

افزایش صحت طبقه‌بندی بیماران دیابتی از لحاظ محدودیت عملکردی با استفاده از ترکیب خطی و غیرخطی بیومارکرها: روش Ramp AUC

لیلی فرجی گاوگانی^۱، پروین سربخش^۲، محمد اصغری جعفرآبادی^۳ مرتضی شمشیرگران^۴

تاریخ دریافت: ۹۷/۸/۲۷

تاریخ پذیرش: ۹۹/۳/۱

چکیده:

سطح زیر منحنی راک یک معیار مرسوم برای ارزیابی عملکرد طبقه‌بندی بیومارکرها است. در عمل یک بیومارکر قدرت طبقه‌بندی محدودی دارد لذا برای بهبود عملکرد طبقه‌بندی، علاقه‌مند به ترکیب مقادیر مربوط به بیومارکرها به صورت خطی و غیرخطی هستیم در این مطالعه ضمن معرفی انواع توابع زیان، به معرفی روش Ramp AUC و برخی ویژگی‌های آن به عنوان یک مدل آماری مبتنی بر سطح زیر منحنی راک پرداخته می‌شود. این مدل جهت ترکیب بیومارکرها به شکل خطی یا غیرخطی باهدف بهبود عملکرد طبقه‌بندی و مینیمم کردن تابع زیان تجربی بر اساس تابع زیان Ramp AUC ارائه شده است. به عنوان مثال کاربردی، در این مطالعه از داده‌های ۳۷۸ بیمار دیابتی مراجعه کننده به مراکز دیابتی اردبیل و تبریز در سال ۱۳۹۴-۱۳۹۳ استفاده شده است. جهت طبقه‌بندی بیماران دیابتی از لحاظ وضعیت محدودیت عملکردی بر مبنای بیومارکهای جمعیت شناختی و بالینی از روش RAUC استفاده گردید. اعتبارسنجی مدل به روش آموزش و آزمایش انجام شد. بر اساس نتایج گروه آزمایش، مقادیر سطح زیر منحنی به دست آمده برای مدل RAUC با ترکیبات خطی از بیومارکرها در قالب هسته خطی برابر ۰/۸۱ و با هسته تابع پایه شعاعی برابر ۱/۰۰ می‌باشد. نتایج بیانگر وجود یک الگوی غیرخطی قوی در داده‌ها می‌باشد به طوری که ترکیبات غیرخطی از بیومارکرها عملکرد طبقه‌بندی بالاتری نسبت به ترکیبات خطی را دارا می‌باشند.

واژه‌های کلیدی: تابع هسته، تابع زیان، سطح زیر منحنی راک، طبقه‌بندی، محدودیت عملکردی.

۱ مقدمه

می‌تواند در نظر گرفتن فضایی از ترکیبات چندجمله‌ای یا اسپلاین (توابع چندجمله‌ای، چند ضابطه‌ای با تعدادی گره‌های به هم پیوسته) برای روابط بین بیومارکرها و انتخاب بسط مناسبی از آن‌ها به عنوان ترکیب بهینه (خطی یا غیرخطی) از بیومارکرها باشد [۱۰]. مدل‌های آماری زیادی مانند رگرسیون لورستیک و ماشین بردار پشتیبان^۶ و مدل جمعی تعمیم یافته^۷ برای ماکزیمم کردن AUC بر مبنای ترکیباتی از بیومارکرها معرفی شده‌اند [۱۱، ۲۳، ۲۶، ۱۷، ۲۸] ولی این مدل‌ها دو محدودیت عمده دارند: اول اینکه اکثر این مدل‌ها به منظور پیدا کردن بهترین ترکیب خطی از بیومارکرها مورد استفاده قرار می‌گیرند که ممکن است در صورت وجود رابطه غیرخطی قوی بین آن‌ها یا در صورت وجود اثر متقابل بین بیومارکرها، کارایی لازم را نداشته باشند. مورد دوم اینکه بسیاری از این مدل‌ها برای پیدا کردن بهترین ترکیبات از

در عمل یک بیومارکر قدرت طبقه‌بندی محدودی دارد لذا برای بهبود عملکرد طبقه‌بندی، علاقه‌مند به ترکیب مقادیر مربوط به بیومارکرها به صورت خطی و غیرخطی هستیم. سطح زیر منحنی راک^۵ یک معیار مرسوم برای ارزیابی عملکرد طبقه‌بندی بیومارکرها می‌باشد. این معیار بیانگر میزان پیشگویی متغیر وابسته توسط مدل به کار برده شده، می‌باشد. مقدار این معیار مابین ۰/۵ تا ۱ می‌باشد که مقدار ۰/۵ نشان دهنده تصادفی بودن مدل، مقدار بالای ۰/۷ بیانگر دقت خوب مدل و مقدار بالای ۰/۹ بیانگر دقت بالای مدل به کار برده شده می‌باشد [۱۹].

بررسی روابط بین بیومارکرها به صورت یک ترکیب خطی، عملکرد طبقه‌بندی را محدود می‌کنند. یک راه حل برای بهبود عملکرد طبقه‌بندی،

^۱ کارشناس ارشد آمار زیستی، دانشکده پزشکی، مرکز تحقیقات پزشکی مبتنی بر شواهد، دانشگاه علوم پزشکی تبریز، تبریز، ایران

^۲ دکتری آمار زیستی، گروه آمار و اپیدمیولوژی دانشگاه علوم پزشکی تبریز، تبریز، ایران

^۳ دانشیار مرکز تحقیقات آموزش علوم پزشکی، دانشگاه علوم پزشکی تبریز، تبریز، ایران

^۴ دکتری اپیدمیولوژی، گروه آمار و اپیدمیولوژی دانشگاه علوم پزشکی تبریز، تبریز، ایران

^۵ Area Under the ROC Curve: AUC

^۶ Support Index Vector Machine

^۷ Generalized Additive Model

بررسی و $\square$ ، اندیس مربوط به نمونه موردها و شاهدها می‌باشد. x_i, x_j بردار مربوط به مقادیر بیومارکرها می‌باشد و β بردار ضرایب و p اندیس مربوط به تمام زوج‌های متشکل از نمونه موردها و شاهدها می‌باشد این تابع زیان تجربی غیر هموار و مشتق ناپذیر است لذا برای مینیمم کردن آن به جای تابع پله‌ای $I(X) - 1$ از تقریب‌های توابع سیگموئید^{۱۲} شامل الگوریتم تفاضل توابع محدب نرمال یا لوژستیک، تابع هینگ^{۱۳} یا تابع زیان درجه ۲ استفاده می‌شود [۱۰].

