





دانشکده مهندسی کامپیوتر و فناوری اطلاعات

گروه هوش مصنوعی

پایان نامه کارشناسی ارشد

تشخیص فعالیت‌های ذهنی در واسطه رایانه - مغز با استفاده از سیگنال EEG

رسول عامری

استاد راهنما

دکتر علی اکبر پویان

استاد مشاور

دکتر وحید ابوالقاسمی

بهمن ماه ۱۳۹۴

دانشگاه صنعتی شاهرود

دانشکده مهندسی کامپیوتر و فناوری اطلاعات

گروه هوش مصنوعی

پایان نامه کارشناسی ارشد آقای / خانم

تحت عنوان:

در تاریخ توسط کمیته تخصصی زیر جهت اخذ مدرک کارشناسی ارشد مورد ارزیابی و
با درجه مورد پذیرش قرار گرفت.

امضاء	اساتید مشاور	امضاء	اساتید راهنما
	نام و نام خانوادگی:		نام و نام خانوادگی:
	نام و نام خانوادگی:		نام و نام خانوادگی:

امضاء	نماینده تحصیلات تکمیلی	امضاء	اساتید داور
	نام و نام خانوادگی:		نام و نام خانوادگی:
			نام و نام خانوادگی:
			نام و نام خانوادگی:
			نام و نام خانوادگی:

تقدیم با تمام وجودم به

پدر، مادر عزیزم

و

برادر و خواهران مهربانم

پاس خدای عالمیان که اوست هدایتگر بی همتا

تشکر فراوان از

آقای دکتر پویان و آقای دکتر ابوالقاسمی

به خاطر زحمات بی دریغشان

و

تمامی آنهایی که مراراً بهمنای زندگی بوده اند.

رسول عامری

زمستان ۱۳۹۴

تعهد نامه

اینجانب **رسول عامری** دانشجوی دوره کارشناسی ارشد رشته **مهندسی کامپیوتر-هوش مصنوعی** دانشکده **مهندسی کامپیوتر و فناوری اطلاعات** دانشگاه صنعتی شاهرود نویسنده پایان نامه **تفحص فعالیت های ذهنی در واسطه رایانه- مغز با استفاده از سیگنال EEG** تحت راهنمایی **دکتر علی اکبر پویان** متعهد می شوم.

- تحقیقات در این پایان نامه توسط اینجانب انجام شده است و از صحت و اصالت برخوردار است.
- در استفاده از نتایج پژوهش های محققان دیگر به مرجع مورد استفاده استناد شده است.
- مطالب مندرج در پایان نامه تاکنون توسط خود یا فرد دیگری برای دریافت هیچ نوع مدرک با امتیازی در هیچ جا ارائه نشده است.
- کلیه حقوق معنوی این اثر متعلق به دانشگاه صنعتی شاهرود می باشد و مقالات مستخرج با نام «دانشگاه صنعتی شاهرود» و یا «Shahrood University of Technology» به چاپ خواهد رسید.
- حقوق معنوی تمام افرادی که در به دست آمدن نتایج اصلی پایان نامه تأثیرگذار بوده اند در مقالات مستخرج از پایان نامه رعایت می گردد.
- در کلیه مراحل انجام این پایان نامه، در مواردی که از موجود زنده (بافتهای آنها) استفاده شده است ضوابط و اصول اخلاقی رعایت شده است.
- در کلیه مراحل انجام این پایان نامه، در مواردی که به حوزه اطلاعات شخصی افراد دسترسی یافته یا استفاده شده است اصل رازداری، ضوابط و اصول اخلاقی انسانی رعایت شده است.

تاریخ

امضای دانشجو

مالکیت نتایج و حق نشر

- کلیه حقوق معنوی این اثر و محصولات آن (مقالات مستخرج، کتاب، برنامه های رایانه ای، نرم افزارها و تجهیزات ساخته شده است) متعلق به دانشگاه صنعتی شاهرود می باشد. این مطلب باید به نحو مقتضی در تولیدات علمی مربوطه ذکر شود .
- استفاده از اطلاعات و نتایج موجود در پایان نامه بدون ذکر مرجع مجاز نمی باشد.

چکیده

رابط رایانه - مغز (BCI) ابزار ارتباطی مناسب بین مغز انسان و دستگاه‌های بیرونی محیط اطراف می‌باشد. دقت و عملکرد دسته‌بندی سیگنال الکتروانسفالوگرام (EEG) نقش موثری در کارایی سیستم BCI ایفا می‌نماید. در این پژوهش دو روش مبتنی بر نمایش تنک (SR) و یادگیری دیکشنری (DL) برای دسته‌بندی فعالیت‌های ذهنی با استفاده از سیگنال EEG ارائه می‌شود. داده‌های مورد استفاده، پایگاه داده دانشگاه کلرادو و دانشگاه برلین می‌باشد که به ترتیب از پنج کلاس و دو کلاس تصور ذهنی حرکت ثبت شده‌اند. برای پیش‌پردازش از فیلتر میان‌گذر و الگوهای مکانی مشترک (CSP) استفاده می‌شود که در زمینه BCI از کارایی زیادی برخوردار است.

در روش پیشنهادی اول یک دسته‌بند بر مبنای نمایش تنک برای دسته‌بندی سیگنال‌های EEG ارائه شده است. در این روش برای ساخت دیکشنری از ویژگی‌های استخراج شده از CSP استفاده شده است. برای بسط دادن آن به یک مسئله چند کلاسه از CSP چند کلاسه استفاده شده است. در روش پیشنهادی دوم دسته‌بندی بر مبنای یادگیری دیکشنری ارائه شده است. در این روش یک دیکشنری تحلیلی برای محاسبه پاسخ تنک و یک دیکشنری مصنوعی برای کمینه‌سازی خطای بازسازی سیگنال آموزش داده می‌شود. با استفاده از دیکشنری تحلیلی ضرایب تنک با پیچیدگی محاسباتی کمتری به دست می‌آیند.

نتایج روش اول با استفاده از اعتبارسنجی بر روی پایگاه داده کلرادو ۹۰٪، ۸۳٫۵٪، ۷۶٫۵٪ و ۷۰٪ به ترتیب برای دسته‌بندی دو، سه، چهار و پنج کلاسه به دست می‌آید. میانگین نتایج بر روی پایگاه داده برلین با استفاده از اعتبارسنجی ۹۷٫۵۶٪ به دست آمد. میانگین نتایج به دست آمده از روش پیشنهادی دوم بر روی پایگاه داده آزمون دانشگاه برلین ۸۰٫۹۸٪ می‌باشد.

کلمات کلیدی: رابط رایانه - مغز، دسته‌بندی سیگنال EEG، نمایش تنک، یادگیری دیکشنری.

لیست مقالات مستخرج از پایان نامه

- Ameri R., Pouyan A., Abolghasemi V, (2015), “EEG Signal Classification Based on Sparse Representation in Brain Computer Interface Applications”, 22nd Iranian Conference on Biomedical Engineering, P 21-24, Tehran, Iran. It was indexed by IEEE.
- Ameri R., Pouyan A., Abolghasemi V, (2015), “Projective dictionary pair learning for EEG signal classification in brain computer interface applications”, At the revision stage in Neurocomputing.

فهرست مطالب

فصل ۱ مقدمه..... ۱

- ۱ - ۱ رابط رایانه - مغز چیست؟..... ۲
- ۲ - ۱ چالش‌های یک سیستم BCI..... ۴
- ۳ - ۱ انگیزه تشخیص فعالیت‌های ذهنی..... ۶
- ۱ - ۳ - ۱ ارتباطات و کنترل..... ۶
- ۲ - ۳ - ۱ نوروفیدبک و توانبخشی..... ۷
- ۳ - ۳ - ۱ پیگیری حالت روانی..... ۷
- ۴ - ۱ ساختار پایان‌نامه..... ۸

فصل ۲ ادبیات موضوع و کارهای مرتبط با رابط رایانه - مغز..... ۱۱

- ۱ - ۲ ساختار مغز..... ۱۲
- ۲ - ۲ سیگنال EEG..... ۱۴
- ۳ - ۲ سیستم BCI..... ۱۵
- ۱ - ۳ - ۲ ثبت سیگنال..... ۱۵
- ۲ - ۳ - ۲ پیش‌پردازش سیگنال..... ۱۶
- ۳ - ۳ - ۲ استخراج ویژگی..... ۱۷
- ۴ - ۳ - ۲ دسته‌بندی..... ۱۷
- ۴ - ۲ پژوهش‌های مرتبط..... ۱۸

فصل ۳ روش‌شناسی..... ۲۷

- ۱ - ۳ معرفی داده‌های موجود..... ۲۹
- ۱ - ۱ - ۳ پایگاه داده دانشگاه کلرادو..... ۲۹
- ۲ - ۱ - ۳ پایگاه داده برلین..... ۳۱
- ۲ - ۳ پیش‌پردازش..... ۳۲
- ۱ - ۲ - ۳ استفاده از الگوهای مکانی مشترک برای انتخاب کانال..... ۳۳
- ۱ - ۲ - ۳ الگوهای مکانی مشترک دو کلاسه..... ۳۴
- ۲ - ۱ - ۲ - ۳ الگوهای مکانی مشترک چند کلاسه..... ۳۶
- ۳ - ۳ نمایش تنک..... ۳۸
- ۱ - ۳ - ۳ مروری بر روش‌های محاسبه ضرایب تنک..... ۴۰
- ۱ - ۱ - ۳ - ۳ روش‌های نرم L1..... ۴۰
- ۲ - ۱ - ۳ - ۳ روش‌های حریص..... ۴۱

۴۳.....	۳ - ۱ - ۳ - ۳ روش‌های آستانه‌گذاری
۴۳.....	۳ - ۱ - ۳ - ۴ روش‌های تقریب نرم L0
۴۳.....	۳ - ۴ - ۴ یادگیری دیکشنری
۴۴.....	۳ - ۴ - ۱ روش K-SVD
۴۵.....	۳ - ۴ - ۲ روش LC-KSVD
۴۶.....	۳ - ۵ روش‌های پیشنهادی
۴۶.....	۳ - ۵ - ۱ روش پیشنهادی دسته‌بندی سیگنال EEG با استفاده نمایش تنک
۵۲.....	۳ - ۵ - ۲ روش پیشنهادی یادگیری دیکشنری برای دسته‌بندی
۵۹.....	۳ - ۵ - ۳ استفاده از الگوریتم ژنتیک برای محاسبه پارامترهای اولیه
۶۰.....	۳ - ۵ - ۱ پارامترهای الگوریتم ژنتیک

فصل ۴ نتایج ۶۳

۶۴.....	۴ - ۱ قسمت‌های مختلف پایگاه داده
۶۶.....	۴ - ۲ محاسبه پارامترهای اولیه
۶۷.....	۴ - ۲ - ۱ پیاده‌سازی الگوریتم ژنتیک برای محاسبه پارامترها
۶۹.....	۴ - ۳ استخراج ویژگی و دسته‌بندی
۷۱.....	۴ - ۴ بررسی نتایج ارزیابی نمایش تنک
۷۱.....	۴ - ۴ - ۱ پیاده‌سازی روی پایگاه داده کلرادو
۷۳.....	۴ - ۴ - ۱ مقایسه روش با سایر دسته‌بندها
۷۴.....	۴ - ۴ - ۲ بررسی عملکرد روش‌های بازسازی تنک
۷۵.....	۴ - ۴ - ۳ مقایسه روش با سایر کارهای انجام شده
۷۷.....	۴ - ۴ - ۲ پیاده‌سازی روی پایگاه داده برلین
۷۷.....	۴ - ۴ - ۱ مقایسه روش با سایر دسته‌بندها
۷۸.....	۴ - ۴ - ۲ بررسی عملکرد استخراج ویژگی‌های مختلف
۷۹.....	۴ - ۴ - ۳ بررسی عملکرد روش‌های بازسازی تنک
۷۹.....	۴ - ۴ - ۳ بررسی ویژگی‌های دیکشنری
۸۱.....	۴ - ۵ بررسی عملکرد روش DPL
۸۱.....	۴ - ۵ - ۱ پیاده‌سازی روش روی پایگاه داده برلین
۸۲.....	۴ - ۵ - ۲ بررسی نتایج برای اندازه‌های مختلف دیکشنری
۸۲.....	۴ - ۵ - ۳ مقایسه روش با سایر کارهای انجام شده
۸۳.....	۴ - ۵ - ۴ مقایسه روش با سایر روش‌های یادگیری دیکشنری
۸۵.....	۴ - ۵ - ۵ بررسی ویژگی‌های دیکشنری

فصل ۵ نتیجه‌گیری و پیشنهاد ۸۷

مراجع ۹۳

فهرست شکل‌ها

- شکل ۱-۱. اجزای یک سیستم BCI [۲]..... ۴
- شکل ۱-۲. مخ و لوب‌های تشکیل‌دهنده آن [۹]..... ۱۳
- شکل ۲-۲. سیستم BCI..... ۱۶
- شکل ۱-۳. مکان الکترودهای استفاده شده در پایگاه داده دانشگاه کلرادو [۲۸]..... ۲۹
- شکل ۲-۳. نمونه سیگنال EEG ثبت شده در پایگاه داده دانشگاه کلرادو..... ۳۱
- شکل ۳-۳. زمان‌بندی انجام ثبت یک آزمایش پایگاه داده برلین..... ۳۲
- شکل ۴-۳. فیلتر میان‌گذر باترورث..... ۳۳
- شکل ۵-۳. الگوریتم CSP چند کلاسه..... ۳۷
- شکل ۶-۳. انواع نرم‌های مختلف برای یک بردار دو بعدی..... ۴۱
- شکل ۷-۳. الگوریتم MP..... ۴۲
- شکل ۸-۳. دسته‌بندی سیگنال با استفاده از نمایش تنک..... ۴۶
- شکل ۹-۳. نحوه ساخت دیکشنری..... ۴۹
- شکل ۱۰-۳. محاسبه پاسخ تنک با استفاده از دیکشنری طراحی شده..... ۵۰
- شکل ۱۱-۳. روش پیشنهادی DPL..... ۵۴
- شکل ۱۲-۳. الگوریتم DPL..... ۵۷
- شکل ۱۳-۳. مراحل الگوریتم ژنتیک..... ۵۹
- شکل ۱۴-۳. نمونه کروموزوم..... ۶۱
- شکل ۱-۴. روند الگوریتم ژنتیک برای حل مسئله..... ۶۸
- شکل ۲-۴. بررسی عملکرد روش‌های به دست آوردن پاسخ تنک..... ۷۵
- شکل ۳-۴. مقایسه روش با دسته‌بندی‌های KNN، SVM، MLP و LDA..... ۷۸

فهرست جدول‌ها

جدول ۱-۱	مقایسه روش‌های ثبت سیگنال مغز [۱]..... ۳
جدول ۱-۲	بازه‌های فرکانسی سیگنال EEG [۱۱]..... ۱۴
جدول ۱-۴	تعداد نمونه‌های آموزشی و آزمون برای پایگاه داده برلین..... ۶۴
جدول ۲-۴	تعداد نمونه‌های پایگاه داده کلرادو..... ۶۵
جدول ۳-۴	نتایج برخی از کروموزوم‌ها..... ۶۸
جدول ۴-۴	نتایج برخی از آزمایش‌ها برای انتخاب بازه زمانی برای فرد aa..... ۶۹
جدول ۵-۴	نتایج روش نمایش تنک روی پایگاه داده کلرادو..... ۷۱
جدول ۶-۴	نتایج حالات مختلف دسته‌بندی دو کلاسه برای فرد اول..... ۷۲
جدول ۷-۴	نتایج حالات مختلف دسته‌بندی سه کلاسه برای فرد اول..... ۷۳
جدول ۸-۴	مقایسه روش پیشنهادی نمایش تنک با سایر دسته‌بندها..... ۷۴
جدول ۹-۴	بررسی عملکرد روش‌های بازسازی تنک..... ۷۴
جدول ۱۰-۴	مقایسه روش دسته‌بندی نمایش تنک با کارهای گذشته..... ۷۶
جدول ۱۱-۴	نتایج روش دسته‌بندی تنک بر روی پایگاه داده برلین..... ۷۷
جدول ۱۲-۴	بررسی عملکرد استخراج ویژگی‌های مختلف..... ۷۸
جدول ۱۳-۴	مقایسه عملکرد روش‌های بازسازی تنک..... ۷۹
جدول ۱۴-۴	مقادیر همبستگی برای پایگاه داده‌های برلین..... ۸۰
جدول ۱۵-۴	بررسی عملکرد فیلتر CSP در ساخت دیکشنری..... ۸۰
جدول ۱۶-۴	نتایج روش DPL روی پایگاه داده برلین..... ۸۲
جدول ۱۷-۴	میزان دقت دسته‌بند برای اندازه‌های مختلف دیکشنری..... ۸۲
جدول ۱۸-۴	مقایسه روش DPL با کارهای گذشته..... ۸۳
جدول ۱۹-۴	مقایسه روش با سایر روش‌های یادگیری دیکشنری..... ۸۴
جدول ۲۰-۴	میزان همبستگی دیکشنری ساخته شده با روش‌های DPL، LC-KSVD1 و LC-KSVD2..... ۸۵

فصل ۱

مقدمه

در سال‌های اخیر، تحقیقات زیادی بر روی تعامل انسان با کامپیوتر^۱ (HCI) با استفاده از الگوهای تعاملی مانند صدا، لمس و مغز انجام شده است. توانایی ثبت پتانسیل الکتریکی از مغز و تبدیل آن به دستوراتی برای ارتباط با محیط اطراف توجه زیادی از محققین را به استفاده از این ویژگی در HCI جلب نموده است. رابط رایانه - مغز^۲ یا BCI یک سیستم ارتباطی بین افراد و محیط اطراف بدون استفاده از واسط عضلانی می‌باشد. از کاربردهای BCI می‌توان بهبود کیفیت زندگی افراد ناتوان حرکتی و کنترل واقعیت مجازی را نام برد.

۱-۱ رابط رایانه - مغز چیست؟

BCI یک سیستم ارتباطی است که به افراد توانایی کنترل ابزار و وسایل محیط اطراف را با استفاده از فعالیت‌های عصبی مغز می‌دهد. فعالیت‌های عصبی مغز با استفاده از روش‌های تهاجمی و غیرتهاجمی^۳ برای استفاده در BCI ثبت می‌شوند. در روش‌های تهاجمی برای ثبت سیگنال مغزی، الکترودها بر روی مغز و در روش‌های غیرتهاجمی الکترودها بر روی سر قرار می‌گیرند. در تحقیقات BCI روش‌های تهاجمی به دلیل ریسک عمل جراحی نسبت به روش‌های غیرتهاجمی از کاربرد کمتری برخوردارند.

از پرکاربردترین این روش‌ها می‌توان الکتروانسفالوگرافی^۴ (EEG)، مغزنگاری مغناطیسی^۵ (MEG)، الکتروکورتیکوگرافی^۶ (ECoG)، تصویرسازی تشدید مغناطیسی کارکردی^۷ (fMRI) را نام برد. خلاصه‌ای از فواید و معایب این روش‌ها در جدول ۱-۱ آورده شده است.

^۱ Human computer interaction.

^۵ Magnetoencephalography.

^۲ Brain computer interface.

^۶ Electrooculography.

^۳ Invasive and non-invasive.

^۷ Functional magnetic resonance imaging.

^۴ Electroencephalography.

در روش ECoG الکترودها با سطح مغز تماس دارند که در نتیجه دقت سیگنال‌های ثبت شده از EEG بالاتر می‌باشد اما به دلیل تهاجمی بودن این روش، کاربرد کمتری نسبت به EEG در کارهای تحقیقاتی BCI دارد. تفکیک‌پذیری فضایی برای سیگنال MEG در محدوده میلی‌متر می‌باشد و دقت فضایی مناسبی دارد. یکی از معایب این سیگنال دامنه کم آن می‌باشد که شرایط ثبت سیگنال را دچار مشکل می‌نماید و باید در مکان‌هایی با پوشش مغناطیسی انجام شود. از طرف دیگر، در کاربردهای بلندمدت، EEG دارای وضوح و دقت زمانی کافی می‌باشد اما تفکیک‌پذیری مکانی آن به تعداد الکترودها بستگی دارد و بیشترین تفکیک‌پذیری آن در محدوده سانتی‌متر می‌باشد. روش ثبت سیگنال EEG با توجه به قیمت ارزان نسبت به سایر روش‌های ثبت سیگنال و همچنین غیرتهاجمی بودن بهترین روش ثبت سیگنال در کارهای تحقیقاتی BCI می‌باشد.

جدول ۱-۱. مقایسه روش‌های ثبت سیگنال مغز [۱]

روش	مزایا	معایب
EEG	ارزان‌قیمت، تفکیک‌پذیری زمانی مناسب غیرتهاجمی، انجام آسان ثبت سیگنال	تفکیک‌پذیری مکانی پایین
ECoG	تفکیک‌پذیری زمانی و مکانی مناسب	بسیار تهاجمی، نیاز به عمل جراحی
MEG	غیرتهاجمی، تفکیک‌پذیری مکانی بسیار بالا تفکیک‌پذیری زمانی مناسب	تجهیزات بسیار گران‌قیمت نیاز به اتاق ایزوله مغناطیسی برای ثبت سیگنال
FMRI	تفکیک‌پذیری مکانی مناسب، غیرتهاجمی	بسیار گران‌قیمت

در شکل ۱-۱ اجزای مختلف یک سیستم BCI آورده شده است. این سیستم متشکل از یک کلاه الکترود مجهز به الکترودهای متصل به یک تقویت‌کننده است که سیگنال به دست آمده را تقویت و آن را از آنالوگ به دیجیتال تبدیل می‌کند.

در این پایان‌نامه سیستم BCI بر اساس EEG مورد بررسی قرار گرفته است که اجزای آن در بخش ۲

- ۳ به طور مفصل بررسی خواهد شد.



شکل ۱-۱. اجزای یک سیستم BCI [۲]

۱-۲ چالش‌های یک سیستم BCI

در سال‌های اخیر با پیشرفت فناوری مشکلاتی مانند اندازه، قابلیت حمل، ثبت سیگنال و هزینه تجهیزات یک سیستم BCI تا حد زیادی بهبود یافته است. این امر منجر به افزایش ورود سیستم‌های BCI به دنیای تجاری شده است. در حال حاضر استفاده از BCI محدود به کارهای آزمایشگاهی و بازی‌ها می‌شود. با این وجود استفاده از BCI‌هایی با زمینه‌های غیر پزشکی در سال‌های آینده بیشتر خواهد شد. برخی از زمینه‌های این BCI‌ها عبارتند از:

- وسایل کنترلی: کنترل کردن وسایلی مانند صندلی چرخدار یا وسایل مصنوعی دیگر [۴], [۳].
- بررسی وضعیت افراد: جمع‌آوری اطلاعاتی در مورد حالات یک شخص از جمله سطح هوشیاری،

احساسات افراد. این برنامه را می‌توان در برنامه‌های کاربردی مانند نظارت بر هوشیاری راننده استفاده نمود [۵].

- سرگرمی و بازی.
 - بهبود شناختی: ارائه آموزش بازخورد عصبی^۱ برای بهبود یا بازسازی حافظه و اجرای فعالیت‌های خاص توسط افراد.
 - سلامت و امنیت: برای نمونه دستگاه‌های دروغ‌سنج با استفاده از سیگنال‌های EEG [۶].
- اگر چه بسیاری از موانع برای کاربردی شدن سیستم‌های BCI برطرف شده است اما چالش‌ها و مشکلاتی نیز در این زمینه وجود دارد که عبارتند از:
- قابلیت استفاده: سهولت استفاده از سیستم BCI یکی از مواردی است که با بهبود آن می‌توان برنامه‌های کاربردی را در دنیای تجاری عرضه نمود. آماده‌سازی یک سیستم BCI و استفاده از ژل بر روی سر یک کار خسته‌کننده برای کاربران می‌باشد.
 - قابل حمل بودن دستگاه: اندازه و وزن دستگاه BCI یک چالش مهم برای استفاده از آن در دنیای واقعی می‌باشد.
 - هزینه: سیستم‌های BCI با ارائه کیفیت بالا برای ثبت سیگنال شامل حذف نویز و مصنوعات^۲ در دهه‌های گذشته بسیار گران بود. در سال‌های اخیر با ظهور شرکت‌هایی در زمینه ساخت دستگاه‌های مورد نیاز BCI هزینه‌ها تا حدودی کاهش یافته است.
 - پردازش سیگنال: سیگنال‌های مغز نسبتاً بسیار ضعیف و پیچیده می‌باشند؛ از این جهت روش‌های پردازش سیگنال برای حذف نویز و مصنوعات مختلف قبل از استفاده از سیگنال ثبت

^۱ Neurofeedback training.

^۲ Artifact.

شده باید انجام شود [۶].

۱-۳ انگیزه تشخیص فعالیت‌های ذهنی

با استفاده از سیستم BCI برای تشخیص فعالیت‌های ذهنی، طیف وسیعی از برنامه‌های کاربردی را می‌توان برای افزایش کیفیت زندگی افراد سالم و افراد دارای معلولیت جسمانی ارائه کرد. در دهه‌های گذشته امکان ترجمه افکار به دستوراتی برای کنترل وسایل اطراف غیر قابل تصور بود؛ اما در حال حاضر با حضور BCI یک واقعیت می‌باشد. نرم‌افزارهای کاربردی برای دسته‌بندی فعالیت‌های ذهنی با استفاده از سیستم BCI در این بخش مورد بررسی قرار خواهد گرفت.

۱-۳-۱ ارتباطات و کنترل

یکی از مهمترین نیازها برای افراد معلول را می‌توان کنترل محیط اطراف و توانایی برقراری ارتباط با دیگران دانست. از سیستم BCI می‌توان برای ترجمه نمودن برخی فعالیت‌های ذهنی به سیگنال‌های کنترلی مانند حرکت صندلی چرخدار، برقراری تماس تلفنی با افراد، کنترل نمودن حرکت مکان‌نما بر روی صفحه مانیتور استفاده نمود. ساخت چنین دستگاه‌هایی می‌تواند به بهبود کیفیت زندگی افراد معلول کمک شایانی نماید. عملکرد چنین سیستم‌هایی به ثبت سیگنال با کیفیت بالا، تجهیزات پردازش و کارهای انتخابی برای آموزش سیستم بستگی دارد. به طور عمده فعالیت‌هایی برای دسته‌بندی انتخاب می‌شود که آموزش آن برای کاربر آسان بوده و همچنین سیستم قادر به دسته‌بندی آن باشد. برخی از فعالیت‌های ذهنی مورد استفاده در کارهای BCI عبارتند از تصور حرکت دست، انجام ضرب ذهنی و غیره [۷]، [۸].

