

الحمد لله
الرحمن الرحيم



دانشکده علوم ریاضی

رشته آمار، گرایش استنباط

رساله دکتری

برآورد در مدل‌های نیمه پارامتری برای تحلیل داده‌های طولی

نگارنده: مژگان تعاونی

استاد راهنما

دکتر محمد آرشی

شهریور ۱۳۹۹

شماره: ۱-۲۱-۳۳

تاریخ: ۹۹/۷/۲۹

ویرایش:

باسمه تعالی



مدیریت تحصیلات تکمیلی

فرم شماره ۱۱: صورت جلسه نهایی دفاع از رساله دکتری (Ph.D)

بدینوسیله گواهی می شود خانم مژگان تعاونی دانشجوی دکتری رشته آمار به شماره دانشجویی ۹۴۰۰۲۲۵ ورودی ۱۳۹۴ در تاریخ ۹۹/۶/۵ از رساله نظری / عملی ☐ خود با عنوان: برآورد در مدل های نیمه پارامتری برای تحلیل داده های طولی دفاع و با اخذ نمره ۱۹.۹ به درجه عالی نائل گردید.

الف) درجه عالی: نمره ۱۹-۲۰ ☒ ب) درجه خیلی خوب: نمره ۱۸/۹۹-۱۷ ☐
ج) درجه خوب: نمره ۱۶/۹۹-۱۵ ☐ د) مردود: کمتر از ۱۵ ☐

ردیف	هیئت داوران	نام و نام خانوادگی	مرتبه علمی	امضاء
۱	دکتر محمد آرشی	استاد راهنما	دانشیار	
۲	دکتر حسین باغبینی	استاد مدعو داخلی	استادیار	
۳	دکتر محمدرضا ربیعی	استاد مدعو داخلی	استادیار	
۴	دکتر مهدی روزبه	استاد مدعو خارجی	دانشیار	
۵	دکتر سید حیدر جعفری	سرپرست (نماینده) تحصیلات تکمیلی دانشکده	استادیار	

مدیر محترم تحصیلات تکمیلی دانشگاه:

ضمن تأیید مراتب فوق مقرر فرمائید اقدامات لازم در خصوص انجام مراحل دانش آموختگی خانم مژگان تعاونی بعمل آید.

نام و نام خانوادگی رئیس دانشکده:

تاریخ و امضاء و مهر دانشکده:



تقدیم به

بزرگوار پدرم،

فداکار مادرم،

مهربان همسرم،

نازنین پسرم،

آنان که دریای بی‌کران فداکاری و عشق‌اند، وجودم برایشان همه رنج بود و وجودشان
برایم همه مهر. هرچه بکوشم، قطره‌ای از دریای بی‌کران مهربانیشان را سپاس نتوانم گفت.
بوسه بر دستان پرمهرتان. اله‌ها به من کمک کن تا بتوانم ادای دین کنم و به خواسته آنان جامه
عمل بپوشانم.

سپاس بی کران پروردگار که هستی مان بخشید؛
و به طریق علم و دانش رهنمونان شد؛
و به همنشینی رهروان علم و دانش مفتخرمان نمود؛
و خوشه چینی از علم و معرفت را روزمان ساخت.

نمی توانم معنایی بالاتر از تقدیر و تشکر بر زبانم جاری سازم و سپاس خود را در وصف استادان خویش آشکار نمایم، که هر چه گویم و سراپم، کم گفته ام. با این وجود بر خود واجب می دانم از استاد فرزانه جناب آقای دکتر محمد آرشی که به عنوان استاد راهنما در مراحل مختلف این رساله همواره با سعه صدر و گشاده رویی در کنار من بودند و در طول مدت تحصیل از راهنمایی های اخلاقی و حمایت های بی دریغ علمی ایشان بهره جسته ام، تشکر و قدردانی نمایم و برای ایشان طول عمر توام با سربلندی را آرزومندم. از اساتید گرامی جناب آقایان دکتر حسین باغیشتی، دکتر محمدرضا ربیعی و دکتر مهدی روزبه که زحمت دآوری رساله اینجانب را قبول فرمودند، قدردانی می نمایم. خالصانه از تمامی معلمان و اساتیدی که در مقاطع مختلف تحصیلی به من علم آموخته و مرا از سرچشمه دانایی سیراب کرده اند، نهایت سپاس را دارم. پروردگارا حسن عاقبت، سلامت و سعادت را برای آنان مقدر نما.

مژگان تعاونی
شهریور ۱۳۹۹

تعهد نامه

اینجانب مژگان تعاونی دانشجوی دکتری رشته آمار دانشکده علوم ریاضی دانشگاه صنعتی شاهرود، نویسنده پایان نامه با عنوان برآورد در مدل های نیمه پارامتری برای تحلیل داده های طولی، تحت راهنمایی دکتر محمد آرشی متعهد می شوم:

- تحقیقات در این پایان نامه توسط اینجانب انجام شده است و از صحت و اصالت برخوردار است.
- در استفاده از نتایج پژوهش های دیگر پژوهشگران، به مرجع مورد استفاده استناد شده است.
- مطالب این پایان نامه، تا کنون توسط خود، یا فرد دیگری برای دریافت هیچ نوع مدرک یا امتیازی در هیچ جا 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 به دست آوردن نتایج اصلی پایان نامه تاثیرگذار بوده اند، در مقالات مستخرج از پایان نامه رعایت می گردد.
- در تمام مراحل انجام این پایان نامه، در مواردی که از موجود زنده (یا بافت های آنها) استفاده شده است، ضوابط و اصول اخلاقی رعایت شده است.
- در تمام مراحل انجام این پایان نامه، در مواردی که به حوزه اطلاعات شخصی افراد دسترسی یافته (یا استفاده شده است)، اصل رازداری و اصول اخلاق انسانی رعایت شده است.

مژگان تعاونی

شهریور ۱۳۹۹

مالکیت نتایج و حق نشر

- تمام حقوق معنوی این اثر و محصولات آن (مقالات مستخرج، کتاب، برنامه های رایانه ای، نرم افزارها و تجهیزات ساخته شده) متعلق به دانشگاه صنعتی شاهرود می باشد. این مطلب باید به نحو مقتضی، در تولیدات علمی مربوطه ذکر شود.
- استفاده از اطلاعات و نتایج موجود در این پایان نامه بدون ذکر منبع مجاز نمی باشد.

چکیده

این رساله به بررسی روش‌های برآورد در مدل رگرسیون نیمه‌پارامتری با اثرات آمیخته به منظور تحلیل داده‌های طولی می‌پردازد. مدل موردنظر ترکیبی از مدل اثرات آمیخته و مدل نیمه‌پارامتری است. ابتدا چند مورد از روش‌های موجود برای برآورد تابع ناپارامتری مدل شامل برآورد هسته‌ای، بازگشتی و اسپلاین معرفی شده و با درنظر گرفتن مساله هم‌خطی بین متغیرهای تبیینی، برآوردگر جدیدی به نام برآوردگر هسته‌ای ریج معرفی می‌شود. در ادامه مساله برآورد و انتخاب متغیر همزمان برای داده‌های گاوسی و غیرگاوسی با بعد بالا درنظر گرفته شده است. مولفه ناپارامتری موجود در مدل با اسپلاین‌های رگرسیونی تقریب زده شده و از طریق بهینه‌سازی تابع هدف مبتنی بر تابع تاوان برآورد و انتخاب متغیر به طور همزمان انجام می‌گیرد. رفتار حدی برآوردگرهای حاصل در چهارچوب داده‌های بعد بالا که در آن تعداد پارامترها متناسب با افزایش حجم نمونه افزایش می‌یابد، مورد مطالعه قرار می‌گیرد. علاوه بر این با درنظر گرفتن پاسخ‌های چندمتغیره یک روش نیرومند بر اساس توزیع تی چندمتغیره، به عنوان جایگزینی برای توزیع نرمال، برای برآوردیابی ارائه می‌شود که نسبت به مقادیر پرت یا مشاهدات غیرمعمول انعطاف‌پذیرتر است. در هر چهارچوب مطالعاتی، به منظور پیاده‌سازی روش‌های برآورد یک الگوریتم تکراری مناسب ارائه گردیده و کارایی روش‌های برآورد پیشنهادی با استفاده از مطالعات شبیه‌سازی و تحلیل داده‌های واقعی بررسی شده است. نتایج بدست آمده نشان می‌دهند که روش‌های پیشنهادی در مقایسه با سایر روش‌های موجود عملکرد مناسبی دارند.

کلمات کلیدی: اسپلاین هموارساز، انتخاب متغیر، برآوردگر ریج، بعد بالا، تاوانیده، توزیع تی چندمتغیره، توزیع مجانبی، داده‌های پرت، داده‌های طولی، دم‌کلفت، بازگشتی، هسته‌ای، معادلات برآورد تعمیم‌یافته، مدل نیمه‌پارامتری با اثرات آمیخته.

فهرست مقالات مستخرج از رساله

۱. تعاونی، م. و آرشی، م. (۱۳۹۹). انتخاب متغیر در مدل‌های خطی-جزئی با اثرات آمیخته برای داده‌های طولی با بعد بالا، مجله علوم آماری، جلد ۱۴، شماره ۲، صفحات ۳۶۷-۳۸۸.
2. Taavoni, M. and Arashi, M. (2018), Semiparametric Mixed Effect Modeling for Longitudinal Data Analysis, The Secound Seminar on Nonparametric Statistics and Its Applications, Allameh Tabataba'i University, Tehran, Iran.
3. Taavoni, M. and Arashi, M. (2018), Semiparametric Ridge Regression for Longitudinal Data, 14th Iranian Statistics Conference, Shahrood University of Technology, Shahrood, Iran.
4. Taavoni, M. and Arashi, M. (2019), Kernel Estimation in Semiparametric Mixed Effect Longitudinal Modeling, *Stat. Papers*, DOI: 10.1007/s00362-019-01125-8.
5. Taavoni, M. and Arashi, M. (2019), Regularization in Generalized Semiparametric Mixed Effects Model for Longitudinal Data, arXiv: 1909.08498v1 [stat.ME].
6. Taavoni, M., Arashi, M., Wang, W.L. and Lin, T.I (2020), Multivariate t Semiparametric Mixed-effects Model for Longitudinal Data with Multiple Characteristics, *J. Statist. Comput. Simul.*, (Revised).
7. Taavoni, M., Arashi, M., Wang, W.L. and Lin, T.I (2020), Estimation in Multivariate t Linear Mixed Models for Longitudinal Data with Multiple Outputs: Application to PBCseq Data Analysis, *Biom. J.*, (Revised).

پیشگفتار

مطالعات طولی یکی از رایج‌ترین و کاربردی‌ترین نوع مطالعات در بسیاری از زمینه‌ها است. داده‌های طولی به مجموعه وسیعی از داده‌های گسسته و پیوسته اطلاق می‌شود که در آن ویژگی‌های چندین آزمودنی به عنوان نمونه‌ای از اعضای جمعیت، در طول زمان اندازه‌گیری می‌شوند. هدف اولیه مطالعات طولی، بررسی تغییرات متغیر پاسخ در طول زمان و نیز استخراج عامل‌های تاثیرگذار بر این تغییرات است. با توجه به انجام اندازه‌گیری‌های مکرر، پاسخ‌های هر آزمودنی دارای همبستگی هستند. جهت رسیدن به استنباط علمی و معتبر، در نظر گرفتن این همبستگی امری ضروری است. به دلیل وجود اندازه‌های تکراری امکان ارزشیابی تغییرات درون هر آزمودنی نیز فراهم می‌شود که این امکان در مطالعات مقطعی که آزمودنی تنها در یک زمان مشخص اندازه‌گیری می‌شود وجود ندارد. مدل رگرسیونی یکی از پرکاربردترین ابزارهای آماری است که به دلیل تنوع زیاد مطالعات، روش‌های گوناگونی برای انجام آن پدید آمده است. از یک منظر رگرسیون را می‌توان به سه بخش پارامتری، ناپارامتری و نیمه پارامتری تقسیم کرد:

- رگرسیون پارامتری بیشتر وقتی مورد استفاده قرار می‌گیرد که اطلاعات تحقیق به قدر کافی زیاد باشد تا آنجایی که بتوان ارتباط بین متغیر پاسخ و تبیینی را به وسیله یک مدل خطی یا قابل تبدیل به خطی بیان کرد. در تحلیل پارامتری داده‌های طولی، رهیافت‌های مختلفی برای در نظر گرفتن همبستگی وجود دارند که شامل مدل‌های حاشیه‌ای، اثرات آمیخته و انتقالی هستند.
- رگرسیون ناپارامتری در مواقعی که نتوان نوع ارتباط بین متغیرها را با یک شکل تابعی معلوم و مشخص بیان کرد، مورد استفاده قرار می‌گیرد، که در آن می‌توان از روش‌های اسپلاین و هسته‌ای استفاده کرد.
- مدل‌های نیمه پارامتری، ترکیبی از مدل‌های پارامتری و ناپارامتری است، به این معنی که بعضی از متغیرها دارای روند معلوم تغییرات نسبت به متغیر پاسخ بوده و برای بعضی از متغیرها نمی‌توان شکل تابعی خاصی را در نظر گرفت.

در چند دهه اخیر، آماردانان در راستای تحلیل هر چه بهتر داده‌های طولی با خلق ترکیبی از مدل‌های ذکر شده، رده دیگری از مدل‌ها را ایجاد کردند. در این میان، مدل اثرات آمیخته از مجموعه مدل‌های پارامتری و مدل‌های نیمه پارامتری مبنای بسیاری از این دسته مدل‌های ترکیبی قرار گرفتند. مدل نیمه پارامتری با اثرات آمیخته که اساس کار این پژوهش است، به عنوان مدلی بسیار منعطف و تنومند که ویژگی‌های خوب اکثر مدل‌ها را دارد مورد توجه بسیار قرار گرفت.

از طرفی، امروزه با افزایش توان محاسباتی کامپیوترها و روش‌های نوین جمع‌آوری اطلاعات، داده‌های با بعد بالا در زمینه‌های مختلف علمی و تحقیقاتی با هزینه نسبتاً کم در اختیار محققین قرار می‌گیرد. در این نوع داده‌ها فقط تعداد اندکی از متغیرهای تبیینی واقعاً با متغیر پاسخ مرتبط هستند و سایر متغیرها تاثیری بر پاسخ ندارند. هنگامی که تعداد متغیرهای تبیینی در مدل زیاد است، تفسیر آن مدل بسیار مشکل بوده و همچنین حضور تعداد زیاد متغیرها باعث ایجاد هم‌خطی بین

متغیرها می‌شود؛ لذا شناسایی متغیرهای موثر بر پاسخ اهمیت بسزایی دارد. یک روش متداول برای مقابله با مشکل هم‌خطی استفاده از برآوردهای اریب مانند ریج است. همچنین فرآیندی که علاوه بر انتخاب متغیر و بهترین معادله رگرسیونی، برآورد ضرایب را نیز انجام دهد، روش مبتنی بر تابع تاوان است که رویکرد نوینی تحت عنوان رگرسیون تاوانیده را گسترش داد. در بسیاری از تحقیقات علمی، با داده‌های رسته‌ای یا داده‌های غیرنرمال، مانند پاسخ‌های دودویی و شمارشی مواجه هستیم که برای مدل‌بندی آن‌ها، مدل‌های خطی تعمیم‌یافته پیشنهاد می‌شوند. مطالعات طولی که متغیر پاسخ آن‌ها غیرنرمال باشد نیز در علوم مختلف جایگاه ویژه‌ای یافته‌اند. مسئله مهم کاربردی دیگر که در تحلیل داده‌های طولی مورد توجه قرار می‌گیرد، تحلیل همزمان یا چند متغیره این داده‌ها است. مدل‌های چندمتغیره آمیخته از پرکاربردترین روش‌های آماری جهت انجام استنباط اندازه‌های تکراری چندمتغیره است. پذیره معمول در برازش این مدل‌ها نرمال بودن توزیع خطاها و مولفه‌های اثرات تصادفی است. در برخورد با برخی داده‌های واقعی این پذیره ممکن است منجر به نتایج نادرست شود به‌ویژه در صورت وجود داده‌های دورافتاده، استنباط را آسیب‌پذیر ساخته و واریانس برآوردها را به شدت تحت تاثیر قرار می‌دهد. از این‌رو در سال‌های اخیر روش‌های پیشرفته آماری با جایگزین کردن توزیع‌های متقارن دم‌سنگین به جای پذیره نرمال بودن در این نوع از مدل‌بندی‌ها منجر به نتایج معتبرتری در مطالعات طولی شده است. علیرغم سازگاری این مدل‌ها با داده‌های چندمتغیره و تنومند بودن در مقابل مقادیر پرت، این مدل‌ها تنها محدود به رابطه خطی بین متغیرهای پاسخ و تبیینی است.

طبق بررسی‌های صورت گرفته، تحقیقات در زمینه انتخاب متغیر داده‌های طولی با بعد بالا در مدل رگرسیون نیمه‌پارامتری با اثرات آمیخته به ویژه در چارچوب خانواده توزیع‌های نمایی تا حد زیادی ناشناخته مانده است. همچنین استفاده از مدل‌های نیمه‌پارامتری برای داده‌های طولی چندمتغیره نسبتاً نادر بوده است. چه بسا مسائلی مانند مسئله هم‌خطی و انتخاب متغیر نیز در این نوع مدل‌بندی‌ها کمتر مورد بحث قرار گرفته‌اند. لذا تحقیقات انجام شده در راستای تمرکز این رساله، که بررسی روش‌های برآورد در مدل رگرسیون نیمه‌پارامتری با اثرات آمیخته است، را در ۴ فصل گنجانده‌ایم که شرح مختصری از محتویات آن‌ها در زیر آمده است.

فصل اول: به بیان مقدمات لازم جهت آشنایی با داده‌های طولی، معرفی ویژگی‌های هر یک از مدل‌های رگرسیونی بیان شده و روش‌های برآورد در آن‌ها اختصاص یافته است. سپس با مبنا قراردادن مدل نیمه‌پارامتری با اثرات آمیخته، برخی از روش‌های برآورد تابع ناپارامتری مدل شامل برآوردهای هسته‌ای، بازگشتی و اسپلاین معرفی می‌شوند.

فصل دوم: ابتدا تاثیر وجود هم‌خطی بر عملکرد برآوردها بررسی شده و برآوردها ریج، به عنوان ابزاری برای مقابله با مشکل هم‌خطی معرفی می‌شود. سپس با معرفی رایج‌ترین روش‌های انتخاب متغیر، برآورد و انتخاب متغیر در مدل نیمه پارامتری با اثرات آمیخته را مبتنی بر رهیافت توابع تاوان، در چارچوب داده‌های بعد بالا که در آن تعداد پارامترها متناسب با افزایش حجم نمونه افزایش می‌یابد، و نحوه انتخاب پارامترهای تابع تاوان مورد بررسی

قرار می‌گیرد.

فصل سوم: به معرفی اجمالی مدل‌های آمیخته تعمیم‌یافته نیمه‌پارامتری پرداخته است. سپس برآورد و انتخاب متغیر در مدل آمیخته تعمیم‌یافته نیمه‌پارامتری را مبتنی بر رهیافت توابع توان برای داده‌های طولی غیرنرمال با بعد بالا و بعد پایین در نظر می‌گیرد. برآوردهای معرفی‌شده صورت بسته ندارند؛ لذا یک الگوریتم تکراری به نام الگوریتم نیوتن رافسون مونت کارلویی توانیده پیشنهاد می‌شود.

فصل چهارم: با در نظر گرفتن پاسخ‌های چندمتغیره، مسئله انتخاب متغیر در مدل‌های چندمتغیره آمیخته با توزیع تی چندمتغیره به عنوان جایگزینی برای توزیع نرمال، می‌پردازد که نسبت به مقادیر پرت یا مشاهدات غیرمعمول انعطاف‌پذیرتر است. سپس نسخه نیمه‌پارامتری مدل مذکور تحت عنوان مدل آمیخته چندمتغیره نیمه‌پارامتری معرفی می‌گردد.

لازم به ذکر است که در هر چهارچوب مطالعاتی، خواص جانبی برآوردهای حاصل بررسی شده و به منظور پیاده‌سازی روش‌های برآورد، یک الگوریتم تکراری مناسب ارائه شده است. کارایی روش‌های برآورد پیشنهادی با استفاده از مطالعات شبیه‌سازی و تحلیل داده‌های واقعی بررسی می‌شود. نتایج بدست آمده نشان می‌دهند که روش‌های پیشنهادی در مقایسه با سایر روش‌های موجود عملکرد مناسبی دارند. همچنین برای برآورد مولفه‌های ناپارامتری مدل، رگرسیون اسپلاین را اتخاذ کرده و در مسائل برآورد و انتخاب متغیر همزمان تابع توان اسکد در نظر گرفته می‌شود. این مجموعه شامل ۲ پیوست است.

پیوست آ: شامل تعاریف و ملزومات استفاده شده در رساله است.

پیوست ب: چند لم مورد نیاز و اثبات قضایای بیان‌شده در فصل‌های مختلف رساله را شامل می‌شود.

فهرست مطالب

ث	فهرست تصاویر
ذ	فهرست جداول
۱	۱ گذری بر مفاهیم داده‌های طولی
۱	۱.۱ تعریف و نمادگذاری
۴	۲.۱ ویژگی‌های داده‌های طولی
۵	۱.۲.۱ همبستگی درون گروهی
۹	۳.۱ مروری بر مدل‌بندی داده‌های طولی
۱۱	۱.۳.۱ مدل‌های پارامتری
۲۱	۲.۳.۱ مدل‌های ناپارامتری
۲۸	۳.۳.۱ مدل‌های نیمه‌پارامتری
۳۲	۴.۳.۱ مدل نیمه‌پارامتری با اثرات آمیخته
۳۳	۴.۱ مثال‌های انگیزشی مورد بحث در این رساله
۳۳	۱.۴.۱ داده‌های $CD4$
۳۴	۲.۴.۱ داده‌های چرخه سلولی مخمر ORF
۳۵	۳.۴.۱ داده‌های بیماری صفراوی
۳۷	۵.۱ روش‌های برآورد در مدل نیمه‌پارامتری با اثرات آمیخته
۳۷	۱.۵.۱ روش برآوردگر هسته‌ای
۳۹	۲.۵.۱ روش برآوردگر بازگشتی
۴۰	۳.۵.۱ روش برآوردگر اسپلاین
۴۱	۴.۵.۱ خواص مجانبی
۴۲	۵.۵.۱ الگوریتم تکراری
۴۳	۶.۵.۱ مطالعات عددی

۴۹	انتخاب متغیر برای داده‌های طولی نرمال	۲
۴۹	انتخاب متغیر در مدل‌های رگرسیونی	۱.۲
۵۲	مسئله هم‌خطی و برآوردگر ریج	۱.۱.۲
۵۵	انتخاب متغیر با معیارهای کلاسیک	۲.۱.۲
۵۷	انتخاب متغیر مبتنی بر تابع تاوان	۳.۱.۲
۷۱	برآوردگر ریج برای داده‌های طولی	۲.۲
۷۲	مطالعات عددی	۱.۲.۲
	انتخاب متغیر در مدل نیمه‌پارامتری با اثرات آمیخته برای داده‌های طولی نرمال	۳.۲
۷۳	با بعد بالا	
۷۶	الگوریتم EM برای بیشینه‌درست‌نمایی توانیده	۱.۳.۲
۷۹	خواص مجانبی	۲.۳.۲
۸۰	مطالعات عددی	۳.۳.۲
۸۵	انتخاب متغیر برای داده‌های طولی در خانواده توزیع‌های نمایی	۳
۸۷	انتخاب متغیر برای داده‌ها با بعد بالا	۱.۳
۸۷	مدل آمیخته نیمه‌پارامتری تعمیم یافته	۱.۱.۳
۸۸	انتخاب متغیر در $GSMM$	۲.۱.۳
۹۳	خواص مجانبی	۳.۱.۳
۹۶	مطالعات عددی	۴.۱.۳
۱۰۲	انتخاب متغیر برای داده‌های با بعد پایین	۲.۳
۱۰۳	انتخاب متغیر برای داده‌های طولی چندمتغیره	۴
۱۰۴	انتخاب متغیر در مدل آمیخته با توزیع t چندمتغیره	۱.۴
۱۰۴	معرفی $MtLMM$	۱.۱.۴
۱۰۶	الگوریتم ECM برای PML	۲.۱.۴
۱۰۸	خواص مجانبی	۳.۱.۴
۱۱۰	مطالعات عددی	۴.۱.۴
۱۱۳	مدل نیمه پارامتری با اثرات آمیخته با توزیع t چندمتغیره	۲.۴
۱۱۳	مدل $MtSMM$ و تقریب اسپلین	۱.۲.۴
۱۱۴	برآورد ML و الگوریتم ECM	۲.۲.۴
۱۱۶	خواص مجانبی	۳.۲.۴
۱۱۷	مطالعات عددی	۴.۲.۴
۱۱۸	پیشنهادهای برای آینده تحقیق	۳.۴
۱۱۸	مدل اثرات آمیخته تعمیم‌یافته نیمه‌پارامتری برای داده‌های چندجمله‌ای	۱.۳.۴

۲۰۳.۴ استفاده از توابع تاوان تجمعی ۱۱۹

۳۰۳.۴ مدل ضرایب متغیر نیمه پارامتری با اثرات آمیخته ۱۱۹

۱۲۷

مراجع

آ تعاریف و ملزومات ۱۳۹

۱.آ معیارهای نیکویی برازش ۱۳۹

۲.آ تعاریف و نامساوی ها ۱۴۱

ب اثبات قضایا ۱۴۷

ب.۱ اثبات قضایای فصل ۱ ۱۴۷

ب.۱.۱ اثبات قضیه ۱.۵.۱ ۱۴۷

ب.۱.۲ اثبات قضیه ۳.۵.۱ ۱۴۸

ب.۲ اثبات قضایای فصل ۲ ۱۴۹

ب.۲.۱ بردار رتبه و ماتریس اطلاع فیشر ۱۴۹

ب.۲.۲ اثبات قضیه ۱.۳.۲ ۱۵۰

ب.۲.۳ اثبات قضیه ۲.۳.۲ ۱۵۲

ب.۳ اثبات قضایای فصل ۳ ۱۵۳

ب.۳.۱ لم ها ۱۵۳

ب.۳.۲ اثبات لم ها ۱۵۴

ب.۳.۳ اثبات قضیه ۱.۱.۳ ۱۶۰

ب.۴.۳ اثبات قضیه ۲.۱.۳ ۱۶۴

ب.۵.۳ اثبات قضیه ۳.۱.۳ ۱۶۶

ب.۴ اثبات قضایای فصل ۴ ۱۷۰

ب.۴.۱ بردار رتبه و ماتریس اطلاع فیشر ۱۷۰

ب.۴.۲ اثبات قضیه ۲.۱.۴ ۱۷۱

ب.۳.۴ اثبات قضیه ۱.۲.۴ ۱۷۲

فهرست تصاویر

۱.۱	نمودار قدرت خواندن کودکان	۳
۲.۱	نمودار ساختگی جهت نمایش تقسیم‌بندی طرح‌های مطالعاتی	۵
۳.۱	نمودار پراکنش تعداد سلول‌های $CD4$ در طول زمان (منحنی‌های خاکستری) به همراه	
۳۴	میانگین هموار شده آن‌ها (منحنی آبی).	
۴.۱	نمودار پراکنش داده‌های $PBCseq$ در طول زمان. سطوح $lbili$ و $albumin$ برای ۱۰۰	
۳۷	بیمار (منحنی‌های خاکستری) به همراه میانگین هموار شده آن‌ها (منحنی آبی).	
۱.۲	منحنی‌های تراز و نواحی محدودیت رگرسیون لاسو (نمودار چپ) و رگرسیون ريج (نمودار راست). ناحیه‌های محدودیت لاسو و ريج به‌ترتیب به‌صورت $s \leq \beta_1 + \beta_2 $ و β_2^2 هستند. بیضی‌ها، منحنی‌های تراز تابع خطای برآورد کمترین توان های دوم هستند.	۵۹
۲.۲	مقایسه نمودارهای توابع تاوان لاسو، اسکد و MCP .	۶۳
۳.۲	مقایسه نمودارهای برآوردهای کمترین توان‌های دوم تاوانیده (PLS) در مقابل برآورد کمترین توان‌های دوم معمولی (OLS) وقتی که ماتریس طرح متعامد است. محور افقی	
۶۴	OLS و محور عمودی PLS است.	
۴.۲	نمودار مشتقات توابع تاوان لاسو، اسکد و MCP .	۶۵
۵.۲	تقریب درجه دو و خطی برای توابع جریمه $L_{0.5}$ و $SCAD$. منحنی آبی: تقریب درجه دو؛ منحنی قرمز: تقریب خطی؛ منحنی مشکی: تابع جریمه.	۷۰
۱.۳	نتایج شبیه‌سازی: نتایج مربوط به $f(t)$ برآورد شده برای حالت‌های $p_n < n$ و $p_n > n$ با حجم نمونه $n = 50$. تصویر بالا-چپ: منحنی $f(t)$ برآورد شده، تصویر بالا-راست: منحنی اریبی، تصویر پایین-چپ: منحنی انحراف استاندارد، تصویر پایین-راست: احتمال پوشش.	۹۸
۲.۳	نتایج شبیه‌سازی: نتایج مربوط به $f(t)$ برآورد شده برای حالت‌های $p_n < n$ و $p_n > n$ در مطالعه شبیه‌سازی با حجم نمونه $n = 100$. تصویر بالا-چپ: منحنی $f(t)$ برآورد شده، تصویر بالا-راست: منحنی اریبی، تصویر پایین-چپ: منحنی انحراف استاندارد، تصویر پایین-راست: احتمال پوشش.	۹۸

۳.۳	نتایج تحلیل داده‌های $CD4$: نتایج مربوط به $f(t)$ برآوردشده تحت مدل $P-GSMM$.
	تصویر بالا-چپ: نمودار پراکنش متغیر پاسخ، تصویر بالا-راست: $\hat{f}(t)$ (آبی) و فواصل
	اطمینان (مشکی)، تصویر پایین-چپ: منحنی‌های انحراف استاندارد مجانبی (آبی)
	و نمونه‌ای (قرمز)، تصویر پایین-راست: منحنی احتمالات پوشش مجانبی (آبی) و
۱۰۱	نمونه‌ای (قرمز)

فهرست جداول

۱۰.۱	مقدار ρ به ازای مقادیر مختلف m و ρ	۹
۲.۱	نتایج شبیه‌سازی: مقادیر MSE ، اریبی (SD) در مدل‌های متفاوت به ازای ρ ها و n های مختلف.	۴۶
۳.۱	نتایج شبیه‌سازی: مقادیر MSE ، اریبی (SD) برای روش‌های متفاوت برآورد به ازای ρ ها و n های مختلف در مدل اثرات آمیخته نیمه‌پارامتری.	۴۷
۴.۱	نتایج تحلیل داده‌های $CD4$: برآورد پارامترهای مدل رگرسیونی (SD) تحت برازش مدل‌های مختلف.	۴۸
۱۰.۲	مقادیر تقریبی $\frac{p}{n}$ برای داده‌های با ابعاد بالا و بسیار بالا به ازای مقادیر مختلف a ، b و n	۵۲
۲.۲	نتایج شبیه‌سازی: مقادیر MSE و اریبی (SD) برآوردگرهای $\hat{\beta}_{RK}$ و $\hat{\beta}_K$ به ازای γ های مختلف.	۷۴
۳.۲	نتایج تحلیل داده‌های $CD4$: برآورد پارامترهای رگرسیونی (SD) تحت برآوردگرهای مختلف.	۷۵
۴.۲	نتایج شبیه‌سازی: مقایسه نتایج انتخاب متغیر تحت روش‌های SMM ، $P-SMM$ و $P-LMM$ برای حالت‌های $p_n > n$ و $p_n < n$	۸۲
۵.۲	نتایج شبیه‌سازی: کارایی روش $P-SMM$ برای حالت‌های $p_n > n$ و $p_n < n$	۸۲
۶.۲	نتایج تحلیل داده‌های $CD4$: برآورد پارامترهای رگرسیونی (SD) تحت برازش سه مدل SMM ، $P-SMM$ و $P-LMM$	۸۴
۱۰.۳	نتایج شبیه‌سازی: مقایسه نتایج انتخاب متغیر روش‌های $GSMM$ ، $P-GSMM$ و $P-GLMM$ برای حالت‌های $p_n > n$ و $p_n < n$	۹۷
۲.۳	نتایج شبیه‌سازی: کارایی روش $P-GSMM$ برای حالت‌های $p_n > n$ و $p_n < n$	۹۷
۳.۳	نتایج تحلیل داده‌های $CD4$: برآورد پارامترهای رگرسیونی (SD) تحت برازش سه مدل $GSMM$ ، $P-GSMM$ و $P-GLMM$	۱۰۰
۴.۳	نتایج تحلیل داده‌های چرخه سلولی مخمر ORF : برآورد پارامترهای رگرسیونی (SD) تحت برازش سه مدل $GSMM$ ، $P-GSMM$ و $P-GLMM$	۱۰۱

۱۰۴	نتایج شبیه‌سازی: مقایسه نتایج انتخاب متغیر روش‌های $MtLMM$ تاوانیده، $MtLMM$
۱۱۲	$MNMM$ تاوانیده و $MNMM$ به ازای n ها و ν های مختلف.
۲۰۴	نتایج شبیه‌سازی: کارایی روش‌های $MtLMM$ تاوانیده و $MNMM$ تاوانیده به ازای
۱۲۱	n ها و ν های مختلف.
۳۰۴	نتایج تحلیل داده‌های $PBCseq$: برآورد پارامترهای رگرسیونی (SD) تحت برازش سه
۱۲۲	مدل $MtLMM$ تاوانیده، $MtLMM$ ، $MNMM$ تاوانیده و $MNMM$
۴۰۴	نتایج شبیه‌سازی: کارایی روش‌های $MtSMM$ ، $MtLMM$ و $MNSMM$ به ازای
۱۲۳	$n = 25$ و ν های مختلف.
۵۰۴	نتایج شبیه‌سازی: کارایی روش‌های $MtSMM$ ، $MtLMM$ و $MNSMM$ به ازای
۱۲۴	$n = 100$ و ν های مختلف.
۶۰۴	نتایج تحلیل داده‌های $PBCseq$: نتایج معیارهای انتخاب مدل تحت برازش مدل‌های
۱۲۵	$MtSMM$ و $MNSMM$ و ساختارهای همبستگی مختلف.
۷۰۴	نتایج تحلیل داده‌های $PBCseq$: برآورد پارامترهای رگرسیونی (SD) تحت برازش
۱۲۶	مدل‌های $MtSMM$ و $MNSMM$

فصل ۱

گذری بر مفاهیم داده‌های طولی

از آنجایی که آشنایی با یک موضوع جدید، مستلزم شناخت مقدمات مربوط به آن است، این فصل به معرفی مقدمات لازم جهت آشنایی با داده‌های طولی^۱ اختصاص یافته است. بدین منظور، ابتدا تعریفی از داده‌های طولی، نمادگذاری متداول، اهداف و لزوم ارائه مبحث تحلیل داده‌های طولی بیان می‌گردد. سپس به معرفی و بررسی مدل‌های رگرسیونی که در تحلیل داده‌های طولی به کار می‌روند پرداخته شده است. مطالب این فصل عمدتاً برگرفته از کتاب تحلیل چندمتغیره‌ی گسسته در مطالعات مقطعی و طولی تألیف گنجعلی و رضایی (۱۳۸۹) و پایان نامه کارشناسی ارشد تعاونی (۱۳۹۲) است. علاقه‌مندان به منظور آشنایی بیشتر با داده‌های طولی و نحوه مدل‌بندی آن‌ها می‌توانند به کتاب‌های دیگل و همکاران (۲۰۰۲)، هدیگر و گیبونز (۲۰۰۶)، وو و ژانگ (۲۰۰۶)، فیتز‌موریس و همکاران (۲۰۰۹)، وربک و مولنبرگ (۲۰۰۹) و فیتز‌موریس و همکاران (۲۰۱۱) مراجعه نمایند.

۱.۱ تعریف و نمادگذاری

مسئله مهم در یک تحلیل آماری، نوع مطالعه و روش گردآوری داده‌ها است. در بسیاری از تحقیقات علمی به‌ویژه در علوم پزشکی به منظور بررسی روند تحقیقات و تاثیر روش‌های مختلف، متغیرهای

¹Longitudinal data

موجود در مطالعه، بیش‌تر از یک بار روی یک آزمودنی^۲ (انسان، حیوان و نمونه‌های آزمایشگاهی) اندازه‌گیری می‌شود؛ مشاهدات حاصل را اندازه‌گیری مکرر^۳ می‌نامند. فاکتور اندازه‌گیری براساس نوع مطالعه انتخاب می‌شود که می‌تواند زمان، مکان یا هر شرط آزمایشی دیگر باشد که اغلب به این فاکتور، فاکتور درون-گروهی^۴ یا درون-فردی^۵ گفته می‌شود. داده‌های مکرری که موقعیت‌های تکرار مشاهدات، نقاط زمانی هستند داده‌های طولی و مطالعاتی را که بر روی این داده‌ها انجام می‌گیرد، مطالعات طولی می‌نامند. می‌توان گفت مطالعات طولی یک شاخه از مطالعات با اندازه‌گیری مکرر هستند که در آن فاکتور اندازه‌گیری زمان است. به عنوان مثال در کودکان مبتلا به بیماری آسم، مشاهدات در ۴ نقطه زمانی یعنی در سنین ۷، ۸، ۹ و ۱۰ سالگی اندازه‌گیری می‌شوند. جهت آشنایی بیش‌تر با ساختار داده‌های طولی در ادامه به معرفی نمادگذاری متداول این داده‌ها می‌پردازیم. با توجه به تعریف داده‌های طولی، نمونه‌ای تصادفی و مستقل به حجم n در نظر بگیرید به‌طوری‌که نمونه i ام دارای n_i مشاهده مکرر در نقطه زمانی است. متغیر پاسخ طولی^۶ به‌صورت

$$y_i = \begin{pmatrix} y_{i1} \\ \vdots \\ y_{in_i} \end{pmatrix}_{n_i \times 1}$$

تعریف می‌شود. بردار متغیرهای تبیینی^۷ وابسته به متغیر پاسخ y_{ij} به‌صورت

$$x_{ij} = \begin{pmatrix} x_{1,ij} \\ \vdots \\ x_{p,ij} \end{pmatrix}_{p \times 1}$$

است که در آن $x_{k,ij}$ مقدار k امین متغیر تبیینی مربوط به آزمودنی i ام در زمان j ام است، به‌طوری‌که $i = 1, 2, \dots, n$ ، $j = 1, 2, \dots, n_i$ و $k = 1, 2, \dots, p$ بوده و n_i و p به‌ترتیب معرف تعداد کل آزمودنی‌ها، تعداد دوره‌های زمانی برای آزمودنی i ام و تعداد متغیرهای تبیینی هستند. در ماتریس طرح مربوط به آزمودنی i ام که به‌صورت

$$X_i = \begin{pmatrix} x_{i1}^\top \\ \vdots \\ x_{in_i}^\top \end{pmatrix} = \begin{pmatrix} x_{1,i1} & \dots & x_{p,i1} \\ \vdots & & \\ x_{1,in_i} & \dots & x_{p,in_i} \end{pmatrix}_{n_i \times p}$$

تعریف می‌شود، ستون‌ها به مشاهدات مکرر هر متغیر در طول زمان مربوط بوده و ردیف‌ها به مشاهدات تمام متغیرها در یک زمان خاص مربوط هستند.

²Subject

³Repeated measurement

⁴Within-group

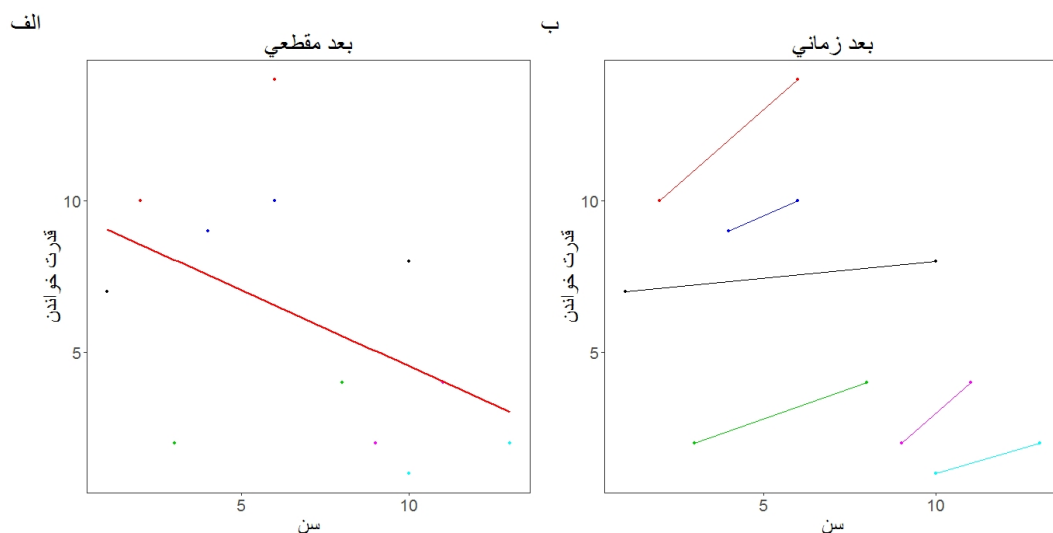
⁵Within-subject

⁶Response

⁷Explanatory

با توجه به تعریف و شیوه نمادگذاری بیان شده، این داده‌ها را می‌توان از دو بعد مورد توجه قرار داد: بعد زمانی و بعد مقطعی. به عبارت دیگر در مطالعات طولی بر خلاف مطالعات مقطعی^۸، علاوه بر بررسی اثر متغیرهای تبیینی^۹ می‌توان اثر زمان و اثر متقابل زمان-مداخله را نیز مورد بررسی قرار داد. به عنوان مثال اگر داده‌هایی را که در زمان z ام جمع‌آوری شده‌اند در نظر بگیریم، در واقع آزمودنی‌های اول تا n ام را تنها در یک مقطع از زمان مورد بررسی قرار داده‌ایم؛ بنابراین برای هر دوره زمانی، داده‌ها مقطعی هستند. اما اگر داده‌های مربوط به آزمودنی i ام را در نظر بگیریم، تنها یک آزمودنی را در زمان‌های متوالی مورد بررسی قرار داده‌ایم؛ بنابراین داده‌ها اثر زمان را نشان می‌دهند. لذا در تحلیل داده‌های طولی با سه هدف اصلی زیر روبرو هستیم:

۱. مطالعه تغییرات آزمودنی‌ها در طول زمان.
 ۲. مطالعه رابطه بین متغیرها (اثر متغیرهای تبیینی بر روی متغیر پاسخ).
 ۳. مطالعه تاثیر زمان بر رابطه بین متغیرها.
- برای درک بهتر موضوع مثالی از دیگل و همکاران (۲۰۰۲) را مورد بررسی قرار می‌دهیم.



شکل ۱۰.۱: نمودار قدرت خواندن کودکان

در شکل ۱۰.۱، داده‌ها از یک مطالعه طولی مربوط به قدرت خواندن کودکان به دست آمده‌اند که در آن شاخص توانایی مطالعه هر کودک، دو بار اندازه‌گیری شده است. بعد مقطعی داده‌ها در قسمت (الف)، کاهش قدرت خواندن کودکان با گذشت زمان را نشان می‌دهد؛ درحالی‌که بعد زمانی داده‌ها در قسمت (ب)، بیان می‌کند که توانایی هر کودک با گذشت زمان افزایش یافته است. در مجموع می‌توان گفت، درحالی‌که توانایی خواندن در کودکان جوان‌تر سطح بالاتری دارد، همه کودکان توانایی خود را با گذشت زمان بهبود می‌بخشند. بنابراین تحلیل داده‌های طولی که با هدف مشاهده تغییر رفتار پدیده‌ها در طول زمان انجام می‌گیرد، نسبت به داده‌های صرفاً مقطعی یا صرفاً زمانی از اهمیت

^۸Cross-sectional

^۹مداخله‌گرها

ویژه‌ای برخوردارند. این ماهیت دو بعدی، لزوم ارائه ترکیبی از ابزارهای آماری را برای تحلیل این داده‌ها آشکار می‌سازد.

۲.۱ ویژگی‌های داده‌های طولی

در مطالعات طولی هر دو متغیر پاسخ و تبیینی می‌توانند گسسته یا پیوسته و متغیرهای تبیینی می‌توانند در طول زمان ثابت یا متغیر باشند. همچنین فاصله نقاط زمانی و تعداد مشاهدات همه آزمودنی‌ها لزوماً برابر نیستند. این داده‌ها با توجه به نوع متغیرهای تبیینی، موقعیت و تعداد تکرارها به چندین دسته تقسیم‌بندی می‌شوند. شرح مختصری از این تقسیم‌بندی‌ها در زیر آمده است.

۱. اگر متغیرهای تبیینی در طول زمان ثابت باشند، یعنی به ازای هر $i = 1, \dots, n$ ، داشته باشیم

$$x_{k,i1} = x_{k,i2} = \dots = x_{k,in_i},$$

به متغیر تبیینی k ام، بین-گروهی^{۱۰} یا زمان-ثابت^{۱۱} گفته می‌شود. به عنوان مثال متغیرهای جنسیت و نژاد جزو این دسته هستند. به متغیر تبیینی k ام، $k = 1, \dots, p$ ، درون-گروهی یا زمان-متغیر^{۱۲} گفته می‌شود، اگر به ازای هر $i = 1, \dots, n$ و $j \neq j'$ داشته باشیم

$$x_{k,ij} \neq x_{k,ij'}.$$

متغیرهایی مانند قد، وزن و سن شامل این موارد هستند.

متغیرهای زمان-ثابت که می‌توانند در طول دوره متغیر باشند اما بر اساس طراحی از قبل تعیین شده مقدار ثابت و معینی گرفته‌اند را طرح-ثابت^{۱۳} می‌نامند. کد شناسایی آزمودنی‌ها و شاخص‌های درمانی شامل این مورد هستند. همچنین به متغیرهای زمان-ثابت، مانند وضعیت اولیه بیمار یا سن بیمار در زمان شروع مطالعه که مقادیر تصادفی هستند، طرح-تصادفی^{۱۴} گفته می‌شود.

۲. در متون مربوط به داده‌های طولی، اگر تمام آزمودنی‌ها دارای تعداد تکرارهای برابر باشند یعنی $n_i = m$ ، آنگاه طرح مطالعاتی را متوازن^{۱۵} و در غیر این صورت نامتوازن می‌نامند. همچنین یک طرح متوازن را کامل می‌گویند اگر همه n آزمودنی در موقعیت‌های زمانی مشابه اندازه‌گیری شده باشند، در غیر این صورت طرح را متوازن ناقص می‌نامند. برای درک بهتر این تقسیم‌بندی‌ها می‌توان به نمودار ساختگی در شکل ۲.۱ توجه نمود. در نمودار (الف) تعداد

¹⁰Between-group

¹¹Time-invariant

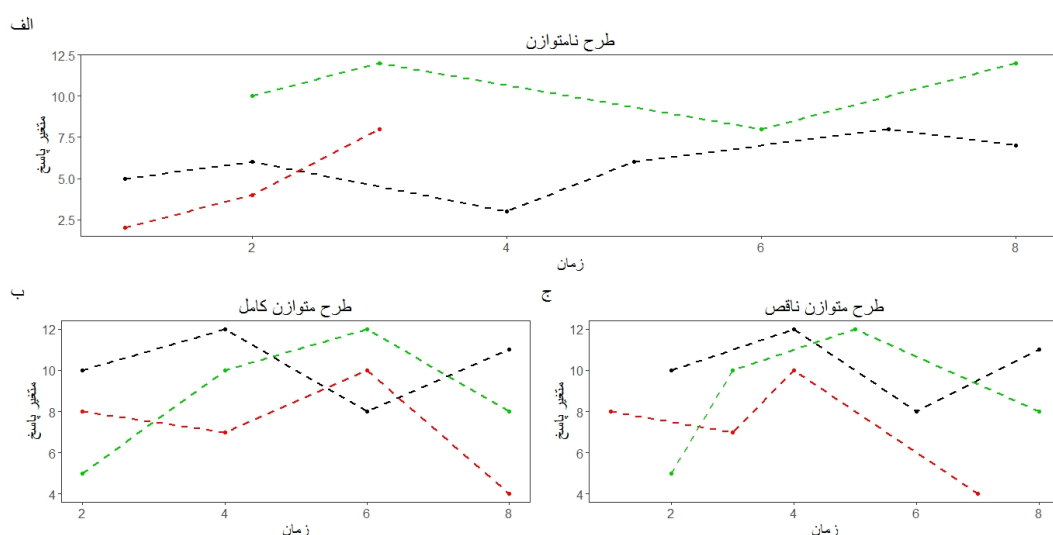
¹²Time-varying

¹³Fixed-design

¹⁴Stochastic-design

¹⁵Balanced

مشاهدات مکرر آزمودنی‌ها متفاوت‌اند بنابراین طرح مطالعاتی نامتوازن است. در نمودار (ب) علاوه بر این که تعداد تکرارهای آزمودنی‌ها برابراند، تمام آزمودنی‌ها در موقعیت‌های زمانی ۲، ۴، ۶ و ۸ اندازه‌گیری شده‌اند؛ لذا طرح متوازن کامل است و نمودار (ج) یک طرح متوازن ناقص را نمایش می‌دهد. یکی از دلایل عمده برخورد با طرح‌های نامتوازن داده‌های گمشده هستند. داده‌های گمشده مسئله‌ای کاملاً رایج در زمینه مطالعات طولی است. روش‌ها برای بررسی داده‌های گمشده به الگوی گم‌شدن و مکانیزم مقادیر گمشده وابسته هستند. شناخت کامل این روش‌ها و بحث درباره آن‌ها، یک زمینه گسترده مطالعاتی را بوجود می‌آورد که مستلزم صرف زمان بیشتری است. لذا در این رساله به این موضوع پرداخته نشده‌است. علاقه‌مندان برای آگاهی بیشتر در خصوص مقادیر گمشده در داده‌های طولی و انواع مدل‌بندی آن‌ها می‌توانند به فیتز‌موریس و همکاران (۲۰۰۹) رجوع کنند. در ادامه با تاکید بر روی ساختار همبستگی در داده‌های طولی، همبستگی درون-گروهی را با جزئیات بیشتری بررسی می‌کنیم.



شکل ۲.۱: نمودار ساختگی جهت نمایش تقسیم‌بندی طرح‌های مطالعاتی

۱.۲.۱ همبستگی درون گروهی

علاوه بر حجم بودن اطلاعات حاصل از داده‌های طولی و ماهیت دوبعدی آن‌ها که پیچیدگی داده‌ها را سبب می‌شوند، ویژگی بسیار مهم دیگری نیز وجود دارد که تحلیل و مدل‌بندی این داده‌ها را با چالش روبرو می‌سازد. برخلاف فرضیات روش‌های آماری متداول که بر پایه استقلال مشاهدات طراحی شده‌اند، در این داده‌ها مجموعه مشاهدات مربوط به هر واحد دارای همبستگی درونی بوده و تنها از مشاهدات سایر واحدها مستقل هستند. با توجه به همبستگی موجود در داده‌های طولی، فرضیات یک رگرسیون خطی را می‌توان به صورت زیر در نظر گرفت:

الف. $(y_1, X_1), \dots, (y_n, X_n)$ دوبه‌دو از هم مستقل‌اند.

ب. $E(y_i|X_i) = X_i\beta$.

ج. $\text{Cov}(y_i|X_i) = \Sigma_{i n_i \times n_i}$ ، که در آن Σ_i تابعی معلوم از متغیرهای تبیینی X_i است.

بدون کاستن از کلیت مسئله، می‌توان X_i ها را تصادفی در نظر گرفت. اگر X_{ij} ها طرح-ثابت باشند با احتمال ۱ مقادیر ثابت x_{ij} را می‌گیرند. اما هنگامی که X_{ij} ها طرح-تصادفی باشند برای سادگی در نوشتار میانگین و کوواریانس شرطی را به صورت $E(y_i)$ و $\text{Cov}(y_i)$ در نظر می‌گیریم. میانگین شرطی پاسخ z ام از نمونه i ام به شرط $X_{i1}, \dots, X_{in_i}$ یک تابع خطی از X_{ij} ها به صورت

$$E(y_{ij}|X_i) = X_{ij}^T \beta = \beta_1 x_{1,ij} + \dots + \beta_p x_{p,ij}$$

است. با استفاده از نماد ماتریسی می‌توان نوشت

$$E \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix}_{N \times 1} = \begin{pmatrix} X_1 \\ \vdots \\ X_n \end{pmatrix}_{N \times p} \times \begin{pmatrix} \beta_1 \\ \vdots \\ \beta_p \end{pmatrix}_{p \times 1}$$

یا

$$E(y)_{N \times 1} = X_{N \times p} \times \beta_{p \times 1}.$$

همچنین

$$\text{Cov}(y)_{N \times N} = \begin{pmatrix} \Sigma_{1 n_1 \times n_1} & 0 & \dots & 0 \\ & \Sigma_{2 n_2 \times n_2} & & 0 \\ & & \ddots & \vdots \\ & & & \Sigma_{n n_n \times n_n} \end{pmatrix}_{N \times N},$$

که در آن $N = \sum_{i=1}^n n_i$.

فرضیه ج این امکان را فراهم می‌کند که همبستگی بین مشاهدات هر نمونه را در مدل در نظر بگیریم. Σ_i می‌تواند تحت تاثیر عوامل مختلفی قرار گرفته باشد. یکی از این عوامل متغیرهای تبیینی هستند، مثلاً گروه‌های جنسیتی یا نژاد می‌توانند عامل همبستگی درون گروهی باشند؛ همچنین Σ_i می‌تواند تابعی از زمان بوده یا تنها از طریق i به n_i ها مرتبط باشد. در بسیاری از موارد تمرکز اصلی روی مطالعه تغییر رفتار متغیر پاسخ است که توسط $X_i\beta$ مدل‌بندی می‌شود و مطالعه روی Σ_i در اولویت‌های بعدی قرار گرفته، یا به عنوان عامل مزاحم در نظر گرفته می‌شود. در چنین حالتی، اگر طرح متوازن بوده یعنی $n_i = m$ و نسبت به n کوچک باشد و همچنین فرض شود که کوواریانس به متغیرهای تبیینی وابسته نیست، آنگاه $\Sigma_i = \Sigma$ را می‌توان یک ماتریس معین مثبت^{۱۶} در نظر گرفت، به‌طوری‌که اعضای آن با دقت بالایی قابل برآورد است. اما اگر m نسبت به n بزرگ بوده یا طرح

¹⁶Positive definite

نامتوازن باشد، آنگاه لازم است که برای Σ_i مدل‌های پارامتری در نظر بگیریم. بنابراین همبستگی درون-گروهی معمولاً به دو روش، شامل مدل‌های همبستگی سریالی و مدل‌های اثرات تصادفی در نظر گرفته می‌شوند. مدل‌های اثرات تصادفی در بخش‌های بعدی به تفصیل معرفی خواهند شد و در ادامه نمونه‌هایی از ساختار Σ بر اساس همبستگی سریالی با واریانس ثابت σ^2 را معرفی خواهیم کرد.

۱. همبستگی تبادلی پذیر^{۱۷}:

$$\Sigma = \sigma^2 \begin{pmatrix} 1 & \rho & \rho & \dots & \rho \\ \rho & 1 & \rho & \dots & \rho \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \rho & \rho & \rho & \dots & 1 \end{pmatrix}.$$

۲. همبستگی غیرساختاری^{۱۸} یا نواری^{۱۹}:

$$\Sigma = \sigma^2 \begin{pmatrix} 1 & \rho_1 & \rho_2 & \dots & \rho_{m-1} \\ \rho_1 & 1 & \rho_1 & \dots & \rho_{m-2} \\ \rho_2 & \rho_1 & 1 & \dots & \rho_{m-3} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \rho_{m-1} & \rho_{m-2} & \rho_{m-3} & \dots & 1 \end{pmatrix}.$$

۳. همبستگی اتو-رگرسیو^{۲۰}:

$$\Sigma = \sigma^2 \begin{pmatrix} 1 & \rho & \rho^2 & \dots & \rho^{m-1} \\ \rho & 1 & \rho & \dots & \rho^{m-2} \\ \rho^2 & \rho & 1 & \dots & \rho^{m-3} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \rho^{m-1} & \rho^{m-2} & \rho^{m-3} & \dots & 1 \end{pmatrix}.$$

همبستگی تبادلی‌پذیر یکی از ساده‌ترین ساختارهاست که در آن مشاهدات هر آزمودنی همبستگی یکسانی دارند، یعنی $\text{corr}(y_{ij}, y_{ik}) = \rho_{jk} \equiv \rho$. در این ساختار موقعیت مشاهدات در همبستگی آن‌ها تاثیر ندارد اما در سایر ساختارها از فاصله بین مشاهدات یا نقاط زمانی برای مدل‌سازی همبستگی استفاده می‌شود. در ساختار نواری همبستگی بین مشاهدات با ترتیب آن‌ها تعیین می‌گردد، یعنی تمام جفت مشاهدات هم‌جوار دارای همبستگی یکسان به صورت $\text{corr}(y_{ij}, y_{ij+k}) = \rho_k$.

¹⁷Exchangeable

¹⁸Unstructured

¹⁹Banded

²⁰Auto-regressive

مشاهداتی که به اندازه k واحد از هم فاصله دارند دارای همبستگی ρ_k می‌باشند. ساختار اتو-رگرسیو بیان شده برای طرح‌های متوازن مناسب است اما برای طرح‌های نامتوازن می‌توان مدل همبستگی را به صورت

$$\text{corr}(y_{ij}, y_{ik}) = \rho^{|t_{ij} - t_{ik}|}$$

در نظر گرفت که در آن $(t_{i1}, t_{i2}, \dots, t_{in_i})$ نقاط زمانی مشاهده شده برای آزمودنی i ام هستند. در این ساختار از یک پارامتر همبستگی مشترک استفاده می‌شود و فاصله نقاط زمانی مشاهدات عامل تفاوت همبستگی در موقعیت‌هاست. یعنی اگر $\rho = 0.8$ و اختلاف نقاط زمانی مشاهدات ۱ واحد است همبستگی برابر $0.8^1 = 0.8$ بوده و اگر این اختلاف به اندازه ۲ واحد باشد همبستگی برابر $0.8^2 = 0.64$ است. در این ساختار با افزایش فاصله بین مشاهدات میزان همبستگی تنزل می‌یابد. علاوه بر ساختارهای ذکر شده، مدل‌های متعدد دیگری نیز وجود دارند که برای آشنایی بیشتر می‌توان به دیگل و همکاران (۲۰۰۲) و وریک و مولنبرگ (۲۰۰۹) مراجعه نمود. در ادامه نشان می‌دهیم همبستگی تاثیر بسزایی در کارایی برآوردگرها دارد.

یک روش ساده برای تحلیل داده‌های همبسته این است که همبستگی را نادیده گرفته و با استفاده از روش‌های کلاسیک به برآورد پارامترها بپردازیم. اما این کار به‌ویژه هنگامی که میزان همبستگی بالاست، موجب کاهش کارایی برآوردگرها می‌شود. به عنوان مثال در یک مدل رگرسیونی خطی ساده، میانگین همه مشاهدات و واریانس آن را در نظر بگیرید.

$$\bar{Y} = \frac{1}{\sum_{i=1}^n n_i} \sum_{i=1}^n \sum_{j=1}^{n_i} y_{ij},$$

$$\text{Var}(\bar{Y}) = \frac{1}{(\sum_{i=1}^n n_i)^2} \sum_{i=1}^n \left[\sum_{j=1}^{n_i} \text{Var}(y_{ij}) + 2 \sum_{j=1}^{n_i-1} \sum_{k=j+1}^{n_i} \text{Cov}(y_{ij}, y_{ik}) \right].$$

اگر واریانس مشاهدات ثابت باشد، $\text{Var}(y_{ij}) = \sigma^2$ ، آنگاه

$$\text{Var}(\bar{Y}) = \frac{\sigma^2}{(\sum_{i=1}^n n_i)^2} \sum_{i=1}^n \left[n_i + \frac{2}{\sigma^2} \sum_{j=1}^{n_i-1} \sum_{k=j+1}^{n_i} \text{Cov}(y_{ij}, y_{ik}) \right].$$

حال اگر فرض کنیم طرح متوازن بوده، $n_i \equiv m$ ، و ساختار همبستگی متبادل پذیر است، $\rho_{jk} = \rho$ ، واریانس میانگین به صورت

$$\text{Var}(\bar{Y}) = \frac{\sigma^2}{nm} [1 + (m-1)\rho].$$

محاسبه می‌شود. به مقدار $[1 + (m-1)\rho]$ عامل تورم واریانس گفته می‌شود؛ زیرا این مقدار واریانس میانگین را افزایش می‌دهد که به دلیل وجود همبستگی درون-گروهی است. برای نشان دادن تاثیر همبستگی روی واریانس میانگین، $[1 + (m-1)\rho]$ به ازای مقادیر مختلف m و ρ در جدول ۱۰.۱ آمده است.

جدول ۱۰۱: مقدار $[1 + (m - 1)\rho]$ به ازای مقادیر مختلف m و ρ

$\rho \backslash m$	۰/۰۰۱	۰/۰۱	۰/۰۲	۰/۰۵	۰/۱
۳	۱/۰۰۱	۱/۰۱	۱/۰۲	۱/۰۵	۱/۱۰
۶	۱/۰۰۴	۱/۰۴	۱/۰۸	۱/۲۰	۱/۴۰
۱۱	۱/۰۰۹	۱/۰۹	۱/۱۸	۱/۴۵	۱/۹۰
۱۰۱	۱/۰۹۹	۱/۹۹	۲/۹۸	۵/۹۵	۱۰/۹۰
۱۰۰۱	۱/۹۹۹	۱۰/۹۹	۲۰/۹۸	۵۰/۹۵	۱۰۰/۹۰

مقادیر جدول ۱۰۱ نشان می‌دهد که حتی اگر همبستگی درون آزمودنی‌ها بسیار ناچیز بوده اما تعداد تکرارهای هر نمونه زیاد باشد، میزان همبستگی تأثیر مهمی بر خطاهای استاندارد دارد. به عنوان مثال، در حالت‌های $(\rho = 0/001, n = 1001)$ ، $(\rho = 0/01, n = 101)$ و $(\rho = 0/1, n = 11)$ مقدار $[1 + (m - 1)\rho]$ تقریباً برابر ۲ است. با توجه به مطالب بیان‌شده، واضح است که جهت رسیدن به استنباط علمی و معتبر، در نظر گرفتن این همبستگی امری ضروری است.

۳.۱ مروری بر مدل‌بندی داده‌های طولی

در بخش‌های قبل دلایلی بیان شد که روش‌های آماری کلاسیک که عمدتاً بر پایه استقلال مشاهدات بنا شده‌اند، به تنهایی برای تحلیل داده‌های طولی مناسب نیستند. بدین منظور آماردانان روش‌های متعددی را معرفی نموده‌اند که تمرکز اصلی این روش‌ها احتساب همبستگی درون آزمودنی‌ها و ارزیابی اثرات متغیرهای تبیینی روی پاسخ‌ها به‌طور همزمان است.

گسترش روش‌های اولیه در زمینه تحلیل داده‌های طولی که متغیر پاسخ آن پیوسته است، به حدود یک قرن پیش باز می‌گردد. پایه این روش‌ها توسط یک ستاره‌شناس انگلیسی به نام آیری (۱۸۶۱) با معرفی مدل تحلیل واریانس اثرات آمیخته^{۲۱} بنا شده است. این مدل که اغلب با عنوان تحلیل واریانس اندازه‌های مکرر تک‌متغیره^{۲۲} شناخته می‌شود دارای تاریخچه‌ای طولانی در زمینه تحلیل این داده‌هاست. از آنجایی که ساختار داده‌های طولی با n آزمودنی و n_i اندازه‌گیری مکرر، مشابه داده‌های جمع‌آوری‌شده در طرح بلوک تصادفی است، طبیعی به نظر می‌رسد که برای تحلیل این داده‌ها از روش‌های تحلیل واریانس استفاده شود. بدین منظور در این نوع مطالعات، آزمودنی‌ها به عنوان بلوک در نظر گرفته شده و مدل به‌صورت

$$y_{ij} = \mathbf{x}_{ij}^T \boldsymbol{\beta} + b_i + \varepsilon_{ij}, \quad i = 1, \dots, n; \quad j = 1, \dots, n_i,$$

^{۲۱}Mixed effects anova

^{۲۲}Univariate repeated measures anova

نوشته می‌شود. در این مدل $b_i \sim N(0, \sigma_b^2)$ و $\varepsilon_{ij} \sim N(0, \sigma_\varepsilon^2)$ به ترتیب عامل تصادفی و جمله خطا بوده که از هم مستقل‌اند. در این مدل، بلوک‌ها یا اثرات طرح به عنوان اثرات تصادفی مورد توجه قرار می‌گیرند نه اثرات ثابت. اثر تصادفی، b_i ، نشان‌دهنده تمام عوامل نامطلوب یا مشاهده‌نشده‌ای است که باعث شده آزمودنی‌ها رفتار یا پاسخ متفاوتی را نشان دهند. نتیجه‌گزینش این مولفه تصادفی در مدل، اعمال همبستگی مثبت به مشاهدات مکرر است به‌طوری‌که پاسخ‌ها دارای واریانس ثابت به صورت $Var(y_{ij}) = \sigma_b^2 + \sigma_\varepsilon^2$ و کوواریانس ثابت به صورت $Cov(y_{ij}, y_{ik}) = \sigma_b^2$ هستند. مشابه همه مدل‌های رگرسیونی ε_{ij} نشان‌دهنده همه مواردی است که در مدل‌بندی در نظر گرفته نشده است. اگر چه روش معرفی‌شده سابقه‌ای طولانی و گسترده در تحلیل داده‌های طولی دارد، فرض‌های محدودکننده و بسیاری از نقص‌های آشکار آن، سودمندی آن را در برنامه‌های کاربردی محدود کرده و این شیوه تنها در مواردی که طرح مطالعاتی بسیار ساده است، می‌تواند مبنایی منطقی برای تحلیل داده‌های طولی فراهم آورد. ویژگی‌های داده‌های طولی از قبیل طرح‌های نامتوازن و ناقص، وجود داده‌های گمشده و پاسخ‌های گسسته، که موجب محدودیت در استفاده از روش‌های تحلیل واریانس شده بود، انگیزه آماردانان را در راستای ارائه تکنیک‌های همه‌جانبه‌ای که دربرگیرنده این ویژگی‌ها باشد برانگیخت. گسترش تحلیل داده‌های طولی گسسته از ۳۰ تا ۳۵ سال پیش آغاز شده و تاکنون پیشرفت‌های قابل‌توجهی داشته است. این روش‌ها عمدتاً بر روی تحلیل رگرسیونی داده‌ها متمرکز شده‌اند. تحلیل رگرسیونی روشی ساده برای بررسی روابط تابعی میان متغیرهاست. این رابطه به شکل یک معادله یا الگویی است که متغیر پاسخ را به یک یا چند متغیر تبیینی مرتبط می‌کند. رگرسیون را بر حسب نوع تابعی که رابطه بین متغیرها را تعیین می‌کند، می‌توان به سه نوع پارامتری، ناپارامتری و نیمه‌پارامتری تقسیم نمود. رگرسیون پارامتری پیش‌تر وقتی مورد استفاده قرار می‌گیرد که اطلاعات تحقیق به قدر کافی زیاد است تا آن‌جایی که بتوان ارتباط بین متغیرهای پاسخ و تبیینی را به وسیله یک مدل خطی یا قابل تبدیل به خطی بیان کرد. این مدل‌ها با وجود سادگی در محاسبات و تفسیر، در مدل‌سازی بسیاری از روابط پیچیده با شکست مواجه می‌شوند؛ زیرا در کاربرد مثال‌های بسیاری را می‌توان یافت که شناسایی و تشخیص اثرات رگرسیون بر اساس فرضیات بیان‌شده امکان‌پذیر نیست. در چنین شرایطی تکنیک‌های مدل‌سازی ناپارامتری مورد نیاز است. رگرسیون ناپارامتری در مواقعی که نتوان نوع ارتباط بین متغیرها را با یک صورت تابعی معلوم و مشخص بیان کرد، مورد استفاده قرار می‌گیرد. در این مدل‌ها تنها فرض می‌شود که تابع مورد نظر یک تابع هموار است که برای برآورد آن می‌توان از روش‌های موضعی مانند اسپلاین و روش‌های هسته‌ای استفاده کرد. برآوردهای رگرسیون ناپارامتری بسیار انعطاف‌پذیر هستند اما با افزایش تعداد متغیرهای تبیینی، دقت آماری آن‌ها به شدت کاهش می‌یابد که این مسئله به مشقت بعد^{۲۳} معروف است. در نتیجه محققین سعی کرده‌اند به مدل‌ها و برآوردهایی بپردازند که انعطاف‌پذیری بیشتری نسبت به رگرسیون پارامتری دارند و در عین حال بر مشکل مشقت بعد غلبه می‌کنند. چنین روش‌هایی معمولاً خصوصیات روش‌های پارامتری و ناپارامتری را با هم ترکیب کرده و به روش‌های نیمه‌پارامتری معروف هستند. در این مدل‌ها فرض می‌شود بعضی از متغیرها دارای روند معلوم تغییرات نسبت به

²³ Curse of dimension

متغیر پاسخ بوده و برای بعضی از متغیرها نمی‌توان صورت تابعی خاصی را در نظر گرفت. در چند دهه اخیر، آماردانان در راستای تحلیل هرچه بهتر داده‌های طولی با خلق ترکیبی از مدل‌های ذکر شده، رده دیگری از مدل‌ها را ایجاد نمودند. در این میان، مدل اثرات آمیخته از مجموعه مدل‌های پارامتری و مدل‌های نیمه‌پارامتری مبنای بسیاری از این دسته مدل‌های ترکیبی قرار گرفتند. مدل نیمه‌پارامتری با اثرات آمیخته که تمرکز اصلی این رساله است، به عنوان مدلی بسیار منعطف و تنومند که ویژگی‌های خوب اکثر مدل‌ها را دارد، مورد توجه بسیار قرار گرفته است. در ادامه ویژگی‌های هر یک از مدل‌های بیان شده و روش‌های برآورد در آن‌ها به تفصیل شرح داده شده است.

۱.۳.۱ مدل‌های پارامتری

هنگامی که علاقه‌مند به تحلیل داده‌های گسسته مانند دودویی، شمارشی، اسمی و ترتیبی باشیم، مدل‌های خطی در تحلیل این داده‌ها ناکارآمد می‌باشند. در این حالت مدل‌های خطی تعمیم‌یافته^{۲۴} (نلدر و ودربرن، ۱۹۷۲)، رده جامعی از مدل‌ها برای تحلیل رگرسیونی داده‌های گسسته و پیوسته که دارای توزیع‌های نمایی باشند، فراهم آورده است؛ به این ترتیب که با اعمال یک تبدیل غیرخطی مناسب روی میانگین متغیر پاسخ، ترکیبی خطی از متغیرهای تبیینی می‌سازد. بسط مدل‌های خطی تعمیم‌یافته به داده‌های همبسته و چندمتغیره، توانایی این مدل‌ها برای تحلیل داده‌های طولی را فراهم نموده است. سه رده گسترده از این مدل‌ها شامل مدل‌های حاشیه‌ای^{۲۵}، مدل‌های اثرات آمیخته^{۲۶} و مدل‌های انتقالی^{۲۷} هستند. پاسخ به این سوال که کدام مدل باید انتخاب شود عمدتاً به سوال یا سوالات تحقیقات بستگی دارد که باید به آن‌ها جواب داده شود. لذا لازم است که با ویژگی‌های هر مدل آشنا شویم.

در مدل‌های حاشیه‌ای، مدل‌بندی بر حسب میانگین پاسخ انجام می‌گیرد و هیچ یک از اثرهای تصادفی یا پاسخ‌های گذشته در آن سهمی ندارند. علاوه بر این میانگین متغیر پاسخ و همبستگی بین اندازه‌های مکرر ثبت شده از یک آزمودنی جداگانه تحلیل می‌شود. وابستگی درون-فردی بین متغیرهای پاسخ با در نظر گرفتن ساختارهای مشخصی برای ماتریس همبستگی منظور می‌شود. انتخاب ساختارهای همبستگی مانند تبادل پذیر، غیرساختاری و اتو-رگرسیو به پیش فرض‌های موجود در مشاهدات و دلایل علمی بستگی دارد. روش دیگر استفاده از مدل اثرات تصادفی است. در این مدل فرض می‌شود همبستگی بین پاسخ‌های تکرار شده برای هر آزمودنی به دلیل ناهمگونی بین آزمودنی‌ها و وجود عوامل اندازه‌گیری نشده و غیرقابل کنترلی مانند عوامل محیطی و ژنتیکی است. بنابراین لازم است این عوامل به عنوان یک عامل تصادفی که از یک آزمودنی به آزمودنی دیگر متفاوت است در نظر گرفته شوند. بنابراین مدل اثرات تصادفی به دلیل دخالت ضرایب رگرسیونی

²⁴Generalized linear models

²⁵Marginal models

²⁶Mixed effects models

²⁷Transition models

تصادفی، متفاوت از مدل حاشیه‌ای عمل می‌کند. معمولاً مدل حاشیه‌ای که از طریق روش‌هایی مرسوم به معادلات برآوردیابی تعمیم‌یافته برازش می‌شود، با عنوان میانگین-جمعیتی^{۲۸} شهرت دارد و مدل اثرات تصادفی به عنوان مدل‌های ویژه-آزمودنی^{۲۹} در منابع ذکر شده است. در مدل‌های انتقالی که به عنوان سومین رویکرد تحلیل داده‌های طولی از آن یاد شده، فرض می‌شود که متغیر پاسخ در یک زمان مشخص به متغیرهای پاسخ در زمان‌های قبلی وابسته است. این تحلیل توسط اصول حاکم بر زنجیرهای مارکوف انجام می‌گیرد و همبستگی‌ها از طریق ارتباط با پاسخ‌های قبلی مورد بررسی قرار می‌گیرند.

۱.۱.۳.۱ مدل‌های حاشیه‌ای

ویژگی اصلی مدل‌های حاشیه‌ای این است که، تاثیر متغیرهای تبیینی روی متغیر پاسخ به‌طور مجزا از همبستگی بین پاسخ‌ها مدل‌بندی می‌شود. به عبارت دیگر، مدل‌های جداگانه‌ای برای میانگین حاشیه‌ای، واریانس حاشیه‌ای و همبستگی حاشیه‌ای در نظر گرفته می‌شوند. به‌طور معمول یک مدل حاشیه‌ای دارای سه مشخصه زیر است:

آ) میانگین حاشیه‌ای، $E(y_{ij}) = \mu_{ij}$ ، از طریق تابع پیوند g به متغیرهای تبیینی x_{ij} وابسته است:

$$g(\mu_{ij}) = x_{ij}^T \beta, \quad i = 1, 2, \dots, n; \quad j = 1, 2, \dots, n_i,$$

که در آن g تابعی مشخص معروف به تابع پیوند مانند لگاریتم $(g(x) = \log(x))$ یا لوجیت $(g(x) = \log(\frac{x}{x-1}))$ است.

ب) واریانس حاشیه‌ای تابعی از میانگین حاشیه‌ای و پارامتر مقیاس ϕ است:

$$\text{var}(y_{ij}|x_{ij}) = \nu(\mu_{ij})\phi,$$

که در آن ν یک تابع مشخص است.

پ) همبستگی بین y_{ij} و y_{ik} ، تابعی از میانگین‌های حاشیه‌ای و پارامترهای دیگر مانند α است:

$$\text{corr}(y_{ij}, y_{ik}) = \rho(\mu_{ij}, \mu_{ik}, \alpha)$$

که در آن ρ ، تابعی معلوم و α ، برداری از پارامترها است.

روش‌های متعددی برای برآورد پارامترهای مدل (β, ϕ, α) وجود دارند، که برآورد ماکزیمم درستنمایی^{۳۰} (ML) و کمترین توان‌های دوم تعمیم‌یافته^{۳۱} (GLS) از جمله متداول‌ترین آن‌هاست. برای داده‌هایی

²⁸Population-average

²⁹Subject-specific

³⁰Maximum likelihood

³¹Generalized least square

با توزیع نرمال این دو روش معادل هستند. در این حالت با تعیین امید و واریانس حاشیه‌ای پاسخ‌ها، توزیع توام پاسخ‌ها نرمال چندمتغیره به دست می‌آید. در نتیجه می‌توان با استفاده از روش MLE ، به برآورد پارامترهای مدل پرداخت. اما هنگامی که با داده‌های غیرنرمال سروکار داریم، محاسبه توزیع توام پاسخ‌ها امکان‌پذیر نبوده و این روش‌ها ناکارآمد هستند. در این حالت استفاده از روش معادلات برآورد تعمیم‌یافته^{۳۲} (GEE) توصیه می‌گردد. این روش که توسط لیانگ و زیگر (۱۹۸۶) پیشنهاد شده است، یک روش برآورد در مدل‌های خطی تعمیم‌یافته به شمار می‌آید. این روش که تعمیمی از روش شبه‌درست‌نمایی (ودرین، ۱۹۷۴) به داده‌های چندمتغیره و همبسته است، دسته‌ای از معادلات برآورد برای پارامترهای مدل رگرسیونی ارائه می‌دهد. با فرض اینکه $\mu_i(\beta)$ و $V_i(\alpha)$ به ترتیب میانگین (گشتاور مرتبه اول) و ماتریس واریانس-کوواریانس (گشتاور مرتبه دوم) برای y_i باشند، معادلات برآورد تعمیم‌یافته به صورت

$$S(\beta, \alpha) = \sum_{i=1}^n S_i(\beta, \alpha) = \sum_{i=1}^n \frac{\partial \mu_i^\top}{\partial \beta} V_i^{-1}(\alpha) (y_i - \mu_i(\beta)) = 0$$

تعریف می‌شوند، که در آن

$$\frac{\partial \mu_i}{\partial \beta} = \begin{pmatrix} \frac{\partial g^{-1}(x_{i1}^\top \beta)}{\partial \beta} \\ \frac{\partial g^{-1}(x_{i2}^\top \beta)}{\partial \beta} \\ \vdots \\ \frac{\partial g^{-1}(x_{in_i}^\top \beta)}{\partial \beta} \end{pmatrix} = \begin{pmatrix} \frac{x_{1,i}}{g'(\mu_{i1})} & \cdots & \frac{x_{p,i}}{g'(\mu_{ip})} \\ \vdots & & \vdots \\ \frac{x_{1,in_i}}{g'(\mu_{in_i})} & \cdots & \frac{x_{p,in_i}}{g'(\mu_{in_i})} \end{pmatrix}_{n_i \times p}$$

و ماتریس واریانس-کوواریانس، $V_i(\alpha)$ ، به صورت

$$V_i(\alpha) = \phi A_i^\top R_i(\alpha) A_i$$

است که در آن، A_i یک ماتریس قطری $n_i \times n_i$ بعدی با عناصر قطری $\nu(\mu_{ij})$ بوده و $A_i^\top$ ریشه دوم ماتریس A_i است. همچنین $R_i(\alpha)$ که به عنوان ماتریس همبستگی عملی^{۳۳} شناخته می‌شود، ساختارهای متعددی دارند که در بخش‌های قبل چند نمونه ذکر شد. در تعیین این معادلات، به گشتاورهای مرتبه اول و دوم پاسخ‌ها و تعیین ماتریس همبستگی برای آزمودنی‌ها نیاز است؛ و درباره توزیع توام پاسخ‌ها هیچ فرضی اختیار نمی‌شود. برای داده‌هایی با توزیع نرمال این روش معادل روش ML یا GLS است. برای حل معادلات فوق از الگوریتم تکراری نیوتن-رافسون استفاده شده و برای برآورد بردار ضرایب رگرسیونی β ، رابطه تکراری به صورت

$$\hat{\beta}^{(r+1)} = \hat{\beta}^{(r)} + \left(\sum_{i=1}^n \frac{\partial \mu_i^\top}{\partial \beta} V_i^{-1}(\hat{\alpha}) \frac{\partial \mu_i}{\partial \beta} \right)^{-1} \left(\sum_{i=1}^n \frac{\partial \mu_i^\top}{\partial \beta} V_i^{-1}(\hat{\alpha}) (y_i - \mu_i) \right)$$

^{۳۲}Generalized estimating equation

^{۳۳}Working correlation matrix

تا رسیدن به همگرایی محاسبه می‌شود. برای برآورد پارامترهای همبستگی (α) می‌توان از روش‌های گشتاوری استفاده کرد. پرنیتس (۱۹۸۸)، دسته‌ای دیگر از معادلات را به‌طور مجزا برای برآورد پارامترهای همبستگی پیشنهاد داده است که به عنوان تعمیم GEE شناخته می‌شوند. زائو و پرنیتس (۱۹۹۰) به برآورد توام پارامترهای رگرسیونی و همبستگی پرداخته‌اند که بسط درجه دوم GEE ^{۳۴} نام گرفته است. هر کدام از این روش‌ها، ویژگی، معایب و محاسن منحصر به فرد خود را دارند. شناخت کامل این روش‌ها و بحث درباره خصوصیات آن‌ها، یک زمینه گسترده مطالعاتی را بوجود می‌آورد که مستلزم صرف زمان بیشتری است. لذا در این رساله، تنها به ذکر چند ویژگی برجسته این روش بسنده می‌کنیم. علاقه‌مندان برای آشنایی بیشتر با روش‌های GEE ، تعمیم و بسط درجه دوم آن می‌توانند به تحقیقات، فیتزموریس و لرد (۱۹۹۳)، سوترادهار و داس (۱۹۹۹) و سوترادهار و فارل (۲۰۰۴) مراجعه نمایند.

این روش برآوردگرهای سازگاری برای ضرایب مدل ارائه می‌دهد که مستقل از انتخاب ساختار همبستگی بوده و $\sqrt{n}(\hat{\beta}_{GEE} - \beta)$ دارای توزیع مجانبی نرمال چند متغیره با میانگین ۰ و ماتریس واریانس-کوواریانس به‌صورت زیر است:

$$\lim_{n \rightarrow \infty} n \left(\sum_{i=1}^n D_i^T V_i^{-1} D_i \right)^{-1} \left\{ \sum_{i=1}^n D_i^T V_i^{-1} \text{Cov}(y_i) V_i^{-1} D_i \right\} \left(\sum_{i=1}^n D_i^T V_i^{-1} D_i \right)^{-1}.$$

اگر $V_i = \text{Cov}(y_i)$ به درستی تشخیص داده شود، آنگاه رابطه فوق به $\lim_{n \rightarrow \infty} n \left(\sum_{i=1}^n D_i^T V_i^{-1} D_i \right)^{-1}$ تبدیل خواهد شد. اما عدم تشخیص درست نوع ساختار همبستگی موجب کاهش کارایی برآوردگرها می‌شود که این نقص در بسط درجه دوم GEE برطرف شده است.

گسترش این شیوه باعث شده است که بسیاری از نرم افزارهای آماری از جمله R و SAS ، به محیط و کدهای اختصاصی این روش مجهز شده و اجازه انتخاب ساختارهای همبستگی مختلف را به کاربران بدهند. در پایان لازم به ذکر است که در این نرم افزارها، از معیار اطلاع تعمیم یافته^{۳۵} (GIC) و معیار اطلاع همبستگی^{۳۶} (CIC) به عنوان معیارهایی برای تشخیص ساختار همبستگی مناسب، تشخیص کارایی برآوردگرها و نیکویی برازش مدل استفاده می‌گردد. این شاخص‌ها که توسط پن (۲۰۰۱) ارائه شده است، تعمیمی از معیار اطلاع آکائیک^{۳۷} (AIC) است. علاقه‌مندان جهت آشنایی بیشتر با این شاخص‌ها می‌توانند به هاردین و هیلب (۲۰۰۳) و زیگلر و ونس (۲۰۱۰) رجوع نمایند.

