

# روش یادگیری گروهی چندوجهی برای کدگشایی



## اشیای دیداری از دادگان fMRI مغزی

اسامه حورانی، نصراله مقدم چرکری\* و سعید جلیلی

گروه کامپیوتر، دانشکده برق و کامپیوتر، دانشگاه تربیت مدرس، تهران، ایران

### چکیده

با توجه به گسترش روزافزون پژوهش‌های علوم شناختی، کدگشایی مغز انسان یک موضوع داغ در حوزه علوم عصب‌شناسی محاسباتی است. در این راستا پژوهش‌های متعددی جهت ارائه روشی کارا و مؤثر برای کدگشایی فعالیت مغز انسان با پردازش دادگان fMRI در حال انجام است. خروجی این روش‌ها به‌طور عمومی معطوف به ارائه یک مدل محاسباتی تعمیم‌یافته است که امکان تشخیص سیگنال مغزی و تعلق آن به عامل محرک (شیء دیداری) را ارائه می‌دهد. دادگان مغزی دارای ابعاد زمانی و فضایی زیادی هستند که سبب افزایش تعداد ویژگی‌ها نیز می‌شود. همچنین استخراج ویژگی‌های مفید از تصاویر مغزی مانند fMRI کاری پیچیده است. این امر سبب طولانی شدن عمل هم‌گرایی در الگوریتم‌های یادگیری برای ایجاد مدل مناسب می‌شود؛ با توجه به چالش‌های یادشده، روش‌های یادگیری گروهی چندوجهی یک پیشنهاد مناسب برای حل مسأله کدگشایی مغز محسوب می‌شود که تلفیقی مناسب بین ویژگی‌های عملکردی متفاوت در دادگان مغزی ایجاد می‌کند. در روش پیشنهادی داده‌های آموزشی بر اساس اطلاعات متقابل در فضای ویژگی خوشه‌بندی می‌شوند، به‌صورتی که فضای ویژگی به چند وجه تفکیک می‌شود؛ سپس روی هر وجه ویژگی یک مدل ماشین بردار پشتیبان به‌صورت موازی آموزش داده می‌شود. در مرحله آزمون، فضای ویژگی دادگان آزمون نیز به‌صورت مشابه دادگان آموزش تقسیم می‌شود؛ و هر بردار ویژگی به مدل مربوطه تخصیص داده خواهد شد. از هر مدل یک بردار احتمالاتی تولید می‌شود و با هم‌جوشی این بردارها ماتریس پروفایل تصمیم‌گیری ساخته خواهد شد؛ در نهایت عملگرهای وزن‌دار مرتب‌شده اعمال می‌شود. جهت بررسی کارایی رویکرد پیشنهادی از رویکرد اعتبار سنجی متقابل با سناریوی درون فردی استفاده شده است. معیارهایی مانند صحت و ماتریس درهم‌ریختگی برای ارزیابی مدل، به‌کار برده شده است. با بررسی عملکرد مدل بر روی هر وجه ویژگی، می‌توان دقت دسته‌بندی با میانگین بیش از ۵۰٪ را به‌دست آورد؛ اما در مدل گروهی نظارتی متوسط صحت تشخیص به بیش از ۹۰ درصد می‌رسد.

واژگان کلیدی: اطلاعات متقابل، تصویربرداری تشدید مغناطیسی عملکردی fMRI، کدگشایی مغز، هم‌جوشی تصمیم‌ها، یادگیری گروهی

## An Ensemble Multiview learning method for visual object decoding from fMRI brain data

Osama Hourani, Nasrollah Moghadam Charkari\* & Saeed Jalili

Department of Computer Engineering, Tarbiat Modares University, Tehran, Iran

### Abstract

In the past two decades, the applications of computational neuroscience have been increasingly growing. Breaking the neural code is a crucial open problem in computational neuroscience. Various research groups attempt to provide an efficient method to decode human brain activity using fMRI data. The output of these methods is a computational model that can assign brain signals to an external stimulus; in this study, visual object recognition has been investigated. The brain decoders are used in many applications, such as the brain-computer interface or detecting specific mental illnesses. In general, brain fMRI data have a high spatial and temporal resolution that increases the number of features of

\* Corresponding author

\* نویسنده عهده‌دار مکاتبات

the problem. Proper feature extraction from brain images is a challenging and time-consuming process. Consequently, the convergence of learning algorithms takes a long time to create an appropriate model. So, breaking down the feature space is highly recommended. We proposed new multi-view learning to solve the brain decoding problem. This approach splits the feature space based on mutual information and finds an appropriate ensemble classification model that detects the related visual object to neural activities in the brain.

The proposed method clusters the feature space based on mutual information and splits it into coherent sub-spaces, views. For each feature view, a support vector machine model is learned in parallel; the used SVM version can generate a vector of probabilities for each class. At the test phase, the feature space of test data is divided similarly to the training data, and each model generates a probabilistic vector for the test instances. Then, these vectors are combined in the decision profile matrix. The decision fusion is employed by the ordered weighted averaging (OWA) approach. The proposed multi-view learning methods achieved higher accuracy rates than the single view model. The main advantage of the MV model is that it can run in parallel, making it counterproductive to deal with the high-dimensional problems based on the divide and conquer strategy. The optimization phase to detect the most acceptable parameters for each model is obtained using the simulated annealing, SA, algorithm. We have employed three real fMRI datasets of the human brain to assess the proposed method, obtained from the Openneuro website. Also, the leave-one-run-out cross-validation approach has been carried out to evaluate the proposed method in the intra-subject scenario. Criteria such as accuracy rate and confusion matrix have been undertaken to analyze the results. The single feature view obtains an accuracy rate of more than 50%. While in the ensemble model, the accuracy rate in most subjects is more than 90%.

**Keywords:** Brain Decoding, Decision Fusion, Ensemble learning, fMRI, Mutual Information

تعمیم یافته است که بتواند فرایند شناختی مغز انسان را کدگشایی کند. برای مثال، پس از برداشتن دادگان مغزی متعلق به فرایند بینایی، همچون تماشای عکس‌های ثابت از اشیا در طبیعت، مدل محاسباتی مغزی ایجاد شده، امکان بازشناسی اشیا را فراهم می‌کند.

در این مقاله، یک روش جدید مبتنی بر یادگیری چندوجهی برای حل مسأله کدگشایی اشیا دیداری در مغز انسان پیشنهاد داده‌ایم. ابتدا اطلاعات متقابل به صورت تقریبی (بین ترکیب‌های ویژگی‌ها مقادیر وکسل‌ها) در دادگان fMRI از یک طرف و برچسب‌ها (دسته‌ها) از طرف به دست آورده و سپس ویژگی‌ها به تعداد مشخصی از خوشه‌ها بر اساس اطلاعات متقابل خوشه‌بندی می‌شوند. به صورتی که هر خوشه از ویژگی‌ها، یک وجه ویژگی مستقل را بیان می‌کند. در ادامه بر روی هر کدام از وجه‌های ویژگی یک دسته‌بند آموزش داده می‌شود؛ در نهایت مرحله دسته‌بندی و هم‌جوشی تصمیمات انجام می‌شود. برای آزمون الگوریتم پیشنهادی از دادگان fMRI واقعی مغز انسان استفاده شده است. به منظور مختصرسازی، روش پیشنهادی را به MMVL نام‌گذاری می‌کنیم که مخفف یادگیری چندوجهی مدرس<sup>۵</sup> است. با توجه به نتایج به دست آمده در این پژوهش، یادگیری چندوجهی نسبت به تک‌وجهی بهبود چشم‌گیری را

<sup>۵</sup> Modares multi-view learning

## ۱- مقدمه

کدگشایی الگوهای فعالیت در مغز انسان یک حوزه مهم در علوم اعصاب محاسباتی<sup>۱</sup> است. شناسایی اشیا دیداری یک نوع کدگشایی الگوهای فعالیت عصبی<sup>۲</sup> محسوب می‌شود. این الگوها هنگام رؤیت اشیا توسط یک فرد در مغز تولید و با کمک دستگاه‌های تصویربرداری مغز تبدیل به سیگنال‌های قابل پردازش می‌شوند. پژوهش‌ها نشان داده است که می‌توان فعالیت‌های مغزی متعلق به محرک‌ها<sup>۳</sup> یا اشیا دیداری را با کمک روش‌های یادگیری ماشین دسته‌بندی کرد [4]–[1].

این پژوهش‌ها مؤید آن است که زمانی که انسان به تصویر یک شیء می‌نگرد، الگوهای فعالیت مغزی او به شکلی وابسته به آن شیء خواهد بود [5]. در این راستا مدل‌های محاسباتی<sup>۴</sup> و هوش مصنوعی می‌توانند راه‌حلی مناسبی برای پیش‌بینی برچسب متعلق به الگوهای فعالیت مغزی باشد. این مدل‌ها در درک نحوه عملکرد ناحیه‌ی بینایی مغز انسان مفید خواهد بود [6] و [7]. طی دو دهه گذشته، پژوهش‌گران تمرکز خود را در ارائه روش‌های کارا برای مدل‌سازی مغز انسان با آنالیز دادگان مغزی داده‌اند. نتایج این روش‌ها یک مدل محاسباتی

<sup>1</sup> Computational Neuroscience

<sup>2</sup> Neural Activity

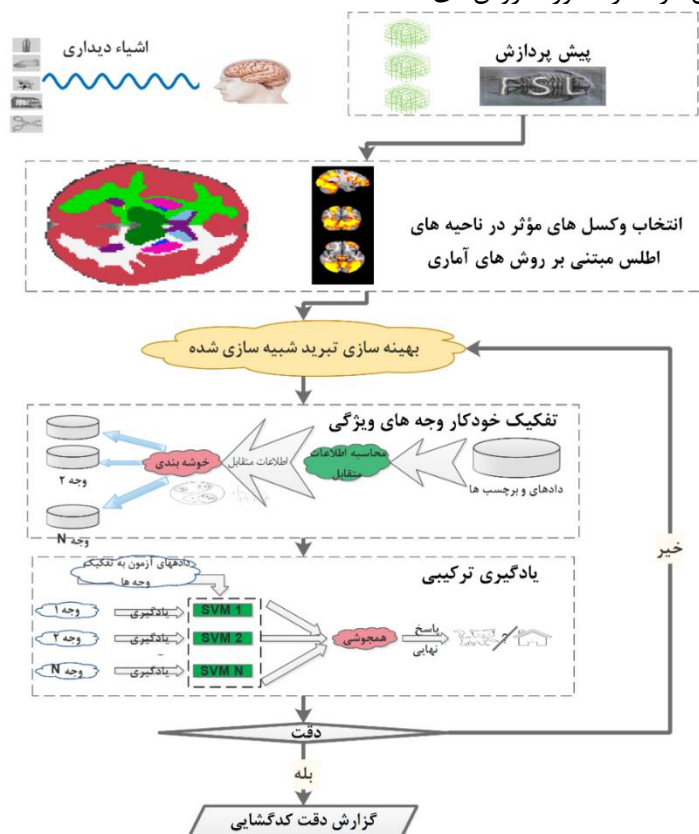
<sup>3</sup> Visual stimulus

<sup>4</sup> Computational Models

کدگذاری مغز انسان با استفاده از یادگیری گروهی بررسی می‌کنیم. قسمت سوم مراحل کلی روش پیشنهادی را مطرح می‌کند. داده‌های مورد استفاده و نتایج تجربی در بخش چهارم ارائه می‌شوند. مقایسه و ارزیابی روش پیشنهادی با آخرین پژوهش‌های کدگذاری مغز انسان در بخش پنجم ارائه می‌شوند؛ در نهایت جمع‌بندی و نتیجه‌گیری بیان می‌شود.

به‌وجود می‌آورد. همچنین در مقایسه با رویکردهای دیگر در کدگذاری مغز انسان، نتایج حاصله بیان‌گر نرخ دقت مناسب در سناریوی درون‌فردی است. هدف اصلی این سناریو ساخت مدل کدگذاری مغز انسان برای هر فرد به‌صورت جداگانه است. مراحل کلی روش پیشنهادی در شکل (۱) نشان داده شده است.