۱- ۳- ۲ نوروفیدبک^۱ و توانبخشی

یکی از کاربردهای دسته‌بندی فعالیت‌های ذهنی در آموزش نوروفیدبک می‌باشد. نوروفیدبک از نمایش به‌هنگام^۲ سیگنال EEG برای بررسی و بهبود عملکرد سیستم عصبی مرکزی^۳ استفاده می‌شود. نوروفیدبک در درمان اضطراب، ناتوانی یادگیری، کمبود توجه و افزایش قدرت تمرکز و توانایی‌های ذهنی در افراد سالم مؤثر می‌باشد. علاوه بر برنامه‌های کاربردی بیان شده برای نوروفیدبک، استفاده از BCI برای توانبخشی اعصاب مغز نیز مطرح شده است. مطالعات انجام شده بر روی بیماران سکته مغزی که از مشکل حرکتی رنج می‌برند نشان می‌دهد که استفاده از سیستم BCI در بهبود مشکل حرکتی مؤثر می‌باشد [۲].

۱- ۳- ۳ پیگیری حالت روانی

در دو بخش گذشته جنبه‌های پزشکی دسته‌بندی فعالیت‌های ذهنی بیان شد. در این بخش استفاده افراد سالم از سیستم BCI مطرح می‌باشد. با افزایش روز افزون کاربران گوشی‌های هوشمند، تبلت‌ها و غیره پیگیری فعالیت بدنی، وضعیت روحی، مصرف کالری برای کاربران می‌تواند مورد توجه قرار گیرد. این فرآیند از طریق حسگرهای ارائه شده در ابزارهای هوشمند امکان‌پذیر می‌باشد.

برنامه‌های کاربردی BCI با تشخیص حالت ذهنی می‌توانند فرد را از میزان تمرکز، سطح استرس، میزان خستگی و غیره مطلع سازند. آگاهی فرد از شرایط روحی و روانی می‌تواند به فرد در برنامه‌ریزی بهتر برای شرایط کنونی کمک نماید. به عنوان مثال شناسایی سطح هیجان، خستگی، احساس، حجم کار و یا سایر متغیرهای مربوط به وضعیت روانی کاربر می‌تواند در طراحی رابط ماشین مطابق با وضعیت روانی

^۱ Neurofeedback.

^۲ Central nervous system.

^۳ Real time.

راننده استفاده شود [۲].

۱-۴ ساختار پایان نامه

این پایان نامه از پنج فصل تشکیل شده است. در فصل اول به معرفی سیستم BCI و چالش های موجود در تحقیقات مربوط به آن پرداختیم. سپس کاربردهای سیستم BCI و دسته بندی فعالیت های ذهنی را بیان نمودیم.

فصل دوم «ادبیات موضوع و کارهای مرتبط» از چهار بخش تشکیل شده است. در بخش اول به بررسی ساختار مغز و در بخش دوم به معرفی مختصر سیگنال EEG پرداخته شده است. در بخش سوم ساختار سیستم BCI و بررسی اجزای تشکیل دهنده معرفی شده است. در بخش پایانی فصل دوم کارهای اخیر انجام شده در زمینه دسته بندی فعالیت های ذهنی مورد تجزیه و تحلیل قرار گرفته است.

در فصل روش دسته بندی سیگنال EEG با استفاده از نمایش تنک^۱ (SR) و یادگیری دیکشنری^۲ (DL) مطرح شده است. در بخش اول فصل سوم به معرفی پایگاه داده های مورد استفاده برای ارزیابی روش پرداخته شده است. در بخش دوم و سوم روش های پیش پردازش و استخراج ویژگی انجام شده بررسی شده است. در بخش سوم نمایش تنک و یادگیری دیکشنری معرفی شده و در پایان دو روش پیشنهادی برای دسته بندی سیگنال EEG ارائه شده است.

در فصل چهارم نتایج روش های پیشنهادی بر روی دو پایگاه داده دانشگاه برلین و کلرادو مورد ارزیابی قرار گرفته است. در فصل پنجم نتیجه گیری و پیشنهاد های برای بهبود روش های پیشنهادی مطرح شده

^۱ Sparse representation.

^۲ Dictionary learning.

است.

فصل ۲

ادبیات موضوع و کارهای

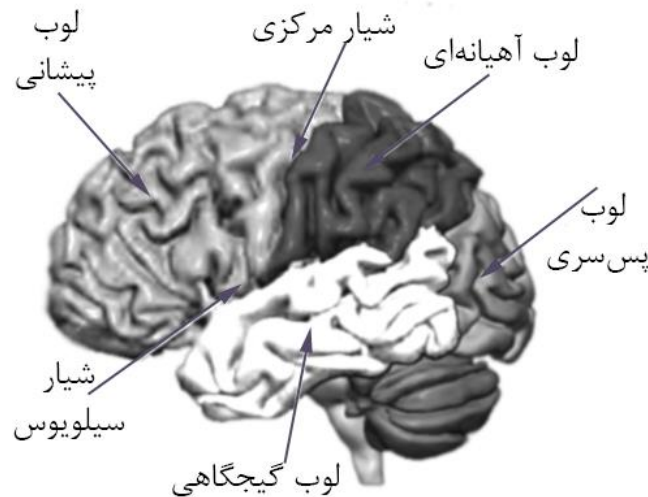
مرتبط با رابط رایانه - مغز

در این فصل پس از بررسی ساختار مغز، سیگنال EEG و چگونگی ثبت سیگنال EEG از مغز معرفی شده است. در بخش سوم ساختار سیستم BCI و اجزای تشکیل دهنده آن معرفی شده است. کارهای مرتبط انجام شده در زمینه دسته‌بندی فعالیت‌های ذهنی در BCI با استفاده از سیگنال EEG در بخش پایانی بیان شده است.

۲-۱ ساختار مغز

مغز مرکز سیستم عصبی بدن و پیچیده‌ترین اندام در بدن انسان است. همچنین این اندام مسئول کنترل متمرکز بر تمام اندام‌های دیگر بدن است. مغز از سه قسمت مخ، مخچه و ساقه مغز تشکیل شده است. مخ بزرگترین قسمت مغز و مسئول عملکردهای حیاتی بدن است. همانطور که در شکل ۲-۱ مشاهده می‌کنید، مخ شامل چهار لوب می‌باشد که عبارتند از:

- لوب پیشانی^۱: این لوب انقباض عضلات ارادی، تمرکز، تصمیم‌گیری و برنامه‌ریزی را اداره می‌کند.



شکل ۱-۲. مخ و لوب‌های تشکیل‌دهنده آن [۹]

- لوب آهیانه‌ای^۲: در این لوب اطلاعات چشایی، لمسی و حرکتی دریافت می‌شود.
 - لوب گیجگاهی^۳: در این لوب اطلاعات حسی مانند شنوایی تفسیر شده و در حافظه ذخیره می‌شود.
 - لوب پس‌سری^۴: این لوب مسئول پردازش اطلاعات بصری می‌باشد.
- اگر چه هر یک از لوب‌ها با یک گروه از حواس و عملکردها در ارتباط است اما برای انجام یک فعالیت با یکدیگر در ارتباط می‌باشند. برای مثال فعالیت «خواندن کتاب»؛ لوب پس‌سری مسئول پردازش اطلاعات بصری مانند حروف، کلمات و شکل‌های موجود در کتاب می‌باشد. لوب پیشانی مسئول فهمیدن کلمات خوانده شده می‌باشد [۹].

^۱ Frontal lobe.

^۲ Parietal lobe.

^۳ Temporal lobe.

^۴ Occipital lobe.

در نگاه دقیق‌تر مغز انسان از ۱۰۰ میلیون نورون^۱ تشکیل شده است. نورون یک سلول است که با استفاده از سیگنال‌های الکتریکی و شیمیایی اطلاعات را دریافت، سپس تحلیل و انتقال دهد. مغز شامل شبکه‌ای از نورون‌ها می‌باشد. در این شبکه، نورون‌ها اطلاعات (حرکتی، حسی و شناختی) را از دیگر نورون‌های شبکه دریافت نموده و به اندام‌های بدن انتقال می‌دهند. این اطلاعات از طریق سیناپس‌ها^۲ منتقل می‌شوند. انتقال اطلاعات از طریق نورون‌ها منجر به ایجاد میدان الکتریکی می‌شود که با استفاده از روش‌هایی مانند EEG و ECoG قابل ثبت می‌باشد. در این پایان‌نامه از سیگنال‌های EEG ثبت شده از فعالیت ذهنی استفاده شده است که در بخش ۲-۲ مورد بررسی قرار خواهد گرفت.

۲-۲ سیگنال EEG

سیگنال EEG در سال ۱۹۲۹ توسط برگر اختراع شد. این سیگنال از پرکاربردترین سیگنال‌های مورد استفاده در سیستم‌های BCI می‌باشد. تفکیک‌پذیری زمانی سیگنال EEG برای سیگنال‌های مغز کمتر از یک میلی‌ثانیه می‌باشد که برتری این سیگنال بر سایر روش‌های ثبت سیگنال می‌باشد. برای افزایش تفکیک‌پذیری مکانی نیز از تعداد الکترود بیشتری در BCI های مدرن (۲۵۶ الکترود) استفاده می‌شود [۱۰].

جدول ۱-۲. بازه‌های فرکانسی سیگنال EEG [۱۱]

باند فرکانس	بازه فرکانس (هرتز)	حالت ذهنی
دلتا ^۳	۰٫۱ - ۴	خواب عمیق یا بیهوشی
تتا ^۴	۴ - ۸	خواب و رؤیا دیدن
آلفا ^۵	۸ - ۱۳	آرام و بدون حرکت و حالت هوشیار
بتا ^۶	۱۳ - ۳۰	آرامش همراه با تفکر

تجزیه	فعالیت‌های حرکتی	۳۰ - ۱۰۰	گاما ^۷
-------	------------------	----------	-------------------

و تحلیل سیگنال‌های EEG در بازه‌های فرکانسی مختلف به دسته‌بندی مناسب‌تر فعالیت‌های ذهنی کمک می‌کند. سیگنال‌های ثبت شده با استفاده از روش‌های پردازش سیگنال مانند تبدیل فوریه سریع^۸ (FFT) و فیلتر میان‌گذر^۹ به چندین بازه فرکانسی تقسیم می‌شوند. در جدول ۱-۲ بازهای فرکانسی مختلف و فعالیت‌های ذهنی انجام شده آنها آورده شده است.

۲ - ۳ سیستم BCI

یک سیستم BCI از تعدادی سنسور برای دریافت سیگنال مغز و اجزای پردازش سیگنال تشکیل شده است که فعالیت‌های مغزی را به یک سری سیگنال‌های ارتباطی و کنترلی برای تعامل با محیط اطراف تبدیل می‌نماید. اجزای اصلی یک سیستم BCI در شکل ۲-۲ آورده شده است. این اجزا عبارتند از:

۲ - ۳ - ۱ ثبت سیگنال

ثبت سیگنال از مغز با دو روش تهاجمی و غیرتهاجمی انجام می‌شود که در فصل قبل مورد بررسی قرار گرفت. ثبت سیگنال با استفاده از الکترودهای قرار گرفته بر روی سر انجام می‌شود. استاندارد ۱۰-۲۰ یکی از پر استفاده‌ترین سیستم‌ها برای قرار گرفتن الکترودها بر روی سر می‌باشد [۱۱]. پس از ثبت

^۱ Neuron.

^۲ Synapses.

^۳ Delta.

^۴ Theta.

^۵ Alpha.

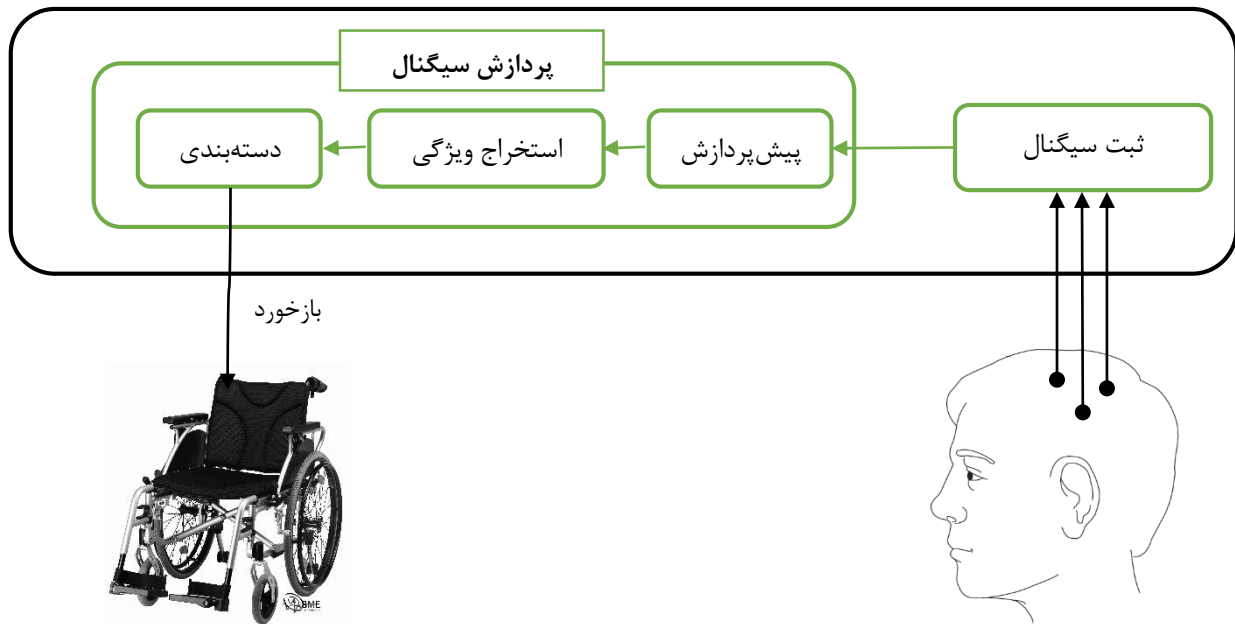
^۶ Beta.

^۷ Gamma.

^۸ Fast Fourier transform.

^۹ Band pass filter.

سیگنال سه مرحله اساسی برای استخراج اطلاعات مهم از سیگنال پیش پردازش، استخراج ویژگی و دسته بندی می باشد که در ادامه توضیح داده خواهد شد.



شکل ۲-۲. سیستم BCI

۲ - ۳ - ۲ پیش پردازش سیگنال

سیگنال EEG ثبت شده آغشته به نویز و مصنوعات می باشد که منشأ غیر از مغز دارند. هدف اصلی این مرحله بهبود کیفیت سیگنال ثبت شده افزایش نسبت سیگنال به نویز^۱ می باشد. منبع نویز و مصنوعات موجود در سیگنال به دو دسته فیزیولوژی و غیر فیزیولوژی^۲ تقسیم می شوند. فیزیولوژی شامل حرکت ماهیچه های بدن^۳ (EMG)، حرکت چشم و پلک زدن^۴ (EOG) و ضربان قلب^۵ (ECG) می باشند. منبع مصنوعات غیر فیزیولوژی نیز می تواند از تجهیزات معیوب و تجهیزات الکتریکی موجود در محیط

^۱ Signal to noise ratio.

^۴ Electrocularographic.

^۲ Physiological and non- physiological.

^۵ Electrocardiogram.

^۳ Electromyographic.

آزمایشگاه باشد. روش‌های متعددی برای حذف نویز و مصنوعات وجود دارد از جمله فیلتر میان‌گذر برای حذف مصنوعات فیزیولوژی و آنالیز مؤلفه‌های مستقل^۱ (ICA) برای حذف مصنوعات غیر فیزیولوژی [۱۲].

۲-۳-۳ استخراج ویژگی

پس از مرحله پیش‌پردازش باید اطلاعات و ویژگی‌های مفید برای دسته‌بندی استخراج شود. هدف از این مرحله، استخراج ویژگی از سیگنال EEG به گونه‌ای که دارای معنا و مفهوم مناسب باشد و به مرحله دسته‌بندی کمک نماید. روش‌های مختلفی برای استخراج ویژگی وجود دارد مانند ویژگی‌های آماری (میانگین، آنتروپی و غیره) در حوزه زمان و ویژگی‌هایی مانند توان باندهای متفاوت (آلفا، بتا، تتا، گاما، دلتا و میو) در حوزه فرکانس.

۲-۳-۴ دسته‌بندی

پس از انجام پیش‌پردازش و استخراج ویژگی مرحله دسته‌بندی می‌باشد. هر چه ویژگی‌ها و دسته‌بند انتخاب شده مناسب‌تر باشند عملکرد سیستم BCI مناسب‌تر می‌باشد. نقش دسته‌بند شناسایی الگوها از داده‌های آموزشی است که بتواند در مرحله آزمون موفقیت‌آمیز عمل نماید. یک دسته‌بند باید سه ویژگی مهم را داشته باشد؛ دقت بالا، کارایی محاسباتی مناسب و پیاده‌سازی در دنیای واقعی. برخی دسته‌بندهای مورد استفاده برای فعالیتهای ذهنی عبارتند از: نزدیکترین همسایه^۲ (KNN)، ماشین بردار پشتیبان^۳ (SVM)، شبکه عصبی مصنوعی^۴ (ANN) [۱۶]–[۱۳].

^۱ Independent component analysis.

^۲ Support vector machine.

^۳ K-nearest neighbor.

^۴ Artificial neural network.

۲-۴ پژوهش‌های مرتبط

در این بخش دسته‌بندی‌های مورد استفاده در کارهای اخیر انجام شده در زمینه BCI مورد بررسی قرار گرفته است. برخی از دسته‌بندی‌ها که به طور وسیعی در این زمینه مورد استفاده قرار گرفته‌اند عبارتند از: ماشین بردار پشتیبان، شبکه عصبی مصنوعی، آنالیز افتراقی خطی و k نزدیکترین همسایه. همچنین برخی از روش‌هایی که از نمایش تنک و یادگیری دیکشنری برای دسته‌بندی سیگنال EEG استفاده نموده‌اند بررسی شده است. در زیر کارهای انجام شده با استفاده از دسته‌بندی ذکر شده و نمایش تنک آورده شده است.

- ماشین بردار پشتیبان

یکی از دسته‌بندی‌های پرکاربرد در تحقیقات BCI ماشین بردار پشتیبان یا SVM می‌باشد. SVM با فرض اینکه داده‌ها به صورت خطی جداپذیر باشند، ابرصفحه‌هایی با حداکثر حاشیه برای جداسازی دسته‌ها به دست می‌آورد. برای مسائل غیرخطی داده‌ها باید به فضای دیگر نگاشت داده شوند یا در همین فضا آنها را باید با کرنل‌های دیگر دسته‌بندی نمود. [۱۷]

در [۱۸] سه فعالیت از پایگاه داده دانشگاه کلرادو را برای دسته‌بندی استفاده نموده‌اند. محققین در این مقاله از روش ICA برای حذف مصنوعات پلک زدن استفاده نموده‌اند. از پارامترهای توان باند و رگرسیون خودکار^۱ (AR) برای استخراج ویژگی و در پایان برای دسته‌بندی از دسته‌بند SVM با کرنل تابع پایه‌ای شعاعی^۲ (RBF) استفاده شده است. بهترین نتیجه به دست آمده در این مقاله ۷۰٪ می‌باشد.

در [۱۹] از ترکیب کرنل‌های مختلف برای ساخت دسته‌بند SVM استفاده شده است. در این مقاله نیز از پایگاه داده دانشگاه کلرادو استفاده شده است. برای پیش‌پردازش و حذف مصنوعات غیر فیزیولوژی،

^۱ Auto regressive.

^۲ Radial basis function.

سیگنال EEG توسط فیلتر میان‌گذر ۰/۱ تا ۶۴ عبور داده شده است. ویژگی آنتروپی برای باندهای مختلف دلتا، تتا، آلفا و بتا به دست آمده از تبدیل موجک استخراج شده است. نتایج به دست آمده بیانگر کارایی مناسب دسته‌بند SVM با چندین کرنل می‌باشد.

در پژوهشی دیگر [۲۰] از سه کاربر برای ثبت سیگنال در محیط آزمایشگاهی استفاده شده است. برای ثبت سیگنال ۱۲ فعالیت ذهنی شامل تصور ذهنی حرکت چپ و راست و غیره در نظر گرفته شده است. نتایج بر مبنای ترکیب پنج فعالیت ذهنی که بیشترین دقت دسته‌بندی را دارند ارائه شده است. برای پیش‌پردازش از فیلتر میان‌گذر ۵ تا ۴۰ هرتز که حاوی اطلاعات مفیدی برای دسته‌بندی فعالیت‌های ذهنی است؛ استفاده شده است. برای استخراج ویژگی از توزیع توان طیفی^۱ (PSD) و برای دسته‌بندی از SVM با کرنل RBF استفاده شده است.

در [۱۴] تحلیلی به منظور انتخاب بهترین الکترودها برای افزایش دقت دسته‌بند انجام شده است. انتخاب بهترین الکترودها نقش مؤثری در عملکرد سیستم BCI دارد. هر چه تعداد الکترودها انتخابی کمتر باشد منجر به کاهش محاسبات و در نتیجه افزایش سرعت سیستم می‌شود. در این مقاله از ۱۲ فعالیت ذهنی مانند حرکت دست و پای چپ و راست، شمارش معکوس و غیره استفاده شده است. داده‌های به دست آمده برای فعالیت‌های ذهنی از ۱۶ کانال می‌باشد. به منظور دریافت سیگنال‌های که سهم مؤثری در فعالیت‌های مغز دارند، فیلتر میان‌گذر ۵ تا ۴۰ هرتزی اعمال می‌شود. برای بهبود سیگنال به دست آمده از فیلتر لاپلاس استفاده شده است. برای استخراج ویژگی نیز مقدار PSD را برای سیگنال‌های ۸ تا ۳۶ هرتز را محاسبه نموده‌اند. برای کلاس‌بندی نیز از SVM با کرنل RBF استفاده شده است.

- شبکه عصبی مصنوعی

شبکه عصبی مصنوعی از دسته‌بندهای پرکاربرد در تحقیقات انجام شده در زمینه BCI می‌باشد. شبکه

^۱ Power spectral density.

عصبی چند لایه یا MLP معروفترین شبکه عصبی مصنوعی می باشد که در این زمینه استفاده شده است. این شبکه دارای دو مد می باشد. مد اول، مد جلورونده^۱ می باشد که شامل ارائه نمودن الگوها توسط نودهای ورودی و انتقال سیگنال ها از طریق شبکه به منظور گرفتن خروجی می باشد. یادگیری مد دوم شبکه می باشد که شامل اصلاح پارامترهای شبکه (وزن ها) برای کاهش فاصله خروجی مطلوب و خروجی محاسبه شده می باشد. [۱۷]

محققین در [۲۱] روشی بر مبنای مدل مخفی مارکوف^۲ (HMM) برای استخراج ویژگی برای دسته بندی سه فعالیت ذهنی ارائه نموده اند. پایگاه داده استفاده شده در این مقاله شامل سیگنال های ثبت شده برای حرکت دست راست، چپ و گفتن کلمه می باشد. بردار ویژگی تولید شده با محاسبه طیف انرژی بر روی سیگنال ورودی به دست می آید. برای کاهش بعد ویژگی های به دست آمده از HMM استفاده شده است. مراحل HMM برای کاهش ویژگی به صورت زیر می باشد.

- نرمال سازی مقادیر بردار ویژگی بین بازه صفر تا یک.
- تعریف نمودن حالت هایی برای مدل مخفی مارکوف در بازه صفر تا یک.
- ماتریس انتقال برای هر ویژگی استخراج شده محاسبه می شود. فرمول ماتریس انتقال به صورت زیر می باشد.

$$\text{ماتریس انتقال} = \frac{\text{تعداد تغییرات از } S_i \text{ به } S_j \text{ در دنباله}}{\text{طول دنباله}}$$

که در آن S_i و S_j به ترتیب بیانگر حالت i و حالت j در دنباله می باشند. پس از اعمال مدل مخفی مارکوف بر روی هر بردار ویژگی، ماتریس انتقال محاسبه می شود. از ماتریس انتقال برای ورودی دسته بند استفاده می شود که به مراتب طول کمتری نسبت به بردار ویژگی استخراج شده برای سیگنال دارد.

^۱ Feed forward.

^۲ Hidden markov model.

برای دسته‌بندی کردن سیگنال ورودی نیز از شبکه عصبی MLP استفاده شده است. نتایج مقاله نشان‌دهنده بهتر بودن روش HMM نسبت به PCA برای کاهش ویژگی‌ها می‌باشد.

در [۲۲] برای دسته‌بندی پنج فعالیت ذهنی پایگاه داده دانشگاه کلرادو از شبکه عصبی استفاده شده است. در این مقاله دو روش مورد بررسی قرار گرفته است:

- روش اول: دسته‌بندی با استفاده از شبکه عصبی
 - روش دوم: دسته‌بندی با استفاده از شبکه عصبی و ورودی به دست آمده از PCA
- نتایج این مقاله بیانگر این موضوع است که پیش‌پردازش روی سیگنال اصلی نقش مهمی در دسته‌بندی سیگنال دارد و روش دوم به مراتب بهتر از روش اول سیگنال‌ها را دسته‌بندی می‌نماید.
- در [۲۳] روشی بر مبنای شبکه عصبی خود سازمانده^۱ (SOM) برای دسته‌بندی سیگنال‌های EEG حرکت دست راست و چپ و نوشتن کلمه ارائه شده است. در این مقاله از PSD به عنوان ویژگی برای دسته‌بند شبکه عصبی SOM استفاده شده است.

- K نزدیکترین همسایه

یکی دیگر از دسته‌بندهای مورد استفاده در زمینه دسته‌بندی سیگنال EEG، دسته‌بند K نزدیکترین همسایه یا K-NN می‌باشد. این دسته‌بند غیرخطی کلاس مربوط به داده آزمون را برابر با کلاس K نزدیکترین داده آموزشی به داده آزمون قرار می‌دهد. این دسته‌بند در کارهای BCI که دارای ابعاد ویژگی زیادی برای دسته‌بندی می‌باشند مناسب نمی‌باشد ولی کارایی مناسبی برای دسته‌بندی داده‌های با ابعاد پایین دارد. برای محاسبه فاصله داده آزمون از داده‌های آموزشی معیارهای فاصله اقلیدسی^۲، چبیشف^۳، ماهالانوبیس^۴ و غیره استفاده می‌شود [۱۷].

