O(m^2d + CN^2)$ است. روش NN جزو سریع‌ترین روش‌ها است؛ اما تطبیق دامنه انجام نمی‌دهد. DAIC در مقایسه با روش TJM دارای پیچیدگی زمانی کمتری بوده، زیرا TJM، دارای یک کرنل به‌نسبه پیچیده و حل تابع هدف آن دشوار و زمان‌بر است.

۴- آزمایش‌ها

در این بخش، آزمایش‌های گسترده‌ای برای مسأله طبقه‌بندی تصویر برای ارزیابی رویکرد DAIC انجام شده است.

۴-۱- معرفی مجموعه داده‌ها

مجموعه داده آفیس و کالت^۱ [13]، اعداد^۲ [14,15] و پای^۳ [16]، یازده مجموعه داده استاندارد هستند که به‌طور گسترده برای ارزیابی الگوریتم‌های دامنه بصری مورد استفاده قرار می‌گیرند.

مجموعه داده اعداد شامل دو دامنه USPS و MNIST است. دامنه USPS حاوی اعداد دست‌نویس پویش‌شده نامه‌های اداره پست آمریکا است که شامل ۷۲۹۱ تصویر آموزشی و ۲۰۰۷ تصویر آزمایشی با اندازه 16×16 پیکسل است. مجموعه داده MNIST حاوی اعداد دست‌نویس پویش‌شده دانش‌آموزان دبیرستانی آمریکا است که شامل شصت‌هزار نمونه آموزشی و ده‌هزار نمونه آزمایشی با اندازه

¹ Office+Caltech-256

² USPS+MNIST

28×28 پیکسل است. این پایگاه داده شامل ده رده مختلف است و به‌منظور آزمایش هر دو دامنه در شرایط یکسان، دامنه USPS_vs_MNIST(U-M) ایجاد شده است که به‌طور تصادفی ۱۸۰۰ نمونه از داده‌های دامنه USPS به‌عنوان داده‌های آموزشی و دوهزار نمونه از داده‌های دامنه MNIST به‌عنوان داده‌های آزمایش به کار گرفته می‌شود. با جابه‌جایی نمونه‌های آموزشی و آزمایشی دامنه MNIST_vs_USPS (M-U) برای یک آزمایش دیگر ایجاد شده است.

مجموعه داده آفیس شامل مجموعه تصاویر از اشیای مختلف است که تصاویر از نظر کیفیت، روشنایی، رنگ و نوع زمینه با هم متفاوت هستند. این مجموعه داده شامل سه دامنه Amazon (تصاویر اشیای دانه‌دیده از سایت‌های تجاری (A))، Webcam (تصاویر اشیای با وضوح پایین که توسط دوربین وب گرفته شده‌اند (W))، DSLR (تصاویر اشیای با وضوح بالا که با استفاده از دوربین‌های دیجیتالی گرفته شده‌اند (D)) و کالتک (C) یک مجموعه داده استاندارد برای تشخیص اشیای است که شامل ۳۰۶۰۷ تصویر و ۲۵۶ رده از تصاویر است. در دامنه‌های مجموعه داده آفیس و کالتک $12 \times 3 \times 4$ آزمایش طراحی شده که از ده رده مشترک بین مجموعه داده آفیس و مجموعه داده کالتک استفاده شده است. در هر یک از آزمایش‌های طراحی شده یکی از مجموعه داده‌ها به‌عنوان دامنه هدف انتخاب می‌شوند؛ به‌عنوان مثال $C \rightarrow A$, $C \rightarrow W$, $(C \rightarrow D, \dots, D \rightarrow W)$.

مجموعه داده پای شامل ۴۱۳۶۸ تصویر از ۶۸ شخص متفاوت با وضوح 32×32 پیکسل است که این تصاویر از لحاظ روشنایی، جهت تصویربرداری و حالت چهره با یکدیگر متفاوت هستند. این مجموعه داده دارای پنج دامنه متفاوت است که هر دامنه نشان‌دهنده یک حالت تصویربرداری است، PIE1 حالت چپ (P1)، PIE2 حالت بالا (P2)، PIE3 حالت پایین (P3)، PIE4 حالت روبه‌رو (P4) و PIE5 حالت راست (P5). در این مجموعه داده، آزمایش بر روی $20 \times 4 \times 5$ مجموعه بین‌دامنه‌ای مختلف طراحی شده است که از بین این پنج دامنه، دو دامنه مختلف به‌عنوان دامنه‌های منبع و هدف انتخاب می‌شوند؛ به‌عنوان مثال $P1 \rightarrow P2$, $P1 \rightarrow P3$, $(\dots, P5 \rightarrow P4)$.

۴-۲- ارزیابی الگوریتم‌ها

در این بخش، عملکرد روش پیشنهادی را با هفت الگوریتم یادگیری ماشین و یادگیری انتقالی بدون نظارت مقایسه

³ PIE

می‌کنیم. روش‌هایی که DAIC با آن‌ها مقایسه شده است، عبارتند از: طبقه‌بند نزدیک‌ترین همسایه (1-NN)، FLDA، LRSR [5]، TSL [17]، DAM [18]، TJM [19]، CDDA [23]، VDA [21]، CLGA [22] و [20].

مهم‌ترین دلیل برتری DAIC نسبت به روش‌های دیگر، استفاده از اطلاعات اولیه دامنه منبع است، که به‌صورت حداکثری به دامنه هدف انتقال می‌یابد. درضمن DAIC اطلاعات خاص داده‌ها (به‌عنوان مثال ساختار منیفولد داده‌ها) را نیز حفظ می‌کند.

(جدول-۱): توضیحات پایگاه داده‌ها

(Table-1): Description of Databases

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
PIE1	3332	1024	68	P1
PIE2	1629	1024	68	P2
PIE3	1632	1024	68	P3
PIE4	3329	1024	68	P4
PIE5	1632	1024	68	P5

اطلاعات جزئی‌تر از پایگاه داده‌ها در جدول (۱) نشان داده شده است.