در ادامه این مقاله به پنج بخش تقسیم می‌شود. در بخش دوم، پیشینه پژوهش را در حوزه روش‌های



(شکل-۱): مراحل کلی روش پیشنهادی  
(Figure 1): General steps of the proposed method

دادگان fMRI از مغز فرد تصویربرداری می‌شود. هر فایل fMRI شامل چند آزمایش (به‌طورمعمول چند دقیقه است) از دادگان فعالیت مغزی (تکراری از بلوک‌های استراحت-فعالیت<sup>۳</sup>) خواهد بود که به اجرا<sup>۴</sup> موسوم است. جهت جلوگیری از نوفه و حفظ پایداری اطلاعات، از هر فرد، آزمایش به دفعات صورت می‌پذیرد. در این پژوهش، وظیفه موردبررسی شناسایی اشیای دیداری تعریف می‌شود، که جزو پارادایم‌های دیداری<sup>۵</sup> در نظر گرفته می‌شود.

## ۲- پیشینه پژوهش

مغز انسان یک دستگاه زیستی پیچیده و پویا است [8] که کدگذاری اطلاعات عصبی از آن، کاری دشوار است [9]. [10]. برای کدگذاری مغز پژوهش‌گران رویکردهایی مبتنی بر وظیفه<sup>۱</sup> طراحی می‌کنند. به‌عنوان مثال، عکس یک شیء دیداری به فرد طی مدت‌زمانی تعریف‌شده نشان داده می‌شود؛ سپس به مدت معین آن فرد استراحت می‌کند، برای مثال به صفحه خالی می‌نگرد. هر دفعه از اجرای این فرایند آزمایش<sup>۲</sup> نامیده می‌شود. هنگام اجرای آزمایش،

<sup>3</sup> Activity- Rest  
<sup>4</sup> Run  
<sup>5</sup> Visual Paradigm

<sup>1</sup> Task Based  
<sup>2</sup> Trail

تصویربرداری تشدید مغناطیسی عملکردی<sup>۱</sup> fMRI از افزایش مصرف اکسیژن در بافت مغزی بهره می‌گیرد که باعث افزایش میزان دی‌اکسید هموگلوبین در گلبول‌های قرمز خون و خواص مغناطیسی آن می‌شود، در نتیجه میزان روشنایی<sup>۲</sup> در تصاویر ناحیه‌های مغزی فعال از دیگر ناحیه‌ها بیشتر خواهد شد [11]؛ این سیگنال به نام BOLD<sup>۳</sup> معروف است که با کمک fMRI ردیابی می‌شود. در این روش تصاویری مکرر از ناحیه‌های مغزی در حال فعالیت و استراحت گرفته می‌شود. دادگان fMRI چهاربعدی هستند که عبارت‌اند از تصاویر سه‌بعدی که شامل سیگنال BOLD هستند. این تصاویر به‌مرور زمان برداشته می‌شوند. تصاویر fMRI از نوع غیرتجمعی هستند که هیچ آسیبی به افراد مورد آزمایش نمی‌رسانند. همچنین ابعاد بسیار بالایی دارند و به‌طور معمول نوفه فراوان در هنگام برداشت تصاویر به اطلاعات مفید تحمیل می‌شود [12].

طیف گسترده‌ای از روش‌های تجزیه و تحلیل الگوهای چندمتغیره<sup>۴</sup> در کدگشایی مغز استفاده می‌شود؛ به‌عنوان مثال روش ماشین بردار پشتیبان<sup>۵</sup> [13]–[15]، طبقه‌بندهای خطی تفکیکی<sup>۶</sup>، نزدیک‌ترین همسایه<sup>۷</sup> [14]، [17]، [16] همچنین جهت مقابله با مشکل ابعاد بالا در کدگشایی مغز انسان روش متنوع انتخاب ویژگی و تقلیل ابعاد ارائه شد، مانند فراهم‌ترازی<sup>۸</sup> [19]، [18]، الگوریتم‌های بهینه‌سازی [21]، [20]، روش‌های مبتنی بر گراف [22] و غیره، این روش‌ها با الگوریتم‌های یادگیری بانظارت مانند ماشین بردار پشتیبان اجرا شده‌اند. بیش‌تر مدل‌های کدگشایی مغز با استفاده از رویکرد یادگیری بانظارت و تک‌وجهی آموزش داده می‌شوند. همچنین، تعدادی کمی از مطالعات با استفاده از روش‌های یادگیری گروهی چندوجهی مسائل کدگشایی مغز را حل کرده‌اند. بیش‌تر مدل‌های کدگشایی مغز با استفاده از رویکرد یادگیری بانظارت و تک‌وجهی آموزش داده می‌شوند. همچنین، تعدادی کمی از مطالعات با استفاده از روش‌های یادگیری گروهی چندوجهی مسائل کدگشایی مغز را حل کرده‌اند [16]–[18].

ایده اصلی یادگیری گروهی این است که چند دسته‌بند پایه یاد گرفته می‌شوند، سپس هم‌جوشی از

آن‌ها به‌دست خواهد آمد. مدل گروهی، دسته‌بندی دقیق‌تری از هر یک از دسته‌بندهای پایه را به‌دست می‌آورد. دو ویژگی مهم در دسته‌بندهای پایه باید وجود داشته باشد: تنوع رفتار دسته‌بندهای پایه و عملکرد مناسب آن‌ها. منظور از عملکرد مناسب این است که هر دسته‌بند باید صحت بالاتر از دسته‌بندهای تصادفی داشته باشد؛ به‌عنوان مثال، میزان صحت در حالت رده‌های دودویی باید بیش از ۵۰٪ باشد [23]–[25]. همچنین منظور از تنوع رفتاری آن است که نتایج دسته‌بندها باید مکمل همدیگر و پاسخ‌های تکراری نداشته باشند. مدل گروهی مناسب، مصالحه عادلانه‌ای بین ویژگی‌های یادشده را فراهم می‌کند. رویکردهای مهم این مجموعه عبارت‌اند از: یادگیری تقویتی، بسته‌بندی<sup>۹</sup>، تقسیم تصادفی و جنگل تصادفی<sup>۱۰</sup>. یک رویکرد یادگیری گروهی شامل سه مرحله است: مرحله نخست، بخش‌بندی فضای یادگیری به چند زیر فضای داده که می‌تواند به یکی از اشکال نمونه‌برداری از داده‌ها یا انتخاب یک زیرمجموعه از ویژگی‌های باشد. مرحله دوم، یادگیری مدل‌های پایه‌ی دسته‌بندی و درنهایت، هم‌جوشی در سطح تصمیم است. تعداد روش‌های یادگیری گروهی چندوجهی برای کدگشایی اشیای دیداری در مقایسه با روش‌های یادگیری تک‌وجهی محدود هستند [26]–[28]. بیش‌تر پژوهش‌ها در این زمینه از روش‌های یادگیری سنتی و تک‌وجهی استفاده کرده‌اند. راه‌حل‌های زیادی در زمینه تفکیک وجه‌های ویژگی مانند CCA<sup>۱۱</sup> مطرح شده است [23]. پس از تفکیک یا بخش‌بندی فضای ویژگی از هر زیر فضای، یک دسته‌بند پایه یاد گرفته‌شده و هم‌جوشی دسته‌بندها در سطح تصمیم صورت می‌پذیرد. به این‌گونه روش‌ها، یادگیری گروهی<sup>۱۲</sup> اطلاق می‌شود [30]، [29]. از ویژگی‌های مهم در یادگیری گروهی این است که بر پایه فرضیات محکم بنا می‌شوند. ادغام نتایج در سطح تصمیم‌گیری در روش یادگیری گروهی، در مقابله با مشکل استفاده از دادگان با ابعاد بالا و پردازش موازی در شناسایی شیء دیداری یک‌راه حل مناسب است.

روش‌های یادگیری چندوجهی و گروهی به‌طور گسترده‌ای در کاربردهای مختلف مانند دسته‌بندی تصاویر طبیعی، متن، صدا، صفحات وب، احساسات صورت انسان و غیره مورد استفاده قرار می‌گیرد. از سوی دیگر، تعداد

<sup>9</sup> Bagging

<sup>10</sup> Random Forest

<sup>11</sup> CCA: Canonical Correlation Analysis

<sup>12</sup> Ensemble Learning

<sup>1</sup> fMRI: Functional Magnetic Resonance Imaging

<sup>2</sup> Intensity

<sup>3</sup> BOLD: Blood oxygen level dependent

<sup>4</sup> MVPA: Multivariate pattern analysis

<sup>5</sup> SVM: Support Vector Machine

<sup>6</sup> Discriminant Analysis

<sup>7</sup> KNN: k nearest neighbor

<sup>8</sup> Hyper-alignment

اما تفاوت اصلی در ترتیب عملیات پیش‌پردازش و همچنین در مراحل یادگیری چندوجهی است که در آن تفکیک خودکار فضای ویژگی صورت می‌پذیرد. روش محاسبه میزان اطلاعات متقابل و خوشه‌بندی فضای ویژگی بر اساس این معیار یکی از نکات متمایز این مقاله می‌باشد. در ادامه، این مراحل به تفصیل ارائه می‌شوند.

### ۱-۳- پیش‌پردازش

در مرحله پیش‌پردازش تصاویر fMRI که حاوی چهار بُعد است به عنوان ورودی دریافت می‌شوند. این ابعاد عبارت‌اند از نقاط با مختصات مکانی سه‌بعدی و زمان. این اطلاعات شامل حجم بالا و دارای نویز فراوان است. در این راستا، خروجی تصاویر fMRI نویز زدایی، نرمال‌سازی شده و در فضای استاندارد MNI152 ثبت می‌شود.

در این مقاله، مرحله نخست پیش‌پردازش تصاویر fMRI با استفاده از کتابخانه FSL انجام داده شده است [39]. FSL یک کتابخانه جامع از ابزارهای آنالیز دادگان تصویربرداری از مغز مشتمل بر سیگنال‌های fMRI، FMRI، MRI و DTI است. این کتابخانه بر روی رایانه‌های شخصی (هر دو محیط لینوکس<sup>۵</sup> و ویندوز<sup>۶</sup> از طریق یک ماشین مجازی) اجرا می‌شود و نصب آن بسیار آسان است. بسیاری از ابزارها را می‌توان از طریق خط فرمان و همچنین به عنوان رابط کاربری گرافیکی (واسطه‌های گرافیکی کاربر) اجرا کرد. برای اطلاعات بیشتر در مورد FSL می‌توان به <https://fsl.fmrib.ox.ac.uk/fsl/fslwiki> مراجعه کرد.

گام‌ها پیش‌پردازش بدین ترتیب هستند، نخست، با کمک Bet بافت مغزی استخراج، سپس، برای تمام تصاویر fMRI ثبت<sup>۷</sup> تصویر مغزی (تبدیل خطی) به فضای استاندارد MNI152 دو میلی‌متری با کمک FLIRT اعمال می‌شود. در ادامه، تجزیه و تحلیل سطح نخست<sup>۸</sup> استاندارد با استفاده از تنظیمات پیش‌فرض در FSL انجام می‌شود. خروجی مرحله پیش‌پردازش، تصاویر پاک‌سازی و نرمال‌سازی شده از دادگان fMRI خواهد بود. جهت انجام این مراحل، از رایانه با کارایی محاسباتی بالا<sup>۹</sup>، برای پردازش تصاویر fMRI و اجرای الگوریتم یادگیری استفاده کرده‌ایم. مشخصات دستگاه اختصاصی عبارت‌اند از: ۱۲۵ گیگابایت حافظه، پردازنده ۱۲ هسته Xeon و سیستم عامل اوبونتو<sup>۱۰</sup> ۱۶/۰۴. مدیریت پردازش‌ها به صورت موازی با استفاده از

محدودی از این روش‌ها برای مشکل کدگذاری مغز پیشنهاد شده است. کنشیوا و رودریگز رویکردهای کلاسیک از یادگیری گروهی برای دسته‌بندی اشیای دیداری با استفاده از مجموعه داده fMRI را بررسی و یک مقایسه جامع بین انواع مختلف روش‌های دسته‌بندی گروهی ارائه کردند. نتایج این یافته‌ها بیان‌گر موفقیت ماشین بردار پشتیبان، جنگل چرخشی و دسته‌بند فضای تصادفی، نسبت به سایر روش‌ها است [29]. یک روش گروهی با استفاده از فضای تصادفی در [31] مورد استفاده قرار گرفته است. انتخاب تصادفی از سیگنال‌های الکترودهای مغزنگاری الکتریکی<sup>۱</sup> برای ایجاد وجه‌های چندگانه مورد استفاده قرار گرفته است. تجزیه و تحلیل مؤلفه اصلی<sup>۲</sup> و تجزیه و تحلیل فیشر<sup>۳</sup>، در مرحله انتخاب ویژگی استفاده می‌شود. سپس، دسته‌بندی متعدد، آموزش داده می‌شود؛ در نهایت، تصمیم‌گیری با استفاده از یک الگوریتم رأی‌گیری ساده در مرحله آزمون صورت می‌پذیرد و میزان صحت به ۶۵٪ می‌رسد [31].

