

الحمد لله
الرحمن الرحيم



دانشکده علوم ریاضی

پایان نامه کارشناسی ارشد آمار ریاضی

رگرسیون جریمه شده در بعد بالا با تابع جریمه لاسوی مشتق پذیر

نگارنده: سمانه رجب‌لو

استاد راهنما

دکتر محمد آرشی

شهریور ۱۳۹۷

شماره:

تاریخ:

باسمه تعالی



مدیریت تحصیلات تکمیلی

فرم شماره (۳) صورتجلسه نهایی دفاع از پایان نامه دوره کارشناسی ارشد

با نام و یاد خداوند متعال، ارزیابی جلسه دفاع از پایان نامه کارشناسی ارشد خانم **سمانه رجبلو** با شماره دانشجویی **۹۵۰۶۵۰۴** رشته **آمار گرایش آمار ریاضی** تحت عنوان: **رگرسیون جریمه شده در بعد بالا** با تابع جریمه لاسوی مشتق پذیر که در تاریخ **۱۳۹۷/۰۶/۱۳** با حضور هیأت محترم داوران در دانشگاه صنعتی شاهرود برگزار گردید به شرح ذیل اعلام می گردد:

☐ مردود ☒ قبول (با درجه: خیالی)			
☐ عملی ☒ نظری نوع تحقیق:			
عضو هیأت داوران	نام و نام خانوادگی	مرتبه علمی	امضاء
۱- استاد راهنمای اول	دکتر محمد آرشی	دانشیار	
۲- استاد راهنمای دوم	-----	-----	-----
۳- استاد مشاور	-----	-----	-----
۴- نماینده تحصیلات تکمیلی	دکتر داود شاهسونی	دانشیار	
۵- استاد ممتحن اول	دکتر حسین باغیشنی	استادیار	
۶- استاد ممتحن دوم	دکتر محمدرضا ربیعی	استادیار	



نام و نام خانوادگی رئیس دانشکده: **دکتر ابراهیم هاشمی**

تاریخ و امضاء و مهر دانشکده:

تبصره: در صورتی که کسی مردود شود حداکثر یکبار دیگر (در مدت مجاز تحصیل) می تواند از پایان نامه خود دفاع نماید (دفاع مجدد نباید زودتر از ۴ ماه برگزار شود).

تقدیم به

آنان که مهر آسمانی‌شان آرام بخش آلام زمینی‌ام است به استوارترین تکیه‌گاهم، دستان
پر مهر پدرم به زیباترین نگاه زندگیم، چشمان زیبای مادرم که هر آنچه آموختم در
مکتب عشق شما آموختم و هر چه بکوشم قطره‌ای از دریای بی‌کران مهربانیتان را
سپاس نتوانم بگویم.
امروز هستی‌ام به امید شماست و فردا کلید باغ بهشت‌م رضای شما باشد که حاصل
تلاشم نسیم گونه غبار خستگی‌تان را بزداید.

بوسه بر دستان پر مهرتان

سپاس‌گزاری

به مصداق (من لم یشکر المخلوق لم یشکر الخالق) بسی شایسته است از استاد فرهیخته و فرزانه جناب آقای دکتر محمد آرشی که با کرامتی چون خورشید، سرزمین دل را روشنی بخشیدند و گلشن سرای علم و دانش را با راهنمایی‌های کارساز و سازنده بارور ساختند تقدیر و تشکر نمایم.

از سرکار خانم دکتر مینا نوروزی‌راد به دلیل یاری‌ها و راهنمایی‌های بی چشم‌داشت ایشان که بسیاری از سختی‌ها را برایم آسان‌تر نمودند، سپاسگزارم. همچنین از پدر و مادر عزیز، دلسوز و مهربانم که آرامش روحی و آسایش فکری فراهم نمودند تا با حمایت‌های همه‌جانبه در محیطی مطلوب، مراتب تحصیلی و نیز پایان‌نامه درسی را به نحو احسن به اتمام برسانم، سپاسگزاری نمایم.

پروردگارا حسن عاقبت، سلامت و سعادت را برای آنان مقدر نما.

سمانه رجب‌لو

شهریور ۱۳۹۷

تعهد نامه

اینجانب **سمانه رجب‌لو** دانشجوی کارشناسی ارشد رشته **آمار** دانشکده **علوم ریاضی** دانشگاه صنعتی شاهرود، نویسنده پایان‌نامه با عنوان **رگرسیون جریمه شده در بعد بالا با تابع جریمه لاسوی مشتق پذیر**، تحت راهنمایی **محمد آرشی** متعهد می‌شوم:

- تحقیقات در این پایان‌نامه توسط اینجانب انجام شده است و از صحت و اصالت برخوردار است.
- در استفاده از نتایج پژوهش‌های دیگر پژوهش‌گران، به مرجع مورد استفاده استناد شده است.
- مطالب این پایان‌نامه، تا کنون توسط خود، یا فرد دیگری برای دریافت هیچ نوع مدرک یا امتیازی در هیچ‌جا ارایه نشده است.
- حقوق معنوی این اثر، به دانشگاه صنعتی شاهرود تعلق دارد، و مقالات مستخرج با نام “دانشگاه صنعتی شاهرود” یا “Shahrood University of Technology” به چاپ خواهد رسید.
- حقوق معنوی تمام افرادی که در به‌دست آوردن نتایج اصلی پایان‌نامه تاثیرگذار بوده‌اند، در مقالات مستخرج از پایان‌نامه رعایت می‌گردد.
- در تمام مراحل انجام این پایان‌نامه، در مواردی که از موجود زنده (یا بافت‌های آنها) استفاده شده است، ضوابط و اصول اخلاقی رعایت شده است.
- در تمام مراحل انجام این پایان‌نامه، در مواردی که به حوزه اطلاعات شخصی افراد دسترسی یافته (یا استفاده شده است)، اصل رازداری و اصول اخلاق انسانی رعایت شده است.

سمانه رجب‌لو

شهریور ۱۳۹۷

مالکیت نتایج و حق نشر

- تمام حقوق معنوی این اثر و محصولات آن (مقالات مستخرج، کتاب، برنامه‌های رایانه‌ای، نرم‌افزارها و تجهیزات ساخته شده) متعلق به دانشگاه صنعتی شاهرود می‌باشد. این مطلب باید به نحو مقتضی، در تولیدات علمی مربوطه ذکر شود.
- استفاده از اطلاعات و نتایج موجود در این پایان‌نامه بدون ذکر منبع مجاز نمی‌باشد.

چکیده

در مسائل تحلیل رگرسیون اگر تعداد متغیرهای توضیحی زیاد باشد، برآورد پارامترها به روش کمترین توان‌های دوم راه مناسبی نبوده و باید از برآوردهای جریمه‌شده استفاده کرد. یکی از موثرترین و کاراترین برآوردهای جریمه‌شده، لاسو است که در صفر مشتق‌پذیر نمی‌باشد و بنابراین توزیع مجانبی آن نرمال نیست. در این پایان‌نامه برآوردهای لاسوی مشتق‌پذیر را مورد بررسی قرار می‌دهیم که دارای توزیع مجانبی نرمال است و خواص خوبی در مقایسه با خیلی از برآوردهای جریمه‌شده دارد. با انجام چند مطالعه شبیه‌سازی و تحلیل داده‌های واقعی، عملکرد این برآوردها و برآوردهای لاسوی مشتق‌پذیر تفاضلی را مورد ارزیابی قرار می‌دهیم.

کلمات کلیدی: برآوردهای لاسوی مشتق‌پذیر، برآوردهای تفاضلی، تابع جریمه، رگرسیون جریمه‌شده در بعد بالا

فهرست مقالات مستخرج از پایان نامه

۱. رجب‌لو، س، آرشی، م. و نوروزی‌راد، م. (۱۳۹۷) کاربرد لاسوی مشتق‌پذیر در تحلیل داده‌های سرطان پرستات، چهاردهمین کنفرانس آمار ایران، شاهرود
۲. Rajabloo, S., Arashi, M. and Norouzirad, M. (2018) Differenced-based differentiable linear models, in progress.

پیش گفتار

در تحلیل مسائل رگرسیونی چنانچه تعداد متغیرهای توضیحی (p) زیاد باشد، تفسیر مدل نهایی ساده نبوده و اغلب در تحلیل مسائل بعد بالای این چینی سعی می‌شود با استفاده از روش‌های انتخاب متغیر، تعداد متغیرها را کم کرده و متغیرهای معنی‌دار را در مدل نگه داشت. برای این منظور به جای استفاده از روش کمترین توان‌های دوم، از صورت جریمه‌شده آن استفاده می‌کنند و برخی از برآوردگرهای جریمه‌شده حاصل این قابلیت را دارند که مقدار p را کاهش دهند. از این میان برآوردگر لاسو که توسط تیبشیرانی (۱۹۹۶) ارائه شد، علاوه بر برآورد پارامترها همزمان تعداد متغیرها را کاهش داده و انتخاب متغیر نیز انجام می‌دهد. اما خاصیت اوراکل را به دلیل مشتق‌پذیر نبودن در صفر ندارد و لذا توزیع مجانبی آن نرمال نیست. حاصلی مشهودی و وینچوتی (۲۰۱۶) برآوردگر جایگزینی معرفی کردند که مشتق‌پذیر است. در این پایان‌نامه پس از مرور انواع مختلف برآوردگرهای جریمه‌شده که انتخاب متغیر انجام می‌دهند، برآوردگر لاسوی مشتق‌پذیر حاصلی مشهودی و وینچوتی را معرفی کرده و خواص آن را از دیدگاه‌های نظری و عددی بررسی می‌کنیم. به عنوان تحقیقی نو، برآوردگر لاسوی مشتق‌پذیر را در مدل تفاضلی، برای تحلیل رگرسیون نیمه‌پارامتری، معرفی کرده و رفتار آن را در قالب یک مثال شبیه‌سازی مورد ارزیابی قرار می‌دهیم. از این مجموعه دو مقاله به‌دست آمده که در فهرست مقالات مستخرج از پایان‌نامه آورده شده‌اند. این مجموعه شامل ۳ فصل و یک ضمیمه است که محتویات آن‌ها به طور مختصر به قرار زیر است. در فصل ۱ به معرفی برآوردگرهای جریمه‌شده می‌پردازیم که در این میان تمرکز بر روی دو برآوردگر معروف ریج و لاسو است. در فصل ۲ برآوردگر لاسوی مشتق‌پذیر معرفی شده، خواص نظری آن را بررسی کرده و کاربرد آن با استفاده از مطالعات عددی مورد ارزیابی و مقایسه قرار گرفته است. مقاله ۱ مستخرج از این پایان‌نامه نتیجه مطالعات این فصل می‌باشد. در فصل ۳، رگرسیونی نیمه‌پارامتری را هدف گرفته و برآوردگر لاسوی مشتق‌پذیر تفاضلی را معرفی می‌کنیم. نشان می‌دهیم برآوردگر ارائه شده دارای توزیع مجانبی نرمال است. نتایج این فصل از نویسنده این مجموعه و جدید می‌باشد. حاصل یافته‌های این فصل به طور خلاصه در مقاله ۲ مستخرج از این پایان‌نامه آمده است. در آخر ملزومات نظری و کاربردی این مجموعه در پیوست آمده است.

فهرست مطالب

م فهرست تصاویر

س فهرست جداول

۱ برآوردگر لاسو ۱

۱.۱ مقدمه ۱

۲.۱ رگرسیون جریمه شده ۵

۱.۲.۱ انواع تابع جریمه ۶

۳.۱ برآوردگر لاسو ۱۳

۱.۳.۱ برآوردگر لاسو در حالت متعامد ۱۶

۲ برآوردگر لاسوی مشتق پذیر ۱۹

۱.۲ مقدمه ۱۹

۲.۲ تابع جریمه مشتق پذیر ۱۹

۳.۲ برآوردگر لاسوی مشتق پذیر ۲۶

۴.۲ الگوریتم برای برآورد پارامترها ۳۰

۵.۲ مطالعه شبیه سازی ۳۱

۳ لاسوی مشتق پذیر در مدل تفاضلی ۳۵

۱.۳ مقدمه ۳۵

۲.۳ رگرسیون ناپارامتری ۳۶

۳.۳ هموارساز خطی ۳۶

۴.۳ هموارساز میانگین متحرک ۳۸

۵.۳ هموارسازهای کرنل ۳۹

۶.۳ مدل های رگرسیون نیمه پارامتری ۴۰

۴۱	برآورد به روش تفاضلی	۷.۳
۴۳	برآوردگر لاسوی مشتق پذیر تفاضلی	۸.۳
۴۴	شبیه سازی	۹.۳
۴۹	تحلیل مثال های واقعی	۴
۴۹	مقدمه	۱.۴
۴۹	تحلیل داده های سرطان پروستات	۲.۴
۵۳	مجموعه داده های PSID	۳.۴
۵۹	تعاریف و قضایای اساسی	آ
۵۹	مرکزی کردن مشاهدات	۱.۰.آ
۶۰	تابع محدب	۱.آ
۶۱	سری تیلور	۱.۱.آ
۶۳	روش اعتبارسنجی متقابل	۲.۱.آ
۶۴	قضیه لایب نیتز	۲.آ
۶۵	معیار AIC و BIC	۱.۲.آ
۶۷	ب گزیده های از برنامه های رایانه ای	
۹۷	مراجع	

فهرست تصاویر

۱.۲	ویژگی‌های کلیدی جریمه DLASSO. (الف) همگرایی سریع به تابع قدرمطلق در مقایسه با تقریب‌های دیگر. (ب) تابع جریمه DLASSO برای s های کوچک شبیه لاسو و به ازای $s = \frac{2}{\sqrt{\pi}}$ شبیه ریج است.	۲۴
۲.۲	نمودار توابع آستانه‌ای به ازای $\lambda = 1$ و مقادیر مختلف s : $s = 0.1$ (تابع لاسو)، $s = 0.5$ ، $s = 1$ (تابع ریج)، $s = 2.0$ (برآوردگر OLS)	۲۷
۳.۲	مطالعه شبیه‌سازی مقایسه برآوردگر لاسو مشتق‌پذیر با روش‌های موجود در چهار سناریوی مختلف	۳۳
۱.۳	تابع داپلر (تابع ناپارامتری مدل در نظر گرفته شده) به ازای $n = 2000$	۴۵
۲.۳	نمودار جعبه‌ای مقایسه مخاطره برآوردگرهای لاسو مشتق‌پذیر	
۴۶	تفاضلی و کمترین توان‌های دوم تعمیم‌یافته	
۳.۳	برآورد تابع داپلر به روش هموارسازی کرنل	۴۷
۱.۴	نمودار برآورد تابع ناپارامتری	۵۸

فهرست جداول

۴۳	۱.۳	ضرایب تفاضلی بهینه
۵۰	۱.۴	جدول توصیف داده‌های سرطان پروستات
	۲.۴	مقایسه برآوردهای لاسوی مشتق‌پذیر، لاسو، Enet، SCAD با $s =$
۵۱		۰/۰۰۱، روی مجموعه داده‌های پروستات.
	۳.۴	مقادیر ضرایب برآوردهای لاسوی مشتق‌پذیر ($s = ۰/۰۰۱$)، لاسو،
۵۱		Enet، SCAD روی مجموعه داده‌های پروستات.
۵۲	۴.۴	مقایسه برآوردهای لاسوی مشتق‌پذیر و ریج
	۵.۴	مقادیر ضرایب برآوردهای لاسوی مشتق‌پذیر ($s = ۱$) و ریج روی
۵۲		مجموعه داده‌های پروستات.
۵۳	۶.۴	مقایسه برآوردهای لاسوی مشتق‌پذیر و LS
	۷.۴	مقادیر ضرایب برآوردهای لاسوی مشتق‌پذیر ($s = ۱۰۰$) و LS
۵۳		روی مجموعه داده‌های پروستات.
۵۷	۸.۴	مقادیر برآوردها و کارایی آن‌ها

فصل ۱

برآوردگر لاسو

۱.۱ مقدمه

مدل‌بندی و تحلیل رگرسیون یکی از روش‌های آماری برای به‌دست آوردن رابطه‌ی بین یک یا چند متغیر توضیحی یا تبیینی با یک یا چند متغیر پاسخ است، که در بیشتر حوزه‌های علوم مورد استفاده قرار می‌گیرد. مدل رگرسیون خطی ساده به صورت زیر تعریف می‌شود:

$$Y = \beta_0 + \beta_1 X + \epsilon, \quad (1.1)$$

که در آن Y متغیر وابسته (پاسخ)، X متغیر توضیحی (تبیینی)، β_0 عرض از مبدأ، β_1 شیب خط (ضریب رگرسیونی) و ϵ مولفه خطای تصادفی است. مرسوم است فرض کنیم $E(\epsilon) = 0$ و $E(\epsilon^2) = \sigma^2$ ، $(\sigma^2 > 0)$. در عمل تعداد متغیرهای توضیحی بیشتر از یکی است؛ به عنوان مثال در تحلیل داده‌های سرطان پروستات علاقه‌مندیم رابطه بین سطح آنتی ژن اختصاصی پروستات^۱ (Ipsa) (پادتن مخصوص پروستات) با هشت متغیر دیگر از جمله لگاریتم حجم سرطان^۲ (lcanvol)، لگاریتم وزن پروستات^۳

^۱Level of prostate specific antigen

^۲Log cancer volume

^۳Log prostate weight

(lweight)، لگاریتم مقدار یاخته پروستات خوش خیم^۴ (lbph)، لگاریتم کپسول نفوذی^۵ (lcp)، سن^۶، میزان گلیسون^۷ (gleason)، درصد گلیسون ۴ یا ۵^۸ (pgg45) و تحاجم بافتی^۹ (svi) را بررسی کنیم. در این مثال مدل (۱.۱) قابل استفاده نبوده و برای مدل بندی متغیر پاسخ با متغیرهای توضیحی عنوان شده می توان از روش رگرسیون چندگانه استفاده کرد.

مدل رگرسیون خطی چندگانه برای p متغیر توضیحی به صورت کلی زیر است

$$\begin{aligned} y_1 &= \beta_0 + \beta_1 x_{11} + \beta_2 x_{12} + \dots + \beta_p x_{1p} + \epsilon_1 \\ y_2 &= \beta_0 + \beta_1 x_{21} + \beta_2 x_{22} + \dots + \beta_p x_{2p} + \epsilon_2 \\ &\vdots \\ y_n &= \beta_0 + \beta_1 x_{n1} + \beta_2 x_{n2} + \dots + \beta_p x_{np} + \epsilon_n, \end{aligned} \quad (2.1)$$

که در آن به ازای y_i متغیر پاسخ i امین مشاهده، β_0 عرض از مبدأ، β_j به ازای j امین پارامتر رگرسیونی، ϵ_i ، i امین مولفه خطای تصادفی است که تغییرات تصادفی y_i را که به وسیله متغیرهای توضیحی مربوطه بیان نمی شود را نشان می دهد. بدیهی است در این حالت x_{ij} متغیر توضیحی مربوط به مشاهده i ام از متغیر توضیحی j است. بدون کاستن از کلیت مسئله فرض می کنیم $\beta_0 = 0$. (برای مشاهده این که چطور با مرکزی کردن مشاهدات می توان اثر عرض از مبدأ را حذف کرد پیوست ۱.۰ را ببینید.)

مدل (۲.۱) با نماد ماتریسی به صورت زیر نمایش داده می شود

$$y = X\beta + \epsilon, \quad (3.1)$$

که در آن

$$y = \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix}, X = \begin{pmatrix} x_{11} & x_{12} & \dots & x_{1p} \\ x_{21} & x_{22} & \dots & x_{2p} \\ \vdots & \vdots & & \vdots \\ x_{n1} & x_{n2} & \dots & x_{np} \end{pmatrix}, \beta = \begin{pmatrix} \beta_1 \\ \beta_2 \\ \vdots \\ \beta_p \end{pmatrix}, \epsilon = \begin{pmatrix} \epsilon_1 \\ \epsilon_2 \\ \vdots \\ \epsilon_n \end{pmatrix}.$$

در حالت کلی فرض می شود

^۴Log benign prostatic hyperplasia

^۵Log capsular penetration

^۶Age

^۷Gleason score

^۸percentage Gleason scores 4 or 5

^۹Seminal vesicle invasion

$$1. E(\epsilon) = 0.$$

$$2. \sigma^2 > 0, E(\epsilon \epsilon^T) = \sigma^2 I.$$

۳. ماتریس X غیرتصادفی است.

در مدل (۳.۱) هدف اولیه برآورد بردار β است. یکی از روش‌های برآورد بردار β روش کمترین توان‌های دوم^{۱۰} (LS) است. می‌توان مسئله را به صورت زیر فرمول‌بندی کرد

$$\hat{\beta} = \underset{\beta \in R^p}{\operatorname{argmin}} \|y - X\beta\|^2. \quad (4.1)$$

روش کمترین توان‌های دوم رهیافتی است که محک زیر را کمینه می‌کند

$$\begin{aligned} S(\beta) &= \|y - X\beta\|^2 \\ &= (y - X\beta)^T (y - X\beta) \\ &= y^T y - 2\beta^T X^T y + \beta^T X^T X \beta. \end{aligned} \quad (5.1)$$

از (۵.۱) نسبت به β مشتق می‌گیریم، داریم

$$\frac{dS(\beta)}{d\beta} = -2X^T(y - X\beta).$$

مشتق S را برابر بردار صفر قرار می‌دهیم و یک سیستم از معادلات نرمال به صورت زیر به دست می‌آوریم:

$$X^T(y - X\beta) = 0.$$

چنانچه $X^T X$ معکوس پذیر باشد، برآوردگر LS بردار β عبارت است از

$$\hat{\beta} = (X^T X)^{-1} X^T y. \quad (6.1)$$

برآوردگر $\hat{\beta}$ نااریب است. این برآوردگر، خواص خوبی داشته و طبق قضیه گوس-مارکف (کاریا و کوراتا^{۱۱}، ۲۰۰۸) در بین تمام برآوردگرهای خطی نااریب به صورت $\hat{\beta}^* = Ay$ که در آن A یک ماتریس معلوم $p \times n$ است و دارای کمترین واریانس است. روش LS ساده است، به آسانی قابل فهم است و برای نوع خاصی از داده‌ها بسیار کاربردی است. اگر مدلی با تعداد متغیرهای کم در مقابل تعداد مشاهدات موجود باشد، روش LS نتایج قابل قبولی ارائه می‌کند. با این وجود، اگر p ، (تعداد متغیرهای توضیحی) به اندازه کافی بزرگ باشد و به ویژه، وقتی بزرگتر از n (تعداد مشاهدات)

^{۱۰}Least squares

^{۱۱}Kariya and Kurata

باشد، برآورد مدل با روش LS بسیار پیچیده است. مدل پیچیده، تعداد متغیرهای زیادی را در برمی گیرد که براساس اینکه روش LS معمولاً ضرایب غیر صفر تولید می کند، تفسیرپذیری دشواری دارد. همچنین، در یک مدل معمولی که انتخاب متغیرها در نظر گرفته نشده است، بیش برآزشی مشاهده می شود. بیش برآزشی، وقتی اتفاق می افتد که یک مدل پیچیده در اختیار داشته باشیم. به عبارت دیگر، تعداد متغیرها در مقایسه با تعداد مشاهدات، بیشتر از حد انتظار وجود داشته باشد که منجر به پیش بینی ضعیفی می شود زیرا برآورد شدیداً تحت تأثیر خطای داده ها قرار می گیرد که این روزها یک مسئله رایج در روش های یادگیری آماری است (فریدمن^{۱۲}، ۲۰۰۱ و همکاران).

یک روش برای مواجهه با بیش برآزشی، استفاده از روش های جریمه شده است. اگر تنها هدف در مدل بندی رگرسیونی برآورد و تفسیر ضرایب رگرسیونی باشد برآوردگر $\hat{\beta}$ ، ابزار مفیدی است اما اگر برخی ضرایب معنی دار نباشند و بخواهیم انتخاب متغیر را نیز همزمان با برآورد ضرایب انجام دهیم، برآوردگر لاسو^{۱۳} (LASSO)، که اولین بار توسط تیشیرانی^{۱۴} (۱۹۹۶) مطرح شد، ابزار مناسب تری است. برآوردگر لاسو انقلاب بزرگی در مدل بندی رگرسیونی برپا کرد و پس از آن در مدل های رگرسیونی زیادی برای برآورد و انتخاب متغیر به طور همزمان توسعه داده شد. برای آگاهی بیشتر در مورد تاریخچه این برآوردگر و کاربردهای آن در ایران به تحقیقات آرست (۱۳۹۵) و نوروزی راد (۱۳۹۶) مراجعه کنید.

برآوردگر لاسو یک برآوردگر جریمه شده است. لذا در این فصل پس از معرفی برآوردگرهای جریمه شده به طور مختصر، برخی از توابع جریمه را نام برده و در ادامه فصل به طور اجمالی توضیحاتی بیشتر در خصوص برآوردگر لاسو، تاریخچه و چگونگی استفاده از آن می پردازیم.

در ادامه، برآوردگر LS پارامتر عرض از مبدأ را در مدل رگرسیون چندگانه می یابیم. برای این منظور فرض کنید تحت مفروضات (۳.۱)، مدل رگرسیون به صورت زیر است

$$y = \beta_0 I_n + X\beta + \epsilon$$

واضح است که مولفه های بردار y را می توان به صورت زیر نوشت

$$y_i = \alpha + \beta_1(x_{i1} - \bar{x}_1) + \beta_2(x_{i2} - \bar{x}_2) + \dots + \beta_p(x_{ip} - \bar{x}_p) + \epsilon_i$$

^{۱۲}Friedman

^{۱۳}Least absolute shrinkage and selection operator

^{۱۴}Tibshirani

که در آن

$$\alpha = \beta_0 + \beta_1 \bar{x}_1 + \dots + \beta_p \bar{x}_p, \quad \bar{x}_j = \frac{1}{n} \sum_{i=1}^n x_{ij}. \quad (7.1)$$

بنابراین

$$(X - \bar{X}) = \begin{pmatrix} x_{11} - \bar{x}_1 & x_{12} - \bar{x}_2 & \dots & x_{1p} - \bar{x}_p \\ x_{21} - \bar{x}_1 & x_{22} - \bar{x}_2 & \dots & x_{2p} - \bar{x}_p \\ \vdots & \vdots & & \vdots \\ x_{n1} - \bar{x}_1 & x_{n2} - \bar{x}_2 & \dots & x_{np} - \bar{x}_p \end{pmatrix} = \begin{pmatrix} (x_1 - \bar{x})^\top \\ (x_2 - \bar{x})^\top \\ \vdots \\ (x_n - \bar{x})^\top \end{pmatrix}$$

که در آن $\bar{x} = (\bar{x}_1, \bar{x}_2, \dots, \bar{x}_p)^\top$ و $x_i = (x_{i1}, x_{i2}, \dots, x_{ip})^\top$. لذا برآوردگر LS پارامتر β_1 به صورت

$$\hat{\beta}_1 = (x_1^\top x_1)^{-1} x_1^\top y \quad (8.1)$$

حاصل می‌شود. واضح است اگر در $E(y) = \beta_0 + \beta_1 x_1 + \dots + \beta_p x_p$ قرار دهیم $x_p = \bar{x}_p, \dots, x_2 = \bar{x}_2, x_1 = \bar{x}_1$ آن‌گاه نتیجه با (7.1) یکسان است. پس α همان $E(y) = \bar{y}$ است و لذا

$$\hat{\alpha} = \bar{y}. \quad (9.1)$$

با استفاده از (9.1) و جایگذاری در رابطه (7.1) برآورد β_0 به صورت زیر حاصل می‌شود

$$\hat{\beta}_0 = \hat{\alpha} - \hat{\beta}_1 \bar{x}_1 - \hat{\beta}_2 \bar{x}_2 - \dots - \hat{\beta}_p \bar{x}_p. \quad (10.1)$$

۲.۱ رگرسیون جریمه‌شده

از آن جایی که در حل معادله (6.1) می‌خواهیم میزان خطا را بر حسب β کمینه کنیم اگر مقادیر β ها بزرگ باشد برآوردگر حاصل، بدون در نظر گرفتن این بزرگی، اختلاف زیادی از β واقعی داشته و خطای برآورد زیاد می‌شود. یک راه حل غلبه بر این مشکل، جریمه کردن مقادیر بزرگ β است؛ لذا به جای کمینه کردن $S(\beta)$ مقدار $S(\beta) + P_\lambda(\beta)$ را کمینه می‌کنیم که در آن $P_\lambda(\beta)$ تابعی بر حسب بزرگی β بوده و پارامتر $\lambda > 0$ میزان بزرگی را مشخص می‌کند. برآوردگر حاصل از کمینه کردن عبارت فوق، برآوردگر جریمه‌شده نام دارد که به صورت زیر به دست می‌آید

$$\hat{\beta}^{Pen} = \underset{\beta \in R^p}{\operatorname{argmin}} \{S(\beta) + \lambda P(\beta)\} \quad (11.1)$$

که در آن $P(\beta)$ همان تابع جریمه $P_\lambda(\beta)$ است با این تفاوت که پارامتر λ کنترل کننده میزان جریمه بوده، به صورت ضربی (برای سادگی در محاسبات) در تابع جریمه اعمال شده است. در ادامه برخی از انواع تابع جریمه $P(\beta)$ که شکل های متفاوتی دارند و در حالت های مختلفی به کار می روند را به عنوان نمونه آورده و برخی از خواص آن ها را بیان کرده و همچنین مرجع مناسبی جهت مطالعه بیشتر ارائه می دهیم.

۱.۲.۱ انواع تابع جریمه

توابع جریمه در رگرسیون جریمه شده انواع مختلفی دارد که با توجه به مسئله مورد بررسی و کاربرد مورد انتظار، رابطه ریاضی آن متفاوت است. اما یک خانواده کلی از توابع جریمه مربوط به رگرسیون جریمه شده بریج است که این توابع جریمه دارای ساختار زیر است

$$P(\beta) = \sum_{i=1}^p |\beta_i|^\gamma, \quad 0 \leq \gamma. \quad (12.1)$$

بدیهی است برآوردگر حاصل در رابطه (۱۱.۱) با استفاده از (۱۲.۱)، برآوردگر بریج نام دارد.

لازم به ذکر است در ادامه این گونه موارد مشابه کلمه برآوردگر بریج، همانند برآوردگر ریج داریم که دلیل نام گذاری، تنها تابع جریمه خاص آن است. آرست (۱۳۹۵) برآوردگر بریج را به طور کامل بررسی کرده است. به ازای $\gamma = 1$ و $\gamma = 2$ به ترتیب برآوردگرهای لاسو و ریج حاصل می شوند. از آن جایی که هدف از این فصل، توضیح بیشتر در مورد برآوردگر لاسو است و برآوردگر ریج نیز شباهت زیادی با آن (به دلیل ساختار رابطه (۱۲.۱)) دارد، از بررسی جزئیات سایر توابع جریمه کاسته و پس از توضیح برآوردگر ریج، برآوردگر لاسو را مورد بررسی قرار می دهیم.