^۱ Self organizing map.

^۲ Chebychev.

^۳ Euclidean.

^۴ Mahalanobis.

تحلیل مجزا ساز خطی یا LDA از یک ابرصفحه برای جداسازی داده‌های کلاس‌های مختلف استفاده می‌کند. LDA یک توزیع نرمال برای داده‌های هر کلاس در نظر می‌گیرد. ابر صفحه جداکننده به گونه‌ای محاسبه می‌شود که نرخ واریانس بین کلاسی به واریانس درون کلاسی را بیشینه نماید؛ در نتیجه بیشترین تفکیک‌پذیری بین کلاس‌ها را تضمین می‌نماید. با توجه به پیچیدگی محاسباتی کم LDA، این دسته‌بند برای BCI های آنلاین مناسب می‌باشد. از طرف دیگر مشکل اصلی خطی بودن آن است که نتایج نامناسبی را برای دسته‌بندی EEG های غیرخطی پیچیده دارد [۱۷].

در پژوهشی دیگر [۲۴] یک سیستم برای دسته‌بندی سیگنال EEG پایگاه داده برلین با استفاده از تجزیه مد تجربی^۱ (EMD) مطرح شده است. پس از به دست آوردن مدهای ذاتی^۲ (IMF) برای هر سیگنال، آنتروپی آن محاسبه شده است و با استفاده از الگوریتم آنالیز مؤلفه‌های اصلی^۳ (PCA) کاهش ویژگی انجام شده است. از آنالیز تفکیک‌کننده خطی^۴ (LDA) برای دسته‌بندی سیگنال‌های EEG استفاده شده است.

در تحقیقی دیگر [۱۶] از روش تبدیل Stockwell برای دسته‌بندی پنج فعالیت ذهنی پایگاه داده کلرادو استفاده شده است. تبدیل Stockwell مزایای تبدیل موجک در مناسب بودن رزولوشن و مزایای تبدیل فوریه در نگه‌داشتن فرکانس‌های محلی را دارا می‌باشد. سیگنال ورودی بدون حذف نویز با استفاده از تبدیل stokwell تجزیه می‌شود. خروجی تبدیل stockwell میزان فرکانس‌های مختلف در بازه‌های زمانی را نشان می‌دهد. با توجه به خروجی این تبدیل، باندهای دلتا، تتا، آلفا، بتا و گاما به راحتی استخراج می‌شوند. برای هر باند میانگین انحراف معیار به عنوان ویژگی در نظر گرفته می‌شود. در این

^۱ Empirical mode decomposition.

^۲ Principal component analysis.

^۳ Intrinsic mode functions.

^۴ Linear discriminant analysis.

مقاله از دو کلاسه‌بند LDA و KNN استفاده شده است. کلاسه‌بند LDA نتایج بهتری را نسبت به KNN دارد.

- نمایش تنک

استفاده از نمایش تنک برای دسته‌بندی در سال‌های اخیر مورد توجه بسیاری قرار گرفته است. تنک بودن بیانگر این موضوع است که داده‌ها را می‌توان به شکل ترکیب خطی تعداد کمی از پایه‌ها که از قبل در نظر گرفته شده بیان نمود. توضیحات بیشتر در بخش ۳ - ۳ آورده شده است.

محققین در [۲۵] یک روش برای جداسازی سیگنال EEG بر مبنای نمایش تنک مطرح نموده‌اند. از پایگاه داده دانشگاه برلین برای ارزیابی روش استفاده شده است. در این مقاله یک روش یادگیری دیکشنری بر مبنای ساخت یک دیکشنری با اندازه کوچکتر و قدرت تفکیک‌پذیری بالاتر ارائه شده است. برای جداسازی هر چه بهتر کلاس‌ها از CSP استفاده شده است.

محققین در [۲۶] با طراحی یک دیکشنری با استفاده از الگوریتم CSP به دسته‌بندی سیگنال EEG پرداخته‌اند که نتایج آن در مقایسه با دسته‌بند LDA بهتر بوده است. در این مقاله از روش جستجوی پایه^۱ (BP) برای محاسبه پاسخ تنک استفاده شده است. پاسخ تنک محاسبه شده برای دسته‌بندی استفاده شده است.

جمع‌بندی:

استفاده از LDA به دلیل پیچیدگی محاسباتی اندک برای سیستم‌های BCI آنلاین مناسب می‌باشد. به همین دلیل از LDA در کارهای زیادی در زمینه BCI برای دسته‌بندی سیگنال EEG استفاده شده است. مشکل اصلی LDA خطی بودن آن است که برای برخی از داده‌های EEG که پیچیده و غیرخطی

^۱ Basis pursuit.

می‌باشند نامناسب است [۱۷].

دسته‌بند KNN حساس به ابعاد ویژگی‌ها می‌باشد به طوری که برای بردار ویژگی‌های با ابعاد زیاد نتایج مطلوبی ارائه نمی‌دهد. از دسته‌بند KNN برای دسته‌بندی ویژگی‌های با ابعاد کم استفاده شده و عملکرد خوب آن نشان داده شده است. از KNN در زمینه BCI به خاطر ابعاد بالای بردار ویژگی کمتر استفاده شده است. از دیگر مسائل مربوط به دسته‌بند KNN می‌توان انتخاب مناسب k و نحوه محاسبه فاصله را نام برد [۱۷].

شبکه عصبی MLP هنگامی که از تعداد نوروں و لایه به اندازه کافی تشکیل شده باشد قادر به تخمین هر تابعی می‌باشد. شبکه عصبی انعطاف‌پذیر می‌باشد و برای دسته‌بندی در زمینه BCI مورد استفاده قرار گرفته‌اند. از مشکلات این دسته‌بند می‌توان به حساس بودن به آموزش بیش از حد^۱ اشاره نمود که برای داده‌های EEG که نویزی و غیر ثابت می‌باشند مشکل‌ساز خواهد شد. همچنین شبکه عصبی نیازمند ساختار مناسبی برای دسته‌بندی می‌باشد [۱۷].

SVM یکی از پرکاربردترین دسته‌بندها در زمینه تحقیقات BCI می‌باشد. معیارهای حداکثر حاشیه و پارامتر تنظیم‌کننده^۲ دلیلی برای بهبود عملکرد SVM برای حل مشکلات داده‌های با ابعاد زیاد و آموزش بیش از حد می‌باشد. این دسته‌بند به تعداد پارامتر اندکی وابسته می‌باشد؛ نوع کرنل و انتخاب پارامتر تنظیم‌کننده. از معایب دسته‌بند SVM می‌توان به زمان اجرا زیاد در مقایسه با سایر دسته‌بندها اشاره نمود [۱۷].

در سال‌های اخیر استفاده از نمایش تنک و یادگیری دیکشنری برای دسته‌بندی سیگنال‌های EEG مورد استفاده قرار گرفته‌اند. ماهیت تنک بودن سیگنال‌های طبیعی می‌تواند به دسته‌بندی کمک نماید و این امر منجر به بهبود عملکرد این روش‌ها می‌شود. در این روش‌ها برای تجزیه سیگنال به اتم‌های

^۱ Over training.

^۲ Regulator.

تشکیل دهنده آن از دیکشنری ساخته شده بر اساس داده‌های آموزشی استفاده می‌شود؛ این روش کارایی مناسب‌تری نسبت به دیکشنری‌های تحلیلی از پیش تعریف شده مانند تبدیل موجک و تبدیل فوریه دارد.

در این پایان‌نامه با استفاده از نمایش تنک و یادگیری دیکشنری روش‌هایی برای دسته‌بندی سیگنال EEG ارائه شده است که در فصل‌های آتی شرح و نتایج آن آورده شده است.

فصل ۳

روش پیشنهادی

در این فصل روش دسته‌بندی سیگنال EEG با استفاده از نمایش تنک و یادگیری دیکشنری مطرح شده است. در بخش اول این فصل به معرفی پایگاه داده‌های مورد استفاده برای ارزیابی روش پرداخته شده است. در بخش دوم پیش‌پردازش انجام شده توضیح داده شده است. در بخش سوم به معرفی نمایش تنک و یادگیری دیکشنری پرداخته شده است. در پایان دو روش پیشنهادی برای دسته‌بندی سیگنال EEG ارائه شده است.

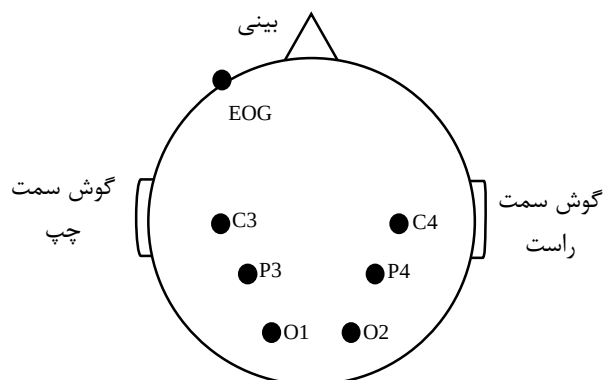
در روش پیشنهادی اول یک دسته‌بند بر مبنای نمایش تنک برای دسته‌بندی سیگنال‌های EEG ارائه شده است. این دسته‌بند، یک دسته‌بند قابل تعمیم برای مسئل چند کلاسه می‌باشد. در این روش برای انتخاب کانال از الگوریتم CSP استفاده شده است. از آنجا که الگوریتم CSP برای مسائل دو کلاسه مطرح شده است؛ برای بسط آن به یک مسئله چند کلاسه از الگوریتم CSP چند کلاسه استفاده شده است. الگوریتم CSP چند کلاسه در بخش ۳ - ۲ - ۱ - ۲ توضیح داده شده است. در این فصل روش‌های محاسبه پاسخ تنک بررسی شده است که در روش پیشنهادی نیز استفاده شده است.

برای بهبود عملکرد دیکشنری ساخته شده از روش‌های یادگیری دیکشنری استفاده می‌شود. اکثر روش‌های یادگیری دیکشنری بر مبنای محاسبه پاسخ تنک می‌باشند. هر چند که برای به دست آوردن پاسخ تنک روش‌های متعددی وجود دارد؛ ولی این روش‌ها غیر خطی هستند و بار محاسباتی برای سیستم ایجاد می‌کنند. در روش پیشنهادی دوم با استفاده از ساخت یک دیکشنری تحلیلی، نیاز به روش‌های محاسبه تنک برطرف شده است. در این روش دو دیکشنری؛ یکی تحلیلی برای محاسبه پاسخ تنک و دیگری مصنوعی برای بازسازی داده اصلی به طور همزمان آموزش دیده و در مرحله آزمون مورد استفاده قرار گرفته شده است. برای ارزیابی دو روش از دو پایگاه داده که به طور گسترده‌ای در زمینه سیگنال EEG مورد استفاده قرار گرفته‌اند استفاده شده است که در این فصل جزئیات آنها آورده شده است.

۳-۱ معرفی داده‌های موجود

۳-۱-۱ پایگاه داده دانشگاه کلرادو

این داده‌ها مطابق استاندارد ۲۰-۱۰ از کانال‌های C3، C4، P3، P4، O1 و O2 همراه با یک کانال EOG توسط گروه Aunon و Keirn ثبت شده است [۲۷]. هر سیگنال به مدت ۱۰ ثانیه و با نرخ نمونه‌برداری ۲۵۰ هرتز ثبت گردیده است. در شکل ۳-۱ مکان الکترودهای استفاده شده در این پایگاه داده بر روی سر آورده شده است. نمونه‌ای از سیگنال‌های ثبت شده توسط این الکترودها در شکل ۳-۲ آورده شده است.



شکل ۳-۱. مکان الکترودهای استفاده شده در پایگاه داده دانشگاه کلرادو [۲۸]

پایگاه داده از هفت نفر فرد سالم برای انجام پنج فعالیت ذهنی به دست آمده است. در ثبت سیگنال‌ها از افراد خواسته شده است که یکی از فعالیت‌های ذهنی زیر را بدون هیچ‌گونه فعالیت فیزیکی یا گفتاری انجام دهند. این پنج فعالیت عبارتند از:

- حالت استراحت^۱: در این حالت از فرد خواسته می‌شود بدون تفکر در رابطه با مسئله‌ای خاص، آرامش خود را حفظ نماید.

^۱ Baseline task.

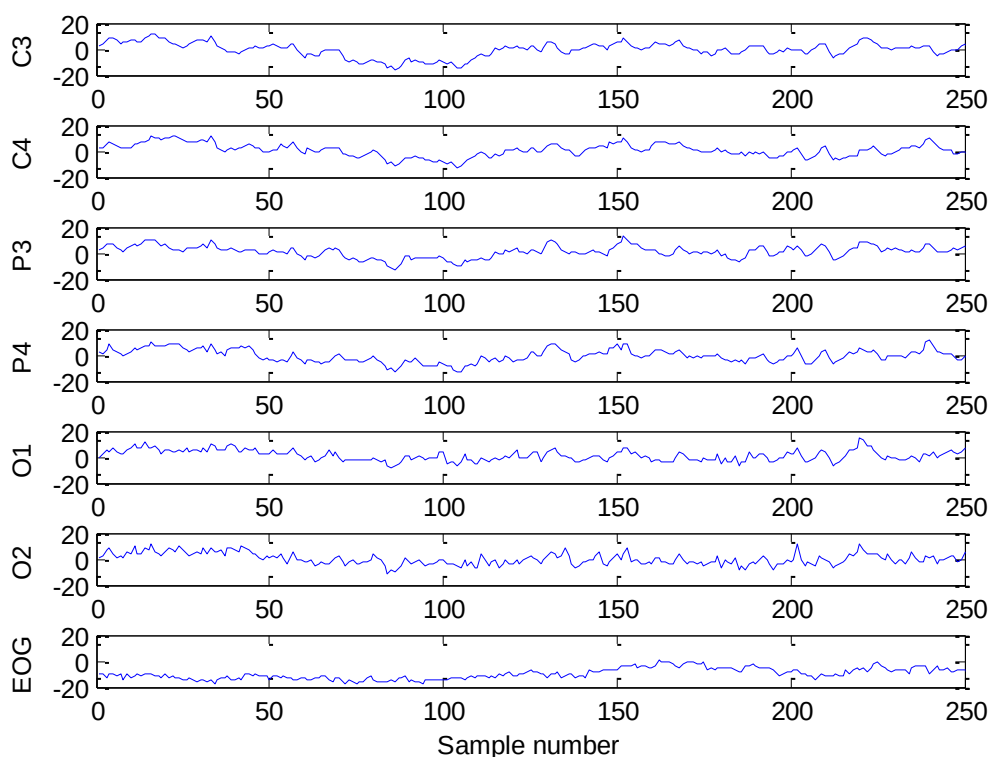
- عمل ضرب ذهنی^۱: از فرد مسئله ضرب ساده‌ای مانند ۱۱×۱۲ پرسیده می‌شود و در این هنگام سیگنال ثبت می‌شود.
- دوران ذهنی یک شی هندسی^۲: در ابتدا به مدت ۳۰ ثانیه یک جسم سه بعدی به فرد نشان داده، سپس از وی خواسته می‌شود که به صورت ذهنی این شی را دوران دهد.
- نامه‌نگاری ذهنی^۳: برای ثبت سیگنال مربوط به این فعالیت؛ از فرد مورد آزمایش خواسته می‌شود که به صورت ذهنی و بدون تکلم به دوست خود نامه بنویسد.
- شمارش ذهنی همراه با تصویر ذهنی^۴: از فرد خواسته می‌شود که یک تخته تصور نماید و شمارش انجام دهد [۲۷].

^۱ Multiplication.

^۲ Mental Letter Composing.

^۳ Geometric figure rotation task.

^۴ Visual Counting.



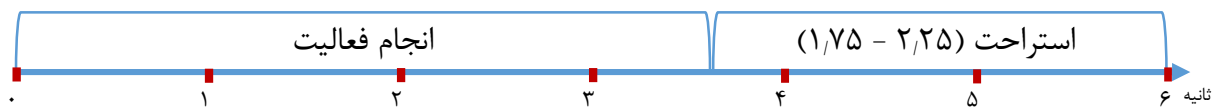
شکل ۳-۲. نمونه سیگنال EEG ثبت شده در پایگاه داده دانشگاه کلرادو

۳-۱-۲ پایگاه داده برلین

پایگاه داده برلین برای مسابقات BCI III ایجاد شده است و در تحقیقات زیادی مورد استفاده قرار گرفته است [۲۳]. این پایگاه داده شامل پنج مجموعه داده مجزای به دست آمده از پنج فرد سالم (av, al, aa, ay و aw) می‌باشد. افراد شرکت کننده در این آزمایش سه فعالیت ذهنی حرکت دست چپ، حرکت دست راست و حرکت پای راست را انجام داده‌اند اما پایگاه داده موجود فقط شامل حرکت دست راست و پای راست می‌باشد. این پایگاه داده با استفاده از تقویت کننده BrainAmp از ۱۱۸ الکترود بر اساس سیستم ۲۰-۱۰ ضبط شده است. سیگنال‌ها از یک فیلتر میان‌گذر ۰/۵ تا ۲۰۰ هرتز عبور داده شده‌اند. سیگنال‌ها با نرخ نمونه‌برداری ۱۰۰۰ هرتز ثبت شده است.

شکل ۳-۳ زمان‌بندی ثبت یک آزمایش پایگاه داده برلین را نشان می‌دهد. در حین آزمایش، فرد بر روی

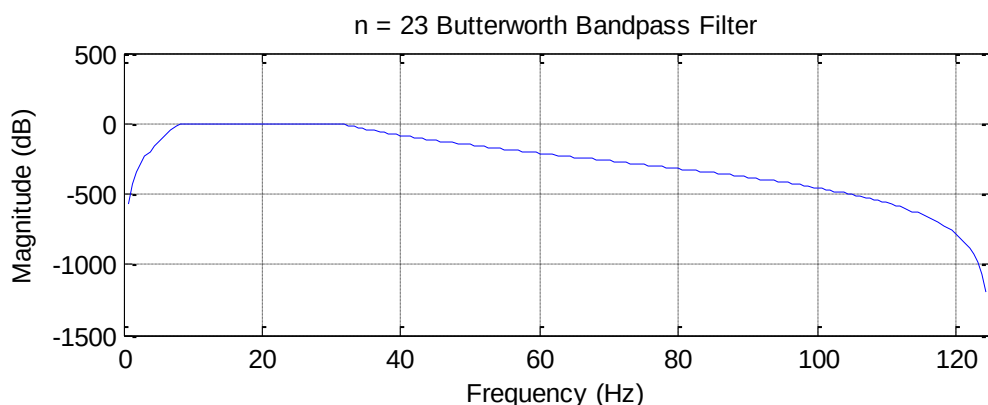
یک صندلی به راحتی می‌نشیند. در مدت ۳/۵ ثانیه، علامت تصویری برای نشان دادن حرکت مورد نظر (حرکت دست چپ، حرکت دست راست و حرکت پای راست) به فرد نمایش داده می‌شود. سپس فرد به مدت ۱/۷۵ تا ۲/۲۵ ثانیه استراحت می‌نماید. برای هر فرد ۱۴۰ داده برای هر کلاس ضبط و ثبت گردیده است [۲۹], [۳۰].



شکل ۳-۳. زمان‌بندی انجام ثبت یک آزمایش پایگاه داده برلین

۲ - ۳ پیش‌پردازش

همانطور که در بخش ۲ - ۳ - ۲ بیان شد، سیگنال EEG آغشته به نویز و مصنوعات می‌باشد که منشأ غیر از مغز دارند. در این پایان‌نامه با استفاده از فیلتر میان‌گذر سعی در بهبود کیفیت سیگنال و حذف مصنوعات سیگنال EEG شده است. برای پایگاه داده کلرادو فیلتر میان‌گذر ۸-۳۱ هرتز و برای پایگاه داده برلین فیلتر میان‌گذر ۴-۴۵ هرتز اعمال شده است. برای طراحی فیلتر میان‌گذر از فیلتر باتروث استفاده شده است. شکل ۴-۳ فیلتر باتروث استفاده شده برای عبور فرکانس‌های ۸-۳۱ هرتز نشان می‌دهد. ریبیل باند گذر استفاده شده ۳ دسیبل (dB) و برای باند حذف ۴۰ dB در نظر گرفته شده است. مرتبه به دست آمده برای این فیلتر ۲۳ می‌باشد.



شکل ۳-۴. فیلتر میان‌گذر باتروورث

۳-۲-۱ استفاده از الگوهای مکانی مشترک^۱ برای انتخاب کانال

تحقیقات انجام شده در زمینه کاهش ابعاد داده (تعداد کانال‌های EEG) نشان می‌دهد که با کاهش تعداد کانال می‌توان دقت سیستم را بهبود داد. کاهش ابعاد منجر به حذف کانال‌های نویزی و نامرتبط با فعالیت ذهنی می‌شود. همچنین با تعداد کانال کمتر بار محاسباتی نیز کاهش و راحتی کاربر افزایش می‌یابد [۳۱].

یکی از روش‌های موفق برای این منظور روش CSP می‌باشد. CSP یک روش تجزیه‌سازی می‌باشد که سیگنال‌های EEG را به دو ماتریس تجزیه می‌نماید. ماتریس اول شامل سیگنال‌های منبع می‌باشد و ماتریس دیگر شامل وزن‌هایی می‌باشد که به ماتریس اول نسبت داده شده است. در برخی از کارها به طور مستقیم از ضرایب CSP کانال‌های مرتبط را انتخاب می‌نمایند [۳۲].

در این بخش به معرفی الگوریتم CSP برای دو کلاس و چند کلاسه پرداخته شده است که در این پایان‌نامه مورد استفاده قرار گرفته است.

^۱ Common spatial pattern.

۳- ۲- ۱- الگوهای مکانی مشترک دو کلاسه

الگوریتم CSP در تحقیقات گسترده‌ای در زمینه BCI مورد استفاده قرار گرفته است. هدف از روش CSP یافتن فیلترهایی است که پس از اعمال آنها بر روی داده‌ها، بیشترین اختلاف از لحاظ واریانس را برای داده‌های کلاس‌های مختلف ایجاد نماید. جزئیات این روش با مثال در ادامه توضیح داده شده است. فرض کنید $X_L \in \mathbb{R}^{N \times T}$ و $X_R \in \mathbb{R}^{N \times T}$ بیانگر ماتریس داده‌های EEG از دو کلاس تصور ذهنی حرکت دست چپ و راست باشند که در آن N تعداد کانال‌های EEG و T تعداد نمونه‌ها برای هر کانال باشد. کواریانس مکانی نرمال شده برای این دو کلاس به صورت معادله ۱-۳ بیان می‌شود.

$$R_L = \frac{X_L X_L^T}{\text{trace}(X_L X_L^T)} \quad 1-3$$

$$R_R = \frac{X_R X_R^T}{\text{trace}(X_R X_R^T)}$$

X^T بیانگر ترانهاده ماتریس X و $\text{trace}(A)$ مجموع عناصر قطر اصلی ماتریس A را محاسبه می‌نماید. میانگین کوواریانس‌های مکانی $\overline{R}_L$ و $\overline{R}_R$ با محاسبه میانگین روی همه نمونه‌های مربوط به هر کلاس به دست می‌آید. ماتریس کواریانس مرکب برای دو کلاس از رابطه ۲-۳ به دست می‌آید.

$$R = \overline{R}_L + \overline{R}_R = U_0 \Sigma U_0^T \quad 2-3$$

در این معادله ماتریس R به $U_0 \Sigma U_0^T$ تجزیه می‌شود به گونه‌ای که ستون‌های U را بردارهای ویژه تشکیل می‌دهند و Σ ماتریس قطری از مقادیر ویژه مربوطه می‌باشد. مقادیر روی قطر اصلی ماتریس Σ به صورت نزولی مرتب شده‌اند. اگر ماتریس سفیدکننده P به صورت زیر تعریف شود:

$$P = \Sigma^{-\frac{1}{2}} U_0^T \quad 3-3$$

این تبدیل می‌تواند نمونه‌ها را به فضایی با بردارهای پایه U نگاشت دهد. انتقال ماتریس کواریانس

میانگین به صورت رابطه ۴-۳ انجام می‌شود.

$$S_L = P \overline{R_L} P^T \quad ۴-۳$$

$$S_R = P \overline{R_R} P^T$$

S_R و S_L بردارهای ویژه مشترک دارند که مجموع مقادیر ویژه متناظر برای آن دو برابر یک می‌باشد.

$$S_L = U \Sigma_L U^T \quad ۵-۳$$

$$S_R = U \Sigma_R U^T$$

$$\Sigma_L + \Sigma_R = I$$

بردارهای ویژه با بزرگترین مقادیر ویژه برای S_L شامل کوچکترین مقادیر ویژه برای S_R می‌باشد و بالعکس؛ بنابراین تصویر انتقال یافته به داده‌ها با استفاده از ماتریس سفیدکننده P متناسب برای دسته‌بندی دو کلاس می‌باشد. ماتریس انتقال W به صورت معادله ۶-۳ بیان می‌شود:

$$W = U^T P \quad ۶-۳$$

با استفاده از ماتریس انتقال W نگاشت داده‌های X به صورت رابطه ۷-۳ محاسبه می‌شود.

$$Z = WX \quad ۷-۳$$

Z بیانگر مؤلفه‌های منابع سیگنال EEG شامل مؤلفه‌های مشترک و مخصوص برای کارهای مختلف می‌باشد. سیگنال اصلی X با استفاده از فرمول ۸-۳ بازسازی می‌شود.

$$X = W^{-1} Z \quad ۸-۳$$

در این فرمول W^{-1} معکوس ماتریس W می‌باشد. ستون‌های ماتریس W^{-1} الگوهای مکانی می‌باشند. اولین و آخرین ستون W^{-1} مهمترین الگوهای مکانی می‌باشند که بیانگر بیشترین واریانس از یک کلاس و کوچکترین واریانس کلاس دیگر می‌باشد [۳۲].

۳-۲-۱ الگوهای مکانی مشترک چند کلاسه

در بخش قبلی الگوریتم CSP برای دو کلاس مورد بررسی قرار داده شد. برای دو کلاس، هدف الگوریتم CSP پیدا کردن ماتریس انتقال W به نحوی که بهره اطلاعاتی بین کلاس‌های مختلف را بیشینه نماید. این موضوع در معادله ۹-۳ آورده شده است.

$$W = \operatorname{argmax}_W \left\{ \frac{W^T \operatorname{cov}(x1) W}{W^T \operatorname{cov}(x2) W} \right\} \quad 9-3$$

که در آن $X1$ و $X2$ دو کلاس مختلف و $\operatorname{cov}(A)$ بیانگر ماتریس کواریانس A می‌باشد.

