



چکیده

در بسیاری از الگوریتم‌های یادگیری ماشین، فرض اولیه بر این اساس است که مجموعه داده آموزشی (دامنه منبع) و مجموعه داده آزمون (دامنه هدف) توزیع یکسانی را به اشتراک می‌گذارند. این در حالی است که در اغلب مسائل دنیای واقعی، به دلیل اختلاف توزیع احتمال بین دامنه منبع و هدف، این فرض نقض می‌شود. برای مقابله با این مشکل، یادگیری انتقالی و تطبیق دامنه، مدل را برای مقابله با داده‌های هدف دارای توزیع متفاوت، تعمیم می‌دهند. در این مقاله ما یک روش تطبیق دامنه با عنوان هم‌ترازی تصویر از طریق یادگیری خصوصیت کرنل شده (IMAKE) را به منظور حفظ اطلاعات عمومی و هندسی دامنه‌های منبع و هدف پیشنهاد می‌دهیم. روش پیشنهادی یک زیرفضای مشترک بین دامنه‌های منبع و هدف جستجو می‌کند تا اختلاف توزیع آنها را به کمینه برساند. IMAKE از هر دو تطبیق توزیع هندسی و عمومی به صورت هم‌زمان بهره می‌برد. روش پیشنهادی دامنه‌های منبع و هدف را به یک زیرفضای کم‌بعد مشترک به صورت بدون نظارت منتقل می‌کند تا اختلاف احتمال توزیع شرطی و حاشیه‌ای داده‌های دامنه منبع و هدف را از طریق بیشینه اختلاف میانگین‌ها کمینه کند و برای تطبیق توزیع هندسی از هم‌ترازی منیفولد بهره می‌گیرد. کارایی روش پیشنهادی با استفاده از پایگاه داده‌های بصری متنوع و استاندارد با ۳۶ آزمایش مورد ارزیابی قرار گرفته است. نتایج به دست آمده، نشان‌دهنده بهبود قابل ملاحظه از عملکرد روش پیشنهادی در مقایسه با جدیدترین روش‌های حوزه یادگیری ماشین و یادگیری انتقالی است.

واژگان کلیدی: طبقه‌بندی تصویر، یادگیری انتقالی، تطبیق دامنه بصری، هم‌ترازی منیفولد، اختلاف توزیع

Image alignment via kernelized feature learning

Elahe Shahrouz & Jafar Tahmoresnezhad*

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

Abstract

Machine learning is an application of artificial intelligence that is able to automatically learn and improve from experience without being explicitly programmed. The primary assumption for most of the machine learning algorithms is that the training set (source domain) and the test set (target domain) follow from the same probability distribution. However, in most of the real-world applications, this assumption is violated since the probability distribution of the source and target domains are different. This issue is known as domain shift. Therefore, transfer learning and domain adaptation generalize the model to face target data with different distribution.

In this paper, we propose a domain adaptation method referred to as IMage Alignment via KERnelized feature learning (IMAKE) in order to preserve the general and geometric information of the source and target domains. IMAKE finds a common subspace across domains to reduce the distribution discrepancy between the source and the target domains. IMAKE adapts both the geometric and the general distributions, simultaneously. Moreover, IMAKE transfers the source and target domains into a shared low dimensional subspace in an unsupervised manner.

Our proposed method minimizes the marginal and conditional probability distribution differences of the source and target data via maximum mean discrepancy and manifold alignment for geometrical distribution adaptation. IMAKE maps the input data into a common latent subspace via manifold alignment as a geometric matching method. Therefore, the samples with the same class labels are collected around their means, and samples with different class are separated, as well. Moreover, IMAKE maintains the source and target domain manifolds to preserve the original data position and

* Corresponding author

* نویسنده عهده‌دار مکاتبات

domain structure. Also, the use of kernels and mapping data into Hilbert space provides more accurate separation between different classes and is suitable for data with complex and unbalanced structures. The proposed method has been evaluated using a variety of benchmark visual databases with 36 experiments. The results indicate the significant improvements of the proposed method performance against other machine learning and transfer learning approaches.

Keyword: Image classification, Transfer learning, Visual domain adaptation, Manifold alignment, Distribution mismatch.

۱- مقدمه

یادگیری ماشین به عنوان یکی از روش‌های تجزیه و تحلیل داده‌ها به ایجاد الگوریتم‌هایی به منظور بهبود توانایی یادگیری سیستم‌ها می‌پردازد. با استفاده از الگوریتم‌های یادگیری ماشین، یک مدل بر روی داده‌های آموزشی^۱ ایجاد می‌شود؛ سپس از مدل ساخته شده جهت برچسب‌گذاری داده‌های آزمون^۲ بهره گرفته می‌شود [1]. در اغلب الگوریتم‌های یادگیری ماشین فرض بر این است که داده‌های دامنه آموزشی و داده‌های دامنه آزمون از یک توزیع یکسان تبعیت می‌کنند، اما در دنیای واقعی، با توجه به عوامل مختلف مانند تفاوت حس‌گرها، تغییرات نور، پس‌زمینه، زوایای دید مختلف و شفافیت داده‌های دامنه هدف، ممکن است داده‌های دامنه آزمون، توزیع متفاوتی از داده‌های برچسب‌دار دامنه منبع داشته باشند [2, 3]. از طرف دیگر روش‌های طبقه‌بندی و رتبه‌بندی در یادگیری ماشین به‌طور معمول مبتنی بر در دسترس بودن مقدار زیادی از داده‌های برچسب‌گذاری شده^۳ برای آموزش یک مدل هستند، اما در بسیاری از کاربردهای دنیای واقعی به تعداد کافی نمونه برچسب‌دار در دسترس نبوده و جمع‌آوری داده‌های برچسب‌دار، بسیار پرهزینه است. این مشکل در بسیاری از برنامه‌های کاربردی در بازیابی اطلاعات، تجارت الکترونیکی، بینایی ماشین و بسیاری از زمینه‌های دیگر به وجود می‌آید. به عنوان نمونه در یک سامانه تشخیص چهره، اگر داده‌های آموزشی مجموعه تصاویری باشند که با دوربین‌های پلیس در روز از افراد مختلف تهیه و برچسب‌گذاری شده و داده‌های آزمون، مجموعه‌ای از تصاویر گرفته شده به وسیله همان دوربین‌ها در شب باشد، اختلاف موجود بین تصاویر آموزشی و آزمون، از نظر کیفیت نور، شفافیت، جهت تابش خورشید، زاویه تصویر و فاصله تصویر گرفته شده از هدف، موجب ایجاد تغییر دامنه^۴ و کاهش بازدهی مدل ساخته شده، در دامنه آزمایش خواهد شد.

یادگیری انتقالی^۵ و انطباق دامنه^۶ روش‌هایی هستند که با استفاده از آن‌ها دانش کسب شده در یک یا چند دامنه مرتبط، به یک دامنه متفاوت ولی مرتبط دیگری منتقل می‌شود تا با وجود اختلاف توزیع بین دامنه‌ها، هزینه یادگیری کاهش یافته و عملکرد الگوریتم‌های یادگیری افزایش یابد. یادگیری انتقالی می‌تواند در دو نوع نیمه نظارت شده^۷ و بدون نظارت^۸ بررسی شود [4]. اگر دامنه منبع^۹ شامل داده‌های آموزشی و دامنه هدف^{۱۰} شامل داده‌های آزمون باشد، در یادگیری انتقالی نیمه نظارت شده، تمام داده‌های دامنه منبع دارای برچسب بوده ولی فقط تعداد محدودی از داده‌های دامنه هدف، دارای برچسب بوده که اغلب این داده‌ها برای ایجاد یک طبقه بند دقیق، مناسب نیستند؛ این در حالی است که اگر تمام داده‌های دامنه منبع دارای برچسب و تمام داده‌های دامنه هدف بدون برچسب باشند، یادگیری انتقالی بدون نظارت است. در بسیاری از کاربردهای دنیای واقعی، به دلیل اینکه به تعداد کافی نمونه برچسب‌دار از دامنه هدف، در دسترس نبوده و همچنین، برچسب‌گذاری دستی آن‌ها بسیار پرهزینه است، از یادگیری انتقالی بدون نظارت استفاده می‌شود [5].

در یادگیری انتقالی بدون نظارت، ابتدا اختلاف توزیع بین دامنه منبع و هدف کاهش یافته، سپس از الگوریتم‌های یادگیری ماشین برای ساختن مدل بر روی داده‌ها استفاده می‌شود. اختلاف توزیع بین دامنه‌های منبع و هدف، شامل اختلاف در توزیع حاشیه‌ای^{۱۱} و توزیع شرطی^{۱۲} است [6]. در شرایطی که دامنه منبع و هدف دارای خصوصیات یکسان باشند، اختلاف در احتمال وقوع مقادیر هر کدام از این خصوصیات در هر دامنه، موجب ایجاد اختلاف توزیع حاشیه‌ای بین دامنه‌ها می‌شود. توزیع شرطی به احتمال پیش‌بینی یک مجموعه برچسب به ازای یک مجموعه داده ورودی گفته می‌شود که معادل مفهوم

⁵ Transfer learning

⁶ Domain adaptation

⁷ Semi-supervised

⁸ Unsupervised

⁹ Source Domain

¹⁰ Target domain

¹¹ Marginal distribution

¹² Conditional distribution

¹ Training data

² Test data

³ Labeled data

⁴ Shift domain

تابع پیش‌بینی است. اگر دو دامنه منبع و هدف دارای اختلاف توزیع شرطی باشند، تابع پیش‌بینی دامنه‌ها باهم متفاوت بوده و مجموعه برچسب متفاوتی به‌ازای داده‌های یکسان از هر دو دامنه، پیش‌بینی خواهد شد. در بحث مربوط به یادگیری انتقالی، هر دو اختلاف توزیع حاشیه‌ای و شرطی دارای اهمیت بوده و مسأله اصلی کاهش هم‌زمان این اختلاف‌ها است.

روش‌های جدید انطباق دامنه که برای به کمینه‌رساندن اختلاف توزیع دامنه منبع و هدف تعریف شده‌اند به سه دسته کلی تقسیم‌بندی می‌شوند [2]: (۱) روش‌های مبتنی بر نمونه^۱: در این روش‌ها با تغییر در وزن نمونه‌های برچسب‌دار دامنه منبع، اختلاف توزیع بین دامنه‌های منبع و هدف کاهش می‌یابد [7، ۲]. روش‌های مبتنی بر مدل^۲: در این روش با تغییر در مدل از طریق تغییر در پارامترهای اصلی یا مشترک، مدل در مقابل اختلاف توزیع‌ها مقاوم می‌شود [3، ۳]. روش‌های مبتنی بر خصوصیت^۳: در این روش‌ها با نگاشت داده‌ها به فضای خصیصه‌ای دیگر، اختلاف توزیع بین داده‌های دامنه منبع و داده‌های دامنه هدف کاهش می‌یابد [8]. که در این مقاله بر روی روش‌های مبتنی بر خصوصیت تمرکز شده است.