• برآوردگر ریج: با انتخاب

$$P(\beta) = \sum_{i=1}^p |\beta_i|^2$$

برآوردگر حاصل در رابطه (۱۱.۱)، برآوردگر ریج نام دارد. در این حالت

$$\begin{aligned} S(\beta) + \lambda P(\beta) &= (\mathbf{y} - \mathbf{X}\beta)^\top (\mathbf{y} - \mathbf{X}\beta) + \lambda \beta^\top \beta \\ &= \mathbf{y}^\top \mathbf{y} - 2\beta^\top \mathbf{X}^\top \mathbf{y} + \beta^\top \mathbf{X}^\top \mathbf{X} \beta + \lambda \beta^\top \beta \end{aligned}$$

با مشتق گیری نسبت به β و مساوی صفر قرار دادن حاصل آن داریم

$$2\mathbf{X}^\top \mathbf{X} \beta + 2\lambda \beta = 2\mathbf{X}^\top \mathbf{y}$$

و در نتیجه برآوردگر ريج با حل رابطه زیر نسبت به β

$$(X^T X + \lambda I_p) \beta = X^T y$$

به صورت

$$\hat{\beta}^R(\lambda) = (X^T X + \lambda I_p)^{-1} X^T y$$

حاصل می‌شود. برآوردگر ريج اولین بار توسط هورل و کنارد^{۱۵} (۱۹۷۰) ارائه شده است. معمولاً از برآوردگر ريج زمانی که هم‌خطی بین ستون‌های ماتریس $X^T X$ وجود دارد، استفاده می‌شود. برای توضیح بیشتر و جلوگیری از ارائه مطالب تکراری، منبع نیرومند (۱۳۸۷) را ببینید. پایان‌نامه‌های زیادی در رابطه با برآوردگر ريج، خواص و کاربردهای آن نوشته شده‌اند که از جمله آن‌ها می‌توان به نجاریان (۱۳۹۰)، برزویی بیدگلی (۱۳۹۳)، و آرست (۱۳۹۵) اشاره کرد. در ادامه به چند خاصیت مهم برآوردگر ريج اشاره می‌کنیم.

امید ریاضی و واریانس برآوردگر ريج $\hat{\beta}^R(\lambda)$ به ترتیب عبارتند از

$$\begin{aligned} E[\hat{\beta}^R(\lambda)] &= (X^T X + \lambda I_p)^{-1} X^T X \beta, \\ \text{Var}[\hat{\beta}^R(\lambda)] &= \sigma^2 (X^T X + \lambda I_p)^{-1} X^T X (X^T X + \lambda I_p)^{-1}. \end{aligned}$$

قضیه ۱.۲.۱. تحت مفروضات مدل (۳.۱)، در صورتی که $X^T X$ معکوس‌پذیر باشد، برآوردگر ريج، تحت نرم L_q ، $q > 0$ یک برآوردگر انقباضی است. به عبارتی داریم

$$\|\hat{\beta}^R(\lambda)\|_q^q \leq \|\hat{\beta}\|_q^q. \quad (13.1)$$

برهان. از آن جایی که $X^T X$ معکوس‌پذیر است، می‌توان نوشت

$$\begin{aligned} \hat{\beta}^R(\lambda) &= (X^T X + \lambda I_p)^{-1} (X^T X) (X^T X)^{-1} X^T y \\ &= (X^T X + \lambda I_p)^{-1} (X^T X) \hat{\beta} \\ &= (I_p + \lambda (X^T X)^{-1})^{-1} \hat{\beta}. \end{aligned} \quad (14.1)$$

حال تجزیه طیفی^{۱۶} $(X^T X)^{-1}$ را به صورت $(X^T X)^{-1} = \Gamma \Lambda \Gamma^T$ در نظر بگیرید که در آن Γ یک ماتریس متعامد در اندازه $p \times p$ شامل بردارهای ویژه و Λ ماتریس قطری شامل مقادیر ویژه $(X^T X)^{-1}$ است. در این صورت، نتیجه می‌شود

$$\|\hat{\beta}^R(\lambda)\|_q^q = \|(I_p + \lambda \Gamma \Lambda \Gamma^T)^{-1} \hat{\beta}\|_q^q$$

^{۱۵} Hoerl and Kennard

^{۱۶} Spectral decomposition

$$\begin{aligned}
&= \|(\mathbf{F}\mathbf{F}^\top + \lambda\mathbf{F}\mathbf{A}\mathbf{F}^\top)^{-1}\hat{\beta}\|_q^q \\
&= \|\mathbf{F}(\mathbf{I}_p + \lambda\mathbf{A})^{-1}\mathbf{F}^\top\hat{\beta}\|_q^q \\
&= (\|\mathbf{F}(\mathbf{I}_p + \lambda\mathbf{A})^{-1}\mathbf{F}^\top\hat{\beta}\|_q^2)^{\frac{q}{2}} \\
&\leq (\|\mathbf{F}^\top\hat{\beta}\|_q^2)^{\frac{q}{2}} \\
&= \|\hat{\beta}\|_q^q
\end{aligned}$$

□

در ادامه نشان می‌دهیم برآوردگر LS در یک مسئله رگرسیون با مشاهدات تجمیع شده^{۱۷}، برآوردگر ريج را نتیجه می‌دهد.

برای این منظور ماتریس گسترش‌یافته X_λ و بردار y_λ به صورت زیر را در نظر بگیرید

$$X_\lambda = \begin{pmatrix} x_{11} & x_{12} & \dots & x_{1p} \\ x_{21} & x_{22} & \dots & x_{2p} \\ \vdots & \vdots & \dots & \vdots \\ x_{n1} & x_{n2} & \dots & x_{np} \\ \sqrt{\lambda} & 0 & \dots & 0 \\ 0 & \sqrt{\lambda} & \dots & 0 \\ \vdots & \vdots & \dots & \vdots \\ 0 & 0 & \dots & \sqrt{\lambda} \end{pmatrix}, \quad y_\lambda = \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \\ 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}.$$

واضح است که

$$X_\lambda = \begin{pmatrix} X \\ \sqrt{\lambda}I_p \end{pmatrix}, \quad y_\lambda = \begin{pmatrix} y \\ 0 \end{pmatrix}.$$

در این حالت برآوردگر LS بردار پارامتر β در مدل

$$y_\lambda = X_\lambda \beta + \epsilon_\lambda$$

(که در آن ϵ_λ به طور مشابه ساخته می‌شود) به صورت زیر محاسبه می‌گردد

$$\begin{aligned}
\hat{\beta} &= \underset{\beta \in \mathbb{R}^p}{\operatorname{argmin}} (\mathbf{y}_\lambda - \mathbf{X}_\lambda \beta)^\top (\mathbf{y}_\lambda - \mathbf{X}_\lambda \beta) \\
&= (\mathbf{X}_\lambda^\top \mathbf{X}_\lambda)^{-1} \mathbf{X}_\lambda^\top \mathbf{y}_\lambda
\end{aligned}$$

^{۱۷}Augmented

$$\begin{aligned}
 &= (\mathbf{X}^\top \mathbf{X} + \lambda \mathbf{I}_p)^{-1} \mathbf{X}^\top \mathbf{y} \\
 &= \hat{\boldsymbol{\beta}}^R(\lambda).
 \end{aligned} \tag{۱۵.۱}$$

در انتها، ساختار برآوردگر ريج را در حالتی که ماتریس $\mathbf{X}$ متعامد است بررسی می‌کنیم. در این حالت، $\mathbf{X}^\top \mathbf{X} = \mathbf{I}_p$ و برآوردگر LS به صورت

$$\begin{aligned}
 \hat{\boldsymbol{\beta}} &= (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y} \\
 &= \mathbf{X}^\top \mathbf{y}
 \end{aligned}$$

است. بنابراین

$$\begin{aligned}
 \hat{\boldsymbol{\beta}}^R(\lambda) &= (\mathbf{X}^\top \mathbf{X} + \lambda \mathbf{I}_p)^{-1} \mathbf{X}^\top \mathbf{y} \\
 &= (\mathbf{I}_p + \lambda \mathbf{I}_p)^{-1} \mathbf{X}^\top \mathbf{y} \\
 &= \frac{1}{(\lambda + 1)} \mathbf{X}^\top \mathbf{y} \\
 &= \frac{1}{(\lambda + 1)} \hat{\boldsymbol{\beta}}
 \end{aligned}$$

لذا همواره

$$\begin{aligned}
 \text{Var}(\hat{\boldsymbol{\beta}}^R(\lambda)) &= \frac{1}{(\lambda + 1)^2} \text{Var}(\hat{\boldsymbol{\beta}}) \\
 &\leq \text{Var}(\hat{\boldsymbol{\beta}})
 \end{aligned} \tag{۱۶.۱}$$

و این نشان می‌دهد در حالت متعامد همواره برآوردگر ريج واریانس کمتری نسبت به برآوردگر LS دارد.

- برآوردگر لاسو: تیبشیرانی (۱۹۹۶) با در نظر گرفتن تابع جریمه $P(\boldsymbol{\beta}) = \sum_{i=1}^p |\beta_i|$ برآوردگر زیر را پیشنهاد کرد

$$\hat{\boldsymbol{\beta}}^{LASSO} = \underset{\boldsymbol{\beta} \in \mathbb{R}^p}{\operatorname{argmin}} \left\{ S(\boldsymbol{\beta}) + \lambda \sum_{i=1}^p |\beta_i| \right\}. \tag{۱۷.۱}$$

برآوردگر (۱۷.۱)، برآوردگر لاسو نام دارد. اگر چه برآوردگر لاسو صورت بسته‌ای، همانند برآوردگر ريج، ندارد ولی می‌تواند علاوه بر برآورد ضرایب $\boldsymbol{\beta}$ ، انتخاب متغیر نیز انجام دهد. در این حالت ضرایب $\hat{\boldsymbol{\beta}}$ به سمت صفر منقبض شده و برخی از آن‌ها مقدار صفر را انتخاب می‌کنند. توضیحات بیشتر در خصوص این برآوردگر در بخش ۳ آورده شده است.

- برآوردگر لاسوی تطبیقی: ژو^{۱۸} (۲۰۰۶) پیشنهاد کرد به جای تابع جریمه قدر مطلق به تنهایی، از تابع جریمه

$$P(\beta) = \hat{W}|\beta| \\ = \sum_{j=1}^p \hat{W}_j |\beta_j|$$

استفاده شود که در آن

$$\hat{W} = \text{diag}(\hat{W}_1, \dots, \hat{W}_p) \quad , \quad |\beta| = (|\beta_1|, \dots, |\beta_p|)^T$$

و

$$\hat{W}_j = \frac{1}{|\hat{\beta}_j|^r}, \quad r > 0$$

که $\hat{\beta}_j$ یک برآوردگر $\sqrt{n}$ -سازگار برای β_j می باشد. می توان از برآوردگرهای LS، لاسو و ریج به عنوان $\hat{\beta}_j$ ، $j = 1, \dots, p$ استفاده کرد. این برآوردگر دارای خاصیت اوراکل^{۱۹} (معجز)، در زیر، است در صورتی که لاسو این خاصیت را ندارد.

– خاصیت ۱: ضرایب غیر صفر را به درستی مشخص می کند. (یعنی زیر مدل درست را تعیین می کند).

– خاصیت ۲: دارای نرخ برآورد بهینه است. به عبارتی ماتریس کوواریانس آن، در حالت حدی ($n \rightarrow \infty$)، با ماتریس کوواریانس برآوردگر زیر مدل برابر است.

برآوردگر لاسوی تطبیقی به صورت زیر است

$$\hat{\beta}^{ALASSO} = \underset{\beta \in \mathbb{R}^p}{\operatorname{argmin}} \left\{ S(\beta) + \lambda \sum_{j=1}^p \hat{W}_j |\beta_j| \right\}.$$

- لاسوی ذوب شده^{۲۰}: در بعضی مواقع ممکن است بتوان یک ترتیب معنی دار بین متغیرها در نظر گرفت. جریمه لاسوی ذوب شده هم نرم L_1 ضرایب و هم تفاضل های پشت سرهم آن را جریمه می کند. از کاربرد این روش می توان در بحث های سری زمانی یا داده های از نوع تصویر اشاره کرد. تیبشیرانی و همکاران

^{۱۸}Zou

^{۱۹}Oracle

^{۲۰}Fused LASSO

(۲۰۰۵) جریمه لاسوی ذوب‌شده را به صورت

$$P_{\lambda}(\beta) = \lambda_1 \sum_{i=1}^p |\beta_i| + \lambda_2 \sum_{i=2}^p |\beta_i - \beta_{i-1}| \quad (18.1)$$

معرفی کردند. برآوردگر لاسوی ذوب‌شده به صورت زیر است:

$$\hat{\beta}^{Fused} = \operatorname{argmin}_{\beta \in \mathbb{R}^p} \{S(\beta) + P_{\lambda}(\beta)\},$$

که $P_{\lambda}(\beta)$ در (۱۸.۱) آمده است.

- برآوردگر اسکد^{۲۱}: در سال ۲۰۰۱، فن و لی^{۲۲}، سه شرط برای یک تابع جریمه خوب ارائه کردند

– شرط ۱ (نااریبی^{۲۳}): برای جلوگیری از اریبی مدل، باید پارامتر برآوردشده به مقادیر بزرگ پارامتر نزدیک باشد.

– شرط ۲ (تنکی^{۲۴}): زمانی که ضرایب کوچک هستند، برآوردگر آن‌ها را صفر می‌کند تا پیچیدگی مدل را کاهش دهد.

– شرط ۳ (پیوستگی^{۲۵}): برآوردگر بدست‌آمده در داده‌ها پیوسته است تا از ناپایداری مدل جلوگیری کند.

همچنین آن‌ها بیان کردند که توابع جریمه محدب این سه شرط را دارا نیستند و اگر مقعر باشد آن‌ها را داراست. برای این منظور تابع جریمه اسکد که مقعر است را معرفی کردند.

فن و لی (۲۰۰۱) تابع جریمه اسکد را به صورت زیر ارائه کردند

$$P_{\lambda}(\beta_j) = \begin{cases} \lambda |\beta_j| & |\beta_j| \leq \lambda \\ -\frac{|\beta_j|^2 - 2a\lambda|\beta_j| + \lambda^2}{2(a-1)} & \lambda < |\beta_j| < a\lambda \\ \frac{(a+1)\lambda^2}{2} & |\beta_j| > a\lambda \end{cases}$$

در عمل مقدار بهینه $a = 3/7$ استفاده می‌شود. تابع جریمه اسکد به طور پیوسته روی $\mathbb{R} - \{0\}$ مشتق‌پذیر است ولی در صفر مشتق‌پذیر نمی‌باشد و دارای شروط

^{۲۱}The Smoothly clipped absolute deviation

^{۲۲}Fan and Li

^{۲۳}Unbiasedness

^{۲۴}Sparsity

^{۲۵}Continuity

سه گانه تابع جریمه خوب است. بنابراین برآوردگر اسکد به صورت زیر بدست می آید

$$\hat{\beta}^{SCAD} = \underset{\beta \in \mathbb{R}^p}{\operatorname{argmin}} \left\{ S(\beta) + \sum_{j=1}^p P_{\lambda}(\beta_j) \right\}.$$

• برآوردگر الاستیکنت^{۲۶}: هرگاه تابع جریمه به صورت

$$P(\beta) = \lambda_1 \sum_{i=1}^n |\beta_i| + \lambda_2 \sum_{i=1}^n \beta_i^2$$

در نظر گرفته شود به آن تابع جریمه الاستیکنت گویند که λ_1 و λ_2 مقادیر نامنفی اختیار می کنند. اگر λ_1 مقدار صفر را اختیار کند، جریمه ریج و اگر مقدار λ_2 را صفر قرار دهیم جریمه لاسو حاصل می شود پس جریمه الاستیکنت ترکیبی از جریمه های ریج و لاسو است. برآوردگر الاستیکنت به صورت زیر به دست می آید

$$\hat{\beta}^{Enet} = \underset{\beta \in \mathbb{R}^p}{\operatorname{argmin}} \left\{ S(\beta) + \lambda_1 \sum_{i=1}^q |\beta_i| + \lambda_2 \sum_{i=1}^q \beta_i^2 \right\}. \quad (19.1)$$

این برآوردگر توسط ژو و هستی^{۲۷} (۲۰۰۵) ارائه شد. این برآوردگر برخی از ضرایب را صفر می کند و برای متغیرهای همبسته اثر یکسانی در نظر می گیرد. بنابراین این روش زمانی که گروه بندی متغیرهای مستقل زیاد است مورد استفاده قرار می گیرد.

• برآوردگر لاسوی گروهی^{۲۸}: یوان و لین^{۲۹} (۲۰۰۶) لاسو گروهی را معرفی کردند. این برآوردگر اجازه می دهد گروهی از متغیرهای همراه هم از مدل حذف شوند. به عبارت ساده تر برخلاف لاسو که انتخاب متغیر تکی صورت می گیرد در این حالت انتخاب متغیر گروهی امکان پذیر است. به عنوان مثال در متغیرهای رسته ای^{۳۰} که سطوح به صورت گروهی از متغیرهای همراه دودویی کد بندی می شود، می توان از لاسوی گروهی برای انتخاب سطوح معنی دار به طور جمعی استفاده کرد. در این حالت تابع جریمه به صورت

$$P(\beta) = \sum_{i=1}^j \|\beta_i\|_{K_i}$$

^{۲۶}Elastic Net

^{۲۷}Zou and Hastie

^{۲۸}Group lasso

^{۲۹}Yuan and Lin

^{۳۰}Categorical variables

است که j تعداد گروه‌ها یا رسته‌ها است و

$$\|\beta_i\|_{K_i} = (\beta_i^\top K_i \beta_i)^{\frac{1}{2}}$$

و K_i یک ماتریس معین مثبت است. برآوردگر لاسوی گروهی به صورت زیر است

$$\hat{\beta} = \operatorname{argmin}_{\beta \in \mathbb{R}^p} \left\{ \left\| y - \sum_{i=1}^p X_i \beta_i \right\|_2^2 + \lambda \sum_{i=1}^p \|\beta_i\|_{K_i} \right\} \quad (20.1)$$

۳.۱ برآوردگر لاسو

اگرچه که لاسو، در ابتدا برای مسائل رگرسیون خطی معرفی شد، این رهیافت در سری‌های زمانی توسط نویسندگان بسیاری به ویژه در حالت مدل اتورگرسیو به کار گرفته شد (وانگ^{۳۱} و همکاران ۲۰۰۷، ناردی و رینالدو^{۳۲}، ۲۰۱۱). از انواع دیگر لاسو می‌توان به لاسوی تطبیقی^{۳۳} (ژو^{۳۴}، ۲۰۰۶) و جزئیات آن (بولمن و ون دگیر^{۳۵}، ۲۰۱۱) اشاره کرد. این روش به عنوان مجموعه‌ای از تکنیک‌هایی برای مدل بعد بالا به کار گرفته می‌شود. واریان^{۳۶} (۲۰۱۴) و فن و همکاران^{۳۷} (۲۰۱۱) به مروری بر این روش‌ها می‌پردازند. هر دوی این مقاله‌ها، تکنیک‌های بعد بالا را خلاصه می‌کند و مثال‌های کاربردی در آن‌ها ارائه شده است. همچنین مثال کاربردی را می‌توان در بای و ان جی^{۳۸} (۲۰۰۸) مشاهده کرد. نویسندگان این مقاله، کارایی چندین روش داده‌های بعد بالا را برای پیش‌بینی سری‌های زمانی مقایسه کردند.

اگر چه رگرسیون ریج، بهبود پیش‌گویی را نتیجه می‌دهد ولی با توجه به حضور همه متغیرها در آن، تفسیرپذیری به‌سادگی امکان‌پذیر نیست. از طرف دیگر، حضور متغیرهایی که با متغیر پاسخ در یک مدل ارتباط ندارند، در فرآیند برآورد، منجر به برآوردگری با واریانس بزرگ می‌شود.

روشی ساده و بدیهی برای حل این مسئله، انتخاب یک زیرمجموعه از متغیرهای توضیحی با اهمیت و تاثیرگذار در مدل می‌باشد. انتخاب این متغیرها می‌تواند بر اساس معنی‌داری آن‌ها در مدل یا برخی از معیارهای اطلاع مانند AIC^{۳۹} (آکائیک، ۱۹۷۳)،

^{۳۱} Wang

^{۳۲} Nardi and Rinaldo

^{۳۳} Adaptive lasso

^{۳۴} Zhou

^{۳۵} Buhlmann and Van de Geer

^{۳۶} Varian

^{۳۷} Fan et al.

^{۳۸} Bay and NG

^{۳۹} Akaike Information Criterion

BIC^{۴۰} (شوارتز^{۴۱}، ۱۹۷۸)، و C_p - مالو^{۴۲} (مالوز، ۱۹۷۳) انجام شود. اما مشکل بزرگی که در هنگام استفاده از این معیارها مطرح می‌شود تعداد کل زیرمجموعه‌هایی از متغیرها است که باید در نظر گرفته شوند، که به‌طور نمایی با افزایش p رشد می‌کند و از نظر محاسباتی برای تعداد متغیرهای بیشتر از ۴۰ یا ۵۰، تقریباً نشدنی است. مشکل دیگر این روش‌ها، ویژگی گسسته بودن آن‌هاست. به عبارت دیگر، تغییر بسیار کوچکی در داده‌ها، می‌تواند برآوردهای کاملاً متفاوتی را نتیجه دهد. در مجموع، روش‌های انتخاب بهترین زیرمجموعه، اغلب ناپایدار و به شدت تغییرپذیر هستند. روش دیگر، برای رویارویی با این مسئله، استفاده از رگرسیون جریمه‌شده است. اگر چه به کاربردن جریمه ريج، در رگرسیون جریمه‌شده، به دستیابی برآوردهایی با واریانس کمتر کمک بسیاری می‌کند اما وقتی مقدار قدرمطلق ضرایب رگرسیونی بزرگ باشد، برآوردگر ريج اریبی زیادی به سمت صفر ایجاد می‌کند و از طرف دیگر، با وجودی که متغیرهای بی‌اهمیت به سمت صفر منقبض شده‌اند، همچنان در مدل حضور دارند که باعث می‌شود تفسیرپذیری مدل به راحتی انجام نشود. در این راستا، تیبشیرانی (۱۹۹۶) روشی جدید برای انتخاب متغیرها پیشنهاد نمود که منجر به مدل‌های دقیق، پایا و صرفه‌جو (تفسیرپذیرتر) می‌شود و نام آن را لاسو نامید.

از دیدگاه آماری، لاسو مخفف عملگر انتخاب‌کننده، منقبض‌کننده و کمینه‌کننده از جنس قدرمطلق است. برآوردگر لاسو، نوع جریمه‌شده برآوردگر LS است که بر اساس ویژگی تنکی نرم L_1 ، در سال‌های اخیر مورد توجه قرار گرفته است. برآوردگر لاسو از حل مسئله بهینه‌سازی زیر به دست می‌آید:

$$\hat{\beta}_n^{LASSO} = \min_{\beta} \left\{ \sum_{i=1}^n (y_i - \sum_{j=1}^p X_{ij}\beta_j)^2 \right\} \quad \text{s.t.} \quad \sum_{j=1}^p |\beta_j| \leq t, \quad (21.1)$$

که در آن t یک مقدار ثابت است. اگر $t = 0$ ، آن‌گاه متغیری در مدل وجود نخواهد داشت در حالی که اگر $t = \infty$ ، آن‌گاه مدل، کامل خواهد بود. فرض کنید $\hat{\beta}^* = (\hat{\beta}_1^*, \dots, \hat{\beta}_p^*)^T$ یک برآوردگر اولیه برای β است که معمولاً این برآوردگر را $\hat{\beta}^{LS}$ در نظر می‌گیرند. اگر $t > \sum_{i=1}^p |\hat{\beta}_i^*|$ ، روش لاسو منجر به برآوردگر LS می‌شود.

اگر $0 < t < \sum_{j=1}^p |\hat{\beta}_j^{LS}|$ باشد، آن‌گاه مسئله بهینه‌سازی (۲۱.۱) معادل با

$$\hat{\beta}^{LASSO} = \operatorname{argmin}_{\beta} \left\{ (y - X\beta)^T (y - X\beta) + \lambda \sum_{j=1}^p |\beta_j| \right\}. \quad (22.1)$$

^{۴۰} Bayesian Information Criterion

^{۴۱} Schwarz

^{۴۲} Mallows

است، که در آن λ پارامتر تنظیم کننده است و سطح تنکی (تعداد پارامترهای صفر) را در برآوردگر لاسو کنترل می کند. وقتی $\lambda \rightarrow 0$ ، $\hat{\beta}^{LASSO} \rightarrow \hat{\beta}^{LS}$ و وقتی $\lambda \rightarrow \infty$ ، $\hat{\beta}^{LASSO} \rightarrow 0$. به عبارت دیگر، هرچه جریمه بزرگتری به کار برده شود، تعداد بیشتری از ضرایب به سمت صفر منقبض می شوند. در حقیقت λ و t رفتار یا رابطه ای معکوس با یکدیگر دارند.

به رگرسیون لاسو، به خاطر شکل تابع جریمه آن، رگرسیون جریمه شده L_1 نیز می گویند. انتخاب این تابع جریمه باعث می شود که برآوردگر لاسو، ضرایب را به سمت صفر منقبض کرده و متغیرهای اضافی را از مدل حذف نماید. در واقع، برآوردگر لاسو، همزمان هم انتخاب متغیر انجام می دهد و هم ضرایب را منقبض می کند. تابع جریمه قدرمطلق باعث می شود که جواب صریحی برای برآوردگر لاسو وجود نداشته باشد. تیبشیرانی (۱۹۹۶) این برآوردگر را از طریق برنامه ریزی درجه دوم به دست آورده بود. افرون و همکاران (۲۰۰۴) نوع دیگری از رگرسیون گام به گام را به نام رگرسیون کمترین زاویه^{۴۳} (LAR) معرفی نمودند که برآوردهای لاسویی را نتیجه می دهد که هزینه محاسباتی آن با برآوردهای کمترین توان های دوم برابر است.

فن و لی (۲۰۰۱) نشان دادند که یک تابع جریمه خوب باید برآوردگری با سه ویژگی مطلوب نالاریبی، تنکی و پیوستگی را نتیجه دهد. جریمه L_1 برآوردهایی را نتیجه می دهد که تنک، نسبت به داده ها پیوسته و اریب هستند زیرا اندازه انقباضی که به ضرایب رگرسیونی کوچک و بزرگ اختصاص می دهد، یکسان است (لازم به ذکر است این مورد در خصوص جریمه L_2 صادق نیست). هم چنین، لاسو برای یک مدل رگرسیونی خطی با p متغیر توضیحی و n مشاهده، حداکثر n متغیر انتخاب می کند و بنابراین اگر متغیرهای بیشتری (بیشتر از n) در مدل معنی دار باشند، توسط لاسو انتخاب نمی شوند. لاسو، نمی تواند متغیرهایی که اثر معنی داری کمی دارند را تشخیص دهد و آن را از مدل حذف می کند. برای داده های هم خط مناسب نیست زیرا در بین گروه متغیرهایی که وابستگی شدیدی دارند، تنها یکی از آن ها را انتخاب می کند. اگر همبستگی بالایی بین متغیرهای توضیحی وجود داشته باشد، کارایی پیش گویی لاسو به خوبی برآوردگر ریج نیست. بزرگترین عیب لاسو این است که خاصیت اوراکل را ندارد.

(۱) زیر مدل درست را مشخص کند به عبارتی، مجموعه $\{j : \hat{\beta}_j \neq 0\}$ را طوری مشخص کند که $\{j : \hat{\beta}_j \neq 0\} = A$ ، که در آن A زیر مجموعه پارامترهای مدل درست است.

^{۴۳}Least Angle Regression

(۲) دارای نرخ برآورد بهینه باشد به عبارتی،

$$\sqrt{n}(\hat{\beta}_{\mathcal{A}}(\delta) - \beta_{\mathcal{A}}) \xrightarrow{\mathcal{D}} \mathcal{N}(\mathbf{0}, \Sigma^*)$$

که در آن Σ^* ماتریس کوواریانس زیرمدل درست است.

۱.۳.۱ برآوردگر لاسو در حالت متعامد

برآوردگر لاسو دارای صورت بسته نیست و نمی‌توان جواب صریحی به صورت تحلیلی، در حالت کلی، برای آن پیدا کرد. اما در حالت متعامد می‌توان رابطه تحلیلی $\hat{\beta}^{LASSO}$ را ارائه کرد که در قالب لم زیر آورده شده است.

لم ۱.۳.۱. تحت مفروضات مدل (۳.۱) فرض کنید X یک ماتریس متعامد باشد. در این صورت

$$\hat{\beta}^{LASSO} = (\hat{\beta}_i^{LASSO}, \quad i = 1, \dots, p)^\top$$

که در آن

$$\hat{\beta}_i^{LASSO} = \text{sign}(\hat{\beta}_i^{LS}) \left(|\hat{\beta}_i^{LS}| - \frac{\lambda}{2} \right)^+, \quad i = 1, \dots, p. \quad (23.1)$$

برهان. طبق فرض داریم

$X^\top X = I_p$ و $(X^\top X)^{-1} = I_p$. لذا $\hat{\beta} = (X^\top X)^{-1} X^\top y = X^\top y$ پس می‌توان نوشت

$$\begin{aligned} \hat{\beta}^{LASSO} &= \min_{\beta \in \mathbb{R}^p} \left\{ S(\beta) + \lambda \sum_{i=1}^p |\beta_i| \right\} \\ &= \min_{\beta \in \mathbb{R}^p} \left\{ y^\top y - y^\top X \beta - \beta^\top X^\top y + \beta^\top X^\top X \beta + \lambda \sum_{i=1}^p |\beta_i| \right\} \\ &\propto \min_{\beta \in \mathbb{R}^p} \left\{ -[\hat{\beta}^{LS}]^\top \beta - \beta^\top \hat{\beta}^{LS} + \beta^\top \beta + \lambda \sum_{i=1}^p |\beta_i| \right\} \\ &= \min_{\beta_1, \dots, \beta_p} \sum_{i=1}^p (-2\hat{\beta}_i^{LS} \beta_i + \beta_i^2 + \lambda |\beta_i|) \\ &= \sum_{i=1}^p \min_{\beta_i \in \mathbb{R}} \left\{ -2\hat{\beta}_i^{LS} \beta_i + \beta_i^2 + \lambda |\beta_i| \right\}. \end{aligned}$$

حال ضرایب رگرسیونی را کمینه می‌کنیم برای این منظور داریم

$$\min_{\beta_i \in \mathbb{R}} \left\{ -2\hat{\beta}_i^{LS} \beta_i + \beta_i^2 + \lambda |\beta_i| \right\} = \begin{cases} \min_{\beta_i \in \mathbb{R}^+} \left\{ -2\hat{\beta}_i^{LS} \beta_i + \beta_i^2 + \lambda \beta_i \right\} & \text{اگر } \beta_i > 0 \\ \min_{\beta_i \in \mathbb{R}^-} \left\{ -2\hat{\beta}_i^{LS} \beta_i + \beta_i^2 - \lambda \beta_i \right\} & \text{اگر } \beta_i < 0 \end{cases} \quad (24.1)$$

حاصل سمت راست عبارت (۲۴.۱) به صورت زیر است

$$\hat{\beta}_i^{LASSO}(\lambda) = \begin{cases} \hat{\beta}_i^{LS} - \frac{1}{4}\lambda & \text{اگر } \beta_i > 0 \\ \hat{\beta}_i^{LS} + \frac{1}{4}\lambda & \text{اگر } \beta_i < 0 \end{cases}$$

□

و اثبات کامل است.

