

## تطبیق دامنه‌های بصری با استفاده از تطبیق خصوصیات و مدل

الهه قولنجی<sup>۱</sup>، دانشجوی کارشناسی ارشد؛ جعفر طهمورث نژاد<sup>۲</sup>، استادیار

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

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

**چکیده:** در اکثر الگوریتم‌های یادگیری ماشین، توزیع احتمالی داده‌های آموزشی و تست (دامنه‌های منبع و هدف) یکسان فرض شده است. این درحالی است که در مسائل دنیای واقعی، اغلب این فرض برقرار نبوده و موجب کاهش بازدهی مدل می‌شود. هدف روش‌های تطبیق دامنه، ایجاد یک مدل تطبیق‌پذیر بر روی داده‌های آموزشی است که دارای عملکرد قابل‌قبولی بر روی داده‌های تست باشد. در این مقاله، یک روش تطبیقی بدون نظارت دومرحله‌ای با بهره‌گیری از روش‌های تطبیق خصوصیات و تطبیق مدل پیشنهاد شده است. در مرحله اول، داده‌های دامنه‌های منبع و هدف به یک فضای مشترک که دارای حداقل اختلاف توزیع حاشیه‌ای و شرطی می‌باشد، نگاشت می‌شوند و سپس از خوشه‌بندی مستقل از دامنه برای ایجاد تفکیک‌پذیری کلاس‌های مختلف در دامنه منبع بهره گرفته می‌شود. در مرحله دوم، یک طبقه‌بند انطباقی با حداقل کردن خطای پیش‌بینی و حداکثر نمودن سازگاری هندسی بین دامنه‌های منبع و هدف ایجاد می‌شود. روش پیشنهادی، بر روی چهار نوع پایگاه‌داده بصری شناخته‌شده با ۳۶ آزمایش طراحی‌شده، مورد ارزیابی قرار گرفته است. نتایج به‌دست‌آمده، نشان‌دهنده بهبود قابل‌ملاحظه از عملکرد روش پیشنهادی در مقایسه با جدیدترین روش‌های حوزه یادگیری ماشین و یادگیری انتقالی است.

**واژه‌های کلیدی:** یادگیری انتقالی، تطبیق دامنه بصری، نمایش خصوصیات، خوشه‌بندی مستقل از دامنه، طبقه‌بند انطباقی.

## Visual Domains Adaptation via Feature and Model Matching

E. Gholenji<sup>1</sup>, MSc Student; J. Tahmoresnezhad<sup>2</sup>, Assistant Professor

1- Faculty of IT & Computer Engineering, Urmia University of Technology, Urmia, Iran, Email: elahe.gholenjy@it.uut.ac.ir

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

**Abstract:** In most machine learning algorithms, the distribution of training and test sets (source and target domains, respectively) are assumed the same. However, this condition is violated in many real world problems and the performance of model degrades as well. The aim of domain adaptation solution is to build an adaptive model on source data to have acceptable performance on target domain. In this paper, we propose an unsupervised two-phases approach which benefits from representation and model adaptation methods. In the first phase, source and target data are projected onto a common subspace on which the marginal and conditional distribution difference is minimized. Moreover, domain invariant clustering is exploited to discriminate between various classes of source data. In the second phase, an adaptation classifier is presented to minimize prediction error rate and maximize manifold adaptability across source and target domains. The proposed approach is evaluated on four visual benchmark datasets according to 36 designed experiments. The obtained results highlight the considerable performance of the proposed approach against other state-of-the-art machine learning and transfer learning methods.

**Keywords:** Transfer learning, Visual domain adaptation, Feature representation, Domain invariant clustering, Adaptive classifier.

تاریخ ارسال مقاله: ۱۳۹۵/۱۲/۲۹

تاریخ اصلاح مقاله: ۱۳۹۶/۰۶/۲۸

تاریخ پذیرش مقاله: ۱۳۹۶/۰۸/۲۶

نام نویسنده مسئول: جعفر طهمورث‌نژاد

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

## ۱- مقدمه

یادگیری ماشین، یکی از شاخه‌های پرکاربرد در زمینه هوش مصنوعی است که به ایجاد الگوریتم‌هایی می‌پردازد که بر اساس آن‌ها، سیستم‌ها توانایی یادگیری داشته باشند. این الگوریتم‌ها به سیستم اجازه استفاده و یادگیری از داده‌ها را برای بهبود عملکردهای مختلف می‌دهند. در اغلب الگوریتم‌های یادگیری ماشین فرض بر این است که داده‌های دامنه آموزشی و دامنه تست از توزیع یکسان تبعیت می‌کنند. درحالی‌که در مسائل بین دامنه‌ای<sup>۱</sup> و واقعی، این شرط برقرار نبوده و مدل طبقه‌بند<sup>۲</sup> ایجاد شده در دامنه آموزشی، صحت پایینی در پیش‌بینی برچسب‌های نمونه‌های دامنه تست خواهد داشت [۱]. برای نمونه، فرض کنید داده‌های آموزشی مجموعه تصاویر گرفته‌شده توسط یک دوربین گوشی موبایل و داده‌های تست، مجموعه‌ای از تصاویر یک دوربین دیجیتال حرفه‌ای باشند. در چنین شرایطی، اختلاف موجود بین تصاویر آموزشی و تست، از نظر کیفیت، نور، شفافیت، زاویه تصویر و خصوصیات دیگر، موجب ایجاد تغییر دامنه<sup>۳</sup> و کاهش بازدهی مدل در دامنه تست خواهد شد. برای حل این مشکل، روش یادگیری انتقالی<sup>۴</sup> با هدف ایجاد تطبیق<sup>۵</sup> بین دامنه‌های منبع و هدف پیشنهاد شده است. در یادگیری انتقالی، با کاهش اختلاف توزیع بین دامنه آموزشی و تست، مدل طبقه‌بند با تغییرات دامنه‌ها سازگاری خواهد یافت.

به‌طور کلی، یادگیری انتقالی یک روش حل مساله در شرایط تغییر دامنه‌ها است. در این روش، دانش کسب‌شده در یک حوزه یا دامنه، به یک حوزه یا دامنه مرتبط منتقل می‌شود تا هزینه یادگیری، کاهش و عملکرد الگوریتم یادگیری بهبود یابد. یادگیری انتقالی، می‌تواند در دو نوع نیمه نظارت‌شده<sup>۶</sup> و بدون نظارت<sup>۷</sup> بررسی شود. اگر دامنه منبع<sup>۸</sup> شامل داده‌های آموزشی و دامنه هدف<sup>۹</sup> شامل داده‌های تست باشد، در یادگیری انتقالی نیمه نظارت‌شده، تمام داده‌های دامنه منبع دارای برچسب<sup>۱۰</sup> می‌باشند ولی فقط تعداد محدودی از داده‌های دامنه هدف، دارای برچسب بوده که اغلب این داده‌ها برای ایجاد یک طبقه‌بند دقیق، مناسب نیستند. این درحالی 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.

در سال‌های اخیر، مساله یادگیری انتقالی کاربردهای زیادی در حوزه‌های مختلف یادگیری ماشین به‌خصوص بینایی ماشین داشته است. از جمله کاربردهای یادگیری انتقالی در این حوزه، آنالیز احساسات، دسته‌بندی متون، تشخیص چهره و تشخیص رویداد می‌باشند. در سیستم‌های بینایی ماشین، مشکل اصلی سیستم، کمبود داده برچسب‌دار در دامنه‌های آموزشی است. در چنین شرایطی، برای

ایجاد مدل طبقه‌بند برای دامنه هدف، راه‌حل، استفاده از داده‌های دامنه‌های مرتبط است [۲]. از آنجایی‌که نمونه‌های موجود در دامنه‌های مرتبط در شرایط متفاوتی جمع‌آوری شده‌اند و از نظر خصوصیات تصویر مانند نور، زاویه تصویر، کیفیت و پس‌زمینه باهم متفاوت می‌باشند، سیستم با مشکل تغییر دامنه مواجه خواهد شد. بروز چنین مشکلی سبب می‌شود که مدل ایجاد شده بر روی دامنه منبع (همان دامنه برچسب‌داری که با دامنه هدف مرتبط است)، کارایی غیرقابل قبولی در طبقه‌بندی نمونه‌های دامنه هدف داشته باشد. بدین ترتیب، چالش اصلی، ایجاد یک طبقه‌بند دقیق و تطبیق‌پذیر است که بتواند با حداقل خطا، نمونه‌های دامنه هدف را برچسب‌گذاری نماید.

اختلاف توزیع بین دامنه‌های منبع و هدف، شامل اختلاف در توزیع حاشیه‌ای و توزیع شرطی است. در شرایطی که دامنه‌های منبع و هدف دارای مجموعه خصوصیات یکسان باشند، اختلاف در احتمال وقوع مقادیر هرکدام از این خصوصیات در هر دامنه، موجب ایجاد اختلاف توزیع حاشیه‌ای بین دامنه‌ها خواهد شد. برای نمونه، در مورد دامنه‌های تصاویر، دامنه‌های منبع و هدف دارای خصوصیات یکسان مانند شرایط نور، کیفیت و پس‌زمینه می‌باشند ولی به علت اینکه احتمال وقوع هرکدام از این خصوصیات در دامنه‌ها متفاوت می‌باشند، اختلاف حاشیه‌ای بین دامنه‌ها ایجاد و تشدید می‌شود. توزیع شرطی به احتمال پیش‌بینی یک مجموعه برچسب به ازای یک مجموعه داده ورودی گفته می‌شود که معادل مفهوم تابع پیش‌بینی است. اگر دو دامنه منبع و هدف، دارای اختلاف توزیع شرطی باشند، تابع پیش‌بینی دامنه‌ها باهم متفاوت بوده و مجموعه برچسب متفاوت به ازای داده‌های یکسان از هر دو دامنه، پیش‌بینی خواهد شد. در بحث مربوط به یادگیری انتقالی در سیستم‌های بینایی ماشین، هر دو اختلاف توزیع شرطی و حاشیه‌ای بین دامنه‌ها دارای اهمیت بوده و مساله اصلی، کاهش هم‌زمان این اختلاف‌ها می‌باشد.

با این حال کاهش اختلاف توزیع بین دامنه‌های منبع و هدف، می‌تواند از طریق ایجاد یک طبقه‌بند انطباقی نیز انجام شود. در یک طبقه‌بند استاندارد، ابعادی که به عنوان تفکیک‌کننده کلاس‌های دامنه منبع انتخاب می‌شوند، به علت اختلاف در ساختار اصلی داده‌ها در دامنه‌های منبع و هدف، به درستی نمی‌توانند تفکیک‌کننده کلاس‌ها در دامنه هدف نیز باشند؛ بنابراین، طبقه‌بند ایجاد شده در دامنه منبع، دارای عملکرد ضعیفی در دامنه هدف است. یک طبقه‌بند انطباقی از طریق کاهش اختلاف توزیع دامنه‌های منبع و هدف، موجب سازگاری بیشتر طبقه‌بند ایجاد شده با ساختار اصلی دامنه هدف می‌شود.

روش پیشنهادی در این مقاله با عنوان تطبیق دامنه‌های بصری از طریق تطبیق خصوصیات و مدل<sup>۱۱</sup> (FMM)، یک روش دومرحله‌ای با بهره‌گیری از تطبیق خصوصیات و طبقه‌بندی انطباقی<sup>۱۲</sup> است. در مرحله اول، سعی شده است با استفاده از تطبیق خصوصیات، یک نمایش جدید از داده‌های دامنه‌های منبع و هدف ایجاد شود که در این نمایش، توزیع حاشیه‌ای و شرطی دامنه‌های آموزشی و تست

دامنه هدف را حداقل نماید [۴، ۵]. در نمایش تنک<sup>۱۵</sup> [۴]، نمونه‌های تست بر اساس تعدادی از نمونه‌های آموزشی که باعث حداقل شدن خطای بازسازی داده شوند، نمایش داده می‌شوند. روش انتخاب لندمارک [۵]، از جمله روش‌های مبتنی بر نمونه است که لندمارک‌ها، به نمونه‌هایی از دامنه منبع گفته می‌شوند که از نظر توزیع بیشترین مشابهت را با نمونه‌های دامنه هدف دارند. در این روش، از روش حداکثر اختلاف میانگین‌ها<sup>۱۶</sup> (MMD) برای محاسبه اختلاف توزیع نمونه‌های دامنه منبع با نمونه‌های دامنه هدف بهره گرفته شده است. سپس، به نمونه‌هایی از دامنه منبع که دارای کمترین اختلاف توزیع با نمونه‌های دامنه هدف باشند، وزن بیشتری اختصاص داده می‌شود. مشکل عمده روش انتخاب لندمارک در این است که در انتخاب لندمارک‌ها، ممکن است بعضی خصوصیات فقط مختص دامنه منبع با هدف باشد که این امر، در انتخاب لندمارک‌ها در نظر گرفته نمی‌شود.

در روش‌های مبتنی بر مدل، هدف پیدا کردن یک طبقه‌بند انطباقی است که این کار توسط انتقال پارامترهای مدل آموزش داده شده بر روی دامنه منبع به دامنه هدف، بدون تغییر فضای خصیصه‌ای انجام می‌گیرد [۶، ۷، ۸]. تمرکز اصلی این روش‌ها، در یادگیری انتقالی نیمه نظارت‌شده است [۶، ۷]. یکی از روش‌های پیشنهاد شده برای یادگیری انتقالی بدون نظارت، روش تنظیم تطبیقی برای یادگیری انتقالی<sup>۱۷</sup> [۸] است که یک طبقه‌بند انطباقی توسط کاهش خطای طبقه‌بند در دامنه منبع، افزایش انطباق هندسی دامنه‌ها در فضای جدید و ایجاد تطبیق در توزیع مشترک بین دامنه‌ها ایجاد می‌کند.

روش‌های مبتنی بر خصوصیت یا تطبیق خصوصیات، فضای خصیصه‌ای را برای ایجاد یک نمایش تطبیق‌پذیر از دامنه‌های منبع و هدف تغییر می‌دهند [۱۳-۹]. سپس در فضای جدید، یک طبقه‌بند استاندارد روی داده‌های دامنه منبع، آموزش داده و به دامنه هدف اعمال می‌شود. روش‌های مبتنی بر خصوصیت، به سه دسته کلی تقسیم می‌شوند:

(۱) روش‌هایی که بر کاهش اختلاف توزیع حاشیه‌ای بین دامنه‌ها تمرکز دارند [۹، ۱۰، ۱۱]. TCA<sup>۱۸</sup> [۱۰] و GFK<sup>۱۹</sup> [۱۱]، نمونه‌های شناخته‌شده از این نوع روش‌ها می‌باشند. در روش TCA در فضای مشترک بین دامنه‌های منبع و هدف، هم‌زمان با کاهش اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف، برای حفظ خصوصیات اصلی داده‌های ورودی، واریانس داده‌های نگاشت‌شده حداکثر می‌شود. برای برقراری این دو شرط، از مولفه‌های انتقال<sup>۲۰</sup> مشترک بین دامنه‌های منبع و هدف که به‌طور هم‌زمان، باعث ایجاد اختلاف توزیع بین دامنه‌ها نشده و ساختار اصلی داده‌ها را حفظ نمایند، برای ایجاد نمایش جدید استفاده می‌شود.