۲ انواع طبقه‌بندی‌ها بر پایه سطح زیر منحنی راک

دو گروه عمده از روش‌ها برای طبقه‌بندی بر پایه سطح زیر منحنی راک وجود دارد. چون هر مشاهده به‌طور مستقیم در محاسبه سطح زیر منحنی راک نقش دارند، در مجموع شامل $\square$ جمله خواهند بود که $\square$ بیانگر اندازه نمونه می‌باشد و به رویکرد تک شاخصه^{۱۴} معروف هستند و تابع زیان استفاده‌شده در خود روش‌ها را بهینه می‌کنند. روش‌های رگرسیون لوژستیک که تابع درستیابی و مدل ماشین بردار پشتیبان که تابع معیار هینگ را بهینه می‌کنند جزء این گروه می‌باشند [۲۲].

گروه دوم به منظور حداکثر کردن تابع سطح زیر منحنی راک تجربی، به عبارتی حداقل کردن تابع زیان سطح زیر منحنی تجربی، مورد استفاده قرار می‌گیرد. در این تابع داده‌ها به دو گروه کنترل و شاهد تقسیم می‌شوند و بر مبنای مقایسه زوج مشاهدات می‌باشد، و در مجموع شامل n^2 جمله خواهند بود. به این گروه رویکرد دو شاخصه یا رویکرد سطح زیر منحنی راک، گفته می‌شود [۱۰] رویکرد سطح زیر منحنی راک نسبت به داده‌های پرت نیرومندتر از رویکرد تک شاخصه می‌باشد. تابع زیان سطح زیر منحنی راک به صورت رابطه (۲) می‌باشد:

$$\frac{1}{n_D n_{\bar{D}}} \sum_{i=1}^{n_D} \sum_{j=1}^{n_{\bar{D}}} [1 - I\{(x_i - x_j)^T \beta > 0\}] \quad (2)$$

$$\equiv \frac{1}{n_D n_{\bar{D}}} \sum_{p=1}^{n_D n_{\bar{D}}} \{1 - I(\eta_p^\Delta > 0)\}$$

$\eta_p^\Delta = (x_i - x_j)^T \beta$ و $n_D n_{\bar{D}}$ برابر تعداد موردها و شاهدها در نمونه تحت بررسی می‌باشد. زمانی که اندازه داده‌های آموزش، بزرگ باشد و سطح زیر منحنی راک به عنوان معیار عملکرد طبقه‌بندی انتخاب شود، رویکرد

الگوریتم گرادیانت استفاده می‌کنند که تابع زیان مدنظر را به صورت موضعی^۸ و نه مطلق بهینه می‌کنند [۱۰] مورد سوم اینکه این روش در مقایسه با روش‌های موجود نسبت به داده‌های پرت حساس نیست و برای تعداد کمی از داده‌ها نیز دارای اعتبار کافی می‌باشد. اخیراً در سال ۲۰۱۶ مدلی به نام RAUC^۹ توسط فونگ^{۱۰} و همکاران برای مرتفع نمودن ایرادات ذکر شده در مورد مدل‌های مذکور معرفی شده است [۱۰]. این مدل با استفاده از تابع زیان RAMP و الگوریتمی غیر از گرادیانت یعنی الگوریتم تفاضل توابع محدب^{۱۱} همچنین روش تابع هسته اقدام به ترکیب بهینه بیومارکرها می‌کند به طوری که توان طبقه‌بندی متغیر پاسخ با استفاده از ترکیبات ساخته‌شده نسبت به وجود یک بیومارکر یا ترکیبات ساده آن‌ها بیشتر باشد و با مطالعات شبیه‌سازی نشان داده‌شده است که در مجموع نسبت به مدل‌های مرسوم کارایی بهتری (طبقه‌بندی صحیح‌تری) دارد [۱۰]. RAUC با استفاده از الگوریتم تفاضل توابع محدب فضای ترکیبات بیومارکرها را جستجو می‌کند که این فضا ترکیبات غیرخطی را نیز شامل شود [۱۰]. هدف مدل RAUC کاهش میزان تابع زیان سطح زیر منحنی راک تجربی می‌باشد [۱۰].

تعریف ۱.۱. در نظریه برآورد نقطه‌ای و علم آمار و بهینه‌سازی، تابع هزینه یا تابع زیان تابعی است که میزان زیان در یک رویداد تصادفی را نشان می‌دهد. اغلب مقدار مورد انتظار برای این تابع زیان، ریسک آماری گفته می‌شود که به منظور مقایسه دو یا چند برآوردگر مورد استفاده قرار می‌گیرد. در چنین مقایسه‌هایی، برآوردگر با کمترین مقدار زیان مورد انتظار معمولاً ترجیح داده می‌شود. [۶].

اکثر روش‌های مرسوم مانند رگرسیون لوژستیک، ماشین بردار پشتیبان و مدل جمعی تعمیم‌یافته برای بهینه کردن تابع زیان از الگوریتم گرادیانت استفاده می‌کنند که منجر به بهینه‌سازی مکانی و نه بهینه‌سازی مطلق در کل فضای مورد بررسی می‌شوند [۱۰].

