

دریافت مقاله: ۱۴۰۳/۰۱/۲۹

پذیرش مقاله: ۱۴۰۳/۰۳/۰۱

نوع مقاله: پژوهشی

تشخیص بیماری‌های مزمن به کمک ماشین بردار پشتیبان دوقلو با قيود نرم

حمیده فدیشه‌ای

دانشجوی کارشناسی ارشد، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد، مشهد، ایران

پست الکترونیکی: ha.fadishehei@mail.um.ac.ir

جلال‌الدین نصیری*

استادیار، گروه ریاضی کاربردی، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد، مشهد، ایران

پست الکترونیکی: Jnasiri@um.ac.ir

سهراب عفتی

استاد، گروه ریاضی کاربردی، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد، مشهد، ایران

پست الکترونیکی: s-effati@um.ac.ir

چکیده

بیماری‌های مزمن از دلایل اصلی مرگ‌ومیر در سال‌های اخیر بوده‌اند. پژوهش‌های تجربی بسیاری برای تشخیص بیماری مزمن در کشورهای مختلف گزارش شده‌اند. در مطالعات انجام شده، روش‌های یادگیری ماشین، نقش به‌سزایی در تشخیص بیماری‌های مزمن داشته‌اند. از جمله این روش‌ها، می‌توان به ماشین بردار پشتیبان استاندارد و دوقلو اشاره کرد. داده‌های بیماری‌های مزمن دارای نوفه ذاتی و نمونه‌های پرت می‌باشند. این مهم باعث افت شدید دقت مدل‌های ماشین بردار پشتیبان استاندارد و دوقلو می‌گردد. در مدل جدید، برای برطرف کردن کاستی بیان شده و رفع حساسیت نسبت به داده‌های نوفه‌ای^۱، به جای دسته‌بندی ماشین بردار پشتیبان دوقلو، ماشین بردار پشتیبان دوقلو با قيود نرم پیشنهاد شده است. این مدل فضای شدنی مسئله بهینه‌سازی درجه دوم را گسترش

می‌دهد. با گسترش فضای شدنی مسئله، به نمونه‌ها اجازه تخطی از ابرصفحات داده می‌شود و تأثیر داده‌های نوفه‌ای و دور افتاده در تشخیص نهایی کاهش می‌یابد. الگوریتم پیشنهادی، بر روی پنج مجموعه داده از داده‌های بالینی پزشکی، اجرا و نتایج آن با روش‌های ماشین بردار پشتیبان، درخت تصمیم، رگرسیون لجستیک و نزدیکترین همسایه مقایسه شده‌اند. میانگین دقت روش پیشنهادی نسبت به بهترین میانگین در روش‌های پیشین، حاکی از بهبود حدود ۷٪ در نتایج تجربی می‌باشد.

واژه‌های کلیدی: یادگیری ماشین، ماشین بردار پشتیبان دوقلو، بهینه‌سازی محدب، محدودیت‌های فازی، تحمل، اطمینان.

مقدمه

در این بخش ابتدا ضرورت انجام پژوهش و سپس

* نویسنده مسئول

بعضی از مفاهیم توضیح داده خواهند شد.

بر اساس گزارش سازمان بهداشت جهانی ویژگی مشترک بیماری‌های مزمن عبارتند از:

- چند دهه زمان لازم است که همه‌گیری آن‌ها به‌طور کامل تثبیت شود.
- ریشه این بیماری‌ها در سنین جوانی است.
- با توجه به زمان طولانی پیش از بروز آن‌ها، فرصت‌های زیادی برای پیش‌گیری از آن‌ها وجود دارد.
- درمان این بیماری‌ها به یک رویکرد نظام‌مند و بلندمدت نیاز دارد.

اگرچه برخی از بیماری‌های مزمن زمینه‌های ارثی دارند یا به دلیل بالا رفتن سن بروز پیدا می‌کنند، پژوهش‌ها نشان داده است که بخش بزرگی از دلایل بروز این بیماری‌ها را می‌توان به تغییر شرایط محیطی، سبک زندگی و رفتار فرد مثل استرس، فعالیت بدنی، سوء تغذیه، مصرف دخانیات و الکل نسبت داد [۶]. تغییرات ذکر شده منجر به چهار تغییر متابولیک/ فیزیولوژیک- فشار خون بالا، اضافه وزن و چاقی، گلوکز خون بالا، کلسترول خون بالا و در نهایت منجر به بیماری‌های مزمن می‌شوند [۷].

با توجه به تأثیر قابل توجهی که بیماری‌های مزمن بر نظام سلامت دارند، طی سال‌های اخیر مطالعات بسیاری به بررسی تأثیر اقتصادی بیماری‌های مزمن پرداخته‌اند. تأثیر معنادار این بیماری‌ها بر تولید ناخالص داخلی و نرخ رشد درآمد سرانه نشان می‌دهد که این اثرات ریشه در تولید اقتصادی در دو سطح فردی و خانواده دارد.

بیماری‌های مزمن به صورت مستقیم و غیرمستقیم بر اقتصاد خانواده‌ها تأثیر می‌گذارند. به‌عنوان مثال بیماری‌های مزمن ۸۶٪ از هزینه‌های سلامت خانوار در سال ۲۰۱۶ ایالات متحده را در بر گرفته است. اما بدیهی است که بار^۲ اقتصادی بیماری‌های مزمن تنها هزینه‌های مربوط به خدمات درمانی نیست.

مهمترین هزینه دیگر، درآمد از دست رفته است که از کاهش بهره‌وری ناشی از بیماری نشئت می‌گیرد. همچنین گزارش سال ۲۰۰۵ سازمان بهداشت جهانی نشان می‌دهد

ضرورت انجام پژوهش

بیماری‌های مزمن به دلیل طولانی بودن دوران ابتلا، عوارض آزاردهنده مختلفی دارند. این بیماری‌ها زمینه‌سازی برای مبتلا شدن به بیماری‌های عفونی دیگر از جمله کرونا می‌شوند. از این‌رو، بیماری‌های مزمن از مهم‌ترین چالش‌های نظام سلامت در جهان هستند. در سال ۲۰۱۹، حدود ۶۳٪ از مرگ و میر جهانی ناشی از بیماری‌های مزمن بوده است. رایج‌ترین این بیماری‌ها، فشار خون، قند خون، و بیماری‌های قلبی هستند [۱]. عوارض ناشی از آن‌ها باعث آسیب رسیدن به قلب، کلیه، مغز و چشم‌ها می‌شود که آثاری مخرب روی کار و زندگی افراد خواهد داشت [۲]. افزون بر موارد ذکر شده این بیماران در برابر بیماری‌های عفونی (مثلاً ویروس کرونا) ایمنی کمتری دارند [۳]. بر پایه گزارش‌های سال ۲۰۱۹، ۴۸٪ افراد مبتلا به این ویروس بیماری مزمن نیز داشته‌اند و احتمال بروز علائم شدید در آن‌ها بیشتر بوده است [۴]. بنابراین تشخیص زودهنگام بیماری کمک قابل توجهی به سلامت فرد و جامعه خواهد داشت.

بنابر تعریف سازمان بهداشت جهانی، بیماری‌های مزمن به آن دسته از بیماری‌ها گفته می‌شود که منجر به آسیب در ساختار یا کاهش عملکرد بدن شده به گونه‌ای که سبب تغییر در زندگی عادی بیمار گردند و طی یک دوره زمانی طولانی ادامه یافته و پایدار باشند [۵].

بیماری‌های مزمن اصلی عبارتند از: بیماری‌های قلبی-عروقی، انواع سرطان، انسداد ریه، نارسایی مزمن کلیه و دیابت. البته بیماری‌های مزمن دیگری از جمله بیماری روانی، اختلالات بینایی، شنوایی، ژنتیک و بیماری‌های مفصلی نیز وجود دارند. از آن‌جا که عوامل محیطی مثل شرایط اجتماعی، سبک زندگی و آلاینده‌های زیست‌محیطی نقش فزاینده‌ای در پیدایش و پیشرفت این بیماری‌ها دارند از آن‌ها با عنوان بیماری‌های مزمن محیطی نیز یاد می‌شود.

برای بررسی ارتباط عوامل خطر بالقوه با دیابت نوع ۲ به کار گرفته شدند. شبکه‌های عصبی با دقت ۸۲/۴٪ بهترین نتیجه را داشتند.

هانگ و همکاران [۱۱] با استفاده از ماشین بردار پشتیبان، جنگل تصادفی و تقویت تطبیقی^۲، ابتلا به بیماری مزمن کلیوی را با ۲۱۴۲ فرد شرکت‌کننده در پژوهش، پیش‌بینی کردند. ماشین بردار پشتیبان با ۸۵/۷٪ بیشترین دقت را به دست آورد.

از آن‌جا که عکاسی از فوندوس شبکیه یک رویکرد غیرتهاجمی برای پیش‌بینی بیماری‌های مزمن می‌باشد، ژانگ و همکاران [۱۲] با روش ذکرشده، پژوهشی مقطعی در مناطق روستایی در مرکز چین انجام دادند. ۱۲۲۲ تصویر برداری شبکیه با بیش از ۵۰ اندازه‌گیری بر روی ۶۲۵ نفر انجام شد. مدل شبکه عصبی پیاده‌سازی شد. پژوهش یادشده، به یک سطح زیر منحنی ROC با $AUC = 0.766$ برای پیش‌بینی فشارخون و دقت ۶۸/۸٪ دست یافت.

العظم و همکاران [۱۳] به مقایسه عملکرد روش‌های یادگیری نظارت‌شده و نیمه نظارت شده در پیش‌بینی سرطان سینه پرداختند. نه الگوریتم، رگرسیون لجستیک، دسته‌بند گوسی نایو بیز، ماشین بردار پشتیبان ساده و با هسته گوسی، درخت تصمیم، جنگل تصادفی، گرادیان افزایشی، دورترین گرادیان افزایشی و نزدیک‌ترین همسایه مورد بررسی قرار گرفتند. برای آموزش و آزمایش مدل‌ها، مجموعه داده‌های سرطان ویسکانسین استفاده شدند. مدل‌های نزدیک‌ترین همسایگی و افزایش گرادیان با دقت ۹۸٪ بهترین پیش‌بینی را داشتند.

شی و همکاران [۱۴] مدل‌های نزدیک‌ترین همسایه، ماشین بردار پشتیبان، جنگل تصادفی، شبکه عصبی و دسته‌بند گوسی بیز ساده و افزایش انطباقی را با اطلاعات پزشکی مربوط به ۱۵۳ بیمار مشکوک به سرطان ریه، مورد آموزش و آزمایش قرار دادند. بیز ساده داده‌های آموزش و آزمایش به دقت ۱۰۰٪ دست یافت. نویسندگان عقیده داشتند که ترکیب روش‌های تشخیص با نشانگرهای

که بیماری‌های مزمن حدود ۴۰٪ از سال‌های مولد را در بازار کار کشورهای نوظهور از بین برده است.