GFK، یکی دیگر از روش‌های کاهش اختلاف حاشیه‌ای است که داده‌های منبع و هدف را به یک زیرفضای جدید که در آن توزیع دامنه‌های منبع و هدف به هم نزدیک هستند، نگاشت داده و مساله را حل می‌کند. در این روش، به دلیل کاهش بعد برای پیدا کردن

تطبیق‌پذیری دقیق‌تری باهم داشته باشند. مرحله دوم، شامل ایجاد یک طبقه‌بند انطباقی است که با ایجاد یک مدل انطباقی بر روی نمایش جدید داده‌های منبع و هدف، سعی دارد خطای پیش‌بینی را در دامنه هدف حداقل سازد. این طبقه‌بند انطباقی، با در نظر گرفتن حداکثر میزان انطباق مدل و حفظ حداکثری شکل هندسی داده‌ها بین دامنه‌های منبع و هدف، مدلی انطباق‌پذیر در راستای جبران اختلاف توزیع دامنه‌ها ایجاد می‌کند.

به‌طور کلی هر کدام از مراحل روش پیشنهادی شامل فرآیندهایی جهت ایجاد تطبیق بین دامنه‌های منبع و هدف می‌باشند. فرآیندهای مرحله اول را می‌توان به‌صورت زیر به اختصار بیان کرد: (۱) نگاشت داده‌های دامنه‌های منبع و هدف به بعد پایین<sup>۲۱</sup> مبتنی بر نظریه کاهش بعد اصلی<sup>۲۲</sup> [۳، ۲] کاهش اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف در بعد پایین، (۳) کاهش اختلاف توزیع شرطی بین دامنه‌های منبع و هدف در بعد پایین، (۴) خوشه‌بندی نمونه‌های دامنه منبع بر اساس برجسب یکسان جهت افزایش کارایی طبقه‌بندی. افزون بر این، فرآیندهای مرحله دوم نیز به شکل زیر اختصار می‌یابند: (۱) حداقل کردن خطای تجربی تابع پیش‌بینی در داده‌های برجسب‌دار دامنه منبع، (۲) حداکثر کردن نرخ سازگاری تابع پیش‌بینی و ساختار هندسی داده‌ها.

کارایی روش پیشنهاد شده 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 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 به تفصیل شرح داده شده است. در بخش چهارم، پایگاه داده‌های مورد ارزیابی این مقاله با جزئیات معرفی شده‌اند. در بخش پنجم، نتایج ارزیابی عملکرد الگوریتم پیشنهادی با روش‌های دیگر یادگیری ماشین و یادگیری انتقالی گزارش شده است. در انتها، مقاله با نتیجه‌گیری و ارائه پیشنهادهایی برای ادامه کار در آینده به اتمام رسیده است.

## ۲- کارهای پیشین

در حوزه یادگیری انتقالی، روش‌های بسیاری پیشنهاد شده است که تمرکز اصلی آن‌ها، بر کاهش اختلاف توزیع بین دامنه‌های منبع و هدف است. به‌طور کلی، روش‌های پیشنهاد شده در حوزه یادگیری انتقالی به سه دسته کلی تقسیم می‌شوند: روش‌های مبتنی بر نمونه، روش‌های مبتنی بر مدل، روش‌های مبتنی بر خصوصیت.

در روش‌های مبتنی بر نمونه، هدف وزن‌دهی مجدد یا انتخاب نمونه‌هایی از دامنه منبع است که اختلاف توزیع بین دامنه منبع و

زیرفضای جدید، داده اصلی به درستی در زیرفضای ایجاد شده نمایش داده نمی‌شود.

روش TJM<sup>۱۲</sup>، 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 کم بعد از دامنه‌های منبع و هدف ایجاد می‌شود که در این نمایش جدید، اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف کاهش می‌یابد. سپس، به نمونه‌هایی از دامنه منبع که دارای کم‌ترین مشابهت از نظر توزیع با نمونه‌های دامنه هدف هستند، وزن کم‌تری در ایجاد مدل طبقه‌بند اختصاص داده می‌شود.

### ۳ روش پیشنهادی

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

#### ۳-۱ هدف تحقیق

در شکل ۱، مراحل روش پیشنهادی را به‌صورت شماتیک نمایش داده شده است. در روش FMM، ابتدا یک الگوریتم تکرارشونده، بر روی هر دو نمونه‌های دامنه منبع و دامنه هدف اعمال می‌شود تا نمایش تطبیق‌پذیری از دامنه‌های منبع و هدف ایجاد شود. این الگوریتم در هر تکرار، اختلاف نمونه‌های دامنه منبع در هر کلاس را نسبت به میانگین کلاس، به حداقل می‌رساند. تکرار الگوریتم باعث می‌شود یک ناحیه تصمیم دقیق‌تری برای هر کلاس ایجاد شود. همچنین، در هر تکرار، برای کاهش اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف، اختلاف میانگین نمونه‌های دامنه هدف با میانگین نمونه‌های دامنه منبع، به حداقل رسیده و همچنین، به‌صورت هم‌زمان، برای کاهش اختلاف توزیع شرطی بین دامنه‌های منبع و هدف، اختلاف میانگین نمونه‌های دامنه هدف با میانگین نمونه‌های دامنه منبع در هر کلاس کاهش می‌یابد. تکرار این عمل موجب افزایش صحت طبقه‌بند در پیش‌بینی برچسب‌های نمونه‌های دامنه هدف می‌شود.

در مرحله دوم از روش FMM، یک طبقه‌بند انطباقی بر روی هر دو داده‌های دامنه‌های منبع و هدف ایجاد می‌شود. طبقه‌بند انطباقی، برای ایجاد یک مدل تطبیق‌پذیر بین دامنه‌ای، ابعاد تفکیک‌کننده کلاس‌های دامنه منبع را با ساختار نمونه‌های دامنه هدف سازگار می‌سازد. این عمل، موجب افزایش صحت طبقه‌بند در پیش‌بینی برچسب نمونه‌های دامنه هدف می‌شود.

#### ۳-۲ تعریف مساله

در این بخش ابتدا دو مفهوم دامنه<sup>۲۶</sup> و وظیفه<sup>۲۷</sup> تعریف شده و در ادامه، بیان مساله مطرح می‌شود.

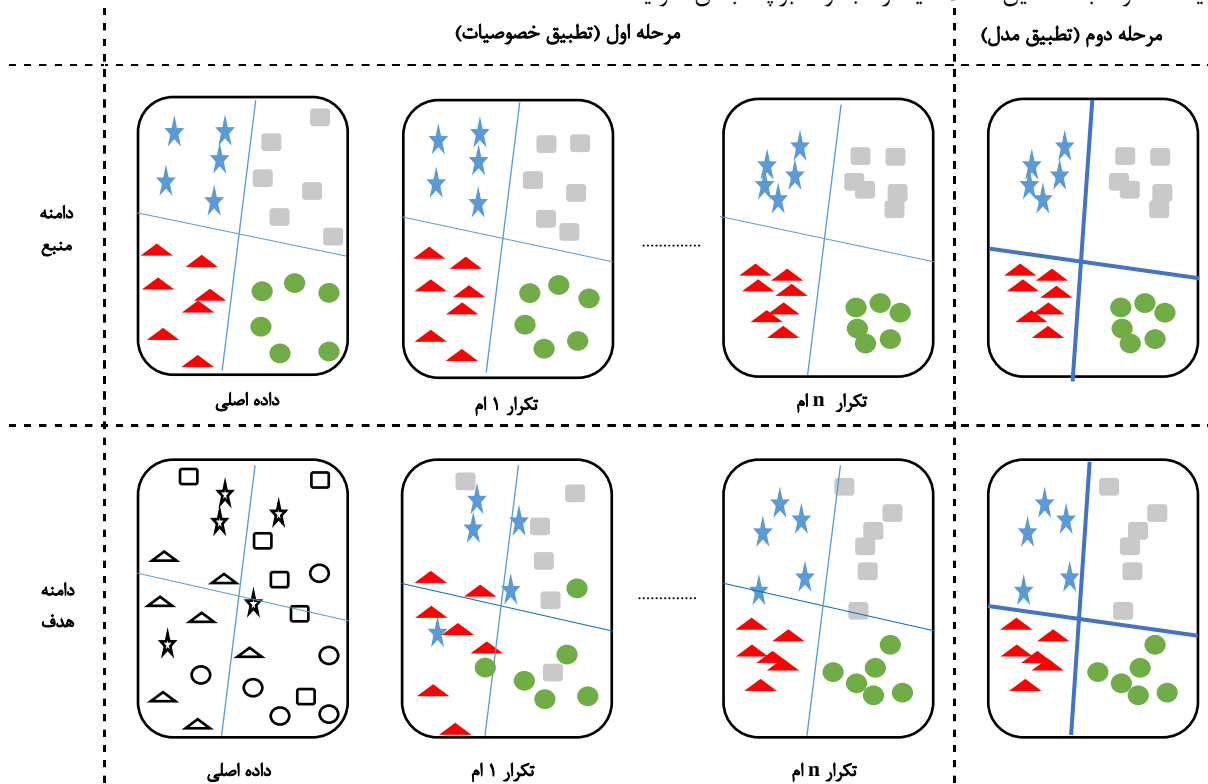
**دامنه.** هر دامنه  $D$  دارای دو عنصر اصلی است: فضای خصیصه‌ای  $X$  و توزیع احتمال حاشیه‌ای  $P(x)$ ، یعنی  $D = \{X, P(x)\}$ . بدین ترتیب دو دامنه زمانی متفاوت هستند که یا فضای خصیصه‌ای آن‌ها و یا توزیع احتمال حاشیه‌ای آن‌ها، باهم اختلاف داشته باشد. برای مثال، اگر  $X_s$  فضای خصیصه‌ای دامنه منبع  $(X_s \in X_s)$ ،  $X_t$  فضای خصیصه‌ای دامنه هدف  $(X_t \in X_t)$  و  $P_s(X_s)$  و  $P_t(X_t)$  به ترتیب، توزیع احتمال حاشیه‌ای دامنه‌های منبع و هدف باشند، دو دامنه وقتی متفاوت هستند که  $P_s(X_s) \neq P_t(X_t)$  یا  $X_s \neq X_t$ . **وظیفه.** برای هر دامنه، یک وظیفه  $T$  وجود دارد که شامل مجموعه برچسب‌های  $Y$  و تابع پیش‌بینی  $f(X)$  است، یعنی

روش‌هایی که بر کاهش اختلاف توزیع شرطی بین دامنه‌های منبع و هدف تمرکز دارند [۱۳]. یکی از این روش‌ها در حوزه یادگیری انتقالی بدون نظارت، روش تطبیق احتمال شرطی بین دامنه‌ها از طریق دسته‌بندی خصوصیات<sup>۲۲</sup> [۱۳] است. در این روش، یک دسته از خصوصیات که به‌طور هم‌زمان، اختلاف توزیع بین دامنه‌های منبع و هدف را کاهش داده و همچنین، مشابهت بین داده‌های برچسب‌دار را افزایش می‌دهند، انتخاب می‌شوند. چالش اصلی در این روش، محاسبه اختلاف توزیع بین دامنه‌های منبع و هدف در شرایطی است که خصوصیات فقط در یکی از دامنه‌های منبع یا هدف تعریف شده باشند. روش‌هایی که به‌طور هم‌زمان بر کاهش اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌ها تمرکز دارند [۱۴، ۱۵، ۱۶]. روش‌های شناخته‌شده در حوزه یادگیری انتقالی بدون نظارت، JDA<sup>۲۳</sup> [۱۵] و VDA<sup>۲۴</sup> [۱۶] می‌باشند. در روش JDA، یک نمایش کم بعد از دامنه‌های منبع و هدف ایجاد می‌شود که دارای حداقل اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌ها است. در VDA، یک نمایش کم بعد از دامنه‌های منبع و هدف ایجاد می‌شود که در آن، علاوه بر کاهش اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌ها، از روش خوشه‌بندی مستقل از دامنه<sup>۲۵</sup> نیز برای ایجاد تفکیک‌پذیری بین کلاس‌های مختلف استفاده شده است. خوشه‌بندی مستقل از دامنه، با حفظ شکل هندسی و آماری بین دامنه‌های منبع و هدف در فضای جدید، باعث افزایش صحت طبقه‌بند در دامنه هدف می‌شود.

در این مقاله، برای حل مساله یادگیری انتقالی بدون نظارت، یک چهارچوب دومرحله‌ای با ترکیبی از دو روش مبتنی بر خصوصیت و مبتنی بر مدل پیشنهاد داده می‌شود. در مرحله اول، داده‌های دامنه منبع و دامنه هدف توسط یک روش تطبیق خصوصیات، به فضای جدید که در آن اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌های منبع و هدف، به‌طور هم‌زمان حداقل شده، نگاشت می‌شوند. همچنین، در این مرحله از خوشه‌بندی مستقل از دامنه برای ایجاد تفکیک‌پذیری بین کلاس‌های مختلف در فضای جدید بهره گرفته می‌شود. سپس در مرحله دوم، یک تابع پیش‌بینی برای طبقه‌بندی در فضای جدید، آموزش داده می‌شود که دو هدف دارد: (۱) طبقه‌بند دارای حداقل خطای پیش‌بینی برچسب در دامنه منبع باشد، (۲) تابع پیش‌بینی با

در توزیع احتمال شرطی آن‌ها است، بدین معنی که  $Y_s \neq Y_t$  یا  $P_s(Y_s | x_s) \neq P_t(Y_t | x_t)$ .

$T = \{Y, f(x)\}$  تابع پیش‌بینی  $f(x)$ ، مجموعه برچسب‌های  $Y$  را به ازای مجموعه نمونه ورودی  $X$  پیش‌بینی می‌کند و می‌توان آن را به صورت توزیع احتمال شرطی  $P(Y|x)$  بیان کرد؛ در این صورت، اگر دو وظیفه متفاوت باشند، این اختلاف یا در مجموعه برچسب آن‌ها و یا



شکل ۱. مراحل روش FMM (نمایش بهتر به صورت رنگی). در روش پیشنهادی، برای پیش‌بینی برچسب‌های نمونه‌های دامنه هدف، یک الگوریتم تکرارشونده بر روی داده‌های دامنه‌های منبع و هدف اعمال می‌شود. در هر تکرار، اختلاف توزیع حاشیه‌ای و شرطی دامنه هدف با دامنه منبع، کاهش یافته و خوشه‌بندی مستقل از دامنه در دامنه منبع ایجاد می‌شود. در ادامه، با اعمال یک طبقه‌بند انطباقی، صحت مدل پیش‌بینی در دامنه هدف افزایش می‌یابد.