اطلاعات متقابل، MI، یک معیار برای تخمین وابستگی آماری بین دو پدیده عددی است [32]، [33]. این معیار در انتخاب ویژگی‌های مؤثر چه در حوزه‌های بیوانفورماتیک چه در حوزه کدگذاری مغز انسان کاربردهای فراوانی دارد [34]–[38]. اطلاعات متقابل بین دو متغیر تصادفی گسسته با استفاده از (رابطه ۱) بیان می‌شود:

$$MI(A, B) = \sum_{a \in A} \sum_{b \in B} p(a, b) \times \log \left( \frac{p(a, b)}{p(a) \cdot p(b)} \right) \quad (1)$$

A و B دو متغیر تصادفی هستند که هر کدام شامل مجموعه از مقادیر اعداد حقیقی یا گسسته و  $p(a, b)$  عبارت از تابع توزیع احتمال مشترک برای A و B است.  $p(a)$  و  $p(b)$  دو توزیع احتمالی حاشیه<sup>۴</sup> متعلق به این دو متغیر تصادفی هستند. نتایج مطالعات در مدل‌های گروهی در کدگذاری مغز بسیار امیدوارکننده است؛ بر این اساس در این پژوهش روشی کارا مبتنی بر تفکیک فضای ویژگی و اطلاعات متقابل به صورت خودکار ارائه می‌دهیم. جزئیات روش پیشنهادی در بخش بعدی ارائه می‌شود.

### ۳- روش پیشنهادی

اگرچه مراحل کلی روش پیشنهادی برای کدگذاری اشیای دیداری در مغز انسان شباهت زیادی با دیگر روش‌ها دارد،

<sup>5</sup> Linux

<sup>6</sup> Windows

<sup>7</sup> Registration

<sup>8</sup> Feat First level analysis

<sup>9</sup> HPC: High Processing Computing

<sup>10</sup> Ubuntu

<sup>1</sup> EEG: electroencephalography

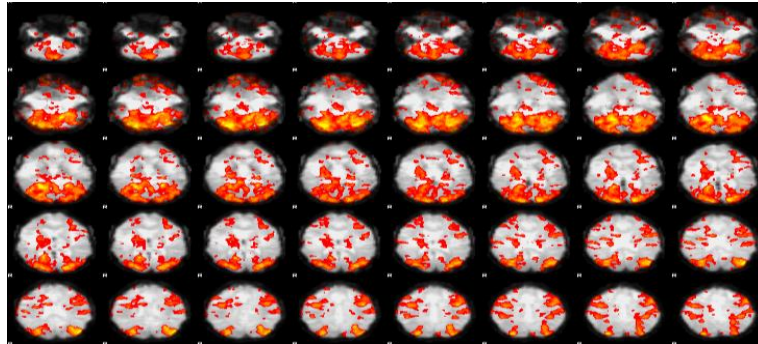
<sup>2</sup> PCA: Principal Component Analysis

<sup>3</sup> FDA: Fisher Discriminant Analysis

<sup>4</sup> Marginal Probability



مجموعه‌ای از صدها هزار وکسل مواجه هستیم. هر وکسل بیان‌گر فعالیت یک نقطه از فضای مغزی است که دارای مختصات سه‌بعدی است. انتخاب وکسل مؤثر چالش جدی در حوزه عصب‌شناسی محاسباتی است. هدف در این مرحله، کاهش ابعاد فضایی وکسل‌ها با انتخاب مؤثرترین آن‌ها به‌جای تمام وکسل‌های مغز است.



(شکل-۲): نتایج موقعیت‌یابی فعالیت مغزی متعلق به تماشای عکس خانه  
(Figure-2): Results of activity localization related to viewing a picture of a house

در الگوریتم انتخاب، از هر ناحیه در اطلس تعداد هشت وکسل در ۹ بازه زمانی استخراج خواهد شد. در این راستا ۷۲ ویژگی برای هر ناحیه به‌دست می‌آید. درواقع، با استفاده از اطلس‌های با تعداد محدود از ناحیه‌ها، مانند اطلس هاروارد-آکسفورد که شامل ۹۶ ناحیه آناتومی است، تعداد کمینه وکسل‌های انتخابی  $96 \times 72$  خواهد شد. نکته حائز اهمیت، انتخاب مصالحه مناسب مابین تعداد مناطق تعریف‌شده در اطلس مغز و میزان اطلاعات مؤثر سیگنال مغزی در هر ناحیه مغزی است. افزایش تعداد ناحیه‌ها سبب بهبود صحت و افزایش پیچیدگی پردازش خواهد بود. به‌طور مشابه با کاهش تعداد نواحی، پیچیدگی پردازش کاسته شده ولی صحت شناسایی پایین می‌آید. همچنین، پیچیدگی زمانی در پردازش مناطق دقیق‌تر، موضوع مهمی است که باید در نظر گرفت. در شکل (۳) جزئیات ساختار بردار ویژگی مورد استفاده نشان داده شده است. با توجه به تعداد ناحیه‌های آناتومی در اطلس مورد استفاده، تعداد وکسل‌های مؤثر در هر ناحیه و تعداد نقطه‌های زمانی در هر وکسل، اندازه بردار ویژگی بزرگ‌تر خواهد شد. در مرحله دوم پیش‌پردازش وکسل‌ها را به یک بردار استاندارد و مرتب‌شده، برای استفاده در مرحله یادگیری تبدیل می‌شوند. گام‌های این مرحله عبارت‌اند از: نرمال‌سازی، قطعه‌بندی<sup>۶</sup>، مسطح‌سازی<sup>۷</sup>، گسسته‌سازی<sup>۸</sup> و درنهایت مقیاس‌گذاری<sup>۹</sup> است که جزئیات بیشتر در مورد این گام‌ها در [38] توضیح داده شده است.

سرور پردازش گرید sun-grid-engine انجام‌شده است که بهره توان پردازشی به‌صورت تقریبی ۸۵٪ بود [40]. متوسط زمان پردازش یک فایل حدود ۲۷ دقیقه بود که هم‌زمان ده فایل به‌صورت موازی پردازش می‌شدند.

## ۲-۳- انتخاب ویژگی‌های مؤثر

انتخاب ویژگی از جمله مراحل حساس و مهم در حوزه یادگیری ماشین است. در حوزه پردازش fMRI با

گفتنی است که تعداد وکسل‌ها در دادگان مغز به‌صورت چشم‌گیر زیاد است؛ به‌ویژه هنگام استفاده از فضای استاندارد دو میلی‌متری تعداد وکسل‌ها به میلیون‌ها عدد می‌رسد؛ ازاین‌رو، ما روشی برای انتخاب وکسل‌های مؤثر به‌کار برده‌ایم که در آن به‌ازای یک محرک خارجی یک وکسل مناسب از هر ناحیه آناتومی در یک اطلس انتخاب می‌شود. مانند هاروارد آکسفورد<sup>۱</sup> [41] که شامل ۹۶ ناحیه از قشر مغز انسان<sup>۲</sup> است یا تالارایخ<sup>۳</sup> [42], [43] که دارای ۹۶۵ ناحیه آناتومی در قشر مغز انسان است. در آغاز وابستگی آماری بین محرک‌ها و سری زمانی وکسل‌ها به‌وسیله مدل رگرسیون خطی تعمیم‌یافته<sup>۴</sup> تعیین می‌شود. برای هر محرک، یک نقشه آماری از مقادیر z-score محاسبه می‌شود. برای هر منطقه آناتومی، وکسلی که دارای حداکثر z-score است انتخاب خواهد شد [44]. نمونه‌ای از نتایج موقعیت‌یابی فعالیت مغزی متعلق به تماشای تصویر خانه در شکل (۲) مشاهده می‌شود. فرایند انتخاب برای هر برجسب در مجموعه داده تکرار می‌شود. به‌این‌ترتیب، تعداد ویژگی‌ها با توجه به نوع محرک‌ها و منطقه‌های آناتومی متغیر است. برای مثال، در مجموعه داده DS105، هشت دسته از اشیا وجود دارد که در آن بازه زمانی نشان دادن تصویر و زمان تکرار<sup>۵</sup>، به‌ترتیب ۲۴ و ۲/۵ ثانیه است.

<sup>1</sup> Harvard-Oxford Cortical Atlas

<sup>2</sup> Human Brain Cortex

<sup>3</sup> Talairach Atlas

<sup>4</sup> GLM: General Linear Model

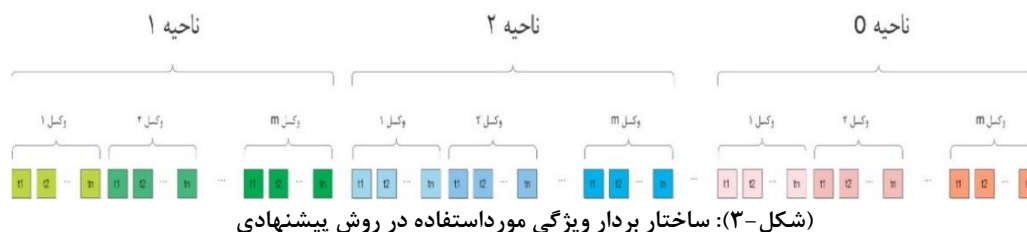
<sup>5</sup> TR: Repetition Time

<sup>6</sup> Segmentation

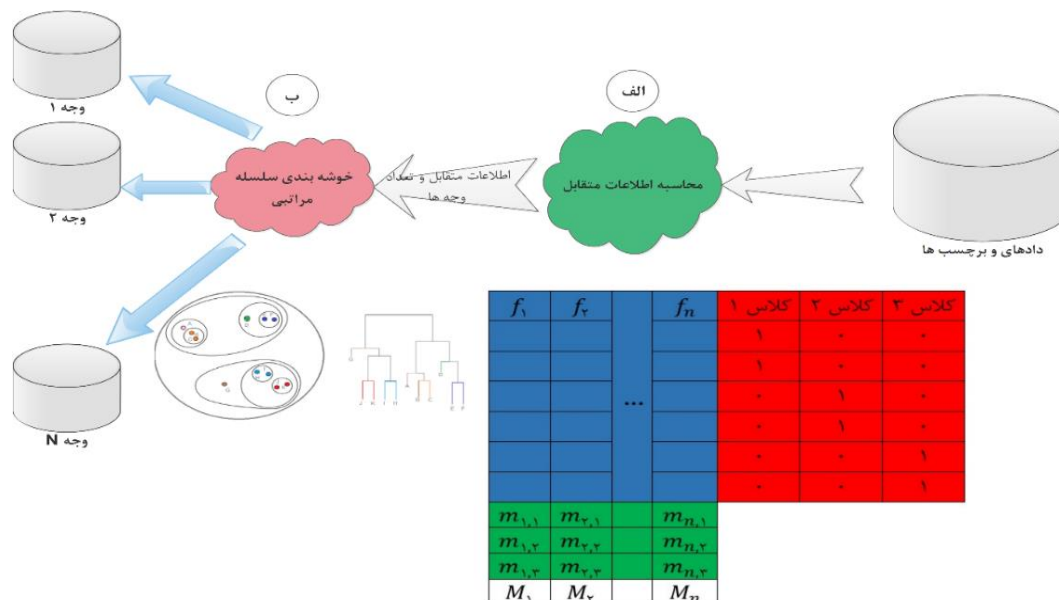
<sup>7</sup> Flattenning

<sup>8</sup> Discretization

<sup>9</sup> Scaling



(Figure-3): The structure of the used feature vector in the proposed method



(شکل-۴): مراحل کلی الگوریتم تفکیک خودکار فضای ویژگی (Figure-4): General steps of feature space automatic decomposition algorithm

در این پژوهش، پارامترهای پنالتی و گاما را برای همه دسته‌بندی‌ها ماشین بردار پشتیبان ثابت در نظر گرفته‌ایم. همچنین هدف الگوریتم، پیدا کردن تعداد مناسبی از وجه‌های ویژگی و عملکرد هم‌جوشی است. میزان خطای مدل دسته‌بندی ترکیبی به‌عنوان تابع برازندگی عمل می‌کند. SA با جستجو در فضای تعداد وجه‌ها و پارامترهای پیش‌فرض در الگوریتم مانند سرعت تبرید<sup>۵</sup>، درجه حرارت نمایی<sup>۶</sup> و تابع پذیرش<sup>۷</sup>، این خطا را کمینه می‌سازد.