^{۳۴}Second order GEE

^{۳۵}Generalized information criterion

^{۳۶}Correlation information criterion

^{۳۷}Akaike information criterion

۲.۱.۳.۱ مدل اثرات آمیخته

مدل اثرات آمیخته یا اثرات تصادفی^{۳۸}، نوع دیگری از تعمیم مدل‌های خطی تعمیم‌یافته به‌شمار می‌آید. در این مدل‌ها عموماً فرض بر آن است که برخی ضرایب رگرسیونی برای آزمودنی‌ها به دلایل متفاوت به خود آزمودنی‌ها وابسته بوده و دارای توزیعی جداگانه می‌باشند. به عبارت دیگر همبستگی بین مشاهدات در داده‌های تکراری را می‌توان ناشی از داشتن یک تاثیر تصادفی یکسان در آن‌ها دانست به‌طوری‌که در مورد یک فرد اثر تصادفی یکسان و در فرد دیگر این اثر متفاوت است. در حقیقت همین ویژگی است که موجب ایجاد همبستگی در اندازه‌های مکرر هر آزمودنی می‌گردد. به عنوان مثال رگرسیون خطی ساده برای رشد کودک را در نظر بگیرید که در آن ضرایب مدل، بیانگر تاثیر وزن تولد و شاخص رشد است. واضح است که کودکان در زمان تولد، دارای وزن و شاخص رشد متفاوت با یکدیگر هستند. بنابراین لازم است این عوامل به عنوان یک عامل تصادفی که از یک آزمودنی به آزمودنی دیگر متفاوت است در نظر گرفته شوند. در این مدل، متغیر پاسخ به‌صورت تابعی از متغیرهای تبیینی با ضرایب ثابت و همچنین زیربرداری از آن‌ها با ضرایب تصادفی ارائه می‌گردد. این مدل که اولین بار توسط هارویل (۱۹۷۷) معرفی شد، برای متغیرهای پاسخ با توزیع نرمال به‌صورت

$$y_{ij} = x_{ij}^T \beta + z_{ij}^T b_i + \varepsilon_{ij}, \quad i = 1, 2, \dots, n; \quad j = 1, 2, \dots, n_i,$$

بیان می‌گردد، که در آن y_{ij} و ε_{ij} به‌ترتیب متغیر پاسخ و خطای آزمودنی i ام در j امین تکرار است. پارامترهای $\beta = (\beta_1, \dots, \beta_p)^T$ و $b_i = (b_{i1}, \dots, b_{iq})^T$ ، به‌ترتیب به عنوان اثرات ثابت و تصادفی شناخته می‌شوند. بردار متغیرهای تبیینی مربوط به اثرات ثابت بوده و $x_{ij} = (x_{1,ij}, \dots, x_{p,ij})^T$ بردار متغیرهای تبیینی مربوط به اثرات تصادفی می‌باشند که با دخالت $z_{ij} = (z_{1,ij}, \dots, z_{q,ij})^T$ برداری از متغیرهای تبیینی مربوط به اثرات تصادفی می‌باشند که با دخالت ضرایب تصادفی b_i ، ماهیت تصادفی به این بردار اعمال شده است. برای هر آزمودنی i ام مدل به‌صورت

$$y_i = X_i \beta + Z_i b_i + \varepsilon_i, \quad i = 1, 2, \dots, n$$

نمایش داده می‌شود، که در آن $y_i = (y_{i1}, \dots, y_{in_i})^T$ و $\varepsilon_i = (\varepsilon_{i1}, \dots, \varepsilon_{in_i})^T$ بردارهایی n_i بعدی هستند. ماتریس طرح اثرات ثابت به‌صورت

$$X_i = \begin{pmatrix} x_{i1}^T \\ \vdots \\ x_{in_i}^T \end{pmatrix} = \begin{pmatrix} x_{1,i1} & x_{2,i1} & \dots & x_{p,i1} \\ x_{1,i2} & x_{2,i2} & \dots & x_{p,i2} \\ \vdots & \vdots & \ddots & \vdots \\ x_{1,in_i} & x_{2,in_i} & \dots & x_{p,in_i} \end{pmatrix}_{n_i \times p}$$

³⁸Random effects model

و ماتریس طرح اثرات تصادفی به صورت

$$Z_i = \begin{pmatrix} z_{i1}^\top \\ \vdots \\ z_{in_i}^\top \end{pmatrix} = \begin{pmatrix} z_{1,i1} & z_{2,i1} & \dots & z_{q,i1} \\ z_{1,i2} & z_{2,i2} & \dots & z_{q,i2} \\ \vdots & \vdots & \ddots & \vdots \\ z_{1,in_i} & z_{2,in_i} & \dots & z_{q,in_i} \end{pmatrix}_{n_i \times q}$$

است. اگر مدل شامل عرض از مبدا ثابت باشد، ستون اول ماتریس X_i برابر با ۱ بوده و اگر مدلی با عرض از مبدا تصادفی داشته باشیم عناصر ستون اول ماتریس Z_i برابر ۱ است. توجه داشته باشید که z_{ij} زیربرداری از x_{ij} است و اغلب Z_i نسبت به ماتریس X_i ستون‌های کمتری دارد یعنی $q \leq p$. در این مدل، فرض می‌شود که بردارهای ε_i و b_i مستقل از هم بوده و دارای توزیع نرمال به صورت $\varepsilon_i \sim \mathcal{N}_{n_i}(\mathbf{0}, \Sigma_i)$ و $b_i \sim \mathcal{N}_q(\mathbf{0}, D_i)$ است. ماتریس کوواریانس اثرات تصادفی است که عناصر روی قطر اصلی نشان‌دهنده واریانس هر عضو از بردار b_i بوده و عناصر غیرقطری، کوواریانس بین اثرات تصادفی است. از آنجایی که در این مدل q اثر تصادفی موجود است، ماتریس D_i یک ماتریس $q \times q$ بعدی متقارن و معین مثبت است که عناصر آن به صورت

$$D_i = \text{Var}(b_i) = \begin{pmatrix} \text{Var}(b_{i1}) & \text{Cov}(b_{i1}, b_{i2}) & \dots & \text{Cov}(b_{i1}, b_{iq}) \\ \text{Cov}(b_{i2}, b_{i1}) & \text{Var}(b_{i2}) & \dots & \text{Cov}(b_{i2}, b_{iq}) \\ \vdots & \vdots & \ddots & \vdots \\ \text{Cov}(b_{iq}, b_{i1}) & \text{Cov}(b_{iq}, b_{i2}) & \dots & \text{Var}(b_{iq}) \end{pmatrix}_{q \times q}$$

است. به طور مشابه ماتریس کوواریانس خطاها را می‌توان به صورت

$$\Sigma_i = \text{Var}(\varepsilon_i) = \begin{pmatrix} \text{Var}(\varepsilon_{i1}) & \text{Cov}(\varepsilon_{i1}, \varepsilon_{i2}) & \dots & \text{Cov}(\varepsilon_{i1}, \varepsilon_{in_i}) \\ \text{Cov}(\varepsilon_{i2}, \varepsilon_{i1}) & \text{Var}(\varepsilon_{i2}) & \dots & \text{Cov}(\varepsilon_{i2}, \varepsilon_{in_i}) \\ \vdots & \vdots & \ddots & \vdots \\ \text{Cov}(\varepsilon_{in_i}, \varepsilon_{i1}) & \text{Cov}(\varepsilon_{in_i}, \varepsilon_{i2}) & \dots & \text{Var}(\varepsilon_{in_i}) \end{pmatrix}_{n_i \times n_i}$$

نمایش داد. نمایش ماتریسی بر پایه همه داده‌ها را می‌توان به صورت

$$y = X\beta + Zb + \varepsilon,$$

$$b \sim \mathcal{N}_{nq}(\mathbf{0}, D), \quad \varepsilon \sim \mathcal{N}_N(\mathbf{0}, \Sigma),$$

نوشت، که در آن

$$\begin{aligned} \mathbf{y} &= (\mathbf{y}_1^\top, \dots, \mathbf{y}_n^\top)^\top, & \boldsymbol{\varepsilon} &= (\boldsymbol{\varepsilon}_1^\top, \dots, \boldsymbol{\varepsilon}_n^\top)^\top, & \mathbf{b} &= (\mathbf{b}_1^\top, \dots, \mathbf{b}_n^\top)^\top, \\ \mathbf{X} &= (\mathbf{X}_1^\top, \dots, \mathbf{X}_n^\top)^\top, & \mathbf{Z} &= \begin{pmatrix} \mathbf{Z}_1 & \mathbf{0} & \dots & \mathbf{0} \\ \mathbf{0} & \mathbf{Z}_2 & \dots & \mathbf{0} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{0} & \mathbf{0} & \dots & \mathbf{Z}_n \end{pmatrix}, \\ \mathbf{D} &= \begin{pmatrix} \mathbf{D}_1 & \mathbf{0} & \dots & \mathbf{0} \\ \mathbf{0} & \mathbf{D}_2 & \dots & \mathbf{0} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{0} & \mathbf{0} & \dots & \mathbf{D}_n \end{pmatrix}, & \boldsymbol{\Sigma} &= \begin{pmatrix} \boldsymbol{\Sigma}_1 & \mathbf{0} & \dots & \mathbf{0} \\ \mathbf{0} & \boldsymbol{\Sigma}_2 & \dots & \mathbf{0} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{0} & \mathbf{0} & \dots & \boldsymbol{\Sigma}_n \end{pmatrix}, \end{aligned}$$

که $N = \sum_{i=1}^n n_i$ در این جا بعدهاى کمیت‌ها به صورت

$$\begin{aligned} \mathbf{y} &\sim N \times 1, & \mathbf{X} &\sim N \times p, & \boldsymbol{\beta} &\sim p \times 1, & \mathbf{Z} &\sim N \times nq, \\ \mathbf{b} &\sim nq \times 1, & \boldsymbol{\varepsilon} &\sim N \times 1, & \mathbf{D} &\sim nq \times nq, & \boldsymbol{\Sigma} &\sim N \times N \end{aligned}$$

هستند. از آنجایی که در داده‌های طولی مشاهدات مکرر از هر آزمودنی دارای همبستگی درونی بوده و از مشاهدات سایر آزمودنی‌ها مستقل‌اند، همبستگی بین پاسخ‌ها به صورت

$$\text{Cov}(\mathbf{y}) = \begin{pmatrix} \text{Cov}(\mathbf{y}_1) & \mathbf{0} & \dots & \mathbf{0} \\ \mathbf{0} & \text{Cov}(\mathbf{y}_2) & \dots & \mathbf{0} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{0} & \mathbf{0} & \dots & \text{Cov}(\mathbf{y}_n) \end{pmatrix}_{N \times N}$$

قابل نمایش است به‌طوری‌که

$$\mathbf{V}_i = \text{Cov}(\mathbf{y}_i) = \mathbf{Z}_i^\top \mathbf{D}_i \mathbf{Z}_i + \boldsymbol{\Sigma}_i.$$

واضح است که همبستگی بین پاسخ‌ها متأثر از تغییرات درون-گروهی $(\mathbf{Z}_i^\top \mathbf{D}_i \mathbf{Z}_i)$ و برون-گروهی $(\boldsymbol{\Sigma}_i)$ است. در این مدل حتی اگر ε_{ij} ‌ها مستقل از هم فرض شوند؛ یعنی $\boldsymbol{\Sigma}_i = \mathbf{I}_i$ ، با این وجود همبستگی بین پاسخ‌ها به دلیل دخالت ضرایب تصادفی $\mathbf{b}_i$ ، قابل بیان است. با اندکی تغییر در ساختار مدل و با فرض معلوم بودن مولفه‌های $\boldsymbol{\Sigma}_i$ و $\mathbf{D}_i$ ، می‌توان پارامترهای ثابت مدل را با روش‌های ساده رگرسیونی برآورد کرد. اگر مدل را به صورت

$$\mathbf{y} = \mathbf{X}^\top \boldsymbol{\beta} + \boldsymbol{\varepsilon}^*, \quad \boldsymbol{\varepsilon}^* = \mathbf{Z}^\top \mathbf{b} + \boldsymbol{\varepsilon}, \quad \boldsymbol{\varepsilon}^* \sim \mathcal{N}_N(\mathbf{0}, \mathbf{V})$$

تغییر دهیم، به‌طوری‌که

$$V = \begin{pmatrix} V_1 & 0 & \dots & 0 \\ 0 & V_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & V_n \end{pmatrix}_{N \times N}$$

آن‌گاه، برآورد پارامترهای ثابت به روش ML یا GLS به‌صورت

$$\hat{\beta} = (X^T V^{-1} X)^{-1} X^T V^{-1} y$$

حاصل می‌شود. برای برآورد پارامترهای تصادفی مدل، می‌توان از بهترین برآوردگر خطی نااریب^{۳۹} ($BLUE$) استفاده کرد. این برآوردگر با محاسبه امید شرطی b_i به شرط y_i و $\hat{\beta}$ به‌صورت

$$\hat{b}_i = E[b_i | y_i] = D Z_i V_i^{-1} (y_i - X_i^T \hat{\beta})$$

حاصل می‌شود. برآوردگرهای حاصل با فرض معلوم بودن مولفه‌های Σ_i و D_i به‌دست آمدند؛ لذا این برآوردگرها در عمل ناکارآمد هستند. به منظور غلبه بر این مسئله استفاده از الگوریتم EM پیشنهاد می‌شود که در بخش‌های بعدی به تفصیل بحث خواهد شد.

تاکنون مطالب بیان شده مربوط به متغیرهای پاسخ با توزیع نرمال بود. برای داده‌هایی که توزیع غیرنرمال دارند از مدل اثرات آمیخته تعمیم‌یافته^{۴۰} استفاده می‌شود. به عنوان مثال پیش‌بینی تعداد خرابی که کمیتی گسسته است، یا زمان انتظار که کمیتی مثبت است را می‌توان به کمک مدل اثرات آمیخته تعمیم‌یافته انجام داد. این مدل از سه جز زیر تشکیل شده است.

(آ) میانگین متغیرهای پاسخ به شرط ضرایب رگرسیونی به‌صورت

$$g[E(y_{ij} | b_i)] = g(\mu_{ij}) = x_{ij}^T \beta + z_{ij}^T b_i$$

است.

(ب) متغیرهای پاسخ $y_{i1}, y_{i2}, \dots, y_{in_i}$ به شرط b_i ، دو به دو از هم مستقل بوده و دارای توزیعی از خانواده توزیع‌های نمایی به‌صورت

$$f(y_{ij} | b_i) = \exp[\phi^{-1}(y_{ij} \theta_{ij} - \psi(\theta_{ij})) + c(y_{ij}, \phi)],$$

هستند، که در آن ϕ پارامتر مقیاس است.

(پ) اگر $g(\cdot)$ تابع پیوند کانونی باشد، آن‌گاه گشتاورهای شرطی $\mu_{ij} = E(y_{ij} | b_i) = \psi'(\theta_{ij})$ و $\nu_{ij} = \text{Var}(y_{ij} | b_i) = \psi''(\theta_{ij})$ در روابط

$$\theta_{ij} = g(\mu_{ij}) = x_{ij}^T \beta + z_{ij}^T b_i, \quad \nu_{ij} = \nu(\mu_{ij}) \phi$$

^{۳۹}Best linear unbiased estimator

^{۴۰}Generalized mixed effects model

صدق می‌کنند.

بنابراین پاسخ‌های هر آزمودنی به شرط ضرایب تصادفی دوبه‌دو از هم مستقل و دارای توزیع‌های نمایی هستند. معمولاً فرض می‌شود بردار ضرایب تصادفی b_i دارای توزیع نرمال با میانگین صفر و ماتریس واریانس-کوواریانس D_i است. با این شرایط محاسبه توزیع توام پاسخ‌ها و استفاده از روش ML برای برآورد ضرایب ثابت β امکان‌پذیر است. بدین ترتیب اگر تابع چگالی y_{ij} به شرط b_i را با نماد $f_{y_{ij}}(y_{ij}|b_i)$ نمایش دهیم، توزیع حاشیه‌ای y_i به صورت

$$f_{y_i}(y_i) = \int \prod_{j=1}^{n_i} f_{y_{ij}}(y_{ij}|b_i) f(b_i) db_i$$

محاسبه می‌شود. با فرض استقلال آزمودنی‌ها، تابع درستنمایی و لگاریتم آن به ترتیب، به صورت

$$L(\beta, \phi) = \prod_{i=1}^n f_{y_i}(y_i) = \prod_{i=1}^n \int \prod_{j=1}^{n_i} f_{y_{ij}}(y_{ij}|b_i) f(b_i) db_i$$

$$l(\beta, \phi) = \sum_{i=1}^n \log\{f_{y_i}(y_i)\} = \sum_{i=1}^n \log\left\{ \int \prod_{j=1}^{n_i} f_{y_{ij}}(y_{ij}|b_i) f(b_i) db_i \right\}$$

است. با حل رابطه $\frac{\partial l(\beta, \phi)}{\partial \beta} = 0$ ، می‌توان برآورد ML پارامترهای مدل را به دست آورد. صورت بسته برآوردگرها به توزیع $f_{y_{ij}}(y_{ij}|b_i)$ وابسته است. این توزیع متناسب با نوع داده‌ها انتخاب می‌شود، به طوری که معمولاً در داده‌های پیوسته، دودویی و شمارشی به ترتیب توزیع‌های نرمال، دوجمله‌ای و پواسن و همچنین برای داده‌های اسمی و ترتیبی توزیع چندجمله‌ای در نظر گرفته می‌شود. معمولاً برآورد و تحلیل درباره b_i ‌ها از طریق روش‌های بیزی انجام می‌گیرد. همچنین به جای روش برآورد ML از روش‌های برآورد ماکزیمم درستنمایی مقید، ماکزیمم درستنمایی تاوانیده، شبه درستنمایی و بسیاری از روش‌های دیگر استفاده می‌شود. با توجه به این که بیان تمام موارد و شرح جزئیات بیش‌تر منجر به دور شدن از بحث این رساله خواهد شد، لذا به مطالب اندک و مختصر شرح داده شده بسنده می‌کنیم. برای آشنایی بیش‌تر با این روش‌ها به تحقیقات استیراتلی و همکاران (۱۹۸۴)، هدیگر و گیبونز (۱۹۹۴)، لی و نلدز (۲۰۰۱)، توتز (۲۰۰۵) و لئون و همکاران (۲۰۰۷) رجوع کنید.

۳.۱.۳.۱ مدل‌های انتقالی

سومین مدل حاصل از تعمیم مدل‌های خطی تعمیم یافته، مدل‌های انتقالی هستند. مدل‌های انتقالی مدل‌هایی هستند که در آن‌ها متغیر پاسخ در زمان j به متغیرهای پاسخ در زمان‌های $j-1, \dots, j-q$ ($q < j-1$) وابسته است. تحت این مدل، همبستگی بین پاسخ‌ها توسط تاثیر مستقیم مقادیر گذشته ($y_{ij-1}, y_{ij-2}, \dots, y_{ij-q}$) روی y_{ij} مورد بررسی قرار می‌گیرد و پاسخ‌های گذشته به عنوان متغیرهای تبیینی در مدل در نظر گرفته می‌شوند. این مدل‌ها که به عنوان مدل‌های اتورگرسیو و مارکوف نیز شناخته می‌شوند، ویژگی‌های عمومی به صورت زیر دارند:

آ) میانگین پاسخ‌ها به شرط پاسخ‌های گذشته به صورت

$$g[E(y_{ij}|\mathbf{H}_{ij})] = g[\mu_{ij}] = \mathbf{x}_{ij}^\top \boldsymbol{\beta} + \sum_{r=1}^s \kappa_r(\mathbf{H}_{ij}, \boldsymbol{\alpha})$$

مدل‌بندی می‌شوند، که در آن، $\mathbf{H}_{ij} = (y_{i1}, y_{i2}, \dots, y_{i(j-1)})^\top$ برداری از مقادیر گذشته y_{ij} ، $\kappa(\cdot)$ تابعی معلوم و $\boldsymbol{\alpha}$ پارامترهای اضافی است. همچنین، g تابع پیوند است.

ب) متغیرهای پاسخ $y_{i1}, y_{i2}, \dots, y_{in_i}$ به شرط $\mathbf{H}_{ij}$ ، دوبه‌دو از هم مستقل بوده و دارای توزیع‌های نمایی به صورت

$$f(y_{ij}|\mathbf{H}_{ij}) = \exp[\phi^{-1}(y_{ij}\theta_{ij} - \psi(\theta_{ij})) + c(y_{ij}, \phi)]$$

است.

پ) اگر $g(\cdot)$ تابع پیوند کانونی باشد، آنگاه گشتاورهای شرطی $\mu_{ij} = E(y_{ij}|\mathbf{H}_{ij}) = \psi'(\theta_{ij})$ و $\nu_{ij} = \text{Var}(y_{ij}|\mathbf{H}_{ij}) = \psi''(\theta_{ij})$ در روابط

$$\theta_{ij} = g(\mu_{ij}) = \mathbf{x}_{ij}^\top \boldsymbol{\beta} + \sum_{r=1}^s \kappa_r(\mathbf{H}_{ij}, \boldsymbol{\alpha}), \quad \nu_{ij} = \nu(\mu_{ij})\phi$$

صدق می‌کنند.

یکی از معمول‌ترین مدل‌های انتقالی، مدل اتورگرسیو مرتبه q است. به عبارت دیگر، $(y_{ij}|\mathbf{H}_{ij})$ به q مشاهده قبلی خود یعنی $y_{ij-1}, y_{ij-2}, \dots, y_{ij-q}$ وابسته است. به عنوان مثال برای پاسخ شمارشی، مدل به صورت

$$\log[E(y_{ij}|\mathbf{H}_{ij})] = \mathbf{x}_{ij}^\top \boldsymbol{\beta} + \sum_{r=1}^q \alpha_r y_{ij-r}, \quad i = 1, \dots, n, \quad j = 1, \dots, n_i$$

بیان می‌گردد. با توجه به این‌که طبق ویژگی دوم مدل‌های انتقالی، $(y_{ij}|\mathbf{H}_{ij})$ ‌ها دوبه‌دو از هم مستقل‌اند، پس از مشخص نمودن تابع $\kappa(\cdot)$ ، می‌توان توزیع توأم پاسخ‌های مربوط به هر آزمودنی را محاسبه نمود. به عنوان مثال برای مدل اتورگرسیو مرتبه q ، توزیع y_i به صورت

$$f_{\mathbf{y}_i}(\mathbf{y}_i) = f(y_{i1}, y_{i2}, \dots, y_{in_i}) \prod_{j=1}^{n_i} f(y_{ij}|y_{ij-1}, y_{ij-2}, \dots, y_{ij-q})$$

ارائه می‌گردد. سپس با فرض استقلال بین آزمودنی‌ها، تابع درستنمایی به صورت

$$\begin{aligned} L(\boldsymbol{\beta}, \boldsymbol{\alpha}) &= \prod_{i=1}^n f_{\mathbf{y}_i}(\mathbf{y}_i) \\ &= \prod_{i=1}^n \left[f(y_{i1}, y_{i2}, \dots, y_{in_i}) \prod_{j=1}^{n_i} f(y_{ij}|y_{ij-1}, y_{ij-2}, \dots, y_{ij-q}) \right] \end{aligned}$$

است. محاسبه توزیع توام q مشاهده اول، یعنی $f(y_{i1}, y_{i2}, \dots, y_{iq})$ نیازمند اعمال پذیره‌های اضافی و پیچیده در مدل است. در چنین شرایطی استفاده از روش ML شرطی پیشنهاد می‌گردد. در این روش تابع درستنمایی به صورت

$$\begin{aligned} L^c(\beta, \alpha) &= \prod_{i=1}^n f(y_{iq+1}, y_{iq+2}, \dots, y_{in_i}) \\ &= \prod_{i=1}^n \prod_{j=q+1}^{n_i} f(y_{ij} | y_{i1}, y_{i2}, \dots, y_{iq}) \end{aligned}$$

محاسبه می‌شود. با شرطی کردن روی q مقدار اولیه از اطلاعات آن‌ها صرف نظر می‌کند. بنابراین کارایی کمتری نسبت به روش درستنمایی عادی خواهد داشت. با این وجود در صورتی که نخواهیم هیچ پذیره اضافی را در مدل اعمال کنیم، این روش مناسب است. علاقه‌مندان جهت آشنایی بیشتر با مباحث مدل‌های انتقالی و برآورد پارامترها، می‌توانند به تحقیقات رضایی و گنجعلی (۱۳۸۸)، مولنبرگ و وریک (۲۰۰۵)، چانگ و همکاران (۲۰۰۵)، سنگال و همکاران (۲۰۰۷)، رضایی و گنجعلی (۲۰۰۹a) و رضایی و گنجعلی (۲۰۰۹b) مراجعه کنند.

۲.۳.۱ مدل‌های ناپارامتری

روش‌های رگرسیون ناپارامتری برای داده‌های مستقل در سه دهه گذشته به خوبی توسعه یافته است که در این زمینه می‌توان به کارهای گرین و سیلورمن (۱۹۹۴) و هاردل و همکاران (۲۰۰۰) اشاره نمود. یک مدل ناپارامتری را می‌توان به صورت

$$E(y|X = x) = f(x)$$

بیان نمود، که در آن $f(\cdot)$ یک تابع هموار نامعلوم است و باید بر اساس یک نمونه از مشاهدات (x_i, y_i) برآورد شود. در این جا منظور از هموار بودن، مشتق‌پذیری تابع $f(\cdot)$ تا یک مرتبه خاص و کوچک است. تقریب این تابع رگرسیونی هموارسازی^{۴۱} نامیده می‌شود. بنابراین یک هموارساز تابعی برای خلاصه‌سازی متغیر پاسخ y_i به عنوان تابعی از متغیر تبیینی x_i است. تاکنون روش‌های متعددی برای هموارسازی معرفی شده‌اند که عبارت‌اند از رگرسیون اسپلاین^{۴۲}، هموارسازی اسپلاین^{۴۳}، اسپلاین توانیده^{۴۴} و هموارسازی چند جمله‌ای موضعی هسته‌ای^{۴۵}. معرفی تمامی این روش‌ها به تنهایی نیازمند رونمایی بحث جدید و گسترده‌ای است که بیان آن‌ها خارج از گنجایش این رساله بوده و منجر به دور شدن از هدف اصلی خواهد شد. لذا در ادامه تنها به مرور اجمالی چند روش متداول خواهیم پرداخت.

⁴¹Smoothing

⁴²Spline regression

⁴³Spline smoothing

⁴⁴Penalized spline

⁴⁵Local polynomial kernel smoothing

یک روش رایج در رگرسیون ناپارامتری میانگین وزنی موضعی است. فرض کنید x_0 مقدار دلخواهی است که x می‌تواند اختیار کند. در این روش برای برآورد تابع $f(\cdot)$ در نقطه x_0 ، به مقادیری از y_i که متناظر با مقادیر x_i نزدیکتر به نقطه x_0 باشند، وزن بیش‌تر و به مقادیر y_i متناظر با نقاط دورتر وزن کمتری نسبت می‌دهیم. وزن‌ها در رگرسیون ناپارامتری می‌توانند بوسیله یک تابع مقارن تک‌مدی حول صفر که مقادیر آن در دو طرف صفر توسط یک پارامتر مقیاس کنترل می‌شود، انتخاب شوند. نخستین نامزدهای آن‌ها برای چنین توابعی که هسته^{۴۶} نامیده شدند، توابع چگالی احتمال بودند. سپس با استفاده از توابع کراندار هسته، برآوردگر

$$\hat{f}(x_0) = \frac{1}{n} \sum_{i=1}^n W_i y_i, \quad W_i = \frac{K_h(x_0, x_i)}{\sum_{i=1}^n K_h(x_0, x_i)},$$

را پیشنهاد دادند، که در آن $K_h(\cdot)$ یک تابع هسته بوده و پارامتر h پهنای باند است. در این برآوردگر که به هموارساز هسته‌ای معروف است، شکل وزن‌ها توسط تابع هسته و مقادیر آن‌ها بر اساس فاصله از نقطه x_0 و پهنای باند h تعیین می‌شود. بزرگ بودن پهنای باند باعث می‌شود که به نقاط نزدیک‌تر وزن‌های بزرگتری داده شود. معمول‌ترین توابع هسته عبارت‌اند از هسته یکنواخت، هسته نرمال و هسته اپانچنیکوف.

یکی دیگر از متداول‌ترین روش‌های هموارسازی، اسپلاین‌ها هستند. اسپلاین توسط شخصی به نام شوئنبرگ در سال ۱۹۴۶ معرفی شد. وی برای نخستین بار قضیه‌ای برای اثبات وجود اسپلاین‌های درون‌یاب بیان کرد و سپس مطالعه چند جمله‌ای‌های قطعه به قطعه برای تابع اسپلاین آغاز شد. فاصله $[a, b]$ به صورت اجتماع دو مجموعه $[a, t] \cup [t, b]$ را در نظر بگیرید. یک اسپلاین از مرتبه یک تابعی است که در هر یک از بازه‌های مذکور یک تابع خطی و در نقطه t نیز پیوسته است. این تابع را می‌توان برای تابع پیوسته $f(\cdot)$ به صورت

$$f(x) = \begin{cases} \beta_0 + \beta_1 x & x \leq t \\ \beta'_0 + \beta'_1 x & x > t \end{cases}$$

نوشت. با اعمال شرط پیوستگی تابع f در نقطه $x = t$ داریم

$$\begin{aligned} \beta_0 + \beta_1 t &= \beta'_0 + \beta'_1 t \\ \beta'_0 &= \beta_0 + \beta_1 t - \beta'_1 t. \end{aligned}$$

با جای‌گذاری β'_0 در ضابطه دوم f داریم

$$\beta'_0 + \beta'_1 x = \beta_0 + \beta_1 t - \beta'_1 t + \beta'_1 x = \beta_0 + \beta_1 t + \beta'_1 (x - t).$$

لذا

$$f(x) = \begin{cases} \beta_0 + \beta_1 x & x \leq t \\ \beta_0 + \beta_1 t + \beta'_1 (x - t) & x > t \end{cases},$$

که می‌توان این تابع دو ضابطه‌ای را به صورت

$$f(x) = \beta_0 + \beta_1 x + \beta_2 (x - t)_+$$

نوشت، که در آن $(x - t)_+ = \max(x - t, 0) = (x - t)I(x > t)$ ، تابع نشانگر و $\beta_2 = \beta'_1 - \beta_1$ تفاوت بین شیب خطوط اول و دوم تابع دو ضابطه‌ای را بیان می‌کند. درواقع می‌توان گفت که تابع f ، ترکیبی خطی از توابع پایه^{۴۷}، ۱، x و $(x - t)_+$ است. اسپلاین مرتبه اول را اسپلاین خطی نیز می‌گویند.

حال فرض کنید که بازه‌ی $[a, b]$ به $K + 1$ زیر بازه به صورت

$$[a, t_1), [t_1, t_2), \dots, [t_K, b]$$

افراز شده باشد، به طوری که $a < t_1 < \dots < t_K < b$. نقاط $\{t_i\}_{i=1}^K$ را گره^{۴۸} می‌نامند. در این صورت اسپلاین مرتبه اول را می‌توان به صورت

$$f(x) = \beta_0 + \beta_1 x + \sum_{k=1}^K \beta_{1+k} (x - t_k)_+$$

بیان کرد که ترکیب خطی از توابع پایه ۱، x ، $(x - t_1)_+$ ، $\dots$ و $(x - t_K)_+$ است. توابع اسپلاین برای مراتب بالاتر به طور مشابه تعریف می‌شود. تابع $f(\cdot)$ را یک اسپلاین مرتبه d نامیم اگر در هر یک از زیر بازه‌های بیان شده یک چندجمله‌ای از مرتبه d باشد به طوری که $f(\cdot)$ و مشتق‌های آن از مراتب ۱ تا $d - 1$ ، در گره‌ها پیوسته است. لذا تابع اسپلاین مرتبه d را می‌توان به صورت

$$f(x) = \beta_0 + \beta_1 x + \beta_2 x^2 + \dots + \beta_d x^d + \sum_{k=1}^K \beta_{d+k} (x - t_k)_+^d$$

بیان نمود، که در آن $(x - t_k)_+^d = (x - t_k)^d I(x > t_k)$ تابع توانی بریده شده^{۴۹} از مرتبه d است. یعنی این تابع ترکیب خطی از یک چندجمله‌ای مرتبه d و توابع توانی بریده شده است. به عبارت دیگر مجموعه

$$1, x, x^2, \dots, x^d, (x - t_1)_+^d, \dots, (x - t_K)_+^d$$

توابع پایه برای فضای اسپلاین‌های مرتبه d با گره‌های $t_1, t_2, \dots, t_K$ و t_K است. اگر $d = 3$ ، آن‌گاه تابع مورد اشاره را یک اسپلاین مکعبی^{۵۰} می‌نامند. اسپلاین مکعبی یکی از متداول‌ترین نوع‌های اسپلاین است که به دلیل انعطاف‌پذیری بالا، در بسیاری از علوم مورد استفاده قرار می‌گیرد.

⁴⁷Basis functions

⁴⁸Knot

⁴⁹Truncated power function

⁵⁰Cubic spline

مدل رگرسیون ناپارامتری

$$y = f(x) + \varepsilon$$

را در نظر بگیرید، که در آن

$$y = (y_1, \dots, y_n)^\top$$

$$f(x) = (f(x_1), \dots, f(x_n))^\top$$

$$\varepsilon = (\varepsilon_1, \dots, \varepsilon_n)^\top.$$

اگر تابع $f(\cdot)$ را با تابع اسپلاین تعریف شده جایگزین کنیم، نتیجه حاصل را رگرسیون اسپلاین نامند، به‌طوری‌که منجر به برازش مدلی به‌صورت

$$y_i = \beta_0 + \beta_1 x_i + \beta_2 x_i^2 + \dots + \beta_d x_i^d + \sum_{k=1}^K \beta_{d+k} (x_i - t_k)_+^d + \varepsilon_i, \quad i = 1, \dots, n$$

می‌شود که در آن ε_i ها در پذیره‌های زیربنایی رگرسیون صدق می‌کنند. با استفاده از نماد ماتریسی، رگرسیون اسپلاین به‌صورت

$$y = B\beta + \varepsilon$$

نمایش داده می‌شود که در آن $\beta = (\beta_0, \beta_1, \dots, \beta_{d+K})^\top$ و B یک ماتریس طرح به‌صورت

$$B = \begin{pmatrix} 1 & x_1 & x_1^2 & \dots & x_1^d & (x_1 - t_1)_+^d & (x_1 - t_2)_+^d & \dots & (x_1 - t_K)_+^d \\ 1 & x_2 & x_2^2 & \dots & x_2^d & (x_2 - t_1)_+^d & (x_2 - t_2)_+^d & \dots & (x_2 - t_K)_+^d \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 1 & x_n & x_n^2 & \dots & x_n^d & (x_n - t_1)_+^d & (x_n - t_2)_+^d & \dots & (x_n - t_K)_+^d \end{pmatrix}_{n \times (d+K+1)}$$

است. واضح است که درایه‌های ستون اول آن برابر ۱، درایه‌های d ستون بعدی مربوط به مشاهدات متغیرهای x ، x^2 ، $\dots$ و x^d و درایه‌های ستون‌های 2 تا $d+K+1$ مربوط به مشاهدات حاصل از متغیرهای $(x - t_1)_+^d$ ، $\dots$ و $(x - t_K)_+^d$ به ازای گره‌های از قبل تعیین شده t_1 ، t_2 ، $\dots$ و t_K هستند. لازم به ذکر است که رگرسیون اسپلاین متعلق به رده رگرسیون خطی است. به عبارت دیگر، تابع $f(\cdot)$ نسبت به پارامترها خطی است و از این رو هزینه‌های محاسباتی برای برآورد منحنی رگرسیون مربوطه ناچیز است. بنابراین برازش یک رگرسیون اسپلاین روی داده‌ها را می‌توان با استفاده از روش‌های مرسوم انجام داد. با به‌کاربردن روش کمترین توان‌های دوم خطا برآوردگر

$$\hat{\beta} = (B^\top B)^{-1} B^\top y$$

حاصل می‌شود. مشکلی که در این برآوردگر وجود دارد هم‌خطی بین پایه‌های توانی بریده‌شده

$$(x - t_1)_+^d, (x - t_2)_+^d, \dots, (x - t_K)_+^d$$

است که باعث عدم وارون‌پذیری یا بدشرطی^{۵۱} ماتریس $B^T B$ خواهد شد که ناپایداری ضرایب رگرسیونی را به همراه دارد. یک روش مناسب برای حل مشکل فوق این است که به جای پایه‌های توانی بریده‌شده از پایه‌های B -اسپلاین استفاده شود. بازه $[a, b]$ و تعداد K گره به صورت

$$a = t_0 < t_1 < \dots < t_K < t_{K+1} = b$$

را در نظر بگیرید. متداول‌ترین روش محاسبه پایه‌های B -اسپلاین مرتبه d ، به وسیله پایه‌های B -اسپلاین مرتبه $d-1$ و براساس رابطه بازگشتی

$$B_{i,d}(x) = \left(\frac{x - t_i}{t_{i+d} - t_i}\right) B_{i,d-1}(x) + \left(\frac{t_{i+d+1} - x}{t_{i+d+1} - t_{i+1}}\right) B_{i+1,d-1}(x), \quad i = -d, -d+1, \dots, K$$

است، که $t_{-d} = t_{-d+1} = \dots = t_{-2} = t_{-1} = t_0 = a$ و $t_{K+1} = b$ بنابراین تعداد پایه‌های B -اسپلاین مرتبه d برابر $d + K + 1$ است. با استفاده از این پایه‌ها می‌توان چندجمله‌ای‌های متعامد به دست آورد. توجه کنید که در این رساله به منظور سهولت در نوشتار از نمادگذاری توابع توانی بریده‌شده استفاده می‌شود اما در محاسبات مربوط به مطالعات شبیه‌سازی و تحلیل داده‌های واقعی از پایه‌های B -اسپلاین استفاده شده است. برای اطلاعات بیشتر درباره این توابع و انواع آن می‌توان به راپرت (۲۰۰۳) مراجعه کرد.

در تحلیل داده‌های طولی به روش ناپارامتری، معمولاً مشاهدات مکرر به صورت تابعی نامعلوم و یک ماهیت واحد و نه به صورت دنباله‌ای از مشاهدات گسسته ارزیابی می‌شوند. یعنی تغییرات پاسخ‌ها در طول زمان را نامعلوم در نظر می‌گیرند، بنابراین ساده‌ترین مدل رگرسیون ناپارامتری برای این داده‌ها به صورت

$$y_i = f(t_i) + \varepsilon_i, \quad i = 1, \dots, n \quad (1.1)$$

بیان می‌شود. با وجود پیشرفت چشمگیر روش‌های رگرسیون ناپارامتری برای داده‌های مستقل، تحولات مهم در این زمینه برای داده‌های طولی تنها در سال‌های اخیر رخ داده است. ویژگی‌های خاص داده‌های طولی، چالش‌های عمده‌ای را در توسعه روش‌های ناپارامتری برای این داده‌ها به وجود آورده است. بنابراین لازم است روش‌های مرسوم جهت تحلیل داده‌های طولی توسعه و تعمیم یابند.

۱.۲.۳.۱ هموارسازی در داده‌های طولی با استفاده از روش هسته

ایده و پایه اصلی این روش بسط تیلور است، که در آن هر تابع هموار را می‌توان به کمک چند جمله‌ای‌هایی از درجه مشخص تقریب زد. حال با استفاده از بسط تیلور حول نقطه زمانی دلخواه t_0 ، تابع $f(\cdot)$ در مدل (۱.۱) را می‌توان به صورت

$$f(t_{ij}) \approx f(t_0) + (t_{ij} - t_0)f^{(1)}(t_0) + \dots + (t_{ij} - t_0)^d \frac{f^{(d)}(t_0)}{d!} = \mathbf{T}^T(t_{ij})\boldsymbol{\beta}$$

⁵¹Ill-condition

تقریب زد، که در آن $\beta_r = \frac{f^{(r)}(t_0)}{r!}$ ، $\beta = (\beta_0, \dots, \beta_d)^\top$ ، $T^\top(t_{ij}) = (1, (t_{ij} - t_0), \dots, (t_{ij} - t_0)^d)^\top$ و $f^{(r)}(t_0)$ مشتق r ام تابع f در نقطه t_0 است. جهت برآورد بردار پارامتر β ، معیار کمترین توان‌های دوم موزون^{۵۲} (WLS) را به صورت

$$\sum_{i=1}^n K_h(t_i - t) \{y_i - T(t_i)\beta\}^2, \quad (2.1)$$

در نظر بگیرید، به طوری که $K_h(\cdot)$ ، $K_h(t_i - t) = (K_h(t_{i1} - t), \dots, K_h(t_{in_i} - t))$ و $T(t_i) = (T^\top(t_{i1}), \dots, T^\top(t_{in_i}))^\top$ با مشتق گیری از رابطه (۲.۱)، برآورد β با حل معادلات برآورد هسته‌ای زیر بدست می‌آید.

$$\sum_{i=1}^n T^\top(t_i) K_h(t_i - t) \{y_i - T^\top(t_i)\beta\} = 0 \quad (3.1)$$

به ازای $d = 1$ ، برآوردگر به صورت $\hat{f}(t_{ij}) = \frac{\sum_{i=1}^n \sum_{j=1}^{n_i} K_h(t_{ij} - t_0) y_{ij}}{\sum_{i=1}^n \sum_{j=1}^{n_i} K_h(t_{ij} - t_0)}$ حاصل می‌شود که همان برآوردگر هسته‌ای یا برآوردگر نادارایا-واتسون است.

به دلیل وجود همبستگی درون گروهی، برآوردهای حاصل از حل معادلات (۳.۱) برای داده‌های طولی ناکارآمد هستند. لین و کارول (۲۰۰۲) با تعمیم نسخه تعمیم یافته این روش به داده‌های غیرنرمال و با اعمال ساختار همبستگی، برآوردگر دیگری را به نام GEE چندجمله‌ای موضعی هسته‌ای^{۵۳} (LPK) ارائه نمودند. بر این اساس یک مدل رگرسیون ناپارامتری تعمیم یافته با سه مشخصه اصلی

- آ. $g[E(y_{ij}|t_{ij})] = g(\mu_{ij}(t_{ij})) = f(t_{ij})$,
- ب. $\text{Var}(y_{ij}|t_{ij}) = \nu(\mu_{ij}(t_{ij}))\phi$,
- پ. $\text{corr}(y_{ij}, y_{ik}) = \rho(\mu_{ij}(t_{ij}), \mu_{ik}(t_{ik}), \alpha)$,

را در نظر بگیرید، که در آن $f(t_{ij}) = T^\top(t_{ij})\beta$ و $T^\top(t_{ij}) = (1, (t_{ij} - t_0), \dots, (t_{ij} - t_0)^d)^\top$ معادلات برآوردیاب تعمیم یافته به صورت

$$\sum_{i=1}^n \frac{\partial \mu_i(t_i)}{\partial \beta} K_h^\top(t_i - t) V_i^{-1}(\alpha) K_h(t_i - t) (y_i - \mu_i(t_i)) = 0 \quad (4.1)$$

ارائه می‌گردد، که در آن

$$\begin{aligned} \mu_i(t_i) &= (\mu_{i1}(t_{i1}), \dots, \mu_{in_i}(t_{in_i}))^\top, \\ \mu_{ij}(t_{ij}) &= g^{-1}(T^\top(t_{ij})\beta), \\ V_i(\alpha) &= \phi A_i^\top R_i(\alpha) A_i, \\ A_i &= \text{diag}(\nu(\mu_{i1}(t_{i1})), \dots, \nu(\mu_{in_i}(t_{in_i}))), \end{aligned}$$

⁵²Weighted least squares

⁵³Local polynomial kernel

$$\frac{\partial \mu_i(t_i)}{\partial \beta} = \begin{pmatrix} \frac{\partial g^{-1}(T^\top(t_{i1})\beta)}{\partial \beta} \\ \frac{\partial g^{-1}(T^\top(t_{i2})\beta)}{\partial \beta} \\ \vdots \\ \frac{\partial g^{-1}(T^\top(t_{in_i})\beta)}{\partial \beta} \end{pmatrix} = \begin{pmatrix} 1 & \frac{(t_{i1}-t_0)}{g'(\mu_{i1})} & \dots & \frac{(t_{i1}-t_0)^d}{g'(\mu_{i1})} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & \frac{(t_{in_i}-t_0)}{g'(\mu_{in_i})} & \dots & \frac{(t_{in_i}-t_0)^d}{g'(\mu_{in_i})} \end{pmatrix}_{n_i \times (d+1)}$$

$R_i(\alpha)$ ماتریس همبستگی عملی و α بردار پارامترهای همبستگی هستند که به تفصیل در بخش ۱۰.۱.۳.۱ معرفی شدند. جواب حاصل از حل معادلات (۴.۱) یعنی $\hat{\beta}$ از طریق الگوریتم تکراری بیان شده در همان بخش به دست آمده و برآوردگر هسته‌ای تابع ناپارامتری به صورت $\hat{f}(t_{ij}) = T_{ij}^\top(t_0)\hat{\beta}$ حاصل می‌شود.

۲.۲.۳.۱ هموارسازی در داده‌های طولی با رگرسیون اسپلاین

همان‌طور که پیش از این نیز بیان شد، رگرسیون اسپلاین با استفاده از تعدادی گره یک رگرسیون پارامتری به کمک توابع پایه تشکیل می‌دهد. رایس و وو (۲۰۰۱) و هوانگ و ژو (۲۰۰۴) این روش را به داده‌های طولی تعمیم دادند. به این صورت که براساس اسپلاین مرتبه d و توابع توانی بریده‌شده، مجموعه توابع پایه در نقاط t_{ij} را به صورت $\{1, t_{ij}, t_{ij}^2, \dots, t_{ij}^d, (t_{ij}-t_1)_+^d, \dots, (t_{ij}-t_K)_+^d\}$ در نظر بگیرید که در مجموعه نقاط گره $\{t_1, \dots, t_K\}$ ایجاد شده‌اند. مشابه آن‌چه در ابتدای این بخش شرح داده‌شد، تابع $f(t_{ij})$ را می‌توان به صورت

$$f(t_{ij}) \approx \beta_0 + \beta_1 t_{ij} + \dots + \beta_d t_{ij}^d + \sum_{k=1}^K (t_{ij} - t_k)_+^d \beta_{d+k} = B^\top(t_{ij})\beta \quad (5.1)$$

تقریب زد، که $B(t_{ij}) = (1, t_{ij}, \dots, t_{ij}^d, (t_{ij}-t_1)_+^d, \dots, (t_{ij}-t_K)_+^d)^\top$ و $\beta = (\beta_0, \beta_1, \dots, \beta_{d+K})^\top$ در این حالت، مدل (۱.۱) را می‌توان به صورت

$$y_i = B(t_i)\beta + \varepsilon_i$$

بازنویسی کرد، به‌طوری‌که $B(t_i) = (B^\top(t_{i1}), \dots, B^\top(t_{in_i}))^\top$ ماتریس طرح $n_i \times (d+K+1)$ بعدی است. با در نظر گرفتن ماتریس کوواریانس V_i و استفاده از روش WLS ، برآورد β با کمینه کردن رابطه

$$\sum_{i=1}^n \{y_i - B(t_i)\beta\}^\top V_i^{-1} \{y_i - B(t_i)\beta\},$$

به صورت

$$\hat{\beta} = (B^\top(t)V^{-1}B(t))^{-1}(B^\top(t)V^{-1}y)$$

حاصل می‌شود، به‌طوری‌که $B(t) = (B(t_1), \dots, B(t_n))^T$ و $V = \text{diag}(V_1, \dots, V_n)$. به دنبال برآورد پارامتر رگرسیونی، برآورد تابع $f(t)$ به روش رگرسیون اسپالین به صورت $\hat{f}(t) = B^T(t)\hat{\beta}$ بدست می‌آید.

۳.۳.۱ مدل‌های نیمه‌پارامتری

همان‌گونه که پیش از این بیان شد، مدل‌های پارامتری در برخورد با تعدد متغیرهای تبیینی، کارایی خوبی دارند اما تحت تاثیر فرضیات پارامتری هستند. رگرسیون ناپارامتری از انعطاف‌پذیری بالایی برخوردار است، اما افزایش بعد فضای متغیرهای تبیینی استفاده از آن را با مشکل مواجه می‌کند. استون (۱۹۸۲) نشان داد که اگر مشتق تابع $f(\cdot)$ تا مرتبه r موجود باشد، سریع‌ترین نرخ هم‌گرایی قابل حصول برآوردگر $\hat{f}(\cdot)$ به $f(\cdot)$ برابر $(\log n/n)^{r/(2r+p)}$ است. این نرخ هم‌گرایی نقش بعد متغیرهای تبیینی را در کارایی برآوردگر نشان می‌دهد. به عبارت دیگر، افزایش p منجر به کاهش عملکرد برآوردگر $\hat{f}(\cdot)$ می‌شود. این مسئله به مشتق بعد معروف است که به دلیل تنگ بودن داده‌ها در فضایی با ابعاد بالاتر اتفاق می‌افتد. برای اطلاعات بیشتر به گیننس (۲۰۱۱) مراجعه کنید.

برای اجتناب از ارببی زیاد مدل‌های پارامتری و مشتق بعد مدل‌های ناپارامتری، مدل‌های نیمه‌پارامتری به عنوان مدل‌هایی سازشگر و میانه‌رو که ویژگی‌های خوب هر دو مدل را دارد، پیشنهاد می‌شود. این مدل ساختار جمعی رگرسیون پارامتری را با انعطاف‌پذیری رگرسیون ناپارامتری ترکیب می‌کند. تاکنون مدل‌های رگرسیون نیمه‌پارامتری زیادی توسط افراد مختلف معرفی شده است، که رایج‌ترین آن‌ها شامل مدل خطی-جزئی^{۵۴}، مدل ضریب متغیر^{۵۵} و مدل تک‌شاخص^{۵۶} است. تمرکز این رساله روی مدل‌های خطی-جزئی است. این مدل تفسیرپذیری مدل‌های خطی و انعطاف‌پذیری مدل‌های ناپارامتری را با هم ترکیب می‌کند و یکی از رایج‌ترین مدل‌های نیمه‌پارامتری است. انگل و همکاران (۱۹۸۶) از نخستین افرادی بودند که از این مدل برای بررسی رابطه بین دما و مصرف برق استفاده کردند. هاردل و همکاران (۲۰۰۰) بررسی‌های دقیقی از این مدل ارائه داده‌اند.

متغیر پاسخ y_{ij} مربوط به واحد i ام در زمان t_{ij} را در نظر بگیرید. فرض کنید میانگین متغیر پاسخ به صورت خطی با متغیرهای تبیینی در ارتباط بوده و از طرف دیگر دارای روندی غیرخطی اما با صورت مشخص نسبت به متغیر زمان است. جهت هم‌ساز کردن مدل‌های رگرسیونی با روند غیرخطی بیان‌شده، دیگل و همکاران (۲۰۰۲) مدل نیمه‌پارامتری خطی-جزئی برای داده‌های طولی را به صورت

$$y_{ij} = \mathbf{x}_{ij}^T \beta + f(t_{ij}) + \varepsilon_{ij}, \quad i = 1, \dots, n; \quad j = 1, \dots, n_i \quad (6.1)$$

ارائه دادند، که در آن $\mathbf{x}_{ij} = (x_{1,ij}, x_{2,ij}, \dots, x_{p,ij})^T$ و $\beta = (\beta_1, \beta_2, \dots, \beta_p)^T$ بردارهایی p بعدی از متغیرهای تبیینی، $f(\cdot)$ تابعی هموار از زمان و ε_{ij} خطای تصادفی نرمال با $E\{\varepsilon_{ij}\} = 0$ است. در

⁵⁴Partially linear models

⁵⁵Varying coefficient model

⁵⁶Single index model

ادامه به مرور کلی و کوتاه بر روش‌های متداول برآورد پارامترهای β و تابع ناپارامتری $f(\cdot)$ خواهیم پرداخت.

۱.۳.۳.۱ برآورد به روش چند جمله‌ای موضعی

در این بخش به معرفی LPK جهت برازش مدل (۶.۱) خواهیم پرداخت. روش اول تعمیمی از اولین روش هسته‌ای اسپیکمن (۱۹۸۸) و بر پایه یک الگوریتم بازگشتی^{۵۷} است. روش دوم بر پایه هموارسازی نیمرخ^{۵۸} و تعمیمی از دومین روش هسته‌ای اسپیکمن (۱۹۸۸) است. هو و همکاران (۲۰۰۴) و وو و لیانگ (۲۰۰۴) این روش‌ها را در تحلیل داده‌های طولی به‌کار بردند و ما به تبعیت از آن‌ها، روش اول را LPK بازگشتی و روش دوم را LPK نیمرخ می‌نامیم. مدل (۶.۱) را در نظر بگیرید. مدل ماتریسی که در آن عناصر بر اساس آزمودنی i ام مشخص و بر اساس نقاط اندازه‌گیری ماتریسی شده‌اند و همچنین مدل ماتریسی کلی به‌ترتیب به‌صورت

$$y_i = X_i\beta + f(t_i) + \varepsilon_i,$$

$$y = X\beta + f(t) + \varepsilon,$$

نمایش داده می‌شوند، که در آن مولفه‌ها به‌صورت

$$\begin{cases} y_i = (y_{i1}, \dots, y_{in_i})^\top \\ X_i = (x_{i1}^\top, \dots, x_{in_i}^\top)^\top \\ f(t_i) = (f(t_{i1}), \dots, f(t_{in_i}))^\top \\ \varepsilon_i = (\varepsilon_{i1}, \dots, \varepsilon_{in_i})^\top \end{cases}, \quad \begin{cases} y = (y_1^\top, \dots, y_n^\top)^\top \\ X = (X_1^\top, \dots, X_n^\top)^\top \\ f(t) = (f^\top(t_1), \dots, f^\top(t_n))^\top \\ \varepsilon = (\varepsilon_1^\top, \dots, \varepsilon_n^\top)^\top \end{cases}.$$

هستند.

روش LPK بازگشتی

مدل نیمه‌پارامتری معرفی‌شده با فرض معلوم بودن β ها، به یک مدل ناپارامتری و با فرض معلوم بودن $f(\cdot)$ ، به یک مدل پارامتری تبدیل می‌شود. براین اساس، الگوریتم بازگشتی دو مرحله‌ای زیر را که توسط هستی و تیشیرانی (۱۹۹۰) معرفی شد، در نظر بگیرید.

الف) به ازای مقدار اولیه $\beta = \beta_0$ ، باقیمانده‌های $r_{ij} = y_{ij} - x_{ij}^\top \beta_0$ را به‌دست آورده و این مقادیر را به عنوان متغیر پاسخ جدید در نظر بگیرید. با در نظر گرفتن مدل ناپارامتری $r_{ij} = f(t_{ij}) + \varepsilon_{ij}$ و با استفاده از روش برآورد LPK ، تابع ناپارامتری به‌صورت $\hat{f} = S(y - X\beta_0)$ برآورد می‌شود، به‌طوری‌که S ماتریس هموارساز است.

^{۵۷}Backfitting algorithm

^{۵۸}Profile smoothing

ب) به ازای $\hat{f}$ به دست آمده در مرحله (الف)، باقیمانده‌های $a_{ij} = y_{ij} - \hat{f}(t_{ij})$ را به دست آورید. با در نظر گرفتن این مقادیر به عنوان پاسخ جدید، و در نظر گرفتن مدل پارامتری $a_{ij} = x_{ij}^\top \beta + \varepsilon_{ij}$ ، پارامترهای β را به روش ML برآورد کنید.

با تکرار دو مرحله فوق، در نهایت برآوردگرهای LPK بازگشتی برای β و f به ترتیب به صورت

$$\hat{\beta} = [X^\top (I_N - S)X]^{-1} [X^\top (I_N - S)y]$$

$$\hat{f} = S(y - X\hat{\beta}).$$

حاصل می‌شوند. شایان ذکر است که ماتریس هموارساز S بسته به نوع انتخاب توابع هسته، صورت‌های مختلفی دارد. علاقه‌مندان جهت آشنایی بیشتر می‌توانند به مراجع ذکر شده در بخش ۲.۳.۱ مراجعه کنند. همچنین در بخش ۴.۳.۱ به ذکر نمونه‌ای از این ماتریس خواهیم پرداخت.

روش LPK نیمرخ

با فرض معلوم بودن β و انتقال جمله پارامتری به طرف اول، مدل را به صورت $y - X\beta = f + \varepsilon$ می‌نویسیم. برآورد تابع ناپارامتری به روش LPK به صورت $\hat{f} = S(y - X\beta)$ خواهد بود. با جای‌گذاری این مقدار به جای f در مدل اولیه و با اندکی محاسبات خواهیم داشت،

$$(I_N - S)y = (I_N - S)X\beta + \varepsilon.$$

در نتیجه برآوردگرهای نیمرخ برای β و f به ترتیب به صورت

$$\hat{\beta} = [X^\top (I_N - S)^\top (I_N - S)X]^{-1} [X^\top (I_N - S)^\top (I_N - S)y],$$

$$\hat{g} = S(y - X^\top \hat{\beta})$$

حاصل می‌شوند.

۲.۳.۳.۱ برآورد به روش رگرسیون اسپلاین

پایه‌های توانی بریده‌شده از مرتبه d در نقاط گره $t_1, \dots, t_K$ را به صورت

$$B(t_{ij}) = (1, t_{ij}, \dots, t_{ij}^d, (t_{ij} - t_1)_+^d, \dots, (t_{ij} - t_K)_+^d)^\top$$

و تقریب تابع $f(t_{ij})$ بر اساس این پایه‌ها را به صورت $f(t_{ij}) = B^\top(t_{ij})\alpha$ در نظر بگیرید، که در آن $\alpha = (\alpha_1, \dots, \alpha_h)^\top$ بردار ضرایب رگرسیونی و $h = 1 + d + K$ است. مدل نیمه‌پارامتری (۶.۱) را می‌توان به صورت

$$y_{ij} = x_{ij}^\top \beta + B^\top(t_{ij})\alpha + \varepsilon_{ij} = \tilde{x}_{ij}^\top \theta + \varepsilon_{ij}, \quad i = 1, \dots, n; \quad j = 1, \dots, n_i$$

بیان نمود، به‌طوری‌که $\tilde{x}_{ij} = (x_{ij}^\top, B^\top(t_{ij}))^\top$ ، $\theta = (\beta^\top, \alpha^\top)^\top$ برآودگر WLS پارامتر θ به‌صورت

$$\hat{\theta} = \left(\sum_{i=1}^n \tilde{X}_i^\top V_i^{-1} \tilde{X}_i \right)^{-1} \sum_{i=1}^n \tilde{X}_i^\top V_i^{-1} y_i,$$

حاصل می‌شود، که در آن $\tilde{X}_i = (\tilde{x}_{i1}^\top, \dots, \tilde{x}_{in_i}^\top)^\top$. بدیهی است که p مولفه اصلی در $\hat{\theta}$ ، برآورد مولفه‌های β یعنی $\hat{\beta}$ هستند.

۳.۳.۳.۱ برآورد به روش اسپلین تاوانیده

مشابه روش رگرسیون اسپلین، این روش نیز $f(t)$ را به کمک پایه‌ها به‌صورت $f(t) = B^\top(t)\alpha$ تقریب می‌زند. این روش میزان ناهمواری $f(t)$ را کنترل می‌کند، لذا برآوردهای β و α از طریق کمینه کردن معیار کمترین توان‌های دوم تاوانیده زیر بدست می‌آید،

$$\sum_{i=1}^n \sum_{j=1}^{n_i} (y_{ij} - x_{ij}^\top \beta - B^\top(t_{ij})\alpha)^2 + \lambda \alpha^\top G \alpha$$

که در آن λ پارامتر هموارسازی و G ماتریس ناهمواری^{۵۹} است و به‌صورت

$$G = \begin{bmatrix} \mathbf{0}_{(d+1) \times (d+1)} & \mathbf{0}_{(d+1) \times K} \\ \mathbf{0}_{K \times (d+1)} & I_K \end{bmatrix}_{h \times h}$$

تعریف می‌شود. صورت ماتریسی معیار فوق را می‌توان به‌صورت

$$\|y - X\beta - B\alpha\|^2 + \lambda \alpha^\top G \alpha$$

بیان نمود، به‌طوری‌که $B = (B_1, \dots, B_n)^\top$ و $B_i = (B^\top(t_{i1}), \dots, B^\top(t_{in_i}))^\top$. با انجام عملیات جبری، کمینه‌کننده‌ها به‌صورت

$$\hat{\beta} = (X^\top P X)^{-1} (X^\top P y),$$

$$\hat{\alpha} = (B^\top Q B + \lambda G)^{-1} (B^\top Q y),$$

$$\hat{f}(t) = B^\top(t) \hat{\beta}$$

بدست می‌آیند، که در آن $P = I_N - B(B^\top B + \lambda G)^{-1} B^\top$ و $Q = I_N - X(X^\top X)^{-1} X^\top$. علاقه‌مندان جهت کسب اطلاعات بیشتر درباره روش اسپلین‌های رگرسیونی و تاوانیده می‌توانند به تحقیقات هی و همکاران (۲۰۰۲)، هی و همکاران (۲۰۰۵) و فیتزموریس و همکاران (۲۰۰۹) مراجعه نمایند.

^{۵۹}Roughness matrix

۴.۳.۱ مدل نیمه‌پارامتری با اثرات آمیخته

در بخش‌های قبل مدل‌هایی که در تحلیل داده‌های طولی استفاده می‌شوند، اعم از مدل‌های پارامتری، ناپارامتری و نیمه‌پارامتری را معرفی نمودیم. در چند دهه اخیر آماردانان در راستای تحلیل هر چه بهتر این داده‌ها، با خلق ترکیبی از مدل‌های بیان‌شده، رده بسیار منعطف از مدل‌ها را ایجاد نمودند. در این میان، ادغام مدل اثرات آمیخته از مجموعه مدل‌های پارامتری با مدل‌های ناپارامتری بسیار مورد توجه قرار گرفت. بر این اساس ما نیز با در نظر گرفتن ترکیب دو مدل ذکر شده به مطالعه هر چه بیشتر مدل نیمه‌پارامتری با اثرات آمیخته پرداختیم. در این مدل اثرات پارامتری و ثابت به مدل‌بندی خطی متغیرهای تبیینی و تابع ناپارامتری به مدل‌بندی غیرخطی نقاط زمانی می‌پردازند. همچنین همبستگی درون موضوعی از طریق اثرات تصادفی، در مدل اعمال شده است. بنابراین یک مدل نیمه‌پارامتری با اثرات آمیخته به صورت

$$y_{ij} = \mathbf{x}_{ij}^T \boldsymbol{\beta} + f(t_{ij}) + \mathbf{z}_{ij}^T \mathbf{b}_i + \varepsilon_{ij}, \quad (7.1)$$

بیان می‌گردد. در مدل فوق y_{ij} ، $(i = 1, \dots, n; j = 1, \dots, n_i)$ ، پاسخ مربوط به آزمودنی i ام در زمان t_{ij} است، که در آن $\boldsymbol{\beta} = (\beta_1, \dots, \beta_p)^T$ بردار $p \times 1$ از ضرایب ثابت رگرسیونی وابسته به متغیرهای تبیینی $\mathbf{x}_{ij}$ ، $f(\cdot)$ تابعی نامعلوم، هموار از زمان و دارای مشتق دوم، $\mathbf{b}_i = (b_{i1}, \dots, b_{iq})^T$ بردار $q \times 1$ بعدی از ضرایب تصادفی، مستقل از هم با میانگین ۰ و ماتریس کوواریانس D_i و وابسته به متغیرهای تبیینی $\mathbf{z}_{ij}$ و خطای تصادفی و مستقل از $\mathbf{b}_i$ ها با $E\{\varepsilon_{ij}\} = 0$ است. بدون کاستن از کلیت مسئله می‌توان فرض کرد که هر t_{ij} اعدادی در فاصله $[0, 1]$ می‌باشند.

مشابه کار فن و همکاران (۲۰۰۷)، فرض می‌کنیم که $\text{Var}\{\varepsilon_{ij}\} = \sigma^2(t_{ij})$ ، یعنی واریانس خطاهای تصادفی تابعی ناپارامتری و هموار نسبت به زمان هستند. از طرف دیگر فرض می‌شود تابع همبستگی بین $\varepsilon_i(t_{ij})$ و $\varepsilon_i(t_{ik})$ دارای صورت تابعی $\text{corr}\{\varepsilon_{ij}, \varepsilon_{ik}\} = \rho(t_{ij}, t_{ik}, \boldsymbol{\theta})$ است، که در آن $\rho(t_{ij}, t_{ik}, \boldsymbol{\theta})$ تابعی معین مثبت از t_{ij} و t_{ik} بوده و $\boldsymbol{\theta}$ برداری از پارامترهای همبستگی است.

مدل ماتریسی که در آن عناصر بر اساس آزمودنی i ام مشخص و بر اساس نقاط اندازه‌گیری ماتریسی شده‌اند به صورت

$$\mathbf{y}_i = \mathbf{X}_i \boldsymbol{\beta} + \mathbf{f}(\mathbf{t}_i) + \mathbf{Z}_i \mathbf{b}_i + \boldsymbol{\varepsilon}_i \quad (8.1)$$

نمایش داده می‌شود، که در آن

$$\begin{cases} \mathbf{y}_i = (y_i(t_{i1}), \dots, y_i(t_{in_i}))^T \sim n_i \times 1 \\ \mathbf{X}_i = (\mathbf{x}_i^T(t_{i1}), \dots, \mathbf{x}_i^T(t_{in_i}))^T \sim n_i \times p \\ \mathbf{f}(\mathbf{t}_i) = (f(t_{i1}), \dots, f(t_{in_i}))^T \sim n_i \times 1 \\ \mathbf{Z}_i = (\mathbf{z}_i^T(t_{i1}), \dots, \mathbf{z}_i^T(t_{in_i}))^T \sim n_i \times q \\ \boldsymbol{\varepsilon}_i = (\varepsilon_i(t_{i1}), \dots, \varepsilon_i(t_{in_i}))^T \sim n_i \times 1. \end{cases}$$

فرض کنید $V_i = Z_i^T D_i Z_i + \Sigma_i$ ماتریس کوواریانس $n_i \times n_i$ بعدی مربوط به پاسخ y_i بوده که در آن D_i ماتریس کوواریانس $q \times q$ بعدی مربوط به b_i ها و Σ_i ماتریس کوواریانس $n_i \times n_i$ بعدی مربوط به ε_i ها است. به تبعیت از لیانگ و زیگر (۱۹۸۶) ما نیز Σ_i را به صورت $\Sigma_i = A_i R_i(\theta) A_i$ در نظر گرفتیم، به طوری که $A_i = \text{diag}(\sigma(t_{i1}), \dots, \sigma(t_{in_i}))$ یک ماتریس قطری $n_i \times n_i$ بعدی با عناصر قطری $\sigma(t_{ij})$ و $R_i(\theta)$ ماتریس همبستگی بین $\varepsilon_i(t_{ij})$ و $\varepsilon_i(t_{ik})$ است که عضو (j, k) ام آن به صورت $\rho(t_{ij}, t_{ik}, \theta)$ تعریف می‌شود.

مدل (۸.۱) را می‌توان به صورت

$$y = X\beta + f(t) + Zb + \varepsilon$$

بازنویسی کرد، به طوری که

$$\begin{cases} y = (y_1^T, \dots, y_n^T)^T \sim N \times 1 \\ X = (X_1^T, \dots, X_n^T)^T \sim N \times p \\ f(t) = (f^T(t_1), \dots, f^T(t_n))^T \sim N \times 1 \\ Z = \text{diag}(Z_1, \dots, Z_n)^T \sim N \times nq \\ \varepsilon = (\varepsilon_1^T, \dots, \varepsilon_n^T)^T \sim N \times 1 \\ b = (b_1^T, \dots, b_n^T)^T \sim nq \times 1 \end{cases}$$

همچنین $V = \text{diag}(V_1, \dots, V_n) = Z^T D Z + \Sigma$ ماتریس کوواریانس $N \times N$ بعدی مربوط به y بوده که در آن $D = \text{diag}(D_1, \dots, D_n)$ ماتریس کوواریانس $nq \times nq$ بعدی مربوط به بردار اثر تصادفی b و $\Sigma = \text{diag}(\Sigma_1, \dots, \Sigma_n)$ ماتریس کوواریانس $N \times N$ بعدی مربوط به ε است. این مدل تعمیم مدل زیگر و دیگل (۱۹۹۴) است که جهت تحلیل داده‌های خود از یک مدل نیمه پارامتری با عرض از مبدأ تصادفی استفاده کردند و تابع ناپارامتری مدل را به روش بازگشتی برآورد نمودند.

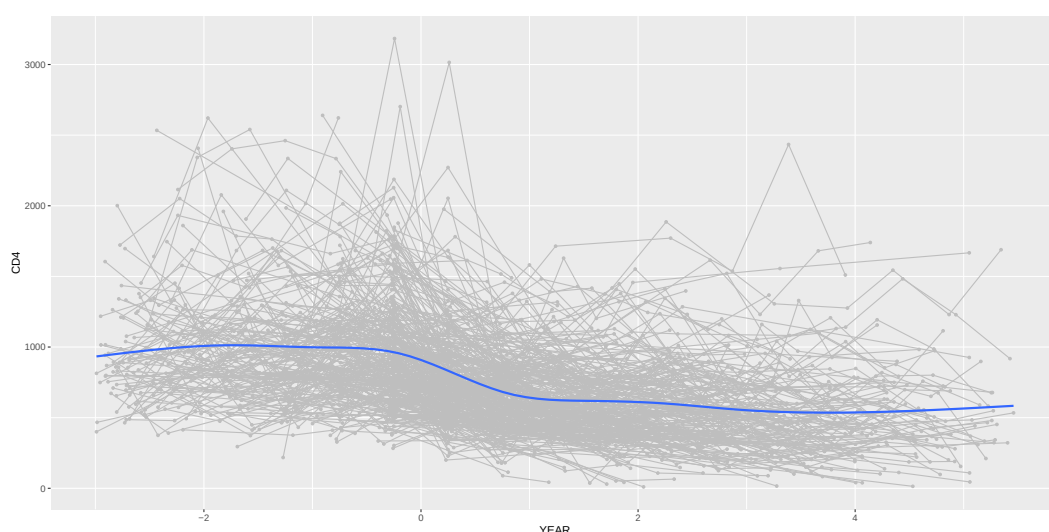
۴.۱ مثال‌های انگیزشی مورد بحث در این رساله

در این بخش مثال‌های واقعی را که در تحلیل مدل‌های ارائه شده در بخش‌های بعدی مورد استفاده قرار می‌گیرند، معرفی می‌کنیم.

۱.۴.۱ داده‌های $CD4$

این داده‌های مربوط به آزمایشات بالینی بیماری ایدز است که تعداد ۳۶۹ مرد مبتلا به HIV را مورد بررسی قرار می‌دهد. در این مثال متغیر پاسخ، تعداد سلول‌های $CD4$ و متغیرهای تبیینی شامل استفاده یا عدم استفاده از مواد مخدر ($DRUG$)، تعداد شرکای جنسی ($SEXP$)، تعداد بسته‌های

سیگار مورد استفاده در روز (*SMOKE*)، میزان افسردگی (*CESD*)، سن (*AGE*) و سال (*YEAR*) هستند. شکل ۳.۱ نمودار پراکنش تعداد سلول‌های $CD4$ در طول زمان یعنی *YEAR* را نشان می‌دهد. به منظور اینکه توزیع داده‌ها به توزیع نرمال نزدیک شود، تبدیل ریشه دوم متغیر پاسخ در نظر گرفته می‌شود. این داده‌ها در بسته *catdata* از نرم‌افزار *R* موجود است. جهت کسب اطلاعات بیشتر در مورد این داده‌ها می‌توان به زیگر و دیگل (۱۹۹۴) و وانگ و همکاران (۲۰۰۵) مراجعه نمود. زیگر و دیگل (۱۹۹۴) دریافتند که داده‌ها دارای همبستگی درونی از نوع ساختار همبستگی تبادلی پذیر با پارامتر همبستگی $\rho = 0.509$ هستند.



شکل ۳.۱: نمودار پراکنش تعداد سلول‌های $CD4$ در طول زمان (منحنی‌های خاکستری) به همراه میانگین هموار شده آن‌ها (منحنی آبی).

۲.۴.۱ داده‌های چرخه سلولی مخمر *ORF*

چرخه سلولی یک فرآیند زندگی کاملاً تنظیم شده است که در آن سلول‌ها رشد می‌کنند، *DNA* هایشان را تکثیر می‌دهند، کروموزوم‌های آن‌ها را جدا می‌کنند و تا جایی که محیط اجازه می‌دهد به تعدادی سلول دختر تقسیم می‌شوند. سلول‌های دختر سلول‌هایی هستند که از تقسیم یک سلول تک واحد حاصل می‌شود. روند چرخه سلولی معمولاً به مراحل G_1 ، S ، G_2 و M تقسیم می‌شوند. به طوری که مرحله G_1 مخفف GAP است. مرحله S مخفف عبارت سنتز (ترکیب یا تجزیه مواد برای تولید ماده یا مواد جدید) است که تکثیر *DNA* در آن رخ می‌دهد. مرحله G_2 مخفف GAP ۲ و مرحله M به معنای مرحله میتوز (تقسیم هسته یک سلول به دو هسته مشابه و روشی برای تکثیر سلولی) است. منظور از بیان ژن یعنی استفاده از آن ژن به منظور تولید *RNA* یا پروتئین است. *RNA* مخفف اسید ریبونوکلیک^{۶۰}، یک ترکیب پیچیده با وزن مولکولی بالاست که در ساختن پروتئین‌های سلولی نقش دارد. سه نوع اصلی *RNA* وجود دارند: *RNA* پیامبر (*mRNA*)، *RNA* ناقل (*tRNA*) و *RNA*

⁶⁰Ribonucleic acid

ریبوزومی ($rRNA$). وقتی ژنی مورد استفاده قرار می‌گیرد، گفته می‌شود آن ژن بیان شده و روشن است و وقتی که یک ژن مورد استفاده قرار نگیرد، می‌گویند آن ژن خاموش است. این که در زمان مشخصی کدام ژن روشن و کدام ژن خاموش است به تنظیم بیان ژن معروف است. عوامل رونویسی^{۶۱} (TFs) که نقش مهمی در تنظیمات بیان ژن دارند پروتئین‌هایی هستند که به دنباله‌های یک DNA خاص متصل می‌شود و از این طریق گردش اطلاعات ژنتیکی از DNA به $mRNA$ را کنترل می‌کند. تعداد این TF ها زیاد است و با علائم اختصاری نمادگذاری می‌شوند که $Mcm1$ ، $Swi6$ ، $Swi4$ ، $Mbp1$ ، $Ace2$ ، $Swi5$ ، $Ndd1$ ، $Fkh2$ ، $Fkh1$ و نمونه‌هایی از انواع این TF ها هستند. برای درک بهتر پدیده اساسی فرآیند چرخه سلولی، مهم است که TF هایی که سطح بیان ژن ژن‌های تنظیم شده از چرخه سلولی را تنظیم می‌کنند، شناسایی شوند. بنابراین نیازمند استفاده از روش‌های انتخاب متغیر هستیم تا TF های کلیدی شناسایی شوند. طبق مطالعات انجام شده توسط سیمون و همکاران (۲۰۰۱)، پروتئین‌های $Swi4$ ، $Swi6$ و $Mbp1$ در مرحله $G1$ ، پروتئین‌های $Ndd1$ ، $Mcm1$ و $Fkh1/2$ در مرحله $G2$ و پروتئین‌های $Swi5$ ، $Ace2$ و $Mcm1$ در مرحله M نقش مهمی دارند. داده بیان ژن چرخه سلولی مخمر^{۶۲} از آزمایشات بیولوژیکی اسپیلمن و همکاران (۱۹۹۸) که روی ژن $CDC15$ مخمر ORF انجام شده، جمع‌آوری شده‌اند. در این داده‌ها سطوح $mRNA$ سراسر ژنوم ۶۱۷۸ مخمر ORF در فواصل ۷ دقیقه و در مجموع به مدت ۱۱۹ دقیقه، شامل دو دوره چرخه سلولی و به تعداد ۱۸ مشاهده، اندازه‌گیری شده است. مجموعه داده‌ای که در این رساله استفاده می‌شود، زیر مجموعه‌ای از بیان ژن‌های چرخه سلولی مخمر ORF است که در آن چرخه سلولی ۲۸۳ ژن تنظیم شده $CDC15$ در ۴ نقطه زمانی و در مرحله $G1$ مشاهده شده‌اند را دربر می‌گیرد. در این داده‌ها متغیر پاسخ y_{ij} لگاریتم سطح بیان ژن مربوط به ژن i ام در زمان j است. متغیرهای تبیینی $x_{k,ij}$ و $k = 1, \dots, 96$ نمره تطابق احتمال اتصال TF k ام به ناحیه پیش‌برنده ژن i ام است. احتمال اتصال با استفاده از یک روش مدل‌سازی ترکیبی بر اساس داده‌های حاصل از آزمایش اتصال $ChIP$ محاسبه می‌شود. این داده‌ها در بسته Ime4 از نرم‌افزار R موجود است. برای آشنایی بیشتر با این مجموعه داده‌ها به وانگ و همکاران (۲۰۱۲) مراجعه کنید.

۳.۴.۱ داده‌های بیماری صفراوی

داده‌های بیماری صفراوی مربوط به مطالعه گروهی ۳۱۲ بیمار است که از ژانویه ۱۹۷۴ تا می ۱۹۸۴ به دلیل بیماری سیروز صفراوی اولیه پی‌آپی^{۶۳} به کلینیک مایو^{۶۴} که یک مرکز پزشکی-دانشگاهی در کشور آمریکا است، مراجعه کرده‌اند. بیماران شرکت‌کننده تا آوریل ۱۹۸۸ به‌طور تصادفی تحت

^{۶۱}Transcription factors

^{۶۲}Yeast cell-cycle gene expression

^{۶۳}Primary Biliary Cirrhosis sequential

^{۶۴}Mayo

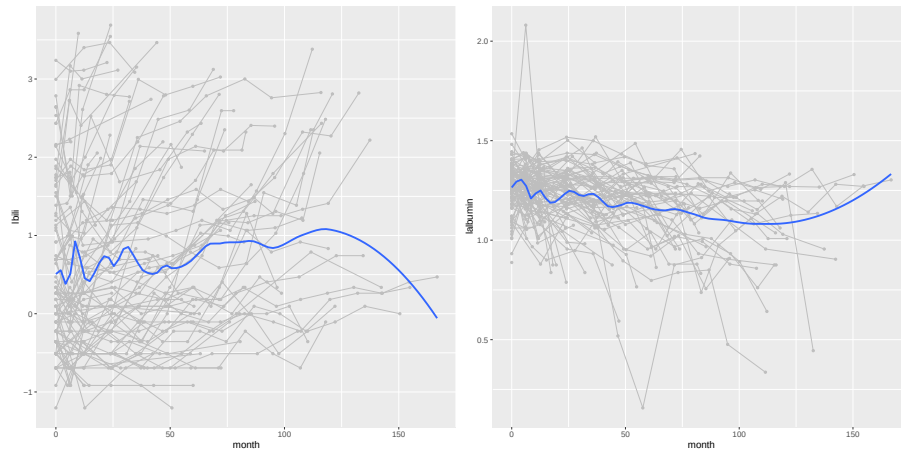
سه نوع درمان شامل D-پنیسیلامین^{۶۵}، دارونما^{۶۶} و آزمایش کور^{۶۷} قرار گرفته‌اند. آزمایش کور (نابینا) یا آزمایش چشمی یک آزمایش علمی است که افراد شرکت‌کننده در آن از داشتن اطلاعاتی که منجر شود به صورت خودآگاه یا ناخودآگاه به سمتی متمایل شوند، منع می‌شوند. به‌طور مثال وقتی از مصرف‌کننده پرسیده شود که طعم و مزه محصولات مختلف را با هم مقایسه کند، نام سازنده محصول باید پوشانده شود. در غیر این صورت معمولاً مصرف‌کننده به سمت محصولی که از قبل نسبت به آن شناخت دارد متمایل می‌شود. به همین شکل وقتی قدرت تأثیر یک دارو مورد آزمایش قرار می‌گیرد، هم بیماران و هم پزشکان شرکت‌کننده در آزمایش از دوز مصرفی بی‌اطلاع هستند. این باعث خواهد شد شانس هرگونه تأثیر تلقینی یا فریب هوشیارانه از بین برود. آزمایش کور ابزار بسیار مهم و کاربردی در شاخه‌های مختلف علوم از جمله علم پزشکی، روانشناسی، عصب‌شناسی و بیولوژی، علوم کاربردی و بازاریابی است. به‌کارگیری روش علمی آزمایش کور در برخی رشته‌ها مانند آزمون دارو ضروری است. مجموعه داده‌های *PBCseq* شامل شماره شناسایی (*ID*) بیمار، پنج متغیر وابسته به زمان (مانند سن و تعداد روزهای تحت درمان)، هشت متغیر رسته‌ای (مانند جنسیت، نوع دارو و وضعیت بیماری) و هفت متغیر پیوسته (مانند لگاریتم بیلی‌روبین و آلبومین) است. این داده‌ها در بسته mixAK از نرم‌افزار *R* و همچنین در آدرس <http://lib.stat.cmu.edu/datasets/pbcseq> قابل دسترسی است. برای توضیحات بیش‌تر درباره داده‌ها به وانگ (۲۰۱۷) مراجعه کنید.

برای بیماران مبتلا به سیروز صفراوی اولیه پیشرفته، پیوند ارتوپدی کبد یک روش جایگزین برای درمان است. بیلی‌روبین و آلبومین دو شاخص اصلی برای ارزیابی وجود یا عدم وجود بیماری‌های کبدی هستند. اگر در صفرا و ادرار سطح بیلی‌روبین بسیار بالاتر از حد استاندارد باشد، نشان‌دهنده وجود بیماری‌های خاصی است. سطح آلبومین بسیار زیاد یا خیلی کم نیز ممکن است برای افراد مضر باشد. اعتقاد بر این است که بین سطح بیلی‌روبین و آلبومین رابطه وجود دارد. بنابراین تحلیل مشترک این دو عامل که به‌طور مکرر در طول زمان جمع‌آوری شده‌اند در تشخیص بیماری‌های کبدی بسیار مفید خواهد بود. وانگ (۲۰۱۷) نشان داد که در این داده‌ها، مشاهدات نامتعارف یا پرت وجود دارد، همچنین اشاره کردند که تحلیل جداگانه این دو عامل باعث از دست رفتن بسیاری از اطلاعات مهم می‌شود. بنابراین وی روی مدل‌سازی توام دو متغیر جمع‌آوری شده در طول زمان، لگاریتم سطح بیلی‌روبین (*lbili*) و لگاریتم سطح آلبومین (*lalbumin*)، در حضور سایر متغیرهای مورد علاقه شامل سن (*Age*)، جنسیت (*Sex*) و نوع درمان (*Drug*) متمرکز شد. شکل ۴.۱ نمودار پراکنش *lbili* و *lalbumin* در طول زمان برای ۱۰۰ بیمار که به‌طور تصادفی انتخاب شده‌اند را نشان می‌دهد. اجرای چنین تحلیلی نیازمند استفاده از ابزار روش‌های چندمتغیره است و از آنجایی که دارای مقادیر پرت است باید روش‌هایی اتخاذ شود که نسبت به داده‌های پرت تنومند باشند. در فصل ۳ به تحلیل این داده‌ها پرداخته می‌شود.

⁶⁵D-penicillamine

⁶⁶placebo

⁶⁷Double-blind



شکل ۴.۱: نمودار پراکنش داده‌های *PBCseq* در طول زمان. سطوح *lbili* و *lalalbumin* برای ۱۰۰ بیمار (منحنی‌های خاکستری) به همراه میانگین هموار شده آن‌ها (منحنی آبی).

۵.۱ روش‌های برآورد در مدل نیمه‌پارامتری با اثرات آمیخته

چالش اصلی در برآورد مدل نیمه‌پارامتری با اثرات آمیخته، نحوه برخورد با مولفه ناپارامتری است که آماردانان بسیاری بسته به نوع مهارت و علاقه خود یکی از روش‌های ناپارامتری معرفی شده را به کار گرفتند. چن و چو (۲۰۰۷) و چن و چو (۲۰۰۹) از روش هموارسازی اسپلاین، چانگ (۲۰۰۴) و لی و چو (۲۰۱۰) از روش اسپلاین‌های توانیده و لیانگ (۲۰۰۹) از روش رگرسیون خطی موضعی استفاده کردند. از آنجایی که در بین روش‌های ناپارامتری معرفی شده روش‌های هسته‌ای و بازگشتی در این مدل مورد توجه قرار نگرفته‌اند، مطالعات خود را روی این روش‌ها به عنوان مطالعه‌ای جدید متمرکز نموده و نتایج حاصل در ادامه گزارش می‌شوند.

۱.۵.۱ روش برآوردگر هسته‌ای

مدل (۸.۱) را در نظر بگیرید و فرض کنید خطاهای تصادفی دارای توزیع نرمال با میانگین صفر هستند. در این صورت توزیع حاشیه‌ای زیر برای پاسخ y_i نتیجه می‌شود:

$$y_i \sim \mathcal{N}_{n_i}(X_i\beta + f(t_i), Z_i^\top D_i Z_i + \Sigma_i). \quad (9.1)$$

اگر امید ریاضی y_{ij} و x_{ij} به شرط نقاط زمانی t_{ij} را به ترتیب به صورت

$$\mu_Y(t_{ij}) = E\{y_{ij}|t = t_{ij}\}, \quad \mu_X(t_{ij}) = E\{x_{ij}|t = t_{ij}\},$$

تعریف کنیم، آنگاه رابطه

$$\mu_Y(t_{ij}) = \mu_X^\top(t_{ij})\beta + f(t_{ij})$$

نتیجه می‌شود. با در نظر گرفتن صورت برداری امیدهای شرطی به صورت

$$\mu_Y(t_i) = (\mu_Y(t_{i1}), \dots, \mu_Y(t_{in_i}))^\top, \quad \mu_X(t_i) = (\mu_X^\top(t_{i1}), \dots, \mu_X^\top(t_{in_i}))^\top,$$

رابطه

$$f(t_i) = \mu_Y(t_i) - \mu_X(t_i)\beta$$

را خواهیم داشت. با جای‌گذاری جمله فوق در رابطه (۹.۱) داریم،

$$y_i - \mu_Y(t_i) \sim \mathcal{N}_{n_i}([X_i - \mu_X(t_i)]\beta, Z_i^\top D_i Z_i + \Sigma_i).$$

به منظور برآورد ضرایب رگرسیونی β ، معیار WLS را به صورت

$$l(\beta) = \sum_{i=1}^n (\tilde{y}_i - \tilde{X}_i^\top \beta)^\top V_i^{-1} (\tilde{y}_i - \tilde{X}_i^\top \beta),$$

در نظر بگیرید که در آن $\tilde{y}_i = y_i - \mu_Y(t_i)$ و $\tilde{X}_i = X_i - \mu_X(t_i)$. معیار فوق به دلیل مجهول بودن توابع $\mu_Y(\cdot)$ و $\mu_X(\cdot)$ قابل استفاده نیست. از این رو ما روش هموارسازی هسته‌ای را در نظر گرفته و برآوردگر آن‌ها را به صورت

$$\hat{\mu}_X(t) = \sum_{i=1}^n \sum_{j=1}^{n_i} \omega_{ij} x_{ij}, \quad \hat{\mu}_Y(t) = \sum_{i=1}^n \sum_{j=1}^{n_i} \omega_{ij} y_{ij},$$

ارائه می‌دهیم، که در آن $\omega_{ij} = K_h(t_{ij} - t) / \sum_{k=1}^n \sum_{l=1}^{n_k} K_h(t_{kl} - t)$ ، $K_h(\cdot) = K(\cdot/h)$ ، h پهنای باند، و $K(\cdot)$ یک تابع هسته است. بنابراین با جای‌گذاری $\hat{\mu}_X(t)$ و $\hat{\mu}_Y(t)$ به ترتیب به جای $\mu_X(t)$ و $\mu_Y(t)$ معیار $l(\beta)$ به صورت

$$\hat{l}(\beta) = \sum_{i=1}^n (\tilde{y}_i - \tilde{X}_i \beta)^\top V_i^{-1} (\tilde{y}_i - \tilde{X}_i \beta) \quad (10.1)$$

تقریب زده می‌شود، که در آن $\tilde{y}_i = y_i - \hat{\mu}_Y(t_i)$ و $\tilde{X}_i = X_i - \hat{\mu}_X(t_i)$ ؛ $\hat{\mu}_Y(t_i) = (\hat{\mu}_Y(t_{i1}), \dots, \hat{\mu}_Y(t_{in_i}))^\top$ و $\hat{\mu}_X(t_i) = (\hat{\mu}_X^\top(t_{i1}), \dots, \hat{\mu}_X^\top(t_{in_i}))^\top$. برآوردگرهای هسته‌ای β و $f(t_i)$ از طریق حل مسئله بهینه‌سازی $\arg \min_{\beta \in \mathbb{R}^p} \hat{l}(\beta)$ ، به ترتیب به صورت

$$\begin{aligned} \hat{\beta}_K &= \left[\sum_{i=1}^n \tilde{X}_i^\top V_i^{-1} \tilde{X}_i \right]^{-1} \left[\sum_{i=1}^n \tilde{X}_i^\top V_i^{-1} \tilde{y}_i \right], \\ \hat{f}(t_i) &= \hat{\mu}_Y(t_i) - \hat{\mu}_X^\top(t_i) \hat{\beta}_K. \end{aligned}$$

حاصل می‌شوند. صورت ماتریسی $\hat{f}(t_i)$ و $\hat{\beta}_K$ عبارت‌اند از:

$$\begin{aligned} \hat{\beta}_K &= \left[X^\top (I_N - S)^\top V^{-1} (I_N - S) X \right]^{-1} \left[X^\top (I_N - S)^\top V^{-1} (I_N - S) y \right], \\ \hat{f} &= S(y - X \hat{\beta}_K). \end{aligned}$$

در روابط فوق، S یک ماتریس هموارساز است. لازم به ذکر است که برآوردگر معرفی شده همان برآوردگر LPK نیمرخی است، با این تفاوت که S به صورت

$$S = \text{diag}(S_1, \dots, S_n), \quad S_i = \text{diag}(K_h(t_{i1} - t), \dots, K_h(t_{in_i} - t))$$

تعریف می‌شود. برآورد پارامترهای تصادفی مدل با محاسبه امید شرطی b_i به شرط y_i و $\hat{\beta}_K$ به صورت

$$\hat{b}_i = E[b_i | y_i] = D_i Z_i V_i^{-1} (y_i - X_i \hat{\beta}_K).$$

حاصل می‌شود.

لازم به ذکر است که برآوردگر هسته‌ای نسبت به پهنای باند حساس است، لذا تعیین پهنای باند در این روش از اهمیت بسزایی برخوردار است. روش‌های متعددی جهت انتخاب پهنای باند مناسب وجود داشته که بیان تمامی آن‌ها خارج از بحث این مجموعه است. لذا جهت آشنایی مختصر، به شرح تنها یک روش به نام روش اعتبارسنجی متقابل^{۶۸} بسنده می‌کنیم. ایده این روش که توسط رادمو (۱۹۸۲) و بومن (۱۹۸۴) ارائه گردید، به این صورت است که پارامتر بهینه h به وسیله رابطه

$$\hat{h} = \arg \min_h CV(h),$$

تعیین می‌گردد که در آن

$$CV(h) = \int \hat{\mu}_X^2(t) dt - \frac{1}{n} \sum_i \hat{\mu}_{(-i)X}(t_{ij}), \quad \hat{\mu}_{(-i)X}(t_{ij}) = \frac{1}{(n-1)h} \sum_{i' \neq i} \sum_j K\left(\frac{t_{i'j} - t}{h}\right).$$

تکنیک دیگری که در تحقیقات جدید از آن برای برآورد β ، f و b استفاده کرده ایم روش بازگشتی است که در بخش بعدی توضیح داده می‌شود.

۲.۵.۱ روش برآوردگر بازگشتی

با در نظر گرفتن رابطه (۹.۱)، روش برآوردگر بازگشتی به این صورت عمل می‌کند که با تکرار دو گام

الف) به ازای مقدار اولیه $\beta = \beta_0$ ، $f(t_i)$ را به صورت $\hat{f}(t_i) = \hat{\mu}_Y(t_i) - \hat{\mu}_X(t_i)$ برآورد کنید.

ب) با اجرای رگرسیون $y_i - \hat{f}(t_i)$ روی X_i ، ضرایب رگرسیونی β را با کمینه کردن معیار زیر برآورد کنید:

$$\hat{l}(\beta) = \sum_{i=1}^n (y_i - \hat{f}(t_i) - X_i \beta)^\top V_i^{-1} (y_i - \hat{f}(t_i) - X_i \beta)$$

برآوردگر بازگشتی β و $f(t_i)$ به ترتیب به صورت

$$\begin{aligned} \hat{\beta}_{BF} &= \left[\sum_{i=1}^n X_i^\top V_i^{-1} \tilde{X}_i \right]^{-1} \left[\sum_{i=1}^n X_i^\top V_i^{-1} \tilde{y}_i \right], \\ \hat{f}(t_i) &= \hat{\mu}_Y(t_i) - \hat{\mu}_X(t_i) \hat{\beta}_{BF} \end{aligned}$$

⁶⁸Cross validation

حاصل می‌شوند. همچنین صورت ماتریسی $\hat{\beta}_{BK}$ و $\hat{f}(t_i)$ عبارت است از:

$$\begin{aligned}\hat{\beta}_{BF} &= \left[X^T V^{-1} (I_N - S) X \right]^{-1} \left[X^T V^{-1} (I_N - S) y \right], \\ \hat{f} &= S(y - X \hat{\beta}_{BF}).\end{aligned}$$

برآورد پارامترهای تصادفی مدل با محاسبه امید شرطی b_i به شرط y_i و $\hat{\beta}_{BF}$ به صورت

$$\hat{b}_i = E[b_i | y_i] = D_i Z_i V_i^{-1} (y_i - X_i^T \hat{\beta}_{BF})$$

حاصل می‌شود.

لازم به ذکر است که تقریب مولفه ناپارامتری به روش اسپلین‌ها بسیار متداول‌تر و محبوب‌تر از روش‌های موضعی قرار گرفته‌اند. زیرا این روش با تعداد کمی از گره‌ها تقریب خوبی را از یک تابع هموار فراهم می‌کند. همچنین با در نظر گرفتن توابع پایه به عنوان شبه‌متغیر، تابعی خطی از مولفه‌های ناپارامتری در اختیار قرار می‌دهد. علیرغم این که در بخش‌های قبل انواع آن را به تفصیل شرح دادیم، به منظور در اختیار داشتن مجموعه‌ای کامل از روش‌های برآورد در مدل نیمه‌پارامتری با اثرات آمیخته، در ادامه به بیان اجمالی این روش می‌پردازیم.