از آن‌جا که بیماری‌های مزمن هزینه‌های بالایی را به فرد تحمیل می‌کنند، ممکن است خانواده‌ها مجبور به افزایش عرضه نیروی کار، کاهش مصرف دیگر کالاها یا دریافت درآمدهای انتقالی از جمله هدایا و کمک‌هایی از دیگر افراد یا نهادهای کشور شوند. شوک هزینه‌ای ناشی از بیماری مزمن باعث کاهش درآمد ناشی از کار و افزایش هزینه‌های پزشکی، افزایش درآمدهای انتقالی و کاهش مصرف خوراکی و غیرخوراکی شده است. بنابراین این بیماری‌ها همچنان چالش مهمی در حوزه بهداشت عمومی در تمام کشورها، از جمله کشورهای با درآمد کم و متوسط، هستند چرا که در این کشورها، بیش از سه چهارم مرگ‌ومیر به دلیل ابتلا به بیماری‌های مزمن رخ می‌دهد.

با توجه به موارد گفته‌شده، وجود یک سیستم هوشمند که بتواند به پیش‌بینی زودهنگام بیماری‌های مزمن دست یابد مورد نیاز است. تشخیص بهنگام بیماری می‌تواند آثار روحی، روانی، اجتماعی و اقتصادی بسیاری بر فرد و جامعه داشته باشد.

تاریخچه‌ای از تشخیص بیماری‌های مزمن

برای تشخیص بیماری‌های مزمن خاص چندین الگوریتم یادگیری ماشین پیشنهاد شده است [۹].

منیرالزمان و همکاران [۱۰] باور داشتند که گسترش مدل‌های پیش‌بینی برای شناسایی دلایل ابتلا به دیابت نوع ۲، می‌تواند به تشخیص بیماری و مداخله زودهنگام و کاهش هزینه‌های پزشکی کمک کند. در پژوهش مذکور ۱۳۸۱۴۶ فرد شرکت داشتند که ۲۰۴۶۷ نفر از آن‌ها به دیابت نوع ۲ مبتلا بودند. برای پیش‌بینی ابتلا به دیابت نوع ۲، روش‌های مختلف یادگیری ماشین شامل ماشین بردار پشتیبان، درخت تصمیم، رگرسیون لجستیک، جنگل تصادفی، دسته‌بند گوسی بیز ساده و از مدل‌های رگرسیون لجستیک و زن‌دار تک متغیره و چند متغیره

زیستی و روش‌های یادگیری ماشین قدرت پیش‌بینی ابتلا به سرطان ریه را بالا می‌برد.

بیمارانی که در مرحله نهایی بیماری کلیوی هستند، مستعد ابتلا به نوعی بیماری قلبی ویژه نیز می‌باشند. مزاتستا و همکاران [۱۵] دو مجموعه داده متفاوت را مورد بررسی قرار دادند. الگوریتم‌های رگرسیون لجستیک، نزدیک‌ترین همسایه، ماشین بردار پشتیبان خطی، ماشین بردار پشتیبان با هسته گوسی و چندجمله‌ای، درخت تصمیم و دسته‌بند گوسی بیز ساده، مرگ و ابتلای به بیماری‌های قلبی عروقی در بیماران دیالیزی را پیش‌بینی کردند. بهترین نتایج مربوط به دسته‌بند ماشین بردار پشتیبان با کرنل گوسی^۵ با دقت ۹۵/۲۵٪ و ۹۲/۱۵٪ بود. شیائو هان و همکاران [۱۶] اولین چارچوب عمومی و کارآمد را برای تشخیص بیماری‌های مزمن ارائه کردند. روش آن‌ها علاوه بر استخراج ویژگی‌های سطح بالا، مشکل عدم تعادل جدی موجود در داده‌های بیماری‌های مزمن و برآزش بیش از حد را نیز حل کرد. ایشان مدل مبتنی بر شبکه عصبی را معرفی کردند که با در نظر گرفتن وزن، نمونه‌ها را تقویت می‌کرد. داده‌هایی که در شبکه عصبی چندجمله‌ای به اشتباه دسته‌بندی شده بودند به طور ویژه مورد توجه قرار گرفتند. مدل پیشنهادی با نه مجموعه داده، شامل داده‌های مقایسه‌ای و نمونه‌های واقعی در ابعاد بالا مورد آموزش و آزمایش قرار گرفت. روش جدید در مقایسه با ماشین بردار پشتیبان، رگرسیون لجستیک، درخت تصمیم، نزدیک‌ترین همسایگی و پرسپترون چند لایه، دارای دقت قابل توجهی بود. مقابله با نوفه ذاتی و عدم تعادل داده‌ها در بیماری‌های مزمن در حال حاضر نیز مورد توجه پژوهشگران می‌باشد.

همان‌طور که ملاحظه گردید، روش ماشین بردار پشتیبان جزو روش‌های تشخیص زودهنگام بیماری‌های مزمن می‌باشد. از طرف دیگر SVM استاندارد، نسبت به نمونه‌های نوفه‌ای و پرت بسیار حساس است. تقویت دسته‌بند ماشین بردار پشتیبان به صورتی که برای داده‌های

نوفه‌ای دقت خوبی را کسب کند، می‌تواند در کارایی سیستم هوشمند تشخیص بیماری مزمن مفید باشد. در این پژوهش، برای دستیابی به هدف بیان شده، مدل مسئله بهینه‌سازی مربوط به ماشین بردار پشتیبان دوقلو، تغییر یافته است و نامساوی‌ها به صورت فازی در نظر گرفته شده‌اند. در نتیجه این تغییر، فضای شدنی مسئله گسترش یافته است. با این روش میزانی از تحمل خطا به الگوریتم، اعمال شده است. الگوریتم جدید، نسبت به داده‌های نوفه‌ای حساسیت کمتری دارد. برای تشخیص میزان کارایی الگوریتم نسبت به روش‌های پیشین، نرخ دقت^۶ روش نوین، محاسبه شده است. جهت ارزیابی، از روش اعتبار سنجی متقابل با $k = 10$ استفاده شده و در نهایت برای مقایسه با روش‌های مشابه، آزمون آماری فریدمن^۶ به کار رفته است.

ادامه ساختار مقاله به این صورت است: بخش دوم مروری بر پیشینه روش‌های ماشین بردار پشتیبان، خواهد بود. در بخش سوم روش پیشنهادی به تفصیل بیان خواهد شد. در بخش چهارم نتایج تجربی روش پیشنهادی، ارزیابی نتایج و در بخش پنجم نتیجه‌گیری و پیشنهادهایی برای پژوهش‌های بعدی ارائه خواهد شد. مراجع مورد استفاده در پژوهش به ترتیبی که در مقاله آمده‌اند در آخرین قسمت مشخص شده‌اند.

مروری بر روش‌های ماشین بردار پشتیبان

پس از الگوریتم‌های شبکه عصبی، در سال ۱۹۹۵ ایده جدیدتری به عنوان ماشین بردار پشتیبان توسط وپنیک و همکاران مطرح شد که در آن از روش‌های بهینه‌سازی مسایل درجه دوم، برای جداسازی داده‌های دو طبقه از یکدیگر استفاده می‌کرد. مبنای کار یافتن دو ابرصفحه موازی با بیشترین فاصله ممکن بود که داده‌های دو طبقه در دوطرف این ابرصفحات قرار گیرند [۱۷]. در طی سال‌ها انواع جدیدتر، دقیق‌تر و کاراتر از ماشین بردار پشتیبان ارائه شد. برخی از این روش‌ها در زیر بیان شده‌اند.

5-Accuracy
6-Friedman test

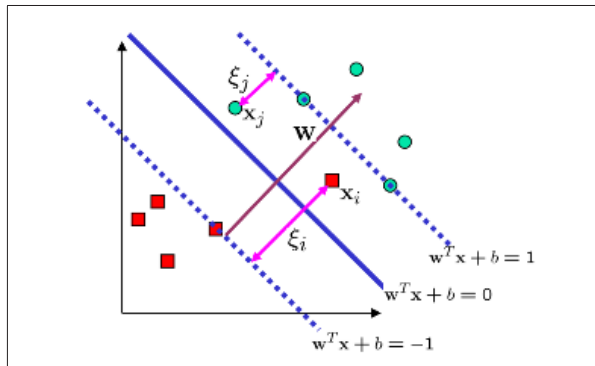
4-Radial Basis Function

$$\begin{aligned} & \text{Min}_{(w, b)} \frac{1}{\tau} \|w\|^\tau \\ & \text{S.t. } Aw + e_1 b \geq e_1, \\ & \quad Bw + e_r b \leq -e_r \end{aligned} \quad (2)$$

در (۲)، e_1 و e_r بردارهای با درایه‌های واحد و متناسب با A و B هستند. در حالتی که داده‌ها در فضای ورودی، به صورت خطی جداپذیر نباشند، مسئله SVM با حاشیه نرم به صورت (۳) مطرح خواهد شد:

$$\begin{aligned} & \text{Min}_{(w, b)} \frac{1}{\tau} \|w\|^\tau + C \sum_{i=1}^m \xi_i \\ & \text{S.t. } y_i(w^T x_i + b) \geq 1 - \xi_i \\ & \quad \xi_i \geq 0, i = 1, 2, \dots, m. \end{aligned} \quad (3)$$

که در آن ξ_i متغیر لغزش برای کاهش تأثیر داده نوفه‌ای x_i است. شکل (۲) نمایان‌گر ماشین بردار پشتیبان با حاشیه نرم می‌باشد.



شکل (۲): ماشین بردار پشتیبان با حاشیه نرم

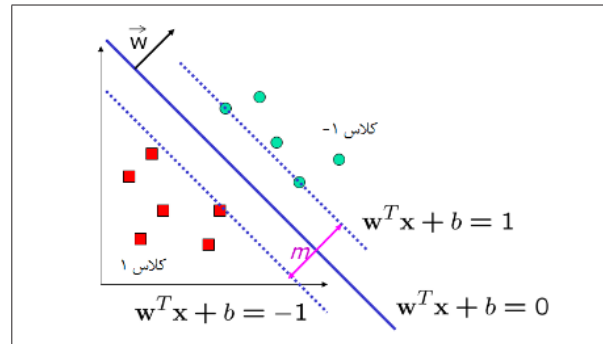
پس از نوشتن تابع لاگرانژ و در نظر گرفتن شرایط مکمل زاید، دوگان مسئله (۳) به صورت زیر است که با حل آن ابرصفحه جداساز معلوم خواهد شد.

$$\begin{aligned} & \text{Min}_{\alpha} \frac{1}{\tau} \sum_{i=1}^m \sum_{j=1}^m \alpha_i \alpha_j y_i y_j x_i^T x_j - \sum_{k=1}^m \alpha_k \\ & \quad \sum_{i=1}^m \alpha_i y_i = 0, \text{ S.t. } \\ & \quad 0 \leq \alpha_i \leq C, i = 1, 2, \dots, m. \end{aligned} \quad (4)$$

که α_i ضرایب لاگرانژ هستند و تابع تصمیم $D(x)$ طبق

ماشین بردار پشتیبان

فرض کنید که $X = \{(x_i, y_i)\}_{i=1}^m$ مجموعه نمونه‌های آموزشی با اندازه m باشد که هر $x_i \in R^n$ یک نمونه یادگیری دارای n ویژگی در فضای ورودی است، $y_i \in \{1, -1\}$ برچسب کلاس داده x_i و $A_{m_1 \times n}$ مجموعه نمونه‌های با برچسب $+1$ ، $B_{m_r \times n}$ مجموعه نمونه‌های با برچسب -1 و $m_1 + m_r = m$ باشند، اگر داده‌ها در فضای ورودی، به صورت خطی قابل جداسازی باشند، می‌توان ابرصفحه‌ای جداساز به فرم $w^T x + b = 0$ یافت که در آن w یک بردار n بعدی و b کمیتی عددی است. ماشین بردار پشتیبان به جای یافتن یک ابرصفحه جداکننده، به دنبال یافتن دو ابرصفحه موازی است که تا حد امکان از یکدیگر دور بوده و داده‌های دو طبقه در دو طرف آنها واقع باشند. شکل (۱) تعبیر هندسی ماشین بردار پشتیبان با حاشیه سخت را نمایش می‌دهد.