بین دامنه‌های منبع و هدف حداقل شده و در کنار آن، تلاش می‌شود که ساختار اصلی داده‌های ورودی نیز حفظ شود. در سال‌های اخیر، روش‌های زیادی برای یادگیری یک نمایش خصیصه تطبیق‌پذیر در یادگیری انتقالی، جهت کاهش اختلاف توزیع بین دامنه‌های منبع و هدف پیشنهاد شده است که از جمله این روش‌ها، روش‌های مبتنی بر کاهش بعد می‌باشند [۳، ۹]. در روش‌های کاهش بعد، داده‌ها از فضای اصلی به یک فضای نگاشت‌شده منتقل می‌شوند با این شرط که هزینه بازنگاشت آن‌ها (بازگرداندن آن‌ها به فضای اصلی<sup>۲۸</sup>) حداقل باشد. از جمله روش‌های شناخته‌شده در حوزه کاهش بعد، می‌توان به روش تحلیل اجزای اصلی (PCA) [۳] اشاره نمود. در ادامه روش PCA با جزئیات بیشتر معرفی می‌شود.

### ۳-۴ کاهش بعد

PCA، یک روش انتقال داده‌ها به یک فضای کم بعد است که ایده اصلی این روش، حفظ ساختار اصلی داده‌ها در فضای جدید می‌باشد. بدین ترتیب، داده‌ها بر روی اجزای اصلی خود نگاشت می‌شوند که

**بیان مساله.** فرض کنید  $n_s$  و  $n_t$  به ترتیب، تعداد نمونه‌های دامنه منبع و هدف باشند. در یادگیری انتقالی بدون نظارت، داده‌های دامنه منبع، داده‌های برچسب‌داری هستند و به صورت  $D_s = \{(x_1, y_1), \dots, (x_{n_s}, y_{n_s})\}$  تعریف می‌شوند. به همین ترتیب، داده‌های دامنه هدف، داده‌های بدون برچسب هستند که به صورت  $D_t = \{(x_{n_s+1}, \dots, x_{n_s+n_t})\}$  تعریف می‌شوند. هدف این مقاله، ارائه یک چهارچوب دومرحله‌ای جهت کاهش اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌های منبع و هدف است. در ادامه، مدل پیشنهادی با جزئیات بیشتر بیان شده است.

### ۳-۴ ایجاد یک نمایش تطبیق‌پذیر

یک مساله 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

$$Mrg(X_s, X_t) = \left\| \frac{1}{n_s} \sum_{i=1}^{n_s} A^T x_i - \frac{1}{n_t} \sum_{j=n_s+1}^{n_s+n_t} A^T x_j \right\|^2 \quad (2)$$

$$= \text{tr}(A^T X M_0 X^T A)$$

ماتریس  $M_0$ ، ماتریس MMD است. اگر  $x_i, x_j \in X_s$  باشد،  $(M_0)_{ij} = \frac{1}{n_s n_t}$ ، اگر  $x_i, x_j \in X_t$  باشد،  $(M_0)_{ij} = \frac{1}{n_s n_t}$  و در غیر این صورت،  $(M_0)_{ij} = \frac{1}{n_s n_t}$  است.

### ۳-۴-۳ کاهش بعد در راستای کاهش اختلاف توزیع شرطی بین دامنه‌ها

در شرایطی که اختلاف توزیع بین دامنه‌های منبع و هدف زیاد باشد، برای بهبود عملکرد مدل یادگیری، تنها کاهش اختلاف توزیع حاشیه‌ای کافی نبوده و برای تطبیق‌پذیری بیشتر مدل، اختلاف توزیع شرطی بین دامنه‌ها نیز باید حداقل شود [۱۶]. برای محاسبه اختلاف توزیع شرطی با بهره‌گیری از روش MMD، مجموع اختلاف میانگین نمونه‌های هر کلاس در دامنه‌های منبع و هدف، به‌صورت زیر محاسبه خواهد شد.

$$Cnd(X_s, X_t) = \left\| \frac{1}{n_s^c} \sum_{x_i \in X_s^c} A^T x_i - \frac{1}{n_t^c} \sum_{x_j \in X_t^c} A^T x_j \right\|^2 \quad (3)$$

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

که  $X_s^c$  و  $X_t^c$ ، به ترتیب، داده‌های دامنه منبع و دامنه هدف در کلاس  $c$  و  $n_s^c$  و  $n_t^c$ ، به ترتیب، تعداد نمونه‌های دامنه منبع و دامنه هدف در کلاس  $c$  می‌باشند. اگر  $x_i, x_j \in X_s^c$  باشد،  $(M_c)_{ij} = \frac{1}{n_s^c n_s^c}$  و اگر  $x_i, x_j \in X_t^c$  باشد، آنگاه  $(M_c)_{ij} = \frac{1}{n_t^c n_t^c}$  و در غیر این صورت،  $(M_c)_{ij} = \frac{1}{n_s^c n_t^c}$  است.

مشکل اصلی در محاسبه اختلاف توزیع شرطی بین دامنه‌های منبع و هدف در یادگیری انتقالی بدون نظارت، عدم وجود داده برچسب‌دار در دامنه هدف می‌باشد. راه‌حل این مشکل، استفاده از یک طبقه‌بند استاندارد برای پیش‌بینی برچسب برای داده‌های دامنه هدف است. این طبقه‌بند استاندارد بر روی دامنه منبع ایجاد و برای پیش‌بینی برچسب‌های داده‌های بدون برچسب دامنه هدف استفاده می‌شود. به علت وجود اختلاف توزیع بین دامنه‌های هدف و منبع، این برچسب‌ها در ابتدای کار، به احتمال زیاد نادرست هستند. به همین دلیل، برای افزایش صحت طبقه‌بند استاندارد در پیش‌بینی برچسب برای داده‌های دامنه هدف، الگوریتم پیشنهادی در مقاله به‌صورت تکرار شونده خواهد بود تا تطبیق‌پذیری بین توزیع‌های شرطی دامنه‌ها در فضای جدید با صحت بیشتری انجام پذیرد.

شامل اجزایی است که دارای پخشش (واریانس) بالایی هستند. از بین اجزای به‌دست‌آمده، اجزایی که دارای حداکثر واریانس باشند، به عنوان اجزای اصلی جهت نگاشت داده‌ها به فضای کم بعد استفاده می‌شوند. 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  $X_t$ ، به ترتیب داده‌های دامنه منبع و هدف باشند، تابع مرکزیت  $H$ ، به‌صورت  $H = I_{n_s + n_t} - \frac{1}{n_s + n_t} \bar{I} \bar{I}^T$  تعریف می‌شود که  $I$  ماتریس همانی  $(I_{n_s + n_t} \in R^{(n_s + n_t) \times (n_s + n_t)})$  و  $\bar{I}$  بردار ستونی از یک‌ها می‌باشد. کوواریانس داده‌ها در فضای اصلی، با بهره‌گیری از تابع مرکزیت، برابر با  $XXH^T$  است که بیانگر اختلاف داده‌ها از میانگین کل نمونه‌ها بوده و استفاده از تابع مرکزیت مانع از پراکندگی داده می‌شود.

هدف PCA، پیدا کردن یک ماتریس نگاشت وارون‌پذیر  $A \in R^{m \times k}$  است که بتواند نمونه‌ها را، از یک فضای اصلی  $m$  بعدی به فضای نگاشت‌شده  $k$  بعدی ( $k < m$ ) با حداکثر واریانس انتقال دهد. تابع بهینه‌سازی PCA به‌صورت رابطه (۱) بیان می‌شود.

$$\max_{A^T A = I} \text{tr}(A^T X H X^T A) \quad (1)$$

که  $A^T X H X^T A$  کوواریانس نمونه‌ها در فضای جدید بوده و عبارت  $AA^T = I$ ، برای برقراری شرط وارون‌پذیری تابع نگاشت به رابطه (۱) اضافه می‌شود.  $\text{tr}$  نشان‌دهنده حاصل جمع عناصر قطر اصلی ماتریس است.

### ۳-۴-۳ کاهش بعد در راستای کاهش اختلاف توزیع حاشیه‌ای بین دامنه‌ها

باوجود این‌که PCA، داده‌ها را به یک فضای کم بعد منتقل می‌کند، ولی تغییر چندانی در اختلاف توزیع بین دامنه‌های منبع و هدف در فضای جدید ایجاد نمی‌شود؛ بنابراین یک مساله مهم، کاهش اختلاف توزیع بین دامنه‌های منبع و هدف در فضای کم بعد است که این اختلاف توزیع شامل اختلاف توزیع حاشیه‌ای و شرطی است. برای محاسبه اختلاف توزیع‌های شرطی و حاشیه‌ای بین دامنه‌های منبع و هدف در فضای جدید، از روش MMD استفاده می‌شود. این روش، یک روش غیر پارامتری برای محاسبه اختلاف توزیع بین دامنه‌ها در مسائل با ابعاد بالا (مانند مسائل بینایی ماشین) است.

در روش MMD برای محاسبه اختلاف توزیع دامنه‌ها، داده‌های دامنه منبع و هدف به فضای هیلبرت<sup>۲۹</sup> نگاشت می‌شوند. برای محاسبه اختلاف توزیع حاشیه‌ای بین دامنه‌ها در فضای اصلی، با استفاده از روش MMD، از اختلاف بین میانگین نمونه‌های دامنه منبع و هدف در فضای هیلبرت استفاده می‌شود. بدین ترتیب، اگر  $A$  تابع نگاشت به فضای جدید در نظر گرفته شود، اختلاف توزیع حاشیه‌ای دامنه‌های منبع و هدف به‌صورت رابطه (۲) تعریف می‌شود [۹، ۱۳، ۱۵].

## ۴ ۳ ۳ خوشه‌بندی مستقل از دامنه

در بخش‌های قبل، چگونگی کاهش اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌های منبع و هدف در فضای کم بعد، توضیح داده شد. با وجود کاهش اختلاف توزیع بین دامنه‌ها، به علت اینکه نمونه‌های دامنه منبع در کلاس‌های مختلف از نظر خصوصیات اصلی (برای نمونه خصوصیات هندسی و آماری) باهم اختلاف دارند، ناحیه تصمیم‌گیری برای هر کلاس در دامنه منبع به خوبی تفکیک نشده و مدل یادگیری ایجاد شده در دامنه منبع، صحت بالایی در پیش‌بینی داده‌های بدون برچسب دامنه هدف نخواهد داشت [۱۷].

برای حل این مشکل، نمونه‌های با برچسب یکسان در دامنه منبع، در یک خوشه قرار داده می‌شوند. برای افزایش صحت مدل طبقه‌بندی، در هر کدام از این خوشه‌ها، فاصله نمونه‌ها از میانگین نمونه‌ها محاسبه و به حداقل خواهد رسید. این عمل از یک سو، باعث حداقل شدن فاصله نمونه‌های هر کلاس از مراکز درون کلاسی (حداکثر شدن واریانس هر کلاس) و از سوی دیگر، باعث حداکثر شدن فاصله مراکز بین کلاسی از هم و افزایش تفکیک‌پذیری بین کلاس‌ها می‌شود.

اگر فاصله نمونه‌های دامنه منبع از میانگین نمونه‌ها در هر کلاس در فضای اصلی به صورت  $S = \sum_{\forall c \in C} \sum_{x_i \in X_c^s} (x_i - \mu^c)^T (x_i - \mu^c)$  در نظر گرفته شود که  $\mu^c$  نشان‌دهنده میانگین نمونه‌های کلاس  $c$  است، این فاصله در فضای جدید به صورت رابطه (۴) تعریف خواهد شد.

$$DIC(X_s, \mu^c) = tr(A^T SA) \quad (4)$$

هدف خوشه‌بندی مستقل از دامنه، پیدا کردن تابع نگاشت  $A$  به صورتی است که در فضای جدید رابطه (۴) حداقل شود.

## ۴ ۳ ۵ مساله بهینه‌سازی

مرحله اول از روش تطبیق خصوصیات و مدل، ایجاد یک نمایش تطبیق‌پذیر از دامنه‌های منبع و هدف است. در نمایش جدید، به صورت هم‌زمان، اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌ها حداقل شده و از روش خوشه‌بندی مستقل از دامنه، برای بهبود صحت مدل طبقه‌بندی در پیش‌بینی برچسب‌ها استفاده شده است. بدین ترتیب، با ترکیب رابطه‌های (۱) تا (۴)، تابع هدف برای مرحله اول، به صورت رابطه (۵) تعریف می‌شود.

$$\min_{A^T X H X^T A = I} \sum_{c=0}^C tr(A^T (X M_c X^T + S) A) + \lambda \|A\|_F^2 \quad (5)$$

که  $\lambda$ ، ضریب تناسب و عبارت دوم، نرم فراینیوس<sup>۳۰</sup> تابع نگاشت  $A$  می‌باشد. برای حل راحت‌تر تابع هدف، به جای محاسبه اختلاف توزیع بین دامنه‌ها در فضای کم بعد، از توزیع دامنه‌ها در فضای کرنلی استفاده می‌شود. تابع نگاشت  $\phi$  به صورت  $\phi: X \rightarrow \phi(X)$  در نظر گرفته می‌شود که داده‌ها را به فضای هیلبرت نگاشت می‌دهد. می‌توان  $V$  را به صورت  $A = V^T \phi(X)$  تعریف کرد که  $V \in R^{(n_s + n_t) \times k}$  است. با تعریف تابع کرنل به صورت  $k(g(x_i), g(x_j)) = \phi(g(x_i))^T \phi(g(x_j))$

و استفاده از تئوری رپرزنتر<sup>۳۱</sup>، تابع هدف مرحله اول به شکل رابطه (۶) بازنویسی می‌شود.

$$\min_{V^T K H K^T V = I} \sum_{c=0}^C tr(V^T (K M_c K^T + S) V) + \lambda \|V\|_F^2 \quad (6)$$

هدف اصلی در این مرحله، پیدا کردن تابع نگاشت  $V$  به صورتی است که اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌های منبع و هدف با حفظ ساختار داده اصلی (تعریف شده در رابطه (۶)) حداقل شود.

