

مقاله پژوهشی

مروری بر روش‌های طبقه‌بندی
در سیگنال Speller-P3۰۰ رابط مغز و کامپیوتر

Doi: 10.30508/KDIP.2022.346627.1042

حمیده جشن^۱ | زینب قاسمی نژاد^۲ | الهه محرم خانی (نویسنده مسئول)^۳ | سید مجتبی صباغ جعفری^۴

۱- کارشناس ارشد، مهندسی کامپیوتر، دانشکده فنی مهندسی، دانشگاه آزاد اسلامی واحد علوم تحقیقات، اهواز، ایران.

۲- کارشناسی ارشد مهندسی کامپیوتر-هوش مصنوعی، دانشکده فنی مهندسی، دانشگاه ولیعصر، رفسنجان، ایران.

۳- کارشناسی ارشد مهندسی کامپیوتر-شبکه‌های کامپیوتری، دانشکده فنی مهندسی، مؤسسه آموزش عالی صائب، ابهر، ایران.

۴- استادیار گروه کامپیوتر، دانشگاه ولیعصر، رفسنجان، ایران.

تاریخ دریافت: ۱۴۰۱/۰۳/۲۵

تاریخ پذیرش: ۱۴۰۱/۰۵/۲۹

صفحه: ۵۲-۶۵

سیستم رابط رایانه‌ای مغز (BCI) بخشی از فناوری عصبی است که فرمان را از مغز انسان به رایانه منتقل می‌کند. BCI در حال حاضر بیشترین رشد را در زمینه تحقیق دارد. برنامه‌های BCI دارای زمینه‌های مختلفی از جمله: پزشکی، آموزش، خودتنظیمی، بازی‌ها و سرگرمی‌ها، تولید، امنیت و همچنین بازاریابی هستند. دستگاه‌های الکترونیکی را می‌توان با استفاده از سیگنال مغزی به نام الکتروانسفالوگرافی (EEG) برای ثبت فعالیت الکتریکی مغز کنترل کرد. موج P3۰۰ اوج مثبت یک پتانسیل مرتبط با رویداد (ERP) است که ۳۰۰ میلی ثانیه رخ می‌دهد توسط EEG ضبط شده است. یک روش عمده در زمینه تحقیق BCI الگوی خاص مبتنی بر P3۰۰ است، محرک‌های که احتمال وقوع کمتری نسبت به سایر محرک‌ها دارند و در نتیجه فرد مورد آزمایش به وقوع غیرمنتظره این محرک‌ها واکنش نشان می‌دهد را با توجه به یک سری محرک‌های استاندارد سریع ارائه شده شناسایی می‌کنند. در این مقاله ما الگوریتم‌های یادگیری ماشین از جمله الگوریتم‌های SVM^۵، SWLDA^۴، LDA^۳، DCPM^۳ و ... را در زمینه استخراج ویژگی و طبقه‌بندی مورد استفاده برای طراحی سیستم‌های رابط مغز و کامپیوتر (BCI) مبتنی بر spell-P3۰۰ را مرور می‌کنیم. و بعد از اینکه روش‌های طبقه‌بندی، مرور و مقایسه شده است. در پایان به جمع‌بندی نهایی مقاله‌مان پرداخته شده است، که نتایج تحقیق مروری مورد بررسی نشان می‌دهد، روش طبقه‌بندی تطبیقی برای نظارت شده و بدون نظارت عملکرد بهتری از طبقه‌بندی استاتیک دارد.

کلمات کلیدی: spell-P3۰۰، رابط مغز و کامپیوتر (BCI)، طبقه‌بندی، استخراج ویژگی، EEG

1- Electroencephalography

2- Event-Related Potential

3- Discriminative Canonical Pattern Matching

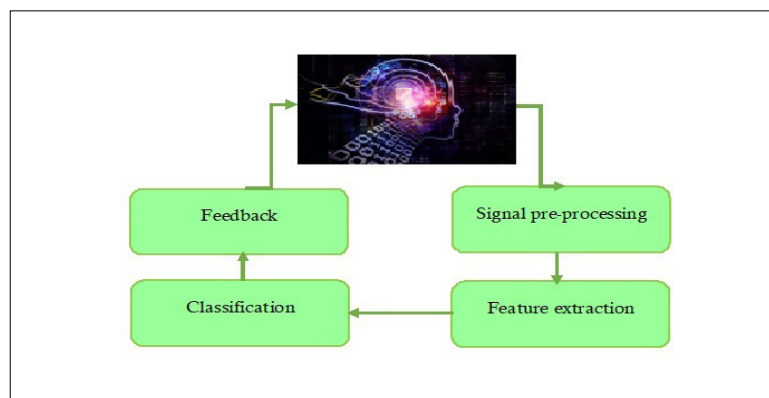
4- Linear Discriminant Analysis

5- Stepwise Linear Discriminant Analysis

۱- مقدمه

با توجه به یک سری محرک‌های استاندارد سریع ارائه شده شناسایی می‌کنند. spell-P300 یک سیستم BCI است که به طور گسترده‌ای مورد استفاده قرار می‌گیرد و به کاربران امکان می‌دهد تا با تمرکز بیشتر ارتباط برقرار کنند. BCI به ویژه spell P300 امواج P300 را طبقه‌بندی می‌کند و ویژگی‌های هدفمند را تشخیص می‌دهد (پاتلیا و پاتل^۲، ۲۰۱۹). یک سیستم BCI شامل مواردی از جمله دریافت سیگنال برای کاهش نویز، پیش پردازش سیگنال برای بهبود سیگنال و کاهش اندازه داده‌ها که شامل استخراج و طبقه‌بندی ویژگی‌ها می‌باشد (سمپانا و میتایم^۳، ۲۰۱۸). در این مقاله به نحوی دریافت، پیش پردازش سیگنال و نحوی انتخاب ویژگی‌ها می‌پردازیم و همچنین الگوریتم‌های یادگیری ماشین از جمله الگوریتم‌های STDA، SVM، DCPM، LDA، SWLDA، BLDA، SKLDA را در زمینه استخراج ویژگی و طبقه‌بندی مورد استفاده برای طراحی سیستم‌های رابط مغز و کامپیوتر (BCI) مبتنی بر spell-P300 را مرور می‌کنیم. ما الگوریتم‌های بکار رفته را مختصر ارائه می‌دهیم و خصوصیات حیاتی آنها را شرح می‌دهیم. بر اساس مطالعات گذشته، ما آنها را از نظر عملکرد مقایسه می‌کنیم و رهنمودهایی را برای انتخاب الگوریتم طبقه‌بندی مناسب برای یک BCI خاص ارائه می‌دهیم. در شکل شماره (۱) به چارچوبی از spell-P300 پرداخته شده است.