فصل ۲

برآوردگر لاسوی مشتق‌پذیر

۱.۲ مقدمه

لازم به ذکر است که تابع جریمه لاسو در نقطه صفر مشتق‌پذیر نبوده و این باعث می‌شود خواص بهینه سه‌گانه یک تابع جریمه را، دارا نباشد. به منظور غلبه بر این مشکل و بهبود تابع جریمه لاسو، حاصلی مشهودی و وینچوتی^۱ (۲۰۱۶) تابع جریمه لاسوی مشتق‌پذیر را معرفی کردند. در این فصل نتایج آن‌ها را مورد بررسی قرار داده و پس از مطالعه خواص برآوردگر لاسوی مشتق‌پذیر، در یک مثال شبیه‌سازی رفتار این برآوردگر را در مقایسه با سایر برآوردگرهای رقیب تحلیل می‌کنیم.

۲.۲ تابع جریمه مشتق‌پذیر

برای این که بتوان تابع جریمه لاسوی مشتق‌پذیر را تعریف کرد، ابتدا باید این نکته را بررسی کنیم که در چه حالتی می‌توانیم مشتق‌پذیری تابع قدر مطلق را مورد مطالعه قرار دهیم. بدیهی است تابع قدر مطلق به خودی خود در نقطه صفر مشتق‌پذیر نبوده اما بعضی تقریب‌های آن، در صفر مشتق‌پذیر هستند.

^۱Haselimashhadi and Vinciotti

اشمیت^۲ و همکاران (۲۰۰۷) تقریب زیر را برای تابع $|x|$ به منظور افزایش قدرت بهینه‌سازی پیشنهاد کردند که در نقطه صفر مشتق‌پذیر است:

$$|x| \approx s \log(2 + e^{-\frac{x}{s}} + e^{\frac{x}{s}}), \quad s \in \mathbb{R}^+.$$

پس از آن رامیرز و سانچز^۳ (۲۰۱۴) بهترین تقریب هموار را برای تابع $|x|$ به صورت زیر ارائه کردند

$$|x| \approx \sqrt{x^2 + s}, \quad s \in \mathbb{R}^+$$

با تمام تلاش‌های فوق، در نهایت حاصلی مشهودی و وینچوتی (۲۰۱۶)، پیشنهاد کردند به عنوان تقریبی هموار از تابع $|x|$ می‌توان از تابع خطا (تعریف ۲.۰.۲) را در پیوست آ ببینید) استفاده نمود. دلیل استفاده از این تقریب سادگی آن در محاسبات بعدی نسبت به دو تقریب فوق است. بنابراین آن‌ها تابع جریمه لاسوی مشتق‌پذیر را به صورت زیر پیشنهاد نمودند:

$$P(\beta) = \sum_{i=1}^p P(\beta_i, s) \quad (1.2)$$

که در آن

$$P(\beta_i, s) = \beta_i \left(\frac{2}{\sqrt{\pi}} \int_0^{\frac{\beta_i}{s}} \exp\{-t\}^2 dt \right), \quad s \in \mathbb{R}^+. \quad (2.2)$$

در ادامه تابع زیر را در نظر بگیرید

$$P(x, s) = x \left(\frac{2}{\sqrt{\pi}} \int_0^{\frac{x}{s}} \exp\{-t\}^2 dt \right), \quad s \in \mathbb{R}^+. \quad (3.2)$$

این تابع جریمه را می‌توان بر حسب تابع خطا به صورت‌های مختلف زیر نوشت

$$\begin{aligned} P(x, s) &= x \operatorname{erf}\left(\frac{x}{s}\right) \\ &= x \left(1 - \operatorname{erfc}\left(\frac{x}{s}\right) \right) \\ &= x \left(2 \Phi\left(\frac{x}{s}, 0, \frac{1}{\sqrt{2}}\right) - 1 \right), \end{aligned}$$

که در آن $\Phi(x, \mu, \sigma)$ تابع توزیع تجمعی نرمال در نقطه x با میانگین μ و انحراف معیار σ است. ویژگی‌های این تابع را در قالب قضیه زیر بیان می‌کنیم.

^۲Schmidt

^۳Ramirez and Sanchez

قضیه ۱.۲.۲. تابع $P(x, s)$ در رابطه (۳.۲) دارای ویژگی‌های زیر است

$$(۱) \text{ به ازای هر } s, P(\circ, s) = \circ.$$

(۲) تابع $P(x, s)$ ، نسبت به x دو بار مشتق‌پذیر بوده که مشتقات آن به صورت زیر است

$$\begin{aligned} \frac{d}{dx}P(x, s) &= \operatorname{erf}\left(\frac{x}{s}\right) + 2\phi\left(\frac{x}{s}, \circ, \frac{1}{\sqrt{\pi}}\right)\frac{x}{s} \\ \frac{d^2}{dx^2}P(x, s) &= \frac{2}{s}\phi\left(\frac{x}{s}, \circ, \frac{1}{\sqrt{\pi}}\right) + \frac{2}{s}\phi\left(\frac{x}{s}, \circ, \frac{1}{\sqrt{\pi}}\right) - \frac{4}{x}\left(\frac{x}{s}\right)^2\phi\left(\frac{x}{s}, \circ, \frac{1}{\sqrt{\pi}}\right) \\ &= 4\phi\left(\frac{x}{s}, \circ, \frac{1}{\sqrt{\pi}}\right)\frac{1}{s}\left(1 - \left(\frac{x}{s}\right)^2\right), \end{aligned}$$

که در آن $\phi(x, \mu, \sigma)$ تابع چگالی توزیع نرمال در نقطه x با میانگین μ و انحراف معیار σ است.

(۳) تابع $P(x, s)$ نسبت به x محدب نیست.

(۴) وقتی $s \rightarrow \circ$ ، تابع با یک نرخ نمایی به $|x|$ به شدت همگرا می‌شود. در حقیقت،

$$||x| - p(x, s)| \leq 2s\phi\left(\frac{x}{s}, \circ, \frac{1}{\sqrt{\pi}}\right) \quad \forall x, s > \circ \quad (۴.۲)$$

(۵) اگر $s = \frac{2}{\sqrt{\pi}}$ ، تابع جریمه رفتاری شبیه تابع جریمه ریج در همسایگی صفر دارد. برای x های کوچک داریم

$$\begin{aligned} \frac{2}{\sqrt{\pi}} \int_{\circ}^{\frac{x}{s}} \exp\{-t^2\} dt &\approx \frac{2}{\sqrt{\pi}} \frac{x}{s} \exp\left\{-\left(\frac{x}{s}\right)^2\right\} \\ &= x^2 \frac{2}{x} \phi\left(\frac{x}{s}, \circ, \frac{1}{\sqrt{\pi}}\right). \end{aligned}$$

برهان. ۱. برای اثبات (۱)، به راحتی مشاهده می‌شود

$$P(\circ, s) = \circ \times \left(\frac{2}{\sqrt{\pi}} \int_{\circ}^{\circ} \exp\{-t^2\} dt \right) = \circ$$

۲. مشتق اول نسبت به x به صورت زیر بدست می‌آید

$$\frac{d}{dx}P(x, s) = \frac{2}{\sqrt{\pi}} \int_{\circ}^{\frac{x}{s}} \exp\{-t^2\} dt + \frac{x}{s} \times \frac{2}{\sqrt{\pi}} \exp\left\{-\left(\frac{x}{s}\right)^2\right\}$$

از رابطه فوق یک بار دیگر مشتق گرفته تا مشتق دوم نسبت به x به صورت زیر حاصل شود

$$\begin{aligned}\frac{d^2}{dx^2}P(x, s) &= \frac{1}{s} \times \frac{2}{\sqrt{\pi}} \exp\left\{-\left(\frac{x}{s}\right)^2\right\} + \frac{1}{s} \times \frac{2}{\sqrt{\pi}} \exp\left\{-\left(\frac{x}{s}\right)^2\right\} \\ &\quad - x \times \frac{2}{s} \times \frac{1}{\sqrt{\pi}} \times 2\left(\frac{x}{s}\right) \times \frac{1}{s} \exp\left\{-\left(\frac{x}{s}\right)^2\right\} \\ &= \frac{1}{s} \times \frac{2}{\sqrt{\pi}} \exp\left\{-\left(\frac{x}{s}\right)^2\right\} + \frac{1}{s} \times \frac{2}{\sqrt{\pi}} \exp\left\{-\left(\frac{x}{s}\right)^2\right\} \\ &\quad - \frac{4}{x} \times \left(\frac{x}{s}\right)^2 \times \frac{1}{\sqrt{\pi}} \exp\left\{-\left(\frac{x}{s}\right)^2\right\}\end{aligned}$$

۳. با توجه به تعریف ۱.۴ در پیوست آ اگر مشتق دوم مثبت نباشد تابع محدب نیست. با توجه به قسمت (۲) قضیه داریم

$$\frac{d^2}{dx^2}P(x, s) = \frac{4}{s\sqrt{\pi}} \exp\left\{-\left(\frac{x}{s}\right)^2\right\} \left[1 - \frac{x^2}{s^2}\right]$$

برای این که $\frac{d^2}{dx^2}p(x, s) > 0$ باید $1 - \frac{x^2}{s^2} > 0$ که نتیجه می‌شود $x^2 < s^2$ و باید به ازای همه مقادیر x و s داشته باشیم $|x| < s$. در حالی که $x \in \mathbb{R}$ است و لذا مشتق دوم همواره مثبت نیست.

۴. با استفاده از آبرامویچ و استاگن^۴ (۲۰۱۲) می‌توان کران‌های زیر را برای تابع خطا نوشت

$$\frac{1}{x + \sqrt{x^2 + 2}} < \exp\{x\}^2 \int_x^\infty \exp\{-t\}^2 dt \leq \frac{1}{x + \sqrt{x^2 + \frac{4}{\pi}}}$$

از طرفی تابع خطا به صورت زیر است

$$\begin{aligned}\frac{1}{x + \sqrt{x^2 + 2}} &< \exp\left\{\frac{x}{s}\right\}^2 \int_{\frac{x}{s}}^\infty \exp\{-t\}^2 dt \leq \frac{1}{x + \sqrt{x^2 + \frac{4}{\pi}}} \\ \frac{2}{\sqrt{\pi}} \frac{e^{-\left(\frac{x}{s}\right)^2}}{\left(\frac{x}{s}\right) + \sqrt{\left(\frac{x}{s}\right)^2 + 2}} &< \operatorname{erfc}\left(\frac{x}{s}\right) \leq \frac{2}{\sqrt{\pi}} \frac{e^{-\left(\frac{x}{s}\right)^2}}{\left(\frac{x}{s}\right) + \sqrt{\left(\frac{x}{s}\right)^2 + \frac{4}{\pi}}}\end{aligned}$$

^۴ Abramowitz and Stegun

با استفاده از کران بالای رابطه فوق به ازای x های مثبت داریم

$$\begin{aligned} \left| x - x \left(1 - \operatorname{erfc} \left(\frac{x}{s} \right) \right) \right| &= \left| x \operatorname{erfc} \left(\frac{x}{s} \right) \right| \\ \left| x \operatorname{erfc} \left(\frac{x}{s} \right) \right| &\leq \frac{\sqrt{x}}{\sqrt{\pi}} \frac{e^{-\left(\frac{x}{s}\right)^2}}{\frac{x}{s} + \sqrt{\left(\frac{x}{s}\right)^2 + \frac{4}{\pi}}} = \frac{\sqrt{x}}{\sqrt{\pi}} \frac{e^{-\left(\frac{x}{s}\right)^2}}{\frac{x}{s} + \sqrt{\left(\frac{x}{s}\right)^2 \left(1 + \frac{4}{\pi} \left(\frac{s}{x} \right)^2 \right)}} \\ &\leq \frac{\sqrt{x}}{\sqrt{\pi}} \frac{e^{-\left(\frac{x}{s}\right)^2}}{\left(\frac{x}{s} \right) + \frac{x}{s} \sqrt{1 + \frac{4}{\pi} \left(\frac{s}{x} \right)^2}} = \frac{\sqrt{s}}{\sqrt{\pi}} e^{-\left(\frac{x}{s}\right)^2} \frac{1}{1 + \sqrt{1 + \frac{4}{\pi} \left(\frac{s}{x} \right)^2}} \\ &\leq \frac{\sqrt{s}}{\sqrt{\pi}} e^{-\left(\frac{x}{s}\right)^2} \\ &= \sqrt{s} \phi \left(\frac{x}{s}, 0, \frac{1}{\sqrt{2}} \right) \end{aligned}$$

برای $x > 0$ روش مشابه‌ای استفاده می‌شود که همان نتیجه را می‌دهد. در نتیجه،

$$|x| - P(x, s) \leq \sqrt{s} \phi \left(\frac{x}{s}, 0, \frac{1}{\sqrt{2}} \right).$$

۵. فرض کنید

$$I = \int_0^{\frac{x}{s}} \exp(-t^2) dt.$$

با استفاده از انتگرال‌گیری جز به جز، فرض می‌کنیم

$$\begin{cases} u = e^{-t^2} \\ dv = dt \end{cases} \Rightarrow \begin{cases} du = -2te^{-t^2} dt \\ v = t \end{cases}$$

با توجه به این که

$$\int e^{-t^2} dt = uv - \int v du$$

می‌توان نوشت

$$I = te^{-t^2} \Big|_0^{\frac{x}{s}} + 2 \int_0^{\frac{x}{s}} t^2 e^{-t^2} dt$$

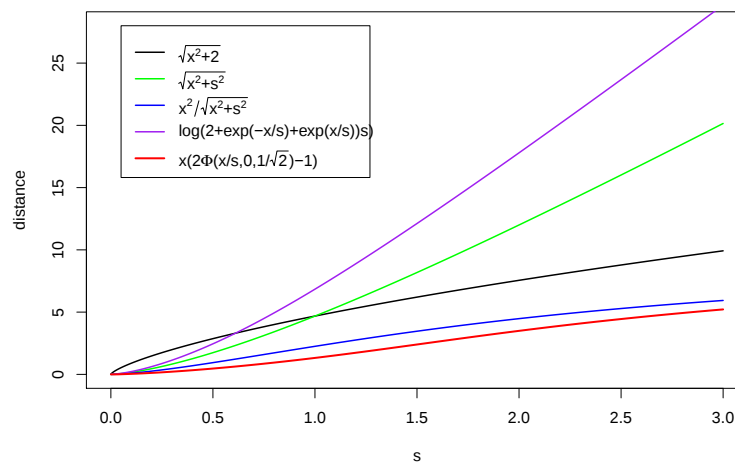
از آنجایی که x کوچک است انتگرال صفر شده و نتیجه حاصل می‌شود

$$I = \frac{x}{s} \exp \left\{ - \left(\frac{x}{s} \right)^2 \right\} + 2 \int_0^{\frac{x}{s}} t^2 e^{-t^2} dt$$

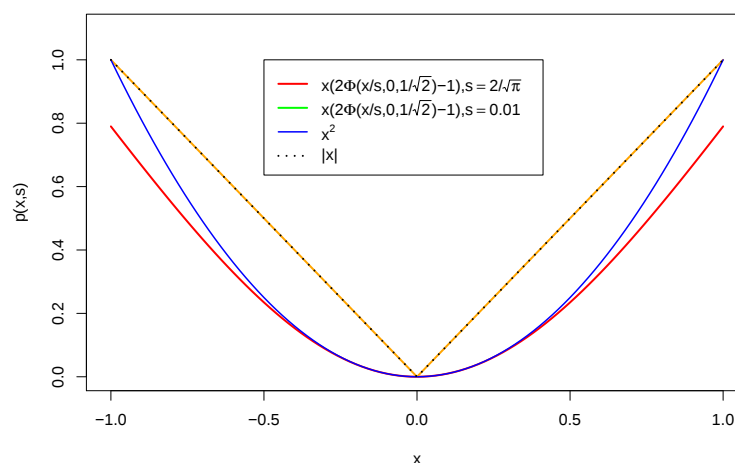
□

سمت راست نامساوی (۴.۲) وقتی $s \rightarrow \infty$ ، با سرعت به نمایی صفر نزدیک می‌شود. شکل ۱.۲ (الف) این نتیجه را به خوبی نشان می‌دهد. همان طور که از شکل مشخص است، تابع پیشنهادی به تابع قدر مطلق سریع‌تر از سایر توابع همگرا می‌شود. نمودار (ب) در شکل ۱.۲ نشان می‌دهد که وقتی s مقدار کوچکی است، تابع پیشنهادی مطابق قدر مطلق رفتار می‌کند و به ازای $s = \frac{2}{\sqrt{\pi}} \approx 1$ ، تابع مانند تابع x^2 در مجاورت $x = 0$ رفتار می‌کند.

(الف) اختلاف بین $p(x, s)$ و $|x|$



(ب) دو مثال از تابع جریمه DLASSO



شکل ۱.۲: ویژگی‌های کلیدی جریمه DLASSO. (الف) همگرایی سریع به تابع قدر مطلق در مقایسه با تقریب‌های دیگر. (ب) تابع جریمه DLASSO برای s های کوچک شبیه لاسو و به ازای $s = \frac{2}{\sqrt{\pi}}$ شبیه ريج است.

نکته نهایی بحث در مورد پیچیدگی محاسباتی است، که تنها مشکل بالقوه پیشنهاد ما است. با این حال، تعدادی از تقریب‌های خوب و سریع که در نوشته‌ها برای ارزیابی ارائه شده‌اند از تابع خطا و توزیع تابع چگالی نرمال هستند. یک نکته خوب استفاده از تقریب تیلور (پیوست ۱.۱.۱) است. به عنوان مثال، تقریب‌ها می‌توانند بر اساس یکی از عبارت‌های زیر باشند

$$\operatorname{erf}(x) = \frac{2x}{\sqrt{\pi}} \sum_{j=0}^{\infty} \frac{(-1)^j x^{2j}}{j!(2j+1)} \quad (5.2)$$

$$\operatorname{erf}(x) = \frac{2xe^{-x^2}}{\sqrt{\pi}} \sum_{j=0}^{\infty} \frac{2^j x^{2j}}{1.3.5 \dots (2j+1)}, \quad (6.2)$$

$$\operatorname{erfc}(x) \approx \frac{e^{-x^2}}{x\sqrt{\pi}} \sum_{j=0}^k (-1)^j \frac{(2j)!}{j!} (2x)^{-2j}. \quad (7.2)$$

برای $|x|$ کوچک، سری (۵.۲) نسبت به سری (۶.۲) سریع‌تر است، زیرا نیازی نیست به صورت نمایی محاسبه شود. اگرچه، سری (۶.۲) برتر از سری (۵.۲) برای $|x|$ است. زیرا شامل هیچ حذفی نمی‌شود. برای $|x|$ بزرگ، هیچ یک از سری‌ها رضایت‌بخش نیستند و در این مورد بهتر است از سری برتر (۷.۲) در انبساط مجانبی برای متمم تابع خطا استفاده شود.

تقریب‌های بسط تیلور، الگوریتم‌های تناوبی سری برای تقریب تابع خطا یا تابع چگالی نرمال دارد. به عنوان مثال، بورجسون و ساندبرگ^۵ (۱۹۷۹)، چویلارد و رول^۶ (۲۰۰۸)، کودی^۷ (۱۹۹۰)، لی^۸ (۱۹۹۲)، اولور^۹ (۲۰۱۰)، پرس^{۱۰} (۲۰۱۰) و لئال^{۱۱} و همکاران (۲۰۱۲) را نگاه کنید. به طور خاص دو تقابل سریع وجود دارند که به صورت زیر هستند:

$$\begin{aligned} \operatorname{erfc}(x) &\approx \tanh\left(\frac{39x}{2\sqrt{\pi}} - \frac{111}{2} \arctan\left(\frac{35x}{111\sqrt{\pi}}\right)\right), \\ \Phi(x, \circ, 1) &\approx \left(\frac{1}{1/9\sqrt{\pi}} \left(\sin\left(\frac{\pi x}{1^\circ}\right) + \sin(x)\right) + \circ/5\right) I_{(|x| \leq 1/513859)} \\ &+ \left(1 - e^{-1/78} + \frac{x}{e^{x+1^\circ}}\right) I_{(x > 1/513859)} + \left(1 - e^{-1/78|x|} - \frac{|x|}{e^{|x|+1^\circ}}\right) I_{(x < -1/513859)} \end{aligned}$$

^۵Borjesson and Sundberg

^۶Chevillard and Revol

^۷Cody

^۸Lee

^۹Olver

^{۱۰}Press

^{۱۱}Leal .

اینها نه تنها سریع بلکه بسیار دقیق هستند، ماکزیمم خطای مطلق آخرین تابع 10^{-4} است. برای آگاهی بیشتر حاصلی مشهودی و وینچوتی (۲۰۱۶) را ببینید.

۳.۲ برآوردگر لاسوی مشتق پذیر

با جای گذاری تابع جریمه لاسوی مشتق پذیر (۱.۲) در معادله (۱۱.۱)، برآوردگر لاسوی مشتق پذیر به صورت زیر بدست می آید

$$\hat{\beta}^{DLASSO} = \operatorname{argmin}_{\beta \in \mathbb{R}^p} \left\{ S(\beta) + \lambda \sum_{j=1}^p P(\beta_j, s) \right\}. \quad (۸.۲)$$

برای درک بهتری از برآوردگر نتیجه شده، مدل زیر را در نظر بگیرید

$$y_i = \beta_i + \epsilon_i \quad i = 1, \dots, n$$

که در آن $\epsilon_i \sim N(0, \sigma^2)$. در این صورت $\hat{\beta}_i^{DLASSO}$ پاسخ دستگاه معادله های زیر است که نسبت به β_i مشتق می گیریم

$$L(\beta) = (y - X\beta)^\top (y - X\beta) + \lambda \sum_{i=1}^p \beta_i \left(2\Phi\left(\frac{\beta_i}{s}, 0, \frac{1}{\sqrt{\lambda}}\right) - 1 \right) \quad s > 0, \lambda \geq 0$$

$$\frac{d}{d\beta_i} L(\beta) = \lambda \left(2\Phi\left(\frac{\beta_i}{s}, 0, \frac{1}{\sqrt{\lambda}}\right) - 1 + 2\left(\frac{\beta_i}{s}\right) \phi\left(\frac{\beta_i}{s}, 0, \frac{1}{\sqrt{\lambda}}\right) \right) - 2(y_i - \beta_i) = 0, i = 1, \dots, n$$

شکل ۲.۲ برآوردگرها را به ازای مقادیر مختلف s نشان می دهد. همان طور که قبلا اشاره شد و با توجه به این شکل، (الف)، (ج) و (د) به ترتیب توابع مشابه لاسو، ریج و جریمه نشده را نشان می دهد.

با توجه به اینکه برای مقادیر کوچک s ($s \rightarrow 0$)، تقریب های بسیار خوبی برای تابع قدر مطلق به دست می آید، انتظار داریم که در این حالت برآوردگرها، ویژگی هایی شبیه برآوردگر لاسو داشته باشند که در قضیه های بعدی اثبات شده است.

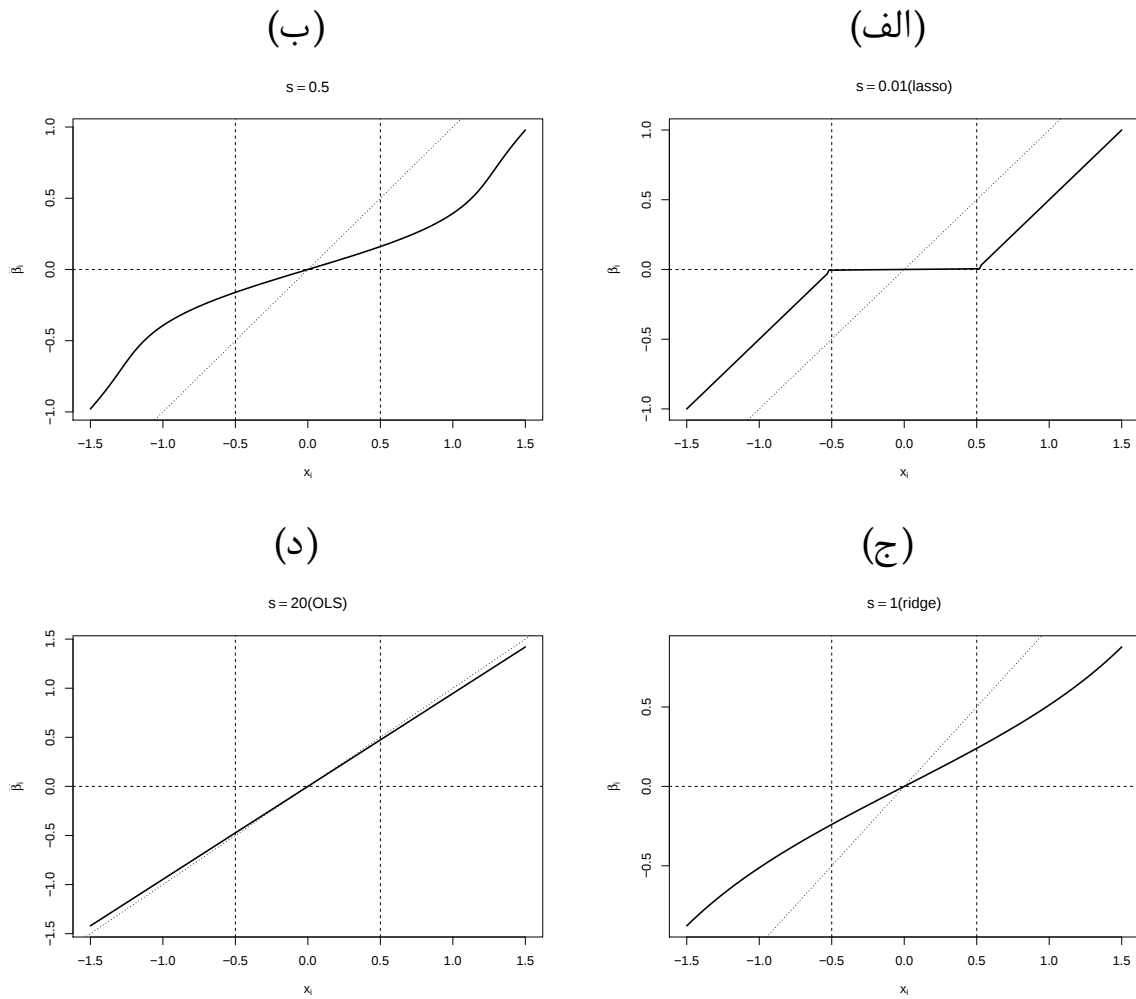
قضیه ۱.۳.۲. برای هر $u \in \mathbb{R}$ و $\lambda \geq 0$ و $s > 0$ ، تعریف کنید

$$k(u, s) = L(\beta + u) - L(\beta) \quad (۹.۲)$$

که در آن $L(\beta)$ تابع هدف در معادله (۸.۲) است. در این صورت

$$\lim_{s \rightarrow 0} k(u, s) = u^\top X_n^\top X_n u - 2u^\top N\left(0, \sigma^2(X_n^\top X_n)\right) + \lambda \sum_{i=1}^m \left(|u_i| I(\beta_i = 0) + u_i \operatorname{sign}(\beta_i + u_i) \right),$$

که در آن $N(0, \sigma^2)$ یک متغیر تصادفی با توزیع نرمال است.



شکل ۲.۲: نمودار توابع آستانه‌ای به‌ازای $\lambda = 1$ و مقادیر مختلف s : $s = 0.01$ (تابع لاسو)، $s = 0.5$ ، $s = 1$ (تابع ریدج)، $s = 20$ (برآوردگر OLS)

برهان. باجای گذاری در $L(\beta)$ داریم،

$$\begin{aligned}
 k(u, s) &= L(\beta + u) - L(\beta) \\
 &= (y - X(\beta + u))^T (y - X(\beta + u)) + \lambda \sum_{i=1}^p (\beta_i + u_i) \left(\mathcal{R}\Phi\left(\frac{(\beta_i + u_i)}{s}, 0, \frac{1}{\sqrt{\lambda}}\right) \right) \\
 &\quad - (y - X\beta)^T (y - X\beta) + \lambda \sum_{i=1}^p \beta_i \left(\mathcal{R}\Phi\left(\frac{\beta_i}{s}, 0, \frac{1}{\sqrt{\lambda}}\right) - 1 \right) \\
 &= (\epsilon - Xu)^T (\epsilon - Xu) - \epsilon^T \epsilon + \mathcal{R}\lambda \sum_{i=1}^p \beta_i \int_{\frac{\beta_i}{s}}^{\frac{\beta_i + u_i}{s}} \frac{1}{\sqrt{\pi}} e^{(-t)^2} dt \\
 &\quad + \lambda \sum_{i=1}^p u_i \left(\mathcal{R}\Phi\left(\frac{\beta_i + u_i}{s}, 0, \frac{1}{\sqrt{\lambda}}\right) - 1 \right)
 \end{aligned}$$

حال اگر $s \rightarrow \circ$ داریم

$$\begin{aligned}
 \lim_{s \rightarrow \circ} k(\mathbf{u}, s) &= \lim_{s \rightarrow \circ} L(\boldsymbol{\beta} + \mathbf{u}) - \lim_{s \rightarrow \circ} L(\boldsymbol{\beta}) \\
 &= (\epsilon - \mathbf{X}\mathbf{u})^\top (\epsilon - \mathbf{X}\mathbf{u}) - \epsilon^\top \epsilon + \lim_{s \rightarrow \circ} \left\{ \gamma \lambda \sum_{i=1}^p \beta_i \int_{\frac{\beta_i}{s}}^{\frac{\beta_i + u_i}{s}} \frac{1}{\sqrt{\pi}} e^{(-t)^\gamma} dt \right. \\
 &\quad \left. + \lambda \sum_{i=1}^p u_i \left(\gamma \Phi\left(\frac{\beta_i + u_i}{s}, \circ, \frac{1}{\sqrt{\gamma}}\right) - 1 \right) \right\} \\
 &= \mathbf{u}^\top \mathbf{X}^\top \mathbf{X} \mathbf{u} - \gamma \mathbf{u}^\top N(\circ, \sigma^\gamma(\mathbf{X}^\top \mathbf{X})) + \lim_{s \rightarrow \circ} \lambda \sum_{i=1}^p \left(\gamma \beta_i \frac{u_i}{s} \phi\left(\frac{u_i}{s}, \circ, \frac{1}{\sqrt{\gamma}}\right) \right. \\
 &\quad \left. + \lim_{s \rightarrow \circ} \left\{ u_i \left(\gamma \Phi\left(\frac{\beta_i + u_i}{s}, \circ, \frac{1}{\sqrt{\gamma}}\right) - 1 \right) \right\} \right) \\
 &= \mathbf{u}^\top \mathbf{X}^\top \mathbf{X} \mathbf{u} - \gamma \mathbf{u}^\top N(\circ, \sigma^\gamma(\mathbf{X}^\top \mathbf{X})) \\
 &\quad + \lim_{s \rightarrow \circ} \lambda \sum_{i=1}^p \begin{cases} u_i \left(\gamma \Phi\left(\frac{u_i}{s}, \circ, \frac{1}{\sqrt{\gamma}}\right) - 1 \right) & \beta_i = \circ \\ \gamma \beta_i \frac{u_i}{s} \phi\left(\frac{u_i}{s}, \circ, \frac{1}{\sqrt{\gamma}}\right) + u_i \left(\gamma \Phi\left(\frac{\beta_i + u_i}{s}, \circ, \frac{1}{\sqrt{\gamma}}\right) - 1 \right) & \beta_i \neq \circ \end{cases} \\
 &= \mathbf{u}^\top \mathbf{X}^\top \mathbf{X} \mathbf{u} - \gamma \mathbf{u}^\top N(\circ, \sigma^\gamma(\mathbf{X}^\top \mathbf{X})) \\
 &\quad + \lambda \sum_{i=1}^p \begin{cases} |u_i| & \beta_i = \circ \\ u_i & \beta_i + u_i > \circ, \\ -u_i & \beta_i + u_i < \circ \end{cases}
 \end{aligned}$$

و داریم

$$k(\mathbf{u}) = \mathbf{u}^\top \mathbf{X}^\top \mathbf{X} \mathbf{u} - \gamma \mathbf{u}^\top N(\circ, \sigma^\gamma(\mathbf{X}^\top \mathbf{X})) + \lambda \sum_{i=1}^p \left(|u_i| I(\beta_i = \circ) + u_i \text{sign}(\beta_i + u_i) \right).$$

□

در قضیه بعد نشان می دهیم که حداقل سرعت s که همگرایی برآوردگرها را به لاسو تضمین می کند، $n^{(-\frac{1}{\gamma} + \epsilon)}$ به ازای هر $\epsilon > \circ$ است.

