





دانشکده علوم ریاضی

پایان نامه کارشناسی ارشد آمار ریاضی

آزمون‌های استوار M در مدل‌های خطی با خطاهای زبرجمعی منفی

نگارنده: سمانه درویشی بهلولی

استاد راهنما

دکتر نگار اقبال

استاد مشاور

دکتر حسین باغیشنی

بهمن ۱۳۹۹

تقدیم با بوسه بر دستان پدرم که راه را
به من نشان داد و تقدیم به مادر عزیزتر از
جانم که چگونه رفتن را به من آموخت.

سپاس‌گزاری...

به مصداق لم یشکر المخلوق، لم یشکر الخالق سپاس بی‌کران نثار پدر و مادرم، که ناتوان شدند تا من توانا شوم، مویشان سپید شد تا من روسفید شوم و عاشقانه سوختند تا گرمابخش وجودم و روشنگر را هم باشند. و هم‌چنین از استاد گرامی سرکار خانم دکتر اقبال و استاد محترم جناب آقای دکتر باغی‌شنی به خاطر زحمات‌ها و راهنمایی‌هایی که در طول این پایان‌نامه برای من کشیده‌اند تقدیر و تشکر می‌کنم.

سمانه درویشی بهلولی

بهمن ۱۳۹۹

تعهد نامه

اینجانب **سمانه درویشی بهلولی** دانشجوی کارشناسی ارشد رشته **آمارریاضی** دانشکده علوم ریاضی دانشگاه صنعتی شاهرود، نویسنده پایان نامه با عنوان **آزمون های استوار M در مدل های خطی با خطاهای زبرجمعی منفی**، تحت راهنمایی **نگار اقبال** 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 پژوهش گران، 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.

سمانه درویشی بهلولی

بهمن 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:

در حالت کلی استواری یک روش آماری را می‌توان به عنوان حساس نبودن آن به انحراف کم از پذیره‌های پایه‌ای تعریف کرد. به عبارت دیگر هرگاه ویژگی‌های یک روش آماری تحت تاثیر انحراف کم از پذیره‌های اساسی قرار نگیرد، آن روش آماری را استوار گوییم. با توجه به اینکه برآوردهای کلاسیک ضرایب رگرسیونی در مدل‌های خطی به شدت تحت تاثیر داده‌های پرت قرار می‌گیرند، برآوردهای استوار مانند M – برآوردها پیشنهاد می‌شوند. همچنین سطح و توان آزمون‌های کلاسیک با حضور داده‌های پرت به شدت کاهش می‌یابند، برای حل این مشکل، آزمون‌های استوار معرفی می‌شوند. برای انجام آزمون استوار پارامترهای رگرسیونی، توزیع مجانبی آماره آزمون مشخص می‌شود و بر اساس پایایی برآوردهای مدل و توان آزمون، روش‌ها مورد ارزیابی قرار می‌گیرند. در این پایان نامه علاوه بر بررسی برخی ویژگی‌های اساسی متغیرهای تصادفی وابسته زبرجمعی منفی، آزمون‌های استوار M در مدل‌های خطی با خطای وابسته زبرجمعی منفی را تحت برخی شرایط مورد مطالعه قرار می‌دهیم و با استفاده از روش‌های شبیه‌سازی مونت کارلو به مقایسه روش‌های کلاسیک و استوار بر اساس پایایی برآوردها و توان آزمون می‌پردازیم.

