

الگوریتم کاهش گروهی در مدل جمعی ناپارامتری توانیده

محمد کاظمی^۱

استادیار، گروه آمار، دانشکده علوم ریاضی، دانشگاه گیلان، رشت، ایران

رسید مقاله: ۲ تیر ۱۴۰۳

پذیرش مقاله: ۲۴ آبان ۱۴۰۳

چکیده

مدل جمعی ناپارامتری یکی از مدل‌های رایج برای مدل‌سازی رابطه بین متغیرها است. در این مقاله، مدل جمعی ناپارامتری بعد بالا را در نظر می‌گیریم که در آن تعداد متغیرهای توضیحی می‌تواند از تعداد مشاهدات بیشتر باشد، اما تعداد متغیرهای توضیحی موثر بر پاسخ نسبت به تعداد مشاهدات، کوچک است. وقتی تعداد متغیرهای توضیحی مدل زیاد باشد، تفسیر مدل مشکل‌تر و هزینه محاسبات افزایش می‌یابد. لذا، شناسایی متغیرهای توضیحی موثر بر پاسخ یا مولفه‌های جمعی غیر صفر در این مدل بسیار مهم است. بدین منظور، ابتدا مولفه‌های جمعی را با استفاده از پایه‌های B - اسپلاین تقریب می‌زنیم. با به کارگیری این تقریب، مساله انتخاب متغیر به انتخاب گروه‌هایی از ضرایب غیر صفر تبدیل می‌شود. سپس از تابع‌های تاوان گروهی برای انتخاب ضرایب غیر صفر استفاده می‌کنیم. این امر معمولاً با مینیم کردن مجموع توان‌های دوم خطا تحت یک شرط محدود کننده انجام می‌شود. مینیم کردن این تابع هدف توانیده مستلزم استفاده از روش‌های بهینه‌سازی است. در این مقاله از الگوریتم کاهش گروهی برای حل مساله مینیم‌سازی فوق استفاده می‌شود. در پایان، عملکرد این الگوریتم تحت سه تابع تاوان مختلف، با مطالعات شبیه‌سازی و تحلیل یک مجموعه داده واقعی بررسی می‌شود.

کلمات کلیدی: انتخاب متغیر گروهی، بهینه‌سازی، داده‌های بعد بالا، مدل جمعی ناپارامتری.

۱ مقدمه

در مطالعه رابطه بین متغیرهای توضیحی و متغیر پاسخ، اطلاعات پیشینی که نشان‌دهنده وجود یک رابطه خطی باشد، به ندرت در دسترس است. بنابراین منطقی به نظر می‌رسد که به جای مدل خطی از مدل‌هایی با انعطاف پذیری بیشتری استفاده کنیم. یکی از این مدل‌ها، مدل جمعی ناپارامتری است که توسط هستی و تیشیرانی در سال ۱۹۹۰ معرفی شد [۱]. در این مدل هیچ فرض محدود کننده قوی برای شکل تابع‌های تعیین کننده اثر متغیرهای توضیحی بر پاسخ در نظر گرفته نمی‌شود؛ بنابراین برای برآورد مولفه‌های جمعی نامعلوم باید از روش‌های ناپارامتری استفاده شود.

^۱ عهده دار مکاتبات

آدرس الکترونیکی: m.kazemi@guilan.ac.ir

از طرفی، ظهور فناوری‌های نوین در دهه اخیر و امکان ذخیره‌سازی داده‌ها در ابعاد بزرگ موجب شده است که مساله مدل‌سازی پدیده‌ها، با تعداد صدها یا هزاران متغیر توضیحی مواجه باشد. اگرچه حضور تعداد زیادی از متغیرهای توضیحی انعطاف‌پذیری مدل را افزایش می‌دهد، اما باعث مشکل شدن تفسیر مدل می‌شود. علاوه بر این، در مدل‌های با بعد بالا، مساله همخطی بسیار شدید است، یعنی هر یک از متغیرهای توضیحی را می‌توان به صورت ترکیب خطی از سایر متغیرها نوشت [۲]. این موضوع منجر به کاهش توان پیشگویی مدل می‌شود. روش پیشنهادی در این حالت، مدل‌سازی با زیرمجموعه‌ای از متغیرهای توضیحی است. در این رهیافت علاقه‌مندیم زیرمجموعه‌ای از متغیرها را طوری انتخاب کنیم که رابطه بین متغیر پاسخ و متغیرهای توضیحی را به خوبی توصیف کند. یافتن زیرمجموعه‌ای مناسب از متغیرهای توضیحی برای ورود به مدل انتخاب متغیر نامیده می‌شود.

برای درک بهتر موضوع انتخاب متغیر، مدل رگرسیون خطی چندگانه $Y = \beta_0 + \beta_1 X_1 + \dots + \beta_p X_p + \varepsilon$ را در نظر بگیرید. در این مدل، برای برآورد متغیر پاسخ Y برحسب متغیرهای توضیحی $X_1, \dots, X_p$ ، به طور معمول از روش معروف مینیمم‌سازی مجموع توان‌های دوم خطا برای یافتن بهینه ضرایب رگرسیونی β_j ها استفاده می‌شود. به عبارت دیگر

$$\hat{\beta} = \arg \min_{\beta} \sum_{i=1}^n (y_i - \beta_0 - \beta_1 x_{i1} - \dots - \beta_p x_{ip})^2 \quad (1)$$

که در آن $\beta = (\beta_0, \beta_1, \dots, \beta_p)'$ بردار ضرایب رگرسیونی است. به منظور انتخاب متغیرهای مهم، باید ضرایب رگرسیونی را طوری برآورد کنیم که برآورد ضرایب متغیرهای توضیحی زاید دقیقاً برابر صفر باشند، یعنی باید مساله مینیمم‌سازی مذکور، به طور مقید انجام شود و بدین وسیله تشخیص متغیرهای مهم و همچنین پیش‌بینی پاسخ در گرو حل مسایل بهینه‌سازی باشد. این قیود بسیار متنوع بوده و برحسب نوع ساختار آن، برنامه‌ریزی خطی، درجه دوم، محدب و غیرمحدب کاربرد پیدا می‌کند. پس از حل این مساله بهینه‌سازی، ضرایب برخی از متغیرهای توضیحی دقیقاً برابر صفر برآورد می‌شوند که منجر به حذف متغیرهای توضیحی متناظر یا به عبارت دیگر، انتخاب متغیر می‌شود. در حالت کلی، برآورد گر کمترین توان‌های دوم توانانیده با مینیمم کردن مجموع توان‌های دوم خطا تحت قید $\sum_{j=1}^p p(|\beta_j|) \leq t$ به دست می‌آید که در آن $p(\cdot)$ تابع توانانیده می‌شود. این مساله مینیمم‌سازی را با

استفاده از لاگرانژ می‌توان به صورت زیر نوشت

$$\hat{\beta} = \arg \min_{\beta} \left\{ \sum_{i=1}^n (y_i - \beta_0 - \sum_{j=1}^p \beta_j x_{ij})^2 + \lambda \sum_{j=1}^p p(|\beta_j|) \right\} \quad (2)$$

که در آن λ پارامتر تنظیم‌کننده نامیده می‌شود. به عنوان نمونه، می‌توان تابع توانان را به شکل $p(|\beta_j|) = |\beta_j|^d$ در نظر گرفت که به آن، توانان بریج گویند [۳]. در ادبیات آماری، این تابع توانان به ازای $d = 1$ موسوم به توانان لاسو^۱ است که در این صورت، حل مساله مینیمم‌سازی (۲) یک مساله بهینه‌سازی محدب است [۴]. در این وضعیت، با

¹ Least absolute shrinkage and selection operator

تغییر مقدار پارامتر λ ممکن است برخی از ضرایب β_j برابر صفر شوند که بدین ترتیب متغیر متناظر آن، یعنی X_j ، از مدل حذف شده و با حذف متغیرهای غیر مرتبط با متغیر پاسخ، انتخاب متغیر انجام می‌شود.