رویکرد سطح زیر منحنی راک برای به حداکثر رساندن سطح زیر منحنی راک تجربی، یا معادل آن به حداقل رساندن تابع زیان سطح زیر منحنی راک، به صورت رابطه (۱) تعریف می‌شود:

$$\frac{1}{n_D n_{\bar{D}}} \sum_{i=1}^{n_D} \sum_{j=1}^{n_{\bar{D}}} [1 - I\{(x_i - x_j)^T \beta > 0\}] \quad (1)$$

$$\equiv \frac{1}{n_D n_{\bar{D}}} \sum_{p=1}^{n_D n_{\bar{D}}} \{1 - I(\eta_p^\Delta > 0)\}$$

$\eta_p^\Delta = (x_i - x_j)^T \beta$ و $n_D n_{\bar{D}}$ برابر تعداد موردها و شاهدها در نمونه تحت

⁸local

⁹Ramp AUC

¹⁰Fong

¹¹Difference of Convex functions Algorithm: DCA

¹²Sigmoid

¹³Hing

¹⁴single-indexed approach

۳ معرفی روش RAUC

RAUC یک روش جدید در پیدا کردن بهترین ترکیب از بیومارکرها مبتنی بر سطح زیر منحنی راک، بر پایه الگوریتمی غیر از گرادیانت یعنی الگوریتم تفاضل توابع محدب و استفاده از روش تابع هسته است که ترکیبات غیر خطی را نیز شامل می‌شود. برای این منظور استفاده از تقریب تابع پله‌ای $1 - I(x)$ در تابع زیان سطح زیر منحنی راک موجود در رابطه (۲)، به صورت تابع شیب‌دار در رابطه (۳) پیشنهاد شده است:

$$h_s(u) = \begin{cases} 1 & \text{اگر } \frac{u}{s} \leq -\frac{1}{4} \\ -\frac{u}{s} + \frac{1}{4} & \text{اگر } -\frac{1}{4} \leq \frac{u}{s} \leq \frac{1}{4} \\ 0 & \text{اگر } \frac{u}{s} \geq \frac{1}{4} \end{cases} \quad (3)$$

$s > 0$ یک پارامتر مقیاس می‌باشد و $\sum_{p=1}^{n_D n_D} \left(\frac{h_s(\eta_p^A)}{n_D n_D} \right)$ را تابع زیان RAMP AUC می‌نامند.

برای هر مقدار $\eta_p^A \neq 0$ مقدار این تابع زیان بین ۰ تا ۱ تغییر می‌کند و در بعضی نقاط که $\square$ به سمت ۰ میل کند، مقدار آن ۰ یا ۱ می‌شود. شکل ۱ روش‌های مختلف طبقه‌بندی بر مبنای سطح زیر منحنی راک بر اساس توابع زیان مختلف را نشان می‌دهد که بیانگر تقریب بهتری از تابع پله‌ای سطح زیر منحنی راک در صورت استفاده از روش RAUC می‌باشد.

تعریف ۱.۳. در صورتی که برای هر یک از پارامترهای آزاد مدل از روی داده‌های مشاهده شده متفاوت، یک مقدار یگانه به دست بیاید آن مدل دارای مسئله شناسایی پذیری^{۱۷} خواهد بود.

مسئله شناسایی پذیری در همه روش‌های مبتنی بر رویکرد سطح زیر منحنی راک وجود دارد. این مشکل در مینیمم تابع زیان رویکرد سطح زیر منحنی راک، به ازای مقدار ثابت $\square$ به صورت $1 - I(x > 0) = 1 - I(cx > 0)$ نمایش داده می‌شود [۱۰].

تاوان تابع زیان RAMP به صورت رابطه (۴) می‌باشد:

$$\sum h(\eta_p^A) + \frac{1}{4} \lambda_n \|\beta\|_2^2 \quad (4)$$

سطح زیر منحنی راک، رویکردی بهتر از رویکرد تک شاخص خواهد داشت [۲۲]. روش‌هایی مانند رگرسیون لوژستیک و ماشین بردار پشتیبان نسبت به داده‌های پرت حساس هستند. نیرومندی رویکرد سطح زیر منحنی راک نتیجه بهینه‌سازی تابع معیار استفاده شده در ارزیابی عملکرد طبقه‌بندی است که به تنظیم پارامتر خاص در مورد نیرومند بودن، نیاز ندارد و تنها از تابع پله‌ای که حالت رتبه‌بندی دارد استفاده شده است. سطح زیر منحنی راک بیشتر از رویکرد تک شاخص مشکل بهینه‌سازی دارد به عنوان مثال رگرسیون لوژستیک که دارای رویکرد تک شاخص می‌باشد تابعی محدب و هموار هست اما رویکرد سطح زیر منحنی راک تجربی این چنین نیست و باید طی مراحل بهینه شود [۱۰].

هم چنین دو گروه اصلی از روش‌ها بر پایه رویکرد AUC وجود دارد. اکثر روش‌ها بر پایه رویکرد سطح زیر منحنی راک، روش‌های سطح زیر منحنی راک هموار^{۱۵} هستند در این روش‌ها از تابع هموار سیگموئید مانند $S_s(x)$ به عنوان تقریبی از تابع پله‌ای $(1 - I(x > 0))$ در تابع زیان سطح زیر منحنی راک استفاده می‌شود. به عنوان مثال تقریب پله‌ای تابع توزیع چگالی نرمال مانند $S_s(x) = \varphi\left(-\frac{x}{s}\right)$ و تقریب تابع لوژستیک مانند $S_s(x) = \frac{1}{1 + \exp\left(\frac{x}{s}\right)}$ [۲۱، ۱۳، ۲۹] استفاده می‌شوند. با تغییر دادن تقریب استفاده شده، تابع معیار نیز تغییر خواهد کرد. هر دو تقریب نرمال و لوژستیک، پارامتر مقیاس $\square$ که بیانگر دقت تقریب می‌باشد، کنترل می‌کنند که $\square$ کوچک‌تر منجر به تقریب بهتر می‌شود اما در این صورت بهینه‌سازی سخت خواهد بود [۸] تقریب‌های هموار دیگری مانند تقریب چندجمله‌ای نیز پیشنهاد شده‌اند [۸]. این گروه به دلیل اینکه از الگوریتم تفاضل توابع محدب برای بهینه‌سازی استفاده می‌کنند بیشتر از گروه دوم کاربرد دارند.

دومین گروه از روش‌های مربوط به رویکرد سطح زیر منحنی راک، مدل SVM-AUC^{۱۶} می‌باشد [۱۶، ۳۰، ۵] در این روش تابع پله‌ای $1 - I(x > 0)$ در تابع زیان سطح زیر منحنی راک به وسیله تابع زیان هینگ که به صورت $(1 - x)_+$ است، جایگزین می‌شود. SVM-AUC تقریب خوبی از تابع زیان سطح زیر منحنی راک نمی‌باشد به خصوص زمانی که در بین داده‌های آموزش، داده‌های پرت وجود داشته باشد، هم چنین از الگوریتم گرادیانت برای بهینه‌سازی استفاده می‌کند.