کلمات کلیدی: آزمون استوار، وابسته زبرجمعی منفی، مدل خطی رگرسیون، شبیه‌سازی‌های مونت کارلو

فهرست مطالب

س	فهرست تصاویر
ف	فهرست جداول
ق	پیشگفتار
۱	۱ مفاهیم اولیه
۱	۱.۱ نابرابری ها
۳	۲.۱ انواع وابستگی
۳	۱.۲.۱ وابسته ربعی منفی
۳	۲.۲.۱ وابستگی پیوندی منفی
۴	۳.۲.۱ چند ویژگی متغیرهای تصادفی پیوندی منفی
۴	۴.۲.۱ مثال هایی از متغیرهای تصادفی پیوندی منفی
۵	۵.۲.۱ وابستگی زبرجمعی منفی
۸	۶.۲.۱ مثال ها و کاربردهایی از وابسته زبرجمعی منفی
۹	۳.۱ ترتیب های غیرتصادفی
۱۱	۲ روش های آماری استوار
۱۱	۱.۲ استواری
۱۳	۲.۲ برآورد استوار پارامتر مکان
۱۶	۳.۲ رگرسیون استوار
۱۸	۱.۳.۲ M-برآوردگر
۲۲	۲.۳.۲ L-برآوردگر
۲۳	۳ آزمون های استوار مدل خطی با خطاهای وابسته زبرجمعی منفی
۲۳	۱.۳ مقدمه

۴۱	۴	ارزیابی برآوردها و آزمون‌های M با مطالعه شبیه‌سازی
۴۱	۱.۴	مثال شبیه‌سازی
۴۷		مراجع
۵۳	آ	کدهای شبیه‌سازی جداول و نمودارهای فصل چهار
۵۳	۱.آ	کدهای مربوط به جدول دو و نمودار شکل چهار
۶۲	۲.آ	کدهای مربوط به جدول یک نمودارهای یک الی سه

فهرست تصاویر

۱.۴	نمودارهای هیستوگرام و منحنی‌های چگالی برازش شده باقی‌مانده‌های مدل رگرسیونی با خطای NSD و مستقل برای حجم نمونه $n = 1000$ به ازای یکی از روش‌های M-برآورد	۴۳
۲.۴	نمودارهای هیستوگرام و منحنی چگالی برازش شده باقی‌مانده‌های M-برآورد با متغیرهای توضیحی مختلف در دو سناریو و روش‌های مختلف M-برآورد به ازای حجم نمونه $n = 1000$	۴۴
۳.۴	مقایسه توابع توزیع NSD فرض شده برای باقی‌مانده‌های رگرسیونی با توزیع تجربی برآورد شده باقی‌مانده‌ها برای حجم نمونه $n = 1000$ و دو سناریو مختلف تولید متغیر توضیحی	۴۵
۴.۴	مقایسه توابع توزیع برازش شده $\hat{\lambda}_n \hat{\sigma}_n^2 M_n$ برای روش‌های M-برآورد مختلف و توزیع کی‌دو مرکزی با دو درجه آزادی به ازای اندازه نمونه $n = 1000$ و سناریوهای تولید متغیر توضیحی	۴۶

فهرست جداول

۱.۲	بعضی از توابعی که برای توزیع وزن مانده‌ها پیشنهاد شده‌اند	۲۱
۱.۴	ارزیابی ضرایب رگرسیون و پارامترهای زیاد	۴۴
۲.۴	اندازه‌های تجربی آزمون‌های M مختلف در مدل رگرسیون خطی با خطاهای NSD – مقادیری که با * مشخص شده‌اند، بیانگر اندازه‌های معنی‌داری اسمی هستند.	۴۵

پیشگفتار

در متون آماری، استوار بودن یک روش، به معنی معتبر و مناسب بودن آن است. به عبارت دیگر، هرگاه ویژگی‌های یک روش آماری، تحت تاثیر انحراف کم از پذیره‌های اساسی قرار نگیرد، آن را روش آماری استوار گوییم. متغیرهای تصادفی زبرجمعی منفی یک رده بسیار وسیع از متغیرهای وابسته هستند که به عنوان مواردی خاص از این رده، می‌توان به دنباله متغیرهای مستقل و متغیرهای پیوندی منفی اشاره نمود. مفهوم متغیرهای تصادفی وابسته زبرجمعی منفی مبتنی بر رده توابع زبرجمعی است. در این پایان‌نامه برخی از روش‌های آماری استوار و آزمون استوار در مدل خطی با جملات خطای وابسته منفی و برخی ویژگی‌های اساسی متغیرهای تصادفی وابسته زبرجمعی منفی را رمورد مطالعه قرار می‌دهیم. همچنین در انتها با مطالعه شبیه‌سازی به ارزیابی نتایج نظری می‌پردازیم. محتوای این پایان‌نامه به صورت زیر تدوین شده است:

- در فصل اول به بیان برخی تعاریف و مفاهیم اولیه که در اثبات نتایج اصلی به آن‌ها نیاز داریم، می‌پردازیم.
- در فصل دوم مفهوم استواری و برخی روش‌های آماری استوار مانند برآورد پارامترهای مدل رگرسیونی استوار معرفی می‌کنیم.
- در فصل سوم مدل خطی آزمون استوار با خطاهای وابسته زبرجمعی منفی را بررسی می‌کنیم.
- در فصل چهارم با استفاده از ابزارهای شبیه‌سازی نتایج نظری استواری فصل سوم را ارزیابی می‌کنیم.
- پیوست پایان‌نامه شامل کدهای نرم‌افزار R برای تولید نتایج شبیه‌سازی می‌باشد.

فصل ۱

مفاهیم اولیه

در این فصل برخی تعاریف و قضیه‌هایی که در فصل‌های بعدی به آن احتیاج داریم بیان می‌کنیم. برای گردآوری عمده مطالب این فصل از گات (۱۹۹۴)، پایان نامه‌های کارشناسی ارشد حیدری نژاد (۱۳۹۷)، حبیبی کم‌نی (۱۳۹۵) و رحمانی (۱۳۹۵) استفاده کرده‌ایم.

۱.۱ نابرابری‌ها

نابرابری‌ها نقش کلیدی در نظریه احتمال دارند. زیرا اگر نتوانیم مقدار دقیق احتمال‌ها و گشتاورها را محاسبه کنیم، با استفاده از نابرابری‌ها می‌توان کران‌هایی برای مقادیر مورد نظر پیدا کرد. یک نابرابری مناسب کلیدی جهت رفع مشکلات محسوب می‌شود. به عبارت دیگر هرچه کران به مقدار واقعی نزدیک‌تر باشد، جواب نهایی قابل اعتمادتر است. در زیر برخی از نابرابری‌های مهم آماری را از گات (۱۹۹۴) معرفی می‌کنیم.

۱- نابرابری مارکف^۱: فرض کنید X متغیر تصادفی که برای عدد حقیقی مثبت r که $E|X|^r < \infty$ آن‌گاه اگر $x > 0$

$$P(|X| > x) \leq \frac{E|X|^r}{x^r}.$$

¹Markov's Inequality

۲- نابرابری هافدینگ^۱: فرض کنید $X_1, X_2, \dots, X_n$ متغیرهای تصادفی مستقل باشند طوری که برای $k = 1, 2, \dots, n$ و $1 \leq n$ ، $P(a_k \leq X_k \leq b_k) = 1$ هم چنین فرض کنید برای $1 \leq n$ ، $S_n = \sum_{k=1}^n X_k$ مجموع‌های جزیی باشد. آن‌گاه برای $x > 0$ و a_n, b_n که اعداد حقیقی ثابت هستند و $a_n \leq b_n$

$$P(S_n - ES_n > x) \leq \exp\left\{\frac{-x^2}{\sum_{k=1}^n (b_k - a_k)^2}\right\}$$

$$P(|S_n - ES_n| > x) \leq 2 \exp\left\{\frac{-x^2}{\sum_{k=1}^n (b_k - a_k)^2}\right\}$$