قضیه ۲.۳.۲. فرض کنید $\boldsymbol{\beta}$ مجموعه ای از ضرایب تنک باشد، $\mathbf{u} \in \mathbb{R}^p$ ، $s_n = s/(n^{\frac{1}{\gamma} + \epsilon}) \rightarrow \circ$ ، $\lambda_n/\sqrt{n} \rightarrow \lambda_\circ \geq \circ$ ، $\epsilon > \circ$ ، در این صورت، $\sqrt{n}(\hat{\boldsymbol{\beta}}^{DLASSO} - \boldsymbol{\beta}) \rightarrow \underset{\mathbf{u}}{\text{argmin}} k(\mathbf{u})$ که در آن

$$k(\mathbf{u}) = -\gamma \mathbf{u}^\top N(\circ, \sigma^\gamma \Sigma) + \mathbf{u}^\top \Sigma \mathbf{u} + \lambda_\circ \sum_{i=1}^p \{u_i \text{sign}(\beta_i) I(\beta_i \neq \circ) + |u_i| I(\beta_i = \circ)\}.$$

برهان. فرض کنید $k_n(\mathbf{u}) = L(\boldsymbol{\beta} + \frac{\mathbf{u}}{\sqrt{n}}) - L(\boldsymbol{\beta})$ در این صورت

$$\begin{aligned}
 k_n(\mathbf{u}) &= L(\boldsymbol{\beta} + \frac{\mathbf{u}}{\sqrt{n}}) - L(\boldsymbol{\beta}) \\
 &= (\mathbf{y} - \mathbf{X}(\boldsymbol{\beta} + \frac{\mathbf{u}}{\sqrt{n}}))^\top (\mathbf{y} - \mathbf{X}(\boldsymbol{\beta} + \frac{\mathbf{u}}{\sqrt{n}})) \\
 &\quad + \lambda_n \sum_{i=1}^p (\beta_i + \frac{u_i}{\sqrt{n}}) \left(\gamma \Phi\left(\frac{(\beta_i + \frac{u_i}{\sqrt{n}})}{s}, \circ, \frac{1}{\sqrt{\gamma}}\right) \right) \\
 &\quad - (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^\top (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) + \lambda_n \sum_{i=1}^p \beta_i \left(\gamma \Phi\left(\frac{\beta_i}{s}, \circ, \frac{1}{\sqrt{\gamma}}\right) - 1 \right) \\
 &= (\boldsymbol{\epsilon} - \mathbf{X} \frac{\mathbf{u}}{\sqrt{n}})^\top (\boldsymbol{\epsilon} - \mathbf{X} \frac{\mathbf{u}}{\sqrt{n}}) - \boldsymbol{\epsilon}^\top \boldsymbol{\epsilon} + \gamma \lambda_n \sum_{i=1}^p \beta_i \int_{\frac{\beta_i}{s}}^{\frac{\beta_i + \frac{u_i}{\sqrt{n}}}{s}} \frac{1}{\sqrt{\pi}} e^{-t^2} dt \\
 &\quad + \lambda_n \sum_{i=1}^p \frac{u_i}{\sqrt{n}} \left(\gamma \Phi\left(\frac{\beta_i + \frac{u_i}{\sqrt{n}}}{s}, \circ, \frac{1}{\sqrt{\gamma}}\right) - 1 \right) \\
 &= \mathbf{u}^\top \frac{\mathbf{X}^\top \mathbf{X}}{n} \mathbf{u} - \gamma \mathbf{u}^\top N(\circ, \sigma^2 \left(\frac{\mathbf{X}^\top \mathbf{X}}{n} \right)) + \gamma \lambda_n \sum_{i=1}^p \beta_i \int_{\frac{\beta_i}{s}}^{\frac{\beta_i + \frac{u_i}{\sqrt{n}}}{s}} \frac{1}{\sqrt{\pi}} e^{-t^2} dt \\
 &\quad + \lambda_n \sum_{i=1}^p \frac{u_i}{\sqrt{n}} \left(\gamma \Phi\left(\frac{\beta_i + \frac{u_i}{\sqrt{n}}}{s}, \circ, \frac{1}{\sqrt{\gamma}}\right) - 1 \right)
 \end{aligned}$$

با توجه به فرضیات اگر $n \rightarrow \infty$ ، داریم

$$\begin{aligned}
 k_n(\mathbf{u}) &\xrightarrow{n \rightarrow \infty} \mathbf{u}^\top \Sigma \mathbf{u} - \gamma \mathbf{u}^\top N(\circ, \sigma^2 \Sigma) + \lim_{n \rightarrow \infty} \gamma \lambda_n \sum_{i=1}^p \beta_i \int_{\frac{\beta_i}{s}}^{\frac{\beta_i + \frac{u_i}{\sqrt{n}}}{s}} \frac{1}{\sqrt{\pi}} e^{-t^2} dt \\
 &\quad + \lim_{n \rightarrow \infty} \lambda_n \sum_{i=1}^p \frac{u_i}{\sqrt{n}} \left(\gamma \Phi\left(\frac{\beta_i + \frac{u_i}{\sqrt{n}}}{s}, \circ, \frac{1}{\sqrt{\gamma}}\right) - 1 \right) \\
 &= \mathbf{u}^\top \Sigma \mathbf{u} - \gamma \mathbf{u}^\top N(\circ, \sigma^2 \Sigma) + \gamma \lambda_\circ \sum_{i=1}^p \beta_i \frac{u_i}{s \sqrt{\pi}} \lim_{n \rightarrow \infty} e^{-\left(\frac{\beta_i + \frac{u_i}{\sqrt{n}}}{s}\right)^2} \\
 &\quad + \lambda_\circ \sum_{i=1}^p u_i \left(\lim_{n \rightarrow \infty} \gamma \Phi\left(\frac{\beta_i + \frac{u_i}{\sqrt{n}}}{s}, \circ, \frac{1}{\sqrt{\gamma}}\right) - 1 \right) \\
 &= \mathbf{u}^\top \Sigma \mathbf{u} - \gamma \mathbf{u}^\top N(\circ, \sigma^2 \Sigma) \\
 &\quad + \lim_{n \rightarrow \infty} \left\{ \begin{aligned} &\lambda_\circ \sum_{i=1}^p u_i \left(\gamma \Phi\left(\frac{\frac{u_i}{\sqrt{n}}}{s}, \circ, \frac{1}{\sqrt{\gamma}}\right) - 1 \right) && \beta_i \in S_\circ \\ &\gamma \lambda_\circ \sum_{i=1}^p \left(\beta_i \frac{u_i}{s \sqrt{\pi}} e^{-\left(\frac{\beta_i + \frac{u_i}{\sqrt{n}}}{s}\right)^2} + u_i \left(\gamma \Phi\left(\frac{\beta_i + \frac{u_i}{\sqrt{n}}}{s}, \circ, \frac{1}{\sqrt{\gamma}}\right) - 1 \right) \right) \beta_i && \beta_i \in S_\circ^c \end{aligned} \right.
 \end{aligned}$$

□

که $S_\circ$ و $S_\circ^c$ به ترتیب مجموعه ضرایب صفر و غیر صفر هستند.

مشابه روش نایت و فو^{۱۲} (۲۰۰۰)، می‌توان نشان داد که این نتایج برآورد تنک پارامترها را تضمین می‌کند. در اثبات قضیه ۲.۳.۲ فرض شده است که $S_n\sqrt{n} \rightarrow 0$. بنابراین، اگر مقدار S_n را کمتر از $\frac{1}{\sqrt{n}}$ انتخاب شود نتایج بسیار شبیه لاسو خواهد بود.

۴.۲ الگوریتم برای برآورد پارامترها

درست‌نمایی جریمه (۸.۲) نسبت به β مشتق‌پذیر است. بنابراین روش‌های بهینه‌سازی استاندارد می‌تواند برای یافتن مینیمم آن استفاده شود. با این وجود، سرعت یافتن پاسخ در این روش‌ها به کندی خواهد بود. در این بخش الگوریتم کارایی بر اساس ویژگی مشتق‌پذیری تابع جریمه لاسوی مشتق‌پذیر ارائه می‌کنیم. بدین منظور، مشابه روش پیشنهادی فن و لی (۲۰۰۱) الگوریتم تکرارشونده زیر پیشنهاد می‌شود

$$\beta^{(k)} = \left(X^T X + \Sigma(\beta^{(k-1)}, \lambda, s) \right)^{-1} X^T y, \quad k = 1, 2, \dots \quad (10.2)$$

در اینجا $\beta^{(0)}$ یک برآورد اولیه برای پارامترهاست که می‌تواند برآورد لاسو باشد و $\Sigma(\beta^{(k-1)}, \lambda, s)$ به صورت زیر تعریف می‌شود

$$\Sigma(\beta^{(k-1)}, \lambda, s) = \lambda \text{Diag} \left[\left(2\Phi\left(\frac{\beta_i^{(k-1)}}{s}, 0, \frac{1}{\sqrt{\lambda}}\right) - 1 \right) + \frac{\beta_i^{(k-1)}}{s} \phi\left(\frac{\beta_i^{(k-1)}}{s}, 0, \frac{1}{\sqrt{\lambda}}\right) / \beta_i^{(k-1)}, i = 1, \dots, r \right].$$

به منظور به دست آوردن این رابطه، از تقریب تیلور مرتبه اول برای تابع جریمه لاسوی مشتق‌پذیر حول $\beta^{(0)}$ استفاده می‌کنیم. به عبارت دیگر،

$$\begin{aligned} \beta \left(2\Phi\left(\frac{\beta}{s}, 0, \frac{1}{\sqrt{\lambda}}\right) - 1 \right) &\approx \beta^{(0)} \left(2\Phi\left(\frac{\beta^{(0)}}{s}, 0, \frac{1}{\sqrt{\lambda}}\right) - 1 \right) \\ &+ \left(2\Phi\left(\frac{\beta^{(0)}}{s}, 0, \frac{1}{\sqrt{\lambda}}\right) - 1 \right. \\ &\left. + \frac{2\beta^{(0)}}{s} \phi\left(\frac{\beta^{(0)}}{s}, 0, \frac{1}{\sqrt{\lambda}}\right) \right) (\beta - \beta^{(0)}). \end{aligned}$$

توجه داشته باشید که مشتق‌پذیر بودن تابع لاسو بدین معناست که نیازی به تقریب درجه دوم موضعی وجود ندارد (فن و لی، ۲۰۰۱). اگر $\beta \approx \beta^{(0)}$ ، می‌توانیم معادله را

^{۱۲} Knight and Fu

به صورت زیر بازنویسی کنیم

$$\begin{aligned} \beta \left(2\Phi\left(\frac{\beta}{s}, \circ, \frac{1}{\sqrt{2}}\right) - 1 \right) &\approx \beta^{(\circ)} \left(2\Phi\left(\frac{\beta^{(\circ)}}{s}, \circ, \frac{1}{\sqrt{2}}\right) - 1 \right) \\ &+ \frac{1}{\beta^{(\circ)}} \left(2\Phi\left(\frac{\beta^{(\circ)}}{s}, \circ, \frac{1}{\sqrt{2}}\right) - 1 \right) \\ &+ \frac{2\beta^{(\circ)}}{s} \phi\left(\frac{\beta^{(\circ)}}{s}, \circ, \frac{1}{\sqrt{2}}\right) (\beta^2 - \beta^{(\circ)2}). \quad (11.2) \end{aligned}$$

با جایگذاری (۱۱.۲) در تابع جریمه (۸.۲) داریم

$$\begin{aligned} (y - X\beta)^\top (y - X\beta) &+ \lambda \sum_{i=1}^p \left[\beta_i^{(\circ)} \left(2\Phi\left(\frac{\beta_i^{(\circ)}}{s}, \circ, \frac{1}{\sqrt{2}}\right) - 1 \right) \right. \\ &+ \frac{1}{\beta_i^{(\circ)}} \left(2\Phi\left(\frac{\beta_i^{(\circ)}}{s}, \circ, \frac{1}{\sqrt{2}}\right) - 1 \right) \\ &\left. + \frac{2\beta_i^{(\circ)}}{s} \phi\left(\frac{\beta_i^{(\circ)}}{s}, \circ, \frac{1}{\sqrt{2}}\right) (\beta_i^2 - \beta_i^{(\circ)2}) \right] \quad (12.2) \end{aligned}$$

در این صورت، مقدار بهینه با روش محاسبات تکراری رگرسیون ریج (۱۰.۲) به دست می‌آید. این الگوریتم در بسته *DLASSO* در نرم‌افزار *R* قابل اجرا است. این بسته اجازه می‌دهد پارامترهای تنظیم‌کننده λ و s را با معیاری که انتخاب مدل رایج مانند *AIC*، *BIC* یا *GCV* انتخاب کنیم.

۵.۲ مطالعه شبیه‌سازی

در این بخش، برای بررسی کارایی تابع جریمه در رگرسیون، از شبیه‌سازی استفاده شده است. مدل زیر را در نظر بگیرید

$$y = X\beta + \sigma\epsilon \quad \epsilon \sim N(\circ, 1)$$

مشابه ژو و هستی (۲۰۰۵)، سه سناریوی مختلف زیر را طراحی می‌کنیم
سناریوی ۱ (استاندارد): بردار پارامتر β را به صورت $(3, 1.5, \circ, \circ, 2, \circ, \circ, \circ)$ قرار می‌دهیم. متغیرهای X را از توزیع نرمال چند متغیره با میانگین صفر و ماتریس همبستگی $Cor(X_i X_j) = \circ/5^{|i-j|}$ و براین اساس به ازای $\sigma^2 = 3$ ، متغیر پاسخ را تولید می‌کنیم.

سناریوی ۲ (β های کوچک): مشابه سناریوی ۱ با این تفاوت که بردار β را به صورت $\beta_j = \circ/5$ به ازای $j = 1, 2, \dots, 8$ در نظر می‌گیریم.

سناریوی ۳ (متغیرهای پیشگو همبسته): در این سناریو فرض می‌کنیم $p = 3$ و ضرایب را به سه گروه تقسیم می‌کنیم. $\beta^{(1)} = (1, 2, 3, 4, 5)$ ، $\beta^{(2)} = (0.5, 0.5, 0.5, 0.5, 0.5)$ و $\beta^{(3)} = (0, 0, 0, 0, 0)$. همبستگی بالای ۰/۹ را بین هر جفت از اولین پنج متغیر در نظر می‌گیریم. به طور مشابه، همبستگی ۰/۵ را در گروه دوم و هیچ وابستگی را در گروه سوم در نظر می‌گیریم. در نهایت $\sigma^2 = 15$.

سناریو ۴ (استاندارد): بردار پارامتر β را به صورت $(3, 1.05, 0, 0, 2, \underbrace{0, 0, 0, \dots}_{295}, 0)$ قرار می‌دهیم. متغیرهای X را از توزیع نرمال چند متغیره با میانگین صفر و ماتریس همبستگی $Cor(X_i X_j) = 0.5^{|i-j|}$ و برای این اساس به ازای $\sigma^2 = 3$ ، متغیر پاسخ را تولید می‌کنیم.

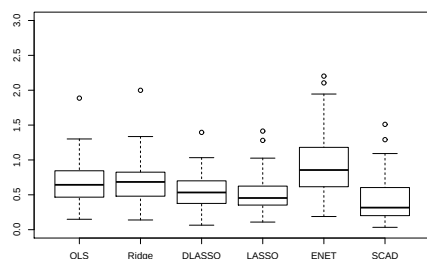
برای هر سه سناریو، ۵۰ مجموعه داده از ۲۴۰ مشاهده تولید می‌کنیم. ۴۰ مشاهده از آن‌ها را به عنوان مجموعه آموزشی و b مشاهده را به عنوان مجموعه آزمون در نظر می‌گیریم. پارامترهای جریمه بر اساس روش اعتبارسنجی متقابل ۱۰ تایی روی MSE محاسبه می‌شوند.

مدل‌های زیر را مقایسه می‌کنیم: لاسوی مشتق‌پذیر (به ازای $s = 0.01$)، لاسوی مشتق‌پذیر (با s و λ محاسبه‌شده براساس CV)، لاسو، ریج، الاستیک‌نت (enet)، اسکد و کمترین توان‌های دوم (LS).

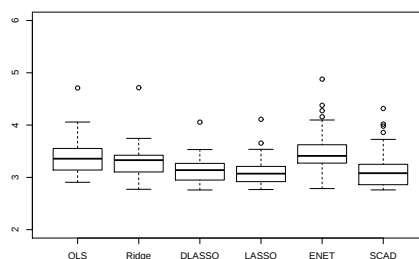
برای محاسبه برآوردگرهای لاسو و الاستیک‌نت از بسته msgps (هیروس، ۲۰۱۲) استفاده می‌کنیم. از بسته ncvmreg (برهنی و هوانگ، ۲۰۱۱) و از بسته glmnet (فریدمن و همکاران، ۲۰۱۰) به ترتیب برای یافتن برآوردگرهای اسکد و ریج استفاده شده است. همچنین برای روش‌های dlasso از بسته DLASSO (حاصلی مشهدی، ۲۰۱۷) استفاده کرده‌ایم.

شکل ۳.۲ نتایج را نشان می‌دهد. برای هر سناریو، توزیع (روی ۵۰ تکرار) از MSE روی مجموعه آزمون که به صورت $\sum_{i=1}^{200} \frac{(y_i - \hat{y}_i)^2}{200}$ تعریف می‌شود و MSE پارامترهای برآوردشده که به صورت $(\hat{\beta} - \beta)^T S_X (\hat{\beta} - \beta)$ تعریف می‌شود که در آن β و $\hat{\beta}$ به ترتیب پارامترهای واقعی و مقدار برآوردشده است.

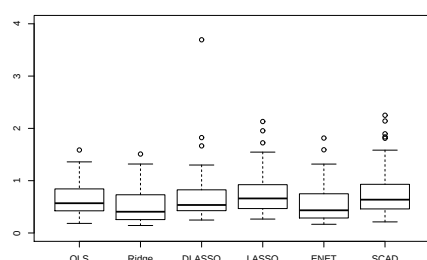
سناریوی ۱: $MSE(\beta)$



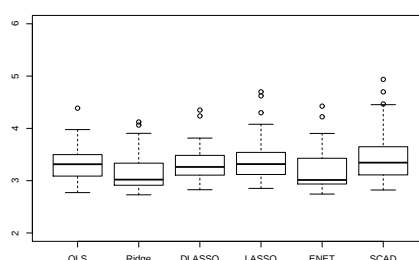
سناریوی ۱: $MSE(\text{Prediction})$



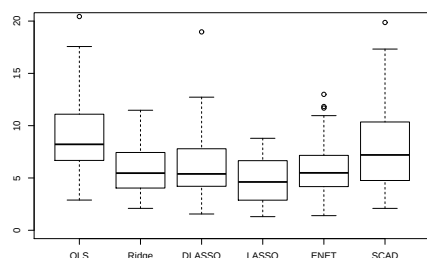
سناریوی ۲: $MSE(\beta)$



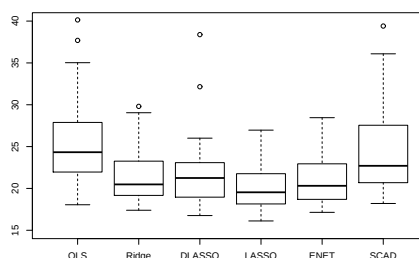
سناریوی ۲: $MSE(\text{Prediction})$



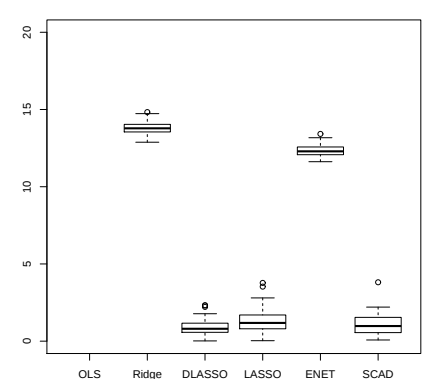
سناریوی ۳: $MSE(\beta)$



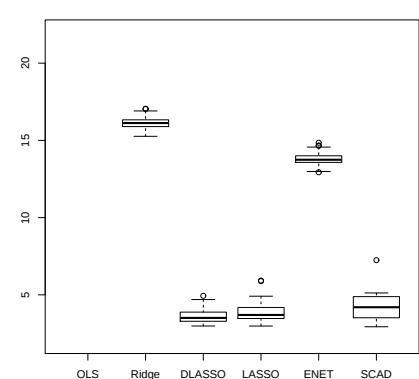
سناریوی ۳: $MSE(\text{Prediction})$



سناریوی ۴: $MSE(\beta)$



سناریوی ۴: $MSE(\text{prediction})$



شکل ۳.۲: مطالعه شبیه‌سازی مقایسه برآوردگر لاسوی مشتق‌پذیر با روش‌های موجود در چهار سناریوی مختلف

در اغلب سناریوها، جریمه لاسوی مشتق‌پذیر بهتر یا مشابه جریمه‌های موجود و بهتر از جریمه ریج رفتار می‌کند که تنها تابع جریمه‌ی مشتق‌پذیر در بین سایر تابع جریمه‌های رایج است و همچنین نسبت به اسکد که تنها جریمه غیر محدب است (به جز سناریو ۱) برتری دارد. اگرچه جریمه لاسو نیز در بعضی اوقات بهتر از لاسوی مشتق‌پذیر است به یاد داشته باشیم که ایرار تابع جریمه لاسو مشتق‌پذیر نبودن آن در صفر بوده که دلیل اصلی پیشنهاد تابع جریمه لاسوی مشتق‌پذیر است.

فصل ۳

لاسوی مشتق‌پذیر در مدل تفاضلی

۱.۳ مقدمه

مدل‌های نیمه‌پارامتری شامل قسمت پارامتری و قسمت نیمه‌پارامتری هستند. این مدل‌ها کاربردهای فراوانی دارند. انگل^۱ و همکاران (۱۹۸۶) اولین‌ها بودند که مدل خطی نیمه‌پارامتری را به کار بردند. در مدل رگرسیونی نیمه‌پارامتری، یاجینو^۲ (۲۰۰۳) بر برآورد پارامتر خطی متمرکز شد و از روش تفاضل برای از بین بردن اریبی ناشی از حضور پارامتر ناپارامتری در مدل استفاده نمود. روش تفاضلی کردن برای حذف اثر تابع ناپارامتری در مدل‌های رگرسیون ناپارامتری چندان جدید نیست (به رایس^۳ (۱۹۸۴) مراجعه شود). در حالی که به کار بردن این ایده در مدل‌های نیمه‌پارامتری نسبتاً جدیدتر hsj (به آن و پاول^۴، ۱۹۹۳ و یاجیو، ۱۹۹۷، ۱۹۹۸، ۱۹۹۹، ۲۰۰۰ و ۲۰۰۱، تاباکان و آکدنیز^۵، ۲۰۱۰، روزبه و همکاران، ۲۰۱۰، آرشی و همکاران، ۲۰۱۱ و روزبه، ۱۳۹۰ مراجعه شود). در این فصل علاقه‌مندیم برآوردگر لاسوی مشتق‌پذیر را با مدل

^۱Engle

^۲Yatchew

^۳Rice

^۴Ahn and Pawell

^۵Tabakan and Akdeniz

رگرسیون نیمه‌پارامتری به دست بیاوریم. مطالب این فصل عمدتاً برگرفته از پایان‌نامه بغلانی (۱۳۹۴) است اما بخش لاسوی مشتق‌پذیر آن جدید بوده و حاصل تحقیقات نویسنده این مجموعه می‌باشد.

۲.۳ رگرسیون ناپارامتری

در مدل رگرسیون ناپارامتری n جفت از مشاهدات $(t_1, y_1), \dots, (t_n, y_n)$ در دست داریم. متغیر پاسخ y توسط معادله رگرسیون ناپارامتری که به t وابسته است به صورت زیر خواهد بود

$$y_i = f(t_i) + \epsilon_i, \quad E(\epsilon_i) = 0, i = 1, \dots, n \quad (1.3)$$

به طوری که $f(\cdot)$ یک تابع ناپارامتری است و t_i ها از هم مستقل هستند. می‌خواهیم تابع $f(\cdot)$ را تحت فرضیات زیر برآورد کنیم:

(۱) t_i ها عضو یک مجموعه فشرده باشند.

(۲) تابع $f(\cdot)$ یک تابع هموار باشد (مشتقات اول و دوم وجود داشته باشند).

برآوردگر $f(t)$ با $\hat{f}_n(t)$ نشان داده می‌شود. همچنین $\hat{f}_n(t)$ به عنوان یک هموارساز شناخته می‌شود. در ادامه برای برآورد $f(t)$ به معرفی هموارساز خطی^۶ و هموارساز میانگین متحرک^۷ و همچنین هموارسازهای کرنل^۸ می‌پردازیم.

۳.۳ هموارساز خطی

یک برآوردگر $\hat{f}_n$ یک هموارساز خطی است اگر برای هر t ، یک بردار $W(t) = (w_1(t), \dots, w_n(t))^T$ وجود داشته باشد به طوری که

$$\hat{f}_n = \sum_{i=1}^n W_i(t) y_i. \quad (2.3)$$

یک بردار از مقادیر برازش‌شده به صورت زیر تعریف می‌شود:

$$f = (\hat{f}_1(t_1), \dots, \hat{f}_n(t_n)). \quad (3.3)$$

^۶ Linear Smoother

^۷ Moving Average Smoother

^۸ Kernel Smoothers

به طوری که $y = (y_1, \dots, y_n)^\top$. بنابراین $f = Wy$ به طوری که W یک ماتریس $n \times n$ است که i امین ردیف آن $W(t_i)^\top$ می باشد. لذا قرار می دهیم

$$W_{ij} = W_j(t_i).$$

اعضای i امین سطر نشان دهنده وزن هایی است که به هر y_i داده می شود. وزن های موجود در همه هموارسازها که ما از آن ها استفاده می کنیم یک خصوصیت دارند و آن اینکه برای تمام t ها داریم

$$\sum_{i=1}^n W_i(t) = 1.$$

این نشان می دهد که هموارساز، منحنی را ثابت نگه می دارد. بنابراین اگر $Y_i = C$ برای تمام i ها باشد، $\hat{f}_n(t) = C$ خواهد بود. برای آگاهی بیشتر بغلانی (۱۳۹۴) را ببینید.

مثال ۱.۳.۳. (رگرسوگرام) فرض کنید $a < t_i < b$ ، $i = 1, 2, \dots, n$ را در m بلوک مساوی تقسیم می کنیم که با $B_1, B_2, \dots, B_m$ نمایش داده می شوند. $\hat{f}_n(t)$ را به صورت زیر تعریف می کنیم

$$\hat{f}_n(t) = \frac{1}{k_j} \sum_{i:t_i \in B_j} y_i, \quad t_i \in B_j$$

به طوری که k_j تعداد نقاط موجود در B_j است. به عبارت دیگر برآورد $\hat{f}_n$ تابعی است که توسط میانگین گیری کردن y_i ها در هر بلوک به دست آمده است. این برآورد، رگرسوگرام^۹ گفته می شود.

اکنون تعریف می کنیم

$$W_i(t) = \begin{cases} \frac{1}{k_j} & \text{اگر } t_i \in B_j \\ 0 & \text{اگر } t_i \notin B_j \end{cases}$$

بنابراین خواهیم داشت $\hat{f}_n(t) = \sum_{i=1}^n y_i W_i(t)$. پس به نظر می رسد که بردار وزن های W_i این گونه باشد

$$W^\top(t) = (0, 0, \dots, \frac{1}{k_j}, \dots, \frac{1}{k_j}, 0, \dots, 0)$$

^۹Regressogram

برای دیدن شکل ماتریس هموارساز W که $n = 9$ و $m = 3$ و $k_1 = k_2 = k_3 = 3$ باشد. همچنین ماتریس هموارساز W به شکل ماتریس مربع خواهد بود. بنابراین

$$W = \frac{1}{3} \times \begin{pmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 \\ 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 1 & 1 & 0 & 0 & 0 \\ 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix}$$

۴.۳ هموارساز میانگین متحرک

در این بخش به معرفی یک هموارساز بسیار ساده به نام میانگین متحرک یا «میانگین پویا» می‌پردازیم. فرض کنید داده‌های $(t_1, y_1), \dots, (t_n, y_n)$ در مدل

$$y = f(t) + \epsilon$$

به ما داده شده باشد. فرض می‌کنیم t اسکالر بوده و داده‌ها به صورت $t_1 \leq \dots \leq t_n$ باشند. به طوری که t_i ها دارای فاصله واحد باشند. برآورد تابع $f(\cdot)$ در t_i به صورت میانگین k مشاهده متوالی مرکزی در t_i است. (برای جلوگیری از ابهام بهتر است که k را فرد بگیریم). به صورت فرمولی تعریف می‌کنیم

$$\hat{f}(t_i) = \frac{1}{k} \sum_{j=\underline{i}}^{\bar{i}} y_j \quad (4.3)$$

به طوری که $\underline{i} = i - (k-1)/2$ و $\bar{i} = i + (k-1)/2$ نشان‌دهنده حد پایین و حد بالای علامت مجموع است. همچنین برآوردگر برابر است با

$$\hat{f}(t_i) = \frac{1}{k} \sum_{j=\underline{i}}^{\bar{i}} \hat{f}(t_j) + \frac{1}{k} \sum_{j=\underline{i}}^{\bar{i}} \epsilon_j \quad (5.3)$$

اگر k (تعداد همسایگی‌های میانگین‌گیری شده) با افزایش n افزایش یابد سپس مطابق قضایای حد مرکزی، قسمت دوم در سمت راست به طور تقریبی نرمال با میانگین صفر و واریانس σ_ϵ^2/k خواهد بود. برای آگاهی بیشتر روزه (۱۳۹۰) را ببینید.