شکل (۱): ماشین بردار پشتیبان با حاشیه سخت

با مدل‌سازی مسئله بالا و حل مسئله بهینه‌سازی (۱)، می‌توان ابرصفحه جداساز را مشخص کرد.

$$\begin{aligned} & \text{Min}_{(w, b)} \frac{1}{\tau} \|w\|^\tau \\ & \text{S.t. } y_i(w^T x_i + b) \geq 1, i = 1, 2, \dots, m. \end{aligned} \quad (1)$$

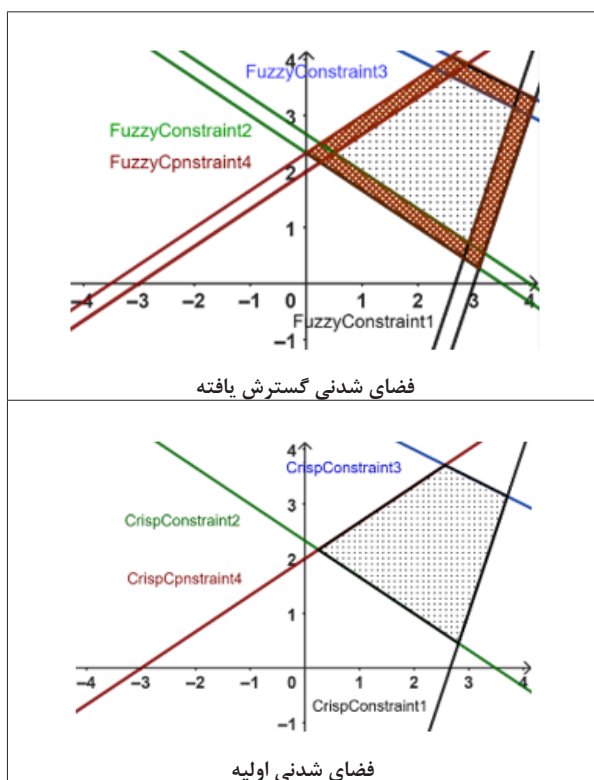
نمایش ماتریسی مسئله (۱) به صورت (۲) خواهد بود:

محدودیت‌ها به جای نامساوی‌های عادی از نامساوی‌های فازی استفاده شده است. این نوع از قیود جدید در فرمول (۶) بازنویسی شده است:

$$y_i(w^T x_i + b) \geq 1 - \xi_i, i = 1, \dots, m. \quad (6)$$

شکل (۴) یک مثال ساده از مسئله بهینه‌سازی (۷) با چهار قید منعطف و گسترش فضای شدنی مسئله را نشان می‌دهد.

$$\begin{aligned} & \text{Minimize } x_1 + 2x_2 \\ & \text{s.t. } \begin{cases} g_1(x_1, x_2) = 2x_1 - x_2 \lesssim 8 \\ g_2(x_1, x_2) = 2x_1 + 2x_2 \gtrsim 7 \\ g_3(x_1, x_2) = x_1 + 2x_2 \lesssim 12 \\ g_4(x_1, x_2) = -x_1 + 2x_2 \lesssim 3 \end{cases} \end{aligned} \quad (7)$$



شکل (۴): گسترش فضای شدنی برای مسئله (۷)

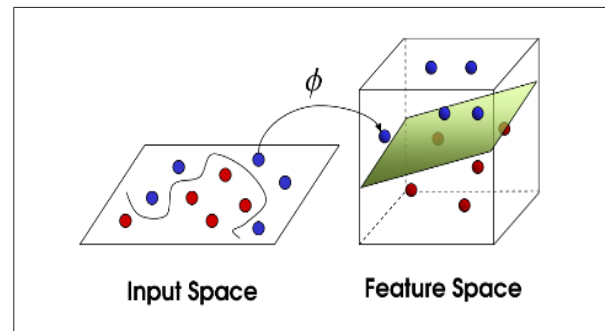
برای به دست آمدن جواب‌های مسئله، باید محدودیت‌های مسئله (۷) برآورده شوند.

علامت $\gtrsim$ به این معنی است که به برخی از نمونه‌های آموزش، اجازه تخطی از حدود هم داده می‌شود. اگر نمونه‌ای اهمیت کمتری داشته باشد، محدودیت نظیر آن

فرمول (۵) موقعیت داده جدید x را نسبت به ابرصفحه مشخص خواهد کرد:

$$\begin{aligned} D(x) &= \text{Sign}(w^T x + b) \\ &= \text{Sign}\left(\sum_{i \in S} \alpha_i y_i x_i^T x + b\right) \end{aligned} \quad (5)$$

در رابطه (۵)، S نشان‌دهنده مجموعه اندیس‌های بردارهای پشتیبان است. اگر نمونه‌ها در فضای ورودی جداپذیر خطی نباشند، نمونه‌ها با تکنیک کرنل مطابق شکل (۳) در فضای ویژگی با ابعاد بیشتر نگاشته می‌شوند. در این فضا، ماهیت خود داده، اهمیتی ندارد بلکه آنچه مهم است فاصله داده‌ها از یکدیگر است. با این روش ماشین بردار پشتیبان یک ابرصفحه جداکننده بهینه برای جداسازی نمونه‌ها پیدا می‌کند.



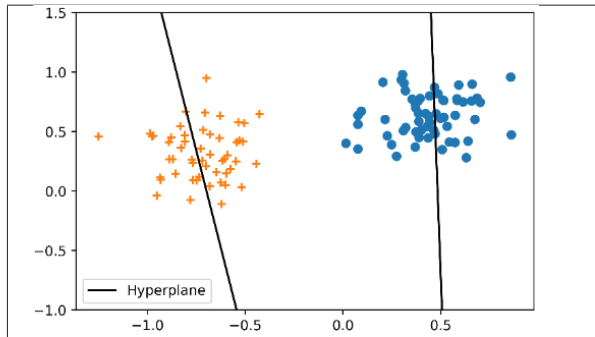
شکل (۳): نگاشت نمونه‌ها با روش کرنل

ماشین بردار پشتیبان با قیود نرم

در فرمول (۳) برای مسئله SVM، نظیر هر یک از نمونه‌ها یک محدودیت به صورت $y_i(w^T x_i + b) \geq 1 - \xi_i$ تعریف شد. بنابراین به تعداد نمونه‌ها محدودیت وجود دارد. مجموعه این محدودیت‌ها، یک فضای شدنی^۸ را به وجود می‌آورد که متغیرهای مسئله یعنی $(w, b, \xi_i) \in R^{m+1+n}$ از این فضا انتخاب می‌شوند. قیدهای به دست آمده در SVM حالت بدون انعطاف دارند. در روش ماشین بردار پشتیبان با قیود نرم، قیدهای مسئله انعطاف بیشتری دارند. در

8-Feasible Region

نصف ابعاد قبل و با مرتبه محاسباتی $\frac{1}{4}$ روش قبل حل می‌شوند. این بالاترین مزیت روش جدید به شمار می‌آید. گسترش‌های قابل توجهی از ماشین بردار پشتیبان دوقلو مطرح [۲۱] و کاربردهای متعددی از این روش ارائه شده است که از جمله آن‌ها می‌توان به تشخیص رفتار انسانی اشاره کرد [۲۲].



شکل (۵): ماشین بردار پشتیبان دوقلو

در توضیح الگوریتم ماشین بردار پشتیبان دوقلو، مجموعه‌ای از نمونه‌ها و اطلاعاتی در مورد آن‌ها موجود است که به هر نمونه یک برچسب اختصاص داده شده و به کمک برچسب‌ها (۱ یا -۱) می‌توان کل نمونه‌ها را به دو گروه دسته‌بندی کرد. بهتر است که توزیع داده‌ها متعادل باشد یعنی با نسبت تقریباً مساوی، با برچسب ۱ یا -۱ دسته‌بندی شده باشند. در روش بردار پشتیبان دوقلو، ابرصفحه‌های جداساز، لزوماً موازی نیستند (مثلاً h_1 و h_2).

$$\begin{aligned} h_1: x^T w_1 + b_1 &= 0 \\ h_2: x^T w_2 + b_2 &= 0 \end{aligned} \quad (9)$$

که در رابطه‌های (۹)، w_1 و w_2 نمایشگر مختصات ابر صفحه و b_1 و b_2 بایاس می‌باشند. h_1 تا حدّ امکان به کلاس +۱ نزدیک و از کلاس -۱ دور است و h_2 تا حدّ ممکن به کلاس -۱ نزدیک و از کلاس +۱ دور می‌باشد. داده‌های با برچسب منفی باید فاصله‌ای مطلوب و حداقلی (مثلاً یک واحد) از h_1 داشته باشند در غیر

نیز کم اهمیت‌تر خواهد بود. چنین داده‌ای می‌تواند تخطی بیشتر از حدی که در رابطه (۶) مشخص شده، داشته باشد. حتی می‌توان داده‌های پرت را با یک اشتباه آگاهانه در دسته مقابل قرار داد. اما برای داده‌های با اهمیت بالا میزان تخطی از حدود کم و ناچیز می‌باشد و جواب مسئله باید در فضای همان قید جستجو شود. با این کار فضای شدنی مسئله گسترش می‌یابد و جواب‌های بهینه می‌توانند از فضای وسیع‌تری انتخاب شوند. روش جداسازی جدید و مسئله بهینه‌سازی نظیر آن به صورت زیر بازنویسی می‌شود:

$$\begin{aligned} \text{Min} \quad & \frac{1}{2} \|w\|^2 + C \sum_{i=1}^m \xi_i \\ & (w, b) \\ \text{S.t.} \quad & y_i(w^T x_i + b) \geq 1 - \xi_i \\ & \xi_i \geq 0, i = 1, 2, \dots, m. \end{aligned} \quad (8)$$

این ساختار، ماشین بردار پشتیبان با قیود نرم نامیده شده است. لازم به ذکر است که در یک مسئله بهینه‌سازی، اگر قیود به شکل فازی بیان گردند، به آن مسئله، برنامه‌ریزی فازی نامتقارن گفته می‌شود. برای حل این گونه مسئله‌ها روش‌هایی پیشنهاد شده است [۱۸]. اولین گام معرفی یک تابع عضویت برای تبدیل کردن روابط فازی به رابطه‌های غیرفازی است. برای توضیحات بیشتر به مرجع [۱۹] مراجعه کنید.