سیگنال EEG (الکتروانسفالوگرام) فعالیت الکتریکی مغز را نشان می‌دهد. ماهیت آن‌ها بسیار تصادفی است و ممکن است حاوی اطلاعات مفیدی در مورد وضعیت مغز باشد. با این حال، به دست آوردن اطلاعات مفید از این سیگنال‌ها فقط با مشاهده آن‌ها بسیار دشوار است. آن‌ها اساساً ماهیتی غیرخطی و غیرثابت دارند. از این رو، می‌توان با استفاده از تکنیک‌های پیشرفته پردازش سیگنال، ویژگی‌های مهمی را برای تشخیص بیماری‌های مختلف استخراج کرد (سوبها، جوزف، آچاریا و لیم^۱، ۲۰۱۰). سیستم رابط رایانه‌ای مغز (BCI) بخشی از فناوری عصبی است که فرمان را از مغز انسان به رایانه منتقل می‌کند. BCI در حال حاضر بیشترین رشد را در زمینه تحقیق دارد. برنامه‌های BCI دارای زمینه‌های مختلفی از جمله پزشکی، آموزش، خودتنظیمی، بازی‌ها و سرگرمی‌ها، تولید، امنیت و همچنین بازاریابی هستند. دستگاه‌های الکترونیکی را می‌توان با استفاده از سیگنال مغزی به نام الکتروانسفالوگرافی (EEG) برای ثبت فعالیت الکتریکی مغز کنترل کرد. موج P300 اوج مثبت یک پتانسیل مرتبط با رویداد (ERP) است که ۳۰ میلی ثانیه رخ می‌دهد توسط EEG ضبط شده است. یک روش عمده در زمینه تحقیق BCI الگوی خاص مبتنی بر spell-P300 است، محرک‌های که احتمال وقوع کمتری نسبت به سایر محرک‌ها دارند و در نتیجه فرد مورد آزمایش به وقوع غیر منتظره این محرک‌ها واکنش نشان می‌دهد را



شکل (۱): چارچوبی از spell-P300

- 1- Subha, Joseph, Acharya & Lim
- 2- Patelia, & Patel
- 3- Sampanna, & Mitaim,

پیش بینی می شود که یک پنجره ۶۶۷ میلی ثانیه، یعنی ۱۶۰ نمونه به اندازه کافی بزرگ هستند تا تمام اطلاعات لازم برای طبقه بندی را بدست آورند. این پنجره ها پنجره هایی هستند که با هم تداخل دارند. (۲) هر سیگنال استخراج شده توسط فیلتر چوبیشف باند گذر مرتبه هشتم از نوع فیلتر شده و فرکانس قطع در ۱۰ و ۲۰ هرتز قرار دارد. تمام نمونه ها در مرحله پیش پردازش به عنوان یک مجموعه ویژگی استفاده می شوند. برای انتخاب بهترین ویژگی ها از مجموعه ویژگی های استخراج شده، PCA در کار اعمال می شود. PCA ویژگی های کم اهمیت را حذف می کند و تعداد کمی از ویژگی های مهم برای طبقه بندی نگهداری می شوند (کاندو و آری، ۲۰۱۸).

انتخاب ویژگی ها

یک مرحله انتخاب ویژگی می تواند بعد از مرحله استخراج ویژگی برای انتخاب زیر مجموعه ای از ویژگی ها با مزایای بالقوه مختلف استفاده شود. اولاً، در میان ویژگی های مختلفی که ممکن است از سیگنال های EEG استخراج شود، برخی از آنها ممکن است زائد باشند یا مربوط به حالات ذهنی مورد هدف BCI نباشند. ثانیاً، تعداد پارامترهایی که طبقه بندی کننده باید بهینه کند با تعداد ویژگی ها ارتباط مثبت دارد. در نتیجه کاهش تعداد ویژگی ها باعث می شود که پارامترهای کمتری توسط طبقه بندی کننده، بهینه شوند. همچنین اثرات احتمالی اضافه کار را کاهش می دهد و بنابراین می تواند عملکرد را به ویژه بهبود بخشد. اگر تعداد نمونه های آموزش کم باشد. ثالثاً، از دیدگاه استخراج دانش، اگر فقط چند ویژگی انتخاب و یا رتبه بندی شوند، مشاهده اینکه کدام ویژگی ها در واقع مربوط به حالات ذهنی هدف قرار می گیرند، آسان تر است. چهارم، یک مدل با ویژگی های سریع تری را برای یک نمونه جدید ایجاد کند، زیرا باید از نظر محاسباتی کارآمدتر باشد. پنجم، جمع آوری و ذخیره سازی داده ها کاهش می یابد. سه رویکرد انتخاب ویژگی شناسایی شده است: روش های فیلتر، روپوش و روش های جاسازی شده. روش های

ادامه مقاله به این صورت می باشد که در بخش های بعدی به ترتیب پیش پردازش داده ها، انتخاب ویژگی ها، طبقه بندی ویژگی های را دارد که هر یک دسته بندی شده و در نهایت از مباحث ذکر شده، نتیجه گیری شده است.

۲- مبانی نظری

پیش پردازش داده ها: عوامل بسیاری در موفقیت یادگیری ماشین (ML) در انجام کار طبقه بندی تأثیر می گذارند. نمایش و کیفیت داده های نمونه قبل از هر چیز مهم است. اگر اطلاعات بی ربط و زاید زیادی وجود داشته باشد یا داده های نویز و غیرقابل اعتماد باشد، کشف دانش در مرحله آموزش دشوارتر است. پیش پردازش داده ها شامل: تمیز کردن داده ها، عادی سازی، تغییر شکل، استخراج و انتخاب ویژگی ها و غیره است. محصول پیش پردازش داده ها مجموعه آموزش نهایی است. خوب است که یک توالی الگوریتم پیش پردازش داده بهترین عملکرد را برای هر مجموعه داده داشته باشد. بنابراین، ما بهترین الگوریتم های شناخته شده را برای هر یک ارائه می دهیم (کوتسیانتیس و کنلوپولس پینتالس^۱، ۲۰۰۶). مرحله پیش پردازش داده ها به طوری که فرد بهترین عملکرد را برای مجموعه داده های خود بدست آورد. پیش پردازش داده ها به منظور جمع بندی توصیفی داده ها، تمیز کردن داده ها، ادغام و تبدیل داده ها، کاهش داده ها، گسسته سازی و خلاصه کردن انجام می شود. برای تمیز کردن داده ها مقادیر از دست رفته را تکمیل می کنیم، داده های نویز را برطرف می کنیم، موارد پرت را شناسایی یا حذف و ناسازگاری ها را برطرف می کنیم. کاهش داده ها حجم داده ها را کاهش می دهد اما نتایج تحلیلی مشابه یا مشابهی را ایجاد می کند. گسسته سازی داده ها بخشی از کاهش داده ها است اما دارای اهمیت ویژه ای است، خصوصاً برای داده های عددی. مرحله پیش پردازش شامل موارد زیر است: (۱) از هر کانال، داده ای با مدت زمان ۰-۶۶۷ میلی ثانیه پس از هر بار فلش زدن استخراج می شود. طبق اطلاعات قبلی در مورد ۲۳۰۰، بعد از ۳۰۰ میلی ثانیه محرک یک اوج مثبت ظاهر می شود. بنابراین،