۳.۵.۱ روش برآوردگر اسپلین

پایه‌های توانی بریده‌شده از مرتبه d در نقاط گره $t_1, \dots, t_K$ را به صورت

$$B(t) = (1, t, \dots, t^d, (t - t_1)_+^d, \dots, (t - t_K)_+^d)^T$$

و تقریب تابع $f(t)$ بر اساس این پایه‌ها را به صورت $f(t) = B^T(t)\alpha$ در نظر بگیرید، که در آن $\alpha = (\alpha_1, \dots, \alpha_h)^T$ بردار ضرایب رگرسیونی و $h = d + K + 1$ است. به این ترتیب مدل (۸.۱) را می‌توان به صورت مدل آمیخته

$$y_i = X_i^T \beta + B_i^T(t_i) \alpha + Z_i^T b_i + \varepsilon_i, \quad (11.1)$$

بیان نمود، که در آن $B_i(t_i) = (B^T(t_{i1}), \dots, B^T(t_{in_i}))^T$. برای سادگی و کاربردی بودن، مدل (۱۱.۱) را به صورت

$$y_i = \tilde{X}_i^T \theta + Z_i^T b_i + \varepsilon_i,$$

بازنویسی می‌کنیم، به طوری که $\tilde{X}_i = (X_i, B_i(t_i))^T$ ماتریس طرح ادغام‌شده $(p + h) \times n_i$ بعدی از ماتریس اثرات ثابت و اثرات اسپلین و $\theta = (\beta^T, \alpha^T)^T$ بردار پارامتر ادغام‌شده $(p + h) \times 1$ بعدی می‌باشند. با در نظر گرفتن توزیع حاشیه‌ای پاسخ‌ها به صورت

$$y_i \sim N_{n_i}(\tilde{X}_i \theta, Z_i^T D_i Z_i + \Sigma_i),$$

برآوردگر WLS پارامتر θ به صورت

$$\hat{\theta}_S = \left(\sum_{i=1}^n \tilde{X}_i^\top V_i^{-1} \tilde{X}_i \right)^{-1} \sum_{i=1}^n \tilde{X}_i^\top V_i^{-1} y_i,$$

حاصل می‌شود. مشابه دو برآوردگر قبل برآورد ضرایب تصادفی مدل با محاسبه امید شرطی b_i به شرط y_i و $\hat{\theta}_S$ به صورت

$$\hat{b}_i = E[b_i | y_i] = D_i Z_i V_i^{-1} (y_i - \tilde{X}_i^\top \hat{\theta}_S)$$

حاصل می‌شود.

۴.۵.۱ خواص جانبی

در این بخش به بیان توزیع جانبی برآوردگرهای هسته‌ای و بازگشتی می‌پردازیم. فرض کنید به ازای هر $n_i = m < \infty$ ، مجموعه سه‌تایی (y_i, X_i, t_i) ، $i = 1, \dots, n$ ، متغیرهایی مستقل و هم‌توزیع باشند؛ نقاط زمانی t_{ij} دارای چگالی حاشیه‌ای $f_j(t)$ باشند؛ مشتق r ام تابعی مانند $d(\cdot)$ را به صورت $d^{(r)}(\cdot)$ تعریف می‌کنیم؛ عضو (j, k) ام ماتریس V^{-1} را به صورت ν^{jk} نشان می‌دهیم؛ تابع چگالی هسته $K(\cdot)$ دارای میانگین 0 و واریانس واحد است؛ یعنی $\int s K(s) ds = 0$ و $\int s^2 K(s) ds = 1$. همچنین فرض کنید $h \propto n^{-\alpha}$ ، $\frac{1}{5} \leq \alpha \leq \frac{1}{3}$ و $\tilde{X} = X + \lim_{n \rightarrow \infty} \frac{\partial \hat{f}(t; \beta)}{\partial \beta}$ شرایط فوق در لین و کارول (۲۰۰۲) نیز استفاده شده‌اند.

قضیه ۱.۵.۱. تحت شرایط بیان‌شده، وقتی $n \rightarrow \infty$ ، آنگاه در رابطه با توزیع جانبی برآوردگر هسته‌ای داریم

$$\sqrt{n} \left\{ \hat{\beta}_K - \beta + h^\top b_K(\beta, f) / 2 \right\} \xrightarrow{d} \mathcal{N}_p(0, V_K),$$

که در آن

$$\begin{aligned} b_K(\beta, f) &= E \left\{ \tilde{X}^\top V^{-1} \tilde{X} \right\}^{-1} E \left\{ \tilde{X}^\top V^{-1} f^{(2)}(t) \right\}, \\ V_K &= E \left\{ \tilde{X}^\top V^{-1} \tilde{X} \right\}^{-1} E \left\{ (W_1 - W_2)^\top V_o (W_1 - W_2) \right\} E \left\{ \tilde{X}^\top V^{-1} \tilde{X} \right\}^{-1}, \\ V_o &= \text{Var}(Y | X, t), \quad W_1 = V^{-1} \tilde{X}, \quad W_2 = (W_{21}, \dots, W_{2m}) \\ W_{2j} &= \frac{\left\{ \sum_{k=1}^m \sum_{l=1}^m E(\tilde{X}_k \nu^{kl} | t_l = t_j) \right\} f_j(t_j)}{\sum_{l=1}^m f_l(t_j)} \end{aligned}$$

اثبات قضیه در پیوست آمده است.

قضیه ۲.۵.۱. تحت شرایط بیان‌شده، وقتی $n \rightarrow \infty$ ، آنگاه در رابطه با توزیع جانبی برآوردگر بازگشتی داریم

$$\sqrt{n} \left\{ \hat{\beta}_{BF} - \beta + h^\top b_{BF}(\beta, f) / 2 \right\} \xrightarrow{d} \mathcal{N}_p(0, V_{BF}),$$

که در آن

$$\begin{aligned} b_{BF}(\beta, f) &= E \left\{ \tilde{X}^\top V^{-1} \tilde{X} \right\}^{-1} E \left\{ X^\top V^{-1} f^{(2)}(t) \right\}, \\ V_{BF} &= E \left\{ \tilde{X}^\top V^{-1} \tilde{X} \right\}^{-1} E \left\{ (W_1^* - W_2^*)^\top V_o (W_1^* - W_2^*) \right\} E \left\{ \tilde{X}^\top V^{-1} \tilde{X} \right\}^{-1}, \\ W_1^* &= V^{-1} X, \quad W_2^* = (W_{21}^*, \dots, W_{2m}^*), \\ W_{2j}^* &= \frac{\left\{ \sum_{k=1}^m \sum_{l=1}^m E(X^k \nu^{kl} | t_l = t_j) \right\} f_j(t_j)}{\sum_{l=1}^m f_l(t_l)}. \end{aligned}$$

اثبات این قضیه مشابه قضیه ۲ در هو و همکاران (۲۰۰۴) است.

قضیه ۳.۵.۱. تحت شرایط نظم بیان شده، ماتریس کوواریانس برآوردگر هسته‌ای حداکثر به بزرگی برآوردگر بازگشتی است. به عبارت دیگر ماتریس $V_{BF} - V_K$ یک ماتریس نیمه‌معین مثبت است.

قضیه فوق به مقایسه عملکرد دو برآوردگر هسته‌ای و بازگشتی می‌پردازد و بیان می‌کند که در داده‌های طولی برآوردگر هسته‌ای همواره بر برآوردگر بازگشتی برتری دارد. اثبات قضیه در پیوست آمده است.

۵.۵.۱ الگوریتم تکراری

در تمامی روش‌های معرفی شده، کارایی برآورد ضرایب رگرسیونی به برآورد تابع کوواریانس وابسته است. بنابراین برآورد نهایی را از طریق به‌روز کردن β ، b_i و $f(t_i)$ نسبت به برآورد ماتریس کوواریانس و با به‌کارگیری یک الگوریتم تکراری بین آن‌ها به‌دست می‌آوریم. ابتدا برآوردگرهای به‌روز شده β ، b_i و $f(t_i)$ نسبت به برآورد ماتریس کوواریانس برای روش هسته‌ای را به‌صورت

$$\begin{aligned} \hat{\beta}_K^{(r)} &= \left(\sum_{i=1}^n \tilde{X}_i^\top V_i^{-1(r-1)} \tilde{X}_i \right)^{-1} \left(\sum_{i=1}^n \tilde{X}_i^\top V_i^{-1(r-1)} \tilde{y}_i \right), \quad (12.1) \\ \hat{f}^{(r)}(t_i) &= \hat{\mu}_Y(t_i) - \hat{\mu}_X(t_i) \hat{\beta}_K^{(r)}, \\ \hat{b}_i^{(r)} &= D_i Z_i V_i^{-1(r-1)} (\tilde{y}_i - \tilde{X}_i \hat{\beta}_K^{(r)}), \end{aligned}$$

در نظر بگیرید. توجه کنید که برای سایر برآوردگرها رابطه فوق به‌طور مشابه تعریف می‌شود و تنها کافی است که به جای $\hat{\beta}_K^{(r)}$ برآوردگر مورد نظر جایگزین شود. الگوریتم پیشنهادی ۱، مراحل را مشخص می‌کند.

الگوریتم ۱ الگوریتم تکراری برای برآوردگر ریج

۱: قرار دهید $r = 0$. با فرض استقلال، برآورد اولیه برای θ ، $\sigma^2(t)$ و D_i ، به صورت $\hat{\theta}^{(r)} = 1$ ، $\hat{\sigma}^{2(r)}(t) = 1$ و $\hat{D}_i^{(r)} = I_q$ است.

۲: قرار دهید $r = r + 1$. بر حسب مقادیر گام قبل، $\hat{f}(t)$ ، $\hat{\beta}$ و $\hat{b}_i$ را به کمک رابطه (۱۲.۱) به‌روز کنید، که در آن

$$\hat{V}_i^{(r-1)} = Z_i^\top D_i^{(r-1)} Z_i + \hat{\Sigma}_i^{(r-1)}, \quad \hat{\Sigma}_i^{(r-1)} = \hat{A}_i^{(r-1)} R_i(\hat{\theta}^{(r-1)}) \hat{A}_i^{(r-1)},$$

$$\hat{A}_i^{(r-1)} = \text{diag}(\hat{\sigma}^{(r-1)}(t_{i1}), \dots, \hat{\sigma}^{(r-1)}(t_{in_i})).$$

۳: مقادیر $\hat{D}_i^{[r]}$ ، $\hat{\sigma}^{2[r]}(t)$ و $\hat{\theta}^{[r]}$ را به صورت

$$\hat{D}_i^{(r)} = n^{-1} \sum_{i=1}^n \left(\hat{b}_i^{(r)} \hat{b}_i^{\top(r)} + [\hat{D}_i^{(r-1)} - \hat{D}_i^{(r-1)} Z_i \hat{V}_i^{(r-1)} Z_i^\top \hat{D}_i^{(r-1)}] \right),$$

$$\hat{\sigma}^{2(r)}(t) = \frac{\sum_{i=1}^n \sum_{j=1}^{n_i} \hat{r}_{ij}^{2(r)} K_{h_1}(t - t_{ij})}{\sum_{i=1}^n \sum_{j=1}^{n_i} K_{h_1}(t - t_{ij})},$$

$$\hat{\theta}^{(r)} = \arg \max_{\theta \in [-1, 1]} \left(-\frac{1}{2} \sum_{i=1}^n \{ \log |R_i(\hat{\theta}^{(r-1)})| + \hat{r}_i^{\top(r)} \hat{A}_i^{-1} R_i^{-1}(\hat{\theta}^{(r-1)}) \hat{A}_i^{-1} \hat{r}_i^{(r)} \} \right),$$

به‌روز کنید، که در آن $\hat{r}_{ij}^{(r)} = y_i(t_{ij}) - x_i(t_{ij})\hat{\beta}^{(r)} - z_i(t_{ij})\hat{b}_i^{(r)} - \hat{f}^{[r]}(t_{ij})$ و آن گام‌های دوم و سوم را تا رسیدن به هم‌گرایی ($|\hat{\beta}_{(r+1)} - \hat{\beta}_{(r)}| < \varepsilon$) تکرار کنید.

۶.۵.۱ مطالعات عددی

۱.۶.۵.۱ مطالعه شبیه‌سازی

در این قسمت به منظور مقایسه برآوردگرهای هسته‌ای، بازگشتی و اسپلین در مدل (۷.۱) برای داده‌های طولی، متغیرهای پاسخ را از مدل

$$y_{ij} = x_{1,ij}\beta_1 + x_{2,ij}\beta_2 + x_{3,ij}\beta_3 + f(t_{ij}) + b_i + \varepsilon_{ij}, \quad i = 1, \dots, n, \quad j = 1, \dots, 4$$

شبیه‌سازی می‌کنیم، که در آن مولفه‌های مدل به صورت

$$\beta^\top = (\beta_1, \beta_2, \beta_3) = (0.5, -1, 1), \quad x_{1,ij}, x_{2,ij}, x_{3,ij} \sim \mathcal{N}(0, 1),$$

$$f(t_{ij}) = 2 \sin(2\pi t_{ij}), \quad t_{ij} \sim \mathcal{U}(0, 1)$$

$$b_i \sim \mathcal{N}(0, \sigma_b^2); \quad \sigma_b^2 = 0.25,$$

$$\varepsilon_{ij} \sim \mathcal{N}(0, \sigma^2(t)), \quad \sigma^2(t) = \exp\left(\frac{t}{\sqrt{2}}\right)$$

تعیین شده‌اند. در این شبیه‌سازی، ساختار همبستگی بین مشاهدات مکرر را از طریق مدل همبستگی اتورگرسیو مرتبه اول $AR(1)$ به صورت $\text{corr}(\varepsilon_{is}, \varepsilon_{it}) = \rho^{|t-s|}$ در آن $s \neq t$ برای

آزمودن تاثیر شدت همبستگی از نوع ضعیف، متوسط و قوی، مقادیر $0/3$ ، $0/6$ و $0/9$ را برای ρ اتخاذ کردیم. همچنین از تابع هسته اپانچنیکوف که به صورت $K(u) = 0/75(1 - u^2)_+$ است، استفاده می‌شود. برای بررسی کارایی برآوردها با افزایش حجم نمونه مقدار n را یک مقدار نسبتاً کوچک، $n = 50$ ، و یک مقدار نسبتاً بزرگ، $n = 100$ ، در نظر گرفتیم.

ابتدا برای مقایسه عملکرد مدل‌ها، مدل اثرات آمیخته، مدل نیمه‌پارامتری و مدل نیمه‌پارامتری با اثرات آمیخته را به مجموعه داده‌های شبیه‌سازی شده برازش دادیم؛ که مولفه ناپارامتری به روش اسپلاین تقریب زده شد. سپس برای مقایسه عملکرد برآوردها سه روش اسپلاین، هسته‌ای و بازگشتی را در مدل نیمه‌پارامتری با اثرات آمیخته به کار بردیم. به منظور بررسی تاثیر اعمال ساختار همبستگی نادرست و همچنین نادیده گرفتن آن، برآوردها را تحت ساختار همبستگی واقعی (AR)، همبستگی تبادلی پذیر (EX) که به صورت $\text{corr}(\varepsilon_i(s), \varepsilon_i(t)) = \rho$ بوده، و با فرض استقلال (IND) به دست آوردیم. دقت برآورد در هر یک از مدل‌ها و همچنین برآوردها را توسط مقادیر اریبی و انحراف استاندارد (SD) پارامترها همچنین توسط معیار کمترین توان‌های دوم خطا (MSE) بر اساس میانگین نتایج 500 تکرار از شبیه‌سازی مورد سنجش قرار دادیم. برای آگاهی در مورد این معیارها ضمیمه آ را ببینید.

جدول ۲.۱ نتایج مربوط به مقادیر MSE ، اریبی و SD تحت برازش سه مدل است که اعداد داخل پرانتز مقدار SD است. از نظر معیار MSE مدل نیمه‌پارامتری با اثرات آمیخته عملکرد بهتری نسبت به دو مدل دیگر دارد؛ حتی زمانی که ساختار همبستگی نادرست انتخاب شود. در این مدل، حالت همبستگی نادرست EX رفتاری بسیار شبیه به حالت همبستگی درست AR دارد اما هنگامی که شدت همبستگی بالاست، AR بر EX پیشی می‌گیرد. حالت IND بدترین عملکرد را داشته است زیرا همبستگی را نادیده گرفته است. همچنین این مدل مقادیر اریبی و SD کمتری نسبت به دو مدل دیگر دارد.

نتایج حاصل از کاربرد برآوردهای مختلف مدل نیمه‌پارامتری با اثرات آمیخته در جدول ۳.۱ گزارش شدند. با توجه به نتایج، هر یک از برآوردها تحت ساختار همبستگی درست AR عملکرد خوبی نسبت به EX یا IND دارند و نادیده گرفتن همبستگی تاثیر نامطلوبی روی مقادیر اریبی، SD و همچنین MSE گذاشته است. روش اسپلاین در تمام جوانب بهتر از دو برآوردها دیگر عمل کرده است. همچنین با مقایسه دو روش موضعی می‌توان دریافت که مقادیر اریبی و SD روش هسته‌ای نسبت به روش بازگشتی کمتر است که این مشاهدات با نتایج نظری قضیه ۳.۵.۱ هم‌خوانی دارد.

۲.۶.۵.۱ تحلیل داده‌های $CD4$

در این بخش به مطالعه رفتار برآوردهای پیشنهادی با استفاده از داده‌های $CD4$ در بخش ۴.۱ می‌پردازیم. داده‌ها با استفاده از سه مدل اثرات آمیخته، نیمه‌پارامتری و نیمه‌پارامتری با اثرات آمیخته تحلیل شدند و برای برآورد جزء ناپارامتری از روش‌های اسپلاین، هسته‌ای و بازگشتی استفاده شد. برای انتخاب پهنای باند در برآوردهای هسته‌ای و بازگشتی، از معیار GCV استفاده شد. بدین

ترتیب برآورد پهنای باند h و h_1 به ترتیب 0.02 و 0.15 به دست آمد. یادآور می‌شویم که h پهنای باند به کار رفته در برآورد $\mu_Y(\cdot)$ و $\mu_X(\cdot)$ و h_1 پهنای باند به کار رفته در برآورد $\sigma^2(\cdot)$ است. هدف این است که ضمن انتخاب بهترین مدل، عملکرد سه برآوردگر معرفی شده مقایسه شوند. جهت رسیدن به این هدف معیارهای SD ، AIC ، BIC و $MSPE$ محاسبه شدند. ضمیمه را ببینید. نتایج حاصل در جدول ۴.۱ ارائه شده است. این نتایج نشان می‌دهد که مدل نیمه‌پارامتری با اثرات آمیخته به طور قابل توجهی دارای معیارهای برازش مدل کمتری نسبت به دو مدل رقیب است. اگرچه دو مدل نیمه‌پارامتری با اثرات آمیخته و مدل نیمه‌پارامتری نزدیک به هم عمل کرده‌اند اما این تفاوت در مدل اثرات آمیخته چشمگیرتر است. برآوردگر بازگشتی همواره مقادیر SD بزرگتری از روش هسته‌ای دارد، بنابراین استفاده از این برآوردگر در تحلیل داده‌های طولی توصیه نمی‌شود. این مورد با نتایج خواص مجانبی برآوردگرها مطابقت دارد. همچنین برآوردگر اسپلاین بهتر از دو روش دیگر عمل کرده است.

جدول ۲.۱: نتایج شبیه‌سازی: مقادیر MSE ، اریبی (SD) در مدل‌های متفاوت به ازای ρ ها و n های مختلف.

پارامتر		$n = 50$			$n = 100$		
		$\rho = 0.3$	$\rho = 0.6$	$\rho = 0.9$	$\rho = 0.3$	$\rho = 0.6$	$\rho = 0.9$
مدل نیمه پارامتری							
AR	β_1	0.2089(0.2561)	0.1630(0.1884)	0.1288(0.1467)	0.1309(0.1456)	0.0922(0.1042)	0.0652(0.0562)
	β_2	0.1461(0.1704)	0.1106(0.1271)	0.0734(0.0644)	0.1466(0.1459)	0.1208(0.1308)	0.1191(0.0680)
	β_3	0.2579(0.3077)	0.1984(0.2343)	0.1930(0.2147)	0.1107(0.1390)	0.0848(0.1048)	0.0556(0.0525)
	MSE	0.2273	0.1286	0.0749	0.0776	0.0481	0.0264
EX	β_1	0.2069(0.2557)	0.1607(0.1886)	0.1259(0.1457)	0.1309(0.1456)	0.0928(0.1068)	0.0651(0.0562)
	β_2	0.1424(0.1670)	0.1159(0.1520)	0.0782(0.1105)	0.1466(0.1459)	0.1204(0.1309)	0.1189(0.0682)
	β_3	0.2588(0.3076)	0.2016(0.2412)	0.1949(0.2162)	0.1107(0.1390)	0.0844(0.1030)	0.0556(0.0533)
	MSE	0.2274	0.1287	0.0775	0.0786	0.0504	0.0263
IND	β_1	0.2163(0.2737)	0.1564(0.1932)	0.1219(0.1422)	0.1247(0.1425)	0.0899(0.1040)	0.0636(0.0596)
	β_2	0.1663(0.2345)	0.1112(0.1390)	0.0670(0.0664)	0.1487(0.1582)	0.1236(0.1287)	0.1105(0.0727)
	β_3	0.2568(0.3061)	0.2160(0.2634)	0.1882(0.2100)	0.1145(0.1408)	0.0886(0.1047)	0.0566(0.0537)
	MSE	0.4950	0.4127	0.4093	0.8439	0.8041	0.8119
مدل اثرات آمیخته							
AR	β_1	0.2040(0.2591)	0.1653(0.2071)	0.1245(0.1442)	0.1284(0.1462)	0.0910(0.1058)	0.0638(0.0591)
	β_2	0.1560(0.2026)	0.1132(0.1386)	0.0671(0.0657)	0.1492(0.1557)	0.1268(0.1313)	0.1106(0.0730)
	β_3	0.2699(0.3266)	0.2095(0.2511)	0.1891(0.2119)	0.1169(0.1445)	0.0908(0.1093)	0.0566(0.0536)
	MSE	0.3401	0.3212	0.3300	0.4466	0.4502	0.4335
EX	β_1	0.2165(0.2736)	0.1562(0.1928)	0.1221(0.1422)	0.1255(0.1431)	0.0939(0.1111)	0.0635(0.0596)
	β_2	0.1590(0.2156)	0.1108(0.1384)	0.0668(0.0661)	0.1533(0.1764)	0.1273(0.1434)	0.1104(0.0727)
	β_3	0.2583(0.3092)	0.2154(0.2631)	0.1877(0.2095)	0.1186(0.1516)	0.0951(0.1314)	0.0566(0.0538)
	MSE	0.3506	0.3240	0.3327	0.4480	0.4505	0.4349
IND	β_1	0.2281(0.1435)	0.1924(0.1084)	0.2039(0.0844)	0.2772(0.0689)	0.2619(0.0517)	0.2580(0.0335)
	β_2	0.3917(0.1535)	0.3599(0.1254)	0.3577(0.1099)	0.6586(0.0803)	0.6349(0.0662)	0.6434(0.0439)
	β_3	0.4813(0.1385)	0.4586(0.1023)	0.4603(0.0940)	0.5637(0.0859)	0.5667(0.0655)	0.5717(0.0378)
	MSE	0.3580	0.3300	0.3351	0.4491	0.4490	0.4287
مدل نیمه پارامتری با اثرات آمیخته							
AR	β_1	0.0685(0.0728)	0.0739(0.0757)	0.0721(0.0728)	0.0492(0.0603)	0.0441(0.0504)	0.0489(0.0586)
	β_2	0.1096(0.0810)	0.1009(0.0814)	0.1054(0.0776)	0.1820(0.0514)	0.1854(0.0502)	0.1804(0.0487)
	β_3	0.0796(0.0807)	0.0745(0.0804)	0.0777(0.0786)	0.0586(0.0545)	0.0582(0.0561)	0.0610(0.0532)
	MSE	0.1975	0.1132	0.0728	0.0757	0.0481	0.0263
EX	β_1	0.0715(0.0743)	0.0751(0.0762)	0.0748(0.0732)	0.0488(0.0604)	0.0443(0.0546)	0.0482(0.0580)
	β_2	0.1089(0.0819)	0.1009(0.0820)	0.1052(0.0783)	0.1823(0.0511)	0.1854(0.0499)	0.1794(0.0483)
	β_3	0.0818(0.0836)	0.0752(0.0808)	0.0781(0.0777)	0.0589(0.0543)	0.0583(0.0556)	0.0617(0.0539)
	MSE	0.1937	0.1227	0.0730	0.0776	0.0481	0.0286
IND	β_1	0.0716(0.0736)	0.0763(0.0767)	0.0748(0.0729)	0.0488(0.0605)	0.0442(0.0544)	0.0480(0.0578)
	β_2	0.1084(0.0822)	0.1012(0.0821)	0.1048(0.0783)	0.1822(0.0515)	0.1849(0.0501)	0.1792(0.0478)
	β_3	0.0824(0.0861)	0.0762(0.0821)	0.0787(0.0772)	0.0594(0.0549)	0.0587(0.0562)	0.0618(0.0540)
	MSE	0.2349	0.1290	0.0860	0.0860	0.0616	0.0286

جدول ۳.۱: نتایج شبیه‌سازی: مقادیر MSE ، اریبی (SD) برای روش‌های متفاوت برآورد به ازای ρ ها و n های مختلف در مدل اثرات آمیخته نیمه پارامتری.

همبستگی	پارامتر	$n = 50$			$n = 100$		
		$\rho = 0.3$	$\rho = 0.6$	$\rho = 0.9$	$\rho = 0.3$	$\rho = 0.6$	$\rho = 0.9$
		روش اسپلین					
AR	β_1	۰/۰۶۸۵(۰/۰۷۲۸)	۰/۰۷۳۹(۰/۰۷۵۷)	۰/۰۷۲۱(۰/۰۷۲۸)	۰/۰۴۹۲(۰/۰۶۰۳)	۰/۰۴۴۱(۰/۰۵۵۴)	۰/۰۴۸۹(۰/۰۵۸۶)
	β_2	۰/۱۰۹۶(۰/۰۸۱۰)	۰/۱۰۰۹(۰/۰۸۱۴)	۰/۱۰۵۴(۰/۰۷۷۶)	۰/۱۸۲۰(۰/۰۵۱۴)	۰/۱۸۵۴(۰/۰۵۰۲)	۰/۱۸۰۴(۰/۰۴۸۷)
	β_3	۰/۰۷۹۶(۰/۰۸۰۷)	۰/۰۷۴۵(۰/۰۸۰۴)	۰/۰۷۷۷(۰/۰۷۸۶)	۰/۰۵۸۶(۰/۰۵۴۵)	۰/۰۵۸۲(۰/۰۵۶۱)	۰/۰۶۱۰(۰/۰۵۳۲)
	MSE	۰/۱۹۷۵	۰/۱۱۳۲	۰/۰۷۲۸	۰/۰۷۵۷	۰/۰۴۸۱	۰/۰۲۶۳
EX	β_1	۰/۰۷۱۵(۰/۰۷۴۳)	۰/۰۷۵۱(۰/۰۷۶۲)	۰/۰۷۴۸(۰/۰۷۳۲)	۰/۰۴۸۸(۰/۰۶۰۴)	۰/۰۴۴۳(۰/۰۵۴۶)	۰/۰۴۸۲(۰/۰۵۸۰)
	β_2	۰/۱۰۸۹(۰/۰۸۱۹)	۰/۱۰۰۹(۰/۰۸۲۰)	۰/۱۰۵۲(۰/۰۷۸۳)	۰/۱۸۲۳(۰/۰۵۱۱)	۰/۱۸۵۴(۰/۰۴۹۹)	۰/۱۷۹۴(۰/۰۴۸۳)
	β_3	۰/۰۸۱۸(۰/۰۸۳۶)	۰/۰۷۵۲(۰/۰۸۰۸)	۰/۰۷۸۱(۰/۰۷۷۷)	۰/۰۵۸۹(۰/۰۵۴۳)	۰/۰۵۸۳(۰/۰۵۵۶)	۰/۰۶۱۷(۰/۰۵۲۹)
	MSE	۰/۱۹۳۷	۰/۱۲۲۷	۰/۰۷۳۰	۰/۰۷۷۶	۰/۰۴۸۱	۰/۰۲۸۶
IND	β_1	۰/۰۷۱۶(۰/۰۷۳۶)	۰/۰۷۶۳(۰/۰۷۶۷)	۰/۰۷۴۸(۰/۰۷۲۹)	۰/۰۴۸۸(۰/۰۶۰۵)	۰/۰۴۴۲(۰/۰۵۴۴)	۰/۰۴۸۰(۰/۰۵۷۸)
	β_2	۰/۱۰۸۴(۰/۰۸۲۲)	۰/۱۰۱۲(۰/۰۸۲۱)	۰/۱۰۴۸(۰/۰۷۸۳)	۰/۱۸۲۲(۰/۰۵۱۵)	۰/۱۸۴۹(۰/۰۵۰۱)	۰/۱۷۹۲(۰/۰۴۷۸)
	β_3	۰/۰۸۳۴(۰/۰۸۶۱)	۰/۰۷۶۲(۰/۰۸۲۱)	۰/۰۷۸۷(۰/۰۷۷۲)	۰/۰۵۹۴(۰/۰۵۴۹)	۰/۰۵۸۷(۰/۰۵۶۲)	۰/۰۶۱۸(۰/۰۵۴۰)
	MSE	۰/۲۳۴۹	۰/۱۲۹۰	۰/۰۸۶۰	۰/۰۸۶۰	۰/۰۶۱۶	۰/۰۲۸۶
روش هستای							
AR	β_1	۰/۰۴۴(۰/۰۲۵۰۶)	۰/۱۶۰۰(۰/۱۸۶۹)	۰/۱۲۵۱(۰/۱۴۴۳)	۰/۱۳۰۰(۰/۱۴۲۲)	۰/۰۹۲۸(۰/۱۰۶۳)	۰/۰۶۵۱(۰/۰۵۶۲)
	β_2	۰/۱۴۳۴(۰/۱۶۵۷)	۰/۱۱۱۱(۰/۱۲۸۵)	۰/۰۷۱۱(۰/۰۶۲۳)	۰/۱۴۴۴(۰/۱۴۵۴)	۰/۱۱۹۱(۰/۱۲۹۰)	۰/۱۱۷۸(۰/۰۶۸۷)
	β_3	۰/۲۵۴۸(۰/۰۳۰۵۰)	۰/۱۹۴۰(۰/۰۲۳۰۲)	۰/۱۹۰۷(۰/۰۲۱۱۸)	۰/۱۰۷۷(۰/۱۳۲۳)	۰/۰۸۳۴(۰/۰۱۰۲۳)	۰/۰۵۵۴(۰/۰۵۲۲)
	MSE	۰/۱۹۰۱	۰/۱۱۰۳	۰/۰۷۴۸	۰/۰۷۴۸	۰/۰۴۷۳	۰/۰۲۸۳
EX	β_1	۰/۰۸۹(۰/۰۲۵۷۳)	۰/۱۶۳۴(۰/۱۹۴۱)	۰/۱۲۵۸(۰/۱۴۵۶)	۰/۱۲۹۱(۰/۱۴۲۲)	۰/۰۹۲۷(۰/۱۰۶۶)	۰/۰۶۹۴(۰/۰۵۵۸)
	β_2	۰/۱۴۶۷(۰/۱۸۵۱)	۰/۱۱۱۳(۰/۱۲۸۸)	۰/۰۷۲۱(۰/۰۶۳۱)	۰/۱۴۶۱(۰/۱۴۶۶)	۰/۱۲۰۳(۰/۱۳۰۶)	۰/۱۱۹۵(۰/۰۶۷۸)
	β_3	۰/۲۶۶۳(۰/۰۳۲۲۷)	۰/۱۹۸۷(۰/۰۲۳۶۰)	۰/۱۹۵۸(۰/۰۲۱۷۲)	۰/۱۰۷۱(۰/۱۳۳۳)	۰/۰۸۴۴(۰/۰۱۰۳۱)	۰/۰۵۵۴(۰/۰۵۲۳)
	MSE	۰/۲۱۱۵	۰/۱۱۵۲	۰/۰۷۷۵	۰/۰۷۵۱	۰/۰۴۸۱	۰/۰۲۸۶
IND	β_1	۰/۰۵۸(۰/۰۲۵۶۲)	۰/۱۸۲۴(۰/۰۳۰۳۹)	۰/۱۶۳۷(۰/۱۵۵۰)	۰/۴۱۳۱(۰/۰۵۰۲۴)	۰/۴۳۳۰(۰/۰۵۷۲۰)	۰/۴۲۳۷(۰/۰۵۳۶۲)
	β_2	۰/۲۸۹۵(۰/۰۳۷۸۲)	۰/۲۴۴۵(۰/۰۳۲۸۶)	۰/۱۳۰۵(۰/۰۳۱۶۶۳)	۰/۳۶۳۷(۰/۰۵۳۴۴)	۰/۴۱۱۷(۰/۰۶۱۸)	۰/۲۰۳۰(۰/۰۲۸۴۹)
	β_3	۰/۳۴۵۷(۰/۰۴۰۸۷)	۰/۲۹۷۷(۰/۰۳۵۴۰)	۰/۲۲۵۰(۰/۰۲۵۸۶)	۰/۴۲۷۱(۰/۰۵۵۷۸)	۰/۵۲۶۵(۰/۰۹۲۹۲)	۰/۴۹۸۱(۰/۰۵۴۶)
	MSE	۰/۴۶۴۶	۰/۳۶۶۰	۰/۱۵۴۷	۰/۹۳۰۴	۲/۴۵۸۸	۱/۵۹۷۴
روش بازگشتی							
AR	β_1	۰/۲۳۲۴(۰/۰۳۰۷۲)	۰/۱۹۹۸(۰/۰۲۴۸۲)	۰/۱۴۹۰(۰/۰۱۵۶۵)	۰/۲۵۴۵(۰/۰۳۳۶۲)	۰/۲۰۶۳(۰/۰۲۸۵۸)	۰/۱۵۴۴(۰/۰۲۱۹۸)
	β_2	۰/۲۴۴۲(۰/۰۳۰۸)	۰/۱۸۵۵(۰/۰۲۷۴۴)	۰/۱۲۴۹(۰/۰۱۹۴۳)	۰/۲۸۹۱(۰/۰۴۰۵۲)	۰/۲۱۰۳(۰/۰۳۱۰۸)	۰/۱۳۷۰(۰/۰۱۵۴۲)
	β_3	۰/۳۰۵۳(۰/۰۳۷۱۹)	۰/۲۴۵۴(۰/۰۳۱۵۹)	۰/۲۳۸۸(۰/۰۲۸۳۸)	۰/۳۰۵۶(۰/۰۴۴۰۵)	۰/۲۶۳۵(۰/۰۴۱۴۶)	۰/۱۷۶۳(۰/۰۳۳۵۹)
	MSE	۰/۳۹۱۴	۰/۲۷۱۳	۰/۱۶۷۰	۰/۵۰۶۹	۰/۳۷۰۸	۰/۲۰۱۶
EX	β_1	۰/۲۰۶۰(۰/۰۲۵۵۹)	۰/۱۹۱۰(۰/۰۲۴۲۷)	۰/۱۶۵۷(۰/۰۱۵۸۷)	۰/۴۰۴۶(۰/۰۴۹۶۰)	۰/۴۲۶۴(۰/۰۵۰۷۶)	۰/۴۰۹۳(۰/۰۴۱۰۲)
	β_2	۰/۲۸۳۶(۰/۰۳۷۱۵)	۰/۲۳۴۷(۰/۰۳۱۵۲)	۰/۱۲۶۲(۰/۰۱۵۲۹)	۰/۳۴۰۹(۰/۰۴۸۳۶)	۰/۳۳۵۱(۰/۰۵۰۷۰)	۰/۲۲۶۸(۰/۰۴۴۴۵)
	β_3	۰/۳۴۲۳(۰/۰۴۰۳۳)	۰/۲۹۱۷(۰/۰۳۴۴۶)	۰/۲۲۴۳(۰/۰۲۵۶۵)	۰/۴۴۰۶(۰/۰۵۸۷۳)	۰/۴۹۲۵(۰/۰۷۲۵۴)	۰/۴۶۳۸(۰/۰۷۳۰۱)
	MSE	۰/۴۵۲۰	۰/۳۶۵۶	۰/۱۴۷۷	۰/۹۰۲۷	۱/۱۳۴۹	۱/۰۲۶۳
IND	β_1	۰/۲۳۶۲(۰/۰۱۴۷۰)	۰/۱۹۲۹(۰/۰۱۰۶۷)	۰/۲۰۳۲(۰/۰۸۱۲)	۰/۲۷۳۴(۰/۰۶۶۷۶)	۰/۲۶۶۱(۰/۰۵۵۲۲)	۰/۲۵۸۶(۰/۰۲۸۲)
	β_2	۰/۴۰۲۹(۰/۰۱۶۹۱)	۰/۳۵۹۴(۰/۰۱۱۸۱)	۰/۳۵۲۷(۰/۰۸۵۵)	۰/۶۶۴۱(۰/۰۷۸۹)	۰/۶۴۳۲(۰/۰۶۵۱)	۰/۶۴۵۰(۰/۰۳۴۳)
	β_3	۰/۴۹۲۵(۰/۰۱۵۳۷)	۰/۴۵۸۶(۰/۰۹۶۴)	۰/۴۵۷۴(۰/۰۸۳۵)	۰/۵۶۹۸(۰/۰۹۰۷)	۰/۵۷۳۹(۰/۰۶۷۵)	۰/۵۷۳۱(۰/۰۳۱۰)
	MSE	۰/۵۲۹۰	۰/۴۰۹۱	۰/۳۹۵۶	۰/۸۵۹۴	۰/۸۲۵۳	۰/۸۱۴۲

جدول ۴.۱: نتایج تحلیل داده‌های $CD4$: برآورد پارامترهای مدل رگرسیونی (SD) تحت برازش مدل‌های مختلف.

	مدل				
	نیمه‌پارامتری با اثرات آمیخته			نیمه‌پارامتری	اثرات آمیخته
	اسپالین	هسته‌ای	بازگشتی	اسپالین	
<i>AGE</i>	۰/۰۸۵۱(۰/۰۰۰۲)	۰/۰۲۱۲(۰/۰۰۲۱)	۰/۰۳۲۹(۰/۰۰۲۱)	۰/۱۵۵۱(۰/۰۵۷۸)	۰/۴۹۱۹(۰/۰۲۳۲)
<i>SMOKE</i>	۰/۷۲۹۰(۰/۰۰۶۰)	۰/۵۹۴۱(۰/۰۰۸۲)	۰/۶۱۱(۰/۰۰۵۱)	۰/۱۵۷۰(۰/۱۷۲۴)	۲/۹۱۴۱(۰/۱۰۳۹)
<i>DRUG</i>	۰/۸۰۱۱(۰/۰۱۴۳)	۰/۵۳۵۰(۰/۰۱۹۴)	۰/۵۴۴(۰/۰۵۲۱)	۰/۱۱۰(۰/۰۵۹۰)	۱/۴۴۵۱(۰/۲۸۴۴)
<i>SEXP</i>	۰/۰۷۱۱(۰/۰۰۱۴)	۰/۰۵۸۵(۰/۰۰۲۴)	۰/۰۶۴۰(۰/۰۰۴۶)	۰/۱۸۹(۰/۰۰۴۴)	۰/۹۷۱۲(۰/۰۲۰۰)
<i>CESD</i>	-۰/۰۵۴۲(۰/۰۰۰۳)	-۰/۰۵۳۱(۰/۰۰۰۹)	-۰/۰۶۲۲(۰/۰۰۲۲)	-۰/۰۱۷۴(۰/۰۰۳۵)	-۰/۰۳۸۴(۰/۰۰۶۹)
<i>AIC</i>	۳۱۳۷/۴۶	۳۲۹۱/۱۴۱	۳۳۱۹/۳۰۵	۳۴۹۱/۱۶۴	۱۰۱۸۷/۰۶
<i>BIC</i>	۳۱۹۶/۱۸	۳۳۱۰/۶۲۷	۳۳۳۸/۷۹۱	۳۵۱۰/۶۴۹	۱۰۲۰۶/۵۵
<i>MSPE</i>	۹۷/۵۲۱۲	۹۸۸۱۹/۴۵	۹۸۳۶۵/۷۲	۱۳۶۰/۲۳۹	۲۰۴/۸۱۱

فصل ۲

انتخاب متغیر برای داده‌های طولی نرمال

۱.۲ انتخاب متغیر در مدل‌های رگرسیونی

مدل رگرسیونی چندگانه

$$y = \beta_0 + \beta_1 x_1 + \dots + \beta_p x_p + \varepsilon = \mathbf{x}^T \boldsymbol{\beta} + \varepsilon \quad (1.2)$$

را در نظر بگیرید، که در آن $\mathbf{x} = (1, x_1, \dots, x_p)^T$ بردار متغیرهای تبیینی، $\boldsymbol{\beta} = (\beta_0, \beta_1, \dots, \beta_p)^T$ بردار ضرایب رگرسیونی و ε خطای تصادفی با میانگین صفر و واریانس ثابت است. در این رساله بدون از دست دادن کلیت، فرض می‌شود عرض از مبدا برابر با صفر است. روش‌های متعددی برای برآورد پارامترهای مدل وجود دارند، که برآورد ماکزیمم درست‌نمایی (MLE) و کمترین توان‌های دوم معمولی (OLS) از جمله معمول‌ترین این روش‌ها هستند. اگر داده‌ها دارای توزیع نرمال باشند این دو روش معادل‌اند. هنگامی که بین متغیرهای تبیینی هم‌خطی وجود دارد، استنباط بر اساس این روش می‌تواند گمراه‌کننده باشد. زیرا وجود هم‌خطی، برآوردهایی با واریانس بزرگ را نتیجه می‌دهد که منجر به کاهش دقت پیشگویی می‌شود. تاکنون روش‌های متعددی برای مقابله با مشکلات هم‌خطی بین متغیرهای تبیینی معرفی شده است. برای اطلاعات بیشتر در این باره به مونت‌گومری و همکاران (۲۰۱۲) و چاترجی و هادی (۲۰۱۲) مراجعه کنید. یک روش متداول برای مقابله با هم‌خطی، استفاده از برآوردهای اریب مانند ریدج^۱ است که توسط هورل و کنارد

^۱Ridge

(۱۹۷۰) معرفی شده است. یکی از عوامل ایجاد هم‌خطی، زیاد بودن تعداد متغیرهای تبیینی در مدل است. در مدل‌های با بعد بالا مسئله هم‌خطی بسیار شدید است، یعنی هر یک از متغیرهای تبیینی را می‌توان به صورت ترکیب خطی از سایر متغیرها نوشت (جیمز و همکاران، ۲۰۱۳). در این حالت روش پیشنهادی، انتخاب زیرمجموعه‌ای از متغیرهای تبیینی برای ورود به مدل است. در این رهیافت علاقه‌مندیم زیرمجموعه‌ای از متغیرها را طوری انتخاب کنیم که رابطه بین متغیر پاسخ و متغیرهای تبیینی را به خوبی توصیف کند. ساختن یک مدل رگرسیونی که تنها شامل زیرمجموعه‌ای از متغیرهای تبیینی است، انتخاب متغیر^۲ نامیده می‌شود. مسئله انتخاب متغیر دو واقعیت ناسازگار را دربر دارد. اولاً علاقه‌مندیم که مدل تا حد امکان متغیرهای تبیینی بیشتری را شامل شود، به نحوی که اطلاعات موجود در این متغیرها بتواند در پیشگویی متغیر پاسخ مؤثر باشد. دوماً می‌خواهیم مدل حتی‌الامکان دارای متغیرهای رگرسیونی کمتری باشد، چون با افزایش تعداد متغیرهای تبیینی در مدل، هزینه محاسبات افزایش می‌یابد و تفسیر مدل دشوارتر می‌شود. علاوه بر این، وجود تعداد زیادی متغیر تبیینی ممکن است باعث بیش‌برازش^۳ شود، یعنی مدل برازش‌شده برای داده‌های آموزشی^۴ که برای ساخت مدل از آن‌ها استفاده شده است، مناسب باشد اما برای پیشگویی مشاهدات آینده بسیار ضعیف باشد. مسئله انتخاب متغیر امری بسیار مهم در تحلیل داده‌هاست. فرآیند به‌دست آوردن یک مدل که سازگاری و توافق بین این دو واقعیت است، انتخاب بهترین معادله رگرسیونی نامیده می‌شود. بدیهی است برای تعیین مدل پیشگو، لازم است پارامترهای مدل برآورد شوند. لذا فرآیندی که علاوه بر انتخاب متغیر و بهترین معادله رگرسیونی، برآورد ضرایب را نیز انجام دهد، موردنظر است. یک روش کارا برای برآورد پارامترها که هم‌زمان انتخاب متغیر را نیز انجام می‌دهد، روش مبتنی بر تابع تاوان است. این روش با اضافه کردن جمله‌ای منظم‌ساز به معادلات برآورد، رویکرد نوینی تحت عنوان رگرسیون تاوانیده را گسترش داد که تاکنون توسط محققین زیادی مورد مطالعه قرار گرفته است. در این میان می‌توان روش‌های رگرسیون بریج^۵ (فرانک و فریدمن، ۱۹۹۳)، لاسو^۶ (تیشیرانی، ۱۹۹۶)، لاسوی سازوار^۷ (زو، ۲۰۰۶)، الاستیک نت^۸ (زو و هستی، ۲۰۰۵)، اسکد^۹ (فن و لی، ۲۰۰۱) و MCP (ژانگ، ۲۰۱۰) را نام برد.

امروزه توسعه و گسترش سریع علوم در زمینه‌های مختلف از جمله ژنومیک، علوم اعصاب، علوم زمین، مالی و غیره، افزایش سرعت ثبت و تحلیل اطلاعات توسط ابررایانه‌ها، آماردانان را با حجم عظیمی از اطلاعات پیچیده در تحلیل مسائل مواجه نموده است. اگرچه محققین در زمینه‌های مختلف اهداف متفاوتی را دنبال می‌کنند، اما با این حال یک موضوع مشترک دارند؛ تحقیق آن‌ها به شدت به استخراج اطلاعات مفید از داده‌های حجیم وابسته است و تعداد متغیرهای تبیینی یعنی بعد متغیرها

^۲Variable selection

^۳Overfitting

^۴Training data set

^۵Bridge

^۶Lasso

^۷Adaptive lasso

^۸Elastic-net

^۹SCAD

در مقایسه با حجم نمونه می‌تواند بزرگ باشد. به‌طور کلی داده‌ها را از نظر بعد می‌توان به دو گروه بعد پایین^{۱۰} و بعد بالا^{۱۱} تقسیم نمود. در داده‌های با بعد پایین، تعداد متغیرها، p ، بسیار کوچک‌تر از تعداد مشاهدات، n ، است. به عنوان مثال مجموعه داده با حجم $n = 10^6$ و بعد $p = 10$ مجموعه داده با بعد پایین است. به‌علاوه در این داده‌ها برای نظریه مجانبی فرض می‌شود که با افزایش n ، مقدار p ثابت می‌ماند. در داده‌های با بعد بالا هم‌چنان p کوچک‌تر از n است ولی مقدار آن بزرگ است یا p از n بزرگ‌تر است. به عنوان مثال $n = 10^6$ و $p = 200$. در این نوع داده‌ها با افزایش n ، مقدار p می‌تواند افزایش یابد. مسئله ”کوچک n و بزرگ p “ چالش‌ها و موهبت‌های فراوانی را پیش‌روی آماردانان قرار داده است و تحلیل این نوع داده‌ها مستلزم به‌کارگیری روش‌های نوین آماری است. یک پذیره اساسی که استنباط پیرامون این نوع داده‌ها را امکان‌پذیر می‌سازد، پذیره تنکی^{۱۲} است که براساس آن، فقط مجموعه کوچکی از متغیرهای تبیینی بر متغیر پاسخ تاثیرگذارند و بقیه متغیرهای تبیینی ارتباطی با متغیر پاسخ ندارند. بنابراین کاهش بعد و انتخاب متغیر نقش مهمی را در تحلیل این نوع داده‌ها ایفا می‌کند. تاکنون روش‌های آماری مختلفی برای انتخاب متغیر در داده‌های با بعد بالا معرفی شده است که از جمله می‌توان لاسو، لاسوی سازوار، اسکد، الاستیک نت، گاروت نامنفی و انتخابگر دن‌زیک^{۱۳} را نام برد. در حالت خاص، ممکن است p با نرخ نمایی رشد کند، به عبارت دیگر رابطه $\log(p) = O(n^a)$ برای برخی مقادیر مثبت a برقرار باشد. داده‌های بیان‌ژنی، ریزآرایه‌ها، داده‌های تابعی، داده‌های مالی، داده‌های پردازش تصویر، داده‌های تصویربرداری پزشکی و طبقه‌بندی تومور نمونه‌هایی از این نوع داده‌ها هستند. به عنوان مثال در مطالعه رابطه بین ویژگی‌های ظاهری^{۱۴} و ویژگی‌های ژنتیکی^{۱۵} از میلیون‌ها SNP ^{۱۶} استفاده می‌کنند، یا در طبقه‌بندی نوعی تومور از داده‌های ریزآرایه استفاده می‌کنند که شامل هزاران عامل ژنتیکی است و هدف شناسایی ژن‌های موثر بر بیماری است. در این مثال‌ها تعداد متغیرهای تبیینی بسیار بزرگ‌تر از تعداد مشاهدات است. برای تمایز این نوع داده‌ها از داده‌های با بعد بالا با نرخ رشد چندجمله‌ای، آن‌ها را داده‌های با بعد بسیار بالا^{۱۷} نامند. لذا عبارت بعد بالا برای نرخ رشد چندجمله‌ای ($p = O(n^b)$) و عبارت بعد بسیار بالا برای نرخ رشد نمایی ($p = \exp(O(n^a))$) استفاده می‌شود. برای درک بهتر مطلب دو حالت $p = n^b$ و $p = \exp(n^a)$ را به ترتیب برای بعد بالا و بعد بسیار بالا در نظر گرفته و نسبت $\frac{p}{n}$ را برای مقادیر مختلف (n, a, b) در جدول ۱۰۲ گزارش کرده‌ایم. با توجه به نتایج جدول مذکور ملاحظه می‌شود که در بعد بسیار بالا، با افزایش n مقدار p به سرعت رشد می‌کند و به سمت بینهایت می‌رود. لازم به ذکر است که تمرکز این رساله روی داده‌ها با بعد بالا است نه بعد بسیار بالا. برای آگاهی بیشتر درباره داده‌های با بعد بالا و چالش‌های تحلیل این نوع داده‌ها فن و لی (۲۰۰۶)، فن و لیو

¹⁰Low dimension¹¹High dimension¹²Sparsity¹³Dantzig selector¹⁴Phenotypes¹⁵Genotypes¹⁶Single-nucleotide polymorphism¹⁷Ultrahigh-dimensional data

(۲۰۰۸)، فن و همکاران (۲۰۰۹)، فن و همکاران (۲۰۱۴)، جان‌استون و تیتزینگتون (۲۰۰۹)، بولمن و فن‌دگیر (۲۰۱۱) و گیراد (۲۰۱۴) مراجعه کنید.

جدول ۱۰.۲: مقادیر تقریبی $\frac{p}{n}$ برای داده‌های با ابعاد بالا و بسیار بالا به ازای مقادیر مختلف n ، a و b .

n	بعد بالا مقادیر مختلف b				بعد بسیار بالا مقادیر مختلف a			
	۱	۱.۵	۲	۲.۵	۳	۵	۷	۱
۱۰	$۱/۰ \times ۱/۰۱$	$۳/۲ \times ۱/۰۱$	$۱/۰ \times ۱/۰۲$	$۳/۲ \times ۱/۰۲$	$۱/۱ \times ۱/۰۱$	$۶/۸ \times ۱/۰۱$	$۱/۸ \times ۱/۰۳$	$۲/۲ \times ۱/۰۴$
۳۰	$۳/۰ \times ۱/۰۱$	$۱/۶ \times ۱/۰۲$	$۹/۰ \times ۱/۰۲$	$۴/۹ \times ۱/۰۳$	$۳/۶ \times ۱/۰۱$	$۴/۴ \times ۱/۰۳$	$۳/۳ \times ۱/۰۸$	$۱/۱ \times ۱/۰۱۳$
۱۰۰	$۱/۰ \times ۱/۰۲$	$۱/۰ \times ۱/۰۳$	$۱/۰ \times ۱/۰۴$	$۱/۰ \times ۱/۰۵$	$۲/۸ \times ۱/۰۲$	$۵/۳ \times ۱/۰۷$	$۲/۶ \times ۱/۰۲۴$	$۲/۷ \times ۱/۰۴۳$
۲۰۰	$۲/۰ \times ۱/۰۲$	$۲/۸ \times ۱/۰۳$	$۴/۰ \times ۱/۰۴$	$۵/۶ \times ۱/۰۵$	$۱/۵ \times ۱/۰۳$	$۸/۱ \times ۱/۰۱۱$	$۶/۲ \times ۱/۰۴۴$	$۷/۲ \times ۱/۰۸۶$
۵۰۰	$۵/۰ \times ۱/۰۲$	$۱/۱ \times ۱/۰۴$	$۲/۵ \times ۱/۰۵$	$۵/۶ \times ۱/۰۶$	$۲/۹ \times ۱/۰۴$	$۱/۳ \times ۱/۰۲۱$	$۷/۲ \times ۱/۰۹۹$	$۱/۴ \times ۱/۰۲۱۷$

در بخش بعد، ابتدا تاثیر وجود هم‌خطی بر عملکرد برآوردگرها و برآوردگر ريج، به عنوان ابزاری برای مقابله با مشکل هم‌خطی بررسی شده و سپس برآوردگر ريج برای داده‌های طولی در مدل نیمه‌پارامتری با اثرات آمیخته معرفی می‌شود. همچنین با معرفی رایج‌ترین روش‌های انتخاب متغیر به مسئله برآورد و انتخاب متغیر همزمان در مدل مذکور پرداخته است.

۱.۱.۲ مسئله هم‌خطی و برآوردگر ريج

فرض کنید $\{y_i, x_i\}_{i=1}^n$ نمونه‌ای تصادفی از مدل (۱.۲) باشد. بنابراین مدل رگرسیون خطی چندگانه را می‌توان با نماد ماتریسی $y = X\beta + \varepsilon$ نوشت، که در آن $y = (y_1, \dots, y_n)^T$ بردار متغیر پاسخ، $X = (x_1^T, \dots, x_n^T)^T$ ماتریس طرح $n \times p$ بعدی با رتبه کامل ستونی، β بردار ضرایب رگرسیونی و $\varepsilon = (\varepsilon_1, \dots, \varepsilon_n)^T$ بردار خطاست. برآوردگر OLS بردار ضرایب به صورت

$$\hat{\beta}_{OLS} = \arg \min_{\beta \in \mathbb{R}^p} \|y - X\beta\|^2 = (X^T X)^{-1} X^T y$$

است. یک مسئله جدی که کاربرد برآوردگرهای معمول را با مشکل مواجه می‌کند وجود هم‌خطی چندگانه^{۱۸} بین متغیرهای تبیینی است، بدین معنی که متغیرهای تبیینی به صورت کامل یا تقریباً کامل وابسته خطی باشند. اگر یکی از متغیرهای تبیینی یک تابع دقیق خطی از یک یا چند متغیر تبیینی باشد، گوئیم رگرسیون دارای هم‌خطی کامل است. هم‌خطی ناقص هنگامی اتفاق می‌افتد که یکی از متغیرهای تبیینی به طور تقریبی یک تابع خطی از یک یا چند متغیر تبیینی دیگر باشد. به عبارت دیگر هم‌خطی کامل هنگامی رخ می‌دهد که حداقل به ازای یک $k = 1, \dots, p$ ، $R_k^2 \approx 1$ باشیم و هم‌خطی ناقص زمانی پیش می‌آید که حداقل به ازای یک k داشته باشیم $R_k^2 \approx 1$ ، به طوری که R_k^2 ضریب تعیین متغیر x_k بر سایر متغیرهای تبیینی در رگرسیون خطی است.

¹⁸Multicollinearity

وجود هم‌خطی تأثیرات جدی بر برآوردها دارد. مثلاً در برآوردگر OLS منجر به برآوردهایی با واریانس بزرگ می‌شود که دقت پیشگویی را کاهش می‌دهد. برای درک بهتر موضوع، مدل رگرسیون خطی با دو متغیر تبیینی را در نظر بگیرید.

$$y = \beta_1 x_1 + \beta_2 x_2 + \varepsilon$$

فرض کنید x_1, x_2 و y استاندارد شده‌اند، به‌طوری‌که $\bar{x}_k = 0$ ، $\sum_{i=1}^n (x_{ik} - \bar{x}_k)^2 = 1$ ، $\bar{y} = 0$ و $\sum_{i=1}^n (y_i - \bar{y})^2 = 1$. در این صورت معادلات نرمال روش OLS به‌صورت

$$(X^T X) \hat{\beta}_{OLS} = X^T y$$

یا به‌طور معادل به‌صورت

$$\begin{bmatrix} 1 & r_{12} \\ r_{12} & 1 \end{bmatrix} \begin{bmatrix} \hat{\beta}_1 \\ \hat{\beta}_2 \end{bmatrix} = \begin{bmatrix} r_{1y} \\ r_{2y} \end{bmatrix}$$

است که r_{12} همبستگی ساده بین x_1 و x_2 ؛ و r_{ky} همبستگی ساده بین y و x_k به ازای $k = 1, 2$ هستند. معکوس $X^T X$ عبارت است از

$$C = (X^T X)^{-1} = \begin{bmatrix} \frac{1}{1-r_{12}^2} & \frac{-r_{12}}{1-r_{12}^2} \\ \frac{-r_{12}}{1-r_{12}^2} & \frac{1}{1-r_{12}^2} \end{bmatrix} \quad (2.2)$$

و برآوردهای ضرایب رگرسیونی به‌صورت

$$\hat{\beta}_1 = \frac{r_{1y} - r_{12}r_{2y}}{1 - r_{12}^2}, \quad \hat{\beta}_2 = \frac{r_{2y} - r_{12}r_{1y}}{1 - r_{12}^2}$$

است. اگر بین x_1 و x_2 هم‌خطی ناقص وجود داشته باشد، آنگاه ضریب همبستگی r_{12} بزرگ خواهد بود. از رابطه (۲.۲) می‌توان دریافت که اگر $|r_{12}| \rightarrow 1$ ، آنگاه به ازای $k = 1, 2$ خواهیم داشت $\text{Var}(\hat{\beta}_k) = C_{kk}\sigma^2 \rightarrow \infty$ و $\text{Cov}(\hat{\beta}_1, \hat{\beta}_2) = C_{12}\sigma^2 \rightarrow \infty$ که مولفه k ام روی قطر اصلی ماتریس C است. بنابراین هم‌خطی زیاد بین x_1 و x_2 منجر به واریانس‌ها و کوواریانس‌های بزرگ برای برآوردهای OLS ضرایب رگرسیونی خواهد شد. برای بیش از دو متغیر تبیینی نیز هم‌خطی اثرات مشابهی ایجاد می‌کند. می‌توان نشان داد که اعضای قطر اصلی ماتریس عبارتند از

$$C_{kk} = \frac{1}{1 - R_k^2}, \quad k = 1, \dots, p,$$

که در آن R_k^2 ضریب تعیین حاصل از رگرسیون x_j بر $p - 1$ متغیر تبیینی دیگر است. اگر هم‌خطی شدید بین x_k و هر زیرمجموعه‌ای از $p - 1$ متغیر دیگر وجود داشته باشد، در این صورت مقدار R_k^2 نزدیک به یک خواهد بود. چون $\text{Var}(\hat{\beta}_k) = C_{kk}\sigma^2 = \sigma^2(1 - R_k^2)^{-1}$ ، هم‌خطی شدید موجب می‌شود که واریانس برآوردهای ضرایب رگرسیون بسیار بزرگ شود. در این صورت $\hat{\beta}_{OLS}$ برآورد دقیقی از β

نیست. به عبارت دیگر، دقت برآورد ضرایب رگرسیونی بسیار کم می‌شود. همچنین هم‌خطی باعث می‌شود که بردار $\hat{\beta}_{OLS}$ دارای اعضایی باشد که از لحاظ قدر مطلق بسیار بزرگ هستند. علاوه بر این، $\hat{\beta}_{OLS}$ ناپایدار است، یعنی با حذف یا اضافه کردن یک متغیر تبیینی یا با یک تغییر کوچک در مقادیر متغیرهای تبیینی، مقادیر $\hat{\beta}_{OLS}$ به‌طور قابل توجهی تغییر می‌کند. در صورت وجود هم‌خطی کامل، ماتریس وارون‌پذیر نیست، زیرا در این حالت پرتبه نیست. بنابراین جواب یکتا برای برآورد ضرایب وجود ندارد. لازم به ذکر است که هم‌خطی کامل معمولاً در عمل رخ نمی‌دهد، زیرا همواره با خطای اندازه‌گیری مواجه هستیم. در عمل هم‌خطی کامل هنگامی رخ می‌دهد که کاربر به اشتباه در مدل رگرسیونی یکی از متغیرها را تابعی از متغیر دیگر قرار دهد، یا وقتی تعداد متغیرهای تبیینی بیشتر از تعداد مشاهدات باشد. برای اطلاعات بیشتر درباره منابع پیدایش هم‌خطی به روزه (۱۳۹۰) مراجعه کنید. بنابراین لازم است بدانیم که چگونه وجود هم‌خطی را کشف کنیم تا مواظب پیامدهای ممکن باشیم. روش‌های مختلفی برای تشخیص هم‌خطی وجود دارند. در این جا به برخی از این روش‌ها اشاره می‌کنیم.

۱. عامل تورم واریانس (VIF): یکی از روش‌های معروف برای تشخیص هم‌خطی استفاده از کمیت VIF است که به‌صورت

$$VIF_k = \frac{1}{1 - R_k^2}, \quad k = 1, \dots, p$$

محاسبه می‌شود. واضح است که اگر x_k رابطه خطی قوی با سایر متغیرهای تبیینی داشته باشد، آنگاه R_k^2 نزدیک به یک بوده و VIF_k بزرگ می‌شود. اگر x_k بر $p - 1$ متغیر تبیینی دیگر عمود باشد، آنگاه R_k^2 صفر بوده و VIF_k برابر یک خواهد شد. بنابراین انحراف VIF_k از مقدار یک، انحراف از متعامد بودن و گرایش به هم‌خطی را نشان می‌دهد. اگر هر یک از مقادیر VIF بیش‌تر از 10 باشد، هم‌خطی بین متغیرهای تبیینی وجود دارد و اگر هر یک از مقادیر آن بیش‌تر از 30 باشد، هم‌خطی بسیار شدید است و ممکن است موجب بروز مشکلات در برآورد شود (گجراتی، ۲۰۰۴). چون $\text{Var}(\hat{\beta}_k) = VIF_k \sigma^2$ ، می‌توان VIF_k را به عنوان عاملی دانست که به وسیله آن واریانس $\hat{\beta}_k$ به دلیل وابستگی خطی بین متغیرهای تبیینی افزایش یافته است. به عبارت دیگر، VIF_k میزان افزایش در واریانس $\hat{\beta}_k$ به واسطه ارتباط خطی x_k با بقیه متغیرهای تبیینی را نسبت به واریانس آن در صورت ناهمبسته بودن متغیر x_k با سایر متغیرهای تبیینی اندازه می‌گیرد. این مطلب دلیل نام‌گذاری این کمیت خاص را نشان می‌دهد.

۲. عدد شرطی^{۱۹}: مقادیر ویژه ماتریس $X^T X$ می‌توانند برای اندازه‌گیری میزان هم‌خطی مورد استفاده قرار گیرند. یک روش متداول برای کشف وجود هم‌خطی محاسبه عدد شرطی به‌صورت

$$\kappa = \sqrt{\frac{\max_{1 \leq k \leq p} \lambda_k}{\min_{1 \leq k \leq p} \lambda_k}}$$

¹⁹Condition number

است که در آن $\lambda_1, \dots, \lambda_p$ مقادیر ویژه ماتریس $X^T X$ هستند. بلسلی و همکاران (۱۹۸۰) پیشنهاد کردند که اگر عدد شرطی برای ماتریس $X^T X$ بزرگ‌تر از 10^6 باشد، آنگاه در ماتریس طرح هم خطی وجود دارد. اگر $10^6 < \kappa < 10^8$ باشد، هم خطی شدید و چنانچه $\kappa > 10^8$ باشد هم خطی بسیار جدی است. برای اطلاعات بیشتر در باره عدد شرطی به آرست (۱۳۹۵) مراجعه کنید.

یکی از روش‌های مقابله با اثرات نامطلوب هم خطی، استفاده از برآوردهای اریب است که رگرسیون ریح یکی از مهم‌ترین و کاراترین روش‌هاست. هورل و کنارد (۱۹۷۰) بر اساس روش منظم‌سازی تیکونوف (۱۹۴۳)، برآوردهای ریح را معرفی کردند. این روش با اضافه کردن جمله منظم‌ساز، $\lambda\beta^T\beta$ ، به معادلات برآورد، به نوعی یک برآوردهای رگرسیونی توانانیده به‌شمار می‌آید. اگر چه برآوردهای حاصل اریب است، با این وجود برآوردهایی حاصل می‌شود که واریانس آن در مقایسه با برآوردهای OLS کمتر است. این کاهش واریانس، کاهش MSE را به دنبال خواهد داشت. در نتیجه برآوردهای ریح در مقایسه با برآوردهای OLS پایدارتر بوده و دقت پیش بینی را بهبود می‌بخشد. توضیحات بیشتر در باره این روش در زیربخش ۱۰.۳.۱.۲ ارائه می‌شود.

هنگامی که تعداد متغیرهای تبیینی بزرگ است، هم خطی می‌تواند جدی باشد. ورود تعداد متغیرهای فراوان به مدل اگرچه انعطاف‌پذیری مدل را افزایش می‌دهد، اما علاوه بر ایجاد هم خطی شدید در ماتریس طرح، می‌تواند باعث کاهش توان پیشگویی و مشکل شدن تفسیر مدل شود. برآوردهای ریح اگرچه توان پیشگویی بالایی دارد، اما مدل پیچیده‌ای را نتیجه می‌دهد، یعنی تعداد پارامترهای آن بزرگ است. تفسیر یک مدل پیچیده بسیار مشکل است. در مدل‌های با بعد بالا، معمولاً تنها تعداد اندکی از متغیرهای تبیینی با متغیر پاسخ مرتبط هستند. بهترین روش در این حالت، حذف متغیرهای زائد و مدل‌بندی با زیرمجموعه‌ای از موثرترین متغیرهای تبیینی است. تاکنون روش‌های مختلف برای انتخاب متغیر معرفی شده‌اند. در ادامه به مهم‌ترین روش‌های انتخاب متغیر می‌پردازیم. این روش‌ها شامل روش‌های انتخاب متغیر کلاسیک و روش‌های مبتنی بر تابع توان هستند.

۲.۱.۲ انتخاب متغیر با معیارهای کلاسیک

تاکنون الگوریتم‌های متعددی برای انتخاب متغیر در مدل‌های خطی معرفی شده است. این الگوریتم‌ها براساس معیارهایی مانند معیار اطلاع آکائیک (AIC)، معیار اطلاع بیزی^{۲۰} (BIC)، اعتبارسنجی متقابل تعمیم‌یافته^{۲۱} (GCV)، ضریب تعیین تعدیل‌شده (R_{adj}^2)، مجموع توان‌های دوم باقیمانده^{۲۲} (RSS) و C_p -مالو^{۲۳}، زیرمجموعه‌های متفاوتی از متغیرهای تبیینی را به عنوان بهترین زیرمجموعه

²⁰Bayesian information criterion

²¹Generalized Cross-Validation

²²Residual sum of squares

²³Mallows's C_p

انتخاب می‌کنند. از جمله این روش‌ها می‌توان انتخاب بهترین زیرمجموعه^{۲۴}، انتخاب پیش‌رو^{۲۵}، روش حذفی پس‌رو^{۲۶}، و روش گام به گام^{۲۷} را نام برد.

در روش انتخاب بهترین زیرمجموعه، تمام زیرمجموعه‌های ممکن را در نظر گرفته و سپس براساس معیارهای فوق، بهترین زیرمجموعه از مدل کامل انتخاب می‌شود. در این حالت، اگر p تعداد کل متغیرهای تبیینی باشد، تعداد 2^p زیرمجموعه باید مورد بررسی قرار گیرند. به عنوان مثال، اگر $p = 10$ ، آنگاه باید ۱۰۲۴ مدل ممکن در نظر گرفته شوند و برای $p = 20$ باید بیش از یک میلیون مدل ممکن بررسی شوند. بنابراین با افزایش تعداد p ، زیرمجموعه‌های ممکن به سرعت افزایش می‌یابد. لذا روش انتخاب بهترین زیرمجموعه از نظر محاسباتی، حتی در صورت استفاده از کامپیوترهای با سرعت بالا، برای $p > 40$ قابل استفاده نیست. در این راستا، دراپر و اسمیت (۱۹۸۱) روش‌های انتخاب پیش‌رو، حذفی پس‌رو و گام به گام را معرفی کردند. روش انتخاب پیش‌رو با مدلی که تنها شامل عرض از مبدا است، شروع می‌شود و سپس در هر گام متغیری که بیش‌ترین تاثیر را در بهبود مدل داشته باشد، انتخاب می‌شود. این رویه تا زمانی ادامه می‌یابد که تاثیر هیچ یک از متغیرهای باقیمانده در بهبود مدل معنی‌دار نباشد. برای اندازه‌گیری میزان تأثیر هر یک از متغیرها در بهبود مدل، از آماره‌هایی مانند R^2 ، RSS ، F و t استفاده می‌شود. در روش حذفی پس‌رو، فرآیند با مدل کامل شروع می‌شود و هر بار یک متغیر از مدل حذف می‌شود. متغیر حذف‌شده همان متغیری است که کمترین تاثیر را در بهبود مدل دارد. به عبارت دیگر در هر گام، متغیر تبیینی با کوچک‌ترین آماره F جزئی برای حذف از مدل نامزد می‌شود. اگر آماره F جزئی متغیر منتخب کمتر از مقدار F^* باشد، متغیر مورد نظر از مدل حذف می‌شود. فرآیند حذف متغیرها زمانی متوقف می‌شود که کوچک‌ترین مقدار آماره F جزئی کوچک‌تر از F^* نباشد، یا همه متغیرهای مدل حذف شده باشند. در روش‌های انتخاب پیش‌رو و حذفی پس‌رو، هنگامی که یک متغیر حذف یا اضافه می‌شود، در مراحل بعدی معنی‌داری این متغیرها مجدداً بررسی نمی‌شود. روش گام به گام تعدیل روش انتخاب پیش‌رو است که در هر گام همه متغیرهای واردشده به مدل در گام‌های قبل، مجدداً ارزیابی می‌شوند. بنابراین یک متغیر تبیینی اضافه شده در گام قبلی ممکن است در گام بعدی از مدل حذف شود. برای اطلاعات بیشتر درباره سایر روش‌های انتخاب متغیر می‌توان به دزیاک و همکاران (۲۰۰۵)، دسبولتس (۲۰۱۸) و دینگ (۲۰۱۸) مراجعه کرد.

اگرچه روش‌های انتخاب زیرمجموعه با معیارهای کلاسیک از ویژگی‌های نمونه‌ای مطلوبی برخوردارند، اما این روش‌ها پایدار نیستند، بدین معنی که یک تغییر بسیار کوچک در داده‌ها ممکن است مدل‌های کاملاً متفاوت را نتیجه دهد. لذا ممکن است این روش‌ها پیشگویی بدی را ایجاد کنند (بريمن، ۱۹۹۶). علاوه بر این، زمانی که p بزرگ باشد، هزینه محاسباتی این روش‌ها بسیار بالاست. برای رفع این چالش‌ها، محققین رگرسیون تاوانیده را معرفی کردند که با انتخاب یک تابع

²⁴ Best subset selection

²⁵ Forward selection

²⁶ Backward elimination

²⁷ Stepwise

تاوان مناسب، دو عمل انتخاب متغیر و برآورد ضرایب را به‌طور همزمان انجام می‌دهد.

۳.۱.۲ انتخاب متغیر مبتنی بر تابع تاوان

در این بخش، به توصیف روش رگرسیون تاوانیده می‌پردازیم. بدیهی است وقتی تعداد متغیرهای تبیینی p زیاد است، اندازه مجموع قدر مطلق ضرایب β_k ، $k = 1, \dots, p$ ، افزایش می‌یابد. از آنجایی که در روش کمترین توان‌های دوم به دنبال کمینه کردن $\|\varepsilon\|^2 = \varepsilon^T \varepsilon$ نسبت به β هستیم، پس چنان‌چه بخواهیم علاوه بر برآورد ضرایب رگرسیونی، تعداد متغیرهای مدل را نیز کاهش دهیم، کافی است بزرگی قدرمطلق ضرایب یا تابعی از آن را نیز به‌طور همزمان کمینه کنیم. بنابراین در رگرسیون تاوانیده، محک $\|\varepsilon\|^2$ را تحت این قید که نباید مجموع تابعی از قدرمطلق ضرایب رگرسیونی مانند $\sum_{k=1}^p P(|\beta_k|)$ بزرگ باشد، کمینه می‌کنند. این یک مسئله بهینه‌سازی مقید^{۲۸} (CO) است. به‌طور دقیق‌تر، علاقه‌مندیم مسئله CO به‌صورت

$$\min_{\beta \in \mathbb{R}^p} \|\varepsilon\|^2 \quad s.t. \quad \sum_{k=1}^p P_\lambda(|\beta_k|) \leq s, \quad (3.2)$$

را حل کنیم که در آن $P_\lambda(|\cdot|)$ تابع تاوان نامیده می‌شود و λ پارامتر تاوان است که میزان تاوان را کنترل می‌کند. یکی از راه‌های حل مسئله CO استفاده از روش ضریب لاگرانژ^{۲۹} است. بنابراین مسئله CO معادل با مسئله

$$\hat{\beta} = \arg \min_{\beta} \left\{ \|\mathbf{y} - \mathbf{X}\beta\|^2 + \lambda \sum_{k=1}^p P(|\beta_k|) \right\} \quad (4.2)$$

است. با دقت در ساختار مسئله (۴.۲) دیده می‌شود که به تابع مجموع توان‌های دوم خطا در رگرسیون معمولی، عبارت تاوان $\left\{ \lambda \sum_{k=1}^p P(|\beta_k|) \right\}$ اضافه شده است. در اینجا، $P_\lambda(|\cdot|) = \lambda P(|\cdot|)$ تابع تاوان و $\lambda > 0$ پارامتر تنظیم‌کننده^{۳۰} است که میزان انقباض را برحسب بزرگی ضرایب رگرسیونی کنترل می‌کند. در این حالت بین λ در (۴.۲) و s در (۳.۲) رابطه‌ای معکوس وجود دارد. به‌طور دقیق‌تر، اگر بخواهیم تعداد متغیرهای بیش‌تری از مدل حذف شوند، کافی است s را کوچک یا λ را بزرگ انتخاب کنیم و برعکس. پارامتر تنظیم‌کننده λ موازنه‌ای بین خطای پیشگویی و پیچیدگی مدل برقرار می‌کند. جواب مسئله بهینه‌سازی (۴.۲) به شدت به مقدار λ وابسته است. تعیین مقدار مناسب برای λ یک مسئله چالش برانگیز است. روش‌های مختلف انتخاب مقدار بهینه λ در زیربخش ۳.۳.۱.۲ داده شده است. بنابراین با انتخاب یک مقدار بهینه برای λ ، مقادیر برخی از پارامترها ممکن است صفر شوند که در این صورت متغیر تبیینی متناظر با آن پارامتر، به عنوان متغیر بی‌تاثیر بر متغیر پاسخ شناخته شده و بدین ترتیب مجموعه متغیرهای مهم و بی‌اهمیت از یکدیگر تشخیص داده می‌شوند و در اصطلاح “انتخاب متغیر” صورت می‌گیرد. در ادامه چند تابع تاوان مهم را معرفی می‌کنیم.

²⁸Constraint optimization

²⁹Lagrange multiplier

³⁰Tuning parameter

۱.۳.۱.۲ انواع تابع تاوان

تاکنون توابع تاوان زیادی توسط افراد مختلف معرفی شده‌اند. در این بخش به معرفی برخی از توابع تاوان معروف می‌پردازیم.

۱. تاوان L_q : با فرض $P(|\beta_k|) = |\beta_k|^q$ در رابطه (۳.۲)، تابع تاوان بریج^{۳۱} حاصل می‌شود. این تابع رگرسیون بریج را نتیجه می‌دهد که توسط فرانک و فریدمن (۱۹۹۳) معرفی شد. برآوردگر بریج به صورت

$$\hat{\beta}^{Bridge} = \arg \min_{\beta \in \mathbb{R}^p} \left\{ \frac{1}{2} \|y - X\beta\|^2 + \lambda \sum_{k=1}^p |\beta_k|^q \right\}, \quad 0 < q \leq 2.$$

تعریف می‌شود. این تابع تاوان، برخی از توابع تاوان مهم را به عنوان یک حالت خاص شامل می‌شود. رگرسیون ریج: برای حالت خاص $q = 2$ ، مدل رگرسیونی حاصل را رگرسیون ریج^{۳۲} گویند و برآورد پارامترها به صورت

$$\begin{aligned} \hat{\beta}^{Ridge} &= \arg \min_{\beta \in \mathbb{R}^p} \left\{ \frac{1}{2} \|y - X\beta\|^2 + \lambda \sum_{k=1}^p |\beta_k|^2 \right\} \\ &= (X^T X + K I_p)^{-1} X^T y, \end{aligned}$$

به دست می‌آید که $K = 2\lambda$ را پارامتر تنظیم‌کننده یا پارامتر ریج گویند. یک مزیت رگرسیون ریج پیاده‌سازی آسان و صورت صریح و بسته جواب آن است. اگر ماتریس طرح X دارای رتبه کامل ستونی نباشد، ماتریس $X^T X$ وارون‌پذیر نخواهد بود و لذا جواب یکتایی برای برآوردگر وجود نخواهد داشت. اما برای هر ماتریس دلخواه X و $K > 0$ ، ماتریس $X^T X + K I_p$ همواره وارون‌پذیر است و در نتیجه پاسخ یکتا برای برآوردگر ریج وجود دارد. برآوردگر ریج اریب است، اما در مقایسه با برآوردگر OLS دارای واریانس کوچکتری است. همچنین همواره λ وجود دارد به طوری که $MSE(\hat{\beta}^{Ridge}) < MSE(\hat{\beta}^{OLS})$ (صالح و همکاران، ۲۰۱۹). بنابراین رگرسیون ریج پیشگویی بهتری از رگرسیون کمترین توان‌های دوم معمولی ارائه می‌دهد. با این وجود برآوردگر ریج اگرچه ضرایب رگرسیون را منقبض می‌کند اما هیچ یک از ضرایب را صفر برآورد نمی‌کند، یعنی این روش توانایی انتخاب متغیر ندارد. این موضوع خللی در پیشگویی ایجاد نمی‌کند، اما با توجه به حضور همه متغیرها در مدل، تفسیر مدل برازش‌شده وقتی p بزرگ است به سادگی امکان‌پذیر نیست. در حالتی که $X^T X = I_p$ ، ماتریس طرح متعامد خواهد بود. در این حالت برآوردگر ریج به صورت

$$\hat{\beta}^{Ridge} = \frac{1}{1 + K} \hat{\beta}^{OLS}$$

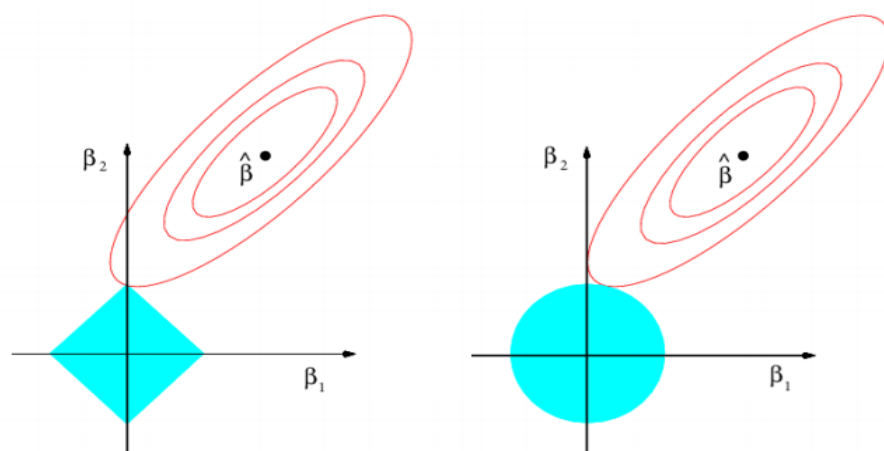
است. این رابطه ماهیت انقباضی رگرسیون تاوانیده را نشان می‌دهد، چون با افزایش K ضرایب به سمت صفر منقبض می‌شوند. هنگامی که $K = 0$ ، برآوردگرهای OLS و ریج معادل‌اند.

³¹Bridge³²Ridge

رگرسیون لاسو: در این حالت مدل رگرسیونی حاصل را که توسط تیشیرانی (۱۹۹۶) معرفی شد، رگرسیون لاسو گویند. از دیدگاه آماری، لاسو مخفف عملگر انتخاب‌کننده، منقبض‌کننده و کمینه‌کننده از جنس قدرمطلق است. در رگرسیون لاسو برآورد ضرایب رگرسیونی به صورت

$$\hat{\beta}^{Lasso} = \arg \min_{\beta \in \mathbb{R}^p} \left\{ \frac{1}{2} \|y - X\beta\|^2 + \lambda \sum_{j=1}^p |\beta_j| \right\}$$

به دست می‌آید که در آن λ پارامتر تنظیم‌کننده است و سطح تنکی (تعداد ضرایب صفر) را کنترل می‌کند. وقتی $\lambda \rightarrow 0$ ، آنگاه $\hat{\beta}^{Lasso} \rightarrow \hat{\beta}^{OLS}$ و وقتی $\lambda \rightarrow \infty$ ، آنگاه $\hat{\beta}^{Lasso} \rightarrow 0$. به عبارت دیگر، هر چه تاوان بزرگ‌تری به کار رود، تعداد بیش‌تری از ضرایب به سمت صفر منقبض می‌شوند. برآورد بدست‌آمده با رگرسیون لاسو تنک است، یعنی برخی از ضرایب صفر برآورد می‌شوند. به عبارت دیگر این روش انتخاب متغیر و برآورد پارامتر را به‌طور همزمان انجام می‌دهد. تابع تاوان قدرمطلق باعث می‌شود که جواب صریحی برای برآوردگر لاسو وجود نداشته باشد. تیشیرانی این برآوردگر را از طریق برنامه‌ریزی درجه دوم^{۳۳} به دست آورده بود. سپس افرون و همکاران (۲۰۰۴) رگرسیون کمترین زاویه^{۳۴} را معرفی نمودند که برآوردهای لاسو را نتیجه می‌دهد و هزینه محاسباتی آن با برآوردهای کمترین توان‌های دوم برابر است.



شکل ۱۰.۲: منحنی‌های تراز و نواحی محدودیت رگرسیون لاسو (نمودار چپ) و رگرسیون ریج (نمودار راست). ناحیه‌های محدودیت لاسو و ریج به ترتیب به صورت $|\beta_1| + |\beta_2| \leq s$ و $\beta_1^2 + \beta_2^2 \leq s$ هستند. بیضی‌ها، منحنی‌های تراز تابع خطای برآورد کمترین توان‌های دوم هستند.

برای روشن‌تر شدن مفهوم انتخاب متغیر توسط تابع تاوان لاسو، حالت خاصی از مدل رگرسیونی را به صورت

$$y = \beta_1 x_1 + \beta_2 x_2 + \varepsilon,$$

در نظر بگیرید. بنابراین مجموع توان‌های دوم خطا $Q(\beta)$ ، تابعی از دو پارامتر β_1 و β_2 است که در

^{۳۳}Quadratic programming

^{۳۴}Least angle regression

آن

$$Q(\beta) = \|y - X\beta\|^2 = \sum_{i=1}^n \varepsilon_i^2 = \sum_{i=1}^n (y_i - \beta_1 x_{i1} - \beta_2 x_{i2})^2.$$

نمودار تابع $Q(\cdot)$ یک رویه درجه دوم به صورت

$$Q(\beta) = a_0 + a_1\beta_1 + a_2\beta_2 + a_3\beta_1^2 + a_4\beta_2^2 + a_5\beta_1\beta_2$$

است که منحنی‌های تراز آن در شکل ۱۰۲ نشان داده شده است. اکنون شرط‌های کمینه‌سازی ريج و لاسو را به صورت

$$\beta_1^2 + \beta_2^2 = s \quad \text{و} \quad |\beta_1| + |\beta_2| = s$$

در نظر بگیرید که شکل هندسی آن‌ها به صورت دایره برای رگرسیون ريج و لوزی برای رگرسیون لاسو نشان داده شده است. با توجه به نمودار سمت چپ شکل ۱۰۲، در رگرسیون لاسو می‌تواند حالتی رخ دهد که یکی از ضرایب صفر شده و دیگری همواره برابر s باشد، در صورتی که در نمودار سمت راست این امکان وجود ندارد، به عبارت دقیق‌تر، در قاب سمت راست، اگر یکی از ضرایب صفر شود، آن‌گاه دیگری می‌تواند $\sqrt{s}$ ، و نه s ، باشد. بنابراین برای لاسو امکان صفر برآورد شدن ضرایب وجود دارد، اما برای ريج این اتفاق هرگز رخ نمی‌دهد. در صورت صفر شدن هر یک از ضرایب در رگرسیون لاسو، متغیر متناظر با آن ضریب به عنوان متغیر بی‌اهمیت شناخته می‌شود. لازم به ذکر است که در هر دو روش لاسو و ريج، با انتخاب مقدار کوچک s ، هر دو پارامتر β_1 و β_2 (در حالت دو بعدی) تا حد ممکن می‌توانند کوچک شوند. بنابراین هر دو روش ريج و لاسو انقباضی هستند. اگر ماتریس طرح متعامد باشد، آن‌گاه می‌توان برآوردگر لاسو را به عنوان تابعی از برآوردگر کمترین توان‌های دوم معمولی β_k به صورت

$$\hat{\beta}_k^{Lasso} = \text{sgn}(z_k)(|z_k| - \lambda)_+, \quad j = 1, \dots, p, \quad (5.2)$$

به دست آورد که در آن $\text{sgn}(x) = I(x \geq 0) - I(x < 0)$ ، $I(A)$ تابع نشانگر، $x_+ = \max\{0, x\}$ و $z_k = X_k^T y$ برآورد OLS ضریب β_k بوده که X_k ، ستون k ام ماتریس طرح X است. برای اطلاعات بیشتر به فن و لی (۲۰۰۱) مراجعه کنید. برآوردگر (۵.۲) را به صورت

$$\hat{\beta}_k^{Lasso} = S(z_k|\lambda), \quad (6.2)$$

نشان می‌دهیم که $S(\cdot|\lambda)$ عملگر آستانه‌ای نرم نامیده می‌شود. رابطه (۶.۲) در بررسی ویژگی‌های مطلوب رگرسیون لاسو مفید است که در پایان این بخش به این موضوع خواهیم پرداخت. رگرسیون لاسو با اینکه در بسیاری از مسائل انتخاب متغیر عملکرد خوبی از خود نشان داده، اما دارای برخی محدودیت‌ها است. به عنوان مثال، لاسو برای یک مدل رگرسیونی خطی با p متغیر پیشگو و n مشاهده، حداکثر n متغیر را انتخاب می‌کند. بنابراین اگر تعداد متغیرهای تبیینی معنی‌دار در مدل بیش‌تر از n باشد، توسط لاسو انتخاب نمی‌شوند. لاسو در بین یک گروه از متغیرهای به شدت همبسته فقط یک متغیر را انتخاب می‌کند (زو و هستی، ۲۰۰۵). برای حالت معمولی $n > p$ ، اگر

همبستگی بالایی بین متغیرهای پیشگو وجود داشته باشد، کارآیی پیشگویی لاسو به خوبی برآوردگر ریح نیست (تیبشیرانی، ۱۹۹۶). رگرسیون لاسو برآوردگری را نتیجه می‌دهد که ارببی آن زیاد است. در حالتی که ماتریس طرح متعامد است، می‌توان نشان داد که برای $k = 1, \dots, p$ داریم

$$\begin{cases} E|\hat{\beta}_k^{Lasso} - \beta_k| = 0, & \beta_k = 0 \\ E|\hat{\beta}_k^{Lasso} - \beta_k| \approx \beta_k, & \beta_k \in [0, \lambda] \\ E|\hat{\beta}_k^{Lasso} - \beta_k| \approx \lambda & |\beta_k| > \lambda \end{cases} \quad (7.2)$$

بنابراین مقدار ارببی برآورد لاسو برای ضرایب بزرگ، نزدیک به λ است (برهنی و هووانگ، ۲۰۱۱). برای از بین بردن این محدودیت‌ها، توابع تاوان متعددی معرفی شدند که در ادامه به برخی از آن‌ها اشاره می‌کنیم. قبل از معرفی این توابع، ابتدا ویژگی پیشگویی را تعریف می‌کنیم.