ماشین بردار پشتیبان دوقلو^۹

با گسترش ایده SVM در سال ۲۰۰۷ توسط جایادوا [۲۰] و همکاران ایده یافتن دو ابرصفحه غیرموازی، مطرح شد که هر یک از آن‌ها تا حد ممکن به داده‌های یک طبقه نزدیک باشند و از طبقه مقابل فاصله مطلوبی بگیرند. این الگوریتم ماشین بردار پشتیبان دوقلو نامیده می‌گردد. شکل (۵) نمایش هندسی ابرصفحه جداساز در فضای دوبعدی است. در این مسایل به جای حل یک مسئله بهینه‌سازی درجه دوم با ابعاد بالا، دو مسئله درجه دوم جدید با ابعاد

این صورت خطای دسته‌بندی ایجاد خواهد شد و لازم است که پارامترهای مشخص‌کننده جداساز طوری تنظیم شوند که داده‌هایی از این دست کمتر رخ دهند. به طور مشابه در مورد داده‌های دسته دوم هم چنین وضعیتی رخ خواهد داد. از این‌رو دو مسئله بهینه‌سازی مطرح خواهد شد و هر یک از این دو مسئله نیاز به کمینه کردن دو مقدار متفاوت دارند؛ کم کردن فاصله ابرصفحه جداساز h_1 از کلاس ۱+ و کم کردن تعداد نمونه‌هایی از کلاس ۱- که به h_1 نزدیک شده‌اند. تمام مباحث گفته شده در مورد جداساز h_2 هم صدق خواهند کرد. دو مسئله بهینه‌سازی به صورت زیر تعریف می‌شوند:

$$\begin{aligned} h_1: K(x^T, C^T)w_1 + b_1 &= \cdot \\ h_2: K(x^T, C^T)w_2 + b_2 &= \cdot \end{aligned} \quad (16)$$

توضیحات بیشتر در مورد داده‌هایی که به‌طور خطی جداناپذیرند، در مرجع [۲۳] ذکر شده‌است.

ماشین بردار پشتیبان دو قلو با قیود نرم^{۱۰}

داده‌ها در یافته‌های بالینی مربوط به بیماری‌های مزمن دارای نوفه ذاتی هستند. برای پیش‌بینی دقیق‌تر بیماری، ناچار از کاهش دادن تأثیر چنین داده‌هایی هستیم. در این بررسی برای کاهش تأثیر نوفه، با ایده گرفتن از RSVM و TWSVM، یک نوع ماشین بردار پشتیبان دو قلو با محدودیت‌های نرم معرفی خواهد شد.

ماشین بردار پشتیبان دو قلو دارای دو مسئله بهینه‌سازی درجه دوم است که در آن قیود حالت بدون انعطاف دارند. اشتراک مناطق محدود به قیدهای هر مسئله، فضای شدنی برای هر یک از دو مسئله بهینه‌سازی را تشکیل می‌دهد. جواب‌های بهینه در این فضا قرار دارند. نوآوری روش پیشنهادی این است که بدون هرگونه تغییر در تابع هدف، فضای شدنی مسئله را گسترش می‌دهد. استفاده از قیود نرم فازی به جای محدودیت‌های عادی، فضای شدنی را وسیع‌تر می‌کند؛ الگوریتم پیشنهادی در فضای شدنی گسترده‌تر، جواب‌های بهینه را می‌یابد. شیوه نوین در مقایسه با روش‌های مشابه عملکرد بهتری داشته است.

این صورت خطای دسته‌بندی ایجاد خواهد شد و لازم است که پارامترهای مشخص‌کننده جداساز طوری تنظیم شوند که داده‌هایی از این دست کمتر رخ دهند. به طور مشابه در مورد داده‌های دسته دوم هم چنین وضعیتی رخ خواهد داد. از این‌رو دو مسئله بهینه‌سازی مطرح خواهد شد و هر یک از این دو مسئله نیاز به کمینه کردن دو مقدار متفاوت دارند؛ کم کردن فاصله ابرصفحه جداساز h_1 از کلاس ۱+ و کم کردن تعداد نمونه‌هایی از کلاس ۱- که به h_1 نزدیک شده‌اند. تمام مباحث گفته شده در مورد جداساز h_2 هم صدق خواهند کرد. دو مسئله بهینه‌سازی به صورت زیر تعریف می‌شوند:

$$\begin{aligned} \text{Min } \frac{1}{\tau} \|Aw_1 + e_1 b_1\|^\tau + C_1 e_1^T q_1 \\ (w_1, b_1) \end{aligned} \quad (10)$$

$$\text{S. t. } -(Bw_1 + e_2 b_1) + q_1 \geq e_2, q_1 \geq \cdot$$

$$\begin{aligned} \text{Min } \frac{1}{\tau} \|Bw_2 + e_2 b_2\|^\tau + C_2 e_2^T q_2 \\ (w_2, b_2) \end{aligned} \quad (11)$$

$$\text{S. t. } (Aw_2 + e_1 b_2) + q_2 \geq e_1, q_2 \geq \cdot$$

که در آن $C_1, C_2, q_1 \in \mathbb{R}^{m_1}, q_2 \in \mathbb{R}^{m_2}$ پارامترهای پنالتی با مقادیر مثبت و سایر متغیرها قبلاً تعریف شده‌اند. برای حل مسئله‌های فوق، بعد از نوشتن دوگان لاگرانژ و

$$\begin{aligned} \text{شرایط مکمل زاید به دو مسئله جدید زیر می‌رسیم:} \\ \text{Min } \frac{1}{\tau} \alpha^T G(H^T H + \epsilon I)^{-1} G^T \alpha - e_1^T \alpha \end{aligned} \quad (12)$$

$$\text{S. t. } \cdot \leq \alpha \leq C_1 e_1$$

$$\begin{aligned} \text{Min } \frac{1}{\tau} \beta^T H(G^T G + \epsilon I)^{-1} H^T \beta - e_2^T \beta \\ \text{S. t. } \cdot \leq \beta \leq C_2 e_2 \end{aligned} \quad (13)$$

پس از حل مسائل و یافتن بردارهای α و β ابر صفحه‌های دو کلاس از روابط زیر به دست می‌آیند:

$$\begin{bmatrix} w_1 \\ b_1 \end{bmatrix} = -(H^T H + \epsilon I)^{-1} G^T \alpha \quad (14)$$

$$\begin{bmatrix} w_2 \\ b_2 \end{bmatrix} = (G^T G + \epsilon I)^{-1} H^T \beta \quad (15)$$

در رابطه (۱۵) ماتریس‌های H و G به‌صورت

طبقه‌بند خطی

در ماشین بردار پشتیبان دو قلو طبق روابط (۱۰ و ۱۱)، برای هر یک از نمونه‌های آزمایش موجود در هر کلاس، محدودیتی مخصوص آن کلاس وجود دارد. اشتراک محدودیت‌ها دو فضای شدنی را می‌سازد. جواب‌های بهینه هر مسئله یعنی $(w_1^*, b_1^*, q_1^*) \in R^{m_1+1+n}$ و $(w_2^*, b_2^*, q_2^*) \in R^{m_2+1+n}$ که مشخص‌کننده ابرصفحات بهینه جداساز هستند، از این دو فضا انتخاب می‌شوند.

روش پیشنهادی جدید با ثابت نگه داشتن تابع هدف و ایجاد تغییر در محدودیت‌ها، به دنبال گسترش دادن فضای شدنی است. برای ساختن فضای شدنی بیشینه با اندازه دلخواه، در محدودیت‌ها از نامساوی‌های فازی استفاده شده است. الگوریتم پیشنهادی برای هر یک از داده‌های آزمایش در دو کلاس، میزانی از تحمل و اطمینان در نظر می‌گیرد تا نمونه‌ها اجازه گذر از محدودیت‌های خود را داشته باشند. هر قدر که میزان تحمل الگوریتم برای گذشتن نمونه از محدودیت بیشتر باشد، محدودیت نظیر آن نرم‌تر خواهد بود. داده‌های کم اهمیت تر اجازه تخطی بیشتری دارند اما داده‌های با اهمیت بالاتر، صرفاً در منطقه محدودیت خود قرار می‌گیرند. داده‌های با اهمیت کمتر می‌توانند فاصله‌ای کمتر از واحد با ابرصفحه مربوط به کلاس مقابل داشته باشند. حتی داده‌های نوفه‌ای می‌توانند با چشم‌پوشی آگاهانه در طبقه مقابل دسته‌بندی شوند. مسئله‌های بهینه‌سازی جدید به صورت زیر هستند:

$$\text{Min } \frac{1}{\tau} \|Aw_1 + e_1 b_1\|^\tau + C_1 e_1^T q_1 \quad (17)$$

$$(w_1, b_1)$$

$$S.t. - (Bw_1 + e_2 b_1) + q_1 \geq e_2, q_1 \geq 0$$

$$\text{Min } \frac{1}{\tau} \|Bw_2 + e_2 b_2\|^\tau + C_2 e_2^T q_2 \quad (18)$$

$$(w_2, b_2)$$

$$S.t. (Aw_2 + e_1 b_2) + q_2 \geq e_1, q_2 \geq 0$$

علامت $\geq$ به این معنی است که برای نمونه‌های یادگیری اجازه تخطی از حدود هم داده می‌شود. استفاده

از نامساوی فازی به این محدودیت‌ها انعطاف بیشتری می‌بخشد. در واقع با این کار فضای شدنی مسئله گسترش می‌یابد و جواب‌های بهینه می‌توانند از فضای بزرگتری انتخاب شوند. ابتدا می‌بایست دو مسئله بهینه‌سازی (۱۷ و ۱۸) که مسئله‌های بهینه‌سازی غیرخطی درجه دوم با تابع هدف محدب و قیود نامساوی فازی هستند، حل شوند. در گام اول برای حل مسئله، محدودیت‌ها را از فرم ماتریسی خارج کرده و برای هر نمونه از هر طبقه منفی، یک محدودیت به صورت (۱۹) در نظر گرفته می‌شود.

$$-(w_1^T x_i + b_1) \geq 1 - q_{1i}, q_{1i} \geq 0 \quad (19)$$

$$i = 1, \dots, m_2.$$

که x_i ، i - امین نمونه از داده‌های با برچسب ۱- و q_{1i} خطای مربوط به این نمونه در اثر تخطی از محدودیت نظیر و نزدیک شدن به h_1 می‌باشد. به‌طور مشابه برای نمونه i - ام از طبقه مثبت، محدودیت به صورت (۲۰) خواهد بود:

$$(w_2^T x_i + b_2) \geq 1 - q_{2i}, q_{2i} \geq 0 \quad (20)$$

$$i = 1, \dots, m_1.$$

که x_i ، i - امین نمونه از داده‌های با برچسب مثبت و q_{2i} خطای مربوط به این نمونه در اثر تخطی از محدودیت نظیر و نزدیک شدن به h_2 می‌باشد. برای نمونه‌های با برچسب مثبت، طرفین محدودیت نظیر در (+۱) و برای نمونه‌های با برچسب منفی، طرفین محدودیت نظیر در (-۱) ضرب شده و در نهایت محدودیت‌ها به صورت (۲۱) ادغام خواهند شد:

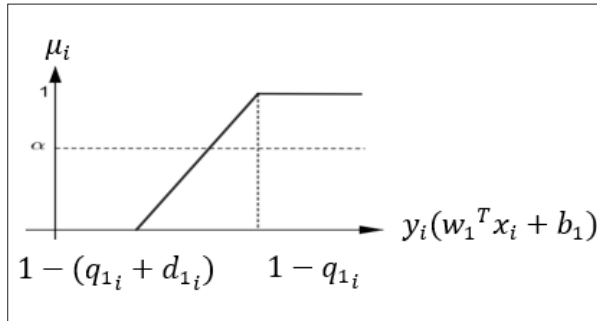
$$y_i(w_1^T x_i + b_1) \geq 1 - q_{1i} \quad (21)$$

$$q_{1i} \geq 0, i = 1, \dots, m_2.$$

$$q_{2i} \geq 0, i = 1, \dots, m_1.$$

و

طبق رابطه (۲۱) هر یک از نمونه‌های آزمایش، یک محدودیت به مسئله تحمیل می‌کنند. مجموعه این محدودیت‌ها فضای شدنی را می‌سازد. متغیرهای مسئله از این فضا انتخاب می‌شوند. هنگام حل مسئله این شرایط

شکل (۶): تابع عضویت μ_i

باید برآورده شوند. یک روش برای حل (۲۱) تبدیل مسئله به حالت غیرفازی است. بنابراین گام دوم تعیین توابع عضویت μ_i, μ'_i برای نامساوی‌های فازی موجود در قیود مسئله است. در روش پیشنهادی هر دو تابع به صورت خطی تعریف شده‌اند. شکل (۶) تابع عضویت برای (۱۷) را نشان می‌دهد و فرمول‌های (۲۲ و ۲۳) بیان‌کننده توابع عضویت μ_i و μ'_i می‌باشند:

که در آن:

$$\mu_i(w_1, b_1, q_1) = \begin{cases} 1, & \text{if } y_i(w_1^T x_i + b_1) \geq 1 - q_{1i} \\ \frac{y_i(w_1^T x_i + b_1) - 1 + q_{1i} + d_{1i}}{d_{1i}}, & \text{if } 1 - (q_{1i} + d_{1i}) \leq y_i(w_1^T x_i + b_1) \leq 1 - q_{1i} \\ 0, & \text{if } y_i(w_1^T x_i + b_1) \leq 1 - (q_{1i} + d_{1i}) \end{cases} \quad (22)$$

و

$$\mu'_i: R^{m_1+1+n} \rightarrow (0, 1], i = 1, 2, \dots, m_1.$$

که در آن

$$\mu'_i(w_\tau, b_\tau, q_\tau) = \begin{cases} 1, & \text{if } y_i(w_\tau^T x_i + b_\tau) \geq 1 - q_{\tau i} \\ \frac{y_i(w_\tau^T x_i + b_\tau) - 1 + q_{\tau i} + d_{\tau i}}{d_{\tau i}}, & \text{if } 1 - (q_{\tau i} + d_{\tau i}) \leq y_i(w_\tau^T x_i + b_\tau) \leq 1 - q_{\tau i} \\ 0, & \text{if } y_i(w_\tau^T x_i + b_\tau) \leq 1 - (q_{\tau i} + d_{\tau i}) \end{cases} \quad (23)$$

اول نزدیک شده و نتوانسته است محدودیت مسئله را برآورده کند پس h_1 قادر به دسته بندی درست چنین نمونه‌ای نیست و درجه‌ی تعلق فازی برای h_1 صفر خواهد بود.

و معنی نامساوی

$$1 - (q_{1i} + d_{1i}) \leq y_i(w_1^T x_i + b_1) \leq 1 - q_{1i}$$

آن است که اگر برای نمونه i -ام بیشینه تحملی به اندازه d_{1i} در نظر بگیریم تا بتواند به ابرصفحه کلاس مقابل نزدیک شود، آنگاه رابطه میان مقدار تحمل و درجه عضویت به صورت خطی است و با نزدیک شدن به آستانه تحمل درجه تعلق فازی هم به عدد ۱ نزدیک می‌شود پس h_1 با درجه تعلق μ_i می‌تواند یک ابرصفحه جداساز برای نمونه مزبور باشد.

برای هر یک از اولین محدودیت‌های رابطه (۲۱) تعریف می‌کنیم:

از این پس در مورد اولین محدودیت در رابطه (۲۱) و با تمرکز بر داده‌های کلاس منفی و با در نظر گرفتن تابع عضویت μ_i محاسبات مربوط به این طبقه به تفصیل بررسی خواهد شد. μ'_i نیز مشابه خواهد بود.

لازم به توضیح است که اگر q_{1i} متغیر لغزش و d_{1i} مقدار تحمل خطا برای نمونه آموزش i -ام و $h_1: x^T w_1 + b_1 = 0$ ابرصفحه جداساز نظیر باشد سه حالت زیر رخ خواهد داد:

اگر $y_i(w_1^T x_i + b_1) \geq 1 - q_{1i}$ یعنی فاصله حداقلی $1 - q_{1i}$ بین نمونه i -ام و ابرصفحه اول حفظ شده، این نمونه توسط h_1 به درستی دسته‌بندی شده است و مقدار تعلق فازی برای h_1 برابر ۱ خواهد بود.

نامساوی $y_i(w_1^T x_i + b_1) \leq 1 - (q_{1i} + d_{1i})$ بیان می‌کند که تحمل مجاز d_{1i} برای نمونه i -ام در نظر گرفته شده است، ولی این نمونه بیش از تحمل مجاز به ابرصفحه

$$\frac{y_i(w_1^T x_i + b_1) - 1 + q_{1i} + d_{1i}}{d_{1i}} \geq \alpha \Rightarrow \quad (31)$$

$$y_i(w_1^T x_i + b_1) \geq 1 - q_{1i} - d_{1i}(1 - \alpha)$$

پارامترهای α و d_{1i} توسط کاربر تعیین می‌شوند، پارامتر α ، که عامل اطمینان نامیده می‌شود، سطحی را مشخص می‌کند که برای سطوح پایین‌تر از آن میزان تعلق فازی نظیر نامساوی موجود در قیود (۲۱)، صفر است. هر چه این مقدار به عدد ۱ نزدیک‌تر باشد، قیود مسئله به TWSVM نزدیک‌تر می‌شود. d_{1i} ها (که در اینجا مقدارشان مساوی d_1 است)، مقدار تحملی هستند که می‌توان به نمونه‌ها نسبت داد. هرچه مقدار تحملی که نسبت می‌دهیم بیشتر باشد، نمونه‌های کلاس منفی، می‌توانند به ابرصفحه جداساز کلاس مثبت، نزدیک‌تر شوند. این بدان معنی است که تعداد بردارهای پشتیبان بیشتر شده و داده‌های بیشتری در تعیین ابرصفحه بهینه نقش خواهند داشت. در نهایت مدل RTWSVM1 عبارت است از:

$$\begin{aligned} \text{Min } & \frac{1}{\gamma} \|Aw_1 + e_1 b_1\|^2 + C_1 e_1^T q_1 \quad (32) \\ \text{s.t. } & -(Bw_1 + e_2 b_1) \geq e_2 - q_1 - d_1(1 - \alpha) \\ & q_{1i} \geq 0, \quad i = 1, \dots, m_2. \end{aligned}$$

گام بعدی حل مسئله بهینه‌سازی غیرخطی درجه دوم با قیدهای غیرفازی^{۱۱} است. تابع هدف اولیه همراه با محدودیت‌هایش، به تابع لاگرانژ تبدیل می‌شوند. بردارهای ضرایب نامنفی لاگرانژ در محدودیت‌ها ضرب شده و به تابع هدف افزوده می‌شوند. دوگان مسئله به صورت (۳۳) بازنویسی خواهد شد:

$$\begin{aligned} Q(w_1, b_1, q_1, \beta_1, \gamma_1) = & \frac{1}{\gamma} \|Aw_1 + e_1 b_1\|^2 + \quad (33) \\ & C_1 e_1^T q_1 - \beta_1^T \{-(Bw_1 + e_2 b_1) + q_1 + \\ & d_1(1 - \alpha) - e_2\} - \gamma_1^T q_1 \end{aligned}$$

که در آن $\beta_1 = (\beta_{11}, \dots, \beta_{1m_2})^T, \gamma_1 = (\gamma_{11}, \dots, \gamma_{1m_2})^T$ ضرایب لاگرانژ نامنفی می‌باشند.

پس از مشتق‌گیری نسبت به متغیرها و یافتن ریشه‌ها خواهیم داشت:

$$\begin{aligned} X_i = \{ & (w_1, b_1, q_1) \in R^{m_2+1+n} \mid y_i(w_1^T x_i + b_1) \\ & \geq 1 - q_{1i}, q_{1i} \geq 0, \quad i = 1, \dots, m_2 \} \end{aligned} \quad (24)$$

آن‌گاه جواب‌های مسئله باید از مجموعه X با تعریف زیر انتخاب شوند:

$$X = \bigcap_{i \in I} X_i, \quad I = \{1, 2, \dots, m_2\} \quad (25)$$

بنابراین مسئله (۱۷) را می‌توان به صورت زیر نوشت:

$$\begin{aligned} \text{Minimize } & Q(w_1, b_1, q_1) = \quad (26) \\ & C_1 e_1^T q_1 \mid (w_1, b_1, q_1) \in X \} + \frac{1}{\gamma} \|Aw_1 + e_1 b_1\|^2 \end{aligned}$$

برای حل (۲۶) از روش آلفا برش استفاده می‌شود. این مسئله در شکل (۶) دیده می‌شود. جواب‌های (۲۶) باید از مجموعه

$$X_\alpha = \{(w_1, b_1, q_1) \in R^{m_2+1+n} \mid \mu_X(w_1, b_1, q_1) \geq \alpha\} \quad (27)$$

انتخاب شوند که:

$$\mu_X(x) = \inf \{\mu_i(x), i \in I\} \quad (28)$$

مجموعه X_α آلفا برش قید i ام است و $\alpha \in (0, 1]$ جواب بهینه مسئله (۱۷) با α داده شده برابر است با:

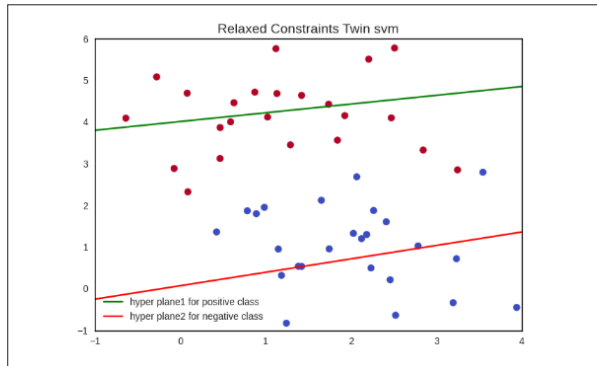
$$\begin{aligned} S(\alpha) = \{ & (w_1, b_1, q_1) \in R^{m_2+1+n} \mid \quad (29) \\ & + C_1 e_1^T q_1 \} = \frac{1}{\gamma} \|Aw_1 + e_1 b_1\|^2 \\ & \{ \text{Min } \frac{1}{\gamma} \|Aw'_1 + e_1 b'_1\|^2 + \\ & C_1 e_1^T q'_1, (w'_1, b'_1, q'_1) \in X_\alpha \} \end{aligned}$$

پس می‌توان نوشت:

$$\begin{aligned} X_\alpha = \bigcap_{i \in I} \{ & (w_1, b_1, q_1) \in R^{m_2+1+n} \mid \quad (30) \\ & y_i(w_1^T x_i + b_1) \geq r_{1i}(\alpha), q_{1i} \geq 0, \\ & i = 1, \dots, m_2 \} \end{aligned}$$

که در آن $r_{1i}(\alpha) = 1 - q_{1i} - d_{1i}(1 - \alpha)$ زیرا:

$$d_j(x) = \frac{|x^T w_j + b_j|}{\|w_j\|} \quad (40)$$



شکل (۷): نمایش هندسی بردار پشتیبان دوقلو با قیود نرم.