### ۳-۴- بخش‌بندی خودکار فضای ویژگی

بخش‌بندی خودکار فضای ویژگی در یادگیری چندوجهی یک گام حیاتی است، به‌ویژه زمانی که تفکیک طبیعی بین وجه‌های ویژگی وجود ندارد. نوآوری اصلی این پژوهش یک روش مبتنی بر خوشه‌بندی اطلاعات متقابل برای بخش‌بندی وجه‌های ویژگی است.

### ۳-۳- بهینه‌سازی

الگوریتم‌های یادگیری به‌صورت عام وابسته به پارامتر هستند، همچون ماشین بردار پشتیبان<sup>۱</sup> که اگر پارامترهای آن خوب تنظیم نشوند (مانند پنالتی<sup>۲</sup> و گاما<sup>۳</sup>)، کارایی مناسبی نشان نخواهد داد. روش پیشنهادی در این مقاله از نوع گروهی است که چندین دسته‌بند ماشین بردار پشتیبان را شامل می‌شود. از طرف دیگر در روش‌های گروهی، مواردی همچون تعداد وجه‌های ویژگی و عملکرد هم‌جوشی از اهمیت بالایی برخوردار هستند. از طرف دیگر، با توجه به اهمیت کاهش فضای جستجو از الگوریتم‌های بهینه‌سازی استفاده می‌کنیم؛ از این‌رو، الگوریتم SA تبرید شبیه‌سازی‌شده<sup>۴</sup> در مرحله بهینه‌سازی پارامترهای الگوریتم پیشنهادی را انتخاب کرده‌ایم. امروزه از SA به‌طور گسترده در حوزه‌های مختلفی در بهینه‌سازی و جستجوی حل مسائل استفاده می‌شود. این الگوریتم را جستجوی سراسری فرامکاشفه‌ای نیز می‌نامند.

- <sup>1</sup> SVM Support vector Machine
- <sup>2</sup> Penalty
- <sup>3</sup> Gamma
- <sup>4</sup> SA: Simulated Annealing

- <sup>5</sup> Annealing Speed
- <sup>6</sup> Exponential Function
- <sup>7</sup> Acceptance

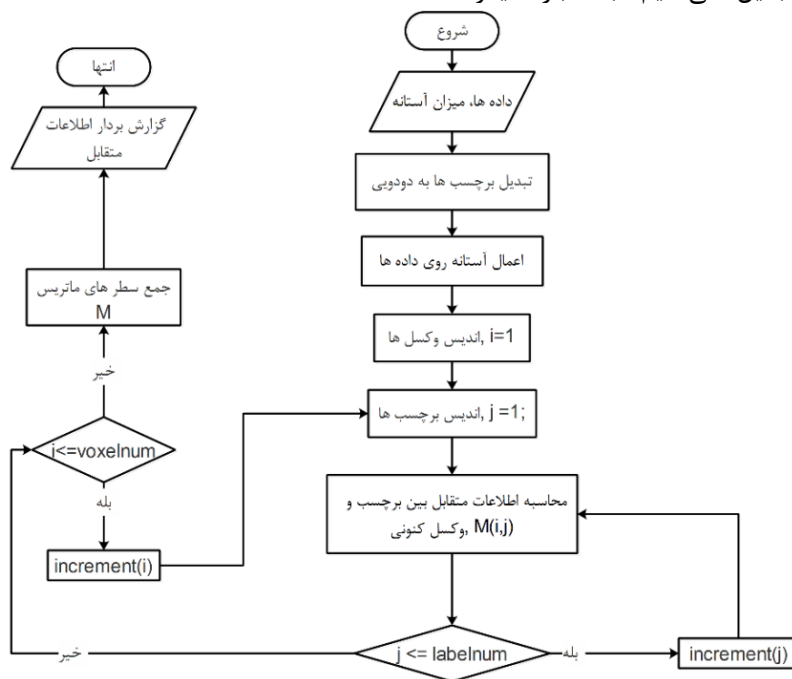
شبه تبدیل برچسب‌ها به یکی-برابر-همه one-versus-all است. یک نمونه از تبدیل برچسب‌های چندرده‌ای به دودویی در جدول (۱) نمایش داده شده است.

(جدول ۱): روند دودویی‌شدن برچسب‌ها

(Table-1): Binarizing process of labels

| برچسب‌ها دودویی | برچسب اصلی |
|-----------------|------------|
| ۰ ۰ ۰ ۱         | ۱          |
| ۰ ۱ ۰ ۰         | ۲          |
| ۰ ۰ ۱ ۰         | ۳          |
| ۰ ۰ ۰ ۱         | ۴          |

در شکل (۴) مراحل کلی الگوریتم بخش‌بندی خودکار فضای ویژگی ارائه شده است. این مراحل عبارت‌اند از محاسبه اطلاعات متقابل و خوشه‌بندی سلسله‌مراتبی. ورودی الگوریتم بخش‌بندی فضای ویژگی عبارت است از: میزان آستانه مقادیر وکسل‌ها؛ تعداد وجه‌های ویژگی؛ داده‌های fMRI مغزی، که شامل مقادیر وکسل‌ها و برچسب‌ها است. ابتدا آستانه روی مقادیر ویژگی‌های ورودی اعمال می‌شود تا این مقادیر به صفر و یک تبدیل شوند؛ سپس مقادیر برچسب‌ها به صورت دودویی کدگذاری می‌شوند. به این صورت که برچسب‌های چندرده‌ای به برچسب‌های دودویی تبدیل می‌کنیم. به عبارت دیگر



(شکل ۵): فلوچارت الگوریتم محاسبه اطلاعات متقابل

(Figure-5): Flowchart of mutual information algorithm

ویژگی‌ها) محاسبه می‌شود (رابطه ۱) تا میزان وابستگی مابین هر جفت از بردارها به دست آید و از آن جهت بخش‌بندی وجه‌های ویژگی استفاده خواهد شد. در ادامه به ازای هر ویژگی که به صورت دودویی است، نرخ اطلاعات متقابل در ترکیب‌های دسته‌ویژگی<sup>۱</sup> حساب می‌شوند. برای مثال در مجموعه داده DS105 برای هر ویژگی هشت مقدار اطلاعات متقابل محاسبه می‌شود. این هشت مقدار با هم جمع و به عنوان اطلاعات متقابل نهایی ویژگی‌ها گزارش داده می‌شوند. نمودار جعبه‌ای الگوریتم محاسبه اطلاعات متقابل در شکل (۵) نشان داده شده است. پس از محاسبه مقادیر اطلاعات متقابل، فاز خوشه‌بندی سلسله‌مراتبی

تعداد سطرهای ماتریس اطلاعات متقابل برابر با تعداد برچسب‌های دودویی و تعداد ستون‌های آن متناظر با ویژگی‌ها است (قسمت الف (شکل). تفاوت اصلی الگوریتم پیشنهادی با دیگر روش‌های مشابه آن است که از یک طرف مقدار آستانه (به عنوان پارامتر ورودی) روی ویژگی‌ها اعمال و از طرف دیگر برچسب‌ها کدگذاری می‌شوند تا به صورت دودویی تبدیل شوند. این فرایند به نوعی تقریب محسوب می‌شود که در سرعت محاسبات و کاهش نوفه احتمالی و نهفته در مقدارهای اطلاعات متقابل نقش مهمی بازی می‌کند. اطلاعات متقابل به عنوان معیار مهم جهت کاهش عدم قطعیت در استخراج خودکار فضای ویژگی استفاده می‌شود. اطلاعات متقابل در الگوریتم پیشنهادی بین هر دو بردار دودویی (برچسب‌ها و

<sup>1</sup> Feature Label Combinations

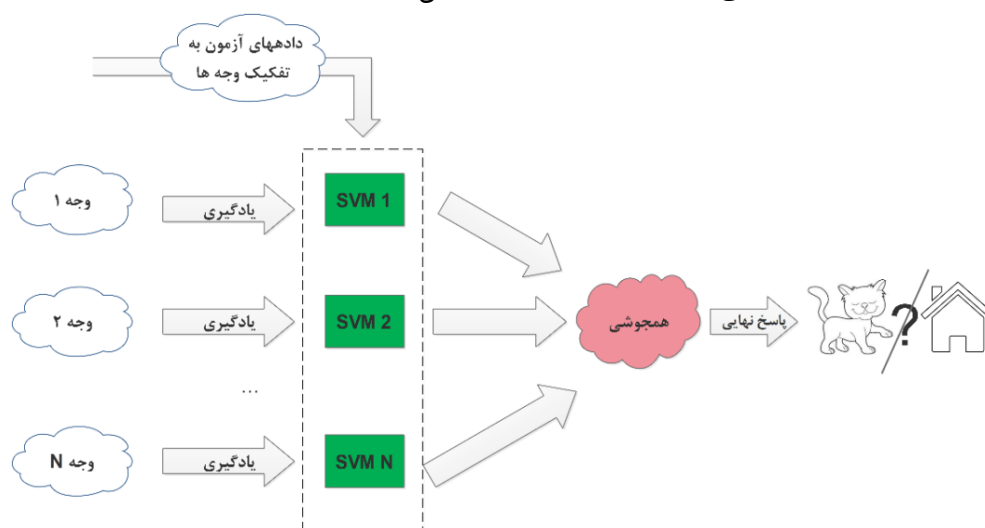


شروع می‌شود؛ بر این اساس تعداد خوشه‌ها و بردار حاوی مقادیر اطلاعات متقابل به‌عنوان ورودی تعریف می‌شوند. خروجی خوشه‌بندی به‌صورت دسته‌هایی از ویژگی‌ها هستند که به هر وجه متعلق خواهد بود (قسمت ب شکل ۴) این کار باعث می‌شود ویژگی‌هایی که مشابه همدیگر هستند، بر اساس اطلاعات متقابل، در یک وجه جداگانه قرار بگیرند و همبستگی بین وجه‌های ویژگی کمتر می‌شود.

### ۵-۳- یادگیری

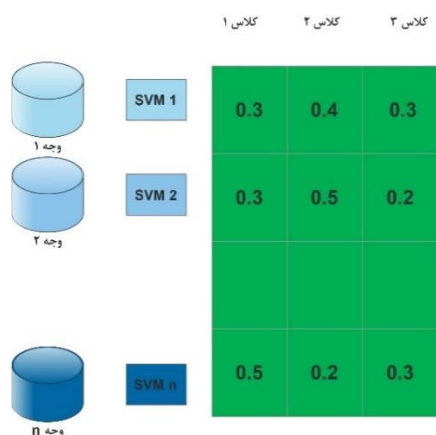
در رویکرد پیشنهادی، چند مدل یادگیری معروف آزمایش کردیم با توجه به آزمایش‌های تجربی، دسته‌بند ماشین

بردار پشتیبان را انتخاب شد که روی هر وجه ویژگی جداگانه آموزش داده می‌شوند، به‌عبارت دیگر تخصص هر دسته‌بند بستگی به وجه ویژگی مربوطه خواهد داشت، به‌عنوان مثال یک دسته‌بند در شناسایی تصاویر صورت انسان دقت بیشتری نسبت به دیگر دسته‌بندها دارد. یکی از ویژگی‌های اصلی مدل مورد استفاده اینکه خروجی آن بایستی به‌صورت یک بردار احتمالاتی باشد. تعداد عناصر این بردار مطابق با تعداد پرچسب‌ها است و هر مقدار از این بردار شامل میزان تعلق نمونه آزمایشی به رده مقابل است. این بردار در مرحله هم‌جوشی نتایج به‌کار برده می‌شود. شمای کلی یادگیری چندی وجهی در شکل (۶) نشان داده شده است.