یکی از محدودیت‌های روش CSP ارائه آن برای مسائل دو کلاسه می‌باشد. به همین منظور تلاش‌های زیادی برای تعمیم آن به یک مسئله چند کلاسه مطرح شده است. در اکثر روش‌های مطرح شده، مسئله چند کلاسه به چند مسئله دو کلاسه تبدیل شده است. روش یک کلاس نسبت به سایر کلاس‌ها^۱ (OVR) ویژگی‌های CSP برای یک کلاس را نسبت به سایر کلاس‌ها مقایسه می‌کند؛ بنابراین برای یک مسئله چهار کلاسه به چهار کلاسه‌بند OVR نیاز می‌باشد که هر کلاسه‌بند ویژگی‌های یک کلاس را نسبت به سایر کلاس‌ها دسته‌بندی می‌کند [۳۳].

روش دیگر ارائه شده برای محاسبه CSP چند کلاسه روش pair-wise می‌باشد که ویژگی‌های CSP را برای هر جفت کلاس موجود محاسبه می‌نماید. به عنوان مثال برای یک مسئله چهار کلاسه به $\binom{4}{2} = 6$ کلاسه‌بند نیاز می‌باشد. ورودی توسط همه کلاسه‌بندهای طراحی شده دسته‌بندی می‌شود و در نهایت کلاس داده‌های ورودی با بیشترین رأی مشخص می‌شود. روش تقسیم و غلبه^۲ رویکرد مشابه با روش OVR دارد اما به اندازه یکی کمتر از تعداد کلاس‌ها ماتریس W داریم؛ بنابراین برای یک مسئله چهار کلاسه به سه کلاسه‌بند نیاز می‌باشد [۳۳].

^۱ One versus rest.

^۲ Divided and conquer.

در این پایان‌نامه از روش ارائه شده در [۳۴] استفاده شده است. الگوریتم کلی این روش در شکل ۵-۳ آورده شده است.

ورودی: ماتریس کواریانس برای کلاس‌ها

گام اول - محاسبه W با استفاده از الگوریتم FFDiag به نحوی که $W^T R W = D_{Ci}$

گام دوم - محاسبه بهره اطلاعاتی مربوط به فیلترهای مکانی به دست آمده

گام سوم - انتخاب L ستون از ماتریس W دارای بیشینه بهره اطلاعاتی

خروجی - ماتریس W با L های انتخاب شده

شکل ۵-۳. الگوریتم CSP چند کلاسه

در [۳۵] محققین روشی ارائه نموده‌اند که در آن در یک مسئله چند کلاسه، تمامی کلاس‌ها به طور مستقیم دسته‌بندی می‌شوند. در این روش فیلترهای فضایی محاسبه شده به طور همزمان، ماتریس کواریانس همه کلاس‌ها را قطری می‌سازد. الگوریتم سریع برای قطری‌سازی مشترک^۱ (FFDiag) یک روش تکراری برای حل مسئله بهینه‌سازی ۱۰-۳ می‌باشد.

$$W = \min_{W \in \mathbb{R}^{N \times N}} \sum_{k=1}^K \sum_{i \neq j} \left((WC^k W^T)_{ij} \right)^2 \quad ۱۰-۳$$

در این فرمول $\{C^1, C^2, \dots, C^k\} \in \mathbb{R}^{N \times N}$ یک ماتریس متقارن می‌باشد. هدف این معادله پیدا نمودن ماتریس انتقال W می‌باشد که ماتریس‌های مورد نظر را قطری نماید.

با استفاده از روش FFDiag ماتریس انتقال W (فیلترهای مکانی) محاسبه می‌شود. پس از محاسبه W ، باید فیلترهای مورد نظر را به گونه‌ای انتخاب نماییم که بهره اطلاعاتی را بین کلاس‌های مختلف را بیشینه نماید؛ بنابراین داریم:

$$W^* = \underset{W}{\operatorname{argmax}} \{I(c, f(x))\} \quad ۱۱-۳$$

^۱ Fast algorithm for joint diagonalization.

در این رابطه $f(x)$ تابع انتقال داده‌ها می‌باشد که برای الگوریتم CSP برابر با $W^T X$ می‌باشد. عبارت $I(c, f(x))$ نیز بهره اطلاعاتی بین کلاس‌های C و تابع $f(x)$ را محاسبه می‌نماید. برای محاسبه تابع I از فرمول بهره اطلاعاتی ۱۲-۳ استفاده شده است.

$$I(C, W^T X) = H(W^T X) - H(W^T X | C) \quad ۱۲-۳$$

که در آن $H(.)$ برابر با آنروپی شانون^۱ می‌باشد. با محاسبه فرمول ۱۲-۳ مقدار $I(c, f(x))$ به صورت زیر به دست خواهد آمد.

$$I(C, W^T X) = - \sum_{i=1}^M P(c_i) \times \log \sqrt{W^T R_{x|c_i} W} - \frac{3}{16} \left(\sum_{i=1}^M P(c_i) \times (W^T R_{x|c_i} W)^2 - 1 \right)^2 \quad ۱۳-۳$$

در این رابطه M تعداد کلاس‌ها مسئله و $P(c_i)$ احتمال کلاس c_i و $R_{x|c_i}$ ماتریس کواریانس x به شرط کلاس c_i می‌باشد؛ بنابراین این روش ابتدا با الگوریتم FFdiag قطری‌سازی را انجام داده و سپس برای مرتب نمودن W ها بر اساس میزان اهمیت با استفاده از الگوریتم ۱۳-۳ عمل می‌نماید.

۳ - ۳ نمایش تنک

در سال‌های اخیر نمایش تنک داده‌ها برای کاربردهایی مانند طبقه‌بندی سیگنال [۳۶]، [۲۶]، شناسایی چهره [۳۷]، فشرده‌سازی [۳۸] مورد استفاده قرار گرفته است. موفقیت نمایش تنک در این کاربردها از آنجا ناشی می‌شود که اکثر سیگنال‌های طبیعی مانند تصویر و صدا با در نظر گرفتن پایه‌های مشخصی دارای نمایش تنک می‌باشند.

^۱ Shannon entropy.

سیگنال $X_n = [x_1, x_2, \dots, x_n] \in \mathbb{R}^n$ را در نظر بگیرید. فرض تنک بودن بیانگر این موضوع است که داده‌ها را می‌توان به شکل ترکیب خطی تعداد کمی از پایه‌ها که از قبل در نظر گرفته شده بیان نمود. اگر تعداد این پایه‌ها را با k نمایش دهیم، این تعداد باید به مراتب از بعد فضای اصلی داده‌ها کوچکتر باشد. ($k \ll n$) پایه‌هایی مورد استفاده برای نمایش داده‌ها را در یک ماتریس قرار می‌دهیم که ماتریس واژه‌نامه یا ماتریس دیکشنری^۱ نامیده می‌شود. هر ستون از ماتریس دیکشنری اتم نامیده می‌شود. اگر تعداد اتم‌های دیکشنری به اندازه بعد فضای برداری داده باشد و پایه‌های کل فضا را پوشش دهد؛ ماتریس کامل^۲ نامیده می‌شود. در این حالت هر داده یک نمایش یکتا توسط اتم‌های دیکشنری خواهد داشت. برای مثال در تبدیل فوریه، پایه‌های تبدیل ثابت و عمود بر هم می‌باشند و هر داده ضریب خاص خود را خواهد داشت. نمایش هر داده با استفاده از ماتریس دیکشنری کامل یکتاست و لزوماً تنک نخواهد بود. اگر تعداد اتم‌های دیکشنری را از حالت کامل بیشتر در نظر بگیریم، ماتریس فراکامل^۳ نامیده می‌شود. مسئله به دست آوردن نمایش تنک، حل دستگاه معادله نامعین زیر می‌باشد.

$$y = Dx \quad ۱۴-۳$$

در این فرمول D نشان‌دهنده دیکشنری و x ترکیب خطی را نشان می‌دهد و y توسط ترکیب خطی اتم‌ها نوشته شده است. این دستگاه بی‌شمار جواب دارد. به دلیل اینکه هدف، یافتن تنک‌ترین پاسخ ممکن می‌باشد باید شروطی در نظر گرفته شود. راه‌حل کلی برای یافتن تنک‌ترین پاسخ در معادله ۱۵-۳ آورده شده است.

$$\min_x \|x\|_0 \text{ subject to } y = Dx \quad ۱۵-۳$$

هدف این معادله پیدا کردن تنک‌ترین پاسخ از میان پاسخ‌های موجود برای معادله ۱۴-۳ می‌باشد.

^۱ Dictionary matrix.

^۲ Over-complete.

^۳ Complete.

معادله ۱۵-۳ را می‌توان به صورت زیر بیان نمود.

$$\min_x \|y - Dx\| \text{ subject to } \|x\|_0 \leq k \quad ۱۶-۳$$

در این معادله عبارت $\|y - Dx\|$ بیانگر خطای بازسازی^۱ با استفاده از نمایش تنک می‌باشد. هدف این معادله کاهش میزان خطای بازسازی با قید میزان تنک بودن x می‌باشد. معادله ۱۶-۳ یک معادله غیر محدب می‌باشد و تنها راه حل کامل برای آن جستجوی کامل می‌باشد که پیچیدگی محاسباتی بالایی دارد و از لحاظ عملی قابل پیاده‌سازی نمی‌باشد.

۳ - ۳ - ۱ روش‌های محاسبه ضرایب تنک

همانطور که در گذشته به آن اشاره شد، معادله ۱۶-۳ یک مسئله غیر محدب NP-hard می‌باشد که با روش‌های جستجوی کامل قابل حل می‌باشد. به منظور به دست آوردن نمایش تنک روش‌های مختلفی ارائه شده است که در ادامه به آنها اشاره می‌کنیم.

۳ - ۳ - ۱ روش‌های نرم $L1$

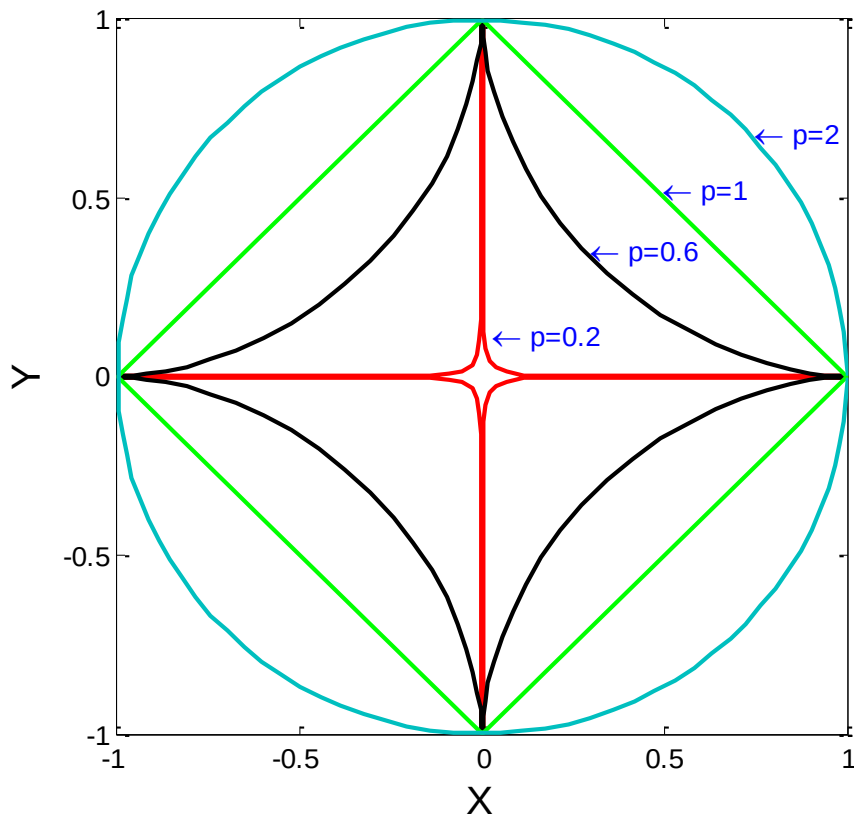
روش بهینه برای یافتن پاسخ تنک، استفاده از نرم کمینه $L0$ می‌باشد که در عمل ممکن نمی‌باشد. در شکل ۶-۳ نرم‌های مختلف برای یک بردار دو بعدی آورده شده است. همانطور که در شکل مشاهده می‌شود کوچکترین نرم محدب نرم $L1$ می‌باشد. به همین دلیل روش‌هایی برای پیدا کردن نمایش تنک بر مبنای نرم $L1$ مطرح شده است. از معروف‌ترین روش‌های این دسته می‌توان روش جستجوی پایه یا BP و LASSO را نام برد. روش BP برای حل معادله ۱۶-۳ از تقریب محدب نرم $L1$ به صورت زیر استفاده می‌نماید.

^۱ Reconstruction.

$$\min_x \|x\|_1 \text{ subject to } y=Dx$$

۱۷-۳

این دسته از الگوریتم‌ها جواب‌های نزدیک به نرم L_0 و حتی با ارضا کردن شرایطی پاسخی برابر با نرم L_0 ارائه می‌کنند. در مقایسه با سایر روش‌ها از سرعت پایین‌تری برخوردارند.



شکل ۳-۶. انواع نرم‌های مختلف برای یک بردار دو بعدی

۳-۳-۱ روش‌های حریص^۱

در این دسته از روش‌ها برای محاسبه پاسخ تنک در هر مرحله به طور حریصانه بهترین انتخاب ممکن را انجام می‌دهند. ساده‌ترین عضو این دسته روش جستجوی تطابق^۲ می‌باشد که در شکل ۳-۷ آورده شده است. شرط توقف را می‌توان خطای بازسازی و یا تعداد تکرار الگوریتم و معیار تعداد اتم‌های مورد

^۱ Greedy.

^۲ Matching pursuit.

استفاده در پاسخ تنک در نظر گرفت.

هدف: محاسبه پاسخ تنک سیگنال y با داشتن دیکشنری D
ورودی: y و D خروجی: x^k
گام اول - مقداردهی اولیه: $r^0=y$ و $x^0=0$ و $k = 1$
گام دوم - محاسبه همبستگی اتم‌ها: $C_i^k=D^T r^{k-1}$
گام سوم - انتخاب بهترین اتم: $i^k=\text{argmax}_i |C_i^k|$
گام چهارم - به‌روزرسانی نمایش تنک: $x_{i^k}^k=x_{i^k}^{k-1}+C_{i^k}^k$
گام پنجم - به‌روزرسانی باقیمانده: $r^k = r^{k-1} - C_{i^k}^k d_{i^k}$
گام ششم - $k = k + 1$
گام هفتم - بررسی شرط توقف

شکل ۷-۳. الگوریتم MP

الگوریتم MP به سبب نیاز به یک جستجوی ساده از سرعت خوبی برخوردار است اما به دلیل حریص بودن تضمینی برای رسیدن این الگوریتم به بهترین پاسخ تنک ممکن وجود ندارد. به منظور بهبود عملکرد MP، الگوریتم‌هایی مانند جستجوی تطابق متعامد^۱ (OMP) و CoSaMp مطرح شده است.

در روش OMP برای به‌روزرسانی نمایش تنک از اتم‌های انتخاب شده تا آن مرحله استفاده می‌شود؛ به عبارت دیگر در گام چهارم شکل ۷-۳ داریم $x^k=D_{\Psi^k}y$ که در آن Ψ حاوی اندیس‌های انتخاب شده تا مرحله k ام می‌باشد. در این روش برای به‌روزرسانی باقیمانده نیز از $r^k = r^{k-1} - Dx^k$ استفاده می‌شود محاسبات این الگوریتم نسبت به روش MP بیشتر می‌باشد اما عملکرد آن را بهبود می‌بخشد.

در روش MP و OMP در هر مرحله تنها یک اتم انتخاب می‌شود و به طور یقین تا پایان الگوریتم در پاسخ تنک حضور خواهد داشت. اگر اتم‌های انتخاب شده اشتباه باشند، اثر آن تا پایان الگوریتم وجود داشته و در نتیجه منجر به پاسخ‌های نادرست می‌شود. روش CoSaMP برای غلبه بر چنین مشکلی

^۱ Orthogonal matching pursuit.

ارائه شده است. این الگوریتم در گام‌های مختلف چندین اتم را انتخاب نموده و تعدادی از این اتم‌ها را برای به‌روزرسانی استفاده می‌نماید. [۳۹]

۳ - ۱ - ۳ روش‌های آستانه‌گذاری

این دسته از روش‌ها مانند روش‌های حریص می‌باشند و برای انتخاب تعداد اتم‌ها از معیار حد آستانه به جای معیار تعداد استفاده می‌شود؛ به عبارت دیگر در هر مرحله اتم‌هایی انتخاب می‌شوند که معیار حد آستانه را داشته باشند در غیر این صورت انتخاب نخواهند شد. از نماینده‌های معروف این دسته می‌توان به روش آستانه‌گذاری سخت تکراری^۱ (IHT) اشاره نمود. [۴۰]

۳ - ۱ - ۴ روش‌های تقریب نرم L0

روش‌های موجود در این دسته بر مبنای تقریب زدن نرم صفر با توابع دیگر شکل گرفته‌اند. در این روش‌ها در هر تکرار به جای کمینه نمودن نرم L0 تابع دیگری که یافتن حداقل آن ساده‌تر است مورد استفاده قرار می‌گیرد. در طول تکرارها، توابع مورد استفاده به سمت L0 نزدیک می‌شوند و در نتیجه جواب نهایی حداقل کننده نرم L0 خواهد بود. نرم صفر هموارشده^۲ (SL0) را می‌توان یکی از معروف‌ترین روش‌های این دسته دانست. [۴۱]

۴ - ۳ یادگیری دیکشنری

در سال‌های اخیر تحقیقات گسترده‌ای بر روی دسته‌بندی با استفاده از نمایش تنک انجام شده است.

^۱ Iterative hard thresholding.

^۲ Smoothed L0.

همانطور که در بخش قبلی به آن اشاره شد؛ مسئله نمایش تنک پیدا نمودن ترکیب تعداد کمی از اتم‌های دیکشنری برای سیگنال ورودی می‌باشد. کارایی نمایش تنک به کیفیت اتم‌های دیکشنری بستگی دارد. روش‌های یادگیری دیکشنری کیفیت اتم‌های دیکشنری را افزایش می‌دهد و یک دیکشنری متناسب با داده‌های ورودی ایجاد می‌نماید. در ادامه این بخش برخی از روش‌های یادگیری دیکشنری را مورد بررسی قرار می‌دهیم.

۳-۴-۱ روش K-SVD

روش K-SVD یک روش یادگیری دیکشنری با استفاده از تجزیه مقدارهای منفرد^۱ (SVD) می‌باشد. این روش تعمیم‌یافته الگوریتم k-means می‌باشد که یک روش تکراری برای بهبود کیفیت اتم‌های دیکشنری ارائه می‌دهد. این الگوریتم با حل معادله ۳-۱۸ یک دیکشنری برای بازسازی بهتر سیگنال ورودی و همچنین تنک بودن ضرایب به دست می‌آورد. با فرض اینکه $D \in \mathbb{R}^{n \times m}$ دیکشنری و $X \in \mathbb{R}^{m \times k}$ ضرایب تنک برای سیگنال ورودی $Y \in \mathbb{R}^{n \times k}$ باشد و پارامتر T میزان تنک بودن را نشان دهد؛ روش K-SVD به حل معادله ۳-۱۸ می‌پردازد.

$$\langle D, X \rangle = \underset{D, X}{\operatorname{argmin}} \|Y - DX\|_2^2 \text{ s.t } \forall i \|X_i\|_0 \leq T \quad ۳-۱۸$$

در این معادله عبارت $\|Y - DX\|$ بیانگر خطای بازسازی سیگنال با استفاده از دیکشنری ساخته شده می‌باشد. [۴۲]

^۱ Singular value decomposition.

۳ - ۴ - ۲ روش LC-KSVD

در [۴۳] با اضافه نمودن عبارت به الگوریتم K-SVD عملکرد آن را بهبود بخشیده‌اند. آنها یک تابع هدف جدید برای بهینه‌سازی بر مبنای الگوریتم K-SVD مطرح نمودند که در زیر آورده شده است.

$$\langle D, A, X \rangle = \underset{D, A, X}{\operatorname{argmin}} \|Y - DX\|_2^2 + \alpha \|Q - AX\|_2^2 \text{ s.t } \forall i \|X_i\|_0 \leq T \quad ۱۹-۳$$

عبارت $\|Y - DX\|$ خطای بازسازی سیگنال اصلی با استفاده از دیکشنری آموزش دیده را نشان می‌دهد. توانایی یک دسته‌بند به میزان افتراق‌پذیر^۱ بودن پاسخ تنک بستگی دارد. Q بیانگر بهینه‌ترین حالت ممکن برای پاسخ تنک می‌باشد و A یک ماتریس انتقال می‌باشد. عبارت $\|Q - AX\|$ میزان افتراق‌پذیر بودن پاسخ تنک را نشان می‌دهد. پارامتر α نیز میزان بازسازی و افتراق‌پذیری را کنترل می‌نماید. همچنین نویسندگان در این مقاله عبارت خطای دسته‌بندی را نیز به تابع هدف اضافه نمودند که در معادله زیر آورده شد است.

$$\langle D, W, A, X \rangle = \underset{D, A, X}{\operatorname{argmin}} \|Y - DX\|_2^2 + \alpha \|Q - AX\|_2^2 + \beta \|H - WX\|_2^2 \quad ۲۰-۳$$

$$\text{s.t } \forall i \|X_i\|_0 \leq T$$

معادله ۲۰-۳ ماتریس H شامل کلاس سیگنال‌های ورودی می‌باشد. برای مثال فرض کنید سیگنال‌های ورودی $Y = [y_1, y_2, \dots, y_6]$ و $D = [d_1, d_2, d_3, d_4]$ باشد و d_1, d_2 و y_1 و y_2 و y_3 مربوط به کلاس یک و d_3, d_4 و y_4 و y_5 و y_6 مربوط به کلاس دو باشند. Q به صورت زیر تعریف می‌شود.

$$Q = \begin{bmatrix} 1 & 1 & 1 & 0 & 0 & 0 \\ 1 & 1 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 1 & 1 \\ 0 & 0 & 0 & 1 & 1 & 1 \end{bmatrix}_{4 \times 6}$$

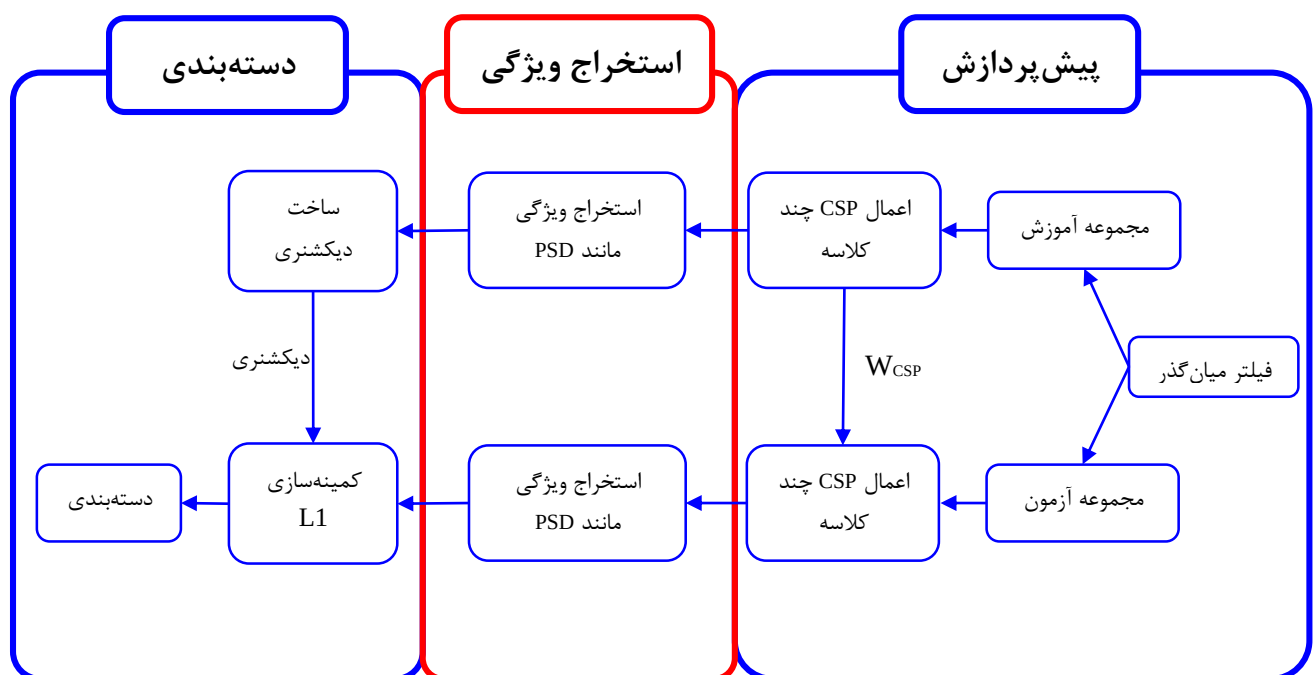
ماتریس H نیز به صورت زیر تعریف می‌شود.

^۱ Discriminability.

$$H = \begin{bmatrix} 1 & 1 & 1 & \vdots & \vdots & \vdots \\ \vdots & \vdots & \vdots & 1 & 1 & 1 \end{bmatrix}_{2 \times 6}$$

که Q بهترین حالت برای پاسخ تنک و H بیانگر کلاس‌های مربوط به سیگنال ورودی می‌باشد.

۳-۵ روش‌های پیشنهادی



در این بخش روش پیشنهادی دسته‌بندی سیگنال EEG با استفاده از نمایش تنک و یادگیری دیکشنری

آورده شده است. در بخش ۳-۵-۱ روش پیشنهادی بر مبنای نمایش تنک و در بخش ۳-۵-۲

روش پیشنهادی یادگیری دیکشنری آورده شده است. ارزیابی و تحلیل این روش‌ها در فصل چهار ارائه

گردیده است.