ویژگی پیشگویی: فرض کنید $\mathcal{A} = \{k : \beta_k \neq 0\}$ مجموعه اندیس ضرایب غیرصفر باشد، به‌طوری که $|\mathcal{A}| = p_0 < p$. در این صورت، مدل درست تنها به زیرمجموعه‌ای از متغیرها (p_0 تا متغیر) وابسته است. منظور از مدل درست مدلی است که تنها شامل ضرایب غیرصفر باشد. براساس فن و لی (۲۰۰۱)، یک روش برآورد ضرایب دارای ویژگی پیشگویی^{۳۵} است، اگر به‌طور مجانبی دارای خواص زیر باشد:

۱. بتواند ضرایب غیرصفر را به درستی شناسایی کند، یعنی اگر $\hat{\mathcal{A}} = \{k : \hat{\beta}_k \neq 0\}$ ، آن‌گاه

$$\lim_{n \rightarrow \infty} \mathbb{P}(\hat{\mathcal{A}} = \mathcal{A}) = 1$$

۲. دارای نرخ برآورد بهینه باشد، به عبارت دیگر

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

که Σ^* ماتریس کوواریانس زیر مدل درست یعنی مدلی برپایه $\mathcal{A}$ ، است.

بنابراین مطلوب است که یک برآوردگر دارای ویژگی پیشگویی باشد. زو (۲۰۰۶) نشان داد که برآوردگر لاسو دارای ویژگی پیشگویی نیست. با توجه به رابطه (۷.۲) ملاحظه می‌شود که مقدار ارببی برآوردگر لاسو توسط λ تعیین می‌شود. بنابراین یک روش کاهش ارببی لاسو، استفاده از تاوان موزون به‌صورت $\lambda_k = \omega_k \lambda$ است. اگر بتوانیم وزن‌ها را طوری انتخاب کنیم که به ضرایب بزرگ‌تر تاوان کوچتری تخصیص داده شود، آن‌گاه با حفظ ویژگی تنکی، ارببی برآورد لاسو کاهش می‌یابد. بنابراین زو (۲۰۰۶) لاسوی سازوار را معرفی کرد و نشان داد که این روش از ویژگی پیشگویی برخوردار است.

۲. تاوان لاسوی سازوار: صورت کلی تابع تاوان لاسوی سازوار مانند لاسو است، اما برخلاف لاسو برای ضرایب مختلف، تاوان‌های مختلفی در نظر می‌گیرد. برآوردگر لاسوی سازوار به‌صورت

$$\hat{\beta}^{aLasso} = \arg \min_{\beta \in \mathbb{R}^p} \left\{ \frac{1}{2} \|\mathbf{y} - \mathbf{X}\beta\|^2 + \lambda \sum_{k=1}^p \omega_k |\beta_k| \right\}$$

^{۳۵}Oracle property

بیان می‌شود، که ω_k وزن β_k است تا ضرایب کوچک را بیش‌تر از ضرایب بزرگ منقبض کند. معمولاً در رگرسیون لاسوی تطبیقی، وزن $\omega_k = \frac{1}{|\hat{\beta}_k^0|^\gamma}$ را انتخاب می‌کنند که $\gamma > 0$ و $\hat{\beta}_k^0$ یک برآوردگر $\sqrt{n}$ -سازگار از β_k مانند برآوردگر OLS ، ريج یا لاسو است.

۳. تاوان اسکد: فن و لی (۲۰۰۱) یک تابع تاوان مقعر به نام اسکد^{۳۶} ($SCAD$) را معرفی کردند و نشان دادند که این تابع تاوان، از همه ویژگی‌های مطلوب یک تابع تاوان برخوردار است (زیر بخش ۲.۳.۱.۲ را ببینید). آن‌ها همچنین ویژگی پیشگویی را نیز برای این روش اثبات کردند. برآوردگر اسکد به صورت

$$\hat{\beta}^{SCAD} = \arg \min_{\beta \in \mathbb{R}^p} \left\{ \frac{1}{2} \|y - X\beta\|^2 + \lambda \sum_{k=1}^p P(|\beta_k|; a, \lambda) \right\}$$

است که $P(\cdot; a, \lambda)$ تابع تاوان اسکد است و به صورت

$$P(\beta; a, \lambda) = \begin{cases} \lambda\beta, & 0 \leq |\beta| \leq \lambda \\ \frac{2a\lambda\beta - \beta^2 - \lambda^2}{2(a-1)}, & \lambda < |\beta| < a\lambda \\ \frac{\lambda^2(a+1)}{2}, & |\beta| \geq a\lambda \end{cases} \quad (۸.۲)$$

تعریف می‌شود که $\lambda > 0$ و $a > 2$ پارامترهای تاوان می‌باشند و اندازه انقباض ضرایب را کنترل می‌کنند. فن و لی (۲۰۰۱) مقدار $a = 3/7$ را پیشنهاد دادند. شکل ۲.۲ نمودار تابع اسکد را به ازای $\lambda = 1$ نشان می‌دهد. با توجه به شکل ۲.۲ و رابطه (۸.۲) ملاحظه می‌شود که در فاصله $|\beta| \leq \lambda$ تابع تاوان اسکد بر لاسو منطبق است و پس از آن تا زمانی که $|\beta| \leq a\lambda$ ، تاوان اسکد یک تابع درجه دوم است. سپس به ازای $|\beta| > a\lambda$ به یک تابع ثابت تبدیل می‌شود. همچنین، با توجه به (۸.۲)، وقتی $a \rightarrow \infty$ ، برآوردگر اسکد با برآوردگر لاسو معادل است، یعنی تابع تاوان اسکد در حالت خاص $a = \infty$ تابع تاوان لاسو را نتیجه می‌دهد.

با فرض متعامد بودن ماتریس طرح، برآوردهای اسکد ضرایب را می‌توان به صورت

$$\hat{\beta}_j^{SCAD} = \begin{cases} S(z_k|\lambda), & |z_k| \leq 2\lambda \\ \frac{a-1}{(a-2)} S(z_k|\frac{a\lambda}{a-1}), & 2\lambda < |z_k| \leq a\lambda \\ |z_k|, & |z_k| > a\lambda \end{cases}$$

نشان داد که z_k برآورد OLS ضریب β_k و $S(\cdot|\lambda)$ عملگر آستانه‌ای نرم در (۵.۲) است.

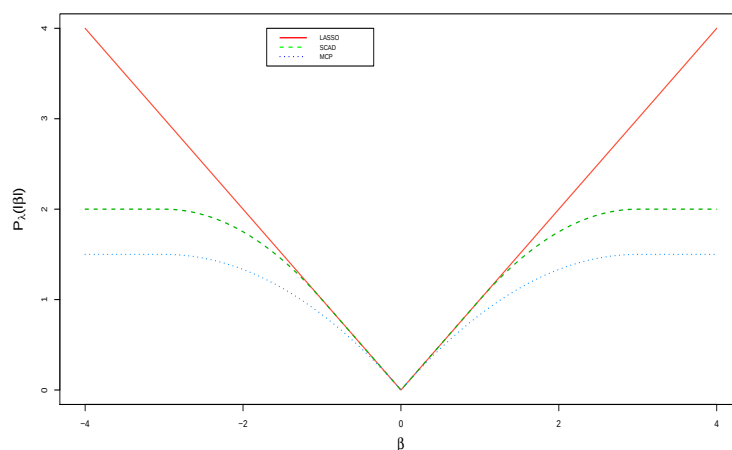
۴. تاوان MCP : ژانگ (۲۰۱۰) تابع تاوان دیگری به نام MCP به صورت

$$P(\beta; \lambda, \gamma) = \lambda \int_0^\beta \left(1 - \frac{x}{\gamma\lambda}\right)_+ dx; \quad \beta \geq 0,$$

معرفی کرد که در آن λ پارامتر تاوان و γ میزان مقعر بودن تابع $P(\cdot)$ را کنترل می‌کند. در این حالت، برآوردگر MCP به صورت

$$\hat{\beta}^{MCP} = \arg \min_{\beta \in \mathbb{R}^p} \left\{ \frac{1}{2} \|y - X\beta\|^2 + \lambda \sum_{j=1}^p P(|\beta_j|; \lambda, \gamma) \right\},$$

³⁶Smoothly clipped absolute deviation



شکل ۲.۲: مقایسه نمودارهای توابع تاوان لاسو، اسکد و MCP .

تعریف می‌شود. تابع تاوان MCP به ازای $\gamma = \infty$ تابع تاوان لاسو را نتیجه می‌دهد. تحت فرض متعامد بودن ماتریس طرح، مشابه برآوردهای لاسو و اسکد، برآوردهای MCP ضرایب را می‌توان به صورت

$$\hat{\beta}_k^{MCP}(\gamma, \lambda) = \begin{cases} \frac{\gamma}{\gamma-1} S(z_k | \lambda), & |z_k| \leq \gamma\lambda \\ z_k, & |z_k| > \gamma\lambda \end{cases}$$

نوشت. شکل ۲.۲ نمودار تابع MCP را به ازای $\lambda = 1$ و $\gamma = 3$ نشان می‌دهد. تابع تاوان MCP در فاصله $|\beta| \leq \gamma\lambda$ یک تابع درجه دوم و پس از آن یک تابع ثابت است.

۲.۳.۱.۲ ویژگی‌های یک تابع تاوان مناسب

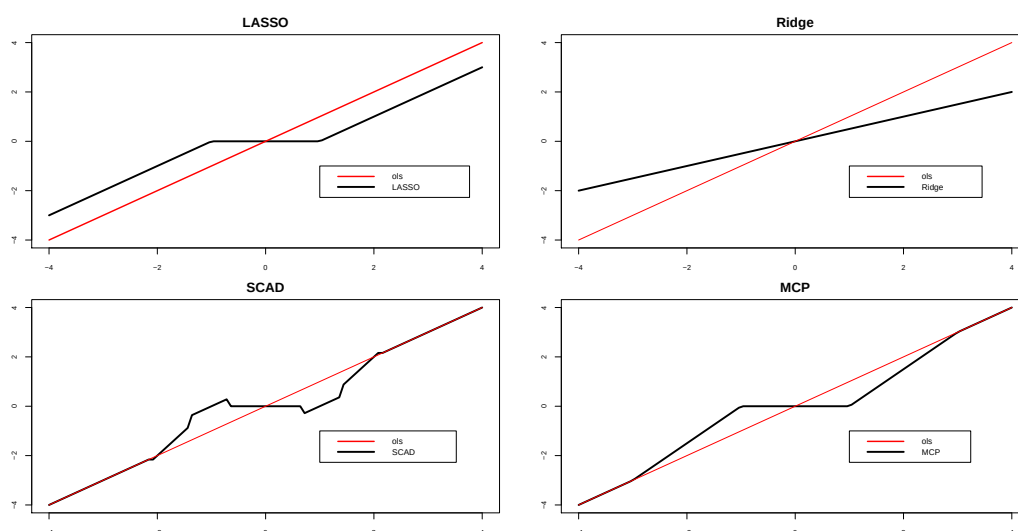
برای استفاده از رگرسیون تاوانیده، سوالی که پیش می‌آید این است که از چه نوع تابع تاوان باید استفاده کرد. فن و لی (۲۰۰۱) نشان دادند که تابع تاوان خوب باید برآوردگری با سه ویژگی مطلوب زیر را نتیجه دهد:

۱. تنکی: برآوردر نتیجه شده باید به‌طور خودکار ضرایب برآورده‌ای که مقدار کوچک دارند را برابر صفر قرار دهد تا متغیرهای مناسب را انتخاب کند. این کار پیچیدگی مدل را کاهش می‌دهد.

۲. نااریبی: برآوردر به‌دست آمده برای ضرایب رگرسیونی که مقادیر واقعی آن‌ها بزرگ است، تقریباً نااریب باشد. این ویژگی اریبی مدل را کاهش می‌دهد.

۳. پیوستگی: برآوردر نتیجه‌شده پیوسته باشد تا بتواند ناپایداری در پیشگویی مدل را کاهش دهد.

شکل ۳.۲ ارتباط برآوردر OLS را با برآوردهای لاسو، ریج، اسکد و MCP با فرض متعامد بودن ماتریس طرح X نشان می‌دهد. محور افقی نشان‌دهنده برآوردر OLS و محور عمودی



شکل ۳.۲: مقایسه نمودارهای برآوردهای کمترین توان‌های دوم تاوانیده (PLS) در مقابل برآورد کمترین توان‌های دوم معمولی (OLS) وقتی که ماتریس طرح متعامد است. محور افقی OLS و محور عمودی PLS است.

برآوردهای کمترین توان‌های دوم تاوانیده است. چون برآوردگر OLS ناریب است، انتظار داریم که برای ضرایب بزرگ، برآورد تاوانیده برابر یا نزدیک به برآورد OLS باشد تا شرط ناریبی برقرار شود. از طرفی برای برقراری شرط تنکی، لازم است روش کمترین توان‌های دوم تاوانیده، ضرایب رگرسیونی کوچک را صفر برآورد کند. همچنین لازم است برآوردگر موردنظر پیوسته باشد.

با توجه به شکل ۳.۲، پیوستگی برآوردگر لاسو بدیهی است. همچنین برای مقادیر کوچک ضرایب، برآورد لاسو برابر صفر است. در نتیجه شرط تنکی برقرار است. اما با بزرگ شدن مقادیر ضرایب، برآورد لاسو از برآورد OLS فاصله می‌گیرد، یعنی شرط ناریبی برقرار نیست. در نتیجه عیب روش لاسو این است که برای مقادیر بزرگ ضرایب، برآوردهای اریب را نتیجه می‌دهد. برآورد ریج در مقایسه با لاسو، علی‌رغم پیوسته بودن، ضرایب کوچک را صفر برآورد نمی‌کند. علاوه بر این، با افزایش قدرمطلق ضرایب، مقدار اریبی به شدت افزایش می‌یابد. بنابراین، تابع تاوان ریج یک تابع مناسب برای انتخاب متغیر نمی‌باشد. با توجه به نمودار برآورد اسکد در مقابل برآورد OLS ، برقراری ویژگی‌های تنکی و پیوستگی واضح است. علاوه بر این، در ابتدا میزان اریبی برآوردگر تاوانیده با اریبی برآوردگر لاسو معادل است، اما با افزایش قدرمطلق ضرایب، اریبی به سمت صفر تنزل می‌یابد. لذا ویژگی ناریبی نیز برقرار است. همچنین، با توجه به شکل، تابع تاوان MCP نیز برآوردگری با ویژگی‌های مطلوب را نتیجه می‌دهد. رفتار تابع تاوان MCP بسیار شبیه اسکد است، با این تفاوت که با فاصله گرفتن ضرایب از صفر، اریبی برآوردگر MCP فوراً به سمت صفر می‌رود، اما اریبی برآوردگر اسکد در یک فاصله معین برابر اریبی لاسو است و سپس به سرعت کاهش می‌یابد. بنابراین از توابع تاوان MCP و اسکد می‌توان به عنوان توابع تاوان مناسب نام برد.

همچنین فن و لی (۲۰۰۱) نشان دادند که تحت شرایط کافی زیر برآوردگر کمترین توان‌های دوم تاوانیده دارای ویژگی‌های مطلوب فوق است:

۱. تنکی: برآوردگر دارای ویژگی تنکی است اگر $\inf_{\beta \geq 0} \{\beta + P'_\lambda(\beta)\} > 0$.

۲. نا اریبی: برآوردگر به دست آمده نااریب است اگر، برای مقادیر بزرگ β ، داشته باشیم $P'_\lambda(\beta) = 0$.

۳. پیوستگی: برآوردگر پیوسته است اگر و فقط اگر $\arg \inf_{\beta \geq 0} \{\beta + P'_\lambda(\beta)\} = 0$.

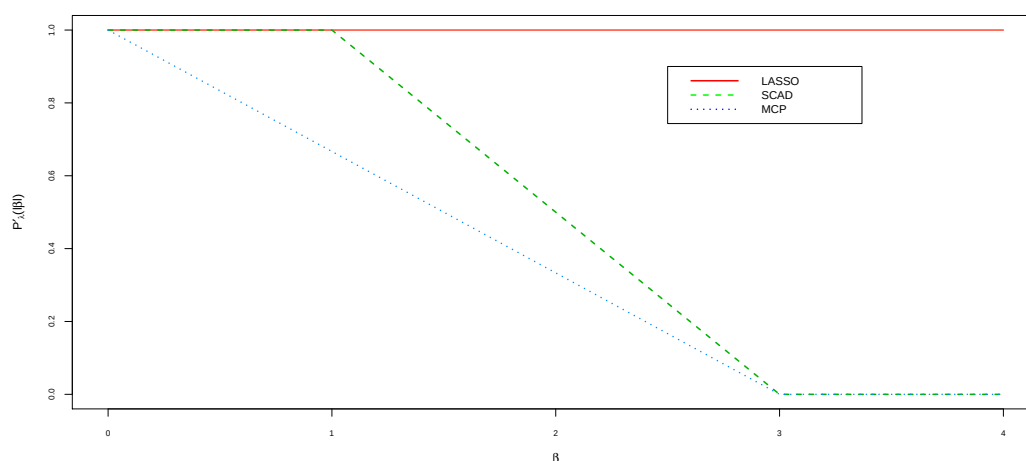
همان‌طور که ملاحظه می‌شود، شرایط کافی برای داشتن ویژگی‌های مطلوب بیان شده توسط فن و لی (۲۰۰۱) وابسته به مشتق تابع تاوان هست. علاوه بر این، شکل مشتق تابع تاوان در حل مسئله بهینه‌سازی تاوانیده تاثیر بسزایی دارد. مشتق تابع تاوان لاسو به صورت $P'_\lambda(\beta) = \lambda$ و مشتق توابع تاوان MCP و اسکد به ترتیب، به صورت

$$P'_\lambda(\beta; a, \lambda) = \lambda \left\{ I(\beta \leq \lambda) + \frac{(a\lambda - \beta)_+}{(a-1)\lambda} I(\beta > \lambda) \right\}, \quad \beta \geq 0$$

و

$$P'_\lambda(\beta; \lambda, \gamma) = \lambda \left(1 - \frac{\beta}{\gamma\lambda} \right), \quad \beta \geq 0$$

به دست می‌آید. بررسی شرایط کافی ویژگی‌های مطلوب توابع تاوان بسیار ساده است. به عنوان مثال، برای تابع تاوان لاسو، $\min_{\beta \geq 0} \{\beta + P'_\lambda(\beta)\} = \lambda > 0$ و $\arg \min_{\beta \geq 0} \{\beta + P'_\lambda(\beta)\} = 0$ ، اما $P'_\lambda(\beta) = \lambda > 0$. بنابراین، تابع تاوان لاسو دارای ویژگی نااریبی نیست.



شکل ۴.۲: نمودار مشتقات توابع تاوان لاسو، اسکد و MCP.

نمودار مشتق توابع تاوان لاسو، اسکد و MCP به ازای $\lambda = 1$ و $a, \gamma = 3$ در شکل ۴.۲ داده شده است. با توجه به شکل می‌بینیم که نرخ تاوان توابع MCP و اسکد در مبدأ با نرخ تاوان تابع لاسو برابر است، اما با افزایش قدر مطلق ضرایب، نرخ تاوان توابع MCP و اسکد به سمت صفر تنزل می‌کند. لازم به ذکر است که نرخ تاوان تابع اسکد در ابتدا با فاصله گرفتن ضرایب از صفر معادل نرخ تاوان لاسو است، سپس با افزایش قدر مطلق ضرایب به سرعت کاهش می‌یابد، اما تابع

تاوان MCP به محض خارج شدن از مبدأ، نرخ تاوان آن فوراً به سمت صفر تنزل می‌یابد. علاوه بر این، شکل نشان می‌دهد که روش MCP مقعر مینیماکس است، یعنی در بین همه توابع تاوان پیوسته مشتق‌پذیر بر بازه $(0, \infty)$ که برای هر $\beta \geq \gamma\lambda$ در دو شرط $P'_\lambda(0_+; \lambda) = \lambda$ و $P'_\lambda(\beta; \lambda) = 0$ صدق می‌کند، MCP بیشینه تقعر زیر را کمینه می‌کند

$$\mathcal{K} = \sup_{0 < \beta_1 < \beta_2} \frac{P'_\lambda(\beta_1; \lambda) - P'_\lambda(\beta_2; \lambda)}{\beta_2 - \beta_1}.$$

مقعر مینیماکس بودن تابع MCP بدین معنی است که در بین همه توابعی که دارای سه ویژگی مطلوب فوق هستند، MCP اندازه بیشینه تقعر را کمینه می‌کند. در شکل مشتق‌های توابع MCP و اسکد در نقاط 0 و $\gamma\lambda = 3$ برابر هستند، اما MCP در این ناحیه دارای تقعر $\frac{1}{3} = \mathcal{K}$ است، در حالی که اسکد دارای ماکزیمم تقعر $\frac{1}{4} = \mathcal{K}$ است. لازم به ذکر است که اغلب توابع تاوانی که شرایط مطلوب تنکی، نارایی و پیوستگی را دارا می‌باشند، توابع غیرمحدب هستند. این موضوع، مسئله بهینه‌سازی را با چالش مواجه می‌کند. برای اطلاعات بیش‌تر به برهنی و هووانگ (۲۰۱۹) مراجعه شود.

۳.۳.۱.۲ انتخاب پارامتر تاوان

انتخاب پارامتر تاوان در رگرسیون تاوانیده نقش بسیار مهمی در کنترل پیچیدگی مدل بازی می‌کند. به عنوان مثال، به ازای $\lambda = 0$ برآوردگر تاوانیده به برآوردگر OLS تبدیل می‌شود و همه متغیرهای تبیینی وارد مدل می‌شوند. اگر $\lambda \rightarrow \infty$ ، جزء دوم تابع (۴.۲) به ∞ همگرا می‌شود، و در نتیجه تمام ضرایب صفر شده و هیچ متغیری وارد مدل نمی‌شود. برای انتخاب بهترین مقدار بین این دو حالت فرین، از معیارهای کلاسیک AIC ، BIC و GCV استفاده می‌شود. به‌طور مشخص، برای λ داده‌شده، برآورد کمترین توان‌های دوم تاوانیده، $\hat{\beta}_\lambda$ را به‌دست آورده و سپس مقدار معیار انتخاب کلاسیک را براساس $\hat{\beta}_\lambda$ محاسبه می‌کنیم. این معیارهای کلاسیک به‌صورت

$$\begin{aligned} AIC_\lambda &= \|y - X\hat{\beta}_\lambda\|^2 + 2df_\lambda\hat{\sigma}^2, \\ BIC_\lambda &= \|y - X\hat{\beta}_\lambda\|^2 + \log(n)\hat{\sigma}^2, \\ GCV_\lambda &= \frac{\frac{1}{n}\|y - X\hat{\beta}_\lambda\|^2}{(1 - df_\lambda/n)^2}, \end{aligned}$$

محاسبه می‌شوند، که $\hat{\sigma}^2 = \frac{RSS_p}{n-p}$ و RSS_p مجموع توان‌های دوم باقیمانده‌ها تحت مدل کامل است. منظور از مدل کامل مدلی است که همه متغیرهای تبیینی در مدل حضور دارند. در معیارهای فوق، درجه آزادی مدل برآورد شده df_λ معمولاً با استفاده از تعداد ضرایب غیرصفر به‌صورت

$$df_\lambda = \sum_{k=1}^p I(\hat{\beta}_{k,\lambda} \neq 0),$$

به دست می‌آید که در آن $I(\cdot)$ تابع نشانگر است. راه دیگر محاسبه df_λ استفاده از ماتریس طرح به صورت

$$df_\lambda = \text{tr} \left(X_\lambda (X_\lambda^\top X_\lambda + n \Sigma_\lambda)^{-1} X_\lambda^\top \right),$$

است که X_λ ماتریس طرح مدل متناظر با λ داده شده و Σ_λ ماتریس بلوکی به صورت

$$\Sigma_\lambda = \text{diag} \left\{ P'_\lambda(|\hat{\beta}_{k,\lambda}|) / |\hat{\beta}_{k,\lambda}| \right\}_{\hat{\beta}_{k,\lambda} \neq 0}$$

بوده و $\hat{\beta}_{k,\lambda}$ ، k امین مولفه از برآورد کمترین توان‌های دوم تاوانیده $\hat{\beta}_\lambda$ است. در پایان، برای یک مجموعه از نقاط داده شده $\lambda_1, \dots, \lambda_M$ ، آن مقداری از λ را به عنوان مقدار بهینه انتخاب می‌کنیم که دارای کمترین مقدار معیار کلاسیک باشد.

یک روش دیگر برای انتخاب پارامتر تاوان استفاده از روش اعتبارسنجی متقابل k تایی^{۳۷} است. در این روش، برای یک λ داده شده، مجموعه مشاهدات را به طور تصادفی به k گروه با اندازه یکسان افراز می‌کنیم. سپس گروه اول را به عنوان داده‌های آزمون^{۳۸} و سایر گروه‌ها را به عنوان داده‌های آموزشی در نظر می‌گیریم. پس از برازش مدل با استفاده از داده‌های آموزشی، مقدار متغیر پاسخ به ازای داده‌های آزمون پیشگویی می‌شود. سپس خطای پیشگویی را به صورت $PE_1 = \sum_{i \in I_1} (y_i - \hat{y}_i)^2$ محاسبه می‌کنیم. که I_1 اندیس مشاهدات گروه اول است. این روند k بار تکرار می‌شود و در هر تکرار یک گروه مختلف به عنوان داده‌های آزمون و بقیه گروه‌ها به عنوان داده‌های آموزشی در نظر گرفته می‌شود. در نتیجه k برآورد خطای آزمون $PE_1, \dots, PE_k$ به دست می‌آید. سپس اعتبارسنجی متقابل با استفاده از رابطه

$$CV(\lambda) = \frac{1}{k} \sum_{i=1}^k PE_i$$

محاسبه می‌شود. آن λ ای انتخاب می‌شود که مقدار $CV(\lambda)$ را کمینه کند. در عمل معمولاً مقدار k برابر ۵ یا ۱۰ در نظر گرفته می‌شود.

۴.۳.۱.۲ الگوریتم محاسباتی روش مبتنی بر تابع تاوان

رگرسیون تاوانیده در بخش‌های قبل را به صورت

$$\ell(\beta) = \frac{1}{2} \|y - X\beta\|^2 + \sum_{k=1}^p P_\lambda(|\beta_k|) \quad (9.2)$$

در نظر بگیرید. با یک تابع تاوان محدب، مانند لاسو، تابع هدف (۹.۲) یک تابع محدب است و لذا برآوردهای رگرسیون تاوانیده را می‌توان با بهینه‌سازی محدب به دست آورد. اما اغلب توابع تاوانی که دارای سه ویژگی مطلوب تنکی، ناریبی و پیوستگی می‌باشند، توابع غیرمحدب بوده و این موضوع باعث می‌شود که تابع هدف (۹.۲) محدب نباشد. می‌توان توابع تاوان غیرمحدب را با برخی توابع

^{۳۷}k-fold cross validation

^{۳۸}Test dataset

محدب تقریب زد و مسئله بهینه‌سازی غیرمحدب را با استفاده از الگوریتم‌های بهینه‌سازی محدب حل کرد. در این قسمت دو تقریب زیر را برای حل مشکل بیان شده، معرفی می‌کنیم.

تقریب درجه دوم موضعی: چون عبارت اول در رابطه (۹.۲) یک تابع محدب است، لذا کافی است تابع تاوان غیر محدب $P_\lambda(|\beta_k|)$ را با یک تابع محدب تقریب کنیم. در این راستا، فن و لی (۲۰۰۱) الگوریتم درجه دوم موضعی^{۳۹} (LQA) را برای تقریب توابع تاوان غیرمحدب ارائه دادند. فرض کنید $\beta^{(\circ)} = (\beta_1^{(\circ)}, \dots, \beta_p^{(\circ)})^\top$ مقدار اولیه‌ای برای β باشد. اگر $\beta_k^{(\circ)}$ نزدیک به صفر باشد، قرار می‌دهیم $\hat{\beta}_k = 0$ ، در غیر این صورت با استفاده از تقریب موضعی داریم

$$[P_\lambda(|\beta_k|)]' = P'_\lambda(|\beta_k|) \frac{\beta_k}{|\beta_k|} \approx P'_\lambda(|\beta_k^{(\circ)}|) \frac{\beta_k}{|\beta_k^{(\circ)}|}, \quad \beta_k \approx \beta_k^{(\circ)}. \quad (10.2)$$

با انتگرال‌گیری از طرفین نسبت به β_k از $\beta_k^{(\circ)}$ تا β_k داریم

$$P_\lambda(|\beta_k|) - P_\lambda(|\beta_k^{(\circ)}|) \approx \frac{1}{2} \frac{P'_\lambda(|\beta_k^{(\circ)}|)}{|\beta_k^{(\circ)}|} \{\beta_k^2 - \beta_k^{(\circ)2}\}$$

سپس با انتقال جمله $P_\lambda(|\beta_k^{(\circ)}|)$ به طرف دوم معادله، عبارت زیر

$$P_\lambda(|\beta_k|) \approx P_\lambda(|\beta_k^{(\circ)}|) + \frac{1}{2} \frac{P'_\lambda(|\beta_k^{(\circ)}|)}{|\beta_k^{(\circ)}|} \{\beta_k^2 - \beta_k^{(\circ)2}\}.$$

به دست می‌آید. بنابراین،

$$\begin{aligned} \ell(\beta) &= \frac{1}{2} \|\mathbf{y} - \mathbf{X}\beta\|^2 + \sum_{k=1}^p P_\lambda(|\beta_k|) \\ &\approx \frac{1}{2} \|\mathbf{y} - \mathbf{X}\beta\|^2 + \sum_{k=1}^p \left\{ P_\lambda(|\beta_k^{(\circ)}|) + \frac{1}{2} \frac{P'_\lambda(|\beta_k^{(\circ)}|)}{|\beta_k^{(\circ)}|} \{\beta_k^2 - \beta_k^{(\circ)2}\} \right\} \\ &= \frac{1}{2} \|\mathbf{y} - \mathbf{X}\beta\|^2 + \sum_{k=1}^p \frac{1}{2} \frac{P'_\lambda(|\beta_k^{(\circ)}|)}{|\beta_k^{(\circ)}|} \beta_k^2 + c, \end{aligned}$$

که ثابت c تابعی از $\beta_k^{(\circ)}$ است. با استفاده از تقریب فوق و حذف جملات ثابت، تابع هدف (۹.۲) به یک تابع درجه دو تبدیل می‌شود. در نتیجه، برآورد کمترین توان‌های دوم تاوانیده را می‌توان به صورت

$$\beta^{(1)} = \left\{ \mathbf{X}^\top \mathbf{X} + \Sigma_\lambda(\beta^{(\circ)}) \right\}^{-1} \mathbf{X}^\top \mathbf{y}$$

به دست آورد که در آن

$$\Sigma_\lambda(\beta^{(\circ)}) = \text{diag} \left\{ \frac{P'_\lambda(|\beta_1^{(\circ)}|)}{|\beta_1^{(\circ)}|}, \dots, \frac{P'_\lambda(|\beta_p^{(\circ)}|)}{|\beta_p^{(\circ)}|} \right\}.$$

^{۳۹}Local quadratic approximation

ساختار فوق شبیه برآوردگر ریج است و با توجه به این واقعیت که رگرسیون ریج نمی‌تواند متغیرهای معنی‌دار را انتخاب کند، ضرایب رگرسیونی برآوردشده را که بسیار نزدیک به صفر می‌باشند، برابر صفر قرار می‌دهیم تا متغیرهای بی‌اهمیت از مدل حذف شوند.

یکی از مزایای تقریب LQA تعمیم‌پذیری آن به هر تابع زیان هموار دیگری غیر از تابع زیان کمترین توان‌های دوم است. با در نظر گرفتن تابع زیان $\ell(\beta)$ تابع هدف تاوانیده به صورت

$$\ell(\beta) + \sum_{k=1}^p P_{\lambda}(|\beta_k|) \quad (11.2)$$

تعریف می‌شود. در این حالت، ابتدا تابع زیان $\ell(\beta)$ را با یک تابع درجه دو به صورت

$$\ell(\beta) = \ell(\beta^{(0)}) + \nabla \ell^T(\beta^{(0)})(\beta - \beta^{(0)}) + \frac{1}{2}(\beta - \beta^{(0)})^T \nabla^2 \ell(\beta^{(0)})(\beta - \beta^{(0)}),$$

تقریب می‌کنیم که

$$\nabla \ell^T(\beta^{(0)}) = \frac{\partial \ell^T(\beta^{(0)})}{\partial \beta}, \quad \nabla^2 \ell(\beta^{(0)}) = \frac{\partial^2 \ell^T(\beta^{(0)})}{\partial \beta \partial \beta^T}.$$

با حذف جمله ثابت، تابع هدف (۱۱.۲) را می‌توان به صورت

$$\nabla \ell^T(\beta^{(0)})(\beta - \beta^{(0)}) + \frac{1}{2}(\beta - \beta^{(0)})^T \nabla^2 \ell(\beta^{(0)})(\beta - \beta^{(0)}) + \frac{1}{2}\beta^T \Sigma_{\lambda}(\beta^{(0)})\beta \quad (12.2)$$

نوشت و سپس با استفاده از الگوریتم نیوتن-رافسون جواب بهینه تابع (۱۱.۲) را به دست آورد. به طور دقیق‌تر، با اجرای مراحل زیر، الگوریتم LQA را می‌توان برای تقریب تابع (۱۱.۲) به کار برد.

- مرحله اول: مقدار اولیه $\beta^{(0)}$ را برای β انتخاب کنید. این مقدار اولیه می‌تواند برآورد OLS باشد.

- مرحله دوم: تابع هدف (۱۱.۲) را براساس (۱۲.۲) تقریب کنید.

- مرحله سوم: الگوریتم نیوتن-رافسون را برای محاسبه مقدار کمینه کننده تابع درجه دو بدست آمده در مرحله دوم به کار ببرید تا مقدار $\beta_k^{(0)}$ به روزسانی شود. سپس اگر $|\beta_k^{(0)}| < \epsilon$ ، متغیر X_k را حذف کنید.

- مرحله چهارم: مراحل دوم و سوم را تا رسیدن به همگرایی تکرار کنید.

فن و لی (۲۰۰۱) نشان دادند که در صورت انتخاب یک مقدار مناسب برای $\beta^{(0)}$ ، الگوریتم LQA پس از تعداد کمی تکرار به همگرایی می‌رسد. علاوه بر این، با یک مقدار اولیه مناسب، کارایی الگوریتم LQA با یک تکرار معادل با تکرار کامل الگوریتم تا رسیدن به همگرایی است. اما ایراد وارد بر این الگوریتم این است که اگر یک متغیر در هر تکرار حذف شود، در تکرار بعدی وارد مدل نمی‌شود و در نتیجه از مدل نهایی حذف می‌شود. بنابراین، این روش مشابه روش انتخاب پیشرو عمل می‌کند.

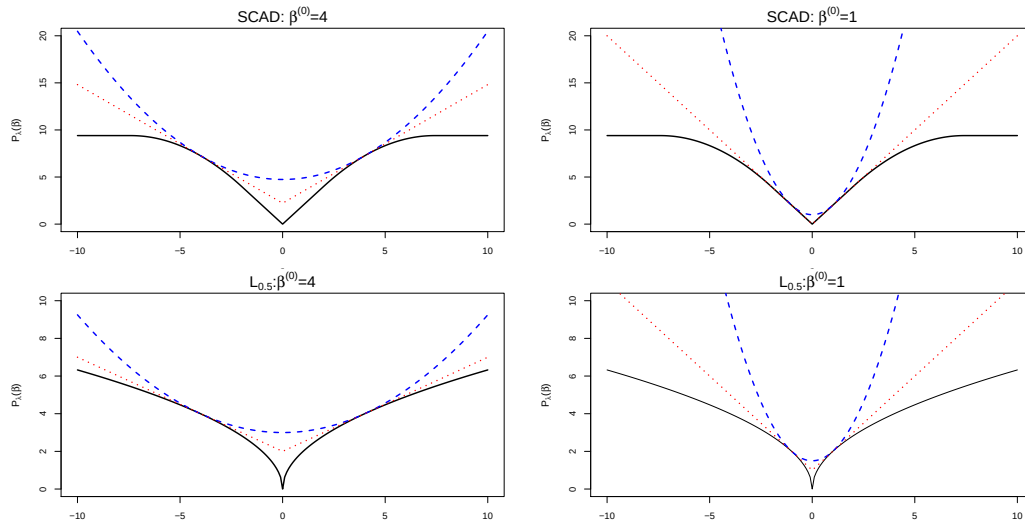
تقریب خطی موضعی: برای غلبه بر مشکل بیان‌شده در تقریب LQA ، ژو و لی (۲۰۰۸) الگوریتم تقریب خطی موضعی^{۴۰} (LLA) را برای تقریب توابع تاوان غیرمحدب پیشنهاد دادند. با استفاده از بسط تیلور $P_\lambda(\cdot)$ در نقطه $\beta_k^{(0)}$ داریم

$$P_\lambda(|\beta_k|) \approx P_\lambda(|\beta_k^{(0)}|) + P'_\lambda(|\beta_k^{(0)}|)\{|\beta_k| - |\beta_k^{(0)}|\}; \quad \beta_k \approx \beta_k^{(0)}.$$

بنابراین می‌توان تابع هدف (۹.۲) را به صورت

$$\begin{aligned} \ell(\beta) &= \frac{1}{2} \|\mathbf{y} - \mathbf{X}\beta\|^2 + \sum_{k=1}^p P_\lambda(|\beta_k|) \\ &\approx \frac{1}{2} \|\mathbf{y} - \mathbf{X}\beta\|^2 + \sum_{k=1}^p \{P_\lambda(|\beta_k^{(0)}|) + P'_\lambda(|\beta_k^{(0)}|)\{|\beta_k| - |\beta_k^{(0)}|\}\} \\ &= \frac{1}{2} \|\mathbf{y} - \mathbf{X}\beta\|^2 + \sum_{k=1}^p \frac{1}{2} P'_\lambda(|\beta_k^{(0)}|)|\beta_k| + c, \end{aligned}$$

تقریب نمود که ثابت c تابعی از $\beta^{(0)}$ است. تقریب LLA تابع هدف تاوانیده غیرمحدب را به تابع هدف لاسوی وزنی تبدیل می‌کند. این تابع هدف تبدیل‌یافته را می‌توان با استفاده از برنامه‌ریزی درجه دو یا الگوریتم کمترین زاویه حل کرد. شکل ۵.۲ تقریب LQA و LLA را برای توابع تاوان $L_{0.5}$ و اسکد نشان می‌دهد. در این شکل برای هر دو تابع $\lambda = 2$ و دو مقدار اولیه مختلف در نظر گرفته شده است.



شکل ۵.۲: تقریب درجه دو و خطی برای توابع جریمه $SCAD$ و $L_{0.5}$. منحنی آبی: تقریب درجه دو؛ منحنی قرمز: تقریب خطی؛ منحنی مشکی: تابع جریمه.

تاکنون، علاوه بر دو تقریب فوق، الگوریتم‌های دیگری نیز برای حل مسئله توابع هدف تاوانیده معرفی شده است. الگوریتم پیش‌برنده^{۴۱} (فو، ۱۹۹۸)، الگوریتم مختصات نزولی^{۴۲} برای لاسوی

⁴⁰Local linear approximation

⁴¹Shooting

⁴²Coordinate descent

گروهی (وو و لانگ ، ۲۰۰۸)، الگوریتم ماکزیم‌سازی شرطی تکراری^{۴۳} (ژانگ و لی ، ۲۰۱۱)، الگوریتم مختصات غیرنزولی برای توابع تاوان غیرمحدب (برهنی و هووانگ ، ۲۰۱۱) و الگوریتم مختصات نزولی گروهی (برهنی و هووانگ ، ۲۰۱۵) نمونه‌هایی از این الگوریتم‌ها می‌باشند.

۲.۲ برآوردگر ريج برای داده‌های طولی

همان‌گونه که در برخورد با داده‌های معمولی و بسیاری از کاربردهای رگرسیون کلاسیک با پدیده هم‌خطی روبرو هستیم، داده‌های طولی نیز از این امر مستثنی نیستند. استفاده از برآوردگر ريج در تحلیل داده‌های طولی پیشینه چندانى نداشته و تنها می‌توان کارهای ژانگ و هورواس (۲۰۰۶)، مالو و همکاران (۲۰۰۸)، الیوت و همکاران (۲۰۱۱) و رحمانی و همکاران (۲۰۱۸) را نام برد. لازم به ذکر است که در هیچ کدام از این موارد، مدل نیمه‌پارامتری با اثرات آمیخته بررسی نشده است و تحقیقات ما را می‌توان به عنوان تعمیمی از نتایج آن‌ها در نظر گرفت. بنابراین ضمن معرفی برآوردگر ريج، رفتار مجانبی برآوردگر معرفی شده ارائه شده است. در پایان نتایج عددی حاصل از مطالعات شبیه‌سازی و تحلیل داده‌های واقعی ارائه می‌گردد.

شرایطی را در نظر بگیرید که بین ستون‌های ماتریس طرح اثرات ثابت داده‌های طولی، هم‌خطی چندگانه وجود داشته باشد. در این حالت استفاده از برآوردگر رگرسیونی ريج برای β پیشنهاد می‌شود. با فرض برآوردگر برای تابع ناپارامتری در رابطه (۸.۱)، صورت تاوانیده رابطه (۱۰.۱) عبارت است از

$$l^*(\beta) = \hat{l}(\beta) + \lambda \beta^\top \beta.$$

در این حالت برآوردگر ريج هسته‌ای موزون به صورت

$$\begin{aligned} \hat{\beta}_{RK} &= \arg \min_{\beta} \{(\tilde{y} - \tilde{X}\beta)^\top V^{-1}(\tilde{y} - \tilde{X}\beta) + \lambda \beta^\top \beta\} \\ &= (\tilde{X}^\top V^{-1} \tilde{X} + \lambda I_{n_i})^{-1} \tilde{X}^\top V^{-1} \tilde{y} \\ &= R(\lambda) \hat{\beta}_K, \end{aligned}$$

حاصل می‌شود، که در آن $R(\lambda) = ((\tilde{X}^\top V^{-1} \tilde{X})^{-1} \lambda + I_{n_i})^{-1}$. واضح است که $f(t_i)$ و b_i مشابه حالت‌های قبل به صورت

$$\begin{aligned} \hat{f}(t_i) &= \hat{\mu}_Y(t_i) - \hat{\mu}_X^\top(t_i) \hat{\beta}_{RK}, \\ \hat{b}_i &= E[b_i | y_i] = D_i Z_i V_i^{-1} (\tilde{y}_i - \tilde{X}_i \hat{\beta}_{RK}) = D_i Z_i V_i^{-1} (y_i - X_i \hat{\beta}_{RK} - \hat{f}(t_i)) \end{aligned}$$

برآورد می‌شوند. به منظور انتخاب مقداری مناسب برای پارامتر تنظیم λ ، فن و لی (۲۰۰۱) معیار

⁴³Iterative conditional maximization

اعتبار سنجی متقابل تعمیم‌یافته^{۴۴} (GCV) را به صورت

$$GCV(\lambda) = \frac{RSS(\lambda)/n}{(1 - \text{tr}(d(\lambda))/n)^2}, \quad \hat{\lambda} = \arg \min_{\lambda \in \mathbb{R}^+} GCV(\lambda)$$

معرفی کردند، که در آن

$$RSS(\lambda) = (\tilde{y} - \tilde{X}\hat{\beta}_R)^\top V^{-1}(\tilde{y} - \tilde{X}\hat{\beta}_R),$$

$$d(\lambda) = \tilde{X}(\tilde{X}^\top V^{-1} \tilde{X} + \lambda I)^{-1} \tilde{X}^\top V^{-1} + Z D Z^\top (I - \tilde{X}[(\tilde{X}^\top V^{-1} \tilde{X} + \lambda I)^{-1} \tilde{X}^\top V^{-1}]).$$

از آنجایی که برآوردگر ريج مضربی از برآوردگر هسته‌ای ($\hat{\beta}_K$) است، به راحتی می‌توان قضیه زیر را از قضیه ۱۰.۵.۱ نتیجه گرفت.

قضیه ۱۰.۲.۲. تحت مفروضات این بخش وقتی $n \rightarrow \infty$ ، آنگاه برآوردگر ريج دارای توزیع مجانبی زیر است.

$$\sqrt{n} \{ \hat{\beta}_R - R(\lambda)(\beta + h^\top b_K(\beta, f)/2) \} \xrightarrow{d} \mathcal{N}_p(0, R(\lambda) V_K R^\top(\lambda)).$$

که در آن مولفه‌های $b_K(\beta, f)$ و V_K در قضیه ۱۰.۵.۱ آمده است.

از آنجایی که برآوردگر معرفی شده در تعامل با ماتریس کوواریانس پاسخ‌ها می‌باشد، به منظور بهبود کارایی برآوردگرها از یک الگوریتم تکراری مشابه الگوریتم ۱ بین برآورد ضرایب رگرسیونی و برآورد ماتریس کوواریانس استفاده می‌کنیم.

۱.۲.۲ مطالعات عددی

۱.۱.۲.۲ مطالعات شبیه‌سازی

در این قسمت یک مطالعه شبیه‌سازی جهت مقایسه برآوردگرهای هسته‌ای و هسته‌ای ريج در شرایطی که متغیرهای تبیینی دارای هم‌خطی باشند، انجام می‌گیرد. داده‌ها از مدل

$$y_i(t_{ij}) = x_{1,i}(t_{ij})\beta_1 + x_{2,i}(t_{ij})\beta_2 + x_{3,i}(t_{ij})\beta_3 + x_{4,i}(t_{ij})\beta_4 + x_{5,i}(t_{ij})\beta_5 \\ + f(t_{ij}) + b_i + \varepsilon_i(t_{ij}),$$

شبیه‌سازی شده‌اند، که در آن $\beta^T = (0.5, 1, 1.5, 2, 0.1, 0.2)$. سایر مولفه‌ها به صورت

$$f(t_{ij}) = 2 \sin(2\pi t_{ij}); \quad t_{ij} \sim \mathcal{U}(0, 1),$$

$$b_i \sim \mathcal{N}(0, \sigma_b^2); \quad \sigma_b^2 = 0.25,$$

$$\varepsilon_{ij} \sim \mathcal{N}(0, \sigma^2(t)); \quad \sigma^2(t) = \exp\left(\frac{t}{1.5}\right).$$

⁴⁴Generalized cross validation

جهت ایجاد شرایط هم خطی، متغیرهای تبیینی به کمک رابطه

$$x_{ik} = (1 - \gamma^2)^{\frac{1}{2}} \omega_{ik} + \gamma^2 \omega_{ip}, \quad i = 1, \dots, n = 50, \quad k = 1, \dots, p = 5,$$

تولید می‌شوند، به طوری که ω_{ik} ها مستقل از هم بوده و دارای توزیع نرمال استاندارد هستند. از آنجایی که میزان همبستگی بین هر دو متغیر تبیینی γ^2 می‌باشد به منظور رسیدن به درجات مختلف همبستگی مقادیر $0.7, 0.8, 0.9, 0.99$ را برای γ در نظر گرفتیم. به منظور بررسی تاثیر اعمال ساختار همبستگی نادرست بر عملکرد برآوردگرها، ابتدا $\hat{\beta}_{RK}$ و $\hat{\beta}_K$ را تحت ساختار همبستگی واقعی $AR(1)$ برآورد کرده و به ترتیب با نماد $RK - C$ و $K - C$ نشان دادیم. سپس برآوردگرها را تحت ساختار همبستگی نادرست تبادل پذیر برآورد کرده و به ترتیب با نماد $RK - I$ و $K - I$ نشان دادیم.

جدول ۲.۲ نتایج مربوط به مقادیر MSE ، اریبی و SD برآوردگرهای مذکور بر اساس ۵۰۰ تکرار از شبیه سازی بیان شده را گزارش می‌کند. از نتایج جداول می‌توان نتیجه گرفت که برآوردگر ریج هسته‌ای، MSE کمتری نسبت به برآوردگر هسته‌ای دارد حتی هنگامی که ساختار همبستگی به درستی تعیین نشود. برآوردگر ریج هسته‌ای از نظر معیارهای اریبی و SD عملکرد بهتری نسبت به برآوردگر هسته‌ای دارد. علاوه بر این، هنگامی که میزان هم خطی بین متغیرهای تبیینی افزایش می‌یابد، کارایی برآوردگر هسته‌ای در مقایسه با برآوردگر ریج هسته‌ای کاهش می‌یابد.

۲.۱.۲.۲ تحلیل داده‌های CD^4

در این بخش به مطالعه رفتار برآوردگرهای هسته‌ای و ریج هسته‌ای معرفی شده با استفاده از داده‌های CD^4 می‌پردازیم.

برآورد پارامترهای رگرسیونی به همراه SD آن‌ها در جدول ۳.۲ ارائه شده است. برای محاسبه مقادیر SD از روش اعتبارسنجی یک طرفه^{۴۵} استفاده شد. این نتایج نشان می‌دهد که روش ریج همواره برآوردهایی با مقادیر SD کمتر را نتیجه می‌دهد.

۳.۲ انتخاب متغیر در مدل نیمه پارامتری با اثرات آمیخته برای داده‌های طولی نرمال با بعد بالا

در بسیاری از مطالعات طولی با تعداد فراوانی از متغیرها روبرو هستیم که عموماً وارد کردن همه آن‌ها در تحلیل موجب انحراف مدل‌های معمول از الگوی واقعی داده‌ها می‌شود. بنابراین مهم است که روش‌های انتخاب متغیر و برآورد برای داده‌های طولی نیز توسعه یابد. پس از آن که فو (۲۰۰۳) و فن و لی (۲۰۰۴) به تعمیم روش مبتنی بر تابع تاوان، به ترتیب در مدل‌های خطی تعمیم یافته و مدل‌های نیمه پارامتری با داده‌های طولی پرداختند، مسئله انتخاب متغیر در زمینه داده‌های طولی با پیشرفت چمشگیری روبرو شد. در این میان می‌توان به کارهای کانتونی و همکاران (۲۰۰۵)، زیاک

⁴⁵Leave-one-out cross-validation

جدول ۲.۲: نتایج شبیه‌سازی: مقادیر MSE و اریبی (SD) برآوردهای $\hat{\beta}_{RK}$ و $\hat{\beta}_K$ به ازای γ ‌های مختلف.

روش‌ها	پارامترها	$\gamma = 0.70$	$\gamma = 0.80$	$\gamma = 0.90$	$\gamma = 0.99$
$RK - C$	β_1	0.316(0.400)	0.376(0.475)	0.514(0.649)	0.1401(0.1748)
	β_2	0.310(0.390)	0.368(0.464)	0.506(0.635)	0.1446(0.1706)
	β_3	0.282(0.355)	0.333(0.422)	0.452(0.577)	0.1359(0.1572)
	β_4	0.306(0.388)	0.364(0.461)	0.498(0.630)	0.1397(0.1745)
	β_5	0.351(0.454)	0.454(0.592)	0.690(0.910)	0.2208(0.2644)
	MSE	0.079	0.118	0.238	0.2103
$RK - I$	β_1	0.337(0.414)	0.402(0.492)	0.553(0.677)	0.1686(0.2059)
	β_2	0.330(0.408)	0.393(0.484)	0.540(0.665)	0.1624(0.2011)
	β_3	0.295(0.375)	0.351(0.446)	0.482(0.613)	0.1453(0.1847)
	β_4	0.362(0.452)	0.431(0.537)	0.593(0.739)	0.1800(0.2236)
	β_5	0.342(0.427)	0.444(0.556)	0.687(0.865)	0.2518(0.3219)
	MSE	0.087	0.128	0.258	0.2704
$K - C$	β_1	0.315(0.400)	0.375(0.476)	0.517(0.656)	0.1596(0.2027)
	β_2	0.313(0.390)	0.373(0.464)	0.513(0.639)	0.1585(0.1975)
	β_3	0.296(0.358)	0.352(0.426)	0.485(0.586)	0.1499(0.1811)
	β_4	0.308(0.389)	0.366(0.463)	0.504(0.638)	0.1558(0.1971)
	β_5	0.356(0.456)	0.463(0.596)	0.720(0.930)	0.2782(0.3608)
	MSE	0.081	0.121	0.249	0.2871
$K - I$	β_1	0.337(0.414)	0.402(0.493)	0.553(0.679)	0.1708(0.2097)
	β_2	0.332(0.407)	0.395(0.485)	0.543(0.667)	0.1679(0.2062)
	β_3	0.296(0.376)	0.353(0.447)	0.485(0.615)	0.1500(0.1902)
	β_4	0.362(0.452)	0.431(0.538)	0.594(0.740)	0.1835(0.2288)
	β_5	0.345(0.428)	0.444(0.557)	0.688(0.868)	0.2633(0.3380)
	MSE	0.087	0.129	0.261	0.2919

(۲۰۰۶)، وانگ و کوای (۲۰۰۹) و سوای و همکاران (۲۰۱۰) در مدل‌های خطی تعمیم‌یافته، باندل و همکاران (۲۰۱۰) در مدل اثرات آمیخته، نی و همکاران (۲۰۱۰) و ما و همکاران (۲۰۱۳) در مدل‌های نیمه‌پارامتری اشاره کرد. رخداد مسئله بعد بالا در داده‌های طولی نیز غیر قابل اجتناب است. با این وجود، مسئله انتخاب متغیر برای داده‌های طولی با بعد بالا به واسطه چالش‌های تحمیل‌شده از جمله وجود همبستگی درون گروهی پیشینه‌چندانی نداشته و تنها می‌توان کارهای سو و همکاران (۲۰۱۳) و وانگ و همکاران (۲۰۱۲) را نام برد. اخیراً کاظمی و همکاران (۱۳۹۷)، در مدل‌های خطی-جزئی با بعد بسیار بالا مسئله انتخاب متغیر را بررسی کرده است. در عین حال، تحقیقات در زمینه انتخاب متغیر داده‌های طولی با بعد بالا در مدل‌های رگرسیونی نیمه‌پارامتری با اثرات آمیخته تا حد زیادی ناشناخته مانده است. با توجه به اهمیت مسئله بیان‌شده در داده‌های طولی با بعد بالا، در این بخش، برآورد و انتخاب متغیر در مدل نیمه‌پارامتری با اثرات آمیخته را مبتنی

جدول ۳.۲: نتایج تحلیل داده‌های $CD4$: برآورد پارامترهای رگرسیونی (SD) تحت برآوردهای مختلف.

	روش			
	ریج اسپلین	ریج هسته‌ای	اسپلین	هسته‌ای
AGE	۰/۰۸۱۱(۰/۰۰۰۱)	۰/۰۲۰۸(۰/۰۰۲۱)	۰/۰۸۵۱(۰/۰۰۰۲)	۰/۰۲۱۲(۰/۰۰۲۱)
SMOKE	۰/۷۸۲۳(۰/۰۰۳۱)	۰/۵۸۰۸(۰/۰۰۸۰)	۰/۷۲۹۰(۰/۰۰۶۰)	۰/۵۹۴۱(۰/۰۰۸۲)
DRUG	۰/۸۴۰۱(۰/۰۱۴۲)	۰/۴۵۸۲(۰/۰۱۶۶)	۰/۸۰۱۱(۰/۰۱۴۳)	۰/۵۳۵۰(۰/۰۱۹۴)
SEXP	۰/۰۷۹۰(۰/۰۰۱۳)	۰/۰۵۹۷(۰/۰۰۲۴)	۰/۰۷۱۱(۰/۰۰۱۴)	۰/۰۵۸۵(۰/۰۰۲۴)
CESD	-۰/۰۵۴۲(۰/۰۰۰۳)	-۰/۰۵۳۱(۰/۰۰۰۹)	-۰/۰۶۸۶(۰/۰۰۰۳)	-۰/۰۵۳۲(۰/۰۰۰۹)

بر رهیافت توابع تاوان، در نظر گرفتیم. روش پیشنهادی به علت وجود توام اثرات تصادفی و مولفه ناپارامتری در مدل، نسبتاً با کارهای انجام شده متفاوت است. از آنجایی که برآوردهای معرفی شده در تعامل با ماتریس کوواریانس پاسخ‌ها و اثرات تصادفی می‌باشند، به منظور بهبود کارایی برآوردهای بیشینه درست‌نمایی توانیده (PML) از الگوریتم تکراری EM استفاده می‌کنیم. علاوه بر این، خواص مجانبی برآوردها در چارچوب داده‌های بعد بالا که در آن تعداد پارامترها متناسب با افزایش حجم نمونه افزایش می‌یابد، و نحوه انتخاب پارامترهای تابع تاوان مورد بررسی قرار می‌گیرد. به منظور بررسی عملکرد روش معرفی شده، مطالعه‌ای شبیه‌سازی و همچنین تحلیل یک مجموعه داده واقعی فراهم آمده است.

فرض کنید یک نمونه تصادفی n تایی از داده‌های طولی با بعد بالا در اختیار داشته باشیم. مدل رگرسیونی نیمه پارامتری با اثرات آمیخته برای داده‌های طولی که در بخش ۴.۳.۱ شرح داده شد را به یاد آورید. همچنین فرض کنید مولفه ناپارامتری مدل به روش اسپلین‌های رگرسیونی شرح داده شده در ۳.۵.۱ تقریب زده شود. بدین ترتیب یک مدل رگرسیون خطی به صورت

$$y_i = \tilde{X}_i \theta + Z_i b_i + \varepsilon_i, \quad (13.2)$$

خواهیم داشت، به طوری که $\tilde{X}_i$ ماتریس طرح ادغام شده $(p_n + h_n) \times n_i$ بعدی از ماتریس اثرات ثابت و اثرات اسپلین و θ بردار پارامتر ادغام شده $1 \times (p_n + h_n)$ بعدی می‌باشند. به طوری که $h_n = d + K_n + 1$ ، d درجه تقریب اسپلین و K_n تعداد گره‌ها در تقریب اسپلین است. حال از طریق بهینه‌سازی تابع هدف مبتنی بر تابع تاوان مانند بیشینه درست‌نمایی توانیده PML به برآورد و انتخاب متغیر می‌پردازیم.

از آنجایی که در مبحث داده‌های با بعد بالا فرض می‌شود که تعداد پارامترها (p) متناسب با افزایش حجم نمونه افزایش می‌یابد، معمولاً p را با نماد p_n نمایش می‌دهند. متعاقباً نمادهای h ، K و λ به صورت h_n ، K_n و λ_n تغییر می‌کنند.

۱.۳.۲ الگوریتم EM برای بیشینه درستنمایی تاوانیده

۱.۱.۳.۲ برآوردگر بیشینه درستنمایی

در روش برآورد ML ، اگر بردار مشاهدات غیرکامل باشد، یعنی برخی از داده‌ها به دلیل سانسور شدن یا مقادیر گمشده در دسترس نباشند، می‌توان از روش‌های تکرارشونده همچون روش نیوتن-رافسون و الگوریتم امیدریاضی بیشینه‌سازی (دمپستر و همکاران، ۱۹۷۷) یا به اختصار الگوریتم EM استفاده کرد. در هر تکرار از الگوریتم دو گام وجود دارد. گام E که امید ریاضی شرطی براساس توزیع شرطی لگاریتم تابع درستنمایی داده‌های کامل به شرط داده‌های مشاهده‌شده را محاسبه می‌کند و گام M که برآورد بیشینه درستنمایی پارامترهای مدل را پیدا می‌کند. ما نیز برای برآورد پارامترهای مدل (13.2) از روش مذکور استفاده می‌کنیم.

فرض کنید $\Theta = (\theta, D, \Sigma)$ مجموعه‌ای از تمام پارامترهای مدل باشد. اگر اثرات تصادفی b_i را به عنوان مقادیر گمشده و مجموعه $\{(y_i, b_i), i = 1, \dots, n\}$ را داده‌های کامل در نظر بگیریم، آنگاه لگاریتم تابع درستنمایی بدون در نظر گرفتن جمله ثابت به صورت

$$\begin{aligned} \ell_c(\Theta) &= \sum_{i=1}^n \ell_c^{[i]}(\Theta) \\ &= \sum_{i=1}^n \frac{1}{2} \left\{ \log |\Sigma_i^{-1}| + \log |D_i^{-1}| - [\varepsilon_i^\top \Sigma_i^{-1} \varepsilon_i + b_i^\top D_i^{-1} b_i] \right\} \end{aligned} \quad (14.2)$$

است. برای ارزیابی امید شرطی (14.2) به شرط داده‌های مشاهده‌شده $y = \{y_i, i = 1, \dots, n\}$ و برآوردهای جاری $\hat{\Theta}^{(r)} = (\hat{\theta}^{(r)}, \hat{D}^{(r)}, \hat{\Sigma}^{(r)})$ ، ما از نتیجه زیر استفاده می‌کنیم.

نتیجه ۱.۳.۲.

$$\begin{aligned} \begin{bmatrix} y_i \\ b_i \end{bmatrix} &\sim \mathcal{N}_{n_i+q} \left(\begin{bmatrix} \tilde{X}_i \theta \\ 0 \end{bmatrix}, \begin{bmatrix} V_i & Z_i D_i \\ D_i Z_i^\top & D_i \end{bmatrix} \right), \\ b_i | y_i &\sim \mathcal{N}_q \left(D_i Z_i^\top V_i^{-1} (y_i - \tilde{X}_i \theta), (D_i^{-1} + Z_i^\top \Sigma_i^{-1} Z_i)^{-1} \right). \end{aligned}$$

مراحل اصلی الگوریتم EM به شرح زیر انجام می‌شود:

گام E : محاسبه تابع $Q(\Theta | \hat{\Theta}^{(r)}) = E(\ell_c(\Theta) | y, \hat{\Theta}^{(r)})$ که مجموع $Q_i(\Theta | \hat{\Theta}^{(r)})$ برای $i = 1, \dots, n$ است و به صورت

$$Q_i(\Theta | \hat{\Theta}^{(r)}) = \frac{1}{2} \left\{ \log |\Sigma_i^{-1}| + \log |D_i^{-1}| - \text{tr}(D_i^{-1} \hat{B}_i^{(r)}) - \text{tr}(\Sigma_i^{-1} \hat{\Psi}_i^{(r)}(\theta)) \right\},$$

تعریف می‌شوند، به طوری که بر اساس نتیجه ۱.۳.۲ داریم:

$$\begin{aligned} \hat{B}_i^{(r)} &= E(b_i b_i^\top | y_i, \hat{\Theta}^{(r)}) = \hat{b}_i^{(r)} \hat{b}_i^{(r)\top} + \hat{V}_{b_i}^{(r)}, \\ \hat{\Psi}_i^{(r)}(\beta) &= E(e_i e_i^\top | y_i, \hat{\Theta}^{(r)}) \\ &= (y_i - \tilde{X}_i \theta - Z_i \hat{b}_i^{(r)})(y_i - \tilde{X}_i \theta - Z_i \hat{b}_i^{(r)})^\top + Z_i \hat{V}_{b_i}^{(r)} Z_i^\top, \end{aligned}$$

$$\begin{aligned}\hat{\mathbf{b}}_i^{(r)} &= E(\mathbf{b}_i | \mathbf{y}_i, \hat{\boldsymbol{\Theta}}^{(r)}) = \hat{\mathbf{D}}_i^{(r)} \mathbf{Z}_i^\top \hat{\mathbf{V}}_i^{(r)-1} (\mathbf{y}_i - \tilde{\mathbf{X}}_i \hat{\boldsymbol{\theta}}^{(r)}), \\ \hat{\mathbf{V}}_{\mathbf{b}_i}^{(r)} &= \text{Cov}(\mathbf{b}_i | \mathbf{y}_i, \hat{\boldsymbol{\Theta}}^{(r)}) = (\hat{\mathbf{D}}_i^{(r)-1} + \mathbf{Z}_i^\top \hat{\boldsymbol{\Sigma}}_i^{(r)-1} \mathbf{Z}_i)^{-1},\end{aligned}$$

$$\cdot \mathbf{e}_i = (\mathbf{y}_i - \tilde{\mathbf{X}}_i \boldsymbol{\theta}^{(r)}) \text{ و}$$

گام M : به روز کردن $\hat{\boldsymbol{\Theta}}^{(r)}$ توسط $Q(\boldsymbol{\Theta} | \hat{\boldsymbol{\Theta}}^{(r)}) = \max_{\boldsymbol{\Theta}} Q(\boldsymbol{\Theta} | \hat{\boldsymbol{\Theta}}^{(r)})$.

۲.۱.۳.۲ برآوردگر PML و الگوریتم EM

به منظور انتخاب متغیرهای مهم و برآورد ضرایب آن‌ها، به لگاریتم تابع درستنمایی (۱۴۰۲)، جمله توان $\sum_{k=1}^p P_{\lambda_n}(|\beta_k|)$ را اضافه می‌کنیم، آنگاه لگاریتم تابع درستنمایی تاوانیده به صورت

$$\ell_c^{\text{Pen}}(\boldsymbol{\Theta}) = \ell_c(\boldsymbol{\Theta}) - n \sum_{k=1}^p P_{\lambda_n}(|\beta_k|). \quad (۱۵.۲)$$

حاصل می‌شود. به طور مشابه تابع Q تاوانیده در گام E را به صورت

$$Q_i^{\text{Pen}}(\boldsymbol{\Theta} | \hat{\boldsymbol{\Theta}}^{(r)}) = Q_i(\boldsymbol{\Theta} | \hat{\boldsymbol{\Theta}}^{(r)}) - n \sum_{k=1}^{p_n} P_{\lambda_n}(|\beta_k|).$$

در نظر می‌گیریم. در ادامه فرض می‌کنیم تابع تاوان مورد استفاده در این جا و کل رساله، تابع تاوان اسکد، رابطه (۸.۲) است و برای حل مسئله بهینه‌سازی در گام M ، تقریب درجه دو موضعی (LQA) را به کار می‌بریم. در مجموع گام M الگوریتم EM به صورت الگوریتم ۲ پیشنهاد می‌شود.

۳.۱.۳.۲ انتخاب پارامترهای تاوان

برای پیاده‌سازی روش پیشنهادشده، باید درباره چندین پارامتر تصمیم‌گیری مناسبی اتخاذ شود. این پارامترها شامل درجه چند جمله‌ای d ، تعداد نقاط گره K_n و تعداد توابع پایه h_n در تقریب B -اسپلاین و پارامترهای a و λ_n در تابع تاوان اسکد هستند.

در تقریب اسپلاین، اگر $d = 3$ ، آنگاه تابع مورد اشاره را اسپلاین مکعبی می‌نامند. اسپلاین مکعبی یکی از متداول‌ترین نوع اسپلاین‌ها است و به دلیل انعطاف‌پذیری بالایی که دارد، در بسیاری از علوم مورد استفاده قرار می‌گیرد، لذا در قسمت مطالعات عددی از این نوع اسپلاین استفاده می‌شود. این نکته که تعداد گره‌ها با افزایش حجم نمونه افزایش یابند اهمیت داشته و از سوی دیگر تعداد زیاد گره‌ها ممکن است واریانس برآوردگرها را افزایش دهد. بنابراین، تعداد گره‌ها باید به درستی انتخاب شود تا بین اریبی و واریانس برآوردگرها تعادل ایجاد گردد. برای سهولت در محاسبات معمولاً گره‌هایی با فواصل برابر و تعداد گره‌ها به صورت $K_n \approx n^{1/(2r+1)}$ انتخاب می‌شود. این استراتژی در بسیاری از متون آماری از جمله هی و همکاران (۲۰۰۲)، چن و چو (۲۰۰۷) و سینها و ستار (۲۰۱۵)

الگوریتم ۲ برآوردگر PML و الگوریتم EM

۱: گام M : به‌روز کردن $\hat{\beta}^{(r)}$ ، $\hat{D}^{(r)}$ و $\hat{\Sigma}^{(r)}$ از طریق بیشینه کردن $Q(\Theta|\hat{\Theta}^{(r)})$ که به‌صورت زیر حاصل می‌شوند:

$$\begin{aligned}\hat{\theta}^{(r+1)} &= \left(\sum_{i=1}^n \tilde{X}_i^\top \hat{\Sigma}_i^{(r)-1} \tilde{X}_i + n E_n(\hat{\theta}^{(r)}) \right)^{-1} \sum_{i=1}^n \tilde{X}_i^\top \hat{\Sigma}_i^{(r)-1} (y_i - Z_i \hat{b}_i^{(r)}), \\ \hat{D}^{(r+1)} &= n^{-1} \sum_{i=1}^n \hat{B}_i^{(r)}, \\ \hat{\Sigma}^{(r+1)} &= n^{-1} \sum_{i=1}^n \hat{\Psi}_i^{(r)}(\beta),\end{aligned}$$

که در آن

$$E_n(\hat{\theta}^{(r)}) = \text{diag} \left\{ \frac{q_{\lambda_n}(|\beta_\lambda|)}{\epsilon + |\beta_\lambda|}, \dots, \frac{q_{\lambda_n}(|\beta_p|)}{\epsilon + |\beta_p|}, \mathbf{0}_{h_n} \right\}, \quad \epsilon = 10^{-\epsilon},$$

و $\mathbf{0}_{h_n}$ برداری h_n بعدی با عناصر صفر است.

۲: گام هم‌گرایی: گام‌های E و M را تا زمان هم‌گرایی تکرار می‌کنیم تا برآوردگر PML ، $\hat{\Theta} = (\hat{\theta}, \hat{D}, \hat{\Sigma})$ ، به‌دست آید.

اجرا شده است. واضح است که $r > 0$ زیرا $r = 0$ تعداد n گره را در نظر می‌گیرد که انتخاب درستی نمی‌باشد. معمولاً r را برابر ۱ در نظر می‌گیرند؛ به‌طور مثال برای نمونه‌هایی به حجم 500 و 1000 ، تعداد گره‌ها به‌ترتیب برابر ۴، ۵ و ۶ خواهد بود. بنابراین اگر اسپلین را مکعبی و تعداد گره‌ها را به‌صورت $K_n \approx n^{1/(2r+1)}$ در نظر بگیریم، تعداد توابع پایه به‌صورت $4 + n^{1/(2r+1)} \approx h_n$ است؛ که برای حجم نمونه‌های 500 ، 1000 و 2000 به‌ترتیب ۸، ۹ و ۱۰ تابع پایه خواهیم داشت. در بخش ۳.۳.۱.۲ به منظور انتخاب مقداری مناسب برای پارامتر تنظیم λ_n ، تکنیک‌های متعدد پیشنهادی توسط افراد مختلف، معرفی شده‌اند. همچنین فن و لی (۲۰۰۱) از دیدگاه بیزی مقدار $a = 3/7$ را به عنوان یک مقدار مناسب برای مسائل مختلف پیشنهاد دادند. در اینجا از معیار اعتبار سنجی متقابل تعمیم‌یافته (GCV) استفاده می‌کنیم که به‌صورت

$$GCV(\lambda_n) = \frac{RSS(\lambda_n)/n}{(1 - \text{tr}(d(\lambda_n))/n)^2}, \quad \hat{\lambda}_n = \arg \min_{\lambda_n \in \text{Re}^+} GCV(\lambda_n),$$

است و در آن

$$\begin{aligned}RSS(\lambda_n) &= (y - \tilde{X}^\top \hat{\theta} - Z^\top \hat{b})^\top \hat{V}^{-1} (y - \tilde{X}^\top \hat{\theta} - Z^\top \hat{b}), \\ d(\lambda_n) &= X(X^\top \hat{V}^{-1} X + \lambda_n I)^{-1} X^\top \hat{V}^{-1} \\ &\quad + Z \hat{D} Z^\top (I - X[(X^\top \hat{V}^{-1} X + \lambda I)^{-1} X^\top \hat{V}^{-1}]).\end{aligned}$$

۲.۳.۲ خواص مجانبی

در این بخش به بیان توزیع مجانبی برآوردهای $\hat{\beta}_n$ و $\hat{f}$ می‌پردازیم. در اکثر متون آماری در حیطه انتخاب متغیر به منظور سادگی در محاسبات، پارامترهای درست $\beta_{n\circ}$ به صورت $\beta_{n\circ} = (\beta_{n\circ 1}^\top, \beta_{n\circ 2}^\top)^\top$ و ماتریس طرح مدل به صورت $X_i = (X_{i(1)}, X_{i(2)})$ تفکیک می‌شوند. بدون کاستن از کلیت مسئله فرض می‌شود $\beta_{n\circ 2} = 0$ و تمام اعضای $\beta_{n\circ 1}$ که برداری s_n^* بعدی است غیرصفر باشد. در اینجا، ضرایب رگرسیونی درست به صورت $\theta_{n\circ} = (\beta_{n\circ 1}^\top, \beta_{n\circ 2}^\top, \alpha_{n\circ}^\top)^\top$ هستند که در آن $\alpha_{n\circ}$ برداری h_n بعدی و وابسته به تابع ناپارامتری f است. به منظور جداسازی فضای پارامتر به فضای صفر و غیرصفر، ما $\theta_{n\circ}$ را به صورت $(\theta_{n\circ 1}^\top, \theta_{n\circ 2}^\top)^\top$ بازتفکیک می‌کنیم، به طوری که $\theta_{n\circ 1} = (\beta_{n\circ 1}^\top, \alpha_{n\circ}^\top)^\top$ یک بردار $(s_n = s_n^* + h_n)$ بعدی است که تمام اعضای آن غیرصفراند. متعاقباً، مقادیر برآورد و ماتریس طرح به ترتیب به صورت $\hat{\theta}_n = (\hat{\theta}_{n1}^\top, \hat{\theta}_{n2}^\top)^\top$ و $\tilde{X}_i = (\tilde{X}_{i(1)}^\top, \tilde{X}_{i(2)}^\top)^\top$ بازتفکیک می‌شوند، که در آن $\hat{\theta}_{n1} = (\hat{\beta}_{n1}^\top, \hat{\alpha}_n^\top)^\top$ ، $\tilde{X}_{i(1)} = (X_{i(1)}^\top, B_i(t_i)^\top)^\top$ ، $\tilde{X}_{i(2)} = X_{i(2)}^\top$ و $\hat{\theta}_{n2} = \hat{\beta}_{n2}$ فرض کنید $\beta_{n\circ, k} \neq 0$ ، $\beta_{n\circ, k} \neq 0$ و $a_n = \max_{1 \leq k \leq p_n} \{P'_{\lambda_n}(|\beta_{n\circ, k}|)\}$ ، $\beta_{n\circ, k} \neq 0$ و $b_n = \max_{1 \leq k \leq p_n} \{P''_{\lambda_n}(|\beta_{n\circ, k}|)\}$ ، $\beta_{n\circ, k} \neq 0$ می‌باشد.

در نظر داشته باشید که اگر گام M الگوریتم EM دارای چندین جواب باشد، تنها دنباله‌ای از برآوردگر سازگار $\hat{\theta}_n$ را در نظر می‌گیریم. دنباله $\hat{\theta}_n$ را سازگار گوئیم اگر $n \rightarrow \infty$ آنگاه $\hat{\beta}_n - \beta_{n\circ} \rightarrow 0$ و $\sup_t |B^\top(t)\hat{\alpha}_n - f_\circ(t)| \xrightarrow{p} 0$.

علاوه بر این، فرض کنید به ازای هر i ، $n_i = m < \infty$ باشد؛ مجموعه سه تایی (y_i, X_i, t_i) ، $i = 1, \dots, n$ متغیرهایی مستقل و هم توزیع باشند و مشتق k ام تابعی مانند $d(\cdot)$ را به صورت $d^{(k)}(\cdot)$ تعریف می‌کنیم. برای مطالعه رفتار حدی برآوردگرها، شرایط نظم زیر لازم می‌باشد:

(۱.A) به ازای هر $k \geq 2$ ، مشتق k ام $f_\circ(t_i)$ کراندار است.

(۲.A) ماتریس کوواریانس برآورد شده $\hat{\Sigma}_i$ در شرط $\|\hat{\Sigma}_i - \Sigma_i\| = O_p(n^{-1/2})$ صدق می‌کند.

(۳.A)

$$b_1 \leq \lambda_{\min}(n^{-1} \sum_{i=1}^n \tilde{X}_i^\top \Sigma_i^{-1} \tilde{X}_i) \leq \lambda_{\max}(n^{-1} \sum_{i=1}^n \tilde{X}_i^\top \Sigma_i^{-1} \tilde{X}_i) \leq b_2,$$

به طوری که b_1 و b_2 مقادیری ثابت و مثبت و $\lambda_{\min}$ و $\lambda_{\max}$ به ترتیب کوچکترین و بزرگترین مقدار ویژه ماتریس هستند.

(۴.A) وقتی $n \rightarrow \infty$ ، آنگاه $\min_{1 \leq k \leq s_n} |\theta_{n\circ k}|/\lambda_n \rightarrow \infty$ ، $\lambda_n \rightarrow 0$ ، $s_n^{-1} = o(1)$ ، $s_n^\top(\log n)^4 = o(n^{-1})$ و $p_n s_n^\top(\log n)^4 = o(n^2 \lambda_n^4)$ و $p_n s_n^\top(\log n)^6 = o(n^2 \lambda_n^2)$ ، $\log(p_n) = o(n \lambda_n^2/(\log n))$ ، $o(n \lambda_n^2)$.

(۵.A) تابع تاوان در شرایط $\liminf_{n \rightarrow \infty} \liminf_{\theta \rightarrow 0^+} p'_{\lambda_n}(\theta)/\lambda_n > 0$ ، $a_n = O(n^{-1/2})$ ، $b_n \rightarrow \infty$ ، $\theta_1, \theta_2 > \lambda_n C$ وجود دارند به طوری که به ازای هر θ_1, θ_2 خواهیم داشت $|P''_{\lambda_n}(\theta_1) - P''_{\lambda_n}(\theta_2)| \leq D|\theta_1 - \theta_2|$.

قضیه ۱.۳.۲. با در نظر گرفتن شرایط بیان شده، وقتی $n \rightarrow \infty$ و $n^{-1}p_n^2 = o(1)$ ، آنگاه $\hat{\theta}_n$ یک جواب الگوریتم EM است، به طوری که

$$(i) \quad \|\hat{\theta}_n - \theta_{n^0}\| = O_p(\sqrt{p_n}(n^{-1/2} + a_n)),$$

$$(ii) \quad \frac{1}{n} \sum_{i=1}^n \sum_{j=1}^{n_i} (\hat{g}(t_{ij}) - g^0(t_{ij}))^2 = O_p(n^{-2m/(2m+1)}).$$

قضیه ۲.۳.۲. با در نظر گرفتن شرایط بیان شده، وقتی $n \rightarrow \infty$ ، اگر $K_n = O_p(n^{1/(2r+1)})$ و $n^{-1}p_n^2 = o(1)$ ، آنگاه

$$(i) \quad \mathbb{P}(\hat{\beta}_{n^2} = \mathbf{0}) \rightarrow 1,$$

$$(ii) \quad \hat{\Theta} \xrightarrow{p} \Theta, \quad \sqrt{n}(\hat{\eta} - \eta) \xrightarrow{d} \mathcal{N}(\mathbf{0}, I_{\eta\eta}^{-1}),$$

$$\sqrt{n}(\hat{\beta}_{n^1} - \beta_{n^0}) \xrightarrow{d} \mathcal{N}_{s_n}(\mathbf{0}, I_{\beta_{n01}\beta_{n01}}^{-1}),$$

که در آن، $n^{-1}J_{\eta\eta} \rightarrow I_{\eta\eta}$ ، $n^{-1}J_{\beta_{n^0}\beta_{n^0}} \rightarrow I_{\beta_{n^0}\beta_{n^0}}$ ، $\eta = (\text{vech}(D)^\top, \text{vech}(\Sigma)^\top)^\top$ ، $J_{\eta\eta}$ و $J_{\beta_{n^0}\beta_{n^0}}$ ماتریس‌های اطلاعات فیشر هستند.

اتبات قضایای ۱.۳.۲ و ۲.۳.۲ و همچنین ماتریس‌های اطلاعات فیشر در قسمت ضمیمه آمده است.

۳.۳.۲ مطالعات عددی

در این بخش، فرایند برازش مدل، برآورد پارامترهای مدل و انتخاب متغیر را به صورت عددی بررسی می‌کنیم. ابتدا مجموعه‌ای از داده‌های شبیه‌سازی شده و یک مورد مطالعاتی با داده‌های واقعی را توسط روش پیشنهادی مقاله یعنی برآوردگر بیشینه درست‌نمایی توانیده در مدل نیمه‌پارامتری با اثرات آمیخته ($P - SMM$) مورد تحلیل قرار می‌دهیم. سپس به منظور بررسی عملکرد روش $P - SMM$ در مسئله برآورد و انتخاب متغیر همزمان، نتایج حاصل از تحلیل داده‌ها بدون رهیافت تابع تاوان یعنی روش برآوردگر بیشینه درست‌نمایی در مدل خطی-جزئی با اثرات آمیخته (SMM) نیز گزارش می‌شود. در نظر داشته باشید که اگر داده‌ها ذاتاً دارای جزء ناپارامتری باشند، مدل نیمه‌پارامتری بهتر از مدل خطی عمل می‌کند. برای نشان دادن این موضوع روش $P - SMM$ را با روش برآوردگر بیشینه درست‌نمایی توانیده در مدل خطی با اثرات آمیخته ($P - LMM$) مقایسه می‌کنیم.

۱.۳.۳.۲ مطالعات شبیه‌سازی

در این قسمت متغیرهای پاسخ را از مدل

$$y_{ij} = \sum_{k=1}^{p_n} x_{k,ij} \beta_k + f(t_{ij}) + b_i + \varepsilon_{ij},$$

شبیه‌سازی می‌کنیم، که در آن مولفه‌های مدل به صورت

$$\begin{aligned} i &= 1, \dots, n, \quad j = 1, \dots, n_i, \quad n = 50, 100, 200, \quad n_i = 5 \\ \beta^\top &= (2, 3, 1/5, 2, 0, \dots, 0), \quad b_i \sim \mathcal{N}(0, \sigma_b^2); \quad \sigma_b^2 = 0.25, \\ \mathbf{X}_{ij}^\top &= (x_{1,ij}, \dots, x_{p_n,ij}), \quad x_{k,ij} \sim \mathcal{U}(-1, 1) \\ f(t_{ij}) &= 2 \sin(2\pi t_{ij}), \quad t_{ij} \sim \mathcal{U}(0, 1). \end{aligned}$$

تعیین شده‌اند. به منظور انتخاب تعداد مولفه‌های پارامتری مدل، نویسندگان پیشنهاد‌های متفاوتی از جمله $p_n = \lfloor \frac{n}{b \log(n)} \rfloor$ و $p_n = \lfloor 4/5 n^{1/4} \rfloor$ ، $p_n = \lfloor \frac{n}{4} \rfloor$ ، $p_n = \lfloor \frac{n}{2} \rfloor$ داشته‌اند که در آن‌ها $b > 1$ و a صحیح مقدار است. موارد بیان شده برای حالت $p_n < n$ است. برای زمانی که $p_n > n$ ، نرخ رشد چند جمله‌ای ($p_n = O(n^b)$) و هنگامی که $p_n \gg n$ (بسیار بزرگتر از n)، نرخ رشد نمایی ($p_n = \exp\{O(n^b)\}$) پیشنهاد شده است که در آن $0 < b < 1$. در حالت p_n بسیار بزرگتر از n اگر چه احتمال انتخاب متغیرهای موثر در مدل افزایش می‌یابد در عین حال متغیرهای بی‌تاثیر بیشتری در مدل وارد خواهند شد که این موضوع بر عملکرد روش‌های برآورد و انتخاب متغیر تاثیر منفی خواهد گذاشت. تحلیل داده‌های با بعد بسیار بالا با چالش‌هایی روبرو است که کاربرد مستقیم روش‌های انتخاب متغیر اشاره شده را با مشکل مواجه می‌کند. یکی از راه حل‌های رفع این مشکل استفاده از روش‌های غربالگری است. از آنجایی که در این رساله به داده‌های با بعد بسیار بالا پرداخته نمی‌شود از توضیحات بیشتر پیرامون این مبحث صرف نظر می‌کنیم و علاقه‌مندان جهت آشنایی بیشتر می‌توانند به کاظمی (۱۳۹۷) مراجعه کنند. به منظور همپوشانی تمامی نرخ‌های پیشنهادی، موارد $(200, 16)$ ، $(100, 14)$ ، $(50, 11)$ برای حالت $p_n < n$ و موارد $(200, 2000)$ ، $(100, 500)$ ، $(30, 100)$ برای حالت $p_n > n$ در نظر گرفته می‌شود.

به منظور مقایسه عملکرد روش پیشنهادی، هر مجموعه از داده‌های شبیه‌سازی شده را توسط سه روش شامل برآوردگر بیشینه درست‌نمایی تاوانیده در مدل خطی-جزئی با اثرات آمیخته ($P - SMM$)، برآوردگر بیشینه درست‌نمایی در مدل خطی-جزئی با اثرات آمیخته (SMM) و برآوردگر بیشینه درست‌نمایی تاوانیده در مدل خطی با اثرات آمیخته ($P - LMM$)، مورد تحلیل قرار دادیم. دقت برآورد هر یک از روش‌ها توسط معیار میانگین کمترین توان‌های دوم خطا (MSE) براساس ۱۰۰ تکرار شبیه‌سازی، مورد سنجش قرار گرفت. ارزیابی عملکرد روش‌ها در مسئله انتخاب متغیر توسط معیارهای (C, I) سنجیده می‌شود به طوری که C میانگین ضرایب صفر که به درستی صفر تخمین زده شده‌اند و I میانگین ضرایب غیر صفر است که به اشتباه صفر شده‌اند. برای ارائه توصیف جامع‌تر از عملکرد مسئله انتخاب متغیر روش‌ها از معیارهای دیگری نیز استفاده کردیم. UF بیانگر نسبت تکرارهایی است که ضرایب تاثیرگذار را به اشتباه صفر در نظر گرفته است، CF بیانگر نسبت تکرارهایی است که مدل را به درستی تعیین کرده‌اند و نسبت تکرارهایی که علاوه بر متغیرهای مهم، متغیرهای بی‌تاثیر را نیز در مدل وارد کرده‌اند با OF نشان دادیم.

جدول ۴.۲، خلاصه‌ای از نتایج مربوط به عملکرد روش‌های $P - LMM$ ، $P - SMM$ و SMM در برآورد و انتخاب مدل برای مقادیر مختلف (n, p_n) است. از نظر دقت برآورد روش SMM

جدول ۴.۲: نتایج شبیه‌سازی: مقایسه نتایج انتخاب متغیر تحت روش‌های SMM ، $P - SMM$ و $P - LMM$ برای حالت‌های $p_n < n$ و $p_n > n$.

method	حالت $p_n < n$						حالت $p_n > n$					
	$(n, p) = (50, 11)$						$(n, p) = (50, 100)$					
	MSE	C(λ)	I(°)	UF	CF	OF	MSE	C(۹۷)	I(°)	UF	CF	OF
SMM	۰/۳۹۶	۰/۰۰	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰	۸/۰۵۰	۰/۱۰	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰
P - LMM	۰/۱۸۳	۶/۲۸	۰/۰۰	۰/۰۰	۰/۵۰	۰/۵۰	۰/۴۵۵	۹۴/۱۹	۰/۰۰	۰/۰۰	۰/۲۳	۰/۷۷
P - SMM	۰/۱۹۱	۶/۲۸	۰/۰۰	۰/۰۰	۰/۵۴	۰/۴۶	۰/۳۸۹	۹۴/۶۵	۰/۰۰	۰/۰۰	۰/۲۷	۰/۷۳
	$(n, p) = (100, 14)$						$(n, p) = (100, 500)$					
	MSE	C(۱۱)	I(°)	UF	CF	OF	MSE	C(۹۹۷)	I(°)	UF	CF	OF
SMM	۰/۲۶۵	۰/۰۲	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰	۱۴۹۹/۱۳۶	۴۷/۰۲	۰/۰۳	۰/۰۳	۰/۰۰	۱/۰۰
P - LMM	۰/۰۹۳	۹/۴۷	۰/۰۰	۰/۰۰	۰/۶۲	۰/۳۸	۰/۰۶۲	۴۹۵/۷۲۰	۰/۰۰	۰/۰۰	۰/۳۰	۰/۷۰
P - SMM	۰/۰۹۶	۹/۵۶	۰/۰۰	۰/۰۰	۰/۶۹	۰/۳۱	۰/۰۳۸	۴۹۶/۲۵۰	۰/۰۰	۰/۰۰	۰/۵۴	۰/۴۶
	$(n, p) = (150, 15)$						$(n, p) = (200, 2000)$					
	MSE	C(۱۲)	I(°)	UF	CF	OF	MSE	C(۱۹۹۷)	I(°)	UF	CF	OF
SMM	۰/۱۳۹	۰/۰۹	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰	۱۲۵/۴۰۶	۱۱۳۷/۶۲	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰
P - LMM	۰/۰۴۳	۱۱/۹۳	۰/۰۰	۰/۰۰	۰/۹۳	۰/۰۷	۰/۰۱۸	۱۹۹۶/۸۹	۰/۰۰	۰/۰۰	۰/۸۹	۰/۱۱
P - SMM	۰/۰۴۳	۱۱/۹۶	۰/۰۰	۰/۰۰	۰/۹۶	۰/۰۴	۰/۰۱۸	۱۹۹۶/۹۳	۰/۰۰	۰/۰۰	۰/۹۲	۰/۰۸

جدول ۵.۲: نتایج شبیه‌سازی: کارایی روش $P - SMM$ برای حالت‌های $p_n < n$ و $p_n > n$.