۳-۴- مفروضات پیاده‌سازی

الگوریتم پیشنهادی DAIC شامل سه پارامتر زیر است که مقادیر بهینه آن‌ها در بخش بعدی آورده شده است: (۱) λ پارامتر تنظیم در رابطه (۱۰)، (۲) η نرخ یادگیری در رابطه (۱۳) و (۳) τ پارامتر تنظیم بین تطبیق خصوصیات و وزن‌دهی به نمونه‌ها است. مقدار بهینه پارامترها در روش DAIC، برای پایگاه داده‌های بصری برای تمامی آزمایش‌ها انجام گرفته و در نهایت از طبقه‌بند استاندارد NN برای پیش‌بینی برچسب‌های داده‌های دامنه هدف استفاده شده است.

۵- نتایج و بحث‌ها

در این بخش، نتایج روش پیشنهادی DAIC و الگوریتم‌های شناخته‌شده حوزه تطبیق دامنه توسط طراحی آزمایش‌های بدون نظارت مورد ارزیابی و مقایسه قرار گرفته است.

۵-۱- ارزیابی نتایج

در این بخش، عملکرد روش پیشنهادی را با هفت الگوریتم یادگیری ماشین و یادگیری انتقالی بدون نظارت مقایسه

می‌کنیم. گفتنی است که FLDA^۳ نشان‌دهنده کاهش اختلاف حاشیه‌ای بین دامنه‌های منبع و هدف با استفاده از الگوریتم کاهش بعد FLDA و واگرایی برگمن بدون وزن‌دهی نمونه‌ها است.

تشخیص شی: دقت طبقه‌بندی DAIC و سایر روش‌ها در مجموعه داده‌های Office + Caltech در جدول (۲) گزارش شده است. به‌منظور تفسیر بهتر، نتایج در شکل (۴) به تصویر کشیده شده است. نتایج حاصل از آزمایش‌ها حاکی از آن است که متوسط بهبود دقت روش DAIC در پایگاه داده آفیس و کالتک ۵۲/۲۱٪ است. علاوه‌براین، بهبود عملکرد DAIC، ۱۶/۰۵٪ بیش از طبقه‌بند NN است، این ثابت می‌کند که DAIC برای طبقه‌بندی تصاویر در دامنه‌های متفاوت می‌تواند مؤثر باشد؛ همچنین متوسط بهبود دقت روش DAIC نسبت به روش FLDA در مجموعه داده آفیس و کالتک ۱۰/۸۳٪ است. DAIC بهبود خوبی (۱/۱۸٪) نسبت به متوسط دقت بهترین الگوریتم مورد مقایسه (یعنی CLGA) دارد.

تشخیص چهره: دقت طبقه‌بندی DAIC و سایر روش‌ها در مجموعه داده‌های PIE در جدول (۳) گزارش شده و به‌منظور تفسیر بهتر، نتایج در شکل (۵) نشان شده است. DAIC دارای (۷/۸۶٪) بهبود از لحاظ متوسط دقت طبقه‌بندی در مقایسه با بهترین الگوریتم مورد مقایسه (یعنی CLGA) است؛ همچنین، DAIC بهبود (۴۵/۲۵٪) را نسبت به NN به دست می‌آورد. DAIC بهبود قابل ملاحظه‌ای (۴۵/۷۱٪) نسبت به متوسط دقت الگوریتم FLDA دارد.

تشخیص اعداد: دقت طبقه‌بندی DAIC و سایر روش‌ها در مجموعه داده‌های USPS+MNIST در جدول (۴) گزارش شده و به‌منظور تفسیر بهتر، نتایج در شکل (۶) آورده شده است. نتایج حاصل از آزمایش‌ها حاکی از متوسط بهبود دقت روش DAIC در مجموعه داده اعداد (۷۵/۷۵٪) است. متوسط بهبود دقت روش DAIC نسبت به روش‌های NN، FLDA و CDDA به ترتیب برابر (۲۰/۴۳٪)، (۱۷/۴۴٪) و (۶/۶۲٪) است. در ادامه روش DAIC به‌تفکیک با تمام روش‌ها مقایسه شده است.

1-NN طبقه‌بند نزدیک‌ترین همسایگی: این الگوریتم بر روی ویژگی‌های اصلی داده‌ها، بدون اعمال هیچ‌گونه انطباقی براساس نزدیک‌ترین همسایگی برچسب داده‌های آزمایش از دامنه هدف را بر طبق داده‌های آموزشی از دامنه منبع پیش‌بینی می‌کند.

دقت (۹/۶۶٪) در پایگاه داده آفیس و کالتک، (۲۷/۹۹٪) در پایگاه داده اعداد و (۳۰/۶۵٪) در پایگاه داده پای است. روش TJM از جمله جدیدترین روش‌های ایجاد تطبیق بین دامنه‌ها توسط ایجاد یک نمایش مشترک بین دامنه‌های منبع و هدف است. روش TJM توسط یک مسأله بهینه‌سازی پیچیده، اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف را کمینه می‌سازد. همچنین TJM به دلیل استفاده از الگوریتم PCA و نادیده گرفتن بخشی از داده ورودی ساختار اصلی داده‌ها را حفظ نمی‌کند و زیرفضای جدید را با بخشی از نمونه‌های دامنه منبع ایجاد می‌کند، درحالی‌که روش DAIC با استفاده از واگرایی برگمن ساختار داده‌های ورودی را حفظ می‌کند. روش DAIC به تفکیک نسبت به روش TJM دارای متوسط بهبود دقت (۷/۸٪) در پایگاه داده آفیس و کالتک، (۱۷/۹۹٪) در پایگاه داده اعداد و (۴۴/۴۸٪) در پایگاه داده پای است.

LRSR با استفاده از یک ماتریس انتقال، نمونه‌های دامنه منبع و هدف را به یک زیرفضای مشترک نگاشت می‌کند که در آن هر نمونه دامنه هدف می‌تواند به طور خطی توسط داده‌های دامنه منبع بازسازی شود. به این ترتیب، اختلاف دامنه‌های منبع و هدف کاهش می‌یابد و با اعمال توأم محدودیت‌های تُنک و سطح پایین روی ماتریس نگاشت ساختار کلی و محلی داده‌ها حفظ می‌شود. مهم‌ترین دلیل برتری روش پیشنهادی نسبت به روش LRSR، کاهش اختلاف توزیع حاشیه‌ای بین دامنه‌ها و انتقال حداکثری اطلاعات تفکیک‌کننده نمونه‌های دامنه منبع برای تفکیک‌پذیری بهتر رده‌های مختلف است. روش DAIC به تفکیک نسبت به روش LRSR دارای متوسط بهبود دقت (۶/۸۲٪) در پایگاه داده آفیس و کالتک، (۱۱/۵۹٪) در پایگاه داده اعداد و (۱۸/۴۸٪) در پایگاه داده پای است.