تابع تاوان لاسو با اینکه در بسیاری از مسایل انتخاب متغیر عملکرد خوبی از خود نشان داده، اما دارای برخی محدودیت‌ها است. به عنوان مثال، لاسو برای یک مدل رگرسیونی خطی با p متغیر پیشگو و n مشاهده، حداکثر n متغیر را انتخاب می‌کند. بنابراین اگر تعداد متغیرهای توضیحی معنی‌دار در مدل بیشتر از n باشد، برخی از آن‌ها توسط لاسو انتخاب نمی‌شوند [۵]. برای غلبه بر این محدودیت‌ها، تابع‌های تاوان متعددی به منظور برآورد و انتخاب متغیر معرفی شده‌اند. فن و لی در سال ۲۰۰۱ یک تابع تاوان غیر محدب به نام SCAD^۱ را معرفی کردند و نشان دادند که این تاوان از سه ویژگی مطلوب نااریبی، تنکی و پیوستگی برخوردار است [۶]. سپس زو در سال ۲۰۰۶ نشان داد برآوردگر لاسو از ویژگی پیشگویی^۲ برخوردار نیست [۷]. او برای رفع این مشکل، تاوان لاسوی تطبیقی را معرفی کرد که دارای این ویژگی است. همچنین ژانگ در سال ۲۰۱۰ تابع تاوان غیر محدب دیگری به نام MCP^۳ را معرفی کرد که تاوان لاسو را به عنوان یک حالت خاص در بر می‌گیرد [۸].

در این مقاله، با هدف نشان دادن نقش به‌سزای مسایل بهینه‌سازی در پیش‌بینی دقیق‌تر مدل‌های رگرسیونی، مساله انتخاب متغیرهای مهم را در یک مدل جمعی ناپارامتری با تابع‌های تاوان مختلف در نظر می‌گیریم. ابتدا به دلیل نامعلوم بودن مولفه‌های جمعی ناپارامتری، از پایه‌های B-اسپلاین برای تقریب آنها استفاده می‌کنیم. با تقریب مولفه‌های جمعی توسط پایه‌های B-اسپلاین، مدل جمعی ناپارامتری به یک مدل خطی با ساختار گروهی کاهش می‌یابد. بنابراین، مساله انتخاب متغیر به انتخاب گروه‌هایی از ضرایب غیرصفر تبدیل می‌شود و لذا باید از تابع‌های تاوان گروهی استفاده کرد که قابلیت انتخاب گروهی را دارند. در نتیجه، انتخاب متغیرهای مهم مستلزم حل یک مساله بهینه‌سازی با ساختار گروهی است.

با یک تابع تاوان محدب، مانند لاسو، مجموع توان‌های دوم خطا یک تابع محدب است و لذا برآورد کمترین توان‌های دوم توانیده را می‌توان با روش‌های بهینه‌سازی محدب به‌دست آورد. اما اغلب تابع‌های تاوان که دارای سه ویژگی مطلوب تنکی، نااریبی و پیوستگی می‌باشند، تابع‌هایی غیر محدب بوده و این موضوع باعث می‌شود که تابع هدف (۲) محدب نباشد. محدب نبودن این تابع، مساله مینیمم‌سازی را با مشکل مواجه می‌کند. در این راستا، دو رویکرد کلی وجود دارد: در رویکرد اول، ابتدا تابع تاوان غیر محدب را با یک تابع محدب تقریب زده و سپس مساله بهینه‌سازی غیر محدب را با استفاده از الگوریتم‌های بهینه‌سازی محدب حل می‌کنند. رویکرد دوم، استفاده از الگوریتم‌های کاهش مختصات است. از جمله الگوریتم‌های موجود برای تقریب تاوان‌های غیر محدب می‌توان به الگوریتم‌های تقریب درجه دوم موضعی و تقریب خطی موضعی اشاره کرد ([۶]، [۹]). اما هر یک از این الگوریتم‌ها دارای برخی نقاط ضعف هستند. ایراد وارد بر الگوریتم تقریب درجه دوم موضعی این است که اگر یک متغیر در هر تکرار حذف شود، در تکرار بعدی وارد مدل نمی‌شود و در نتیجه از مدل نهایی حذف می‌شود. بنابراین، این روش مشابه روش انتخاب پیشرو عمل می‌کند. برای غلبه بر این مشکل، زو و لی در سال ۲۰۰۸

¹ Smoothly Clipped Absolute Deviation

² Oracle property

³ Minimax concave penalty

الگوریتم تقریب خطی موضعی را برای تقریب تابع‌های تاوان غیر محدب پیشنهاد دادند و برتری این الگوریتم را نسبت به الگوریتم تقریب درجه دوم موضعی نشان دادند [۹]. اما هیچ یک از این دو الگوریتم در مدل‌های رگرسیونی بعد بالا کارایی لازم را ندارند. در مقابل، رویکرد کاهش مختصات برای برازش مدل‌های رگرسیونی توانیده بعد بالا بسیار مناسب است. تاکنون انواع الگوریتم‌های کاهش مختصات در مدل‌های توانیده مورد استفاده قرار گرفته است که از جمله آنها می‌توان به [۱۴-۱۰] اشاره کرد. این الگوریتم‌ها در ابتدا برای تابع‌های تاوان محدب مانند لاسو معرفی شدند، اما سپس تعمیم‌هایی از آنها برای تاوان‌های SCAD و MCP نیز معرفی شدند. به عنوان نمونه، بره‌نی و هوآنگک در سال ۲۰۱۱ با به‌کارگیری رویکرد کاهش مختصات، الگوریتم تقریب خطی موضعی را بهبود دادند و نشان دادند که چگونه این رویکرد می‌تواند به طور کارا در برازش مدل‌های رگرسیونی با تاوان‌های غیرمحدب مانند SCAD و MCP موثر بوده و در مسایل با بعد بسیار بالا نیز قابل استفاده است [۱۵]. سپس بره‌نی و هوآنگک در سال ۲۰۱۵ این الگوریتم را به حالت گروهی تعمیم داده و الگوریتم کاهش گروهی را برای حل مسایل بهینه‌سازی در مدل‌های با ساختار گروهی معرفی کردند [۱۶].

مهم‌ترین مزیت الگوریتم ارایه‌شده در [۱۶] کارایی محاسباتی آن است، زیرا اجرای آن نیاز به هیچ بهینه‌سازی عددی پیچیده یا محاسبات ماتریسی دشوار مانند وارون‌گیری ندارد و تنها شامل تعداد اندکی عملیات حسابی ساده است. لازم به ذکر است که پیش از این الگوریتم دیگری توسط شی در سال ۲۰۱۲ برای انتخاب متغیر گروهی ارایه شده بود [۱۷] که مراحل آن بسیار شبیه الگوریتم کاهش گروهی در [۱۶] است. با این وجود، برتری الگوریتم ارایه شده در [۱۶] این است که برای حالت تک گروهی، این الگوریتم در هر مرحله جواب‌های دقیقی را تولید می‌کند. در مقابل، روش ارایه‌شده در شی [۱۷] نیاز به چندین تکرار برای همگرا شدن به همان جواب دارد، به عبارت دیگر، این الگوریتم پایدار نیست.