1- Machine Learning
2- Kotsiantis, Kanellopoulos, & Pintelas
3- Kundu, & Ari

روش‌های جاسازی شده انتخاب ویژگی‌ها و ارزیابی را در یک فرآیند منحصر به فرد ادغام می‌کنند، به عنوان مثال در یک درخت تصمیم‌گیری یا یک پرسپترون چند لایه با آسیب سلول بهینه (لاتی، کانگدو، لکویر، لامرچ و امالیدی، ۲۰۰۷). برای انتخاب یک طبقه بندی مناسب برای یک سیستم BCI داده شده، درک ویژگی‌های مهم، خصوصیات آنها و رابطه آنها ضروری است. ورودی BCI عموماً با حجم بالایی از ویژگی‌ها از جمله مقادیر دامنه سیگنال‌های EEG، توان باند، مقادیر چگالی طیفی قدرت، پارامترهای خود رگرسیون و خود رگرسیون تطبیقی، ویژگی‌های فرکانس زمان. روش‌های انتخاب ویژگی‌ها به محققان کمک می‌کند تا ویژگی‌های قابل توجهی را از بین تعداد زیاد ویژگی‌ها انتخاب و مطالعه کنند تا کل فرایند را با سهولت انجام دهند. روند انتخاب ویژگی با استخراج کاملاً متفاوت است. فرآیند انتخاب ویژگی‌ها به محققان کمک می‌کند زیرمجموعه مهم ویژگی را که از طریق روش‌های مختلف استخراج ویژگی استخراج شده است، شناسایی کنند. یک انتخاب و فرآیند استخراج ویژگی خوب به انتخاب یک طبقه بندی مناسب نیز کمک می‌کند. در بیشتر موارد، انتخاب ویژگی نقشی حیاتی در طبقه بندی سیگنال مبتنی بر EEG ایفا می‌کند. انتخاب زیرمجموعه‌ای از ویژگی‌ها که برای مسئله طبقه بندی بسیار مفید هستند، اغلب دقت طبقه بندی را بر روی داده‌های جدید افزایش می‌دهد (چاکلادار و چاکرابرتی، ۲۰۱۹).

استخراج ویژگی با استفاده از الگوی فضایی مشترک (CSP)
الگوی فضایی مشترک (CSP) یک روش ریاضی است که در پردازش سیگنال برای جداسازی سیگنال چند متغیره به زیرمولفه‌های افزودنی استفاده می‌شود که حداکثر اختلاف واریانس بین دو پنجره را دارند. این الگوهای فضایی، به معنای حداقل مربعات، حداکثر برای واریانس EEG از یک جمعیت و حداقل برای واریانس در جمعیت دیگر حساب می‌شوند، بنابراین به نظر می‌رسد برای تفکیک کمی بین EEG‌های منفرد در دو جمعیت، بهینه

جایگزین بسیاری برای هر روش ارائه شده است. روش‌های فیلتر، مستقل از طبقه بندی مورد استفاده، به معیارهای ارتباط بین هر ویژگی و کلاس هدف متکی هستند. ضریب تعیین، که مربع برآورد ضریب همبستگی پیرسون است، می‌تواند به عنوان یک معیار رتبه بندی ویژگی استفاده شود. ضریب تعیین همچنین می‌تواند برای یک مسئله دو طبقه استفاده شود، به عنوان مثال کلاس‌های که با برچسب -۱ یا +۱ مشخص می‌شوند. هر چند ضریب همبستگی فقط می‌تواند وابستگی‌های خطی بین ویژگی‌ها و کلاس‌ها را تشخیص دهد. برای بهره‌برداری از روابط غیرخطی، یک راه حل ساده این است که پیش پردازش غیرخطی را انجام دهید، مانند پیش پردازش درجه دو یا ثبت ویژگی‌ها. به عنوان مثال معیارهای رتبه بندی بر اساس تئوری اطلاعات همچنین می‌تواند از اطلاعات متقابل بین هر ویژگی و متغیر هدف استفاده کند. بسیاری از رویکردهای انتخاب ویژگی فیلتر نیاز به تخمین چگالی احتمال و چگالی توام ویژگی‌ها و برچسب کلاس داده‌ها دارند. یک راه حل این است که ویژگی‌ها و برچسب‌های کلاس را تفکیک کنید. راه حل دیگر تقریب چگالی آنها با یک روش غیرپارامتریک مانند؛ پنجره‌های Parzen است. اگر تراکم‌ها با یک توزیع نرمال تخمین زده شوند، نتیجه بدست آمده توسط اطلاعات متقابل مشابه نتیجه حاصل از ضریب همبستگی خواهد بود. رویکردهای فیلتر با توجه به تعداد ویژگی‌ها دارای پیچیدگی خطی هستند. با این حال، این ممکن است منجر به انتخاب ویژگی‌های زاید شود. بسته بندی و روش‌های جاسازی شده این مشکل را با هزینه زمان طولانی تر محاسبه می‌کند. این رویکردها از یک طبقه بندی کننده برای بدست آوردن زیرمجموعه‌ای از ویژگی‌ها استفاده می‌کنند. روش‌های بسته بندی زیرمجموعه‌ای از ویژگی‌ها را انتخاب می‌کنند، آن را به عنوان ورودی به یک طبقه بندی کننده برای آموزش ارائه می‌دهند، عملکرد حاصل را مشاهده می‌کنند و جستجو را بر اساس معیار توقف متوقف می‌کنند یا در صورت عدم رضایت ملاک زیرمجموعه جدیدی را پیشنهاد می‌کنند.

1- Lotte, Congedo, Lécuyer, Lamarche, & Arnaldi

2- Chakladar & Chakraborty

3- Common Spatial Patterns

متغیرهای دستکاری شده تجربی است، که در روش‌های رگرسیون یا طبقه‌بندی رسمی شده است. ICA روش محبوب دیگری برای استخراج ویژگی‌ها است. ICA فرض می‌کند که سیگنال EEG مشاهده شده ترکیبی از چندین سیگنال منبع مستقل است (چاکلادرو چاکرابرتی، ۲۰۱۹). بنابراین ICA یک روش اکتشافی یا داده محور است: ما به راحتی می‌توانیم برخی سیستم‌ها یا پدیده‌ها را بدون طراحی شرایط آزمایشی مختلف اندازه‌گیری کنیم. وقتی فرضیه‌های مناسب در دسترس نیستند، یا خیلی محدود یا ساده انگاشته می‌شوند، می‌توان از ICA برای بررسی ساختار داده‌ها استفاده کرد (هیوانن^۴، ۲۰۱۳).

استخراج ویژگی با استفاده از مولفه اتورگرسیون (AR)
روش AR برای استخراج ویژگی با توجه به دامنه زمان مفید است. مدل AR برای دستیابی به فرکانس بهتر و بدون مشکل نشت طیفی از روش پارامتریک استفاده می‌کند. از این رو عملکرد بهتری نسبت به رویکرد غیرپارامتری دارد. برآورد طیفی AR ترجیح می‌دهد از تبدیل فوریه استفاده کند، اما در سیگنال ناپایدار عملکرد ضعیفی دارد. برای غلبه بر این مشکل، AR تطبیقی چند متغیره^۵ (MVAAR) را اجرا کرده است (چاکلادرو چاکرابرتی، ۲۰۱۹).

استخراج ویژگی با استفاده از تبدیل موجک مداوم (CWT)
CWT یک تکنیک قدرتمند برای استخراج ویژگی‌های ارزشمند از سیگنال است، و همچنین به عنوان ابزاری برای شناسایی الگو استفاده می‌شود. CWT با استفاده از برخی خصوصیات الگوی موجک عملکرد بهتری نسبت به روش تطبیق الگوی کلاسیک دارد. موجک‌ها برای تجزیه و تحلیل سیگنال گذرا، جایی که خواص طیفی سیگنال در زمان تغییر می‌کنند، مناسب هستند (چاکلادرو چاکرابرتی، ۲۰۱۹).