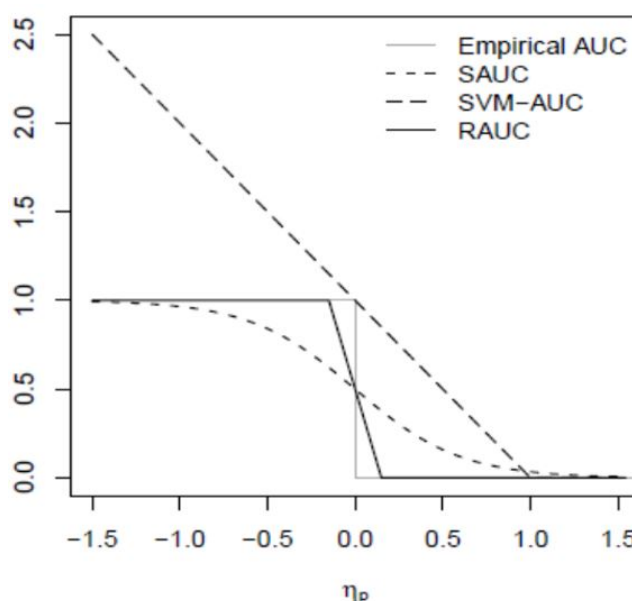
مینیمم RAUC دارای مشکل شناسایی پذیری می‌باشد زیرا به ازای $h_s(x) \geq h_s(cx)$ ، $c > 1$ یک روش پیشنهادی جدید برای حل این مسئله، جایگذاری $s = 1$ و جریمه $\|\beta\|_2^2$ می‌باشد. در نهایت حداقل مقدار

¹⁵Smoothed Area Under the ROC Curve: SAUC

¹⁶Support Vector Machine-AUC

¹⁷identifiability

¹⁸adjust



شکل ۱: نمودار انواع توابع زیان مورد استفاده در سطح زیر منحنی راک [۱۰]

است و به صورت رابطه (۵) می‌توان نوشت:

$$h(x) = h_1(x) - h_2(x) \quad (5)$$

$$h_1(x) = \left(\frac{1}{2} - x\right)_+, h_2(x) = \left(-\frac{1}{2} - x\right)_+$$

این نحوه نوشتن، امکان استفاده از الگوریتم تفاضل توابع محدب را به وجود می‌آورد [۱۸]. اساس الگوریتم تفاضل توابع محدب در این مسئله، استفاده از یک سری مسائل بهینه‌سازی بر حسب توابع محدب برای بهینه کردن یک تابع غیر محدب است و این به وسیله تقریب تکراری h_2 با بسط تیلور مرتبه اول به دست می‌آید و به تقریبی برای h_1 هم نیازی نیست زیرا تفاوت بین یک تابع محدب و یک تابع خطی یک تابع محدب است. شکل ۲ نشان می‌دهد تابع شیب‌دار h به صورت تفاضل دو تابع محدب h_1 و h_2 نوشته می‌شود.

□ گام اول: یک مقدار اولیه برای η_p^Δ در نظر می‌گیریم و آن را با $\eta_p^{(\Delta,0)}$ نشان می‌دهیم. برای مثال اگر $\hat{\theta}^{rob}$ را به عنوان برآورد اولیه رگرسیون لورستیک نیرومند^{۱۹} به کار ببریم رابطه (۶) را خواهیم داشت:

$$(1, \hat{\theta}^{rob})^T (x_{ip} - x_{jp})^T \rightarrow \eta_p^{\Delta,0} \quad (6)$$

□ گام دوم: حل معادله موجود در رابطه (۷) می‌باشد:

$$\hat{\beta} = \arg \min \sum_{p=1}^{n_D n_D} \left\{ h_1(\eta_p^\Delta) - h_2(\eta_p^\Delta, \eta_p^{\Delta,0}) + \frac{1}{2} \lambda_n \|\beta\|_2^2 \right\} \quad (7)$$

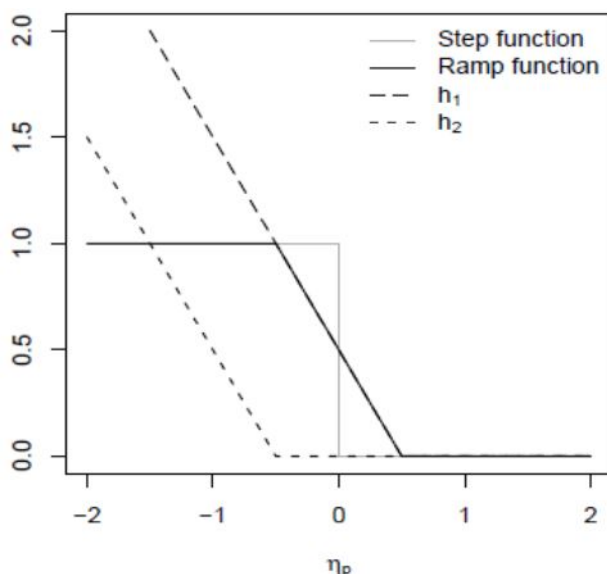
این در مینیم کردن تابع زیان RAUC مؤثر است، همچنین باعث محدودیت و ثبات مقدار $\|\beta\|_2^2$ می‌شود. λ_n به عنوان یک اهرم برای تعدیل^{۱۸} کردن $\|\beta\|_2^2$ به کار برده می‌شود به طوری که وقتی λ_n کوچک باشد $\|\beta\|_2^2$ مقدار بزرگی در نظر گرفته می‌شود و نسبت اشباع افزایش می‌یابد بنابراین اگر λ_n به اندازه کافی کوچک در نظر گرفته شود به طوری که حداقل ۷۵٪ از $h(\eta_p^\Delta)$ برابر ۰ یا ۱ باشد عملکرد طبقه‌بندی رضایت‌بخش و مشکل بهینه‌سازی حل خواهد شد [۲۷].