حالت $p_n < n$						حالت $p_n > n$					
(n, p_n)		β_1	β_2	β_3	β_4	(n, p_n)		β_1	β_2	β_3	β_4
$(50, 11)$	Bias	۰/۰۳۹	۰/۰۰۸	۰/۰۰۳	۰/۰۰۵	$(50, 100)$	Bias	۰/۰۸۰	۰/۱۰۸	۰/۰۴۶	۰/۰۵۸
	SD۱	۰/۱۰۵	۰/۱۰۷	۰/۱۰۶	۰/۱۰۳		SD۱	۰/۰۹۳	۰/۰۸۷	۰/۰۸۰	۰/۰۸۱
	SD۲	۰/۰۷۷	۰/۰۸۲	۰/۰۸۰	۰/۰۸۰		SD۲	۰/۰۸۱	۰/۰۷۹	۰/۰۷۷	۰/۰۷۷
	CP	۰/۹۵	۰/۹۷	۰/۹۵	۰/۹۱		CP	۰/۹۶	۰/۹۶	۰/۹۹	۰/۹۳
$(100, 14)$	Bias	۰/۰۱۳	۰/۰۴۹	۰/۰۱۱	۰/۰۱۷	$(100, 500)$	Bias	۰/۰۴۱	۰/۰۳۶	۰/۰۱۹	۰/۰۲۹
	SD۱	۰/۰۷۵	۰/۰۷۶	۰/۰۷۳	۰/۰۷۷		SD۱	۰/۰۷۸	۰/۰۷۹	۰/۰۷۷	۰/۰۷۸
	SD۲	۰/۰۵۵	۰/۰۵۵	۰/۰۵۵	۰/۰۵۵		SD۲	۰/۰۵۵	۰/۰۵۴	۰/۰۵۵	۰/۰۵۴
	CP	۰/۹۱	۰/۹۴	۰/۹۵	۰/۹۴		CP	۰/۹۳	۰/۹۳	۰/۹۵	۰/۹۴
$(200, 16)$	Bias	۰/۰۳۷	۰/۰۰۳	۰/۰۱۳	۰/۰۱۴	$(200, 2000)$	Bias	۰/۰۱۶	۰/۰۴۲	۰/۰۱۳	۰/۰۰۵
	SD۱	۰/۰۵۷	۰/۰۵۷	۰/۰۵۸	۰/۰۵۹		SD۱	۰/۰۷۶	۰/۰۷۷	۰/۰۷۵	۰/۰۷۹
	SD۲	۰/۰۳۹	۰/۰۳۹	۰/۰۴۰	۰/۰۴۰		SD۲	۰/۰۵۵	۰/۰۵۵	۰/۰۵۵	۰/۰۵۶
	CP	۰/۹۶	۰/۹۵	۰/۹۵	۰/۹۵		CP	۰/۹۳	۰/۹۶	۰/۹۲	۰/۹۴

MSE بیش‌تری نسبت به دو روش دیگر دارد، در حالی‌که دو روش $P - LMM$ و $P - SMM$ بسیار نزدیک به هم عمل می‌کنند. با این وجود MSE روش پیشنهادی همواره کمتر از $P - LMM$ است. از نظر انتخاب مدل مشاهده می‌کنیم که SMM انتخاب مدل انجام نمی‌دهد اما روش‌های $P - LMM$ و $P - SMM$ تمامی متغیرهای مهم را انتخاب کرده است، زیرا معیار I آن‌ها صفر است. همچنین روش پیشنهادی به دلیل داشتن مقادیر بالا در نرخ‌های C و CF و مقادیر کمتر در OF ، نسبت به رقیب خود یعنی $P - LMM$ عملکرد بهتری داشته است. با توجه به معیارهای CF و OF ، هنگامی که $p_n > n$ نسبت به حالت $p_n < n$ روش‌ها در انتخاب مدل عملکرد ضعیفی دارند و ضرایب صفر تمایل بیش‌تری دارند که در مدل گنجانده شوند، اما با افزایش حجم نمونه هر دو

معیار بهبود می‌یابند. برای بررسی بیش‌تر عملکرد روش پیشنهادی، نتایج مربوط به اریبی ($Bias$)، انحراف استاندارد مجانبی (SD_1)، انحراف استاندارد نمونه‌ای (SD_2) و احتمال پوشش نمونه‌ای^{۴۶} (CP) برای فاصله اطمینان ۹۵٪ در جدول ۵.۲ گزارش شده است. نتایج نشان می‌دهد که خطاهای استاندارد نمونه‌ای و مجانبی نزدیک به هم برآورد شده‌اند و احتمال پوشش در اغلب موارد نزدیک به ۹۵٪ است که این‌ها بیانگر عملکرد خوب روش است. به عبارت دیگر برآوردها رفتار مجانبی خوبی را نشان داده‌اند که این امر با نتایج قضایای حدی بیان‌شده مطابقت دارد.

۲.۳.۳.۲ تحلیل داده‌های CD_4

برای بهره‌مندی از انعطاف بالای مدل‌های نیمه پارامتری ما متغیر $YEAR$ را به عنوان جزء ناپارامتری و سایر متغیرها را به صورت خطی وارد مدل کردیم. از آنجایی‌که اثرات متقابل نیز ممکن است در مدل اهمیت داشته باشند، ما تمامی اثرات متقابل متغیرها را نیز در مدل وارد کرده و از روش پیشنهادی برای انتخاب متغیرهای مهم استفاده کردیم. به منظور مقایسه عملکرد روش $P - SMM$ با دو روش دیگر یعنی $P - LMM$ و SMM از معیار انحراف استاندارد مجانبی SD استفاده کردیم. همچنین برای شناسایی بهترین مدل پشتیبانی‌شده توسط داده‌ها، معیار اطلاع آکائیک (AIC) و معیار اطلاع بیزی (BIC) را محاسبه کردیم. این معیارها در قسمت ضمیمه معرفی شده‌اند. نتایج حاصل در جدول ۶.۲ ارائه شده است. این نتایج نشان می‌دهد که برآوردها در مقایسه با دو روش دیگر دارای مقادیر SD کمتری می‌باشند. همچنین معیارهای AIC و BIC نیز عملکرد بهتر روش پیشنهادی را تایید می‌کنند. تحت مدل $P - SMM$ ، متغیرهای $SMOKE$ ، $SEXP$ ، $CESD$ و $SMOKE * SEXP$ به عنوان متغیرهای مهم شناسایی شدند.

نتایج نشان می‌دهند که مدل در مقایسه با سایر مدل‌های رقیب، در برآورد ضرایب غیر صفر و انتخاب مدل عملکرد بهتری دارد. به عبارت دیگر وقتی داده‌ها دارای روند غیرخطی هستند، مدل‌های خطی-جزئی نسبت به مدل‌های خطی کارآمدتر هستند. نتایج در هر دو حالت $p_n > n$ و $p_n < n$ برقرار است.

⁴⁶Coverage probability

جدول ۶.۲: نتایج تحلیل داده‌های $CD4$: برآورد پارامترهای رگرسیونی (SD) تحت برازش سه مدل $P - LMM$ و SMM ، $P - SMM$.

متغیرها	SMM $\hat{\beta}(SD)$	$P - LMM$ $\hat{\beta}(SD)$	$P - SMM$ $\hat{\beta}(SD)$
AGE	۰/۰۸۴ (۰/۰۰۴)	-۰/۶۶۱ (۰/۰۰۵)	۰ (۰)
$SMOKE$	۰/۷۳۰ (۰/۰۲۴)	۲/۴۹۸ (۰/۰۳۰)	۰/۴۱۷ (۰/۰۱۴)
$DRUG$	۰/۸۱۱ (۰/۰۴۷)	۶/۶۰۸ (۰/۰۳۸)	۰ (۰)
$SEXP$	۰/۱۷۰ (۰/۰۰۹)	۱/۱۴۵ (۰/۰۰۸)	۰/۰۸۴ (۰/۰۰۶)
$CESD$	-۰/۰۶۹ (۰/۰۰۴)	۰ (۰)	-۰/۰۵۴ (۰/۰۰۳)
$AGE * SMOKE$	۰/۰۱۲ (۰/۰۰۱)	۰/۰۱۶ (۰/۰۰۱)	۰ (۰)
$AGE * DRUG$	-۰/۰۲۰ (۰/۰۰۴)	۰/۰۷۰ (۰/۰۰۵)	۰ (۰)
$AGE * SEXP$	-۰/۰۱۲ (۰/۰۰۱)	۰ (۰)	۰ (۰)
$AGE * CESD$	-۰/۰۰۵ (۰/۰۰۰۲)	۰ (۰)	۰ (۰)
$SMOKE * DRUG$	-۰/۲۷۸ (۰/۰۲۲)	-۱/۱۶۶ (۰/۰۳۱)	۰ (۰)
$SMOKE * SEXP$	۰/۰۳۸ (۰/۰۰۳)	۰ (۰)	۰/۰۴۲ (۰/۰۰۳)
$SMOKE * CESD$	-۰/۰۰۵ (۰/۰۰۱)	۰ (۰)	۰ (۰)
$DRUG * SEXP$	-۰/۰۷۹ (۰/۰۰۹)	-۰/۵۵۸ (۰/۰۱۰)	۰ (۰)
$DRUG * CESD$	۰/۰۲۰ (۰/۰۰۳)	۰ (۰)	۰ (۰)
$SEXP * CESD$	-۰/۰۰۱ (۰/۰۰۰۳)	۰ (۰)	۰ (۰)
$\ell_{\max}$	-۱۶۵۶۴/۶۲	-۱۷۳۰۶/۱۹	-۱۶۵۵۳/۷۲
AIC	۳۳۱۵۹/۲۴	۳۴۶۴۲/۳۷	۳۳۱۳۷/۴۵
BIC	۳۳۲۱۷/۹۰	۳۴۷۰۱/۰۳	۳۳۱۹۶/۱۱

فصل ۳

انتخاب متغیر برای داده‌های طولی در خانواده توزیع‌های نمایی

در بسیاری از تحقیقات علمی، با داده‌های رسته‌ای^۱ یا داده‌های غیرنرمال، مانند پاسخ‌های دودویی و شمارشی مواجه هستیم. نلدر و ودربرن (۱۹۷۲) اولین کسانی بودند که مدل‌های خطی تعمیم یافته^۲ (GLM) را معرفی کردند و مک‌کلاچ و نلدر (۱۹۸۹) مدل‌های GLM را برای مدل‌بندی متغیرهای پاسخ غیرنرمال پیشنهاد دادند. در این مدل‌ها با پذیرفتن استقلال مشاهدات، با استفاده از یک تابع پیوند بین میانگین مشاهدات و متغیرهای کمکی ارتباط خطی برقرار می‌شود. هنگامی که بین مشاهدات همبستگی وجود دارد، تعمیمی از مدل‌های مذکور به نام مدل‌های آمیخته خطی تعمیم‌یافته ($GLMM$) استفاده می‌شود که توسط بریسلو و کلیتون (۱۹۹۳) ارائه شد؛ که در آن پیشگوه‌های مدل علاوه بر اثرات ثابت شامل اثرات تصادفی نیز هستند. در $GLMM$ ‌ها، تابع درست‌نمایی برخلاف مدل‌های خطی به دلیل غیرنرمال بودن متغیرهای پاسخ و وجود اثرات تصادفی، صورت بسته‌ای نداشته و برآورد پارامترها به راحتی امکان‌پذیر نیست. لذا در اکثر مقالات با فرض نرمال بودن اثرات تصادفی به ارائه راه‌حلی برای برآورد پارامترهای مدل و اثرات تصادفی با بیشینه کردن توابع درست‌نمایی، شبه درست‌نمایی^۳ و درست‌نمایی سلسله مراتبی^۴ به روش‌های عددی پرداخته

^۱Categorical

^۲Generalized linear models

^۳Quasi likelihood

^۴Hierarchical likelihood

شده است. از جمله مک‌کلاچ (۱۹۹۷)، با بکار بردن روش‌های عددی مانند بیشینه‌سازی امیدریاضی مونت‌کارلویی^۵ ($MCEM$)، نیوتون رافسون مونت‌کارلویی^۶ ($MCNR$) و بیشینه درستنمایی شبیه‌سازی شده^۷ (SML) برآورد بیشینه درستنمایی پارامترها را در مدل $GLMM$ به دست آورد. همچنین در مدل‌های GLM ، پن و تامپسون (۲۰۰۷)، برآورد شبه مونت‌کارلو^۸ و لین (۲۰۰۷) برآورد پارامترهای GLM پواسنی را با ترکیب درستنمایی نما^۹ و شبه درستنمایی توانیده مورد مطالعه قرار دادند. ناتاراجان و گس (۲۰۰۰) با بسط روش‌های بیزی در تحلیل مدل‌های GLM ، نحوه انتخاب پیشین در این مدل‌ها را مطالعه کردند.

مطالعات طولی که متغیر پاسخ آن‌ها غیرنرمال باشد نیز در علوم مختلف جایگاه ویژه‌ای یافته‌اند. استفاده از $GLMM$ در تحلیل داده‌های طولی (زیگر و کریم، ۱۹۹۱) بسیار مفید واقع شده‌اند، زیرا تفاوت بین گروه‌ها را می‌توان با استفاده از اثرات تصادفی مدل‌بندی کرد. با این حال این مدل‌ها به دلیل اعمال فرضیات پارامتری اثرات ثابت، در تحلیل داده‌های طولی که متغیر پاسخ آن‌ها دارای رفتارهای پیچیده در طول زمان است، دارای محدودیت‌هایی هستند. برای از بین بردن محدودیت $GLMM$ برای مدل‌سازی روند غیرخطی زمان، مدل آمیخته نیمه‌پارامتری تعمیم‌یافته^{۱۰} ($GSMM$)، برای تحلیل داده‌های طولی استفاده می‌شود که در آن همبستگی درون آزمودنی‌ها با استفاده از اثرات تصادفی و مدل‌سازی اثر زمان توسط یک تابع هموار در نظر گرفته می‌شود. فن و همکاران (۲۰۰۷)، چن و چو (۲۰۰۷)، چن و چو (۲۰۰۹)، لیانگ (۲۰۰۹) و کروم و همکاران (۲۰۱۶) از مدل $GSMM$ در تحلیل داده‌های طولی استفاده کرده‌اند.

اگرچه انتخاب مدل برای داده‌های طولی نرمال تا اندازه‌ای گسترش یافته است، تحقیقات در مورد انتخاب مدل برای داده‌های طولی غیرنرمال در چارچوب GLM تا حد زیادی ناشناخته مانده است. در این میان پن (۲۰۰۱) معیار اطلاع شبه‌درستنمایی^{۱۱} (QIC)، کانتونی و همکاران (۲۰۰۵) معیار C_p -مالو و وانگ و کوای (۲۰۰۹) معیار BIC را برای انتخاب متغیر داده‌های طولی در مدل‌های GLM تعمیم دادند. این روش‌ها انتخاب متغیر را بر اساس معیارهای کلاسیک انجام می‌دهند. همان‌طور که پیش از این نیز بیان شد هنگامی که بعد داده‌ها بالاست این روش‌ها از نظر محاسباتی کارایی ندارند. در زمینه انتخاب متغیر داده‌های طولی براساس روش مبتنی بر تابع تاوان، فو (۲۰۰۳) رگرسیون بریج و لاسو، سوای و همکاران (۲۰۱۰) روش غربالگری مستقل و زیاک (۲۰۰۶) روش اسکد را در مدل‌های خطی تعمیم یافته پیشنهاد دادند. در تمامی تحقیقات فوق که تعمیم GLM توانیده به داده‌های طولی است، بعد مدل ثابت است. در زمینه داده‌های طولی با بعد بالا، وانگ و همکاران (۲۰۱۲) و سو و همکاران (۲۰۱۳) روش GEE توانیده با تابع تاوان

⁵Monte Carlo Expectation Maximization

⁶Monte carlo Newton Raphson

⁷Simulated Maximum Likelihood

⁸Quasi Monte Carlo

⁹Pseudo Likelihood

¹⁰Generalized semiparametric mixed models

¹¹Quasi-likelihood information criterion

اسکد را پیشنهاد دادند.

در این فصل با تمرکز روی مدل‌های GSM ، به معرفی اجمالی مدل‌های آمیخته نیمه‌پارامتری تعمیم‌یافته پرداخته می‌شود. سپس برآورد و انتخاب متغیر مبتنی بر رهیافت تابع تاوان برای داده‌های طولی خانواده توزیع‌های نمایی با بعد بالا و بعد پایین در نظر گرفته می‌شود. برآوردگرهای معرفی شده صورت بسته ندارند؛ لذا با بکارگیری روش سو و همکاران (۲۰۱۳)، برآوردگر GEE تاوانیده را معرفی کرده و با استفاده از یک الگوریتم تکراری به نام الگوریتم نیوتن رافسون مونت کارلویی تاوانیده، به طور همزمان برآورد و انتخاب متغیر انجام می‌شود. علاوه بر این، خواص جانبی برآوردگرها در چارچوب داده‌های بعد بالا که در آن تعداد پارامترها متناسب با افزایش حجم نمونه افزایش می‌یابد، مورد بررسی قرار می‌گیرد. به منظور بررسی عملکرد روش معرفی شده، مطالعه‌ای شبیه‌سازی و همچنین تحلیل مجموعه داده‌های واقعی فراهم آمده است.

۱.۳ انتخاب متغیر برای داده‌ها با بعد بالا

۱.۱.۳ مدل آمیخته نیمه‌پارامتری تعمیم یافته

یک نمونه تصادفی n تایی از داده‌های طولی با بعد بالا را در نظر بگیرید. فرض کنید متغیرهای پاسخ به شرط اثرات تصادفی از هم مستقل و هر $y_{ij}|u_i$ دارای خانواده توزیع‌های نمایی باشد، به طوری که تابع چگالی احتمال به صورت

$$p(y_{ij}|u_i, \beta_n, \phi) = \exp \left[\phi^{-1} \{y_{ij}\theta_{ij} - b(\theta_{ij})\} + c(y_{ij}, \phi) \right], \quad (1.3)$$

است، که در آن θ_{ij} پارامتر کانونی، ϕ پارامتر پراکندگی، $a(\cdot)$ و $b(\cdot)$ توابع معلوم هستند. همچنین امید و واریانس شرطی به ترتیب به صورت $\mu_{ij} = E(y_{ij}|u_i) = b'(\theta_{ij})$ و $\nu_{ij} = \text{var}(y_{ij}|u_i) = \phi b''(\theta_{ij})$ هستند، به طوری که $b'(\theta) = \frac{\partial b(\theta)}{\partial \theta}$ و $b''(\theta) = \frac{\partial^2 b(\theta)}{\partial \theta^2}$. اکنون با فرض استقلال شرطی مشاهدات، با استفاده از تابع پیوند معلوم g ، داریم

$$g(\mu_{ij}) \triangleq \eta_{ij} = \mathbf{X}_{ij}^\top \beta_n + \mathbf{Z}_{ij}^\top \mathbf{u}_i + f(t_{ij}), \quad i = 1, \dots, n; \quad j = 1, \dots, n_i. \quad (2.3)$$

در مدل فوق y_{ij} پاسخ مربوط به آزمودنی i ام در زمان t_{ij} است، که در آن $\beta = (\beta_1, \dots, \beta_{p_n})^\top$ برداری $1 \times p_n$ بعدی از ضرایب ثابت رگرسیونی وابسته به متغیرهای تبیینی x_{ij} ؛ $f(\cdot)$ تابعی نامعلوم، هموار و دارای مشتق دوم؛ $\mathbf{u}_i = (u_{i1}, \dots, u_{iq})^\top$ برداری $1 \times q$ بعدی از ضرایب تصادفی، مستقل از هم و وابسته به متغیرهای تبیینی z_{ij} و دارای توزیعی به صورت

$$\mathbf{u}_i \sim f_u(\mathbf{u}_i|\Sigma), \quad (3.3)$$

هستند. مجموع روابط (۱.۳)–(۳.۳) به عنوان مدل آمیخته نیمه‌پارامتری تعمیم‌یافته (GSM) شناخته می‌شود.

حال فرض کنید مولفه ناپارامتری مدل به روش اسپلین‌های رگرسیونی شرح داده‌شده در ۳.۵.۱ تقریب زده شود. بدین ترتیب رابطه (۲.۳) یک مدل رگرسیونی به صورت

$$g(\mu_{ij}) \triangleq \eta_{ij} = \mathbf{X}_{ij}^\top \boldsymbol{\beta}_n + \mathbf{Z}_{ij}^\top \mathbf{u}_i + \mathbf{B}(t_{ij}) \boldsymbol{\alpha}_n, \quad i = 1, \dots, n; j = 1, \dots, n_i. \quad (4.3)$$

را نتیجه می‌دهد. برای سادگی در نوشتار، مدل را می‌توان به صورت

$$g(\mu_{ij}) \triangleq \eta_{ij} = \mathbf{D}_{ij}^\top \boldsymbol{\theta}_n + \mathbf{Z}_{ij}^\top \mathbf{u}_i, \quad i = 1, \dots, n; j = 1, \dots, n_i,$$

بازنویسی کرد، به طوری که $\mathbf{D}_{ij} = (\mathbf{X}_{ij}^\top, \mathbf{B}_j(t_i)^\top)^\top$ ماتریس طرح ادغام‌شده $(p_n + h_n) \times n_i$ بعدی از ماتریس اثرات ثابت و اثرات اسپلین و $\boldsymbol{\theta}_n = (\boldsymbol{\beta}_n^\top, \boldsymbol{\alpha}_n^\top)^\top$ بردار پارامتر ادغام‌شده $(p_n + h_n) \times 1$ بعدی هستند. همچنین $1, h_n = d + K_n + 1, d$ درجه تقریب اسپلین و K_n تعداد گره‌ها در تقریب اسپلین است.

لذا به طور خلاصه یک $GSMM$ را به صورت

$$\begin{aligned} p(y_{ij} | \mathbf{u}_i, \boldsymbol{\theta}_n, \phi) &= \exp \left[\phi^{-1} \{ y_{ij} \theta_{ij} - b(\theta_{ij}) \} + c(y_{ij}, \phi) \right], \\ \mathbf{u}_i &\sim f_u(\mathbf{u}_i | \boldsymbol{\Sigma}), \\ \mu_{ij} &= E(y_{ij} | \mathbf{u}_i), \\ g(\mu_{ij}) \triangleq \eta_{ij} &= \mathbf{D}_{ij}^\top \boldsymbol{\theta}_n + \mathbf{Z}_{ij}^\top \mathbf{u}_i, \quad i = 1, \dots, n; j = 1, \dots, n_i, \end{aligned} \quad (5.3)$$

خواهیم داشت. اکنون با توجه به مدل (۵.۳) که یک $GLMM$ به شمار می‌آید، می‌توان تمام تکنیک‌های برآورد در $GLMM$ را برای این مدل نیز بکار برد.

۲.۱.۳ انتخاب متغیر در $GSMM$

۱.۲.۱.۳ الگوریتم $MCNR$

یک روش معمول برای برآورد پارامترها در $GLMM$ ها، روش بیشینه درستنمایی است. برای مدل (۵.۳)، تابع درستنمایی به صورت

$$L(\boldsymbol{\theta}_n, \boldsymbol{\Sigma}, \phi) = \prod_{i=1}^n \int_{\mathbb{R}^q} p_{\mathbf{y}_i | \mathbf{u}_i}(\mathbf{y}_i | \mathbf{u}_i, \boldsymbol{\theta}_n, \phi) p_{\mathbf{u}}(\mathbf{u}_i | \boldsymbol{\Sigma}) d\mathbf{u}_i \quad (6.3)$$

است که در آن $\mathbf{u} = (\mathbf{u}_1, \dots, \mathbf{u}_n)$ ، $\mathbf{y}_i = (y_{i1}, \dots, y_{in_i})^\top$ و

$$p_{\mathbf{y}_i | \mathbf{u}_i}(\mathbf{y}_i | \mathbf{u}_i, \boldsymbol{\theta}_n, \phi) = \prod_{j=1}^{n_i} p(y_{ij} | \mathbf{u}_i, \boldsymbol{\theta}_n, \phi).$$

تابع درستنمایی (۵.۳) دارای صورت بسته نیست و بعد انتگرال‌گیری برابر با تعداد اثرات تصادفی مدل است. لذا برای به دست آوردن برآورد ML از الگوریتم بیشینه‌سازی امیدریاضی مونت‌کارلویی (مک‌کلاچ، ۱۹۹۷) استفاده می‌کنیم. بدین ترتیب که اثرات تصادفی $\mathbf{u}_i$ را به عنوان مقادیر گم‌شده

و (y_i, u_i) را داده‌های کامل در نظر می‌گیرد. سپس لگاریتم تابع درستنمایی برای داده‌های کامل به صورت

$$\ell(\theta_n, \Sigma, \phi) = \sum_{i=1}^n \ln p_{y_i|u_i}(y_i|u_i, \theta_n, \phi) + \sum_{i=1}^n \ln p_{u_i}(u_i|\Sigma). \quad (7.3)$$

ارائه می‌شود. در نظر گرفتن u_i به عنوان مقادیر گمشده این مزیت را دارد که در گام-M، برای بیشینه‌سازی نسبت به پارامترهای θ_n و ϕ کافی است که تنها امید عبارت اول در رابطه (7.3) یعنی $E[\ln p_{y_{ij}|u_i}(y_{ij}|u_i, \theta_n)|y_{ij}]$ را بیشینه کنیم و برای بیشینه‌سازی نسبت به پارامترهای Σ تنها از امید عبارت دوم رابطه (7.3) یعنی $E[\ln p_{u_i}(u_i|\Sigma)|y_{ij}]$ استفاده می‌شود. بنابراین معادلات درستنمایی برای θ_n و Σ ، به ترتیب به صورت

$$\begin{aligned} E\left[\frac{\partial \ln p_{y_{ij}|u_i}(y_{ij}|u_i, \theta_n)}{\partial \theta_n} | y_{ij}\right] &= 0, \\ E\left[\frac{\partial \ln p_{u_i}(u_i|\Sigma)}{\partial \Sigma} | y_{ij}\right] &= 0, \end{aligned}$$

حاصل می‌شود. ملاحظه می‌شود که عبارت اول در رابطه فوق مشابه بیشینه‌سازی در مدل خطی تعمیم‌یافته است. بنابراین در این مرحله منطقی است که برای متناسب‌سازی با مدل خطی تعمیم‌یافته، رویکردی نظیر روش نیوتن رافسون را به کار ببریم. در این شرایط مک‌کلاچ (۱۹۹۷) الگوریتم نیوتن رافسون مونت کارلویی^{۱۲} ($MCNR$) را پیشنهاد داد. با بسط دادن این الگوریتم به داده‌های طولی، معادلات برآورد به صورت

$$E_{u|y} \left[n^{-1} \sum_{i=1}^n \frac{\partial \mu_i(\theta_n, u_i)}{\partial \theta_n^\top} V_i^{-1}(\theta_n, u_i) (y_i - \mu_i(\theta_n, u_i)) \right] = 0, \quad (8.3)$$

ارائه می‌گردد، که در آن $\mu_i(\theta_n, u_i) = (\mu_{i1}, \dots, \mu_{in_i})^\top$ و $V_i(\theta_n, u_i)$ ماتریس کوواریانس $y_i|u_i$ است. الگوریتم $MCNR$ معرفی شده را می‌توان تعمیمی از روش GEE از بخش ۱.۱.۳.۱ در نظر گرفت. لذا مشابه روش GEE ، ماتریس کوواریانس $V_i(\theta_n, u_i)$ را براساس ماتریس همبستگی برداری از پارامترهای همبستگی است، مشخص می‌کنیم. در فصل ۱ برای $R(\rho)$ ساختارهای متعددی معرفی شد که به ازای یک ساختار خاص می‌توان پارامترهای همبستگی را به روش‌های متفاوت برآورد کرد. با در نظر گرفتن برآورد ماتریس کوواریانس به صورت $\hat{R} \equiv R(\hat{\rho})$ ، رابطه (8.3) به صورت

$$E_{u|y} \left[n^{-1} \sum_{i=1}^n D_i^\top A_i^{-1}(\theta_n, u_i) \hat{R}^{-1} A_i^{-1}(\theta_n, u_i) (y_i - \mu_i(\theta_n, u_i)) \right] = 0, \quad (9.3)$$

درمی‌آید، که در آن $D_i = (D_{i1}^\top, \dots, D_{in_i}^\top)^\top$. برآوردگر $\hat{\theta}_n$ با حل معادلات (9.3) به دست می‌آید. در ادامه برای سادگی در نوشتار و تفسیر فرض می‌کنیم $\phi = 1$ و $n_i = m < \infty$. حال از طریق بهینه‌سازی تابع هدف مبتنی بر تابع تاوان به برآورد و انتخاب متغیر می‌پردازیم.

¹²Monte carlo newton raphson

۲.۲.۱.۳ الگوریتم MCNR تاوانیده

به منظور انتخاب متغیرهای مهم و برآورد ضرایب آن‌ها، به لگاریتم تابع درستنمایی (۷.۳)، جمله تاوان $\sum_{k=1}^p P_{\lambda_n}(|\beta_k|)$ را اضافه می‌کنیم، آنگاه لگاریتم تابع درستنمایی تاوانیده به صورت

$$\ell^p(\beta_n, \alpha_n, D, \phi) = \sum_{i=1}^n \ln p_{y_i|u_i}(y_i|u_i, \theta_n) + \sum_{i=1}^n p_{u_i}(u_i|\Sigma) - n \sum_{k=1}^{p_n} p_{\lambda_n}(|\beta_{nk}|), \quad (۱۰.۳)$$

به دست می‌آید. از آن جایی که ضرایب θ_n با جملات اول و سوم عبارت فوق وابسته است، معادلات برآورد تاوانیده به صورت

$$U_n(\theta_n) = S_n(\theta_n) - q_{\lambda_n}(|\beta_n|) \text{sgn}(\beta_n),$$

پیشنهاد می‌شود، به طوری که

$$S_n(\theta_n) = E_{u|y} \left[\sum_{i=1}^n D_i^\top A_i^{-1}(\theta_n, u_i) \hat{R}^{-1} A_i^{-1}(\theta_n, u_i) (y_i - \mu_i(\theta_n, u_i)) \right],$$

$\text{sgn}(\beta_n) =$ برداری $p_n \times 1$ بعدی از توابع تاوان، $q_{\lambda_n}(|\beta_n|) = (q_{\lambda_n}(|\beta_{n1}|), \dots, q_{\lambda_n}(|\beta_{np_n}|))^\top$ و $(\text{sgn}(\beta_{n1}), \dots, \text{sgn}(\beta_{np_n}))^\top$ هستند.

در نظر داشته باشید که برای سادگی، فرض می‌کنیم قسمت ناپارامتری سهم قابل توجهی در مدل داشته باشد و معادلات برآورد تاوانیده تنها مولفه‌های کوچک β_n را منقبض می‌کند و نه α_n ها را. در اینجا تابع تاوان اسکد با تقریب درجه دو موضعی (LQA) را به کار می‌بریم. برآوردگر θ_n باید جواب معادله $U_n(\theta_n) = 0$ باشد، اما از آن جایی که $U_n(\theta_n)$ دارای نقاط ناپیوسته است، ممکن است جواب دقیق برای $U_n(\theta_n) = 0$ وجود نداشته باشد. لذا برآوردگر $\hat{\theta}_n$ را به عنوان جواب تقریبی معادلات برآورد می‌شناسیم، به عبارت دیگر به ازای دنباله $a_n \rightarrow 0$ داشته باشیم $U_n(\hat{\theta}_n) = o(a_n)$. اکنون برای حل مسئله بهینه سازی تاوانیده $U_n(\hat{\theta}_n) = o(a_n)$ از الگوریتم تکراری نیوتن رافسون استفاده شده و رابطه بازگشتی به صورت

$$\hat{\theta}_n^{(m+1)} = \hat{\theta}_n^{(m)} + \left\{ H_n(\hat{\theta}_n^{(m)}) + n E_n(\hat{\theta}_n^{(m)}) \right\}^{-1} \times \left\{ S_n(\hat{\theta}_n^{(m)}) + n E_n(\hat{\theta}_n^{(m)}) \hat{\theta}_n^{(m)} \right\}, \quad (۱۱.۳)$$

حاصل می‌شود، به طوری که

$$H_n(\hat{\theta}_n^{(m)}) = E_{u|y} \left[\sum_{i=1}^n D_i^\top A_i^{-1}(\theta_n, u_i) \hat{R}^{-1} A_i^{-1}(\theta_n, u_i) D_i \right],$$

$$E_n(\hat{\theta}_n^{(m)}) = \text{diag} \left\{ \frac{q_{\lambda_n}(|\beta_{n1}|)}{\epsilon + |\beta_{n1}|}, \dots, \frac{q_{\lambda_n}(|\beta_{np_n}|)}{\epsilon + |\beta_{np_n}|}, 0_{h_n} \right\}, \quad \epsilon = 10^{-6}$$

و 0_{h_n} بردار h_n بعدی با عناصر ۰ است.

عبارت‌های امید ریاضی در رابطه (۱۱.۳) صورت بسته ندارند، زیرا توزیع شرطی $u_i|y_i$ به توزیع حاشیه‌ای y_i وابسته است که محاسبه صریح آن کار ساده‌ای نیست. مشابه [۷۸] از یک روش جایگزین

استفاده می‌کنیم بدین ترتیب که با استفاده از الگوریتم متروپولیس تانر (۱۹۹۳) نمونه‌هایی تصادفی از توزیع شرطی $u_i|y_i$ تولید می‌کنیم که دیگر به تعیین توزیع حاشیه‌ای y_i نیازی نیست. الگوریتم متروپولیس، یک روش زنجیر مارکوف مونت کارلویی^{۱۳} برای به‌دست آوردن نمونه‌های تصادفی از یک توزیع احتمالی است که نمونه‌برداری مستقیم از آن دشوار می‌باشد. این نمونه‌ها را می‌توان برای برآورد یک توزیع یا برای محاسبه برخی انتگرال‌ها (به‌طور مثال یک امید ریاضی) استفاده کرد. در این الگوریتم یک توزیع پیشنهادی متناسب با مسئله مورد نظر انتخاب می‌شود به طوری که مشاهدات جدید از آن ساخته می‌شوند. توزیع پیشنهادی باید به گونه‌ای ساخته شود که نمونه‌گیری از آن ساده باشد. سپس با تعیین یک تابع پذیرش، احتمال پذیرش مشاهدات جدید در مقابل مشاهدات گذشته محاسبه می‌شود. حال فرض کنید U مشاهدات قبلی از توزیع شرطی $U|y$ باشد و مقادیر جدید u_k^* را برای مولفه k ام $U^* = (u_1, \dots, u_{k-1}, u_k^*, u_{k+1}, \dots, u_{nq})$ تولید می‌کنیم. با در نظر گرفتن p_u به عنوان توزیع پیشنهادی، U^* را با احتمال

$$\alpha_k(U, U^*) = \min\left\{1, \frac{p_{u|y}(U^*|y, \theta_n, D)p_u(U|D)}{p_{u|y}(U|y, \theta_n, D)p_u(U^*|D)}\right\}. \quad (12.3)$$

می‌پذیریم. در صورت رد مقادیر جدید U^* ، مقادیر قبلی U باقی می‌مانند. عبارت دوم در رابطه (۱۲.۳)، را می‌توان به صورت

$$\begin{aligned} \frac{p_{u|y}(U^*|y, \theta_n, D)p_u(U|D)}{p_{u|y}(U|y, \theta_n, D)p_u(U^*|D)} &= \frac{p_{y|u}(y|U^*, \theta_n)}{p_{y|u}(y|U, \theta_n)} \\ &= \frac{\prod_{i=1}^n p_{y_i|u}(y_i|U^*, \theta_n)}{\prod_{i=1}^n p_{y_i|u}(y_i|U, \theta_n)}, \end{aligned}$$

نوشت. محاسبه تابع پذیرش $\alpha_k(U, U^*)$ تنها شامل توزیع شرطی $y|u$ است که صورت بسته آن قابل محاسبه است. در پایان الگوریتم MCNR به صورت الگوریتم ۳ پیشنهاد می‌شود. این روش برآوردگرهای سازگاری برای ضرایب مدل ارائه می‌دهد که مستقل از انتخاب ساختار همبستگی بوده و $\sqrt{n}(\hat{\theta} - \theta)$ دارای توزیع نرمال مجانبی چندمتغیره با میانگین 0 و ماتریس واریانس-کوواریانس به صورت

$$\text{Cov}(\hat{\theta}_n) \approx [H_n(\hat{\theta}_n, u_i) + nE_n(\hat{\theta}_n)]^{-1} M_n(\hat{\theta}_n, u_i) [H_n(\hat{\theta}_n, u_i) + nE_n(\hat{\theta}_n, u_i)]^{-1},$$

است، به طوری که

$$M_n(\hat{\theta}_n, u_i) = \sum_{i=1}^n D_i^\top A_i^{-1/2}(\hat{\theta}_n, u_i) \hat{R}^{-1} [\epsilon_i(\hat{\theta}_n, u_i) \epsilon_i^\top(\hat{\theta}_n, u_i)] \hat{R}^{-1} A_i^{-1/2}(\hat{\theta}_n, u_i) D_i^\top,$$

و

$$\epsilon_i(\theta_n, u_i) = (\epsilon_{i1}(\theta_n, u_i), \dots, \epsilon_{in_i}(\theta_n, u_i))^\top = A_i^{-1/2}(\theta_n, u_i)(Y_i - \mu_i(\theta_n, u_i)).$$

¹³Markov chain Monte Carlo

الگوریتم ۳ الگوریتم $MCNR$ تاوانیده

گام ۱. فرض کنید $m_k = 0$ برای θ_n و Σ مقادیر اولیه به صورت θ_n^0 و Σ^0 انتخاب کنید.
 گام ۲. با استفاده از الگوریتم متروپولیس، N مشاهده به صورت $U^{(1)}, \dots, U^{(N)}$ از توزیع $p_{u|y}(u|y, \theta_n^{(m_k)}, \Sigma^{(m_k)})$ تولید کنید. از این مشاهدات برای برآورد مونت کارلویی امید ریاضی‌ها استفاده کنید. یعنی

(a) $\theta_n^{(m_k+1)}$ را به صورت

$$\begin{aligned} \theta_n^{(m_k+1)} = & \theta_n^{(m_k)} + \left\{ \frac{1}{N} \sum_{k=1}^N \left[H_n(\hat{\theta}_n^{(m_k)}, U^{(k)}) \right] + n E_n(\hat{\beta}_n^{(m_k)}) \right\}^{-1} \\ & \times \left\{ \frac{1}{N} \sum_{k=1}^N \left[S_n(\hat{\theta}_n^{(m_k)}, U^{(k)}) \right] - n E_n(\hat{\beta}_n^{(m_k)}) \hat{\beta}_n^{(m_k)} \right\}, \end{aligned}$$

محاسبه کنید، به طوری که

$$\begin{aligned} H_n(\hat{\theta}_n^{(m_k)}, U^{(k)}) &= \sum_{i=1}^n D_i^\top A_i^{-\frac{1}{2}}(\theta_n^{(m_k)}, U_i^{(k)}) \hat{R}^{-1} A_i^{-\frac{1}{2}}(\theta_n^{(m_k)}, U_i^{(k)}) D_i, \\ S_n(\hat{\theta}_n^{(m_k)}, U^{(k)}) &= \sum_{i=1}^n D_i^\top A_i^{-\frac{1}{2}}(\theta_n^{(m_k)}, U_i^{(k)}) \hat{R}^{-1} A_i^{-\frac{1}{2}}(\theta_n^{(m_k)}, U_i^{(k)}) (y_i - \mu_i(\theta_n^{(m_k)}, U_i^{(k)})). \end{aligned}$$

(b) $\Sigma^{(m_k+1)}$ را از طریق بیشینه‌سازی عبارت $\frac{1}{N} \sum_{k=1}^N \ln f_u(U^{(k)}|\Sigma)$ به دست آورید.

(c) قرار دهید $m_k = m_{k+1}$.

گام ۳. گام ۲ را تا رسیدن به همگرایی تکرار کنید. $\theta_n^{(m_k+1)}$ و $\Sigma^{(m_k+1)}$ برآوردگر $MCNR$ تاوانیده برای θ_n و Σ هستند.

۳.۲.۱.۳ انتخاب پارامترهای تاوان

برای پیاده‌سازی روش پیشنهادشده، باید درباره چندین پارامتر تصمیم‌گیری مناسبی اتخاذ شود. این پارامترها شامل درجه چند جمله‌ای d ، تعداد نقاط گره L_n و تعداد توابع پایه h_n در تقریب B -اسپلاین و پارامترهای a و λ_n در تابع تاوان اسکد هستند. تصمیم‌گیری درباره d ، K_n ، h_n و a مشابه روند بیان‌شده در بخش ۳.۱.۳.۲ است. همچنین برای پارامتر تاوان λ_n از معیار GCV معرفی‌شده در همان بخش که به‌صورت

$$GCV_{\lambda_n} = \frac{RSS(\lambda_n)/n}{(1 - d(\lambda_n)/n)^2},$$

است، استفاده می‌کنیم با این تفات که مولفه‌های آن به‌صورت

$$RSS(\lambda_n) = \frac{1}{N} \sum_{k=1}^N \left[\sum_{i=1}^n (y_i - \mu_i(\hat{\theta}_n, U_i^{(k)}))^T W_i^{-1} (y_i - \mu_i(\hat{\theta}_n, U_i^{(k)})) \right] \quad (13.3)$$

و

$$d(\lambda_n) = \text{tr} \left[\left\{ \frac{1}{N} \sum_{k=1}^N [H_n(\hat{\theta}_n, U^{(k)})] + nE_n(\hat{\theta}_n) \right\}^{-1} \times \left\{ \frac{1}{N} \sum_{k=1}^N [H_n(\hat{\theta}_n, U^{(k)})] \right\} \right]$$

تعیین می‌شوند. W_i در (۱۳.۳) ماتریس کوواریانس $n_i \times n_i$ برای پاسخ‌های y_i است و به‌صورت $W_i = E_{u|y}(\text{Var}(y_i|u_i)) + \text{Var}_{u|y}(E(y_i|u_i))$ محاسبه می‌شود، به‌طوری‌که

$$E_{u|y}(\text{var}(y_i|u_i)) = \frac{1}{N} \sum_{k=1}^N [V_i(\hat{\theta}_n, U_i^{(k)})],$$

$$\text{Var}_{u|y}(E(y_i|u_i)) = \frac{1}{N} \sum_{k=1}^N [\mu_i(\hat{\theta}_n, U_i^{(k)})]^2 - \left[\frac{1}{N} \sum_{k=1}^N [\mu_i(\hat{\theta}_n, U_i^{(k)})] \right]^2.$$

۳.۱.۳ خواص جانبی

در این بخش رفتار جانبی برآوردهای $\hat{\beta}_n$ و $\hat{f}$ تحت الگوریتم $MCNR$ پیشنهادشده، مطالعه می‌شوند. توجه داشته باشید که در این بخش نیز، تعاریف، شرایط و تفکیک‌های معرفی‌شده در بخش ۲.۳.۲ برقرار است. همچنین شرایط نظم زیر را در نظر بگیرید:

(۱.A) به ازای هر $k \geq 2$ ، مشتق k ام $f_\circ(t_i)$ کراندار است.

(۲.A) $1 \leq j \leq m$ ، $1 \leq i \leq n$ ، X_{ij} به‌طور یکنواخت کراندار است.

(۳.A) این شرط مشابه شرط (۳.A) در بخش ۲.۳.۲ است.

(۴.A) ماتریس همبستگی واقعی $R_\circ$ در دارای مقادیر ویژه کراندار است؛ ماتریس همبستگی برآوردشده $\bar{R}$ در شرط $\|\hat{R}^{-1} - \bar{R}^{-1}\| = O_p(n^{-1/2})$ صدق می‌کند به‌طوری‌که $\bar{R}$ یک ماتریس معین مثبت و دارای مقادیر ویژه کراندار است و لازم نیست که $\bar{R}$ برابر با $R_\circ$ باشد.

(۵.۸) فرض کنید

$$\epsilon_i(\theta_n, u_i) = (\epsilon_{i1}(\theta_n, u_i), \dots, \epsilon_{in_i}(\theta_n, u_i))^T = A_i^{-1/2}(\theta_n, u_i)(Y_i - \mu_i(\theta_n, u_i)).$$

مقدار ثابت و متناهی $M_1 > 0$ وجود دارد که به ازای تمام i ها و بعضی از $\delta > 0$ ها داریم

$$E(\|\epsilon_i(\theta_{n^0}, u_i)\|^{2+\delta}) \leq M_1;$$

مقادیر ثابت و مثبت M_2 و M_3 وجود دارند که

$$E[\exp(M_2|\epsilon_{ij}(\theta_{n^0}, u_i)|)|X_i] \leq M_3.$$

(۶.۸) فرض کنید $B_n = \{\theta_n : \|\theta_n - \theta_{n^0}\| \leq \Delta\sqrt{p_n/n}\}$ آنگاه $\mu^*(D_{ij}^T \theta_n)$ و $\mu^{**}(D_{ij}^T \theta_n)$ به طور یکنواخت کراندار هستند.

(۷.۸) وقتی $n \rightarrow \infty$ ، آنگاه $s_n(\log n)^2 = s_n^3 n^{-1} = o(1)$ ، $\lambda_n \rightarrow 0$ ، $\min_{1 \leq k \leq s_n} |\theta_{n^0 k}|/\lambda_n \rightarrow \infty$ ، $p_n s_n^3 (\log n)^4 = o(n^2 \lambda_n^4)$ و $p_n s_n^4 (\log n)^6 = o(n^2 \lambda_n^2)$ ، $\log(p_n) = o(n \lambda_n^2 / (\log n))$ ، $o(n \lambda_n^2)$

(۸.۸) تابع تاوان در شرایط $b_n \rightarrow \infty$ ، $a_n = O(n^{-1/2})$ ، $\liminf_{n \rightarrow \infty} \liminf_{\theta \rightarrow 0^+} p'_{\lambda_n}(\theta)/\lambda_n > 0$ صدق می‌کند و همچنین مقادیر ثابت C و D وجود دارند به طوری که به ازای هر $\theta_1, \theta_2 > \lambda_n C$ خواهیم داشت $|p''_{\lambda_n}(\theta_1) - p''_{\lambda_n}(\theta_2)| \leq D|\theta_1 - \theta_2|$

حال معادلات برآوردیاب زیر را در نظر بگیرید،

$$\bar{S}_n(\theta) = E_{u|y} \left[\sum_{i=1}^n D_i^T A_i^{-1/2}(\theta, u_i) \bar{R}^{-1} A_i^{-1/2}(\theta, u_i) (y_i - \mu_i(\theta, u_i)) \right].$$

فرض کنید $\bar{M}_n(\theta_n)$ ماتریس کوواریانس $\bar{S}_n(\theta)$ است، پس

$$\bar{M}_n(\theta) = E_{u|y} \left[\sum_{i=1}^n D_i^T A_i^{-1/2}(\theta, u_i) \bar{R}^{-1} R_0 \bar{R}^{-1} A_i^{-1/2}(\theta, u_i) D_i \right].$$

برای مطالعه رفتار حدی برآوردها به بیان نتیجه‌های دیگری نیاز داریم که در قالب لم در زیر بخش ب.۱.۳ از پیوست ارائه شده است. با استفاده از قضیه ۱۲.۷ در اسچوماکر (۱۹۸۱)، $f_0(t)$ توسط $B(t)\alpha_0$ تقریب زده می‌شود، پس داریم

$$\eta_{ij}(\theta_0) = g(\mu_{ij}(\theta_0)) = X_{ij}^T \beta_0 + B(t_{ij})\alpha_0 + Z_{ij}^T u_i, \quad \theta_0 = (\beta_0^T, \alpha_0^T)_{(p_n+h_n) \times 1}^T.$$

در ادامه قضایای ۱.۱.۳-۳.۱.۳ بیان می‌شوند که به ترتیب موجود بودن، سازگاری و نرمال بودن برآوردهای تاوانیده را هنگامی که $p_n \rightarrow \infty$ نشان می‌دهند.

قضیه ۱.۱.۳. تحت شرایط بیان شده، برآوردها $MCNR$ تاوانیده $\hat{\theta} = (\hat{\theta}_1^T, \hat{\theta}_2^T)^T$ وجود دارد به طوری که

$$(i) \quad \mathbb{P}_n(|U_{nk}(\hat{\theta}_n)| = 0, k = 1, \dots, s_n^*, (s_n^* + 1), \dots, (s_n = s_n^* + h_n)) \rightarrow 1,$$

$$(ii) \quad \mathbb{P}_n \left(|U_{nk}(\hat{\theta}_n)| \leq \frac{\lambda_n}{\log n}, k = (s_n^* + h_n + 1), \dots, p_n \right) \rightarrow 1,$$

که در آن

$$U_{nk}(\hat{\theta}_n) = \begin{cases} S_{nk}(\hat{\theta}_n) - n \frac{q_{\lambda_n}(|\hat{\beta}_{nk}|)}{\epsilon + |\hat{\beta}_{nk}|} \hat{\beta}_{nk} & k = 1, \dots, s_n, \\ S_{nk}(\hat{\theta}_n) & k = (s_n + 1), \dots, p_n \end{cases},$$

و $S_{nk}(\hat{\theta}_n)$ مولفه k ام $S_n(\hat{\theta}_n)$ است.

قضیه ۲.۱.۳. تحت شرایط بیان شده، وقتی $n^{-1} p_n^2 = o(1)$ آنگاه $U_n(\theta_n) = o(1)$ دارای جواب $\hat{\theta}_n$ است به طوری که

$$(i) \quad \|\hat{\theta}_n - \theta_{n^0}\| = O_p(\sqrt{p_n/n}),$$

$$(ii) \quad \frac{1}{n} \sum_{i=1}^n \sum_{j=1}^{n_i} (\hat{f}(t_{ij}) - f_{\circ}(t_{ij}))^2 = O_p(n^{-2r/(2r+1)}).$$

قضیه ۳.۱.۳. تحت شرایط بیان شده، وقتی $L_n = O_p(n^{1/(2r+1)})$ و $n^{-1} p_n^3 = o(1)$ آنگاه به ازای هر $\xi_n \in \mathcal{R}^{p_n}$ که $\|\xi_n\| = 1$ داریم

$$(i) \quad \mathbb{P}_n(\hat{\beta}_{n^2} = \mathbf{0}) \rightarrow 1,$$

$$(ii) \quad \xi_n^\top \overline{M}_n^{*-1/2}(\beta_{n^0}) \overline{H}_n^*(\beta_{n^0})(\hat{\beta}_{n^1} - \beta_{n^0}) \xrightarrow{d} \mathcal{N}_{p_n}(\mathbf{0}, 1),$$

به طوری که

$$\begin{aligned} \overline{M}_n^* &= E_{u|y} \left[\sum_{i=1}^n X_i^{*\top} A_i^{\frac{1}{2}}(\theta_n, u_i) \overline{R}^{-1} R_{\circ} \overline{R}^{-1} A_i^{\frac{1}{2}}(\theta_n, u_i) X_i^* \right], \\ \overline{H}_n^* &= E_{u|y} \left[\sum_{i=1}^n X_i^{*\top} A_i^{\frac{1}{2}}(\theta_n, u_i) \overline{R}^{-1} A_i^{\frac{1}{2}}(\theta_n, u_i) X_i^* \right], \\ X_i^* &= (I - P) X_i \\ P &= B(B^\top \Omega B)^{-1} B^\top \Omega \\ \Omega &= \text{diag}\{\Omega_i\} \\ \Omega_i &= E_{u|y} \left[A_i^{\frac{1}{2}}(\theta_n, u_i) \overline{R}^{-1} A_i^{\frac{1}{2}}(\theta_n, u_i) \right] \end{aligned}$$

قضیه ۳.۱.۳ نشان می‌دهد که تحت شرایط نظم بیان شده، برآوردگر پیشنهادی می‌تواند ضریب متغیرهای بی‌تاثیر در مدل را با احتمال نزدیک به ۱ به سمت صفر منقبض کند. همچنین ضرایب متغیرهای مهم در مدل، دارای توزیع نرمال مجانبی هستند. اثبات قضایای ۳.۱.۳-۱.۱.۳ در قسمت پیوست ارائه شده است.

۴.۱.۳ مطالعات عددی

در این بخش، فرآیند برازش مدل شامل برآورد پارامترهای مدل و انتخاب متغیر را بر اساس مطالعات شبیه‌سازی و تحلیل داده‌های $CD4$ بخش ۴.۱ بررسی می‌کنیم. ابتدا مجموعه‌ای از داده‌های شبیه‌سازی شده و یک مورد مطالعاتی با داده‌های واقعی را توسط روش پیشنهادی مقاله یعنی برآوردگر $MCNR$ تاوانیده در مدل نیمه‌پارامتری با اثرات آمیخته تعمیم‌یافته $(P - GSMM)$ مورد تحلیل قرار می‌دهیم. سپس به منظور بررسی عملکرد روش $P - GSMM$ در مسئله برآورد و انتخاب متغیر همزمان، نتایج حاصل از تحلیل داده‌ها بدون رهیافت تابع تاوان یعنی روش برآوردگر $MCNR$ در مدل آمیخته خطی تعمیم‌یافته گزارش می‌شود. در نظر داشته باشید که اگر داده‌ها ذاتاً دارای جزء ناپارامتری باشند، مدل نیمه‌پارامتری بهتر از مدل خطی عمل می‌کند. برای نشان دادن این موضوع روش $P - GSMM$ را با روش تاوانیده در مدل آمیخته خطی تعمیم‌یافته $(P - GLMM)$ مقایسه می‌کنیم.

۱.۴.۱.۳ مطالعه شبیه‌سازی

در این قسمت متغیرهای پاسخ را از مدل پواسن با عرض از مبدا تصادفی و به صورت

$$y_{ij}|b_i \sim \text{Pois}(\mu_{ij}), \quad i = 1, \dots, n, \quad j = 1, \dots, n_i,$$

$$\eta_{ij} = \log(\mu_{ij}) = \sum_{k=1}^{p_n} x_{k,ij} \beta_k + f(t_{ij}) + b_i$$

تولید می‌کنیم که در آن مولفه‌های مدل به صورت

$$\beta = (-1, -1, 2, 0, \dots, 0), \quad b_i \sim \mathcal{N}(0, \sigma_b^2); \quad \sigma_b^2 = 0.25,$$

$$\mathbf{X}_{ij}^\top = (x_{1,ij}, \dots, x_{p_n,ij}) \sim \mathcal{U}(-1, 1)$$

$$f(t_{ij}) = 2 \sin(2\pi t_{ij}), \quad t_{ij} \sim \mathcal{U}(0, 1),$$

تعیین شده‌اند. تعداد متغیرهای مدل، p_n ، با افزایش حجم نمونه افزایش می‌یابد اما بعد متغیرهای تاثیرگذار در مدل ثابت و برابر ۳ است. با توجه به مطالب بیان شده در بخش ۱.۳.۳.۲، برای (n, p_n) موارد $(200, 16)$ ، $(100, 14)$ ، $(50, 11)$ را برای حالت $p_n < n$ و برای حالت $p_n > n$ موارد $(200, 2000)$ ، $(100, 500)$ ، $(30, 100)$ در نظر گرفته شده است.

برای مقایسه عملکرد روش‌های برآورد، از معیارهای مشابه بخش ۱.۳.۳.۲ از فصل ۲ استفاده شد. نتایج مربوط به این معیارها برای سه روش برآورد بیان شده و به ازای مقادیر مختلف (n, p_n) در جداول ۱.۳ و ۲.۳ گزارش شده است. مشابه استدلالی که در بخش ۱.۳.۳.۲ آمده است، روش $P - GSMM$ ، از نظر دقت برآورد و همچنین انتخاب متغیر بهتر از دو مدل دیگر عمل می‌کند. نتایج جدول ۲.۳ نیز گواهی بر صحت قضایای حدی بیان شده است.

همچنین رفتار تابع $f(t)$ از طریق رسم نمودارهایی بررسی شده است. این نمودارها شامل منحنی‌های $\hat{f}(t)$ ، اربیی، انحراف استاندارد و احتمال پوشش هستند. شکل‌های ۱.۳ و ۲.۳ به

جدول ۱۰۳: نتایج شبیه‌سازی: مقایسه نتایج انتخاب متغیر روش‌های $P - GSMM$ ، $P - GLMM$ و $P - GLMM$ برای حالت‌های $p_n < n$ و $p_n > n$.

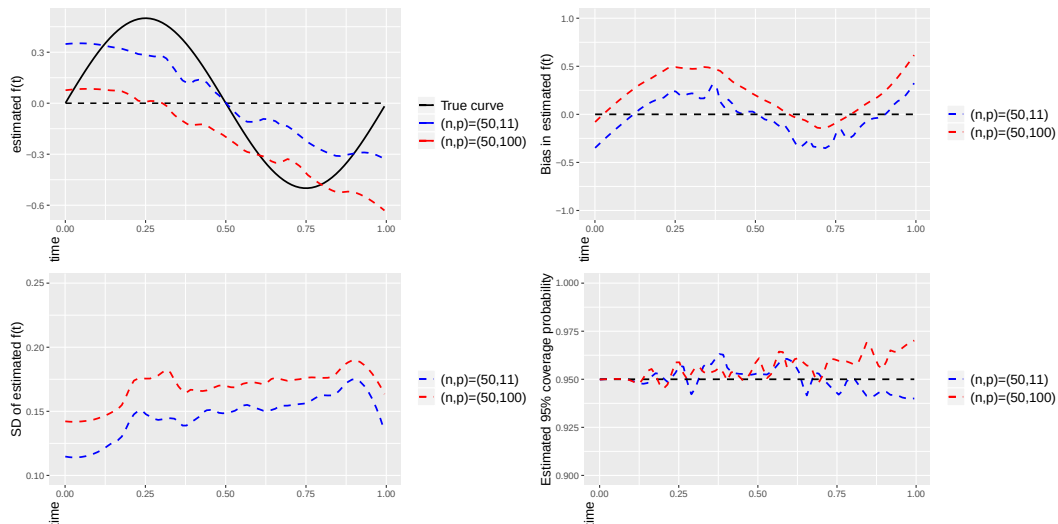
method	حالت $p_n < n$						حالت $p_n > n$					
	$(n, p) = (50, 11)$						$(n, p) = (50, 100)$					
	MSE	C(λ)	I(°)	UF	CF	OF	MSE	C(۹۷)	I(°)	UF	CF	OF
GSMM	۰/۱۱۶	۰/۰۹	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰	۶۸۰۲۸	۰/۰۷۴	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰
P - GLMM	۰/۰۶۰	۶/۵۴	۰/۰۰	۰/۰۰	۰/۱۳	۰/۸۷	۰/۴۳۵	۹۶/۴۸	۰/۰۰	۰/۰۰	۰/۵۵	۰/۴۵
P - GSMM	۰/۰۵۲	۷/۵۹	۰/۰۰	۰/۰۰	۰/۶۴	۰/۳۶	۰/۳۹۱	۹۶/۴۱	۰/۰۰	۰/۰۰	۰/۶۰	۰/۴۰
	$(n, p) = (100, 14)$						$(n, p) = (100, 500)$					
	MSE	C(۱۱)	I(°)	UF	CF	OF	MSE	C(۹۷)	I(°)	UF	CF	OF
GSMM	۰/۰۷۲	۰/۱۶	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰	۱۴۹۹/۱۳۶	۴۷/۰۲	۰/۰۳	۰/۰۳	۰/۰۰	۱/۰۰
P - GLMM	۰/۰۴۱	۱۰/۵۲	۰/۰۰	۰/۰۰	۰/۷۷	۰/۲۳	۰/۰۶۲	۴۹۵/۷۲۰	۰/۰۰	۰/۰۰	۰/۸۹	۰/۱۱
P - GSMM	۰/۰۳۶	۱۰/۷۰	۰/۰۰	۰/۰۰	۰/۹۳	۰/۰۷	۰/۰۳۸	۴۹۶/۲۵۰	۰/۰۰	۰/۰۰	۰/۹۲	۰/۰۸
	$(n, p) = (150, 15)$						$(n, p) = (200, 2000)$					
	MSE	C(۱۲)	I(°)	UF	CF	OF	MSE	C(۱۹۹۷)	I(°)	UF	CF	OF
GSMM	۰/۰۶۰	۰/۲۶	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰	۱۲۵/۴۰۶	۱۱۳۷/۶۲	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰
P - GLMM	۰/۰۴۴	۱۱/۲۵	۰/۰۰	۰/۰۰	۰/۹۲	۰/۰۸	۰/۰۱۸	۱۹۹۶/۸۹	۰/۰۰	۰/۰۰	۰/۳۰	۰/۷۰
P - GSMM	۰/۰۴۵	۱۱/۸۷	۰/۰۰	۰/۰۰	۰/۹۶	۰/۰۴	۰/۰۱۸	۱۹۹۶/۹۳	۰/۰۰	۰/۰۰	۰/۵۴	۰/۴۶

جدول ۲۰۳: نتایج شبیه‌سازی: کارایی روش $P - GSMM$ برای حالت‌های $p_n < n$ و $p_n > n$.

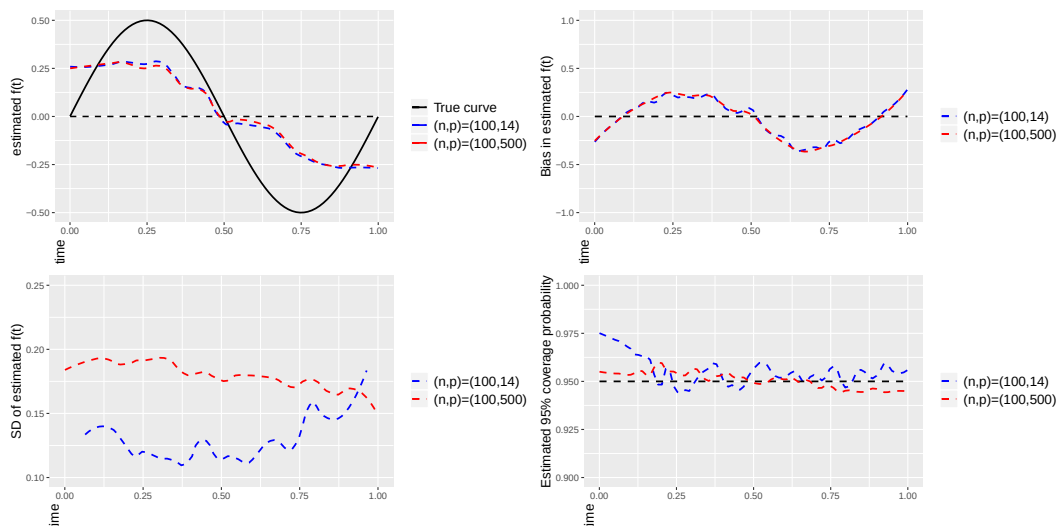
حالت $p_n < n$				حالت $p_n > n$					
(n, p_n)		β_1	β_2	β_3	(n, p_n)		β_1	β_2	β_3
$(50, 11)$	<i>Bias</i>	۰/۰۴۷	۰/۰۹۶	۰/۰۶۹	$(50, 100)$	<i>Bias</i>	۰/۱۴۷	۰/۳۴۴	۰/۳۶۵
	<i>SD1</i>	۰/۰۹۲	۰/۰۸۵	۰/۱۱۳		<i>SD1</i>	۰/۰۸۴	۰/۰۷۱	۰/۰۹۴
	<i>SD2</i>	۰/۰۹۷	۰/۰۹۴	۰/۰۹۷		<i>SD2</i>	۰/۱۶۲	۰/۱۷۸	۰/۱۴۸
	<i>CP</i>	۰/۹۶	۰/۹۵	۰/۹۲		<i>CP</i>	۰/۹۵	۰/۹۶	۰/۹۷
$(100, 14)$	<i>Bias</i>	۰/۰۷۶	۰/۱۰۳	۰/۰۵۳	$(100, 500)$	<i>Bias</i>	۰/۰۷۸	۰/۰۶۱	۰/۰۱۷
	<i>SD1</i>	۰/۰۷۱	۰/۰۶۷	۰/۰۸۴		<i>SD1</i>	۰/۰۴۹	۰/۰۴۴	۰/۰۶۶
	<i>SD2</i>	۰/۰۷۲	۰/۰۶۷	۰/۰۷۸		<i>SD2</i>	۰/۰۷۰	۰/۰۷۲	۰/۰۸۶
	<i>CP</i>	۰/۹۶	۰/۹۵	۰/۹۶		<i>CP</i>	۰/۹۴	۰/۹۵	۰/۹۶
$(150, 15)$	<i>Bias</i>	۰/۱۰۱	۰/۱۲۴	۰/۰۹۹	$(200, 2000)$	<i>Bias</i>	۰/۰۴۹	۰/۰۷۲	۰/۰۲۸
	<i>SD1</i>	۰/۰۶۰	۰/۰۵۹	۰/۰۷۰		<i>SD1</i>	۰/۰۳۹	۰/۰۴۰	۰/۰۵۴
	<i>SD2</i>	۰/۰۵۲	۰/۰۵۹	۰/۰۵۹		<i>SD2</i>	۰/۰۵۱	۰/۰۵۲	۰/۰۵۸
	<i>CP</i>	۰/۹۷	۰/۹۶	۰/۹۵		<i>CP</i>	۰/۹۴	۰/۹۲	۰/۹۴

ترتیب نتایج مربوط به حجم نمونه ۵۰ و ۱۰۰ هستند. در هر دو تصویر نتایج برای حالت‌های $p_n < n$ و $p_n > n$ رسم شده‌اند. مشاهده می‌شود که منحنی‌های برآورد شده تقریباً دارای تابع سینوسی هستند و این امر در حجم نمونه بیشتر یعنی شکل ۲۰۳ مشهودتر است. منحنی اریبی برآوردگرها حول خط صفر در نوسان هستند و برای حجم نمونه ۱۰۰ میزان اریبی کمتر از حجم نمونه ۵۰ است. همچنین برای حجم نمونه ۵۰ میزان اریبی حالت $p_n < n$ کمتر از $p_n > n$ است در حالی که در حجم نمونه بالا

میزان اریبی در دو حالت رفتار یکسانی را نشان می‌دهد. در هر دو شکل مقادیر انحراف استاندارد حالت $p_n < n$ کمتر از $p_n > n$ است. همچنین منحنی احتمال پوشش در هر دو شکل حول خط ۹۵٪ در نوسان است.



شکل ۱.۳: نتایج شبیه‌سازی: نتایج مربوط به $f(t)$ برآوردشده برای حالت‌های $p_n < n$ و $p_n > n$ با حجم نمونه $n = 50$. تصویر بالا-چپ: منحنی $f(t)$ برآوردشده، تصویر بالا-راست: منحنی اریبی، تصویر پایین-چپ: منحنی انحراف استاندارد، تصویر پایین-راست: احتمال پوشش.



شکل ۲.۳: نتایج شبیه‌سازی: نتایج مربوط به $f(t)$ برآوردشده برای حالت‌های $p_n < n$ و $p_n > n$ در مطالعه شبیه‌سازی با حجم نمونه $n = 100$. تصویر بالا-چپ: منحنی $f(t)$ برآوردشده، تصویر بالا-راست: منحنی اریبی، تصویر پایین-چپ: منحنی انحراف استاندارد، تصویر پایین-راست: احتمال پوشش.

۲.۴.۱.۳ مثال واقعی

تحلیل داده‌های $CD4$

در این بخش به مطالعه رفتار برآوردگرهای پیشنهادی با استفاده از داده‌های $CD4$ می‌پردازیم. مشابه بخش ۳.۳.۲ علاوه بر اثرات اصلی، اثرات متقابل را نیز وارد مدل کردیم. با این تفاوت که در آن بخش، تبدیل ریشه دوم متغیر پاسخ (تعداد سلول‌های $CD4$) در نظر گرفته شد تا توزیع داده‌ها به توزیع نرمال نزدیک باشد؛ اما در اینجا بدون تبدیل متغیرهای پاسخ، مدل رگرسیون پواسن را در نظر می‌گیریم. متغیر $YEAR$ به عنوان جزء ناپارامتری و سایر متغیرها به صورت خطی وارد مدل می‌شوند. بدین ترتیب یک مدل آمیخته نیمه‌پارامتری تعمیم‌یافته به داده‌ها برازش داده می‌شود. به منظور مقایسه عملکرد روش $P-GSMM$ با دو روش دیگر یعنی $P-GLMM$ و $GSMM$ از معیار انحراف استاندارد جانبی (SD) استفاده و برای شناسایی بهترین مدل پشتیبانی‌شده توسط داده‌ها، معیار اطلاع آکائیک (AIC) و معیار اطلاع بیزی (BIC) محاسبه می‌شود. نتایج حاصل در جدول ۳.۳ گزارش شده است. روش $P-GSMM$ مقادیر SD کمتری در مقایسه با دو روش دیگر یعنی $P-GLMM$ و $GSMM$ دارد. همچنین معیارهای AIC و BIC مدل پیشنهادی در مقایسه با دو مدل رقیب دیگر مقادیر کمتری دارند که نشان‌دهنده عملکرد بهتر آن است. تحت مدل $P-GSMM$ ، متغیرهای $SMOKE$ ، $DRUGS$ ، $SEXP$ ، $SOMKE * DRUG$ و $DRUG * SEXP$ به عنوان متغیرهای مهم شناسایی شده‌اند. انتخاب متغیر تحت مدل $P-GLMM$ کمی متفاوت بوده است و علاوه بر متغیرهای ذکر شده، متغیرهای $AGE * SMOKE$ ، $AGE * DRUG$ و $SMOKE * SEXP$ نیز وارد مدل شده‌اند. در این تحلیل برخی از اثرات متقابل نیز به عنوان عوامل مهم وارد مدل شده‌اند که در تحقیقات قبلی وانگ و همکاران (۲۰۰۵) و هووانگ و همکاران (۲۰۰۷) نادیده گرفته شده بودند.

نتایج مربوط به برآورد قسمت ناپارامتری مدل توسط نمودارهایی در شکل ۳.۳ ارائه شده است. این شکل شامل نمودار تابع ناپارامتری برآورد شده $f(t)$ به همراه نمودارهای فواصل اطمینان ۹۵٪، منحنی‌های انحراف استاندارد جانبی و نمونه‌ای، و منحنی‌های احتمال پوشش جانبی و نمونه‌ای می‌باشند. مشاهده می‌شود که با گذشت زمان تابع ناپارامتری روند نزولی دارد. همچنین در نقاط مرزی بین انحراف استاندارد جانبی و انحراف استاندارد نمونه‌ای اختلاف چشمگیری وجود دارد. در این نقاط انحراف استاندارد جانبی به شدت کاهش می‌یابد، این امر باعث می‌شود که احتمالات پوشش از ۹۵٪ به ۹۲٪ تنزل یابند.

تحلیل داده‌های چرخه سلولی مخمر بیان‌ژن ORF

در این بخش به داده‌های چرخه سلولی مخمر بیان‌ژن ORF از بخش ۲.۴.۱، مدل آمیخته نیمه‌پارامتری به صورت

$$y_{ij} = \sum_{k=1}^{96} x_{ij}^{(k)} + f(t_{ij}) + b_i; \quad i = 1, \dots, 283; j = 1, \dots, 4$$

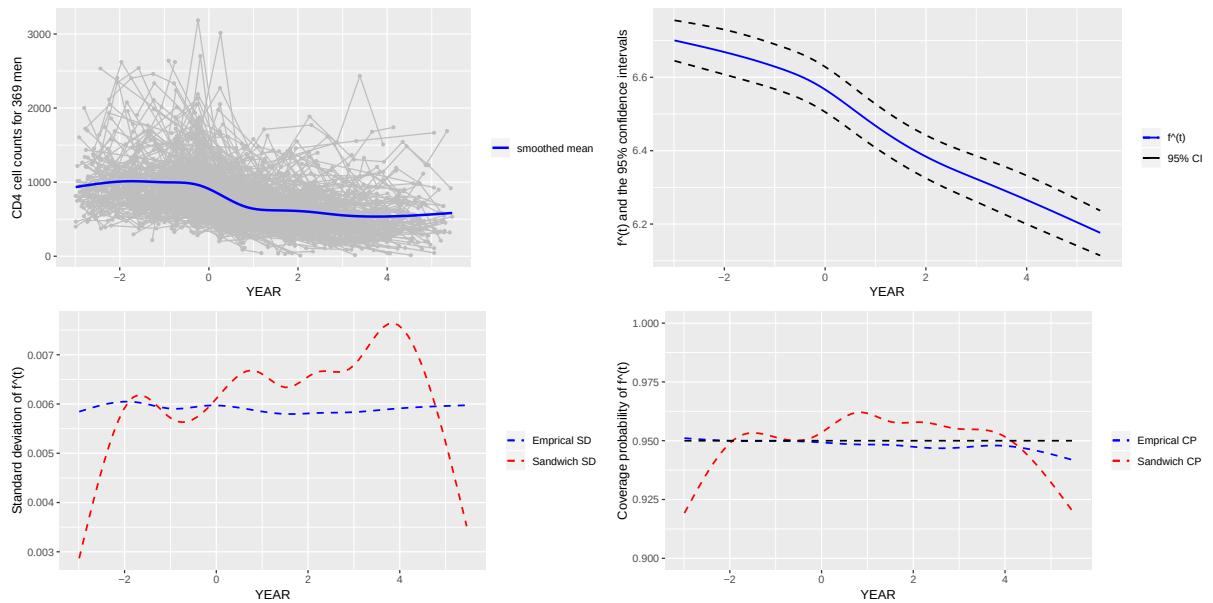
را برازش می‌دهیم. هدف شناسایی TF ‌هایی است که ممکن است با الگوهای بیان‌ژن ۲۸۳ ژن

جدول ۳.۳: نتایج تحلیل داده‌های $CD4$: برآورد پارامترهای رگرسیونی (SD) تحت برازش سه مدل $P - GLMM$ و $GSMM$ ، $P - GSMM$.

متغیرها	$GSMM$ $\hat{\beta}(SE)$	$P - GLMM$ $\hat{\beta}(SE)$	$P - GSMM$ $\hat{\beta}(SE)$
AGE	۰/۰۷۳ (۰/۰۳۹)	-۰/۰۹۲ (۰/۰۵۱)	۰ (۰)
$SMOKE$	۰/۱۸۸ (۰/۱۷۹)	۰/۸۸۸ (۰/۱۹۲)	۰/۰۷۹ (۰/۰۴۵)
$DRUG$	۰/۱۳۰ (۰/۱۴۳)	۶/۰۶۸ (۰/۱۲۵)	۰/۱۴۲ (۰/۰۷۴)
$SEXP$	-۰/۰۴۹ (۰/۰۳۱)	۰/۶۷۲ (۰/۰۳۰)	۰/۰۱۷ (۰/۰۱۲)
$CESD$	-۰/۰۰۱ (۰/۰۱۱)	۰ (۰)	۰ (۰)
$AGE * SMOKE$	۰/۰۰۲ (۰/۰۱۴)	۰/۰۱۴ (۰/۰۰۴)	۰ (۰)
$AGE * DRUG$	-۰/۰۳۴ (۰/۰۲۴)	۰/۰۳۲ (۰/۰۳۵)	۰ (۰)
$AGE * SEXP$	-۰/۰۰۹ (۰/۰۰۳)	۰ (۰)	۰ (۰)
$AGE * CESD$	۰/۰۰۱ (۰/۰۰۲)	۰ (۰)	۰ (۰)
$SMOKE * DRUG$	۰/۰۰۹ (۰/۰۵۴)	-۰/۵۸۴ (۰/۱۵۰)	-۰/۰۱۴ (۰/۰۳۸)
$SMOKE * SEXP$	-۰/۰۱۰ (۰/۰۱۲)	-۰/۰۳۴ (۰/۰۱۰)	۰ (۰)
$SMOKE * CESD$	-۰/۰۰۶ (۰/۰۰۹)	۰ (۰)	۰ (۰)
$DRUG * SEXP$	-۰/۰۲۵ (۰/۰۱۹)	-۰/۵۹۸ (۰/۰۴۱)	-۰/۰۲۲ (۰/۰۱۲)
$DRUG * CESD$	۰/۰۰۶ (۰/۰۰۶)	۰ (۰)	۰ (۰)
$SEXP * CESD$	۰/۰۰۱ (۰/۰۰۳)	۰ (۰)	۰ (۰)
$\ell_{\max}$	۸۴۶۳۰۰۷	۷۵۲۹۱۵۸	۸۶۲۴۴۲۹
AIC	-۱۶۹۲۵۹۸۳	-۱۵۰۵۸۲۸۶	-۱۷۲۴۸۸۲۷
BIC	-۱۶۹۲۵۹۲۴	-۱۵۰۵۸۲۲۸	-۱۷۲۴۸۷۶۹

تنظیم‌شده از چرخه سلولی مرتبط باشد. بنابراین یک روش تاوانیده را بر اساس مدل $P - GSMM$ به کار می‌بریم. همچنین با نادیده گرفتن مولفه ناپارامتری $f(t_{ij})$ ، مدل $P - GLMM$ را نیز استفاده می‌کنیم.

جدول ۴.۳، خلاصه‌ای از نتایج مربوط به TF ‌های برگزیده‌شده بر اساس مدل‌های $P - GSMM$ و $P - GLMM$ است که نشان می‌دهد در مجموع به تعداد ۱۳ و ۱۶ تا از TF ‌ها مربوط به فرآیند چرخه سلولی مخمر به ترتیب توسط $P - GSMM$ و $P - GLMM$ شناسایی می‌شوند. تعداد TF ‌های انتخاب‌شده که در دو مدل مشترک باشند کم بوده و شامل $GAT3$ ، $MBP1$ ، $MSN4$ ، $PHD1$ ،



شکل ۳.۳: نتایج تحلیل داده‌های $CD4$: نتایج مربوط به $f(t)$ برآورد شده تحت مدل $P-GSMM$. تصویر بالا-چپ: نمودار پراکنش متغیر پاسخ، تصویر بالا-راست: $\hat{f}(t)$ (آبی) و فواصل اطمینان (مشکی)، تصویر پایین-چپ: منحنی‌های انحراف استاندارد مجانبی (آبی) و نمونه‌ای (قرمز)، تصویر پایین-راست: منحنی احتمالات پوشش مجانبی (آبی) و نمونه‌ای (قرمز).

جدول ۴.۳: نتایج تحلیل داده‌های چرخه سلولی مخمر ORF : برآورد پارامترهای رگرسیونی (SD) تحت برازش سه مدل $P-GSMM$ ، $P-GLMM$ و $P-GSMM$.

متغیرها	$P-GLMM$ $\hat{\beta}(SE)$	$P-GSMM$ $\hat{\beta}(SE)$	متغیرها ادامه	$P-GLMM$ $\hat{\beta}(SE)$	$P-GSMM$ $\hat{\beta}(SE)$
$ARGA1$	۰/۰۲۲ (۰/۰۱۹)	۰ (۰)	$PHD1$	۰/۰۶۵ (۰/۰۲۷)	-۰/۰۱۹ (۰/۰۰۶)
$DOT6$	۰/۰۱۸ (۰/۰۱۷)	۰ (۰)	$RAP1$	۰/۰۵۳ (۰/۰۲۷)	۰ (۰)
$FKH1$	۰ (۰)	۰/۰۰۳ (۰/۰۰۵)	$RGM1$	۰ (۰)	-۰/۰۲۲ (۰/۰۱۳)
$FKH2$	۰ (۰)	۰/۱۶۶ (۰/۰۰۸)	$RLM1$	۰ (۰)	-۰/۰۰۲ (۰/۰۰۴)
$GAT1$	-۰/۰۰۳ (۰/۰۰۷)	۰ (۰)	$RME1$	۰/۰۷۲ (۰/۰۲۸)	۰ (۰)
$GAT3$	۰/۰۱۲ (۰/۰۱۴)	-۰/۰۲۲۳ (۰/۰۱۲)	$SMP1$	۰/۰۴۵ (۰/۰۲۴)	-۰/۰۱۵ (۰/۰۰۶)
$MBP1$	۰/۱۴۷ (۰/۰۳۵)	-۰/۱۴۷۷ (۰/۰۰۷)	$STB1$	۰ (۰)	-۰/۰۰۸ (۰/۰۰۵)
$MIG1$	-۰/۰۰۳ (۰/۰۰۷)	۰ (۰)	$STP1$	۰/۰۰۲ (۰/۰۰۵)	۰ (۰)
$MSN4$	۰/۰۶۰ (۰/۰۲۷)	-۰/۰۰۸ (۰/۰۰۶)	$SWI4$	۰/۰۷۶ (۰/۰۳۰)	-۰/۰۰۷ (۰/۰۰۶)
$NDD1$	۰ (۰)	۰/۰۸۴ (۰/۰۰۸)	$SWI6$	۰/۱۱۵۱ (۰/۰۳۴)	-۰/۰۲۰ (۰/۰۰۷)
$PDR1$	۰/۰۲۲۸ (۰/۰۱۷)	۰ (۰)	$YAP5$	۰/۰۰۷ (۰/۰۱۱)	۰ (۰)
$\ell_{\max}$	-5.71×10^{14}	-1.58×10^{14}			
AIC	1.14×10^{15}	3.17×10^{14}			
BIC	1.14×10^{15}	3.17×10^{14}			

$SWI6$ و $SWI4$ ، $SMP1$ است. در آزمایشات بیولوژیکی پروتئین‌های $MBP1$ ، $SWI4$ و $SWI6$ به عنوان سه TF مهم در مرحله $G1$ از چرخه سلولی شناخته می‌شوند که این TF ها در هر دو مدل وارد شده‌اند. با این حال، معیارهای انتخاب مدل، از جمله مقادیر خطاهای استاندارد و معیارهای AIC و BIC عملکرد بهتر مدل پیشنهادی ما را تأیید می‌کنند.

۲.۳ انتخاب متغیر برای داده‌های با بعد پایین

فرآیند انتخاب متغیر در مدل آمیخته نیمه‌پارامتری تعمیم‌یافته برای داده‌های طولی غیرنرمال با بعد پایین حالت خاصی از بعد بالا است، با این تفاوت که بعد مدل p_n ثابت است و با افزایش حجم نمونه تغییر نمی‌کند. لذا در این بخش تنها به بیان قضایای حدی مربوط به آن می‌پردازیم. تمام شرایط نظم بیان‌شده در بخش ۳.۱.۳ در این جا نیز برقرار است با این تفاوت که تعداد کل متغیرها p_n و تعداد متغیرهای مهم s_n ثابت بوده و به p و s تغییر می‌یابند.

قضیه ۱.۲.۳. تحت شرایط بیان‌شده در بخش ۳.۱.۳، اگر تعداد گره‌ها به صورت $L_n = O_p(n^{1/(2r+1)})$ باشد، آن‌گاه به ازای هر $\xi_n \in \mathcal{R}^p$ که $\|\xi_n\| = 1$ ، داریم

$$\mathbb{P}(\hat{\beta}_{n2} = \circ) \rightarrow 1,$$

$$\xi_n^\top \overline{M}_n^{*-1/2}(\beta_{n\circ}) \overline{H}_n^*(\beta_{n\circ})(\hat{\beta}_{n1} - \beta_{n\circ1}) \xrightarrow{d} \mathcal{N}(\circ, 1),$$

که در آن‌ها

$$\overline{M}_n^* = E_{u|y} \left[\sum_{i=1}^n X_i^{*\top} A_i^{\dagger}(\theta_n, u_i) \overline{R}^{-1} R_{\circ} \overline{R}^{-1} A_i^{\dagger}(\theta_n, u_i) X_i^* \right],$$

$$\overline{H}_n^* = E_{u|y} \left[\sum_{i=1}^n X_i^{*\top} A_i^{\dagger}(\theta_n, u_i) \overline{R}^{-1} A_i^{\dagger}(\theta_n, u_i) X_i^* \right],$$

$$X_i^* = (I - P)X_i$$

$$P = B(B^\top \Omega B)^{-1} B^\top \Omega$$

$$\Omega = \text{diag}\{\Omega_i\}$$

$$\Omega_i = E_{u|y} \left[A_i^{\dagger}(\theta_n, u_i) \overline{R}^{-1} A_i^{\dagger}(\theta_n, u_i) \right].$$

فصل ۴

انتخاب متغیر برای داده‌های طولی چندمتغیره

در تحقیقات میدانی معمولاً چند متغیر پاسخ به طور همزمان اندازه‌گیری می‌شوند که این متغیرها با یکدیگر همبستگی داشته و داده‌های حاصل را داده‌های چندمتغیره می‌نامند. در مطالعات طولی نیز داشتن داده‌هایی با بیش از یک متغیر پاسخ که به طور مکرر در یک موضوع خاص و برای یک دوره زمانی خاص اندازه‌گیری شوند، منجر به داده‌های طولی چندمتغیره می‌شود. روش‌های مدل‌سازی ویژه‌ای برای رویارویی با داده‌های طولی چندمتغیره توسعه یافته‌اند. در ابتدا شاه و همکاران (۱۹۹۷) مدل آمیخته خطی چندمتغیره^۱ (*MLMM*) را معرفی نمود و برای برآورد پارامترها از الگوریتم *EM* استفاده کرد. سپس ساموئل و همکاران (۱۹۹۹)، سونگ و همکاران (۲۰۰۲)، روی (۲۰۰۶) و وانگ و فن (۲۰۱۰) به تحقیقات بیشتر پیرامون این مدل پرداختند. در تحقیقات ذکرشده با پذیرش نرمال بودن اثرات آمیخته و خطاهای مدل به ارائه راه‌حلی برای برآورد پارامترهای مدل و اثرات آمیخته با بیشینه کردن تابع درستنمایی پرداخته شده است. در عمل به دلیل غیر قابل مشاهده بودن اثرات آمیخته، بررسی نرمال بودن آن‌ها مقدور نیست و پذیرش ناصحیح این فرض می‌تواند بر روی دقت برآورد پارامترها و پیشگوها تاثیر سوء داشته باشد. اگرچه بررسی هر توزیع دیگری نیز امکان‌پذیر نیست اما اگر مطالعاتی داشته باشیم که برای اثرات آمیخته توزیع نرمال در نظر گرفته شده باشد، می‌توان از برخی توزیع‌های رقیب برای بهبود نتایج مدل نرمال، استفاده کرد. اما مهم‌تر هنگامی

¹Multivariate linear mixed model

که داده‌ها دارای مقادیر پرت باشند، پذیرش توزیع نرمال برای خطاهای مدل می‌تواند گمراه‌کننده باشد. برای مدل‌سازی تنومند داده‌های طولی در حضور مقادیر پرت یا مشاهدات غیرمعمول، اتخاذ یک رده انعطاف‌پذیرتر از توزیع‌ها به عنوان جایگزینی برای توزیع نرمال پیشنهاد شده است. در این راستا پینه‌رو و همکاران (۲۰۰۱) با در نظر گرفتن توزیع t برای اثرات تصادفی و خطاهای مدل، به تعمیم LMM پرداخته و مدلی تحت عنوان مدل آمیخته خطی $tMLMM$ را پیشنهاد دادند. توسعه بیشتر این روش توسط روزا و همکاران (۲۰۰۴)، لین و لی (۲۰۰۶)، لین و لی (۲۰۰۷) و سونگ و همکاران (۲۰۰۷) انجام گرفته است. با این حال روش آن‌ها به داده‌هایی با یک متغیر پاسخ محدود می‌شود. برای استنباط تنومند در برابر مقادیر پرت برای داده‌های طولی چند متغیره، وانگ و فن (۲۰۱۱) نسخه چندمتغیره مدل $tLMM$ را تحت عنوان مدل آمیخته خطی چندمتغیره ($MtLMM$) گسترش دادند. مقالات وانگ (۲۰۱۳)، وانگ و فن (۲۰۱۲) و وانگ و لین (۲۰۱۵) از جمله تحقیقات جامعی هستند که استراتژی‌های محاسباتی و برنامه‌های کاربردی $MtLMM$ را پوشش می‌دهند.

علیرغم سازگاری $MtLMM$ با داده‌های چندمتغیره و تنومند بودن در مقابل مقادیر پرت، این مدل‌ها تنها محدود به رابطه خطی بین متغیرهای پاسخ و متغیرهای تبیینی هستند. اخیراً وانگ و لین (۲۰۱۴) و وانگ و کاسترو (۲۰۱۸) به مطالعه مدل آمیخته غیرخطی t چندمتغیره ($MtNLMM$) پرداخته‌اند که تعمیم مدل آمیخته غیرخطی چندمتغیره مارشال و همکاران (۲۰۰۶) به شمار می‌آید. اگرچه $MtNLMM$ مانند $MNLMM$ می‌تواند طیف گسترده‌تری از کاربردهای عملی را نسبت به مدل‌های خطی فراهم کند، اما به طور قابل توجهی پیچیده‌تر و از نظر محاسباتی سنگین‌تر هستند. تا کنون استفاده از مدل‌های نیمه‌پارامتری برای داده‌های طولی چندمتغیره و مسئله انتخاب متغیر در $MtLMM$ نسبتاً نادر بوده است. بر این اساس، در این فصل ابتدا به مسئله انتخاب متغیر در $MtLMM$ پرداخته می‌شود. سپس نسخه نیمه‌پارامتری $MtLMM$ تحت عنوان مدل آمیخته نیمه‌پارامتری چندمتغیره ($MtSMM$) معرفی می‌شود.