۳-۵-۱ روش پیشنهادی دسته‌بندی سیگنال EEG با استفاده از نمایش تنک

همانطور که در بخش‌های قبلی به آن اشاره شد؛ نمایش تنک در دسته‌بندی تصاویر، دسته‌بندی سیگنال

استفاده شده است. در این بخش پس از بررسی و ارائه ساخت دیکشنری با استفاده از فیلتر CSP به چگونگی دسته‌بندی سیگنال EEG پرداخته شده است. روش پیشنهادی در شکل ۳-۸ آورده شده است. برای پیش‌پردازش از فیلتر میان‌گذر ۵ تا ۴۰ هرتز استفاده شده است تا اطلاعات مفید باندهای مختلف (تتا، آلفا، بتا و گاما) را داشته باشیم. پس از اعمال فیلتر میان‌گذر بر روی سیگنال‌های ورودی، با اعمال فیلتر CSP اطلاعات مفید را بین کلاس‌های مختلف بیشینه می‌نماییم. در مرحله بعدی پس از ساخت دیکشنری با استفاده از روش کمینه‌سازی L1 دسته‌بندی را انجام می‌دهیم.

• مرحله اول: پیش‌پردازش و اعمال فیلتر CSP

روش پیشنهادی را روی پایگاه داده دانشگاه برلین شرح می‌دهیم. این پایگاه داده شامل سیگنال‌های ثبت شده از تصور ذهنی حرکت دست چپ و دست راست می‌باشد. جزئیات پایگاه داده در بخش ۳ - ۱ آورده شده است. فرض کنید $X_R \in R^{C \times T}$ مربوط به کلاس تصور ذهنی حرکت دست راست و $X_L \in R^{C \times T}$ مربوط به کلاس تصور ذهنی حرکت دست چپ باشد که C تعداد کانال‌های EEG و T تعداد نمونه‌ها باشد.

پس از اعمال فیلتر میان‌گذر روی سیگنال‌های ورودی؛ $W \in R^{C \times C}$ را از اعمال فیلتر CSP محاسبه می‌نماییم. W به گونه‌ای محاسبه می‌شود که واریانس بین دو کلاس را افزایش دهد. معادله ۳-۲۱ مربوط به محاسبه W می‌باشد.

$$W = \max_W \left(\frac{W^T C_R W}{W^T C_L W} \right) \quad ۲۱-۳$$

در این معادله $C_L = X_L X_L^T$ و $C_R = X_R X_R^T$. (جزئیات بیشتر روش CSP برای دو کلاس و چند کلاس در بخش ۳ - ۲ - ۱ آورده شده است)

پایگاه داده برلین شامل ۱۱۸ کانال EEG می‌باشد؛ بنابراین $W \in R^{C \times C} = \{W_1, W_2, \dots, W_{118}\}$ می‌باشد

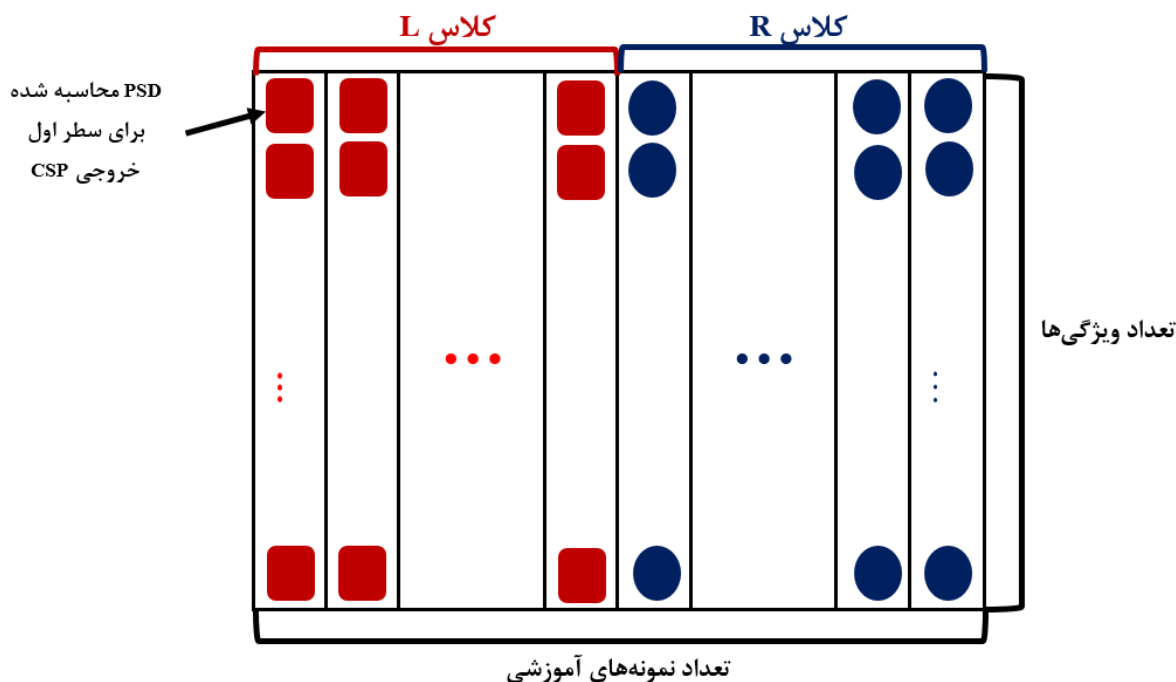
که برخی از این W_i ($1 \leq i \leq 118$) ها باید انتخاب شوند. در این روش از تعداد $2n$ فیلتر CSP استفاده شده است. n ستون ابتدایی ماتریس W و n ستون انتهایی آن انتخاب شده است. این شیوه انتخاب تعداد CSP ها واریانس بین کلاس ها را افزایش و از طرف دیگر با انتخاب n ستون انتهایی واریانس درون کلاسی را کاهش می دهد. پس از اعمال فیلتر CSP و انتخاب فیلترهای مورد نظر W_{CSP} به صورت زیر تعریف می شود.

$$W_{CSP} = [W_1, \dots, W_n, W_{C-n+1}, \dots, W_{2n}] \quad 22-3$$

انتخاب تعداد فیلترهای CSP بستگی به افراد مورد آزمایش و تعداد داده های آموزشی دارد. در بخش ۳ - ۵ - ۳ یک روش بر مبنای الگوریتم ژنتیک برای به دست آوردن تعداد بهینه فیلتر CSP ارائه شده است. نحوه اعمال W_{CSP} روی داده های ورودی X_L و X_R به صورت زیر می باشد.

$$\begin{aligned} X_R^{CSP} \in R^{2n \times T} &= W_{CSP}^T X_R \\ X_L^{CSP} \in R^{2n \times T} &= W_{CSP}^T X_L \end{aligned} \quad 23-3$$

پس از اعمال فیلتر CSP بر روی داده های ورودی مرحله ساخت دیکشنری می باشد. دیکشنری طراحی شده در شکل ۳-۹ آورده شده است. در مرحله بعدی چگونگی ساخت دیکشنری ارائه داده شده است.



شکل ۳-۹. نحوه ساخت دیکشنری

• مرحله دوم: ساخت دیکشنری

خروجی مرحله قبل؛ X_L^{CSP} و X_R^{CSP} می‌باشد که باید با استفاده از آنها یا ویژگی‌های استخراج شده از آنها دیکشنری مورد نظر طراحی شود. همانطور که در بخش ۲ - ۳ - ۳ بیان شد، می‌توان ویژگی‌های گوناگونی نظیر PSD، واریانس، آنتروپی و غیره را برای سیگنال در نظر گرفت. دیکشنری طراحی شده در این قسمت از مقدار PSD محاسبه شده برای هر سطر از X_L^{CSP} و X_R^{CSP} ساخته شده است. مقدار PSD را می‌توان برای باندهای مهم در نظر گرفت که در کارهای متعددی مورد استفاده قرار گرفته است. PSD برای سیگنال $x(t)$ از رابطه زیر محاسبه می‌شود.

$$PSD = \int_{-\infty}^{+\infty} S_x(f) df \quad ۲۴-۳$$

که در این فرمول $S_x(f)$ به صورت زیر محاسبه می‌شود.

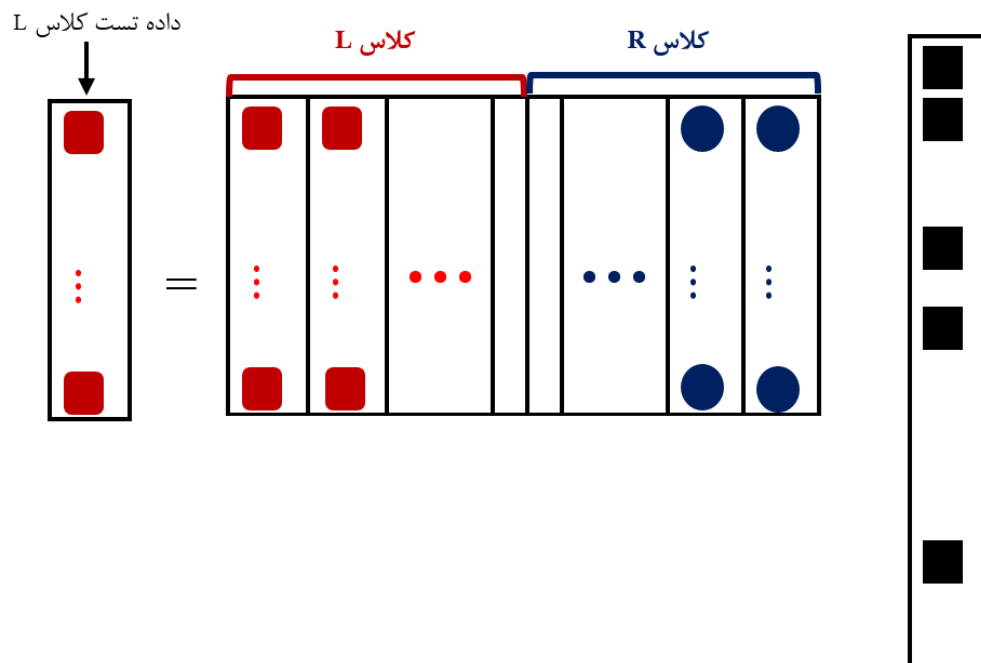
$$S_x(f) = \lim_{T \rightarrow \infty} \left\{ \frac{1}{2T} \left| \int_{-T}^T x(t) e^{-j2\pi f t} dt \right|^2 \right\} \quad ۲۵-۳$$

همچنین می‌توان بهترین بازه فرکانسی برای استخراج ویژگی را به وسیله روش‌هایی مانند الگوریتم ژنتیک (بخش ۳ - ۵ - ۳) انتخاب نمود. پس از استخراج ویژگی برای هر سطر از ماتریس‌های X_R^{CSP} و X_L^{CSP} دیکشنری مانند آنچه در شکل ۹-۳ مشاهده می‌کنید؛ طراحی شده است. هر ستون از دیکشنری (اتم) ویژگی‌های استخراج شده برای یک داده آموزشی را نشان می‌دهد و هر عضو از آن بیانگر مقدار PSD محاسبه شده از هر سطر ماتریس X_R^{CSP} و X_L^{CSP} می‌باشد.

دیکشنری ایجاد شده شامل دو بخش D_L و D_R مربوط به کلاس‌های تصور ذهنی حرکت دست چپ و راست می‌باشد. اگر تعداد کل آزمایش‌های موجود را برابر N_t فرض نماییم؛ ابعاد ماتریس ایجاد شده برابر است با $2n \times N_t$. $2n$ برابر است با ویژگی‌های استخراج شده از هر داده آموزشی)

• مرحله سوم: دسته‌بندی

پس از انجام دو مرحله گذشته دیکشنری ایجاد شده ($D \in R^{2n \times N_t}$) برای دسته‌بندی مورد استفاده قرار می‌گیرد. شکل ۱۰-۳ نحوه استفاده از دیکشنری برای به دست آوردن را نشان می‌دهد.



شکل ۱۰-۳. محاسبه پاسخ تنک با استفاده از دیکشنری طراحی شده

برای توضیح مرحله دسته‌بندی، فرض کنید $Y \in R^{C \times T}$ داده آزمون مربوط به کلاس تصور ذهنی حرکت دست راست باشد. $y_{csp} \in R^{2n \times T}$ را با استفاده از W_{CSP} محاسبه شده در مرحله آموزشی به صورت زیر محاسبه می‌شود.

$$y_{CSP} = W_{CSP}^T \times Y \quad ۲۶-۳$$

پس از محاسبه مقدار PSD برای هر سطر از y_{csp} ، $y \in R^{2n \times 1}$ به دست می‌آید. تعداد سطرهای آرایه y با تعداد سطرهای ماتریس دیکشنری برابر است. اکنون با استفاده از دیکشنری طراحی شده و سیگنال y ، نمایش تنک را می‌توان برای سیگنال y محاسبه نمود. روش‌های به دست آوردن نمایش تنک در بخش ۳-۳-۱ توضیح داده شده است. در روش پیشنهادی با توجه به دقت روش‌های نرم $L1$ از روش BP برای محاسبه پاسخ تنک استفاده شده است؛ بنابراین برای به دست آوردن نمایش تنک x از معادله ۳-۱۷ استفاده شده است.

پس از محاسبه پاسخ تنک x ، برای دسته‌بندی از معادله زیر استفاده شده است.

$$\text{Class}(y) = \underset{i=R,L}{\operatorname{argmin}} \operatorname{Error}_i(y) \quad ۲۷-۳$$

کلاس i با اندیس کمترین میزان تابع Error است. تابع $\operatorname{Error}(y)$ به صورت زیر تعریف می‌شود.

$$\operatorname{Error}_{i=R,L}(y) = \left\| y - D_i \delta_i(x) \right\|_2 \quad ۲۸-۳$$

تابع $\delta_i(x)$ ضرایب پاسخ تنک مربوط به کلاس‌های مخالف i را برابر صفر قرار می‌دهد؛ به عبارت دیگر برای هر کلاس به صورت جداگانه مقدار خطای بازسازی داده آزمون را با استفاده از اتم‌های مربوط به همان کلاس محاسبه می‌نماید.

توجه شود که روش پیشنهادی برای بیش از دو کلاس قابل پیاده‌سازی است و به منظور سادگی در ارائه و فهم بهتر روش، در این قسمت برای دو کلاس توضیح داده شده است. پیاده‌سازی این روش برای بیش

از دو کلاس نیز با استفاده از CSP چند کلاسه و طراحی دیکشنری انجام شده است. نتایج و تحلیل‌های مربوط به این روش در بخش ۴ - ۳ مورد بررسی قرار گرفته است.

۳ - ۵ - ۲ روش پیشنهادی یادگیری دیکشنری برای دسته‌بندی

در بخش گذشته، روش دسته‌بندی با استفاده از نمایش تنک بیان شد. همانطور که مشاهده کردید، دیکشنری از ویژگی‌های به دست آمده از داده‌های آموزشی ساخته شد و سپس برای دسته‌بندی از نمایش تنک محاسبه شده توسط الگوریتم BP استفاده شد. در این بخش روش پیشنهادی برای دسته‌بندی داده EEG بر مبنای یادگیری دیکشنری مطرح شده است. در ابتدا مقدمه‌ای در مورد دیکشنری‌های مصنوعی^۱ و تحلیلی^۲ بیان شده و سپس روش پیشنهادی مطرح گردیده است.

دیکشنری نقش مهمی در فرایند بازسازی سیگنال با استفاده از نمایش تنک ایفا می‌کند. برخی از دیکشنری‌های تحلیلی از پیش تعریف شده (دیکشنری تبدیل موجک، گابور) بازسازی سیگنال را با استفاده از ضرب ساده مهیا می‌نمایند. ویژگی مناسب این نوع دیکشنری‌ها به دست آوردن پاسخ در زمان مناسب می‌باشد اما این نوع دیکشنری‌ها در مدل‌های پیچیده غیر مؤثر عمل می‌نمایند. برای مثال دیکشنری تبدیل FFT شامل انواع اتم‌های متعامد سینوسی می‌باشد؛ این نوع دیکشنری کارایی مناسبی برای سیگنال‌های با ماهیت سینوسی داراست اما اگر ماهیت سیگنال سینوسی نباشد این نوع دیکشنری فاقد کارایی لازم می‌باشد.

در سال‌های اخیر کارهای متعددی برای ساخت دیکشنری‌های مصنوعی انجام شده است که در بخش ۳ - ۴ به برخی از آنها اشاره شد. اکثر روش‌های یادگیری دیکشنری برای به دست آوردن پاسخ تنک از نرم‌های L_p ($0 \leq p \leq 1$) استفاده می‌نمایند که بار محاسباتی بیشتری نسبت به دیکشنری‌های تحلیلی

^۱ Synthesis.

^۲ Analysis.

دارند. از طرف دیگر، دیکشنری‌های مصنوعی برای مدل‌های پیچیده از دیکشنری تحلیلی بهتر عمل می‌نماید زیرا در اغلب روش‌های ایجاد دیکشنری از داده‌های آموزشی همان مدل استفاده می‌شود. همچنین در ساخت این نوع دیکشنری می‌توان پارامترهایی نظیر افتراق‌پذیری بین کلاس‌های مختلف و میزان خطای دسته‌بند را در نظر گرفت.

در این بخش روش پیشنهادی مطرح شده برای ساخت و یادگیری جفت دیکشنری^۱ (DPL) به صورت همزمان مطرح شده است. یک دیکشنری مصنوعی حاوی برخی ویژگی‌های داده‌های آموزشی با داشتن ویژگی‌های افتراق‌پذیری و دیکشنری تحلیلی برای محاسبه پاسخ تنک مورد استفاده قرار می‌گیرد. روش پیشنهادی در شکل ۳-۱۱ نشان داده شده است.

مراحل این روش پیشنهادی در ادامه شرح شده است.

- مرحله اول: پیش‌پردازش، اعمال فیلتر CSP و ساخت دیکشنری

این مرحله همانند مراحل اول و دوم بخش ۳ - ۵ - ۱ می‌باشد.

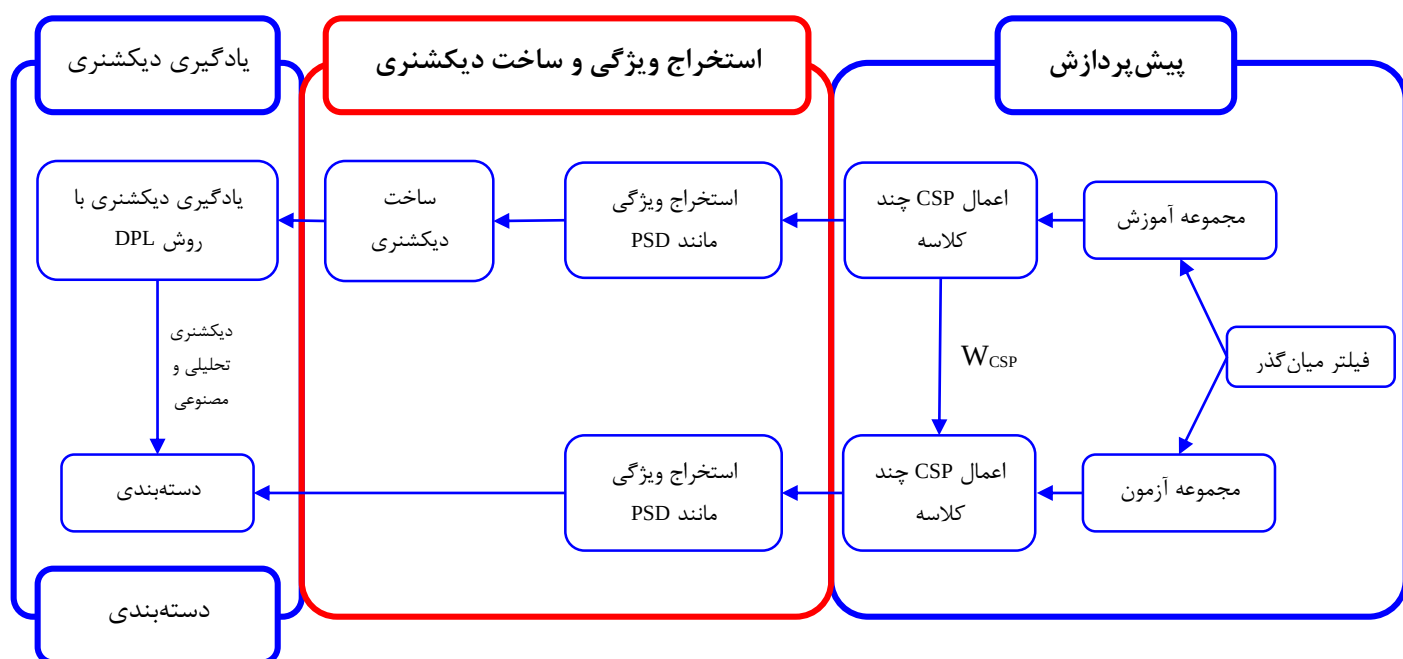
- مرحله دوم: یادگیری جفت دیکشنری

اغلب روش‌های یادگیری دیکشنری از معادله ۳-۲۹ زیر برای ساخت و آموزش دیکشنری استفاده می‌نمایند.

$$\min_{D, X} \|Y - DX\|_F^2 + \lambda \|X\|_p \quad ۳-۲۹$$

در این معادله $Y \in \mathbb{R}^{n \times k}$, $D \in \mathbb{R}^{n \times m}$, $X \in \mathbb{R}^{m \times k}$ به ترتیب بیانگر پاسخ تنک، دیکشنری مصنوعی و سیگنال ورودی می‌باشند. $\|X\|_p$ نیز بیان‌کننده نرم p برای محاسبه X و عبارت $\|Y - DX\|$ میزان

^۱ Dictionary pair learning.



خطای بازسازی سیگنال اصلی می باشد. پارامتر λ عدد ثابت برای کنترل میزان تنک بودن سیگنال می باشد.

اغلب روش های یادگیری دیکشنری از نرم L_0 و L_1 در مرحله آموزش و آزمون استفاده می نمایند که برای سیستم بار محاسباتی دارد. در روش DPL با استفاده از یادگیری دیکشنری تحلیلی برای محاسبه پاسخ تنک از ضرب خطی به جای روش های غیر خطی استفاده شده است و همچنین عملکرد سیستم در میزان دقت بهبود یافته است.

فرض نماییم سیگنال ورودی $Y = [y_1, \dots, y_q, \dots, y_Q]$ می باشد که در آن $y_q \in \mathbb{R}^{n \times k}$ سیگنال ورودی از کلاس q باشد و در مجموع Q تا کلاس داشته باشیم. اگر بتوانیم یک دیکشنری تحلیلی $P \in \mathbb{R}^{m \times n}$ ایجاد نماییم که برای محاسبه پاسخ تنک X بتوانیم از رابطه $X = PY$ استفاده نماییم در کارایی سیستم تأثیر مثبت دارد. با در نظر گرفتن دو دیکشنری تحلیلی و مصنوعی معادله ۳-۲۹ به شکل زیر تبدیل می شود.

$$\{P^*, D^*\} = \underset{P, D}{\operatorname{argmin}} \|Y - DPY\|_F^2 + \Psi(D, P, Y, H) \quad 3-30$$

که تابع Ψ در آن می تواند شامل عبارت هایی برای ساخت دیکشنری نظیر میزان افتراق پذیری و برای

ساخت پاسخ تنک مانند میزان تنک بودن باشد. H شامل کلاس‌های مربوط به سیگنال ورودی Y می‌باشد. دیکشنری P برای محاسبه پاسخ تنک و دیکشنری D برای بازسازی سیگنال ورودی استفاده می‌شود.

برای افزودن ویژگی‌هایی مانند افتراق‌پذیر بودن دیکشنری‌ها و تنک بودن پاسخ X باید تابع Ψ را به طور مناسب طراحی نماییم. برای این منظور ساختار P و D را به صورت جداگانه برای هر کلاس در نظر می‌گیریم؛ به عبارت دیگر $D = \{D_1, \dots, D_q, \dots, D_Q\}$ و $P = \{P_1, \dots, P_q, \dots, P_Q\}$ از دیکشنری‌های مصنوعی و تحلیلی ایجاد شده برای هر کلاس تشکیل شده‌اند به طوری که $P_q \in \mathbb{R}^{m \times n}$ و $D_q \in \mathbb{R}^{n \times m}$ زیرمجموعه‌های دیکشنری‌های مصنوعی و تحلیلی برای کلاس q می‌باشند.

دیکشنری P باید به گونه‌ی تعریف شود که هنگام اعمال P_q روی سیگنال Y پاسخ تنک به دست آمده برای کلاس‌های غیر q برابر صفر شود؛ به عبارت دیگر دیکشنری P_q فقط برای سیگنال‌های کلاس q آموزش دیده باشد و برای سایر کلاس‌ها مقدار صفر را برگرداند. رابطه ۳-۳۱ بیانگر اعمال چنین شرایطی برای دیکشنری P می‌باشد.

$$P_q Y_{q'} \approx 0; q \neq q' \text{ and } 1 < q, q' < Q \quad 3-31$$

از آنجا که دیکشنری P و D را به قسمت‌هایی متناسب با کلاس‌های مسئله در نظر گرفته شده است، خطای بازسازی برای هر کلاس به صورت جداگانه باید محاسبه شود. این خطا به صورت زیر محاسبه می‌شود.

$$\min_{P,D} \sum_{i=1}^Q \|Y_i - D_i P_i Y_i\|_F^2 \quad 3-32$$

با در نظر گرفتن معادله ۳-۳۱ برای ساخت تابع Ψ و معادله ۳-۳۲ برای خطای بازسازی، معادله ۳-۳۰ را می‌توان به صورت زیر می‌توان نوشت.

$$\{P^*, D^*\} = \underset{P, D}{\operatorname{argmin}} \sum_{i=1}^Q \|Y_i - D_i P_i Y_i\|_F^2 + \lambda \|P_i \bar{Y}_i\|_2^2, \text{ s.t. } \|d_j\| \leq 1 \quad 33-3$$

این معادله یک معادله غیر محدب می باشد؛ با در نظر گرفتن متغیر $A = PY$ معادله ۳۳-۳ را می توان به صورت زیر نوشت:

$$\{P^*, D^*, A^*\} = \underset{P, D, A}{\operatorname{argmin}} \sum_{i=1}^Q \|Y_i - D_i A_i\|_F^2 + \gamma \|P_i Y_i - A_i\|_F^2 + \lambda \|P_i \bar{Y}_i\|_F^2, \quad 34-3$$

$$\text{s.t. } \|d_j\| \leq 1$$

حل معادله ۳۴-۳ در [۴۴] انجام شده است. برای حل معادله ۳۴-۳ گام های زیر انجام شده است.