(شکل-۶): شمای فرایند یادگیری گروهی چندوجهی  
(Figure-6): Scheme of the multi-view ensemble learning process

(۷) این ماتریس را تحت نام پروفایل تصمیم‌گیری<sup>۲</sup> نشان می‌دهد.



(شکل-۷): نمونه‌ای از ماتریس پروفایل تصمیم‌گیری  
(Figure-7): example of decision profile matrix

### ۳-۶- هم‌جوشی

مزیت اصلی در هم‌جوشی دسته‌بندها دستیابی به دقت بالاتر نسبت به حالت تک‌وجهی از طریق اهمیت‌دادن به جواب‌های درست در دسته‌بندهای گوناگون<sup>۱</sup> است. حساب‌کردن وزن تعلق نمونه آزمایشی، به هرکدام از رده‌های مسأله (هشت شیء دیداری) در این مرحله صورت می‌پذیرد. در هر وجه ویژگی، یک دسته‌بند به‌صورت جداگانه یاد گرفته می‌شود؛ سپس برای تصمیم‌گیری در مرحله بعدی از ماتریس پروفایل تصمیم‌گیری استفاده می‌کنیم. اگر تعدادی از دسته‌بندها داشته باشیم آنگاه خروجی هر یک به‌صورت بردار میزان تعلق برای آن دسته‌بند تعریف می‌شود. درنهایت خروجی این دسته‌بندها به‌صورت یک ماتریس خواهد بود. شکل

<sup>2</sup> DP: Decision Profile Matrix

<sup>1</sup> Diverse Classifiers

در روش میانگین، وزن‌های مرتب‌شده فقط به ستون‌های ماتریس که رده را نشان می‌دهند، توجه می‌شود. به عنوان مثال، سطرها متعلق به نوع دسته‌بند و ستون‌های متعلق به رده هستند که مجموع وزن‌های هر سطر باید یک باشد. هر وزن میزان تعلق نمونه به رده (ستون) را نشان می‌دهد. روش‌های خطی مختلف برای هم‌جوشی نتایج با کمک ماتریس پروفایل تصمیم‌گیری به کار می‌رود. در کار حاضر از ترکیب این روش‌ها در مرحله هم‌جوشی استفاده شده است. جدول (۲) مجموعه‌ای از روش‌های به کار گرفته‌شده در مرحله هم‌جوشی در این پژوهش را نشان می‌دهد.

(جدول ۲): جزئیات عملگرهای هم‌جوشی تصمیم‌های

مورد استفاده در این پژوهش

(Table-2): Details of used decision fusion operators

| نام روش ترکیب | نام انگلیسی | فرمول                                   |
|---------------|-------------|-----------------------------------------|
| بیشینه        | Maximum     | $Y_i = \max_j d_{ji}$                   |
| کمینه         | Minimum     | $Y_i = \min_j d_{ji}$                   |
| میانگین       | mean        | $Y_i = \frac{1}{L} \sum_{j=1}^L d_{ji}$ |
| میانه         | Median      | $Y_i = \text{median}_j d_{ji}$          |
| ضرب           | Products    | $Y_i = \prod d_{ji}$                    |

## ۴- نتایج

در این قسمت نتایج دسته‌بندی متعلق به مدل پیشنهادی چندوجهی MMVL ارائه می‌شود که با استفاده از دو اطلس با تعداد ناحیه متفاوت (کم ۹۶ ناحیه هاروارد اکسفورد و زیاد ۹۶۵ ناحیه تالارایخ) اجرا شده. ابتدا مجموعه‌های داده مورد استفاده معرفی می‌شوند؛ در نهایت صحت تشخیص برای مجموعه‌های داده DS105، DS107 و DS116 مورد بررسی قرار می‌گیرند.

### ۴-۱- معرفی مجموعه‌های داده

سه مجموعه داده تصاویر fMRI برداشته شده از سوژه‌های انسانی، پردازش شده‌اند و در اعتبارسنجی روش پیشنهادی مورد بررسی قرار گرفته‌اند. این دادگان از تارنمای openfMRI بارگیری شده‌اند. مجموعه نخست در مطالعات شناسایی اشیای دیداری در مغز انسان استفاده می‌شود [45]. مجموعه دوم برای موقعیت‌یابی برخی

فعالیت مغزی متعلق به تفکیک بین اشیای دیداری و نوشته‌های متنی است [46]. مجموعه سوم شامل دادگان وظیفه‌های دیداری و شنیداری odd-ball است که فقط فعالیت دیداری مورد استفاده قرار گرفت [47]. گفتنی است که وجه مشترک هر سه مجموعه داده در فعالیت متعلق به شناسایی اشیای دیداری است. مجموعه داده DS105 یا مجموعه داده شناسایی اشیای دیداری شامل شش نفر، پنج زن و یک مرد است. برای هر فرد در این مجموعه داده دوازده اجرا وجود دارد که در دوازده فایل تقسیم شدند. هر اجرا شامل ۱۲۰ حجم (۱۲۰ نقطه زمانی<sup>۱</sup>). به ازای هر ۲/۵ ثانیه یک حجم<sup>۲</sup> نشان داده می‌شود، آن‌گاه مدت زمان هر اجر سیصد ثانیه است که بر اساس رابطه (۲) محاسبه می‌شود:

$$\text{overall}_{runtime} = vn \times TR \quad (2)$$

پارمترها عبارت‌اند از  $TR^3$  زمان تکرار و  $vn^4$  تعداد حجم‌ها یا عکس‌های سه‌بعدی از سیگنال BOLD. عکس‌های اشیای (محرک‌ها) مورد استفاده در وظیفه مجموعه داده DS105 عبارت‌اند از خانه؛ عکس درهم شده؛ گربه؛ کفش؛ بطری؛ قیچی؛ صندلی و صورت انسان است. مجموعه داده DS107 شامل تصاویر fMRI متعلق به (۲۳ مرد و ۲۲ زن) آشنا به یک زبان، در مطالعه [46] معرفی می‌شود. افراد دارای میانگین سن ۲۵ سال و بین ۱۹ تا ۳۸ سال هستند، هیچ کدام مشکل بینایی ندارد. برچسب‌های این مجموعه داده چهار گروه هستند. تصاویر کلمات انگلیسی؛ اشیای دیداری؛ اشیای درهم شده و رشته حرفی نامفهوم. مجموعه داده DS116 شامل هفده فرد است (شش زن و ۱۱ مرد، میانگین سن ۲۷.۷ سال و دامنه سنی ۲۰ تا ۴۰ سال) سناریوی بینایی و شنوایی مشابه برای شرکت‌کنندگان در سه اجرا انجام شده است. محرک‌ها عبارت‌اند از دایره بزرگ قرمز رنگ و دایره کوچک سبز رنگ که هر دو با پس‌زمینه خاکستری هستند. در این مقاله تنها سناریوی دیداری مورد استفاده قرار گرفته است که بقیه روش‌های مورد مقایسه در جدول (۵) از سناریوهای دیداری و شنوایی استفاده کردند. مقایسه مختصر مجموعه‌های دادگان پردازش شده در این مقاله در جدول (۳) نشان داده شده است.

<sup>1</sup> Time point

<sup>2</sup> Volume

<sup>3</sup> TR: time to repetition

<sup>4</sup> Vn: volume number

(جدول-۳): مجموعه داده‌های fMRI مورد بررسی در این مقاله

(Table-3) The details of the studied fMRI datasets in this paper

| نام مجموعه داده | تعداد افراد | تعداد اجراها | نوع وظیفه                    | توضیح فرایند مورد بررسی                                                                                   | روش تصویربرداری |
|-----------------|-------------|--------------|------------------------------|-----------------------------------------------------------------------------------------------------------|-----------------|
| DS105           | ۶           | ۱۲           | یک-عقب <sup>۱</sup>          | هشت نوع از تصاویر اشیاء: خانه، تصویر به هم خورده <sup>۲</sup> ، گربه، کفش، بطری، قیچی، صندلی و صورت انسان | fMRI, MRI       |
| DS107           | ۴۹          | ۲            | یک-عقب                       | چهار گروه از اشیاء و کلمات: کلمات نوشته شده، اشیاء، اشیاء درهم، رشته‌های از حروف غیر صدadar               | fMRI, MRI       |
| DS116           | ۱۷          | ۲            | مبتنی بر رویداد <sup>۳</sup> | فرایند oddball دیداری شنیداری                                                                             | Fmri, MRI       |

## ۲-۴- معیار ارزیابی

در زمینه یادگیری ماشین، در حل مسأله حاضر معیار صحت از (رابطه ۳) به دست می‌آید.

$$accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad j = 1 \dots n \quad (3)$$

برای ارزیابی عملکرد یک دسته‌بند چندرده برای یکی از برچسب‌ها چون label-i معیارهای گفته شده به صورت زیر بیان می‌شود:

مثبت صحیح  $TP_i = C(i, i)$

منفی صحیح  $TN_i = \sum_{j \neq i} C(j, j)$

مثبت کاذب  $FP_i = \sum_{i \neq j} C(i, j)$

منفی کاذب  $FN_i = \sum_{i \neq j} C(j, i)$

به طوری که مثبت صحیح<sup>۴</sup> نمونه از کلاس مثبت که به درستی مثبت تشخیص داده می‌شود. مثبت کاذب<sup>۵</sup>: نمونه از رده منفی که به اشتباه مثبت تشخیص داده می‌شود. منفی صحیح<sup>۶</sup>: نمونه از رده منفی که به درستی منفی تشخیص داده شود. منفی کاذب<sup>۷</sup>: نمونه از رده مثبت که به اشتباه منفی تشخیص داده شود.

برای ارزیابی رویکرد پیشنهادی از ماتریس درهم‌ریختگی استفاده می‌کنیم. اگر ماتریس درهم‌ریختگی در حالت چندرده با حرف C نشان بدهیم، به طوری که هر سطر متعلق به تعداد برچسب‌های پیش‌بینی شده است و در مقابل ستون‌ها متعلق به جواب‌های درست هستند. چند معیار به تفکیک برچسب‌ها می‌توان بیان کرد. در این مقاله، معیارهای صحت حساسیت و تشخیص مورد استفاده قرار گرفته‌اند.

میزان صحت متعلق به رده i برابر است با (رابطه ۴)، و معیار حساسیت متعلق به رده i برابر است با (رابطه ۵)، معیار ویژگی متعلق به رده i با (رابطه ۶) بیان می‌شوند.

$$Accuracy_i = \frac{TP_i}{TP_i + TN_i + FP_i + FN_i} \quad (4)$$

$$Sensitivity_i = \frac{TP_i}{TP_i + FN_i} \quad (5)$$

$$Specificity_i = \frac{TN_i}{TN_i + FP_i} \quad (6)$$

در نهایت صحت تجمیعی به صورت زیر تعریف می‌شود:

$$Overall ACC = \frac{\sum_i^n TP_i}{\sum_i^n TP_i + \sum_i^n TN_i + \sum_i^n FP_i + \sum_i^n FN_i} \quad (7)$$

سناریوی اعتبارسنجی متقابل<sup>۸</sup> یک اجرا بیرون<sup>۹</sup> به کار برده شده است که با رویکرد یک نمونه بیرون<sup>۱۰</sup> LOOCV شباهت دارد؛ اما تنها تفاوت این است که نمونه‌های آزمون باید متعلق به یک اجرا باشد و نمونه‌های اجراها در زمان تولید داده‌های آزمون به صورت تصادفی انتخاب نمی‌شوند. ضمن آنکه امکان پخش شدن نمونه‌های یک اجرا در داده‌های آزمون و یادگیری وجود ندارد. به این صورت که داده‌های فرد را بر اساس اجراها به زیرمجموعه<sup>۱۱</sup> داده تجزیه می‌کند، در هر تکرار یکی از زیرمجموعه‌ها به عنوان داده‌های آزمون استفاده می‌شود و بقیه به عنوان آموزش مورد استفاده قرار می‌گیرند. آموزش و آزمون مدل‌ها تکرار می‌شود تا زمانی که همه زیرمجموعه‌ها در آزمایش یکبار استفاده شوند.