استخراج ویژگی با استفاده از تبدیل موجک گسسته

باشد (کولز، لازار و ژوو^۱، ۱۹۹۰)، CSP یک روش استخراج ویژگی است که سیگنال‌های EEG چند کاناله را به فضایی فرعی ارائه می‌دهد، جایی که تفاوت بین کلاس‌ها برجسته می‌شود و شباهت‌ها به حداقل می‌رسد. CSP را می‌توان به زیر فضایی ثابت (SCSP^۲) قاعده‌مند کرد، و این دقت طبقه‌بندی را افزایش می‌دهد. در روش‌های فعلی، CSP براساس برخی از رویکردهای ابتکاری در طبقه‌بندی در سطح چند کلاس اجرا شده است (چاکلادرو چاکرابرتی، ۲۰۱۹).

استخراج ویژگی با استفاده از تجزیه و تحلیل کامپوننت مستقل (ICA^۳)

تجزیه و تحلیل مولفه‌های مستقل به عنوان یک روش اساسی برای تجزیه و تحلیل داده‌های چند متغیری ایجاد شده است. این روش تجزیه خطی (تبدیل)، مانند روش‌های کلاسیک برای تجزیه و تحلیل عامل و تجزیه و تحلیل مولفه‌های اصلی (PCA) داده‌ها را می‌آموزد. با این حال، ICA در بسیاری از مواردی که روش‌های کلاسیک با شکست روبرو می‌شود، می‌تواند اجزای اساسی و منابع مخلوط شده در داده‌های مشاهده شده را پیدا کند. ICA سعی دارد با استفاده از برخی فرض‌های ساده از خصوصیات آماری آنها، اجزای اصلی یا منابع را پیدا کند. برخلاف روش‌های دیگر، فرضیه‌های اساسی مستقل از یکدیگر فرض می‌شوند، که اگر با فرایندهای فیزیکی مشخص مطابقت داشته باشند واقع‌بینانه است. با این حال، آنچه ICA را از PCA و تحلیل عاملی متمایز می‌کند، استفاده از ساختار غیر گاوسی داده‌ها است که برای بازیابی مولفه‌های اساسی ایجاد کننده داده‌ها بسیار مهم است. ICA یک روش بدون نظارت است به این معنا که داده‌های ورودی را به شکل یک ماتریس داده واحد می‌گیرد. لازم نیست "خروجی" مورد نظر سیستم را بدانید، یا اندازه‌گیری‌ها را به شرایط مختلف تقسیم کنید. این در تضاد کامل با روش‌های علمی کلاسیک مبتنی بر برخی

- 1- Koles, Lazar, & Zhou
- 2- Stationary Common Spatial Patterns
- 3- Independent Component Analysis
- 4- Hyvärinen
- 5- Autoregressive Component
- 6- Multivariate Adaptive
- 7- Continuous Wavelet Transform

(DWT)

در تجزیه و تحلیل عددی و تحلیل عملکردی، تبدیل موجک گسسته (DWT) هر تبدیل موجک است که از آن موجک به طور مجزا نمونه برداری می شود. مانند سایر تبدیل های موجک، مزیت کلیدی آن نسبت به تبدیل های فوریه وضوح زمانی است: هم اطلاعات فرکانس و هم مکان (موقعیت به موقع) را به دست می آورد. تبدیل موجک گسسته دارای کاربردهای زیادی در علوم، مهندسی، ریاضیات و علوم کامپیوتر است. از همه مهم تر، این برای کدگذاری سیگنال استفاده می شود، تا سیگنال گسسته را به شکل زائدتر نشان دهد، که اغلب به عنوان پیش شرط فشرده سازی داده است (نصیر و ساسانی، ۲۰۱۹). اشکال اصلی CWT این است که تولید افزونه و پیچیدگی می کند زیرا شامل تجزیه و تحلیل سیگنال در تعداد بسیار زیادی از فرکانس ها پس از تغییر موجک مادر است. DWT موجک مادر را در مقادیر گسسته منشعب و تغییر می دهد. وضوح چندگانه داده خام $x(n)$ از EEG به دو فیلتر دیجیتال $g(n)$ و $h(n)$ تقسیم شده و نمونه ها را به دو قسمت تقسیم می کند. موجک گسسته $g(n)$ یک فیلتر عبور بالا گذر است، در حالی که تصویر آینه $h(n)$ یک فیلتر پایین گذر است (هیوانن، ۲۰۱۳).

استخراج ویژگی با استفاده از تجزیه و تحلیل مولفه اصلی (PCA^۳)

تجزیه و تحلیل مولفه های اصلی (PCA) برای رتبه بندی الکترودها بر اساس اهمیت آنها استفاده می شود. PCA به عنوان روش انتخاب الکترودها یا کانال استفاده می شود که به عنوان کاهش ابعاد وابسته به موضوع عمل می کند. علاوه بر این، از آنجا که سیگنال های به دست آمده دارای نویز هستند، این نویزها با استفاده از PCA برداشته شده و ویژگی های استخراج شده کاهش بعد می یابد. بنابراین، می توان تعداد الکترودهای لازم برای طبقه بندی سیگنال های مغزی را کاهش داد بدون اینکه عملکرد قابل توجه طبقه بندی را از دست بدهد. ویژگی های جدید

کاهش یافته در طبقه بندی تحلیل خطی بیزین^۴، ارائه می شوند (سلیم، واحد و کاداه، ۲۰۰۹). در بیشتر برنامه های پردازش سیگنال، PCA به عنوان یک روش کاهش بعد داده استفاده می شود. مهمترین ویژگی را از مجموعه داده های بعد بالا استخراج می کند. PCA یک پیش بینی خطی بدون نظارت است که پراکندگی همه نمونه ها را به حداکثر می رساند (کاندو و آری، ۲۰۱۸).

PCA روشی برای استخراج ویژگی ها پس از کاهش بعد داده است. PCA داده های ورودی را در یک فضای ویژه k بعدی به بردارهای ویژه k تایی تبدیل می کند. تمام نقاط داده (X) در راستای اولین بردار ویژه (v) پیش بینی می شوند تا واریانس حاصل (Z) حداکثر باشد. مقادیر ویژه قابل توجهی با استفاده از آزمون ANOVA انتخاب می شوند، سپس با استفاده از چندین طبقه بندی نظارت شده آموزش و آزمایش می شوند. PCA همچنین مصنوعات را برای کاهش بعد سیگنال حذف می کند (هیوانن، ۲۰۱۳).

طبقه بندی ویژگی ها

پس از استخراج ویژگی های مورد نیاز از سیگنال EEG، آن ویژگی های مفید باید با استفاده از یک طبقه بندی کننده مناسب طبقه بندی مناسب شوند. روش طبقه بندی ویژگی ها دارای دو بخش است، (۱) رگرسیون و (۲) طبقه بندی، اما روش طبقه بندی مناسب ترین روش مورد استفاده در روزهای اخیر است. الگوریتم های رگرسیون برای ورودی های کمی مفید هستند در حالی که الگوریتم های طبقه بندی عمدتاً ورودی های طبقه ای را اداره می کنند. الگوریتم های رگرسیون بیشتر وقتی استفاده می شوند که بخواهیم قصد کاربر را بر اساس ویژگی های استخراج شده از سیگنال EEG پیش بینی کنیم، در حالی که الگوریتم های طبقه بندی برای طبقه بندی ویژگی های مستقل به کلاس مناسب استفاده می شوند. الگوریتم های طبقه بندی بیشتر برای دو مسئله کلاس استفاده می شود، زیرا گاهی اوقات "طبقه بندی باینری" نامیده می شود. بنابراین، رویکرد طبقه بندی ممکن است برای کاربردهای دو هدف

- 1- Discrete Wavelet Transform
- 2- Nasir, Cool, & Sassani,
- 3- Principal Component Analysis
- 4- Bayesian Linear Discriminant Analysis
- 5- Selim, Wahed, & Kadah

SVM ها در تحقیقات BCI استفاده شده‌اند زیرا این یک روش قدرتمند برای تشخیص الگو، به ویژه برای مشکلات با ابعاد بالا است (چاکلادرو چاکرابرتی، ۲۰۱۹).