اغلب نوشتن هموارساز میانگین متحرک (وسایر هموارسازها) به صورت ماتریسی مناسب خواهد بود. S را ماتریس هموارساز قرار می‌دهیم که به صورت زیر تعریف می‌شود

$$S_{(n-k+1)} = \begin{pmatrix} \frac{1}{k} & \cdots & \frac{1}{k} & 0 & 0 & \cdots & 0 \\ 0 & \frac{1}{k} & \cdots & \frac{1}{k} & 0 & \cdots & 0 \\ \vdots & \cdots & \vdots & \cdots & \vdots & \cdots & \vdots \\ 0 & \cdots & 0 & \frac{1}{k} & \cdots & \frac{1}{k} & 0 \\ 0 & \cdots & 0 & 0 & \frac{1}{k} & \cdots & \frac{1}{k} \end{pmatrix}$$

بنابراین می‌توان رابطه (۱.۳) را به شکل بردار-ماتریس به صورت زیر نوشت

$$\hat{y} = \hat{f}(t) = Sy \quad (۶.۳)$$

به طوری که t ، Y ، $\hat{Y}$ و $\hat{f}(t)$ بردار هستند.

۵.۳ هموارسازهای کرنل

اکنون برآوردگر کرنل (هسته) f را در نقطه t_0 به صورت زیر در نظر بگیرید

$$\hat{f}(t_0) = \sum_{i=1}^n w_i(t_0) y_i \quad (۷.۳)$$

در اینجا تابع رگرسیون را در نقطه t_0 به عنوان یک وزن از مجموع y_i ها برآورد می‌کنیم به طوری که وزن‌های $w_i(t_0)$ به t_0 بستگی دارد.

یک راه مناسب برای ساخت مفهوم متوسط وزنی، به کار بردن یک تابع کیفی مرکزی در صفر است که در هر دو جهت با یک نرخ کاهش می‌یابد و به وسیله یک پارامتر مقیاس کنترل می‌شود. این توابع عموماً به عنوان کرنل شناخته می‌شوند و توابع چگالی احتمال هستند. $k(\cdot)$ را یک تابع کراندار قرار می‌دهیم که انتگرال آن ۱ و در اطراف ۰ متقارن است. وزن‌ها را به صورت زیر تعریف می‌کنیم

$$w_i(t_0) = \frac{\frac{1}{hn} k\left(\frac{t_i - t_0}{h}\right)}{\frac{1}{hn} \sum_{i=1}^n k\left(\frac{t_i - t_0}{h}\right)} \quad (۸.۳)$$

شکل وزن‌ها توسط $k(\cdot)$ تعیین می‌گردد و بزرگی مقدار آن‌ها توسط h کنترل می‌شود که به عنوان پهنای باند شناخته می‌شود. انتخاب پهنای باند h ، مقدار همواری را کنترل می‌کند. پهنای باند کم برآورد بسیار ناهنجاری می‌دهد، در حالی که پهنای باند بزرگتر

برآورد هموارتری به دست می‌آورد. هر چه مشاهدات از نقطه t_0 دور تر شوند، پارامتر h وزن بزرگتری به آن‌ها در برآورد $f(t_0)$ می‌دهد. با استفاده از رابطه (۸.۳) برآوردگر تابع رگرسیون ناپارامتری، در ابتدا توسط نادارای^{۱۰} (۱۹۶۴) و واتسون^{۱۱} (۱۹۶۴)، پیشنهاد گردید که به صورت زیر آمده است

$$\hat{f}(t_0) = \frac{\frac{1}{hn} \sum_{i=1}^n k\left(\frac{t_i - t_0}{h}\right) y_i}{\frac{1}{hn} \sum_{i=1}^n k\left(\frac{t_i - t_0}{h}\right)}. \quad (9.3)$$

عموما انتخاب کرنل از انتخاب پهنای باند اهمیت کمتری دارد. ساده‌ترین کرنل، کرنل یکنواخت است (همچنین به کرنل مستطیلی و جعبه‌ای نیز می‌توان اشاره کرد)، به طوری که مقداری به اندازه $\frac{1}{2}$ در $[-1, 1]$ و در جاهای دیگر ۰ به دست می‌دهد. اما کرنل نرمال و سایر کرنل‌ها نیز به طور گسترده مورد استفاده قرار می‌گیرند. (به وند و جونز^{۱۲} ۱۹۹۵، برای مشاهده انواع بیشتری از هموارسازهای کرنل مراجعه شود). در بسیاری از موارد به طور رایج از میانگین متحرک استفاده می‌شود. کرنل یکنواخت به آسانی میانگین مشاهداتی که در محدوده $t_0 \pm h$ قرار می‌گیرند را به دست می‌دهد.

۶.۳ مدل‌های رگرسیون نیمه‌پارامتری

رگرسیون نیمه‌پارامتری می‌تواند به عنوان یک نمونه دیگری از رگرسیون پارامتری و همچنین رفتار دو مدل پارامتری و ناپارامتری باهم، مورد بحث قرار گیرد. داشتن هر دو اجزای پارامتری و ناپارامتری در مدل آن را به یک مدل نیمه‌پارامتری تبدیل می‌کند.

مجموعه داده‌های $(y_1, x_1, t_1), \dots, (y_n, x_n, t_n)$ از مدل رگرسیون نیمه‌پارامتری تبعیت می‌کند، هرگاه

$$y_i = x_i^\top \beta + f(t_i) + \epsilon_i, \quad i = 1, \dots, n \quad (10.3)$$

به طوری که $x_i^\top = (x_{i1}, x_{i2}, \dots, x_{ip})$ برای $i = 1, 2, \dots, n$ و $\beta = (\beta_0, \beta_1, \dots, \beta_p)^\top$ یک بردار از پارامترهای نامعلوم است. به طوری که $f(\cdot)$ یک تابع مجهول و هموار است و t ها عضو یک مجموعه فشرده و به صورت $t_1 \leq t_2 \leq \dots \leq t_n$ مرتب شده‌اند به گونه‌ای که n تعداد مشاهدات در نمونه است. فرض می‌کنیم ϵ یک بردار n تایی با توزیع نامشخص که دارای $E(\epsilon) = 0$ و $E(\epsilon\epsilon^\top) = \sigma^2 V$ می‌باشد. مدل‌های رگرسیون

^{۱۰}Nadaraya

^{۱۱}Watson

^{۱۲}Wand and Jones

نیمه پارامتری نسبت به مدل های خطی استاندارد انعطاف پذیرترند، به این دلیل که آن ها هم دارای قسمت پارامتری و هم ناپارامتری هستند. در واقع استفاده از این مدل زمانی مفید است که متغیر وابسته y به طور خطی با متغیر مستقل x و به طور غیر خطی با متغیر مستقل t رابطه داشته باشد.

۷.۳ برآورد به روش تفاضلی

در برآورد تفاضلی، ابتدا روندی را که توسط تابع $f(\cdot)$ در داده ها به وجود می آید را حذف نموده و سپس به برآورد پارامترهای مدل می پردازیم. بنابراین در این روش نیازی به برآورد تابع $f(\cdot)$ نداشته و به طور مستقیم پارامترها را برآورد می کنیم. به طور خلاصه، با استفاده از روش تفاضلی می توان:

(۱) بدون نیاز به استفاده از تکنیک های برازش ناپارامتری، واریانس مانده ها را برآورد کرد.

(۲) آزمون فرضیه در مورد خطی بودن تابع ناپارامتری را انجام داد.

(۳) قسمت خطی مدل های نیمه پارامتری را برآورد کرد.

(۴) آزمون فرضیه در مورد خطی بودن تابع ناپارامتری در مدل های نیمه پارامتری انجام داد.

(۵) می توان آزمون فرضیه برابر بودن دو تابع در دو مدل مختلف در مدل های خطی نیمه پارامتری را انجام داد.

برای به دست آوردن شکلی از برآوردگر تفاضلی ابتدا مدل (۱۰.۳) را به صورت ماتریس و بردار به صورت زیر بازنویسی می کنیم

$$y = X\beta + f(t) + \epsilon \quad (11.3)$$

به طوری که $X = (x_1, \dots, x_n)^T$ و $x_i^T = (x_{i1}, x_{i2}, \dots, x_{ip})$, $i = 1, 2, \dots, n$ که X یک ماتریس $n \times p$ می باشد و داشته باشیم $f(t) = (f(t_1), \dots, f(t_n))^T$ و $y = (y_1, \dots, y_n)^T$ و $\epsilon = (\epsilon_1, \dots, \epsilon_n)^T$. $d = (d_0, \dots, d_m)$ را یک بردار $m + 1$ تایی در نظر می گیریم، به گونه ای که m مرتبه تفاضل و $d_0, \dots, d_m$ وزن های تفاضل با شرایط زیر باشند

$$\sum_{j=0}^m d_j = 0, \quad \sum_{j=0}^m d_j^2 = 1 \quad (12.3)$$

اکنون ماتریس تفاضلی D با ابعاد $(n-m) \times n$ که عناصر آن از رابطه (۱۲.۳) تبعیت می‌کنند را به صورت زیر تعریف می‌کنیم

$$D = \begin{pmatrix} d_0 & d_1 & \dots & d_m & 0 & 0 & \dots & 0 \\ 0 & d_0 & d_1 & \dots & d_m & 0 & \dots & 0 \\ \vdots & \vdots & \vdots & \vdots & \dots & \vdots & \dots & \vdots \\ 0 & 0 & \dots & 0 & d_0 & d_1 & \dots & d_m \end{pmatrix}$$

به عنوان مثال اگر وزن‌ها را بصورت زیر تعریف کنیم، آن‌گاه شرایط (۱۲.۳) برقرار است:

$$d_0 = \sqrt{\frac{m}{m+1}}, d_1 = d_2 = \dots, d_m = -\sqrt{\frac{1}{m(m+1)}}$$

هال و همکارانش (۱۹۹۰) مقدار بهینه ضرایب تفاضلی را با استفاده از روش‌های عددی تا مرتبه $m = 10$ به دست آوردند. این ضرایب در جدول (۱.۳) آورده شده‌اند. در عین حال روش مشخصی برای تعیین مقدار m پیشنهاد نشده است. برای آگاهی بیشتر روزه (۱۳۹۴) را ببینید.

با استفاده از ماتریس تفاضلی در مدل (۱۱.۳) اجازه برآورد مستقیم اثر پارامتر را می‌دهد. در مجموع داریم

$$Dy = DX\beta + Df(t) + D\epsilon \quad (13.3)$$

ماتریس تفاضلی D در مدل (۱۳.۳) اثر تابع ناپارامتری را در نمونه‌های بزرگ از بین می‌برد. بنابراین مدل (۱۳.۳) را می‌توان به شکل زیر باز نویسی نمود

$$Dy = DX\beta + D\epsilon$$

یا

$$\tilde{y} = \tilde{X}\beta + \tilde{\epsilon} \quad (14.3)$$

به گونه‌ای که $\tilde{y} = Dy$ ، $\tilde{X} = DX$ و $\tilde{\epsilon} = D\epsilon$. بنابراین، $\tilde{\epsilon}$ یک بردار $(n-m)$ تایی دارای توزیع نامعلوم با $E(\tilde{\epsilon}) = 0$ و $E(\tilde{\epsilon}\tilde{\epsilon}^T) = \sigma^2 V_D$ است به طوری که داریم $V_D = DVD^T \neq I_{n-m}$ برآوردگر تفاضلی کمترین توان‌های دوم تعمیم‌یافته^{۱۳} (DGLSE) به صورت زیر به دست می‌آید:

$$\begin{aligned} \hat{\beta}_{GD}^{LS} &= \underset{\beta \in \mathbb{R}^p}{\operatorname{argmin}} (\tilde{y} - \tilde{X}\beta)^T V_D^{-1} (\tilde{y} - \tilde{X}\beta) \\ &= C_D^{-1} \tilde{X}^T V_D^{-1} \tilde{y} \end{aligned} \quad (15.3)$$

به گونه‌ای که $C_D = \tilde{X}^T V_D^{-1} \tilde{X}$.

^{۱۳} Difference-based generalized least squares estimator

جدول ۱.۳: ضرایب تفاضلی بهینه

m	۱	۲	۳	۴	۵
d_0	۰/۷۰۷۱	۰/۸۰۹۰	۰/۸۵۸۲	۰/۸۸۷۳	۰/۹۰۶۴
d_1	-۰/۷۰۷۱	-۰/۵۰۰۰	-۰/۳۸۳۲	-۰/۳۰۹۹	-۰/۲۶۰۰
d_2		-۰/۳۰۹۰	-۰/۲۸۰۹	-۰/۲۴۶۴	-۰/۲۱۹۷
d_3			-۰/۱۹۴۲	-۰/۱۹۰۱	-۰/۱۷۷۴
d_4				-۰/۱۴۰۹	-۰/۱۴۲۰
d_5					-۰/۱۱۰۳
m	۶	۷	۸	۹	۱۰
d_0	۰/۹۲۰۰	۰/۹۳۰۲	۰/۹۳۸۰	۰/۹۴۴۳	۰/۹۴۹۴
d_1	-۰/۲۲۳۸	-۰/۱۹۶۵	-۰/۱۷۵۱	-۰/۱۵۷۸	-۰/۱۴۳۷
d_2	-۰/۱۹۲۵	-۰/۱۷۲۸	-۰/۱۵۶۵	-۰/۱۴۲۹	-۰/۱۳۱۴
d_3	-۰/۱۶۳۵	-۰/۱۵۰۶	-۰/۱۳۸۹	-۰/۱۲۸۷	-۰/۱۱۹۷
d_4	-۰/۱۳۶۹	-۰/۱۲۹۹	-۰/۱۲۲۴	-۰/۱۱۵۲	-۰/۱۰۸۵
d_5	-۰/۱۱۲۶	-۰/۱۱۰۷	-۰/۱۰۶۹	-۰/۱۰۲۵	-۰/۰۹۷۸
d_6	-۰/۰۹۰۶	-۰/۰۹۳۰	-۰/۰۹۲۵	-۰/۰۹۰۵	-۰/۰۸۷۷
d_7		-۰/۰۷۶۸	-۰/۰۷۹۱	-۰/۰۷۹۲	-۰/۰۷۸۲
d_8			-۰/۰۶۶۶	-۰/۰۶۸۷	-۰/۰۶۹۱
d_9				-۰/۰۵۸۸	-۰/۰۶۰۶
d_{10}					-۰/۰۵۲۷

۸.۳ برآوردگر لاسوی مشتق‌پذیر تفاضلی

در این بخش برآوردگر لاسوی مشتق‌پذیر تفاضلی^{۱۴} (DDLASSO) را معرفی کرده و خواص آن را بررسی می‌کنیم. تحت مفروضات مدل (۱۴.۳) برآوردگر DDLASSO به صورت زیر به دست می‌آید:

$$\begin{aligned}
 \hat{\beta}^{DLASSO} &= \operatorname{argmin}_{\beta \in \mathbb{R}^p} \left\{ (\tilde{y} - \tilde{X}\beta)^\top V_D^{-1} (\tilde{y} - \tilde{X}\beta) + \lambda \sum_{j=1}^p P(\beta_j, s) \right\} \\
 &= C_D^{-1} \tilde{X}^\top V_D^{-1} \tilde{y} + \lambda \sum_{j=1}^p P(\beta_j, s).
 \end{aligned} \tag{۱۶.۳}$$

^{۱۴} Differenced-based DLASSO

قضیه زیر توزیع مجانبی برآوردگر $\hat{\beta}^{DLASSO}$ را ارائه می‌کنیم. از آن جایی که اثبات این قضیه مشابه اثبات قضیه ۲.۳.۲ است از آوردن جزئیات آن خودداری می‌کنیم.

قضیه ۱.۸.۳. فرض کنید $u \in \mathbb{R}^p$ ، $S_n^* = \frac{s}{(n-m)^{\frac{1}{2}+\epsilon}} \rightarrow 0$ ، $\epsilon > 0$ ، $\frac{\lambda_n}{\sqrt{n-m}} \rightarrow \lambda_0 \geq 0$ ، $\frac{\tilde{X}^\top \tilde{X}}{n-m} \rightarrow \Sigma^*$ در این صورت

$$\sqrt{n-m}(\hat{\beta}^{DDLASSO} - \beta) \rightarrow \underset{\beta \in \mathbb{R}^p}{\operatorname{argmin}} k^*(u)$$

که در آن

$$k^*(u) = -2u^\top N(0, r\Sigma^*) + u^\top \Sigma^* u + \lambda_0 \sum_{i=1}^p \left\{ u_i \operatorname{Sign}(\beta_i) \mathbf{I}(\beta_i \neq 0) + |u_i| \mathbf{I}(\beta_i = 0) \right\}.$$

بدیهی است الگوریتم دستیابی به برآوردگر $\hat{\beta}^{DDLASSO}$ مشابه بخش ۴.۲ بوده با این تفاوت که X به $\tilde{X}$ تبدیل می‌شود.

۹.۳ شبیه‌سازی

در این قسمت، در قالب یک مثال شبیه‌سازی برآوردگر لاسوی مشتق‌پذیر را در مدل تفاضلی به‌دست می‌آوریم. در این مطالعه، متغیر پاسخ را از مدل زیر به حجم $n = 100$ با 10^3 تکرار تولید می‌کنیم

$$y_i = x_{1i}\beta_1 + x_{2i}\beta_2 + x_{3i}\beta_3 + x_{4i}\beta_4 + x_{5i}\beta_5 + f(t_i) + \epsilon_i, i = 1, \dots, n \quad (17.3)$$

که در آن

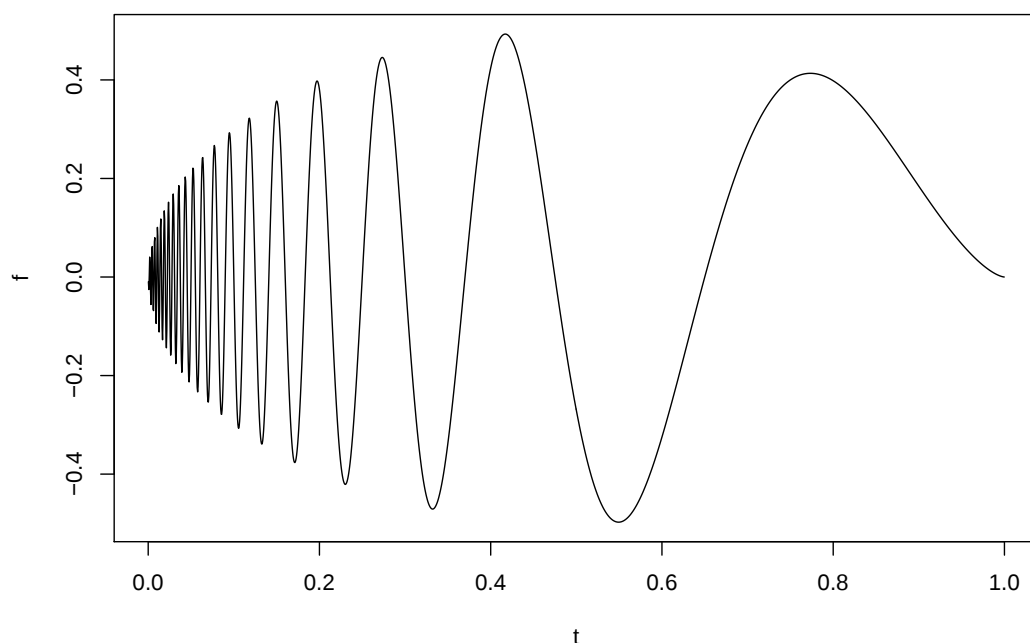
$$\beta = (1/5, 2, 3, -5, 4)^\top$$

$$x_i \sim N_\Delta(\mu, \Sigma_x), \mu = \begin{pmatrix} -2 \\ 1 \\ 2 \\ 3 \\ 2 \end{pmatrix}, \Sigma_x = \begin{pmatrix} 0/81 & 0/4 & 0/3 & -0/2 & -0/1 \\ 0/4 & 2/25 & 0/4 & 0/3 & -0/2 \\ 0/3 & 0/4 & 1 & 0/4 & 0/3 \\ -0/2 & 0/3 & 0/4 & 0/64 & 0/4 \\ -0/1 & -0/2 & 0/3 & 0/4 & 0/49 \end{pmatrix}$$

$$f(t_i) = \sqrt{t_i(1-t_i)} \sin\left(\frac{2/1\pi}{t_i + 0/05}\right), t_i = (i - 0/5)/n$$

$$\epsilon_i \sim N(0, \sigma_\epsilon^\top V), \sigma_\epsilon^\top = 0/01, v_{ij} = \left(\frac{1}{n}\right)^{|i-j|+1}$$

تابع ناپارامتری در نظر گرفته شده، تابع داپلر^{۱۵} نامیده می‌شود. تابع داپلر به دلیل شکل موجی و ناهمگن بودنش (شکل ۱.۳) به سختی برآورد می‌شود و به همین دلیل استفاده از آن برای نشان دادن کارایی روش تفاضلی و روش هموارسازی به کار رفته می‌باشد.



شکل ۱.۳: تابع داپلر (تابع ناپارامتری مدل در نظر گرفته شده) به ازای $n = 2000$

شبیه‌سازی مونت کارلو به ازای $N = 1000$ انجام شد و برآوردگرهای پیشنهادی به دست آورده شد. قسمت پارامتری مدل (۱۷.۳)، به عبارتی β با استفاده از ضرایب تفاضلی مرتبه چهارم، $d_0 = 0.8873$ ، $d_1 = -0.3099$ ، $d_2 = -0.2465$ ، $d_3 = -0.1901$ و $d_4 = -0.1409$ برآورد می‌شود. آن‌گاه با جایگزینی برآورد به دست آمده برای پارامتر β ، تابع ناپارامتری $f(\cdot)$ را با استفاده از تکنیک هموارسازی کرنل و بکار بردن کرنل نرمال، برآورد می‌کنیم. لذا ماتریس تفاضلی از مرتبه $100 \times (4 - 100)$ بصورت زیر

^{۱۵}Doppler function

خواهد بود

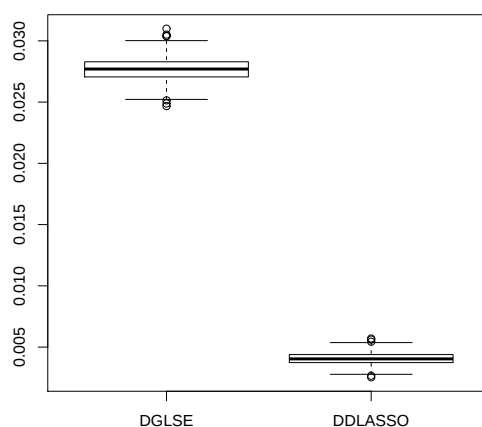
$$D = \begin{pmatrix} d_0 & d_1 & d_2 & d_3 & d_4 & 0 & 0 & \dots & 0 \\ 0 & d_0 & d_1 & d_2 & d_3 & d_4 & 0 & \dots & 0 \\ \vdots & \ddots & & & & & & \ddots & \\ 0 & 0 & \dots & 0 & d_0 & d_1 & d_2 & d_3 & d_4 \end{pmatrix}$$

مخاطره برآوردگرهای پیشنهادی به صورت زیر به دست می‌آید.

$$R(\hat{\beta}; \beta) = \frac{1}{N} \sum_{i=1}^N \left(\hat{\beta}^{(i)} - \beta \right)^\top \left(\hat{\beta}^{(i)} - \beta \right),$$

که $\hat{\beta}^{(i)}$ برآوردگرهای به دست آمده در تکرار m ام می‌باشد.

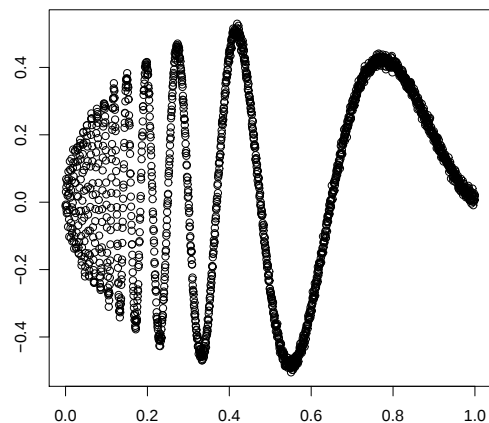
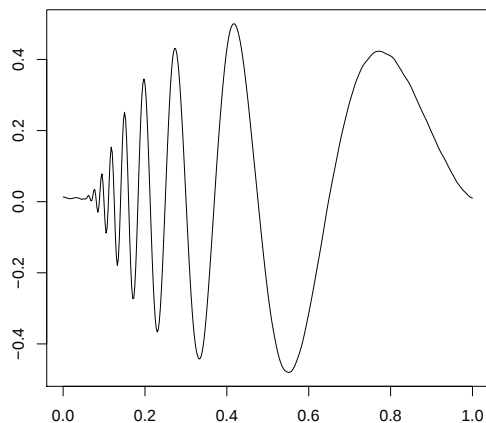
شکل ۲.۳ مقایسه مخاطره برآوردگر لاسوی مشتق‌پذیر تفاضلی را با برآوردگر تفاضلی نشان می‌دهد. بر اساس این نمودار مشخص می‌شود که استفاده از برآوردگر لاسوی مشتق‌پذیر در روش تفاضلی باعث کاهش مخاطره می‌گردد.



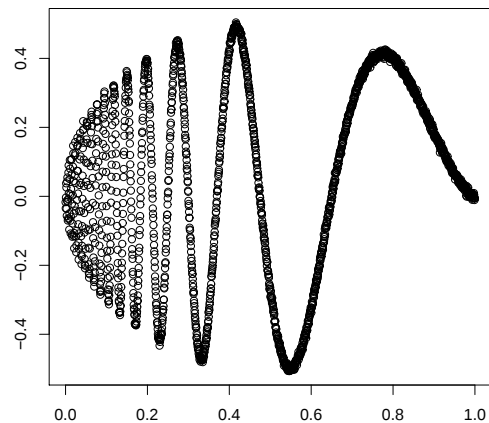
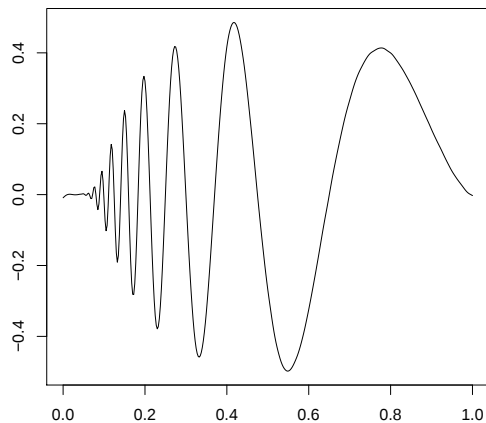
شکل ۲.۳: نمودار جعبه‌ای مقایسه مخاطره برآوردگرهای لاسوی مشتق‌پذیر تفاضلی و کمترین توان‌های دوم تعمیم‌یافته

در شکل ۳.۳ (قاب)، نمودار مانده‌های برآوردشده با روش هموارسازی کرنل پس از برآورد قسمت خطی به وسیله برآوردگر تفاضلی تعمیم یافته یعنی $y - X\hat{\beta}_{GD}^{LS}$ در مقابل t_i ها و منحنی ناپارامتری برازش شده به این مانده‌ها را نشان می‌دهد. نمودارهای قاب پایین با جایگزینی $\hat{\beta}_{GD}^{LS}$ با $\hat{\beta}^{DDLASSO}$ به دست آمده است.

$$\hat{\beta}_{GD}^{LS}$$



$$\hat{\beta}^{DDLASSO}$$



شکل ۳.۳: برآورد تابع داپلر به روش هموار سازی کرنل، بالا سمت راست: مانده‌ها پس از برآورد قسمت خطی با استفاده از $\hat{\beta}_{GD}^{LS}$ ، بالا سمت چپ: برازش منحنی هموار به مانده‌ها با استفاده از روش هموار سازی کرنل و به کار بردن نرمال با پهنای باند بهینه $h = 0.05$ ، پایین مانند قسمت بالاست که در آن به جای $\hat{\beta}_{GD}^{LS}$ از $\hat{\beta}^{DDLASSO}$ استفاده شده است.

فصل ۴

تحلیل مثال‌های واقعی

۱.۴ مقدمه

در این فصل، در دو مجموعه داده واقعی، کاربردی از روش‌های بیان‌شده در این پایان‌نامه را بررسی می‌کنیم. ابتدا، در مجموعه داده‌های پروستات، کارایی روش لاسوی مشتق‌پذیر را با سایر برآوردها مورد مطالعه قرار می‌دهیم. سپس در مجموعه داده‌های مطالعه پنلی روی درآمدهای پویا^۱ PSID (امروز، ۱۹۸۷)، برآوردها لاسوی مشتق‌پذیر تفاضلی را با برآوردهای لاسوی تفاضلی و نیز برآوردها تفاضلی، مورد بررسی قرار می‌دهیم.

۲.۴ تحلیل داده‌های سرطان پروستات

مجموعه داده‌های سرطان پروستات، ارتباط بین سطح آنتی‌ژن اختصاصی پروستات (پادتن مخصوص پروستات) با اندازه‌های بالینی ریشه غده پروستات ۹۷ مرد، که اشعه رادیکال پرستاکتومی را دریافت کرده‌اند، و هشت متغیر زیر را بررسی می‌کند: لگاریتم حجم سرطان (lcanvol)، لگاریتم وزن پروستات (lweight)، لگاریتم مقدار یاخته

^۱Panel study on income dynamics

پروستات خوش‌خیم (lbph)، لگاریتم خصوصیات کپسول نفوذی (lcp)، سن، نمره گلیسون (gleason)، درصد گلیسون ۴ یا ۵ (pgg45)، تهاجم بافتی هسته‌ای (svi). همه متغیرها استاندارد شده‌اند تا میانگین صفر و واریانس یک را داشته باشند و متغیر پاسخ مرکزی شده است که میانگین آن در این صورت صفر است. توصیفی از این داده‌ها در جدول ۱.۴ ارائه شده است.