۳- نابرابری C_r ^۲: فرض کنید $\{X_n; n \geq 1\}$ دنباله‌ای از متغیرهای تصادفی مستقل باشند. اگر عدد حقیقی مثبت $r > 0$ وجود داشته باشد که به ازای $i = 1, 2, \dots, n$ ، $E|X_i|^r < \infty$ باشد، آن‌گاه

$$E\left|\sum_{i=1}^n X_i\right|^r \leq C_r \sum_{i=1}^n E|X_i|^r$$

که C_r مقادیر زیر را اختیار می‌کند

$$C_r = \begin{cases} 1 & \text{اگر } r \leq 1 \\ 2^{r-1} & \text{اگر } r \geq 1 \end{cases}$$

۴- نابرابری هولدر^۳: فرض کنید برای هر عدد حقیقی $q > 1$ و $p > 1$ که

$$p^{-1} + q^{-1} = 1$$

اگر $E|X|^p < \infty$ و $E|Y|^q < \infty$ آن‌گاه

$$E|XY| \leq (E|X|^p)^{\frac{1}{p}} (E|Y|^q)^{\frac{1}{q}}.$$

۵- نابرابری کوشی-شوارتز^۴: در نابرابری هولدر اگر $p = q = 2$ قرار دهیم آن‌گاه نابرابری کوشی شوارتز به صورت زیر به دست می‌آید:

$$E|XY| \leq (E|X|^2)^{\frac{1}{2}} (E|Y|^2)^{\frac{1}{2}}.$$

۶- نابرابری مینکوفسکی^۵: فرض کنید X و Y متغیرهای تصادفی باشند و برای هر عدد حقیقی $p \geq 1$ ، $E|X|^p < \infty$ و $E|Y|^p < \infty$ باشد. آن‌گاه

$$(E|X + Y|^p)^{\frac{1}{p}} \leq (E|X|^p)^{\frac{1}{p}} + (E|Y|^p)^{\frac{1}{p}}.$$

¹Hoeffding's Inequality

² C_r Inequality

³Holder's Inequality

⁴Cauchy-Schwarz Inequality

⁵Mincowski Inequality

۲.۱ انواع وابستگی

در این بخش به معرفی چند نوع از وابستگی‌ها که رده‌ی بزرگی از دنباله متغیرهای تصادفی را در بر می‌گیرند می‌پردازیم.

۱.۲.۱ وابسته ربعی منفی

تعریف ۱.۲.۱. (لهمن ۱۹۶۶) متغیرهای تصادفی X و Y دارای وابستگی ربعی منفی^۱ (NQD) هستند اگر برای هر x و y حقیقی داشته باشیم:

$$P(X \leq x, Y \leq y) \leq P(X \leq x) P(Y \leq y)$$

یا به طور معادل

$$F(x, y) \leq F_1(x) F_2(y).$$

که F تابع توزیع توام و F_1, F_2 توابع توزیع حاشیه‌ای هستند، اگر علامت نابرابری روابط بالا را برعکس کنیم به وابستگی ربعی مثبت^۲ (PQD) تبدیل می‌شود. شرایط تعریف شده برای NQD نتیجه می‌دهد که $Cov(X, Y) \leq 0$ است.

۲.۲.۱ وابستگی پیوندی منفی

نتایج مختلفی در آمار و احتمال تحت فرض برخی از متغیرهای تصادفی دارای وابستگی پیوندی منفی^۳ (NA) هستند توسط آلام و سکنا (۱۹۸۱) معرفی شده‌اند و جاج و پروشان (۱۹۸۳) آنها را به دقت مورد مطالعه قرار داده‌اند.

تعریف ۲.۲.۱. یک دنباله متناهی از متغیرهای تصادفی $\{X_i, 1 \leq i \leq n\}$ ، وابسته پیوندی منفی (NA) نامیده می‌شود اگر برای هر جفت زیرمجموعه جدا از هم A_1 و A_2 از $\{1, 2, \dots, n\}$ بتوان نوشت

$$Cov(f_1(X_i, i \in A_1), f_2(X_j, j \in A_2)) \leq 0$$

هرگاه دو تابع f_1 و f_2 نسبت به هر مولفه صعودی باشند و کواریانس تعریف شده در رابطه بالا وجود داشته باشد.