طبقه‌بندی برپایه تطبیق الگوی متمایز (DCPM)

یک الگوریتم جدید، به عنوان مثال تطبیق الگوی متمایز ((DCPM، برای شناسایی ERP های کوچک ایجاد شده است. از ویژگی تقارن EEG در فضا برای سرکوب نویز حالت مشترک پس زمینه و سپس شناخت الگوهای متعارف ERP استفاده می‌کند. DCPM در طبقه‌بندی پتانسیل‌های برانگیخته نامتقارن بصری کوچک (aVEP) با دامنه حتی به اندازه ۵٪ میکروولت عملکرد خوبی دارد. روش تطبیق الگوی متمایز روش DCPM شامل سه قسمت عمده است: (۱) ساخت الگوهای فضایی افتراقی (DSP^3) ، (۲) ساخت الگوهای CCA^۴، (۳) تطبیق الگو. روش الگوی فضایی افتراقی ماتریس طرح‌ریزی را می‌سازد، که می‌تواند به عنوان یک فیلتر فضایی در نظر گرفته شود و تفکیک‌پذیری دو نمونه پس از فیلتر کردن با ماتریس W محاسبه می‌شود، همچنین همبستگی توسط الگوریتم CCA محاسبه می‌شود (ایکسائو، ایکسو، وانگ، جانگ و مینگ، ۲۰۱۹).

رگرسیون خطی (LREG)

الگوریتم‌های رگرسیون خطی (LREG) بیشتر بسته به تعداد متغیرهای مستقل و نوع رابطه بین متغیرهای مستقل و وابسته متفاوت هستند. الگوریتم‌های رگرسیون خطی (LREG) بیشتر به تعداد متغیرهای مستقل و نوع رابطه بین متغیرهای مستقل و وابسته بستگی دارند. رگرسیون خطی یک مدل خطی است، به عنوان مثال مدلی که یک رابطه خطی را بین متغیرهای ورودی (x) و متغیر خروجی منفرد (y) در نظر می‌گیرد. به طور دقیق‌تر، y را می‌توان از ترکیب خطی متغیرهای ورودی (x) محاسبه کرد (چاکلادرو چاکرابرتی، ۲۰۱۹).

مفیدتر باشد. الگوریتم‌های رگرسیون برای چندین هدف بهتر هستند. در اینجا چندین طبقه‌بندی ویژگی از جمله خطی، غیرخطی و گرافیکی مورد بحث قرار گرفته است. ما همچنین مختصراً درباره خصوصیات مهم و کاربرد چندین رویکرد مدل که به آن دسته تعلق دارند، بحث می‌کنیم (چاکلادرو چاکرابرتی، ۲۰۱۹).

طبقه‌بندی برپایه ماشین‌های بردار پشتیبان (SVM)

الگوریتم‌های مبتنی بر ماشین بردار پشتیبان (SVM) نمونه‌هایی از طبقه‌بندی کننده‌ها هستند که می‌توانند برای مطالعه سیگنال‌های EEG استفاده شوند. مزیت استفاده از الگوریتم‌های طبقه‌بندی، بهبود کاهش خطا در تشخیص یک دسته‌بند متمایز کننده است که می‌تواند صفحه بهینه‌ای که دسته‌های گوناگون داده‌ها را از هم جدا می‌سازد، پیدا کند. یک ماشین بردار پشتیبانی از یک ابرصفحه متمایز برای تشخیص کلاس‌ها استفاده می‌کند. در SVM، ابرصفحه‌ای (خط) را انتخاب می‌کنیم که حاشیه اطمینان بیشتری داشته باشد، یعنی نزدیک‌ترین فاصله را از نقاط آموزش داشته باشد. با به حداکثر رساندن حاشیه قابلیت‌های تعمیم‌پذیری افزایش می‌یابد. این طبقه‌بندی از یک پارامتر منظم‌سازی C استفاده می‌کند که مقادیر کوچک C نقاط نزدیک به حاشیه را نادیده می‌گیرد و مرز حاشیه را افزایش می‌دهد، در حالی که مقادیر بزرگ C تمام نقاط را در نظر می‌گیرد، و مرز را کاهش می‌دهد. با استفاده از پارامتر C می‌توان تصمیم گرفت که چقدر خطاهای مجموعه آموزش وجود داشته باشد. ماشین‌های بردار پشتیبان طبقه‌بندی را با استفاده از مرزهای تصمیم خطی امکان‌پذیر می‌کنند و به عنوان SVM های خطی شناخته می‌شوند. SVM ها چندین مزیت دارند. در واقع، به حداکثر رساندن حاشیه اطمینان و اصطلاحاً تنظیم صفحه (خط) بهینه برای دسته‌بندی داده‌ها، SVM دارای خواص تعمیم‌پذیری خوبی است تا نسبت به آموزش بیش از حد (بیش برآزش) و لعنت ابعاد حساسیت نداشته باشد.

- 1- curse of dimensionality
- 2- miniature aVEP-speller
- 3- Discriminative Spatial Patterns
- 4- Canonical Correlation Analysis
- 5- Xiao, Xu, Wang, Jung, & Ming
- 6- Linear Regression

۱۵/۰ باشد از مدل حذف می‌شود. این فرایند تشخیص تا زمانی تکرار می‌شود که مدل تفکیک شامل ویژگی‌های کافی باشد یا هیچ ویژگی دیگری شرایط ورود و حذف را نداشته باشد (ایکسائو و همکاران، ۲۰۱۹).

طبقه‌بندی با تابع تفکیک خطی نیمه نظارتی (sLDA^۳) و (SKLDA^۳)

به عنوان یک روش معمول برای جبران بایاس سیستماتیک ماتریس‌های کوواریانس برآورد شده و پارامتر انقباض برای فضاهای ویژگی با ابعاد بالا، SKLDA عملکرد برتر خود را هنگام استفاده از نمونه‌های آموزش ناکافی نشان داده است. از طریق تنظیم مقادیر ویژه خیلی بزرگ ماتریس کوواریانس به میانگین مقادیر ویژه، LDA سنتی را بهبود می‌بخشد (ایکسائو و همکاران، ۲۰۱۹). انقباض LDA (sLDA)، یک طبقه‌بندی‌کننده استاندارد LDA است که در آن ماتریس‌های کوواریانس مرتبط با کلاس مورد استفاده در بهینه‌سازی آن با استفاده از جمع‌شدگی منظم می‌شوند (بلنکرتز، لیم، تریدر، هافه و مولر^۴، ۲۰۱۱). در واقع ماتریس‌های کوواریانس برآورد شده از داده‌های کم، دارای مقادیر ویژه فشرده تر از توزیع واقعی داده‌ها هستند، که منجر به تخمین کوواریانس ضعیف می‌شود. این را می‌توان با کوچک کردن ماتریس‌های کوواریانس تنظیم حل کرد (لدویت و وولف^۵، ۲۰۰۴) نشان داده شده است که طبقه‌بندی sLDA برتر از طبقه‌بندی استاندارد LDA است.