## ۴ ۳ ۴ ایجاد یک طبقه‌بند انطباقی

در بسیاری از موارد، تنها حداقل کردن اختلاف توزیع بین دامنه‌های منبع و هدف، برای ایجاد یک طبقه‌بند دارای صحت بالا در پیش‌بینی داده‌های بدون برچسب کافی نیست. در این شرایط، حتی با ایجاد تفکیک بین کلاس‌های دامنه منبع نیز، نمی‌توان یک طبقه‌بند تطبیق‌پذیر ایجاد کرد، زیرا ابعاد تفکیک‌کننده دامنه منبع و دامنه هدف نسبت به هم اختلاف قابل‌توجهی دارند. به بیان دیگر، ابعاد تفکیک‌کننده ایجاد شده توسط طبقه‌بند در دامنه منبع، به دلیل عدم سازگاری ساختار داده‌های دامنه منبع و دامنه هدف، دارای خطای زیادی در ایجاد تفکیک در دامنه هدف می‌باشند. بدین ترتیب، در مرحله دوم، هدف ایجاد یک طبقه‌بند انطباقی، برای افزایش تطبیق‌پذیری دامنه‌های منبع و هدف است که دارای حداقل خطای پیش‌بینی برچسب دامنه هدف، ساختار هندسی دامنه هدف را به خوبی نشان می‌دهند، بنابراین طبقه‌بند انطباقی، با بهره‌گیری از داده‌های برچسب‌دار دامنه منبع و داده‌های بدون برچسب دامنه هدف به صورت هم‌زمان، باعث تطبیق بهتر طبقه‌بند با اختلاف بین دامنه‌ها می‌شود. در ادامه، طبقه‌بند انطباقی با جزئیات بیشتری معرفی می‌شود.

۴ ۳ ۴ ۱ یادگیری براساس حداقل کردن ریسک تجربی<sup>۳۲</sup>

برای ایجاد طبقه‌بند انطباقی، در ابتدا، یک طبقه‌بند استاندارد  $f$  روی دامنه منبع برای پیش‌بینی داده‌های بدون برچسب دامنه هدف ایجاد می‌شود. این طبقه‌بند، باید دارای حداقل خطای پیش‌بینی در دامنه منبع باشد. اگر  $g$  تابع نگاشت به نمایش جدید باشد،  $l$ ، تابع خطای<sup>۳۳</sup> پیش‌بینی طبقه‌بند انطباقی در دامنه منبع بوده و به شکل رابطه (۷) تعریف می‌شود.

$$l(f(g(x_i)), y_i) = \sum_{i=1}^{n_s + n_t} R_{ii} (y_i - f(g(x_i)))^2 \quad (7)$$

که  $R$ ، یک ماتریس قطری است و اگر  $x_i$ ، یک نمونه از دامنه منبع باشد،  $R_{ii} = 1$  می‌باشد و در غیر این صورت،  $R_{ii} = 0$  است. رابطه (۷)، نشان‌دهنده مجموع مربع خطای برچسب اصلی و برچسب پیش‌بینی شده توسط تابع پیش‌بینی در داده‌های دامنه منبع می‌باشد.

## ۳-۴-۴- یادگیری براساس حفظ ساختار داده

برای ایجاد سازگاری بین طبقه‌بند انطباقی و ساختار هندسی داده‌ها در فضای جدید، از تعریف فرض منیفلد<sup>۳۴</sup> استفاده می‌شود. براساس فرض منیفلد، اگر دو نمونه از نظر هندسی در توزیع حاشیه‌ای دامنه‌های منبع و هدف، دارای اختلاف کمی با یکدیگر باشند، توزیع شرطی دو نمونه (برچسب‌های دو نمونه) نیز دارای اختلاف کمی با یکدیگر می‌باشد [۱۸]؛ بنابراین می‌توان از ساختار داده‌های دامنه‌های منبع و هدف در توزیع حاشیه‌ای، برای ایجاد یک طبقه‌بند تطبیق‌پذیر بین دامنه‌ها استفاده کرد.

برای محاسبه فاصله بین نمونه‌ها، در ابتدا داده‌ها توسط گراف نزدیک‌ترین همسایه<sup>۳۵</sup> که دارای  $n_s + n_t$  رأس است، مدل می‌شوند. هر رأس، مربوط به یک داده از مجموعه داده‌های دامنه‌های منبع و هدف است. برای هر داده،  $p$  نزدیک‌ترین همسایه پیدا شده و به یکدیگر متصل می‌شوند. ماتریس  $W$ ، برای محاسبه وزن ارتباطی بین هر دو نمونه  $x_i$  و  $x_j$  متصل به هم در گراف، به شکل رابطه (۸) تعریف می‌شود.

$$W_{ij} = e^{-\frac{|(x_i - x_j)|^2}{\delta}} \quad (8)$$

تابع  $M_f$  برای ایجاد حداکثر سازگاری بین طبقه‌بند  $f$  و توزیع هندسی دامنه‌های منبع و هدف  $(P_s, P_t)$  به شکل رابطه (۹) تعریف می‌شود.

$$M_f(P_s, P_t) = \sum_{i,j=1}^{n_s+n_t} (f(x_i) - f(x_j))^2 W_{ij} \\ = \sum_{i,j=1}^{n_s+n_t} f(x_i) \bar{L}_{ij} f(x_j) f(x_j) \quad (9)$$

که  $\bar{L}$ ، ماتریس لاپلاسین نرمال شده است. اگر  $D$  یک ماتریس قطری بوده و به شکل رابطه (۱۰) تعریف شود.

$$D_{ii} = \sum_{j=1}^{n_s+n_t} W_{ij} \quad (10)$$

هر عنصر  $D_{ii}$  نشانگر وزن ارتباطی رأس  $i$  با تمام رؤس است. بدین ترتیب  $L = D - W$ ، ماتریس لاپلاسین<sup>۳۶</sup> نرمال نشده است که هر عنصر قطری  $L_{ii}$  در ماتریس، وزن ارتباطی رأس  $i$  را با تمام رؤس به جز خودش نشان می‌دهد. شکل نرمال شده ماتریس لاپلاسین به شکل رابطه (۱۱) تعریف می‌شود [۱۹].

$$\bar{L} = I - D^{-\frac{1}{2}} W D^{-\frac{1}{2}} \quad (11)$$

در نهایت، مساله بهینه‌سازی برای مرحله دوم روش پیشنهادی به شکل رابطه (۱۲) ایجاد می‌شود.

$$\min_{f \in U} \sum_{i=1}^{n_s} (f(g(x_i)), y_i) + \sigma f^2 + \gamma M_f(P_s, P_t) \quad (12)$$

که  $U$ ، یک مجموعه از طبقه‌بندها و  $f^2$  نرم مربع  $f$  در  $U$ ،  $\sigma$  و  $\gamma$  پارامترهای نسبت<sup>۳۷</sup>  $M_f$ ، محاسبه‌کننده عدم سازگاری بین تابع پیش‌بینی و توزیع بین دامنه‌های منبع و هدف است.

برای حل راحت‌تر رابطه (۱۲) می‌توان از تعریف کرنل استفاده کرد. تابع پیش‌بینی  $f$  به صورت  $f(g(x_i)) = w^T \phi(g(x_i))$  تعریف می‌شود که  $w$  پارامترهای طبقه‌بندی و  $\phi$ ، تابع نگاشتی است که داده‌ها را از فضای جدید به فضای هیلبرت نگاشت می‌دهد. با تعریف تابع کرنل  $k$  به صورت  $k(g(x_i), g(x_j)) = \phi(g(x_i))^T \phi(g(x_j))$  و استفاده از تئوری رگرزنتر، تابع پیش‌بینی به شکل رابطه (۱۳) تعریف می‌شود [۲۰].

$$f(g(x)) = \sum_{i=1}^{n_s+n_t} \alpha_i k(g(x_i), g(x)) \quad (13)$$

در پایان تابع هدف طبقه‌بندی انطباقی (رابطه (۱۲)) به شکل رابطه (۱۴) بازنویسی می‌شود.

$$\alpha = \argmin_{\alpha \in R^{n_s+n_t}} (Y - \alpha^T) \\ + \argmin_{\alpha \in R^{n_s+n_t}} \text{tr}(\gamma \alpha^T \bar{K} L K \alpha + \sigma \alpha^T K \alpha) \quad (14)$$

بعد از مشتق‌گیری از معادله بالا، پارامترهای طبقه‌بندی  $\alpha$  به شکل رابطه (۱۵) تعریف می‌شوند.

$$\alpha = (\sigma I + (R + \gamma \bar{L}) K)^{-1} R Y^T \quad (15)$$

پس از به دست آوردن پارامترهای طبقه‌بندی، تابع پیش‌بینی براساس معادله (۱۲) ساخته شده و برای پیش‌بینی داده‌های بدون برچسب دامنه هدف استفاده می‌شود.

## ۳-۵- تحلیل زمان اجرا

در این بخش پیچیدگی زمانی روش FMM که دارای یک الگوریتم دومرحله‌ای است، مورد بررسی قرار می‌گیرد. اگر  $m$ ، تعداد ابعاد اصلی داده‌ها در دامنه‌های منبع و هدف و  $C$  تعداد کل کلاس‌ها باشد، پیچیدگی زمانی در مرحله اول (رابطه (۶)) به شرح زیر است: محاسبه ماتریس  $M_0$  در زمان  $O((n_s + n_t)^2)$ ، محاسبه ماتریس  $H$ ،  $K$  و  $I$  با پیچیدگی  $O((n_s + n_t)^2)$ ، محاسبه ماتریس  $S$  با پیچیدگی  $O(n_s * m)$ ، پیش‌بینی برچسب‌های دامنه هدف در زمان  $O((n_s + n_t) * m)$  و به‌روزرسانی ماتریس  $M_c$  در هر تکرار، با پیچیدگی  $O(C(n_s + n_t)^2)$  انجام می‌گیرد. همچنین، حل رابطه بهینه مرحله اول (رابطه (۶)) توسط روش تجزیه مقادیر ویژه<sup>۳۸</sup>، پیچیدگی زمانی  $O(m^2)$  دارد. به‌طور کلی، پیچیدگی زمانی مرحله اول به صورت  $O((n_s + n_t)^2 + (n_s + n_t) * m + m^2)$  محاسبه می‌شود.

در مرحله دوم (رابطه (۱۲))، محاسبه ماتریس‌های  $K$  و  $L$  دارای پیچیدگی زمانی  $O((n_s + n_t)^2)$  بوده و محاسبه پارامترهای طبقه‌بند انطباقی (رابطه (۱۲)) دارای پیچیدگی زمانی  $O((n_s + n_t)^2)$  است؛ بنابراین، پیچیدگی زمانی مرحله دوم  $O((n_s + n_t)^2)$  می‌باشد.



۲۰۰۰ نمونه از داده‌های دامنه MNIST را به عنوان داده‌های آموزشی و ۱۸۰۰ نمونه از داده‌های دامنه USPS را به عنوان داده‌های تست شامل می‌شود. به‌طور مشابه، دامنه USPS\_vs\_MNIST (U\_M)، با جابجایی نمونه‌های آموزشی و تست دامنه MNIST\_vs\_USPS، ایجاد شده است. در هر دو پایگاه داده ایجاد شده، نمونه‌های آموزشی و تست به اندازه‌های ۱۶\*۱۶ پیکسلی تغییر کرده است و به‌صورت بردارهای خصیصه کدشده توسط پیکسل‌های سیاه و سفید (تصاویر سیاه سفید)، نمایش داده می‌شوند.

پایگاه داده کوپل [۲۵]، شامل ۱۴۴۰ تصویر سیاه و سفید از ۲۰ شی با زمینه سیاه در زاویه‌های مختلف است که هر تصویر در اندازه ۳۲\*۳۲ پیکسلی نمایش داده می‌شود. تصاویر با اختلاف ۵ درجه‌ای نسبت به همدیگر تصویربرداری شده‌اند و بدین ترتیب، ۷۲ تصویر در ۳۶۰ درجه جمع‌آوری شده است. پایگاه داده کوپل، شامل دو دامنه کوپل ۱ و کوپل ۲ است که کوپل ۱، مجموعه تصاویر اشیاء گرفته‌شده در زاویه‌های [۰، ۸۵] و [۱۸۰، ۲۶۵] (ربع اول و سوم) و کوپل ۲، مجموعه تصاویر اشیاء گرفته‌شده در زاویه‌های [۹۰، ۱۷۵] و [۲۷۰، ۳۵۵] (ربع دوم و چهارم) هستند. بدین ترتیب، وجود اختلاف توزیع بین دامنه‌های کوپل ۱ و کوپل ۲، مشهود است. دامنه COIL1\_vs\_COIL2 (C1\_C2) که ۷۲۰ نمونه از دامنه کوپل ۱ را به عنوان داده آموزشی و ۷۲۰ نمونه از دامنه کوپل ۲ را به عنوان داده تست شامل شده است، به عنوان آزمایش اول از پایگاه داده کوپل در نظر گرفته می‌شود. به‌طور مشابه، دامنه COIL2\_vs\_COIL1 (C2\_C1)، با جابجایی نمونه‌های آموزشی و تست دامنه COIL1\_vs\_COIL2 ایجاد شده و به عنوان آزمایش دوم در نظر گرفته می‌شود.

پایگاه داده پای [۲۶]، پایگاه داده شناخته‌شده‌ای در زمینه تشخیص چهره است. این پایگاه داده، شامل ۴۱۳۶۸ تصویر از ۶۸ شخص مختلف در اندازه‌های ۳۲\*۳۲ پیکسلی و در شرایط روشنایی و حالت تصویربرداری مختلف می‌باشد. ۵ دامنه در این پایگاه داده وجود دارد که هر کدام مربوط به یک حالت تصویربرداری است: پای ۱ (حالت چپ (P1)، پای ۲ (حالت بالا (P2)، پای ۳ (حالت پایین (P3)، پای ۴ (حالت روبرو (P4)، پای ۵ (حالت راست (P5). در این دامنه‌ها، تصاویر در شرایط مختلف گرفته شده و بنابراین اختلاف توزیع بین هر دو دامنه وجود دارد. در این پایگاه داده، آزمایش بر روی ۲۰ مجموعه بین دامنه‌ای مختلف طراحی شده است که از بین ۵ دامنه، دو دامنه مختلف به عنوان دامنه‌های منبع و هدف انتخاب می‌شوند. به‌طور کلی، کارایی الگوریتم پیشنهادی بر روی ۳۶ مجموعه تصاویر بین دامنه‌ای مختلف مورد ارزیابی قرار گرفته است که در بخش بعدی گزارش می‌شود.

#### ۴.۴ ارزیابی الگوریتم‌ها

روش‌هایی که الگوریتم FMM با آن‌ها مقایسه شده است، عبارتند از: طبقه‌بند نزدیک‌ترین همسایه (NN)، تحلیل اجزای اصلی (PCA) [۳]،

به‌طور کلی، با توجه به اینکه  $n_s + n_t \ll m$  است، پیچیدگی زمانی روش FMM به‌صورت رابطه (۱۶) تعریف می‌شود.

$$T = O(C(n_s + n_t)^2) \quad (16)$$

#### ۴.۴-تنظیمات اولیه محیط آزمایش

در این بخش، آزمایش‌های انجام گرفته برای ارزیابی روش پیشنهادی به تفصیل بیان می‌شود.