CDDA از جمله روش‌های شناخته‌شده در حوزه یادگیری انتقالی بدون نظارت است که با ایجاد یک نمایش کم‌بعد از دامنه‌های منبع و هدف، به طور هم‌زمان اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌ها را کاهش می‌دهند. به دلیل خصوصیات متفاوت داده‌های آموزشی و آزمون در دامنه‌های منبع و هدف، طبقه‌بند ایجادشده در نمایش جدید توسط روش CDDA نمی‌تواند با صحت بالایی برچسب داده‌های دامنه هدف را پیش‌بینی کند. دلیل برتری DAIC نسبت به CDDA این است که DAIC علاوه بر اینکه ساختار داده‌ها را حفظ می‌کند، اطلاعات تفکیک‌کننده اولیه‌های ورودی را نیز در نظر می‌گیرد که این باعث کاهش اختلاف

FLDA یک روش انتقال داده‌ها به یک فضای کم‌بعد است که ایده اصلی آن حفظ ساختار رده‌ها برای طبقه‌بندی است. از آنجا که FLDA اختلاف توزیع بین دامنه‌ها را در نظر نمی‌گیرد، نمی‌تواند عملکرد خوبی در برابر مسائل تطبیق دامنه داشته باشد. با این وجود FLDA عملکرد بهتری نسبت به الگوریتم NN دارد.

TSL یکی از روش‌های تطبیق دامنه بدون نظارت است که اختلاف توزیع حاشیه‌ای دامنه منبع و هدف را براساس تخمین تراکم هسته کاهش می‌دهد. TSL به اندازه مجموعه داده‌ها حساس است در صورتی‌که دامنه هدف حاوی داده‌های کم باشد، تخمین تراکم هسته نمی‌تواند توصیف دقیقی از داده‌ها داشته باشد؛ علاوه بر این برای داده‌هایی در مقیاس بالا (مانند مجموعه داده TSL (PIE) دچار مشکل هم‌گرایی می‌شود. DAIC بر روی هر دو مجموعه داده کوچک و بزرگ به خوبی کار می‌کند و دارای بهبود قابل توجه در مقایسه با TSL است. مهم‌ترین علت برتری DAIC نسبت به TSL استفاده از تطبیق خصوصیات برای کاهش اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف است. روش DAIC به تفکیک نسبت به روش TSL دارای متوسط بهبود دقت (۱۵/۱۲٪) در پایگاه داده آفیس و کالتک، (۲۰/۵۸٪) در پایگاه داده اعداد و (۴۲/۵۵٪) در پایگاه داده پای است.

FIDOS یک چارچوب جدید براساس تعمیم الگوریتم FLDA است. FIDOS یک زیرفضای کم‌بعد را ایجاد می‌کند که اختلاف توزیع بین دامنه‌ها را کاهش می‌دهد. FIDOS مشابه DAIC یک رویکرد مبتنی بر FLDA است، اما فقط برای مجموعه داده‌های مرتبط عملکرد خوبی دارد؛ درواقع برای مسائل چند دامنه‌ای به دلیل استخراج ویژگی‌های بیشتر، بهتر عمل می‌کند. این درحالی‌که DAIC برای مسائلی که دارای اختلاف توزیع زیادی هستند، پیشنهاد شده است. روش DAIC به تفکیک نسبت به روش FIDOS دارای متوسط بهبود دقت (۱۲/۳۷٪) در پایگاه داده آفیس و کالتک، و (۴۴/۰۶٪) در پایگاه داده پای است.

DAM، یک روش تطبیق دامنه بدون نظارت است که به طور مستقیم یک طبقه‌بند سازگار با دامنه هدف را به دست می‌آورد. DAM از طبقه‌بندهای پیش‌آمोخته با استفاده از نمونه‌های برچسب‌دار دامنه منبع، استفاده می‌کند. علاوه بر این DAM براساس فرضیه مسطح‌سازی یک شرایط تنظیمی را به وجود می‌آورد که طبقه‌بند دامنه هدف مجبور به تصمیم‌گیری مشابه با طبقه‌بند پیش‌آمोخته می‌شود. روش DAIC به تفکیک نسبت به روش DAM دارای متوسط بهبود

می‌کند. تعداد نمونه‌های انتخاب شده با توجه به τ در جدول (۵) گزارش شده است.

به عنوان مثال هنگامی که $\tau=100$ است در مجموعه داده اعداد هیچ نمونه‌ای از دامنه منبع انتخاب نمی‌شود؛ اما زمانی که τ مقدار کوچکی دارد، ماتریس A تنگی کمتری دارد و بنابراین نمونه‌های منبعی که توزیع نزدیک تری با نمونه‌های هدف دارند، انتخاب می‌شوند.

(جدول-۵): تأثیر پارامتر تنظیم تنگی

τ بر تعداد نمونه‌های انتخاب شده از دامنه منبع

(Table-5): Influence of tensile setting parameter τ on the number of samples selected from the source domain

τ	0.01	0.1	1	10	100
U-M#	890	430	98	39	0
P1-P5#	1800	1800	1800	780	40
C-A#	1068	646	494	73	2

۳-۵- بررسی ضرورت کاهش اختلاف توزیع حاشیه‌ای

در روش پیشنهادی DAIC برای کاهش اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف از روش غیر پارامتری واگرایی برگمن استفاده شده است. واگرایی برگمن قادر است اطلاعات به دست آمده از نمونه‌های دامنه منبع را، با به کمینه رساندن اختلاف توزیع حاشیه‌ای بین دامنه‌ها به نمونه‌های دامنه هدف منتقل کند که باعث تفکیک پذیری بهتر رده‌های مختلف در فضای جدید می‌شود. در صورتی که در تابع هدف (رابطه ۱۲) پارامتر λ را برابر صفر در نظر بگیریم (طبق نتایج گزارش شده در جداول (۲) الی (۴)، فیلد FLDA) دقت مدل به دلیل وجود اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف به شدت کاهش می‌یابد. در شرایطی که دامنه‌های منبع و هدف دارای مجموعه خصوصیات یکسان باشند، اختلاف در احتمال وقوع مقادیر هر کدام از این خصوصیات در هر دامنه، موجب ایجاد اختلاف توزیع حاشیه‌ای بین دامنه‌ها خواهد شد.