هم‌ترازی منیفلد^۴ [9, 10] یکی از روش‌های کاهش بعد است. این روش یک چارچوب هندسی برای ساخت فضایی کم‌بعد مانند فضای نهان^۵ فراهم می‌کند. ایده کلی این رویکرد این است که دامنه‌های مختلف را به یک فضای نهان نگاشت کند؛ درحالی‌که هم‌زمان نمونه‌های یکسان در کنار هم قرار گرفته و هندسه محلی (رابطه همسایگی) هر دامنه حفظ شود. هم‌ترازی منیفلد هم از داده‌های برچسب‌دار و هم از داده‌های بدون برچسب بهره می‌برد. استفاده از داده‌های بدون برچسب برای مسئله انطباق دامنه، به دلیل کمبود داده‌های برچسب‌دار مفید خواهد بود. بدین ترتیب، در فضای جدید دامنه‌های ورودی با خصوصیات یکسانی تعریف می‌شوند؛ بنابراین هم‌ترازی منیفلد می‌تواند با انواع مختلف روش‌های یادگیری انتقالی ترکیب شده و برای حل چالش‌های انتقال دانش در مسائل دنیای واقعی مورد استفاده قرار گیرد [11].

روش پیشنهادی در این مقاله با عنوان (IMAKE)^۶ به دنبال یافتن یک نمایش جدید از نمونه‌هاست که در این

نمایش از تطبیق توزیع‌های عمومی و هندسی و خوشه‌بندی هندسی-طبقه‌ای بهره گرفته می‌شود. در مرحله تطبیق خصوصیات عمومی، یک نمایش جدید از داده‌های دامنه منبع و هدف ایجاد می‌شود که در آن، اختلاف توزیع شرطی و حاشیه‌ای به کمینه می‌رسد تا داده‌ها تطبیق‌پذیری دقیق‌تری باهم داشته باشند. در مرحله تطبیق توزیع هندسی، با حفظ توپولوژی اولیه دامنه‌ها، داده‌ها به یک فضای نهان نگاشت شده به‌طوری‌که در آن فضا داده‌های متعلق به یک طبقه یکسان نزدیک به هم قرار گرفته و طبقه‌های متفاوت از هم فاصله می‌گیرند. در این مرحله، استفاده از برچسب‌های موجود در دامنه منبع و هدف و انطباق داده‌هایی با برچسب‌های یکسان و تفکیک داده‌هایی با برچسب متفاوت، موجب خوشه‌بندی مناسب داده‌ها می‌شود. در مرحله نهایی، برای تطبیق توزیع هندسی بهتر دامنه‌هایی که ساختارهای هندسی بسیار متفاوتی باهم دارند، از کرنل استفاده شده است تا دامنه‌های مختلف به فضای کرنل هیلبرت نگاشت شده و در آن فضا باهم تطبیق شوند. این مرحله نقش بسیار مهمی در انطباق دامنه‌هایی دارد که دارای ساختارهای هندسی نامتوازن از هم هستند.

کارایی روش پیشنهادی در این مقاله، بر روی پایگاه داده‌های شناخته‌شده بصری تحت شرایط مختلف مورد آزمایش قرار گرفته و نتایج آن‌ها، با جدیدترین روش‌ها در حوزه یادگیری انتقالی مقایسه شده است. نتایج حاصل، نشان‌دهنده برتری قابل‌ملاحظه الگوریتم پیشنهادی نسبت به سایر روش‌های شناخته‌شده در حوزه یادگیری انتقالی است.

در ادامه، مقاله به‌صورت زیر سازمان‌دهی شده است. در بخش دوم مقاله، مروری بر کارهای پیشین در این حوزه گنجانده شده است. در بخش سوم، روش پیشنهادی به‌تفصیل شرح داده شده است. در بخش چهارم، پایگاه داده‌های مورد ارزیابی در این مقاله با جزئیات معرفی شده‌اند. در بخش پنجم، نتایج عملکرد الگوریتم پیشنهادی با روش‌های دیگر یادگیری ماشین و یادگیری انتقالی گزارش شده است. در انتها، مقاله با نتیجه‌گیری و ارائه پیشنهادهایی برای کار در آینده به اتمام رسیده است.

۲- کارهای پیشین

در حوزه یادگیری ماشین روش‌های بسیاری در زمینه تطبیق دامنه‌ها پیشنهاد شده که تمرکز بیشتر آن‌ها بر کاهش اختلاف توزیع بین دامنه‌های منبع و هدف است. همان‌طور که در بخش قبلی اشاره شد، روش‌های تطبیق

¹ Sample based

² Model based

³ Feature based

⁴ Manifold Alignment

⁵ Latent space

⁶ IMake Alignment via Kernalized feature learning

دامنه به سه دسته کلی تقسیم می‌شوند: روش‌های مبتنی بر نمونه، روش‌های مبتنی بر مدل و روش‌های مبتنی بر خصوصیت.

رویکردهای مبتنی بر نمونه بر اساس وزن‌دهی مجدد و یا انتخاب نمونه‌هایی که تفاوت بین توزیع دامنه‌های منبع و هدف را کمینه می‌سازد، عمل می‌کنند. از جمله روش‌های انتخاب نمونه، روش انتخاب لندمارک است [7, 12]. لندمارک، نمونه‌هایی از دامنه منبع هستند که دارای کمینه اختلاف توزیع حاشیه‌ای با نمونه‌های دامنه هدف هستند. TIT [13] یکی از الگوریتم‌های شاخص در این دسته است. عملکرد TIT به این صورت است که وزن نمونه‌های محوری را افزایش داده و وزن نمونه‌های خارج از محدوده را کاهش می‌دهد. این نوع انتخاب لندمارک مبتنی بر بهینه‌سازی گراف است. در این روش با کاهش بردارهای ویژگی به رئوس گراف، عملیات ریاضی نسبت به روش‌های دیگر بسیار ساده‌تر می‌شود.

در روش‌های مبتنی بر مدل، هدف پیدا کردن یک طبقه‌بند انطباقی است که این کار به وسیله انتقال پارامترهای آموزش داده شده در دامنه منبع به دامنه هدف، بدون تغییر فضای خصیصه‌ای انجام می‌شود [14-16]. JDAC¹ [17] یک روش تطبیق مدل پیشنهاد شده برای مسائل تطبیق دامنه بدون نظارت است. این روش یک طبقه‌بند انطباقی بر روی هر دو دامنه منبع و هدف ایجاد می‌کند به طوری که نخست اختلاف توزیع‌های شرطی و حاشیه‌ای را به کمینه رسانده، دوم بهینه انطباق با توزیع هندسی نمونه‌ها را در هر دو دامنه فراهم کرده و سوم دارای کمینه خطای پیش‌بینی برچسب نمونه‌های دامنه منبع باشد.

روش‌های مبتنی بر خصوصیت یا تطبیق خصوصیات، فضای خصیصه‌ای را برای ایجاد یک نمایش تطبیق‌پذیر از دامنه‌های منبع و هدف تغییر داده، سپس در فضای جدید، یک طبقه‌بند استاندارد روی داده‌های دامنه منبع، آموزش داده و روی داده‌های دامنه هدف اعمال می‌کنند [18, 19]. روش TJM² [18] اختلاف توزیع دامنه‌ها را با انطباق هم‌زمان خصوصیات و وزن‌دهی مجدد نمونه‌ها در یک فضای با بعد کاهش یافته انطباق می‌دهد. در روش TJM زیرفضای به دست آمده هم در برابر اختلاف توزیع‌ها و هم در برابر نمونه‌های نامرتبب مقاوم است. روش GFK³ [19] جابه‌جایی دامنه‌ها را با استفاده از

یک پارچه‌سازی بی‌نهایت زیرفضا که تغییرات در هندسه و آمار خصوصیات را نشان می‌دهند، مدل می‌کند. روش‌های مبتنی بر خصوصیت به دو دسته کلی تقسیم می‌شوند:

(۱) روش‌های تطبیق خصوصیات مبتنی بر داده. در این روش‌ها، دامنه‌های منبع و هدف با استخراج خصوصیات مشترک بین دامنه‌ها، به یک زیرفضای مشترک دارای کمینه اختلاف توزیع و بیشینه حفظ ساختار اولیه‌شان نگاشت می‌شوند. از جمله روش‌های شناخته شده در یادگیری بدون نظارت CLGA⁴ [20] به تطبیق ساختار هندسی نمونه‌های دامنه منبع و هدف بر اساس اصل منیفولد⁵ می‌پردازد. این روش علاوه بر تطبیق توزیع هندسی دامنه‌ها، موجب بهبود تطبیق توزیع هندسی در کلاس‌های مختلف و افزایش طبقه‌بندی می‌شود. روش دیگری با عنوان VDA⁶ [33] است. در این روش نمایش کم‌بعد از دامنه‌های منبع و هدف ایجاد شده که در آن، علاوه بر کاهش اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌ها، از روش خوشه‌بندی مستقل از دامنه نیز برای تفکیک‌پذیری بین طبقه‌های مختلف استفاده می‌شود. روش JDA⁷ [5] به طور هم‌زمان توزیع احتمال‌های حاشیه‌ای و شرطی را در یک فرایند کاهش بعد تطبیق می‌دهد. بدین ترتیب، JDA یک نمایش خصوصیات جدیدی ایجاد می‌کند که در مقابل اختلاف توزیع‌های اساسی مؤثر و قوی است.

(۲) روش تطبیق خصوصیات مبتنی بر زیرفضا. این روش‌ها، برخلاف روش‌های مبتنی بر داده، به دنبال یافتن یک زیر فضای مشترک تطبیق‌پذیر نبوده، بلکه از خصوصیات مرتبط در هر دامنه برای ایجاد زیر فضاهایی مجزا برای دامنه‌ها بهره می‌برند. از جمله این روش‌ها، روش تطبیق هم‌زمان آماری و هندسی برای تطبیق دامنه‌های بصری (JGSA)⁸ [21] است. این روش زیرفضاهای مجزا برای هر دامنه پیدا می‌کند. در ایجاد این زیرفضاها، علاوه بر بهره‌گیری از اطلاعات خصوصیات مختص هر دامنه، از خصوصیات مشترک بین دامنه‌ای نیز بهره گرفته می‌شود. روش LRSR⁹ [2] از یک ماتریس انتقال برای انتقال دادن دامنه‌های منبع و هدف به یک زیرفضای مشترک استفاده می‌کند که هر نمونه هدف می‌تواند با ترکیبی از نمونه‌های منبع نشان داده شود، به طوری که نمونه‌ها از دامنه‌های مختلف

⁴ Coupled local-global adaptation

⁵ Manifold

⁶ Visual domain adaptation

⁷ Joint distribution adaptation

⁸ Joint geometrical and statistical alignment

⁹ Low-rank and sparse representation

¹ Joint distribution based adaptive classifier

² Transfer joint matching

³ Geodesic flow kernel

می‌توانند در هم آمیخته شوند. روش ¹CDDA [3] یک نمایش خصوصیات نهان ایجاد می‌کند که دارای دو خصوصیات زیر است. اختلاف توزیع دامنه‌های منبع و هدف در قالب اختلاف توزیع‌های حاشیه‌ای و شرطی اندازه‌گیری می‌شود که موجب کاهش اختلاف نمونه‌های دو دامنه می‌شود. همچنین از یک نیروی دافعه برای بیشینه‌کردن اختلاف هر برجسب مرتبط با زیردامنه‌ها استفاده می‌شود.