#### ۴.۴-معرفی مجموعه داده‌ها

کارایی روش پیشنهاد شده در این مقاله، بر روی چهار نوع پایگاه داده مختلف ارزیابی شده است: (۱) آفیس و کالتک، (۲) اعداد USPS و MNIST، (۳) کوپل، (۴) چهره (پای).

پایگاه داده آفیس [۲۱]، شامل مجموعه تصاویر از اشیاء مختلف است که از ۳ دامنه آمازون (A)، وبکم (W) و DSLR (D) جمع‌آوری شده است که تصاویر از نظر کیفیت، روشنایی، رنگ و نوع زمینه باهم متفاوت هستند. دامنه آمازون، متشکل از تصاویر اشیاء داندلود شده از سایت‌های تجاری (برای نمونه amazon.com) است که این تصاویر با زمینه سفید بوده و اشیاء در مرکز آن‌ها قرار گرفته و در شرایط نورپردازی استودیو تصویربرداری شده‌اند. دامنه وبکم، شامل تصاویر اشیاء با وضوح پایین، گرفته‌شده توسط دوربین وب و دامنه DSLR، شامل تصاویر اشیاء با وضوح بالا، گرفته‌شده توسط دوربین‌های دیجیتالی حرفه‌ای هستند. در دامنه‌های آفیس، ابتدا خصوصیات توسط روش استخراج خصیصه<sup>۳۹</sup> SURF [۲۱] به‌دست آمده و سپس این خصوصیات، در هیستوگرام ۸۰۰ خانه‌ای با کدبوک<sup>۴۰</sup>های محاسبه‌شده، توسط اعمال روش k-means بر روی نمونه‌های دامنه آمازون، مقداردهی شده و مورد استفاده قرار می‌گیرند. کالتک ۲۵۶ (C) [۲۲]، پایگاه داده استاندارد برای تشخیص اشیاء شامل ۳۰۶۰۷ تصویر و ۲۵۶ کلاس از تصاویر است. در دامنه‌های پایگاه داده آفیس و دامنه کالتک، ۱۲ آزمایش طراحی شده است که از ۱۰ کلاس مشترک بین پایگاه داده آفیس و پایگاه داده کالتک استفاده می‌کنند. در هر یک از آزمایش‌های طراحی شده، یکی از مجموعه داده‌ها (برای مثال وبکم)، به عنوان دامنه منبع و یکی دیگر از مجموعه داده‌ها (برای مثال آمازون)، به عنوان دامنه هدف انتخاب می‌شوند.

پایگاه داده اعداد شامل دو دامنه USPS [۲۳] و MNIST [۲۴] است. دامنه USPS، شامل اعداد دست‌نویس اسکن‌شده از نامه‌های اداره پست آمریکا در اندازه‌های ۱۶\*۱۶ پیکسلی است که ۷۲۶۱ داده آموزشی و ۲۰۰۷ داده تست دارد. دامنه MNIST، شامل اعداد دست‌نویس جمع‌آوری شده از دانش‌آموزان دبیرستانی آمریکا و کارمندان سازمان‌های مالیاتی آمریکا در اندازه ۲۸\*۲۸ پیکسلی است که ۶۰۰۰ داده آموزشی و ۱۰۰۰۰ داده تست دارد. برای اینکه دو دامنه در شرایط یکسان مورد آزمایش قرار بگیرند، دامنه MNIST\_vs\_USPS (M\_U) ایجاد شده است که به‌صورت تصادفی،

## ۵- نتایج و بحث‌ها

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

### ۵.۱ ارزیابی نتایج

جدول ۲ تا ۴، نشان‌دهنده نتایج به‌دست‌آمده از روش FMM و الگوریتم‌های مورد مقایسه به ترتیب، بر روی پایگاه داده آفیس و کالتک، پایگاه داده پای و پایگاه داده‌های اعداد و کوئل است. در پایگاه داده آفیس و کالتک، FMM دارای ۳/۹۱٪ متوسط بهبود صحت نسبت به بهترین الگوریتم مورد مقایسه و دارای ۲۱/۵۷٪ متوسط بهبود صحت نسبت به الگوریتم استاندارد NN بوده و در ۱۰ پایگاه داده از ۱۲ پایگاه داده صحت FMM بهتر شده است. در مورد پایگاه داده‌های اعداد و کوئل، FMM به ترتیب دارای ۶/۳۶٪ و ۰/۷۶٪ نسبت به بهترین الگوریتم مورد مقایسه و همچنین، به ترتیب ۲۰/۱۵٪ و ۱۶/۱۸٪ نسبت به الگوریتم استاندارد NN بهبود عملکرد دارد. همچنین، FMM در هر چهار دامنه از پایگاه داده‌های اعداد و کوئل، عملکرد بهتری از خود نشان می‌دهد. در پایگاه داده پای، متوسط بهبود صحت FMM، ۷/۵۶٪ نسبت به بهترین الگوریتم مورد مقایسه و ۳۴/۸۶٪ نسبت به الگوریتم استاندارد NN است و همچنین FMM در ۱۸ پایگاه داده از ۲۰ پایگاه داده عملکرد بهتری از خود نشان می‌دهد.

به‌طور کلی می‌توان نتیجه گرفت الگوریتم پیشنهادی، در هر چهار نوع پایگاه داده بصری، تطبیق‌پذیری بهتری بین دامنه‌های منبع و هدف ایجاد می‌کند. در ادامه، عملکرد FMM نسبت به هریک از روش‌های مورد مقایسه به تفکیک مورد بحث قرار گرفته است.

روش PCA، یک روش کاهش بعد است که یک فضای مشترک کم بعد بین دامنه‌ها ایجاد می‌کند. این روش به کاهش اختلاف توزیع بین دامنه‌ها کمک زیادی نمی‌کند و به همین دلیل، دارای صحت بسیار پایینی نسبت به الگوریتم‌های شناخته‌شده در حوزه یادگیری انتقالی می‌باشد. متوسط بهبود صحت روش FMM نسبت به روش PCA در پایگاه داده آفیس و کالتک ۱۳/۲۹٪، در پایگاه داده اعداد ۱۹/۸۸٪، در پایگاه داده کوئل ۱۵٪ و در پایگاه داده پای، ۴۴/۷۱٪ می‌باشد.

TCA، یک روش تطبیق خصوصیات است که با استفاده از اجزای انتقال، یک نمایش مشترک بین دامنه‌های منبع و هدف ایجاد می‌کند که در نمایش جدید، اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف کاهش می‌یابد. در این روش به اختلاف توزیع شرطی بین دامنه‌های منبع و هدف توجهی نشده و همچنین از داده‌های برچسب‌دار دامنه منبع در ایجاد فضای جدید استفاده نمی‌شود. درحالی‌که در روش FMM، اختلاف توزیع حاشیه‌ای و شرطی به‌طور هم‌زمان حداقل شده و از برچسب‌های نمونه‌های دامنه منبع برای ایجاد تفکیک‌پذیری بهتر بین کلاس‌ها استفاده می‌شود. به همین دلیل، عملکرد روش FMM نسبت به روش TCA، بهتر بوده و متوسط بهبود صحت روش FMM نسبت به روش TCA در پایگاه داده آفیس و

TCA [۱۰]، GFK [۱۱]، TJM [۱۲]، JDA [۱۵]، VDA [۱۶]. به دلیل این‌که تمامی این روش‌ها، روش‌های کاهش بعد می‌باشند، از طبقه‌بند استاندارد نزدیک‌ترین همسایه (NN) برای ایجاد یک طبقه‌بند بر روی داده‌های دامنه منبع جهت پیش‌بینی برچسب داده‌های دامنه هدف، استفاده کرده‌اند. دلیل دیگر استفاده از طبقه‌بند NN در الگوریتم‌های مورد مقایسه، عدم نیاز به تنظیم پارامترهای اعتبارسنجی متقابل<sup>۴</sup> است. طبقه‌بند نزدیک‌ترین همسایه، در ابتدا فاصله اقلیدسی بین هر نمونه از دامنه هدف را نسبت به نمونه‌های دامنه منبع محاسبه می‌کند، سپس، با توجه به این‌که درجه همسایگی یک در نظر گرفته شده است، برچسب نزدیک‌ترین نمونه از دامنه منبع، به عنوان برچسب هر نمونه از دامنه هدف اختصاص داده می‌شود [۲۷]. عملکرد روش پیشنهادی FMM، با بهترین نتایج گزارش‌شده از الگوریتم‌های مورد مقایسه، مورد ارزیابی قرار گرفته است.

### ۴.۴ مفروضات پیاده‌سازی

برای مقایسه روش پیشنهادی با الگوریتم‌های شناخته‌شده در حوزه یادگیری انتقالی، صحت طبقه‌بند بر روی داده‌های دامنه هدف محاسبه می‌شود. این صحت توسط محاسبه خطای پیش‌بینی در دامنه هدف به‌صورت رابطه (۱۷) تعریف می‌شود.

$$Accuracy = \frac{|\{x : x \in D_t \wedge f(x) = y(x)\}|}{n_t} \quad (17)$$

که  $D_t$  دامنه هدف،  $f(x)$  تابع پیش‌بینی به‌دست‌آمده،  $y(x)$  برچسب واقعی داده و  $n_t$  تعداد داده‌های دامنه هدف می‌باشد. در روش پیشنهادی ۴ پارامتر مختلف وجود دارد: (۱)  $\lambda$ : پارامتر نسبت در رابطه (۶)، (۲)  $k$ : تعداد ابعاد فضای جدید، (۳)  $\gamma$ : پارامتر نسبت در رابطه (۱۲)، (۴)  $\delta$ : پارامتر نسبت در رابطه (۱۲). روش FMM با مقادیر مختلف پارامترها در ۳۶ پایگاه داده مورد آزمایش قرار گرفته است. مقدار بهینه پارامترها برای پایگاه داده‌های مختلف در جدول ۱ نشان داده شده است. همچنین، تعداد بهینه تکرار الگوریتم در مرحله اول، ۱۰ در نظر گرفته شده و پارامتر  $p$  (تعداد نزدیک‌ترین همسایه هر نمونه در گراف همسایگی) برای تمام پایگاه داده‌ها از بین تعداد [۶۴، ۲]، تعداد بهینه ۵ به‌دست‌آمده است. به‌علاوه، پیاده‌سازی روش پیشنهادی FMM، توسط نرم افزار Matlab انجام گرفته است.

جدول ۱. مقدار بهینه پارامترها برای ۴ پایگاه داده بصری

| پای   | کوئل   | اعداد | آفیس و کالتک |           |
|-------|--------|-------|--------------|-----------|
| ۱۶۰   | ۲۰     | ۱۲۰   | ۲۰           | $k$       |
| ۰/۰۱  | ۰/۰۱   | ۱     | ۰/۱          | $\lambda$ |
| ۰/۰۵  | ۱      | ۰/۰۵  | ۰/۵          | $\gamma$  |
| ۰/۰۰۵ | ۰/۰۰۰۱ | ۰/۰۰۱ | ۰/۱          | $\delta$  |

کالتک ۹/۹۱٪، در پایگاه داده اعداد ۲۱/۸٪، در پایگاه داده کوپل ۱۲/۲۳٪ و در پایگاه داده پای، ۳۳/۸۱٪ می‌باشد.

در روش GFK، زیرفضاهایی از دامنه‌های منبع و هدف ایجاد می‌شود که در این زیرفضاها، اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف به حداقل می‌رسد. در این روش، به علت اینکه در ایجاد زیرفضاها ابعاد اصلی کاهش می‌یابد، داده اصلی به درستی در زیرفضای جدید نمایش داده نمی‌شود. درحالی‌که در مرحله اول روش FMM،

جدول ۲. صحت (%) طبقه‌بند در پایگاه داده آفیس و کالتک

| FMM   | VDA   | JDA   | TJM   | TCA   | GFK   | PCA   | NN    |         |
|-------|-------|-------|-------|-------|-------|-------|-------|---------|
| ۵۵/۶۴ | ۴۶/۱۴ | ۴۴/۷۸ | ۴۶/۷۶ | ۴۵/۸۲ | ۴۱/۰۲ | ۳۶/۹۵ | ۲۳/۷  | C-A     |
| ۵۲/۵۴ | ۴۶/۱  | ۴۱/۶۹ | ۳۹/۹۸ | ۳۰/۵۱ | ۴۰/۶۸ | ۳۲/۵۴ | ۲۵/۷۶ | C-W     |
| ۵۹/۲۴ | ۵۱/۵۹ | ۴۵/۲۲ | ۴۴/۵۹ | ۳۵/۶۷ | ۳۸/۸۵ | ۳۸/۲۲ | ۲۵/۴۸ | C-D     |
| ۴۷/۰۲ | ۴۲/۲۱ | ۳۹/۳۶ | ۳۹/۴۵ | ۴۰/۰۷ | ۴۰/۲۵ | ۳۴/۷۳ | ۲۶    | A-C     |
| ۵۶/۶۱ | ۵۱/۱۹ | ۳۷/۹۷ | ۴۲/۰۳ | ۳۵/۲۵ | ۳۸/۹۸ | ۳۵/۵۹ | ۲۹/۸۳ | A-W     |
| ۵۰/۳۲ | ۴۸/۴۱ | ۳۹/۴۹ | ۴۵/۲۲ | ۳۴/۳۹ | ۳۶/۳۱ | ۲۷/۳۹ | ۲۵/۴۸ | A-D     |
| ۳۱/۵۲ | ۲۷/۶  | ۳۱/۱۷ | ۳۰/۱۹ | ۲۹/۹۲ | ۳۰/۷۲ | ۲۶/۳۶ | ۱۹/۸۶ | W-C     |
| ۲۶/۳  | ۲۶/۱  | ۳۲/۷۸ | ۲۹/۹۶ | ۲۸/۸۱ | ۲۹/۷۵ | ۲۹/۳۵ | ۲۲/۹۶ | W-A     |
| ۹۲/۳۶ | ۸۹/۱۸ | ۸۹/۱۷ | ۸۹/۱۷ | ۸۵/۹۹ | ۸۰/۸۹ | ۷۷/۰۷ | ۵۹/۲۴ | W-D     |
| ۳۱/۰۸ | ۳۱/۲۶ | ۳۱/۵۲ | ۳۱/۴۳ | ۳۲/۰۶ | ۳۰/۲۸ | ۲۹/۶۵ | ۲۶/۲۷ | D-C     |
| ۴۰/۸۱ | ۳۷/۶۸ | ۳۳/۰۹ | ۳۲/۷۸ | ۳۱/۴۲ | ۳۲/۰۵ | ۳۲/۰۵ | ۲۸/۵  | D-A     |
| ۹۱/۸۶ | ۹۰/۸۵ | ۸۹/۴۹ | ۸۵/۴۲ | ۸۶/۴۴ | ۷۵/۵۹ | ۷۵/۹۳ | ۶۳/۳۹ | D-W     |
| ۵۲/۹۴ | ۴۹/۰۳ | ۴۶/۳۱ | ۴۶/۴۲ | ۴۳/۰۳ | ۴۲/۹۵ | ۳۹/۶۵ | ۳۱/۳۷ | میانگین |