روش‌های کاهش بُعد به تنهایی پاسخ‌گوی مسائل شیف‌د دامنه‌ها نیستند و به دلیل نادیده گرفتن اطلاعات اولیه در مورد اختلاف توزیع بین نمونه‌های دامنه منبع و هدف، مدل روی نمونه‌های دامنه منبع دچار زیادآموزی شده و عملکرد مدل به شدت کاهش می‌یابد. برای اجتناب از مشکل زیادآموزی و بهبود عملکرد مدل یادگیری، ما اطلاعات اولیه را

توزیع حاشیه‌ای و شرطی بین دامنه‌ها می‌شود. روش DAIC به تفکیک نسبت به روش CDDA دارای متوسط بهبود دقت (۹۹/۴٪) در پایگاه داده آفیس و کالتک، (۶۶/۶٪) در پایگاه داده اعداد، و (۹۱/۱۶٪) در پایگاه داده پای است.

روش VDA علاوه بر تطبیق توزیع‌های حاشیه‌ای و شرطی شرطی از یک خوشه‌بندی مستقل از دامنه برای بهبود صحت طبقه‌بند استفاده می‌کند؛ با این حال، به دلیل ویژگی‌های متفاوت داده‌های آموزشی و آزمون در دامنه‌های منبع و هدف، صحت بالایی در پیش‌بینی داده‌های دامنه هدف را ندارد.

روش CLGA یک رویکرد تطبیق دامنه بدون نظارت چند منبعی است که به بررسی نحوه بهره‌گیری مؤثر و هم‌زمان از اطلاعات رده‌ها و دامنه‌ها می‌پردازد. روش DAIC به تفکیک نسبت به روش CLGA دارای متوسط بهبود دقت (۹۸/۱۳٪) در پایگاه داده آفیس و کالتک و (۸۶/۷٪) در پایگاه داده پای است.

۲-۵- ارزیابی پارامترها

برای به دست آوردن تنظیمات مدل، روش DAIC با مقادیر مختلف پارامترها مورد ارزیابی قرار می‌گیرد. در روش‌های یادگیری ماشین، یک اصل مهم در فرآیند بهینه‌سازی انتخاب نرخ یادگیری مناسب است ($\eta(k)$ در رابطه (۱۳)). اگر مقدار پایینی برای $\eta(k)$ انتخاب شود، سرعت هم‌گرایی بسیار کاهش می‌یابد، در حالی که اگر $\eta(k)$ مقدار بالایی داشته باشد، فرآیند هم‌گرایی دچار جهش شده و از محدوده مورد انتظار فراتر می‌رود و الگوریتم در تعداد تکرار معقول همگرا نمی‌شود. در این چارچوب پیشنهادی به طور تجربی $\eta(k) = \eta(0)/k$ در نظر گرفته شده که نرخ یادگیری با افزایش تعداد تکرار کاهش یابد. جهت محاسبه مقدار بهینه W از بیست تکرار برای به دست آوردن نتایج الگوریتم استفاده شده است؛ در بیشتر موارد DAIC در بیست تکرار اولیه به هم‌گرایی رسیده و تعداد تکرار بیشتر، تأثیری در افزایش دقت الگوریتم ندارد. λ پارامتر تنظیم بین واگرایی برگمن (D_W) و تابع هدف $F(W)$ است و در محدوده [۰، ۰/۰۰۰۰۱] بررسی می‌شود، مقدار بهینه پارامتر λ برای تمامی پایگاه داده‌ها برابر ۰/۱ است. علاوه بر این پارامتر τ در محدوده [۰/۰۱ ۱۰۰] در نظر گرفته می‌شود. τ یکی از پارامترهای تأثیرگذار در DAIC است، در واقع ضرایب ماتریس A را کنترل و تعداد نمونه‌های منبع مرتبط را تعیین

بین دامنه‌ها، دقت پایینی در تمام پایگاه‌داده‌های مورد آزمایش دارد.

با استفاده از واگرایی برگمن در نظر می‌گیریم. به همین ترتیب براساس نتایج نشان داده‌شده در شکل (۷) نمودار FLDA به دلیل عدم بهره‌گیری از اطلاعات تفکیک‌کننده اولیه داده‌های دامنه منبع و نادیده گرفتن اختلاف توزیع حاشیه‌ای

(جدول-۲): دقت (%) طبقه‌بند در مجموعه‌داده آفیس و کالتک

(Table-2): Classification accuracy (%) on Office and Caltech-256 datasets.

Dataset	NN	FLDA	TJM	TSL	DAM	FIDOS	LRSR	VDA	CLGA	CDDA	FLDA _m	DAIC
C-A	23.7	40.22	46.76	32.78	42.69	44.78	51.25	46.14	48.02	48.33	52.3	53.03
C-W	25.76	40.11	39.98	44.07	34.58	38.64	38.64	46.1	42.37	44.75	41.8	48.81
C-D	25.48	39.99	44.59	49.04	33.12	42.68	47.13	51.59	49.04	48.41	38.86	47.13
A-C	26	41.36	39.45	21.1	35.35	39.63	43.37	42.21	42.30	42.12	41.85	41.5
A-W	29.83	41.65	42.03	24.34	31.86	37.29	36.61	51.19	41.36	41.69	31.39	51.86
A-D	25.48	40.89	45.22	49.04	36.31	32.48	38.85	48.41	36.31	97.58	40.95	50.32
W-C	19.86	40	30.19	22.53	33.84	25.29	29.83	27.6	32.95	31.97	33.39	34.55
W-A	22.96	42.90	29.96	27.24	37.58	31.32	34.13	26.1	34.75	37.27	37.68	42.90
W-D	59.24	41.52	89.17	61.78	80.87	70.7	82.80	89.18	92.36	87.9	60.43	90.45
D-C	26.27	43.21	31.43	22.71	32.41	27.69	31.61	31.26	33.66	34.64	33.04	31.97
D-A	28.5	42.56	32.78	19.83	34.34	28.6	33.19	37.68	35.99	33.51	32.43	41.75
D-W	63.39	42.91	85.42	60.68	77.63	58.98	77.29	90.85	89.83	90.51	69.97	92.88
Average	31.37	41.38	44.41	37.09	42.55	39.84	45.39	49.03	48.23	47.22	42.84	52.21