۱.۴ انتخاب متغیر در مدل آمیخته با توزیع t چندمتغیره

۱.۱.۴ معرفی $MtLMM$

n نمونه تصادفی از داده‌های طولی را در نظر بگیرید، به طوری که نمونه i ام دارای r متغیر پاسخ و s_i مشاهده مکرر در طول زمان باشد. فرض کنید

$$Y_i = \begin{pmatrix} y_{i1,1} & y_{i2,1} & \cdots & y_{ir,1} \\ y_{i1,2} & y_{i2,2} & \cdots & y_{ir,2} \\ \vdots & \vdots & \ddots & \vdots \\ y_{i1,s_i} & y_{i2,s_i} & \cdots & y_{ir,s_i} \end{pmatrix} = [y_{i1} : \cdots : y_{ir}] = [y_{i1}^\top : \cdots : y_{i,s_i}^\top]^\top$$

ماتریس متغیرهای پاسخ $s_i \times r$ بعدی برای نمونه i ام به ازای $i = 1, \dots, n$ باشد، به طوری که برداری ستونی از مشاهدات مکرر متغیر پاسخ j ام، $j = 1, \dots, r$ ، در طول زمان است. $y_{ij,k} = (y_{ij,1}, \dots, y_{ij,s_i})^T$ بردار سطری از متغیرهای پاسخ مشاهده شده در موقعیت زمانی k ام، $k = 1, \dots, s_i$ ، است. همچنین فرض کنید $E_i = [e_{i1} : \dots : e_{ir}] = [e_{i1}^T : \dots : e_{ir}^T]^T$ ماتریس $s_i \times r$ بعدی از خطاهای درون-گروهی، $e_{ij} = (e_{ij,1}, \dots, e_{ij,n_i})^T$ بردار ستونی خطاهای هم‌ارز با y_{ij} و $e_{i,k} = (e_{i1,k}, \dots, e_{ir,k})$ بردار سطری خطاهای هم‌ارز با $y_{i,k}$ هستند. برای سهولت در نمادگذاری از عملگر بردار ساز^۲ (vec) استفاده می‌شود و بردارهای پاسخ و خطای مدل به ترتیب به صورت $y_i = \text{vec}(Y_i)$ و $\varepsilon_i = \text{vec}(E_i)$ با بعد $n_i = s_i r$ تعریف می‌شوند.

برای نمونه i ام $MtLMM$ به صورت

$$y_i = X_i \beta + Z_i b_i + \varepsilon_i, \quad (1.4)$$

بیان می‌شود و فرض می‌کنیم ε_i و b_i دارای توزیع توام t چندمتغیره به صورت زیر است:

$$\begin{bmatrix} b_i \\ \varepsilon_i \end{bmatrix} \sim \mathcal{T}_{q+n_i} \left(\begin{bmatrix} 0 \\ 0 \end{bmatrix}, \begin{bmatrix} D & 0 \\ 0 & R_i \end{bmatrix}, \nu \right). \quad (2.4)$$

در مدل (۱.۴)، X_i و Z_i به صورت $X_i = \text{diag}\{X_{i1}, \dots, X_{ir}\}$ و $Z_i = \text{diag}\{Z_{i1}, \dots, Z_{ir}\}$ تعریف می‌شوند، به طوری که X_{ij} یک ماتریس از اثرات ثابت با بعد $s_i \times p_j$ و هم‌ارز با پاسخ j ام از نمونه i است و Z_{ij} که معمولاً توسط زیرمجموعه‌ای از X_{ij} ها ساخته می‌شود، یک ماتریس طرح $s_i \times q_j$ بعدی مربوط به اثرات تصادفی است. ساختار قطری بلوکی X_i و Z_i به واسطه تعیین ماتریس طرح مجزا برای هر متغیر پاسخ، به تحلیل‌گران اجازه می‌دهد که به راحتی روابط بین متغیرهای پاسخ و تبیینی را که ممکن است در فواصل زمانی نابرابر جمع‌آوری شده باشند، به یکدیگر پیوند دهند. اگر بعد کلی اثرات ثابت با $p = \sum_{j=1}^r p_j$ و بعد کلی اثرات تصادفی با $q = \sum_{j=1}^r q_j$ مشخص شوند، آنگاه $\beta = (\beta_1^T, \dots, \beta_r^T)^T$ یک بردار $p \times 1$ بعدی از اثرات ثابت است و β_j زیربرداری $1 \times p_j$ بعدی است که برای توصیف مشخصات پاسخ j ام مورد استفاده قرار می‌گیرد. $b_i = (b_{i1}^T, \dots, b_{ir}^T)^T$ یک بردار $q \times 1$ بعدی از اثرات تصادفی مشاهده نشده است و b_{ij} زیربرداری $1 \times q_j$ بعدی هم‌ارز با y_{ij} است. در فرضیات توزیعی رابطه (۲.۴)، D و R_i به ترتیب ماتریس‌های کوواریانس مربوط به اثرات تصادفی و خطاهای مدل هستند. درجه آزادی ν یک پارامتر تنظیم‌کننده است که میزان سنگینی دم توزیع را تعیین می‌کند. علاوه بر این، $D = [D_{jj'}]$ یک ماتریس $q \times q$ بعدی متقارن و معین مثبت با عناصر $D_{jj'}$ است. به ازای $j = j'$ ساختار کوواریانس اثرات تصادفی برای پاسخ j ام است و به ازای $j \neq j'$ ساختار کوواریانس اثرات تصادفی برای جفت متغیر پاسخ j و j' است. به منظور شناسایی بهتر ماتریس کوواریانس R_i ، فرض کنید به ازای $j = 1, \dots, r$ ، داشته باشیم $e_{ij} \sim \mathcal{T}_{s_i}(\circ, \sigma_{jj} \Omega_i, \nu)$ و به ازای $k = 1, \dots, s_i$ ، داشته باشیم $e_{i,k} \sim \mathcal{T}_r(\circ, \Sigma, \nu)$ به طوری که $\Sigma = [\sigma_{jj'}]$ واریانس‌ها و کوواریانس‌های بین r متغیر پاسخ را توصیف می‌کند و Ω_i

^۲Vectorial operation

ماتریس همبستگی وابسته به موقعیت‌های زمانی است که همبستگی سریالی بین s_i موقعیت زمانی را نشان می‌دهد. بر این اساس، ماتریس خطای درون-گروهی E_i ، دارای توزیع t ماتریسی کیریا (۲۰۰۶) بوده و بردار خطای ε_i دارای توزیع t چندمتغیره با پارامتر مکان $\circ$ و پارامتر مقیاس R_i است که توسط ساختار ضرب کرونکر^۳ (KP) به صورت $R_i = \Sigma \otimes \Omega_i$ تعریف می‌شود. ساختار KP کمک می‌کند تا R_i را دقیق‌تر برآورد کنیم. برای غلبه بر مسئله عدم شناساپذیری ناشی از جواب‌های غیریکتا برای Σ و Ω_i در برآورد R_i با ساختار KP ، لازم است که Ω_i را یک ماتریس همبستگی بدانیم نه یک ماتریس کوواریانس. برای این که Ω_i دقیق‌تر برآورد شود، می‌توان بر اساس ویژگی‌های داده‌های موجود یک ساختار پارامتری برای آن در نظر گرفت که تابعی از پارامترهای همبستگی ϕ و نقاط زمانی t_i به صورت $\Omega_i = \Omega_i(\phi, t_i)$ باشد. تحت مدل (۱۰۴) به راحتی می‌توان نتیجه گرفت که

$$\Lambda_i = Z_i D Z_i^T + R_i \text{ و } y_i \sim \mathcal{T}_{n_i}(X_i \beta, \Lambda_i, \nu)$$

به تبعیت از وانگ و فن (۲۰۱۱)، مجموعه‌ای از پارامترهای مقیاس τ_i را بکار گرفته و به کمک آن‌ها مدل $MtLMM$ در رابطه (۱۰۴) را توسط یک مدل سلسله مراتبی به صورت

$$y_i | (b_i, \tau_i) \sim \mathcal{N}_{n_i}(X_i \beta + Z_i b_i, \tau_i^{-1} R_i), \quad (3.4)$$

$$b_i | \tau_i \sim \mathcal{N}_q(\circ, \tau_i^{-1} D),$$

$$\tau_i \sim \mathcal{G}(\nu/2, \nu/2).$$

بیان می‌کنیم. بر اساس فرضیات رابطه (۲.۴)، $b_i | \tau_i$ و $\varepsilon_i | \tau_i$ از هم مستقل‌اند. با انتگرال‌گیری مدل (۳.۴) نسبت به b_i ، مدل را می‌توان به صورت

$$y_i | \tau_i \sim \mathcal{N}_{n_i}(X_i \beta, \tau_i^{-1} \Lambda_i), \quad \tau_i \sim \mathcal{G}(\nu/2, \nu/2). \quad (4.4)$$

نوشت. حال از طریق بهینه‌سازی تابع هدف مبتنی بر تابع تاوان مانند PML به برآورد و انتخاب متغیر می‌پردازیم.

۲.۱.۴ الگوریتم ECM برای PML

در بخش ۱.۳.۲ از فصل ۲، توضیحات مختصری پیرامون برآورد ML پارامترها و الگوریتم EM ارائه شد. در این بخش نیز برای برآورد پارامترهای مدل از روش مذکور استفاده می‌شود. اما در مدل $MtLMM$ ، گام M الگوریتم EM از نظر محاسباتی پیچیده است و صورت بسته‌ای برای برآوردگرها نمی‌توان یافت. برای حل این مسئله، الگوریتم دیگری به نام الگوریتم بیشینه‌سازی امید ریاضی شرطی^۴ (ECM) استفاده می‌شود. الگوریتم ECM که توسط مینگ و رابین (۱۹۹۳) معرفی شد مانند الگوریتم EM است با این تفاوت که گام M آن توسط گام بیشینه‌سازی شرطی (CM) که از نظر محاسباتی ساده‌تر است، جایگزین می‌شود.

³Kronecker product

⁴Expectation conditional maximization

۱.۲.۱.۴ برآوردگر بیشینه درستنمایی

فرض کنید $\Theta = (\beta, \text{vech}(D), \text{vech}(\Sigma), \phi, \nu)$ مجموعه‌ای از تمام پارامترهای مدل باشد. اگر اثرات تصادفی b_i را به عنوان مقادیر گمشده و مجموعه $\{(y_i, b_i), i = 1, \dots, n\}$ را داده‌های کامل در نظر بگیریم، آنگاه لگاریتم تابع درستنمایی کامل بدون در نظر گرفتن جمله ثابت به صورت

$$\ell_c(\Theta) = \sum_{i=1}^n \ell_c^{[i]}(\Theta) = \sum_{i=1}^n \frac{1}{\nu} \left\{ \log |R_i^{-1}| + \log |D^{-1}| - \tau_i [\varepsilon_i^\top R_i^{-1} \varepsilon_i + b_i^\top D^{-1} b_i] + \nu \log(\frac{\nu}{\nu}) - \nu \log \frac{\nu}{\nu} + \nu(\log \tau_i - \tau_i) \right\}. \quad (5.4)$$

است. برای ارزیابی امید شرطی (۵.۴) به شرط داده‌های مشاهده شده $y = \{y_i, i = 1, \dots, n\}$ و برآوردهای جاری $\hat{\Theta}^{(r)} = (\hat{\beta}^{(r)}, \text{vech}(\hat{D}^{(r)}), \text{vech}(\hat{\Sigma}^{(r)}), \hat{\phi}^{(r)}, \hat{\nu}^{(r)})$ از قضیه زیر استفاده می‌کنیم.

قضیه ۱.۱.۴. (وانگ و فن، ۲۰۱۱) تحت مدل سلسله مراتبی (۳.۴)، داریم

$$\begin{aligned} \begin{bmatrix} y_i \\ b_i \end{bmatrix} | \tau_i &\sim \mathcal{N}_{n_i+q} \left(\begin{bmatrix} X_i \beta \\ 0 \end{bmatrix}, \tau_i^{-1} \begin{bmatrix} \Lambda_i & Z_i D \\ D Z_i^\top & D \end{bmatrix} \right), \\ b_i | (y_i, \tau_i) &\sim \mathcal{N}_q \left(D Z_i^\top \Lambda_i^{-1} (y_i - X_i \beta), \tau_i^{-1} (D^{-1} + Z_i^\top R_i^{-1} Z_i)^{-1} \right), \\ \tau_i | y_i &\sim \mathcal{G} \left((\nu + n_i)/2, (\nu + \Delta_i)/2 \right), \end{aligned}$$

که در آن $\Delta_i = \Delta_i(\beta, D, \Sigma, \phi, \nu) = \varepsilon_i^\top \Lambda_i^{-1} \varepsilon_i$ فاصله ماهالانوبیس^۵ بین y_i و $X_i \beta$ است و $\varepsilon_i = y_i - X_i \beta$.

مراحل اصلی الگوریتم EM به شرح زیر انجام می‌شود:

گام E: محاسبه تابع $Q(\Theta | \hat{\Theta}^{(r)}) = E(\ell_c(\Theta) | y, \hat{\Theta}^{(r)})$ که مجموع $Q_i(\Theta | \hat{\Theta}^{(r)}) = E(\ell_c^{[i]}(\Theta) | y_i, \hat{\Theta}^{(r)})$ برای $i = 1, \dots, n$ است و به صورت

$$\begin{aligned} Q_i(\Theta | \hat{\Theta}^{(r)}) &= \frac{1}{\nu} \left\{ \log |R_i^{-1}| + \log |D^{-1}| - \text{tr}(D^{-1} \hat{B}_i^{(r)}) - \text{tr}(R_i^{-1} \hat{\Psi}_i^{(r)}(\theta)) \right. \\ &\quad \left. + \nu(\log(\nu/2) + \hat{\kappa}_i^{(r)} - \hat{\tau}_i^{(r)}) \right\} - \log(\nu/2), \end{aligned}$$

تعریف می‌شود، به طوری که

$$\begin{aligned} \hat{\tau}_i^{(r)} &= E(\tau_i | y_i, \hat{\Theta}^{(r)}) = (\hat{\nu}^{(r)} + n_i) / (\hat{\nu}^{(r)} + \hat{\Delta}_i^{(r)}), \\ \hat{\kappa}_i^{(r)} &= E(\log \tau_i | y_i, \hat{\Theta}^{(r)}) = \mathcal{D}_g \left(\frac{\hat{\nu}^{(r)} + n_i}{2} \right) - \log \left(\frac{\hat{\nu}^{(r)} + \hat{\Delta}_i^{(r)}}{2} \right), \\ \hat{B}_i^{(r)} &= E(\tau_i b_i b_i^\top | y_i, \hat{\Theta}^{(r)}) = \hat{\tau}_i^{(r)} \hat{b}_i^{(r)} \hat{b}_i^{(r)\top} + \hat{V}_{b_i}^{(r)}, \\ \hat{\Psi}_i^{(r)}(\beta) &= E(\tau_i \varepsilon_i \varepsilon_i^\top | y_i, \hat{\Theta}^{(r)}) \\ &= \hat{\tau}_i^{(r)} (y_i - X_i \beta - Z_i \hat{b}_i^{(r)}) (y_i - X_i \beta - Z_i \hat{b}_i^{(r)})^\top + Z_i \hat{V}_{b_i}^{(r)} Z_i^\top, \end{aligned}$$

⁵Mahalanobis distance

$$\begin{aligned}\hat{\Delta}_i^{(r)} &= (\mathbf{y}_i - \mathbf{X}_i \hat{\boldsymbol{\beta}}^{(r)})^\top \hat{\boldsymbol{\Lambda}}_i^{(r)-1} (\mathbf{y}_i - \mathbf{X}_i \hat{\boldsymbol{\beta}}^{(r)}), \\ \hat{\boldsymbol{\Lambda}}_i^{(r)} &= \mathbf{Z}_i \hat{\mathbf{D}}^{(r)} \mathbf{Z}_i^\top + \hat{\mathbf{R}}_i^{(r)}, \\ \hat{\mathbf{R}}_i^{(r)} &= \hat{\boldsymbol{\Sigma}}^{(r)} \otimes \boldsymbol{\Omega}_i(\hat{\phi}^{(r)}), \\ \hat{\mathbf{b}}_i^{(r)} &= E(\mathbf{b}_i | \mathbf{y}_i, \hat{\boldsymbol{\Theta}}^{(r)}) = \hat{\mathbf{D}}^{(r)} \mathbf{Z}_i^\top \hat{\boldsymbol{\Lambda}}_i^{(r)-1} (\mathbf{y}_i - \mathbf{X}_i \hat{\boldsymbol{\beta}}^{(r)}), \\ \hat{\mathbf{V}}_{\mathbf{b}_i}^{(r)} &= \hat{\tau}_i^{(r)} \text{cov}(\mathbf{b}_i | \mathbf{y}_i, \hat{\boldsymbol{\Theta}}^{(r)}) = (\hat{\mathbf{D}}^{(r)-1} + \mathbf{Z}_i^\top \hat{\mathbf{R}}_i^{(r)-1} \mathbf{Z}_i)^{-1}.\end{aligned}$$

گام M: به روز کردن $\hat{\boldsymbol{\Theta}}^{(r)}$ توسط $\hat{\boldsymbol{\Theta}}^{(r+1)} = \max_{\boldsymbol{\Theta}} Q(\boldsymbol{\Theta} | \hat{\boldsymbol{\Theta}}^{(r)})$.

۲.۲.۱.۴ برآوردگر PML و الگوریتم ECM

به منظور انتخاب متغیرهای مهم و برآورد ضرایب آن‌ها، به لگاریتم تابع درستنمایی (۵.۴)، جمله تاوان $\sum_{k=1}^p P_{\lambda_n}(|\beta_k|)$ را اضافه می‌کنیم، آنگاه لگاریتم تابع درستنمایی توانیده به صورت

$$\ell_c^{\text{Pen}}(\boldsymbol{\Theta}) = \ell_c(\boldsymbol{\Theta}) - n \sum_{k=1}^p P_{\lambda_n}(|\beta_k|). \quad (۶.۴)$$

حاصل می‌شود. به طور مشابه تابع Q توانیده در گام E را به صورت

$$Q_i^{\text{Pen}}(\boldsymbol{\Theta} | \hat{\boldsymbol{\Theta}}^{(r)}) = Q_i(\boldsymbol{\Theta} | \hat{\boldsymbol{\Theta}}^{(r)}) - n \sum_{k=1}^{p_n} P_{\lambda_n}(|\beta_k|)$$

در نظر می‌گیریم. در این جا از تابع تاوان اسکد استفاده می‌کنیم و برای حل مسئله بهینه سازی در گام M، تقریب درجه دو موضعی را به کار می‌بریم. بیشینه سازی گام M از نظر محاسباتی پیچیده است، لذا از گام CM استفاده می‌کنیم. در این گام فرض می‌شود $\hat{\mathbf{e}}_{il}^{(r)} = \mathbf{y}_{il} - \mathbf{X}_{il} \boldsymbol{\theta}_l - \mathbf{Z}_{il} \hat{\mathbf{b}}_{il}^{(r)}$ و $\hat{\mathbf{e}}_{is}^{(r)} = \mathbf{y}_{is} - \mathbf{X}_{is} \boldsymbol{\theta}_s - \mathbf{Z}_{is} \hat{\mathbf{b}}_{is}^{(r)}$ شامل عناصر $q_l \times 1$ یک زیربردار $q_l \times 1$ بعدی و شامل عناصر $(\sum_{j=1}^{l-1} q_j + 1)$ تا $(\sum_{j=1}^l q_j)$ ام از بردار $\hat{\mathbf{b}}_i^{(r)}$ است. در این حالت داریم $[\hat{\boldsymbol{\psi}}_{ils}^{(r)}(\boldsymbol{\theta})] = \hat{\boldsymbol{\Psi}}^{(r)}(\boldsymbol{\theta})$ که در آن $\hat{\boldsymbol{\psi}}_{ils}^{(r)}(\boldsymbol{\theta}) = E(\tau_i \mathbf{e}_{il} \mathbf{e}_{is}^\top | \mathbf{y}_i, \hat{\boldsymbol{\Theta}}^{(r)}) = \hat{\tau}_i^{(r)} \hat{\mathbf{e}}_{il}^{(r)} \hat{\mathbf{e}}_{is}^{(r)\top} - \mathbf{Z}_{il} \hat{\mathbf{V}}_{\mathbf{b}_{ils}}^{(r)} \mathbf{Z}_{is}^\top$ یک زیرماتریس $q_l \times q_s$ بعدی شامل سطرهای $(\sum_{j=1}^{l-1} q_j + 1)$ تا $(\sum_{j=1}^l q_j)$ بعد s_i است که در آن $\hat{\mathbf{V}}_{\mathbf{b}_{ils}}^{(r)}$ یک زیرماتریس $(\sum_{j=1}^s q_j + 1)$ تا $(\sum_{j=1}^s q_j)$ ام ماتریس $\hat{\mathbf{V}}_{\mathbf{b}_i}^{(r)}$ به ازای $l, s = 1, \dots, r$ است. در مجموع گام CM الگوریتم ECM به صورت الگوریتم ۴ پیشنهاد می‌شود.

۳.۱.۴ خواص جانبی

در این بخش خواص جانبی برآوردگر توانیده معرفی شده، مطالعه می‌شود. موقعیتی را در نظر بگیرید که مقادیر درست $\boldsymbol{\beta} = (\boldsymbol{\beta}_1^\top, \boldsymbol{\beta}_2^\top)^\top$ و ماتریس طرح هم‌ارز با آن به صورت $\mathbf{X}_i = (\mathbf{X}_{i(1)}, \mathbf{X}_{i(2)})$ تفکیک شوند. $\boldsymbol{\beta}_1$ بردار s بعدی است که همه مقادیر آن غیر صفراند و $\boldsymbol{\beta}_2 = \mathbf{0}$. در

الگوریتم ۴ برآوردگر PML و الگوریتم ECM

۱: گام CM-۱: ثابت نگه داشتن $\phi = \hat{\phi}^{(r)}$ و $\nu = \hat{\nu}^{(r)}$ ، به روز کردن $\hat{\beta}^{(r)}$ ، $\hat{D}^{(r)}$ و $\hat{\Sigma}^{(r)}$ از طریق بیشینه کردن $Q(\Theta|\hat{\Theta}^{(r)})$ که به صورت زیر حاصل می شوند:

$$\begin{aligned}\hat{\beta}^{(r+1)} &= \left(\sum_{i=1}^n \hat{\tau}_i^{(r)} \mathbf{X}_i^\top \hat{\mathbf{R}}_i^{(r)-1} \mathbf{X}_i + n \mathbf{E}_n(\hat{\beta}^{(r)}) \right)^{-1} \sum_{i=1}^n \hat{\tau}_i^{(r)} \mathbf{X}_i^\top \hat{\mathbf{R}}_i^{(r)-1} (\mathbf{y}_i - \mathbf{Z}_i \hat{\mathbf{b}}_i^{(r)}), \\ \hat{D}^{(r+1)} &= n^{-1} \sum_{i=1}^n \hat{\mathbf{B}}_i^{(r)}, \\ \hat{\sigma}_{lm}^{(r+1)} &= \begin{cases} (\sum_{i=1}^n s_i)^{-1} \sum_{i=1}^n \text{tr}(\Omega_i^{-1}(\hat{\phi}^{(r)}) \hat{\psi}_{ils}^{(r)}(\hat{\beta}^{(r+1)})) & \text{for } l = s, \\ (\sum_{i=1}^n s_i)^{-1} \sum_{i=1}^n \text{tr}(\Omega_i^{-1}(\hat{\phi}^{(r)}) (\hat{\psi}_{ils}^{(r)}(\hat{\beta}^{(r+1)}) + \hat{\psi}_{isl}^{(r)}(\hat{\beta}^{(r+1)}))) & \text{for } l \neq s, \end{cases} \end{aligned}$$

که در آن

$$\mathbf{E}_n(\hat{\theta}^{(r)}) = \text{diag} \left\{ \frac{q_{\lambda_n}(|\beta_1|)}{\epsilon + |\beta_1|}, \dots, \frac{q_{\lambda_n}(|\beta_p|)}{\epsilon + |\beta_p|} \right\}, \quad \epsilon = 10^{-6}.$$

۲: گام CM-۲: بر اساس $\hat{\beta}^{(r+1)}$ ، $\hat{D}^{(r+1)}$ و $\hat{\Sigma}^{(r+1)}$ ، محاسبه شده در گام قبل، $(\hat{\phi}^{(r+1)}, \hat{\nu}^{(r+1)})$ از طریق بیشینه کردن تابع Q به دست می آید. به عبارت دیگر

$$\begin{aligned}(\hat{\phi}^{(r+1)}, \hat{\nu}^{(r+1)}) &= \arg \max_{(\phi, \nu)} \left\{ \sum_{i=1}^n \left(r \log |\Omega_i^{-1}(\phi)| - \text{tr}((\hat{\Sigma}^{(r+1)})^{-1} \otimes \Omega_i^{-1}(\phi)) \times \hat{\Psi}_i^{(r+1/2)} \right) \right. \\ &\quad \left. + \nu \log(\frac{\nu}{\gamma}) - \gamma \log \Gamma(\frac{\nu}{\gamma}) + \nu(\hat{\kappa}_i^{(r)} - \hat{\tau}_i^{(r)}) \right\}, \end{aligned}$$

به طوری که $\hat{\Psi}_i^{(r+1/2)} = \hat{\Psi}_i^{(r)}(\hat{\beta}^{(r+1)})$. از آنجایی که با برابر صفر قرار دادن مشتق اول عبارت Q نسبت به ϕ و ν ، صورت بسته ای نتیجه نمی دهد؛ برای به روز کردن $(\hat{\phi}^{(r+1)}, \hat{\nu}^{(r+1)})$ ، باید از روش های عددی استفاده کرد که در این رساله از پکیج nlminb در نرم افزار R استفاده شده است.

۳: گام همگرایی: گام های E و CM را تا زمان همگرایی به مقداری ثابت تکرار می کنیم تا برآوردگر PML بدست آید.

نتیجه مقادیر برآوردگرها نیز به صورت $\hat{\beta} = (\hat{\beta}_1^\top, \hat{\beta}_\gamma^\top)^\top$ تفکیک می شوند. فرض کنید $\eta = (\omega^\top, \nu)^\top$ که در آن $\omega = (\text{vec}(\mathbf{D})^\top, \text{vec}(\Sigma)^\top, \phi^\top)^\top$. همچنین $J_{\beta_1 \beta_1}$ ، $J_{\beta_1 \beta_\gamma}$ ، $J_{\beta_\gamma \beta_\gamma}$ و $J_{\eta \eta} \rightarrow \mathbf{I}_{\eta \eta}$ به ترتیب ماتریس های اطلاع فیشر مربوط به β_1 ، β_γ و η هستند. برای مطالعه رفتار حدی برآوردگرها، شرایط نظم زیر را در نظر می گیریم:

(۱.۸) اگر $n \rightarrow \infty$ ، آنگاه

$$\begin{aligned} n^{-1} J_{\beta\beta} &\rightarrow I_{\beta\beta}, & n^{-1} J_{\beta_1\beta_1} &\rightarrow I_{\beta_1\beta_1} \\ n^{-1} J_{\beta_2\beta_2} &\rightarrow I_{\beta_2\beta_2}, & n^{-1} J_{\eta\eta} &\rightarrow I_{\eta\eta} \end{aligned}$$

(۲.۸) به ازای بردار $a \neq 0$ که $pr \times 1$ بعدی است داریم

$$\frac{\max_i \{a^\top X_i^\top \Lambda_i^{-1} X_i a\}}{a^\top (\sum_{i=1}^n X_i^\top \Lambda_i^{-1} X_i) a} \rightarrow 0. \quad (۷.۴)$$

(۳.۸) مقادیر ثابت و مثبت C_1 و C_2 وجود دارند، به طوری که

$$C_1 < \lambda_{\min} \left(\frac{1}{n} \sum_{i=1}^n X_i^\top \Lambda_i^{-1} X_i \right) < \lambda_{\max} \left(\frac{1}{n} \sum_{i=1}^n X_i^\top \Lambda_i^{-1} X_i \right) < C_2,$$

که در آن $\lambda_{\max}(A)$ و $\lambda_{\min}(A)$ به ترتیب کوچکترین و بزرگترین مقدار ویژه ماتریس A هستند.

قضیه ۲.۱.۴. تحت مفروضات این بخش و شرایط نظم (۱.۸)-(۳.۸) اگر

$$\liminf_{n \rightarrow \infty} \liminf_{\lambda_n \rightarrow 0^+} \sqrt{n} P'_{\lambda_n} \rightarrow \infty,$$

آنگاه $\hat{\beta} = (\hat{\beta}_1^\top, \hat{\beta}_2^\top)^\top$ یک جواب الگوریتم ECM است، به طوری که

$$\mathbb{P}_n(\hat{\beta}_2 = 0) \rightarrow 0,$$

$$\begin{aligned} \hat{\Theta} &\xrightarrow{P} \Theta, & \sqrt{n}(\hat{\beta}_1 - \beta_1) &\xrightarrow{d} \mathcal{N}_s(0, I_{\beta_1\beta_1}^{-1}), \\ \sqrt{n}(\hat{\eta} - \eta) &\xrightarrow{d} \mathcal{N}(0, I_{\eta\eta}^{-1}). \end{aligned}$$

قضیه ۲.۱.۴ بیان می‌کند که با احتمال نزدیک به یک، ضرایب صفر واقعی یعنی متغیرهای بی‌تأثیر شناسایی می‌شوند و ضرایب غیر صفر کارآیی بالایی دارند. علاوه بر این با مطالعه رفتار مجانبی برآوردهای $\hat{\eta}$ و $\hat{\beta}$ ، ابزار اولیه برای اجرای سایر استنباط‌های آماری مانند فاصله اطمینان و آزمون‌های فرض را فراهم می‌آورد. برای این منظور لازم است ماتریس کوواریانس مجموعه برآوردهای $\hat{\Theta}$ شناسایی شوند. از آنجا که $\hat{\Theta}$ سازگار و دارای توزیع نرمال مجانبی است، می‌توان معکوس ماتریس اطلاع فیشر در نقطه $\Theta = \hat{\Theta}$ را به عنوان تقریبی از ماتریس کوواریانس مجانبی برای $\hat{\Theta}$ در نظر گرفت. ماتریس‌های اطلاع فیشر و اثبات قضیه ۲.۱.۴ در قسمت ضمیمه آمده است.

۴.۱.۴ مطالعات عددی

در این بخش با طراحی یک مطالعه شبیه‌سازی شده و تحلیل داده‌های واقعی نشان داده می‌شود که روش معرفی‌شده علاوه بر انتخاب متغیرهای مهم به دلیل استفاده از توزیع t به جای توزیع نرمال، در مقابل داده‌های پرت تنومند است. بدین منظور عملکرد برآوردها تحت رویکردهای $MtLMM$ تاوانیده، $MNMM$ ، $MtLMM$ و $MNMM$ را با هم مقایسه خواهیم کرد.

۱.۴.۱.۴ مطالعه شبیه‌سازی

در این بخش با طراحی یک مطالعه شبیه‌سازی شده، نشان داده می‌شود که روش معرفی شده در انتخاب متغیرهای مهم از توزیع t از کارایی خوبی برخوردار است.

متغیرهای پاسخ y_i را به ازای $r = 2$ متغیر پاسخ از مدل (۱.۴) با فرضیات (۲.۴) تولید می‌کنیم. تعداد تکرارها برای هر نمونه را ثابت و برابر $s_i = 5$ در نظر می‌گیریم. به‌طوری‌که مولفه‌های مدل به‌صورت

$$X_{i1}, X_{i2} \sim U(-0.5, 0.5)$$

$$D = \begin{bmatrix} 1 & 0.25 \\ 0.25 & 1 \end{bmatrix}, \quad \Sigma = \begin{bmatrix} 2 & 0.75 \\ 0.75 & 1 \end{bmatrix}, \quad \text{و} \quad \Omega = I_{s_i}$$

تعیین شده‌اند. مقادیر فرضی برای پارامترهای مدل به‌صورت $\beta = (\beta_{11}, \dots, \beta_{1p}, \beta_{21}, \dots, \beta_{2p})^T$ بوده که در آن $p = 10$ ، $(\beta_{11}, \beta_{12}, \beta_{13}, \beta_{21}, \beta_{22}, \beta_{23}) = (-1, -1, 2, -1, -1, 2)$ و سایر اعضای بردار پارامترها برابر صفر هستند. برای در نظر گرفتن داده‌هایی با دم سنگین و همچنین داده‌هایی نزدیک به توزیع نرمال، مقادیر درجه آزادی را به ترتیب $\nu = 5$ و $\nu = 50$ در نظر گرفتیم. برای بررسی خواص نمونه‌ای برآورد ML ، اندازه حجم نمونه به‌صورت $n = 25$ برای نمونه نسبتاً کوچک و $n = 100$ برای نمونه نسبتاً بزرگ تعیین شده است. هر مجموعه از داده‌های شبیه‌سازی شده توسط رهیافت‌های $MtLMM$ تاوانیده، $MtLMM$ ، $MNMM$ تاوانیده و $MNMM$ و به ازای ترکیب‌های مختلف ν و n مورد تحلیل قرار گرفتند. به منظور سنجش دقت برآوردگرها از معیارهای AIC ، BIC و MSE بر اساس میانگین‌گیری ۱۰۰ تکرار از شبیه‌سازی، استفاده می‌شود. برای ارزیابی عملکرد روش‌ها در مساله انتخاب متغیر مشابه بخش ۱.۳.۳.۲ از معیارهای C ، I ، UF ، CF و OF استفاده می‌کنیم. جدول ۱.۴، خلاصه‌ای از نتایج مربوط به عملکرد چهار روش ذکر شده در برآورد و انتخاب مدل برای مقادیر مختلف (n, ν) است. برای بررسی بیشتر عملکرد روش‌های رقیب $MtLMM$ تاوانیده و $MNMM$ تاوانیده، نتایج مربوط به برآورد (Est) ، SD و MSE برای هر یک از پارامترها در جدول ۲.۴ گزارش شده است.

از نظر انتخاب مدل، همان‌طور که انتظار می‌رود مشاهده می‌کنیم که $MtLMM$ و $MNMM$ انتخاب متغیر انجام نمی‌دهند. علاوه بر این، $MtLMM$ تاوانیده و $MNMM$ تاوانیده با موفقیت همه متغیرها با ضرایب غیر صفر را انتخاب می‌کنند زیرا معیار I آن‌ها صفر است. اما بدیهی است که رویکرد پیشنهادی با تفاوت بسیار کم بهتر از $MNMM$ تاوانیده است زیرا C بزرگ‌تر است. در $MtLMM$ تاوانیده، احتمال شناسایی مدل درست حداقل ۷۰٪ است و این نرخ با افزایش حجم نمونه بهبود می‌یابد. بنابراین رفتار مجانبی برآوردگرها به خوبی تایید می‌شود. از نظر دقت برآورد، $MtLMM$ تاوانیده نسبت به سایر روش‌ها به طور قابل توجهی منجر به SD و Mse کمتری شده است. این موضوع در مورد معیارهای AIC ، BIC و MSE نیز صدق می‌کند. این برتری در حالت $\nu = 5$ که دم توزیع کلفت‌تر است، یا به عبارتی دارای مقادیر پرت بیشتری است، مشهودتر است. هنگامی که توزیع داده‌ها به نرمال نزدیک‌تر است، $\nu = 50$ ، یا حجم نمونه بزرگ است، $n = 100$ ،

۱۱۲ انتخاب متغیر برای داده‌های طولی چندمتغیره

هر دو روش تقریباً عملکرد یکسانی داشته‌اند. همچنین دریافتیم که مقادیر SD و Mse با افزایش حجم نمونه کاهش می‌یابند، لذا رفتار مجانبی خوب برآوردهای ML را تأیید می‌کند.

جدول ۱۰۴: نتایج شبیه‌سازی: مقایسه نتایج انتخاب متغیر روش‌های $MtLMM$ تاوانیده، $MtLMM$ ، $MNMM$ تاوانیده و $MNMM$ به ازای n ها و ν های مختلف.

$n = 25$										
	ν	MSE	ℓ	AIC	BIC	C	I	UF	CF	OF
$MtLMM$ تاوانیده	۵	۰/۴۵۲	-۸۹/۸۵	۲۳۳/۷۰	۲۶۶/۶۱	۱۳/۷۵	۰/۰۰	۰/۰۰	۰/۸۳	۰/۱۷
	۵۰	۰/۴۸۰	-۹۰/۱۵	۲۳۴/۳۰	۲۶۷/۲۱	۱۳/۶۶	۰/۰۰	۰/۰۰	۰/۷۳	۰/۲۷
$MtLMM$	۵	۰/۸۶۷	-۷۸/۶۹	۲۱۱/۳۹	۲۴۴/۲۹	۲/۷۵۰	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰
	۵۰	۰/۸۳۵	-۸۴/۷۷	۲۲۳/۵۵	۲۵۶/۴۶	۲/۸۹	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰
$MNMM$ تاوانیده	۵	۰/۷۱۱	-۲۴۶/۰۹	۵۴۶/۱۹	۵۷۹/۱۰	۱۳/۶۰	۰/۰۰	۰/۰۰	۰/۷۲	۰/۲۸
	۵۰	۰/۵۵۰	-۱۸۱/۵۳	۴۱۷/۰۷	۴۴۹/۹۸	۱۳/۵۷	۰/۰۰	۰/۰۰	۰/۶۹	۰/۳۱
$MNMM$	۵	۲/۵۱۰	-۲۴۰/۸۹	۵۳۵/۷۷	۵۶۸/۶۸	۱/۹۷	۰/۰۰	۰/۰۱	۰/۰۰	۱/۰۰
	۵۰	۱/۶۶۱	-۱۷۷/۱۳	۴۰۸/۲۸	۴۴۱/۱۸	۲/۱۵	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰
$n = 100$										
$MtLMM$ تاوانیده	۵	۰/۱۰۱	-۴۷۰/۷۳	۹۹۵/۴۶	۱۰۶۵/۸۰	۱۳/۹۷	۰/۰۰	۰/۰۰	۰/۹۷	۰/۰۳
	۵۰	۰/۰۸۳	-۵۲۴/۵۲	۱۱۰۳/۰۴	۱۱۷۳/۳۸	۱۳/۹۴	۰/۰۰	۰/۰۰	۰/۹۴	۰/۰۶
$MtLMM$	۵	۰/۱۸۹	-۴۹۰/۹۲	۱۰۳۵/۸۴	۱۱۰۶/۱۸	۵/۵۹	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰
	۵۰	۰/۱۸۳	-۵۴۳/۷۹۶	۱۱۴۱/۵۹۳	۱۲۱۱/۹۳۲	۵/۶۷	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰
$MNMM$ تاوانیده	۵	۰/۲۱۹	-۱۰۵۱/۱۴	۲۱۵۶/۲۹	۲۲۲۶/۶۳	۱۳/۹۱	۰/۰۱	۰/۰۱	۰/۹۰	۰/۰۹
	۵۰	۰/۰۷۹	-۷۷۵/۸۳	۱۶۰۵/۶۵	۱۶۷۵/۹۹	۱۳/۸۸	۰/۰۰	۰/۰۰	۰/۸۹	۰/۱۱
$MNMM$	۵	۰/۵۴۹	-۱۰۴۳/۰۲۸	۲۱۴۰/۰۵	۲۲۱۰/۳۹	۳/۹۴	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰
	۵۰	۰/۳۴۶	-۷۶۸/۶۰	۱۵۹۱/۲۰	۱۶۶۱/۵۴	۴/۵۹	۰/۰۰	۰/۰۰	۰/۰۰	۱/۰۰

۲.۴.۱.۴ تحلیل داده‌های بیماری صفراوی

در این بخش داده‌های بیماری صفراوی را توسط روش‌های $MtLMM$ تاوانیده، $MtLMM$ ، $MNMM$ تاوانیده و $MNMM$ تحلیل کرده و نتایج با هم مقایسه می‌شوند. $y_i = (y_{i1}^T, y_{i2}^T)^T$ متغیرهای پاسخ برای بیمار i ام است، به‌طوری‌که y_{i1} و y_{i2} به ترتیب بیانگر سطوح $lbili$ و $lalalbumin$ است. هدف از تحلیل این داده‌ها تعیین ارتباط بین این دو متغیر پاسخ و متغیرهای تبیینی Age ، Sex و $Drug$ است. بنابراین ماتریس طرح برای اثرات ثابت $\beta = (\beta_{10}, \beta_{11}, \beta_{12}, \beta_{13}, \beta_{14}, \beta_{20}, \beta_{21}, \beta_{22}, \beta_{23}, \beta_{24})$ به‌صورت

$$X_i = I_2 \otimes [I_{s_i} : t_i : sex_i I_{s_i} : drug_i I_{s_i} : age_i I_{s_i}],$$

تعیین می‌شوند، به‌طوری‌که I_{s_i} یک بردار $1 \times s_i$ بعدی با عناصر ۱، $t_i = (t_{i1}, \dots, t_{is_i})^T$ با $t_{i,k} = month_{i,k}/12$ ، sex_i شاخص جنسیت (۰ برای مرد و ۱ برای زن)، $drug_i$ شاخص دارو (۰ برای دارونما و ۱ برای D -پنسیلامین)، age_i سن بیمار i ام هستند. ماتریس طرح برای عرض از مبدا و شیب تصادفی به‌صورت $Z_i = I_2 \otimes [I_{s_i} : t_i]$ است. برای سادگی در محاسبات

فرض می‌کنیم $\Omega_i = I$. برای مقایسه عملکرد روش $MLMM$ تاوانیده با سه روش دیگر، خطای استاندارد برآوردها از طریق روش اعتبار سنجی متقابل یک طرفه محاسبه می‌شود. همچنین برای شناسایی بهترین مدل پشتیبانی شده توسط داده‌ها، از معیارهای AIC ، BIC و $MSPE$ (میانگین توان‌های دوم خطای پیشگویی) استفاده می‌شود. جدول ۳.۴، نتایج شامل برآورد پارامترها، خطاهای استاندارد و معیارهای AIC ، BIC و $MSPE$ را تحت برازش روش‌های مذکور ارائه می‌دهد. نتایج نشان می‌دهند که پارامترها تحت روش‌های $MtLMM$ تاوانیده و $MNMM$ تاوانیده، خطاهای استاندارد کمتری نسبت به $MtLMM$ و $MNMM$ دارند. اگر چه از نظر معیار $MSPE$ تفاوت عملکرد روش‌ها چشمگیر نیست با این وجود معیارهای AIC ، BIC مدل پیشنهادی در مقایسه با سه مدل دیگر مقدار کمتری دارند. براساس مدل $MtLMM$ تاوانیده، نتیجه می‌شود که $Drug$ اثر معنی‌داری روی $lbili$ ندارد و $albumin$ تحت تاثیر Sex و $Drug$ نیست.

۲.۴ مدل نیمه پارامتری با اثرات آمیخته با توزیع t چندمتغیره

در این بخش ابتدا به معرفی مدل آمیخته نیمه پارامتری t چندمتغیره ($MtSMM$) پرداخته می‌شود. مولفه ناپارامتری مدل به روش اسپلین‌های رگرسیونی شرح داده شده در بخش ۳.۵.۱ تقریب زده می‌شود. سپس از الگوریتم تکراری ECM برای برآورد پارامترها استفاده می‌کنیم. برای بررسی کارایی مدل پیشنهادی، نتایج حاصل از تحلیل مجموعه داده شبیه‌سازی شده و یک مجموعه داده واقعی گزارش می‌شود.

۱.۲.۴ مدل $MtSMM$ و تقریب اسپلین

n نمونه تصادفی از داده‌های طولی با مقادیر پرت را در نظر بگیرید، به‌طوری‌که نمونه i ام دارای r متغیر پاسخ و s_i مشاهده مکرر در طول زمان باشد. مدل $MtSMM$ برای نمونه i ام به‌صورت

$$y_i = X_i\beta + f(t_i) + Z_ib_i + \varepsilon_i, \quad (۸.۴)$$

بیان می‌شود، که در آن $g_j(t_i) = (g_j(t_{i,1}), \dots, g_j(t_{i,s_i}))^\top$ ، $f(t_i) = (f_1(t_i)^\top, \dots, f_r(t_i)^\top)^\top$ و $g_j(t_{i,k})$ یک تابع هموار و مشتق‌پذیر برای پاسخ j ام اندازه‌گیری شده در زمان $t_{i,k}$ است. فرضیات توزیعی مدل به‌صورت رابطه (۲.۴) است. سایر مولفه‌های مدل مشابه بخش ۱.۴ تعیین می‌شوند. برای تقریب تابع ناپارامتری $g_j(t_{i,k})$ به روش اسپلین‌های رگرسیونی فرض کنید $t_{i,1} = t_{ij}^{(1)} < t_{ij}^{(L_j)} = t_{i,s_i}$... افزای برای بازه $[t_{i,1}, t_{i,s_i}]$ باشد که هم‌ارز با متغیر پاسخ j از نمونه i ام است. l امین گره به ازای $l = 1, \dots, L_j$ است. اگر برای تقریب $g_j(t_{i,k})$ از تابع اسپلین مرتبه $d_j (\geq 1)$ استفاده کنیم، آن‌گاه خواهیم داشت

$$g_j(t_{i,k}) = \alpha_{j,0} + \alpha_{j,1}t_{i,k} + \dots + \alpha_{j,d_j}t_{i,k}^{d_j} + \sum_{l=1}^{L_j} \alpha_{(d_j+1)+l}(t_{i,k} - t_{ij}^{(l)})_+^{d_j} = B_j(t_{i,k})^\top \alpha_j,$$

که در آن $B_j(t_{i,k}) = (1, t_{i,k}, \dots, t_{i,k}^{d_j}, (t_{i,k} - t_{ij}^{(1)})_+^{d_j}, \dots, (t_{i,k} - t_{ij}^{(L_j)})_+^{d_j})$ برداری $1 \times h_j$ بعدی از توابع پایه بوده و $\alpha_j = (\alpha_{j,0}, \dots, \alpha_{j,d_j}, \alpha_{j,d_j+1}, \dots, \alpha_{j,h_j})^\top$ برداری از ضرایب اسپلاین با بعد $h_j = d_j + 1 + L_j$ است. بنابراین مدل نیمه‌پارامتری (۸.۴) به یک مدل خطی به صورت

$$y_i = X_i \beta + B(t_i) \alpha + Z_i b_i + \varepsilon_i, \quad (9.4)$$

تبدیل می‌شود. که $B(t_i) = \text{diag}(B_1(t_i), \dots, B_r(t_i))$ و $B_j(t_i) = (B_j(t_{i,1}), \dots, B_j(t_{i,s_i}))^\top$ یک ماتریس طرح $s_i \times h_i$ بعدی هم‌ارز با بخش اسپلاینی مدل است. اگر بعد کلی اثرات اسپلاین را با $h = \sum_{j=1}^r h_j$ نشان دهیم، آنگاه $\alpha = (\alpha_1^\top, \dots, \alpha_r^\top)^\top$ یک بردار $1 \times h$ بعدی از اثرات اسپلاین است. برای سادگی، مدل (۹.۴) را می‌توان به صورت

$$y_i = \tilde{X}_i \theta + Z_i b_i + \varepsilon_i, \quad (10.4)$$

نوشت، به طوری که $\tilde{X}_i = \text{diag}\{\tilde{X}_{i1}, \dots, \tilde{X}_{ir}\}$ و $\tilde{X}_{ij} = (X_{ij}, B_j(t_i))$ ماتریس طرح $s_i \times (p_j + h_j)$ بعدی برای پاسخ j از نمونه i ام است. $\theta = (\beta^\top, \alpha^\top)^\top$ بردار پارامترهای رگرسیونی ترکیب شده با بعد $1 \times (p_j + h_j)$ است.

طبق راهکار ارائه شده در بخش ۱۰.۴، مدل (۱۰.۴) را می‌توان به صورت مدل سلسله مراتبی زیر

$$y_i | (b_i, \tau_i) \sim \mathcal{N}_{n_i}(\tilde{X}_i \beta + Z_i b_i, \tau_i^{-1} R_i), \quad (11.4)$$

$$b_i | \tau_i \sim \mathcal{N}_q(0, \tau_i^{-1} D),$$

$$\tau_i \sim \mathcal{G}(\nu/2, \nu/2).$$

نوشت. حال به برآورد ML پارامترها بر اساس الگوریتم ECM می‌پردازیم. این روش کاملاً مشابه الگوریتم بکار رفته در بخش است با این تفاوت که

۲.۲.۴ برآورد ML و الگوریتم ECM

فرض کنید $\Theta = (\theta, \text{vech}(D), \text{vech}(\Sigma), \phi, \nu)$ مجموعه‌ای از تمام پارامترهای مدل باشد. اگر اثرات تصادفی b_i را به عنوان مقادیر گمشده و مجموعه $\{(y_i, b_i), i = 1, \dots, n\}$ را داده‌های کامل در نظر بگیریم، آنگاه لگاریتم تابع درستنمایی شرطی بدون در نظر گرفتن جمله ثابت به صورت

$$\ell_c(\Theta) = \sum_{i=1}^n \ell_c^{[i]}(\Theta) = \sum_{i=1}^n \frac{1}{2} \left\{ \log |R_i^{-1}| + \log |D^{-1}| - \tau_i [\varepsilon_i^\top R_i^{-1} \varepsilon_i + b_i^\top D^{-1} b_i] + \nu \log \left(\frac{\nu}{2} \right) - \nu \log \frac{\nu}{2} + \nu (\log \tau_i - \tau_i) \right\}. \quad (12.4)$$

است. برای ارزیابی امید شرطی (۱۲.۴) به شرط داده‌های مشاهده شده $y = \{y_i, i = 1, \dots, n\}$ و برآوردهای جاری $\hat{\Theta}^{(r)} = (\hat{\theta}^{(r)}, \text{vech}(\hat{D}^{(r)}), \text{vech}(\hat{\Sigma}^{(r)}), \hat{\phi}^{(r)}, \hat{\nu}^{(r)})$ ، از قضیه زیر استفاده می‌کنیم.

نتیجه ۱.۲.۴. (وانگ و فن، ۲۰۱۱). تحت مدل سلسله مراتبی (۳.۴)، داریم

$$\begin{aligned} \begin{bmatrix} \mathbf{y}_i \\ \mathbf{b}_i \end{bmatrix} | \tau_i &\sim \mathcal{N}_{n_i+q} \left(\begin{bmatrix} \tilde{\mathbf{X}}_i \boldsymbol{\theta} \\ \mathbf{0} \end{bmatrix}, \tau_i^{-1} \begin{bmatrix} \boldsymbol{\Lambda}_i & \mathbf{Z}_i \mathbf{D} \\ \mathbf{D} \mathbf{Z}_i^\top & \mathbf{D} \end{bmatrix} \right), \\ \mathbf{b}_i | (\mathbf{y}_i, \tau_i) &\sim \mathcal{N}_q \left(\mathbf{D} \mathbf{Z}_i^\top \boldsymbol{\Lambda}_i^{-1} (\mathbf{y}_i - \tilde{\mathbf{X}}_i \boldsymbol{\theta}), \tau_i^{-1} (\mathbf{D}^{-1} + \mathbf{Z}_i^\top \mathbf{R}_i^{-1} \mathbf{Z}_i)^{-1} \right), \\ \tau_i | \mathbf{y}_i &\sim \mathcal{G} \left((\nu + n_i)/2, (\nu + \Delta_i)/2 \right), \end{aligned}$$

به طوری که $\tilde{\epsilon}_i = \tilde{\epsilon}_i^\top \boldsymbol{\Lambda}_i^{-1} \tilde{\epsilon}_i$ فاصله مایلانویس بین $\mathbf{y}_i$ و $\tilde{\mathbf{X}}_i \boldsymbol{\theta}$ است و $\tilde{\epsilon}_i = \mathbf{y}_i - \tilde{\mathbf{X}}_i \boldsymbol{\theta}$

مراحل اصلی الگوریتم EM به شرح زیر انجام می شود:

گام E: محاسبه تابع $Q(\boldsymbol{\Theta} | \hat{\boldsymbol{\Theta}}^{(r)}) = E(\ell_c(\boldsymbol{\Theta}) | \mathbf{y}, \hat{\boldsymbol{\Theta}}^{(r)})$ که مجموع $Q_i(\boldsymbol{\Theta} | \hat{\boldsymbol{\Theta}}^{(r)}) = E(\ell_c^{[i]}(\boldsymbol{\Theta}) | \mathbf{y}_i, \hat{\boldsymbol{\Theta}}^{(r)})$ برای $i = 1, \dots, n$ است و به صورت

$$\begin{aligned} Q_i(\boldsymbol{\Theta} | \hat{\boldsymbol{\Theta}}^{(r)}) &= \frac{1}{2} \left\{ \log |\mathbf{R}_i^{-1}| + \log |\mathbf{D}^{-1}| - \text{tr}(\mathbf{D}^{-1} \hat{\mathbf{B}}_i^{(r)}) - \text{tr}(\mathbf{R}_i^{-1} \hat{\boldsymbol{\Psi}}_i^{(r)}(\boldsymbol{\theta})) \right. \\ &\quad \left. + \nu (\log(\nu/2) + \hat{\kappa}_i^{(r)} - \hat{\tau}_i^{(r)}) \right\} - \log(\nu/2), \end{aligned}$$

تعریف می شوند، به طوری که

$$\begin{aligned} \hat{\tau}_i^{(r)} &= E(\tau_i | \mathbf{y}_i, \hat{\boldsymbol{\Theta}}^{(r)}) = (\hat{\nu}^{(r)} + n_i) / (\hat{\nu}^{(r)} + \hat{\Delta}_i^{(r)}), \\ \hat{\kappa}_i^{(r)} &= E(\log \tau_i | \mathbf{y}_i, \hat{\boldsymbol{\Theta}}^{(r)}) = \mathcal{D}_g \left(\frac{\hat{\nu}^{(r)} + n_i}{2} \right) - \log \left(\frac{\hat{\nu}^{(r)} + \hat{\Delta}_i^{(r)}}{2} \right), \\ \hat{\mathbf{B}}_i^{(r)} &= E(\tau_i \mathbf{b}_i \mathbf{b}_i^\top | \mathbf{y}_i, \hat{\boldsymbol{\Theta}}^{(r)}) = \hat{\tau}_i^{(r)} \hat{\mathbf{b}}_i^{(r)} \hat{\mathbf{b}}_i^{(r)\top} + \hat{\mathbf{V}}_{\mathbf{b}_i}^{(r)}, \\ \hat{\boldsymbol{\Psi}}_i^{(r)}(\boldsymbol{\theta}) &= E(\tau_i \mathbf{e}_i \mathbf{e}_i^\top | \mathbf{y}_i, \hat{\boldsymbol{\Theta}}^{(r)}) \\ &= \hat{\tau}_i^{(r)} (\mathbf{y}_i - \tilde{\mathbf{X}}_i \boldsymbol{\theta} - \mathbf{Z}_i \hat{\mathbf{b}}_i^{(r)}) (\mathbf{y}_i - \tilde{\mathbf{X}}_i \boldsymbol{\theta} - \mathbf{Z}_i \hat{\mathbf{b}}_i^{(r)})^\top + \mathbf{Z}_i \hat{\mathbf{V}}_{\mathbf{b}_i}^{(r)} \mathbf{Z}_i^\top, \end{aligned}$$

و

$$\begin{aligned} \hat{\Delta}_i^{(r)} &= (\mathbf{y}_i - \tilde{\mathbf{X}}_i \hat{\boldsymbol{\theta}}^{(r)})^\top \hat{\boldsymbol{\Lambda}}_i^{(r)-1} (\mathbf{y}_i - \tilde{\mathbf{X}}_i \hat{\boldsymbol{\theta}}^{(r)}), \\ \hat{\boldsymbol{\Lambda}}_i^{(r)} &= \mathbf{Z}_i \hat{\mathbf{D}}^{(r)} \mathbf{Z}_i^\top + \hat{\mathbf{R}}_i^{(r)}, \\ \hat{\mathbf{R}}_i^{(r)} &= \hat{\boldsymbol{\Sigma}}^{(r)} \otimes \boldsymbol{\Omega}_i(\hat{\phi}^{(r)}), \\ \hat{\mathbf{b}}_i^{(r)} &= E(\mathbf{b}_i | \mathbf{y}_i, \hat{\boldsymbol{\Theta}}^{(r)}) = \hat{\mathbf{D}}^{(r)} \mathbf{Z}_i^\top \hat{\boldsymbol{\Lambda}}_i^{(r)-1} (\mathbf{y}_i - \tilde{\mathbf{X}}_i \hat{\boldsymbol{\theta}}^{(r)}), \\ \hat{\mathbf{V}}_{\mathbf{b}_i}^{(r)} &= \hat{\tau}_i^{(r)} \text{cov}(\mathbf{b}_i | \mathbf{y}_i, \hat{\boldsymbol{\Theta}}^{(r)}) = (\hat{\mathbf{D}}^{(r)-1} + \mathbf{Z}_i^\top \hat{\mathbf{R}}_i^{(r)-1} \mathbf{Z}_i)^{-1}. \end{aligned}$$

گام M: به روز کردن $\hat{\boldsymbol{\Theta}}^{(r)}$ توسط $\hat{\boldsymbol{\Theta}}^{(r+1)} = \max_{\boldsymbol{\Theta}} Q(\boldsymbol{\Theta} | \hat{\boldsymbol{\Theta}}^{(r)})$.

در مجموع گام CM الگوریتم ECM به صورت الگوریتم ۵ پیشنهاد می شود.

الگوریتم ۵ برآوردگر PML و الگوریتم ECM

۱: گام CM-۱: ثابت نگه داشتن $\phi = \hat{\phi}^{(r)}$ و $\nu = \hat{\nu}^{(r)}$ ، به روز کردن $\hat{\theta}^{(r)}$ ، $\hat{D}^{(r)}$ و $\hat{\Sigma}^{(r)}$ از طریق پیشینه کردن $Q(\Theta|\hat{\Theta}^{(r)})$ که به صورت زیر حاصل می شوند:

$$\begin{aligned}\hat{\beta}^{(r+1)} &= \left(\sum_{i=1}^n \hat{\tau}_i^{(r)} \tilde{X}_i^\top \hat{R}_i^{(r)-1} \tilde{X}_i \right)^{-1} \sum_{i=1}^n \hat{\tau}_i^{(r)} \tilde{X}_i^\top \hat{R}_i^{(r)-1} (y_i - Z_i \hat{b}_i^{(r)}), \\ \hat{D}^{(r+1)} &= n^{-1} \sum_{i=1}^n \hat{B}_i^{(r)}, \\ \hat{\sigma}_{lm}^{(r+1)} &= \begin{cases} (\sum_{i=1}^n s_i)^{-1} \sum_{i=1}^n \text{tr}(\Omega_i^{-1}(\hat{\phi}^{(r)}) \hat{\psi}_{ils}^{(r)}(\hat{\theta}^{(r+1)})) & \text{for } l = s, \\ (\sum_{i=1}^n s_i)^{-1} \sum_{i=1}^n \text{tr}(\Omega_i^{-1}(\hat{\phi}^{(r)}) (\hat{\psi}_{ils}^{(r)}(\hat{\theta}^{(r+1)}) + \hat{\psi}_{isl}^{(r)}(\hat{\theta}^{(r+1)}))) & \text{for } l \neq s, \end{cases}\end{aligned}$$

۲: گام CM-۲: بر اساس $\hat{\theta}^{(r+1)}$ ، $\hat{D}^{(r+1)}$ و $\hat{\Sigma}^{(r+1)}$ محاسبه شده در گام قبل، $(\hat{\phi}^{(r+1)}, \hat{\nu}^{(r+1)})$ از طریق پیشینه کردن تابع Q به دست می آید. به عبارت دیگر

$$\begin{aligned}(\hat{\phi}^{(r+1)}, \hat{\nu}^{(r+1)}) &= \arg \max_{(\phi, \nu)} \left\{ \sum_{i=1}^n \left(r \log |\Omega_i^{-1}(\phi)| - \text{tr}((\hat{\Sigma}^{(r+1)})^{-1} \otimes \Omega_i^{-1}(\phi)) \times \hat{\Psi}_i^{(r+1/2)} \right) \right. \\ &\quad \left. + \nu \log(\frac{\nu}{r}) - 2 \log \Gamma(\frac{\nu}{r}) + \nu (\hat{\kappa}_i^{(r)} - \hat{\tau}_i^{(r)}) \right\},\end{aligned}$$

به طوری که $\hat{\Psi}_i^{(r+1/2)} = \hat{\Psi}_i^{(r)}(\hat{\beta}^{(r+1)})$.

۳: گام همگرایی: گام های E و CM را تا زمان همگرایی به مقداری ثابت تکرار می کنیم تا برآوردگر ML بدست آید.

۳.۲.۴ خواص جانبی

برای مطالعه رفتار حدى برآوردگرهای $\hat{\Xi} = (\hat{\beta}^\top, \omega^\top, \hat{\nu})$ و $\hat{g}(t_i) = B(t_i)\hat{\alpha}$ ، مجموعه ای از شرایط نظم مشابه بخش ۱.۴ مورد نیاز است. در این جا $\omega = (\text{vech}(D)^\top, \text{vech}(\Sigma)^\top, \phi^\top)^\top$. اگر $\hat{\beta} \xrightarrow{p} \beta$ و هنگامی که $n \rightarrow \infty$ ، داشته باشیم $\sup_t |B(t_i)\hat{\alpha} - g(t_i)| \xrightarrow{p} 0$ ، آن گاه می گوییم برآوردگرهای $(\hat{\beta}, \hat{g}(t_i))$ سازگاراند. تحت شرایط نظم، قضیه ۱.۲.۴ را بیان می کنیم که رفتار حدى برآوردگرهای $\hat{\Xi}$ و $\hat{g}(t_i)$ را نشان می دهد.

قضیه ۱.۲.۴. اگر $L_j \approx N^{1/(2m_j+1)}$ آن گاه

$$(i) : \quad \frac{1}{n} \sum_{i=1}^n \sum_{k=1}^{s_i} (\hat{g}_j(t_{i,k}) - g_j(t_{i,k}))^2 = O_p(N^{2m_j/(2m_j+1)}), \quad j = 1, \dots, r,$$

$$(ii) : \quad \hat{\Xi} \xrightarrow{p} \Xi, \quad \sqrt{n}(\hat{\beta} - \beta) \xrightarrow{d} \mathcal{N}(0, I_{\beta\beta}^{-1}), \quad \sqrt{n}(\boldsymbol{\eta} - \boldsymbol{\eta}) \xrightarrow{d} \mathcal{N}(0, I_{\eta\eta}^{-1}).$$

که در آن $N = \sum_{i=1}^n s_i$ تعداد کل مشاهدات بوده و ماتریس های $I_{\beta\beta}$ و $I_{\eta\eta}$ در شرط (۱.A) معرفی شده اند.

۴.۲.۴ مطالعات عددی

در این بخش، فرایند برازش مدل و برآورد پارامترها را به صورت عددی بررسی می‌کنیم. ابتدا مجموعه‌ای از داده‌های شبیه‌سازی شده و یک مورد مطالعاتی با داده‌های واقعی را توسط روش پیشنهادی یعنی برآوردگر ML در مدل نیمه پارامتری با اثرات آمیخته t چندمتغیره ($MtSMM$) مورد تحلیل قرار می‌دهیم.

۱.۴.۲.۴ مطالعه شبیه‌سازی

متغیرهای پاسخ y_i را به ازای $r = 2$ متغیر پاسخ از مدل (۸.۴) و با فرضیات (۲.۴) تولید می‌کنیم. به طوری که

$$\begin{aligned} g_1(t_{i,k}) &= 2/5 - 1/5 t_{i,k} + 0/4 t_{i,k}^2, \\ g_2(t_{i,k}) &= 2 \sin(2\pi t_{i,k}) + 2 \cos(2\pi t_{i,k}), \end{aligned}$$

و پارامترهای مدل به صورت $\beta = (\beta_{11}, \beta_{12}, \beta_{21}, \beta_{22})^T = (1, 2, -2, 4)^T$ تعیین می‌شوند. سایر مولفه‌های مدل مشابه مطالعه شبیه سازی در بخش ۱.۴ است. برای $g_1(t_{i,k})$ تقریب اسپلاین درجه دو، $d_1 = 2$ ، و برای $g_2(t_{i,k})$ تقریب اسپلاین مکعبی، $d_2 = 3$ ، را در نظر می‌گیریم. مطابق دلایل بیان شده در بخش (۲.۴) حجم نمونه و مقادیر درجه آزادی به ترتیب با مقادیر $n = 50, 100$ و $\nu = 5, 50$ تعیین می‌شوند. دقت برآورد با مقادیر اریبی، SD و MSE پارامترهای برآورد شده تحت برازش مدل‌های مذکور اندازه‌گیری می‌شود. علاوه بر این، انتخاب مدل با محاسبه معیارهای AIC و BIC انجام می‌شود. نتایج برای حالت‌های $n = 25$ و $n = 100$ به ترتیب در جداول ۴.۴ و ۵.۴ گزارش شده است. همان‌طور که مشاهده می‌شود، $MtSMM$ نسبت به $MNSMM$ مقادیر MSE ، SD کمتری دارد. این قضیه در مورد معیارهای AIC و BIC نیز صادق است. این رفتار در حالت $(n, \nu) = (25, 5)$ مشهودتر است. هنگامی که داده‌ها به توزیع نرمال نزدیک‌ترند ($\nu = 50$)، هر دو مدل برآوردها و انحراف‌های استاندارد تقریباً مشابهی را نتیجه می‌دهند. همچنین دریافتیم که مقادیر SD و MSE با افزایش اندازه نمونه کاهش می‌یابند که تأیید کننده ویژگی‌های جانبی خوب برآوردهای ML هستند.

۲.۴.۲.۴ تحلیل داده‌های بیماری صفراوی

در این بخش داده‌های $PBCseq$ تحلیل می‌شوند. شکل ۴.۱ نشان می‌دهد رفتار دو متغیر در طول زمان متفاوت از هم است. بنابراین به نظر می‌رسد که برای مطالعه تغییر رفتار آن‌ها در طول زمان، استفاده از طرح ناپارامتری مجزا برای هر یک از متغیرها، مناسب باشد. طرح باید به گونه‌ای باشد که تاثیر مقادیر پرت موجود در داده‌ها را نیز کنترل کند. لذا استفاده از مدل $MtSMM$ پیشنهاد می‌شود. استفاده از این مدل این اجازه را می‌دهد که علاوه بر مطالعه ارتباط بین بلیروبین و آلبومین، نحوه

تغییرات و ارتباط این دو متغیر در طول زمان را بررسی کنیم. برای بررسی تنومند بودن این روش در حضور مشاهدات پرت، علاوه بر مدل پیشنهادی، مدل $MNSMM$ را نیز به داده‌ها برازش داده و نتایج با هم مقایسه می‌شوند. تنضیمات مدل مشابه بخش ۱۰۴ است. در اینجا برای ماتریس Ω ، سه نوع اختیار همبستگی شامل همبستگی استقلال (UNC)، همبستگی تبادلی پذیر (EX) و همبستگی اتورگرسیو مرتبه یک ($AR(1)$) را در نظر می‌گیریم. در مجموع با ۶ مدل نامزد روبرو هستیم.

جدول ۶۰۴ خلاصه‌ای از برازش ۶ مدل نامزد را نشان می‌دهد که شامل مقادیر BIC ، AIC و $MSPE$ است. با توجه به تمام معیارها، استفاده از ساختار همبستگی $AR(1)$ ، برای هر دو مدل $MtSMM$ و $MNSMM$ مناسب بوده است. جدول ۷۰۴ برآورد پارامترهای رگرسیونی به همراه انحراف‌های استاندارد اثرات ثابت برای $MtSMM$ و $MNSMM$ با همبستگی $AR(1)$ را ارائه می‌دهد. نتایج بیان می‌کنند که برآورد اثرات ثابت در دو مدل $MtSMM$ و $MNSMM$ تقریباً یکسان است اما $MtSMM$ در مقایسه با $MNSMM$ ، انحراف استاندارد بسیار کمتری را نتیجه می‌دهد، که بیانگر عملکرد مناسب $MtSMM$ است.

۳.۴ پیشنهادات برای آینده تحقیق

مدل اثرات آمیخته نیمه‌پارامتری با جمع‌آوری ویژگی‌های مثبت هر سه دسته مدل‌های رگرسیونی پارامتری، ناپارامتری و نیمه‌پارامتری در یک مدل واحد، از انعطاف بسیار بالایی در تحلیل داده‌های طولی برخوردار است. همان‌گونه که نتایج گزارش شده در مطالعات شبیه‌سازی و تحلیل داده‌های واقعی، برتری این مدل نسبت به مدل‌های معمول و کارایی بالای برآوردگرهای معرفی‌شده را نشان می‌دهد، بر این اساس، تحلیل داده‌های طولی در چارچوب مدل معرفی‌شده زمینه‌ای نوپا و ابزاری کارآمد برای فعالیت‌های آتی آماردانان فراهم خواهد آورد. ما نیز با تکیه بر این مهم تصمیم داریم، مطالعات و فعالیت‌های خود را در این راستا قرار داده و این مدل را در بسیاری از مسائل و استنباط‌های آماری دیگر وارد کنیم. در ادامه پیشنهادات خود را به طور مختصر ارائه خواهیم نمود.

۱.۳.۴ مدل اثرات آمیخته تعمیم‌یافته نیمه‌پارامتری برای داده‌های چندجمله‌ای

با بررسی و مطالعه اجمالی مباحث این رساله، می‌توان پی برد که یکی از اولین موضوعات کاربردی که می‌توان به عنوان مطالعه بعدی در نظر گرفت، ادغام بخش‌های ۱۰۴ و ۲۰۴ و پرداختن به مسئله انتخاب متغیر در مدل نیمه‌پارامتری با اثرات آمیخته با توزیع t چندمتغیره است. همچنین فرض کنید که داده‌های چندجمله‌ای مورد بحث در این فصل، ترکیبی از انواع داده‌های رسته‌ای باشند که هر کدام نیازمند به‌کارگیری تابع ربط منحصر به خود است. در این حالت می‌توان یک مدل اثرات

آمیخته تعمیم یافته نیمه پارامتری را به صورت

$$g_j(\mu_{ij}) = \eta_{ij} = \mathbf{X}_{ij}^\top \boldsymbol{\beta} + \mathbf{Z}_{ij}^\top \mathbf{u}_i + f(t_{ij})$$

در نظر گرفت. اگرچه این مدل می تواند طیف گسترده تری از کاربردهای عملی را دربر بگیرد، اما به طور قابل توجهی پیچیده تر و از نظر محاسباتی سنگین تر خواهد بود. لذا نیازمند تحقیقات جامعی پیرامون استراتژی های محاسباتی و برنامه های کاربردی است.

۲.۳.۴ استفاده از توابع تاوان تجمعی

در این رساله تمرکز اصلی روی برآورد و انتخاب متغیر جزء پارامتری مدل ها بود. با این حال می توان مشابه کار سوای (۲۰۰۹)، تکنیک پیشنهادی را برای جزء ناپارامتری مدل ها نیز به کار برد. به عنوان مثال در فصل ۲ با به کارگیری جملات تاوان تجمعی، مسئله برآورد و انتخاب متغیر هم زمان در رابطه (۱۵.۲) به حل مسئله بیشینه سازی لگاریتم تابع درست نمایی تاوانیده، به صورت

$$\ell_c^{\text{Pen}}(\boldsymbol{\Theta}) = \ell_c(\boldsymbol{\Theta}) - n \sum_{k=1}^p P_{\lambda_n}^{(1)}(|\beta_k|) - n \sum_{k=1}^p P_{\lambda_n}^{(2)}(|\alpha_k|)$$

تبدیل می شود، که $P_{\lambda_n}^{(1)}(\cdot)$ ضرایب پارامتری مدل و $P_{\lambda_n}^{(2)}(\cdot)$ ضرایب ناپارامتری مدل را کنترل می کنند. جزئیات نظری مسئله بیان شده نیازمند توجه ویژه ای است و می توان به عنوان موضوع جالبی برای تحقیقات آینده در نظر گرفت.

۳.۳.۴ مدل ضرایب متغیر نیمه پارامتری با اثرات آمیخته

همان گونه که در فصل اول این رساله بیان شد، در تحلیل داده های طولی با سه هدف اصلی زیر روبرو هستیم:

۱. مطالعه تغییرات آزمودنی ها در طول زمان.
 ۲. مطالعه رابطه بین متغیرها (اثر متغیرهای تبیینی بر روی متغیر پاسخ).
 ۳. مطالعه تاثیر زمان بر رابطه بین متغیرها.
- مدل اثرات آمیخته، تنها قادر به برآورد نقاطی است که مشاهده واقعی هم ارز با آن موجود باشد؛ بنابراین اثر متغیرهای تبیینی بر روی متغیر پاسخ را می توان دریافت. از آنجایی که، برآورد متغیر پاسخ در مدل اثرات آمیخته نیمه پارامتری به دلیل وجود تابع ناپارامتری نسبت به زمان، به صورت یک منحنی هموار می باشد، هدف مطالعه تغییرات آزمودنی ها در طول زمان محقق می گردد. حال اگر متغیرهایی تبیینی داشته باشیم که در هر نقطه زمانی تاثیر متفاوتی روی پاسخ ها داشته باشند و علاقه مند به مطالعه تاثیر زمان بر رابطه بین متغیرها باشیم، مدل اثرات آمیخته نیمه پارامتری با ضرایب متغیر را می توان به صورت

$$\mathbf{y}(t) = \mathbf{X}(t)^\top \boldsymbol{\beta}(t) + \mathbf{W}(t)^\top \boldsymbol{\alpha} + \mathbf{Z}(t)^\top \mathbf{b} + \varepsilon(t),$$

در نظر گرفت، به طوری که پارامتر $\beta(t) = (\beta_1(t), \dots, \beta_p(t))^T \in \mathbb{R}^p$ تابعی نامعلوم از زمان است. این ویژگی باعث می‌شود که اطلاعات حجیمی از تغییر رفتار پدیده‌ها و نحوه تاثیر متغیرها روی یکدیگر در زمان‌های مختلف به دست آیند. با در نظر گرفتن هر یک از روش‌های ناپارامتری معرفی شده برای $\beta(t)$ ، می‌توان مساله برآورد در این نوع مدل‌های بسیار منعطف را حل نمود.

اگرچه تاکنون تحقیقات بسیار متعددی در زمینه مدل‌بندی داده‌های طولی با ضرایب متغیر از جمله هوور و همکاران (۱۹۹۸)، وو و یانگ (۲۰۰۰)، زو و ژو (۲۰۰۷)، تانگ و چنگ (۲۰۰۸)، نوه و پارک (۲۰۱۰)، تانگ و همکاران (۲۰۱۳)، یانگ و همکاران (۲۰۱۴) و چنگ و همکاران (۲۰۱۴) انجام گرفته است؛ با این وجود در هیچ‌کدام از این تحقیقات مدل پیشنهادی این مجموعه مورد بررسی قرار نگرفته است. از این رو پیشنهاد می‌شود ضمن بررسی ویژگی‌ها و اهمیت کاربردی مدل مذکور به مباحث و موضوعات زیر نیز پرداخته شود.

۱. مساله برآورد و انتخاب متغیر همزمان.
۲. بررسی نظریه‌های مرتبط با داده‌های بعد بالا و بعد بسیار بالا.
۳. مطالعه نسخه تعمیم‌یافته مدل مذکور جهت انطباق با داده‌های رسته‌ای.

جدول ۲.۴: نتایج شبیه‌سازی: کارایی روش‌های $MtLMM$ تاوانیده و $MNMM$ تاوانیده به ازای n ها و ν های مختلف.

		$n = 25$				$n = 100$			
		$\nu = 5$		$\nu = 50$		$\nu = 5$		$\nu = 50$	
		$MtLMM$	$MNMM$	$MtLMM$	$MNMM$	$MtLMM$	$MNMM$	$MtLMM$	$MNMM$
β_{11}	Est	-۱.۰۲۰	-۱.۰۳۵	-۱.۰۶۷	-۱.۰۲۷	-۰.۹۷۷	-۰.۹۵۶	-۱.۰۰۰	-۱.۰۰۱
	SE	۰.۲۸۹	۰.۲۸۶	۰.۳۲۹	۰.۳۳۱	۰.۱۸۰	۰.۲۷۰	۰.۱۴۵	۰.۱۴۵
	MSE	۰.۰۸۳	۰.۰۸۲	۰.۱۰۸	۰.۱۱۲	۰.۰۳۳	۰.۰۷۴	۰.۰۲۱	۰.۰۲۱
β_{12}	Est	-۱.۰۱۰	-۱.۰۲۷	-۱.۰۰۷	-۱.۰۱۳	-۰.۹۸۶	-۰.۹۵۸	-۰.۹۹۵	-۱.۰۰۰
	SE	۰.۲۸۲	۰.۳۴۲	۰.۲۶۰	۰.۲۵۸	۰.۱۴۰	۰.۲۵۴	۰.۱۳۸	۰.۱۳۵
	MSE	۰.۰۷۹	۰.۱۱۶	۰.۰۶۷	۰.۰۶۶	۰.۰۲۰	۰.۰۶۶	۰.۰۱۹	۰.۰۱۸
β_{13}	Est	۲.۰۶۰	۲.۰۴۹	۲.۰۳۳	۲.۰۳۶	۱.۹۸۲	۱.۹۳۷	۲.۰۲۰	۲.۰۲۲
	SE	۰.۲۴۷	۰.۳۲۵	۰.۲۶۷	۰.۲۹۱	۰.۱۵۲	۰.۱۸۰	۰.۱۲۵	۰.۱۲۰
	MSE	۰.۰۶۴	۰.۱۰۶	۰.۰۷۲	۰.۰۸۵	۰.۰۲۳	۰.۰۳۶	۰.۰۱۶	۰.۰۱۶
β_{21}	Est	-۰.۹۷۹	-۱.۰۲۲	-۰.۹۸۷	-۰.۹۸۰	-۰.۹۹۱	-۰.۹۸۰	-۰.۹۷۹	-۰.۹۸۱
	SE	۰.۲۰۰	۰.۲۱۸	۰.۲۰۵	۰.۲۲۰	۰.۱۰۴	۰.۱۳۳	۰.۱۰۰	۰.۰۹۲
	MSE	۰.۰۴۰	۰.۰۴۷	۰.۰۴۲	۰.۰۴۸	۰.۰۱۱	۰.۰۱۸	۰.۰۱۰	۰.۰۰۹
β_{22}	Est	-۱.۰۱۲	-۰.۹۸۹	-۰.۹۹۰	-۱.۰۰۰	-۰.۹۹۹	-۱.۰۰۶	-۱.۰۰۳	-۱.۰۰۲
	SE	۰.۱۹۸	۰.۲۵۷	۰.۲۰۷	۰.۱۹۷	۰.۰۸۷	۰.۱۱۸	۰.۰۹۱	۰.۰۸۳
	MSE	۰.۰۳۹	۰.۰۶۵	۰.۰۳۸	۰.۰۴۲	۰.۰۰۸	۰.۰۱۴	۰.۰۰۸	۰.۰۰۷
β_{23}	Est	۲.۰۰۳	۲.۰۱۸	۱.۹۶۶	۱.۹۵۸	۲.۰۰۹	۲.۰۰۴	۱.۹۹۱	۱.۹۹۱
	SE	۰.۱۷۶	۰.۱۹۴	۰.۱۸۰	۰.۲۰۳	۰.۰۸۵	۰.۱۰۰	۰.۰۸۵	۰.۰۸۵
	MSE	۰.۰۳۱	۰.۰۳۷	۰.۰۳۴	۰.۰۴۲	۰.۰۰۷	۰.۰۱۰	۰.۰۰۷	۰.۰۰۷
d_{11}	Est	۱.۱۷۸	۳.۴۴۷	۱.۱۸۵	۲.۲۰۵	۱.۲۸۹	۳.۵۵۱	۱.۳۹۳	۲.۲۲۵
	SE	۰.۴۰۲	۱.۳۵۰	۰.۷۲۲	۰.۴۴۹	۰.۲۶۹	۰.۷۹۴	۰.۲۶۸	۰.۳۴۶
	MSE	۰.۸۳۵	۳.۸۸۶	۰.۵۵۷	۰.۸۶۴	۰.۵۷۷	۳.۰۳۰	۰.۴۴۰	۰.۱۶۹
d_{12}	Est	۰.۱۴۲	۰.۳۳۰	۰.۱۶۱	۰.۳۱۸	۰.۱۷۳	۰.۴۴۶	۰.۲۰۳	۰.۳۱۱
	SE	۰.۲۴۷	۰.۹۶۰	۰.۵۰۱	۰.۲۸۲	۰.۱۶۱	۰.۵۹۹	۰.۱۵۶	۰.۲۴۹
	MSE	۰.۰۷۲	۰.۹۱۳	۰.۲۵۳	۰.۰۸۷	۰.۰۳۲	۰.۳۹۳	۰.۰۲۶	۰.۰۶۵
d_{22}	Est	۱.۱۵۰	۳.۴۳۴	۱.۲۳۵	۲.۲۴۱	۱.۲۱۰	۳.۴۲۲	۱.۳۵۶	۲.۱۵۸
	SE	۰.۳۷۰	۱.۳۶۰	۰.۶۶۰	۰.۵۴۱	۰.۲۷۴	۰.۸۵۳	۰.۲۹۱	۰.۳۵۲
	MSE	۰.۸۵۷	۳.۸۷۴	۰.۴۸۸	۰.۸۷۵	۰.۶۹۹	۲.۷۴۲	۰.۴۹۹	۰.۱۴۷
σ_{11}	Est	۰.۶۴۶	۲.۱۷۲	۰.۶۵۹	۱.۳۳۷	۰.۷۳۷	۲.۳۱۲	۰.۸۴۴	۱.۴۴۷
	SE	۰.۱۵۹	۰.۵۵۲	۰.۲۱۰	۰.۲۱۲	۰.۱۴۳	۰.۳۹۲	۰.۱۳۹	۰.۱۲۲
	MSE	۱.۸۵۹	۰.۳۲۹	۰.۴۸۴	۱.۸۴۳	۱.۶۱۶	۰.۲۴۹	۱.۳۵۵	۰.۳۲۲
σ_{12}	Est	۰.۲۵۲	۰.۸۶۰	۰.۲۶۲	۰.۵۴۰	۰.۲۷۷	۰.۸۸۴	۰.۳۲۴	۰.۵۵۲
	SE	۰.۰۶۶	۰.۲۵۹	۰.۱۲۵	۰.۰۹۵	۰.۰۵۸	۰.۱۹۳	۰.۰۵۹	۰.۰۶۸
	MSE	۰.۲۵۲	۰.۰۷۸	۰.۲۴۷	۰.۰۶۰	۰.۲۲۷	۰.۰۵۵	۰.۱۸۵	۰.۰۴۴
σ_{22}	Est	۰.۳۰۸	۱.۰۴۹	۰.۳۲۶	۰.۶۷۴	۰.۳۶۳	۱.۱۴۵	۰.۴۱۴	۰.۷۱۲
	SE	۰.۰۷۹	۰.۲۷۸	۰.۱۱۰	۰.۰۹۴	۰.۰۶۷	۰.۱۷۸	۰.۰۶۶	۰.۰۶۰
	MSE	۰.۴۸۶	۰.۰۷۸	۰.۴۶۳	۰.۱۱۸	۰.۴۱۱	۰.۰۵۲	۰.۳۴۸	۰.۰۸۷
ν	Est	۱.۹۴۶	-	3.17×10^6	-	۲.۱۵۰	-	۳.۹۷۱	-
	SE	۰.۳۵۲	-	2.02×10^7	-	۰.۴۲۵	-	۰.۹۰۲	-
	MSE	۹.۴۵۱	-	4.14×10^{14}	-	۸.۳۰۱	-	۲۱۱۹.۴۵	-

جدول ۳.۴: نتایج تحلیل داده‌های $PBCseq$: برآورد پارامترهای رگرسیونی (SD) تحت برازش سه مدل $MtLMM$ تاوانیده، $MNMM$ ، $MtLMM$ تاوانیده و $MNMM$.

پارامترها	$MtLMM$ تاوانیده	$MtLMM$	$MNMM$ تاوانیده	$MNMM$
β_{10} (Intercept)	۱/۲۷۹ (۰/۰۲۰)	۱/۶۴۵ (۰/۱۲۱)	۱/۲۳۶ (۰/۰۲۰)	۱/۴۹۷ (۰/۱۸۰)
β_{11} (Time)	۰/۰۲۴ (۸۸ × ۱۰ ^{-۴})	۰/۰۰۴ (۰/۰۱۱)	۰/۰۱۷۱ (۸۹ × ۱۰ ^{-۴})	۰/۰۱۷ (۰/۰۲۰)
β_{12} (Sex)	-۰/۳۷۲ (۰/۰۱۰)	-۰/۵۵۸ (۰/۰۴۰)	-۰/۳۳۱ (۰/۰۱۱)	-۰/۵۳۰ (۰/۰۶۴)
β_{13} (Drug)	۰/۰۰۰ (۰/۰۰۰)	-۰/۰۰۸ (۰/۰۲۰)	۰/۰۰۰ (۰/۰۰۰)	-۰/۰۲۳ (۰/۰۱۸)
β_{14} (Age)	-۰/۰۱۰ (۰/۰۰۶)	-۰/۰۱۱ (۰/۰۰۱)	-۰/۰۰۸ (۰/۰۰۷)	-۰/۰۱۰ (۰/۰۰۲)
اثرات ثابت				
β_{20} (Intercept)	۱/۰۱۳ (۰/۰۰۲)	۱/۴۰۲ (۰/۰۸۶)	۱/۳۴۳ (۰/۰۰۲)	۱/۳۱۲ (۰/۰۱۷)
β_{21} (Time)	-۰/۰۱۵ (۱/۱ × ۱۰ ^{-۴})	-۰/۰۲۰ (۰/۰۰۶)	-۰/۰۱۳ (۱/۲ × ۱۰ ^{-۴})	-۰/۰۱۴ (۰/۰۰۱)
β_{22} (Sex)	۰/۰۰۰ (۱/۴ × ۱۰ ^{-۴})	-۰/۰۲۴ (۰/۰۳۵)	-۰/۰۰۰ (۱/۵ × ۱۰ ^{-۴})	-۰/۰۰۱ (۰/۰۰۵)
β_{23} (Drug)	۰/۰۰۰ (۸/۴ × ۱۰ ^{-۴})	۰/۰۲۸ (۰/۰۱۶)	۰/۰۰۰ (۸۸ × ۱۰ ^{-۴})	۰/۰۲۵ (۰/۰۰۳)
β_{24} (Age)	۰/۰۰۶ (۳/۶ × ۱۰ ^{-۵})	-۰/۰۰۰۲ (۰/۰۰۱)	-۰/۰۰۰۲ (۳/۶ × ۱۰ ^{-۵})	-۰/۰۰۱ (۲ × ۱۰ ^{-۴})
d_{11}	۰/۸۷۹	۲/۱۲۵	۱/۰۴۰	۱/۱۱۴
d_{21}	۰/۰۲۵	۰/۰۰۰	۰/۰۲۸	-۰/۳۴۳
d_{22}	۰/۰۷۴	۰/۲۱۱	۰/۰۸۲	۲/۶۳۱
d_{31}	-۰/۰۳۸	-۰/۱۰۰	-۰/۰۵۹	-۰/۰۶۵
d_{32}	-۰/۰۰۳	-۰/۰۰۱	-۰/۰۰۱	۰/۰۱۸
اثرات تصادفی				
d_{43}	۰/۰۱۳	۰/۰۲۲	۰/۰۱۲	۰/۰۱۳
d_{41}	-۰/۰۱۱	-۰/۰۲۳	-۰/۰۰۷	۰/۰۱۶
d_{42}	-۰/۰۰۹	-۰/۰۲۴	-۰/۰۰۹	-۰/۱۵۷
d_{43}	۰/۰۰۱	-۰/۰۰۲	-۰/۰۰۱	-۰/۰۱۵
d_{44}	۰/۰۰۳	۰/۰۱۰	۰/۰۰۲	۰/۰۳۷
σ_{11}	۰/۰۴۶	۰/۱۲۴	۰/۰۶۶	۰/۱۰۲
σ_{21}	-۰/۰۰۱	-۰/۰۰۱	-۰/۰۰۰۵	-۰/۰۰۲
σ_{22}	۰/۰۰۳	۰/۰۱۲	۰/۰۰۶۱	۰/۰۱۹
ν	۱/۰۸۲	۱/۰۰۰	-	-
خطاها				
ℓ_{max}	۸۶۷۵۴۰	۵۴۳۳۴۴	۶۵۰۱/۰۵	۵۳۴۶۲۷
AIC	-۱۷۳۰۲/۸۱	-۱۰۸۱۸/۸۸	-۱۲۹۵۶/۱۱	-۱۰۶۴۶/۵۴
BIC	-۱۷۲۱۲/۹۸	-۱۰۷۲۹/۰۵	-۱۲۸۷۰/۰۲	-۱۰۵۶۰/۴۵
$MSPE$	۷/۷۴	۶/۸۶	۷/۷۵	۷/۸۴

جدول ۴.۴: نتایج شبیه‌سازی: کارایی روش‌های $MtSMM$ ، $MtLMM$ و $MNSMM$ به ازای $n = ۲۵$ و ν های مختلف.