طبقه‌بند غیرخطی

مسایل یادگیری ماشین در دنیای واقعی غالباً به صورت خطی جدپذیر نیستند. روش پیشنهادی RT-WSVM برای جداسازی مسایل غیرخطی، نمونه‌ها را به فضای ویژگی با ابعاد بالاتر نگاشت می‌کند. بدین‌منظور دو ابرسطح^{۱۲} به صورت روابط (۱۶) تعریف می‌شوند، که در آن $C = [A \ B]$ و K تابع کرنل دلخواه می‌باشد. در نسخه غیرخطی، مسایل بهینه‌سازی اصلی برای کلاس مثبت و منفی به ترتیب در روابط (۴۱) و (۴۲) تعریف شده است.

$$\text{Min } \frac{1}{\tau} \|K(A, C^T)w_1 + e_1 b_1\|^2 + C_1 e_1^T q_1 \quad (41)$$

$$S. t. \quad - (K(B, C^T)w_1 + e_1 b_1) + q_1 \geq e_1, \quad q_1 \geq 0$$

$$\text{Min } \frac{1}{\tau} \|K(B, C^T)w_2 + e_2 b_2\|^2 + C_2 e_2^T q_2 \quad (42)$$

$$S. t. \quad (K(A, C^T)w_2 + e_2 b_2) + q_2 \geq e_2, \quad q_2 \geq 0.$$

برای مسئله (۴۱)، مشابه حالت خطی، تابع عضویت فازی μ_i تعریف شده و به کمک روش آلفا برش، نامساوی‌های فازی به غیرفازی تبدیل می‌شوند.

رابطه (۴۳) مسئله را در حالت غیرخطی و با قیود غیر فازی نشان می‌دهد:

$$\frac{\partial Q(w_1, b_1, q_1, \beta_1, \gamma_1)}{\partial w_1} = 0 \Rightarrow \quad (34)$$

$$A^T (Aw_1 + e_1 b_1) + B^T \beta_1 = 0$$

$$\frac{\partial Q(w_1, b_1, q_1, \beta_1, \gamma_1)}{\partial b_1} = 0 \Rightarrow \quad (35)$$

$$e_1^T (Aw_1 + e_1 b_1) + e_1^T \beta_1 = 0$$

$$\frac{\partial Q(w_1, b_1, q_1, \beta_1, \gamma_1)}{\partial q_1} = 0 \Rightarrow \quad (36)$$

$$C_1 e_1^T - \beta_1^T - \gamma_1^T = 0$$

به‌علاوه باید شرایط مکمل زاید نیز برقرار شوند:

$$\beta_1^T \{-(Bw_1 + e_2 b_1) + q_1 + d_1(1 - \gamma_1^T q_1) - e_2\} = 0 \quad (37)$$

از ترکیب روابط (۳۴ و ۳۵) خواهیم داشت:

$$U_1 = -(H^T H + \epsilon I)^{-1} G^T \beta_1 \quad (38)$$

$$U_1 = [w_1 \ b_1]^T H = [A \ e_1], \quad G = [B \ e_2]. \quad \text{که در آن:}$$

با جای‌گذاری (۳۸) در (۳۳) و ساده کردن روابط و استفاده از رابطه (۳۶)، مسئله نهایی به صورت زیر خواهد بود:

$$\begin{aligned} \text{Min } & \frac{1}{\tau} \beta_1^T G (H^T H + \epsilon I)^{-1} G^T \beta_1 \\ & - \beta_1^T \{d_1(1 - \alpha) - e_2\} \\ S. t. & \quad 0 \leq \beta_1 \leq C_1 e_2 \end{aligned} \quad (39)$$

مسئله مشابه تمام روابط (۲۴ تا ۳۹) برای داده‌های کلاس دیگر نیز برقرار است. با حل (۳۹) ضرایب لاگرانژ به دست می‌آیند. دو ابرصفحه غیرموازی، مشخص می‌شوند. فاصله عمودی نمونه‌های جدید از هر ابر صفحه محاسبه می‌گردد. کمترین فاصله، برچسب داده آزمایش را مشخص می‌کند:

$$\text{class}(x) = \arg \text{Min}_{j=1,2} (d_j(x)),$$

که در آن $d_j(x)$ از دستور (۴۰) محاسبه می‌شود:

جدول (۱): مشخصات مجموعه داده برای ارزیابی روش RTWSVM

مجموعه داده	تعداد نمونه‌ها و ویژگی	تعداد نمونه‌های مثبت	تعداد نمونه‌های منفی	نسبت منفی/مثبت	داده از دست‌رفته
CKD	۲۴×۴۰۰	۲۵۰	۱۵۰	۱/۷	ندارد
Hep	۱۹×۱۵۵	۱۲۳	۳۲	۳/۸	دارد
Heart	۱۳×۲۷۰	۱۲۰	۱۵۰	۰/۸۰	ندارد
PIDD	۸×۷۶۸	۲۶۸	۵۰۰	۰/۵۴	ندارد
WDBC	۳۰×۵۶۹	۲۱۲	۳۵۷	۰/۵۹	دارد

پنج مجموعه داده CKD، Hep، Heart-statlog، PIDD و WDBC که به ترتیب مرتبط با بیماری کلیه، هپاتیت، بیماری قلبی، دیابت و سرطان سینه هستند، مورد بررسی قرار گرفته‌اند. داده‌های توصیفی^{۱۳} از دست‌رفته با اولین مُد در ستون مزبور و داده‌های عددی با صفر جایگزین شده‌اند. عدم تعادل بین نمونه‌ها در مجموعه داده‌های CKD و Hep مشهود است. جدول (۱) مجموعه داده را توصیف می‌کند.

نحوه پیاده‌سازی و اجرای الگوریتم‌ها

تمام روش‌ها در زبان برنامه‌نویسی پایتون و در محیط رایگان برخط google colab پیاده‌سازی شده‌اند. از کتابخانه cvxopt برای بهینه‌سازی و از کتابخانه‌های numpy، pandas و matplotlib جهت محاسبات ریاضی، ساختن چارچوبی برای داده‌ها و رسم نمودارها استفاده شده است.

از آن‌جا که در دنیای واقعی داده‌ها جدایی‌پذیر خطی نیستند، تابع هسته گوسی که قدرت تعمیم‌پذیری بهتری دارد به کار رفته است:

$$K(x, x') = \exp\left(-\frac{\|x - x'\|^2}{2\gamma^2}\right) \quad (۴۷)$$

دقت روش RTWSVM به انتخاب پارامترهای بهینه بسیار وابسته است. بدین‌منظور از روش جستجوی شبکه‌ای برای پیدا کردن پارامتر بهینه استفاده شده است. پارامترهای پناستی c_1 و c_2 از مجموعه $\{2^i | i = -7, -6, \dots, 7\}$ و پارامترهای α_1 ، α_2 و d_1 که به ترتیب، برای سطح اطمینان برش آلفا و ضریب تحمل الگوریتم، مورد استفاده بوده‌اند از مجموعه $[0.95, 0.75, 0.25]$ و پارامتر

$$\text{Min } \frac{1}{\gamma} \|K(A, C^T)w_1 + e_1 b_1\|^2 + C_1 e_1^T q_1 \quad (۴۳)$$

$$S. t. \quad -(K(B, C^T)w_1 + e_1 b_1) \geq e_2 - q_1 - d_1(1 - \alpha) \\ q_{1i} \geq 0, \quad i = 1, \dots, m_2.$$

مسئله دوگان (۴۳) به صورت (۴۴) خواهد بود:

$$Q(w_1, b_1, q_1, \beta_1, \gamma_1) = \frac{1}{\gamma} \|K(A, C^T)w_1 + e_1 b_1\|^2 + C_1 e_1^T q_1 - \beta_1^T \{-(K(B, C^T)w_1 + e_1 b_1) + q_1 + d_1(1 - \alpha) - e_2\} - \gamma_1^T q_1 \quad (۴۴)$$

پس از مشتق‌گیری و برقراری شرایط مکمل زاید:

$$U_1 = -(S^T S + \epsilon I)^{-1} R^T \beta_1 \quad (۴۵)$$

$$S = [K(A, C^T) e_1] \quad R = [K(B, C^T) e_2]$$

$$U_1 = [w_1 \ b_1]^T$$

و مشابه نسخه خطی کلاس نمونه جدید x با محاسبه فاصله عمودی آن از دو ابرسطح مشخص می‌شود، به‌طوری‌که تابع تصمیم برای نسخه غیرخطی به‌صورت زیر تعریف می‌گردد:

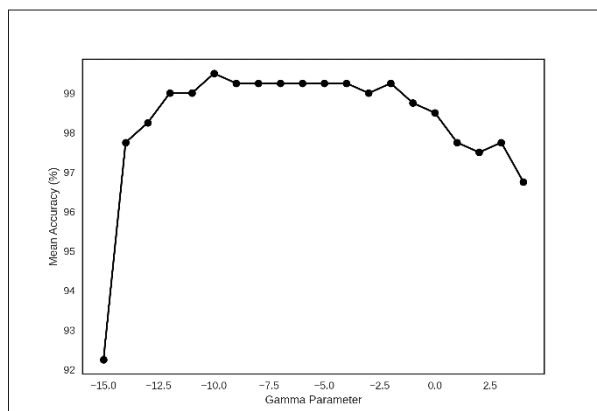
$$\text{class}(x) = \arg \min_{j=1,2} (d_j(x)),$$

که در آن $d_j(x)$ از دستور (۴۶) محاسبه می‌شود:

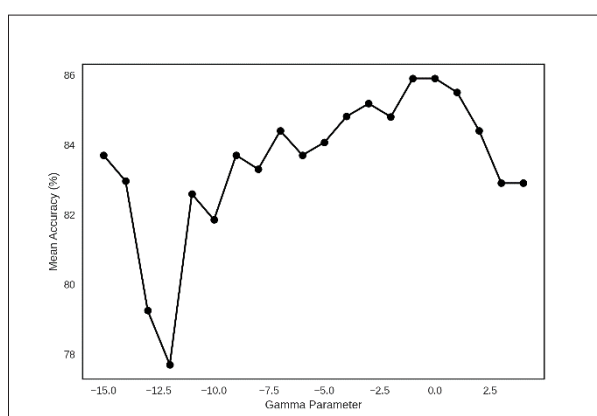
$$d_j(x) = \frac{|K(x, C^T)w_j + b_j|}{\|w_j\|} \quad (۴۶)$$