## ۳-۴- نتایج به تفکیک مجموعه‌های داده

در این بخش نتایج صحت برای سناریوی درون فردی<sup>۱۲</sup> گزارش شده است. از آنجا که سناریو درون فردی مورد

<sup>8</sup> Cross-validation

<sup>9</sup> Leave-one-run out

<sup>10</sup> LOOCV: Leave-one-out cross validation

<sup>11</sup> fold

<sup>12</sup> Intra-subject

<sup>1</sup> One-back Task

<sup>2</sup> Scrambled

<sup>3</sup> Event Based

<sup>4</sup> TP: true positive

<sup>5</sup> FP: false positive

<sup>6</sup> TN: true negative

<sup>7</sup> FN: false negative

بررسی قرار می‌گیرد، مدل کدگذاری پیشنهادی بر اساس افراد به صورت جداگانه آموزش داده و آزمون می‌شود. در این راستا، برای هر فرد پارامترهای جداگانه بهینه با استفاده از الگوریتم SA به دست خواهد آمد. در مرحله بهینه‌سازی یک بازه برای جستجوی مقادیر هر پارامتر در نظر گرفته شد. پارامتر آستانه در بازه  $0/8-0/3$  و تعداد وجه‌ها در محدوده  $3-50$  جستجو می‌شوند.

### ۱-۳-۴- نتایج مجموعه داده DS105

از آنجایی که مجموعه داده DS105 برچسب‌گذاری شده است. پس یادگیری گروهی از نوع بانظارت است. میزان صحت برای همه افراد در جدول (۴) نمایش داده شده است. مرحله هم‌جوشی با استفاده از روش میانگین وزن‌دار مرتب‌شده<sup>۱</sup> انجام شده است. عمل‌گر کمینه و میانگین در مرحله هم‌جوشی، بهترین نتایج را داشته‌اند. بر اساس معیار صحت نشان داده شده است که الگوریتم یادگیری چندوجهی در حالت درون فردی عملکرد بهتری نسبت به تک‌وجهی دارد. تفاوت چشم‌گیری بین نتایج الگوریتم یادگیری تک‌وجهی و چندوجهی و روش [48] در ((جدول دیده می‌شود. گفتنی است که توزیع آماری نتایج نرمال نیست. در الگوریتم MMVL-965 صحت تشخیص برای سه فرد آخر به یکصد درصد رسید.

(جدول ۴): مقایسه میزان صحت روش‌های یادگیری

چندوجهی با یادگیری تک‌وجهی مجموعه داده DS105

(Table-4): comparison of accuracy rate in multiview and single view learning methods in dataset DS105

| # | SVM<br>تک‌وجهی | روش [48] | MMVL-96 | MMVL-965 |
|---|----------------|----------|---------|----------|
| ۱ | ۰۶±۹۵          | ۰۸±۹۴    | ۰۵±۹۸   | ۰۸±۹۵    |
| ۲ | ۰۸±۸۹          | ۱۳±۸۱    | ۰۶±۹۶   | ۱۰±۹۳    |
| ۳ | ۰۶±۹۳          | ۰۸±۹۴    | ۰۶±۹۶   | ۰۵±۹۷    |
| ۴ | ۰۸±۹۲          | ۱۱±۸۴    | ۰۵±۹۸   | ۱۰۰      |
| ۵ | ۰۹±۹۴          | ۷۷±۳۴    | ۰۴±۹۹   | ۱۰۰      |
| ۶ | ۰۸±۹۶          | ۰۸±۹۲    | ۱۰۰     | ۱۰۰      |

ماتریس درهم‌ریختگی<sup>۳</sup> برای آزمون مدل چندوجهی در شکل (۸) نشان داده شده است. برچسب‌ها در خط افقی هستند و ستون پاسخ‌های مدل پیشنهادی را به صورت تجمیعی<sup>۴</sup> نشان می‌دهد. هرچقدر مقدار در

<sup>۱</sup> OWA: Ordered Weighted Averaging

<sup>۲</sup> شماره فرد

<sup>۳</sup> Confusion matrix

<sup>۴</sup> Accumulative

قطر اصلی ماتریس درهم‌ریختگی بزرگ‌تر باشند؛ یعنی دسته‌بندی از صحت بالاتری برخوردار است. ستون اضافه در سمت راست ماتریس متعلق به میزان حساسیت<sup>۵</sup> است. همچنین سطر اضافه متعلق به معیار تشخیص<sup>۶</sup> است. با توجه به شکل (۸) بیشترین صحت متعلق به سه رده خانه، گربه و صندلی بوده و کمترین عملکرد متعلق به قیچی است. صحت تشخیص مدل پیشنهادی به صورت تجمیعی نزدیک به ۹۸ درصد است. این برتری الگوریتم یادگیری گروهی را در مقابل یادگیری تک‌وجهی نشان می‌دهد.

### ۲-۳-۴- نتایج مجموعه داده DS107

نتایج دسته‌بندی متعلق به مجموعه داده DS107 در شکل (۹) به صورت نمودار جعبه‌ای<sup>۷</sup> نمایش داده شده است. محور افقی شماره‌سریال افراد و محور عمودی میزان صحت تشخیص را نشان می‌دهد. دایره‌ها در وسط جعبه‌ها در نمودار، میانه<sup>۸</sup> صحت تشخیص برای هر فرد را ارائه می‌دهد. میانه در ۱۸ فرد در بازه بالاتر از ۹۸ درصد قرار دارد و تنها پنج فرد میانه صحت تشخیص بین ۹۰ و ۸۸ درصد را دارند. بقیه بالاتر از ۹۲ هستند. به دلیل تعداد زیاد افراد و بیان بهتر نتایج، نمودار جعبه‌ای به کار برده شده است. برخلاف اینکه برخی جعبه‌های طولانی که انحراف معیار زیاد در صحت تشخیص الگوریتم یادگیری را نشان می‌دهند، همچنان صحت بیشتر افراد بالاتر از ۹۲ درصد است. ماتریس درهم‌ریختگی تجمیعی متعلق به آزمون مدل پیشنهادی چندوجهی برای مجموعه داده DS107 در شکل (۱۰) ملاحظه می‌شود. بیشترین خطاهای تشخیص متعلق به اشیای دیداری و کمترین خطا متعلق به رده حروف نامفهوم بوده است. شباهت زیاد بین هر دو برچسب و تغییرات در الگوهای فعالیت مغز مرتبط می‌تواند سبب کاهش صحت تشخیص مدل گروهی تفسیر کند.

### ۳-۳-۴- نتایج مجموعه داده DS116

ماتریس درهم‌ریختگی تجمیعی وابسته به روش پیشنهادی برای مجموعه داده DS116 در شکل (۱۱) نمایش شده است. صحت تجمیعی برای همه افراد ۹۷/۷

<sup>۵</sup> Sensitivity

<sup>۶</sup> Specificity

<sup>۷</sup> Box plot

<sup>۸</sup> median

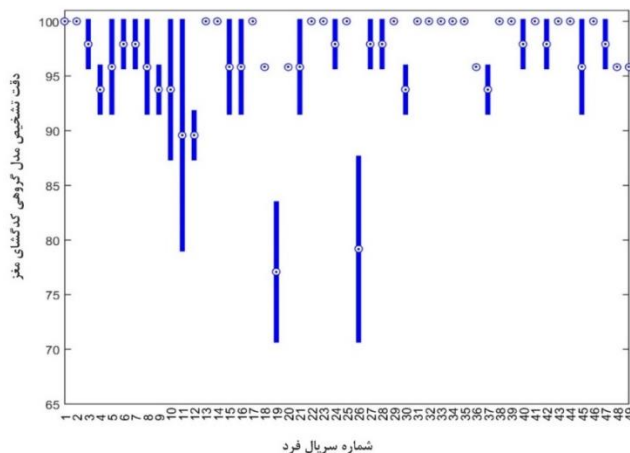
شناسایی اشیای دیداری تمرکز می‌کند، فعالیت مورد بررسی تنها oddball دیداری بوده و فعالیت شنوایی مورد بررسی قرار نگرفته است.

در صد است. بیشترین میزان حساسیت متعلق به دسته دایره قرمز بوده و بیشترین میزان تشخیص مربوط به دسته دایره سبز است. به دلیل اینکه این پژوهش روی

|              |                 |               |                 |               |              |               |               |               |              |               |
|--------------|-----------------|---------------|-----------------|---------------|--------------|---------------|---------------|---------------|--------------|---------------|
| کلاسهای پاسخ | خانه            | 70<br>12.3%   | 0<br>0.0%       | 0<br>0.0%     | 0<br>0.0%    | 0<br>0.0%     | 0<br>0.0%     | 0<br>0.0%     | 0<br>0.0%    | 100%<br>0.0%  |
|              | عکس<br>درهم شده | 0<br>0.0%     | 70<br>12.3%     | 0<br>0.0%     | 0<br>0.0%    | 0<br>0.0%     | 0<br>0.0%     | 0<br>0.0%     | 0<br>0.0%    | 100%<br>0.0%  |
|              | گربه            | 0<br>0.0%     | 0<br>0.0%       | 70<br>12.3%   | 0<br>0.0%    | 0<br>0.0%     | 1<br>0.2%     | 0<br>0.0%     | 0<br>0.0%    | 98.6%<br>1.4% |
|              | کفش             | 0<br>0.0%     | 0<br>0.0%       | 0<br>0.0%     | 71<br>12.5%  | 1<br>0.2%     | 1<br>0.2%     | 1<br>0.2%     | 0<br>0.0%    | 95.9%<br>4.1% |
|              | بتری            | 0<br>0.0%     | 0<br>0.0%       | 0<br>0.0%     | 0<br>0.0%    | 67<br>11.8%   | 0<br>0.0%     | 0<br>0.0%     | 0<br>0.0%    | 100%<br>0.0%  |
|              | قیچی            | 0<br>0.0%     | 0<br>0.0%       | 0<br>0.0%     | 0<br>0.0%    | 1<br>0.2%     | 69<br>12.1%   | 0<br>0.0%     | 0<br>0.0%    | 98.6%<br>1.4% |
|              | صندلی           | 0<br>0.0%     | 0<br>0.0%       | 0<br>0.0%     | 0<br>0.0%    | 0<br>0.0%     | 0<br>0.0%     | 69<br>12.1%   | 0<br>0.0%    | 100%<br>0.0%  |
|              | صورت            | 1<br>0.2%     | 1<br>0.2%       | 1<br>0.2%     | 0<br>0.0%    | 2<br>0.4%     | 0<br>0.0%     | 1<br>0.2%     | 71<br>12.5%  | 92.2%<br>7.8% |
|              |                 | 98.6%<br>1.4% | 98.6%<br>1.4%   | 98.6%<br>1.4% | 100%<br>0.0% | 94.4%<br>5.6% | 97.2%<br>2.8% | 97.2%<br>2.8% | 100%<br>0.0% | 98.1%<br>1.9% |
|              |                 | خانه          | عکس<br>درهم شده | گربه          | کفش          | بتری          | قیچی          | صندلی         | صورت         |               |

(شکل-۸): ماتریس درهم‌ریختگی تجمیعی متعلق به روش پیشنهادی برای مجموعه داده DS105

(Figure-8): accumulative confusion matrix of the proposed method for dataset DS105



(شکل-۹): نتایج صحت تشخیص روش پیشنهادی به ازای هر فرد مجموعه داده DS107

(Figure-9): accuracy results of the proposed method for each subject in dataset DS107

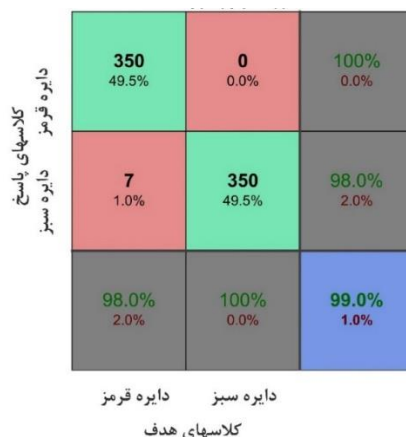
|              |                   |               |               |                   |                 |               |
|--------------|-------------------|---------------|---------------|-------------------|-----------------|---------------|
| کلاسهای پاسخ | کلمات             | 543<br>23.1%  | 20<br>0.9%    | 0<br>0.0%         | 0<br>0.0%       | 96.4%<br>3.6% |
|              | اشیاء             | 38<br>1.6%    | 563<br>24.0%  | 2<br>0.1%         | 1<br>0.0%       | 93.2%<br>6.8% |
|              | اشیاء<br>درهم شده | 5<br>0.2%     | 3<br>0.1%     | 583<br>24.8%      | 4<br>0.2%       | 98.0%<br>2.0% |
|              | حروف<br>نامفهوم   | 2<br>0.1%     | 2<br>0.1%     | 3<br>0.1%         | 580<br>24.7%    | 98.8%<br>1.2% |
|              |                   | 92.3%<br>7.7% | 95.7%<br>4.3% | 99.1%<br>0.9%     | 99.1%<br>0.9%   | 96.6%<br>3.4% |
|              |                   | کلمات         | اشیاء         | اشیاء<br>درهم شده | حروف<br>نامفهوم |               |

کلاسهای هدف

(شکل-۱۰): ماتریس درهم‌ریختگی روش پیشنهادی برای مجموعه داده DS107

(Figure-10): accumulative confusion matrix of the proposed method in DS107 dataset





(شکل-۱۱): ماتریس درهم‌ریختگی روش پیشنهادی برای مجموعه داده DS116  
(Figure-11): accumulative confusion matrix of the proposed method in DS116 dataset.