۴ بهینه‌سازی و ترکیب غیرخطی بیومارکرها

تابع زیان RAUC یک تابع غیر محدب و غیر هموار است برای حل مشکل غیر محدب بودن و پیدا کردن بهترین ترکیب، تابع شیب‌دار h به صورت تفاضل دو تابع محدب h_1 و h_2 نوشته می‌شود که در شکل ۲ نشان داده شده

فرایند بهینه‌سازی الگوریتم تفاضل توابع محدب به شرح زیر می‌باشد:

¹⁹Robust Logistic Regression



شکل ۲: تفاضل دو تابع محدب به عنوان برآوردی از تابع RAUC [۱۰]

روش لاگرانژ است. روش لاگرانژ اولیه در رابطه (۹) نشان داده شده است:

$$L_p = \sum_{p=1}^{n_D n_D} \left\{ \varepsilon_p - h_2 \left(\eta_p^{\Delta,0} \right) \eta_p^{\Delta} \right\} + \frac{1}{\gamma} \lambda_n \|\beta\|_2^2 - \sum_p \alpha_p \left(\varepsilon_p + \eta_p^{\Delta} - \frac{1}{\gamma} \right) - \sum_p \gamma_p \varepsilon_p \quad (9)$$

که α و γ ، لاگرانژ چندگانه غیر منفی می‌باشند و مربوط به دو مجموعه از محدودیت‌های موجود در مسئله بهینه‌سازی می‌باشند. اگر مشتق مرتبه اول لاگرانژ را به دست آورده و برابر قرار دهیم، مقادیر قابل قبول برای متغیرهای فضای اولیه ε و β را به صورت رابطه (۱۰) به دست خواهیم آورد:

$$\beta = \frac{1}{\lambda_n} \sum_{p=1}^{n_D n_D} \left\{ h_2 \left(\eta_p^{\Delta,0} \right) + \alpha_p \right\} x_p^{\Delta}, \quad 1 = \alpha_p + \gamma_p \quad (10)$$

همچنین اگر مقادیر به دست آمده را در L_p جایگذاری کنیم شرایط سخت مربوط به ε_p از بین خواهد رفت. در بعضی مواقع در حل جبری ممکن است با دو برابر مشکل موجود مواجه شویم، مثلاً در مواردی که $0 \leq \alpha_p \leq 1$ باشد رابطه (۱۱) را خواهیم داشت:

$$\min_{\alpha} \alpha^T Q \alpha - b^T \alpha \quad (11)$$

در عبارت بالا Q ماتریس توان دوم به صورت $[p, p']^{th}$ با اجزاء $\langle x_p^{\Delta}, x_{p'}^{\Delta} \rangle / \lambda_n$ و $b = 1 - 2 \times Q \times h_2'(\eta^{\Delta,0})$ می‌باشد. برای یافتن بهترین ترکیب و پیدا کردن طبقه مربوطه به یک مشاهده‌ی جدید [۱۰] از رابطه (۱۲) استفاده می‌شود:

$$\eta_{new} = \sum_{p=1}^{n_D n_D} \left\{ h_2 \left(\eta_p^{\Delta,0} \right) + \alpha_p \right\} \langle x_{new}, x_p^{\Delta} \rangle \quad (12)$$

به طوری که به ازای تمام مقادیر x مشتق مرتبه اول $h_2(x)$ به جز در نقطه $x = -\frac{1}{2}$ که مقدار آن برابر صفر می‌باشد، وجود داشته باشد، $\hat{h}_2(\eta_p^{\Delta}, \eta_p^{\Delta,0}) = h_2'(\eta_p^{\Delta,0}) \eta_p^{\Delta} + h_2(\eta_p^{\Delta,0})$ و $h_2'(x) = -I(x < -\frac{1}{2})$ می‌باشد.

□ گام سوم: اگر مجموعه $\hat{\beta}(x_{ip} - x_{jp})^T$ به $\eta_p^{\Delta,0}$ میل کند آنگاه به گام دوم بازمی‌گردیم و این مراحل را تکرار می‌کنیم تا زمانی که تغییرات ایجاد شده در پناهی تابع زیان RAUC کمتر از یک مقدار ثابت از قبل تعیین شده باشد.

در گام دوم از الگوریتم فوق، یک مشکل بهینه‌سازی محدب حل می‌شود. به عنوان مثال h_1 و $\hat{h}_2$ توابعی غیر هموار هستند که با استفاده از روش استاندارد و معرفی متغیرهای ε_p به جای h_1 ، به یک مسئله بهینه‌سازی هموار تبدیل می‌شود.

در مواردی که $\varepsilon_p \geq \frac{1}{2} - \eta_p^{\Delta}$ و $\varepsilon_p \geq 0$ منجر به رابطه (۸) خواهد شد:

$$\min_{\varepsilon, \beta} \sum_{p=1}^{n_D n_D} \left\{ \varepsilon_p - h_2 \left(\eta_p^{\Delta,0} \right) \right\} + \frac{1}{2} \lambda_n \|\beta\|_2^2 \quad (8)$$

و این زمانی روی خواهد داد که محدودیت‌های مربوط به ε و β را در نظر بگیریم. چون در حل مشکل بهینه‌سازی، کار کردن با محدودیت‌ها سخت می‌باشد [۴] هینگ به دنبال پیدا کردن حل مشکل دوگانه (β, ε) از طریق

غیرخطی تحت عنوان فضای ویژگی، داده‌ها به فضایی با ابعاد بیشتر انتقال داده می‌شود تا در این فضا به صورت خطی از هم جدا شوند و در فضای جدید با جایگزینی x_k با $\varphi(x_k)$ مرز بهینه را می‌توان محاسبه نمود. فضای ویژگی دارای ابعاد بزرگ می‌باشد به خاطر همین دارای محاسبات زیاد و پرهزینه خواهد بود برای غلبه بر این مشکل از تابع هسته استفاده می‌شود که مبتنی بر ضرب داخلی می‌باشند [۳، ۱۵]. هسته‌های متفاوتی وجود دارد که انتخاب هسته مناسب بستگی به ماهیت داده‌ها دارد. در حالت کلی دو نوع هسته داریم: هسته خطی و هسته غیرخطی که شامل سیگموئید چندجمله‌ای و تابع شعاعی پایه^{۲۲} می‌باشد که متداول‌ترین آن‌ها تابع شعاعی پایه می‌باشد [۹] انواع توابع هسته همراه با الگوریتم‌های مربوطه در جدول ۱ نشان داده شده است:

جدول ۱. انواع تابع هسته

الگوریتم	تابع هسته
$K(x_i, x_j) = \langle x_i, x_j \rangle$	خطی
$K(x_i, x_j) = (r + \gamma \langle x_i, x_j \rangle)^d$	چندجمله‌ای
$K(x_i, x_j) = \exp(-\gamma \ x_i - x_j\ ^2)$	تابع شعاعی پایه
$K(x_i, x_j) = \tanh(\gamma X_i^T X_j + r)$	سیگموئید

۶ مثال کاربردی

مثال ۱.۶. مثال کاربردی مورداستفاده در این مطالعه ترکیبی از دو مطالعه مقطعی - تحلیلی صورت گرفته در شمال غرب ایران در سال ۱۳۹۴-۱۳۹۳ می‌باشد جزییات بیشتر در مورد این مطالعات در چندین مطالعه ارائه شده‌اند [۲۴، ۲۵، ۲۶]. در مطالعه حاضر با ترکیب داده‌های دو مطالعه و حذف موارد دارای داده گمشده برای متغیرهای موردنیاز، تعداد ۳۷۸ نفر با داده‌های کامل برای فرایند مدل‌سازی باقی ماندند. سن بیماران دیابتی بین ۲۵ و ۷۰ سال با میانگین ۵۶/۱۶ سال (با انحراف معیار ۸/۷۹ سال) و میانگین طول مدت ابتلا ۸/۹۱ سال (با انحراف معیار ۶/۱۱ سال) بود. از بین ۳۷۸ بیمار دیابتی، ۲۵۵ نفر (۶۷/۴۶ درصد) دارای محدودیت عملکردی بودند. از ۲۴۹ زن تحت مطالعه نیز ۷۲/۳ درصد دارای محدودیت عملکردی بودند درحالی‌که این نسبت در مردان ۵۸/۱ درصد بود و بین زنان و مردان تفاوت معنی‌داری از لحاظ احتمال ابتلا به محدودیت عملکردی مشاهده شده است ($p\text{-value} = 0.05$).

مدل $RAUC$ برای بررسی نحوه ارتباط بین بیومارکرهای جنسیت، سن، طول مدت ابتلا، نمایه توده بدن، HDL ، کلسترول، تری گلیسرید، قند خون ناشتا، فشارخون سیستولیک و دیاستولیک و هموگلوبین $A1C$ به عنوان عوامل خطر مرتبط و محدودیت عملکردی به عنوان پیامد به کار گرفته شده است. برای تحلیل بر اساس مدل $RAUC$ از هسته تابع پایه شعاعی به عنوان

$\eta(\Delta, \alpha)$ و α از آخرین تکرار الگوریتم تفاضل توابع محدب به دست می‌آید، طبق اصطلاحات ماشین بردار پشتیبان [۷] دو عبارت x_p^Δ و $h_2(\eta_p^{\Delta,0}) + \alpha_p = 0$ را جفت غیر پشتیبان^{۲۰} نامیده می‌شود زیرا که در امتیازدهی شرکت نمی‌کنند.

۵ استفاده از فضای ویژگی برای جداسازی ترکیبات غیرخطی

گاهی ممکن است بهترین طبقه‌بندی بین ترکیبات خطی از متغیرهای بیومارکر ورودی صورت نگیرد و نتوانیم داده‌ها را به صورت خطی جدا کنیم در این صورت آن‌ها را به فضایی با ابعاد بیشتر نگاشت می‌دهیم تا بتوان آن‌ها را در این فضا به صورت خطی جدا کرد بدین منظور برای به دست آوردن بهترین ترکیب از فضای ویژگی استفاده می‌شود. فضای ویژگی برای مورد i با φ_i نشان می‌دهند. در این صورت مشکل دوگانه موجود در رابطه (۱۲) با استفاده از Q که یک ماتریس توان دوم به صورت $[p, p]^{th}$ با اجزاء $\langle x_p^\Delta, x_p^\Delta \rangle / \lambda_n$ می‌باشد می‌توان حل کرد که $\langle \varphi_p^\Delta, \varphi_p^\Delta \rangle$ را به صورت رابطه (۱۳) نوشت:

$$\langle \varphi_p^\Delta, \varphi_p^\Delta \rangle = \langle \varphi_{ip}, \varphi_{ip} \rangle + \langle \varphi_{jp}, \varphi_{jp} \rangle - \langle \varphi_{jp}, \varphi_{jp'} \rangle - \langle \varphi_{ip}, \varphi_{jp'} \rangle \quad (13)$$

مشکل دوگانگی موجود در مرحله سوم را با معرفی $\hat{\eta}^\Delta$ در رابطه (۱۴) می‌توان حل کرد، به طوری که نیازی به برگشتن به مرحله دوم نباشد:

$$\hat{\eta}^\Delta = (\varphi_1^\Delta, \dots, \varphi_p^\Delta) \hat{\beta} = Q \times \{h'_2(\eta^{\Delta,0}) + \hat{\alpha}\} \quad (14)$$

φ_i ممکن است بی نهایت-بعدی باشد و $\hat{\eta}^\Delta$ معمولاً نماینده یک-بعدی می‌باشد. بنابراین نمره یک مشاهده جدید را می‌توان بر اساس رابطه (۱۵) به دست آورد:

$$\begin{aligned} \eta^{new} &= \sum_{p=1}^{n_D n_D} \{h'_2(\eta_p^{\Delta,0}) + \alpha_p\} \langle \varphi(x_{new}), \varphi_p^\Delta \rangle \\ &= \sum_{p=1}^{n_D n_D} \{h'_2(\eta_p^{\Delta,0}) + \alpha_p\} \\ &\quad \times \{ \langle \varphi(x_{new}), \varphi_{ip} \rangle - \langle \varphi(x_{new}), \varphi_{jp} \rangle \} \quad (15) \end{aligned}$$

$\eta_p^{\Delta,0}$ و α از آخرین تکرار الگوریتم تفاضل توابع محدب به دست می‌آیند. این رویکرد از نگاشت از بردار ورودی در یک فضای ویژگی بدون نیاز به مشخص کردن دقیق مقدار نگاشت شده در ادبیات ماشین بردار پشتیبان، تکنیک هسته‌ای^{۲۱} نامیده می‌شود [۱، ۱۲].