۱- ثابت فرض نمودن P و D و به روز رسانی A :

$$A^* = \underset{A}{\operatorname{argmin}} \sum_{i=1}^Q \|Y_i - D_i A_i\|_F^2 + \gamma \|P_i Y_i - A_i\|_F^2 \quad 35-3$$

معادله ۳۵-۳ یک مسئله Standard least square می باشد و با حل آن A^* برابر می شود با:

$$A^* = (D_i^T D_i + \gamma I)^{-1} (\gamma P_i Y_i + D_i^T D_i) \quad 36-3$$

۲- ثابت فرض نمودن A و به روز رسانی P و D :

$$P^* = \underset{P}{\operatorname{argmin}} \sum_{i=1}^Q \gamma \|P_i Y_i - A_i\|_F^2 + \lambda \|P_i \bar{Y}_i\|_F^2 \quad 37-3$$

$$D^* = \underset{D}{\operatorname{argmin}} \sum_{i=1}^Q \|Y_i - D_i A_i\|_F^2 \text{ s.t. } \|d_j\| \leq 1$$

پس از حل آن P^* برابر است با:

$$P_i^* = \gamma A_i Y_i^T (\gamma Y_i Y_i^T + \lambda \bar{Y}_i \bar{Y}_i^T + 10^{-4} \times I)^{-1} \quad 38-3$$

برای به دست آوردن D^* از الگوریتم ADMM شده است [۴۵]. با در نظر گرفتن $D = S$ مقدار بهینه

برای D به صورت زیر محاسبه می‌شود.

$$D^* = \underset{D, S}{\operatorname{argmin}} \sum_{i=1}^Q \|Y_i - D_i A_i\|_F^2 \text{ s.t } D=S, \|d_j\| \leq 1 \quad ۳۹-۳$$

برای حل معادله ۳۹-۳ با استفاده از الگوریتم ADMM گام‌های زیر را باید انجام داد.

$$D^{(r+1)} = \min_D \sum_{i=1}^Q \|Y_i - D_i A_i\|_F^2 + \rho \|D_i - S_i^{(r)} + T_i^{(r)}\|_F^2$$

$$S^{(r+1)} = \min_S \sum_{i=1}^Q \rho \|D_i^{(r+1)} - S_i^{(r)} + T_i^{(r)}\|_F^2, \text{ s.t } \|S_i\| \leq 1 \quad ۴۰-۳$$

$$T^{(r+1)} = T^{(r)} + D_i^{(r+1)} - S_i^{(r+1)}$$

خلاصه الگوریتم DPL در شکل ۱۲-۳ آورده شده است. خروجی الگوریتم DPL دو دیکشنری تحلیلی و مصنوعی آموزش دیده P و D می‌باشد. از این دو دیکشنری برای مرحله دسته‌بندی استفاده می‌شود.

هدف: محاسبه دیکشنری تحلیلی P و دیکشنری مصنوعی D
ورودی: $Y = [y_1, \dots, y_q, \dots, y_Q]$ خروجی: D و P
گام اول - مقداردهی اولیه: دیکشنری P با اعداد تصادفی و D با ورودی Y
گام دوم - مراحل زیر را تا همگرایی انجام بده
گام سوم: $t = t + 1$
گام چهارم: برای $i = 1$ تا Q مراحل زیر را انجام بده
گام پنجم: به‌روزرسانی $A_i^{(t)}$ با استفاده از فرمول ۳۶-۳
گام ششم: به‌روزرسانی $P_i^{(t)}$ با استفاده از فرمول ۳۸-۳
گام هفتم: به‌روزرسانی $D_i^{(t)}$ با استفاده از فرمول ۴۰-۳
گام هشتم: به گام چهارم برو
گام نهم: به گام دوم برو
خروجی: D^* و P^*

شکل ۱۲-۳. الگوریتم DPL

در این مرحله با استفاده از الگوریتم DPL دو دیکشنری تحلیلی و مصنوعی به دست آمد. دیکشنری مصنوعی D برای بازسازی سیگنال اصلی و دیکشنری تحلیلی P برای به دست آوردن پاسخ تنک مورد

استفاده قرار گرفت. دیکشنری تحلیلی P با توان افتراق پذیری بین کلاس‌های مختلف آموزش دیده است که نقش مؤثری را در دسته‌بندی ایفا می‌نماید. در مرحله بعدی نحوه دسته‌بندی با استفاده از روش DPL شرح داده شده است.

- مرحله سوم: دسته‌بندی

فرض نماییم $Y_t \in \mathbb{R}^{C \times T}$ داده ورودی مرحله آزمون باشد که دارای C کانال EEG و تعداد نمونه برابر T باشد. همچنین در نظر بگیرید که داده‌ها متعلق به دو کلاس تصور ذهنی حرکت راست (R) و چپ (L) باشد. برای دسته‌بندی Y_t گام‌های زیر باید انجام شود:

- اعمال مرحله پیش‌پردازش روی سیگنال Y_t (فیلتر میان‌گذر)
- اعمال فیلتر CSP روی خروجی گام قبلی با استفاده از W محاسبه شده در مرحله آموزش
- محاسبه بردار ویژگی $Y_t \in \mathbb{R}^{n \times 1}$.
- کلاس داده‌ها با استفاده از P^* و D^* به دست آمده از اعمال روش DPL به صورت زیر مشخص می‌شود

$$\text{class}(Y_f) = \underset{i}{\operatorname{argmin}} \|Y_f - D_i P_i Y_f\|_F^2 \quad ۴۱-۳$$

از آنجا که در اینجا برای داده‌ها دو کلاس R و L را در نظر گرفته شده است؛ در نتیجه داریم:

$$P = \{P_L, P_R\}; \text{ s.t } P_L, P_R \in \mathbb{R}^{m \times n}$$

$$D = \{D_L, D_R\}; \text{ s.t } D_L, D_R \in \mathbb{R}^{n \times m}$$

با استفاده از فرمول ۴۱-۳ کلاس داده ورودی برابر است با:

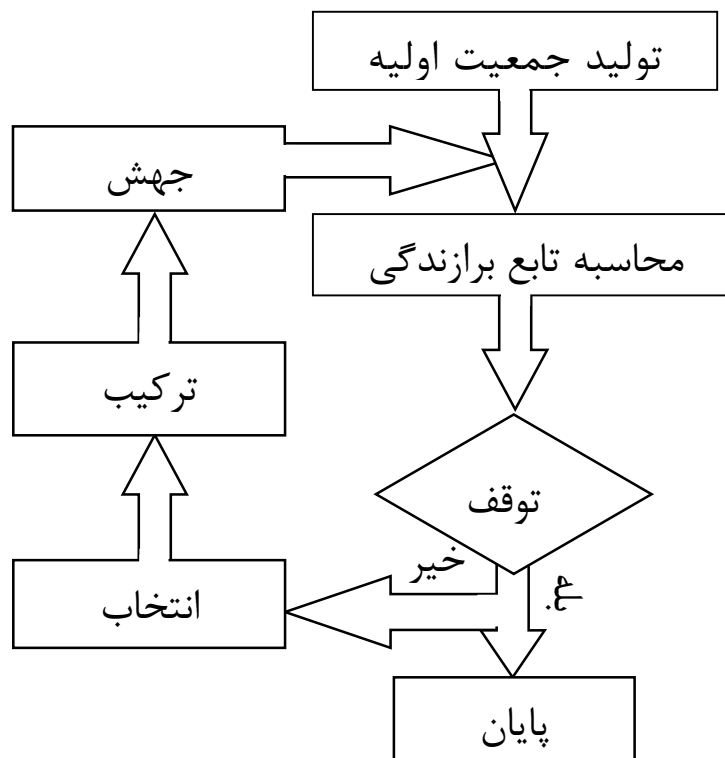
$$\text{class}(Y_f) = \underset{i}{\operatorname{argmin}} \|Y_f - D_i P_i Y_f\|_F^2$$

اگر فرض نماییم کلاس Y_t مربوط به کلاس L باشد؛ با در نظر گرفتن معادله ۴۱-۳ داریم $P_R Y_f \approx 0$.

در نتیجه مقدار باقیمانده $\|Y_f - D_R P_R Y_f\|_F^2$ بزرگتر از $\|Y_f - D_L P_L Y_f\|_F^2$ می‌باشد و کلاس Y_t برابر با L تشخیص داده خواهد شد.

۳-۵-۳ استفاده از الگوریتم ژنتیک^۱ برای محاسبه پارامترهای اولیه

الگوریتم ژنتیک (GA) یکی از قدرتمندترین الگوریتم‌های انتخاب پارامترهای بهینه می‌باشد و کاربردهای وسیعی در همه زمینه‌های علوم و فناوری دارد. این الگوریتم برای فضاهای جستجوی بزرگ که جستجو در آن در عمل غیر ممکن است، مناسب می‌باشد. این الگوریتم روش مناسبی برای یافتن جواب مسائل بهینه‌سازی می‌باشد که الهام گرفته از تولید نسل می‌باشد. مراحل الگوریتم ژنتیک در شکل ۳-۱۳ آورده شده است. [۴۶]



شکل ۳-۱۳. مراحل الگوریتم ژنتیک

^۱ Genetic algorithm.

همانطور که در شکل ۳-۱۳ مشاهده می‌کنید؛ الگوریتم ژنتیک با یک جمعیت اولیه شروع به کار می‌کند. هر فرد در جمعیت، یک کروموزوم نامیده می‌شود که یک جواب برای مسئله بهینه‌سازی می‌باشد. الگوریتم ژنتیک برخی از کروموزوم‌ها را برای مرحله بعدی انتخاب نموده و آن‌ها را با یکدیگر ترکیب نموده و برخی از آنها جهش یافته و مجموع این کروموزوم‌ها جمعیت نسل بعد را تشکیل می‌دهند.

انتخاب تعداد فیلتر CSP نقش مستقیمی در نتیجه دسته‌بندی دارد. همچنین بازه فرکانسی انتخابی برای استخراج ویژگی از اهمیت زیادی برخوردار است. این بازه باید به نحوی انتخاب شود که شامل اطلاعات فعالیت ذهنی انجام شده باشد. انتخاب مناسبترین ترکیب بازه فرکانسی و تعداد فیلتر CSP برای به دست آوردن بهترین نتیجه دسته‌بند را می‌توان با آزمون و خطا، الگوریتم‌های تکاملی و جستجوهای هیورستیک انجام داد. در این پایان‌نامه جهت یافتن پارامترهای اولیه نظیر تعداد فیلتر CSP و انتخاب باندهای فرکانسی برای استخراج ویژگی از الگوریتم ژنتیک استفاده شده است.

۳ - ۵ - ۳ پارامترهای الگوریتم ژنتیک

برای حل یک مسئله با الگوریتم ژنتیک باید یک سری پارامتر را برای آن مشخص نمود. پارامترهایی که برای حل این مسئله در نظر گرفته شده‌اند در زیر آمده است.

• کروموزوم

برای حل این مسئله خاص، کروموزوم به صورت عدد باینری در نظر گرفته شده است که شامل سه قسمت می‌باشد. قسمت اول بیانگر تعداد فیلتر CSP، قسمت دوم عدد کوچکتر بازه فرکانسی و قسمت سوم عدد بزرگتر بازه فرکانسی را در خود جای داده است. نمونه‌ای از کروموزوم در شکل ۳-۱۴ آورده شده است. این کروموزوم نشان‌دهنده تعداد فیلتر CSP برابر با ۶ و باند فرکانسی مورد استفاده شامل بازه ۸ تا ۱۳ هرتز می‌باشد.

0	1	0	1	0	1	0	0	1	0	0	0	0	0	1	1	0
---	---	---	---	---	---	---	---	---	---	---	---	---	---	---	---	---

فرکانس بالا: ۱۳

فرکانس پایین: ۸

تعداد CSP: ۶

شکل ۳-۱۴. نمونه کروموزوم

• تابع برازش^۱

تابع برازش مورد استفاده در برای این مسئله میزان دقت دسته‌بندی می‌باشد.

عملگر انتخاب^۲ و عملگر ترکیب^۳ و عملگر جهش^۴ و شرط توقف برای این مسئله در بخش ۴ - ۲ آورده شده است.

در این فصل روش دسته‌بندی سیگنال EEG با استفاده از نمایش تنک و یادگیری دیکشنری مطرح شد. در بخش اول جزئیات پایگاه داده‌های برلین و کلرادو آورده شد. در بخش دوم پیش‌پردازش استفاده شده برای روش پیشنهادی معرفی و توضیح داده شد. از فیلتر میان‌گذر برای حذف مصنوعات استفاده شد. از فیلتر CSP دو کلاسه و چند کلاسه برای افزایش بهره اطلاعاتی بین کلاس‌های مختلف استفاده شد. استفاده از فیلتر CSP کیفیت اتم‌های دیکشنری را افزایش می‌دهد.

روش‌های پیشنهادی بر مبنای نمایش تنک و یادگیری دیکشنری مطرح گردید و جزئیات آن مورد بررسی قرار گرفت. روش‌های محاسبه نمایش تنک و همچنین چند روش یادگیری دیکشنری در این فصل آورده شد. نتایج دو روش پیشنهادی در فصل چهار آورده شده است.

^۱ Fitness function.

^۲ Cross over.

^۳ Selection.

^۴ Mutation.

فصل ۴

نتایج

در این فصل نتایج به دست آمده از روش‌های پیشنهادی بر روی دو پایگاه داده برلین و کلرادو مورد بحث و بررسی قرار داده شده است. در بخش اول قسمت‌های مختلف دو پایگاه داده و استفاده از اعتبارسنجی برای ارزیابی نتایج توضیح داده است. در بخش دوم پارامترهای اولیه سیستم معرفی و نحوه محاسبه آنها با استفاده از الگوریتم ژنتیک بررسی شده است. در دو بخش پایانی نتایج به دست آمده از روش‌های نمایش تنک و DPL به ترتیب آورده شده است و به مقایسه و بررسی عملکرد روش‌های پیشنهادی پرداخته شده است.

۴-۱ قسمت‌های مختلف پایگاه داده

برای ارزیابی روش‌های دسته‌بندی از یک پایگاه داده شامل داده‌های آموزشی و آزمون استفاده می‌شود. این دو مجموعه باید کاملاً از یکدیگر مجزا باشند. از طرف دیگر برای تنظیم پارامترهای اولیه مورد نیاز سیستم، از مجموعه داده اعتبارسنجی^۱ استفاده می‌شود زیرا از داده‌های مجموعه آزمون نمی‌توان برای آموزش سیستم یا تنظیم پارامترهای سیستم استفاده نمود.

در پایگاه داده برلین مجموعه داده آموزش و آزمون برای هر فرد مشخص شده است که جدول ۴-۱ بیانگر تعداد نمونه‌های آموزشی و آزمون برای فرد می‌باشد.

جدول ۴-۱. تعداد نمونه‌های آموزشی و آزمون برای پایگاه داده برلین

افراد	aa	al	av	aw	ay
تعداد نمونه‌های آموزشی	۱۶۸	۲۲۴	۸۴	۵۶	۲۸

^۱ Validation.

تعداد نمونه‌های آزمون	۱۱۲	۵۶	۱۹۶	۲۲۴	۲۵۲
-----------------------	-----	----	-----	-----	-----

در این پایگاه داده برای تنظیم نمودن پارامترهای اولیه از مجموعه اعتبارسنجی که شامل قسمتی از داده‌های آموزشی می‌باشد؛ استفاده شده است.

در پایگاه داده کلرادو مجموعه آزمون و آموزشی مشخص نشده است و همچنین تعداد داده‌های این مجموعه محدود می‌باشد؛ بنابراین برای ارزیابی عملکرد سیستم از اعتبارسنجی متقابل^۱ استفاده شده است. تعداد نمونه‌های این پایگاه داده برای هر فرد در جدول ۲-۴ آورده شده است.

جدول ۲-۴. تعداد نمونه‌های پایگاه داده کلرادو

افراد	فرد ۱	فرد ۲	فرد ۳	فرد ۴	فرد ۵	فرد ۶	فرد ۷
تعداد نمونه‌ها	۵۰	۲۵	۵۰	۵۰	۵۰	۵۰	۲۵

پارامترهای مربوط به این پایگاه داده با استفاده از روش سعی و خطا به دست آمده است. برای ارزیابی نتایج آن از اعتبارسنجی متقابل استفاده شده است. میزان دقت دسته‌بند از رابطه ۱-۴ محاسبه شده است.

$$۱-۴ \quad \text{دقت دسته‌بند} = \frac{\text{تعداد نمونه‌هایی که درست دسته‌بندی شده‌اند}}{\text{تعداد کل نمونه‌ها}}$$

• اعتبارسنجی متقابل

هنگامی که پایگاه داده شامل تعداد محدودی داده باشد نمی‌توانیم دسته‌بندی مناسبی را برای مجموعه آزمون و آموزشی داشته باشیم. همچنین هنگامی که داده‌های آموزشی و آزمون مشخص نباشد؛ انتخاب مجموعه نماینده مناسب برای آموزش و آزمون دشوار می‌باشد. در چنین مواقعی از روش اعتبارسنجی

^۱ Cross validation.

متقابل استفاده می‌شود که با استفاده از آن می‌توان همه داده‌ها را برای آموزش و آزمون در نظر گرفت. اعتبارسنجی متقابل معمولاً به صورت تکرار مراحل آموزش و آزمون روی نمونه‌های موجود در پایگاه داده به نحوی انجام می‌شود که در هر تکرار مجموعه در نظر گرفته شده برای آموزش و آزمون با تکرارهای دیگر متفاوت است. از انواع اعتبارسنجی می‌توان به اعتبارسنجی با k زیر نمونه^۱، اعتبارسنجی دسته‌ای با یک نمونه بیرون^۲ (LOO) و زیر نمونه‌برداری تصادفی تکراری^۳ اشاره نمود.

۴ - ۲ محاسبه پارامترهای اولیه

برخی پارامترهای اولیه نقش مهم و مؤثری در میزان دقت سیستم ایفا می‌نمایند. در این سیستم می‌توان به برخی پارامترها مانند تعداد فیلتر CSP بازه فرکانسی مورد نظر، بازه زمانی مورد استفاده برای استخراج ویژگی و نمونه‌برداری از نمونه‌ها نام برد. برای ارزیابی عملکرد این پارامترها از مجموعه اعتبارسنجی استفاده شده است. در این بخش عملکرد پارامترهای زیر بر روی پایگاه داده برلین مورد بررسی قرار گرفته است.

- تعداد فیلتر CSP

- بازه فرکانسی

- بازه زمانی استفاده از نمونه‌ها

برای محاسبه پارامترهای بازه فرکانسی و تعداد فیلتر CSP از الگوریتم ژنتیک و برای بازه زمانی مورد

^۱ K-fold cross validation.

^۲ Repeated random sub sampling.

^۳ Leave one out cross validation.

استفاده از روش سعی و خطا استفاده شده است.

۴ - ۲ - ۱ پیاده‌سازی الگوریتم ژنتیک برای محاسبه پارامترها

الگوریتم ژنتیک روش مناسبی برای یافتن جواب مسائل بهینه‌سازی می‌باشد. در این پایان‌نامه، الگوریتم ژنتیک یک بار برای به دست آوردن پارامترهای اولیه سیستم مورد اجرا قرار گرفته است. پارامترهای در نظر گرفته شده برای الگوریتم ژنتیک در ادامه توضیح داده شده است.

الف - کروموزوم

پایگاه داده برلین شامل ۱۱۸ کانال EEG با فرکانس نمونه‌برداری ۱۰۰ هرتز می‌باشد. موارد زیر برای تعریف نمودن یک کروموزوم باید در نظر گرفته شود.

- تعداد فیلتر CSP بین ۱ تا ۱۱۸ باشد.
- بازه فرکانسی باید بین صفر تا ۵۰ هرتز تعریف شود. باید به این نکته توجه نمود که برای تعریف کروموزوم فرکانس پایین از فرکانس بالا کوچکتر باشد. (با توجه به تغییر مقادیر کروموزوم در

ترکیب و جهش)

ب- جمعیت اولیه

جمعیت اولیه در نظر گرفته شده برای این مسئله شامل ۵۰ کروموزوم می‌باشد که به صورت تصادفی تولید شده‌اند.

ج- عملگرهای مورد استفاده برای ترکیب، جهش و انتخاب

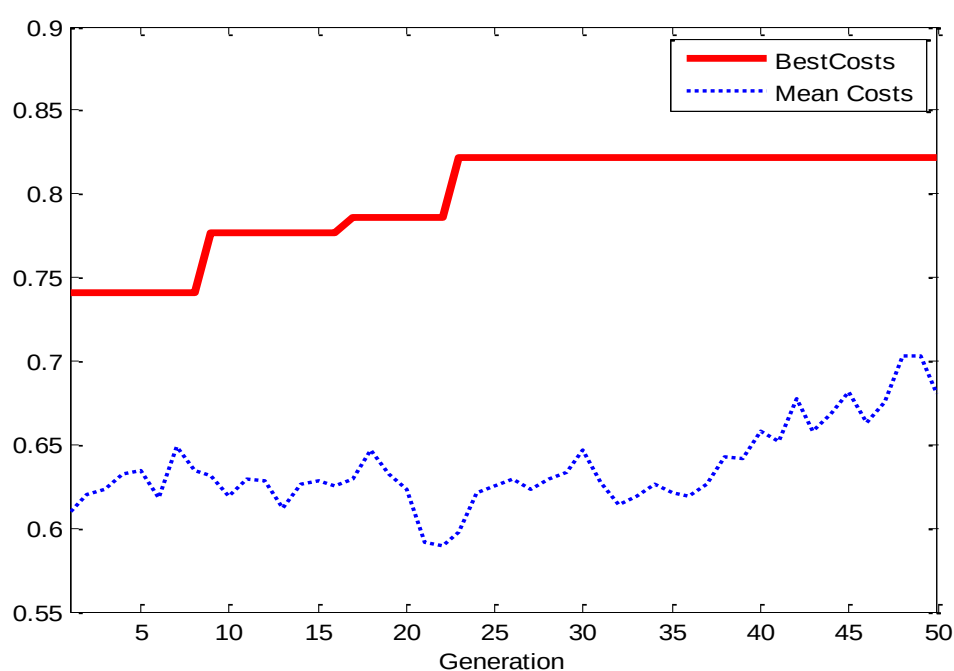
انتخاب کروموزوم‌ها برای تولید نسل آینده با استفاده از عملگر انتخاب نخبگان انجام شده است. برای ترکیب کروموزوم‌ها نیز از ترکیب یک نقطه‌ای و عملگر جهش نیز مقدار یک زن را به صورت تصادفی

تغییر می‌دهد.

ح- شرط توقف

شرایط در نظر گرفته شده برای پایان الگوریتم ژنتیک، یکی تعداد تکرار مشخص (۵۰ تکرار) و دیگری تغییر نیافتن بهترین کروموزوم پس از ۳۰ تکرار می‌باشد.

شکل ۱-۴ روند نتایج به دست آمده از الگوریتم ژنتیک را برای فرد «aa» نشان می‌دهد.



شکل ۱-۴. روند الگوریتم ژنتیک برای حل مسئله

همانطور که در شکل ۱-۴ مشاهده می‌شود بیشترین دقت برای مجموعه اعتبارسنجی نزدیک ۸۲٪ می‌باشد. جدول ۳-۴ نتایج برخی از کروموزوم‌ها را نشان می‌دهد.

جدول ۳-۴. نتایج برخی از کروموزوم‌ها

تعداد فیلتر CSP	بازه فرکانس (HZ)	دقت دسته‌بند (%)
۸	۱۳-۳۰	۷۴٪

۸۲٪	۱۲-۲۱	۴
۷۸٪	۱۳-۲۹	۴

با توجه به عملکرد الگوریتم ژنتیک روی داده‌های اعتبارسنجی تعداد فیلتر CSP برای فرد aa برابر ۸ در نظر گرفته شده است. همچنین بازه فرکانسی مورد نظر برای اعمال فیلتر میان‌گذر و استخراج ویژگی بین ۱۲ تا ۲۱ در نظر گرفته شده است. همانطور که گفته شد برای در نظر گرفتن بازه زمانی مناسب از آزمایش‌های مختلف استفاده شد که در جدول ۴-۴ برخی از نتایج آن آورده شده است. برای در نظر گرفتن بازه زمانی، یک تا دو ثانیه پس از علامت تصور حرکت ذهنی در نظر گرفته شده است.

جدول ۴-۴. نتایج برخی از آزمایش‌ها برای انتخاب بازه زمانی برای فرد aa

بازه زمانی (ثانیه) پس از نمایش علامت	میزان دقت دسته‌بند (%)
[۱ ۲]	۸۲٪
[۱ ۲٫۵]	۶۹٫۸۴٪
[۱٫۲ ۲]	۶۷٫۳۴٪
[۱٫۲ ۲٫۵]	۷۶٫۷۸٪
[۱ ۳]	۷۵٪

۴ - ۳ استخراج ویژگی و دسته‌بندی

در این بخش روش‌های استخراج ویژگی و دسته‌بندی استفاده شده در آزمایش‌های بخش‌های بعدی به صورت خلاصه توضیح داده شده است.

استخراج ویژگی گام مهم قبل از دسته‌بندی است که تأثیر زیادی در نحوه عملکرد سیستم BCI دارد. در این پایان‌نامه برای استخراج ویژگی از PSD استفاده شده است. مقایسه استخراج ویژگی با استفاده از ویژگی‌های آماری مانند میانگین و واریانس صورت گرفته است.

میانگین $\bar{x}$ برای سیگنال x به طول N به صورت فرمول زیر محاسبه می‌شود.

$$\bar{x} = \frac{1}{N} \sum_{i=1}^N x_i = \frac{1}{N} (x_1 + x_2 + \dots + x_N) \quad 2-4$$

مقدار واریانس برای سیگنال x با استفاده از معادله زیر به دست می‌آید.

$$S_{N-1}^2 = \frac{1}{N-1} \sum_{i=1}^N (x_i - \bar{x})^2 \quad 3-4$$

بنابراین برای مقایسه از دو ویژگی میانگین و واریانس استفاده شده است. این دو مقادیر مانند PSD برای دو پایگاه داده موجود به شکل زیر محاسبه شده‌اند:

- محاسبه شده از سیگنال بازه زمانی [۱,۲] پس از نمایش علامت برای پایگاه داده برلین

- محاسبه شده از سیگنال پنجره‌های به طول ۲۵۰ برای پایگاه داده کلرادو.