در ادامه، ساختار مقاله به صورت زیر است: در بخش ۲ مدل جمعی ناپارامتری معرفی می‌شود. در بخش ۳ به موضوع انتخاب متغیر گروهی و معرفی روش‌های متداول انتخاب گروهی پرداخته می‌شود. در بخش ۴ با معرفی الگوریتم کاهش گروهی به حل مساله بهینه‌سازی برای انتخاب متغیر می‌پردازیم. در بخش ۵ با مطالعات عددی، عملکرد الگوریتم مذکور در انتخاب متغیرهای مهم را تحت تابع‌های تاوان مختلف بررسی می‌کنیم. در بخش ۶ تحلیل یک مجموعه داده واقعی ارایه می‌شود.

۲ مدل جمعی ناپارامتری

مدل جمعی ناپارامتری به صورت زیر تعریف می‌شود:

$$Y = \sum_{j=1}^p f_j(X_j) + \varepsilon \quad (3)$$

که Y متغیر پاسخ مرکزی شده، $\mathbf{X} = (X_1, \dots, X_p)^T$ متغیرهای توضیحی، f_j ها توابع یک متغیره نامعلوم هموار با $E[f_j(X_j)] = 0$ و ε خطای تصادفی با میانگین صفر است. همچنین فرض می‌شود که برخی از f_j ها برابر صفر هستند. هدف، پیدا کردن متغیرهای توضیحی مهم یا به عبارت دیگر، یافتن مولفه‌های غیر صفر در مدل فوق و

برآورد آنها است. از آنجا که برخلاف مدل خطی، تابع‌های f_j نامعلوم هستند، ابتدا باید این تابع‌ها را با استفاده از پایه‌های B-اسپلاین تقریب بزنیم. در این مقاله، از اسپلاین‌های چندجمله‌ای برای تقریب f_j ها استفاده می‌شود. بدون از دست دادن کلیت، فرض می‌کنیم که X_j مقادیر خود را در بازه $[0, 1]$ اختیار می‌کند و $\tau_0 = 0 < \tau_1 < \dots < \tau_{K'} < 1 = \tau_{K'+1}$ افزایش از $[0, 1]$ به زیر بازه‌های (τ_k, τ_{k+1}) با K' گره داخلی با فواصل یکسان باشد. یک اسپلاین چندجمله‌ای از درجه q با گره‌های $\tau_k, k = 1, \dots, K'$ ، یک چندجمله‌ای تکه‌ای از درجه $q-1$ است که مشتق آن تا مرتبه $q-2$ موجود و پیوسته باشد. مجموعه اسپلاین‌ها با یک دنباله ثابت از گره‌ها، دارای پایه‌های B-اسپلاین نرمال شده $\{B_1(x), B_2(x), \dots, B_{K'+q}(x)\}$ می‌باشند. برای اطلاعات بیشتر درباره توابع B-اسپلاین به [۱۸] و [۱۹] مراجعه کنید. بنابراین می‌توانیم مولفه‌های f_j را به صورت زیر تقریب بزنیم

$$f_j(x) \approx \sum_{k=1}^K \beta_{jk} B_{jk}(x) = \boldsymbol{\beta}_j^T \mathbf{B}_j(x), \quad j = 1, \dots, p$$

که $\mathbf{B}_j(x) = (B_{j1}(x), \dots, B_{jK}(x))^T$ پایه‌های B-اسپلاین، $\boldsymbol{\beta}_j = (\beta_{j1}, \dots, \beta_{jK})^T$ بردار ضرایب نامعلوم و $K = K' + q$ تعداد پایه‌ها می‌باشند. فرض کنید $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n)$ ، n مشاهده از $(\mathbf{X}, Y)$ باشند که $\mathbf{a} \in \mathbb{R}^m$ باشد. بنابراین، تابع هدف $\|\mathbf{a}\| = \left(\sum_{i=1}^m |a_i|^2 \right)^{\frac{1}{2}} = (\mathbf{a}'\mathbf{a})^{\frac{1}{2}}$ و $\mathbf{x}_i = (x_{i1}, \dots, x_{ip})^T$ توانیده را می‌توانیم به صورت زیر بنویسیم

$$Q(\boldsymbol{\beta}) = \frac{1}{2n} \sum_{i=1}^n \left(y_i - \sum_{j=1}^p \boldsymbol{\beta}_j^T \mathbf{B}_j(x_{ij}) \right)^2 + \sum_{j=1}^p p_\lambda(\|\boldsymbol{\beta}_j\|) \quad (۴)$$

که $p_\lambda(\cdot)$ تابع جریمه و λ پارامتر تنظیم کننده است که میزان انقباض ضرایب رگرسیونی را کنترل می‌کند. با استفاده از نمایش ماتریسی می‌توان تابع $Q(\boldsymbol{\beta})$ را به صورت زیر ساده کرد

$$Q(\boldsymbol{\beta}) = \frac{1}{2n} \|\mathbf{y} - \mathbf{Z}\boldsymbol{\beta}\|^2 + \sum_{j=1}^p p_\lambda(\|\boldsymbol{\beta}_j\|) \quad (۵)$$

که $\mathbf{Z}_{(n \times pK)} = (\mathbf{Z}_1, \dots, \mathbf{Z}_p)$ ، $\boldsymbol{\beta} = (\boldsymbol{\beta}_1^T, \dots, \boldsymbol{\beta}_p^T)^T$ ، $\mathbf{y} = (y_1, \dots, y_n)^T$ است و $\mathbf{Z}_j$ به صورت زیر است:

$$\mathbf{Z}_j = \begin{pmatrix} B_{j1}(x_{1j}) & B_{j2}(x_{1j}) & \dots & B_{jK}(x_{1j}) \\ B_{j1}(x_{2j}) & B_{j2}(x_{2j}) & \dots & B_{jK}(x_{2j}) \\ \vdots & \vdots & \ddots & \vdots \\ B_{j1}(x_{nj}) & B_{j2}(x_{nj}) & \dots & B_{jK}(x_{nj}) \end{pmatrix}$$

با تقریب f_j ها توسط پایه‌های B-اسپلاین، ماتریس طرح $\mathbf{Z}$ ساختار گروهی پیدا می‌کند، یعنی K تابع پایه‌ای اول برای تقریب f_1 ، K تابع پایه‌ای دوم برای تقریب f_2 و به همین ترتیب، K تابع پایه‌ای p ام برای تقریب f_p به کار می‌روند. بنابراین برای برقراری $f_j(x_j) = 0$ ، باید همه مؤلفه‌های بردار $\boldsymbol{\beta}_j = (\beta_{j1}, \dots, \beta_{jK})^T$ صفر شوند، یعنی باید از روش‌های انتخاب متغیر گروهی استفاده کرد. از جمله تابع‌های توان رایج برای انتخاب متغیر

گروهی می‌توان تاوان لاسوی گروهی، تاوان SCAD گروهی و MCP گروهی را نام برد. این تابع‌های تاوان در [۲۳-۲۰، ۱۱] مورد بررسی قرار گرفته‌اند. همچنین یک روش بیزی موسوم به SSSL¹ برای انتخاب متغیر گروهی براساس تابع تاوان لاسوی گروهی توسط بای و همکاران در سال ۲۰۲۲ ارایه شده است [۲۴]. در بخش بعد ابتدا به معرفی این تابع‌ها و سپس الگوریتم کاهش گروهی برای حل مساله بهینه‌سازی می‌پردازیم.