جدول ۱.۴: جدول توصیف داده‌های سرطان پروستات

متغیرها	کمینه	چارک اول	میانه	میانگین	چارک سوم	بیشینه
lcanvol	-۱/۳۴۷۱	۰/۵۱۲۸	۱/۴۴۶۹	۱/۳۵۰۰	۲/۱۲۷۰	۳/۸۲۱۰
lweight	۲/۳۷۵	۳/۳۷۶	۳/۶۲۳	۳/۶۲۹	۳/۸۷۶	۴/۷۸۰
lbph	-۱/۳۸۶۳	-۱/۳۸۶۳	۰/۳۰۰۱	۰/۱۰۰۴	۱/۵۵۸۱	۲/۳۲۶۳
lcp	-۱/۳۸۶۳	-۱/۳۸۶۳	-۰/۷۹۸۵	-۰/۱۷۹۴	۱/۱۷۸۷	۲/۹۰۴۲
age	۴۱/۰۰	۶۰/۰۰	۶۵/۰۰	۶۳/۸۷	۶۸/۰۰	۷۹/۰۰
gleason	۶/۰۰۰	۶/۰۰۰	۷/۰۰۰	۶/۷۵۳	۷/۰۰۰	۹/۰۰۰
pgg45	۰/۰۰	۰/۰۰	۱۵/۰۰	۲۴/۳۸	۴۰/۰۰	۱۰۰/۰۰
svi	۰/۰۰۰۰	۰/۰۰۰۰	۰/۰۰۰۰	۰/۲۱۶۵	۰/۰۰۰۰	۱/۰۰۰۰
lpsa	-۰/۴۳۰۸	۱/۷۳۱۷	۲/۵۹۱۵	۲/۴۷۸۴	۳/۰۵۶۴	۵/۵۸۲۹

روش‌های لاسو، ریج، SCAD، LS، Enet و لاسوی مشتق‌پذیر روی داده‌ها به کار گرفته شده‌اند تا تصویر واضحی از شباهت برآوردگر لاسوی مشتق‌پذیر با سایر برآوردگرهای جریمه‌شده به دست آید. پارامترهای تنظیم‌کننده با BIC انتخاب شده‌اند. مدل‌ها براساس BIC و AIC مقایسه شده‌اند. تعریف AIC و BIC را در پیوست ۱.۲ ببینید. کارایی برآوردگر لاسوی مشتق‌پذیر به ازای مقدار ثابت $s = ۰/۰۰۱$ با برآوردگرهای لاسو، Enet و SCAD در جدول ۲.۴ مقایسه شده‌اند. برآوردگرهای لاسوی مشتق‌پذیر، لاسو و Enet، هر سه، پنج متغیر انتخاب می‌کنند و بر اساس معیارهای AIC و BIC بر برآوردگر SCAD برتری دارند. جدول ۳.۴ مقدار ضرایب این برآوردگرها را نشان می‌دهد.

اگر برای لاسوی مشتق‌پذیر، s را ثابت و برابر ۱ قرار دهیم، انتظار داریم نتیجه شبیه به برآوردگر ریج باشد. جدول ۴.۴ این ادعا را تایید می‌کند. لازم به ذکر است که هدف از ارائه مقادیر AIC و BIC انتخاب مدل و مقایسه دو روش لاسوی مشتق‌پذیر و ریج نبوده، بلکه هدف این است که نشان دهیم چنانچه پارامتر s مقدار یک را انتخاب کند نتایج لاسوی مشتق‌پذیر مشابه نتایج ریج است. همانطور که از مقادیر AIC و BIC

جدول ۲.۴: مقایسه برآوردهای لاسوی مشتق‌پذیر، لاسو، Enet، SCAD با $s = 0.001$ روی مجموعه داده‌های پروستات.

روش	دقت	AIC	BIC	درجه آزادی	متغیرهای مهم
لاسوی مشتق‌پذیر	$s = 0.001$	۲۰۷/۶	۲۱۶/۱	۵	،pgg45،lweight،lbph،lcavol svi
لاسو	—	۲۰۶/۷	۲۱۵/۳	۵	،pgg45،lweight،lbph،lcavol svi
Enet	$\alpha = 0.001$	۲۰۶/۸	۲۱۵/۳	۵	،pgg45،lweight،lbph،lcavol svi
SCAD	—	۲۱۴/۶	۲۳۱	۴	svi ،lweight،lbph،lcavol

جدول ۳.۴: مقادیر ضرایب برآوردهای لاسوی مشتق‌پذیر ($s = 0.001$)، لاسو، Enet، SCAD روی مجموعه داده‌های پروستات.

	روش			
	لاسوی مشتق‌پذیر	لاسو	Enet	SCAD
lcavol	۰/۵۲۷۳	۰/۵۹۵۵	۰/۵۲۴۶	۰/۶۰۴۱
lweight	۰/۱۴۹۳	۰/۱۴۴۲	۰/۱۴۵۷	۰/۱۹۲۹
age	۰/۰۰۰۰	۰/۰۰۰۰	۰/۰۰۰۰	۰/۰۰۰۰
lbph	۰/۰۶۴۸	۰/۱۳۲۶	۰/۰۵۹۳	۰/۱۳۵۲
svi	۰/۲۰۳۱	۰/۲۷۰۹	۰/۱۹۹۹	۰/۲۷۱۷
lcp	۰/۰۰۰۰	۰/۰۰۰۰	۰/۰۰۰۰	۰/۰۰۰۰
gleason	۰/۰۰۰۰	۰/۰۰۰۰	۰/۰۰۰۰	۰/۰۰۰۰
pgg45	۰/۰۳۷۱	۰/۱۰۶۹	۰/۰۳۴۰	۰/۰۰۰۰

نتیجه می‌شود به ازای $s = 1$ لاسوی مشتق‌پذیر مشابه ريج عمل می‌کند. در جدول ۵.۴ مقدار ضرایب آمده‌اند.

جدول ۴.۴: مقایسه برآوردگر لاسوی مشتق‌پذیر و ريج

متغیرهای مهم	درجه آزادی	BIC	AIC	دقت	روش
همه متغیرها	۸	۲۲۷	۲۰۷/۴۴	۱	لاسوی مشتق‌پذیر
همه متغیرها	۸	۲۲۶	۲۰۷/۱	–	ريج

وقتی در برآوردگر لاسوی مشتق‌پذیر، s را ثابت و برابر ۱۰۰ قرار دهیم انتظار داریم که پاسخ شبیه برآوردگر LS باشد. جدول ۶.۴ انتظار ما را تایید می‌کند و توضیحات مربوط به مقادیر AIC و BIC جدول ۴.۴ در این‌جا نیز صادق است. جدول ۷.۴ مقدار ضرایب را نشان می‌دهد.

مقدار BIC در جدول ۷.۴ نشان می‌دهد که لاسوی مشتق‌پذیر بهتر از برآوردگر LS است، که این موضوع احتمالاً به دلیل تاثیر مقدار کمی پارامتر تنظیم‌کننده برای $s = 100$ است.

جدول ۵.۴: مقادیر ضرایب برآوردهای لاسوی مشتق‌پذیر ($s = 1$) و ريج روی مجموعه داده‌های پروستات.

روش		
ريج	لاسوی مشتق‌پذیر	
۰/۴۶۶۴	۰/۴۶۵۲	lcavol
۰/۱۸۹۷	۰/۱۷۹۷	lweight
–۰/۰۹۵	–۰/۰۷۲۷	age
۰/۱۱۸۴	۰/۱۰۵۷	lbph
۰/۲۴۵۶	۰/۲۲۲۹	svi
–۰/۰۰۴۰	–۰/۰۰۸۵	lcp
۰/۰۳۹۰	۰/۰۲۸۶	gleason
۰/۰۸۲۹	۰/۰۷۱۲	pgg45

جدول ۶.۴: مقایسه برآوردگر لاسوی مشتق‌پذیر و LS

روش	دقت	AIC	BIC	درجه آزادی	متغیرهای مهم
لاسوی مشتق‌پذیر	۱۰۰	۲۰۴	۲۲۸	۸	همه متغیرها
LS	—	۲۰۲	۲۳۳	۸	همه متغیرها

جدول ۷.۴: مقادیر ضرایب برآوردگرهای لاسوی مشتق‌پذیر ($s = 100$) و LS روی مجموعه داده‌های پروستات.

روش		
لاسوی مشتق‌پذیر	LS	
۰/۴۵۱۰	۰/۵۹۹۳	lcavol
۰/۱۸۰۹	۰/۱۹۵۵	lweight
—۰/۰۷۱۴	—۰/۱۲۶۵	age
۰/۱۰۵۲	۰/۱۳۴۵	lbph
۰/۲۲۴۲	۰/۲۷۴۷	svi
—۰/۰۰۱۹	—۰/۱۲۷۷	lcp
۰/۰۲۸۲	۰/۰۲۸۲	gleason
۰/۰۷۰۹	۰/۱۱۰۵	pgg45

۳.۴ مجموعه داده‌های PSID

امروز (۱۹۸۷) از نمونه‌ای در مورد مطالعه پنلی روی درآمدها پویا استفاده نمود تا به‌طور سیستماتیک چندین نظریه و فرضیه‌های آماری را که تا آن زمان در مورد مدل‌های تجربی درآمد زنان بررسی شده بود را مطالعه کند. مجموعه داده‌های PSID به‌طور رایگان از سایت <https://ideas.repec.org/p/boc/bocins/mroz.html> قابل دسترسی است.

این داده‌ها از بین زنان سفیدپوست متاهل آمریکایی بین ۳۰ تا ۶۰ ساله در سال ۱۹۷۵ جمع‌آوری شده است که ۷۵۳ مشاهده و ۱۹ متغیر به‌شرح زیر دارد:

- INLF (اگر برابر ۱ باشد، بدین معنی است که در سال ۱۹۷۵، نیروی کار اجباری بوده است)

- HOURS (ساعت‌های کاری در سال ۱۹۷۵)

- k5 (کودک زیر ۶ سال)
- k618 (کودکان ۶ تا ۱۸ سال)
- AGE (سن زنان بر حسب سال)
- EDUC (سال‌های تحصیل در مدرسه)
- WAGE (برآورد حقوق ساعتی از درآمد)
- REPWAGE (نرخ حقوق گزارش شده در سال ۱۹۷۶)
- HUSHRS (ساعات کاری شوهر در سال ۱۹۷۵)
- HUSAGE (سن شوهر)
- HUSEDUC (سال‌های تحصیل در مدرسه شوهر)
- HUSWAGE (حقوق ساعتی شوهر در سال ۱۹۷۵)
- FAMINC (درآمد خانواده در سال ۱۹۷۵)
- MTR (نرخ مالیات زنان)
- MOTHEduc (سال‌های تحصیل در مدرسه مادران)
- FATHEDUC (سال‌های تحصیل در مدرسه پدران)
- UNEM (نرخ بیکاری در کشور)
- CITY (اگر در شهر SMSA زندگی کند، مقدار ۱ را می‌پذیرد)
- EXPER (تجربه بازار کار)
- NWIFEINC (برابر با $(FAMINC - WAGE)/1000$ است)

مشابه امرز (۱۹۸۷)، متغیر HOURS، ساعات کاری زنان در سال ۱۹۷۵ را به‌عنوان متغیر پاسخ در نظر می‌گیریم. مانند رحیم و همکاران (۲۰۱۲)، با توجه به ماهیت متغیر پاسخ، فقط بخشی از داده‌ها مربوط به زنان شاغل به کار اجباری را در نظر می‌گیریم. بنابراین ۲۴۸ مشاهده در مجموعه داده جدید داریم. مدل پیشنهادی شامل متغیرهای زیر است:

AGE ●

NWIFEINC ●

K5 ●

k618 ●

WC ● (اگر مقدار متغیر EDUC از ۱۲ بیشتر باشد، مقدار ۱ می‌گیرد.)

textsche ● (اگر مقدار متغیر HUSEDUC از ۱۲ بیشتر باشد، مقدار ۱ می‌گیرد.)

UNEM ●

EXPER ●

MTR ●

آن‌ها مدل زیر را پیشنهاد کردند

$$\begin{aligned} hours = & wc\beta_1 + g(nwifeinc) + mtr\beta_2 + exper\beta_3 + unem\beta_4 + k5\beta_5 + age\beta_6 \\ & + k618\beta_7 + hc\beta_8 + \epsilon \end{aligned} \quad (1.4)$$

در اینجا متغیر NWIFEINC مولفه ناپارامتری است. هال و همکاران (۱۹۹۰) مقادیر بهینه ماتریس D را در روش تفاضلی با روش‌های عددی به‌دست آوردند. در این جا از روش تفاضلی مرتبه $m = 6$ استفاده می‌شود. بر اساس مقادیر پیشنهادی هال و همکاران (۱۹۹۰)،

$$d = (0 : 0.0906; 0 : 0.1926; 0 : 0.1369; 0 : 0.1635; 0 : 0.1925; 0 : 0.2238; 0 : 0.9200) \quad (2.4)$$

و در نتیجه ماتریس D ساخته می‌شود.

در ادامه، روش اعتبارسنجی متقابل K تایی برای به‌دست آوردن برآوردی از خطاهای پیش‌گویی مدل استفاده می‌شود. در این روش، مجموعه داده را به‌صورت تصادفی به K زیرمجموعه نسبتاً مساوی تقسیم می‌کنیم. یکی از مجموعه‌ها را کنار می‌گذاریم، یعنی $\{(X_D^{\text{test}}, y_D^{\text{test}})\}$ ، که آن را مجموعه آزمون می‌نامیم و از $K-1$ مجموعه باقیمانده که آن‌ها را مجموعه آموزشی می‌نامیم برای برآزش مدل استفاده می‌کنیم. برآوردگر به‌دست آمده را $\hat{\beta}^{\text{train}}$ می‌نامیم. مدل برآزش شده برای پیش‌گویی متغیر پاسخ در

مجموعه داده آزمون به کار می‌رود. سرانجام خطاهای پیش‌گویی بر اساس توان دوم خطای مقادیر مشاهده‌شده از مقادیر پیش‌گویی‌شده به دست می‌آید. یعنی

$$PE^k = \|y_{Dk}^{\text{test}} - \hat{y}_{Dk}^{\text{test}}\|^2; \quad k = 1, \dots, K,$$

که در آن $\hat{y}_{Dk}^{\text{test}} = X_{Dk}^{\text{test}} \hat{\beta}_{Dk}^{\text{train}}$

برآیند برای همه K زیرمجموعه تکرار می‌شود و خطاهای پیش‌گویی محاسبه می‌شوند. برای اثر دادن تغییرات تصادفی اعتبارسنجی متقابل، فرآیند N بار تکرار می‌شود و متوسط خطای پیش‌گویی^۲ (APE) به صورت زیر محاسبه می‌شود

$$APE = \frac{1}{N} \sum_{i=1}^N \left(\frac{1}{K} \sum_{k=1}^K PE_i^k \right),$$

که PE_i^k خطای پیش‌گویی k امین مجموعه آزمون در i امین تکرار است. روش اعتبارسنجی متقابل به طور مفصل‌تری در پیوست ۲.۱.۲ آمده است. کارایی برآوردگر دلخواه $\hat{\beta}^*$ نسبت به برآوردگر لاسوی مشتق‌پذیر تفاضلی، $\hat{\beta}^{\text{D-LASSO}}$ ، با فرمول کارایی^۳ (Eff) به صورت زیر به دست می‌آید.

$$\text{Eff}(\hat{\beta}^*; \hat{\beta}^{\text{D-LASSO}}) = \frac{APE(\hat{\beta}^{\text{D-LASSO}})}{APE(\hat{\beta}^*)}.$$

اگر مقدار کارایی از ۱ بیشتر باشد، آن‌گاه $\hat{\beta}^{\text{D-LASSO}}$ بهتر از برآوردگر $\hat{\beta}^*$ رفتار می‌کند.

جدول ۸.۴ برآوردگرها و کارایی آن‌ها را به‌ازای $N = 5000$ نشان می‌دهد. برای برآورد قسمت پارامتری، روش تفاضلی مرتبه ششم ($m = 6$) به کار برده شد. اکنون از روش کرنل با پهنای $h_n = 1/2$ برای برآورد متغیر hours استفاده می‌شود. به منظور برآورد مولفه ناپارامتری، ابتدا پارامتر را با روش تفاضلی برآورد می‌کنیم و سپس رهیافت کرنل را برای برازش $Z_i(\hat{\beta}^*) = \text{hours}_i - x_{Di}^\top \hat{\beta}^*$ روی nwifcinc استفاده می‌کنیم که $x_i = (\text{exper}, \text{mtr}, \text{wc}, \text{age}, \text{unem}, \text{k5}, \text{k6}, \text{k18}, \text{hc})$ و $\hat{\beta}^*$ یکی از برآوردگرهای تفاضلی، لاسوی تفاضلی و لاسوی مشتق‌پذیر تفاضلی است (شکل ۱.۴).

همان‌طور که در جدول ۸.۴ دیده می‌شود، با توجه به این که مقدار Eff برای لاسوی مشتق‌پذیر تفاضلی $1/12$ است، این برآوردگر کاراتر از هر دو برآوردگر تفاضلی و لاسوی تفاضلی است. این نتیجه یکی از مهم‌ترین یافته‌های این پایان‌نامه می‌باشد که نشان می‌دهد در مدل نیمه‌پارامتری لاسوی مشتق‌پذیر تفاضلی برآوردگری کاراتر

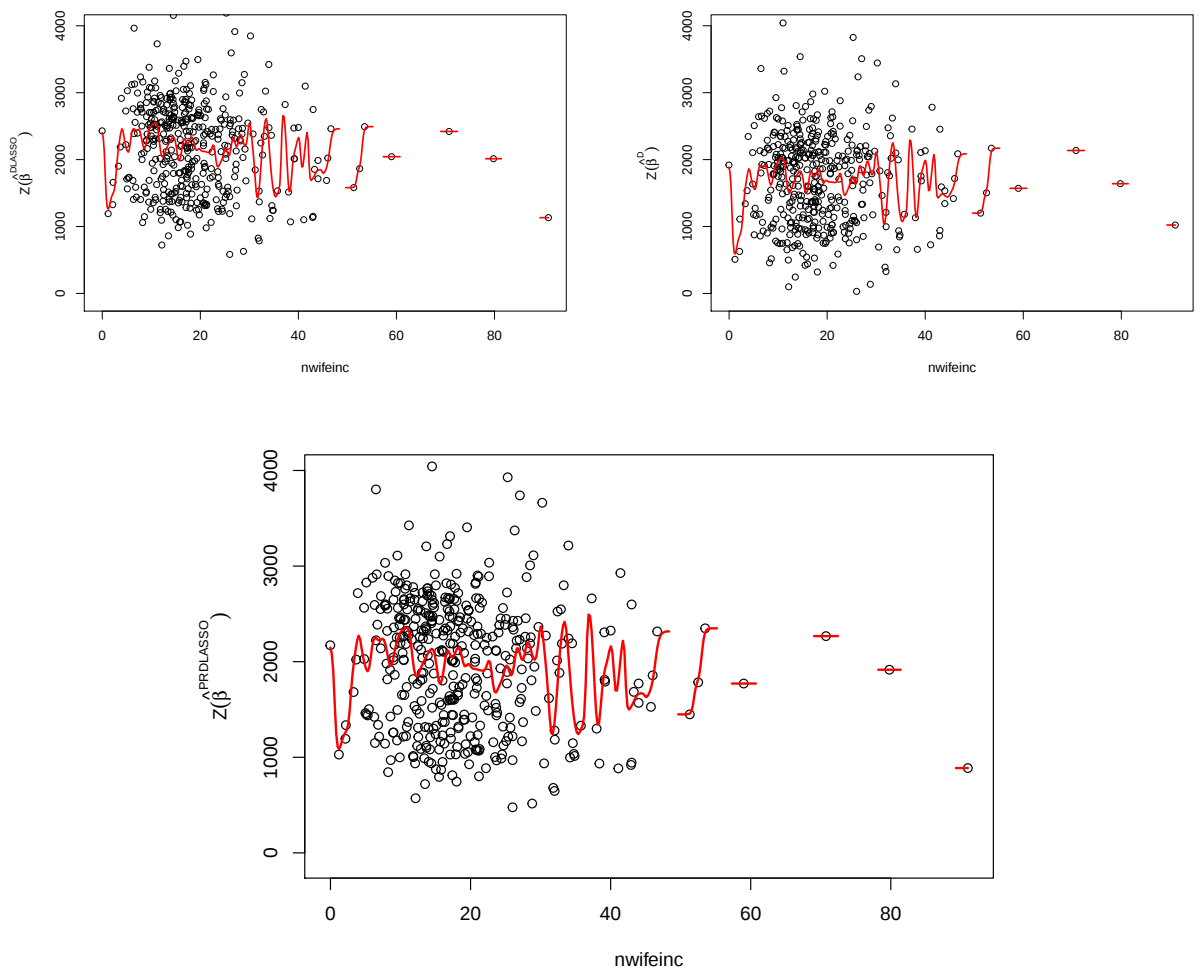
^۲ Average prediction error

^۳ Efficiency

جدول ۸.۴: مقادیر برآوردگرها و کارایی آنها

متغیرها	تفاضلی	لاسوی تفاضلی	لاسوی مشتق‌پذیر تفاضلی
exper	۳۰/۲۶	۲۳/۵۷	۳۱/۸۷
mtr	-۴۱۶/۰۰	-۱۲۱۳/۵۰	-۷۰۷/۱۸
wc	-۶۸/۸۴	-۴۵/۵۶	-۱۱۰/۴۳
age	-۸/۳۶	-۵/۹۲	-۸/۷۵
unem	-۹/۸۵	-۹/۱۹	-۴/۱۰
k۵	-۱۲۸/۱	-۱۳۲/۱۱	-۱۳۹/۴۲۸۴
k۶۱۸	-۴۳/۰۶	-۲۵/۵۶	-۳۶/۱۸
hc	۲۲/۲۲	۰/۰۰	۲۴۴۳
Eff	۱/۰۰	۱/۰۵	۱/۱۲

نسبت به سایر برآوردگرهای جریمه‌شده برای انتخاب متغیر بر اساس متوسط خطای پیشگویی مدل است.



شکل ۱.۴: نمودار برآورد تابع ناپارامتری

پیوست آ

تعاریف و قضایای اساسی

آ.۱۰ مرکزی کردن مشاهدات

در این قسمت نشان می‌دهیم با مرکزی کردن مشاهدات در مدل رگرسیون خطی ساده، می‌توان اثر پارامتر عرض از مبدأ β_0 را حذف کرد. مدل (۱.۱) را در نظر بگیرید. با کم کردن میانگین مشاهدات از طرفین مدل رگرسیون خطی داریم

$$(y_i - \bar{y}) + \bar{y} = \beta_0 + \beta(x_i - \bar{x}) + \beta\bar{x} + \epsilon_i.$$

در نتیجه

$$(y_i - \bar{y}) = (\beta_0 + \beta\bar{x} - \bar{y}) + \beta(x_i - \bar{x}) + \epsilon_i. \quad (1.A)$$

همچنین با میانگین‌گیری از معادله (۱.۱) داریم

$$\bar{y} = \beta_0 + \beta\bar{x}. \quad (2.A)$$

با جای‌گذاری معادله (۲.آ) در (۱.آ) نتیجه می‌شود

$$(y_i - \bar{y}) = (\beta_0 + \beta\bar{x} - \beta_0 - \beta\bar{x}) + \beta(x_i - \bar{x}) + \epsilon_i$$

$$= \beta(x_i - \bar{x}) + \epsilon_i.$$

فرض کنید $y_i^* = (y_i - \bar{y})$ و $x_i^* = (x_i - \bar{x})$. در این صورت مدل نهایی با رابطه زیر داده می‌شود

$$y_i^* = \beta x_i^* + \epsilon_i \quad (۳.آ)$$

همانطور که در مدل (۳.آ) مشاهده می‌شود اثر β_0 از مدل حذف شده است. به طور مشابه می‌توان در مدل رگرسیون چندگانه نیز اثر β_0 را حذف کرد.

تعریف ۱.۰. (نرم- L_p). فرض کنید $x \in \mathbb{R}^1$ یک بردار و $p \geq 1$ یک عدد صحیح باشد. سپس

$$\|x\|_p = \left(\sum_{i=1}^n |x_i|^p \right)^{\frac{1}{p}} \quad (۴.آ)$$

نرم- L_p از بردار x نامیده می‌شود. واضح است

$$\|x\|_p^p = \sum_{i=1}^n |x_i|^p \quad (۵.آ)$$

تعریف ۲.۰. (تابع خطا). تابع خطا به صورت زیر تعریف می‌شود

$$\operatorname{erf} z = \frac{2}{\sqrt{\pi}} \int_0^z \exp\{-t^2\} dt$$

و اگر $z = \frac{x}{s}$

$$\operatorname{erfc} z = \frac{2}{\sqrt{\pi}} \int_z^\infty \exp\{-t^2\} dt$$

از رامیز و سانچزا (۲۰۱۴) رابطه زیر را داریم

$$\frac{1}{x + \sqrt{x^2 + \frac{4}{\pi}}} < \exp\left\{\left(\frac{x}{s}\right)^2\right\} \int_{\frac{x}{s}}^\infty \exp\{-t^2\} dt \leq \frac{1}{x + \sqrt{x^2 + \frac{4}{\pi}}}$$

$$\frac{2}{\sqrt{\pi}} \frac{e^{-\left(\frac{x}{s}\right)^2}}{\left(\frac{x}{s}\right) + \sqrt{\left(\frac{x}{s}\right)^2 + \frac{4}{\pi}}} < \operatorname{erfc}\left(\frac{x}{s}\right) \leq \frac{2}{\sqrt{\pi}} \frac{e^{-\left(\frac{x}{s}\right)^2}}{\left(\frac{x}{s}\right) + \sqrt{\left(\frac{x}{s}\right)^2 + \frac{4}{\pi}}}$$

۱.آ تابع محدب

فرض کنیم $-\infty \leq a \leq b \leq \infty$ ، تابع $f : (a, b) \rightarrow \mathbb{R}$ را محدب می‌گوییم در صورتی که به ازای هر دو عدد $x_1, x_2 \in (a, b)$ و هر t که $0 \leq t \leq 1$ ، داشته باشیم:

$$\forall x_1, x_2 \in X, \forall t \in [0, 1]: \quad f(tx_1 + (1-t)x_2) \leq tf(x_1) + (1-t)f(x_2).$$

اگر در تعریف بالا تساوی را برداریم آن گاه f را اکیدا محدب می نامیم.

$$\forall x_1 \neq x_2 \in X, \forall t \in (0, 1) : \quad f(tx_1 + (1-t)x_2) \leq tf(x_1) + (1-t)f(x_2).$$

خاصیت‌ها:

$$R(x_1, x_2) = \frac{f(x_1) - f(x_2)}{x_1 - x_2} \quad (6. \bar{A})$$

$$\forall x_1, x_2 \in C : f\left(\frac{x_1 + x_2}{2}\right) \leq \frac{f(x_1) + f(x_2)}{2} \quad (7. \bar{A})$$

$$f(x) \geq f(y) + f'(y)(x - y) \quad (8. \bar{A})$$

$$f(tx) = f(tx + (1-t)0) \leq tf(x) + (1-t)f(0) \leq tf(x), \forall t \in [0, 1]. \quad (9. \bar{A})$$

$$\begin{aligned} f(a) + f(b) &= f\left((a+b)\frac{a}{a+b}\right) + f\left((a+b)\frac{b}{a+b}\right) \\ &\leq \frac{a}{a+b}f(a+b) + \frac{b}{a+b}f(a+b) = f(a+b) \end{aligned} \quad (10. \bar{A})$$

همچنین چنانچه تابع $f : (a, b) \rightarrow \mathbb{R}$ مشتق پذیر باشد، می توان نشان داد این تابع محدب است اگر و تنها اگر f' صعودی باشد. چنانچه f دارای مشتق دوم باشد، این بدان معناست که تابع f محدب است اگر و تنها اگر $f'' \geq 0$.

۱.۱. سری تیلور

تعریف ۱.۱.۱. $(\limsup)$. دنباله X_n را در نظر بگیرید. در این صورت

$$\begin{aligned} \limsup_{n \rightarrow \infty} X_n &= \lim_{n \rightarrow \infty} \left(\sup_{m \geq n} X_m \right) \\ &= \lim_{n \rightarrow \infty} \left(\sup X_n \right) \\ &= \inf_{n \geq 0} \sup_{m \geq n} X_m \\ &= \inf \sup X_m : m \geq n : n \geq 0 \end{aligned}$$

یا به طور معادل، می توان فرض کرد دنباله J_n به صورت زیر تعریف شده باشد.

$$\begin{aligned} J_n &= \sup \left\{ X_m : m \in \{n, n+1, n+2, \dots\} \right\} \\ &= \bigcup_{m=n}^{\infty} X_m = X_n \cup X_{n+1} \cup X_{n+2} \cup \dots \end{aligned}$$

دنباله J_n یک دنباله ناصعودی ($J_n \supseteq J_{n+1}$) است، زیرا J_{n+1} اجتماع مجموعه‌های کوچکتر از J_n است. بزرگترین کران پایین روی این دنباله از اجتماع X_n ها برابر است با

$$\begin{aligned}\limsup_{n \rightarrow \infty} X_n &= \inf \{ \sup \{ X_m : m \in \{n, n+1, n+2, \dots\} \} : n \in \{1, 2, \dots\} \} \\ &= \bigcap_{n=1}^{\infty} (\bigcup_{m=n}^{\infty} X_m)\end{aligned}$$

تعریف ۲.۱.آ. فرض کنید X_n و Y_n دو دنباله از متغیرهای تصادفی باشند. اگر $\limsup_{n \rightarrow \infty} |\frac{X_n}{Y_n}| < \infty$ به عبارت دیگر اگر $n \rightarrow \infty$ ، عبارت $|\frac{X_n}{Y_n}|$ کران دار باشد می‌نویسیم $X_n = O(Y_n)$.

تعریف ۳.۱.آ. فرض کنید X_n و Y_n دو دنباله از متغیرهای تصادفی باشند. اگر

$$\lim_{n \rightarrow \infty} \left| \frac{X_n}{Y_n} \right| = 0$$

می‌نویسیم $X_n = o(Y_n)$.

در ریاضیات علامت O بزرگ رفتار حدی یک تابع را وقتی آرگومان‌های آن به یک عدد خاص یا به بی نهایت میل می‌کند را بیان می‌کند علامت O بزرگ به کاربر این اجازه را می‌دهد که تابع را ساده کند تا بر روی نرخ رشد متمرکز شود؛ بنابراین توابع مختلف با نرخ رشد یکسان می‌توانند دارای یک علامت O مشابه باشند. تابع زیر را بر حسب n در نظر بگیرید

$$T(n) = 4n^2 - 5n + 7$$

اگر تعداد ورودی این مسأله به بی نهایت میل کند اندازه جمله n^2 بسیار نزدیک تر از دیگر جمله‌ها خواهد بود. در این صورت گفته می‌شود $T(n) = O(n^2)$ یا $T(n) \in O(n^2)$. در ریاضیات معمولاً برای نشان دادن این که یک سری هندسی متناهی تا چه اندازه به تابع مورد نظر نزدیک است خصوصاً در سری تیلور ناقص از این علامت استفاده می‌شود.