نتایج و ارزیابی پایگاه داده

داده‌های مورد استفاده در این تحقیق، از پایگاه‌های داده UCI و Kaggle استخراج شده‌اند. به دلیل اهمیت موضوع

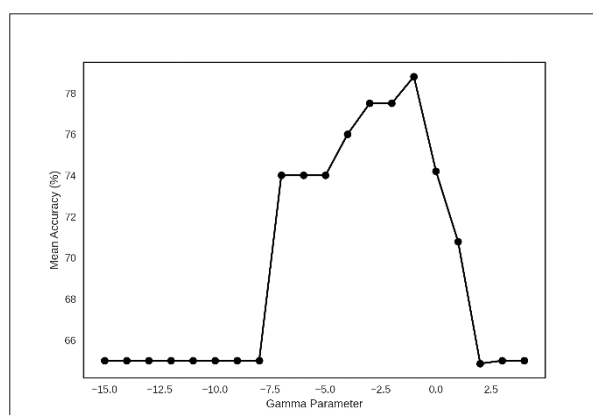


شکل (۱۰): تاثیر پارامتر گاما بر دقت الگوریتم در پایگاه داده WDBC



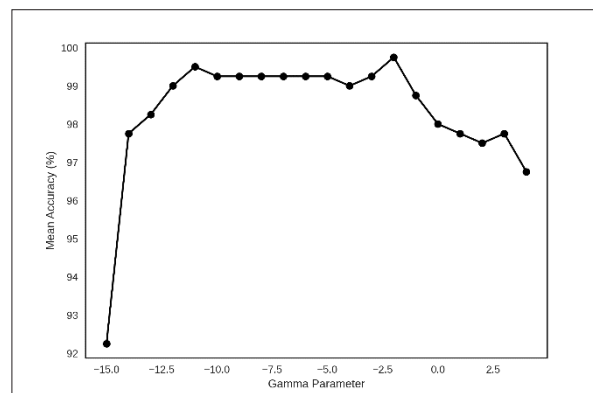
شکل (۱۱): تاثیر پارامتر گاما بر دقت الگوریتم در پایگاه داده Heart-statlog

تابع هسته γ نیز از $\{2^i | i = -15, -14, \dots, 5\}$ انتخاب شده‌اند. برای تعیین میزان دقت روش، اعتبارسنجی متقابل به کار رفته است؛ الگوریتم ارائه شده با ۹۰٪ داده‌ها ($k=10$) آموزش دیده و با ۱۰٪ باقیمانده آزمایش شده، این مراحل ده بار تکرار شده‌اند. پس از ۱۰ بار تکرار، میانگین و انحراف معیار دقت روش پیشنهادی به دست آمده‌اند. شکل‌های (۸-۱۱) بیانگر تأثیر تغییرات پارامتر گاما بر میزان دقت روش پیشنهادی در پایگاه‌های داده CKD، Heart-statlog، PIDD و WDBC می‌باشند.



شکل (۸): تاثیر پارامتر گاما بر دقت الگوریتم در پایگاه داده PIDD

به منظور بزرگنمایی نمودار رابطه بین گاما و دقت الگوریتم، نمودار دقت، بر اساس لگاریتم گاما در مبنای ۲ رسم شده است. بیشترین حساسیت نسبت به تغییرات گاما در مجموعه داده Heart-statlog قابل مشاهده است اما WDBC و CKD نسبت به تغییرات مقاوم‌تر بوده‌اند.



شکل (۹): تاثیر پارامتر گاما بر دقت الگوریتم در پایگاه داده CKD

نتایج تجربی و تجزیه و تحلیل آماری بر روی داده‌های دنیای واقعی

با هدف تحلیل عملکرد، روش پیشنهادی، همراه با SVM ساده، رگرسیون لجستیک، درخت تصمیم و نزدیک‌ترین همسایه، را بر روی پنج مجموعه داده‌های جدول (۱)، مقایسه شده‌اند. نتایج روش جدید هم با نتایج به دست آمده از مرجع [۲۴] و هم با پیاده‌سازی دوباره مدل‌های ذکر شده، بر روی داده‌های بیماری‌های مزمن، گزارش شده‌اند. به جز نتایج مجموعه داده هیپاتیت که دچار عدم تعادل محسوس می‌باشد، اعداد به دست آمده از پیاده‌سازی مجدد و جدول گزارش شده در مرجع [۲۴] تقریباً یکسان هستند که اعتبار پژوهش را تایید می‌کند. نتایج مقایسه دقت و رتبه الگوریتم پیشنهادی با روش‌های پیشین در جدول‌های (۲-۵) قابل

جدول (۲): مقایسه دقت دسته‌بندی‌های پیشین و روش پیشنهادی برگرفته از مرجع [۲۴]

دقت	SVM	LR	DT	KNN	RTWSVM
CKD	۹۸/۷۵	۹۸/۷۵	۹۷/۵۰	۹۵/۰۰	۹۹/۵۰
Hep	۸۰/۰۰	۸۳/۰۰	۸۳/۰۰	۸۰/۰۰	۱۰۰/۰۰
PIDD	۷۷/۹۲	۷۹/۸۷	۷۴/۶۸	۷۶/۶۲	۷۸/۸۱
WDBC	۹۷/۱۴	۹۷/۱۴	۹۵/۰۰	۹۷/۱۴	۹۸/۳۹
میانگین	۸۸/۲۲	۸۹/۶۹	۸۷/۵۴	۸۷/۱۹	۹۷/۱۴

جدول (۳): مقایسه رتبه دسته‌بندی‌های پیشین و روش پیشنهادی برگرفته از مرجع [۲۴]

رتبه	SVM	LR	DT	KNN	RTWSVM
CKD	۳	۲	۴	۵	۱
Hep	۴	۲	۲	۴	۱
PIDD	۳	۱	۵	۴	۲
WDBC	۳	۳	۵	۳	۱
میانگین	۳/۳۷۵	۲/۱۲۵	۴/۱۲۵	۴/۱۲۵	۱/۲۵

جدول (۴): مقایسه دقت دسته‌بندی‌های پیشین و روش پیشنهادی

دقت	SVM	LR	DT	KNN	RTWSVM
CKD	۹۷/۰۰±۱/۸۷	۹۸/۰۰±۱/۵۰	۹۵/۲۵±۳/۰۵	۹۷/۲۵±۲/۰۷	۹۹/۵۰±۱/۰۰
Heart	۸۳/۳۳±۸/۹۵	۸۳/۷۰±۸/۶۳	۷۴/۴۴±۷/۴۹	۸۱/۱۱±۷/۴۹	۸۵/۹۲±۷/۵۵
Hep	۰±۱۰۰	۰±۱۰۰	۰±۱۰۰	۸۸/۴۱±۷/۷۱	۰±۱۰۰
PIDD	۷۷/۵۹±۴/۱۷	۷۷/۱۹±۴/۶۸	۶۹/۵۳±۵/۱۶	۷۳/۸۳±۳/۷۵	۷۸/۸۱±۵/۵۹
WDBC	۹۷/۵۴±۲/۹۵	۹۵/۷۸±۳/۴۳	۹۱/۷۵±۴/۳۶	۹۶/۶۶±۲/۴۱	۹۸/۳۹±۱/۴۸
میانگین	۹۱/۰۹۲	۹۰/۹۳۴	۸۶/۱۹۴	۸۷/۴۵۲	۹۲/۵۲۴

مشاهده‌اند؛ RTWSVM عملکرد بهتری دارد.

آزمون آماری ساده و ناپارامتریک فریدمن از توزیع χ^2 با $(k-1)$ درجه آزادی که در آن k تعداد مدل‌ها و N مجموعه داده مورد استفاده قرار گرفته است. به منظور انجام آزمون آماری، میانگین رتبه چهار روش بر اساس دقت محاسبه شده و در جدول (۲ و ۳) نشان داده شده است. ابتدا بر اساس فرضیه صفر، تمام روش‌ها را یکسان در نظر می‌گیریم و سپس آزمون فریدمن از رابطه (۴۸) محاسبه می‌شود.

$$\chi_F^2 = \frac{12N}{k(k+1)} \left[\sum_j R_j^2 - \frac{k(k+1)^2}{4} \right] \quad (48)$$

در رابطه (۴۸)، $R_j = \frac{1}{N} \sum_i r_i^j$ ، میانگین رتبه روش j ام و r_i^j نشان‌دهنده رتبه j امین روش بر روی i امین مجموعه داده از N مجموعه داده، است:

$$\chi_F^2 = \frac{12 \times 5}{5(5+1)} \left[2/7^2 + 7 + 4/5^2 + 3/8^2 + 1/3^2 - \frac{5(5+1)^2}{4} \right] = 11/92$$

سپس مقدار F_F بر اساس فرمول (۴۹) محاسبه می‌شود.

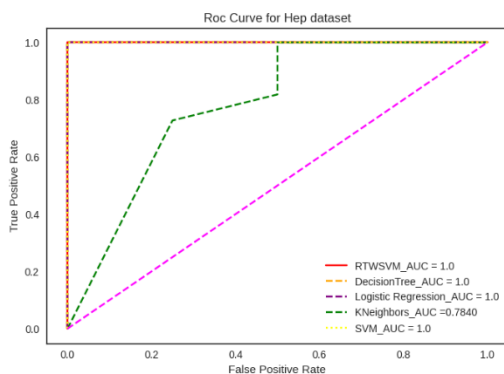
$$F_F = \frac{(N-1)\chi_F^2}{N(k-1) - \chi_F^2} = \frac{4 \times 11/92}{5 \times 4 - 11/92} = 5/9 \quad (49)$$

که در آن $k-1 = 4$ ، $k = 5$ ، $(k-1)(N-1) = 16$ و F_F از توزیع F با $(k-1)(N-1)$ درجه آزادی پیروی می‌کند. با در نظر گرفتن سطح اهمیت $\alpha = 0/05$ ، مقدار بحرانی $F(4,16) = 3/01$ به دست می‌آید و از آنجا که مقدار F_F از آن بیشتر است، فرضیه صفر رد می‌شود.