پارامترها		$n = ۲۵$					
		$\nu = ۵$			$\nu = ۵۰$		
		$MtSMM$	$MtLMM$	$MNSMM$	$MtSMM$	$MtLMM$	$MNSMM$
β_{11}	$BIAS$	۰٫۶۴۱	۷٫۹۶۶	۰٫۶۴۵	۰٫۷۴۶	۷٫۹۰۶	۰٫۷۴۶
	SD	۰٫۳۶۶	۰٫۴۸۹	۰٫۴۱۸	۰٫۳۶۸	۰٫۳۹۳	۰٫۳۶۸
	MSE	۰٫۵۴۳	۶۳٫۷۰۱	۰٫۵۸۹	۰٫۶۹۰	۶۲٫۶۵۸	۰٫۶۹۰
β_{12}	$BIAS$	۰٫۰۳۰	۰٫۵۵۵	۰٫۰۶۴	۰٫۰۵۸	۰٫۶۴۸	۰٫۰۶۴
	SD	۰٫۲۲۴	۰٫۷۷۸	۰٫۲۶۵	۰٫۲۷۱	۰٫۷۳۰	۰٫۲۷۰
	MSE	۰٫۰۵۰	۰٫۹۰۸	۰٫۰۷۴	۰٫۰۷۶	۰٫۹۴۷	۰٫۰۷۶
β_{21}	$BIAS$	۰٫۹۸۷	۱٫۰۷۴	۰٫۹۵۵	۰٫۸۹۲	۱٫۰۶۷	۰٫۸۷۹
	SD	۰٫۳۷۷	۰٫۴۱۷	۰٫۳۷۷	۰٫۳۷۵	۰٫۳۹۱	۰٫۳۷۳
	MSE	۱٫۱۱۶	۱٫۳۲۴	۱٫۰۵۳	۰٫۹۳۵	۱٫۲۸۹	۰٫۹۱۰
β_{22}	$BIAS$	۰٫۳۱۲	۱٫۲۹۲	۰٫۲۱۵	۰٫۱۷۳	۱٫۲۸۷	۰٫۱۵۱
	SD	۰٫۴۸۴	۰٫۵۱۲	۰٫۴۱۴	۰٫۳۱۴	۰٫۴۳۵	۰٫۳۰۰
	MSE	۰٫۳۲۹	۱٫۹۳۰	۰٫۲۱۶	۰٫۱۲۸	۱٫۸۴۴	۰٫۱۱۲
d_{11}	$BIAS$	۰٫۸۴۹	۰٫۹۵۴	۰٫۹۶۵	۰٫۱۳۱	۰٫۹۵۷	۰٫۱۱۰
	SD	۰٫۳۷۹	۰٫۲۶۱	۱٫۳۲۵	۰٫۶۴۴	۰٫۲۵۹	۰٫۷۳۷
	MSE	۰٫۸۶۲	۰٫۹۷۷	۲٫۶۷۰	۰٫۴۲۷	۰٫۹۸۳	۰٫۵۵۰
d_{12}	$BIAS$	۰٫۰۳۶	۰٫۱۶۰	۰٫۳۰۳	۰٫۰۵۰	۰٫۱۹۴	۰٫۰۸۶
	SD	۰٫۳۳۸	۰٫۱۶۵	۰٫۹۹۰	۰٫۵۰۲	۰٫۱۶۶	۰٫۵۶۲
	MSE	۰٫۱۱۴	۰٫۰۵۳	۱٫۰۶۲	۰٫۲۵۱	۰٫۰۶۵	۰٫۳۲۰
d_{22}	$BIAS$	۰٫۵۵۳	۱٫۶۵۰	۱٫۷۰۳	۰٫۰۴۳	۱٫۷۳۴	۰٫۳۰۲
	SD	۰٫۵۴۶	۰٫۱۵۲	۲٫۰۷۰	۰٫۸۲۸	۰٫۱۰۰	۰٫۹۳۹
	MSE	۰٫۶۰۱	۲٫۷۴۴	۷٫۱۴۳	۰٫۶۸۰	۳٫۰۱۷	۰٫۹۶۴
σ_{11}	$BIAS$	۱٫۲۲۶	۳٫۱۳۵	۰٫۱۸۸	۰٫۷۷۰	۳٫۲۴۷	۰٫۵۹۵
	SD	۰٫۱۷۷	۰٫۵۱۱	۰٫۷۳۴	۰٫۱۹۰	۰٫۶۱۴	۰٫۲۲۳
	MSE	۱٫۵۳۳	۱۰٫۰۸۵	۰٫۵۶۹	۰٫۶۲۷	۱۰٫۹۱۳	۰٫۴۰۳
σ_{12}	$BIAS$	۰٫۴۵۳	۰٫۷۶۸	۰٫۰۸۳	۰٫۲۲۷	۰٫۷۷۰	۰٫۱۵۶
	SD	۰٫۱۵۷	۰٫۱۲۷	۰٫۴۳۲	۰٫۲۴۰	۰٫۰۹۶	۰٫۲۷۰
	MSE	۰٫۲۳۰	۰٫۶۰۶	۰٫۱۹۲	۰٫۱۰۸	۰٫۶۰۲	۰٫۱۰۰
σ_{22}	$BIAS$	۰٫۳۹۸	۰٫۶۳۹	۲٫۱۶۸	۱٫۵۵۶	۰٫۶۴۶	۱٫۸۸۳
	SD	۰٫۲۲۲	۰٫۰۵۷	۰٫۵۰۷	۰٫۳۱۰	۰٫۰۶۰	۴۰٫۳۴۴
	MSE	۰٫۲۰۷	۰٫۴۱۱	۴٫۹۵۳	۲٫۵۱۵	۰٫۴۲۰	۳٫۶۶۱
ν	$BIAS$	۰٫۰۲۳	۳٫۲۶۲	–	۱٫۱۵۷	۴۸٫۹۹۹	–
	SD	۰٫۱۳۲	۰٫۴۳۹	–	۰٫۹۷۹	۰٫۰۰۸	–
	MSE	۰٫۰۱۸	۱۰۸	–	۲٫۲۸۸	۲۴۰۰٫۸۸۴	–
ℓ_{max}		–۲۱۵٫۳۲۹	–۲۳۲٫۳۷۴	–۳۲۷٫۲۵۳	–۲۶۵٫۹۲۳	–۲۲۸٫۷۵۷	–۲۸۷٫۰۸۹
AIC		۴۵۲٫۶۵۸	۴۸۶٫۷۴۸	۶۷۴٫۵۰۶	۵۵۳٫۸۴۶	۴۷۹٫۵۱۵	۵۹۴٫۱۷۷
BIC		۴۶۶٫۰۶۵	۵۰۰٫۱۵۶	۶۸۶٫۶۹۵	۵۶۷٫۲۵۴	۴۹۲٫۹۲۲	۶۰۶٫۳۶۶

جدول ۵.۴: نتایج شبیه‌سازی: کارایی روش‌های $MtSMM$ ، $MtLMM$ و $MNSMM$ به ازای $n = ۱۰۰$ و ν های مختلف.

Parameters		$n = ۱۰۰$					
		$\nu = ۵$			$\nu = ۵۰$		
		$MtSMM$	$MtLMM$	$MNSMM$	$MtSMM$	$MtLMM$	$MNSMM$
β_{11}	BIAS	۰/۷۲۴	۸/۱۸۷	۰/۷۲۷	۰/۷۴۰	۸/۱۹۴	۰/۷۴۰
	SD	۰/۱۹۰	۰/۲۵۵	۰/۲۱۹	۰/۱۶۲	۰/۱۷۶	۰/۱۶۳
	MSE	۰/۵۶۰	۶۷/۰۸۵	۰/۵۷۶	۰/۵۷۳	۶۷/۱۷۳	۰/۵۷۳۴
β_{12}	BIAS	۰/۰۳۳	۰/۸۶۹	۰/۰۲۳	۰/۰۰۶	۰/۸۲۵	۰/۰۰۶
	SD	۰/۱۲۰	۰/۲۷۲	۰/۱۳۸	۰/۱۰۰	۰/۲۴۷	۰/۱۰۳
	MSE	۰/۰۱۵	۰/۸۲۸	۰/۰۱۹	۰/۰۱۰	۰/۷۴۲	۰/۰۱۱
β_{21}	BIAS	۰/۰۱۸	۰/۱۶۲	۰/۰۱۴	۰/۰۶۱	۰/۱۷۹	۰/۰۵۶
	SD	۰/۲۲۳	۰/۲۳۹	۰/۲۳۹	۰/۱۴۸	۰/۲۱۴	۰/۱۴۷
	MSE	۰/۰۵۰	۰/۰۸۳	۰/۰۵۷	۰/۰۲۵	۰/۰۷۷	۰/۰۲۵
β_{22}	BIAS	۰/۷۰۸	۰/۵۶۸	۰/۶۲۲	۰/۶۴۳	۰/۵۰۹	۰/۶۳۲
	SD	۰/۲۳۰	۰/۲۸۶	۰/۲۱۸	۰/۱۵۵	۰/۲۱۱	۰/۱۵۷
	MSE	۰/۵۵۲	۰/۴۰۴	۰/۴۳۴	۰/۴۳۷	۰/۳۰۳	۰/۴۲۴
d_{11}	BIAS	۰/۵۶۲	۱/۰۲۸	۱/۵۷۰	۰/۱۱۴	۱/۰۴۴	۰/۲۱۲
	SD	۰/۲۷۷	۰/۱۶۷	۰/۹۷۵	۰/۳۳۹	۰/۱۵۸	۰/۳۹۴
	MSE	۰/۳۹۲	۱/۰۸۴	۳/۴۰۴	۰/۱۲۷	۱/۱۱۴	۰/۱۹۹
d_{12}	BIAS	۰/۰۴۴	۰/۲۴۴	۰/۲۵۳	۰/۰۲۱	۰/۲۳۷	۰/۰۲۱
	SD	۰/۲۰۸	۰/۰۸۷	۰/۷۴۷	۰/۲۷۱	۰/۰۷۹	۰/۳۱۱
	MSE	۰/۰۴۵	۰/۰۶۷	۰/۶۱۶	۰/۰۷۳	۰/۰۶۲	۰/۰۹۶
d_{22}	BIAS	۰/۳۳۲	۱/۶۳۱	۱/۶۹۷	۰/۱۰۲	۱/۶۸۳	۰/۴۵۸
	SD	۰/۲۹۳	۰/۰۸۰	۰/۸۶۱	۰/۳۵۱	۰/۰۵۸	۰/۴۰۳
	MSE	۰/۱۹۵	۲/۶۶۸	۳/۶۱۲	۰/۱۳۳	۲/۸۳۶	۰/۳۷۰
σ_{11}	BIAS	۱/۱۳۷	۴/۰۰۹	۴/۰/۲۹۸	۰/۷۶۸	۴/۳۲۰	۰/۵۴۲
	SD	۰/۱۲۳	۰/۶۲۳	۰/۴۰۳	۰/۱۰۶	۰/۵۸۰	۰/۱۲۳
	MSE	۱/۳۰۸	۱۶/۴۵۲	۰/۲۵۰	۰/۶۰۱	۱۸/۹۹۸	۰/۳۰۹
σ_{12}	BIAS	۰/۴۲۳	۰/۷۴۱	۰/۱۲۱	۰/۲۴۷	۰/۷۸۰	۰/۱۵۶
	SD	۰/۰۷۹	۰/۰۵۸	۰/۲۲۳	۰/۱۱۴	۰/۰۵۶	۰/۱۳۰
	MSE	۰/۱۸۵	۰/۵۵۲	۰/۰۶۴	۰/۰۷۴	۰/۶۱۱	۰/۰۴۱
σ_{22}	BIAS	۰/۸۰۷۸	۰/۵۶۹	۲/۶۹۵	۱/۹۲۱	۰/۵۶۷	۲/۳۹۱
	SD	۰/۲۰۷	۰/۰۵۴	۰/۲۶۰	۰/۱۸۳	۰/۰۵۲	۰/۱۷۵
	MSE	۰/۶۹۵	۰/۳۲۶	۷/۳۲۷	۳/۷۲۳	۰/۳۲۵	۵/۷۴۵
ν	BIAS	۰/۰۱۴	۳/۲۹۵	–	۰/۷۶۰	۴۸/۹۹۸	–
	SD	۰/۱۴۳	۰/۰۸۰	–	۰/۹۷۶	۰/۰۰۶	–
	MSE	۰/۰۲۰	۱۰/۸۶۴	–	۱/۵۲۰	۲۴/۰۰/۸۴۶	–
ℓ_{max}		–۹۹۴/۸۰۸	–۱۰۳۴/۴۹۱	–۱۴۰۰/۳۷۷	–۱۱۳۵/۹۴۸	–۱۰۳۸۵۰۰	–۱۲۲۶/۶۳۱
AIC		۲۰۱۱/۶۱۷	۲۰۹۰/۹۸۲	۲۸۲۰/۷۵۳	۲۲۹۳/۸۹۶	۲۰۹۸/۹۹۹	۲۴۷۳/۲۶۲
BIC		۲۰۴۰/۲۷۴	۲۱۱۹/۶۳۹	۲۸۴۶/۸۰۵	۲۳۲۲/۵۵۲	۲۱۲۷/۶۵۶	۲۴۹۹/۳۱۴

جدول ۶.۴: نتایج تحلیل داده‌های $PBCseq$: نتایج معیارهای انتخاب مدل تحت برازش مدل‌های $MtSMM$ و $MNSMM$ و ساختارهای همبستگی مختلف.

مدل	Ω_i	Criteria			
		ℓ_{max}	AIC	BIC	$MSPE$
$MtSMM$	UNC	۵۰۹۰/۴۹	-۱۰۱۴۲/۹۸	-۱۰۰۵۴/۶۳	۶/۴۰
	EX	۵۲۵۱/۲۴	-۱۰۴۶۲/۴۸	-۱۰۳۷۶/۱۳	۶/۳۸
	$AR(1)$	۵۳۵۸/۴۱	-۱۰۶۷۲/۸۲	-۱۰۵۹۰/۴۷	۶/۱۳
$MNSMM$	UNC	۳۸۱۷/۹۴	-۷۵۹۳/۸۸	-۷۵۱۵/۲۸	۷/۲۴
	EX	۳۹۳۸/۵۰	-۷۸۳۵/۰۰	-۷۷۵۶/۴۰	۷/۱۸
	$AR(1)$	۴۰۱۸/۸۸	-۷۹۹۵/۷۷	-۷۹۱۷/۱۶	۷/۱۰

جدول ۷.۴: نتایج تحلیل داده‌های PBC_{seq} : برآورد پارامترهای رگرسیونی (SD) تحت برازش مدل‌های $MtSMM$ و $MNSMM$.

پارامترها	$MtSMM$	$MNSMM$
$\beta_{10}(Intercept)$	۱/۴۳۲(۰/۰۱۸)	۱/۴۲۸(۰/۱۱۲)
$\beta_{11}(Sex)$	-۰/۵۴۲(۰/۰۱۰)	-۰/۵۳۸(۰/۰۴۲)
$\beta_{12}(Drug)$	-۰/۰۱۵(۰/۰۰۷)	-۰/۰۱۹(۰/۰۱۳)
$\beta_{13}(Age)$	-۰/۰۱۰(۳/۳ $\times 10^{-4}$)	-۰/۰۱۰۲(۰/۰۰۰۲)
اثرات ثابت	$\beta_{20}(Intercept)$	۱/۱۷۱(۰/۰۰۲)
	$\beta_{21}(Sex)$	-۰/۰۰۰۴(۰/۰۰۰۱)
	$\beta_{22}(Drug)$	۰/۰۱۰(۸/۹ $\times 10^{-4}$)
	$\beta_{23}(Age)$	-۰/۰۰۰۲(۴/۱ $\times 10^{-5}$)
اثرات تصادفی	d_{11}	۰/۸۰۷
	d_{21}	-۰/۰۱۲
	d_{22}	۰/۳۱۷
	d_{31}	-۰/۰۲۲
	d_{32}	-۰/۰۰۳
	d_{33}	۰/۳۱۰
	d_{41}	-۰/۰۰۷
	d_{42}	-۰/۰۰۵
	d_{43}	-۰/۰۹۵
	d_{44}	۰/۲۵۶
	σ_{11}	-۰/۰۸۵
	σ_{21}	-۰/۰۰۳
خطاها	σ_{22}	-۰/۲۲۲
	ν	۳/۱۲۰
		-

مراجع

- [۱] آرست، م. (۱۳۹۵)، مقایسه رفتار برخی برآوردگرهای انقباضی بریج در مدل رگرسیون چندگانه، پایان نامه کارشناسی ارشد، دانشگاه صنعتی شاهرود.
- [۲] تعاونی، م. (۱۳۹۲)، استفاده از مدل های مختلف رگرسیونی در تحلیل داده های طولی، پایان نامه کارشناسی ارشد، دانشگاه گلستان.
- [۳] رضایی قهرودی، م.، گنجعلی، م. و هرنندی، ف. (۱۳۸۸)، مدل انتقالی برای تحلیل داده های طولی با پاسخ های دو متغیره ترتیبی و اسمی، پژوهش های آماری ایران.
- [۴] روزبه، م. (۱۳۹۰)، برآورد در مدل های خطی جزئی، رساله دکتری، دانشگاه فردوسی مشهد.
- [۵] کاظمی، م.، شاهسونی، د. و آرشی، م. (۱۳۹۷)، انتخاب متغیر و تشخیص ساختار در بعد بالا برای مدل های جمعی خطی-جزیی، مجله علوم آماری، ۱۲، ۴۸۵-۵۱۲.
- [۶] گنجعلی، م. و رضایی قهرودی، ز. (۱۳۸۹)، تحلیل چند متغیره ی گسسته در مطالعات مقطعی و طولی، پژوهشکده آمار: تهران.
- [7] Airy, G.B. (1861), *On the Algebraical and Numerical Theory of Errors of Observation and the Combination of Observations*, Macmillan, London.
- [8] Akaike, H. (1973), *Information Theory and an Extension of the Maximum Likelihood Principle* In 2nd Int. Symp. on Information Theory. (Edited by B. N. Petrov and F. Csaki), Akademiai Kiado, Budapest, 267-281.
- [9] Belsley, D.A., Kuh, E. and Welsch, R.E. (1980), *Regression Diagnostics*, John Wiley, New York.
- [10] Bondell, H.D., Krishna, A. and Ghosh, S.K. (2010), Joint Variable Selection for Fixed and Random Effects in Linear Mixed-effects Models, *Biometrics*, 66, 1069-1077.
- [11] Bowman, A. (1984), An Alternative Method of Cross-Validation for the Smoothing of Density Estimates, *Biometrika*, 71, 353-360.

-
- [12] Breiman, L. (1996), Heuristics of Instability and Stabilization in Model Selection, *Annal. Statist.*, 24(6), 2350–2383.
- [13] Breheny, P., and Huang, J.(2011), Coordinate Descent Algorithms for Nonconvex Penalized Regression, with Applications to Biological Feature Selection, *Annal. Appl. Statist.*, 5, 232–253.
- [14] Breheny, P., and Huang, J. (2015), Group Descent Algorithms for Nonconvex Penalized Linear and Logistic Regression Models with Grouped Predictors, *Statist. Comp.*, 25(2), 173–187.
- [15] Breheny, P., and Huang, J. (2019), *High-Dimensional Regression Modeling: Methodology, Applications, and Software*, Chapman and Hall.
- [16] Breslow, N.E. and Clayton, D.G. (1993), Approximate Inference in Generalized Linear Mixed Models, *J. Amer. Statist. Assoc.*, 88(421), 9–25.
- [17] Buhlmann, P., and Van De Geer, S. (2011), *Statistics for High-Dimensional Data: Methods, Theory and Applications*, Springer Science and Business Media.
- [18] Cantoni, E., Mills Flemming, J. and Ronchetti, E. (2005), Variable Selection for Marginal Longitudinal Generalized Linear Models, *Biometrics*, 61(2), 507–514.
- [19] Chatterjee, S., and Hadi, A.S. (2012), *Regression Analysis by Example*, Wiley, New York.
- [20] Cheng, M.Y., Honda, T., Li, J. and Peng, H. (2014), Nonparametric Independence Screening and Structural Identification for Ultra-high Dimensional Longitudinal Data, *Annal. Statist.*, 42, 1819–1849.
- [21] Chung, H., Park, Y. and Lanza, S.T. (2005), Latent Transition Analysis with Covariates: Pubertal Timing and Substance Use Behaviours in Adolescent Females, *Statist. Med.*, 24, 2895–2910.
- [22] Dempster AP, Laird NM, Rubin DB. (1977), Maximum Likelihood from Incomplete Data Via the EM Algorithm. *J. Roy. Statist. Soc. B*, 39, 1–38.
- [23] Diggle, P.J., Heagerty, P., Liang, K.Y. and Zeger, S.L. (2002), *Analysis of Longitudinal Data*, Oxford University Press, New York.
- [24] Dziak, J.J. (2006), *Penalized Quadratic Inference Functions for Variable Selection in Longitudinal Research*, Ph.D Thesis, the Pennsylvania State University.

- [25] Drapper, N.R., and Smith, H. (1981), *Applied Regression Analysis*, Wiley.
- [26] Dziak, J., Li, R., and Collins, L. (2005), *Critical Review and Comparison of Variable Selection Procedures for Linear Regression*, State College, Pennsylvania State University.
- [27] Desboulets, L. (2018), A Review on Variable Selection in Regression Analysis, *Econometrics*, 6(4), ; doi: 10.3390/econometrics6040045.
- [28] Ding, J., Tarokh, V., and Yang, Y. (2018), Model Selection Techniques: An Overview. *IEEE Sig. Proc. Mag.*, 35(6), 16–34.
- [29] Efron, B., Hastie, T., Johnstone, I. and Tibshirani, R. (2004), Least Angle Regression, *Annal. Statist.*, 32, 407–451.
- [30] Eliot, M., Ferguson, J., Reilly M.P. and Foulkes, A.S. (2011), Ridge Regression for Longitudinal Biomarker Data, *Int. J. Biostat.*, 7, Article 37.
- [31] Engle, R.F., Granger, C.W.J., Rice, J. and Weiss, A. (1986). Semiparametric Estimates of the Relation Between Weather and Electricity Sales, *J. Amer. Statist. Assoc.*, 81, 310–320.
- [32] Fan, J.Q., Huangm, T. and Li, R. (2007), Analysis of Longitudinal Data with Semiparametric Estimation of Covariance Function, *J. Amer. Statist. Assoc.*, 102, 632–641.
- [33] Fan, J.Q. and Li, R. (2001), Variable Selection via Nonconcave Penalized Likelihood and its Oracle Properties, *J. Amer. Statist. Assoc.*, 96, 1348–1360.
- [34] Fan, J.Q. and Li, R. (2004), New Estimation and Model Selection Procedures for Semiparametric Modeling in Longitudinal Data Analysis, *J. Amer. Statist. Assoc.*, 99, 710–723.
- [35] Fan, J.Q, and Li, R. (2006), Statistical Challenges with High Dimensionality: Feature Selection in Knowledge Discovery, arXiv preprint math/0602133.
- [36] Fan, J.Q, and Lv, J. (2008), Sure Independence Screening for Ultrahigh Dimensional Feature Space, *J. Roy. Statist. Soc.*, B, 70(5), 849–911.
- [37] Fan, J.Q, Ma, Y., and Dai, W. (2014), Nonparametric Independence Screening in Sparse Ultra-high-dimensional Varying Coefficient Models, *J. Amer. Statist. Assoc.*, 109(507), 1270–1284.
- [38] Fan, J.Q, Samworth, R., and Wu, Y. (2009), Ultrahigh Dimensional Feature Selection: Beyond the Linear Model, *J. Mach. Learn. Res.*, 10, 2013–2038.

-
- [39] Fitzmaurice, G., Davidian, M., Verbeke, G. and Molenberghs, G. (2009), *Longitudinal Data Analysis*, Chapman and Hall, London.
- [40] Fitzmaurice, G.M. and Laird, N.M. (1993), A Likelihood-Based Method for Analysing Longitudinal Binary Responses, *Biometrika*, 80, 141- 151.
- [41] Fitzmaurice, G.M., Laird, N.M. and Ware, J.H. (2011), *Applied Longitudinal Analysis*, Wiley, New York.
- [42] Frank, I.E. and Friedman, J.H. (1993), A statistical View of Some Chemometrics Regression Tools (with discussion), *Technometrics*, 35, 109–148.
- [43] Fu, W.J. (1998), Penalized Regressions: the Bridge Versus the Lasso, *J. Comp. Graph. Gstatist.*, 7(3), 397–416.
- [44] Fu, W.J. (2003), Penalized Estimating Equations, *Biometrics*, 59, 126–132.
- [45] Geenens, G. (2011). Curse of Dimensionality and Related Issues in Nonparametric Functional Regression, *Statist. Surveys*, 5, 30–43.
- [46] Giraud, C. (2014), *Introduction to High-dimensional Statistics*, Chapman and Hall.
- [47] Green, P.J. and Silverman, B.W. (1994), *Nonparametric Regression and Generalized Linear Models: A Roughness Penalty Approach*, Chapman and Hall, London.
- [48] Gujarati, D.N. (2004), *Basic Econometrics*, Tata McGraw- Hill.
- [49] Hardin, J.W. and Hilbe, J.M. (2003), *Generalized Estimating Equations*, Chapman and Hall, New York.
- [50] Hardle, W., Liang, H. and Gao, J.T. (2000), *Partially Linear Models*, Springer, Heidelberg.
- [51] Harville, D.A. (1977), Maximum Likelihood Approaches to Variance Component Estimation and to Related Problems, *J. Amer. Statist. Assoc.*, 72, 320–338.
- [52] Hastie, T. and Tibshirani, R. (1990), *Generalized Additive Models*, Chapman and Hall, London.
- [53] Hedeker, D. and Gibbons, R.D. (1994), A Random-Effects Ordinal Regression Model for Multilevel Analysis, *Biometrics*, 50, 933–944.
- [54] Hedeker, D. and Gibbons, R.D. (2006), *Longitudinal Data Analysis*, Wiley, New York.

- [55] He, X, Fung, W.K., Zhu, Z.Y. (2005), Robust Estimation in Generalized Partial Linear Models for Clustered Data, *J. Amer. Statist. Assoc.*, 100, 1176–1184
- [56] He, X.M., Zhu, Z.Y. and Fung, W.K. (2002), Estimation in a Semiparametric Model for Longitudinal Data with Unspecified Dependence Structure, *Biometrika*, 89, 579–590.
- [57] Hoerl, A. and Kennard, R. (1970b), Ridge Regression: Biased Estimation for Nonorthogonal Problems, *Technometrics*, 12, 55–67.
- [58] Hoover, D.R., Rice, J.A., Wu, C.O. and Yang, L.P. (1998), Nonparametric Smoothing Estimates of Timevarying Coefficient Models with Longitudinal Data, *Biometrika*, 85, 809–822.
- [59] Huang, J., Wu, C., and Zhou, L. (2004), Polynomial Spline Estimation and Inference for Varying Coefficient Models with Longitudinal Data, *Statist. Sinica*, 14, 763–788.
- [60] Huang, J.Z., Zhang, L., and Zhou, L. (2007), Efficient Estimation in Marginal Partially Linear Models for Longitudinal/Clustered Data using Splines, *Scand. J. Stat.*, 34, 451–477.
- [61] Hu, Z., Wang, N. and Carroll, J. R. (2004), Profile-Kernel Versus Backfitting in the Partially Linear Models for Longitudinal/Clustered Data, *Biometrika*, 91, 251–262.
- [62] James, G., Witten, D., Hastie, T., and Tibshirani, R. (2013), An Introduction to Statistical Learning, springer, New York.
- [63] Johnstone, I.M., and Titterington, D.M. (2009), Statistical Challenges of Highdimensional Data, *Philos. Trans. R. Soc. B A*, 367, 4237–4253.
- [64] Kibria, B.M.G. (2006), The Matrix-t Distribution and its Applications in Predictive Inference, *J. Multivariate Anal.* 97, 785–795.
- [65] Kurum, E., Li, R., Shiffman, S. and Yao, W. (2016), Time- varying Coefficient Models for Joint Modeling Binary and Continues Outcome in Longitudinal Data, *Stat. Sinica*, 29, 979–1000
- [66] Lee, Y. and Nelder, J.A. (2001), Hierarchical Generalized Linear Models: A Synthesis of Generalized Linear Models, Random-Effect Models and Structured Dispersions, *Biometrika*, 88, 987–1006.
- [67] Leon, A.C., Hedeker, D. and Teres, J.J. (2007), Bias Reduction in Effectiveness Analyses of Longitudinal Ordinal Doses with a Mixed-Effects Propensity Adjustment, *Statist. Med.*, 26, 110-123.

-
- [68] Liang, H. (2009), Generalized Partially Linear Mixed-Effects Models Incorporating Mis-measured Covariates, *Ann. Inst. Statist. Math.*, 61, 27–46.
- [69] Liang, K.Y. and Zeger, S.L. (1986), Longitudinal Data Analysis Using Generalized Linear Models, *Biometrika*, 73, 13–22.
- [70] Li, Z. and Zhu, L. (2010), On Variance Components in Semiparametric Mixed Models for Longitudinal Data, *Scan. J. Statist.*, 37, 442–457.
- [71] Lin, X. (2007), Estimation using Penalized Quasilikelihood and Quasi-pseudo-likelihood in Poisson Mixed Models, *Lifetime Data Analysis*, 13, 533–544.
- [72] Lin, X. and Carroll, R.J. (2000), Nonparametric Function Estimation for Clustered Data when the Predictor is Measured without/with Error, *J. Amer. Statist. Assoc.*, 95, 520–534.
- [73] Lin, T.I. and Lee, J.C. (2006), A Robust Approach to t Linear Mixed Models Applied to Multiple Sclerosis Data, *Statist. Med.* 25, 1397–1412.
- [74] Lin, T.I. and Lee, J.C. (2007), Bayesian Analysis of Hierarchical Linear Mixed Modeling using the Multivariate t Distribution, *J. Statist. Plann. Inference.* 137, 484–495.
- [75] Ma, S., Song, Q. and Wang, L. (2013), Simultaneous Variable Selection and Estimation in Semiparametric Modeling of Longitudinal/Clustered Data, *Bernoulli*, 19(1), 252–274.
- [76] Malo, N., Libiger, O. and Schork, N. (2008), Accommodating Linkage Disequilibrium in Genetic-Association Analyses via Ridge Regression, *Amer. J. Hum. Genet.*, 82, 375–85.
- [77] Marshall, G., De la Cruz-Mesia, R., Baron, A.E., Rutledge, J.H. and Zerbe, G.O. (2006), Non-linear Random Effects Model for Multivariate Responses with Missing Data, *Statist. Med.*, 25, 2817–2830.
- [78] Mcculloch, C.E. (1997), Maximum Likelihood Algorithms for Generalized Linear Mixed Models, *J. Amer. Statist. Assoc.*, 92, 162–170.
- [79] McCullagh, P. and Nelder, J.A. (1989), *Generalized Linear Models*, Chapman and Hall, London, second edition.
- [80] Meng, X.L. and Rubin, D.B. (1993), Maximum Likelihood Estimation via the ECM Algorithm: A General Framework. *Biometrika*, 80, 267–278.

- [81] Molenberghs, G. and Verbeke, G. (2005), *Models for Discrete Longitudinal Data*, Springer, New York.
- [82] Montgomery, D.C., Peck, E.A. and Vining, G.G. (2012), *Introduction to Linear Regression Analysis*, Wiley, New York.
- [83] Natarajan, R. and Kass, R.E. (2000), Reference Bayesian methods for generalized linear mixed models, *Amer. Statist.*, 95, 227–237.
- [84] Nelder, J.A. and Wedderburn, R.W.M. (1972), Generalized Linear Models, *J. Roy. Statist. Soc.*, 135, 370–384.
- [85] Ni, X., Zhang, D. and Zhang, H.H. (2010), Variable Selection for Semiparametric Mixed Models in Longitudinal Studies, *Biometrics*, 66, 79–88.
- [86] Noh, H. and Park, B., (2010), Sparse Varying Coefficient Models for Longitudinal Data, *Statist. Sinica*, 20, 1183–1202.
- [87] Pan, W. (2001), Akaike's Information Criterion in Generalized Estimating Equations, *Biometrics*, 57, 120–125.
- [88] Pan, W. (2002), Goodness-of-Fit Tests for GEE with Correlated Binary Data, *Scand. J. Statist.*, 29, 101–110.
- [89] Pan, J. and Thompson, R. (2007), Quasi-Monte Carlo Estimation in Generalized Linear Mixed Models, *Comput. Statist. Data Anal.*, 51, 5765–5775.
- [90] Pinheiro, J.C., Liu, C.H. and Wu, Y.N. (2001), Efficient Algorithms for Robust Estimation in Linear Mixed Effects Model using the Multivariate t Distribution, *J. Comput. Graph. Statist.*, 10, 249–276.
- [91] Prentice, R.L. (1998), Correlated Binary Regression with Covariates Specific to Each Binary Observation, *Biometrics*, 44, 1033–1048.
- [92] Qin, G.Y. and Zhu, Z.Y. (2007), Robust Estimation in Generalized Semiparametric Mixed Models for Longitudinal Data, *J. Mult. Anal.*, 98, 1658–1683.
- [93] Qin, G.Y. and Zhu, Z.Y., (2009), Robustified Maximum Likelihood Estimation in Generalized Partial Linear Mixed Model for Longitudinal Data, *Biometrics*, 65, 52–59.

-
- [94] Rahmani, M., Arashi, M., Mamode Khan, N. and Sunecher, Y. (2018), Improved Mixed Model for Longitudinal Data Analysis Using Shrinkage m Method, *Math. Sciences*, 12, 305–312.
 - [95] Rezaee, Z., and Ganjali, M. (2009), Testing Homogeneity in Markov Models for Analyzing Longitudinal Ordinal Response Data with Random Dropout, *J. Statist. Theo. App.*, 8, 125–139.
 - [96] Rezaee, Z., Ganjali, M., and Berridge, D. (2009), A Transition Model for Ordinal Response Data with Random Dropout: An Application to the Fluvoxamine Data, *J. Biopharm. Statist.*, 19, 658–671.
 - [97] Rice, J.A. and Wu, C.O. (2001), Nonparametric Mixed Effects Models for Unequally Sampled Noisy Curves, *Biometrics*, 57, 253–259.
 - [98] Rosa, G.J.M., Gianola, D. and Padovani, C.R. (2004), Bayesian Longitudinal Data Analysis with Mixed Models and Thick-tailed Distributions using MCMC, *J. Appl. Stat.* 31, 855–873.
 - [99] Roy, A. (2006), Estimating Correlation Coefficient Between two Variables with Repeated Observations using Mixed Effects Model, *Biom J.*, 48, 286–301.
 - [100] Rudemo, M. (1982), Empirical Choice of Histograms and Kernel Density Estimators, *Scand. J. Statist.*, 9, 65–78.
 - [101] Ruppert, D., Wand, M.P. and Carroll, R.J. (2003), *Semiparametric Regression*, Cambridge University Press, Cambridge.
 - [102] Saleh, A.K.Md.E., Arashi, M. and Kibria, B.M.G. (2019), *Theory of Ridge Regression Estimation with Applications*. Wiley, New York.
 - [103] Sammel, M., Lin, X. and Ryan, L. (1999), Multivariate Linear Mixed Models for Multiple Outcomes, *Statist. Med.*, 18, 2479–2492.
 - [104] Schumaker, L.L. (1981), *Spline Functions*, Wiley, New York.
 - [105] Schwarz, G. (1978), Estimating the Dimension of a Model, *Ann. Statist.*, 6, 461–464.
 - [106] Sengul, T.K., Stoffer, D.S. and Day, N.L. (2007), A Residuals-Based Transition Model for Longitudinal Analysis with Estimation in the Presence of Missing Data, *Statist. Med.*, 26, 3330–3341.

- [107] Shah, A., Laird, N. and Schoenfeld, D. (1997), A Random-effects Model for Multiple Characteristics with Possibly Missing Data, *J. Roy. Statist. Soc.*, 92, 775–779.
- [108] Simon, I., Barnett, J., Hannett, N., Harbison, C. T., Rinaldi, N.J., Volkert, T.L., Wyrick, J.J., Zeitlinger, J., Gifford, D.K., Jaakola, T. S., and Young, R.A. (2001), Serial Regulation of Transcriptional Regulators in the Yeast Cell Cycle, *Cell*, 106, 697–708.
- [109] Sinha, S.K. and Sattar, A. (2015), Inference in Semi-Parametric Spline Mixed Models for Longitudinal Data, *METRON*, 73, 377–395.
- [110] Song, X., Davidian, M. and Tsiatis, A.A. (2002), An Estimator for the Proportional Hazards Model with Multiple Longitudinal Covariates Measured with Error, *Biostatistics*, 3, 511–528.
- [111] Song, P.X.K., Zhang, P. and Qu, A. (2007), Maximum Likelihood Inference in Robust Linear Mixed-Effects Models using Multivariate t Distributions, *Statist. Sinica*, 17, 929–943.
- [112] Speckman, P. (1988), Kernel Smoothing in Partial Linear Models, *J. Roy. Statist. Soc.*, B, 50, 413–436.
- [113] Spellman, P.T., Sherlock, G., Zhang, M.Q., Iyer, V.R., Anders, K., Eisen, M.B., Brown, P.O., Botstein, D., and Futch, B. (1998), Comprehensive Identification of Cell Cycle-Regulated Genes of the Yeast *Saccharomyces Cerevisiae* by Microarray Hybridization, *Molecular Biology of Cell*, 9, 3273–3297.
- [114] Stiratelli, R., Laird, N.M. and Ware, J.H. (1984), Random Effects Models for Serial Observations with Binary Response, *Biometrics*, 40, 961–971.
- [115] Stone, C. (1982), Optimal Global Rates of Convergence for Nonparametric Regression, *Ann. Statist.*, 10, 1040–1053.
- [116] Sutradhar, B.C. and Das, K. (1999), On the Efficiency of Regression Estimators in Generalized Linear Models for Longitudinal Data, *Biometrika*, 86, 459–465.
- [117] Sutradhar, B.C. and Farrell, P.J. (2004), Analyzing Multivariate Longitudinal Binary Data: A Generalized Estimating Equations Approach, *Can. J. Statist.*, 32, 39–55.
- [118] Tanner, M.A. (1993), *Tools for Statistical Inference: Methods for the Exploration*, Springer.
- [119] Tang, Q.G. and Cheng, L.S. (2008), M-Estimation and B-spline Approximation for Varying Coefficient Models with Longitudinal Data, *J. Nonp. Statist.*, 20, 611–625.

-
- [120] Tang, Y., Wang, H.J. and Zhu, Z. (2013), Variable Selection in Quantile Varying Coefficient Models with Longitudinal Data, *Comp. Statist. Data Anal.*, 57, 435–449.
- [121] Tibshirani, R. (1996), Regression Shrinkage and Selection via the Lasso. *J. Roy. Statist. Soc.*, B, 58, 267–288.
- [122] Tikhonov, A. (1943), On the Stability of Inverse Problems, *Proceedings of the USSR Academy of Sciences*, 39, 267–88.
- [123] Tutz, G. (2005), Modelling of Repeated Ordered Measurements by Isotonic Sequential Regression, *Statist. Model.*, 5, 269- -87.
- [124] Verbeke, G. and Molenberghs, G. (2009), *Linear Mixed Models for Longitudinal Data*, Springer, New York.
- [125] Wang, W.L. (2013), Multivariate t Linear Mixed Models for Irregularly Observed Multiple Repeated Measures with Missing Outcomes, *Biom. J.*, 55, 554–571.
- [126] Wang N., Carroll R. and Lin X.H. (2005), Efficient Semiparametric Marginal Estimation for Longitudinal/Clustered Data, *J. Amer. Statist. Assoc.*, 100, 147–157.
- [127] Wang, W.L. and Castro, L.M. (2018), Bayesian Inference on Multivariate- t Nonlinear Mixed-effects Models for Multiple Longitudinal Data with Missing Values, *Stat. Interface.*, 11, 251–264.
- [128] Wang, W.L. and Fan, T.H. (2010), ECM-based maximum likelihood inference for multivariate linear mixed models with autoregressive errors, *Computat. Stat. Data Anal.* 54, 1328–1341.
- [129] Wang, W.L. (2017), Mixture of Multivariate t Linear Mixed Models for Multi-Outcome Longitudinal Data with Heterogeneity, *Statist. Sinica*, 27, 733–760.
- [130] Wang, W.L. and Fan, (2011), T.H. Estimation in Multivariate t Linear Mixed Models for Multiple Longitudinal Data, *Statist. Sinica*, 21, 1857–1880.
- [131] Wang, W.L. and Fan, T.H. (2012), Bayesian Analysis of Multivariate t Linear Mixed Models using a Combination of IBF and Gibbs Samplers, *J Multivariate Anal.*, 105, 300–310.
- [132] Wang, W.L. and Lin, T.I. (2015), Bayesian Analysis of Multivariate t Linear Mixed Models with Missing Responses at Random, *J. Stat. Comput. Simul.*, 85, 3594–3612.

- [133] Wang, W.L. and Lin, T.I. (2014), Multivariate t Nonlinear Mixed-effects Models for Multi-outcome Longitudinal Data with Missing Values, *Statist. Med.*, 33, 3029–3046.
- [134] Wang, L. and Qu, A. (2009), Consistent Model Selection and Data-Driven Smooth Tests for Longitudinal Data in the Estimating Equations Approach, *J. Roy. Statist. Soc., B*, 71, 177–190.
- [135] Wang, L., Zhou, J. and Qu, A. (2012), Penalized Generalized Estimating Equations for High-Dimensional Longitudinal Data Analysis, *Biometrics*, 68(2), 353–360.
- [136] Wedderburn, R.W.M. (1974), Quasi-Likelihood Functions, Generalized Linear Models, and the Gauss-Newton Method, *Biometrika*, 61, 439–47.
- [137] Wu, C.O. and Chiang, C.T. (2000), Kernel Smoothing on Varying Coefficient Models with Longitudinal Dependent Variable, *Statist. Sinica*, 10, 433–456.
- [138] Wu, T.T., and Lange, K. (2008), Coordinate Descent Algorithms for Lasso Penalized Regression, *Annal. Appl. Statist.*, 2(1), 224–244.
- [139] Wu, H. and Liang, H. (2004), Backfitting Random Varying-Coefficient Models with Time-Dependent Smoothing Covariates, *Scand. J. Statist.*, 31, 3–19.
- [140] Wu, H. and Zhang, T. (2006), *Nonparametric Regression Methods for Longitudinal Data Analysis*, Wiley, New York.
- [141] Xu, P.R., Fu, W. and Zhu, L.X. (2013), Shrinkage Estimation Analysis of Correlated Binary Data with a Diverging Number of Parameters, *Science China Math.*, 56, 359–377.
- [142] Xue, L. (2009), Consistent Variable Selection in Additive Models. *Statist. Sinica*, 19, 1281–1296.
- [143] Xue, L., Qu, A. and Zhou, J. (2010), Consistent Model Selection for Marginal Generalized Additive Model for Correlated Data. *J. Amer. Statist. Assoc.*, 105, 1518–1530.
- [144] Xue, L.G. and Zhu, L.X. (2007), Empirical Likelihood for a Varying Coefficient Model with Longitudinal Data, *J. Amer. Statist. Assoc.*, 102, 642–654.
- [145] Xue, L., Qu, A., and Zhou, J. (2010), Consistent model selection for marginal generalized additive model for correlated data, *J. Amer. Statist. Assoc.*, 105, 1518–1530.

-
- [146] Yang, Y., Li, G. and Peng, H. (2014), Empirical Likelihood of Varying Coefficient Errors-in-Variables Models with Longitudinal Data, *J. Mult. Anal.*, 127, 1–18.
- [147] Zeger, S.L. and Diggle, P.J. (1994), Semi-Parametric Models for Longitudinal Data with Application to CD4 Cell Numbers in HIV Seroconverters. *Biometrics*, 50, 689–99.
- [148] Zeger, S.L. and Karim, M.R. (1991), Generalized Linear Models with Random Effects: A Gibbs sampling approach, *J. Amer. Statist. Assoc.*, 86, 79–86.
- [149] Zhang, B. and Horvath, S. (2006), Ridge Regression Based Hybrid Genetic Algorithms for Multi-Locus Quantitative Trait Mapping, *Int. J. Bioinformatics Res. App.*, 1, 261–272.
- [150] Zhang, C.H. (2010). Nearly Unbiased Variable Selection Under Minimax Concave Penalty, *Annal. statist.*, 38(2), 894–942.
- [151] Zhang, D. (2004), Generalized Linear Mixed Models with Varying Coefficients for Longitudinal Data, *Biometrics*, 60, 8–15.
- [152] Zhang, Y., and Li, R. (2011), Iterative Conditional Maximization Algorithm for Nonconcave Penalized Likelihood, *Nonpar. Statist. Mixture Models*, 336–351.
- [153] Zhao, L.P. and Prentice, R.L. (1990), Correlated Binary Regression Using a Quadratic Exponential Model, *Biometrika*, 77, 642–648.
- [154] Ziegler, A. and Vens, M. (2010). Generalized Estimating Equations: Notes on the Choice of the Working Correlation Matrix, *Methods Inf. Med.*, 49, 421–425.
- [155] Zou, H. (2006), The Adaptive Lasso and its Oracle Properties, *J. Amer. Statist. Assoc.*, 101, 1418–1429.
- [156] Zou, H. and Hastie, T. (2005), Regularization and Variable Selection via the Elasticnet, *J. Roy. Statist. Soc., B*, 67, 301–320.
- [157] Zou, H., and Li, R. (2008), One-step Sparse Estimates in Nonconcave Penalized Likelihood Models. *Annal. Statist.*, 36(4), 1509–1533.

پیوست آ

تعاریف و ملزومات

در این پیوست در بخش‌های مختلف ملزومات و تعاریف استفاده شده در فصل‌های این رساله، گردآوری شده‌اند.

۱.۰.۱ معیارهای نیکویی برازش

در رساله حاضر، سنجش نیکویی برازش مدل‌های مختلف توسط معیارهای متعدد انجام می‌شود؛ از جمله معیار اطلاع آکائیک (AIC)، معیار اطلاع بیزی (BIC) و میانگین توان‌های دوم خطای پیشگویی^۱ ($MSPE$). همچنین برای تشخیص کارایی برآوردهای مختلف از معیارهایی مانند اریبی برآوردها ($Bias$)، انحراف استاندارد برآوردها (SD)، میانگین توان دوم خطای (Mse) برآوردها و میانگین توان دوم خطا (MSE) برای بردار برآوردها استفاده می‌شود. در ادامه هر یک از این معیارها را تعریف می‌کنیم.

تعریف ۱.۰.۱.۰ AIC : این معیار که توسط آکائیک (۱۹۷۳) معرفی شد به صورت

$$AIC = 2m - 2\ell_{\max}$$

تعریف می‌شود که در آن $\ell_{\max}$ مقدار بیشینه تابع درستنمایی است و m تعداد پارامترهای مدل است.

^۱Mean squared prediction error

تعریف ۲.۱.۰. BIC : این معیار که توسط شوارتز (۱۹۷۸) معرفی شد به صورت

$$BIC = m \log n - 2\ell_{\max}$$

تعریف می‌شود که در آن $\ell_{\max}$ و m مشابه تعریف AIC بوده و n حجم نمونه است.

تعریف ۳.۱.۰. $MSPE$: این معیار مدل را بر اساس کارایی آن در پیش‌بینی مقدارهای جدید (مشاهده‌نشده) ارزیابی می‌کند. به عبارت دیگر بر داده‌هایی تکیه دارد که مشاهده شده‌اند ولی در هنگام ساختن مدل به کار گرفته نمی‌شوند. این داده‌ها به منظور بررسی و سنجش کارایی مدل برای پیش‌بینی داده‌های جدید به کار می‌روند. بدین منظور داده‌ها به دو زیرمجموعه مکمل افراز می‌شوند. بر روی یکی از آن زیرمجموعه‌ها (داده‌های آموزشی) تحلیل انجام می‌شود و بر اساس نتایج حاصل، داده‌های مجموعه دیگر (داده‌های آزمایش) پیشگویی می‌شوند. برای افزایش دقت، عمل پیشگویی چندین بار با افرازهای مختلف انجام و از خطاهای پیشگویی میانگین گرفته می‌شود. اگر فرض کنیم داده‌ها به S زیرمجموعه افراز شده‌اند. از این S زیرمجموعه، هر بار یکی به عنوان داده‌های آزمایش و $S - 1$ تای دیگر برای آموزش بکار می‌روند. این روال S بار تکرار می‌شود و همه داده‌ها یک بار برای آموزش و یک بار برای آزمایش به کار می‌روند. $MSPE$ به صورت

$$MSPE = \frac{1}{S} \sum_{s=1}^S (\mathbf{Y}_{test}^{(-s)} - \hat{\mathbf{Y}}_{test}^{(-s)})^\top (\mathbf{Y}_{test}^{(-s)} - \hat{\mathbf{Y}}_{test}^{(-s)}),$$

تعریف می‌شود که در آن $\mathbf{Y}_{test}^{(-s)}$ داده‌های آزمایشی در افراز s ام و $\hat{\mathbf{Y}}_{test}^{(-i)}$ مقادیر پیشگویی بر اساس نتایج کنار گذاشتن داده‌های افراز s ام هستند.

تعریف ۴.۱.۰. $Bias$: اریبی یک برآوردگر همان اختلاف بین امید ریاضی آن برآوردگر و مقدار واقعی پارامتر برآورد شده است. اریبی برآوردگر $\hat{\beta}$ به صورت

$$Bias(\hat{\beta}) = E(\hat{\beta}) - \beta$$

تعریف می‌شود. در مطالعات شبیه‌سازی اگر S مجموعه تولید شده داشته باشیم، $E(\hat{\beta})$ به صورت

$$\hat{E}(\hat{\beta}) = \frac{1}{S} \sum_{s=1}^S \hat{\beta}^{(s)}$$

تقریب می‌شود که $\hat{\beta}^{(s)}$ برآورد β بر اساس مجموعه تولید شده s ام است.

تعریف ۵.۱.۰. SD : در مطالعات شبیه‌سازی مقاله حاضر، SD برآوردگر $\hat{\beta}$ بر اساس S تکرار از شبیه‌سازی به صورت

$$SD = \sqrt{\frac{\sum_{s=1}^S (\hat{\beta}^s - E(\hat{\beta}))^2}{S - 1}}$$

تعریف می‌شود که در آن $\hat{E}(\hat{\beta}) = \frac{1}{S} \sum_{s=1}^S \hat{\beta}^{(s)}$ و $\hat{\beta}^s$ برآورد β در s -امین مجموعه داده تولید شده است.

در مجموعه داده‌های واقعی، مشابه روش بیان شده در تعریف $MSPE$ ، داده‌ها به S زیرمجموعه افراز می‌شوند. هر بار با کنار گذاشتن داده‌های آزمایش، برآوردها را بر اساس داده‌های آموزشی به دست می‌آوریم. بدین ترتیب تعداد S برآورد برای پارامتر مدل خواهیم داشت. با محاسبه انحراف استاندارد مجموعه برآوردهای حاصل، معیار SD مورد نظر به دست می‌آید که معمولاً با عنوان انحراف استاندارد نمونه‌ای از آن نام می‌بریم.

تعریف ۶.۱.۰. Mse : در مطالعات شبیه‌سازی مقاله حاضر، Mse برآوردگر $\hat{\beta}$ بر اساس S تکرار از شبیه‌سازی به صورت

$$Mse = \frac{1}{S} \sum_{s=1}^S (\hat{\beta}^s - \beta)^2$$

تعریف می‌شود که در آن $\hat{\beta}^s$ برآورد β در s -امین مجموعه داده تولید شده است.

تعریف ۷.۱.۰. MSE : در مطالعات شبیه‌سازی مقاله حاضر، MSE بر اساس S تکرار از شبیه‌سازی به صورت

$$MSE = \frac{1}{S} \sum_{s=1}^S \|\hat{\beta}^s - \beta\| = \frac{1}{S} \sum_{s=1}^S (\hat{\beta}^s - \beta)^\top (\hat{\beta}^s - \beta)$$

تعریف می‌شود که در آن $\hat{\beta}^s$ برآورد بردار پارامتر β در s -امین مجموعه داده تولید شده است.

۲.۰ تعاریف و نامساوی‌ها

تعریف ۱.۲.۰. نرم اقلیدسی ($\|\cdot\|$): یک تابع حقیقی است که بر فضای برداری $x = (x_1, \dots, x_n)^\top$ به صورت

$$\|x\| = \left(\sum_{i=1}^n |x_i|^2 \right)^{1/2}$$

تعریف می‌شود.

تعریف ۲.۲.۰. پایه: در جبر خطی منظور از یک پایه برای یک فضای برداری، مجموعه‌ای از بردارهای موجود در آن فضا بوده که دارای دو خاصیت زیر است:

۱. هیچ کدام از بردارهای آن مجموعه را نتوان به صورت ترکیب خطی از بقیه بردارها نوشت، به عبارت دیگر، بردارها به طور خطی از هم مستقل باشند.
۲. هر بردار دیگر در آن فضای برداری را بتوان از ترکیب خطی این بردارها به دست آورد.

تعریف ۳.۲.آ. همگرایی در توزیع: دنباله متغیرهای تصادفی $\{X_n; n \geq 1\}$ در توزیع به متغیر تصادفی X همگراست اگر برای هر $x \in \mathbb{R}$ که تابع F در آن پیوسته است، داشته باشیم

$$\lim_{n \rightarrow \infty} F_n(x) = F(x),$$

که در آن F_n تابع توزیع تجمعی $\{X_n; n \geq 1\}$ و F تابع توزیع تجمعی X است. این نوع همگرایی را به صورت $X_n \xrightarrow{d} X$ نشان می‌دهند.

تعریف ۴.۲.آ. همگرایی در احتمال: دنباله متغیرهای تصادفی $\{X_n\}$ در احتمال به متغیر تصادفی X همگراست اگر برای هر $\varepsilon > 0$ ، داشته باشیم

$$\lim_{n \rightarrow \infty} P(|X_n - X| > \varepsilon) = 0.$$

این نوع همگرایی را به صورت $X_n \xrightarrow{P} X$ نشان می‌دهند.

تعریف ۵.۲.آ. سازگاری: دنباله برآوردگرهای $\{T_n; n \geq 1\}$ برای $a(\theta)$ سازگار است، اگر $\{T_n\}$ در احتمال به $a(\theta)$ همگرا باشد.

تعریف ۶.۲.آ. گویم برآوردگر $\hat{\theta}_n$ یک برآوردگر $\sqrt{n}$ -سازگار برای θ است، اگر رابطه

$$\hat{\theta}_n - \theta = O_p(n^{-1/2}),$$

برقرار باشد.

تعریف ۷.۲.آ. گویم $f(n)$ از مرتبه $g(n)$ است یا $f(n) = O(g(n))$ ، اگر وقتی $n \rightarrow \infty$ ، آنگاه $\frac{f(n)}{g(n)}$ متناهی باشد. به عبارت دیگر، عدد حقیقی مثبت M و عدد حقیقی n_0 موجود باشند به طوری که

$$\forall n > n_0, \quad |f(n)| \leq M|g(n)|.$$

یعنی برای مقادیر به اندازه کافی بزرگ n ، قدرمطلق $f(n)$ حداکثر M برابر قدرمطلق $g(n)$ باشد.

تعریف ۸.۲.آ. گویم $f(n)$ از مرتبه کوچک‌تر $g(n)$ است یا $f(n) = o(g(n))$ ، اگر وقتی $n \rightarrow \infty$ ، آنگاه داشته باشیم

$$\lim_{n \rightarrow \infty} \frac{f(n)}{g(n)} = 0.$$

به عبارت دیگر $g(n)$ بسیار سریع‌تر از $f(n)$ رشد می‌کند. بنابراین o حالت قوی‌تر O است. هر عبارتی که $o(g)$ باشد، $O(g)$ نیز هست، اما عکس آن برقرار نیست.

تعریف ۹.۲.آ. فرض کنید $\{X_n\}$ یک دنباله از متغیرهای تصادفی و $\{a_n\}$ یک دنباله از ثابت‌های متناظر باشد. نماد $X_n = O_p(a_n)$ بدین معنی است که مقادیر $\frac{X_n}{a_n}$ در احتمال کراندار است، یعنی برای هر $\varepsilon > 0$ ، اعداد متناهی $M, N > 0$ وجود دارند به طوری که

$$\forall n > N, \quad \mathbb{P}\left(\left|\frac{X_n}{a_n}\right| > M\right) < \varepsilon.$$

تعریف ۱۰.۲.۰. بسط تیلور: فرض کنید که تابع $f(x)$ بر بازه $[a, b]$ تعریف شده و $n+1$ بار مشتق‌پذیر باشد. نقطه $x_0 \in (a, b)$ را در نظر بگیرید. برای هر $x \in (a, b)$ که $x = x_0 + \Delta x$ ، عددی مانند ξ بین دو نقطه x و x_0 وجود دارد، به‌طوری‌که تساوی

$$f(x_0 + \Delta x) = f(x_0) + \Delta x f'(x_0) + \frac{\Delta x^2}{2!} f''(x_0) + \dots + \frac{\Delta x^n}{n!} f^{(n)}(x_0) + R_n(x_0 + \Delta x),$$

برقرار است. که در آن

$$R_n(x_0 + \Delta x) = \frac{\Delta x^{n+1}}{(n+1)!} f^{(n+1)}(\xi), \quad x_0 \leq \xi \leq x.$$

در این روابط منظور از $f^{(n)}$ مشتق مرتبه n ام است. در عمل معمولاً تعداد اندکی از جملات اول سری برای تقریب استفاده می‌شود. در این حالت اگر مقدار Δx به قدر کافی کوچک باشد، تقریب مناسبی از تابع در همسایگی x_0 حاصل می‌شود. در صورتی که تقریب تابع با سری تیلور برای n جمله اول انجام شود، مانده سری یا خطای تقریب متناسب با Δx^n خواهد بود و آن را با $O(\Delta x^n)$ نشان می‌دهند. یعنی

$$f(x_0 + \Delta x) = f(x_0) + \Delta x f'(x_0) + \frac{\Delta x^2}{2!} f''(x_0) + \dots + \frac{\Delta x^n}{n!} f^{(n)}(x_0) + O(\Delta x^n).$$

تعریف ۱۱.۲.۰. عملگر بردارسازی^۲: در ریاضیات، به ویژه در جبر خطی و نظریه ماتریس، بردارسازی یک ماتریس یک تبدیل خطی است که ماتریس را به یک بردار ستونی تبدیل می‌کند. به عبارت دیگر بردارسازی ماتریس A با بعد $m \times n$ ، که با نماد $\text{vec}(A)$ نمایش داده می‌شود، یک بردار ستونی $mn \times 1$ بعدی است که با انباشتن ستون‌های ماتریس A بر روی یک‌دیگر به‌دست می‌آید. به $\text{vec}(\cdot)$ عملگر بردارسازی گویند و به صورت

$$\text{vec}(A) = (a_{1,1}, \dots, a_{m,1}, a_{1,2}, \dots, a_{m,2}, \dots, a_{1,n}, \dots, a_{m,n})^\top$$

تعریف می‌شود، به‌طوری‌که $a_{i,j}$ عضو مربوط به سطر i ام و ستون j ام ماتریس A است.

تعریف ۱۲.۲.۰. عملگر نیمه‌بردارسازی^۳: برای ماتریس متقارن A ، بردار $\text{vec}(A)$ حاوی اطلاعات غیرضروری است؛ زیرا به ازای $i \neq j$ عضو a_{ij} برابر عضو a_{ji} است. اگر A یک ماتریس متقارن $n \times n$ بعدی باشد، عملگر نیمه‌بردارسازی، $\text{vech}(A)$ ، با بردارسازی مثلث پایین ماتریس A یک بردار $\frac{n(n+1)}{2} \times 1$ بعدی را نتیجه می‌دهد. به عبارت دیگر

$$\text{vech}(A) = (a_{1,1}, \dots, a_{n,1}, a_{2,2}, \dots, a_{n,2}, \dots, a_{n-1,n-1}, a_{n,n-1}, a_{n,n})^\top.$$

²Vectorization

³Half-vectorization

تعریف ۱۳.۲. ضرب کرونکر: ضرب کرونکر ماتریس‌های $A \in \mathbb{R}^{m_A \times n_A}$ و $B \in \mathbb{R}^{m_B \times n_B}$ به صورت $A \otimes B$ نمایش داده می‌شود و برابر است با:

$$A \otimes B = \begin{pmatrix} a_{11}B & \dots & a_{1n_A}B \\ \vdots & \ddots & \vdots \\ a_{m_A1}B & \dots & a_{m_An_A}B \end{pmatrix}$$

هر $a_{ij}B$ یک بلوک در سائز $m_B \times n_B$ است، بنابراین $A \otimes B$ یک ماتریس $m_A m_B \times n_A n_B$ بعدی است.

تعریف ۱۴.۲. قضیه لیندبرگ^۴: متغیرهای تصادفی X_{ij} را در نظر بگیرید که یک آرایه مثلی به صورت

$$\begin{matrix} X_{11} & X_{12} & \dots & X_{1n_1} \\ X_{21} & X_{22} & \dots & \dots & X_{2n_2} \\ X_{31} & X_{32} & \dots & \dots & \dots & X_{3n_3} \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \end{matrix}$$

را تشکیل می‌دهند. فرض می‌شود که X_{ij} ها در شروط زیر که شروط آرایه مثلی نامیده می‌شود، صدق می‌کنند.

۱. برای هر i, n_i متغیر تصادفی $X_{i1}, X_{i2}, \dots, X_{in_i}$ در سطر i ام، دو به دو از هم مستقل‌اند.

۲. به ازای تمام i ها و j ها، $E(X_{ij}) = 0$.

۳. به ازای تمام i ها، $\sum_j E(X_{ij}^2) = 1$.

دنباله تصادفی که از مجموع عناصر سطری آرایه مثلی به دست می‌آید را به صورت $S_i = \sum_j X_{ij}$ تعریف می‌کنیم. فرض می‌کنیم $L(S_i)$ تابع توزیع این دنباله باشد. به علاوه فرض می‌کنیم X_{ij} ها در شرط لیندبرگ

$$\forall \epsilon > 0, \quad \lim_{i \rightarrow \infty} \sum_{j=1}^{n_i} E(X_{ij}^2 I(|X_{ij}| > \epsilon)) = 0$$

صدق می‌کنند. آنگاه داریم

$$\lim_{i \rightarrow \infty} L(S_i) \rightarrow \mathcal{N}(0, 1).$$

تعریف ۱۵.۲. نامساوی مارکف: فرض کنید X_n یک متغیر تصادفی با مقادیر مثبت باشد، یعنی $\mathbb{P}(X_n \geq 0) = 1$. آنگاه این نامساوی بیان می‌کند که

$$\mathbb{P}(X_n \geq c) \leq \frac{E(X_n)}{c}.$$

این نامساوی یک کران بالا برای مقدار احتمال ارائه می‌دهد.

⁴Lindeberg

تعریف ۱۶.۲. نامساوی جنسن^۵: فرض کنید g یک تابع محدب در بازه (a, b) باشد و امید ریاضی متغیر تصادفی X_n با ویژگی $\mathbb{P}(X_n \in (a, b)) = 1$ وجود داشته باشد. آنگاه نامساوی جنسن به صورت زیر برقرار است.

$$g(E(X_n)) \leq E(g(X_n)).$$

تعریف ۱۷.۲. نامساوی کوشی-شوارتز^۶: این نامساوی بیان می‌کند که برای هر دو متغیر تصادفی X_n و Y_n داریم

$$E^2(X_n Y_n) \leq E(X_n^2) E(Y_n^2).$$

هم‌چنین با گرفتن ریشه دوم طرفین نامساوی فوق داریم

$$|E(X_n Y_n)| \leq E^{1/2}(X_n^2) E^{1/2}(Y_n^2).$$

تعریف ۱۸.۲. الگوریتم نیوتن-رافسون: در آنالیز عددی این روش یک الگوریتم ریشه‌یابی است که تقریب‌های خوبی در نزدیکی ریشه یک تابع می‌زند. در پایه‌ای‌ترین حالت، الگوریتم نیوتن برای یک تابعی چون f با متغیر x و با مشتق f' به همراه حدس اولیه x_0 به کار می‌رود. شکل عمومی الگوریتم به شرح زیر می‌باشد:

$$x_1 = x_0 - \frac{f(x_0)}{f'(x_0)}$$

اگر تابع حدس کافی و دقیقی را برآورد سازد و هم‌چنین حدس اولیه نزدیک به ریشه تابع مفروض باشد آنگاه x_1 تقریب بهتری نسبت به x_0 به حساب می‌آید. چرا که با احتساب همگرایی جواب‌ها، هر تقریب نسبت به تقریب قبل از خودش از دقت بالاتری برخوردار بوده و به ریشه تابع نزدیک‌تر است.

⁵Jensen

⁶Cauchy-Schwarz

پیوست ب

اثبات قضایا

در این پیوست اثبات قضایایی که در متن رساله (فصل‌های گذشته) آمده، به تفصیل آورده شده است.

ب.۱ اثبات قضایای فصل ۱

ب.۱.۱ اثبات قضیه ۱.۵.۱

برای برآوردگر هسته‌ای، $\hat{\beta}_K$ ، داریم

$$\sqrt{n}(\hat{\beta}_K - \beta) = D_n^{-1}\{\sqrt{n}C_n\} + o_p(1),$$

به طوری که

$$D_n = \frac{1}{n} \sum_{i=1}^n \tilde{X}_i^\top V_i^{-1} \tilde{X}_i, \quad \text{و} \quad C_n = \frac{1}{n} \sum_{i=1}^n \tilde{X}_i^\top V_i^{-1} [Y_i - (X_i^\top \beta + \hat{f}(t_i))].$$

فرض کنید $D = \lim_{n \rightarrow \infty} D_n = E\{\tilde{X}^\top V^{-1} \tilde{X}\}$. با اندکی محاسبات ساده می‌توان نشان داد
 $C_n = C_{1n} - C_{2n} + o_p(1)$ که در آن C_{1n} و C_{2n} به ترتیب به صورت

$$C_{1n} = \frac{1}{n} \sum_{i=1}^n \tilde{X}_i^\top V_i^{-1} [Y_i - \mu_i] = \frac{1}{n} \sum_{i=1}^n W_{1i}^\top [Y_i - \mu_i]$$

و

$$C_{\Psi n} = \frac{1}{n} \sum_{i=1}^n \tilde{X}_i^\top V_i^{-1} [\hat{f}(t_i, \beta) - f(t_i)]$$

تعریف می‌شوند، به طوری که $\mu_i = X_i^\top \beta + f(t_i)$ و $W_{\Psi i} = V_i^{-1} \tilde{X}_i$. مشابه اثبات صورت گرفته در لین و کارول (۲۰۰۲)، می‌توان نشان داد که

$$C_{\Psi n} = \frac{1}{n} \sum_{i=1}^n W_{\Psi i}^\top [Y_i - \mu_i] + \frac{h^{\frac{1}{2}}}{\sqrt{n}} E\{X^\top V^{-1} f''(t)\} + o_p(1),$$

به طوری که $W_{\Psi i} = \{W_{\Psi i1}, \dots, W_{\Psi im}\}^\top$ و

$$W_{\Psi ij} = \frac{\{\sum_{k=1}^m \sum_{l=1}^m E(\tilde{X}_k V^{kl} | t_l = t_{ij})\} f_j(t_{ij})}{\sum_{l=1}^m f_l(t_{ij})}.$$

بنابراین

$$\sqrt{n}(\hat{\beta}_K - \beta) = D^{-1} \frac{1}{\sqrt{n}} \sum_{i=1}^n (W_{\Psi i} - W_{\Psi i})(Y_i - \mu_i) + \sqrt{nh^{\frac{1}{2}}} b_K(\beta, f)/\sqrt{2} + o_p(1),$$

که $b_K(\beta, g) = D^{-1} E\{\tilde{X}^\top V^{-1} g''(t)\}$ با این فرض که $V_o = \text{Var}(Y|X, t)$ و همچنین $V_K = D^{-1} E\{(W_{\Psi 1} - W_{\Psi 1})^\top V_o (W_{\Psi 1} - W_{\Psi 1})\} D^{-1}$ ، نتیجه موردنظر حاصل می‌شود.

۲.۱.۱ اثبات قضیه ۳.۵.۱

می‌خواهیم ثابت کنیم $V_{BF} - V_K$ یک ماتریس نیمه معین مثبت است. جزء اصلی ماتریس V_{BF} به صورت $\varepsilon^* = Z^\top b + \varepsilon$ است که $X^\top V^{-1} (I - S) \varepsilon^*$ به صورت $X^\top (I - S)^\top V^{-1} (I - S) \varepsilon^*$ و یا به عبارتی به صورت $\tilde{X}^\top V^{-1} (I - S) \varepsilon^*$ است. برای مقایسه دو ماتریس V_{BF} و V_K کافی است که اجزای اصلی آن‌ها را مقایسه کنیم. اکنون نشان می‌دهیم که

$$\text{Cov}\{X^\top V^{-1} (I - S) \varepsilon^*\} \geq \text{Cov}\{\tilde{X}^\top V^{-1} (I - S) \varepsilon^*\}.$$

برای برآوردگر بازگشتی داریم

$$\begin{aligned} \text{Cov}\{X^\top V^{-1} (I - S) \varepsilon^*\} &= E\{X^\top V^{-1} (I - S) V_o (I - S)^\top V^{-1} X | t\} \\ &= \mu_X^\top(t) V^{-1} (I - S) V_o (I - S)^\top V^{-1} \mu_X(t) \\ &\quad + \text{tr}[V^{-1} (I - S) V_o (I - S)^\top V^{-1} \text{Cov}(X|t)]. \end{aligned} \quad (۱.ب)$$

جمله اول عبارت فوق نیمه معین مثبت است زیرا $V^{-1} (I - S) V_o (I - S)^\top V^{-1}$ یک ماتریس نیمه معین مثبت است. همچنین عبارت $\mu_X(t)$ مقداری غیرصفر است. از طرفی برای برآوردگر

هسته‌ای داریم

$$\begin{aligned} \text{Cov}\{\tilde{X}^\top V^{-1}(I-S)\varepsilon^*\} &= E\left\{\tilde{X}^\top V^{-1}(I-S)V_o(I-S)^\top V^{-1}\tilde{X}|t\right\} \quad (\text{ب.۲}) \\ &= E(\tilde{X}|t)^\top V^{-1}(I-S)V_o(I-S)^\top V^{-1}E(\tilde{X}|t) \\ &\quad + \text{tr}\left[V^{-1}(I-S)V_o(I-S)^\top V^{-1}\text{tr}(\tilde{X}|t)\right]. \end{aligned}$$

با توجه به این که

$$E(\tilde{X}_i|t_i) = E\{X_i - \mu_X(t_i)|t_i\} = 0, \quad \text{Cov}(\tilde{X}_i|t_i) = \text{Cov}\{X_i - \mu_X(t_i)|t_i\} = \text{Cov}(X_i|t_i).$$

بنابراین جمله اول در عبارت (ب.۲) برابر ۰ بوده و جملات دوم در عبارات (ب.۱) و (ب.۲) برابر هستند. لذا نتیجه می‌گیریم که $\text{Cov}\{X^\top V^{-1}(I-S)\varepsilon^*\} \geq \text{Cov}\{\tilde{X}^\top V^{-1}(I-S)\varepsilon^*\}$ ، و نتیجه مورد نظر حاصل می‌شود.

ب.۲ اثبات قضایای فصل ۲

ب.۲.۱ بردار رتبه و ماتریس اطلاع فشر

مشتق اول جزئی از لگاریتم تابع درستنمایی تاوانیده (۱۵.۲) نسبت به هر یک از اعضای پارامتر Θ را بردار رتبه (S_Θ) می‌گویند که اعضای آن شامل

$$\begin{aligned} S_\theta &= \sum_{i=1}^n \tilde{X}_i^\top V_i^{-1}(y_i - \tilde{X}_i\theta) - nE_n(\theta), \\ [S_\eta]_l &= -\frac{1}{\gamma} \sum_{i=1}^n \left\{ \text{tr}(V_i^{-1}\dot{V}_{il}) - (y_i - \tilde{X}_i\theta)^\top V_i^{-1}\dot{V}_{il}V_i^{-1}(y_i - \tilde{X}_i\theta) \right\}, \end{aligned}$$

است، که در آن $g = \frac{q(q+1)}{\gamma} + \frac{m(m+1)}{\gamma}$ ، $l = 1, \dots, g$

$$\dot{V}_{il} = \frac{\partial \nabla_i(D, \Sigma)}{\partial \eta_l} = \begin{cases} Z_i \frac{\partial D}{\partial \eta_l} Z_i^\top & \text{اگر } \eta_l = \text{vec}(D) \\ \frac{\partial \Sigma}{\partial \eta_l} & \text{اگر } \eta_l = \text{vec}(\Sigma) \end{cases}$$

و

$$\Delta_i(D, \Sigma) = (y_i - \tilde{X}_i\theta - Z_i b_i)^\top V_i^{-1}(y_i - \tilde{X}_i\theta - Z_i b_i).$$

همچنین امید مشتق دوم جزئی از لگاریتم تابع درستنمایی تاوانیده (۱۵.۲) نسبت به هر یک از اعضای پارامتر Θ را ماتریس اطلاع فشر (J_Θ) می‌گویند که اعضای آن به صورت

$$J_{\Theta\Theta} = \begin{bmatrix} J_{\theta\theta} & J_{\theta\eta} \\ J_{\eta\theta}^\top & J_{\eta\eta} \end{bmatrix},$$

تعیین می‌شود، به طوری که

$$\begin{aligned} J_{\theta\eta} &= 0, \\ J_{\theta\theta} &= \sum_{i=1}^n \tilde{X}_i^\top V_i^{-1} \tilde{X}_i - n \frac{\partial E_n(\theta)}{\partial \theta}, \\ [J_{\eta\eta}]_{ls} &= \frac{1}{2} \sum_{i=1}^n \left(\text{tr}(V_i^{-1} \overset{\circ}{V}_{il} V_i^{-1} \overset{\circ}{V}_{is}) - \text{tr}(V_i^{-1} \overset{\circ}{V}_{il}) \text{tr}(V_i^{-1} \overset{\circ}{V}_{is}) \right), \end{aligned}$$

و $l, s = 1, \dots, g$. به علاوه، سایر ماتریس‌های اطلاع مورد نیاز در اثبات قضایا به صورت‌های

$$J_{\theta\theta} = \begin{bmatrix} J_{\theta_1\theta_1} & J_{\theta_1\theta_2} \\ J_{\theta_2\theta_1}^\top & J_{\theta_2\theta_2} \end{bmatrix}, \quad J_{\theta_1\theta_1} = \begin{bmatrix} J_{\beta_1\beta_1} & J_{\beta_1\alpha} \\ J_{\alpha\beta_1}^\top & J_{\alpha\alpha} \end{bmatrix}$$

قابل تفکیک هستند و اعضای آن‌ها به صورت

$$\begin{aligned} J_{\theta_1\theta_1} &= \sum_{i=1}^n \tilde{X}_{i1}^\top V_{i1}^{-1} \tilde{X}_{i1} - n \frac{\partial E_n(\theta_1)}{\partial \theta_1}, \\ J_{\theta_2\theta_2} &= \sum_{i=1}^n \tilde{X}_{i2}^\top V_{i2}^{-1} \tilde{X}_{i2} - n \frac{\partial E_n(\theta_2)}{\partial \theta_2}, \\ J_{\theta_1\theta_2} &= \sum_{i=1}^n \tilde{X}_{i1}^\top V_{i1}^{-1} \tilde{X}_{i2}, \\ J_{\beta_1\beta_1} &= \sum_{i=1}^n X_{i1}^\top V_{i1}^{-1} X_{i1} - n \frac{\partial E_n(\beta_1)}{\partial \beta_1}, \\ J_{\beta_1\alpha} &= \sum_{i=1}^n X_{i1}^\top V_{i1}^{-1} B_i, \\ J_{\alpha\alpha} &= \sum_{i=1}^n B_i^\top V_{i1}^{-1} B_i; \end{aligned}$$

تعیین می‌شوند، که در آن‌ها $V_{i11}, V_{i12}, V_{i22}$ اعضای ماتریس

$$V_i = \begin{bmatrix} V_{i11} & V_{i12} \\ V_{i12} & V_{i22} \end{bmatrix}$$

هستند.

ب.۲.۲ اثبات قضیه ۱.۳.۲

فرض کنید $\alpha_n = \sqrt{p_n}(n^{-1/2} + a_n)$ و $\|U\| = C$ ، به طوری که C یک مقدار ثابت و به اندازه کافی بزرگ است. می‌خواهیم نشان دهیم که به ازای هر $\xi > 0$ یک مقدار ثابت C وجود دارد که

$$\mathbb{P}\left\{\inf_{\|U\|=C} \ell_c^{Pen}(\theta_{n^0} + \alpha_n U) \geq \ell_c^{Pen}(\theta_{n^0})\right\} \geq 1 - \xi.$$

یعنی با احتمال نزدیک به ۱، یک جواب محلی برای $\ell_c^{Pen}(\theta_{n^\circ})$ در مجموعه $\{\theta_{n^\circ} + \alpha_n U : \|U\| \leq C\}$ وجود دارد. با استفاده از $p_{\lambda_n}(\circ) = \circ$ داریم

$$\begin{aligned} D_n(U) &= \ell_c^{Pen}(\theta_{n^\circ} + \alpha_n U) - \ell_c^{Pen}(\theta_{n^\circ}) \\ &\geq \ell_c(\theta_{n^\circ} + \alpha_n U) - \ell_c(\theta_{n^\circ}) + n \sum_{i=1}^n \{p_{\lambda_n}(|\theta_{n^\circ, k} + \alpha_n U_k|) - p_{\lambda_n}(|\theta_{n^\circ, k}|)\} \\ &\equiv I + II. \end{aligned}$$

با اندکی محاسبات ساده داریم

$$\begin{aligned} I &= \frac{n\alpha_n^2}{2} U^\top \hat{\Gamma}_n U - \alpha_n U^\top \sum_{i=1}^n \tilde{X}_i \Sigma_i^{-1} (Y_i - \tilde{X}_i \theta - Z_i b_i) \\ &\equiv I_1 - I_2, \end{aligned}$$

که در آن $\hat{\Gamma}_n = n^{-1} \sum_{i=1}^n \tilde{X}_i^\top \hat{\Sigma}_i^{-1} \tilde{X}_i$ ، $\Gamma_n = n^{-1} \sum_{i=1}^n \tilde{X}_i^\top \Sigma_i^{-1} \tilde{X}_i$ و

$$\begin{aligned} I_1 &= \frac{n\alpha_n^2}{2} U^\top (\Gamma_n - \hat{\Gamma}_n + \hat{\Gamma}_n) U \\ &= \frac{n\alpha_n^2}{2} U^\top \hat{\Gamma}_n U + O_p(n^{-1/2}) n\alpha_n^2 \|U\|^2, \end{aligned}$$

$$I_2 = \alpha_n U^\top \sum_{i=1}^n \tilde{X}_i \Sigma_i^{-1} \epsilon_i.$$

با استفاده از شرط (۳.۸) می‌توان نشان داد که

$$\begin{aligned} |\alpha_n U^\top \sum_{i=1}^n \tilde{X}_i \Sigma_i^{-1} \epsilon_i| &\leq \alpha_n \|\alpha_n\| \|U\| \\ &= O_p(\alpha_n \sqrt{np_n}) \|U\| = O_p(\alpha_n^2 n) \|U\|. \end{aligned}$$

بنابراین $I_2 = O_p(\alpha_n^2 n) \|U\|$. مشابه اثبات فن و لی (۲۰۰۱) می‌توان نشان داد که

$$II = O_p(\sqrt{s_n} n \alpha_n a_n) \|U\| + O_p(n \alpha_n^2 b_n) \|U\|.$$

بر اساس نتایج بدست آمده، شرط (۳.۸) و اینکه $\|U\|$ به اندازه کافی بزرگ باشد، می‌توان دریافت که عبارت‌های I_1 و II تحت الشعاع عبارت I_1 می‌باشند که این عبارت مقداری مثبت است. بنابراین قسمت (i) قضیه ۱.۳.۲ اثبات می‌شود.

برای اثبات قسمت (ii) قضیه ۱.۳.۲، با استفاده از قضیه ۱۲.۷ در اسچوماکر (۱۹۸۱) تحت شرط (۱.۸) مقدار ثابتی مانند C وجود دارد به طوری که

$$\sup_{t \in [0, 1]} |f_\circ(t) - B(t)\alpha_\circ| \leq CK_n^{-m}.$$

حال با انتخاب $K_n \approx n^{1/(2r+1)}$ داریم

$$\begin{aligned} (\hat{f}(t) - f_\circ(t))^2 &= (B(t)\hat{\alpha}_n - f_\circ(t))^2 \simeq (B(t)\alpha_{n^\circ} - f_\circ(t))^2 \\ &\leq |B(t)\alpha_{n^\circ} - f_\circ(t)| |B(t)\alpha_{n^\circ} - f_\circ(t)| \\ &\leq \sup_{t \in [\circ, 1]} |B(t)\alpha_{n^\circ} - f_\circ(t)| \sup_{t \in [\circ, 1]} |B(t)\alpha_{n^\circ} - f_\circ(t)| \\ &\leq CL_n^{-r} CL_n^{-r} = C^2 n^{-2r/(2r+1)} = O_p(n^{-2r/(2r+1)}). \end{aligned}$$

بنابراین قسمت (ii) قضیه ۱.۳.۲ اثبات می‌شود.

۳.۲.۲ اثبات قضیه ۲.۳.۲

برای اثبات قسمت (i)، باید نشان داد که به ازای هر θ_{n1} ای که $\|\theta_{n1} - \theta_{n^\circ 1}\| = O_p(\sqrt{p_n/n})$ و هر مقدار ثابت C ، $(\theta_{n1}^\top, 0)^\top$ جواب بهینه تابع درستنمایی تاوانیده است. به عبارت دیگر

$$\ell_c^{Pen}\{(\theta_{n1}^\top, 0)^\top\} = \max_{\|\theta_{n2}\| \leq C\sqrt{p_n/n}} \ell_c^{Pen}\{(\theta_{n1}^\top, \theta_{n2}^\top)^\top\}.$$

برای اثبات عبارت فوق کافی است ثابت کنیم که

$$\begin{aligned} \frac{\partial \ell_c^{Pen}\{(\theta_n)\}}{\partial \theta_{n,k}}, &< 0 \quad \forall \circ < \theta_{n,k} < \xi_n; \\ \frac{\partial \ell_c^{Pen}\{(\theta_n)\}}{\partial \theta_{n,k}}, &> 0 \quad \forall -\xi_n < \theta_{n,k} < \circ. \end{aligned}$$

به ازای هر $\theta_{n,k} \neq \circ$ و $k = s_n + 1, \dots, p_n$ برقرار است، به طوری که $\xi_n = C\sqrt{p_n/n}$. واضح است که

$$\frac{\partial \ell_c^{Pen}\{(\theta_n)\}}{\partial \theta_{n,k}} = \frac{\partial \ell_c\{(\theta_n)\}}{\partial \theta_{n,k}} - np'_{\lambda_n}(|\theta_{n,k}|) \operatorname{sgn}(\theta_{n,k}),$$

و مشابه فن و لی (۲۰۰۱) می‌توان نشان که

$$\frac{\partial \ell_c\{(\theta_n)\}}{\partial \theta_{n,k}} = O_p(\sqrt{np_n}).$$

بنابراین می‌توان عبارت

$$\frac{\partial \ell_c^{Pen}\{(\theta_n)\}}{\partial \theta_{n,k}} = n\lambda_n \left\{ -\frac{p'_{\lambda_n}(|\theta_{n,k}|)}{\lambda_n} \operatorname{sgn}(\theta_{n,k}) + O_p\left(\frac{\sqrt{p_n/n}}{\lambda_n}\right) \right\}$$

را نتیجه گرفت. از آنجایی که $\frac{\sqrt{p_n/n}}{\lambda_n} \rightarrow \circ$ و $\liminf_{\theta \rightarrow \circ_+} p'_{\lambda_n}(\theta)/\lambda_n > \circ$ علامت $\theta_{n,k}$ تعیین‌کننده علامت $\frac{\partial \ell_c^{Pen}\{(\theta_n)\}}{\partial \theta_{n,k}}$ است. بنابراین اثبات کامل می‌شود.

بخش همگرایی در احتمال و همگرایی در توزیع مستقیماً از طریق قانون ضعیف اعداد بزرگ، قضیه حد مرکزی و قضیه اسلاتسکی اثبات می‌شود.

ب.۳ اثبات قضایای فصل ۳

فرض کنید C یک مقدار ثابت و مثبت باشد، که در جاهای مختلف این بخش ممکن است مقادیر مختلفی داشته باشد. در این بخش $\bar{S}_n(\theta)$ را به صورت $\bar{S}_n(\theta) = (\bar{S}_{n1}(\theta), \dots, \bar{S}_{np_n}(\theta))^\top$ تعریف می‌شود به طوری که $\bar{S}_{nk}(\theta_n) = e_k^\top \bar{S}_n(\theta)$ و e_k یک بردار p_n بعدی است که عضو k ام آن برابر ۱ بوده و سایر اعضا برابر ۰ هستند. همچنین مشتق $\bar{S}_{nk}(\theta_n)$ نسبت به θ_n به صورت $\bar{D}_{nk}(\theta_n) = \frac{\partial}{\partial \theta_n^\top} \bar{S}_{nk}(\theta_n)$ تعریف می‌شود و داریم

$$\bar{D}_{nk}(\theta_n) = \bar{H}_{nk}(\theta_n) + \bar{E}_{nk}(\theta_n) + \bar{G}_{nk}(\theta_n),$$

که در آن $\bar{H}_{nk}$ ، $\bar{E}_{nk}$ و $\bar{G}_{nk}$ عضو k ام $\bar{H}_n$ ، $\bar{E}_n$ و $\bar{G}_n$ هستند که به طور مشابه از مشتق $\bar{S}_n(\theta)$ نسبت به θ_n استخراج شده‌اند. در ادامه لم‌هایی را که در اثبات قضایا استفاده می‌شوند، بیان می‌کنیم.

ب.۱.۳ لم‌ها

ایده اصلی برای اثبات سازگاری برآوردگرها این است که $S_n(\theta)$ را توسط $\bar{S}_n(\theta)$ تقریب بزنیم؛ زیرا ارزیابی گشتاورهای آن آسان‌تر است. لم ب.۱.۳ صحت این تقریب را تعیین می‌کند، همچنین این لم نقش مهمی در اثبات نرمال مجانبی بودن توزیع برآوردگرها در قضیه ۳.۱.۳ دارد.

لم ب.۱.۳. تحت شرایط نظم، اگر $n^{-1}p_n^2 = o(1)$ ، آنگاه

$$\|U_n(\theta_{n^0}) - \bar{U}_n(\theta_{n^0})\| = \|S_n(\theta_{n^0}) - \bar{S}_n(\theta_{n^0})\| = O_p(p_n).$$

برای سادگی در یافتن بسط تیلور تابع برآوردیابی $S_n(\theta_n)$ ، از $\bar{D}_n(\theta_n) = \frac{\partial}{\partial \theta_n^\top} \bar{S}_n(\theta_n)$ تقریب تابع گرادیانت، $D_n(\theta_n) = \frac{\partial}{\partial \theta_n^\top} S_n(\theta_n)$ ، استفاده می‌شود. لم ب.۲.۳ نمایش کاربردی برای $\bar{D}_n(\theta_n)$ ارائه می‌دهد.

لم ب.۲.۳

$$\bar{D}_n(\theta_n) = \bar{H}_n(\theta_n) + \bar{E}_n(\theta_n) + \bar{G}_n(\theta_n), \quad (\text{ب.۳})$$

به طوری که

$$\begin{aligned} \bar{H}_n(\theta_n) &= -n^{-1} E_{u|y} \left[\sum_{i=1}^n D_i^\top A_i^{\downarrow}(\theta_n, u_i) \bar{R}^{-1} A_i^{\downarrow}(\theta_n, u_i) D_i \right], \\ \bar{E}_n(\theta_n) &= -(Yn)^{-1} E_{u|y} \left[\sum_{i=1}^n D_i^\top A_i^{\downarrow}(\theta_n, u_i) \bar{R}^{-1} A_i^{\downarrow}(\theta_n, u_i) C_i(\theta_n, u_i) F_i(\theta_n, u_i) D_i \right], \\ \bar{G}_n(\theta_n) &= (Yn)^{-1} E_{u|y} \left[\sum_{i=1}^n D_i^\top A_i^{\downarrow}(\theta_n, u_i) F_i(\theta_n, u_i) J_i(\theta_n, u_i) D_i \right], \end{aligned}$$

$$\begin{aligned} C_i(\theta_n, \mathbf{u}_i) &= \text{diag} \left(y_{i1} - \mu_{i1}(\theta_n, \mathbf{u}_i), \dots, y_{im} - \mu_{im}(\theta_n, \mathbf{u}_i) \right), \\ F_i(\theta_n, \mathbf{u}_i) &= \text{diag} \left(\mu_i^\top (D_i^\top \theta_n + Z_i^\top \mathbf{u}_i), \dots, \mu_i^\top (D_{im}^\top \theta_n + Z_{im}^\top \mathbf{u}_i) \right), \\ J_i(\theta_n, \mathbf{u}_i) &= \text{diag} \left(\overline{R}_i^{-1} A_i^\top (\theta_n, \mathbf{u}_i) (y_i - \mu_i(\theta_n, \mathbf{u}_i)) \right), \end{aligned}$$

در عبارت‌های فوق، هر سه نماد $\text{diag}(a)$ و $\text{diag}(a_1, \dots, a_m)$ ، $a = (a_1, \dots, a_{n_i})^\top$ یک ماتریس قطری $m \times m$ بعدی است که عناصر قطر اصلی آن $(a_1, \dots, a_m)$ هستند.