برای مقایسه روش‌های پیشنهادی با سایر دسته‌بندها؛ دسته‌بندهای KNN، SVM، LDA و MLP پیاده‌سازی شده است. به طور مختصر این چهار دسته‌بند را مورد بررسی قرار می‌دهیم.

در دسته‌بند KNN، برای محاسبه فاصله داده آزمون از داده‌های آموزشی معیارهای فاصله اقلیدسی^۱، چبیشف^۲، ماحالانوبیس^۳ و غیره استفاده می‌شود که در این مطالعه از فاصله اقلیدسی استفاده شده است.

^۱ Euclidean.

^۲ Mahalanobis.

^۳ Chebychev.

فاصله اقلیدسی فاصله معمولی دو نقطه است که توسط قضیه فیثاغورث به دست می‌آید. فاصله اقلیدسی دو نقطه $p = (p_1, p_2, \dots, p_n)$ و $q = (q_1, q_2, \dots, q_n)$ در فضای اقلیدسی n بعدی به صورت زیر محاسبه می‌شود.

$$\text{distance} = \sqrt{\sum_{i=1}^n (p_i - q_i)^2} \quad 4-4$$

۴ - ۴ بررسی نتایج ارزیابی نمایش تنک

در این قسمت به بررسی روش پیشنهادی دسته‌بندی با استفاده از نمایش تنک پرداخته شده است. برای ارزیابی روش، نتایج روی دو پایگاه داده دانشگاه کلرادو و برلین مورد بررسی قرار گرفته است. همچنین عملکرد برخی از پارامترها مانند به دست آوردن ضرایب تنک با استفاده از نرم‌های مختلف و انتخاب تعداد CSP نیز مورد تحلیل قرار گرفته است. برای مقایسه روش دسته‌بندی، سایر دسته‌بندها مانند LDA، KNN، SVM و MLP پیاده‌سازی شده و نتایج آن آورده شده است.

۴ - ۴ - ۱ پیاده‌سازی روی پایگاه داده کلرادو

نتایج دسته‌بندی با استفاده از روش نمایش تنک روی پایگاه داده کلرادو در جدول ۴-۵ نشان داده شده است. نتایج با استفاده از اعتبارسنجی با ۱۰ زیر نمونه محاسبه شده است.

جدول ۴-۵. نتایج روش نمایش تنک روی پایگاه داده کلرادو

تعداد کلاس				فرد
پنج کلاسه	چهار کلاسه	سه کلاسه	دو کلاسه	

فرد اول	۹۰٪	۸۳٫۵٪	۷۶٫۵٪	۷۰٪
فرد دوم	۷۴٪	۵۹٫۳۳٪	۵۰٪	۴۴٪
فرد سوم	۷۳٪	۵۷٪	۴۷٪	۴۰٪
فرد چهارم	۵۲٫۵٪	۴۳٫۶٪	۳۵٫۳۳٪	۳۰٫۵٪
فرد پنجم	۵۰٪	۴۲٫۶٪	۳۱٪	۲۷٫۶٪
فرد ششم	۷۵٪	۵۹٫۶۷٪	۵۰٪	۴۴٪
فرد هفتم	۸۳٪	۷۳٫۳۳٪	۶۶٪	۶۰٪

نتایج این روش برای پایگاه داده فرد اول مناسب می‌باشد و برای افراد دیگر در حد متوسط می‌باشد که یکی از دلایل آن را می‌توان استخراج ویژگی نامناسب برای افراد دیگر دانست. همانطور که اشاره شد این پایگاه داده شامل ۵ فعالیت ذهنی می‌باشد. نتایج آورده شده در جدول ۴-۵ میانگین تمامی حالات برای دو کلاسه، سه کلاسه و چهار کلاسه می‌باشد. جدول ۴-۶ نتایج همه حالات برای دسته‌بندی دو کلاس از پایگاه داده کلرادو را نشان می‌دهد. نتایج با استفاده از اعتبارسنجی با ۱۰ زیر نمونه‌برداری به دست آمده است. کلاس‌های حالت استراحت، ضرب ذهنی، نامه‌نگاری ذهنی، دوران ذهنی یک شی و شمارش ذهنی به ترتیب با B، M، L، R و C نشان داده شده است. با توجه به نتایج میانگین دسته‌بندی برای حالت استراحت با سایر کلاس‌ها برابر با ۸۷٫۵٪ می‌باشد. میانگین دسته‌بندی کلاس‌های ضرب ذهنی، نامه‌نگاری ذهنی، دوران و شمارش ذهنی با دیگر کلاس‌ها به ترتیب برابر با ۹۲٫۵٪، ۹۲٫۵٪، ۹۱٫۲۵٪ و ۸۶٫۲۵٪ می‌باشد.

جدول ۴-۶. نتایج حالات مختلف دسته‌بندی دو کلاسه برای فرد اول

فعالیت‌های ذهنی	درصد درستی	فعالیت‌های ذهنی	درصد درستی
B ,M	۷۵٪	M ,R	۹۵٪

۱۰۰٪	M ,C	۱۰۰٪	B ,L
۹۰٪	L ,R	۹۵٪	B, R
۸۰٪	L ,C	۸۰٪	B ,C
۸۵٪	R ,C	۱۰۰٪	M ,L

در جدول ۷-۴ نتایج حالات مختلف دسته‌بندی سه کلاس برای فرد اول آورده شده است. بهترین دسته‌بندی مربوط به کلاس‌های ضرب و دوران و نام‌نگاری ذهنی و همچنین حالت استراحت، دوران و نام‌نگاری ذهنی می‌باشد.

جدول ۷-۴. نتایج حالات مختلف دسته‌بندی سه کلاس برای فرد اول

درصد درستی	فعالیت‌های ذهنی	درصد درستی	فعالیت‌های ذهنی
۷۶٫۶۷٪	B, R, C	۸۳٫۳۳٪	B ,M, L
۹۳٫۳۳٪	M, L, R	۸۰٪	B ,M, R
۸۶٫۶۷٪	M, L, C	۷۰٪	B, M, C
۹۰٪	M, R, C	۹۳٫۳۳٪	B ,L, R
۷۶٫۶۷٪	L, R, C	۷۶٫۶۷٪	B ,L, C

۴ - ۴ - ۱ - ۱ مقایسه روش با سایر دسته‌بندها

نتایج روی پایگاه داده فرد اول با سایر دسته‌بندها مقایسه شده است. نتایج در جدول ۸-۴ آورده شده است. دسته‌بندی با استفاده از نمایش تنک به مراتب از سایر دسته‌بندها نتایج بهتری را ارائه داده است. همچنین این دسته‌بند در مقایسه با دسته‌بند MLP زمان کمتری را برای اجرا دارد و با دسته‌بندهای دیگر مشابه می‌باشند. عملکرد دسته‌بند SVM از دسته‌بندهای KNN، MLP و LDA مناسب‌تر بوده و

دقت بهتری را برای دسته‌بندی انجام داده است.

جدول ۴-۸. مقایسه روش پیشنهادی نمایش تنک با سایر دسته‌بندها

دسته‌بند	دو کلاسه	سه کلاسه	چهار کلاسه	پنج کلاسه
روش پیشنهادی	۹۰٪	۸۳٫۵٪	۷۶٫۵٪	۷۰٪
KNN	۸۷٫۵٪	۷۷٪	۶۵٪	۵۲٪
SVM	۸۹٫۵٪	۷۹٫۳۳٪	۶۹٫۵٪	۶۲٪
MLP	۸۶٪	۶۵٫۶٪	۵۷٪	۵۲٪
LDA	۸۸٫۵٪	۷۲٪	۵۶٫۵٪	۴۲٪

۴ - ۴ - ۱ - ۲ بررسی عملکرد روش‌های بازسازی تنک

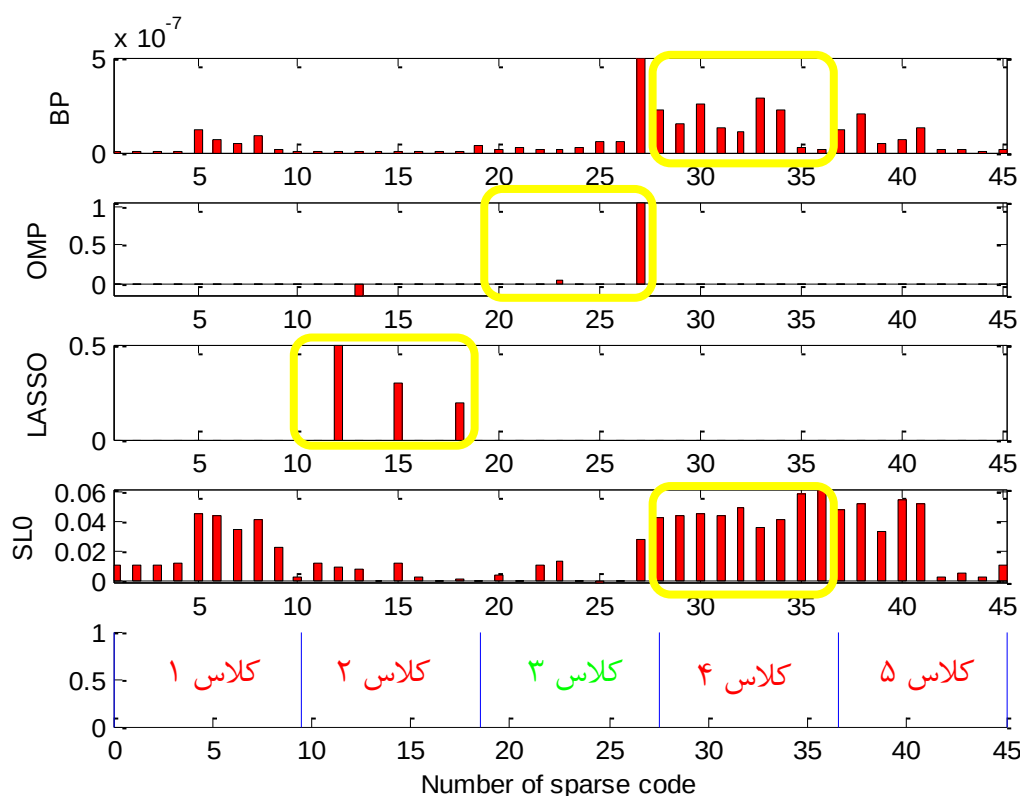
برای ارزیابی دسته‌بندی با استفاده از نمایش تنک که از الگوریتم‌های مختلف برای یافتن پاسخ تنک استفاده می‌کند، میانگین دقت دسته‌بندی بر روی پایگاه داده فرد اول در جدول ۴-۹ آورده شده است. الگوریتم‌های BP و LASOO از خانواده محاسبه پاسخ تنک با استفاده از نرم یک، الگوریتم OMP از دسته الگوریتم‌های حریص و SL0 نیز از دسته تخمین نرم صفر برای محاسبه پاسخ تنک می‌باشند. نتایج به دست آمده نشان می‌دهد که الگوریتم OMP نسبت به الگوریتم‌های دیگر دارای نرخ دقت بالاتری می‌باشد. همچنین این الگوریتم زمان اجرای کمتری را نسبت به سایر روش‌های دیگر نیاز دارد.

جدول ۴-۹. بررسی عملکرد روش‌های بازسازی تنک

	دو کلاسه	سه کلاسه	چهار کلاسه	پنج کلاسه	زمان اجرا (ثانیه)
روش BP	۹۰٪	۸۳٫۵٪	۷۶٫۵٪	۷۰٪	۳۲٫۹۵

روش OMP	۹۲٪	۸۶٪	۸۰٪	۷۴٪	۹,۷۸
روش LASSO	۸۹,۵٪	۸۲,۳۳٪	۷۶,۵٪	۷۲٪	۱۰,۵۳
روش SL0	۸۶,۵٪	۷۴,۶۷٪	۶۵,۵٪	۶۰٪	۱۰,۱۵

برای بررسی عملکرد روش‌های BP، OMP، LASSO و SL0 مثالی روی یک نمونه مربوط به کلاس ۳ در شکل ۲-۴ آورده شده است. خروجی پاسخ تنک مربوط به الگوریتم‌های ذکر شده آورده شده است. کلاس این نمونه توسط روش‌های BP و SL0 برابر با چهار و به وسیله روش LASSO دو پیش‌بینی شده است. کلاس این نمونه توسط روش OMP به درستی تشخیص داده شده است.



شکل ۲-۴. بررسی عملکرد روش‌های به دست آوردن پاسخ تنک

۴ - ۱ - ۳ مقایسه روش با سایر کارهای انجام شده

در این بخش نتایج روش دسته‌بندی نمایش تنک با دیگر کارهای انجام شده روی پایگاه داده کلرادو مقایسه شده است. Zhang و همکاران [۴۷] با محاسبه توان طیفی برای باندهای دلتا، تتا، بتا، آلفا و گاما و همچنین بازه فرکانسی بیشتر از گاما، نرخ عدم تقارن توان طیفی محاسبه شده برای باندهای مختلف بین کانال‌های سمت راست و چپ سر را به عنوان ویژگی انتخاب نموده است. برای دسته‌بندی نیز آنالیز افتراقی خطی^۱ استفاده شده است. Xiaoo و همکاران در [۱۹] با ارائه یک دسته‌بند SVM چند هسته‌ای به دسته‌بندی سیگنال EEG پرداخته‌اند. در این مقاله از تبدیل موجک بسته‌ای برای استخراج باندها استفاده شده است. ویژگی انتخاب شده برای دسته‌بندی نیز میزان آنتروپی برای هر بسته موجک در نظر گرفته شده است. نتایج مقایسه روش پیشنهادی با دو کار مطرح شده در جدول ۴-۱۰ آورده شده است.

جدول ۴-۱۰. مقایسه روش دسته‌بندی نمایش تنک با کارهای گذشته

	دو کلاسه	سه کلاسه	چهار کلاسه	پنج کلاسه
روش پیشنهادی	۹۰٪	۸۳٫۵٪	۷۶٫۵٪	۷۰٪
Zhang و همکاران [۴۷]	۸۵٫۹٪	۷۵٫۳٪	۶۶٫۶٪	۶۴٫۴٪
Xiaoo و همکاران [۱۹]	۹۸٫۲۴٪	۷۳٫۹۵٪	۷۳٫۱۴٪	۶۵٫۱۵٪

نتایج آورده شده بیانگر این موضوع است که روش پیشنهادی در مقایسه با روش Zhang در دسته‌بندی‌های مختلف نتایج بهتری را ارائه داده است. روش پیشنهادی در مقایسه با روش Xiaoo غیر از دسته‌بندی دو کلاسه، برای دسته‌بندی دو، سه و چند کلاسه نتایج بهتری را ارائه می‌دهد.

^۱ Fisher discriminant analysis.

۴ - ۴ - ۲ پیاده‌سازی روی پایگاه داده برلین

در این بخش نتایج اعمال روش را بر روی پایگاه داده برلین مورد بررسی قرار می‌دهیم. نتایج به دست آمده از این روش بر روی پایگاه داده با اعتبارسنجی LOO در جدول ۴-۱۱ آورده شده است. برای این آزمایش موارد زیر اعمال شده است:

- ویژگی استخراجی: واریانس و PSD

- استفاده از روش BP برای به دست آوردن ضرایب تنک

جدول ۴-۱۱. نتایج روش دسته‌بندی تنک بر روی پایگاه داده برلین

فرد	aa	Al	av	aw	ay	میانگین
درصد	۹۸٫۲۱٪	۹۸٫۹۲٪	۹۴٫۲۸٪	۹۶٫۴۲٪	۱۰۰٪	۹۷٫۵۶٪
درستی						

۴ - ۴ - ۲ - ۱ مقایسه روش با سایر دسته‌بندها



شکل ۳-۴. مقایسه روش با دسته‌بندی‌های KNN، SVM، MLP و LDA

۴ - ۲ - ۲ - بررسی عملکرد استخراج ویژگی‌های مختلف

در جدول ۱۲-۴ نتایج عملکرد استخراج ویژگی‌های مختلف مانند میانگین و واریانس آورده شده است. همانطور که مشاهده می‌شود ترکیب دو ویژگی PSD و واریانس منجر به بهترین عملکرد دسته‌بندی می‌شود.

جدول ۱۲-۴. بررسی عملکرد استخراج ویژگی‌های مختلف

فرد	aa	al	av	aw	ay	میانگین
PSD + واریانس	۹۸٫۲۱٪	۹۸٫۹۲٪	۹۴٫۲۸٪	۹۶٫۴۲٪	۱۰۰٪	۹۷٫۵۶٪
PSD	۹۸٫۹۳٪	۱۰۰٪	۹۵٫۷۱٪	۹۷٫۸۶٪	۹۱٫۷۹٪	۹۶٫۸۵٪
واریانس	۹۸٫۲۱٪	۹۸٫۹۲٪	۹۳٫۹۲٪	۹۶٫۴۲٪	۹۶٫۴۲٪	۹۶٫۷۷٪
میانگین	۹۷٫۱۴٪	۹۴٫۶۴٪	۹۲٫۸۸٪	۹۴٫۶۴٪	۹۷٫۵٪	۹۴٫۵۶٪

۴ - ۲ - ۳ بررسی عملکرد روش‌های بازسازی تنک

برای ارزیابی دسته‌بندی با استفاده از نمایش تنک که از الگوریتم‌های مختلف برای یافتن پاسخ تنک استفاده می‌کند، میانگین دقت دسته‌بندی بر روی پایگاه داده فرد اول در جدول ۴-۱۳ آورده شده است. با توجه به نتایج مشاهده می‌شود که روش‌های SL0 و LASSO و BP عملکردی تقریباً مشابه دارند و نسبت به OMP بهتر عمل نموده‌اند.

جدول ۴-۱۳. مقایسه عملکرد روش‌های بازسازی تنک

میانگین	ay	aw	av	al	aa	
۹۷/۵۶٪	۱۰۰٪	۹۶/۴۲٪	۹۴/۲۸٪	۹۸/۹۲٪	۹۸/۲۱٪	روش BP
۹۵/۹۹٪	۹۸/۹۲٪	۹۵/۷۱٪	۹۰/۷۱٪	۹۸/۵۷٪	۹۶/۰۷٪	روش OMP
۹۷/۶۳٪	۹۹/۶۴٪	۹۶/۷۸٪	۹۳/۹۲٪	۹۸/۹۲٪	۹۸/۹۲٪	روش LASSO
۹۷/۵٪	۹۹/۶۴٪	۹۷/۵٪	۹۵٪	۹۸/۵۷٪	۹۷/۱۴٪	روش SL0

۴ - ۴ - ۳ بررسی ویژگی‌های دیکشنری

یکی از مهمترین ویژگی‌های یک دیکشنری را می‌توان میزان عدم وابستگی اتم‌های ایجاد کننده دیکشنری دانست. هر چه اتم‌های دیکشنری وابستگی کمتری به هم داشته باشند (در حالت ایده‌آل ضرب داخلی آنها صفر است) یعنی که اتم‌ها دیکشنری ویژگی‌های متفاوت‌تری را پوشش می‌دهند؛ به عبارت دیگر حداقل وابستگی بین اتم‌های دیکشنری منجر به پوشش هر چه بهتر و بیشتر فضای داده‌ها می‌باشد. میزان همبستگی بین اتم‌های دیکشنری را می‌توان با استفاده از ضریب همبستگی محاسبه نمود. ضریب همبستگی برای هر جفت از اتم‌های دیکشنری با استفاده از رابطه ۴-۵ محاسبه می‌شود.

$$R(x,y) = \frac{\text{cov}(x,y)}{\sqrt{\text{cov}(x,x) \times \text{cov}(y,y)}} \quad ۵-۴$$

در این رابطه x و y دو اتم غیر مشترک از دیکشنری می‌باشند و تابع $\text{cov}(x,y)$ کواریانس بین دو اتم x و y را محاسبه می‌نماید. مقدار کوچک برای R نشان‌دهنده میزان همبستگی کم بین دو اتم x و y می‌باشد. R برابر با یک بیانگر بیشترین میزان همبستگی بین اتم‌ها و R برابر صفر عمود بودن دو اتم را بر هم نشان می‌دهد.

جدول ۴-۱۴. مقادیر همبستگی برای پایگاه داده‌های برلین

پایگاه داده	aa	al	av	aw	ay
میانگین R	۰٫۲۴	۰٫۳۳	۰٫۲۸	۰٫۲۶	۰٫۲۸
مینیمم R	$۴/۳۵ \times ۱۰^{-۶}$	$۳/۳۹ \times ۱۰^{-۶}$	$۱/۴۴ \times ۱۰^{-۶}$	$۱/۸۳ \times ۱۰^{-۵}$	$۱/۲۶ \times ۱۰^{-۵}$
ماکزیمم R	۰٫۹۳	۰٫۹۵	۰٫۹۹	۰٫۹۸	۰٫۹۲

- بررسی عملکرد فیلتر CSP در ساخت دیکشنری

در این قسمت به بررسی نقش الگوریتم CSP در دیکشنری ساخته شده پرداخته شده است. همانطور که ذکر شد؛ میزان همبستگی اتم‌های دیکشنری در کیفیت آن نقش مؤثری دارد. در جدول ۴-۱۵ میانگین مقادیر همبستگی بین اتم‌های دیکشنری با استفاده از CSP و عدو استفاده از آن آورده شده است.

جدول ۴-۱۵. بررسی عملکرد فیلتر CSP در ساخت دیکشنری

پایگاه داده	aa	al	av	aw	ay
میانگین R با استفاده از CSP	۰٫۲۴	۰٫۳۳	۰٫۲۸	۰٫۲۶	۰٫۲۸

۰٫۵۱	۰٫۵۱	۰٫۲۶	۰٫۵۲	۰٫۵۱	مینیمم R با بدون استفاده از CSP
------	------	------	------	------	---------------------------------

۴-۵ بررسی عملکرد روش DPL

در این بخش به بررسی عملکرد روش DPL که در فصل ۳ بخش ۳-۵-۲ توضیح داده شده است روی پایگاه داده برلین می‌پردازیم. آزمایش‌ها بر روی کامپیوتر با مشخصات Intel core i3 با پردازنده ۲٫۲۷ گیگا هرتز و حافظه ۶ گیگا بایت با نرم‌افزار MATLAB R2014a انجام شده است.

۴-۵-۱ پیاده‌سازی روش روی پایگاه داده برلین

برای این آزمایش از ۱۱۸ کانال و نمونه‌های بین یک تا دو ثانیه پس از علامت انجام تصور ذهنی استفاده شده است. همچنین از فیلتر میان‌گذر ۱۱ تا ۳۶ هرتز برای حذف مصنوعات و نویزها استفاده شده است که این بازه شامل اطلاعات مفید دو باند آلفا و بتا می‌باشد. فیلتر CSP برای افزایش اطلاعات دوطرفه بین کلاس‌ها قابل اعمال است. به دست آوردن تعداد فیلتر مناسب برای این روش در بخش ۴-۲-۱ با استفاده از الگوریتم ژنتیک محاسبه شد.

روش DPL روی پایگاه داده al که بیشترین داده آموزشی را دارد بهترین عملکرد را دارد. میانگین دقت دسته‌بندی با استفاده از روش DPL برابر با ۷۸٫۲۹٪ می‌باشد. نتایج به دست آمده از روش DPL در جدول ۴-۱۶ آورده شده است. نتایج روی پایگاه داده آموزشی و آزمون مشخص شده برای این پایگاه داده محاسبه شده است.

جدول ۴-۱۶. نتایج روش DPL روی پایگاه داده برلین

ay	aw	av	al	aa	پایگاه داده
۶۷/۸۵٪	۷۵/۴۴٪	۶۴/۲۸٪	۱۰۰٪	۸۳/۹۲٪	درصد درستی

۴ - ۵ - ۲ بررسی نتایج برای اندازه‌های مختلف دیکشنری

در جدول ۴-۱۷ میزان دقت دسته‌بند به ازای اندازه‌های مختلف دیکشنری آورده شده است. همانطور که در جدول ۴-۱۷ مشاهده می‌شود؛ میزان دقت به دست آمده برای اندازه‌های کوچک دیکشنری نیز مناسب می‌باشد و این عملکرد را می‌توان یکی از پارامترهای خوب برای دیکشنری دانست؛ به عبارت دیگر الگوریتم به ابعاد انتخابی دیکشنری حساس نمی‌باشد

جدول ۴-۱۷. میزان دقت دسته‌بند برای اندازه‌های مختلف دیکشنری

اندازه دیکشنری							پایگاه داده aa
۴۰	۲۰	۱۸	۱۶	۱۴	۱۲	۱۰	
۸۳/۹۲٪	۸۳/۹۲٪	۸۳/۹۲٪	۸۳/۹۲٪	۸۰/۳۵٪	۷۶/۷۸٪	۷۵/۸۹٪	میزان دقت

۴ - ۵ - ۳ مقایسه روش با سایر کارهای انجام شده

در این بخش به مقایسه نتایج روش DPL با کارهای انجام شده روی پایگاه داده برلین پرداخته شده است. Talukdar و همکاران در [۲۰] با استفاده از آنتروپی محاسبه شده از IMP به دسته‌بندی سیگنال EEG پرداخته‌اند. در این مقاله از روش PCA برای کاهش ویژگی‌ها استفاده شده و از دسته‌بند LDA برای دسته‌بندی ویژگی‌های استخراج شده استفاده شده است. نویسندگان در [۲۱] یک الگوریتم بر

مبنای بیز برای دسته‌بندی سیگنال‌های EEG پایگاه داده دانشگاه برلین استفاده شده است. از SVM برای دسته‌بندی ویژگی‌های هر باند استفاده شده و در پایان با توجه به اهمیت هر باند خروجی نهایی مشخص می‌شود. Zhou و همکاران در [۲۲] یک روش برای جداسازی سیگنال EEG بر مبنای نمایش تنک مطرح شده است. در این مقاله یک روش یادگیری دیکشنری برای ساخت یک دیکشنری با اندازه کوچکتر و قدرت تفکیک‌پذیری بالاتر ارائه شده است. برای جداسازی هر چه بهتر کلاس‌ها از الگوریتم الگوی مکانی مشترک (CSP) زیادی استفاده شده است.

نتایج مقایسه روش DPL با روش‌های مطرح شده در این قسمت در جدول ۴-۱۸ آورده شده است. نتایج به دست آمده برای دو فرد aa و al به مراتب از سایر روش‌ها بهتر بوده است. نتیجه برای فرد al برابر با ۱۰۰٪ با استفاده از روش DPL به دست آمده است. به طور میانگین روش ارائه شده با روش ارائه شده توسط Talukdar برابر و از دو روش دیگر بیشتر می‌باشد.