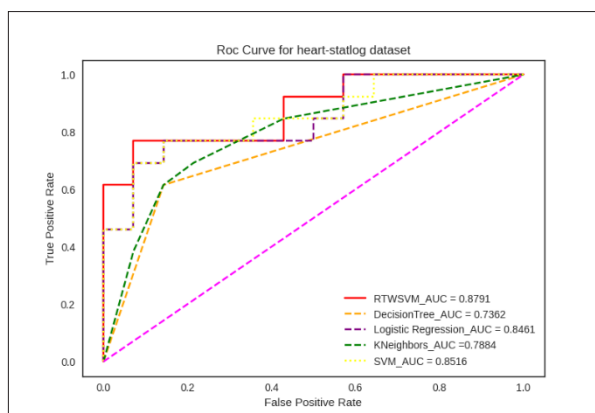
یکی از روش‌های بررسی و ارزیابی عملکرد دسته‌بندی دودویی، نمودار مشخصه عملکرد یا به اختصار منحنی ROC است. کارایی الگوریتم‌های دسته‌بندی دودویی معمولاً به وسیله شاخص‌هایی مثل حساسیت^{۱۴} و صراحت^{۱۵} سنجیده می‌شوند. اما در نمودار ROC هر دوی این شاخص‌ها ترکیب شده و به صورت یک منحنی نمایش داده می‌شوند. منطقه زیر منحنی^{۱۶} اندازه‌گیری، بیان‌گر توانایی طبقه‌بندی‌کننده برای تمایز بین کلاس‌ها است و به عنوان خلاصه منحنی ROC استفاده می‌شود. هر چه AUC بالاتر باشد، عملکرد مدل در تشخیص کلاس‌های مثبت و

14-Sensitivity
15-Specificity
16-Area Under the Curve

شکل (۱۴): نمودار مشخصه عملکرد بر مجموعه داده PIDD



شکل (۱۵): نمودار مشخصه عملکرد بر مجموعه داده Hep



شکل (۱۶): نمودار مشخصه عملکرد بر مجموعه داده Heart-statlog

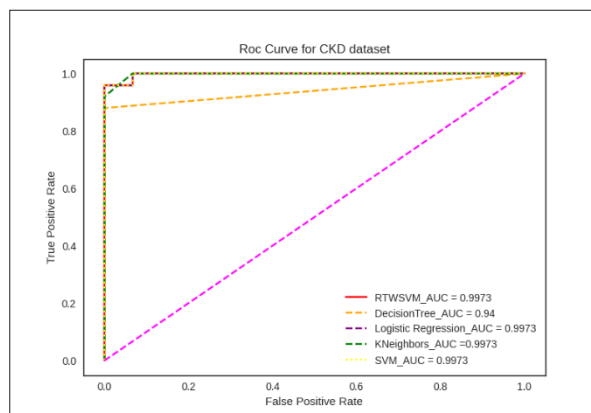
نتیجه‌گیری و پیشنهادهای آتی

گسترش فضای شدنی، عدم تأثیر منفی بر تابع هدف، کاهش اثر داده‌های پرت، سرعت بیشتر و کارایی بالاتر از ویژگی‌های مثبت روش جدید است. به‌کارگیری الگوریتم‌های جستجو مانند الگوریتم‌های تکاملی، می‌تواند یافتن مقدار بهینه برای پارامترهای زیاد موجود را سرعت بخشد. همچنین می‌توان با در نظر گرفتن یک ماتریس وزنی ضرایب تحمل، به هر یک از نمونه‌ها توجه ویژه داشت. با این کار تأثیر داده‌های نوفه کاهش خواهد یافت.

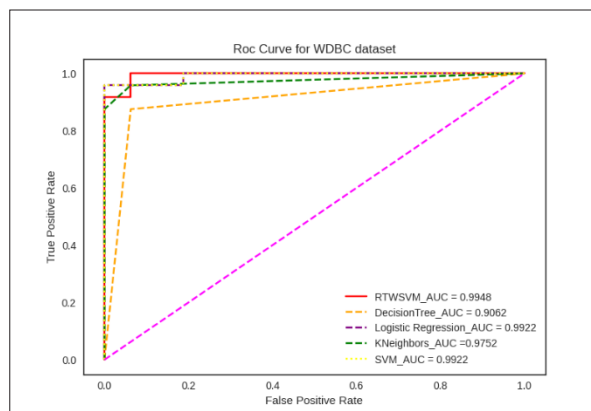
مراجع

[1] Souza-Pereira, L., Pombo, N., Ouhbi, S., Felizardo, V. and Garcia, N., 2020. «Clinical decision support systems for chronic diseases: A systematic literature review», Computer

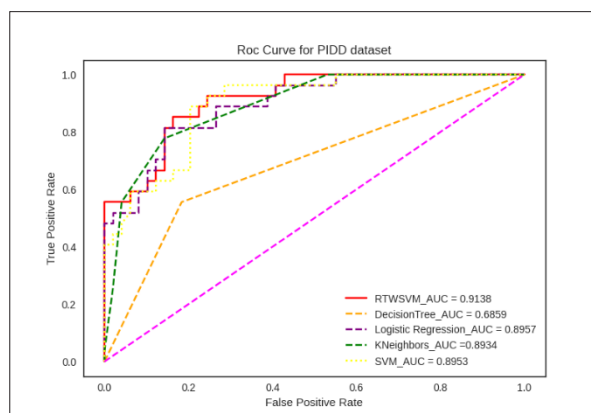
منفی بهتر است. وقتی $AUC = 1$ ، طبقه‌بندی‌کننده می‌تواند کاملاً بین تمام نقاط کلاس مثبت و منفی به درستی تمایز قائل شود و اگر $AUC = 0$ ، طبقه‌بندی‌کننده همه منفی‌ها را به‌عنوان مثبت و همه مثبت‌ها را به‌عنوان منفی پیش‌بینی خواهد کرد. عملکرد روش پیشنهادی در شکل‌های (۱۲)–(۱۶) در قالب نمودارهای ROC آمده است.



شکل (۱۲): نمودار مشخصه عملکرد بر مجموعه داده CKD



شکل (۱۳): نمودار مشخصه عملکرد بر مجموعه داده WDBC



- machine learning methods», *Translational oncology*, Vol. 14(1), pp. 100907.
- [15] Mezzatesta, S., Torino, C., De Meo, P., Fiumara, G. and Vilasi, A., 2019. "A machine learning-based approach for predicting the outbreak of cardiovascular diseases in patients on dialysis», *Computer methods and programs in biomedicine*, Vol. 177, pp. 9-15.
- [16] Yuan, X., Chen, S., Sun, C. and Yuwen, L., 2022. "A novel early diagnostic framework for chronic diseases with class imbalance», *Scientific Reports*, Vol. 12(1), pp. 8614.
- [17] Tanveer, M., Rajani, T., Rašogi, R., Shao, Y.H. and Ganai, M.A., 2022. Comprehensive review on twin support vector machines. *Annals of Operations Research*, pp. 1-46.
- [۱۸] ناجی عظیمی. زهرا، تابستان ۹۵، آشنایی با برنامه‌ریزی خطی فازی، ویراسته وحیدیان کامباد، علی، ویرایش اول، مشهد، انتشارات دانشگاه فردوسی مشهد.
- [۱۹] سبزه‌کار، مصطفی، آذر ۸۸، بررسی تأثیر توجه مضاعف به نمونه‌های یادگیری با استفاده از قیود بهینه‌سازی طبقه‌بندی ماشین بردار پشتیبان، پایان‌نامه کارشناسی ارشد، گروه مهندسی کامپیوتر، دانشکده مهندسی، دانشگاه فردوسی مشهد، صفحات ۸۲-۹۰.
- [20] Jayadeva, Khemchandani, R. and Chandra, S., 2018. "Twin support vector machines: models, extensions and applications», Springer Publishing Company, Incorporated.
- [21] Nasiri, J.A. and Mir, A.M., 2020. "An enhanced KNN-based twin support vector machine with stable learning rules», *Neural computing and applications*, Vol. 32(16), pp. 12949-12969.
- [22] Mozafari, K., Nasiri, J.A., Charkari, N.M. and Jalili, S., 2011, December. Action recognition by local space-time features and least square twin SVM (LS-TSVM). In 2011 first international conference on informatics and computational intelligence (pp. 287-292). IEEE.
- [23] Raschka, S. and Mirjalili, V. 2017, Python machine learning second edition. Birmingham, England: Packt Publishing.
- [24] Yuan, X., Chen, S., Sun, C. and Yuwen, L., 2022, "A novel early diagnostic framework for chronic diseases with class imbalance», *Scientific Reports*, Vol. 12(1), pp. 8614.
- Methods and Programs in Biomedicine, Vol. 195, pp. 105565.
- [2] Alkenani, A. H., Li, Y., Xu, Y. & Zhang, Q., 2021. "Predicting Alzheimer's disease from spoken and written language using fusion-based stacked generalization", *Journal of Biomedical Informatics*, Vol. 118, pp. 103803.
- [3] Yuan, X., Chen, S., Yuwen, L., An, S., Mei, S. & Chen, T., 2020. "An improved SEIR model for reconstructing the dynamic transmission of COVID-19", In *Proceedings of IEEE International Conference on Bioinformatics and Biomedicine (BIBM)*, pp. 2320 – 2327.
- [4] Guo, Y. et al., 2021. "A review of wearable and unobtrusive sensing technologies for chronic disease management", *Computers in Biology and Medicine*, Vol. 129, pp. 104163.
- [5] Wang, Haidong, et al., 2016. "Global, regional, and national life expectancy, all-cause mortality, and cause-specific mortality for 249 causes of death, 1980–2015: a systematic analysis for the Global Burden of Disease Study 2015", *The lancet* 388, Vol. 10053, pp. 1459-1544.
- [6] Beyer, K.M., Namin, S., 2023. "Chronic environmental diseases: burdens, causes, and response», In *Biological and environmental hazards, risks, and disasters*, Elsevier, pp. 223-249.
- [7] Organization, W. H. WHO Noncommunicable diseases <https://www.who.int/westernpacific/health-topics/noncommunicable-diseases>
- [۸] نورالدینی، صدرالدین، ۱۴۰۰. «اثر اقتصادی بیماری‌های مزمن بر هزینه و درآمد خانوار ایرانی»، نشریه علمی (فصلنامه)، پژوهش‌ها و سیاست‌های اقتصادی شماره ۹، صفحه ۲۰۷ – ۲۴۱.
- [9] Delpino, F.M., Costa, Â.K., Farias, S.R., Chiavegatto Filho, A.D.P., Arcêncio, R.A. and Nunes, B.P., 2022. "Machine learning for predicting chronic diseases: a systematic review», *Public Health*, Vol. 205, pp. 14-25.
- [10] Maniruzzaman, M., Rahman, M.J., Ahammed, B. and Abedin, M.M., 2020. "Classification and prediction of diabetes disease using machine learning paradigm», *Health information science and systems*, Vol. 8, pp. 1-14.
- [11] Huang, J., Huth, C., Covic, M., Troll, M., Adam, J., Zukunft, S., Prehn, C., Wang, L., Nano, J., Scheerer, M.F. and Neschen, S., 2020. "Machine learning approaches reveal metabolic signatures of incident chronic kidney disease in individuals with prediabetes and type 2 diabetes», *Diabetes*, Vol. 69(12), pp. 2756-2765.
- [12] Zhang, L., Yuan, M., An, Z., Zhao, X., Wu, H., Li, H., Wang, Y., Sun, B., Li, H., Ding, S. and Zeng, X., 2020. "Prediction of hypertension, hyperglycemia and dyslipidemia from retinal fundus photographs via deep learning: A cross-sectional study of chronic diseases in central China», *PloS one*, Vol. 15(5), p.e0233166.
- [13] Al-Azzam, N. and Shatnawi, I., 2021. "Comparing supervised and semi-supervised machine learning models on diagnosing breast cancer», *Annals of Medicine and Surgery*, Vol. 62, pp. 53-64.
- [14] Xie, Y., Meng, W.Y., Li, R.Z., Wang, Y.W., Qian, X., Chan, C., Yu, Z.F., Fan, X.X., Pan, H.D., Xie, C. and Wu, Q.B., 2021. "Early lung cancer diagnostic biomarker discovery by