روش پیشنهادی در این مقاله در دسته روش‌های تطبیق خصوصیات مبتنی بر داده قرار دارد که از دو دیدگاه تطبیق توزیع عمومی و تطبیق توزیع هندسی مورد بررسی قرار گرفته است. گفتنی است که بسیاری از مطالعات انجام‌شده در این حوزه، هرکدام از دسته‌بندی تطبیق توزیع عمومی و تطبیق توزیع هندسی را به‌صورت جداگانه مورد مطالعه قرار داده‌اند؛ اما در این مقاله، به‌طور هم‌زمان هر دو دسته‌بندی مورد توجه قرار گرفته است به‌طوری‌که در تطبیق توزیع عمومی، ابتدا برای به کمینه‌رساندن اختلاف توزیع شرطی و حاشیه‌ای، یک نمایش جدید از داده‌های دامنه منبع و هدف ایجاد می‌شود تا داده‌ها تطبیق‌پذیری دقیق‌تری باهم داشته باشند؛ سپس در تطبیق توزیع هندسی، با حفظ ساختار اولیه دامنه‌ها، داده‌ها به یک فضای نهان نگاشت می‌شوند تا در آن فضا، بین طبقه‌های مختلف، تفکیک‌پذیری ایجاد شود. در مرحله نهایی برای تطبیق توزیع هندسی بهتر با بهره‌گیری از کرنل، داده‌ها به فضای کرنل هیلبرت نگاشت شده و در آن فضا با یکدیگر تطبیق می‌شوند.

۳- روش پیشنهادی

در این بخش، روش IMAKE برای حل مسأله یادگیری انتقالی بدون نظارت، با جزئیات بیشتر توضیح داده می‌شود.

۳-۱- هدف پژوهش

این مقاله به دنبال آن است که روشی برای دسته‌بندی تصاویر ارائه کند تا کاربران با اطمینان از برجسب نسبت داده‌شده به هر تصویر، بتوانند در کاربردهای مختلف از آن تصویر استفاده کنند. بدین ترتیب، با استفاده از روش‌های تطبیق خصوصیات، یک نمایش کم‌بعد از داده‌های دامنه‌های منبع و هدف ایجاد شده که در این نمایش، اختلاف توزیع شرطی و حاشیه‌ای بین دامنه‌های منبع و

هدف، کمینه می‌شود؛ علاوه بر آن، با بهره‌گیری از تطبیق توزیع هندسی و همچنین خوشه‌بندی هندسی-کلاسی، داده‌ها به‌گونه‌ای به فضای نهان نگاشت می‌یابند که فاصله داده‌های متعلق به طبقه‌های یکسان کاهش یافته و اختلاف بین طبقه‌های مختلف افزایش یابد. در نتیجه، مدل در پیش‌بینی برجسب داده‌های دامنه هدف عملکرد بهتری خواهد داشت. همچنین به دلیل وجود داده‌هایی با هندسه‌های متفاوت و اغلب نامتوازن از همدیگر، روش خوشه‌بندی مطرح‌شده به‌تنهایی برای تراز کردن این داده‌ها به‌طوری‌که باعث ایجاد تفکیک‌پذیری قابل قبولی بین طبقه‌های مختلف شود، کافی نیست. به همین جهت برای بهبود مدل ساخته‌شده بر روی داده‌های با شرایط یادشده از کرنل به جهت تطبیق توزیع هندسی دقیق‌تر بهره گرفته‌شده است. هدف از این پژوهش ارائه راه‌حلی است که بتواند بر محدودیت‌های الگوریتم‌های کلاسیک یادگیری ماشین غلبه کرده و بازدهی روش‌های یادگیری انتقالی را افزایش دهد.

۳-۲- تعریف مسأله

در این بخش دو مفهوم دامنه و وظیفه را معرفی کرده و در ادامه به تشریح مسئله می‌پردازیم.

دامنه. هر دامنه D شامل دو مفهوم کلی فضای خصیصه‌ای X و توزیع احتمال حاشیه‌ای $P(x)$ برای هر $x \in X$ است؛ یعنی $D = \{X, P(x)\}$. بدین ترتیب، اگر دو دامنه متفاوت باشند، ممکن است فضای خصیصه‌ای مختلف و یا توزیع احتمال حاشیه‌ای متفاوت از یکدیگر داشته باشند. به‌طور دقیق‌تر، اگر X_s فضای خصیصه‌ای دامنه منبع، X_t فضای خصیصه‌ای دامنه هدف و $P_s(X_s)$ و $P_t(X_t)$ به ترتیب، توزیع احتمال حاشیه‌ای دامنه‌های منبع و هدف باشند (برای هر نمونه $x_s \in X_s$ و $x_t \in X_t$)، دو دامنه زمانی متفاوت هستند که $X_s \neq X_t$ یا $P_s(X_s) \neq P_t(X_t)$.

وظیفه. برای هر دامنه D ، یک وظیفه T شامل مجموعه برجسب‌های Y و تابع پیش‌بینی $f(x)$ وجود دارد که به صورت $T = \{Y, f(x)\}$ نشان داده می‌شود. تابع پیش‌بینی $f(x)$ ، بازای مجموعه نمونه ورودی X مجموعه برجسب‌های Y را پیش‌بینی می‌کند که احتمال شرطی آن به صورت $P(Y|x)$ تعریف می‌شود؛ بنابراین، اگر دو وظیفه متفاوت باشند، ممکن است، مجموعه برجسب‌های مختلفی داشته باشند یا توزیع احتمالی شرطی آن‌ها

¹ Close yet distinctive domain adaptation

متفاوت از یکدیگر باشد، بدین معنی که $Y_s \neq Y_t$ یا $P_s(Y|x) \neq P_t(Y|x)$.

مسئله یافتن توابع، نگاشتی است که به طور هم زمان داده های دامنه منبع و هدف را به فضای نهان نگاشت کند به طوری که داده ها در آن فضا، تطبیق توزیع هندسی دقیق تری داشته باشند و طبقه های مختلف به خوبی از هم تفکیک شوند. همچنین موجب تطبیق توزیع عمومی داده ها و افزایش کارایی طبقه بند شود.

۳-۳-۳-۳-۳ تطبیق توزیع هندسی

در این مقاله برای تطبیق توزیع هندسی از روش هم ترازی منیفلد بهره گرفته شده است. هم ترازی منیفلد یکی از روش های کاهش بعد محسوب می شود، به طوری که یک چارچوب هندسی برای ساخت فضای کم بعد فراهم می کند. این روش ابتدا هر مجموعه داده ورودی متعلق به دامنه های منبع و هدف را به شکل گراف مدل می کند، که به ساختارهای هندسی حاصل از این مرحله، منیفلد گفته می شود؛ سپس همه منیفلد های حاصل، هم زمان به یک فضای نهان کم بعد نگاشت می شوند. فرض کنید K مجموعه داده ورودی داریم که نمونه های این مجموعه داده ها متعلق به c طبقه مختلف هستند. اگر $X_k = (x_k^1, \dots, x_k^{m_k})$ نشان دهنده k امین مجموعه داده ورودی باشد، در حالی که i امین نمونه x_k^i دارای p_k برچسب است. X_k را می توان به شکل یک ماتریس $p_k \times m_k$ نشان داد. اگر برچسب های مجموعه داده X_k به صورت $V_k = (v_k^1, \dots, v_k^{l_k})$ نشان داده شود، آن گاه اگر X_k مربوط به دامنه منبع باشد l_k مقداری بزرگ و اگر X_k مربوط به دامنه هدف باشد l_k مقدار کوچکی خواهد بود. در اینجا فرض می شود $X_1, \dots, X_k$ مجموعه های گسسته هستند.

هدف این مسئله به دست آوردن K تابع نگاشت $f_1, \dots, f_k$ است، تا K مجموعه داده ورودی را به فضای نهان d بعدی نگاشت کند. به طوری که (۱) هندسه و ساختار اولیه هر دامنه ورودی حفظ شود، (۲) داده های متعلق به طبقه یکسان، در فضای کم بعد، نزدیک به هم قرار بگیرند، (۳) داده هایی که متعلق به طبقه های متفاوتی هستند، به خوبی از یکدیگر تفکیک شوند. برای دستیابی به این هدف، ابتدا یک نمایش واحد، حاصل از ترکیب منیفلد های مجموعه داده های ورودی ایجاد می کنیم. هر منیفلد با استفاده از یک ماتریس لاپلاسیان^۱ نمایش داده می شود که این ماتریس از روی گراف هر مجموعه داده و با معیار اتصال نمونه های نزدیک به هم ساخته می شود.

¹ Laplacian matrix

اطلاعات برچسب دار، نقش مهمی در ساختن این گراف های همسایگی ایفا می کنند؛ به این ترتیب نمونه هایی با برچسب یکسان در همسایگی یکدیگر قرار گرفته و نمونه هایی با برچسب غیر یکسان از هم فاصله می گیرند. در ادامه، منیفلد حاصل از پیوست تک تک منیفلد ها، دارای خصیصه های زائد خواهد بود، به همین جهت برای از بین بردن این خصیصه های تکراری، منیفلد پیوستی به فضای کم بعد نگاشت می شود. در نهایت، مسئله حاصل از طریق تجزیه مقادیر ویژه ی تعمیم یافته حل می شود.

۳-۳-۳-۱-۳-۳ نمادها

ابتدا ماتریس شباهت W_s ، ماتریس عدم شباهت W_d ، ماتریس های قطری D_s, D_d و ماتریس های لاپلاسیان L_s و L_d را تعریف می کنیم.

ماتریس شباهت $W_s = \begin{pmatrix} W_s^{1,1} & \dots & W_s^{1,K} \\ \dots & \dots & \dots \\ W_s^{K,1} & \dots & W_s^{K,K} \end{pmatrix}$ دارای ابعاد $(m_1 + \dots + m_K) \times (m_1 + \dots + m_K)$ می باشد. در اینجا ماتریس $W_s^{a,b}$ دارای ابعاد $m_a \times m_b$ است و زمانی که $W_s^{a,b}(i,j) = 1$ است که x_a^i و x_b^j متعلق به طبقه یکسان باشند، در غیر این صورت $W_s^{a,b}(i,j) = 0$ خواهد بود. ماتریس قطری $D_s(i,i) = \sum_j W_s(i,j)$ جمع سطری های ماتریس مربوطه حاصل می شود و ماتریس گراف لاپلاسیان ترکیبی به صورت $L_s = D_s - W_s$ قابل محاسبه است.

ماتریس عدم شباهت $W_d = \begin{pmatrix} W_d^{1,1} & \dots & W_d^{1,K} \\ \dots & \dots & \dots \\ W_d^{K,1} & \dots & W_d^{K,K} \end{pmatrix}$ نیز دارای ابعاد $(m_1 + \dots + m_K) \times (m_1 + \dots + m_K)$ است. در این ماتریس $W_d^{a,b}$ دارای ابعاد $m_a \times m_b$ است. زمانی که $W_d^{a,b}(i,j) = 1$ است که x_a^i و x_b^j از دو طبقه متفاوت باشند و در غیر این صورت $W_d^{a,b}(i,j) = 0$ خواهد بود. ماتریس قطری $D_d(i,i) = \sum_j W_d(i,j)$ جمع سطری های ماتریس مربوطه حاصل و ماتریس گراف لاپلاسیان ترکیبی به صورت $L_d = D_d - W_d$ محاسبه می شود.

برای نشان دادن توپولوژی هر دامنه، سه ماتریس D_k, W_k و L_k تعریف می شود که $W_k(i,j)$ نشان دهنده شباهت بین x_k^i و x_k^j است و از طریق $e^{-\|x_k^i - x_k^j\|^2}$ محاسبه می شود. همچنین ماتریس قطری $D_k(i,i) = \sum_j W_k(i,j)$ و ماتریس گراف لاپلاسیان ترکیبی برابر $L_k = D_k - W_k$ است.