۳ روش‌های انتخاب متغیر گروهی

در رویکرد انتخاب متغیر گروهی، هنگام برآورد و انتخاب متغیرهای مهم، هر گروه از متغیرها را به عنوان یک واحد در نظر می‌گیرند و این عمل فرآیند انتخاب متغیر را تسهیل می‌کند. در نتیجه متغیرهای یک گروه، یا همگی از مدل حذف شده و یا همگی در مدل باقی می‌مانند. ساختار گروهی متغیرها می‌تواند دلایل و اهداف متفاوتی داشته باشد. به عنوان مثال، در تحلیل داده‌های بیان ژنی، ژن‌های متعلق به یک مسیر بیولوژیکی یکسان می‌توانند در یک گروه قرار بگیرند. همچنین در مطالعات پیوند ژنتیکی، نشانگرهای ژنتیکی از ژن یکسان می‌توانند به عنوان یک گروه در نظر گرفته شوند. مطلوب است ساختار گروهی را در تحلیل این نوع داده‌ها مورد توجه قرار دهیم.

ابتدا تابع تاوان لاسوی گروهی توسط یوآن و لین [۱۱] به صورت زیر معرفی شد

$$P(\beta) = \lambda \sum_{j=1}^p c_j \|\beta_j\|$$

که λ پارامتر جریمه و c_j ‌ها برای تعدیل اندازه گروه‌ها استفاده می‌شوند. یک انتخاب مناسب $c_j = \sqrt{d_j}$ است که d_j اندازه گروه j ام می‌باشد. با توجه به یکسان بودن اندازه تمام گروه‌ها، $c_j = \sqrt{K}$ در نظر گرفته شده است. همان‌طور که در بخش ۱ بیان شد، تاوان لاسو از ویژگی پیشگویی برخوردار نیست. دلیل این مشکل آن است که تاوان لاسو برای تمام ضرایب مدل میزان انقباض یکسانی را در نظر می‌گیرد. این امر باعث می‌شود که مدل انتخاب‌شده توسط لاسو معمولاً بزرگ‌تر از مدل واقعی باشد، زیرا قادر به تشخیص ضرایب کوچک از ضرایب صفر نیست. تابع تاوان لاسوی گروهی هم رفتاری شبیه لاسو داد و گروه‌های بی‌اهمیت زیادی را وارد مدل می‌کند. لذا، منطقی به نظر می‌رسد که از تاوان‌های مناسب‌تری مانند SCAD گروهی و MCP گروهی استفاده کنیم. این دو تابع تاوان همه ویژگی‌های مطلوب یک تابع تاوان مناسب را دارند. در صورت استفاده از تاوان‌های SCAD گروهی و MCP گروهی، عبارت تاوان در رابطه (۵) دارای شکل کلی زیر است:

$$P(\beta) = \sum_{j=1}^p \rho_{\gamma}(\|\beta_j\|; c_j \lambda)$$

که در آن، برای تابع تاوان MCP گروهی، $\gamma > 1$ و $\rho_{\gamma}(x; \lambda) = \lambda \int_0^{|x|} (1 - \frac{t}{\lambda \gamma})_+ dt$ و $x_+ = xI\{x \geq 0\}$ هنگامی که $\gamma \rightarrow \infty$ ، MCP گروهی به لاسوی گروهی همگرا می‌شود. همچنین، برای تابع تاوان SCAD گروهی $\gamma > 2$ و $\rho_{\gamma}(x; \lambda) = \lambda \int_0^{|x|} \min\{1, \frac{(\gamma - \frac{t}{\lambda})_+}{\gamma - 1}\} dt$ است.

¹ Spike-and-Slab Group Lasso

۴ حل مساله بهینه‌سازی

برای مینیمم کردن تابع هدف (۵)، ابتدا یک حالت خاص از این مساله را در نظر می‌گیریم، یعنی فرض می‌کنیم که ماتریس‌های Z_j متعامد باشند. تحت این فرض برآوردگرهای صریحی برای ضرایب به‌دست می‌آید. در ادامه با معرفی الگوریتم کاهش گروهی، از آنجا که این الگوریتم هر بار تابع هدف را نسبت به یک گروه تکی بهینه می‌کند، این برآوردگرها در الگوریتم کاهش گروهی برای بهینه‌سازی استفاده خواهند شد.

۴-۱ گروه‌های متعامد

در این بخش، فرض می‌کنیم که گروه‌ها متعامد هستند به طوری که $Z_j'Z_k = 0, k \neq j$ ، $Z_j'Z_j = nI_K$. در این حالت مساله به برآورد p مدل رگرسیونی K متغیره تک گروهی به‌صورت $y = Z_j\beta_j + \varepsilon, j = 1, \dots, p$ می‌یابد. فرض کنید $\tilde{\beta}_j = \frac{1}{n}Z_j'y$ برآورد کمترین توان‌های دوم معمولی برای بردار ضرایب β_j باشد. بدون از دست دادن کلیت فرض می‌کنیم $c_j = 1$ از طرفی چون $Z_j'Z_j = nI_K$ ، بنابراین

$$\frac{1}{n}\|y - Z_j\beta_j\|^2 = \|\tilde{\beta}_j - \beta_j\|^2 + \frac{1}{n}\|y\|^2 - \|\tilde{\beta}_j\|^2.$$

در نتیجه، تابع هدف توانیده برای مدل $y = Z_j\beta_j + \varepsilon$ را می‌توان به‌صورت زیر نوشت

$$\frac{1}{2n}\|y - Z_j\beta_j\|^2 + \rho(\|\beta_j\|; \lambda, \gamma) = \frac{1}{2}\|\tilde{\beta}_j - \beta_j\|^2 + \rho(\|\beta_j\|; \lambda, \gamma).$$

برای برآورد کمترین توان‌های دوم توانیده بردار β_j ، ابتدا فرض کنید $H(\tilde{\beta}_j; \lambda) = u(\|\tilde{\beta}_j\|; \lambda) \frac{\tilde{\beta}_j}{\|\tilde{\beta}_j\|}$ که در

آن $u(.,.)$ به صورت زیر تعریف می‌شود

$$u(t; \lambda) = \begin{cases} t - \lambda & t > \lambda \\ 0 & |t| \leq \lambda \\ t + \lambda & t < -\lambda. \end{cases}$$

اکنون با در نظر گرفتن $\rho(.,.)$ برابر تابع‌های تاوان لاسو، SCAD و MCP، می‌توان برآوردگرهای لاسوی گروهی، SCAD گروهی و MCP گروهی را به‌صورت زیر به‌دست آورد

$$\text{Group LASSO: } \hat{\beta}_j = H(\tilde{\beta}_j; \lambda) \quad (۶)$$

$$\text{Group SCAD: } \hat{\beta}_j = F_s(\tilde{\beta}_j; \lambda, \gamma) = \begin{cases} \frac{\gamma}{\gamma-1} H(\tilde{\beta}_j; \lambda) & \|\tilde{\beta}_j\| \leq \gamma\lambda \\ \tilde{\beta}_j & \|\tilde{\beta}_j\| > \gamma\lambda \end{cases} \quad (۷)$$

$$\text{Group MCP: } \hat{\beta}_j = F_M(\tilde{\beta}_j; \lambda, \gamma) = \begin{cases} H(\tilde{\beta}_j; \lambda) & \|\tilde{\beta}_j\| \leq \gamma \\ \frac{\gamma - 1}{\gamma - 2} H(\tilde{\beta}_j; \frac{\gamma \lambda}{\gamma - 1}), & \gamma \lambda < \|\tilde{\beta}_j\| \leq \gamma \lambda \\ \tilde{\beta}_j & \|\tilde{\beta}_j\| > \gamma \lambda. \end{cases} \quad (8)$$

بر آوردگرهای فوق به عنوان بلوک‌های ساختمان الگوریتم کاهش گروهی در بخش ۴-۲ به کار می‌روند.