(جدول-۳): دقت (%) طبقه‌بند در مجموعه‌داده پای

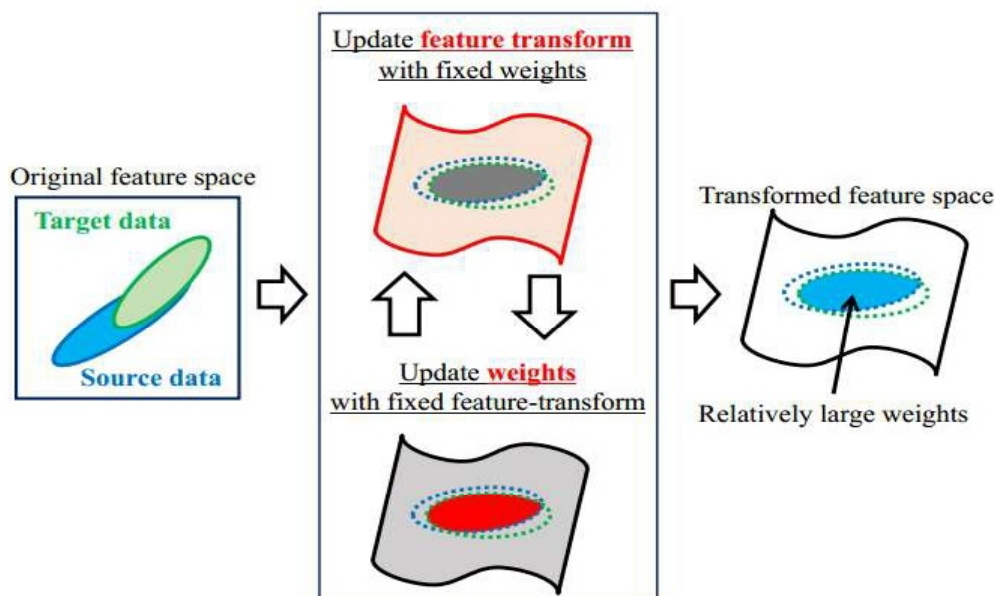
(Table-3): Classification accuracy (%) on PIE datasets.

Dataset	NN	FLDA	TJM	TSL	DAM	FIDOS	LRSR	VDA	CLGA	CDDA	FLDA _m	DAIC
P1_P2	26.09	33.89	23.87	33.46	46.65	20.87	65.87	40.09	67.83	60.22	37.63	82.63
P1_P3	26.59	33.56	28.86	33.58	45.04	27.51	64.09	45.22	63.85	58.7	47.55	78.06
P1_P4	30.67	32.93	43.37	35.27	68.52	39.62	82.03	71.64	88.95	83.48	69.18	92.67
P1_P5	16.67	38.79	19.3	26.78	28.74	22.79	54.90	33.33	61.76	54.17	36.09	64.95
P2_P1	24.49	35.29	26.14	45.47	35.32	27.79	45.54	55.58	71.40	62.33	38.6	76.65
P2_P3	46.63	34.78	37.93	42.59	35.78	34.8	53.49	57.23	72.98	64.64	44.91	77.7
P2_P4	54.07	35.17	50.53	45.63	72.33	46.59	71.43	68.82	86.24	79.9	69.84	89.67
P2_P5	26.53	32.41	21.63	31.43	35.11	27.63	47.97	35.23	51.23	44	34.01	64.22
P3_P1	21.37	37.36	28.66	42.44	41	33.19	52.49	46.64	70.17	58.46	44.48	71.58
P3_P2	41.01	37.03	35.97	37.02	59.36	36.28	55.56	53.04	73.48	59.73	40.7	82.26
P3_P4	46.53	38.45	51.97	41.12	72.33	53.08	77.50	62.75	89.31	77.2	69.42	88.22
P3_P5	26.23	32.59	25.31	30.27	40.32	35.29	54.11	39.22	55.51	47.24	46.81	72.49
P4_P1	32.95	34.53	45.71	36.49	69.36	51.35	81.54	71.55	89.56	83.1	72.09	93.88
P4_P2	62.68	35.21	57.58	39.29	77.29	51.87	58.39	75.51	92.94	82.26	65.56	95.27
P4_P3	73.22	34.96	71.63	43.01	76.72	66.05	82.23	80.21	93.08	86.64	72.79	93.81
P4_P5	37.19	34.80	30.94	32.11	61.15	42.65	72.61	53.06	71.63	58.33	59.25	84.99
P5_P1	18.49	31.81	27.13	37.39	30.52	21.13	52.19	49.01	57.68	48.02	39.83	59.03
P5_P2	24.1	29.28	22.65	34.62	37.14	22.9	49.41	38.67	55.43	45.61	34.62	71.09
P5_P3	28.31	34.06	28.86	35.11	41.11	26.41	58.45	36.4	58.03	52.02	50.67	74.45
P5_P4	31.24	29.12	32.59	46.08	39.14	31.12	64.31	46.11	71.85	55.99	62.06	86.57
Average	34.76	34.30	35.53	37.46	49.5	35.95	61.53	52.97	72.15	63.1	51.81	80.01

(جدول-۴): دقت (%) طبقه‌بندی در مجموعه داده اعداد

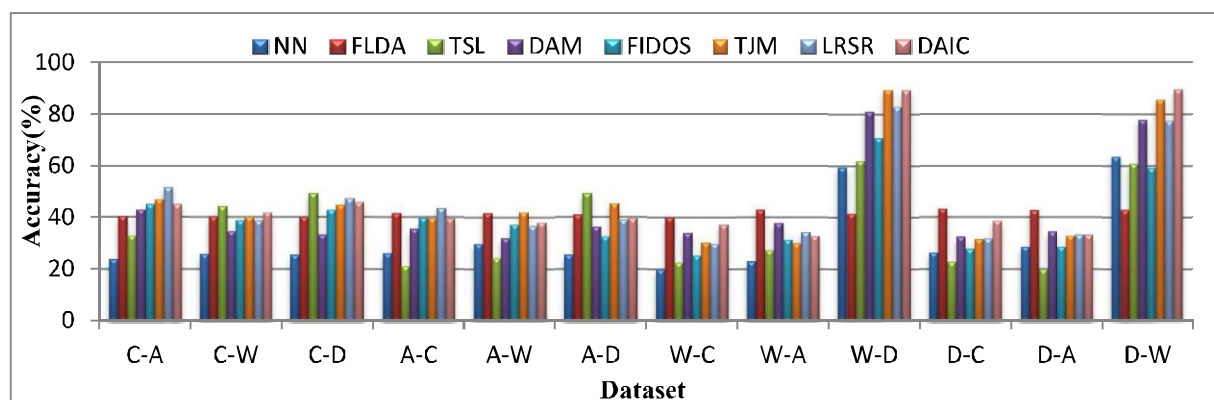
(Table-4): Classification accuracy (%) on USPS+ MNIST datasets