همچنین توابع L و Z به صورت زیر تعریف می شوند:

هدف به‌دست‌آوردن K تابع نگاشت $f_1, \dots, f_k$ است، تا K مجموعه‌داده ورودی را به فضای نهان d بعدی نگاشت کند؛ در نتیجه مسئله به‌صورت زیر خواهد بود:

$$\begin{aligned} \{f_1, \dots, f_k\} \\ = \operatorname{argmin}_{f_1, \dots, f_k} (C(f_1, \dots, f_k)) \\ = \operatorname{argmin}_{f_1, \dots, f_k} \left(\frac{A + C}{B} \right) \end{aligned} \quad (5)$$

که اگر $\gamma = (f_1^T, \dots, f_k^T)^T$ ماتریسی با ابعاد $(p_1 + \dots + p_k) \times d$ بوده و نمایش‌دهنده K تابع نگاشتی باشد که می‌خواهیم به‌دست آوریم. مسئله، کمینه‌کردن تابع هدف به این صورت است که فضای تعبیه‌شده‌ای که تابع هدف $C(f_1, \dots, f_k)$ را کمینه می‌کند از مسأله بردارهای ویژه مرتبط با کوچک‌ترین مقادیر ویژه غیرصفر، از تجزیه مقادیر ویژه تعمیم یافته $\lambda Z L_d Z^T x = \lambda Z L_s Z^T x$ به‌دست می‌آید.

با در نظر گرفتن $A = \operatorname{tr}(\gamma^T Z L_s Z^T \gamma)$ ، $B = \operatorname{tr}(\gamma^T Z L_d Z^T \gamma)$ ، $C = \operatorname{tr}(\gamma^T Z \mu L Z^T \gamma)$ در نتیجه داریم:

$$\begin{aligned} \operatorname{argmin}_{f_1, \dots, f_k} (C(f_1, \dots, f_k)) = \\ \operatorname{argmin}_{f_1, \dots, f_k} \left(\frac{\gamma^T Z (L_s + \mu L) Z^T \gamma}{\gamma^T Z L_d Z^T \gamma} \right) \end{aligned} \quad (6)$$

در نهایت مسئله بردارهای ویژه، متعلق به کمترین مقادیر ویژه تعمیم‌یافته، به‌صورت رابطه (۷) خواهد بود:

$$Z(\mu L + (1 - \delta)L_s)Z^T x = \lambda Z \delta L_d Z^T x \quad (7)$$

بدین ترتیب با استفاده از توابع نگاشت حاصل شده، داده‌ها به فضای نهان نگاشت می‌یابند و در آن فضا داده‌های طبقه‌های یکسان در نزدیکی هم قرار می‌گیرند و داده‌های با طبقه‌های متفاوت به‌خوبی از همدیگر فاصله می‌گیرند.

۳-۴- تطبیق توزیع عمومی

یک مسئله مهم در یادگیری انتقالی، یافتن یک راه‌حل برای کاهش اختلاف توزیع بین دامنه‌های منبع و هدف است، یک روش کاهش اختلاف توزیع بین دامنه‌های منبع و هدف، ایجاد یک نمایش مشترک بین دامنه‌ها است که در نمایش جدید، به‌طور هم‌زمان، اختلاف توزیع بین دامنه‌های منبع و هدف کمینه‌شده و در کنار آن تلاش می‌شود که ساختار اصلی داده‌های ورودی نیز حفظ شود. از جمله روش‌های پیشنهادشده در سال‌های اخیر برای کاهش اختلاف توزیع منبع و هدف، روش‌های مبتنی بر کاهش بُعد هستند [22، 23]. در این روش‌ها، داده‌ها از فضای اصلی به یک فضای نگاشت‌شده منتقل می‌شوند با

که $L = \begin{pmatrix} L_1 & 0 & \dots & 0 \\ \dots & \dots & \dots & \dots \\ 0 & \dots & 0 & L_k \end{pmatrix}$ با ابعاد $(m_1 + \dots + m_k) \times (m_1 + \dots + m_k)$ و $Z = \begin{pmatrix} X_1 & 0 & \dots & 0 \\ \dots & \dots & \dots & \dots \\ 0 & \dots & 0 & X_k \end{pmatrix}$ دارای ابعاد $(p_1 + \dots + p_k) \times (m_1 + \dots + m_k)$ است.

۳-۳-۲- تابع هزینه

سه پارامتر عددی A ، B ، C برای استفاده در تابع هزینه به شکل زیر تعریف می‌شوند:

$$\begin{aligned} A = 0.5 \sum_{a=1}^K \sum_{b=1}^K \sum_{i=1}^{m_a} \sum_{j=1}^{m_b} \|f_a^T x_a^i - f_b^T x_b^j\|^2 W_s^{a,b}(i, j) \end{aligned} \quad (1)$$

که اگر x_a^i و x_b^j از یک طبقه یکسان باشند، اما در فضای تعبیه‌شده در موقعیت‌های دور از هم قرار گرفته باشند، آن‌گاه مقدار A بزرگ خواهد بود. در این حالت با کمینه‌کردن مقدار A این نمونه متعلق به طبقه یکسان، در فضای جدید در موقعیت‌های مشابهی قرار خواهند گرفت:

$$\begin{aligned} B = 0.5 \sum_{a=1}^K \sum_{b=1}^K \sum_{i=1}^{m_a} \sum_{j=1}^{m_b} \|f_a^T x_a^i - f_b^T x_b^j\|^2 W_d^{a,b}(i, j) \end{aligned} \quad (2)$$

که اگر x_a^i و x_b^j متعلق به دو طبقه متفاوت باشند، اما در فضای تعبیه‌شده نزدیک به هم قرار گرفته باشند، آن‌گاه مقدار B کوچک خواهد بود؛ بنابراین با بیشینه‌کردن مقدار B نمونه‌های متعلق به طبقه‌های مختلف، در فضای جدید از هم، فاصله می‌گیرند:

$$\begin{aligned} C = 0.5 \mu \sum_{k=1}^K \sum_{i=1}^{m_k} \sum_{j=1}^{m_k} \|f_k^T x_k^i - f_k^T x_k^j\|^2 W_k(i, j) \end{aligned} \quad (3)$$

که اگر x_k^i و x_k^j متعلق به یک دامنه یکسان باشند، آن‌گاه مقدار $W_k(i, j)$ بزرگی خواهد بود. حال اگر در فضای تعبیه‌شده $f_k^T x_k^i$ و $f_k^T x_k^j$ در فضای جدید به‌خوبی از هم فاصله گرفته باشند، C مقدار بزرگی خواهد داشت؛ بنابراین با کمینه‌کردن مقدار C توپولوژی هر دامنه حفظ می‌شود. با توجه به مطالب مطرح‌شده تابع هزینه‌ای که باید کمینه شود به‌صورت زیر است:

$$C(f_1, \dots, f_k) = (A + C)/B \quad (4)$$

این شرط که هزینه بازنگاشت آن‌ها (بازگرداندن آن‌ها به فضای اصلی) کمینه باشد. از جمله روش‌های کاهش بعد، می‌توان به روش تحلیل اجزای اصلی (PCA) اشاره کرد [22]. در ادامه روش PCA با جزئیات بیشتر معرفی می‌شود.

۳-۴-۱- کاهش بعد

PCA، یک روش انتقال داده به یک فضای کم‌بعد است که ایده اصلی این روش، حفظ ساختار اصلی داده‌ها در فضای جدید است. بدین ترتیب، داده‌ها بر روی اجزای اصلی خود نگاشت می‌شوند که شامل اجزایی است که دارای پخشش (واریانس) بالایی هستند. از بین اجزای به‌دست‌آمده، اجزایی که دارای بیشینه واریانس باشند، به‌عنوان اجزای اصلی جهت نگاشت داده‌ها به فضای کم‌بعد استفاده می‌شوند. با فرض تعریف تابع نگاشت $A \in R^{m \times d}$ برای نگاشت نمونه‌ها به زیرفضای جدید، تابع هدف PCA به‌صورت رابطه (۸) تعریف می‌شود:

$$\max_{A^T A = I} \text{tr}(AX^T H X A^T) \quad (۸)$$

اگر $X = [X_s; X_t] \in R^{m \times (n_s + n_t)}$ کل داده‌های ورودی با m بعد در فضای اصلی و X_s و X_t به‌ترتیب داده‌های دامنه منبع و هدف هستند و تابع مرکزیت H ، به‌صورت $H = I_{n_s + n_t} - \frac{1}{n_s + n_t} \vec{1} \vec{1}^T$ ماتریس همانی و $\vec{1}$ برداری از یک‌هاست. کوواریانس داده‌ها در فضای اصلی، با بهره‌گیری از تابع مرکزیت، برابر با $X^T H X$ است که نشان‌دهنده اختلاف داده‌ها از میانگین کل نمونه‌ها بوده و استفاده از تابع مرکزیت مانع از پراکندگی داده می‌شود. همچنین در رابطه (۱) $A X^T H X A^T$ کوواریانس نمونه‌ها در فضای جدید بوده و عبارت $A A^T$ برای برقراری شرط وارون‌پذیری تابع نگاشت به این رابطه اضافه می‌شود. $\text{tr}(\cdot)$ نشان‌دهنده حاصل جمع عناصر قطر اصلی ماتریس است.

۳-۴-۲- کاهش بعد در راستای کاهش اختلاف توزیع حاشیه‌ای

یک مسئله مهم، کاهش اختلاف توزیع بین دامنه‌های منبع و هدف در فضای کم‌بعد است که این اختلاف توزیع شامل اختلاف توزیع حاشیه‌ای و شرطی است. در روش PCA، تنها به کاهش بعد فضای اصلی پرداخته شده و به کاهش اختلاف توزیع توجهی نشده است. برای محاسبه اختلاف توزیع بین دامنه‌های منبع و هدف از روش

^۱MMD استفاده می‌شود. در روش MMD به‌عنوان یک روش غیر پارامتری، برای محاسبه اختلاف توزیع دامنه‌ها، داده‌ها به فضای هیلبرت نگاشت می‌شوند. در این روش برای محاسبه اختلاف توزیع حاشیه‌ای در فضای اصلی، از اختلاف بین میانگین نمونه‌های دامنه منبع و هدف در فضای هیلبرت استفاده می‌شود. با فرض A به‌عنوان تابع نگاشت به فضای جدید، اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف با استفاده از رابطه (۹) تعریف می‌شود:

$$\begin{aligned} Mrg(X_s, X_t) &= \left\| \frac{1}{n_s} \sum_{i=1}^{n_s} A^T x_i - \frac{1}{n_t} \sum_{j=n_s+1}^{n_s+n_t} A^T x_j \right\|^2 \\ &= \text{tr}(A^T X M_0 X^T A) \end{aligned} \quad (۹)$$

که ماتریس M_0 ، ماتریس ضرایب MMD است. همچنین n_s و n_t به ترتیب، تعداد نمونه‌های دامنه منبع و هدف می‌باشند. اگر $x_i, x_j \in X_s$ باشد، $(M_0)_{ij} = \frac{1}{n_s n_s}$ ، اگر $x_i, x_j \in X_t$ باشد، $(M_0)_{ij} = \frac{1}{n_t n_t}$ و درغیراین صورت، $(M_0)_{ij} = -\frac{1}{n_s n_t}$ است.