۴-۲ الگوریتم کاهش گروهی

الگوریتم‌های کاهش هر بار تابع هدف را نسبت به یک پارامتر تکی بهینه می‌کنند و این چرخه به‌طور مکرر برای سایر پارامترها تا رسیدن به همگرایی ادامه می‌یابد. به‌طور مشابه، الگوریتم کاهش گروهی هر بار تابع هدف را نسبت به یک گروه تکی بهینه می‌کند و این عمل بهینه‌سازی برای سایر گروه‌ها تا رسیدن به همگرایی تکرار می‌شود. این الگوریتم به‌ویژه برای برازش مدل‌های لاسوی گروهی، SCAD گروهی و MCP گروهی مناسب می‌باشند، چون همان‌طور که دیدیم، هر سه برآوردگر مذکور در یک مدل تک-گروهی دارای شکل بسته می‌باشند.

هر مرحله از الگوریتم کاهش گروهی شامل بهینه‌سازی مجموع توان‌های دوم خطا نسبت به ضرایب گروه

j می‌باشد. لذا تعریف کنید

$$Q_j(\beta_j) = \frac{1}{2n} \left\| y - \sum_{k \neq j} Z_k \tilde{\beta}_k - Z_j \beta_j \right\|^2 + \rho(\|\beta_j\|; c_j \lambda, \gamma)$$

که $\tilde{\beta}_k$ آخرین مقدار به روز شده β_k است. حال $\tilde{y}_j = \sum_{k \neq j} Z_k \tilde{\beta}_k$ و $\tilde{s}_j = \frac{1}{n} Z_j' (y - \tilde{y}_j)$ را در نظر بگیرید. دقت کنید که $\tilde{y}_j$ همان مقادیر برازش شده بدون در نظر گرفتن گروه j ام و $\tilde{s}_j$ برآورد کمترین توان‌های دوم معمولی β_j پس از تعدیل اثر سایر گروه‌ها روی متغیر پاسخ است. بنابراین، مشابه رگرسیون کمترین توان‌های دوم معمولی، مقدار β_j که مینیمم‌کننده $Q_j(\beta_j)$ باشد، برابر با مقدار β_j ای است که از رگرسیون کردن Z_j روی باقیمانده‌های جزئی، یعنی $y - \tilde{y}_j$ ، به‌دست می‌آید. با توجه به بخش ۴-۱ می‌دانیم که مینیمم‌کننده $Q_j(\beta_j)$ برابر $F(\tilde{s}_j; \lambda, \gamma)$ است که $F(\cdot)$ با توجه به تابع تاوان، یکی از جواب‌های (۶) تا (۸) می‌باشد.

فرض کنید $\tilde{\beta}^{(i)} = (\tilde{\beta}_1^{(i)}, \dots, \tilde{\beta}_p^{(i)})^T$ مقادیر اولیه ضرایب و t نشان‌دهنده تکرار باشد. الگوریتم کاهش گروهی در الگوریتم ۱ خلاصه شده است. در این الگوریتم با قرار دادن $F(\cdot; \lambda, \gamma)$ برابر یکی از جواب‌های $\hat{\beta}_j$ در روابط (۶) تا (۸)، برآوردهای لاسوی گروهی، SCAD گروهی و MCP گروهی را برای بردار ضرایب β در مدل (۵) به‌دست می‌آوریم. در مرحله ۴ الگوریتم ۱، همگرایی زمانی رخ می‌دهد که $\|\tilde{\beta}_j^{(t+1)} - \tilde{\beta}_j^{(t)}\| < \varepsilon$. در این الگوریتم با قرار دادن $F(\cdot; \lambda, \gamma)$ برابر یکی از جواب‌های $\hat{\beta}_j$ در روابط (۶) تا (۸)، برآوردهای لاسوی گروهی، SCAD گروهی و MCP گروهی را برای بردار ضرایب β در مدل (۵) به‌دست می‌آوریم. در مرحله ۴ الگوریتم ۱، همگرایی زمانی رخ می‌دهد که $\|\tilde{\beta}_j^{(t+1)} - \tilde{\beta}_j^{(t)}\| < \varepsilon$.

الگوریتم ۱: الگوریتم کاهش گروهی

- مرحله ۱: قرار دهید $t = 0$ و بردار اولیه باقیمانده‌ها را به صورت $\mathbf{r} = \mathbf{y} - \tilde{\mathbf{y}}$ در نظر بگیرید که

$$\tilde{\mathbf{y}} = \sum_{j=1}^p \mathbf{Z}_j \tilde{\boldsymbol{\beta}}_j^{(0)}$$

- مرحله ۲: برای $j = 1, \dots, p$ محاسبات زیر را انجام دهید:

$$\tilde{\mathbf{s}}_j = \frac{1}{n} \mathbf{Z}_j' \mathbf{r} + \tilde{\boldsymbol{\beta}}_j^{(t)} \quad \text{آ.}$$

$$\tilde{\boldsymbol{\beta}}_j^{(t+1)} \leftarrow F(\tilde{\mathbf{s}}_j; \lambda, \gamma) \quad \text{ب.}$$

$$\mathbf{r} \leftarrow \mathbf{r} - \mathbf{Z}_j (\tilde{\boldsymbol{\beta}}_j^{(t+1)} - \tilde{\boldsymbol{\beta}}_j^{(t)}) \quad \text{پ.}$$

- مرحله ۳: به روزرسانی $t \leftarrow t + 1$.

- مرحله ۴: مراحل ۲ و ۳ را تا رسیدن به همگرایی تکرار کنید.

این الگوریتم از دو مزیت عمده برخوردار است: الف) از نظر هزینه محاسباتی بسیار کارا است. ب) چون در هر مرحله به روزرسانی، تابع هدف نسبت به $\boldsymbol{\beta}_j$ مینیمم می‌شود، این الگوریتم دارای ویژگی نزولی^۱ است؛ یعنی الگوریتم با هر تکرار، تابع هدف را کاهش می‌دهد. از طرفی، برهنی و هوآنگ [۱۶] نشان دادند که $Q_j(\boldsymbol{\beta}_j)$ یک تابع اکیدا محدب نسبت به $\boldsymbol{\beta}_j$ است. بنابراین، ویژگی نزولی این الگوریتم با توجه به اکیدا محدب بودن $Q_j(\boldsymbol{\beta}_j)$ ، همگرایی الگوریتم را نتیجه می‌دهد.