Dataset	NN	FLDA	TJM	TSL	DAM	LRSR	VDA	CDDA	FLDA _m	DAIC
U-M	44.7	73.51	52.25	50.74	42.69	54.51	62.95	62.05	60.54	68.50
M-U	65.94	64.89	63.28	59.6	52.83	73.82	74.72	76.22	73.21	83.00
Average	55.32	58.31	57.76	55.17	47.76	64.16	68.84	69.13	66.87	75.75



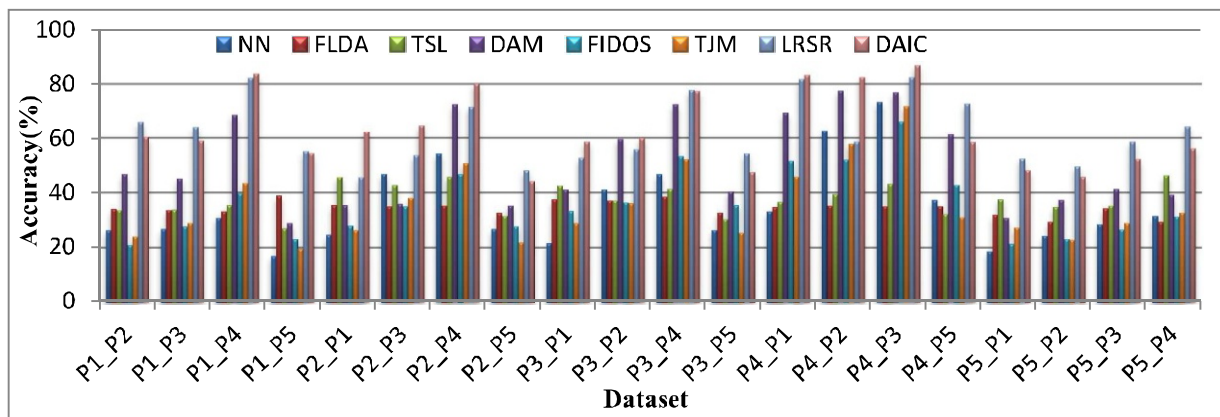
(شکل-۳): روش پیشنهادی DAIC داده‌های منبع و هدف را که در فضای اصلی دارای توزیع متفاوت هستند با استفاده از ماتریس تبدیل FLDA و واگرایی برگمن به یک زیرفضای جدید نگاشت می‌کند تا اختلاف توزیع آن‌ها را کاهش دهد، هم‌زمان با این کار با استفاده از یک ماتریس وزن تأثیر نمونه‌های نامرتب دامنه منبع را کم می‌کند. از آنجایی که پیدا کردن هم‌زمان ماتریس تبدیل و ماتریس وزن کار دشواری است DAIC ابتدا در مرحله نخست ماتریس وزن را ثابت در نظر می‌گیرد و ماتریس تبدیل بهینه را به دست می‌آورد و در مرحله دوم ماتریس تبدیل به دست آمده را ثابت در نظر می‌گیرد و ماتریس وزن بهینه را به دست می‌آورد.

(Figure-3): An overview of the proposed method. DAIC projects the source and target data into a shared low dimensional subspace based on Bregman divergence minimization and FLDA criteria and at the same time, the use of a weight matrix reduces the impact of unrelated instances on the source domain. Since it is mathematically difficult to optimize them simultaneously, we adopt an interactive approach in which the optimization of feature transform and that of instance weight are conducted alternately.



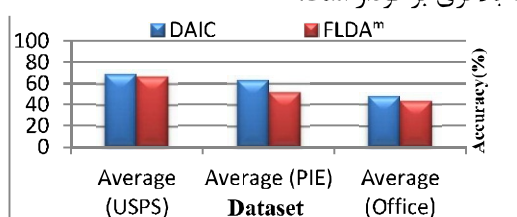
(شکل-۴): صحت طبقه‌بندی در مجموعه داده‌های آفیس و کالتک با استفاده از روش‌های مختلف NN, FLDA, TSL, DAM, FIDOS, TJM, LRSR و DAIC

(Figure-4): Classification accuracy (%) on Office and Caltech-256 datasets using different methods NN, FLDA, TSL, DAM, FIDOS, TJM, LRSR and DAIC.



شکل-۵: صحت طبقه‌بندی در مجموعه داده‌های پای با استفاده از روش‌های مختلف NN, FLDA, TSL, DAM, FIDOS, TJM, LRSR و DAIC (Figure-5): Classification accuracy (%) on PIE datasets using different methods NN, FLDA, TSL, DAM, FIDOS, TJM, LRSR and DAIC.

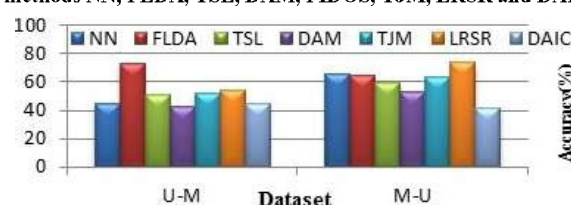
(۸) نمودار DAIC در مقایسه با $FLDA^m$ به دلیل بهره‌گیری هم‌زمان از تطبیق خصوصیات و روش وزن‌دهی به نمونه‌ها از دقت بالاتری برخوردار است.