وقتی داده‌ها در فضای ورودی به صورت خطی تفکیک‌پذیر نباشند، یعنی کلاس داده‌ها هم‌پوشانی داشته باشند در این صورت با استفاده از توابع

²¹Kernel trick²²Radial Basis Function) RBF

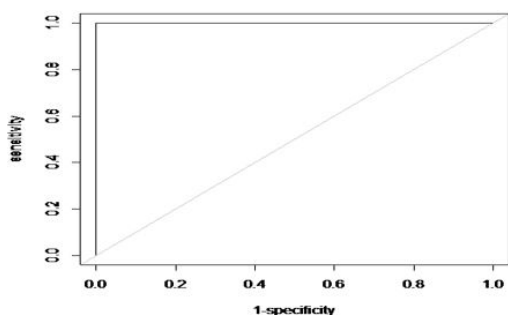
۷ بحث و نتیجه گیری

ارزیابی توانایی طبقه بندی در تشخیص یا پیشگویی صحیح طبقه ضروری است. معیارهای مختلفی برای ارزیابی عملکرد طبقه بندی وجود دارد ولی در طبقه بندی های دو طبقه ای بیشتر سطح زیر منحنی راک توصیه شده است. علاقه مندی به استفاده از سطح زیر منحنی راک برای ارزیابی عملکرد طبقه بندی منجر به طراحی الگوریتم هایی شده است که هدف آن ها به حداکثر رساندن سطح زیر منحنی راک می باشد [۱۴].

سطح زیر منحنی راک اغلب برای اندازه گیری عملکرد برآوردگر در مسائل مربوط به طبقه بندی دو سطحی استفاده می شود. این معیار بیانگر این موضوع می باشد که مدل به کار برده شده، به چه اندازه متغیر وابسته را می تواند به خوبی پیشگویی کند.

در پژوهش حاضر میزان درستی طبقه بندی افراد دیابتی از لحاظ محدودیت عملکردی، بر مبنای بیومارکرهایی مانند جنسیت، سن، طول مدت ابتلا، و بیومارکرهای کیفیت مراقبت شامل فشارخون، قند خون، نمایه توده بدن، کلسترول و هموگلوبین A1C تری گلیسیرید مورد بررسی قرار گرفت و با توجه به نتایج حاصل از این پژوهش و مقدار سطح زیر منحنی راک به دست آمده از مدل RAUC با تابع هسته غیرخطی پایه شعاعی که بیان گر یک الگوی غیرخطی بین داده ها می باشد می توان بیان کرد که ترکیبات غیرخطی عملکرد طبقه بندی یا پیشگویی بالاتری نسبت به ترکیبات خطی را دارا می باشند.

با توجه به اهمیت طبقه بندی صحیح مطالعات پزشکی، تشخیص و تعیین ارتباط پیچیده و غیرخطی بین بیومارکرها، بر اساس روش هایی که برای یافتن چنین ترکیبات بهینه استفاده می گردد مدل RAUC می تواند مفید باشد.



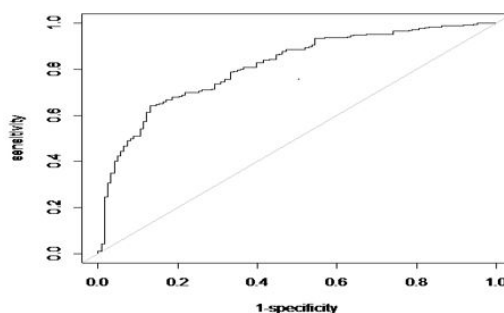
شکل ۴:

نمودار سطح زیر منحنی راک مربوط به مدل RAUC با تابع هسته غیرخطی تابع شعاعی پایه

بر کاربردترین و مهم ترین هسته غیرخطی به منظور بررسی ترکیبات غیرخطی و تابع هسته خطی به منظور بررسی ترکیبات خطی استفاده شده است. برای مقایسه قدرت تشخیص و افتراق این مدل با دو تابع هسته متفاوت استفاده شده، از سطح زیر منحنی راک استفاده شده است. تحلیل های لازم برای برازش مدل RAUC از نرم افزار R 3.3.1 با بسته نرم افزاری *aucom* استفاده گردیده است. به منظور اعتبارسنجی نتایج از روش آزمون و آزمایش مورد استفاده قرار گرفته است. داده ها به طور تصادفی به دو گروه، داده آموزشی (۷۰ درصد داده ها) و داده آزمایش (۳۰ درصد داده ها) تقسیم شدند. مقدار احتمال ابتلا به محدودیت عملکردی پس از برازش مدل RAUC بر روی داده های گروه آزمایش، برای داده های گروه آموزش محاسبه گردید و با استفاده از مقدار احتمال به دست آمده، شاخص سطح زیر منحنی راک برای تعیین صحت طبقه بندی بیماران دیابتی مبتلا به محدودیت عملکردی محاسبه گردید. مدل RAUC با استفاده از تابع هسته خطی و تابع شعاع پایه بر روی داده های آزمایش برازش داده شده است. مقادیر سطح زیر منحنی به دست آمده از نتایج حاصل از داده های آزمایش، برای مدل RAUC با تابع هسته خطی برابر ۰.۸۱ به دست آمده که همین مقدار بر اساس تابع شعاعی برابر ۱/۰۰ می باشد لذا با توجه به نتایج به دست آمده، هسته تابع پایه شعاعی بیشترین مقدار سطح زیر منحنی راک را نسبت به هسته خطی دارا می باشد.

نمودارهای مربوط به سطح زیر منحنی راک برای مدل RAUC با هسته خطی و هسته غیرخطی تابع شعاعی پایه در شکل های ۳ و ۴ نشان داده شده است.

نتایج حاصل از این پژوهش نشان دهنده این موضوع می باشد که ترکیبات غیرخطی عملکرد طبقه بندی یا پیشگویی بالاتری نسبت به ترکیبات خطی را دارا می باشند، این بیان گر وجود یک الگوی غیرخطی در داده های این مطالعه می باشد.