۳-۴-۳- کاهش بُعد در راستای کاهش اختلاف توزیع شرطی

یادگیری انتقالی به‌صورت انتقال‌دهنده، نوعی از یادگیری است که در آن، از نمونه‌های مشاهده‌شده از یک مجموعه آموزشی خاص، برای یک مجموعه آزمون خاص، مدل ساخته می‌شود. این نوع یادگیری برخلاف دیگر مدل‌های یادگیری است که در آن‌ها مدل از نمونه‌های آموزشی مشاهده شده با قواعد کلی، برای اعمال به داده‌های تست ساخته می‌شود. در یادگیری انتقالی انتقال‌دهنده، نمونه‌های برچسب دار آموزشی و نمونه‌های بدون برچسب آزمون در دسترس هستند. با این حال، با وجود اینکه برچسب داده‌های آزمون در دسترس نیستند، اما الگوها و اطلاعات اضافی موجود در داده می‌تواند کمک شایانی در ساخت مدل موردنظر داشته باشد. بدین ترتیب، برای استفاده بیشینه‌ای از پتانسیل موجود در داده‌های تست، از شبه‌برچسب‌ها برای برچسب‌گذاری داده‌های تست استفاده می‌شود. شبه‌برچسب‌ها عملاً برچسب‌هایی هستند که خیلی از صحت آنها اطمینان نداریم، ولی می‌توانیم با تکرار و انطباق دامنه‌های آموزشی و آزمون، آنها را در هر مرحله نسبت به مرحله قبل بهبود ببخشیم.

^۱ Maximum mean discrepancy (MMD)

در شرایطی که اختلاف توزیع بین دامنه‌های منبع و هدف زیاد باشد، برای عملکرد بهتر مدل یادگیری و تطبیق‌پذیری بیشتر مدل، اختلاف توزیع شرطی بین دامنه‌ها نیز باید کمینه شود. برای محاسبه اختلاف توزیع شرطی با بهره‌گیری از روش MMD، مجموع اختلاف میانگین نمونه‌های هر طبقه در دامنه‌های منبع و هدف، به‌صورت زیر محاسبه می‌شود:

$$Cnd(X_s, X_t) = \left\| \frac{1}{n_s^c} \sum_{x_i \in X_s^c} A^T x_i - \frac{1}{n_t^c} \sum_{x_j \in X_t^c} A^T x_j \right\|^2 \quad (10)$$

$$= \text{tr}(A^T X M_c X^T A)$$

که X_s^c و X_t^c به‌ترتیب، داده‌های دامنه منبع و دامنه هدف در طبقه c و n_s^c و n_t^c به‌ترتیب، تعداد نمونه‌های دامنه منبع و هدف در طبقه c هستند. اگر $x_i, x_j \in X_s^c$ باشد، $(M_c)_{ij} = \frac{1}{n_s^c n_t^c}$ و اگر $x_i, x_j \in X_t^c$ باشد، آن‌گاه $(M_c)_{ij} = -\frac{1}{n_s^c n_t^c}$ در غیر این صورت، $(M_c)_{ij} = 0$ است. با اضافه کردن روابط (۹) و (۱۰) به تابع هزینه در رابطه (۶) خواهیم داشت:

$$\text{argmin}_{f_1, \dots, f_k} (C(f_1, \dots, f_k)) = \text{argmin}_{f_1, \dots, f_k} \left(\frac{\gamma^T Z (\delta L + \mu L_s + M_0 + M_c) Z^T \gamma}{\gamma^T Z L_d Z^T \gamma} \right) \quad (11)$$

و درنهایت با استفاده از روش ضریب لاگرانژ رابطه (۱۲) حاصل می‌شود.

$$Z(\delta L + \mu L_s + M_0 + M_c) Z^T x = \lambda Z L_d Z^T x \quad (12)$$

در رابطه حاصل‌شده در ضمن هم‌ترازی داده‌های منبع و هدف در فضای نهان، با کاهش بُعد در راستای کاهش اختلاف توزیع حاشیه‌ای و شرطی، داده‌های طبقه‌های مختلف دارای تفکیک‌پذیری دقیق‌تری خواهند بود.

۳-۵- کرنل سازی

در کاربردهای یادگیری ماشین، منظور از کرنل، به‌طورعمومی تابع (نیمه) معین مثبتی است که می‌توان آن را به‌عنوان یک معیار سنجش شباهت بین دو نمونه در نظر گرفت. درواقع، شباهت میان زوج نمونه‌های آموزشی $S = \{x_i\}_{i=1}^n$ به‌وسیله یک مقدار حقیقی به‌دست‌آمده از توابع کرنل مشخص می‌شود. تابع کرنل به‌صورت $K_{ij} = \langle \phi(x_i), \phi(x_j) \rangle$ تعریف می‌شود که در آن $\phi: X \rightarrow H$ نگاشت خطی از فضای اولیه به فضای H (فضای هیلبرت) است و $\langle u, v \rangle$ ویژگی ناشی از کرنل

نشان‌دهنده ضرب داخلی بین دو بردار داده دلخواه u و v است؛ بنابراین محاسبه تابع کرنل را می‌توان به‌مثابه انجام عملیات ضرب نقطه‌ای دو داده در فضای هیلبرت (فضای ویژگی) متناظر با آن کرنل در نظر گرفت و این‌یکی از مهم‌ترین خصوصیات توابع کرنل است که سبب معرفی ترفند کرنل شده است. فرض کنید تابع کرنل k داده‌شده باشد، به ماتریس K که اندازه آن $n \times n$ بوده و هر المان آن با رابطه $k_{i,j} = k(x_i, x_j)$ به‌دست می‌آید، ماتریس کرنل حاصل از تابع کرنل k گفته می‌شود؛ بنابراین با توجه به ترفند کرنل نیازی به دانستن تابع نگاشت ϕ نبوده و انتخاب تابع کرنل برای نگاشت داده‌ها کافی است.

با کرنلایز کردن رابطه (۱۲)، ابتدا هر دامنه منبع با استفاده از تابع نگاشت ϕ جداگانه به فضای هیلبرت کرنل برده می‌شوند؛ سپس تمام دامنه‌ها تحت روش هم‌ترازی منیفلد به فضای نهان به‌صورت هم‌زمان نگاشت می‌شوند و در آن فضا باهم انطباق یافته و خوشه‌بندی هندسی-کلاسی روی داده‌های نگاشت شده، جهت تفکیک‌پذیری بهتر بین طبقه‌ای، اعمال می‌شود. به همین منظور در اینجا، ابتدا یک تئوری را مطرح می‌کنیم.

تئوری ۱. مجموع مستقیم فضاهای هیلبرت: اگر فرض کنیم دو فضای هیلبرت H_1 و H_2 وجود داشته باشند، به‌طوری‌که دو نمونه $\{x, y\}$ به‌صورت $x \in H_1$ و $y \in H_2$ باشد، دراین‌صورت فضای هیلبرت H وجود دارد که $\langle \{x_1, y_1\}, \{x_2, y_2\} \rangle = \langle x_1, x_2 \rangle_{H_1} + \langle y_1, y_2 \rangle_{H_2}$ به این قضیه مجموع مستقیم فضاهای هیلبرت گفته می‌شود و با $H = H_1 \oplus H_2$ نشان داده می‌شود. این قضیه را می‌توان به D فضای هیلبرت گسترش داد که به‌صورت $H = \bigoplus_{i=1}^D H_i$ می‌توان نشان داد.

حال ما ابزار ضروری برای کرنلایز کردن رابطه (۱۲) را داریم. ابتدا D دامنه مختلف به D فضای هیلبرت مختلف با ابعاد $D, \dots$ نگاشت می‌شوند $\phi_i(\cdot): x \mapsto \phi_i(x) \in H_i$ حال با جای‌گذاری همه نمونه‌ها با بردارهای ویژگی نگاشت‌شده آنها مسأله به‌صورت زیر تغییر می‌کند:

$$\Phi(\delta L + \mu L + M_0 + M_c) \Phi^T U = \lambda \Phi L_d \Phi^T U \quad (13)$$

که Φ یک ماتریس قطری شامل ماتریس داده‌ها به‌صورت $\Phi_i = [\phi_i(x_1), \dots, \phi_i(x_{n_i})]^T$ است و U ماتریسی است که هر سطر آن شامل بردارهای ویژه متعلق به هر دامنه جداگانه در فضای هیلبرت است $U = [u_1, u_2, \dots, u_H]^T$ به‌طوری‌که $H = \sum_{i=1}^D H_i$ با توجه به این‌که بردارهای ویژه

u_i دارای ابعاد نامحدودی هستند؛ بنابراین محاسبه مستقیم آن‌ها امکان‌پذیر نیست. به همین جهت از تئوری نمایش راز^۱ [24] به‌عنوان جایگزین استفاده می‌کنیم؛ بنابراین بردارهای ویژه می‌توانند به‌صورت ترکیب خطی از نمونه‌های نگاشت یافته نمایش داده شوند، به این صورت که $u_i = \Phi_i \alpha_i$ و ماتریسی که از این نگاشت حاصل خواهد شد برابر $U = \Phi \Lambda$ است، درنهایت مسأله به شکل زیر تبدیل می‌شود:

$$\Phi(\delta L + \mu L + M_0 + M_c) \Phi^T \Phi \Lambda = \lambda \Phi L_d \Phi^T \Phi \Lambda \quad (14)$$

حال با ضرب تقدیمی طرفین در Φ^T و جای‌گذاری ضرایب نقطه‌ای با ماتریس کرنل $K_i = \Phi_i^T \Phi_i$ ، رابطه نهایی حاصل می‌شود:

$$K(\delta L + \mu L + M_0 + M_c) K \Lambda = \lambda K L_d K \Lambda \quad (15)$$

که در رابطه (۱۵) ماتریس قطری K شامل ماتریس‌های کرنل K_i است.

هدف ما در این روش، یافتن توابع نگاشتی است که با انتقال داده‌های منبع و هدف به فضای نهان با استفاده از این توابع نگاشت، منیفولدهای حاصل از داده‌های هر دو دامنه منبع و هدف دارای تطبیق‌پذیری قابل قبولی بین کلاس‌های مختلف باشند. در رابطه نهایی حاصل شده، توابع نگاشت به‌دست‌آمده داده‌ها را به فضای نهان نگاشت می‌کنند؛ به‌طوری‌که در فضای جدید حاصل شده داده‌ها دارای تفکیک‌پذیری مناسبی بین طبقه‌های مختلف است و نمونه‌های متعلق به کلاس‌های یکسان در مجاورت هم قرار می‌گیرند.

۴- تنظیمات اولیه محیط آزمایش

در این بخش عملکرد و کارایی روش پیشنهادی IMAKE در مقایسه با الگوریتم‌های شناخته‌شده در حوزه یادگیری انتقالی به تفصیل بیان می‌شود.

۴-۱- معرفی مجموعه داده‌ها

کارایی روش پیشنهاد شده در این مقاله بر روی چهار پایگاه داده بصری شناخته‌شده مورد ارزیابی قرار گرفته است: (۱) آفیس [25] و کالتک [26]، (۲) اعداد، (۳) کوپل [27]، (۴) چهره (پای) [28].