تعریف ۴.۱.آ. فرض کنید f یک تابع حقیقی-مقدار روی R باشد و همچنین $x \in R$. اگر برای بعضی δ ها تابع f دارای مشتق‌های پیوسته تا مرتبه‌ی p در بازه $(x - \delta, x + \delta)$ باشد، آن گاه برای هر دنباله‌ی α_n همگرا به صفر، داریم

$$\begin{aligned}f(x + \alpha_n) &= \sum_{j=1}^p \left(\frac{\alpha_n^j}{j!} \right) f^{(j)}(x) + o(\alpha_n^p) \\ &= \sum_{j=1}^p \left(\frac{\alpha_n^j}{j!} \right) f^{(j)}(x) + O(\alpha_n^{p+1})\end{aligned}$$

۲.۱.۱ روش اعتبارسنجی متقابل

بدیهی است که روش های جریمه شده، به پارامتر تنظیم کننده λ وابسته هستند پس انتخاب این پارامتر برای به کارگیری تابع جریمه از اهمیت زیادی برخوردار است. روش های زیادی برای یافتن آن وجود دارد اما روشی که بیشتر مورد استفاده قرار می گیرد روش اعتبارسنجی متقابل^۱ (CV) است. این روش برای تخمین بهترین مقدار λ به کار می رود. در این روش علاقه مندیم پارامتر λ را طوری پیدا کنیم که میانگین توان دوم خطای پیشگویی^۲ (MPE) کمترین مقدار را داشته باشد. لازم به ذکر است در مدل (۳.۱) معیار MPE به صورت زیر محاسبه می شود

$$MPE = E(\mathbf{y} - \hat{\mathbf{y}})^\top (\mathbf{y} - \hat{\mathbf{y}})$$

که در آن $\hat{\mathbf{y}} = \mathbf{X}\hat{\beta}$. کار اعتبارسنجی متقابل به صورت زیر است. ابتدا داده ها را به k زیرگروه مساوی تقسیم می کنیم. یک گروه به عنوان برازش $\hat{\beta}$ و $k-1$ گروه به عنوان ارزیابی $\hat{y} = \hat{\beta}x$ در نظر گرفته می شود که خطای مقدار برازش شده مشاهدات محاسبه می شود. این روند روی هر قسمت تکرار می شود و در آخر میانگین این خطاها محاسبه می شود. معمولاً k را ۵، ۱۰ و n در نظر می گیرند. در این حالت معیار CV به صورت زیر تعریف می شود

$$CV(\lambda) = \frac{1}{n} \sum_{i=1}^k (y_i - \hat{y}_{\lambda}^{-i})^2$$

که در آن $\hat{Y}_{\lambda}^{-i}$ بیانگر مقدار برازش شده پس از کنار گذاشتن i امین گروه مشاهدات است. در این صورت λ مورد نظر متناظر با λ یی خواهد بود که این معیار را مینیمم کند.

محاسبه $CV(\lambda)$ بسیار وقت گیر است. کروان و واهبا (۱۹۷۹) معیار اعتبارسنجی تعمیم یافته^۳ (GCV) را به عنوان جایگزینی برای $CV(\lambda)$ به صورت زیر معرفی کردند

$$GCV(\lambda) = \frac{(\mathbf{y} - \hat{\mathbf{Y}}_{\lambda})^\top (\mathbf{Y} - \hat{\mathbf{Y}}_{\lambda})}{n(1 - \frac{tr(A(\lambda))}{n})^2}$$

که در آن، $\hat{Y}_{\lambda}$ بردار مقادیر برازش شده تحت ماتریس تصویر $A(\lambda)$ است. انتخاب پارامتر جریمه بخش مهمی از برازش مدل در رگرسیون جریمه شده است. روش های بسیاری در این مورد وجود دارد اما مهمترین روش های کاربردی بر اساس این است که برآوردگر

^۱Cross-validation

^۲Mean prediction error

^۳Generalized cross-validation

جریمه شده چه قدر خوب می تواند مقدار واقعی متغیر پاسخ را پیش گویی کند. منصفانه نیست از داده ها دوبار استفاده شود، یک بار برای برازش مدل و دوباره از همان داده ها برای برآورد دقت پیشگویی. این کار، منجر به بیش برازشی می شود.

ایده ای که به ذهن می رسد، تقسیم بندی داده ها به دو مجموعه است. از یک مجموعه (مجموعه آموزش) برای برازش برآوردگر جریمه شده (مثلاً β^*) استفاده شده و از مجموعه دیگر (مجموعه آزمون) برای محاسبه اینکه کمیت $X\beta^*$ چه قدر خوب توانسته است مشاهدات مجموعه دیگر را پیش گویی کند. اما به ندرت تعداد داده ها آن قدر زیاد است که بتوان آن ها را به دو قسمت تقسیم کرد به طوری که هر یک از مجموعه ها به تنهایی برای انتخاب پارامتر جریمه کافی باشند. از این رو، روش اعتبارسنجی متقابل پیشنهاد می شود.

در اعتبارسنجی متقابل، داده ها به طور تصادفی به K زیرمجموعه تقسیم می شود. یک زیرمجموعه کنار گذاشته می شود که نام مجموعه آزمون را به خود می گیرد، $K-1$ مجموعه باقیمانده، مجموعه آموزشی نامیده و برای برازش مدل استفاده می گردد. سپس مدل برازش شده به منظور پیش گویی متغیرهای پاسخ در مجموعه داده آزمون به کار برده می شود. سرانجام، خطاهای پیش گویی با میانگین توان دوم انحراف مقادیر مشاهده شده از مقادیر پیش گویی در مجموعه داده آزمون به دست می آیند. معمولاً $K = 5, 10, n$ انتخاب می شود.

۲.۱ قضیه لایب نیتز

مشتق انتگرال: هرگاه بخواهیم از تابعی شامل انتگرال، مشتق بگیریم باید از قاعده زیر استفاده کنیم که مستقیماً فرمول محاسبه مشتق انتگرال را به ما می دهد. مشتق انتگرال $\int_{m(x)}^{n(x)} h(t) dt$ برابر است با:

$$\left(\int_{m(x)}^{n(x)} h(t) dt \right)' = m'(x).h(m(x)) - n'(x).h(n(x)) \quad (11.A)$$

یعنی مشتق تابعی که شامل انتگرالی با یک تابع تک متغیره مانند $h(t)$ و کران هایی بر حسب x مانند $m(x)$ (کران بالا) و $n(x)$ (کران پایین) باشد برابر است با مشتق کران بالا ضربدر تابع داخل انتگرال که کران بالا به جای متغیر آن جایگزین شده است منهای مشتق کران پایین ضربدر تابع داخل انتگرال که کران پایین به جای متغیر آن جایگزین شده باشد.

مثال ۱.۲.آ. مشتق انتگرال

$$f(x) = \int_{Ln(x)}^{x^2+2x} \frac{2t}{3-t} dt \quad (12.آ)$$

را بیابید.

حل.

$$\begin{aligned} f(x) &= \int_{Ln(x)}^{x^2+2x} \frac{2t}{3-t} dt \\ \Rightarrow f'(x) &= (x^2+2) \cdot \frac{2(x^2+2x)}{3-(x^2+2x)} - \frac{1}{x} \cdot \frac{2Ln(x)}{3-Ln(x)} \end{aligned}$$

□

قضیه ۱.۲.آ (قضیه انتگرال گیری جزء به جزء). فرض کنید u و v توابعی دیفرانسیل پذیر باشد در این صورت دیفرانسیل تابع uv به صورت زیر محاسبه می شود

$$d(uv) = u dv + v du$$

که می توان آن را به صورت زیر نوشت

$$u dv = d(uv) - v du$$

پس از انتگرال گیری از آن داریم

$$\int u dv = uv - \int v du \quad (13.آ)$$

در فرمول انتگرال گیری جزء به جزء، انتگرال $\int u dv$ بر حسب انتگرال $\int v du$ بیان می شود. با انتخاب مناسب u و v ممکن است محاسبه انتگرال دوم ساده تر از محاسبه انتگرال اول باشد. اهمیت این فرمول در همین است. وقتی با انتگرالی رو به رو می شویم که نمی توانیم آن را حل کنیم، باید آن را به انتگرال دیگری تبدیل کنیم تا بتوانیم آن را حل کنیم.

۱.۲.آ معیار AIC و BIC

معیار آکائیک (AIC) به صورت زیر تعریف می شود

$$AIC = -2l + 2p$$

که در آن، l لگاریتم تابع درست نمایی و p تعداد ضریب رگرسیونی موجود در مدل می‌باشد. هر مدل زیرمجموعه با AIC کمتر، بر سایر مدل‌های زیر مجموعه برتری دارد. پس از این معیار، معیارهای دیگری پدید آمده اند که معروف‌ترین آنها BIC است و به صورت زیر تعریف می‌شود

$$BIC = -2l + p \ln(n)$$

که در آن، n حجم نمونه است.

پیوست ب

گزیده‌ای از برنامه‌های رایانه‌ای

نمودار ۱.۲

```
library(philentropy)
library(latex2exp)
X=seq(-2,2,0.05)
s = seq(0.0001,3,length.out = 1000)

F1 = function(x,s)      sqrt((x^2)+s)
F2 = function(x,s)      sqrt((x^2)+(s^2))
F3 = function(x,s)      (x^2)/sqrt((x^2)+(s^2))
F4 = function(x,s) log(2+exp(-x/s)+exp(x/s))*s
F5 = function(x,s)      x*(2*pnorm(x/s,0,1/sqrt(2),lower.tail = TRUE)-1)

D1 <- D2 <- D3 <- D4 <- D5 <- numeric(length(s))
for(i in 1:length(D1)){
```

```

D1[i] = distance(rbind(abs(X),F1(X,s[i])), method = "euclidean")
D2[i] = distance(rbind(abs(X),F2(X,s[i])), method = "euclidean")
D3[i] = distance(rbind(abs(X),F3(X,s[i])), method = "euclidean")
D4[i] = distance(rbind(abs(X),F4(X,s[i])), method = "euclidean")
D5[i] = distance(rbind(abs(X),F5(X,s[i])), method = "euclidean")
}

plot(s,D1,type = "l",ylim = c(0,28),lwd = 1.5 , col = "black",ylab = "distance")
lines(s,D2,type = "l",lwd = 1.5 , col = "green")
lines(s,D3,type = "l",lwd = 1.5 , col = "blue")
lines(s,D4,type = "l",lwd = 1.5 , col = "purple")
lines(s,D5,type = "l",lwd = 2 , col = "red")

expression.legend = c(TeX("$\\sqrt{x^2+2}$"),TeX("$\\sqrt{x^2+s^2}$"),
TeX("$x^2/\\sqrt{x^2+s^2}$"),TeX("$\\log(2+exp(-x/s)+exp(x/s))s$"),
TeX("$x(2 \\Phi(x/s,0,1/\\sqrt{2})-1)$"))
legend(0.05,28,expression.legend , col = c("black","green","blue","purple","red"),
lwd = c(1.5,1.5,1.5,1.5,2), lty = 1)

x = seq(-1,1,BY = 0.05)
x2= function(x) x^2
s1 = 0.01
f1 = function(x) x*(2*pnorm(x/s1,0,1/sqrt(2))-1)
s2=2/sqrt(pi)
f2=function(x) x*(2*pnorm(x/s2,0,1/sqrt(2))-1)

curve(f2,-1,1,col = "red",lwd = 2,xlab = "x" ,ylab = "p(x,s)",ylim = c(0,1.1))
curve(f1, -1,1,col = "orange" , lwd = 2, add= TRUE)
curve(x2, -1,1,col = "blue" , lwd = 1.5,add = TRUE)
curve(abs(x),-1,1,col="black",lwd=2,add=TRUE , lty = 3)

expression.legend = c(TeX("$x(2 \\Phi(x/s,0,1/\\sqrt{2})-1)$"), s = 2/\\sqrt{\\pi}$"),
TeX("$x(2 \\Phi(x/s,0,1/\\sqrt{2})-1)$", s = 0.01$"),TeX("$x^2$"),TeX("$|x|$"))

```

```
legend(-0.5,1,col=c("red","green","blue","black"), expression.legend,
lty = c(1,1,1,3),lwd=c(2,2,1.5,2),text.width = 0.7)
```

نمودار ۲.۲

```
library(latex2exp)
thresholding.function <- function(beta,s = 0.01){
.5*(2*pnorm(beta/s,0,1/sqrt(2))-1+2*(beta/s)*dnorm(beta/s,0,1/sqrt(2)))
+ beta }
f.inv <- inverse(thresholding.function, -1,1)

y.vec = seq(-1.5,1.5,0.01)

D = numeric(length(y.vec))
for(i in 1:length(y.vec))
D[i] = f.inv(y.vec[i])
plot(y.vec,D, type = "l",lwd = 2 , xlab = TeX("$x_i$") ,
ylab = TeX("$\\hat{\\beta}_i$"))
title(TeX("$ s = 0.01$ (lasso) $"))
abline(a = 0 , b = 1 , lty = 3)
abline(h = 0, lty = 2)
abline(v = -0.5, lty = 2)
abline(v = 0.5, lty = 2)

thresholding.function <- function(beta,s = 0.5){
.5*(2*pnorm(beta/s,0,1/sqrt(2))-1+2*(beta/s)*dnorm(beta/s,0,1/sqrt(2))) + beta
}
f.inv <- inverse(thresholding.function, -1,1)

y.vec = seq(-1.5,1.5,0.01)

D = numeric(length(y.vec))
for(i in 1:length(y.vec))
```

```
D[i] = f.inv(y.vec[i])
plot(y.vec,D, type = "l",lwd = 2 , xlab = TeX("$x_i$") ,
     ylab = TeX("$\\hat{\\beta}_i$"))
title(TeX("$ s = 0.5 $"))
abline(a = 0 , b = 1 , lty = 3)
abline(h = 0, lty = 2)
abline(v = -0.5, lty = 2)
abline(v = 0.5, lty = 2)

thresholding.function <- function(beta,s = 1){
  .5*(2*pnorm(beta/s,0,1/sqrt(2))-1+2*(beta/s)*dnorm(beta/s,0,1/sqrt(2)))
  + beta }
f.inv <- inverse(thresholding.function, -1,1)

y.vec = seq(-1.5,1.5,0.01)

D = numeric(length(y.vec))
for(i in 1:length(y.vec))
D[i] = f.inv(y.vec[i])
plot(y.vec,D, type = "l",lwd = 2 , xlab = TeX("$x_i$") ,
     ylab = TeX("$\\hat{\\beta}_i$"))
title(TeX("$ s = 1 (ridge) $"))
abline(a = 0 , b = 1 , lty = 3)
abline(h = 0, lty = 2)
abline(v = -0.5, lty = 2)
abline(v = 0.5, lty = 2)

thresholding.function <- function(beta,s = 20){
  .5*(2*pnorm(beta/s,0,1/sqrt(2))-1+2*(beta/s)*dnorm(beta/s,0,1/sqrt(2))) + beta
}
f.inv <- inverse(thresholding.function, -2,2)
```

```
y.vec = seq(-1.5,1.5,0.01)

D = numeric(length(y.vec))
for(i in 1:length(y.vec))
D[i] = f.inv(y.vec[i])
plot(y.vec,D, type = "l",lwd = 2 , xlab = TeX("$x_i$") ,
     ylab = TeX("$\\hat{\\beta}_i$"))
title(TeX("$ s = 20$ (OLS) $"))
abline(a = 0 , b = 1 , lty = 3)
abline(h = 0, lty = 2)
abline(v = -0.5, lty = 2)
abline(v = 0.5, lty = 2)
```

نمودار ۳.۲

```
library(MASS)
library(DLASSO)
library(ncvreg)
library(glmnet)
library(msgps)
library(rms)
library(lasso2)

#Senario 1
sigm2 = 3
s = sqrt(sigm2)
p = 8
beta = matrix(c(3,1.5,0,0,2,0,0,0),nr = 8, nc = 1)
Sigma = matrix(0,nr = p, nc = p)

for(i in 1:p){
  for(j in 1:p){
    Sigma[i,j] = 0.5^(abs(i-j))
```

```
}
}
```

```
n = 240 #number of observations
mu = rep(0,p)
X = mvrnorm(n , mu = mu, Sigma = Sigma)
X.train = X[1:40,]
X.test = X[41:240,]

N.Sim = 50
DLASSO.mat <- OLS.mat <- ridge.mat <- lasso.mat <-enet.mat <-
SCAD.mat <- matrix(0,nr = p, nc = N.Sim)
for(i.sim in 1:N.Sim){
e = rnorm(n,mean = 0, sd = 1) #error
e.train = e[1:40]
e.test = e[41:240]
y.train = X.train%%beta + s*e.train
y.train <- c(y.train)
y.test = X.test%%beta + s*e.test
dLasso = dlasso(x=X.train,y=y.train,adp = FALSE , s = 0.01)
DLASSO.mat[,i.sim] = coef(dLasso)[,4]
ridge=cv.glmnet(x=X.train,y=y.train,alpha = 0)
ridge.mat[,i.sim] = coef(ridge,s = ridge$lambda.min)[-1]
lasso=fit <- msgps(X.train,y.train)
lasso.mat[,i.sim]=coef(lasso)[-1,4]
enet=msgps(X.train,y.train,penalty="enet",alpha=0.5)
enet.mat[,i.sim]=coef(enet)[-1,4]
scad = cv.ncvreg(X.train,y.train,penalty = "SCAD")
SCAD.mat[,i.sim]<-ncvreg(X.train,y.train,penalty = "SCAD",
lambda = scad$lambda.min)$beta[-1]
ols = lm(y.train ~ X.train -1)
OLS.mat[,i.sim] = coef(ols)
```

```
}
```

```
y.ols <- y.scad <- y.dlasso <- y.enet <- y.lasso <-  
  y.ridge <- matrix(0, nr = 200, nc = N.Sim)  
for(i.test in 1:N.Sim){  
  y.ols[,i.test] <- X.test %*% OLS.mat[,i.test]  
  y.scad[,i.test] <- X.test %*% SCAD.mat[,i.test]  
  y.dlasso[,i.test] <- X.test %*% DLASSO.mat[,i.test]  
  y.enet[,i.test] <- X.test %*% enet.mat[,i.test]  
  y.lasso[,i.test] <- X.test %*% lasso.mat[,i.test]  
  y.ridge[,i.test] <- X.test %*% ridge.mat[,i.test]  
}
```

```
MSE <- matrix(0, nr = N.Sim, nc = 6)  
for(i.mse in 1:N.Sim){  
  MSE[i.mse,1] <- sum((y.ols[,i.mse]-y.test)^2)/200 #OLS  
  MSE[i.mse,2] <- sum((y.ridge[,i.mse]-y.test)^2)/200 #ridge  
  MSE[i.mse,3] <- sum((y.dlasso[,i.mse]-y.test)^2)/200 #DLASSO  
  MSE[i.mse,4] <- sum((y.lasso[,i.mse]-y.test)^2)/200 #lasso  
  MSE[i.mse,5] <- sum((y.enet[,i.mse]-y.test)^2)/200 #enet  
  MSE[i.mse,6] <- sum((y.scad[,i.mse]-y.test)^2)/200 #scad  
}  
colnames(MSE)<- c("OLS","Ridge","DLASSO","LASSO","ENET","SCAD")  
boxplot(MSE,ylim = c(2,6))
```

```
MSE.beta <- matrix(0, nr = N.Sim, nc = 6)  
S.X <- cov(X)  
for(i.mse.beta in 1:N.Sim){  
  MSE.beta[i.mse.beta,1] <- t(OLS.mat[,i.mse.beta]-beta)%*%S.X%*%  
    (OLS.mat[,i.mse.beta]-beta)#OLS  
  MSE.beta[i.mse.beta,2] <- t(ridge.mat[,i.mse.beta]-beta)%*%S.X%*%  
    (ridge.mat[,i.mse.beta]-beta) #ridge
```

```

MSE.beta[i.mse.beta,3] <- t(DLASSO.mat[,i.mse.beta]-beta)%*%S.X%*%
(DLASSO.mat[,i.mse.beta]-beta) #DLASSO
MSE.beta[i.mse.beta,4] <- t(lasso.mat[,i.mse.beta]-beta)%*%S.X%*%
(lasso.mat[,i.mse.beta]-beta) #lasso
MSE.beta[i.mse.beta,5] <- t(enet.mat[,i.mse.beta]-beta)%*%S.X%*%
(enet.mat[,i.mse.beta]-beta) #enet
MSE.beta[i.mse.beta,6] <- t(SCAD.mat[,i.mse.beta]-beta)%*%S.X%*%
(SCAD.mat[,i.mse.beta]-beta) #scad
}
colnames(MSE.beta)<- c("OLS","Ridge","DLASSO","LASSO","ENET","SCAD")
boxplot(MSE.beta,ylim = c(0,3))

#Senario 2
sigm2 = 3
s = sqrt(sigm2)
p = 8
beta2=rep(0.5,8)
Sigma = matrix(0,nr = p, nc = p)

for(i in 1:p){
  for(j in 1:p){
    Sigma[i,j] = 0.5^(abs(i-j))
  }
}

n = 240 #number of observations
mu = rep(0,p)
X = mvrnorm(n , mu = mu, Sigma = Sigma)
X.train = X[1:40,]
X.test = X[41:240,]

N.Sim = 50

```

```

DLASSO.mat <- OLS.mat <- ridge.mat <- lasso.mat <-enet.mat <-SCAD.mat <-
  matrix(0,nr = p, nc = N.Sim)
for(i.sim in 1:N.Sim){
e = rnorm(n,mean = 0, sd = 1) #error
e.train = e[1:40]
e.test = e[41:240]
y.train = X.train%*%beta2 + s*e.train
y.train <- c(y.train)
y.test = X.test%*%beta2 + s*e.test
dLasso = dlasso(x=X.train,y=y.train,adp = FALSE , s = 0.01)
DLASSO.mat[,i.sim] = coef(dLasso)[,4]
ridge=cv.glmnet(x=X.train,y=y.train,alpha = 0)
ridge.mat[,i.sim] = coef(ridge,s = ridge$lambda.min)[-1]
lasso=fit <- msgps(X.train,y.train)
lasso.mat[,i.sim]=coef(lasso)[-1,4]
enet=msgps(X.train,y.train,penalty="enet",alpha=0.5)
enet.mat[,i.sim]=coef(enet)[-1,4]
scad = cv.ncvreg(X.train,y.train,penalty = "SCAD")
SCAD.mat[,i.sim]<-ncvreg(X.train,y.train,penalty = "SCAD",
lambda = scad$lambda.min)$beta[-1]
ols = lm(y.train ~ X.train -1)
OLS.mat[,i.sim] = coef(ols)
}

```

```

y.ols <- y.scad <- y.dlasso <- y.enet <- y.lasso <- y.ridge <-
  matrix(0, nr = 200, nc = N.Sim)
for(i.test in 1:N.Sim){
y.ols[,i.test] <- X.test %*% OLS.mat[,i.test]
y.scad[,i.test] <- X.test %*% SCAD.mat[,i.test]
y.dlasso[,i.test] <- X.test %*% DLASSO.mat[,i.test]
y.enet[,i.test] <- X.test %*% enet.mat[,i.test]
y.lasso[,i.test] <- X.test %*% lasso.mat[,i.test]

```

```

y.ridge[,i.test] <- X.test %*% ridge.mat[,i.test]
}

MSE <- matrix(0, nr = N.Sim, nc = 6)
for(i.mse in 1:N.Sim){
MSE[i.mse,1] <- sum((y.ols[,i.mse]-y.test)^2)/200 #OLS
MSE[i.mse,2] <- sum((y.ridge[,i.mse]-y.test)^2)/200 #ridge
MSE[i.mse,3] <- sum((y.dlasso[,i.mse]-y.test)^2)/200 #DLASSO
MSE[i.mse,4] <- sum((y.lasso[,i.mse]-y.test)^2)/200 #lasso
MSE[i.mse,5] <- sum((y.enet[,i.mse]-y.test)^2)/200 #enet
MSE[i.mse,6] <- sum((y.scad[,i.mse]-y.test)^2)/200 #scad
}

colnames(MSE)<- c("OLS","Ridge","DLASSO","LASSO","ENET","SCAD")
boxplot(MSE,ylim = c(2,6))

MSE.beta <- matrix(0, nr = N.Sim, nc = 6)
S.X <- cov(X)
for(i.mse.beta in 1:N.Sim){
MSE.beta[i.mse.beta,1] <- t(OLS.mat[,i.mse.beta]-beta2)%*%S.X%*%
(OLS.mat[,i.mse.beta]-beta2) #OLS
MSE.beta[i.mse.beta,2] <- t(ridge.mat[,i.mse.beta]-beta2)%*%S.X%*%
(ridge.mat[,i.mse.beta]-beta2) #ridge
MSE.beta[i.mse.beta,3] <- t(DLASSO.mat[,i.mse.beta]-beta2)%*%S.X%*%
(DLASSO.mat[,i.mse.beta]-beta2) #DLASSO
MSE.beta[i.mse.beta,4] <- t(lasso.mat[,i.mse.beta]-beta2)%*%S.X%*%
(lasso.mat[,i.mse.beta]-beta2) #lasso
MSE.beta[i.mse.beta,5] <- t(enet.mat[,i.mse.beta]-beta2)%*%S.X%*%
(enet.mat[,i.mse.beta]-beta2) #enet
MSE.beta[i.mse.beta,6] <- t(SCAD.mat[,i.mse.beta]-beta2)%*%S.X%*%
(SCAD.mat[,i.mse.beta]-beta2) #scad
}

colnames(MSE.beta)<- c("OLS","Ridge","DLASSO","LASSO","ENET","SCAD")

```

```
boxplot(MSE.beta,ylim = c(0,4))
#####
#Senario 3
sigma3=15
s=sqrt(sigma3)
p=15
beta1=c(1,2,3,4,5)
beta2=rep(0.5,5)
beta3=rep(0,5)
beta=matrix(c(beta1,beta2,beta3),nr=p,nc=1)
Sigma = matrix(0,nr = p, nc = p)

for(i in 1:p){
  for(j in 1:p){
    Sigma[i,j] = 0.5^(abs(i-j))
  }
}

n = 240 #number of observations
mu = rep(0,p)
X = mvrnorm(n , mu = mu, Sigma = Sigma)
X.train = X[1:40,]
X.test = X[41:240,]

N.Sim = 50
DLASSO.mat <- OLS.mat <- ridge.mat <- lasso.mat <-enet.mat <-SCAD.mat <-
matrix(0,nr = p, nc = N.Sim)
for(i.sim in 1:N.Sim){
  e = rnorm(n,mean = 0, sd = 1) #error
  e.train = e[1:40]
  e.test = e[41:240]
  y.train = X.train%*%beta + s*e.train
```

```

y.train <- c(y.train)
y.test = X.test%*%beta + s*e.test
dLasso = dlasso(x=X.train,y=y.train,adp = FALSE , s = 0.01)
DLASSO.mat[,i.sim] = coef(dLasso)[,4]
ridge=cv.glmnet(x=X.train,y=y.train,alpha = 0)
ridge.mat[,i.sim] = coef(ridge,s = ridge$lambda.min)[-1]
lasso=fit <- msgps(X.train,y.train)
lasso.mat[,i.sim]=coef(lasso)[-1,4]
enet=msgps(X.train,y.train,penalty="enet",alpha=0.5)
enet.mat[,i.sim]=coef(enet)[-1,4]
scad = cv.ncvreg(X.train,y.train,penalty = "SCAD")
SCAD.mat[,i.sim]<-ncvreg(X.train,y.train,penalty = "SCAD",
lambda = scad$lambda.min)$beta[-1]
ols = lm(y.train ~ X.train -1)
OLS.mat[,i.sim] = coef(ols)
}

```

```

y.ols <- y.scad <- y.dlasso <- y.enet <- y.lasso <-
y.ridge <- matrix(0, nr = 200, nc = N.Sim)
for(i.test in 1:N.Sim){
y.ols[,i.test] <- X.test %*% OLS.mat[,i.test]
y.scad[,i.test] <- X.test %*% SCAD.mat[,i.test]
y.dlasso[,i.test] <- X.test %*% DLASSO.mat[,i.test]
y.enet[,i.test] <- X.test %*% enet.mat[,i.test]
y.lasso[,i.test] <- X.test %*% lasso.mat[,i.test]
y.ridge[,i.test] <- X.test %*% ridge.mat[,i.test]
}

```

```

MSE <- matrix(0, nr = N.Sim, nc = 6)
for(i.mse in 1:N.Sim){
MSE[i.mse,1] <- sum((y.ols[,i.mse]-y.test)^2)/200 #OLS
MSE[i.mse,2] <- sum((y.ridge[,i.mse]-y.test)^2)/200 #ridge

```

```

MSE[i.mse,3] <- sum((y.dlasso[,i.mse]-y.test)^2)/200 #DLASSO
MSE[i.mse,4] <- sum((y.lasso[,i.mse]-y.test)^2)/200 #lasso
MSE[i.mse,5] <- sum((y.enet[,i.mse]-y.test)^2)/200 #enet
MSE[i.mse,6] <- sum((y.scad[,i.mse]-y.test)^2)/200 #scad
}

colnames(MSE)<- c("OLS","Ridge","DLASSO","LASSO","ENET","SCAD")
boxplot(MSE,ylim = c(15,40))

MSE.beta <- matrix(0, nr = N.Sim, nc = 6)
S.X <- cov(X)
for(i.mse.beta in 1:N.Sim){
MSE.beta[i.mse.beta,1] <- t(OLS.mat[,i.mse.beta]-beta)%*%S.X%*%
(OLS.mat[,i.mse.beta]-beta) #OLS
MSE.beta[i.mse.beta,2] <- t(ridge.mat[,i.mse.beta]-beta)%*%S.X%*%
(ridge.mat[,i.mse.beta]-beta) #ridge
MSE.beta[i.mse.beta,3] <- t(DLASSO.mat[,i.mse.beta]-beta)%*%S.X%*%
(DLASSO.mat[,i.mse.beta]-beta) #DLASSO
MSE.beta[i.mse.beta,4] <- t(lasso.mat[,i.mse.beta]-beta)%*%S.X%*%
(lasso.mat[,i.mse.beta]-beta) #lasso
MSE.beta[i.mse.beta,5] <- t(enet.mat[,i.mse.beta]-beta)%*%S.X%*%
(enet.mat[,i.mse.beta]-beta) #enet
MSE.beta[i.mse.beta,6] <- t(SCAD.mat[,i.mse.beta]-beta)%*%S.X%*%
(SCAD.mat[,i.mse.beta]-beta) #scad
}

colnames(MSE.beta)<- c("OLS","Ridge","DLASSO","LASSO","ENET","SCAD")
boxplot(MSE.beta,ylim = c(0,20))

#####
#Senario 4
sigm2 = 3
s = sqrt(sigm2)
p = 300
u=rep(c(3,1.5,0,0,2,0),c(1,1,1,1,1,295))