لازم به ذکر است که در صورت استفاده از تاوان لاسوی گروهی، که تابع هدف تاوانیده (۵) محدب است، قضیه فوق همگرایی الگوریتم به مینیمم سراسری را ثابت می‌کند. اما برای SCAD گروهی و MCP گروهی که تابع هدف (۵) مجموع دو تابع محدب و غیر محدب است، همگرایی به یک مینیمم موضعی امکان‌پذیر است.

۵ مطالعات شبیه‌سازی

در این بخش، با استفاده از شبیه‌سازی، عملکرد الگوریتم کاهش گروهی را در حل مساله بهینه‌سازی مربوط به انتخاب متغیرهای مهم تحت سه تابع تاوان مختلف بررسی می‌کنیم. در این مثال فرض می‌کنیم که حجم نمونه برابر $n = 200$ و تعداد متغیرهای توضیحی برابر $p = 100, 500, 1000$ باشند. همچنین X_i ها متغیرهای تصادفی مستقل یکنواخت (۰,۱) و خطاها دارای توزیع نرمال استاندارد فرض شده‌اند. دو مدل جمعی ناپارامتری زیر را در نظر بگیرید:

$$\bullet \quad \text{مدل (۱): داده‌ها از مدل } Y = \sum_{j=1}^p f_j(X_j) + \sqrt{1.74} \varepsilon \text{ تولید شده‌اند که در آن}$$

$$f_1(x) = 5x, \quad f_7(x) = 3(2x-1)^2, \quad f_9(x) = \frac{4 \sin(2\pi x)}{2 - \sin(2\pi x)},$$

$$f_9(x) = 6(0.1 \sin(2\pi x) + 0.2 \cos(2\pi x) + 0.3 \sin(2\pi x)^2 + 0.4 \cos(2\pi x)^2 + 0.5 \sin(2\pi x)^3).$$

¹ Descent property

• مدل (۲): داده ها از مدل $Y = \sum_{j=1}^p f_j(X_j) + \varepsilon$ تولید شده‌اند که در آن

$$\begin{aligned} f_1(x) &= \frac{2(e^{-1.5x} - e^{-1.0})}{1 - e^{-1.0}} - 1, & f_2(x) &= \frac{-2(e^{-1.5x} - e^{-1.0})}{1 - e^{-1.0}} + 1, & f_3(x) &= 2x^2 - 1, \\ f_4(x) &= -2x^2 + 1, & f_5(x) &= 8(x - 0.5)^2 - 1, & f_6(x) &= -8(x - 0.5)^2 + 1. \end{aligned}$$

در مدل (۱)، متغیرهای $(X_1, X_2, X_3, X_4, X_5, X_6)$ و در مدل (۲)، متغیرهای $(X_1, X_2, X_3, X_4, X_5, X_6)$ متغیرهای توضیحی موثر بر پاسخ هستند و بقیه متغیرهای توضیحی در پیشگویی پاسخ هیچ کمکی نمی‌کنند. هدف، پیدا کردن متغیرهای توضیحی مهم، یا بطور معادل شناسایی مولفه‌های جمعی غیر صفر است. ابتدا برای تقریب f_j ها، از پایه‌های B-اسپلاین مکعبی با ۴ گره، با فاصله‌های یکسان، استفاده می‌شود. برای تابع تاوان MCP گروهی مقدار پارامتر $\gamma = 3$ ، برای SCAD گروهی مقدار $\gamma = 4$ و برای انتخاب پارامتر λ از روش اعتبارسنجی متقابل تعمیم یافته (GCV) استفاده شده است. معیار نتایج شبیه‌سازی پس از ۵۰۰ بار تکرار در جدول ۱ خلاصه شده است. در این جدول، NS میانگین تعداد مولفه‌های غیر صفر انتخاب شده، NST میانگین تعداد مولفه‌های غیر صفر انتخاب شده که در مدل واقعی غیر صفر هستند، IN درصد دفعاتی که همه مولفه‌های غیر صفر در مدل واقعی، توسط این روش به عنوان مولفه‌های غیر صفر تشخیص داده شده‌اند و RMSE ریشه میانگین مربع خطا در ۵۰۰ بار تکرار، گزارش شده است. مقدار RMSE از رابطه زیر محاسبه می‌شود:

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (\mu_i - \hat{\mu}_i)^2}$$

که در آن μ_i میانگین واقعی و $\hat{\mu}_i$ میانگین برآورد شده Y_i به شرط معلوم بودن X_i می‌باشد. در جدول ۱ مشاهده می‌شود که الگوریتم کاهش گروهی در حل مساله بهینه‌سازی (۵) تحت تابع تاوان لاسوی گروهی، SCAD گروهی و MCP گروهی عملکرد مناسبی دارد. در مدل ۱، تحت هر یک از سه تابع تاوان، در ۱۰۰٪ دفعات متغیرهای مهم (X_1, X_2, X_3, X_4) به درستی انتخاب می‌شوند. با این وجود، تابع تاوان لاسویگروهی در مقایسه با دو تابع تاوان دیگر، تمایل به انتخاب مدل‌های بزرگ‌تر دارد. براساس معیار RMSE می‌توان دید که مقدار خطای مدل برازش شده با تابع تاوان لاسوی گروهی در مقایسه با دو مدل دیگر اندکی بیشتر، و خطای متناظر با تابع تاوان MCP گروهی نسبت به دو مدل دیگر کمتر است، اما تفاوت خطای این مدل‌ها بسیار ناچیز است.

در مدل (۲)، تحت هر یک از تابع‌های تاوان، متوسط تعداد متغیرهای مهم انتخاب شده نزدیک به مقدار واقعی، یعنی ۶ است. در این مدل نیز تابع تاوان لاسویگروهی گرایش به انتخاب مدل‌های با اندازه بزرگ‌تر را دارد و در مقابل، تابع تاوان MCP گروهی مدل کوچک‌تری را نتیجه می‌دهد. اگرچه برای $p = 100$ مناسبیت مدل‌های برازش شده با هریک از سه تابع تاوان را می‌توان نتیجه گرفت، اما با افزایش p ، کفایت مدل جمعی ناپارامتری برازش شده به کمک الگوریتم کاهش گروهی کاهش می‌یابد.

جدول ۱. نتایج شبیه‌سازی برای بررسی عملکرد الگوریتم کاهش گروهی در برازش مدل جمعی ناپارامتری با تابع‌های تاوان لاسوی گروهی، SCAD گروهی و MCP گروهی. NS میانگین تعداد مولفه‌های غیر صفر انتخاب شده، NST میانگین تعداد مولفه‌های غیر صفر انتخاب شده که به درستی انتخاب شده‌اند، IN درصد دفعاتی که همه مولفه‌های غیر صفر در مدل واقعی، به درستی تشخیص داده شده‌اند و RMSE خطای مدل در ۵۰۰ بار تکرار. (اعداد داخل پرانتز انحراف معیار را نشان می‌دهند.)