پایگاه داده آفیس و کالتک که شامل تصاویر اشیای جمع‌آوری شده از دامنه‌های وبکم، آمازون، دی اس ال آر و

کالتک است که تصاویر در این دامنه‌ها از نظر شرایط نور و پس‌زمینه با یک‌دیگر تفاوت قابل‌توجهی دارند. دامنه آمازون، شامل تصاویر اشیای داندودشده از سایت‌های تجاری است. زمینه این تصاویر سفید بوده و اشیاء در مرکز آن‌ها قرار دارند و در شرایط نورپردازی استودیو تصویربرداری شده‌اند. دامنه وبکم، شامل تصاویر اشیاء با کیفیت پایین است که با دوربین وب گرفته شده و دامنه دی اس ال آر، شامل تصاویر اشیاء با کیفیت بالا است که با دوربین‌های دیجیتالی گرفته شده‌اند. مجموعه داده کالتک شامل ۳۰۶۰۷ تصویر و ۲۵۶ طبقه است که تصاویر این دامنه از وب سایت گوگل جمع‌آوری شده است. با این حال، در آزمایش‌هایی که ما طراحی کرده‌ایم از ۱۰ طبقه مشترک بین پایگاه داده‌های آفیس و کالتک استفاده می‌شود. برای این پایگاه داده، دوازده آزمایش طراحی شده است که در هر یک از آزمایش‌ها یکی از مجموعه داده‌ها، به‌عنوان دامنه منبع و یکی دیگر از مجموعه داده‌ها به‌عنوان دامنه هدف انتخاب می‌شوند.

پایگاه داده اعداد که شامل دو دامنه از تصاویر اعداد دست‌نویس انگلیسی صفر تا نه است. این پایگاه داده شامل دو مجموعه داده USPS [29] و MNIST [30] است. دامنه USPS، شامل حدود نه‌هزار تصویر با اندازه ۱۶×۱۶ پیکسلی است که ۷۲۶۱ داده آموزشی و ۲۰۰۷ داده آزمون دارد. دامنه MNIST، شامل هفتاد هزار تصویر با اندازه ۲۸×۲۸ پیکسلی است که شصت هزار داده آموزشی و ده هزار داده آزمون دارد. این مجموعه داده شامل ده طبقه مختلف است که با طراحی دو آزمایش، در یکی از آن‌ها، USPS دامنه منبع و MNIST دامنه هدف (U-M) و در آزمایش دیگر بالعکس (M-U) است، که کارایی الگوریتم در هر دو حالت مختلف مورد ارزیابی قرار گرفته است.

پایگاه داده کوپل، شامل ۱۴۴۰ تصویر سیاه و سفید از بیست شیء با زمینه سیاه در زاویه‌های مختلف است که هر تصویر در اندازه ۳۲×۳۲ پیکسلی نمایش داده می‌شود. پایگاه داده کوپل شامل دو دامنه کوپل یک و کوپل دو است که کوپل یک، مجموعه تصاویر اشیای گرفته‌شده در زاویه‌های [۰، ۸۵] و [۲۶۵، ۱۸۰] (ربع اول و سوم) و کوپل ۲، مجموعه تصاویر گرفته‌شده در زاویه‌های [۹۰، ۱۷۵] و [۲۷۰، ۳۵۵] (ربع دوم و چهارم) هستند. بدین ترتیب، وجود اختلاف توزیع بین دامنه‌های کوپل یک و کوپل دو، مشهود است. دامنه COIL1_vs_COIL2 (C1_C2) که ۷۲۰ نمونه از دامنه کوپل یک را به‌عنوان

¹ Riesz representation theory

² DSLR

برای مقایسه روش پیشنهادی با الگوریتم‌های شناخته‌شده در حوزه یادگیری انتقالی، دقت طبقه‌بند بر روی داده‌های دامنه هدف محاسبه می‌شود. این دقت به‌وسیله محاسبه خطای پیش‌بینی در دامنه هدف به‌صورت رابطه زیر تعریف می‌شود:

$$\text{Accuracy} = \frac{|\{x: x \in D_t \text{ and } f(x) = y(x)\}|}{n_t} \quad (16)$$

که D_t دامنه هدف، $f(x)$ تابع پیش‌بینی به‌دست‌آمده، $y(x)$ برچسب واقعی داده و n_t تعداد داده‌های دامنه هدف است. در روش پیشنهادی چهار پارامتر مختلف وجود دارد: $\lambda(1)$: پارامتر نسبت در رابطه $(\gamma, 2)$: تعداد ابعاد فضای جدید، $\mu(3)$: پارامتر نسبت در رابطه $(\gamma, 4)$: پارامتر نسبت در رابطه (γ) .

(جدول ۱-): مقدار بهینه پارامترها برای ۴ پایگاه‌داده بصری

(Table-1): Optimal value of parameters for 4 visual databases

پای	کویل	اعداد	آفیس و کالتک	
25	20	120	20	k
0.1	0.1	0.01	0.1	λ
0.9	0.1	0.2	0.5	μ
0.1	0.1	0.9	0.1	δ

روش IMAKE با مقادیر مختلف پارامترها در ۳۶ پایگاه داده همچون روش‌های [32, 33, 34, 35] مورد آزمایش قرار گرفته است. مقدار بهینه پارامترها برای پایگاه داده‌های مختلف در جدول (۱) نشان داده شده است. باتوجه به اینکه مدل پیشنهادی دارای چهار پارامتر است، یافتن بهینه پارامترها به‌صورت هم‌زمان ممکن نیست، اما با ثابت گرفتن سه پارامتر و بررسی نقطه بهینه یک پارامتر به‌صورت تکرارشونده می‌توان نقاط بهینه جواب را به‌دست دهد. این رویکرد در اغلب روش‌های یادگیری انتقالی مورد استفاده قرار می‌گیرد و به همین دلیل مقایسه نتایج را منصفانه می‌کند. درخصوص برچسب‌های نمونه‌های آزمون، در ابتدا شبه برچسب‌ها از طریق مدل اولیه‌ای که بر روی داده‌های آموزشی ایجاد می‌شود، به‌دست می‌آیند و سپس در هر مرحله، مدل دقیق‌ترشده و می‌تواند برچسب‌های درست‌تری را پیشنهاد کند. بدین ترتیب الگوریتم پیشنهادی می‌تواند به شکل تکرارشونده مقادیر بهینه هر پارامتر را به‌دست آورد. پیاده‌سازی روش پیشنهادی، به‌وسیله نرم‌افزار Matlab انجام گرفته است.

۵- نتایج و بحث‌ها

در این بخش، عملکرد روش IMAKE و الگوریتم‌های شناخته‌شده در حوزه یادگیری انتقالی مورد مقایسه و تحلیل قرار می‌گیرد.

داده آموزشی و ۷۲۰ نمونه از کویل دو را به‌عنوان داده تست شامل شده است، به‌عنوان آزمایش نخست از پایگاه داده کویل در نظر گرفته می‌شود. به‌طور مشابه COIL2_vs_COIL1 (C2_C1)، با جایه‌جایی نمونه‌های آموزشی و آزمون دامنه (C2_C1) ایجاد شده و به‌عنوان آزمایش دوم در نظر گرفته می‌شود.

مجموعه‌داده پای، از مجموعه پایگاه‌داده‌های شناخته‌شده برای مسأله تطبیق دامنه‌های بصری در زمینه تشخیص چهره است که شامل ۴۱۳۶۸ تصویر چهره از ۶۸ فرد مختلف است. پنج دامنه در این پایگاه داده وجود دارد که هر کدام مربوط به یک حالت تصویربرداری است: پای یک (حالت چپ) (P1)، پای دو (حالت بالا) (P2)، پای سه (حالت پایین) (P3)، پای چهار (حالت روبه‌رو) (P4)، پای پنج (حالت راست) (P5). در مجموع بیست آزمایش بین‌دامنه‌ای بر روی پایگاه داده پای قابل طراحی است که از بین پنج دامنه، دو دامنه مختلف به‌عنوان دامنه‌های منبع و هدف انتخاب می‌شوند. به‌طورکلی، کارایی الگوریتم پیشنهادی بر روی ۳۶ مجموعه تصاویر بین دامنه‌ای مختلف مورد ارزیابی قرار گرفته است که در بخش بعدی گزارش می‌شود.

۴-۲- ارزیابی الگوریتم

روش‌هایی که الگو ریت IMAKE با آن‌ها مقایسه شده است، عبارتند از:

طبقه‌بند نزدیک‌ترین همسایه (NN) [31]، GFK [19]، JDA [5]، TJM [18]، LRSR [2]، CDDA [3]، CLGA [20]. عملکرد روش پیشنهادی با بهترین نتایج گزارش شده از الگوریتم‌های مورد مقایسه، مورد ارزیابی قرار گرفته است. تمامی این الگوریتم‌های نامبرده شده به جز NN جزء روش‌های کاهش بُعد هستند که با استفاده از طبقه‌بند استاندارد NN روی داده‌های دامنه منبع برای پیش‌بینی برچسب داده‌های دامنه هدف مورد آزمایش قرار می‌گیرند. طبقه‌بند NN، در ابتدا فاصله اقلیدسی بین هر نمونه از دامنه هدف را نسبت به نمونه‌های دامنه منبع محاسبه کرده سپس با توجه به درجه همسایگی، برچسب نزدیک‌ترین نمونه از دامنه منبع را به‌عنوان برچسب هر نمونه از دامنه هدف اختصاص می‌دهد. عملکرد روش پیشنهادی IMAKE، با بهترین نتایج گزارش‌شده از الگوریتم‌های مورد مقایسه، مورد ارزیابی قرار گرفته است.

۴-۳- مفروضات پیاده‌سازی

(جدول-۲): صحت (%) طبقه‌بند در پایگاه‌داده آفیس و کالتک

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

IMAKE	CLGA (2018)	CDDA (2017)	LRSR (2016)	TJM (2014)	JDA (2013)	GFK (2012)	NN	Dataset
51.98	48.02	48.33	51.25	46.76	44.78	41.02	23.7	C-A
47.12	42.37	44.75	38.64	39.98	41.69	40.68	25.76	C-W
51.59	49.04	48.41	47.13	44.59	45.22	38.85	25.48	C-D
41.76	42.30	42.12	43.37	39.45	39.36	40.25	26	A-C
46.1	41.36	41.69	36.61	42.03	37.97	38.98	29.83	A-W
37.58	36.31	37.58	38.85	45.22	39.49	36.31	25.48	A-D
32.98	32.95	31.97	29.83	30.19	31.17	30.72	19.86	W-C
40.92	34.75	37.27	34.13	29.96	32.78	29.75	22.96	W-A
88.54	92.36	87.9	82.80	89.17	89.17	80.89	59.24	W-D
35.35	33.66	34.64	31.61	31.43	31.52	30.28	26.27	D-C
39.77	89.83	33.51	33.19	32.78	33.09	32.05	28.5	D-A
90.58	35.99	90.51	77.29	85.42	89.49	75.59	63.39	D-W
50.36	48.23	48.22	45.39	44.41	46.31	42.95	31.37	میانگین