استفاده می‌کند که برچسب‌ها به دودویی کدگذاری (مانند یکی-مقابل-همه) و ویژگی‌ها بر اساس اعمال آستانه به مقادیر صفر و یک تبدیل می‌شوند. مقادیر بهینه آستانه و تعداد وجه‌های ویژگی به صورت خودکار با کمک الگوریتم جستجوی فرامکاشف‌ای SA تعیین می‌شوند.

## ۵- مقایسه و ارزیابی

مقایسه نتایج روش پیشنهادی با آخرین روش‌های کدگذاری اشیای دیداری، در جدول (۵) آورده شده است. فرق اصلی روش پیشنهادی با دیگر روش‌ها، در روند تفکیک فضای ویژگی با کمک اطلاعات متقابل بین ترکیب‌های برچسب-ویژگی است. روش پیشنهادی از محاسبه اطلاعات متقابل مبتنی بر رویکرد تقریبی

(جدول-۵): مقایسه آخرین روش‌های کدگذاری مغز انسان با روش پیشنهادی در حالت درون فردی

(Table-5): Comparison of the latest decoding methods of human brain with the proposed method in intra-subject case

| روش‌ها و مجموعه‌های داده           | سال        | DS105      | DS107      | DS116 <sup>۱</sup> |
|------------------------------------|------------|------------|------------|--------------------|
| L1 SVM [13], [50]                  | ۲۰۱۵، ۱۹۹۸ | ۸۵/۳±۲۹/۴۹ | ۸۱/۳±۲۵/۶۲ | ۶۹/۳±۲۴/۲۸         |
| L1 SVM + HA [15], [18]             | ۲۰۱۶، ۲۰۱۱ | ۸۷/۲±۰۳/۸۷ | ۸۴/۱±۰۱/۵۶ | ۷۴/۱±۶۲/۸۴         |
| L1 SVM + KHA [19]                  | ۲۰۱۲       | ۹۰/۲±۰۵/۳۹ | ۸۶/۱±۶۸/۷۱ | ۸۰/۲±۵۱/۱۲         |
| Graph Based [22]                   | ۲۰۱۶       | ۹۰/۱±۸۲/۸۷ | ۸۵/۱±۶۲/۹۵ | ۷۸/۲±۹۱/۰۴         |
| PSO-SVM [20]                       | ۲۰۱۶       | ۷۷/۱±۹۱/۰۳ | ۸۹/۲±۲۱/۳۳ | ۷۶/۱±۱۴/۴۹         |
| HHPSO-SVM [20]                     | ۲۰۱۶       | ۹۴/۱±۴۶/۲۳ | ۸۹/۱±۹۱/۶۷ | ۹۶/۰±۰۳/۵۶         |
| <sup>۲</sup> NSGA-II Kao [21]      | ۲۰۱۲       | ۹۵/۰±۳۲/۴۴ | ۸۷/۰±۷۷/۲۸ | ۸۴/۰±۱۶/۷۳         |
| بردار ویژگی روش پیشنهادی+SVM       | ۲۰۲۰       | ±۵/۸۹ ۹/۸۴ | ۸±۸۸       | ۹±۸۹               |
| (MOMVP) [49] linear kernel         | ۲۰۱۹       | ۹۳/۰±۶۱/۵۷ | ۹۲/۰±۸۳/۵۷ | ۹۰/۰±۹۳/۷۱         |
| (MOMVP) [49] Gaussian kernel       | ۲۰۱۹       | ۹۸/۰±۳۴/۲۹ | ۹۶/۰±۷۹/۵۹ | ۹۷/۱±۰۹/۳۳         |
| روش پیشنهادی MMVL-96 <sup>۳</sup>  | ۲۰۲۰       | ۹۷/۴±۷۱/۸۷ | ۹۴/۶±۳/۸۱  | ۹۷/۴±۲/۳           |
| روش پیشنهادی MMVL-965 <sup>۴</sup> | ۲۰۲۰       | ۹۸/۱±۱/۹   | ۹۶/۳±۶/۴   | ۱±۹۹               |

روش‌های دیگر به‌ازای هر مجموعه داده از یک ناحیه مورد توجه خاص استفاده کرده‌اند؛ اما روش پیشنهادی از مراحل انتخاب ویژگی‌های (وکسل‌ها) مؤثر در بخش ۳-۲

همچنین انتخاب وکسل‌های مورد استفاده در روش پیشنهادی از دیگر روش‌ها متفاوت است. در

<sup>۱</sup> در این مجموعه داده تنها از سناریوی دیداری استفاده کردیم که بقیه روش‌ها از سناریوی دیداری و شنوایی استفاده کردند. روش پیشنهادی بین دو شیء دیداری دایره سبز رنگ و دایره قرمز رنگ دسته‌بندی می‌کند؛ درحالی‌که بقیه روش‌ها بین حالت دیداری و حالت شنوایی تفکیک می‌کنند.

<sup>۲</sup> Non-dominant Search Algorithm

<sup>۳</sup> در حالت اطلس هاروارد اکسفورد-۹۶ ناحیه آناتومی قشره مغز

<sup>۴</sup> در حالت اطلس تالاریخ شامل ۹۶۵ ناحیه آناتومی قشره مغز

تک‌وجهی به‌صورت چشم‌گیری افزایش می‌یابد. در روش MMVL با انتخاب تعداد کافی وکسل‌ها صحت تشخیص به‌طور چشم‌گیری افزایش می‌یابد و در مقابل دیگر روش‌ها عملکرد مناسبی نشان می‌دهد. پیشنهاد روش تخمین وزن‌های دسته‌بندها در فرایند هم‌جوشی تصمیم‌های دسته‌بندها جزو کارهای آینده پیش‌بینی می‌شود. همچنین بهبود روش محاسبه اطلاعات متقابل در انتخاب و تفکیک فضای ویژگی پراهمیت تلقی می‌شود.

## 7- References

## ۷- مراجع

- [1] J. V. Haxby, M. I. Gobbini, M. L. Furey, A. Ishai, J. L. Schouten, and P. Pietrini, "Distributed and Overlapping Representations of Faces and Objects in Ventral Temporal Cortex," *Science* (80-. ), vol. 293, no. 5539, pp. 2425–2430, 2001.
- [2] T. Naselaris, K. N. Kay, S. Nishimoto, and J. L. Gallant, "Encoding and decoding in fMRI," *Neuroimage*, vol. 56, no. 2, pp. 400–410, 2011.
- [3] H. Uchida et al., "Visual Image Reconstruction from Human Brain Activity using a Combination of Multiscale Local Image Decoders," *Neuron*, vol. 60, no. 5, pp. 915–929, 2008.
- [4] ali mahloojifar, "Classification of Parkinson Disease Based on Inter - and Intra -Regional Biomarkers of the Brain Motor Network Using Resting State fMRI Data," *Signal Data Process.*, vol. 11, no. 2, 2015.
- [۴] قاسمی مهدیه، محلوچی فر علی. طبقه‌بندی بیماری پارکینسون بر مبنای شاخص‌های درون-ناحیه‌ای و بین-ناحیه‌ای شبکه حرکتی مغز با استفاده از داده‌گان fMRI حالت استراحت. پردازش علائم و داده‌ها. ۱۳۹۳؛ ۱۱ (۲): ۲۹-۱۵
- [5] P. Dayan and L. F. Abbott, *Theoretical neuroscience: computational and mathematical modeling of neural systems*. Computational Neuroscience Series, 2001.
- [6] J. L. Gallant and K. N. Kay, "I can see what you see.," *Nat. Neurosci.*, vol. 12, no. 3, p. 245, 2009.
- [7] N. Branin, "Reading Minds," *Sci. Am. Mind*, vol. 20, no. 6, pp. 8–8, 2009.
- [8] O. Sporns, "Network analysis, complexity, and brain function," *Complexity*, vol. 8, no. 1, pp. 56–60, 2002.
- [9] J. V. Haxby, "Multivariate pattern analysis of fMRI: The early beginnings," *Neuroimage*, vol. 62, no. 2, pp. 852–855, Aug. 2012.

بررسی‌شده، استفاده می‌کند. این مراحل اندازه فضای وکسل ورودی در روش پیشنهادی را به‌صورت قابل‌توجه کاهش می‌دهد. همچنین سرعت یادگیری را به کمتر از ده ثانیه می‌رساند. در مقابل زمان یادگیری در دیگر روش‌ها به‌صورت دقیق گزارش نشده است زمان پردازش روش MOMVP به‌عنوان واحد زمان فرض شده است و زمان پردازش روش‌های دیگر به‌صورت کسری از زمان پردازش MOMVP گزارش شده است [49]. هرچند صحت تشخیص روش پیشنهادی MMVL-96 پایین‌تر از روش MOMVP با هسته گاوسی<sup>۱</sup> است؛ اما صحت تشخیص روش MMVL-965 به‌صورت چشم‌گیری برای همه مجموعه‌های داده بیشتر شده و با روش MOMVP از لحاظ آماری تفاوت چندانی ندارد در این رابطه یک مقایسه آماری بر اساس آزمون ANOVA انجام شده است و عدم تفاوت معناداری بین دو روش در DS105 و DS107 ثابت شد. گزارش مقایسه آماری و زمان یادگیری و آزمون روش MMVL در پیوست الف آورده شده است که سرعت یادگیری و آزمون قابل‌مقایسه است. به‌عبارت‌دیگر در MMVL سعی شده است که مصالحه مناسبی میان صحت و زمان اجرا داشته باشیم. گفتنی است که با افزایش تعداد وکسل‌ها در بردار ویژگی، روش MMVL به نتایج بهتری دست می‌یابد.