طبقه‌بندی برپایه تحلیل متمایز فضایی-زمانی (STDA^۶)

(STDA) می‌تواند به عنوان یک توسعه چند بعدی از LDA در نظر گرفته شود و دقت آزمایشی ۳۰٪s حدود ۶۱٪ به دست آورد (ژانگ، ژوو، ژوآ، جین، وانگ و کیچوکی^۷، ۲۰۳۱). بر خلاف نقاط زمانی و کانال‌های مکانی متصل به هم که در LDA و سه روش پیشرفته فوق‌الذکر در نظر گرفته می‌شدند، STDA به جای بردار نمونه ویژگی، یک نمونه

طبقه‌بندی با تابع تفکیک خطی (LDA)

هدف LDA استفاده از ابرصفحه برای تفکیک داده‌های نشان دهنده طبقات مختلف است. برای یک مسئله دو کلاس، کلاس بردار ویژگی به این بستگی دارد که ناقل از کدام طرف بیش از صفحه باشد. LDA توزیع عادی داده‌ها را با ماتریس کوواریانس برابر برای هر دو کلاس فرض می‌کند. ابرصفحه جداکننده با جستجوی فزاینده بدست می‌آید که فاصله بین دو کلاس را به حداکثر می‌رساند و واریانس بین کلاس را به حداقل می‌رساند. این روش نیاز محاسباتی بسیار کمی دارد که آن را برای سیستم‌های BCI مناسب می‌کند. علاوه بر این، استفاده از این طبقه‌بندی ساده است و به طور کلی نتایج خوبی را ارائه می‌دهد (فواد، لیب، مبروک، شاروی، و سید^۱، ۲۰۲۰). LDA به عنوان یکی از محبوب‌ترین الگوریتم‌ها برای کاربردهای BCI، از ابرصفحه‌ها برای جداسازی دو کلاس داده استفاده می‌کند. اگرچه LDA به دلیل نیاز محاسباتی بسیار کم در سیستم BCI آنلاین قابل اجرا است، اما در طبقه‌بندی تک‌محور با تعداد کمی نمونه آموزش عملکرد ضعیفی دارد. بنابراین، برخی از نسخه‌های پیشرفته LDA برای حل این مشکل ارائه شده است (ایکسائو و همکاران، ۲۰۱۹).

طبقه‌بندی برپایه تطبیق الگوی متمایز (SWLDA)

SWLDA به دلیل تأثیر در کاهش ابعاد ویژگی، همیشه در مورد حجم نمونه کوچک به کار رفته است. ترکیبی از تجزیه و تحلیل گام‌های رو به جلو و روبه عقب برای انتخاب ویژگی‌های مناسب در مدل تفکیک استفاده می‌شود. SWLDA از رگرسیون حداقل مربعات معمولی برای وزن ویژگی‌های ورودی استفاده می‌کند و سپس برچسب‌های دو کلاس را پیش‌بینی می‌کند. فرایند تشخیص بدون ویژگی اولیه، آغاز می‌شود. مهم‌ترین ویژگی ورودی که مقدار p کمتر از ۰/۰۵ است، به مدل اضافه خواهد شد، و سپس کم اهمیت‌ترین ویژگی ورودی که مقدار p بیش از

- 1- Fouad, Labib, Mabrouk, Sharawy, & Sayed
- 2- Shrinkage Linear Discriminant Analysis
- 3- Shrinkage Linear Discriminant Analysis
- 4- Blankertz, Lemm, Treder, Haufe, & Müller
- 5- Ledoit & Wolf,
- 6- Spatial-Temporal Discriminant Analysis
- 7- Zhang, Zhou, Zhao, Jin, Wang, & Cichocki

توزیع پیش‌بینی برای هر نمونه آزمون که به عنوان نمره ردیف یا ستون مرتبط در نظر گرفته شده، محاسبه شده است (سلیم و همکاران، ۲۰۰۹).

طبقه‌بندی کننده جنگل تصادفی (RF^۳)

طبقه‌بندی کننده‌های تصادفی جنگل (RF) مجموعه‌ای از چندین طبقه‌بندی کننده درخت تصمیم هستند. ایده پشت این طبقه‌بندی این است که به طور تصادفی زیرمجموعه‌ای از ویژگی‌های موجود را انتخاب کنید و یک طبقه‌بندی درخت تصمیم‌گیری را روی آنها آموزش دهید، سپس فرایند را با زیرمجموعه‌های بسیاری از ویژگی‌های تصادفی تکرار کنید تا درختان تصمیم‌گیری زیادی تولید شود، از این رو جنگل تصادفی نامیده می‌شود. تصمیم نهایی با ترکیب خروجی‌های همه درختان تصمیم گرفته می‌شود. از آنجا که هر درخت فقط از زیرمجموعه‌ای از ویژگی‌ها استفاده می‌کند، حساسیت کمتری در لعنت ابعاد دارد و بنابراین برای اثربخشی به داده‌های آموزشی کمتری نیاز دارد. خارج از تحقیقات BCI، در میان طبقه‌بندی کننده‌های مختلف و در مورد مشکلات مختلف طبقه‌بندی و دامنه‌ها، الگوریتم‌های جنگل تصادفی اغلب در میان دقیق‌ترین طبقه‌بندی‌ها از جمله مشکلات مجموعه داده‌های آموزش کوچک شناخته می‌شوند. از RF حتی به صورت آنلاین هم برای BCI مبتنی بر ERP و هم برای تصاویر موتور BCI با موفقیت استفاده شد. آنها از طراحی مبتنی بر طبقه‌بندی LDA برای تصاویر موتور بهتر عمل نکردند (لاتی و همکاران، ۲۰۰۷).

در ادامه در جدول شماره (۱) کارهای انجام شده در زمینه تکنیک‌های طبقه‌بندی مورد استفاده در سینگال sepper-p300 رابط مغزو کامپیوتر نشان داده شده است:

ویژگی مکانی-زمانی جدید را ایجاد می‌کند. STDA با ساختن نمونه‌های ویژگی دو طرفه، تجزیه و تحلیل تفکیک را به طور مشترک در ابعاد مکانی و زمانی اتخاذ می‌کند و سپس ویژگی‌ها را به صورت متناوب و مکانی بهینه می‌کند. این روش ابعاد ویژگی را کاهش می‌دهد و از این رو صحت طبقه‌بندی را حتی با چند نمونه آموزش بهبود می‌بخشد. در همین حال، این روش به منظور افزایش قابلیت تعمیم طبقه‌بندی آموزش دیده کمک می‌کند (ایکسانوو و همکاران، ۲۰۱۹).