(جدول-۳): صحت (%) طبقه‌بند در پایگاه‌داده پای

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

IMAKE	CLGA (2018)	CDDA (2017)	LRSR (2016)	TJM (2014)	JDA (2013)	GFK (2012)	NN	Dataset
75.32	67.83	60.22	65.87	23.87	58.81	26.15	26.09	P1-P2
78.55	63.85	58.7	64.09	28.86	54.23	27.27	26.59	P1-P3
94.02	88.95	83.48	82.03	43.37	84.5	31.15	30.67	P1-P4
72.73	61.76	54.17	54.90	19.3	49.75	17.59	16.67	P1-P5
76.59	71.40	62.33	45.54	26.14	57.62	25.24	24.49	P2-P1
77.7	72.98	64.64	53.49	37.93	62.93	47.37	46.63	P2-P3
91.41	86.24	79.9	71.43	50.53	75.82	54.25	54.07	P2-P4
67.16	51.23	44	47.97	21.63	39.89	27.08	26.53	P2-P5
82.68	70.17	58.46	52.49	28.66	50.96	21.82	21.37	P3-P1
74.4	73.48	59.73	55.56	35.97	57.95	43.16	41.01	P3-P2
93.3	89.31	77.2	77.50	51.97	68.45	46.41	46.53	P3-P4
71.08	55.51	47.24	54.11	25.31	39.95	26.78	26.23	P3-P5
95.85	89.56	83.1	81.54	45.71	80.58	34.24	32.95	P4-P1
96.19	92.94	82.26	58.39	57.58	82.63	62.92	62.68	P4-P2
95.96	93.08	86.64	82.23	71.63	87.25	73.35	73.22	P4-P3
82.84	71.63	58.33	72.61	30.94	54.66	37.38	37.19	P4-P5
76.02	57.68	48.02	52.19	27.13	46.46	20.35	18.49	P5-P1
64.95	55.43	45.61	49.41	22.65	42.05	24.62	24.1	P5-P2
71.94	58.03	52.02	58.45	28.86	53.31	28.49	28.31	P5-P3
81.65	71.85	55.99	64.31	32.59	57.01	31.33	31.24	P5-P4
81.02	72.15	63.1	61.53	35.53	60.24	35.35	34.76	میانگین

(جدول-۴): صحت (%) طبقه‌بند در پایگاه‌داده اعداد

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

IMAKE	CLGA (2018)	CDDA (2017)	LRSR (2016)	TJM (2014)	JDA (2013)	GFK (2012)	NN	Dataset
88.78	58.35	62.05	54.51	52.25	59.65	46.45	44.7	U-M
88.78	71.28	76.22	73.82	63.28	67.28	67.22	65.94	M-U
88.78	64.81	69.13	64.16	57.76	63.47	56.84	55.32	میانگین

(جدول-۵): صحت (%) طبقه‌بند در پایگاه‌داده کویل

(Table-5): classification accuracy (%) on COIL dataset

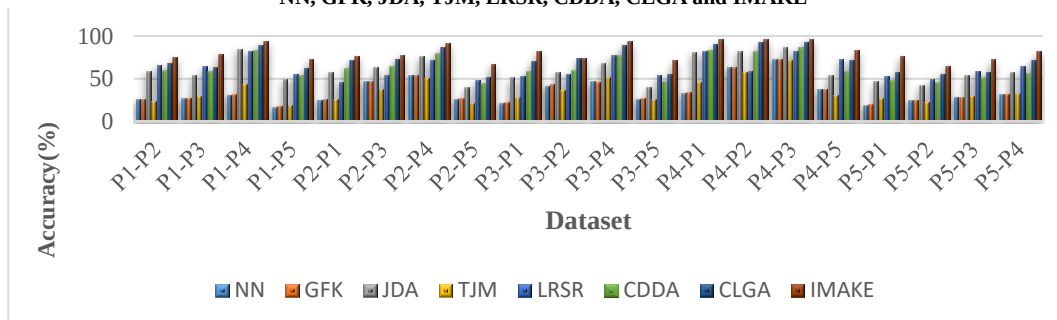
IMAKE	CLGA (2018)	CDDA (2017)	LRSR (2016)	TJM (2014)	JDA (2013)	GFK (2012)	NN	Dataset
100	96.81	91.53	88.61	91.67	89.31	72.5	83.61	C1-C2
100	91.11	93.89	89.17	91.53	88.47	74.17	82.78	C2-C1
100	93.96	92.71	88.89	91.6	88.89	73.34	83.2	میانگین



(شکل-۱): دقت طبقه‌بندی در پایگاه‌داده‌های آفیس و کالتک با استفاده از روش‌های مختلف

IMAKE, CLGA, CDDA, LRSR, TJM, JDA, GFK, NN

(Figure-1): Classification accuracy (%) on Office and Caltech-256 datasets using different methods
NN, GFK, JDA, TJM, LRSR, CDDA, CLGA and IMAKE



(شکل-۲): دقت طبقه‌بندی در پایگاه‌داده پای با استفاده از روش‌های مختلف

IMAKE, CLGA, CDDA, LRSR, TJM, JDA, GFK, NN

(Figure-2): Classification accuracy (%) on PIE datasets using different methods
NN, GFK, JDA, TJM, LRSR, CDDA, CLGA and IMAKE

۵-۱- ارزیابی نتایج

می‌کند. در ادامه عملکرد IMAKE نسبت به هر یک از روش‌های مورد مقایسه به تفکیک مورد بحث قرار گرفته است.

در روش GFK، که از جمله روش‌های تطبیق دامنه مبتنی بر ویژگی است، زیرفضاهایی از دامنه‌های منبع و هدف ایجاد می‌شود که در این زیرفضاها، اختلاف توزیع حاشیه‌ای بین داده‌های دامنه منبع و هدف به کمینه می‌رسد. در این روش، به علت اینکه در ایجاد زیرفضاها ابعاد اصلی کاهش می‌یابد، داده اصلی به درستی در زیرفضای جدید نمایش داده نمی‌شود. این در حالی است که در روش IMAKE، به طور هم‌زمان اختلاف توزیع شرطی و حاشیه‌ای بین دامنه‌های منبع و هدف کاهش می‌یابد و اطلاعات هندسی نمونه‌ها به وسیله روش لاپلاسی حفظ می‌شود. بدین ترتیب، متوسط دقت روش پیشنهادی در مقایسه با روش GFK، در پایگاه‌داده‌های آفیس و کالتک و پای به ترتیب $7/4\%$ و $45/67\%$ و همچنین در دو پایگاه‌داده اعداد و کوئل به ترتیب $31/94\%$ و $26/66\%$ بهبود عملکرد داشته است.

روش‌های TJM و JDA، از جمله روش‌های ایجاد تطبیق بین دامنه‌ها به وسیله ایجاد یک نمایش مشترک بین دامنه‌های منبع و هدف هستند. روش TJM به وسیله یک مسئله بهینه‌سازی پیچیده، اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف را کمینه می‌کند. روش

جدول‌های (۲ تا ۵)، نشان‌دهنده نتایج به دست آمده از روش IMAKE و الگوریتم‌های مورد مقایسه به ترتیب، بر روی پایگاه‌داده آفیس و کالتک، پایگاه‌داده پای و پایگاه‌داده‌های اعداد و کوئل است. در پایگاه‌داده آفیس و کالتک، IMAKE دارای $2/12\%$ متوسط بهبود نسبت به بهترین الگوریتم مورد مقایسه و دارای $18/98\%$ متوسط بهبود صحت نسبت به الگوریتم استاندارد NN بوده و در هشت پایگاه‌داده از دوازده پایگاه‌داده صحت IMAKE بهتر شده است. در مورد پایگاه‌داده‌های اعداد و کوئل، روش مورد بررسی به ترتیب دارای $19/65\%$ و $6/04\%$ نسبت به بهترین الگوریتم مورد مقایسه و همچنین به ترتیب $33/46\%$ و $16/8\%$ نسبت به الگوریتم استاندارد NN بهبود عملکرد دارد. همچنین IMAKE در هر چهار دامنه اعداد و کوئل، عملکرد بهتری از خود نشان می‌دهد. در پایگاه داده پای، متوسط بهبود صحت IMAKE، نسبت به بهترین الگوریتم مورد مقایسه و $8/88\%$ نسبت به الگوریتم استاندارد NN است و همچنین IMAKE در هر بیست دامنه از این پایگاه‌داده عملکرد بهتری نشان می‌دهد.

به طور کلی می‌توان نتیجه گرفت الگوریتم پیشنهادی، در هر چهار نوع پایگاه داده بصری، تطبیق‌پذیری بهتری بین دامنه‌های منبع و هدف ایجاد

JDA، به طور همزمان، اختلاف توزیع حاشیه‌ای و شرطی را کاهش می‌دهد و به همین دلیل، دارای عملکرد بهتری نسبت به روش TJM است. این در حالی است که تطبیق همزمان توزیع شرطی و حاشیه‌ای و هندسی IMAKE، در مقایسه با TJM به طور متوسط به بهبود $3/93\%$ و $45/49\%$ به ترتیب نسبت به پایگاه داده آفیس و کالتک و پای دست یافته و در پایگاه داده‌های اعداد و کوئل به ترتیب و به طور متوسط $31/01\%$ و $8/4\%$ بهبود یافته است. متوسط دقت IMAKE نسبت به JDA، بر روی پایگاه داده‌های آفیس و کالتک و پای به طور متوسط دارای $4/04\%$ و $28/54\%$ بهبود عملکرد و همچنین در پایگاه داده‌های اعداد و کوئل به ترتیب دارای $31/25\%$ و $11/11\%$ بهبود قابل ملاحظه‌ای است.

LRSR، با وجود کمینه‌سازی خطای بازسازی مجدد نمونه‌ها، به کاهش اختلاف بین دامنه‌ای توجهی نکرده است. این در حالی است که IMAKE، با کاهش همزمان اختلاف توزیع‌های عمومی و هندسی، موجب کاهش اختلاف توزیع نسبت به روش LRSR شده است. روش مطرح شده در این مقاله نسبت به روش LRSR در پایگاه داده آفیس و کالتک و پای به ترتیب دارای $4/96\%$ و $17/49\%$ بهبود عملکرد و در دو پایگاه داده اعداد و کوئل به ترتیب دارای $42/62\%$ و $15/62\%$ بهبود قابل توجه است.

CDDA، علاوه بر تطبیق توزیع آماری دامنه‌های منبع و هدف، طبقه‌های مختلف در زیرفضای کم بعد را از یکدیگر تفکیک می‌کند. IMAKE، با تطبیق توزیع عمومی با کاهش اختلاف توزیع دو دامنه منبع و هدف از تطبیق توزیع هندسی دامنه‌های منبع و هدف نیز بهره گرفته است. در مقایسه با روش CDDA، به طور متوسط بر روی پایگاه داده آفیس و کالتک و پایگاه داده پای به ترتیب $3/13\%$ و $17/92\%$ و در پایگاه داده‌های اعداد و کوئل دارای $19/65\%$ و $17/29\%$ بهبود عملکرد است.

CLGA، علاوه بر کاهش اختلاف توزیع حاشیه‌ای و شرطی از تفکیک پذیری هندسی برای تطبیق دو دامنه استفاده کرده است، این در حالی است که IMAKE علاوه بر تطبیق پذیری عمومی، در تطبیق پذیری هندسی خود علاوه بر استفاده از ماتریس شباهت و ماتریس عدم شباهت جهت تفکیک داده‌های موجود در طبقه‌های مختلف از ماتریس لاپلاس جهت حفظ ساختار اولیه مجموعه داده بهره می‌برد. به همین جهت IMAKE دارای عملکرد بهتری در هر چهار پایگاه داده مورد بررسی بوده و به ترتیب در پایگاه داده آفیس و کالتک و پای دارای

$2/12\%$ و $8/88\%$ بهبود عملکرد و همچنین در پایگاه داده‌های اعداد و کوئل به ترتیب و به طور متوسط دارای $23/97\%$ و $3/07\%$ عملکرد بهتر نسبت به روش یاد شده است.