شکل-۸: مقایسه دقت (%) مدل پیشنهادی با توجه به کاهش

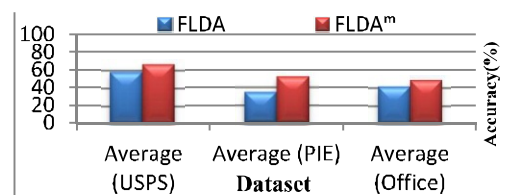
اختلاف توزیع حاشیه‌ای و وزن‌دهی به نمونه‌ها

(Figure-8): Comparison of the accuracy (%) of the proposed model with respect to the reduction of marginal distribution and sample reweighting



شکل-۶: صحت طبقه‌بندی در مجموعه داده‌های اعداد با استفاده از روش‌های مختلف NN, FLDA, TSL, DAM, TJM, LRSR و DAIC

(Figure-6): Classification accuracy (%) on Digit datasets using different methods NN, FLDA, TSL, DAM, TJM, LRSR and DAIC.



شکل-۷: مقایسه دقت (%) مدل پیشنهادی با توجه به کاهش

اختلاف توزیع حاشیه‌ای و روش FLDA

(Figure-7): Comparison of the accuracy (%) of the proposed model with regard to the reduction of the marginal gap and the FLDA method

۴-۵- بررسی ضرورت وزن‌دهی مجدد به نمونه‌ها

در شرایطی که اختلاف توزیع دامنه‌های منبع و هدف زیاد باشد، فقط با کاهش اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف، نمی‌توان باعث ایجاد تطبیق‌پذیری طبقه‌بند ایجادشده بر روی دامنه منبع با دامنه هدف شد (جدول ۲ الی ۴ فیلد $FLDA^m$). در چنین شرایطی، حتی با ایجاد تفکیک‌پذیری بین رده‌ها، به علت اینکه ساختار داده‌های دامنه هدف با داده‌های دامنه منبع متفاوت است، یک طبقه‌بند تطبیق‌پذیر بین دامنه‌های ایجاد نخواهد شد. به همین دلیل روش DAIC، به نمونه‌هایی از دامنه منبع که دارای کمترین مشابهت از نظر توزیع با نمونه‌های دامنه هدف هستند، وزن کمتری در ایجاد مدل طبقه‌بند اختصاص می‌دهد. در شکل

۶- نتیجه‌گیری و کارهای آینده

پردازش تصویر به عنوان ابزار قدرتمند برای درک تصاویر توسط رایانه است که توانایی استخراج اطلاعات و ویژگی‌های تصاویر را دارد. یادگیری ماشین یکی از شاخه‌های پرکاربرد هوش مصنوعی بوده که در زمینه پردازش تصویر به دنبال یافتن الگوریتم‌هایی برای طبقه‌بندی تصاویر با دقت بالا است. در این مقاله، روش تطبیق دامنه بدون نظارت DAIC برای ایجاد تطبیق در دامنه‌های بصری و کاهش اختلاف توزیع حاشیه‌ای بین دامنه‌ها پیشنهاد شده است. روش پیشنهادی، در ابتدا با استفاده از واگرایی برگمن یک نمایش تطبیق‌پذیر مشترک بین دامنه‌های منبع و هدف ایجاد می‌کند که اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف را کاهش می‌دهد. درواقع این روش با استفاده از اطلاعات تفکیک‌کننده اولیه داده‌های ورودی، مدلی را ایجاد می‌کند که مقاومت بیشتری در برابر اختلاف توزیع داده‌های منبع و هدف داشته باشد؛ سپس به طور هم‌زمان، به نمونه‌هایی از دامنه منبع که دارای کمترین مشابهت از نظر توزیع با نمونه‌های دامنه هدف

work for transfer learning”, *IEEE Trans. Knowl. Data Eng.*, vol. 26, pp. 1076–1089, 2013.

- [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] J. Tahmoresnezhad and S. Hashemi, “Visual domain adaptation via transfer feature learning”, *KnowlInf Syst.*, vol. 50, no. 2, pp. 585–605, 2016.
- [10] M. Long, J. Wang, G. Ding, J. Sun, and P. Yu, “Transfer feature learning with joint distribution adaptation”, in *Proc IEEE International Conference on Computer Vision*, pp. 2200–2207, 2013.
- [11] 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*, vol. PP, no. 99, pp. 1–1, 2016.
- [12] R. A. Fisher, “The use of multiple measurements in taxonomic problems”, *Annals of eugenics*, vol. 7, no. 2, pp. 179–188, 1936.
- [13] K. Sacnko, 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] J. J. Hull, “A database for handwritten text recognition research”, *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 16, no. 5, pp. 550–554, 1994.
- [15] 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.
- [16] 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.
- [17] Cuong V. Dinh, Robert P.W. Duin, Ignacio Piqueras-Salazar, and Marco Loog. Fidos: A generalized fisher based feature extraction method for domain shift. *Pattern Recognition*, 46(9):2510–2518, 2013.
- [18] 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, pp. 504–518, 2012.
- [19] M. Long, J. Wang, G. Ding, J. Sun, and P. S. Yu, “Transfer joint matching for unsupervised domain adaptation”, in *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, IEEE, pp. 1410–1417, 2014.

هستند، وزن کمتری در ایجاد مدل طبقه‌بند اختصاص داده می‌شود.

DAIC با افزایش تطبیق‌پذیری بین ساختار دامنه‌های منبع و هدف، باعث افزایش دقت طبقه‌بند در پیش‌بینی برچسب برای داده‌های بدون برچسب دامنه هدف می‌شود. روش DAIC بر روی ۳۴ مجموعه داده بصری مورد آزمایش قرار گرفته که این مجموعه داده‌ها، دارای اختلاف توزیع قابل توجهی نسبت به یکدیگر هستند. نتایج به دست آمده، نشان‌دهنده بهبود قابل ملاحظه‌ای از کارایی روش DAIC نسبت به جدیدترین روش‌های حوزه یادگیری ماشین و یادگیری انتقالی بر روی دامنه‌های مختلف است. برای ادامه کار، ما در حال برنامه‌ریزی جهت گسترش DAIC برای کاهش هم‌زمان توزیع اختلاف‌های حاشیه‌ای و شرطی و علاوه بر این در تلاش جهت توسعه روش DAIC برای سامانه‌های چندمنبعی هستیم. در این راستا، به دنبال انتقال دانش از چند منبع مرتبط به یک منبع هدف هستیم تا بتوانیم صحت پیش‌بینی برچسب دامنه هدف را هرچه بیشتر ارتقا بخشیم.