لم ب.۳.۳. تحت شرایط نظم، اگر $n^{-1}p_n^2 = o(1)$ ، به ازای هر $\Delta > 0$ و $b_n \in \mathcal{R}^{p_n}$ داریم

$$\sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \sup_{\|b_n\|=1} \left| b_n^\top [D_n(\theta_n) - \overline{D}_n(\theta_n)] b_n \right| = O_p(\sqrt{np_n}).$$

ماتریس $D_n(\theta_n) - \overline{D}_n(\theta_n)$ یک ماتریس متقارن است. لم فوق گزاره زیر را نتیجه می‌دهد.

نتیجه ب.۱.۳.

$$\begin{aligned} \sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \left| \lambda_{\min} [D_n(\theta_n) - \overline{D}_n(\theta_n)] \right| &= O_p(\sqrt{np_n}), \\ \sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \left| \lambda_{\max} [D_n(\theta_n) - \overline{D}_n(\theta_n)] \right| &= O_p(\sqrt{np_n}). \end{aligned}$$

لم ب.۴.۳. تحت شرایط نظم، اگر $n^{-1}p_n^2 = o(1)$ ، آن‌گاه به ازای هر $\Delta > 0$ و $b_n \in \mathcal{R}^{p_n}$ داریم

$$\sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \sup_{\|b_n\|=1} \left| b_n^\top [\overline{D}_n(\theta_n) - \overline{H}_n(\theta_n)] b_n \right| = O_p(\sqrt{np_n}).$$

لم ب.۵.۳. تحت شرایط نظم، اگر $n^{-1}p_n^2 = o(1)$ ، آن‌گاه به ازای هر $\Delta > 0$ و $b_n \in \mathcal{R}^{p_n}$ داریم

$$\sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \sup_{\|b_n\|=1} \left| b_n^\top [\overline{H}_n(\theta_n) - \overline{H}_n(\theta_{n^0})] b_n \right| = O_p(\sqrt{np_n}).$$

لم ب.۶.۳. تحت شرایط نظم، اگر $n^{-1}p_n^2 = o(1)$ ، آن‌گاه به ازای هر $\xi_n \in \mathcal{R}^{p_n}$ که $\|\xi_n\| = 1$ ، داریم

$$\xi_n^\top \overline{M}_n^{-1}(\theta_{n^0}) \overline{S}_n(\theta_{n^0}) \xrightarrow{d} \mathcal{N}_{p_n}(0, 1).$$

ب.۲.۳ اثبات لم‌ها

اثبات لم ب.۱.۳

فرض کنید $Q = \{q_{j_1, j_2}\}_{1 \leq j_1, j_2 \leq m}$ نمایانگر ماتریس $\hat{R}^{-1} - \overline{R}^{-1}$ باشد. آن‌گاه

$$\begin{aligned} S_n(\theta_{n^0}) - \overline{S}_n(\theta_{n^0}) &= \sum_{i=1}^n \sum_{j_1=1}^m \sum_{j_2=1}^m q_{j_1, j_2} A_{ij_1}^{1/2}(\theta_{n^0}, \mathbf{u}_i) A_{ij_2}^{-1/2}(\theta_{n^0}, \mathbf{u}_i) \\ &\quad \times (y_{ij_1} - \mu_{ij_2}(\theta_{n^0}, \mathbf{u}_i)) D_{ij_1} \\ &= \sum_{j_1=1}^m \sum_{j_2=1}^m q_{j_1, j_2} \left[\sum_{i=1}^n A_{ij_1}^{1/2}(\theta_{n^0}, \mathbf{u}_i) \epsilon_{ij_2}(\theta_{n^0}, \mathbf{u}_i) D_{ij_1} \right], \end{aligned}$$

$$\begin{aligned}
 & \text{به طوری که } \epsilon_{ij_2}(\theta_{n^0}, \mathbf{u}_i) = \mathbf{A}_{ij_2}^{-1/2}(\theta_{n^0}, \mathbf{u}_i)(y_{ij_1} - \mu_{ij_2}(\theta_{n^0}, \mathbf{u}_i)) \text{ در نظر داشته باشید که} \\
 \mathbb{E} \left[\left\| \sum_{i=1}^n \mathbf{A}_{ij_1}^{-1/2}(\theta_{n^0}, \mathbf{u}_i) \epsilon_{ij_2}(\theta_{n^0}, \mathbf{u}_i) \mathbf{D}_{ij_1} \right\|^2 \right] &= \sum_{i=1}^n \mathbf{A}_{ij_1}^{-1/2}(\theta_{n^0}, \mathbf{u}_i) \mathbb{E}[\epsilon_{ij_2}^2(\theta_{n^0}, \mathbf{u}_i)] \mathbf{D}_{ij_1}^\top \mathbf{D}_{ij_1} \\
 &\leq \sum_{i=1}^n \mathbf{D}_{ij_1}^\top \mathbf{D}_{ij_1} = O(np_n).
 \end{aligned}$$

بنابراین $\|\sum_{i=1}^n \mathbf{A}_{ij_1}^{-1/2}(\theta_{n^0}, \mathbf{u}_i) \epsilon_{ij_2}(\theta_{n^0}, \mathbf{u}_i)\| = O_p(\sqrt{np_n})$ بر اساس شرط (۴.۵)، به ازای هر $1 \leq j_1$ و $j_2 \leq m$ داریم $q_{j_1, j_2} = O_p(\sqrt{p_n/n})$ ، لذا اثبات کامل می‌شود.

اثبات لم ب.۲.۳

اثبات مشابه پن (۲۰۰۲) است.

اثبات لم ب.۳.۳

فرض کنید $H_n(\theta_n)$ ، $E_n(\theta_n)$ و $G_n(\theta_n)$ به ترتیب مشابه $\bar{H}_n(\theta_n)$ ، $\bar{E}_n(\theta_n)$ و $\bar{G}_n(\theta_n)$ تعریف شوند با این تفاوت که $\bar{R}$ توسط $\hat{R}$ جایگزین می‌شود. بر اساس لم ب.۲.۳، کافی است که نتایج زیر را اثبات کنیم.

$$\sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \sup_{\|\mathbf{b}_n\|=1} \left| \mathbf{b}_n^\top [H_n(\theta_n) - \bar{H}_n(\theta_n)] \mathbf{b}_n \right| = O_p(\sqrt{np_n}), \quad (۴.ب)$$

$$\sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \sup_{\|\mathbf{b}_n\|=1} \left| \mathbf{b}_n^\top [E_n(\theta_n) - \bar{E}_n(\theta_n)] \mathbf{b}_n \right| = O_p(\sqrt{np_n}), \quad (۵.ب)$$

$$\sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \sup_{\|\mathbf{b}_n\|=1} \left| \mathbf{b}_n^\top [G_n(\theta_n) - \bar{G}_n(\theta_n)] \mathbf{b}_n \right| = O_p(\sqrt{np_n}). \quad (۶.ب)$$

حال داریم

$$\begin{aligned}
 \left| \mathbf{b}_n^\top [H_n(\theta_n) - \bar{H}_n(\theta_n)] \mathbf{b}_n \right| &= \left| \sum_{i=1}^n \mathbf{b}_n^\top \mathbf{D}_i^\top \mathbf{A}_i^{-1/2}(\theta_n, \mathbf{u}_i) [\hat{R}^{-1} - \bar{R}^{-1}] \mathbf{A}_i^{-1/2}(\theta_n, \mathbf{u}_i) \mathbf{D}_i \mathbf{b}_n \right| \\
 &\leq \|\hat{R}^{-1} - \bar{R}^{-1}\| \lambda_{\max}(\mathbf{A}_i(\theta_n, \mathbf{u}_i)) \lambda_{\max}\left(\sum_{i=1}^n \mathbf{D}_i^\top \mathbf{D}_i\right) \|\mathbf{b}_n\|^2.
 \end{aligned}$$

تحت شرط (۴.۵)، عبارت (۴.ب) اثبات می‌شود. سپس در نظر داشته باشید که

$$\begin{aligned}
 \left| \mathbf{b}_n^\top [E_n(\theta_n) - \bar{E}_n(\theta_n)] \mathbf{b}_n \right| &= \frac{1}{2} \left| \sum_{i=1}^n \sum_{j=1}^m (1 - 2\mu_{ij}(\theta_n, \mathbf{u}_i)) \epsilon_{ij}(\theta_n, \mathbf{u}_i) \mathbf{b}_n^\top \mathbf{D}_i^\top \mathbf{A}_i^{-1/2}(\theta_n, \mathbf{u}_i) \right. \\
 &\quad \left. \times [\hat{R}^{-1} - \bar{R}^{-1}] \mathbf{e}_j \mathbf{e}_j^\top \mathbf{D}_i \mathbf{b}_n \right| \\
 &\leq \sum_{i=1}^n \sum_{j=1}^m \mathbf{A}_i^{-1/2}(\theta_n, \mathbf{u}_i) \left| \mathbf{b}_n^\top \mathbf{D}_i^\top \mathbf{A}_i^{-1/2}(\theta_n, \mathbf{u}_i) [\hat{R}^{-1} - \bar{R}^{-1}] \mathbf{e}_j \right| \cdot |\mathbf{e}_j^\top \mathbf{D}_i \mathbf{b}_n| \\
 &\leq \sum_{i=1}^n \sum_{j=1}^m \mathbf{A}_i^{-1/2}(\theta_n, \mathbf{u}_i) \|\hat{R}^{-1} - \bar{R}^{-1}\| \cdot \|\mathbf{A}_i^{-1/2}(\theta_n, \mathbf{u}_i)\| \cdot \|\mathbf{D}_i \mathbf{b}_n\|^2.
 \end{aligned}$$

تحت برقراری شرط (۴.A)، داریم

$$\begin{aligned} & \sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \sup_{\|b_n\|=1} \left| b_n^\top [E_n(\theta_n) - \bar{E}_n(\theta_n)] b_n \right| \\ & \leq C \|\hat{R}^{-1} - \bar{R}^{-1}\| \cdot \sum_{i=1}^n \sum_{j=1}^m \sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} A_{ij}^{-1/2}(\theta_n, u_i) \sup_{\|b_n\|=1} \|D_i b_n\|^2 \\ & = O_p(\sqrt{p_n/n}) O(n) = O_p(\sqrt{np_n}). \end{aligned}$$

بنابراین عبارت (ب.۵) نیز اثبات می‌شود. عبارت (ب.۶) مشابه دو عبارت قبل برقرار است.

اثبات لم ب.۴.۳

بر اساس لم ب.۲.۳ تنها کافی است نشان دهیم که

$$\sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \sup_{\|b_n\|=1} \left| b_n^\top \bar{E}_n(\theta_n) b_n \right| = O_p(\sqrt{np_n}), \quad (\text{ب.۷})$$

و

$$\sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \sup_{\|b_n\|=1} \left| b_n^\top \bar{G}_n(\theta_n) b_n \right| = O_p(\sqrt{np_n}). \quad (\text{ب.۸})$$

در نظر داشته باشید که می‌توان $\bar{E}_n(\theta_n)$ را به صورت

$$\begin{aligned} \bar{E}_n(\theta_n) &= \frac{1}{\sqrt{n}} \sum_{i=1}^n \sum_{j=1}^m (\mathbb{I} - \Psi_{\mu_{ij}}(\theta_{n^0}, u_i)) \epsilon_{ij}(\theta_{n^0}, u_i) D_i^\top A_i^{-1/2}(\theta_{n^0}, u_i) \bar{R}^{-1} e_j e_j^\top D_i \\ &+ \frac{1}{\sqrt{n}} \sum_{i=1}^n \sum_{j=1}^m (\mathbb{I} - \Psi_{\mu_{ij}}(\theta_{n^0}, u_i)) \epsilon_{ij}(\theta_{n^0}, u_i) D_i^\top \left[A_i^{-1/2}(\theta_n, u_i) \right. \\ &\quad \left. - A_i^{-1/2}(\theta_{n^0}, u_i) \right] \bar{R}^{-1} e_j e_j^\top D_i \\ &+ \frac{1}{\sqrt{n}} \sum_{i=1}^n \sum_{j=1}^m \left[(\mathbb{I} - \Psi_{\mu_{ij}}(\theta_n, u_i)) A_{ij}^{-1/2}(\theta_n, u_i) - (\mathbb{I} - \Psi_{\mu_{ij}}(\theta_{n^0}, u_i)) \right. \\ &\quad \left. \times A_{ij}^{-1/2}(\theta_{n^0}, u_i) \right] (y_{ij} - \mu_{ij}(\theta_{n^0}, u_i)) D_i^\top A_i^{-1/2}(\theta_n, u_i) \bar{R}^{-1} e_j e_j^\top D_i \\ &+ \frac{1}{\sqrt{n}} \sum_{i=1}^n \sum_{j=1}^m (\mathbb{I} - \Psi_{\mu_{ij}}(\theta_n, u_i)) A_{ij}^{-1/2}(\theta_n, u_i) (\mu_{ij}(\theta_{n^0}, u_i) - \mu_{ij}(\theta_n, u_i)) \\ &\quad \times D_i^\top A_i^{-1/2}(\theta_n, u_i) \bar{R}^{-1} e_j e_j^\top D_i \\ &= \bar{E}_{\setminus n}(\theta_{n^0}) + \sum_{k=\setminus}^{\setminus} \bar{E}_{kn}(\theta_n) \end{aligned}$$

تجزیه کرد. بنابراین برای عبارت (ب.۷) کافی است نشان دهیم که

$$\sup_{\|b_n\|=1} \left| b_n^\top \bar{E}_{\setminus n}(\theta_{n^0}) b_n \right| = O_p(\sqrt{np_n})$$

$$\sup_{\|\theta_n - \theta_{n^\circ}\| \leq \Delta \sqrt{p_n/n}} \sup_{\|b_n\|=1} |b_n^\top \bar{E}_{kn}(\theta_n) b_n| = O_p(\sqrt{np_n}).$$

ابتدا نشان می‌دهیم که $\sup_{\|b_n\|=1} |b_n^\top \bar{E}_n(\theta_{n^\circ}) b_n| = O_p(\sqrt{np_n})$ برقرار است. اگر نشان دهیم که $\|\bar{E}_n(\theta_{n^\circ})\| = O_p(\sqrt{np_n})$ ، به طوری که $\|\bar{E}_n(\theta_{n^\circ})\| = \sqrt{\text{tr}(\bar{E}_n(\theta_{n^\circ}) \bar{E}_n^\top(\theta_{n^\circ}))}$ ، آنگاه نتیجه حاصل می‌شود. حال داریم

$$\begin{aligned} & E[\|\bar{E}_n(\theta_{n^\circ})\|^2] \\ &= \frac{1}{4} \sum_{i=1}^n \sum_{j_1=1}^m \sum_{j_2=1}^m (1 - 2\mu_{ij_1}(\theta_{n^\circ}, u_i))(1 - 2\mu_{ij_2}(\theta_{n^\circ}, u_i)) \\ & \quad \times E[\epsilon_{ij_1}(\theta_{n^\circ}, u_i) \epsilon_{ij_2}(\theta_{n^\circ}, u_i)] \\ & \quad \times \text{tr}[D_i^\top A_i^{1/2}(\theta_{n^\circ}, u_i) \bar{R}^{-1} e_{j_1} e_{j_2}^\top D_i D_i^\top e_{j_2} e_{j_1}^\top \bar{R}^{-1} A_i^{1/2}(\theta_{n^\circ}, u_i) D_i] \\ &\leq C \sum_{i=1}^n \sum_{j_1=1}^m \sum_{j_2=1}^m |e_{j_1}^\top D_i D_i^\top e_{j_2} e_{j_2}^\top \bar{R}^{-1} A_i^{1/2}(\theta_{n^\circ}, u_i) D_i D_i^\top A_i^{1/2}(\theta_{n^\circ}, u_i) \bar{R}^{-1} e_{j_1}| \\ &\leq C \sum_{i=1}^n \sum_{j_1=1}^m \sum_{j_2=1}^m \|e_{j_1}^\top D_i\| \cdot \|D_i^\top e_{j_2}\| \cdot \|e_{j_2}^\top \bar{R}^{-1} A_i^{1/2}(\theta_{n^\circ}, u_i) D_i\| \\ & \quad \times \|D_i^\top A_i^{1/2}(\theta_{n^\circ}, u_i) \bar{R}^{-1} e_{j_1}\|. \end{aligned}$$

در نظر داشته باشید که $\|e_{j_2}^\top \bar{R}^{-1} A_i^{1/2}(\theta_{n^\circ}, u_i) D_i\| \leq \|D_i^\top e_{j_2}\| = \|D_{ij_2}\|$ ، $\|e_{j_1}^\top D_i\| = \|D_{ij_1}\|$ و $\|D_i^\top A_i^{1/2}(\theta_{n^\circ}, u_i) \bar{R}^{-1} e_{j_1}\| \leq C(\text{tr}(D_i D_i^\top))^{1/2}$ و $C(\text{tr}(D_i D_i^\top))^{1/2}$ بنابرین براساس شرایط (۲.۸) و (۳.۸)، داریم

$$\begin{aligned} E[\|\bar{E}_n(\theta_{n^\circ})\|^2] &\leq C \sum_{i=1}^n \sum_{j_1=1}^m \sum_{j_2=1}^m \|D_{ij_1}\| \cdot \|D_{ij_2}\| \text{tr}(D_i D_i^\top) \\ &\leq C \max_{i,j} \|D_{ij}\|^2 \text{tr}\left(\sum_{i=1}^n D_i D_i^\top\right) = O(np_n^2), \end{aligned}$$

که $\sup_{\|b_n\|=1} |b_n^\top \bar{E}_n(\theta_{n^\circ}) b_n| = O_p(\sqrt{np_n})$ سپس داریم

$$\begin{aligned} |b_n^\top \bar{E}_n(\theta_{n^\circ}) b_n| &= \left| \frac{1}{4} \sum_{i=1}^n \sum_{j=1}^m (1 - 2\mu_{ij}(\theta_{n^\circ}, u_i)) \epsilon_{ij}^{1/2}(\theta_{n^\circ}, u_i) b_n^\top D_i^\top [A_i^{1/2}(\theta_n, u_i) \right. \\ & \quad \left. - A_i^{1/2}(\theta_{n^\circ}, u_i)] \bar{R}^{-1} e_j e_j^\top D_i b_n \right| \\ &\leq C \sum_{i=1}^n \sum_{j=1}^m |b_n^\top D_i^\top [A_i^{1/2}(\theta_n, u_i) - A_i^{1/2}(\theta_{n^\circ}, u_i)] \bar{R}^{-1} e_j| \cdot |e_j^\top D_i b_n| \\ &\leq C \sum_{i=1}^n \sum_{j=1}^m \|D_i b_n\|^2 \lambda_{\max}(\bar{R}^{-1}) \max_j |A_{ij}^{1/2}(\theta_n, u_i) - A_{ij}^{1/2}(\theta_{n^\circ}, u_i)|. \end{aligned}$$

بین θ_n و θ_{n° ، θ_n^* می‌توان یافت که

$$\begin{aligned} A_{ij}^{1/2}(\theta_n, u_i) - A_{ij}^{1/2}(\theta_{n^\circ}, u_i) &= \frac{1}{4} A_{ij}^{1/2}(\theta_n^*, u_i) (1 - 2\mu_{ij}(\theta_n^*, u_i)) D_{ij}^\top (\theta_n - \theta_{n^\circ}) \\ &\leq C \|D_{ij}\| \cdot \|\theta_n - \theta_{n^\circ}\|. \end{aligned}$$

بنابراین

$$\begin{aligned} & \sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \sup_{\|b_n\|=1} \left| b_n^\top \bar{E}_{\gamma_n}(\theta_n) b_n \right| \\ & \leq C \max_{ij} \|D_{ij}\| \sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \|(\theta_n - \theta_{n^0})\| \cdot \lambda_{\max} \left(\sum_{i=1}^n D_i D_i^\top \right) \\ & = O(\sqrt{p_n}) O(\sqrt{p_n/n}) O(n) = O(\sqrt{n} p_n). \end{aligned}$$

به طور مشابه می‌توان نشان داد که $k = O_p(\sqrt{n} p_n)$, $\sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \sup_{\|b_n\|=1} \left| b_n^\top \bar{E}_{kn}(\theta_n) b_n \right| = O_p(\sqrt{n} p_n)$, $k = 3, 4$. لذا عبارت (ب.۷) اثبات می‌شود. برقراری عبارت (ب.۸) از طریق تکنیک‌های مشابه نتیجه می‌شود.

اثبات لم ب.۵.۳

$$\begin{aligned} & b_n^\top [\bar{H}_n(\theta_n) - \bar{H}_n(\theta_{n^0})] b_n \\ & = \left| \sum_{i=1}^n b_n^\top D_i^\top [A_i^{1/\gamma}(\theta_n) \bar{R}^{-1} A_i^{1/\gamma}(\theta_n) - A_i^{1/\gamma}(\theta_{n^0}) \bar{R}^{-1} A_i^{1/\gamma}(\theta_{n^0})] D_i b_n \right| \\ & \leq \left| \sum_{i=1}^n b_n^\top D_i^\top [A_i^{1/\gamma}(\theta_n) \bar{R}^{-1} A_i^{1/\gamma}(\theta_n) - A_i^{1/\gamma}(\theta_n) \bar{R}^{-1} A_i^{1/\gamma}(\theta_{n^0})] D_i b_n \right| \\ & \quad + \left| \sum_{i=1}^n b_n^\top D_i^\top [A_i^{1/\gamma}(\theta_n) \bar{R}^{-1} A_i^{1/\gamma}(\theta_n) - A_i^{1/\gamma}(\theta_{n^0}) \bar{R}^{-1} A_i^{1/\gamma}(\theta_{n^0})] D_i b_n \right| \\ & = I_{n1} + I_{n2}. \end{aligned}$$

بر اساس نامساوی کوشی-شوارتز می‌توان

$$\begin{aligned} I_{n1} & \leq \sum_{i=1}^n \left| b_n^\top D_i^\top A_i^{1/\gamma}(\theta_n) \bar{R}^{-1} [A_i^{1/\gamma}(\theta_n) - A_i^{1/\gamma}(\theta_{n^0})] D_i b_n \right| \\ & \leq \sum_{i=1}^n \|b_n^\top D_i^\top A_i^{1/\gamma}(\theta_n) \bar{R}^{-1}\| \cdot \|\bar{R}^{-1}\| \| [A_i^{1/\gamma}(\theta_n) - A_i^{1/\gamma}(\theta_{n^0})] D_i b_n \|, \end{aligned}$$

را نتیجه گرفت. سپس از عبارت‌های

$$\begin{aligned} \|b_n^\top D_i^\top A_i^{1/\gamma}(\theta_n, u_i) \bar{R}^{-1}\| & = \left[b_n^\top D_i^\top A_i^{1/\gamma}(\theta_n) \bar{R}^{-1} A_i^{1/\gamma}(\theta_n, u_i) D_i b_n \right] \\ & \leq \lambda_{\max}^{1/\gamma}(\bar{R}^{-1}) \lambda_{\max}^{1/\gamma}(A_i(\theta_n, u_i)) \|D_i b_n\|, \end{aligned}$$

$$\begin{aligned} & \|\bar{R}^{-1}\| \| [A_i^{1/\gamma}(\theta_n) - A_i^{1/\gamma}(\theta_{n^0})] D_i b_n \| \\ & \leq \lambda_{\max}^{1/\gamma}(\bar{R}^{-1}) \lambda_{\max}^{1/\gamma}([A_i^{1/\gamma}(\theta_n) - A_i^{1/\gamma}(\theta_{n^0})]) \|D_i b_n\|, \end{aligned}$$

می‌توان نشان داد که

$$\begin{aligned} I_{n1} &\leq C \max_{i,j} \left| A_{ij}^{1/2}(\theta_n) - A_{ij}^{1/2}(\theta_{n^0}) \right| \sum_{i=1}^n \|D_i b_n\|^2 \\ &\leq C \max_{i,j} \|D_{ij}\| \cdot \|\theta_n - \theta_{n^0}\| \cdot \lambda_{\max} \left(\sum_{i=1}^n D_i^\top D_i \right) \|b_n\|^2. \end{aligned}$$

بنابراین خواهیم داشت

$$\sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \sup_{\|b_n\|=1} I_{n1} = O_p(\sqrt{np_n}).$$

به طور مشابه می‌توان نشان داد که $\sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \sup_{\|b_n\|=1} I_{n2} = O_p(\sqrt{np_n})$ و لم اثبات می‌شود.

اثبات لم ب.۳.۶

عبارت مورد نظر که باید توزیع مجانبی آن را یافت به صورت $\alpha_n^\top \bar{M}_n^{-1/2}(\theta_{n^0}) \bar{S}_n(\theta_{n^0}) = \sum_{i=1}^n Z_{ni}$ قابل بیان است، به طوری که $Z_{ni} = \alpha_n^\top \bar{M}_n^{-1/2}(\theta_{n^0}) D_i^\top A_i^{1/2}(\theta_{n^0}, u_i) \bar{R}^{-1} \epsilon_i(\theta_{n^0}, u_i)$ از آنجایی که $\bar{M}_n(\theta_{n^0}) = \text{Cov}(\bar{S}_n(\theta_{n^0}))$ داریم $\text{Var}(\alpha_n^\top \bar{M}_n^{-1/2}(\theta_{n^0}) \bar{S}_n(\theta_{n^0})) = 1$ برای نشان دادن توزیع نرمال مجانبی کافی است که شرط لیندبرگ را بررسی کنیم، به این صورت که به ازای هر $\epsilon > 0$ داشته باشیم $\sum_{i=1}^n E[Z_{ni}^2 I(|Z_{ni}| > \epsilon)] \rightarrow 0$ بر اساس نامساوی کوشی-شوارتز داریم

$$\begin{aligned} Z_{ni}^2 &\leq \|\alpha_n^\top \bar{M}_n^{-1/2}(\theta_{n^0}) D_i^\top A_i^{1/2}(\theta_{n^0}, u_i) \bar{R}^{-1}\|^2 \cdot \|\epsilon_i(\theta_{n^0}, u_i)\|^2 \\ &\leq \lambda_{\max}(\bar{R}^{-2}) \lambda_{\max}(A_i(\theta_{n^0}, u_i)) (\alpha_n^\top \bar{M}_n^{-1/2}(\theta_{n^0}) D_i^\top D_i \bar{M}_n^{-1/2}(\theta_{n^0}) \alpha_n) \|\epsilon_i(\theta_{n^0}, u_i)\|^2 \\ &\leq C \gamma_{ni} \|\epsilon_i(\theta_{n^0}, u_i)\|^2, \end{aligned}$$

به طوری که $\gamma_{ni} = \alpha_n^\top \bar{M}_n^{-1/2}(\theta_{n^0}) D_i^\top D_i \bar{M}_n^{-1/2}(\theta_{n^0}) \alpha_n$ سپس نشان خواهیم داد هنگامی که $n \rightarrow \infty$ آن‌گاه $\max_{1 \leq i \leq n} \gamma_{ni} \rightarrow 0$ در نظر داشته باشید که $\gamma_{ni} \leq \lambda_{\max}(D_i^\top D_i) \lambda_{\min}^{-1}(\bar{M}_n(\theta_{n^0}))$ از آنجایی که $\bar{M}_n(\theta_{n^0})$ متقارن است، برای ارزیابی $\lambda_{\min}(\bar{M}_n(\theta_{n^0}))$ ، به ازای هر $b_n \in \mathbb{R}^{p_n}$ داریم

$$\begin{aligned} b_n^\top \bar{M}_n(\theta_{n^0}) b_n &\geq \lambda_{\min}(R_0) \lambda_{\min}(\bar{R}^{-2}) \sum_{i=1}^n \lambda_{\min}(A_i(\theta_{n^0}, u_i)) b_n^\top D_i^\top D_i b_n \\ &\geq C b_n^\top \left(\sum_{i=1}^n D_i^\top D_i \right) b_n \geq C \lambda_{\min} \left(\sum_{i=1}^n D_i^\top D_i \right) \|b_n\|^2. \end{aligned}$$

بنابراین نامساوی $\inf_{\|b_n\|=1} |b_n^\top \bar{M}_n(\theta_{n^0}) b_n| \geq C \lambda_{\min} \left(\sum_{i=1}^n D_i^\top D_i \right)$ را نتیجه می‌گیریم و این نشان می‌دهد که $\lambda_{\min}(\bar{M}_n(\theta_{n^0})) \geq C \lambda_{\min} \left(\sum_{i=1}^n D_i^\top D_i \right)$ حال می‌توان نامساوی

$$\gamma_{ni} \leq \frac{\lambda_{\max}(D_i^\top D_i)}{C \lambda_{\min} \left(\sum_{i=1}^n D_i^\top D_i \right)} \leq \frac{\text{tr}(D_i^\top D_i)}{C \lambda_{\min} \left(\sum_{i=1}^n D_i^\top D_i \right)} = \frac{\sum_{j=1}^m D_{ij}^\top D_{ij}}{C \lambda_{\min} \left(\sum_{i=1}^n D_i^\top D_i \right)},$$

و به دنبال آن $\max_{1 \leq i \leq n} \gamma_{ni} \leq O(n^{-1} p_n) = o(1)$ را نتیجه گرفت. بنابراین داریم

$$\sum_{i=1}^n E[Z_{ni}^2 I(|Z_{ni}| > \epsilon)] \leq \sum_{i=1}^n C \gamma_{ni} E \left[\|\epsilon_i(\theta_{n^0}, u_i)\|^2 I \left\{ \|\epsilon_i(\theta_{n^0}, u_i)\|^2 > \frac{\epsilon^2}{C \gamma_{ni}} \right\} \right].$$

بر اساس شرط (۳.A)، عبارت $\|\epsilon_i(\theta_{n^0}, u_i)\|^2$ به طور یکنواخت کراندار است، لذا به ازای هر $\epsilon > 0$ و $\delta > 0$ ، یک عدد ثابت و مثبت مانند C وجود دارد که $I \left\{ \|\epsilon_i(\theta_{n^0}, u_i)\|^2 > \frac{\epsilon^2}{C \gamma_{ni}} \right\} = 0$ این تضمین می‌کند که

$$\sum_{i=1}^n C \gamma_{ni} E \left[\|\epsilon_i(\theta_{n^0}, u_i)\|^2 I \left\{ \|\epsilon_i(\theta_{n^0}, u_i)\|^2 > \frac{\epsilon^2}{C \gamma_{ni}} \right\} \right] \rightarrow 0.$$

لذا شرط لیندبرگ برقرار است.

۳.۳.۱.۳ اثبات قضیه

قسمت (i): بر اساس تعریف $\hat{\theta}_n$ بدیهی است که به ازای $k = 1, \dots, s_n$ ، $S_{nk}(\hat{\theta}_n, u_i) = 0$ ، کافی است ثابت کنیم که $\mathbb{P}(|\hat{\theta}_{nk}| \geq a \lambda_n, k = 1, \dots, s_n) \rightarrow 0$ ، یعنی با احتمال نزدیک به ۱، تابع تاوان برابر ۰ است. می‌دانیم که

$$\min_{1 \leq k \leq s_n} |\hat{\theta}_{nk}| \geq \min_{1 \leq k \leq s_n} |\hat{\theta}_{n^0 k}| - \max_{1 \leq k \leq s_n} |\hat{\theta}_{n^0 k} - \hat{\theta}_{nk}| \geq \min_{1 \leq k \leq s_n} |\hat{\theta}_{n^0 k}| - \|\theta_{n^0} - \hat{\theta}_{n^0}\|.$$

همچنین فرض کنید

$$\|\theta_{n^0} - \hat{\theta}_{n^0}\| = \sqrt{s_n/n}. \quad (۹.ب)$$

از آنجایی که $\min_{1 \leq k \leq s_n} |\hat{\theta}_{n^0 k}|/\lambda \rightarrow \infty$ و $\|\theta_{n^0} - \hat{\theta}_{n^0}\| = o(\lambda_n)$ می‌توان نتیجه گرفت که

$$\mathbb{P} \left(\min_{1 \leq k \leq s_n} |\hat{\theta}_{n^0 k}| - \|\theta_{n^0} - \hat{\theta}_{n^0}\| \geq a \lambda_n \right) = \mathbb{P}(\|\theta_{n^0} - \hat{\theta}_{n^0}\| \leq \min_{1 \leq k \leq s_n} |\hat{\theta}_{n^0 k}| - a \lambda_n) \rightarrow 1.$$

بنابراین هنگامی که $n \rightarrow \infty$ ، داریم $\mathbb{P}(\min_{1 \leq k \leq s_n} |\hat{\theta}_{nk}| \geq a \lambda_n) \rightarrow 1$.

قسمت (ii): به ازای هر $k = s_n + 1, \dots, p_n$ ، داریم $\hat{\theta}_{nk} = 0$ ، لذا $q_{\lambda_n}(|\theta_{nk}|) \operatorname{sgn}(\theta_{nk}) = 0$ ، کافی است نشان دهیم که

$$\mathbb{P} \left(\max_{s_n+1 \leq k \leq p_n} |S_{nk}(\hat{\theta})| \leq \frac{\lambda_n}{\log n} \right) \rightarrow 1. \quad (۱۰.ب)$$

به عبارت دیگر کافی است نشان دهیم که

$$\mathbb{P} \left(\max_{s_n+1 \leq k \leq p_n} |S_{nk}(\hat{\theta}_n) - \bar{S}_{nk}(\hat{\theta}_n)| > \frac{\lambda_n}{\sqrt{2} \log n} \right) \rightarrow 0, \quad (۱۱.ب)$$

و

$$\mathbb{P} \left(\max_{s_n+1 \leq k \leq p_n} |\bar{S}_{nk}(\hat{\theta}_n)| > \frac{\lambda_n}{\sqrt{2} \log n} \right) \rightarrow 0. \quad (۱۲.ب)$$

عبارت سمت چپ در (ب.۱۱)، توسط عبارت

$$\begin{aligned}
 & \mathbb{P} \left(\max_{s_n+1 \leq k \leq p_n} n^{-1} E_{u|y} \left[\sum_{i=1}^n \left| e_k D_i^\top A_i^{-\frac{1}{\gamma}}(\theta_n, u_i) [\hat{R}^{-1} - \bar{R}^{-1}] A_i^{-\frac{1}{\gamma}}(\theta_n, u_i) \right. \right. \right. \\
 & \quad \left. \left. \left. \times (y_i - \mu_i(\theta_n, u_i)) \right] \right] > \frac{\lambda_n}{\sqrt{\log n}} \right) \\
 & \leq \mathbb{P} \left(\max_{s_n+1 \leq k \leq p_n} n^{-1} E_{u|y} \left[\sum_{i=1}^n \|e_k D_i^\top A_i^{-\frac{1}{\gamma}}(\theta_n, u_i)\| \cdot \|\hat{R}^{-1} - \bar{R}^{-1}\| \cdot \|A_i^{-\frac{1}{\gamma}}(\theta_n, u_i)\| \right. \right. \\
 & \quad \left. \left. \left. \times (y_i - \mu_i(\theta_n, u_i)) \right] \right] > \frac{\lambda_n}{\sqrt{\log n}} \right) \\
 & \leq \mathbb{P} \left(\|\hat{R}^{-1} - \bar{R}^{-1}\| n^{-1} E_{u|y} \left[\sum_{i=1}^n \left(\max_{s_n+1 \leq k \leq p_n} \|e_k D_i^\top A_i^{-\frac{1}{\gamma}}(\theta_n, u_i)\| \right) \|A_i^{-\frac{1}{\gamma}}(\theta_n, u_i)\| \right. \right. \right. \\
 & \quad \left. \left. \left. \times (y_i - \mu_i(\theta_n, u_i)) \right] \right] > \frac{\lambda_n}{\sqrt{\log n}} \right) \\
 & \leq \mathbb{P} \left(n^{-1} E_{u|y} \left[\sum_{i=1}^n \|\epsilon_i(\theta_n, u_i)\| \right] > \frac{\lambda_n \sqrt{n}}{\sqrt{\log n}} \right) \\
 & \leq C \frac{n^{-1} E_{u|y} \left[\sum_{i=1}^n E(\|\epsilon_i(\theta_n, u_i)\|) \right] \sqrt{s_n \log n}}{\lambda_n \sqrt{n}} = O\left(\frac{\sqrt{s_n \log n}}{\lambda_n \sqrt{n}}\right) = o(1),
 \end{aligned}$$

از بالا کراندار است، به طوری که نامساوی سوم از شرط‌های (A.۲) و (A.۴) نتیجه می‌شود. همچنین شرط (A.۴) بیان می‌کند که $\sqrt{s_n \log n} / \lambda_n \sqrt{n} \rightarrow 0$. برای اثبات عبارت (ب.۱۲)، بسط تیلور عبارت $\bar{S}_{nk}(\hat{\theta}_n, u_i)$ را به صورت

$$\begin{aligned}
 \bar{S}_{nk}(\hat{\theta}_n, u_i) &= \bar{S}_{nk}(\theta_{n^\circ}, u_i) + \frac{\partial \bar{S}_{nk}(\theta_n, u_i)}{\partial \theta_n^\top} (\hat{\theta}_n - \theta_{n^\circ}) \\
 &+ (\hat{\theta}_n - \theta_{n^\circ})^\top \frac{\partial^2 \bar{S}_{nk}(\theta_n^*, u_i)}{\partial \theta_n \partial \theta_n^\top} (\hat{\theta}_n - \theta_{n^\circ}), \tag{ب.۱۳}
 \end{aligned}$$

در نظر می‌گیریم، که θ_n^* بین θ_{n° و $\hat{\theta}_n$ است. فرض کنید $D_k(\theta_n, u_i) = \frac{\partial \bar{S}_{nk}(\theta_n, u_i)}{\partial \theta_n^\top}$ و $\nabla_k(\theta_n, u_i) = \frac{\partial^2 \bar{S}_{nk}(\theta_n, u_i)}{\partial \theta_n \partial \theta_n^\top}$. $D_k(\theta_n, u_i)$ زیربرداری است که شامل s_n عضو اول بردار $D_k(\theta_n, u_i)$ است. همچنین $\nabla_k(\theta_n, u_i)$ یک ماتریس $s_n \times s_n$ بعدی است که همه اعضای s_n سطر و ستون اول ماتریس $\nabla_k(\theta_n, u_i)$ را شامل می‌شود. از آنجایی که $(\hat{\theta}_n - \theta_{n^\circ})^\top = ((\hat{\theta}_{n1} - \theta_{n^\circ 1})^\top, 0^\top)^\top$ ، عبارت (ب.۱۳) را می‌توان به صورت

$$\begin{aligned}
 \bar{S}_{nk}(\hat{\theta}_n, u_i) &= \bar{S}_{nk}(\theta_{n^\circ}, u_i) + D_{k1}(\theta_n, u_i)(\hat{\theta}_{n1} - \theta_{n^\circ 1}) \\
 &+ (\hat{\theta}_{n1} - \theta_{n^\circ 1})^\top \nabla_{k1}(\theta_n^*, u_i)(\hat{\theta}_{n1} - \theta_{n^\circ 1}),
 \end{aligned}$$

بیان نمود. حال داریم

$$\begin{aligned}
 & \mathbb{P} \left(\max_{s_n+1 \leq k \leq p_n} |\bar{S}_{nk}(\hat{\theta}_n)| > \frac{\lambda_n}{\sqrt{\log n}} \right) \\
 \leq & \mathbb{P} \left(\max_{s_n+1 \leq k \leq p_n} |\bar{S}_{nk}(\hat{\theta}_n)| > \frac{\lambda_n}{\sqrt{\log n}} \right) \\
 & + \mathbb{P} \left(\max_{s_n+1 \leq k \leq p_n} |D_{k1}(\theta_n, u_i)(\hat{\theta}_{n1} - \theta_{n1\circ})| > \frac{\lambda_n}{\sqrt{\log n}} \right) \\
 & + \mathbb{P} \left(\max_{s_n+1 \leq k \leq p_n} |(\hat{\theta}_{n1} - \theta_{n1\circ})^\top \nabla_{k1}(\theta_n^*, u_i)(\hat{\theta}_{n1} - \theta_{n1\circ})| > \frac{\lambda_n}{\sqrt{\log n}} \right) \\
 = & I_{n1} + I_{n2} + I_{n3}.
 \end{aligned}$$

اگر نشان دهیم $I_{ni} = o(1)$ که $i = 1, 2, 3$ ، آنگاه رابطه (ب.۱۲) برقرار خواهد بود. ابتدا نشان می‌دهیم که $I_{n1} = o(1)$. بدیهی است که $\mathbb{P} \left(|\bar{S}_{nk}(\hat{\theta})| > \frac{\lambda_n}{\sqrt{\log n}} \right) \leq \sum_{k=s_n+1}^{p_n} \mathbb{P} \left(|\bar{S}_{nk}(\hat{\theta})| > \frac{\lambda_n}{\sqrt{\log n}} \right)$. حال می‌توان نوشت $Z_i = e_k D_i^\top A_i^{\frac{1}{2}}(\theta_{n\circ}, u_i) \bar{R}^{-1} \epsilon_i(\theta_{n\circ}, u_i)$ که به طوری که $\bar{S}_{nk}(\hat{\theta}_{n\circ}, u_i) = n^{-1} \sum_{i=1}^n Z_i$ متغیرهای تصادفی مستقل با میانگین ۰ هستند. به ازای هر $l \geq 2$ و ثابت‌های $M > 0$ و $\delta > 0$ داریم

$$\begin{aligned}
 \mathbb{E} |Z_i|^l & \leq \mathbb{E} \left[\|e_k D_i^\top A_i^{\frac{1}{2}}(\theta_{n\circ}, u_i) \bar{R}^{-1}\|^l \|\epsilon_i(\theta_{n\circ}, u_i)\|^l \right] \\
 & \leq C^l \mathbb{E} \left[\|\epsilon_i(\theta_{n\circ}, u_i)\|^l \right] = C^l \mathbb{E} \left[\left(\sum_{j=1}^m \epsilon_{ij}^2(\theta_{n\circ}, u_i) \right)^{l/2} \right] \\
 & \leq C^l m^{l/2-1} \sum_{j=1}^m \mathbb{E} |\epsilon_{ij}(\theta_{n\circ}, u_i)|^l \leq C^l m^{l/2-1} \sum_{j=1}^m l! M_{\sqrt{}}^{-l} \mathbb{E} \left(\exp(M_{\sqrt{}} |\epsilon_{ij}(\theta_{n\circ}, u_i)|) \right) \\
 & \leq C^l m^{l/2-1} m l! M_{\sqrt{}} \leq l! M^{l-2} \delta / 2.
 \end{aligned}$$

در عبارت فوق نامساوی دوم از شرایط (۲.A) و (۳.A) نتیجه می‌شود. نامساوی سوم را می‌توان از نتیجه نامساوی جنسن که بیان می‌کند برای $m \geq 1$ و $p \geq 1$ داریم $\sum_{i=1}^m |a_i|^p \leq m^{p-1} \sum_{i=1}^m |a_i|^p$ ، نتیجه گرفت. نامساوی چهارم از بسط تیلور و نامساوی آخر از شرط (۵.A) نتیجه می‌شوند. بنابراین Z_i ها در شرط نامساوی برنشتاین صدق می‌کنند و داریم

$$\begin{aligned}
 \mathbb{P}(|\bar{S}_{nk}(\hat{\theta}_{n\circ})| > \frac{\lambda_n}{\sqrt{\log n}}) & \leq 2 \exp \left[-\frac{1}{2} \frac{n^2 \lambda_n^2 / (36(\log n))^2}{n\delta + M^* n \lambda_n / (\sqrt{\log n})} \right] \\
 & \leq 2 \exp \left[-C \frac{n \lambda_n^2}{(\log n)^2} \right].
 \end{aligned}$$

بر اساس شرط (۷.A) $\log p_n = o(n \lambda_n^2 / (\log n)^2)$ و $n \lambda_n^2 / (\log n)^2 \rightarrow \infty$ برقرار هستند، بنابراین داریم

$$I_{n1} \leq 2 \exp \left[\log p_n - C \frac{n \lambda_n^2}{(\log n)^2} \right] = o(1),$$

و $I_{n1} = o(1)$ اثبات می‌شود.

سپس نشان خواهیم داد که $I_{n2} = o(1)$. فرض کنید $H_{nk1} = (H_{nk1}, \dots, H_{nks_n})^\top$ زیربردارى باشد که s_n عضو اولیه بردار H_{nk} را شامل می‌شود. E_{nk1} و G_{nk1} به طور مشابه تعریف می‌شوند. حال داریم

$$\begin{aligned} I_{n2} &= \mathbb{P}\left(\max_{s_n+1 \leq k \leq p_n} |D_{k1}(\theta_{n0}, u_i)(\hat{\theta}_{n1} - \theta_{n10})| > \frac{\lambda_n}{\epsilon \log n}\right) \\ &= \mathbb{P}\left(\max_{s_n+1 \leq k \leq p_n} |D_{k1}(\theta_{n0}, u_i)(\hat{\theta}_{n1} - \theta_{n10})| > \frac{\lambda_n}{\epsilon \log n}, \|\hat{\theta}_{n1} - \theta_{n10}\| \leq \sqrt{s_n/n} \log n\right) \\ &\quad + \mathbb{P}(\|\hat{\theta}_{n1} - \theta_{n10}\| > \sqrt{s_n/n} \log n) \\ &\leq \mathbb{P}\left(\max_{s_n+1 \leq k \leq p_n} \|D_{k1}(\theta_{n0}, u_i)\| > \frac{\lambda_n \sqrt{n}}{\epsilon \sqrt{s_n} (\log n)^2}\right) + o(1) \\ &\leq \mathbb{P}\left(\max_{s_n+1 \leq k \leq p_n} \|H_{nk1}(\theta_{n0}, u_i)\| > \frac{\lambda_n \sqrt{n}}{\epsilon \sqrt{s_n} (\log n)^2}\right) \\ &\quad + \mathbb{P}\left(\max_{s_n+1 \leq k \leq p_n} \|E_{nk1}(\theta_{n0}, u_i)\| > \frac{\lambda_n \sqrt{n}}{\epsilon \sqrt{s_n} (\log n)^2}\right) \\ &\quad + \mathbb{P}\left(\max_{s_n+1 \leq k \leq p_n} \|G_{nk1}(\theta_{n0}, u_i)\| > \frac{\lambda_n \sqrt{n}}{\epsilon \sqrt{s_n} (\log n)^2}\right) + o(1) \\ &= I_{n21} + I_{n22} + I_{n23} + o(1). \end{aligned}$$

برای ارزیابی I_{n21} ، تحت شرایط (۵.A)، (۴.A) و (۶.A) و اینکه $|H_{nkk'}(\theta_{n0}, u_i)|$ به طور یکنواخت کراندار است، می‌توان نشان داد که

$$\begin{aligned} I_{n21} &\leq \mathbb{P}\left(\max_{s_n+1 \leq k \leq p_n} \|H_{nk1}(\theta_{n0}, u_i)\|^2 > C \frac{n \lambda_n^2}{s_n (\log n)^4}\right) \\ &\leq \mathbb{P}\left(\max_{s_n+1 \leq k \leq p_n} \left| \|H_{nk1}(\theta_{n0}, u_i)\|^2 - \mathbb{E} \|H_{nk1}(\theta_{n0}, u_i)\|^2 \right| \right. \\ &\quad \left. + \max_{s_n+1 \leq k \leq p_n} \|H_{nk1}(\theta_{n0}, u_i)\|^2 > C \frac{n \lambda_n^2}{s_n (\log n)^4}\right). \end{aligned}$$

بنابراین داریم $\max_{s_n+1 \leq k \leq p_n} \|H_{nk1}(\theta_{n0}, u_i)\|^2 = \max_{s_n+1 \leq k \leq p_n} \mathbb{E} (\sum_{k'=1}^{s_n} H_{nkk'}(\theta_{n0}, u_i)) \leq C s_n$ آنجایی که تحت شرط (۷.A)، $s_n^2 (\log n)^4 = o(n \lambda_n^2)$ و $p_n s_n^3 (\log n)^4 / (n^2 \lambda_n^4) = o(1)$ برقرار است، داریم

$$\begin{aligned} I_{n21} &\leq \mathbb{P}\left(\max_{s_n+1 \leq k \leq p_n} \left| \|H_{nk1}(\theta_{n0}, u_i)\|^2 - \mathbb{E} \|H_{nk1}(\theta_{n0}, u_i)\|^2 \right| > \frac{C}{2} \frac{n \lambda_n^2}{s_n (\log n)^4}\right) \\ &\leq \sum_{k=s_n+1}^{p_n} \mathbb{P}\left(\left| \|H_{nk1}(\theta_{n0}, u_i)\|^2 - \mathbb{E} \|H_{nk1}(\theta_{n0}, u_i)\|^2 \right| > \frac{C}{2} \frac{n \lambda_n^2}{s_n (\log n)^4}\right) \\ &\leq C \sum_{k=s_n+1}^{p_n} \frac{\mathbb{E} [\sum_{k'=1}^{s_n} (H_{nkk'}^2(\theta_{n0}, u_i) - \mathbb{E} (H_{nkk'}^2(\theta_{n0}, u_i)))]^2 s_n^2 (\log n)^4}{n^2 \lambda_n^4} \\ &= O(p_n s_n^3 (\log n)^4 / (n^2 \lambda_n^4)) = o(1), \end{aligned}$$

به طوری که نامساوی سوم از نامساوی مارکوف نتیجه می‌شود. به طور مشابه می‌توان نشان داد که $I_{n22} = o(1)$ و $I_{n23} = o(1)$. لذا عبارت $I_{n2} = o(1)$ اثبات می‌شود.

در پایان اثبات می‌کنیم که $I_{n3} = o(1)$. نا مساوی‌های زیر را در نظر بگیرید

$$\begin{aligned} I_{n3} &\leq \mathbb{P}\left(\max_{s_n+1 \leq k \leq p_n} |(\hat{\theta}_{n1} - \theta_{n1\circ})^\top \nabla_{k1}(\theta_n^*, u_i)(\hat{\theta}_{n1} - \theta_{n1\circ})| > \frac{\lambda_n}{\epsilon \log n}\right) \\ &\leq \mathbb{P}\left(\max_{s_n+1 \leq k \leq p_n} |(\hat{\theta}_{n1} - \theta_{n1\circ})^\top \nabla_{k1}(\theta_n^*, u_i)(\hat{\theta}_{n1} - \theta_{n1\circ})| > \frac{\lambda_n}{\epsilon \log n}, \right. \\ &\quad \left. \|\hat{\theta}_{n1} - \theta_{n1\circ}\| \leq \sqrt{s_n/n \log n}\right) + \mathbb{P}\left(\|\hat{\theta}_{n1} - \theta_{n1\circ}\| > \sqrt{s_n/n \log n}\right) \\ &\leq \sum_{k=s_n+1}^{p_n} \mathbb{P}\left(\text{tr}(\nabla_{k1}(\theta_n^*, u_i)) > \frac{n\lambda_n}{s_n(\log n)^3}\right) + o(1) \\ &\leq C \sum_{k=s_n+1}^{p_n} \frac{\mathbb{E}[\text{tr}(\nabla_{k1}(\theta_n^*, u_i))^2] s_n^2 (\log n)^\epsilon}{n^2 \lambda_n^2} + o(1), \end{aligned}$$

به طوری که نامساوی سوم از عبارت (ب.۹) و نامساوی آخر از نامساوی مارکوف نتیجه می‌شوند. تحت شرایط (۲.۸)، (۴.۸)، (۵.۸) و (۶.۸) داریم

$$\mathbb{E}[\text{tr}(\nabla_{k1}(\theta_n^*, u_i))^2] = \mathbb{E}\left[\sum_{k'=1}^{s_n} \frac{\partial^2 \bar{S}_{nk}}{\partial \beta_{nk'}^2}(\theta_n^*, u_i)\right]^2 \leq C s_n^2.$$

از آنجایی که تحت شرط (۷.۸) $I_{n3} = O(p_n s_n^2 (\log n)^\epsilon)$ برقرار است لذا $I_{n3} = (n^2 \lambda_n^2) = o(1)$ $(n^2 \lambda_n^2) + o(1) = o(1)$

با جمع‌آوری همه شرایط مورد نیاز، قضیه اثبات می‌شود.

ب.۴.۳ اثبات قضیه ۲.۱.۳

قسمت (i): کافی است نشان دهیم، به ازای هر $\epsilon > 0$ ، یک ثابت $\Delta > 0$ وجود دارد که برای n به اندازه کافی بزرگ، داریم

$$\mathbb{P}\left(\sup_{\|\theta_n - \theta_{n\circ}\| = \Delta \sqrt{p_n/n}} (\theta_n - \theta_{n\circ})^\top U_n(\theta_n) < 0\right) \geq 1 - \epsilon.$$

لذا لازم است که علامت $(\theta_n - \theta_{n\circ})^\top U_n(\theta_n)$ را روی $\{\theta_n : \|\theta_n - \theta_{n\circ}\| = \Delta \sqrt{p_n/n}\}$ بررسی کنیم. به ازای θ_n^* بین θ_n and $\theta_{n\circ}$ ، یعنی به ازای $0 < t < 1$ داشته باشیم $\theta_n^* = t\theta_n + (1-t)\theta_{n\circ}$ ، داریم

$$\begin{aligned} (\theta_n - \theta_{n\circ})^\top U_n(\theta_n) &= (\theta_n - \theta_{n\circ})^\top U_n(\theta_{n\circ}) - (\theta_n - \theta_{n\circ})^\top D_n(\theta_n^*)(\theta_n - \theta_{n\circ}) \\ &= J_{n1} + J_{n2}. \end{aligned}$$

برای J_{n1} می‌توان نوشت

$$J_{n1} = (\theta_n - \theta_{n\circ})^\top \bar{U}_n(\theta_{n\circ}) + (\theta_n - \theta_{n\circ})^\top [U_n(\theta_{n\circ}) - \bar{U}_n(\theta_{n\circ})] = J_{n11} + J_{n12}.$$

بر اساس نامساوی کوشی شوارتز داریم $\|\bar{U}_n(\theta_{n^\circ})\| = \Delta \sqrt{p_n/n} \|\bar{U}_n(\theta_{n^\circ})\|$ و $|J_{n11}| \leq \|\theta_n - \theta_{n^\circ}\| \cdot \|\bar{U}_n(\theta_{n^\circ})\|$ بنابراین

$$\begin{aligned} & \mathbb{E} \left[\|\bar{U}_n(\theta_{n^\circ})\|^2 \right] \\ &= \mathbb{E} \left[\sum_{i=1}^n \epsilon_i^\top(\theta_{n^\circ}, \mathbf{u}_i) \bar{\mathbf{R}}^{-1} \mathbf{A}_i^\top(\theta_{n^\circ}, \mathbf{u}_i) \mathbf{D}_i \mathbf{D}_i^\top \mathbf{A}_i^\top(\theta_{n^\circ}, \mathbf{u}_i) \bar{\mathbf{R}}^{-1} \epsilon_i(\theta_{n^\circ}, \mathbf{u}_i) \right] \\ &\leq \sum_{i=1}^n \lambda_{\max}(\mathbf{D}_i \mathbf{D}_i^\top) \lambda_{\max}(\mathbf{A}_i(\theta_{n^\circ}, \mathbf{u}_i)) \lambda_{\max}(\bar{\mathbf{R}}^{-1}) \mathbb{E} \left[\epsilon_i^\top(\theta_{n^\circ}, \mathbf{u}_i) \epsilon_i(\theta_{n^\circ}, \mathbf{u}_i) \right] \\ &\leq C \operatorname{tr} \left(\sum_{i=1}^n \mathbf{D}_i \mathbf{D}_i^\top \right) = C \sum_{i=1}^n \sum_{j=1}^{n_i} \mathbf{D}_{ij}^\top \mathbf{D}_{ij} = O(np_n). \end{aligned}$$

از عبارات فوق $\|\bar{U}_n(\theta_{n^\circ})\| = O_p(\sqrt{np_n})$ و به دنبال آن $|J_{n11}| = \Delta O_p(p_n)$ را می‌توان نتیجه گرفت. برای عبارت J_{n12} بر اساس لم ب.۱.۳ می‌توان نشان داد که

$$|J_{n12}| \leq \|\theta_n - \theta_{n^\circ}\| \cdot \|U_n(\theta_{n^\circ}) - \bar{U}_n(\theta_{n^\circ})\| = \Delta \sqrt{p_n/n} O_p(p_n) = \Delta o_p(p_n).$$

بنابراین تساوی $|J_{n1}| = \Delta O_p(p_n)$ برقرار است. در ادامه J_{n2} را می‌توان به صورت

$$\begin{aligned} J_{n2} &= -(\theta_n - \theta_{n^\circ})^\top \bar{\mathbf{D}}_n(\theta_n^*)(\theta_n - \theta_{n^\circ}) - (\theta_n - \theta_{n^\circ})^\top [\mathbf{D}_n(\theta_n^*) - \bar{\mathbf{D}}_n(\theta_n^*)](\theta_n - \theta_{n^\circ}) \\ &= J_{n21} + J_{n22}, \end{aligned}$$

نشان داد. ابتدا بر اساس لم ب.۳.۳ می‌توان نشان داد که

$$\begin{aligned} |J_{n22}| &\leq \max \left(|\lambda_{\max}(\mathbf{D}_n(\theta_n^*) - \bar{\mathbf{D}}_n(\theta_n^*))|, |\lambda_{\min}(\mathbf{D}_n(\theta_n^*) - \bar{\mathbf{D}}_n(\theta_n^*))| \right) \times \|\theta_n - \theta_{n^\circ}\|^2 \\ &= O_p(\sqrt{np_n}) \Delta^2 \frac{p_n}{n} = \Delta^2 o_p(p_n). \end{aligned}$$

به عبارت دیگر

$$\begin{aligned} J_{n21} &= -(\theta_n - \theta_{n^\circ})^\top \bar{\mathbf{H}}_n(\theta_{n^\circ})(\theta_n - \theta_{n^\circ}) \\ &\quad - (\theta_n - \theta_{n^\circ})^\top [\mathbf{H}_n(\theta_n^*) - \bar{\mathbf{H}}_n(\theta_{n^\circ})](\theta_n - \theta_{n^\circ}) \\ &\quad - (\theta_n - \theta_{n^\circ})^\top [\bar{\mathbf{D}}_n(\theta_n^*) - \bar{\mathbf{H}}_n(\theta_n^*)](\theta_n^*)(\theta_n - \theta_{n^\circ}) \\ &= J_{n21}^a + J_{n21}^b + J_{n21}^c. \end{aligned}$$

بر اساس لم ب.۵.۳، عبارت $J_{n21}^b = \Delta^2 o_p(p_n)$ ؛ و بر اساس لم ب.۴.۳، عبارت $J_{n21}^c = \Delta^2 o_p(p_n)$ نتیجه می‌شود. لذا کافی است که عبارت J_{n21}^a را بررسی کنیم. تحت شرط (۳.۴) داریم

$$\begin{aligned} J_{n21}^a &= -(\theta_n - \theta_{n^\circ})^\top \left[\sum_{i=1}^n \mathbf{D}_i^\top \mathbf{A}_i^\top(\theta_{n^\circ}, \mathbf{u}_i) \bar{\mathbf{R}}^{-1} \mathbf{A}_i^\top(\theta_{n^\circ}, \mathbf{u}_i) \mathbf{D}_i \right] (\theta_n - \theta_{n^\circ}) \\ &\leq \lambda_{\min}(\bar{\mathbf{R}}^{-1}) \min_i \lambda_{\min}(\mathbf{A}_i(\theta_{n^\circ}, \mathbf{u}_i)) \lambda_{\min} \left(\sum_{i=1}^n \mathbf{D}_i^\top \mathbf{D}_i \right) \|\theta_n - \theta_{n^\circ}\|^2 \leq -C \Delta^2 p_n. \end{aligned}$$

بنابراین در فضای پارامتر $(\theta_n - \theta_{n^0})^\top U_n(\theta_n)$ ، عبارت $(\theta_n - \theta_{n^0})^\top U_n(\theta_n)$ در احتمال به عبارت $J_{n1} + J_{n2}^a$ همگراست که مقداری منفی است. لذا نتیجه اثبات می‌شود. قسمت (ii): با استفاده از قضیه ۱۲.۷ در اسچوماکر (۱۹۸۱) تحت شرط (۱.۵) مقدار ثابتی مانند C وجود دارد به طوری که

$$\sup_{t \in [\circ, 1]} |f_\circ(t) - B(t)\alpha_\circ| \leq CK_n^{-m}.$$

حال با انتخاب $K_n \approx n^{1/(2r+1)}$ داریم

$$\begin{aligned} (\hat{f}(t) - f_\circ(t))^2 &= (B(t)\hat{\alpha}_n - f_\circ(t))^2 \simeq (B(t)\alpha_{n^0} - f_\circ(t))^2 \\ &\leq |B(t)\alpha_{n^0} - f_\circ(t)| |B(t)\alpha_{n^0} - f_\circ(t)| \\ &\leq \sup_{t \in [\circ, 1]} |B(t)\alpha_{n^0} - f_\circ(t)| \sup_{t \in [\circ, 1]} |B(t)\alpha_{n^0} - f_\circ(t)| \\ &\leq CL_n^{-r} CL_n^{-r} = C^2 n^{-2r/(2r+1)} = O_p(n^{-2r/(2r+1)}), \end{aligned}$$

و اثبات کامل می‌شود.

۵.۳.۱.۳ اثبات قضیه ۳.۱.۳

قسمت (i): ابتدا نشان می‌دهیم که $\hat{\theta}_n$ در شرط تنکی، $\hat{\theta}_n = \circ$ صدق می‌کند. برای این که نشان دهیم به ازای $\hat{\theta}_n = \circ$ جواب بهینه بدست می‌آید، کافی است نشان دهیم هنگامی که $n \rightarrow \infty$ برای هر $\hat{\theta}_n$ و $\hat{\theta}_{n^0}$ که به ترتیب در شرایط $\|\hat{\theta}_n - \hat{\theta}_{n^0}\| = O_p(\sqrt{p_n/n})$ و $\|\hat{\theta}_n\| \leq C(\sqrt{p_n/n})$ صدق می‌کنند، عبارت $U_{nk}(\theta)$ برای $\beta_k \in (-C(\sqrt{p_n/n}), C(\sqrt{p_n/n}))$ که $k = s_n + 1, \dots, p_n$ با احتمال همگرا به ۱ دارای علامت‌های متفاوت هستند. عبارت $U_{nk}(\theta)$ را به صورت

$$S_{nk}(\hat{\theta}) = O_p(\sqrt{p_n/n}),$$

در نظر بگیرید، بنابراین می‌توان نشان داد که

$$U_{nk}(\theta) = n\lambda_n \left\{ \frac{q_{\lambda_n}(|\beta_k|)}{\lambda_n} \text{sgn}(\beta_k) + O_p(\sqrt{p_n/n}) \right\}.$$

از آنجایی که $\sqrt{p_n/n}/\lambda_n \rightarrow \circ$ و $\liminf_{n \rightarrow \infty} \liminf_{\beta_k \rightarrow \circ+} \frac{q_{\lambda_n}(|\beta_k|)}{\lambda_n} > \circ$ ، می‌توان نشان داد که علامت β_k تعیین کننده علامت $U_{nk}(\theta)$ است. لذا نتیجه مورد نظر حاصل می‌شود. قسمت (ii): ابتدا نشان می‌دهیم که $(\hat{\theta}_n - \theta_{n^0})^\top \bar{H}_n(\theta_{n^0}) \xrightarrow{d} \mathcal{N}_{s_n}(\circ, 1)$ می‌توان نشان داد

$$\begin{aligned} &\xi_n^\top \bar{M}_n^{-1/2}(\theta_{n^0}) \bar{S}_n(\theta_{n^0}) \\ &= \xi_n^\top \bar{M}_n^{-1/2}(\theta_{n^0}) S_n(\theta_{n^0}) + \xi_n^\top \bar{M}_n^{-1/2}(\theta_{n^0}) [\bar{S}_n(\theta_{n^0}) - S_n(\theta_{n^0})] \\ &= \xi_n^\top \bar{M}_n^{-1/2}(\theta_{n^0}) D_n(\theta_n^*)(\hat{\theta}_n - \theta_{n^0}) + \xi_n^\top \bar{M}_n^{-1/2}(\theta_{n^0}) [\bar{S}_n(\theta_{n^0}) - S_n(\theta_{n^0})] \end{aligned}$$

$$\begin{aligned}
 &= \xi_n^\top \overline{M}_n^{-1/2}(\theta_{n^\circ}) \overline{H}_n(\theta_{n^\circ}) (\widehat{\theta}_n - \theta_{n^\circ}) \\
 &\quad + \xi_n^\top \overline{M}_n^{-1/2}(\theta_{n^\circ}) [D_n(\theta_n^*) - \overline{H}_n(\theta_{n^\circ})] (\widehat{\theta}_n - \theta_{n^\circ}) \\
 &\quad + \xi_n^\top \overline{M}_n^{-1/2}(\theta_{n^\circ}) [\overline{S}_n(\theta_{n^\circ}) - S_n(\theta_{n^\circ})],
 \end{aligned}$$

به طوری که برای رسیدن به تساوی دوم فرض می‌شود $S_n(\widehat{\theta}_n) = 0$ و برای θ_n^* بین $\widehat{\theta}_n$ و θ_{n° و بر اساس بسط تیلور داریم $S_n(\theta_{n^\circ}) = D_n(\theta_n^*) (\widehat{\theta}_n - \theta_{n^\circ})$. بر اساس لم ب.۳.۶، داریم $\xi_n^\top \overline{M}_n^{-1/2}(\theta_{n^\circ}) \overline{S}_n(\theta_{n^\circ}) \xrightarrow{D} N_{s_n}(0, 1)$ بنابراین برای اثبات قضیه کافی است نشان دهیم که به ازای هر $\Delta > 0$ ، عبارت‌های

$$\sup_{\|\theta_n - \theta_{n^\circ}\| \leq \Delta \sqrt{p_n/n}} |\xi_n^\top \overline{M}_n^{-1/2}(\theta_{n^\circ}) [D_n(\theta_n) - \overline{H}_n(\theta_{n^\circ})] (\widehat{\theta}_n - \theta_{n^\circ})| = o_p(1) \quad (\text{ب.۱۴})$$

و

$$|\xi_n^\top \overline{M}_n^{-1/2}(\theta_{n^\circ}) [\overline{S}_n(\theta_{n^\circ}) - S_n(\theta_{n^\circ})]| = o_p(1) \quad (\text{ب.۱۵})$$

برقرار است. ابتدا عبارت (ب.۱۵) را اثبات می‌کنیم. در نظر بگیرید که

$$\begin{aligned}
 &[\xi_n^\top \overline{M}_n^{-1/2}(\theta_{n^\circ}) [\overline{S}_n(\theta_{n^\circ}) - S_n(\theta_{n^\circ})]]^2 \\
 &= \xi_n^\top \overline{M}_n^{-1/2}(\theta_{n^\circ}) [\overline{S}_n(\theta_{n^\circ}) - S_n(\theta_{n^\circ})] [\overline{S}_n(\theta_{n^\circ}) - S_n(\theta_{n^\circ})]^\top \overline{M}_n^{-1/2} \xi_n \\
 &\leq \lambda_{\max}(\overline{M}_n^{-1}(\theta_{n^\circ})) \lambda_{\max}([\overline{S}_n(\theta_{n^\circ}) - S_n(\theta_{n^\circ})] [\overline{S}_n(\theta_{n^\circ}) - S_n(\theta_{n^\circ})]^\top) \\
 &\leq \frac{\|\overline{S}_n(\theta_{n^\circ}) - S_n(\theta_{n^\circ})\|^2}{\lambda_{\min}(\overline{M}_n(\theta_{n^\circ}))} \\
 &\leq \frac{\|\overline{S}_n(\theta_{n^\circ}) - S_n(\theta_{n^\circ})\|^2}{C \lambda_{\min}(\sum_{i=1}^n \mathbf{X}_i^\top \mathbf{X}_i)} = O_p(p_n^2/n) = o_p(1).
 \end{aligned}$$

بر اساس لم ب.۳.۱۰ و این که

$$\lambda_{\min}(\overline{M}_n(\theta_{n^\circ})) \geq C \lambda_{\min}(\sum_{i=1}^n \mathbf{X}_i^\top \mathbf{X}_i), \quad (\text{ب.۱۶})$$

می‌توان عبارت (ب.۱۵) را مشابه لم ب.۳.۶ اثبات کرد. حال به بررسی عبارت (ب.۱۴) می‌پردازیم. تساوی‌های و نامساوی‌های

$$\begin{aligned}
 &\sup_{\|\theta_n - \theta_{n^\circ}\| \leq \Delta \sqrt{p_n/n}} |\xi_n^\top \overline{M}_n^{-1/2}(\theta_{n^\circ}) [D_n(\theta_n) - \overline{H}_n(\theta_{n^\circ})] (\widehat{\theta}_n - \theta_{n^\circ})| \\
 &\leq \sup_{\|\theta_n - \theta_{n^\circ}\| \leq \Delta \sqrt{p_n/n}} |\xi_n^\top \overline{M}_n^{-1/2}(\theta_{n^\circ}) [D_n(\theta_n) - \overline{D}_n(\theta_{n^\circ})] (\widehat{\theta}_n - \theta_{n^\circ})| \\
 &\quad + \sup_{\|\theta_n - \theta_{n^\circ}\| \leq \Delta \sqrt{p_n/n}} |\xi_n^\top \overline{M}_n^{-1/2}(\theta_{n^\circ}) [\overline{D}_n(\theta_n) - \overline{H}_n(\theta_{n^\circ})] (\widehat{\theta}_n - \theta_{n^\circ})| \\
 &\quad + \sup_{\|\theta_n - \theta_{n^\circ}\| \leq \Delta \sqrt{p_n/n}} |\xi_n^\top \overline{M}_n^{-1/2}(\theta_{n^\circ}) [\overline{H}_n(\theta_n) - \overline{H}_n(\theta_{n^\circ})] (\widehat{\theta}_n - \theta_{n^\circ})| \\
 &= K_{n1} + K_{n2} + K_{n3},
 \end{aligned}$$

را در نظر بگیرید. براساس نامساوی کوشی شوارتز و نتیجه ب.۱۰.۳، داریم

$$\begin{aligned}
 K_{n1} &\leq \sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} [\xi_n^\top \bar{M}_n^{-1/2}(\theta_{n^0})(D_n(\theta_n) - \bar{D}_n(\theta_n))^2 \\
 &\quad \times \bar{M}_n^{-1/2}(\theta_{n^0})\xi_n]^{1/2} \|\hat{\theta}_n - \theta_{n^0}\| \\
 &\leq \sup_{\|\theta_n - \theta_{n^0}\| \leq \Delta \sqrt{p_n/n}} \max(|\lambda_{\min}(D_n(\theta_n) - \bar{D}_n(\theta_n))|, |\lambda_{\max}(D_n(\theta_n) - \bar{D}_n(\theta_n))|) \\
 &\quad \times \lambda_{\min}^{-1/2}(\bar{M}_n(\theta_{n^0})) O_p(p_n^{1/2} n^{-1/2}) \\
 &= O_p(\sqrt{n} p_n) O(n^{-1/2}) O_p(p_n^{1/2} n^{-1/2}) = O_p(n^{-1/2} p_n^{3/2}) = o_p(1).
 \end{aligned}$$

بنابراین $K_{n1} = o_p(1)$. با استدلال مشابه لم‌های ب.۴.۳ و ب.۵.۳، می‌توان نشان داد که $K_{n2} = o_p(1)$ و $K_{n3} = o_p(1)$. لذا عبارت (ب.۱۴) اثبات می‌شود. تا این مرحله نشان دادیم که

$$\xi_n^\top \bar{M}_n^{-1/2}(\theta_{n^0}) \bar{H}_n(\theta_{n^0})(\hat{\theta}_{n1} - \theta_{n^01}) \xrightarrow{D} \mathcal{N}_{s_n}(\circ, 1),$$

بنابراین $(\hat{\theta}_{n1} - \theta_{n^01}) \xrightarrow{D} \mathcal{N}_{s_n}(\circ, \bar{H}_n^{-1}(\theta_{n^0}) \bar{M}_n(\theta_{n^0}) \bar{H}_n^{-1}(\theta_{n^0}))$ بدیهی است که

$$(\hat{\theta}_{n1} - \theta_{n^01}) = \begin{pmatrix} \hat{\beta}_{n1} - \beta_{n^01} \\ \hat{\alpha}_{n1} - \alpha_{n^01} \end{pmatrix} \xrightarrow{D} N_2 \left(\begin{pmatrix} 0_{s_n} \\ 0_{h_n} \end{pmatrix}, \begin{pmatrix} Q_1 & Q_2 \\ Q_3 & Q_4 \end{pmatrix} \right).$$

در نتیجه برای اثبات قضیه کافی است نشان دهیم که $Q_1 = \bar{H}_n^{*-1}(\beta_{n^0}) \bar{M}_n^*(\beta_{n^0}) \bar{H}_n^{*-1}(\beta_{n^0})$ را به صورت

$$\begin{aligned}
 \bar{H}_n &= \begin{pmatrix} X^\top \\ B^\top \end{pmatrix} \Omega \begin{pmatrix} X & B \end{pmatrix} = \begin{pmatrix} X^\top \Omega X & X^\top \Omega B \\ B^\top \Omega X & B^\top \Omega B \end{pmatrix}, \\
 \bar{M}_n &= \begin{pmatrix} X^\top \\ B^\top \end{pmatrix} \Omega_\circ \begin{pmatrix} X & B \end{pmatrix} = \begin{pmatrix} X^\top \Omega_\circ X & X^\top \Omega_\circ B \\ B^\top \Omega_\circ X & B^\top \Omega_\circ B \end{pmatrix},
 \end{aligned}$$

بازتعریف می‌کنیم، به طوری که $\Omega_\circ = \text{diag}\{\Omega_{\circ i}\}$ و $\Omega_{\circ i} = E_{u|y} [A_i^\top(\theta_n, u_i) \bar{R}^{-1} R_\circ \bar{R}^{-1} A_i^\top(\theta_n, u_i)]$ این نشان می‌دهد که

$$\bar{H}_n^{-1} = \begin{pmatrix} (X^\top \Omega X)^{-1} + (X^\top \Omega X)^{-1} B^\top \Omega X Q B^\top \Omega X (X^\top \Omega X)^{-1}, & Q^* \\ Q^{*\top}, & Q \end{pmatrix},$$

به طوری که $Q^* = -(X^\top \Omega X)^{-1} X^\top \Omega B Q$ و $Q = (B^\top \Omega B - B^\top \Omega X (X^\top \Omega X)^{-1} X^\top \Omega B)^{-1}$

با کمی محاسبات جبری می‌توان Q_1 را به صورت

$$\begin{aligned} Q_1 = & (I - Q^* B^\top \Omega X)(X^\top \Omega X)^{-1} \\ & \times \left[X^\top \Omega \circ X + X^\top \Omega \circ B Q^{*\top} (X^\top \Omega X)(I - Q^* B^\top \Omega X)^{-1} \right. \\ & + (X^\top \Omega X)(I - Q^* B^\top \Omega X)^{-1} Q^* B^\top \Omega \circ X \\ & \left. + (X^\top \Omega X)(I - Q^* B^\top \Omega X)^{-1} Q^* B^\top \Omega \circ B Q^{*\top} (X^\top \Omega X)(I - Q^* B^\top \Omega X)^{-1} \right] \\ & \times (I - Q^* B^\top \Omega X)(X^\top \Omega X)^{-1}, \end{aligned}$$

بیان کرد. بنابراین کافی است نشان دهیم که

$$(I - Q^* B^\top \Omega X)(X^\top \Omega X)^{-1} = \left(X^\top \Omega X (I - Q^* B^\top \Omega X)^{-1} \right)^{-1} = \overline{H}_n^{*-1} \quad (\text{ب.۱۷})$$

و

$$\begin{aligned} X^\top \Omega \circ X + X^\top \Omega \circ B Q^{*\top} (X^\top \Omega X)(I - Q^* B^\top \Omega X)^{-1} & \quad (\text{ب.۱۸}) \\ + (X^\top \Omega X)(I - Q^* B^\top \Omega X)^{-1} Q^* B^\top \Omega \circ X \\ + (X^\top \Omega X)(I - Q^* B^\top \Omega X)^{-1} Q^* B^\top \Omega \circ B Q^{*\top} (X^\top \Omega X)(I - Q^* B^\top \Omega X)^{-1} \\ = \overline{M}_n^*. \end{aligned}$$

هر دو عبارت (ب.۱۷) و (ب.۱۸) شامل جمله $(I - Q^* B^\top \Omega X)^{-1}$ هستند که

$$(I - Q^* B^\top \Omega X)^{-1} = I - (X^\top \Omega X)^{-1} X^\top \Omega B (B^\top \Omega B)^{-1} B^\top \Omega X.$$

با استفاده از

$$X^\top \Omega X (I - Q^* B^\top \Omega X)^{-1} = X^\top \Omega X - X^\top \Omega B (B^\top \Omega B)^{-1} B^\top \Omega X = \overline{H}_n^*,$$

عبارت (ب.۱۷) اثبات می‌شود. همچنین با استفاده از

$$X^\top \Omega X (I - Q^* B^\top \Omega X)^{-1} Q^* = -X^\top \Omega B (B^\top \Omega B)^{-1},$$

عبارت (ب.۱۸) را می‌توان به صورت

$$\begin{aligned} X^\top \Omega \circ X - X^\top \Omega \circ B (B^\top \Omega B)^{-1} B^\top \Omega X \\ - X^\top \Omega B (B^\top \Omega B)^{-1} B^\top \Omega \circ X \\ + X^\top \Omega B (B^\top \Omega B)^{-1} B^\top \Omega \circ B (B^\top \Omega B)^{-1} B^\top \Omega X = \overline{M}_n^*, \end{aligned}$$

بیان کرد. لذا (ب.۱۸) نیز برقرار بوده و اثبات قضیه کامل می‌شود.

ب.۴ اثبات قضایای فصل ۴

ب.۴.۱ بردار رتبه و ماتریس اطلاع فشر

مشتق اول جزئی از لگاریتم تابع درستنمایی تاوانیده نسبت به هر یک از اعضای پارامتر Θ را بردار رتبه (S_Θ) می‌گویند که اعضای آن شامل

$$\begin{aligned} S_\beta &= \sum_{i=1}^n \left(\frac{\nu + n_i}{\nu + \Delta_i} \right) \mathbf{X}_i^\top \Lambda_i^{-1} (\mathbf{y}_i - \mathbf{X}_i \beta), \\ [S_\omega]_l &= -\frac{1}{\nu} \sum_{i=1}^n \left\{ \text{tr}(\Lambda_i^{-1} \mathring{\Lambda}_{il}) - \left(\frac{\nu + n_i}{\nu + \Delta_i} \right) (\mathbf{y}_i - \mathbf{X}_i \beta)^\top \Lambda_i^{-1} \mathring{\Lambda}_{il} \Lambda_i^{-1} (\mathbf{y}_i - \mathbf{X}_i \beta) \right\}, \\ s_\nu &= \frac{1}{\nu} \sum_{i=1}^n \left\{ \mathcal{D}_g \left(\frac{\nu + n_i}{\nu} \right) - \mathcal{D}_g \left(\frac{\nu}{\nu} \right) - \frac{n_i}{\nu} - \log \left(1 + \frac{\Delta_i}{\nu} \right) + \frac{(\nu + n_i) \Delta_i}{(\nu + \Delta_i) \nu} \right\}, \end{aligned}$$

است، که در آن $\mathcal{D}_g(x) = \partial \Gamma(x) / \partial x$ ، $g = q(q+1)/2 + r(r+1)/2 + 2$ ، $l = 1, \dots, g$ بوده و

$$\mathring{\Lambda}_{il} = \frac{\partial \Delta_i(\mathbf{D}, \Sigma, \phi)}{\partial \omega_l} = \begin{cases} \mathbf{Z}_i \frac{\partial \mathbf{D}}{\partial \omega_l} \mathbf{Z}_i^\top & \text{if } \omega_l = \text{vec}(\mathbf{D}), \\ \frac{\partial \Sigma}{\partial \omega_l} \otimes \mathbf{\Omega}_i & \text{if } \omega_l = \text{vec}(\Sigma), \\ \Sigma \otimes \frac{\partial \mathbf{\Omega}_i}{\partial \omega_l} & \text{if } \omega_l = \phi. \end{cases}$$

همچنین امید مشتق دوم جزئی از لگاریتم تابع درستنمایی تاوانیده نسبت به هر یک از اعضای پارامتر Θ را ماتریس اطلاع فشر (J_Θ) می‌گویند که اعضای آن به صورت

$$J_{\Theta\Theta} = \begin{bmatrix} J_{\beta\beta} & J_{\beta\eta} \\ J_{\eta\beta}^\top & J_{\eta\eta} \end{bmatrix},$$

تعیین می‌شود، به طوری که

$$J_{\beta\beta} = \sum_{i=1}^n \frac{(\nu + n_i)}{(\nu + n_i + 2)} \mathbf{X}_i^\top \Lambda_i^{-1} \mathbf{X}_i, \quad J_{\beta\eta} = \mathbf{m}^\circ, \quad \text{and} \quad J_{\eta\eta} = \begin{bmatrix} J_{\omega\omega} & J_{\omega\nu} \\ J_{\nu\omega}^\top & J_{\nu\nu} \end{bmatrix}$$

و

$$\begin{aligned} [J_{\omega\omega}]_{ls} &= \frac{1}{\nu} \sum_{i=1}^n \frac{1}{\nu + n_i + 2} \left((\nu + n_i) \text{tr}(\Lambda_i^{-1} \mathring{\Lambda}_{il} \Lambda_i^{-1} \mathring{\Lambda}_{is}) - \text{tr}(\Lambda_i^{-1} \mathring{\Lambda}_{il}) \text{tr}(\Lambda_i^{-1} \mathring{\Lambda}_{is}) \right), \\ [J_{\omega\nu}]_l &= -\sum_{i=1}^n \frac{1}{(\nu + n_i)(\nu + n_i + 2)} \text{tr}(\Lambda_i^{-1} \mathring{\Lambda}_{il}), \\ J_{\nu\nu} &= \frac{1}{\nu} \sum_{i=1}^n \left(\mathcal{T}_g \left(\frac{\nu}{\nu} \right) - \mathcal{T}_g \left(\frac{\nu + n_i}{\nu} \right) - \frac{2n_i(\nu + n_i + 2)}{\nu(\nu + n_i)(\nu + n_i + 2)} \right), \end{aligned}$$

که در آن $T_g(x) = \partial^2 \log \Gamma(x) / \partial x^2$ و $l, s = 1, \dots, g$ به علاوه $J_{\beta\beta}$ را می‌توان به صورت

$$J_{\beta\beta} = \begin{bmatrix} J_{\beta_1\beta_1} & J_{\beta_1\beta_2} \\ J_{\beta_2\beta_1}^\top & J_{\beta_2\beta_2} \end{bmatrix},$$

تفکیک کرد، به طوری که

$$\begin{aligned} J_{\beta_1\beta_1} &= \sum_{i=1}^n \frac{(\nu+n_i)}{(\nu+n_i+2)} X_{i1}^\top \Lambda_{i1}^{-1} X_{i1}, \\ J_{\beta_2\beta_2} &= \sum_{i=1}^n \frac{(\nu+n_i)}{(\nu+n_i+2)} X_{i2}^\top \Lambda_{i2}^{-1} X_{i2}, \\ J_{\beta_1\beta_2} &= \sum_{i=1}^n \frac{(\nu+n_i)}{(\nu+n_i+2)} X_{i1}^\top \Lambda_{i1}^{-1} X_{i2}; \end{aligned}$$

و عضوهای $\Lambda_{i12}, \Lambda_{i22}, \Lambda_{i11}$

$$\Lambda_i = \begin{bmatrix} \Lambda_{i11} & \Lambda_{i12} \\ \Lambda_{i12}^\top & \Lambda_{i22} \end{bmatrix}.$$

هستند.

ب. ۲.۴ اثبات قضیه ۲.۱.۴

مشابه فن و لی (۲۰۰۱)، می‌توان نشان داد که با احتمال همگرا به ۱، $\hat{\beta}$ در شرط

$$\|\hat{\beta} - \beta\| = O_p(n^{-1/2} + a_n).$$

صدق می‌کند. این نشان می‌دهد که نرخ همگرایی $\hat{\beta}$ به پارامتر تاوان λ_n وابسته است. برای اینکه نشان دهیم $\hat{\beta}$ یک برآوردگر $\sqrt{n}$ -سازگار است، فرض می‌کنم $a_n = O_p(n^{-1/2})$ و پارامتر تاوان λ_n را مقداری بسیار کوچک در نظر می‌گیریم. ابتدا نشان می‌دهیم که $\hat{\beta}$ در شرط تنگی، $\hat{\beta}_2 = 0$ ، صدق می‌کند. فرض کنید که U_k و S_k ، به ترتیب مشتق اول ℓ_c^{Pen} و ℓ_c نسبت به β_k هستند. برای این که نشان دهیم به ازای $\hat{\beta}_2 = 0$ ، جواب بهینه بدست می‌آید، کافی است نشان دهیم هنگامی که $n \rightarrow \infty$ ، برای هر β_1 که در شرط $\|\hat{\beta}_1 - \beta_1\| = O_p(n^{-1/2} + a_n)$ صدق می‌کند، عبارت $U_k(\beta)$ برای $\beta_k \in (-C(n^{-1/2} + a_n), C(n^{-1/2} + a_n))$ که $k = s+1, \dots, p$ با احتمال همگرا به ۱ دارای علامت‌های متفاوت هستند. عبارت $U_k(\beta)$ را به صورت

$$U_k(\beta) = S_k(\hat{\beta}) + nq_{\lambda_n}(|\beta_k|) \text{sgn}(\beta_k),$$

در نظر بگیرید. می‌توان نشان داد $S_k(\hat{\beta}) = O_p(n^{-1/2} + a_n)$. بنابراین داریم

$$U_k(\beta) = n\lambda_n \left\{ \frac{q_{\lambda_n}(|\beta_k|)}{\lambda_n} \text{sgn}(\beta_k) + O_p(n^{-1/2} + a_n) \right\}.$$

از آنجایی که $O_p(n^{-1/2} + a_n) \rightarrow 0$ و $\liminf_{n \rightarrow \infty} \liminf_{\beta_k \rightarrow 0^+} \frac{q_{\lambda_n}(|\beta_k|)}{\lambda_n} > 0$ ، می‌توان نشان داد که علامت β_k تعیین کننده علامت $U_{nk}(\theta)$ است. لذا نتیجه مورد نظر حاصل می‌شود.

ب.۳.۴ اثبات قضیه ۱.۲.۴

با استفاده از قضیه ۱۲.۷ در اسچوماکر (۱۹۸۱) می‌توان $g_j(t_i)$ را توسط $B_j(t_i)\alpha_j$ تقریب زد. سپس با انتخاب $K_j \approx N^{1/(2r_j+1)}$ داریم

$$\begin{aligned} (\hat{g}_j(t) - g_j(t))^2 &= |\hat{g}_j(t) - g_j(t)| |\hat{g}_j(t) - g_j(t)| \\ &\leq \sup_{t \in [\circ, 1]} |\hat{g}_j(t) - g_j(t)| \sup_{t \in [\circ, 1]} |\hat{g}_j(t) - g_j(t)| \\ &\leq C_j L_j^{-m_j} C_j L_j^{-m_j} = C_j^2 N^{2m_j/(2m_j+1)} = O_p(N^{2m_j/(2m_j+1)}), \end{aligned}$$

که بخش (i) قضیه اثبات می‌شود. قسمت (ii) قضیه مستقیماً از طریق قانون ضعیف اعداد بزرگ، قضیه حد مرکزی و قضیه اسلاتسکی اثبات می‌شود.

Abstract

This thesis considers the estimation methods in semiparametric mixed effects regression models for longitudinal data analysis. The model is a combination of mixed effects and semiparametric models. The nonparametric function of the model is estimated using kernel, backfitting and spline methods, then consider the multicollinearity problem and introduce a new estimator called the ridge estimator. Also considers the problem of simultaneous variable selection and estimation in the generalized semiparametric mixed effects model for gaussian and non-gaussian high dimensional data. A penalization type of generalized estimating equation method is proposed while using regression spline to approximate the nonparametric component. In addition, for analyzing multivariate longitudinal data in the presence of atypical observations or outliers, a robust method based on multivariate t-distribution is offered. In each study framework, the asymptotic behavior of proposed estimators is studied and for practical implementation, an appropriate iterative algorithm is developed. The performance of the proposed methods are compared through both simulation examples and real data sets.

Keywords: Smoothing Spline, Variable selection, Ridge estimator, High dimension, Penalized, Multivariate t-distribution, Asymptotic distribution, Outliers, Longitudinal data, Heavy-tailed distribution, Backfitting, Kernel, Generalized estimating equations, Semiparametric mixed effects model.



Shahrood University of Technology

Faculty of Mathematical Sciences

PhD Thesis in: Statistics

Estimation in Semiparametric Modeling for Longitudinal Data Analysis

By: Mozhgan Taavoni

Supervisor

Mohammad Arashi

August 2020