مدل	p	LASSO				SCAD				MCP			
		NS	NST	IN	RMSE	NS	NST	IN	RMSE	NS	NST	IN	RMSE
مدل ۱	۱۰۰	۲۳/۶۲ (۹/۲۴)	۴ (۰)	۱۰۰	۵/۱۳	۱۳/۷۳ (۷/۷۷)	۴ (۰)	۱۰۰	۵/۰۷	۵/۷۴ (۲/۳۲)	۴ (۰)	۱۰۰	۵/۹۰
	۵۰۰	۳۵/۴۶ (۱۶/۸۵)	۴ (۰)	۱۰۰	۵/۱۸	۲۱/۰۵ (۱۳/۷۳)	۴ (۰)	۱۰۰	۵/۱۰	۷/۰۲ (۳/۵۹)	۴ (۰)	۱۰۰	۵/۰۸
	۱۰۰۰	۴۰/۸۲ (۱۷/۵۸)	۴ (۰)	۱۰۰	۵/۱۵	۳۲/۳۷ (۱۵/۷۵)	۴ (۰)	۱۰۰	۵/۰۹	۷/۶۸ (۳/۷۰)	۴ (۰)	۱۰۰	۵/۰۷
مدل ۲	۱۰۰	۲۷/۹۷ (۱۰/۳۷)	۵/۹۷ (۰/۱۸)	۹۷	۰/۷۴	۲۵/۶۷ (۶/۴۷)	۵/۹۷ (۰/۱۷)	۹۷	۰/۶۴	۱۰/۵ (۲/۶۴)	۵/۹۱ (۰/۳۱)	۹۱	۰/۶۱
	۵۰۰	۴۱/۰۹ (۱۷/۳۷)	۵/۷۷ (۰/۴۵)	۷۹	۰/۸۳	۴۵/۱۸ (۸/۶۶)	۵/۸۹ (۰/۳۳)	۹۱	۰/۷۱	۱۲/۸۵ (۳/۳۷)	۵/۶۳ (۰/۵۷)	۶۸	۰/۶۷
	۱۰۰۰	۴۵/۸۶ (۲۲/۴۰)	۵/۶۰ (۰/۶۴)	۶۷	۰/۹۰	۵۰/۵۱ (۱۱/۰۶)	۵/۷۴ (۰/۵۲)	۷۹	۰/۷۶	۱۴/۱۷ (۳/۲۵)	۵/۴۰ (۰/۷۷)	۵۵	۰/۷۱

۶ مثال کاربردی

در این بخش، برای بررسی کارایی الگوریتم کاهش گروهی در حل مساله بهینه‌سازی مدل‌بندی داده‌های بعد بالا، مجموعه داده‌های مربوط به تولید ریوفلاوین (ویتامین B۲) در باسیلوس سوبتیلیس را در نظر می‌گیریم که توسط بالمن و همکاران در سال ۲۰۱۴ مورد استفاده قرار گرفته [۲۵] و در بسته "hdi" نرم افزار R در دسترس است. در این داده‌ها، لگاریتم نرخ تولید ریوفلاوین به عنوان متغیر پاسخ و لگاریتم سطح بیان $p = 4088$ به عنوان متغیرهای توضیحی در نظر گرفته می‌شوند. این مجموعه داده شامل اطلاعات مربوط به $n = 71$ است. هدف، پیدا کردن ژن‌های مؤثر در پیش‌بینی نرخ تولید ریوفلاوین است.

برای مدل‌سازی رابطه بین نرخ تولید ریوفلاوین و ۴۰۸۸ ژن، از مدل جمعی ناپارامتری (۳) استفاده می‌شود. برای برازش این مدل، تابع‌های تاوان لاسوی گروهی، SCAD گروهی و MCP گروهی را به کار می‌بریم. همچنین متغیرهای توضیحی را طوری مقیاس‌بندی می‌کنیم که مقادیر آنها بین ۰ و ۱ باشند و از اسپلاین مکعبی با ۴ گره برای برآورد مؤلفه‌های جمعی استفاده می‌شود. جدول ۲ تعداد ژن‌های انتخاب شده، مجموع توان‌های دوم باقیمانده‌ها (RSS) و ضریب تعیین (R^2) متناظر با هر روش را نشان می‌دهد که این مقادیر، به ترتیب، به صورت

$$RSS = \sum_{i=1}^n (y_i - \hat{y}_i)^2 \quad \text{و} \quad R^2 = 1 - \frac{RSS}{S_{yy}}$$

محاسبه می‌شوند. برای یافتن مقدار بهینه پارامتر λ از معیارهای اعتبارسنجی متقابل تعمیم یافته (GCV)، معیار اطلاع بیزی (BIC) و معیار اطلاع آکائیک (AIC) استفاده شده است. این معیارها به صورت زیر محاسبه می‌شوند:

$$AIC = L + 2df, \quad BIC = 2L + \log(n)df, \quad GCV = \frac{2L}{\left(1 - \frac{df}{n}\right)^2},$$

که L مجموع توان‌های دوم خطا (یا لگاریتم تابع درستنمایی)، n تعداد مشاهدات و df تعداد ضرایب غیرصفر مدل به ازای λ داده شده است

با توجه به جدول ۲، برای تابع تاوان MCP، مقدار بهینه λ براساس معیارهای مختلف، مدل یکسانی را نتیجه می‌دهد. ژن‌های انتخاب شده در این مدل عبارت از YWRE_at، YVFK_at، YURN_at، YTGD_at، YJZB_at، YFHI_at، YCKE_at، XKDI_at هستند. اگرچه مقادیر RSS و R^2 مناسب هر یک از مدل‌های برازش شده با هر سه تابع تاوان را تایید می‌کنند، اما تابع تاوان لاسوی گروهی تعداد ژن‌های بیشتری را انتخاب می‌کند و تابع تاوان MCP گروهی با انتخاب ۸ ژن، مدل مقرون به صرفه‌تری را نتیجه می‌دهد. به طور خلاصه، با توجه به اندازه مدل و معیارهای نیکویی برازش، تابع تاوان MCP گروهی مدل بهتری را به داده‌ها برازش می‌دهد.

جدول ۲. نتایج تحلیل داده‌های واقعی. NS تعداد ژن‌های انتخاب شده، RSS مجموع توان‌های دوم باقیمانده‌ها، R^2 ضریب تعیین

	MCP			SCAD			LASSO		
	NS	RSS	R^2	NS	RSS	R^2	NS	RSS	R^2
GCV	۸	۰/۱۰۲	۹۹/۸۲	۲۱	۰/۳۲۸	۹۹/۴۴	۳۴	۱/۵۷	۹۷/۳۴
BIC	۸	۰/۱۰۲	۹۹/۸۲	۱۹	۰/۴۶۶	۹۹/۲۱	۱۴	۱۲/۴۰	۷۹/۰۸
AIC	۸	۰/۱۰۲	۹۹/۸۲	۲۱	۰/۳۲۸	۹۹/۴۴	۳۸	۰/۷۳۹	۹۸/۷۵

۷ بحث و نتیجه‌گیری

در این مقاله، موضوع برآورد و انتخاب متغیر در مدل جمعی ناپارامتری بعد بالا با استفاده از رگرسیون توانانیده مورد بررسی قرار گرفت. تحقق این امر مستلزم استفاده از روش‌های بهینه‌سازی برای مینیمم کردن مجموع توان‌های دوم خطای توانانیده است. در این راستا، از الگوریتم کاهش گروهی برای حل مساله بهینه‌سازی استفاده شد. در نهایت عملکرد این الگوریتم در قالب یک مثال شبیه‌سازی و تحلیل یک مجموعه داده واقعی با سه تابع تاوان لاسوی گروهی، SCAD گروهی و MCP گروهی مورد ارزیابی قرار گرفت. نتایج شبیه‌سازی عملکرد مناسب الگوریتم کاهش را در حل این مساله بهینه‌سازی تحت سه تابع تاوان مذکور نشان می‌دهد. در خصوص مقایسه مدل‌های برازش داده شده با این سه تابع تاوان گروهی، اگرچه اختلاف خطای این مدل‌ها ناچیز است، اما تابع تاوان لاسوی گروهی متغیرهای توضیحی بیشتری را وارد مدل می‌کند. این در حالی است که تاوان‌های SCAD گروهی و MCP گروهی منتج به مدل‌های کوچک‌تری می‌شوند. لذا می‌توان نتیجه گرفت که مدل جمعی ناپارامتری برازش شده با تابع تاوان MCP گروهی و حل مساله بهینه‌سازی به کمک الگوریتم کاهش گروهی، با انتخاب متغیرهای توضیحی کمتر، مدل مقرون به صرفه‌تری را نتیجه می‌دهد.