روش IMAKE، علاوه بر کمینه کردن اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌های منبع و هدف در فضای جدید، با بهره‌گیری از توزیع هندسی داده‌ها، ابعاد تفکیک کننده طبقه‌ها در دامنه‌های منبع و هدف را با هم تطبیق داده و در نتیجه، صحت طبقه‌بند برای پیش‌بینی برچسب‌های داده‌های آزمون بهبود می‌یابد. برای جلوگیری از نگاشت ویژگی‌ها به بعدهای نامربوط از کمینه کردن واریانس داده‌های هدف در فضای نگاشت شده استفاده می‌شود. همچنین از آنجایی که برچسب‌های داده‌های منبع در دسترس هستند، از اطلاعات برچسب‌ها برای محدود کردن نمایش داده‌های منبع در بازنمایی جدید استفاده می‌شود. علاوه بر این، اختلاف بین دامنه‌ها با نزدیک هم آوردن زیرفضاهای منبع و هدف کاهش می‌یابد. بدین ترتیب، زیرفضای حاصل شده می‌تواند از تفکیک پذیری قابل قبولی برای طبقه‌بندی نمونه‌ها برخوردار باشد.

همان‌طور که از نتایج مشخص است، عیب کلی روش‌های زیرفضا محور اختلاف بین توزیع‌ها است که به طور صریح کاهش نمی‌یابد. با این حال، روش‌های داده محور، اختلاف توزیع را به صورت صریح کاهش می‌دهند. به طور کلی، یک تبدیل یک پارچه وجود ندارد که هم اختلاف توزیع را کاهش دهد و هم ویژگی‌های داده اصلی را حفظ کند؛ بنابراین روش پیشنهادی ما در مقایسه با هر دوی روش‌های داده محور و زیرفضا محور عملکرد بهتری از خود نشان داده است. همچنین ما، نسخه‌های اصلی و کرنل شده روش خود را بر روی مجموعه داده‌های مختلف اعمال کردیم. نتایج حاصل شده نشان می‌دهد که به طور میانگین روش پیشنهادی ما در حالات اصلی و کرنل شده عملکرد مشابهی دارد.

۶- نتیجه‌گیری و کارهای آینده

روش پیشنهاد شده در این مقاله، یک روش تطبیق دامنه بدون نظارت بوده که با هدف بهبود کارایی مدل‌های یادگیری در حوزه پردازش تصویر ارائه شده است. در این روش ابتدا با بهره‌گیری از هم‌ترازی منیفولد به عنوان یک روش تطبیق توزیع هندسی، نمونه‌های دامنه منبع و هدف به یک زیرفضای نهان مشترک نگاشت می‌شوند به طوری که داده‌های متعلق به طبقه‌های یکسان نزدیک

- Computer Science, journal article June 03 2019.
- [7] ا. قولنجی ج. طهمورث نژاد، "تطبیق دامنه‌های بصری با استفاده از تطبیق خصوصیات و مدل"، *مجله مهندسی برق دانشگاه تبریز*، ۱۳۹۸
- [8] J. Tahmoresnezhad and S. Hashemi, "A generalized kernel-based random k-samplesets method for transfer learning," *Iranian Journal of Science and Technology Transactions of Electrical Engineering*, vol. 39, no. E2, pp. 193-207, 2015.
- [9] M. Zandifar, S. Noori Saray, and J. Tahmoresnezhad, "Domain adaptation via Bregman divergence minimization", *Scientia Iranica*, vol. 10, no. 1, pp. 1-12, 2021.
- [10] C. Wang and S. Mahadevan, "Manifold Alignment Preserving Global Geometry," in *IJCAI*, pp. 1743-1749, 2013.
- [11] M. Mardani, and J. Tahmoresnezhad, "Cross-and multiple-domains visual transfer learning via iterative Fischer linear discriminant analysis", *Knowledge and Information Systems*, vol. 63, no. 8, pp. 2157-2188, 2021.
- [12] Y. Wu and Q. Ji, "Facial landmark detection: A literature survey," *International Journal of Computer Vision*, vol. 127, no. 2, pp. 115-142, 2019.
- [13] J. Li, K. Lu, Z. Huang, L. Zhu, and H. T. Shen, "Transfer independently together: a generalized framework for domain adaptation," *IEEE transactions on cybernetics*, no. 99, pp. 1-12, 2018.
- [14] M. Vrigkas, E. Kazakos, C. Nikou, and I. A. Kakadiaris, "Human activity recognition using robust adaptive privileged probabilistic learning", *Pattern Analysis and Applications*, vol. 24, no. 3, pp. 915-932, 2021.
- [15] L. Bruzzone and M. Marconcini, "Domain adaptation problems: A DASVM classification technique and a circular validation strategy," *IEEE transactions on pattern analysis and machine intelligence*, vol. 32, no. 5, pp. 770-787, 2010.
- [16] M. Long, J. Wang, G. Ding, S. J. Pan, and S. Y. Philip, "Adaptation regularization: A general framework for transfer learning," *IEEE Transactions on Knowledge and Data Engineering*, vol. 26, no. 5, pp. 1076-1089, 2014.
- [17] Y. Yang, X. Kong, B. Wang, K. Ren, and Y. Guo, "Steganalysis on Internet Images via Domain Adaptive Classifier," *Neurocomputing*, 2019.
- [18] M. Long, J. Wang, G. Ding, J. Sun, and P. S. Yu, "Transfer joint matching for unsupervised domain adaptation," in *Proceedings of the IEEE conference on*

به هم قرار گرفته و طبقه‌های مختلف از هم فاصله بگیرند در نتیجه تفکیک‌پذیری قابل قبولی بین طبقه‌های مختلف ایجاد می‌شود. همچنین این روش با حفظ منیفلد دامنه منبع و هدف موجب می‌شود تا جایگاه داده‌ها و ساختار اولیه دامنه‌ها تا حد امکان حفظ شود و سپس در این فضای نهان اختلاف توزیع حاشیه‌ای و شرطی بین نمونه‌ها کاهش می‌یابد. همچنین استفاده از کرنل و نگاشت داده‌ها به فضای هیلبرت تفکیک‌پذیری دقیق‌تری بین طبقه‌های مختلف ایجاد شده و برای داده‌هایی با ساختارهای پیچیده و نامتوازن مناسب است. روش پیشنهادی بر روی پایگاه‌داده‌های بصری شناخته‌شده مورد ارزیابی قرار گرفته و نتایج آزمایش‌ها بیان‌گر برتری IMAKE نسبت به روش‌های مطرح‌شده در حوزه تطبیق دامنه است.

به‌عنوان کارهای آینده، در نظر داریم از این مدل به‌عنوان طبقه‌بندی برای برچسب‌گذاری تصاویر در مسائل برخط استفاده کرده و همچنین این روش را برای مسائل ناهمگون^۱، بهینه‌سازی خواهیم کرد. بدین ترتیب کاربرد این روش در حوزه وسیع‌تری ارزیابی خواهد شد.

7- References

۷- مراجع

- [1] S. Noori Saray, and J. Tahmoresnezhad, "Iterative joint classifier and domain adaptation for visual transfer learning", *International Journal of Machine Learning and Cybernetics*, vol.13, no. 4, pp.947-961, 2022.
- [2] 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.
- [3] L. Luo, X. Wang, S. Hu, C. Wang, Y. Tang, and L. Chen, "Close yet distinctive domain adaptation," *arXiv preprint arXiv:1704.04235*, 2017.
- [4] D. Tuia and G. Camps-Valls, "Kernel manifold alignment for domain adaptation," *PloS one*, vol. 11, no. 2, p. e0148655, 2016.
- [5] M. Long, J. Wang, G. Ding, J. Sun, and S. Y. Philip, "Transfer feature learning with joint distribution adaptation," in *2013 IEEE International Conference on Computer Vision*, 2013: IEEE, pp. 2200-2207.
- [6] S. Rezaei and J. Tahmoresnezhad, "Discriminative and domain invariant subspace alignment for visual tasks," *Iran Journal of*

¹ Heterogeneous

- [33] J. Tahmoresnezhad and S. Hashemi, "Visual domain adaptation via transfer feature learning," *Knowledge and Information Systems*, vol. 50, no. 2, pp. 585-605, 2017.
- [34] Z. Khorshidpour, J. Tahmoresnezhad, S. Hashemi, and A. Hamzeh, "Domain invariant feature extraction against evasion attack," *International Journal of Machine Learning and Cybernetics*, vol. 9, no. 12, pp. 2093-2104, 2018.
- [35] M. Ahmadvand and J. Tahmoresnezhad, "Metric transfer learning via geometric knowledge embedding," *Applied Intelligence*, vol. 51, no 2, pp. 921-934, 2021.



جعفر طهمورث‌نژاد مدرک دکترای مهندسی کامپیوتر-هوش مصنوعی خود را در سال ۱۳۹۴ از دانشگاه شیراز دریافت کرد. ایشان در راستای فعالیت‌های علمی خود، در حال حاضر

به‌عنوان دانشیار دانشگاه صنعتی ارومیه درحال فعالیت هستند. 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



الهه شهروز دانشجوی کارشناسی ارشد مهندسی فناوری اطلاعات، مدرک کارشناسی خود را در رشته مهندسی فناوری اطلاعات در سال ۱۳۹۴ دریافت کرده‌اند.

نشانی رایانامه ایشان عبارت است از:

elahe_oskuie@it.uut.ac.ir

- [19] B. Gong, Y. Shi, F. Sha, and K. Grauman, "Geodesic flow kernel for unsupervised domain adaptation," in *2012 IEEE Conference on Computer Vision and Pattern Recognition*, 2012: IEEE, pp. 2066-2073.
- [20] J. Liu, J. Li, and K. Lu, "Coupled local-global adaptation for multi-source transfer learning," *Neurocomputing*, vol. 275, pp. 247-254, 2018.
- [21] J. Zhang, W. Li, and P. Ogunbona, "Joint geometrical and statistical alignment for visual domain adaptation," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2017, pp. 1859-1867.
- [22] I. Jolliffe, *Principal component analysis*. Springer, 2011.
- [23] 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.
- [24] F. Riesz and B. S.-. Nagy, "Functional Analysis, Frederick Ungar Pub," Co., New York, 1955.
- [25] K. Saenko, B. Kulis, M. Fritz, and T. Darrell, "Adapting visual category models to new domains," in *European conference on computer vision*, 2010: Springer, pp. 213-226.
- [26] G. Griffin, A. Holub, and P. Perona, "Caltech-256 object category dataset," 2007.
- [27] S. A. Nene, S. K. Nayar, and H. Murase, "Columbia object image library (coil-20)," 1996.
- [28] T. Sim, S. Baker, and M. Bsat, "The CMU pose, illumination, and expression (PIE) database," in *Proceedings of Fifth IEEE International Conference on Automatic Face Gesture Recognition*, 2002: IEEE, pp. 53-58.
- [29] J. J. Hull, "A database for handwritten text recognition research," *IEEE Transactions on pattern analysis and machine intelligence*, vol. 16, no. 5, pp. 550-554, 1994.
- [30] Y. LeCun, L. Bottou, Y. Bengio, and P. Haffner, "Gradient-based learning applied to document recognition," *Proceedings of the IEEE*, vol. 86, no. 11, pp. 2278-2324, 1998.
- [31] T. Cover and Peter Hart, "Nearest neighbor pattern classification," *IEEE transactions on information theory*, vol. 13, no.1, pp. 21-27, 1967.
- [32] S. Rezaei, J. Tahmoresnezhad, and V. Solouk, "A transductive transfer learning approach for image classification," *International Journal of Machine Learning*