```

```
beta = matrix(u,nr = 300, nc = 1)
Sigma = matrix(0,nr = p, nc = p)

for(i in 1:p){
  for(j in 1:p){
    Sigma[i,j] = 0.5^(abs(i-j))
  }
}

n = 240 #number of observations
mu = rep(0,p)
X = mvrnorm(n , mu = mu, Sigma = Sigma)
X.train = X[1:40,]
X.test = X[41:240,]

N.Sim = 50
DLASSO.mat <- OLS.mat <- ridge.mat <- lasso.mat <-enet.mat <-
SCAD.mat <- matrix(0,nr = p, nc = N.Sim)
for(i.sim in 1:N.Sim){
  e = rnorm(n,mean = 0, sd = 1) #error
  e.train = e[1:40]
  e.test = e[41:240]
  y.train = X.train%*%beta + s*e.train
  y.train <- c(y.train)
  y.test = X.test%*%beta + s*e.test
  dLasso = dlasso(x=X.train,y=y.train,adp = FALSE , s = 0.01)
  DLASSO.mat[,i.sim] = coef(dLasso)[,4]
  ridge=cv.glmnet(x=X.train,y=y.train,alpha = 0)
  ridge.mat[,i.sim] = coef(ridge,s = ridge$lambda.min)[-1]
  lasso=fit <- msgps(X.train,y.train)
  lasso.mat[,i.sim]=coef(lasso)[-1,4]
  enet=msgps(X.train,y.train,penalty="enet",alpha=0.5)
```

```
enet.mat[,i.sim]=coef(enet)[-1,4]
scad = cv.ncvreg(X.train,y.train,penalty = "SCAD")
SCAD.mat[,i.sim]<-ncvreg(X.train,y.train,penalty = "SCAD",
lambda = scad$lambda.min)$beta[-1]
ols = lm(y.train ~ X.train -1)
OLS.mat[,i.sim] = coef(ols)
}

y.ols <- y.scad <- y.dlasso <- y.enet <- y.lasso <-
y.ridge <- matrix(0, nr = 200, nc = N.Sim)
for(i.test in 1:N.Sim){

y.ols[,i.test] <- X.test %*% OLS.mat[,i.test]
y.scad[,i.test] <- X.test %*% SCAD.mat[,i.test]
y.dlasso[,i.test] <- X.test %*% DLASSO.mat[,i.test]
y.enet[,i.test] <- X.test %*% enet.mat[,i.test]
y.lasso[,i.test] <- X.test %*% lasso.mat[,i.test]
y.ridge[,i.test] <- X.test %*% ridge.mat[,i.test]
}

MSE <- matrix(0, nr = N.Sim, nc = 6)
for(i.mse in 1:N.Sim){
MSE[i.mse,1] <- sum((y.ols[,i.mse]-y.test)^2)/200 #OLS
MSE[i.mse,2] <- sum((y.ridge[,i.mse]-y.test)^2)/200 #ridge
MSE[i.mse,3] <- sum((y.dlasso[,i.mse]-y.test)^2)/200 #DLASSO
MSE[i.mse,4] <- sum((y.lasso[,i.mse]-y.test)^2)/200 #lasso
MSE[i.mse,5] <- sum((y.enet[,i.mse]-y.test)^2)/200 #enet
MSE[i.mse,6] <- sum((y.scad[,i.mse]-y.test)^2)/200 #scad
}

colnames(MSE)<- c("OLS","Ridge","DLASSO","LASSO","ENET","SCAD")
boxplot(MSE,ylim = c(2,22))
```

```

MSE.beta <- matrix(0, nr = N.Sim, nc = 6)
S.X <- cov(X)
for(i.mse.beta in 1:N.Sim){
MSE.beta[i.mse.beta,1] <- t(OLS.mat[,i.mse.beta]-beta)%*%S.X%*%
(OLS.mat[,i.mse.beta]-beta)#OLS
MSE.beta[i.mse.beta,2] <- t(ridge.mat[,i.mse.beta]-beta)%*%S.X%*%
(ridge.mat[,i.mse.beta]-beta) #ridge
MSE.beta[i.mse.beta,3] <- t(DLASSO.mat[,i.mse.beta]-beta)%*%S.X%*%
(DLASSO.mat[,i.mse.beta]-beta) #DLASSO
MSE.beta[i.mse.beta,4] <- t(lasso.mat[,i.mse.beta]-beta)%*%S.X%*%
(lasso.mat[,i.mse.beta]-beta) #lasso
MSE.beta[i.mse.beta,5] <- t(enet.mat[,i.mse.beta]-beta)%*%S.X%*%
(enet.mat[,i.mse.beta]-beta) #enet
MSE.beta[i.mse.beta,6] <- t(SCAD.mat[,i.mse.beta]-beta)%*%S.X%*%
(SCAD.mat[,i.mse.beta]-beta) #scad
}
colnames(MSE.beta)<- c("OLS","Ridge","DLASSO","LASSO","ENET","SCAD")
boxplot(MSE.beta,ylim = c(0,20))

```

شکل ۱.۳

```

n = 2000
t = numeric(n)
for(i in 1:n){
t[i] = (i - 0.5)/n
}
f = function(x)
sqrt(x*(1-x))*sin((2.1*pi)/(x+0.05))
f <- Vectorize(f)
f.t = matrix(f(t),nc = 1)

plot(t,f.t,type = "l",xlab = "t", ylab = "f")

```

```
library(mvtnorm)
library(DLASSO)
library(KernSmooth)

n = 100
p = 5

N.sim = 1000

beta = matrix(c(1.5,2,3,-5,4),nc = 1)
mu = c(-2,1,2,3,2)
Sigma = matrix(c(0.81,0.4,0.3,-0.2,-0.1,0.4,2.25,0.4,0.3,-0.2,
0.3,0.4,1,0.4,0.3,-0.2,0.3,0.4,0.64,0.4,
-0.1,-0.2,0.3,0.4,0.49), nr = 5 , nc = 5 , byrow = TRUE)
X = rmvnorm(n,mean = mu, sigma = Sigma)
t = numeric(n)
for(i in 1:n){
t[i] = (i - 0.5)/n }
f = function(x)
sqrt(x*(1-x))*sin((2.1*pi)/(x+0.05))
f <- Vectorize(f)
f.t = matrix(f(t),nc = 1)

V = matrix(0, nr = n, nc = n)
for(i in 1:n){
for(j in 1:n){
V[i,j] = (1/n)^(abs(i-j)+1) } }

sigma.e = 0.01
S.V = sigma.e * V
```

```

mu.e = rep(0,n)

d = c(0.8873,-0.3099,-0.2464,-0.1901,-0.1409)
m = length(d) -1
I.nm <- diag(1,n-m,n-m)
D = matrix(0,nr = (n-m), nc = n)
for(i in 1:(n-m)){
D[i,i:(i+m)] <- d }

DLS.mat <- DDLASSO.mat <- matrix(0,nr = p,nc = N.sim)
e.DLS <- e.DDLASSO <- matrix(0,nr = n, nc = N.sim)
for(i.sim in 1:N.sim){
e = matrix(rmvnorm(1,mu.e,S.V),nc = 1)
y = X%%beta + f.t + e

y.D = D %% y
X.D = D %% X

V.D = D %% S.V %% t(D)
C.D = t(X.D)%%solve(V.D)%%X.D
DLS.mat[,i.sim] = solve(C.D)%%t(X.D)%%solve(V.D)%%y.D

dLasso = dlasso(x=X.D,y=y.D,adp = FALSE , s = 0.01)
DDLASSO.mat[,i.sim] = coef(dLasso)[,4]

e.DLS[,i.sim] <- y - X%%DLS.mat[,i.sim]
e.DDLASSO[,i.sim] <- y - X%%DDLASSO.mat[,i.sim] }

R.DLS <- R.DDLASSO <- numeric(N.sim)
for(i.sim in 1:N.sim){
R.DLS[i.sim] <- t(DLS.mat[,i.sim]-beta)%%(DLS.mat[,i.sim]-beta)
R.DDLASSO[i.sim] <- t(DDLASSO.mat[,i.sim]-beta)%%(DDLASSO.mat[,i.sim]-beta)}

```

```
R<-cbind(R.DLS,R.DDLASSO)
colnames(R)<- c("DGLSE","DDLASSO")
boxplot(R)
```

شکل ۳.۳

```
library(mvtnorm)
library(DLASSO)
library(KernSmooth)
n = 2000
p = 5

beta = matrix(c(1.5,2,3,-5,4),nc = 1)
mu = c(-2,1,2,3,2)
Sigma = matrix(c(0.81,0.4,0.3,-0.2,-0.1,
0.4,2.25,0.4,0.3,-0.2,
0.3,0.4,1,0.4,0.3,
-0.2,0.3,0.4,0.64,0.4,
-0.1,-0.2,0.3,0.4,0.49),
nr = 5 , nc = 5 , byrow = TRUE)
X = rmvnorm(n,mean = mu, sigma = Sigma)
t = numeric(n)
for(i in 1:n){
t[i] = (i - 0.5)/n }
f = function(x)
sqrt(x*(1-x))*sin((2.1*pi)/(x+0.05))
f <- Vectorize(f)
f.t = matrix(f(t),nc = 1)

V = matrix(0, nr = n, nc = n)
for(i in 1:n){
```

```
for(j in 1:n){
V[i,j] = (1/n)^(abs(i-j)+1) } }

sigma.e = 0.01
S.V = sigma.e * V
mu.e = rep(0,n)

d = c(0.8873,-0.3099,-0.2464,-0.1901,-0.1409)
m = length(d) -1
I.nm <- diag(1,n-m,n-m)
D = matrix(0,nr = (n-m), nc = n)
for(i in 1:(n-m)){
D[i,i:(i+m)] <- d }

e = matrix(rmvnorm(1,mu.e,S.V),nc = 1)
y = X%%beta + f.t + e

y.D = D %% y
X.D = D %% X

V.D = D %% S.V %% t(D)
C.D = t(X.D)%%solve(V.D)%%X.D
DLS = solve(C.D)%%t(X.D)%%solve(V.D)%%y.D

dLasso = dlasso(x=X.D,y=y.D,adp = FALSE , s = 0.01)
DDLASSO = coef(dLasso)[,4]

e.DLS <- y - X%%DLS.mat[,i.sim]
e.DDLASSO <- y - X%%DDLASSO.mat[,i.sim]

x <- as.matrix(seq(0,1,length = n))
```

```
y.DLS <- as.matrix(e.DLS[,1])
plot(x,y.DLS, xlab = " ", ylab = " ")

fit.DLS <- locpoly(x,y.DLS,bandwidth = dpill(x,y),degree = 1)
plot(fit.DLS, type = "l", xlab = " ", ylab = " ")

y.DDLASSO <- as.matrix(e.DDLASSO[,1])
plot(x,y.DDLASSO, xlab = " ", ylab = " ")

fit.DDLASSO <- locpoly(x,y.DDLASSO,bandwidth = dpill(x,y),degree = 1)
plot(fit.DDLASSO, type = "l", xlab = " ", ylab = " ")
```

جدول ۱.۴

```
library(lasso2)
data("Prostate")
summary(Prostate)
```

جدول های ۲.۴ و ۳.۴

```
library(lasso2)
library(DLASSO)
library(glmnet)
library(msgps)
library(ncvreg)

data("Prostate")
summary(Prostate)

attach(Prostate)
prostate = as.data.frame(Prostate)
prostate[,1:8] = scale(prostate[,1:8])
```

```

prostate[,9] = scale(prostate[,9],scale = FALSE)

x = as.matrix(prostate[,-9])
y = as.matrix(prostate[,9])

dLasso = dlasso(x,y, adp = FALSE, s = 0.001)
AIC.dlasso = dLasso[which(dLasso[,3] == min(dLasso[,3])),3]
BIC.dlasso = dLasso[which(dLasso[,5] == min(dLasso[,5])),5]
AIC.dlasso
BIC.dlasso
DLASSO = coef(dLasso)[,2]
DLASSO

fit.lasso <- glmnet(x, y, family="gaussian", alpha=1)
fit.lasso.cv <- cv.glmnet(x, y, type.measure="mse", alpha=1,
family="gaussian")
lasso = coef(fit.lasso,s = fit.lasso.cv$lambda.min)
lasso

tLL <- fit.lasso$nulldev - deviance(fit.lasso)
k <- fit.lasso$df
n <- fit.lasso$noobs
AICc <- -tLL+2*k+2*k*(k+1)/(n-k-1)
min(AICc)

BIC<-log(n)*k - tLL
min(BIC)

y = as.vector(y)
enet=msgps(x,y,penalty="enet",alpha=0.001)
summary(enet)

```

```
enet=coef(enet)[-1,4]
enet

scad = cv.ncvreg(x,y,penalty = "SCAD")
SCAD<-ncvreg(x,y,penalty = "SCAD",lambda = scad$lambda.min)$beta[-1]
summary(scad)
SCAD
```

جدول‌های ۴.۴ و ۵.۴

```
library(lasso2)
library(DLASSO)
library(glmnet)
library(msgps)
library(ncvreg)

data("Prostate")
summary(Prostate)

attach(Prostate)

prostate = as.data.frame(Prostate)
prostate[,1:8] = scale(prostate[,1:8])
prostate[,9] = scale(prostate[,9],scale = FALSE)

x = as.matrix(prostate[, -9])
y = as.matrix(prostate[,9])
dLasso = dlasso(x,y, adp = FALSE, s = 1)
AIC.dlasso = dLasso[which(dLasso[,3] == min(dLasso[,3])),3]
BIC.dlasso = dLasso[which(dLasso[,5] == min(dLasso[,5])),5]
AIC.dlasso
BIC.dlasso
DLASSO = coef(dLasso)[,2]
DLASSO
```

```

ridge=cv.glmnet(x,y,alpha = 0)
ridge = coef(ridge,s = ridge$lambda.min)[-1]

tLL <- ridge$nulldev - deviance(ridge)
k <- fit.lasso$df
n <- fit.lasso$nobs
AICc <- -tLL+2*k+2*k*(k+1)/(n-k-1)
min(AICc)

BIC<-log(n)*k - tLL
min(BIC)

ridge

```

جدول‌های ۴.۴ و ۵.۴

```

library(lasso2)
library(DLASSO)
library(glmnet)
library(msgps)
library(ncvreg)

data("Prostate")
summary(Prostate)

attach(Prostate)
prostate = as.data.frame(Prostate)
prostate[,1:8] = scale(prostate[,1:8])
prostate[,9] = scale(prostate[,9],scale = FALSE)

x = as.matrix(prostate[,-9])
y = as.matrix(prostate[,9])

```

```
dLasso = dlasso(x,y, adp = FALSE, s = 100)
AIC.dlasso = dLasso[which(dLasso[,3] == min(dLasso[,3])),3]
BIC.dlasso = dLasso[which(dLasso[,5] == min(dLasso[,5])),5]
AIC.dlasso
BIC.dlasso
DLASSO = coef(dLasso)[,2]
DLASSO
```

```
ols = lm(y ~ x -1)
ols
AIC(ols)
BIC(ols)
```

جدول ٨.٤

```
setwd("E:/R")
library(haven)
library(latex2exp)
library(DLASSO)
library(msgps)

devtools::install_github('stefano-meschiari/latex2exp')
mroz <- read_dta("mroz.dta")
attach(mroz)

data = as.data.frame(subset(mroz, inlf == 1))
data$wc <- ifelse(data$educ > 12,1,0)
data$hc <- ifelse(data$huseduc > 12,1,0)
attach(data)

data.new <- subset(data, select= c(hours,age,nwifeinc,kidslt6,kidsge6,
exper,unem,mtr,wc,hc))
```

```

n <- nrow(data.new)
p <- ncol(data.new)-2

d = c(0.9302,-0.1965,-0.1728,-0.1506,-0.1299,-0.1107,-0.0930,-0.0768)

m = length(d) -1
I.nm <- diag(1,n-m,n-m)
D = matrix(0,nr = (n-m), nc = n)
for(i in 1:(n-m)){
  D[i,i:(i+m)] <- d
}
y = data.new$hours
X = as.matrix(data.new[,-c(1,3)])
y.D = D %*% y
X.D = D %*% X
colnames(X.D) <- c("age","k5","k618","exper","unem","mtr","wc","hc")
colnames(y.D) <- "hours"

B.D = solve(t(X.D)%*%M %*% X.D)%*%t(X.D)%*%M %*%y.D

B.LS0 = msgps(X.D,y.D,alpha = 0)
B.LS0=coef(B.LS0)[-1,4]

dLasso = dlasso(X.D,y.D, adp = FALSE, s = 0.001)
B.DLS0 = coef(dLasso)[,2]

d = c(0.9200,-0.2238,-0.1925,-0.1635,-0.1369,-0.1926,-0.0906)
# d = c(0.9302,-0.1965,0.1728,-0.1506,-0.1299,-0.1107,-0.0930,-0.0768)
m = length(d) -1

```

```

alpha=0.05
N.cv<-500
k<-10 #k-fold cross validation
PE.D.mat=PE.LSO.mat=PE.DLSO.mat=matrix(0,nc=k,nr=N.cv)
for(j in 1:N.cv){
  data.new$id<-rep(1:k,length=n)
  data.new$id<-sample(data.new$id)
  list <- 1:k
  for(i.cv in 1:k){
    train<-subset(data.new, id %in% list[-i.cv])
    test<-subset(data.new, id %in% list[i.cv])
    X.train = as.matrix(train[,-c(1,3,11)])
    y.train<-as.matrix(train[,1])
    n.train <- nrow(X.train)
    I.train.nm <- diag(1,n.train-m,n.train-m)
    D.train = matrix(0,nr = (n.train-m), nc = n.train)
    for(i in 1:(n.train-m)){
      D.train[i,i:(i+m)] <- d
    }
    y.train.D = D.train %*% y.train
    X.train.D = D.train %*% X.train
    colnames(X.train.D) <- c("age","k5","k618","exper","unem","mtr","wc","hc")
    colnames(y.train.D) <- "hours"

    X.test<-as.matrix(test[,-c(1,3,11)])
    y.test<-as.matrix(test[,1])
    n.test <- nrow(X.test)
    D.test = matrix(0,nr = (n.test-m), nc = n.test)
    for(i in 1:(n.test-m)){
      D.test[i,i:(i+m)] <- d
    }
    y.test.D = D.test %*% y.test
  }
}

```

```

X.test.D = D.test %% X.test
colnames(X.test.D) <- c("age", "k5", "k618", "exper", "unem", "mtr", "wc", "hc")
colnames(y.test.D) <- "hours"

p = ncol(X.train.D)
M.train = I.train.nm
- X.train.D %% solve(t(X.train.D)%%X.train.D)%%t(X.train.D)

B.D = solve(t(X.train.D)%%M.train %% X.train.D)%%t(X.train.D)%%M.train
%%y.train.D

PE.D.mat[j,i.cv]<-drop(t(X.test.D%%B.D -y.test.D)%%(X.test.D%%B.D -y.test.D))

y.train.D <- as.vector(y.train.D)
B.LSO = msgps(X.train.D,y.train.D,alpha = 0)
B.LSO=coef(B.LSO)[-1,4]

PE.LSO.mat[j,i.cv]<-t(X.test.D%%B.LSO -y.test.D)
%%(X.test.D%%B.LSO -y.test.D)

dLasso = dlasso(X.train.D,y.train.D, adp = FALSE, s = 0.001)
B.DLSO = coef(dLasso)[,2]
PE.DLSO.mat[j,i.cv]<-t(X.test.D%%B.DLSO -y.test.D)%%(X.test.D%%B.DLSO - y.test.D)
}
}
PE.D=apply(PE.D.mat,1,mean)
PE.LSO.mat=apply(PE.LSO.mat,1,mean)
PE.DLSO=apply(PE.DLSO.mat,1,mean)

```

```

result1=cbind(mean(PE.D),mean(PE.LSO),mean(PE.DLSO))
result2=cbind(sd(PE.D),sd(PE.LSO),sd(PE.DLSO))/sqrt(N.cv)
result=rbind(result1,result2)
colnames(result)<-c("FM","SM","PT","JS","PR")
rownames(result)<-c("mean","sd")
result[,1]/result

```


مراجع

- [۱] الوین، ا. (۱۹۳۴) (ترجمه حسنعلی آذرنوش و ابوالقاسم بزرگ نیا)، مدل‌های خطی برای آمار، دانشگاه فردوسی مشهد (۱۳۸۲) شماره ۳۶۴.
- [۲] آرست، م. (۱۳۹۵). مقایسه رفتار برخی برآوردگرهای انقباضی بریج در مدل رگرسیون چندگانه، پایان‌نامه کارشناسی ارشد، دانشگاه صنعتی شاهرود.
- [۳] برزویی بیدگلی، م. (۱۳۹۳). بررسی رفتار برآوردگر جک‌نایف ریج، پایان‌نامه کارشناسی ارشد، دانشگاه صنعتی شاهرود.
- [۴] بغلانی، م. (۱۳۹۴). برآوردگرهای تفاضلی استوار در مدل‌های رگرسیون نیمه پارامتریک، پایان‌نامه کارشناسی ارشد، دانشگاه سمنان.
- [۵] توماس، ج.، فینی، ر. (۱۹۸۸) (ترجمه مهدی بهزاد، سیامک کاظمی، علی کافی)، حساب دیفرانسیل و انتگرال و هندسه تحلیلی تهران، مرکز نشر دانشگاهی (۱۳۷۳-۱۳۷۴).
- [۶] روزبه، م. (۱۳۹۰) برآورد در مدل‌های خطی جزئی، رساله دکتری، دانشگاه فردوسی مشهد.
- [۷] گلپایگانی، م. (۱۳۹۵) مدل رگرسیون نرخ خطر کاکس جریمه‌شده، پایان‌نامه کارشناسی ارشد، دانشگاه صنعتی شاهرود.
- [۸] میرزائی، م. (۱۳۹۴)، مطالعه روشهای تعیین پهنای باند در برآورد کرنل، پایان‌نامه کارشناسی ارشد دانشگاه صنعتی شاهرود.
- [۹] نجاریان، س. (۱۳۹۰). بررسی رفتار برآوردگرهای رگرسیونی ریج در مدل‌های خطی منفرد، پایان‌نامه کارشناسی ارشد، دانشگاه صنعتی شاهرود.
- [۱۰] نوروزی‌راد، م. (۱۳۹۶) برآوردگرهای بهبودیافته در برخی مدل‌های رگرسیونی خطی جریمه‌شده، رساله دکتری، دانشگاه صنعتی شاهرود.

[۱۱] نیرومند، ح. ع. (۱۳۸۷) تحلیل رگرسیون خطی: ابزاری برای تحقیق، نشر ارسلان، مشهد.

- [12] Abramowitz, M. and Stegun, I. A. (1965), **Handbook of Mathematical Functions: with Formulas, Graphs, and Mathematical Tables**, Dover Publications. New York.
- [13] Ahn, H. and Powell, J. (1993), **Semiparametric Estimation of Censored Selection Models with a Nonparametric Selection Mechanism**, Journal of Econometrics. **58**, 3-29.
- [14] Akaike, H. (1973), **Information Theory and an Extension of the Maximum Likelihood Principle**, Second International Symposium on Information Theory, 267-281.
- [15] AlNasser H. (2014), **On Ridge Regression and Least Absolute Shrinkage and Selection Operator**. University of Victoria.
- [16] Arashi, M., Roozbeh, M. and Niroumand, H. (2012), **A Note on Stein-type Shrinkage Estimator in Partial Linear Models**, Statistics. **46**, 673-685.
- [17] Bai J, and Ng S. (2000). Large dimensional factor analysis. **Foundations and Trends in Econometrics**, **3**(2), 89–163.
- [18] Borjesson, P. and Sundberg, C. E. (1979). Simple approximations of the error function $Q(x)$ for communications applications, **IEEE Transactions on Communications**, **27**(3), 639-643.
- [19] Buhlmann, P. and van de Geer, S. (2011). **Statistics for High-Dimensional Data: Methods, Theory and Applications**, Springer Series in Statistics. Berlin.
- [20] Chevillard, S. and Revol, N. (2008), Computation of the error function erf in arbitrary precision with correct rounding. **JD Bruguera et M. Daumas (editeurs): RNC**, **8**, 27-36.
- [21] Cody, W. J. (1990). Performance evaluation of programs for the error and complementary error functions, **ACM Transactions on Mathematical Software**, **16**(1), 29-37.

- [22] Efron, B and Hastie, T and Johnstone, I and Tibshirani, R and others. (2004), Least angle regression, **The Annals of statistics**, **32**(2), 407-499.
- [23] Engle, R.F., Granger, C.W.J., Rice, J. and Weiss, A. (1986), **Semi parametric estimates of the relation between ewather and electricity sales**, Journal of American statistical Association. **81** , 310-320.
- [24] Fan, J. and Li, R. (2001). **Variable Selection Via Non-Concave Penalized Likelihood and its Oracle Properties**, Journal of the American Statistical Association, **96**(456), 1348-1360.
- [25] Fan, J., Lv, J. and Qi, L., (2011). Sparse high-dimensional models in economics, **Annual Review of Economics**, **3**(1), 291–317.
- [26] Fan, J. and Peng, H. (2004), Nonconcave penalized likelihood with a diverging number of parameters, **The Annals of Statistics**, **32**(3), 928-961.
- [27] Hall, P., Kay, J. W. and Titterington, D. M., (1990). On estimation of noise variance in two-dimensional signal processing. **Advances in Applied Probability**, **23**, 476-495.
- [28] Haselimashhadi, H and Vinciotti, V. (2016), A Differentiable Alternative to the Lasso Penalty, **arXiv preprint arXiv:1609.04985**.
- [29] Hoerl, A. E. and Kennard, R. W., (1970). Ridge regression: Biased estimation for nonorthogonal problems, **Technometrics**, **12**(1), 55-67.
- [30] Kariya, T. and H. Kurata (2004). **Generalized Least Squares**. Wiley series in Probability and Statistics. New York.
- [31] Knight, K. and Fu, W. (2000) Asymptotics for Lasso-type estimators. **The Annals of Statistics** , **28**, 1356-1378.
- [32] Leal, H. V., Sheissa, R. C., Nino, U. F., Reyes, A. S., and Orea, J. S. (2012) High accurate simple approximation of normal distribution integral. **Mathematical Problems in Engineering**, 1-23.

-
- [33] Lee, C. C. (1992). On Laplace continued fraction for the normal integral, **Annals of the Institute of Statistical Mathematics**, **44**(1), 107-120.
- [34] Mac Nally, R.(2000), Regression and model-building in conservation biology, biogeography and ecology: the distinction between—and reconciliation of—predictive and explanatory models, **Biodiversity and Conservation**, **9**(5),655-671.
- [35] Mallows, C. L.(1973), Some comments on C_p , **Technometrics**, **15**(4), 661-675.
- [36] Montgomery, D. C.(2017), **Design and Analysis of Experiments**, John wiley, New Jersey.
- [37] Mroz, T. A., (1987). The sensitivity of an empirical model of married women's hours of work to economic and statistical assumptions. **Econometrica**, **55**(4), 765-799.
- [38] Nadaraya, E.A.(1964), On estimating regression, **Theory of Probability and its Applications**. **9** , 141-142.
- [39] Nardi, Y. and Rinaldo, A. (2011). **Autoregressive process modeling via the lasso procedure**, Journal of Multivariate Analysis, **102**(3), 528-549.
- [40] Press, W.H. (1992). **Numerical Recipes in C: The Art of Scientific Computing**. Number v. 4. Cambridge University Press.
- [41] Olver, F.W.J. (2010). National Institute of Standards, and Technology (U.S.). **NIST Handbook of Mathematical Functions**. Cambridge University Press.
- [42] Raheem, S. M. E., Ahmed, S. E. and Doksum, K. A. (2012). Absolute penalty and shrinkage estimation in partially linear models, **Computational Statistics and Data Analysis**, **56**, 874-891.
- [43] Rice, J. (1984), Bandwidth Choice for Nonparametric Regression, **Annals of Mathematical Statistics**, **12** , 1215-1230.
- [44] Roozbe, M., Arashi, M. and Niroumand, H. (2010), **Ridge Regression Methodology in Partial Linear Models with Correlated Errors**, Journal of Statistical Computation and Simulation. **81** , 517-528.

- [45] Schmidt, M., Fung, G., and Rosales, R. (2007). Fast optimization methods for l1 regularization: A comparative study and two new approaches. **In Machine Learning: ECML 2007**, pages 286-297.
- [46] Schwarz, G. (1978), Estimating the dimension of a model, **The Annals of Statistics**, **6**(2), 461-464.
- [47] Stackelberg, H. v. and Heinrich. (1952), **Theory of the Market Economy** Oxford University Press.
- [48] Stamey, T. A., Kabalin, J. N., McNeal, J. E. and Johnstone, I. M. and Freiha, F. and Redwine, E. A. and Yang, N.(1989). **Prostate Specific Antigen in the Diagnosis and Treatment of Adenocarcinoma of the Prostate. II. Radical Prostatectomy Treated Patients**, The Journal of Urology, **141**(5), 1076–1083.
- [49] Tabakan, G. and Akdeniz, F.(2010), **Difference-based Ridge Estimator of Parameters in Partial Linear Model**, Statical Papers. **51**, 357-368.
- [50] Tibshirani R.(1996). **Regression Shrinkage and Selection Via the Lasso**, Journal of the Royal Statistical Society. Series B (Methodological), 267-288.
- [51] Tibshirani, R., Saunders, M., Rosset, S., Zhu, J., and Knight, K. (2005). **Sparsity and Smoothness Via the Fused Lasso**, Journal of the Royal Statistical Society: Series B (Statistical Methodology),**67**(1), 91-108.
- [52] Wand, M.P. and Jones, M.C.(1995), **Kernel Smoothing**, Chapman and Hill. London.
- [53] Wang, H., Li, G., and Jiang, G. (2007). **Robust Regression Shrinkage and consistent variable Selection Through the LAD-LASSO**, Journal of Business and Economic Statistics, **25**(3), 347-355.
- [54] Watson, G.S. (1964), **Smooth regression analysis**, Journal of Sankhya. **26**, 35-72.
- [55] van Wieringen, Wessel N. (2015), Lecture notes on ridge regression, **arXiv preprint arXiv:1509.09169**.

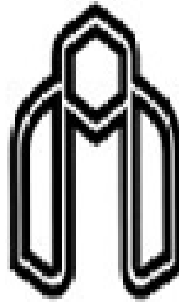
-
- [56] Variyan, H. R. (2014). **Big Data: New tricks for Econometrics**, Journal of Economic Perspectives, **28**(2), 3-28.
- [57] Yatchew, A. (1997), **An Elementary Estimator of the Partial Linear Model**, Journal of Economics Letters. **57**, 135-143.
- [58] Yatchew, A. (1998), **Nonparametric Regression Techniques in Economics**, Journal of Economic Literature. **36** , 669-721.
- [59] Yatchew, A. (1999), **An Elementary Nonparametric Differencing Test of Equality of Regression Functions**, Journal of Economics Letters. **62**, 271-278.
- [60] Yatchew, A. (2000), **Scale Economies in Electricity Distribution**, Journal of Applied Economics. **15**, 187-210.
- [61] Yatchew, A. (2003), **Semiparametric Regression for the Applied Econometricians**, Cambridge University Press, Cambridge.
- [62] Yuan, M. and Lin, Y. (2006). **Model selection and Estimation in Regression With Grouped Variables**, Journal of the Royal Statistical Society: Series B (Statistical Methodology), **68**(1), 49-67.
- [63] Zou, H. (2006), **The adaptive lasso and its oracle properties**, Journal of the American Statistical Association, **101**(476), 1418-1429.
- [64] Zou, H., and Hastie, T. (2005). **Regularization and Variable Selection Via the Elastic Net**, Journal of the Royal Statistical Society: Series B (Statistical Methodology), **67**(2), 301-320.

Abstract

In regression analysis if the number of explanatory variables is large, estimation of parameters with least squares method is not relevant and one must use penalized estimators.

One of the most efficient penalized estimators is the LASSO. This estimator is not differentiable at zero and hence not normally distributed. In this dissertation we study the performance of a differentiable LASSO which has asymptotic normal distribution. With some simulation studies and real data analyses, we evaluate the performance of this estimator along with a difference-based differentiable LASSO.

Keywords: Differentiable LASSO Estimator, Difference Estimator, Penalty Function, High Dimensional Penalized Regression



Shahrood University of Technology

Faculty of Mathematical Sciences

M.Sc. Thesis in Mathematical Statistics

High Dimensional Penalized Regression With Differentiable LASSO

By: Samane Rajabloo

Supervisor

Mohammad Arashi

September 2018