همان‌طور که در مثال شبیه‌سازی دیدیم، با افزایش تعداد متغیرهای توضیحی p نسبت به اندازه نمونه n ، عملکرد روش‌های پیشنهادی تحت تاثیر قرار می‌گیرد. هنگامی که p نسبت به n بسیار بزرگ باشد، توصیه می‌شود

که ابتدا با استفاده از یک روش غربالگری مناسب بعد مساله را کاهش داده و سپس از رگرسیون توانیده برای برآورد و انتخاب متغیر استفاده شود. برای اطلاعات بیشتر در مورد انواع روش‌های غربالگری به فن [۲۶-۳۲] مراجعه شود.

تقدیر و تشکر

نویسنده مقاله از داوران محترم که نظرات ارزشمند ایشان باعث بهبود مطالب ارایه شده در این مقاله گردید، کمال تشکر و قدردانی را دارد. این پژوهش در قالب طرح پژوهشی به شماره ۸۲۰۵۴-۱۵ پ و با استفاده از اعتبارات پژوهش و فناوری دانشگاه گیلان انجام شده است.

منابع

- [1] Hastie, T., Tibshirani, R. (1990). Generalized Additive Models, Chapman and Hall, New York.
- [2] James, G., Witten, D., Hastie, T., Tibshirani, R. (2013). An Introduction to Statistical Learning, New York: springer.
- [3] Frank, L. E., Friedman, J. H. (1993). A statistical view of some chemometrics regression tools, *Technometrics*, 35(2), 109-135.
- [4] Tibshirani, R. (1996). Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society. Series B (Methodological)*, 267-288.
- [5] Zou, H., Hastie, T. (2005). Regularization and variable selection via the elastic net. *Journal of the Royal Statistical Society, Series B*, 67(2), 301-320.
- [6] Fan, J., Li, R. (2001). Variable selection via nonconcave penalized likelihood and its oracle properties. *Journal of American Statistical Association*, 96, 1348-1360.
- [7] Zou, H. (2006). The adaptive lasso and its oracle properties, *Journal of the American Statistical Association*, 101, 1418-1429.
- [8] Zhang, C. H. (2010). Nearly unbiased variable selection under the minimax concave penalty. *The Annals of Statistics*, 38, 894-942.
- [9] Zou, H., Li, R. (2008). One-step sparse estimates in nonconcave penalized likelihood models. *Annals of Statistics*, 36, 1509-1533.
- [10] Fu, W. (1998). Penalized regressions: the bridge versus the lasso. *Journal of Computational and Graphical Statistics*, 7, 397-416.
- [11] Yuan, M., and Lin, Y. (2006). Model selection and estimation in regression with grouped variables. *Journal of the Royal Statistical Society Series B*, 68, 49-67.
- [12] Friedman, J., Hastie, T., Hofling, H., Tibshirani, R. (2007). Pathwise coordinate optimization. *Annals of Applied Statistics*, 1, 302-332.
- [13] Wu, T., Lange, K., (2008). Coordinate descent algorithms for lasso penalized regression. *Annals of Applied Statistics*, 2 224-244.
- [14] Friedman, J., Hastie, T., Tibshirani, R. (2010). Regularization paths for generalized linear models via coordinate descent. *Journal of Statistical Software*, 33 1-22.
- [15] Breheny, P., Huang, J. (2011). Coordinate descent algorithms for nonconvex penalized regression, with applications to biological feature selection. *Annals of Applied Statistics*, 5, 232-253.
- [16] Breheny, P., Huang, J. (2015). Group descent algorithms for nonconvex penalized linear and logistic regression models with grouped predictors. *Statistics and Computing*, 25(2), 173-187.
- [17] She, Y. (2012). An iterative algorithm for fitting nonconvex penalized generalized linear models with grouped predictors. *Computational Statistics & Data Analysis*, 56(10) 2976-2990.
- [18] De Boor, C. (2001). A Practical Guide to Splines. Revised edition. *Applied Mathematical Sciences*. New York: Springer-Verlag.
- [19] Ruppert, D., Wand, M. P., Carroll, R. J. (2003). Semiparametric regression, Cambridge University Press, Cambridge.
- [20] Bakin, S. (1999). Adaptive regression and model selection in data mining problems. PhD Thesis. Australian National University, Canberra.

- [21] Wang, L., Chen, G. and Li, H. (2007). Group SCAD regression analysis for microarray time course gene expression data. *Bioinformatics*, 23(12), 1494-1486.
- [22] Wang, L., Li, H., Huang, J. Z. (2008). Variable selection in nonparametric varying-coefficient models for analysis of repeated measurements. *Journal of the American Statistical Association*, 103, 1556-1569.
- [23] Huang, J., Breheny, P., Ma, S. (2012). A selective review of group selection in high-dimensional models. *Statistical Science*, 27 481-499.
- [24] Bai, R., Moran, G. E., Antonelli, J. L., Chen, Y., Boland, M. R. (2022). Spike-and-slab group lassos for grouped regression and sparse generalized additive models. *Journal of the American Statistical Association*, 117(537) 184-197.
- [25] Bühlmann, P., Kalisch, M., Meier, L. (2014). High-dimensional statistics with a view towards applications in biology. *Annual Review of Statistics and its Applications*, 1, 255-278.
- [26] Fan, J., Lv, J. (2008). Sure independence screening for ultrahigh dimensional feature space. *Journal of Machine Learning Research*, 70 (5), 849-911.
- [27] Li, R.Z., Zhong, W., Zhu, L.P. (2012). Feature screening via distance correlation learning. *Journal of the American Statistical Association*, 107, 1129-1139.
- [28] Li, R., Zhong, W., Zhu, L. (2012). Feature screening via distance correlation learning. *Journal of the American Statistical Association*, 107(499), 1129-1139.
- [29] Kazemi, M., Shahsavani, D., Arashi, M. (2018). A sure independence screening procedure for ultrahigh dimensional partially linear additive models. *Journal of Applied Statistics*, 1-19.
- [30] Kazemi, M. (2020). Partial correlation screening for varying coefficient models. *Journal of Mathematical Modeling*, 8(4), 363-376.
- [31] Kazemi, M. (2024). A Robust variable screening approach with application to gene expression data, *REVSTAT-Statistical Journal*. Retrieved from <https://revstat.ine.pt/index.php/REVSTAT/article/view/613>
- [32] Kazemi, M. (2024). Support vector machine in ultrahigh-dimensional feature space. *Journal of Statistical Computation and Simulation*, 94(3), 517-535.