جدول ۴-۱۸. مقایسه روش DPL با کارهای گذشته

روش	aa	al	av	Aw	ay
Talukdar و همکاران [۲۴]	۷۲٫۶۴٪	۷۵٫۱۷٪	۷۳٫۳۶٪	۸۸٫۵۱٪	۸۳٫۵۹٪
Suk و همکاران [۴۸]	۷۹٫۴۶٪	۹۴٫۶۴٪	۵۷٫۶۵٪	۹۱٫۹۶٪	۵۳٫۳۵٪
Zhou و همکاران [۲۵]	۷۳٫۲۰٪	۹۴٫۶۰٪	-	-	-
روش پیشنهادی	۸۳٫۹۲٪	۱۰۰٪	۶۴٫۲۸٪	۷۵٫۴۴٪	۶۷٫۸۵٪

۴ - ۵ - ۴ مقایسه روش با سایر روش‌های یادگیری دیکشنری

در بخش ۳ - ۴ برخی از روش‌های یادگیری دیکشنری مطرح شد. در این بخش به مقایسه روش DPL

با روش‌های LC-KSVD1 و LC-KSVD2 پرداخته شده است. مقایسه به لحاظ زمانی و میزان دقت دسته‌بند انجام شده است.

با توجه به نتایج آورده شده در جدول ۴-۱۹ از مزایای روش DPL نسبت به دو روش دیگر می‌توان به پارامتر زمان اجرای آموزشی و آزمون کمتر اشاره نمود. پارامتر زمان اجرای کمتر نقش مؤثر و مهمی در سیستم BCI ایفا می‌کند. همچنین روش DPL به مراتب نتایج مناسب‌تری را نسبت به روش LCKSVD1 ارائه کرده است و نتایج آن از روش یادگیری دیکشنری LC-KSVD2 نیز بیشتر می‌باشد.

جدول ۴-۱۹. مقایسه روش با سایر روش‌های یادگیری دیکشنری

روش یادگیری دیکشنری	روش DPL	روش LC-KSVD1	روش LC-KSVD2
Aa	میزان دقت (%)	۸۳٫۹۲٪	۷۱٫۴۲٪
	زمان آموزش (ثانیه)	۰٫۳۷۳۳	۴٫۴۸
	زمان آزمون (ثانیه)	۰٫۰۰۰۲۸	۰٫۰۰۰۶۷
al	میزان دقت (%)	۱۰۰٪	۶۲٫۵٪
	زمان آموزش (ثانیه)	۰٫۴۲	۶٫۲۳
	زمان آزمون (ثانیه)	۰٫۰۰۰۲۶	۰٫۰۰۰۳۳
av	میزان دقت (%)	۶۴٫۲۸٪	۴۶٫۴۲٪
	زمان آموزش (ثانیه)	۰٫۳۵	۱٫۸۸
	زمان آزمون (ثانیه)	۰٫۰۰۰۳	۰٫۰۰۰۹۸
aw	میزان دقت (%)	۷۵٫۴۴٪	۴۷٫۶۷٪
	زمان آموزش (ثانیه)	۰٫۲۶	۱٫۱۸
	زمان آزمون (ثانیه)	۰٫۰۰۰۳	۰٫۰۰۰۸۷
ay	میزان دقت (%)	۶۷٫۸۶٪	۴۹٫۶٪

زمان آموزش (ثانیه)	۰/۱۱	۰/۶۳	۰/۶۳
زمان آزمون (ثانیه)	۰/۰۰۰۷	۰/۰۱۰۴	۰/۰۰۹۷

۴ - ۵ - ۵ بررسی ویژگی‌های دیکشنری

در جدول ۴-۲۰ میزان همبستگی دیکشنری به دست آمده از سه روش DPL، LC-KSVD1 و LC-KSVD2 آورده شده است. مقادیر مینیمم، ماکزیمم و میانگین ماتریس همبستگی R در این جدول ذکر شده است. با در نظر گرفتن مقادیر مشاهده می‌شود که به طور میانگین اتم‌های دیکشنری ساخته شده توسط روش DPL نسبت به دو روش دیگر دارای همبستگی پایین‌تری می‌باشند؛ در نتیجه کیفیت دیکشنری ساخته شده به وسیله این روش بالاتر از روش‌های دیگر مطرح شده می‌باشد.

جدول ۴-۲۰. میزان همبستگی دیکشنری ساخته شده با روش‌های DPL، LC-KSVD1 و LC-KSVD2

روش	ماکزیمم R	مینیمم R	میانگین R
LC-KSVD1	۰/۹۹	۰/۰۰۰۳۶	۰/۵۳
LC-KSVD2	۰/۹۹	۰/۰۰۰۱۰۲	۰/۵۲
DPL	۰/۹۴	۰/۰۰۰۱۵	۰/۳۱

فصل ۵

نتیجه‌گیری و پیشنهادها

در این پایان نامه تشخیص فعالیت های ذهنی در BCI با استفاده از سیگنال EEG مورد بررسی قرار گرفت. BCI یک سیستم ارتباطی است که به افراد توانایی کنترل ابزار و وسایل محیط اطراف را با استفاده از فعالیت های عصبی مغز می دهد. برنامه های کاربردی وسیعی را می توان برای افزایش کیفیت زندگی افراد سالم و افراد دارای معلولیت جسمانی با استفاده از سیستم BCI ارائه نمود. سیگنال EEG به دلیل ارزان قیمت بودن، تفکیک پذیری زمانی مناسب و غیرتهاجمی بودن یکی از پرکاربردترین روش های ثبت سیگنال مغز در تحقیقات BCI می باشد. مراحل اصلی یک سیستم BCI عبارتند از ثبت سیگنال، پیش پردازش، استخراج ویژگی و دسته بندی. یکی از چالش های سیستم BCI پردازش سیگنال می باشد که در عملکرد آن نقش بسزایی دارد. هر چه این مرحله دارای دقت و سرعت بالاتری باشد؛ عملکرد سیستم BCI را می تواند بهبود دهد و امکان استفاده از آن را در دنیای واقعی آسانتر نماید. در این پایان نامه روش های دسته بندی جدیدی بر مبنای نمایش تنک و یادگیری دیکشنری برای سیگنال های EEG ارائه شد و مورد بررسی قرار گرفت.

در فصل اول چالش ها و انگیزه های استفاده از یک سیستم BCI مورد بررسی قرار گرفت. در این فصل کاربردهای سیستم BCI و تشخیص فعالیت های ذهنی نیز مطرح شد. در فصل دوم ساختار مغز و قسمت های مختلف آن به صورت مختصر آورده شد و در ادامه این فصل سیگنال EEG، اجزای یک سیستم BCI و کارهای گذشته انجام شده مرتبط با سیستم BCI مورد بررسی قرار گرفت. در فصل سوم روش پیشنهادی ارائه شد و در فصل چهارم نتایج آن مورد ارزیابی قرار گرفت.

نتایج الگوریتم های ارائه شده بر روی دو پایگاه داده برلین و کلرادو مورد ارزیابی قرار گرفت. پایگاه داده کلرادو شامل پنج فعالیت ذهنی استراحت، ضرب ذهنی، نوشتن نامه ذهنی، شمارش ذهنی و چرخش ذهنی یک شکل سه بعدی می باشد. پایگاه داده برلین شامل تصور ذهنی حرکت دست چپ و راست ثبت شده از پنج نفر می باشد.

برای به دست آوردن پارامترهای اولیه از الگوریتم ژنتیک استفاده شد. برای الگوریتم ژنتیک و محاسبه

پارامترهای اولیه از داده‌های آموزشی و اعتبارسنجی استفاده شد و با توجه با اینکه برای هر پایگاه داده تنها یک بار انجام شد؛ مشکل زمانی برای مرحله آزمون وجود نداشت. برای پیش‌پردازش و حذف نویز از فیلتر میان‌گذر استفاده شد. سپس از فیلتر CSP برای افزایش بهره اطلاعاتی بین کلاس‌های مختلف استفاده گردید. برای اعمال فیلتر CSP بر روی چند کلاس از الگوریتم JAD استفاده شد و سپس بهره اطلاعاتی برای فیلترهای به دست آمده مورد محاسبه قرار گرفت. بررسی عملکرد CSP در فصل چهار بیانگر مناسب بودن آن برای دسته‌بندی سیگنال EEG می‌باشد. همچنین اعمال CSP کیفیت دیکشنری ساخته شده را افزایش داد. پس از اعمال CSP، استخراج ویژگی انجام شد و دیکشنری مورد نظر برای کلاس‌های مختلف ساخته شد. سپس دو رویکرد برای ادامه راه مطرح شد. اولین روش پیشنهادی استفاده از نمایش تنک برای دسته‌بندی بدون یادگیری دیکشنری بود. روش پیشنهادی دیگر برای دسته‌بندی استفاده از یادگیری دیکشنری DPL بود.

روش پیشنهادی دسته‌بندی با استفاده از نمایش تنک؛ دیکشنری ساخته شده از ویژگی‌های استخراجی را برای مرحله آزمون استفاده می‌نماید. برای محاسبه پاسخ تنک از روش‌هایی مانند BP، OMP، LASSO و SL0 می‌توان استفاده نمود که در فصل چهارم عملکرد آنها مورد بررسی قرار گرفت. نتایج این روش پیشنهادی روی پایگاه داده کلرادو برای دسته‌بندی پنج فعالیت ذهنی و همه حالات دو، سه و چهار کلاسه انجام شد. دسته‌بندی ارائه شده بر مبنای پاسخ تنک در مقایسه با دسته‌بندهای SVM، KNN، LDA و MLP نتایج مناسبتری را ارائه می‌دهد. همچنین این روش با اعتبارسنجی یک نمونه بیرون روی پایگاه داده برلین اعمال شد و کارایی آن مورد بررسی قرار گرفت. نتایج به دست آمده بیانگر مناسب بودن روش برای دسته‌بندی سیگنال‌های EEG می‌باشد.

روش پیشنهادی دیگر ارائه شده بر مبنای یادگیری دیکشنری بود. دیکشنری نقش مهمی در فرایند بازسازی سیگنال با استفاده از نمایش تنک ایفا می‌کند. یادگیری دیکشنری کیفیت اتم‌های دیکشنری را افزایش داده و در نتیجه با بهبود دیکشنری می‌توان نتایج مناسبی برای دسته‌بندی حاصل شود. روش

ارائه شده بر اساس یادگیری دو دیکشنری همزمان می‌باشد که یک دیکشنری مصنوعی حاوی برخی ویژگی‌های داده‌های آموزشی با داشتن ویژگی‌های افتراق‌پذیری و دیکشنری تحلیلی برای محاسبه پاسخ تنک مورد استفاده قرار می‌گیرد. استفاده از دیکشنری تحلیلی برای محاسبه پاسخ تنک، نیاز به استفاده از الگوریتم‌های غیرخطی مانند BP، OMP را رفع نموده و پاسخ تنک را با استفاده از یک ضرب محاسبه می‌نماید. ارزیابی این روش بر روی پایگاه داده برلین انجام شد و نتایج بیانگر بهتر بودن روش نسبت به سایر روش‌های یادگیری دیکشنری (LC-KSVD1 و LCKSVD2) از جنبه دقت و زمان بود. همچنین کیفیت دیکشنری ایجاد شده توسط این روش نسبت به سایر روش‌ها مناسب‌تر بود.

برای ادامه کار پیشنهاد‌های زیر قابل انجام می‌باشد:

- در این پژوهش تکنیک یادگیری دیکشنری برای دسته‌بندی استفاده شد. ساخت دیکشنری با استفاده از تابع هدف انجام شد. برای بهبود کیفیت اتم‌های دیکشنری مرتبط با دسته‌بندی می‌توان پارامترهایی نظیر کمینه کردن خطای دسته‌بندی را به تابع هدف اضافه نمود.
- در این پژوهش برای مرحله استخراج ویژگی از پارامتر PSD استفاده شد و نتایج مناسب به دست آمد اما برای برخی از پایگاه داده‌ها دقت دسته‌بندی قابل بهبود می‌باشد. می‌توان برای ادامه کار استخراج ویژگی‌های مختلفی مانند AR و غیره را استفاده نمود.
- روش DPL برای پایگاه داده‌های دارای داده آموزشی زیاد جواب مناسبی ارائه می‌دهد ولی پاسخ آن برای پایگاه داده با داده آموزشی کم قابل بهبود می‌باشد. می‌توان با بررسی و پژوهش در این زمینه نواقص آن را برطرف نمود.
- برای تعریف تابع هدف برخی از ضرایب و پارامترها مانند اندازه دیکشنری با استفاده از سعی و خطا تنظیم شده است. برای ادامه کار پیشنهاد می‌شود که با توجه به اندازه داده‌های آموزشی و میزان شباهت اتم‌های دیکشنری برخی پارامترها را به صورت سازگار محاسبه نمود.

- برای بهبود عملکرد دیکشنری و جداسازی هر چه بهتر کلاس‌ها روش‌هایی را می‌توان برای حذف اتم‌های مشترک در کلاس‌های مختلف ارائه نمود.

مراجع

- [1] J. C. Lee and D. S. Tan, (2006), "Using a low-cost electroencephalograph for task classification in HCI research", Proceedings of the 19th annual ACM symposium on User interface software and technology, pp. 81–90, New York, USA.
- [2] M. Nadin Sergio, Albeverio, V. Jentsch, H. Kantz, B. Graimann, B. Z. Allison, and G. Pfurtscheller, (2010), "**Brain-computer interfaces: Revolutionizing human-computer interaction**", Vol. 1, Springer-Verlag Berlin Heidelberg, Berlin, pp. 21–45.
- [3] R. Chai, S. H. Ling, G. P. Hunter, Y. Tran, and H. T. Nguyen, (2013), "Classification of wheelchair commands using brain computer interface: comparison between able-bodied persons and patients with tetraplegia", Engineering in Medicine and Biology Society (EMBC), 35th Annual International Conference of the IEEE , pp. 989–992. Osaka, Japan.
- [4] G. Pfurtscheller, R. Leeb, C. Keinrath, D. Friedman, C. Neuper, C. Guger, and M. Slater, (2006), "Walking from thought", **Brain Res.**, **1071**, **1**, pp. 145–152.
- [5] A. Nijholt and D. F. D. Tan, (2008), "Brain-computer interfacing for intelligent systems", **Intell. Syst. IEEE**, **23**, **3**, pp. 72–79.

- [6] J. van Erp, F. Lotte, and M. Tangermann, (2012), “Brain-Computer Interfaces: Beyond Medical Applications” **Computer (Long. Beach. Calif)**, **45, 4**, pp. **26–34**.
- [7] T. M. Vaughan, D. J. McFarland, G. Schalk, W. A. Sarnacki, L. Robinson, and J. R. Wolpaw, (2001), “EEG-based brain-computer interface: development of a speller application”, **Society for Neuroscience Abstracts**, **1, 26**.
- [8] I. Iturrate, J. M. Antelis, A. Kübler, and J. Minguez, (2009), “A noninvasive brain-actuated wheelchair based on a P300 neurophysiological protocol and automated navigation”, **Robot. IEEE Trans.**, **25, 3**, pp. **614–627**.
- [9] W. L. Nowinski, (2011), Introduction to brain anatomy, pp. 5–40, In: "**Biomechanics of the Brain**", K. Miller, Springer New York, New York.
- [10] A. Vallabhaneni, T. Wang, and B. He, (2005), Brain—computer interface, pp. 85–121, In: "**Neural Engineering**", Springer US, USA.
- [11] H. H. Jasper, (1958), “The ten twenty electrode system of the international federation”, **Electroencephalogr. Clin. Neurophysiol.**, **10, 1**, pp. **371–375**.
- [12] I. Manuel and B. Núñez, (2010), “EEG Artifact Detection”, pp. 1-33.
- [13] P. Rajasekaran, R. Kottaimalai, M. P. Rajasekaran, V. Selvam, and B. Kannapiran, (2013), “EEG Signal Classification using Principal Component Analysis with Neural Network in Brain Computer Interface Applications”, *Emerg. Trends Comput. Commun. Nanotechnol. (ICE-CCN)*, pp. 227–231, Tirunelveli, India.
- [14] E. Hortal, I. Eduardo, D. Planelles, E. Ianez, A. Ubeda, A. Costa, and J. M. Azorin, (2014), “Selection of the Best Mental Tasks for a SVM-Based BCI System”, *Syst. Man Cybern. (SMC)*, pp. 1483–1488, San Diego, CA.
- [15] E. Hortal, D. Planelles, A. Costa, E. Iáñez, A. Úbeda, J. M. Azorín, and E. Fernández, (2015), “SVM-based Brain–Machine Interface for controlling a robot arm through four mental tasks”, **Neurocomputing**, **151, 1**, pp. **116–121**.

- [16] V. Vijean, M. Hariharan, A. Saidatul, and S. Yaacob, (2011), “Mental tasks classifications using S-transform for BCI applications”, *Sustainable Utilization and Development in Engineering and Technology*, pp. 69–73, Semenyih, Malaysia.
- [17] F. Lotte, M. Congedo, L. Anatole, F. Lamarche, B. A. A, and A. Lécuyer, (2007), “A review of classification algorithms for EEG-based brain–computer interfaces”, **J. Neural Eng.**, **4**, pp 24.
- [18] C. W. Anderson and Z. Sijercic, (1996), “Classification of EEG signals from four subjects during five mental tasks”, *Solving engineering problems with neural networks: proceedings of the conference on engineering applications in neural networks (EANN’96)*, pp. 407–414.
- [19] X. Li, X. Chen, Y. Yan, W. Wei, and Z. Wang, (2014), “Classification of EEG Signals Using a Multiple Kernel Learning Support Vector Machine”, **Sensors**, **14**, 7, pp. 12784–12802.
- [20] E. Hortal, D. Planelles, A. Ubeda, A. Costa, and J. M. Azorin, (2014), “Brain-Machine Interface system to differentiate between five mental tasks”, *Systems Conference (SysCon)*, pp. 172–175, Ottawa, Canada.
- [21] S. Nasehi and H. Pourghassem, (2013), “Mental task classification based on HMM and BPNN”, *Communication Systems and Network Technologies (CSNT)*, pp. 210–214, Gwalior, India.
- [22] R. Kottaimalai, M. P. Rajasekaran, V. Selvam, and B. Kannapiran, (2013), “EEG signal classification using principal component analysis with neural network in brain computer interface applications”, *Emerging Trends in Computing, Communication and Nanotechnology (ICE-CCN)*, pp. 227–231, Tirunelveli, India.
- [23] M. N. Bawane and K. M. Bhurchandi, (2014), “Classification of Mental Task Based on EEG Processing Using Self Organising Feature Map”, *Intelligent Human-Machine Systems and Cybernetics (IHMSC)*, pp. 240–244, Hangzhou, China.

- [24] M. T. F. Talukdar, S. K. Sakib, C. Shahnaz, S. A. Fattah, T. F. Talukdar, and T. Foysal, (2014), “Motor imagery EEG signal classification scheme based on entropy of intrinsic mode function”, *Electrical and Computer Engineering (ICECE)*, pp. 703–706, Dhaka, Bangladesh.
- [25] W. Zhou, Y. Yang, and Z. Yu, (2012), “Discriminative dictionary learning for EEG signal classification in Brain-computer interface”, *Control Automation Robotics & Vision (ICARCV)*, pp. 1582–1585, Guangzhou, China.
- [26] Y. Shin, S. Lee, J. Lee, and H.-N. Lee, (2012), “Sparse representation-based classification scheme for motor imagery-based brain–computer interface systems”, **J. Neural Eng.**, **9**, 5, p. 56002.
- [27] Z. Keirn and J. I. Aunon, (1990), “A new mode of communication between man and his surroundings”, **Biomed. Eng. IEEE Trans.**, **37**, 12, pp. 1209–1214.
- [28] Z. M. Lwin and M. M. Thaw, (2015), “Mental Tasks Classification from Electroencephalogram (EEG) Signal Using Gabor Based Matching Pursuit (MP),” **Int. J. Comput. Sci. Technol.**, **6**, 1, pp. 22–26.
- [29] B. Blankertz, (2015), “Data Set IVa for the BCI Competition III”. [Online]. Available: http://www.bbci.de/competition/iii/desc_IVa.html. [Accessed: 10-Sep-2015].
- [30] B. Blankertz, K. R. Muller, D. J. Krusienski, G. Schalk, J. R. Wolpaw, A. Schlogl, G. Pfurtscheller, J. D. R. Millan, M. Schroder, and N. Birbaumer, (2006), “The BCI Competition III: Validating Alternative Approaches to Actual BCI Problems”, **IEEE Trans. Neural Syst. Rehabil. Eng.**, **14**, 2, pp. 153–159.
- [31] M. Arvaneh, C. Guan, K. K. Ang, and C. Quek, (2011), “Optimizing the channel selection and classification accuracy in EEG-based BCI,” **Biomed. Eng. IEEE Trans.**, **58**, 6, pp. 1865–1873.
- [32] Y. Wang, S. Gao, X. Gao, R. Iru, K. Lq, P. Dvhg, U. Frpsxwhu, F. Vljqdo, W. Fodvvli, and L. Q. J. Wkh, (2006), “Common spatial pattern method for channel

- selection in motor imagery based brain-computer interface”, Engineering in Medicine and Biology Society (IEEE-EMBS), pp. 5392–5395. Shanghai, China.
- [33] Z. Y. Chin, K. K. Ang, C. Wang, C. Guan, and H. Zhang, (2009), “Multi-class filter bank common spatial pattern for four-class motor imagery BCI”, Engineering in Medicine and Biology Society (EMBC), pp. 571–574, Minneapolis, USA.
- [34] M. Grosse-Wentrup and M. Buss, (2008), “Multiclass common spatial patterns and information theoretic feature extraction”, **Biomed. Eng. IEEE Trans.**, **55**, **8**, pp. **1991–2000**.
- [35] A. Ziehe, P. Laskov, G. Nolte, and K.-R. Müller, (2004), “A fast algorithm for joint diagonalization with non-orthogonal transformations and its application to blind source separation”, **J. Mach. Learn. Res.**, **5**, pp. **777–800**.
- [36] K. Huang and S. Aviyente, (2007), “Sparse representation for signal classification”, **Adv. Neural Inf. Process. Syst.**, **19**, p. **609**.
- [37] Q. Zhang and B. Li, (2010), “Discriminative K-SVD for dictionary learning in face recognition”, Computer Vision and Pattern Recognition (CVPR), pp. 2691–2698, San Francisco, USA.
- [38] O. Bryt and M. Elad, (2008), “Compression of facial images using the K-SVD algorithm”, **J. Vis. Commun. Image Represent.**, **19**, **4**, pp. **270–282**.
- [39] D. Needell, J. A. Tropp, and S. J. S. J. Andriole, (2010), “Cosamp: iterative signal recovery from incomplete and inaccurate samples”, **Commun. ACM**, **53**, **12**, pp. **93–100**.
- [40] T. Blumensath and M. E. Davies, (2008), “Iterative thresholding for sparse approximations”, **J. Fourier Anal. Appl.**, **14**, **5–6**, pp. **629–654**.
- [41] M. Babaie-Zadeh, H. Mohimani, and C. Jutten, (2008), “A fast approach for overcomplete sparse decomposition based on smoothed L0 norm”, **IEEE Trans. Signal Process**, **57**, **1**, pp. **289–301**.

- [42] M. Aharon, M. Elad, and A. Bruckstein, (2006), “K-SVD: An Algorithm for Designing Overcomplete Dictionaries for Sparse Representation”, **IEEE Trans. Signal Process.**, **54**, **11**, pp. **4311–4322**.
- [43] Z. Jiang, Z. Lin, and L. S. Davis, (2011), “Learning A Discriminative Dictionary for Sparse Coding via Label Consistent”, Computer Vision and Pattern Recognition (CVPR), pp. 1697–1704, Providence, Island.
- [44] S. Gu, L. Zhang, W. Zuo, and X. Feng, (2014), Projective dictionary pair learning for pattern classification, pp. 793–801, In: "**Advances in Neural Information Processing Systems**", Z. Ghahramani and M. Welling and C. Cortes and N.D. Lawrence and K.Q. Weinberger, Curran Associates.
- [45] T. O. M. Goldstein, B. O. Donoghue, S. Setzer, and R. Baraniuk, (2014), “Fast alternating direction optimization methods”, **SIAM J. Imaging Sci.**, **7**, **3**, pp. **1588–1623**.
- [46] R. L. Haupt, S. E. Haupt, and A. J. Wiley, (2004), "**Practical genetic algorithms**", vol. 1, John Wiley & Sons, Canada, pp. 28–47.
- [47] L. Zhang, W. He, C. He, and P. Wang, (2010), “Improving mental task classification by adding high frequency band information”, **J. Med. Syst.**, **34**, **1**, pp. **51–60**.
- [48] H.-I. Suk and S.-W. Lee, (2013), “A novel Bayesian framework for discriminative feature extraction in brain-computer interfaces”, **IEEE Trans. Pattern Anal. Mach. Intell.**, **35**, **2**, pp. **286–299**.

Abstract

Brain-Computer Interface (BCI) is a useful communication tool between the human's brain and external devices. Accurate and effective classification of Electroencephalography (EEG) signals plays an important role in performance of BCI applications. In this research, two EEG signal classification approaches based on sparse representation and dictionary learning is proposed to classify mental tasks. We used Colorado and Berlin university databases that includes two and five classes respectively. Band pass filtering and common spatial pattern (CSP), which are effective in BCI, were applied to preprocess data.

In first proposed method a classifier based on sparse representation was presented. In this method feature extracted from CSP were used to build a dictionary and to apply to a multiclass problem, multiclass CSP was used. In second proposed method classification based on dictionary learning was proposed. In this method an analytical dictionary (to calculate sparse response) and a synthetic dictionary (to minimize reconstruction error of signal) are trained.

The main advantage of using analytical dictionary is to achieve the sparse coefficient with lower computational complexity.

Result of first method by using validation on Colorado database are 90%, 83.5%, 76.5% and 70% to classify two, three, four and five classes, respectively. The average results on Berlin database by using validation is 97.56%. Result of second proposed method on Berlin test database is 80.98%.

Keywords: Brain computer interface, EEG signal classification, sparse representation, dictionary learning.



Shahrood University of Technology
Faculty of Computer Engineering and information technology
Artificial intelligence group

Mental task recognition in BCI using EEG data

Rasool Ameri

Supervisor:
Dr. Ali Akbar Pouyan

Associate Supervisor:
Dr. Vahid Abolghasemi

February 2016