شکل ۳:

نمودار سطح زیر منحنی راک برای مدل RAUC با هسته خطی

مراجع

- [1] Aizerman, M. (1964). Theoretical foundations of the potential function method in pattern recognition learning. *Automation and remote control*, **25**, 821–37.
- [2] Ataei, J., Shamshirgaran, S., Iranparvar Alamdari, M. and Safaeian, A. (2015). Evaluation of Diabetes Quality of Care Based on a Care Scoring System among People Referring to Diabetes Clinic in Ardabil, 2014. *Journal of Ardabil University of Medical Sciences*, **15**(2), 207–19.
- [3] Bellotti, T. and Crook, J. (2009). Support vector machines for credit scoring and discovery of significant features. *Expert Systems with Applications*, **36**(2), 3302–8.
- [4] Boyd, S. and Vandenberghe, L. (2004). *Convex Optimization*, Cambridge University Press.
- [5] Brefeld, U. and Scheffer, T. (2005). AUC maximizing support vector learning. *Proceedings of the ICML 2005 workshop on ROC Analysis in Machine Learning*, Citeseer semanticscholar.
- [6] Bühlmann, P. and Van De Geer, S. (2011). *Statistics for high-dimensional data: methods, theory and applications*. Springer Heidelberg Dordrecht London. New York.
- [7] Burges, C.J. (1998). A tutorial on support vector machines for pattern recognition. *Data mining and knowledge discovery*, **2**(2), 121–67.
- [8] Calders, T. and Jaroszewicz, S. (2007). Efficient AUC optimization for classification, *Knowledge Discovery in Databases: PKDD 2007*, 42–53.
- [9] Daugherty, C., Dewhurst, W. and Spertus, J. (1998). Health related quality of life outcomes in clinical research. *American Journal of Epidemiology*, **51**(7), 569–75.
- [10] Fong, Y., Yin, S. and Huang, Y. (2016). Combining biomarkers linearly and nonlinearly for classification using the area under the ROC curve. *Statistics in medicine*, **35**(21), 3792–809.
- [11] Freund, Y. and Schapire, R. (1996). Experiments with a new boosting algorithm, in *Proceedings of the Thirteenth International Conference on Machine Learning*, 148–156, Morgan Kaufmann.
- [12] Hastie, T., Tibshirani, R. and Friedman, J. (2009). *The Elements of Statistical Learning: Data Mining, Inference, and Prediction*, Second Edition, Springer Series in Statistics, Springer. New York.
- [13] Herschtal, A. and Raskutti, B. (2004). Optimising area under the ROC curve using gradient descent, in *Proceedings of the Twenty-First International Conference on Machine Learning*, 49, ACM.
- [14] Holst, A. (2008). Efficient AUC maximization with regularized least-squares. *Tenth Scandinavian Conference on Artificial Intelligence*, SCAI 2008, IOS Press. Amsterdam.
- [15] Huang, C.L., Chen, M.C. and Wang, C.J. (2007). Credit scoring with a data mining approach based on support vector machines. *Expert systems with applications*, **33**(4), 847–56.
- [16] Joachims, T. (1998). Making large-scale SVM learning practical. *Technical Report, SFB 475: Komplexitätsreduktion in Multivariaten Datenstrukturen, Universität Dortmund*.

- [17] Komori, O. (2011). A boosting method for maximization of the area under the ROC curve. *Annals of the Institute of Statistical Mathematics*, **63**(5), 961–79.
- [18] Le Thi Hoai, A. and Tao, P.D. (1997). Solving a class of linearly constrained indefinite quadratic problems by DC algorithms. *Journal of global optimization*, **11**(3), 253–85.
- [19] Linkie, M., Smith, R.J. and Leader-Williams, N. (2004). Mapping and predicting deforestation patterns in the lowlands of Sumatra. *Biodiversity & Conservation*, **13**(10), 1809–18.
- [20] Lloyd, C.J. (1998). Using smoothed receiver operating characteristic curves to summarize and compare diagnostic systems. *Journal of the American Statistical Association*, **93**(444), 1356–64.
- [21] Ma, S. and Huang, J. (2007) Combining multiple markers for classification using ROC. *Biometrics*, **63**(3), 751–7.
- [22] Pepe, M.S., Cai, T. and Longton, G. (2006). Combining predictors for classification using the area under the receiver operating characteristic curve. *Biometrics*, **62**(1), 221–9.
- [23] Pepe, M.S. and Thompson, M.L. (2000). Combining diagnostic test results to increase accuracy. *Biostatistics*, **1**(2), 123–40.
- [24] Shamshirgaran, S.M., Ataei, A., Malek, A., Iranparvar-Alamdari, M. and Aminisani, N. (2017). Quality of sleep and its determinants among people with type 2 diabetes mellitus in Northwest of Iran. *World J Diabetes*, **8**(7), 358–364.
- [25] Shamshirgaran, S.M., Mamaghanian, A., Aliasgarzadeh, A., Aiminisani, N., Iranparvar-Alamdari, M. and Ataei, J. (2017). Age differences in diabetes-related complications and glycemic control. Under press. *BMC Endocrine Disorders*, **17**, 25.
- [26] Tutz, G. and Binder, H. (2006). Generalized Additive Modeling with Implicit Variable Selection by Likelihood-Based Boosting. *Biometrics*, **62**(4), 961–71.
- [27] Vexler, A., Liu, A., Schisterman, E.F. and Wu, C. (2006). Note on distribution-free estimation of maximum linear separation of two multivariate distributions. *Nonparametric Statistics*, **18**(2), 145–58.
- [28] Xu-hui, W., Ping, S., Li, C. and Ye, W. (2009). A ROC curve method for performance evaluation of support vector machine with optimization strategy. *Computer Science-Technology and Applications*, 2009 IFCSTA'09 International Forum on; IEEE.
- [29] Yan, L., Dodier, R.H., Mozer, M. and Wolniewicz, R.H. (2003). Optimizing classifier performance via an approximation to the Wilcoxon-Mann-Whitney statistic. *Proceedings of the 20th International Conference on Machine Learning (ICML-03)*.
- [30] Yu, L. and Yao, X. (2013). A total least squares proximal support vector classifier for credit risk evaluation. *Soft Computing*, **17**(4), 643–50.