جدول ۳. صحت (%) طبقه‌بند در پایگاه داده پای

| FMM   | VDA   | JDA   | TJM   | TCA   | GFK   | PCA   | NN    |         |
|-------|-------|-------|-------|-------|-------|-------|-------|---------|
| ۷۱/۱۵ | ۷۳/۴۸ | ۵۸/۸۱ | ۲۳/۸۷ | ۴۰/۷۶ | ۲۶/۱۵ | ۲۴/۸  | ۲۶/۰۹ | P1_P2   |
| ۶۶/۱۸ | ۶۲/۹۲ | ۵۴/۲۳ | ۲۸/۸۶ | ۴۱/۷۹ | ۲۷/۲۷ | ۲۵/۱۸ | ۲۶/۵۹ | P1_P3   |
| ۹۴/۲۳ | ۹۰/۵۱ | ۸۴/۵  | ۴۳/۳۷ | ۵۹/۶۳ | ۳۱/۱۵ | ۲۹/۲۶ | ۳۰/۶۷ | P1_P4   |
| ۶۹/۳  | ۵۷/۲۹ | ۴۹/۷۵ | ۱۹/۳  | ۲۹/۳۵ | ۱۷/۵۹ | ۱۶/۳  | ۱۶/۶۷ | P1_P5   |
| ۷۶/۵۶ | ۷۰/۰۲ | ۵۷/۶۲ | ۲۶/۱۴ | ۴۱/۸۱ | ۲۵/۲۴ | ۲۴/۲۲ | ۲۴/۴۹ | P2_P1   |
| ۷۱/۵۱ | ۷۳/۰۴ | ۶۲/۹۳ | ۳۷/۹۳ | ۵۱/۴۷ | ۴۷/۳۷ | ۴۵/۵۳ | ۴۶/۶۳ | P2_P3   |
| ۹۱/۲  | ۸۴/۲۹ | ۷۵/۸۲ | ۵۰/۵۳ | ۶۴/۷۳ | ۵۴/۲۵ | ۵۳/۳۵ | ۵۴/۰۷ | P2_P4   |
| ۶۷/۴۶ | ۵۴/۶۶ | ۳۹/۸۹ | ۲۱/۶۳ | ۳۳/۷  | ۲۷/۰۸ | ۲۵/۴۳ | ۲۶/۵۳ | P2_P5   |
| ۸۰/۸۵ | ۶۷/۳۵ | ۵۰/۹۶ | ۲۸/۶۶ | ۳۴/۶۹ | ۲۱/۸۲ | ۲/۹۵  | ۲۱/۳۷ | P3_P1   |
| ۷۷/۹  | ۷۰/۴۱ | ۵۷/۹۵ | ۳۵/۹۷ | ۴۷/۷  | ۴۳/۱۶ | ۴۰/۴۵ | ۴۱/۰۱ | P3_P2   |
| ۹۲/۱۹ | ۸۴/۴۷ | ۶۸/۴۵ | ۵۱/۹۷ | ۵۶/۲۳ | ۴۶/۴۱ | ۴۶/۱۴ | ۴۶/۵۳ | P3_P4   |
| ۶۳/۳  | ۵۲/۳۹ | ۳۹/۹۵ | ۲۵/۳۱ | ۳۳/۱۵ | ۲۶/۷۸ | ۲۵/۳۱ | ۲۶/۲۳ | P3_P5   |
| ۹۶/۰۴ | ۹۱/۶  | ۸۰/۵۸ | ۴۵/۷۱ | ۵۵/۶۴ | ۳۴/۲۴ | ۳۱/۹۶ | ۳۲/۹۵ | P4_P1   |
| ۹۴/۹  | ۹۱/۴۷ | ۸۲/۶۳ | ۵۷/۵۸ | ۶۷/۸۳ | ۶۲/۹۲ | ۶۰/۹۶ | ۶۲/۶۸ | P4_P2   |
| ۹۲/۷۷ | ۹۰/۹۳ | ۸۷/۲۵ | ۷۱/۶۳ | ۷۵/۸۶ | ۷۳/۳۵ | ۷۲/۱۸ | ۷۳/۲۲ | P4_P3   |
| ۷۳/۷۱ | ۶۳/۳۶ | ۵۴/۶۶ | ۳۰/۹۴ | ۴۰/۲۶ | ۳۷/۳۸ | ۳۵/۱۱ | ۳۷/۱۹ | P4_P5   |
| ۷۰/۰۸ | ۵۵/۷  | ۴۶/۴۶ | ۳۷/۱۳ | ۲۶/۹۸ | ۲۰/۳۵ | ۱۸/۸۵ | ۱۸/۴۹ | P5_P1   |
| ۷۳/۳  | ۶۱/۵۷ | ۴۲/۰۵ | ۲۲/۶۵ | ۲۹/۹  | ۲۴/۶۲ | ۲۳/۳۹ | ۲۴/۱۹ | P5_P2   |
| ۶۶/۱۸ | ۵۵/۵۸ | ۵۳/۳۱ | ۲۸/۸۶ | ۲۹/۹  | ۲۸/۴۹ | ۲۷/۲۱ | ۲۸/۳۱ | P5_P3   |
| ۸۲/۴۳ | ۶۸/۸۲ | ۵۷/۰۱ | ۳۲/۵۹ | ۳۳/۶۴ | ۳۱/۳۳ | ۳۰/۳۴ | ۳۱/۲۴ | P5_P4   |
| ۷۸/۵۶ | ۷۱    | ۶۰/۲۴ | ۳۵/۵۳ | ۴۴/۷۵ | ۳۵/۳۵ | ۳۳/۸۵ | ۳۴/۷۶ | میانگین |

جدول ۴. صحت (%) طبقه‌بند در پایگاه داده‌های اعداد و کوپل

| FMM   | VDA   | JDA   | TJM   | TCA   | GFK   | PCA   | NN    |         |
|-------|-------|-------|-------|-------|-------|-------|-------|---------|
| ۷۱/۱  | ۶۲/۹۵ | ۵۹/۶۵ | ۵۲/۲۵ | ۵۱/۰۵ | ۴۶/۴۵ | ۴۴/۹۵ | ۴۴/۷  | U_M     |
| ۷۹/۸۳ | ۷۴/۷۲ | ۶۷/۲۸ | ۶۳/۲۸ | ۵۶/۲۸ | ۶۷/۲۲ | ۶۶/۲۲ | ۶۵/۹۴ | M_U     |
| ۷۵/۴۷ | ۶۸/۸۴ | ۶۳/۴۷ | ۵۷/۷۷ | ۵۳/۶۷ | ۵۶/۸۴ | ۵۵/۵۹ | ۵۵/۳۲ | میانگین |

|       |       |       |       |       |       |       |       |         |
|-------|-------|-------|-------|-------|-------|-------|-------|---------|
| ۹۹/۷۲ | ۹۹/۳۱ | ۸۹/۳۱ | ۹۱/۶۷ | ۸۸/۴۷ | ۷۲/۵  | ۸۴/۷۲ | ۸۳/۶۱ | C1_C2   |
| ۹۹/۰۴ | ۹۷/۹۲ | ۸۸/۴۷ | ۹۱/۵۳ | ۸۵/۸۳ | ۷۴/۱۷ | ۸۴/۰۳ | ۸۲/۷۸ | C2_C1   |
| ۹۹/۳۸ | ۹۸/۶۲ | ۸۸/۸۹ | ۹۱/۶  | ۸۷/۱۵ | ۷۳/۳۴ | ۸۴/۳۸ | ۸۳/۲  | میانگین |

گرفته می‌شوند. نتایج به‌دست‌آمده از روش FMM با مقادیر مختلف پارامترهای  $k$  و  $\lambda$  با روش‌های JDA، TJM و VDA در پایگاه‌داده‌های آفیس و کالکت، اعداد، کوئل و پای مقایسه شده است. شکل‌های ۳ و ۴، نشان‌دهنده نتایج به‌دست‌آمده برای مقادیر مختلف  $k$  و  $\lambda$  برای روش‌های TJM، JDA، VDA و FMM در پایگاه‌داده‌های A-C، C1-C2، U-M و P1-P3 هستند. همچنین، شکل‌های ۵ و ۶، نشان‌دهنده نتایج به‌دست‌آمده از روش FMM برای مقادیر مختلف  $\gamma$  و  $\delta$  در پایگاه‌داده‌های A-C، C1-C2، U-M و P1-P3 می‌باشند. در مورد پارامتر  $p$ ، در مقادیر بالا به علت اینکه نمونه‌های کاملاً متفاوت از هم در گراف به یکدیگر متصل می‌شوند، صحت طبقه‌بند در تمام پایگاه‌داده‌ها کاهش می‌یابد. همچنین در مقادیر پایین پارامتر  $p$ ، ساختار نمونه‌ها به‌طور دقیق محاسبه و حفظ نخواهد شد. بدین‌ترتیب، در تمامی پایگاه‌داده‌ها، بهترین مقدار برای پارامتر  $p$ ، تعداد ۵ به‌دست‌آمده است.

نتایج به‌دست‌آمده از آزمایش‌های انجام‌گرفته بر روی پایگاه‌داده آفیس و کالکت، نشان می‌دهند پایگاه داده‌های مختلف در آفیس و کالکت دارای رفتار متفاوتی نسبت به مقادیر مختلف پارامتر  $k$  می‌باشند. صحت به‌دست‌آمده برای پارامتر  $k$  در هر روش، نشان‌دهنده صحت روش برای بازسازی داده‌ها در فضای جدید است؛ به‌عبارت‌دیگر، نشان‌دهنده صحت انتقال دانش از دامنه منبع به دامنه هدف در فضای جدید می‌باشد. برای به دست آوردن مقدار محدوده بهینه پارامتر، محدوده‌ای انتخاب می‌شود که دارای بهینه‌ترین صحت بر روی تعداد بیشتری از پایگاه داده‌ها باشد. برای بیشتر پایگاه داده‌ها محدوده بهینه پارامتر  $k$  در پایگاه داده آفیس و کالکت، ۲۰ بعد است (برای نمونه، عملکرد پایگاه داده C-A نسبت به مقادیر مختلف پارامتر  $k$  در شکل ۳ (C-A) نشان داده شده است). باوجوداینکه در مورد برخی پایگاه‌داده‌ها در مقادیر پایین پارامتر  $k$ ، صحت پایینی به‌دست‌آمده است، هدف، پیدا کردن محدوده بهینه کلی برای پایگاه داده آفیس و کالکت است. همچنین نتایج نشان می‌دهند در بیشتر پایگاه داده‌های آفیس و کالکت مقادیر پایین  $\lambda$ ، به علت اینکه موجب ایجاد یک مساله بهینه بدیهی می‌شوند، دارای صحت پایینی هستند. به‌طورکلی، برای بیشتر پایگاه‌داده‌های آفیس و کالکت محدوده بهینه پارامتر  $\lambda$ ، [۱۰، ۰/۵] است. برای نمونه، عملکرد پایگاه داده C-A نسبت به مقادیر مختلف پارامتر  $\lambda$  در شکل ۴ (C-A) نشان داده شده است.

عملکرد پایگاه داده C-A نسبت به مقادیر مختلف پارامترهای  $\delta$  و  $\gamma$ ، به ترتیب در شکل‌های ۵ (نمودار C-A) و ۶ (نمودار C-A) نشان داده شده است. به‌طورکلی، پایگاه داده آفیس و کالکت در مقادیر مختلف پارامتر  $\delta$ ، دچار کاهش صحت در مقادیر پایین پارامتر مورد نظر می‌شود. از آنجا که مقادیر پایین پارامتر  $\delta$ ، باعث می‌شود مدل دچار بیش‌برازش<sup>۲۲</sup> شود، مدل تطبیق‌پذیری بین دامنه‌های منبع و هدف

روش‌های TJM، JDA و VDA، از جمله جدیدترین روش‌های ایجاد تطبیق بین دامنه‌ها توسط ایجاد یک نمایش مشترک بین دامنه‌های منبع و هدف هستند. روش TJM توسط یک مساله بهینه‌سازی پیچیده، اختلاف توزیع حاشیه‌ای بین دامنه‌های منبع و هدف را حداقل می‌کند. روش JDA، به‌طور هم‌زمان، اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌های منبع و هدف را کاهش می‌دهد و به همین دلیل، دارای عملکرد بهتری نسبت به روش TJM می‌باشد. روش VDA، برپایه روش JDA بوده که از خوشه‌بندی مستقل از دامنه برای بهبود صحت طبقه‌بند استفاده می‌کند. بااین‌حال، به دلیل خصوصیات متفاوت داده‌های آموزشی و تست در دامنه‌های منبع و هدف، طبقه‌بند ایجاد شده در نمایش جدید توسط روش‌های TJM، JDA و VDA، نمی‌تواند با صحت بالایی برچسب داده‌های دامنه هدف را پیش‌بینی کند. روش FMM، علاوه بر حداقل کردن اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌های منبع و هدف در فضای جدید، با بهره‌گیری از توزیع هندسی داده‌ها، ابعاد تفکیک‌کننده کلاس‌ها در دامنه‌های منبع و هدف را باهم تطبیق داده و در نتیجه، صحت طبقه‌بند برای پیش‌بینی برچسب‌های داده‌های تست بهبود می‌یابد.

متوسط بهبود صحت روش FMM نسبت به روش TJM در پایگاه‌داده آفیس و کالکت ۵۲/۶٪، در پایگاه‌داده اعداد ۱۷/۷٪، در پایگاه‌داده کوئل ۷۸/۷٪ و در پایگاه‌داده پای ۴۳/۰۳٪ است. همچنین، متوسط بهبود صحت روش FMM نسبت به روش JDA در پایگاه‌داده آفیس و کالکت ۶۳/۶٪، در پایگاه‌داده اعداد ۱۲٪، در پایگاه‌داده کوئل ۴۹/۱۰٪ و در پایگاه‌داده پای ۱۸/۳۲٪ و همچنین، نسبت به روش VDA در پایگاه‌داده آفیس و کالکت ۳/۹۱٪، در پایگاه‌داده اعداد ۶۳/۶٪، در پایگاه‌داده کوئل ۷۶/۰٪ و در پایگاه‌داده پای ۷/۵۶٪ است.