طبقه‌بندی برپایه تحلیل تمایز خطی بیزی و فیشر (BLDA^۱) و (FLDA^۲)

بر اساس BLDA، LDA ماتریس کوواریانس نمونه‌های آموزشی را با استفاده از دانش قبلی آزمایشات می‌سازد. BLDA یک الگوریتم احتمالی مبتنی بر رگرسیون بیزی است و در شرایطی که مجموعه‌های آموزشی کمی داشته یا نویز صوتی زیادی داشته باشد، عملکرد بهتری نسبت به LDA دارد (ایکسانوو و همکاران، ۲۰۱۹). تجزیه و تحلیل خطی بیزی (BLDA) توسعه یافته تجزیه و تحلیل خطی فیشر (FLDA) است، که یک روش ساده اما کارآمد برای یادگیری ماشین است. به طور خاص، چارچوبی از یادگیری ماشین بیزی، اصطلاحاً چارچوب شواهد، در FLDA اعمال می‌شود. نسخه بیزی FLDA از نظر صحت طبقه‌بندی از FLDA ساده بهتر عمل می‌کند. در این الگوریتم اهداف رگرسیون از برجسب‌های کلاس مجموعه داده‌های آموزش به منظور اعمال FLDA از طریق رگرسیون محاسبه شد. سپس یک برآورد تکراری از روش پارامترها استفاده شد. برای پیش‌بینی مرحله آزمون، هر نمونه آزمون مرتبط با یک ردیف یا ستون خاص متعلق به کلاسی است که میانگین

1- Bayesian Linear Discriminant Analysis

2- Fisher Linear Discriminant Analysis

3- Random Forest

جدول شماره (۱): بررسی تکنیک‌های طبقه‌بندی مورد استفاده در در سینگال sepiller-p300 رابط مغز و کامپیوتر

محققین	عنوان	شرح	نتایج
لاتی و همکاران (۲۰۰۷)	مقایسه روشهای ارائه شده در زمینه طبقه بندی BCI در ده سال اخیر	در این مقاله، مقالات سال‌های ۲۰۰۷-۲۰۱۷ جهت طبقه‌بندی سیستم‌های BCI مورد بررسی و ارزیابی قرار گرفت. رویکردهای مورد بررسی در چهار دسته طبقه بندی شد: طبقه بندی تطبیقی، طبقه بندی ماتریسی و transfer learning و یادگیری عمیق	نتایج تحقیق نشان می‌دهد، روش طبقه بندی تطبیقی برای نظارت شده و بدون نظارت عملکرد بهتری از طبقه بند استاتیک دارد. طبقه بند ماتریس و transfer learning قابلیت اعتماد بالایی دارد و transfer learning برای زمانی که داده آموزشی کمی در دسترس باشد عملکرد بهتری دارد و جنگل تصادفی و LDA عملکرد خوبی دارد. یادگیری عمیق عملکرد بالایی در BCI نداشته است.
سلیم و همکاران (۲۰۰۹)	الگوریتم‌های یادگیری مورد استفاده (GAT)، (FLDA)، (SVM)، (BLDA)،	با استفاده از PCA جهت استخراج ویژگی موجب کاهش ابعاد شده است که موجب بهبود الگوریتم طبقه بندی شده است.	تجزیه و تحلیل تشخیصی خطی (LDA)، (BLDA SVM) بالاترین دقت را برای هر ۳ داشتند
لی، کیو، لی، جی، لیو، ویو (۲۰۱۵)	استفاده از الگوریتم SVM	پس از استخراج ویژگی، از svm گروهی برای طبقه بندی سیگنال‌های EEG استفاده می‌شود. برای ایجاد تفاوت در طبقه بندی فرعی، از آموزش‌های متفاوت جهت آموزش زیر طبقه بندی‌های مختلف استفاده می‌شود (۵ جلسه آموزش و ۱ جلسه تست). دقت از میانگین نتایج پنج جلسه محاسبه می‌شود.	نتایج بدست آمده، نشان می‌دهد الگوریتم پیشنهادی نسبت به الگوریتم svm برای زمانی که الگو چهره آشنا باشد دقت بالاتری دارد.
کاندو و آری (۲۰۱۷)	بکارگیری از الگوریتم‌های SVM و ESVM	طبقه بندی گروهی، به هر طبقه بند یک امتیاز اختصاص داده می‌شود. از تکنیک‌های متفاوتی، مانند روش‌های نرمال سازی min و max نرمال سازی نمره Z-score، متوسط و انحراف مطلق متوسط در امتیازدهی استفاده می‌شود که بهبود طبقه بندی شده است.	الگوریتم ESVM نسبت به SVM عملکرد بهتر دارد.
کاندو و آری (۲۰۱۸)	بکارگیری الگوریتم‌های WESVM و PCA	از ترکیب الگوریتم WESVM و WELDA جهت طبقه بندی استفاده می‌شود، زمانی که ابعاد داده کم باشد از LDA به علت عملکرد بهتر در داده با ابعاد کم استفاده می‌شود و در ابعاد داده بالا از EWSVM استفاده می‌شود.	نتایج تحقیق نشان می‌دهد که کاهش ابعاد توسط PCA باعث بهبود دقت کلاس بندی توسط WESVM می‌شود.

جدول شماره (۱): بررسی تکنیک‌های طبقه‌بندی مورد استفاده در در سینگال sepper-p۳۰۰ رابط مغز و کامپیوتر			
ایکسانو و همکاران (۲۰۱۹)	مقایسه الگوریتم‌های طبقه‌بندی SWLDA, BLDA, SKLDA STDA, DCPM	در این مطالعه، الگوریتم DCPM را برای طبقه‌بندی sepper-p۳۰۰ در مقایسه با پنج طبقه‌بندی کننده دیگر، LDA, SWLDA, BLDA, SKLDA و STDA عملکرد به مراتب بهتری داشت. مزیت DCPM نسبت به سایر روش‌ها برای زمانی است که نمونه آموزش به اندازه کافی در دسترس نباشد، این الگوریتم نتایج بهتری ارائه می‌دهد.	روش DCPM نسبت به سایر روش‌های داری دقت دارد.
فواد، و همکاران (۲۰۲۰)	مقایسه الگوریتم‌های مختلف کلاس‌بندی	این کار در مورد عملکرد الگوریتم‌های مختلف یادگیری ماشین بحث می‌کند ابتدا یک مرحله پیش پردازش برای استخراج ویژگی مهم قبل از استفاده از الگوریتم‌های یادگیری ماشین معرفی می‌شود. الگوریتم‌های ارائه شده عبارتند از: (LDA SVM I), LDA SVM II, SVM III, SVM IV, (LREG), (BLDA)	نتایج تحقیق نشان می‌دهد، الگوریتم SVMIV و الگوریتم BLDA بالاترین عملکرد را دارند.
مووالا، مورالس، مارتینز، آلچین و تامپسون (۲۰۲۰)	استفاده از الگوریتم‌های LS, SWLAD, SAE	در روش پیشنهادی در بررسی میران تاخیر در الگوریتم کلاس‌بندی پیشنهادی بررسی می‌شود. در این مقاله ضرایب رگوسیون از رابطه بین VCBL و دقت بدست می‌آید.	نتایج تحقیق نشان می‌دهد، زمانی که ویژگی‌ها کم باشد الگوریتم طبقه‌بندی ۵ عملکرد بهتری دارد و با تعداد ویژگی بالا الگوریتم SWLAD و SAE عملکرد بهتری دارد ولی به طور کلی عملکرد الگوریتم SWLD نسبت به دو روش دیگر بهتر است.
بیانچی، لیتی، لیوزی، پیکالی و سالواتوره (۲۰۲۱)	ترکیب الگوریتم SVM الگوریتم بهینه سازی OSBF	ابتدا، برای هر شرکت کننده مقدار نمرات بهینه با حل یک مسئله بهینه سازی بر روی داده‌های آموزشی وی تعیین می‌شود. در مرحله دوم، با حل یک نسخه اصلاح شده از بهینه سازی، یک روش کارآمد متوقف سازی اولیه پیاده‌سازی می‌شود و در پایان حداکثر ارزش برای تابع تصمیم گیری تعیین می‌شود. استفاده از مسئله MILP برای اختصاص "نمره اطمینان" به طبقه‌بندی هر یک محرک در هر تکرار سبب افزایش دقت کلاس‌بندی شده است.	نتایج نشان می‌دهد که الگوریتم پیشنهادی موجب افزایش دقت و نرخ انتقال در یک محیط بدون توقف و با توقف می‌شود