## ۶- جمع‌بندی

در این مقاله، رویکرد نوین یادگیری چندوجهی برای حل مشکل کدگذاری اشیای دیداری از تصاویر مغزی را پیشنهاد دادیم. تفکیک فضای ویژگی مبتنی بر خوشه‌بندی سلسله‌مراتبی از اطلاعات متقابل ویژگی‌ها و برچسب‌ها انجام می‌شود. در رویکرد موجود، هم‌جوشی در سطح تصمیم‌گیری صورت می‌گیرد و در مقایسه با رویکردهای هم‌جوشی در سطح ویژگی، از سرعت و سادگی بیشتری برخوردار است. تفکیک خودکار فضای ویژگی مبتنی بر رویکرد تقریبی محاسبه اطلاعات متقابل یک نوآوری در این مقاله محسوب می‌شود؛ و انتخاب روش هم‌جوشی مناسب برای دسته‌بندهای حاصل از وجه‌های متفاوت مسأله با کمک روش میانگین، وزن‌های مرتب‌شده با عملگرهای استاندارد در فرایند هم‌جوشی تصمیمات مورد استفاده قرار گرفت. نشان داده شده که صحت تشخیص یادگیری گروهی چندوجهی در مقایسه با

<sup>1</sup> Gaussian Kernel

- [22] D. E. Osher, R. R. Saxe, K. Koldewyn, J. D. E. Gabrieli, N. Kanwisher, and Z. M. Saygin, "Structural Connectivity Fingerprints Predict Cortical Selectivity for Multiple Visual Categories across Cortex," *Cereb. Cortex*, vol. 26, no. 4, pp. 1668–1683, 2016.
- [23] S. Sun, "A survey of multi-view machine learning," *Neural Comput. Appl.*, vol. 23, no. 7–8, pp. 2031–2038, 2013.
- [24] R. Polikar, "Ensemble based systems in decision making," *Circuits Syst. Mag. IEEE*, vol. 6, no. 3, pp. 21–44, 2006.
- [25] O. Sagi and L. Rokach, "Ensemble learning: A survey," *Wiley Interdiscip. Rev. Data Min. Knowl. Discov.*, vol. 8, no. 4, p. e1249, 2018.
- [26] M. N. Hebart, K. GÅrgen, and J.-D. Haynes, "The Decoding Toolbox (TDT): a versatile software package for multivariate analyses of functional imaging data," *Front. Neuroinform.*, vol. 8, p. 88, Jan. 2015.
- [27] K. Seeliger *et al.*, "Convolutional neural network-based encoding and decoding of visual object recognition in space and time," *Neuroimage*, no. July, pp. 1–14, 2017.
- [28] M. P. Eckstein *et al.*, "Neural decoding of collective wisdom with multi-brain computing," *Neuroimage*, vol. 59, no. 1, pp. 94–108, 2012.
- [29] L. I. Kuncheva and J. J. Rodríguez, "Classifier ensembles for fMRI data analysis: an experiment," *Magn. Reson. Imaging*, vol. 28, no. 4, pp. 583–593, 2010.
- [30] W. J. Faithfull, J. J. Rodríguez, and L. I. Kuncheva, "Combining univariate approaches for ensemble change detection in multivariate data," *Inf. Fusion*, vol. 45, no. January 2018, pp. 202–214, 2019.
- [31] S. Sun, C. Zhang, and Y. Lu, "The random electrode selection ensemble for EEG signal classification," *Pattern Recognit.*, vol. 41, no. 5, pp. 1680–1692, May 2008.
- [32] T. M. Cover and J. A. Thomas, *Elements of information theory*. John Wiley & Sons, 2012.
- [33] S. Guia csu, *Information Theory with Applications*. New York : McGraw-Hill, 1977.
- [34] Y. Saeys, I. Inza, and P. Larranaga, "A review of feature selection techniques in bioinformatics," *Bioinformatics*, vol. 23, no. 19, pp. 2507–2517, Oct. 2007.
- [35] C. A. Chou, K. Kampa, S. H. Mehta, R. F. Tungaraza, W. A. Chaovalitwongse, and T. J. Grabowski, "Voxel selection framework in multi-voxel pattern analysis of fMRI data for prediction of neural response to visual stimuli," *IEEE Trans. Med. Imaging*, vol. 33, no. 4, pp. 925–934, Apr. 2014.
- [10] J. V. Haxby, A. C. Connolly, and J. S. Guntupalli, "Decoding Neural Representational Spaces Using Multivariate Pattern Analysis," *Annu. Rev. Neurosci.*, vol. 37, no. 1, pp. 435–456, 2014.
- [11] D. J. Heeger and D. Ress, "WHAT DOES fMRI TELL US ABOUT NEURONAL ACTIVITY?," *Nat. Rev. Neurosci.*, vol. 3, no. February, pp. 142–151, 2002.
- [12] K. Smith, "Brain imaging: fMRI 2.0," *Nature*, vol. 484, no. 7392, pp. 24–26, 2012.
- [13] P. S. Bradley and O. L. Mangasarian, "Feature selection via concave minimization and support vector machines.," in *ICML*, 1998, vol. 98, pp. 82–90.
- [14] H. R. Holger Mohr, Uta Wolfenstelle, Steffi Frimmel, H. Mohr, U. Wolfensteller, S. Frimmel, H. Ruge, and H. R. Holger Mohr, Uta Wolfenstelle, Steffi Frimmel, "Sparse regularization techniques provide novel insights into outcome integration processes," *Neuroimage*, vol. 104, pp. 163–176, 2015.
- [15] J. V. Haxby *et al.*, "A common, high-dimensional model of the representational space in human ventral temporal cortex," *Neuron*, vol. 72, no. 2, pp. 404–416, 2011.
- [16] D. D. Cox and R. L. Savoy, "Functional magnetic resonance imaging (fMRI) 'brain reading': Detecting and classifying distributed patterns of fMRI activity in human visual cortex," *Neuroimage*, vol. 19, no. 2, pp. 261–270, Jun. 2003.
- [17] T. Naselaris, R. J. Prenger, K. N. Kay, M. Oliver, and J. L. Gallant, "Bayesian Reconstruction of Natural Images from Human Brain Activity," *Neuron*, vol. 63, no. 6, pp. 902–915, 2009.
- [18] J. S. Guntupalli, M. Hanke, Y. O. Halchenko, A. C. Connolly, P. J. Ramadge, and J. V Haxby, "A model of representational spaces in human cortex," *Cereb. cortex*, vol. 26, no. 6, pp. 2919–2934, 2016.
- [19] A. Lorbort and P. J. Ramadge, "Kernel hyperalignment," in *Advances in Neural Information Processing Systems*, 2012, pp. 1790–1798.
- [20] X. Ma, C.-A. Chou, H. Sayama, and W. A. Chaovalitwongse, "Brain response pattern identification of fMRI data using a particle swarm optimization-based approach," *Brain Informatics*, vol. 3, no. 3, pp. 181–192, 2016.
- [21] M.-H. Kao, A. Mandal, and J. Stufken, "Constrained multiobjective designs for functional magnetic resonance imaging experiments via a modified non-dominated sorting genetic algorithm," *J. R. Stat. Soc. Ser. c (applied Stat.)*, vol. 61, no. 4, pp. 515–534, 2012.

Cortical Attention Networks and the Brainstem," *J. Neurosci.*, vol. 33, no. 49, pp. 19212–19222, 2013.

- [48] C. A. Chou, K. Kampa, S. H. Mehta, R. F. Tunga, W. A. Chaovalitwongse, and T. J. Grabowski, "Voxel selection framework in multi-voxel pattern analysis of fMRI data for prediction of neural response to visual stimuli," *IEEE Trans. Med. Imaging*, vol. 33, no. 4, pp. 925–934, Apr. 2014.
- [49] M. Yousefnezhad and D. Zhang, "Multi-Objective Cognitive Model: a Supervised Approach for Multi-subject fMRI Analysis," *Neuroinformatics*, vol. 17, no. 2, pp. 197–210, 2019.
- [50] H. R. Holger Mohr, Uta Wolfenstette, Steffi Frimmel, "Sparse regularization techniques provide novel insights into outcome integration processes," *Neuroimage*, vol. 104, pp. 163–176, 2015.
- [36] C. Cabral, M. Silveira, and P. Figueiredo, "Automatic classification of cognitive states," *1st Port. Meet. Biomed. Eng. ENBENG 2011*, no. February, 2011.
- [37] V. Gómez-Verdejo, M. Martínez-Ramón, J. Florensa-Vila, and A. Oliviero, "Analysis of fMRI time series with mutual information," *Med. Image Anal.*, vol. 16, no. 2, pp. 451–458, 2012.
- [38] O. Hourani, N. M. Charkari, and S. Jalili, "Voxel selection framework based on meta-heuristic search and mutual information for brain decoding," *Int. J. Imaging Syst. Technol.*, vol. 29, no. 4, pp. 663–676, Jun. 2019.
- [39] M. Jenkinson, C. F. Beckmann, T. E. J. J. Behrens, M. W. Woolrich, and S. M. Smith, "Fsl," *Neuroimage*, vol. 62, no. 2, pp. 782–790, 2012.
- [40] R. Buyya, G. M. Mohay, P. Roe, and IEEE Computer Society., "Sun Grid Engine: towards creating a compute power grid," in *Proceedings of the 1st International Symposium on Cluster Computing and the Grid*, 2001, p. 704.
- [41] R. S. Desikan *et al.*, "An automated labeling system for subdividing the human cerebral cortex on MRI scans into gyral based regions of interest," *Neuroimage*, vol. 31, no. 3, pp. 968–980, 2006.
- [42] J. L. Lancaster *et al.*, "Automated Talairach atlas labels for functional brain mapping," *Hum. Brain Mapp.*, vol. 10, no. 3, pp. 120–131, 2000.
- [43] P. Talairach, J. & Tournoux, "Co-planar stereotaxic atlas of the human brain," *Clin. Neurol. Neurosurg.*, vol. 91, no. 3, pp. 277–278, 1989.
- [44] M. W. Woolrich, B. D. Rippey, J. M. Brady, and S. M. Smith, "Temporal Autocorrelation in Univariate Linear Modelling of FMRI Data," *Neuroimage*, vol. 14, no. 6, pp. 1370–1386, 2001.
- [45] J. V. Haxby, M. I. Gobbini, M. L. Furey, A. Ishai, J. L. Schouten, and P. Pietrini, "Distributed and Overlapping Representations of Face and Objects in Ventral Temporal Cortex," *Science (80-. )*, vol. 293, no. 5539, pp. 2425–2430, 2001.
- [46] K. J. Duncan, C. Pattamadilok, I. Knierim, and J. T. Devlin, "Consistency and variability in functional localisers," *Neuroimage*, vol. 46, no. 4, pp. 1018–1026, 2009.
- [47] J. M. Walz, R. I. Goldman, M. Carapezza, J. Muraskin, T. R. Brown, and P. Sajda, "Simultaneous EEG-fMRI Reveals Temporal Evolution of Coupling between Supramodal

#### اسامه حورانی دوره کارشناسی خود را

در سال ۱۳۸۹ در رشته مهندسی رایانه نرم‌افزار در دانشگاه اصفهان گذراند. مدرک کارشناسی ارشد و دکترای خود را در رشته مهندسی رایانه نرم‌افزار از



دانشگاه تربیت مدرس به‌ترتیب در سال‌های ۱۳۹۲ و ۱۳۹۸ دریافت کرد. ایشان در حال حاضر پژوهشگر و مدرس در دانشگاه سیستم بلوچستان است. زمینه پژوهشی مورد علاقه ایشان پردازش تصاویر مغز انسان، بینایی ماشین و یادگیری ماشین است. نشانی رایانه‌ای ایشان عبارت است از:

hourani@modares.ac.ir

#### نصراله مقدم چرکری کارشناسی خود

را در علوم رایانه از دانشگاه شهید بهشتی تهران در سال ۱۳۶۵ دریافت کرد. تحصیلات کارشناسی ارشد و دکترای خود را در رشته‌های مهندسی



رایانه و مهندسی سامانه‌های اطلاعاتی در دانشگاه یاماناشی ژاپن به‌ترتیب در سال‌های ۱۳۷۱ و ۱۳۷۴ به پایان رساند. ایشان در حال حاضر دانشیار دانشکده مهندسی برق و کامپیوتر دانشگاه تربیت مدرس هستند.

زمینه‌های پژوهشی مورد علاقه ایشان آنالیز و بازیابی تصویر و ویدیو، بیوانفورماتیک و شبکه‌های پیچیده هستند. نشانی رایانامه ایشان عبارت است از:

**charkari@modares.ac.ir**



**سعید جلیلی** ایشان کارشناسی خود را در مهندسی رایانه از دانشگاه شهید بهشتی در سال ۱۳۵۸ اخذ و تحصیلات کارشناسی ارشد را در علوم رایانه از دانشگاه صنعتی شریف دریافت کرد.

ایشان دوره دکترای خود را در علوم رایانه از دانشگاه برادفورد انگلستان در سال ۱۳۷۰ دریافت کرد. اکنون ایشان دانشیار دانشکده برق و رایانه دانشگاه تربیت مدرس هستند. علایق پژوهشی ایشان طراحی جستجو بنیان معماری سامانه‌های نرم‌افزاری و بیوانفورماتیک هستند.

نشانی رایانامه ایشان عبارت است از:

**sjalili@modares.ac.ir**