برای بررسی همگرایی روش FMM در تکرارهای مختلف، یک آزمایش دیگر نیز ترتیب داده شده است. از آنجایی که روش FMM یک روش تکرارشونده است، بررسی شده است که در پایگاه‌داده‌های مختلف میزان تکرار بهینه الگوریتم چقدر می‌باشد. شکل ۲، الگوریتم FMM را با روش‌های TJM، JDA و VDA در ۲۰ تکرار مقایسه می‌کند. چنانچه از نتایج بر می‌آید، در بیشتر موارد، FMM در ۱۰ تکرار اولیه به همگرایی رسیده است و تعداد تکرار بیشتر، تاثیری در افزایش صحت الگوریتم ندارد. از این رو، تمام نتایج گزارش شده برای FMM حاصل انجام ۱۰ تکرار بر روی الگوریتم می‌باشند.

#### ۴.۵ ارزیابی پارامترها

برای به دست آوردن تنظیمات مدل، روش FMM با مقادیر مختلف پارامترها مورد ارزیابی قرار می‌گیرد. پارامتر  $k$ ، در محدوده [۲۲۰، ۲۰] و پارامترهای  $\lambda$ ،  $\delta$  و  $\gamma$  در محدوده [۰/۰۰۰۰۱، ۱۰] در نظر

در پایگاه داده کوپل، نتایج به دست آمده نشان دهنده حساسیت به مقادیر بالای پارامتر  $\lambda$  است. از آنجایی که براساس نتایج، با مقادیر بالای پارامتر  $\lambda$  در پایگاه داده کوپل، صحت پیش‌بینی برچسب کاهش می‌یابد، یک مدل تطبیق‌پذیر بین دامنه‌های منبع و هدف ایجاد نخواهد شد. براساس شکل ۴ (C1-C2)، محدوده بهینه در پایگاه داده کوپل برای پارامتر  $\lambda$ ،  $[0.001, 0.1]$  است.

همان‌طور که در شکل ۵ نمودار C1-C2، نشان داده شده است، محدوده بهینه در پایگاه داده کوپل برای پارامتر  $\delta$ ،  $[0.001, 0.0001]$  می‌باشد، به دلیل اینکه در این محدوده، طبقه‌بند دارای بهترین صحت با پیش‌بینی دقیق برچسب برای تقریباً تمام داده‌های دامنه هدف می‌باشد. همچنین، براساس شکل ۶ نمودار C1-C2، برای پایگاه داده کوپل در مقادیر پایین پارامتر  $\gamma$ ، به دلیل نادیده گرفتن اطلاعات داده‌های برچسب‌دار دامنه منبع و در مقادیر بالای پارامتر  $\gamma$ ، به دلیل نادیده گرفتن اطلاعات داده‌های بدون برچسب دامنه هدف، در ایجاد طبقه‌بند انطباقی، صحت مدل کاهش می‌یابد. محدوده بهینه در پایگاه داده کوپل برای پارامتر  $\gamma$ ،  $[0.5, 1]$  است که در این شرایط مدل طبقه‌بند بهترین عملکرد و بالاترین صحت پیش‌بینی را دارد.

همچنین، براساس نتایج به دست آمده، محدوده‌های بهینه برای پارامترهای  $k$ ،  $\lambda$ ،  $\delta$  و  $\gamma$  در تعداد بیشتری از پایگاه داده‌های پای، به ترتیب،  $[0.18, 0.16]$ ،  $[0.01, 0.01]$ ،  $[0.05, 0.05]$  و  $[0.5, 0.5]$  می‌باشند. برای نمونه، رفتار پایگاه داده P1-P2 نسبت به مقادیر مختلف پارامترهای  $k$  و  $\lambda$  به ترتیب، در شکل ۳ (P1-P2) و شکل ۴ (P1-P2) و همچنین، نسبت به مقادیر مختلف پارامترهای  $\delta$  و  $\gamma$  به ترتیب، در شکل‌های ۵ و ۶ نمودار P1-P2 نشان داده شده است.

#### ۴-۵ ارزیابی زمانی

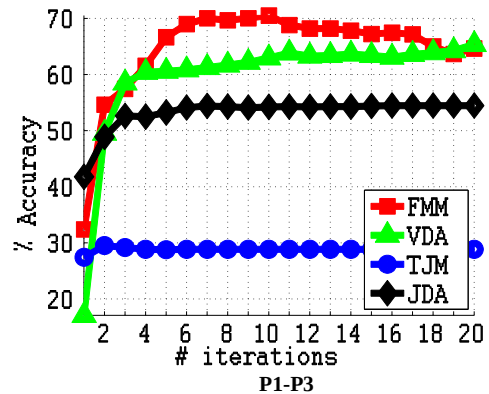
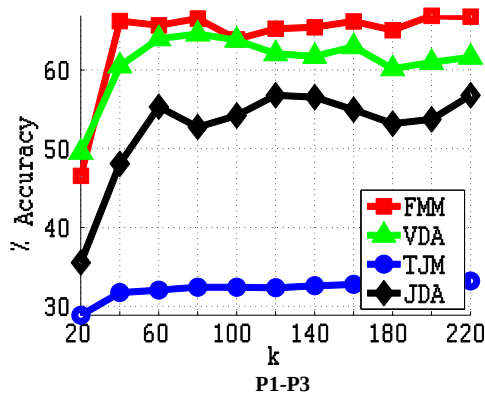
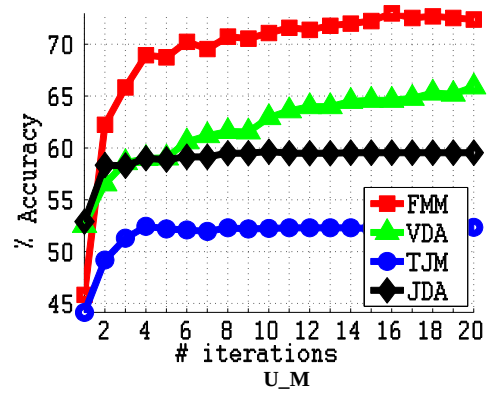
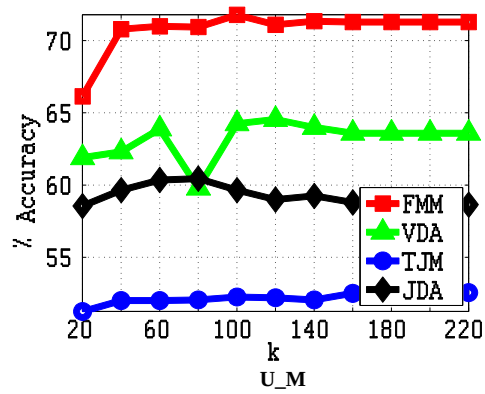
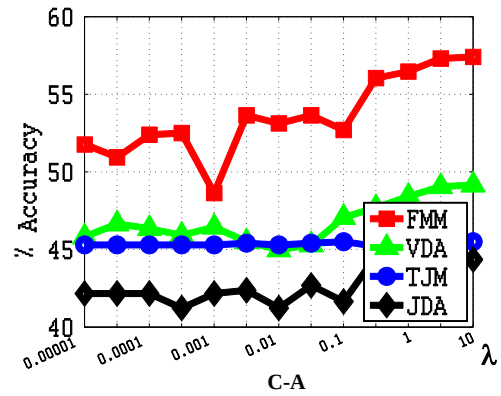
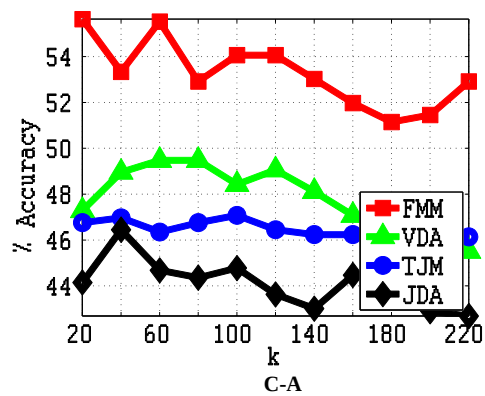
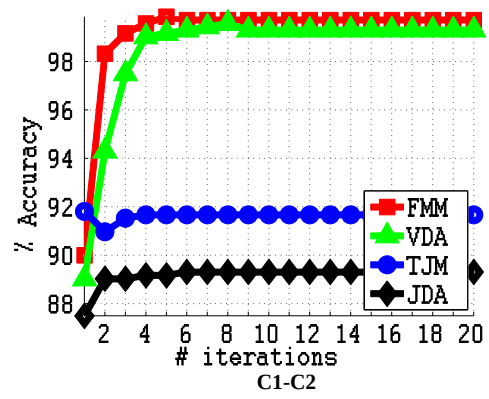
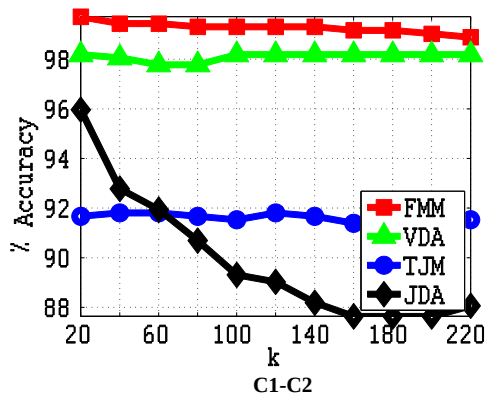
در این بخش، پیچیدگی زمانی روش پیشنهادی FMM نسبت به روش‌های مورد مقایسه، مورد ارزیابی قرار می‌گیرد. اگر  $k$ ،  $m$ ،  $n$  و  $C$  به ترتیب تعداد کل داده‌ها، اندازه ابعاد فضای اصلی، اندازه ابعاد فضای جدید و تعداد کلاس‌ها باشند، روش NN (با در نظر گرفتن تعداد همسایگی برابر با ۱ برای نمونه‌ها)، دارای پیچیدگی زمانی  $O(nm)$  می‌باشد. پیچیدگی زمانی روش PCA، وابسته به تعداد ابعاد فضای جدید بوده و به صورت  $O(k^2n + k^3)$  تعریف می‌شود. روش TCA، با وابستگی نسبت به تعداد داده‌ها، با پیچیدگی  $O(mn^2)$  انجام می‌گیرد. به همین ترتیب، روش TJM دارای پیچیدگی  $O(km^2 + Cn^2)$  می‌باشد.

ایجاد نمی‌شود. همچنین، به دلیل اینکه مقادیر بالای پارامتر  $\gamma$ ، باعث نادیده گرفتن اطلاعات داده‌های برچسب‌دار دامنه منبع در ایجاد طبقه‌بند انطباقی می‌شوند، پایگاه داده آفیس و کالتک دچار کاهش صحت در مقادیر بالای پارامتر  $\gamma$  می‌باشد. محدوده بهینه در پایگاه داده آفیس و کالتک برای پارامتر  $\delta$ ،  $[0.1, 0.5]$  و برای پارامتر  $\gamma$ ،  $[0.0001, 0.5]$  است.

نتایج به دست آمده برای مقادیر مختلف پارامتر  $k$  در پایگاه داده اعداد نشان می‌دهد دو پایگاه داده U\_M و M\_U، رفتار متفاوتی در مورد مقادیر مختلف پارامتر  $k$  دارند. U\_M دارای حساسیت کمی نسبت به مقادیر مختلف  $k$  است (همان‌طور که در شکل ۳ (U-M) نشان داده شده است) در حالی که M\_U در مقادیر پایین پارامتر  $k$ ، صحت پایینی در بازسازی داده‌ها در فضای جدید دارد. بهترین محدوده برای پارامتر  $k$  در پایگاه داده اعداد،  $[100, 160]$  است. همچنین، نتایج به دست آمده نشان می‌دهند پایگاه داده اعداد در مقادیر بالای پارامتر  $\lambda$ ، به دلیل این که یک مدل تطبیق‌پذیر بین دامنه‌های منبع و هدف ایجاد نمی‌شود، دارای صحت پایینی می‌باشد و همچنین، در مقادیر پایین پارامتر  $\lambda$ ، به دلیل یافتن یک مساله بهینه بدیهی، مدل دچار کاهش صحت شده است. محدوده بهینه پایگاه داده اعداد برای پارامتر  $\lambda$ ،  $[0.005, 1]$  است (برای نمونه، عملکرد پایگاه داده U-M نسبت به مقادیر مختلف پارامتر  $\lambda$  در شکل ۴ (U-M) نشان داده شده است).

همچنین، براساس نتایج نشان داده شده در شکل ۵ نمودار U-M، در مقادیر بالای پارامتر  $\delta$  در پایگاه داده اعداد صحت مدل کاهش می‌یابد. دلیل این مساله این است که با در نظر گرفتن مقادیر بالا برای پارامتر  $\delta$ ، به دلیل این که پیچیدگی مدل افزایش یافته و تأثیر سایر عوامل در ایجاد طبقه‌بند انطباقی نادیده گرفته شده، ساختار اصلی داده‌ها در ایجاد تطبیق بین دامنه‌های منبع و هدف استفاده نمی‌شود. از این رو، محدوده بهینه پایگاه داده اعداد برای پارامتر  $\delta$ ،  $[0.1, 0.005]$  می‌باشد. به همین ترتیب براساس نتایج نشان داده شده در شکل ۶ نمودار U-M، پایگاه داده اعداد، صحت پایینی در مقادیر پایین پارامتر  $\gamma$  دارد. مقدار پایین پارامتر  $\gamma$ ، باعث عدم بهره‌گیری از اطلاعات داده‌های بدون برچسب دامنه هدف در ایجاد طبقه‌بند انطباقی می‌شود. محدوده بهینه پایگاه داده اعداد برای پارامتر  $\gamma$ ،  $[0.1, 0.5]$  است.

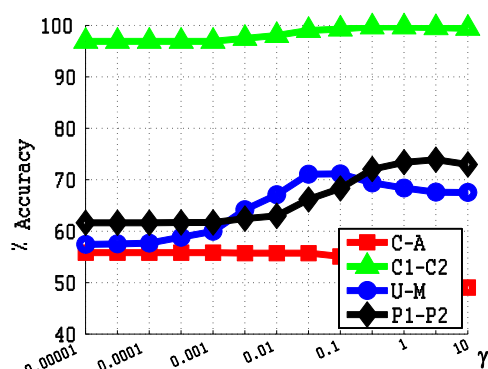
نتایج در پایگاه داده کوپل، نشان می‌دهد در مقادیر بالای پارامتر  $k$  نمی‌توان با صحت مناسبی فضای تطبیق‌پذیر بین دامنه‌های منبع و هدف ایجاد کرد، به عبارت دیگر نمی‌توان با صحت بالایی دانش را از دامنه منبع به دامنه هدف انتقال داد. محدوده بهینه پایگاه داده کوپل برای پارامتر  $k$ ، ۲۰ بعد می‌باشد (براساس شکل ۳ (C1-C2)) که در این محدوده، مدل طبقه‌بند تنها برای ۳ نمونه از ۷۲۰ نمونه تست در پایگاه داده C1\_C2 و ۷ نمونه از ۷۲۰ نمونه تست در پایگاه داده C2\_C1، برچسب نادرستی پیش‌بینی کرده است. در چنین شرایطی، روش پیشنهادی FMM توانسته است به طور تقریباً کاملی، دامنه‌های منبع و هدف را با یکدیگر تطابق دهد.