7- References

۷- مراجع

- [1] S. J. Pan and Q. Yang, “A survey on transfer learning”, *IEEE Transactions on Knowledge and Data Engineering*, vol. 22, no. 10, pp. 1345–1359, 2010.
- [2] L. Shao, F. Zhu, and X. Li, “Transfer learning for visual categorization: A survey”, *IEEE transactions on neural networks and learning systems*, vol. 26, no. 5, pp. 1019–1034, 2015.
- [3] S. Ben-David, J. Blitzer, K. Crammer, A. Kulesza, F. Pereira, and J. W. Vaughan, “A theory of learning from different domains”, *Machine learning*, vol. 79, no. 1–2, pp. 151–175, 2010.
- [4] 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.
- [5] S. Si, D. Tao, and B. Geng, “Bregman divergence-based regularization for transfer subspace learning”, *IEEE TKDE*, 2010.
- [6] L. M. Bregman, “The relaxation method of finding the common points of convex sets and its application to the solution of problems in convex programming”, *USSR Computational Mathematics and Mathematical Physics*, vol. 7, pp. 200–217, 1967.
- [7] M. Long, J. Wang, G. Ding, S. J. Pan and P. Yu, “Adaptation regularization: a general frame-

- [32] Z. Ding, M. Shao, and Y. Fu, "Deep low-rank coding for transfer learning," in *Proceedings of the 24th International Conference on Artificial Intelligence*. AAAI Press, 2015, pp. 3453–3459.
- [33] M. Long, J. Wang, G. Ding, J. Sun, and P. Yu, "Transfer feature learning with joint distribution adaptation," in *Proceedings of the IEEE International Conference on Computer Vision*, 2013, pp. 2200–2207.
- [34] H. Li, T. Jiang, and K. Zhang, "Efficient and robust feature extraction by maximum margin criterion," *IEEE Transactions on Neural Networks*, vol. 17, no. 1, pp. 157–165, 2006.
- [35] M. Long, J. Wang, G. Ding, S. J. Pan, and P. S. Yu, "Adaptation regularization: A general framework for transfer learning," *IEEE Transactions on Knowledge and Data Engineering*, vol. 26, no. 5, pp. 1076–1089, 2014.
- [36] Y. Xu, S. J. Pan, H. Xiong, Q. Wu, R. Luo, H. Min, and H. Song, "A unified framework for metric transfer learning," *IEEE Trans. Knowl. Data Eng.*, vol. 29, no. 6, pp. 1158–1171, 2017.
- [37] S. Z. Seyyedsalehi, S. A. Seyyedsalehi, "Improving nonlinear manifold separator model to the face recognition by a single image of per person", *Signal and Data Processing*, vol. 12, no. 1, pp. 3–16, 2015.
- [38] S. Ahmadkhani, and P. Adibi, "Supervised Probability Component Analysis Mixture Model in a Lossless Dimensionality Reduction Framework for Face Recognition", *Signal and Data Processing*, vol. 12, no. 4, pp. 53–65, 2016.
- [20] Y. Xu, X. Fang, J. Wu, X. Li, and D. Zhang, "Discriminative transfer subspace learning via low-rank and sparse representation," *IEEE Transactions on Image Processing*, vol. 25, no. 2, pp. 850–863, 2016.
- [21] J. Tahmoresnezhad and S. Hashemi, "Visual domain adaptation via transfer feature learning," *Knowledge and Information Systems*, vol. 50, no. 2, pp. 585–605, 2017.
- [22] J. Liu, J. Li, and K. Lu, "Coupled local-global adaptation for multi-source transfer learning," *Neurocomputing*, vol. 275, pp. 247–254, 2018.
- [23] L. Luo, X. Wang, S. Hu, C. Wang, Y. Tang, and L. Chen, "Close yet distinctive domain adaptation," *arXiv preprint arXiv: 1704.04235*, 2017.
- [24] W. Dai, Q. Yang, G.-R. Xue, and Y. Yu, "Boosting for transfer learning," in *Proceedings of the 24th international conference on Machine learning*, 2007, pp. 193–200.
- [25] Y. Tsuboi, H. Kashima, S. Hido, S. Bickel, and M. Sugiyama, "Direct density ratio estimation for large-scale covariate shift adaptation," *Information and Media Technologies*, vol. 4, no. 2, pp. 529–546, 2009.
- [26] J. Quionero-Candela, M. Sugiyama, A. Schwaighofer, and N. D. Lawrence, "Dataset Shift in Machine Learning", The MIT Press, 2009.
- [27] S. J. Pan, X. Ni, J.-T. Sun, Q. Yang, and Z. Chen, "Cross-domain sentiment classification via spectral feature alignment," in *Proceedings of the 19th international conference on World Wide Web. ACM*, pp. 751–760, 2010.
- [28] S. J. Pan, I. W. Tsang, J. T. Kwok, and Q. Yang, "Domain adaptation via transfer component analysis," *IEEE Transactions on Neural Networks*, vol. 22, no. 2, pp. 199 – 210, 2011.
- [29] X. Shi, Q. Liu, W. Fan, and P. S. Yu, "Transfer across completely different feature spaces via spectral embedding," *IEEE Transactions on Knowledge and Data Engineering*, vol. 25, no. 4, pp. 906–918, 2013.
- [30] R. Aljundi, R. Emonet, D. Muselet, and M. Sebban, "Landmarks-based kernelized subspace alignment for unsupervised domain adaptation," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 56–63, 2015.
- [31] M. Gheisari and M. S. Baghshah, "Joint predictive model and representation learning for visual domain adaptation," *Engineering Applications of Artificial Intelligence*, vol. 58, pp. 157–170, 2017.



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 و علایق پژوهشی

ایشان شامل حوزه‌های یادگیری ماشین، یادگیری انتقالی، داده‌کاوی و سامانه‌های امنیتی است.
نشانی رایانامه ایشان عبارت است از:
j.tahmores @it.uut.ac.ir



مژده زندی‌فر دانشجوی کارشناسی ارشد
مهندسی فناوری اطلاعات، مدرک
کارشناسی خود را در رشته مهندسی
کامپیوتر-ترجم‌افزار در سال ۹۴ دریافت
کردند.

نشانی رایانامه ایشان عبارت است از:
mozhddeh.zandifar @it.uut.ac.ir