۳- نتیجه‌گیری

در این مقاله، انواع طبقه‌بندی $sepller-p300$ بررسی شده است. الگوریتم‌های یادگیری ماشین از جمله الگوریتم‌های $DCPM$ ، $LDAT$ ، $SWLDA$ ، $BLDA$ ، $SKLDA$ ، $DCMP$ و $STDA$ و الگوریتم‌های تطبیقی و وزن‌دهی را در زمینه طبقه‌بندی برای طراحی سیستم رابط مغز و کامپیوتر (BCI) مبتنی بر $sepller-p300$ مرور شده است. براساس مقالات مورد بررسی $DCPM$ نسبت به سایر روش‌ها برای زمانی است که نمونه آموزش به اندازه کافی در دسترس نباشد، نتایج بهتری ارائه می‌دهد. زمانی که ویژگی‌ها کم باشد الگوریتم طبقه‌بندی Is عملکرد بهتری دارد و با تعداد ویژگی بالا الگوریتم $SWLAD$ و SAE عملکرد بهتری دارد ولی به طور کلی عملکرد الگوریتم $SWLD$ نسبت به دو روش

دیگر بهتر است. کاهش ابعاد توسط PCA باعث بهبود دقت کلاس‌بندی توسط $WESVM$ می‌شود. SVM گروهی برای زمانی که الگو چهره آشنا باشد دقت بالاتری دارد از تکنیک‌های متفاوتی، مانند روش‌های نرمال‌سازی min و max نرمال‌سازی نمره Z -score، در دقت طبقه‌بندی SVM تاثیر بالایی دارد. نتایج تحقیق مروری مورد بررسی نشان می‌دهد، روش طبقه‌بندی تطبیقی برای نظارت شده و بدون نظارت عملکرد بهتری از طبقه‌بندی استاتیک دارد. طبقه‌بندی ماتریس و $transfer learning$ قابلیت اعتماد بالایی دارد و $transfer learning$ برای زمانی که داده آموزشی کمی در دسترس باشد عملکرد بهتری دارد و جنگل تصادفی و LDA عملکرد خوبی دارد. یادگیری عمیق عملکرد بالایی در BCI نداشته است.

منابع:

- 1-Bianchi, L., Liti, C., Liuzzi, G., Piccialli, V., & Salvatore, C. (2021). Improving P300 Speller performance by means of optimization and machine learning. *Annals of Operations Research*, 1-39.
- 2-Blankertz, B., Lemm, S., Treder, M., Haufe, S., & Müller, K. R. (2011). Single-trial analysis and classification of ERP components—a tutorial. *NeuroImage*, 56(2), 814-825.
- 3-Chakladar, D. D., & Chakraborty, S. (2019). Feature extraction and classification in brain-computer interfacing: Future research issues and challenges. *Natural Computing for Unsupervised Learning*, 101-131.
- 4-Fouad, I. A., Labib, F. E. Z. M., Mabrouk, M. S., Sharawy, A. A., & Sayed, A. Y. (2020). Improving the performance of P300 BCI system using different methods. *Network Modeling Analysis in Health Informatics and Bioinformatics*, 9(1), 1-13.
- 5-Hyvärinen, A. (2013). Independent component analysis: recent advances. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 371(1984), 20110534.
- 6-Koles, Z. J., Lazar, M. S., & Zhou, S. Z. (1990). Spatial patterns underlying population differences in the background EEG. *Brain topography*, 2(4), 275-284.
- 7-Kotsiantis, S. B., Kanellopoulos, D., & Pintelas, P. E. (2006). Data preprocessing for supervised learning. *International journal of computer science*, 1(2), 111-117.
- 8-Kundu, S., & Ari, S. (2017, July). Score normalization of ensemble SVMs for brain-computer interface P300 speller. In *2017 8th International Conference on Computing, Communication and Networking Technologies (ICCCNT)* (pp. 1-5). IEEE.
- 9-Kundu, S., & Ari, S. (2017, July). Score normalization of ensemble SVMs for brain-computer interface P300 speller. In *2017 8th International Conference on Computing, Communication and Networking Technologies (ICCCNT)* (pp. 1-5). IEEE.
- 10-Ledoit, O., & Wolf, M. (2004). A well-conditioned estimator for large-dimensional covariance matrices. *Journal of multivariate analysis*, 88(2), 365-411.
- 11-Li, Q., Li, J., Liu, S., & Wu, Y. (2015, January). Improving the performance of P300-speller with familiar face paradigm using support Vector machine ensemble. In *2015 International Conference on Network and Information Systems for Computers* (pp. 606-610). IEEE.
- 12-Lotte, F., Congedo, M., Lécuyer, A., Lamarche, F., & Arnaldi, B. (2007). A review of classification algorithms for EEG-based brain-computer interfaces. *Journal of neural engineering*, 4(2), R1.
- 13-Mowla, M. R., Gonzalez-Morales, J. D., Rico-Martinez, J., Ulichnie, D. A., & Thompson, D. E. (2020). A comparison of classification techniques to predict brain-computer interfaces accuracy using classifier-

based latency estimation. *Brain Sciences*, 10(10), 734.

14-Nasir, V., Cool, J., & Sassani, F. (2019). Intelligent machining monitoring using sound signal processed with the wavelet method and a self-organizing neural network. *IEEE Robotics and Automation Letters*, 4(4), 3449-3456.

15-Patel, V., & Patel, M. S. (2019, July). Brain computer interface: Applications and p300 speller overview. In *2019 10th International Conference on Computing, Communication and Networking Technologies (ICCCNT)* (pp. 1-5). IEEE.

16-Sampanna, R., & Mitaim, S. (2018). Noise benefits in the array of brain-computer interface classification systems. *Informatics in Medicine Unlocked*, 12, 88-97.

17-Selim, A. E., Wahed, M. A., & Kadah, Y. M. (2009, March). Machine learning methodologies in P300 speller Brain-Computer Interface systems. In *2009 National Radio Science Conference* (pp. 1-9). IEEE.

18-Subha, D. P., Joseph, P. K., Acharya U, R., & Lim, C. M. (2010). EEG signal analysis: a survey. *Journal of medical systems*, 34(2), 195-212.

19-Xiao, X., Xu, M., Wang, Y., Jung, T. P., & Ming, D. (2019, July). A comparison of classification methods for recognizing single-trial P300 in brain-computer interfaces. In *2019 41st Annual International Conference of the IEEE Engineering in Medicine and Biology Society (EMBC)* (pp. 3032-3035). IEEE.

20-Zhang, Y., Zhou, G., Zhao, Q., Jin, J., Wang, X., & Cichocki, A. (2013). Spatial-temporal discriminant analysis for ERP-based brain-computer interface. *IEEE Transactions on Neural Systems and Rehabilitation Engineering*, 21(2), 233-243.

©Authors, Published by Journal of Intelligent Knowledge Exploration and Processing. This is an open-access paper distributed under the CC BY (license <http://creativecommons.org/licenses/by/4.0/>).