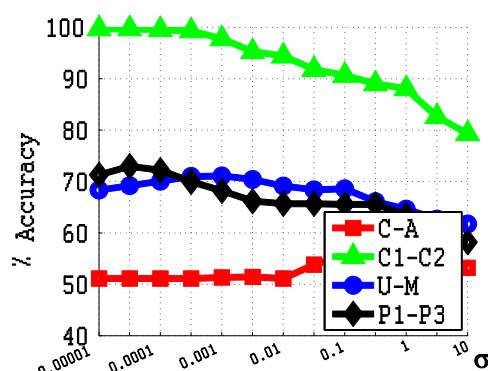
شکل ۳. ارزیابی صحت روش‌های TJM، JDA، VDA و FMM در پایگاه داده‌های C1-C2، U-M، A-C، P1-P3 با مقادیر مختلف پارامتر  $k$

شکل ۲. ارزیابی صحت روش‌های TJM، JDA، VDA و FMM در پایگاه داده‌های C1-C2، U-M، A-C، P1-P3 در ۲۰ تکرار

شکل ۴. ارزیابی صحت روش‌های TJM، JDA، VDA و FMM در پایگاه داده‌های C1-C2، U-M، C-A، P1-P3 با مقادیر مختلف پارامتر  $\lambda$



شکل ۵. ارزیابی صحت پایگاه داده‌های C1-C2، U-M، C-A، P1-P3 با مقادیر مختلف پارامتر  $\gamma$

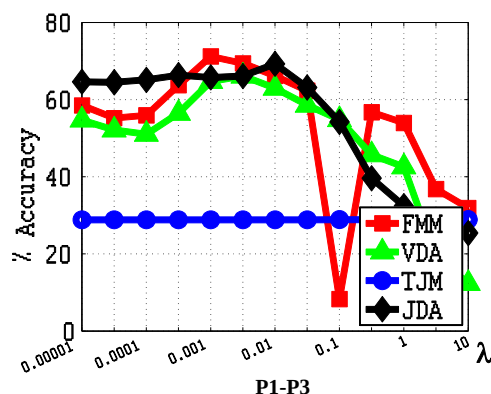
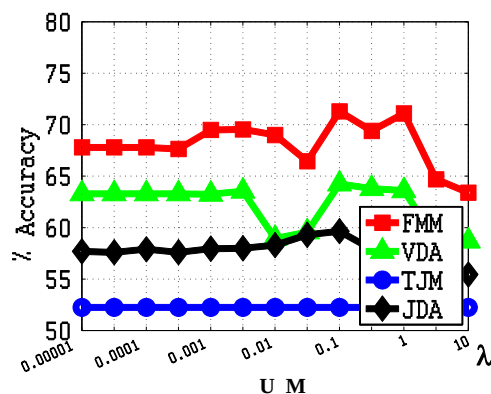
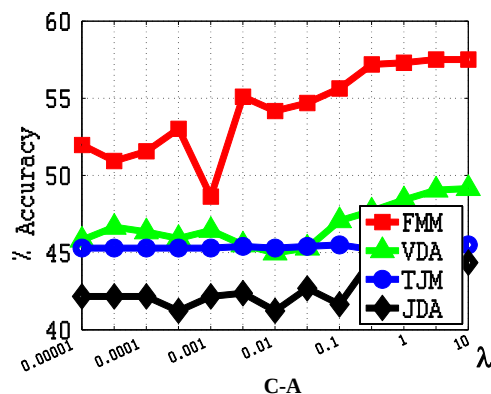
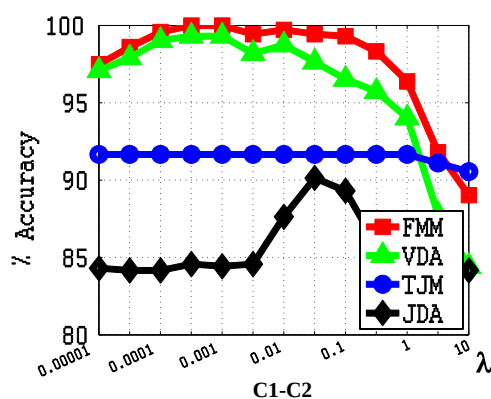


شکل ۶. ارزیابی صحت پایگاه داده‌های C1-C2، U-M، C-A، P1-P3 با مقادیر مختلف پارامتر  $\delta$

روش‌های JDA و VDA نیز هر دو دارای پیچیدگی زمانی یکسان  $O(mn + m^2 + Cn^2)$  هستند. به‌طور کلی، با فرض  $n \leq m > k \leq C$ ، روش‌های PCA و NN، سریع‌ترین روش‌ها می‌باشند (ولی تطبیق دامنه انجام نمی‌دهند). پیچیدگی زمانی روش TCA از FMM بیشتر بوده و پیچیدگی زمانی روش‌های TJM، JDA و VDA به‌صورت کلی  $O(Cn^2)$  می‌باشد که برابر با پیچیدگی زمانی روش FMM است.

## ۶ نتیجه‌گیری و کارهای آتی

در این مقاله، روش یادگیری انتقالی بدون نظارت، با عنوان تطبیق خصوصیات و مدل (FMM) پیشنهاد شد. یک چهارچوب دومرحله‌ای با ترکیبی از روش‌های تطبیق خصوصیات و تطبیق مدل است. در مرحله اول روش FMM، با استفاده از روش‌های تطبیق خصوصیات، یک فضای جدید برای نمایش دامنه‌های منبع و هدف ایجاد می‌شود. در نمایش جدید، اختلاف توزیع شرطی و حاشیه‌ای بین دامنه‌های منبع و هدف به‌صورت هم‌زمان، به حداقل رسانده می‌شود. همچنین، برای افزایش صحت طبقه‌بندی در نمایش جدید، خوشه‌بندی مستقل از دامنه در دامنه منبع اعمال می‌شود.



- [8] M. Long, J. Wang, G. Ding, S. J. Pan and P. Yu, "Adaptation regularization: a general framework for transfer learning", IEEE Trans. Knowl. Data Eng, vol. 26, pp. 1076-1089, 2013.
- [9] S. J. Pan, I. W. Tsang, J. T. Kwok and Q. Yang, "Domain adaptation via transfer component analysis", IEEE Trans. Neural Netw, vol. 22, pp. 199-210, 2011.
- [10] S. Si, D. Tao and B. Geng, "Bregman divergence-based regularization for transfer subspace learning", IEEE Trans Knowl Data Eng, vol. 22, no. 7, pp. 929-942, 2010.
- [11] B. Gong, Y. Shi, F. Sha and K. Grauman, "Geodesic flow kernel for unsupervised domain adaptation", Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, pp. 2066-2073, 2012.
- [12] M. Long, J. Wang, G. Ding, J. Sun and P. S. Yu, "Transfer joint matching for unsupervised domain adaptation", IEEE conference on computer vision and pattern recognition, pp. 1410-1417, 2014.
- [13] S. Satpal and S. Sarawagi, "Domain adaptation of conditional probability models via feature subsetting", Proceedings of PKDD, vol. 4702, pp. 224-235, 2007.
- [14] B. Quanz, J. Huan and M. Mishra, "Knowledge transfer with low-quality data: a feature extraction issue", IEEE Trans Knowl Data Eng, vol. 24, no. 10, pp. 1789-1802, 2012.
- [15] M. Long, J. Wang, G. Ding, J. Sun and S. Yu Philip, "Transfer feature learning with joint distribution adaptation", IEEE international conference on computer vision, pp. 2200-2207, 2013.
- [16] J. Tahmoresnezhad and S. Hashemi, "Visual domain adaptation via transfer feature learning", KnowlInf Syst, vol. 50, no. 2, pp. 585-605, 2016.
- [17] 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.
- [18] M. Belkin, P. Niyogi, V. Sindhwani, "Manifold regularization: a geometric framework for learning from labeled and unlabeled examples", J. Mach. Learn. Res, vpl. 7, pp. 2399-2434, 2006.
- [19] U. von Luxburg, "A tutorial on spectral clustering", Stat. Comput, vol. 17, no. 4, pp. 395-416, 2007.
- [20] B. Schölkopf, R. Herbrich and A. J. Smola, "A generalized representer theorem", Proceedings of the Conference on Computational Learning Theory, pp. 416-426, 2001.
- [21] K. Saenko, B. Kulis, M. Fritz and T. Darrell, "Adapting visual category models to new domains", Proceedings of the European Conference on Computer Vision, pp. 213-226, 2010.
- [22] G. Griffin, A. Holub and P. Perona, "Caltech-256 object category dataset", Technical Report 7694, 2007.
- [23] J. J. Hull, "A database for handwritten text recognition research", IEEE Trans. Pattern Anal. Mach. Intell, vol. 16, no. 5, pp. 550-554, 1994.
- [24] 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.
- [25] S. A. Nene, S. K. Nayar and H. Murase, "Columbia object image library (COIL-20)", Technical Report CUCS, 1996.
- [26] 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.

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

اما در شرایطی که اختلاف توزیع دامنه‌های منبع و هدف زیاد باشد، فقط با کاهش اختلاف توزیع حاشیه‌ای و شرطی بین دامنه‌های منبع و هدف، نمی‌توان باعث ایجاد تطبیق‌پذیری طبقه‌بند ایجاد شده بر روی دامنه منبع با دامنه هدف شد. در چنین شرایطی، حتی با ایجاد تفکیک‌پذیری بین کلاس‌ها، به علت اینکه ساختار داده‌های دامنه هدف با داده‌های دامنه منبع متفاوت است، یک طبقه‌بند تطبیق‌پذیر بین دامنه‌ای ایجاد نخواهد شد. به همین دلیل در مرحله دوم روش FMM، یک طبقه‌بند انطباقی بر روی هر دو نمونه‌های برچسب‌دار دامنه منبع و نمونه‌های بدون برچسب دامنه هدف در نمایش جدید ایجاد می‌شود. طبقه‌بند انطباقی با افزایش تطبیق‌پذیری بین ساختار دامنه‌های منبع و هدف، باعث افزایش صحت طبقه‌بند در پیش‌بینی برچسب برای داده‌های بدون برچسب دامنه هدف می‌شود. روش FMM، بر روی ۳۶ پایگاه‌داده بصری مورد آزمایش قرار گرفته است که این پایگاه‌داده‌ها، دارای اختلاف توزیع قابل‌توجهی نسبت به یکدیگر هستند. نتایج به‌دست‌آمده، نشان‌دهنده بهبود قابل‌ملاحظه‌ای از کارایی روش FMM نسبت به جدیدترین روش‌های حوزه یادگیری ماشین و یادگیری انتقالی بر روی دامنه‌های مختلف می‌باشد.

برای ادامه راه، ما در حال برنامه ریزی جهت توسعه روش FMM برای سیستم‌های چندمنبعی هستیم. در این راستا، به دنبال انتقال دانش از چند منبع مرتبط به یک منبع هدف می‌باشیم تا بتوانیم صحت پیش‌بینی برچسب دامنه هدف را هرچه بیشتر ارتقا بخشیم.

## مراجع

- [1] B. Gong, K. Grauman and F. Sha, "Reshaping visual datasets for domain adaptation", Proceedings of the Advances in Neural Information Processing Systems, vol. 26, pp. 1286-1294, 2013.
- [2] B. Gong, K. Grauman and F. Sha, "Learning kernels for unsupervised domain adaptation with applications to visual object recognition", Int J Comput Vision vol. 109, pp. 3-27, 2014.
- [3] Jolliffe I, *Principal component analysis*, Wiley, vol. 2, pp. 433-459, 2002.
- [۴] طاهره زارع بیدکی، محمدتقی صادقی، «بهینه‌سازی وزن‌ها در کرنل مرکب برای طبقه‌بند مبتنی بر نمایش تنک کرنلی»، مجله مهندسی برق دانشگاه تبریز، جلد ۴۷، شماره ۳، صفحات ۱۰۵۹-۱۰۷۲، ۱۳۹۶.
- [5] B. Gong, K. Grauman and F. Sha, "Connecting the dots with landmarks: discriminatively learning domain-invariant features for unsupervised domain adaptation", Proceedings of the International Conference on Machine Learning, vol. 28, no. 1, pp. 222-230, 2013.
- [6] L. Duan L, I. W. Tsang, D. Xu and S. J. Maybank, "Domain transfer SVM for video concept detection", IEEE Conference on computer vision and pattern recognition, pp. 1375-1381, 2009.
- [7] L. Bruzzone and M. Marconcini, "Domain adaptation problems: a DASVM classification technique and a circular validation strategy", IEEE Trans Pattern Anal Mach Intell, vol. 32, no. 5, pp. 770-787, 2010.

## زیرنویس‌ها

<sup>3</sup> Domain shift

<sup>4</sup> Transfer learning

<sup>5</sup> Adaptation

<sup>6</sup> Semi-supervised

<sup>7</sup> Unsupervised

<sup>8</sup> Source domain

<sup>1</sup> Cross-domain

<sup>2</sup> Classifier



- 
- <sup>9</sup> Target domain
  - <sup>10</sup> Label
  - <sup>11</sup> Feature and model matching
  - <sup>12</sup> Adaptive classifier
  - <sup>13</sup> Low-dimensional
  - <sup>14</sup> Principal component analysis
  - <sup>15</sup> Sparse representation
  - <sup>16</sup> Maximum mean discrepancy (MMD)
  - <sup>17</sup> Adaptation regularization for transfer learning
  - <sup>18</sup> Transfer component analysis
  - <sup>19</sup> Geodesic flow kernel
  - <sup>20</sup> Transfer component
  - <sup>21</sup> Transfer joint matching
  - <sup>22</sup> Domain adaptation of conditional probability models via feature subsetting
  - <sup>23</sup> Transfer feature learning with joint distribution adaptation
  - <sup>24</sup> Visual domain adaptation via transfer feature learning
  - <sup>25</sup> Domain invariant clustering
  - <sup>26</sup> Domain
  - <sup>27</sup> Task
  - <sup>28</sup> Reconstruction error
  - <sup>29</sup> Hilbert
  - <sup>30</sup> Frobenius norm
  - <sup>31</sup> Representer theory
  - <sup>32</sup> Empirical risk
  - <sup>33</sup> Loss function
  - <sup>34</sup> Manifold assumption
  - <sup>35</sup> Nearest neighbor graph
  - <sup>36</sup> Laplacian matrix
  - <sup>37</sup> Regularization parameters
  - <sup>38</sup> Eigen decomposition
  - <sup>39</sup> Feature extraction
  - <sup>40</sup> Codebook
  - <sup>41</sup> Cross validation
  - <sup>42</sup> Overfitting