نکته ۱.۲.۱. (جاج و پروشان ۱۹۸۳). اگر توابع f_1 و f_2 هر دو نزولی باشند نابرابری رابطه $Cov(f_1(X_i, i \in A_1), f_2(X_j, j \in A_2)) \leq 0$ نیز برقرار است. یکی از مزیت‌های وابستگی NA

¹Negative Quadrant Dependence

²Positive Quadrant Dependence

³ Negatively Associated

این است که توابع f_1 و f_2 در رابطه بالا می‌توانند صورتهای مختلفی داشته باشند و این نتیجه به بسیاری از نابرابری‌های چند متغیره مفید منجر می‌شود.
بدون از دست دادن کلیت می‌توان فرض کرد $A_1 \cup A_2 = \{1, 2, \dots, n\}$. هم‌چنین یک دنباله نامتناهی پیوندی منفی است اگر هر زیر دنباله متناهی آن پیوندی منفی باشد.

۳.۲.۱ چند ویژگی متغیرهای تصادفی پیوندی منفی

چند ویژگی برای متغیرهای تصادفی NA توسط جاج و پروشان (۱۹۸۳) بیان شده‌اند که در زیر به آنها اشاره می‌کنیم:

۱- هر جفت از متغیرهای تصادفی NA ، NQD می‌باشد.

۲- فرض کنید $A_1, A_2, \dots, A_m$ زیرمجموعه‌های مجزا از $\{1, 2, \dots, n\}$ و هم‌چنین توابع $f_1, f_2, \dots, f_m$ نسبت به هر مولفه صعودی باشند، در این حالت NA بودن $X_1, X_2, \dots, X_n$ دلالت بر این دارد که

$$E \prod_{i=1}^m f_i(X_j, j \in A_i) \leq \prod_{i=1}^m E f_i(X_j, j \in A_i).$$

۳- یک نتیجه بدیهی از ویژگی ۲ این است که برای زیرمجموعه‌های جدا از هم A_1 و A_2 از $\{1, 2, \dots, n\}$ و هم‌چنین مقادیر حقیقی $x_1, x_2, \dots, x_n$ می‌توان نوشت

$$P(X_1 > x_1, X_2 > x_2, \dots, X_n > x_n) \leq \prod_{i=1}^n P(X_i > x_i)$$

$$P(X_1 \leq x_1, X_2 \leq x_2, \dots, X_n \leq x_n) \leq \prod_{i=1}^n P(X_i \leq x_i)$$

۴- یک زیرمجموعه دوتایی یا بیشتر از متغیرهای تصادفی NA ، NA است.

۵- یک مجموعه از متغیرهای تصادفی مستقل، NA است.

۶- توابع صعودی تعریف شده روی زیرمجموعه‌های جدا از هم یک مجموعه از متغیرهای تصادفی NA ، NA هستند.

۴.۲.۱ مثال‌هایی از متغیرهای تصادفی پیوندی منفی

جاج و پروشان (۱۹۸۳) نشان دادند برای دو مثال زیر ویژگی NA می‌تواند برقرار باشد:
(۱) فرض کنید $Z = (Z_1, Z_2, \dots, Z_k)$ بردار تصادفی دارای توزیع چند جمله‌ای است که تنها دارای یک مشاهده است. بنابراین فقط یک Z_i برابر با ۱ می‌باشد و بقیه اعضا صفر هستند. به عبارتی $\sum_{i=1}^k Z_i = 1$ است. پس با افزایش یک Z_i بقیه یا ثابت می‌مانند یا کاهش پیدا می‌کنند.

پس کوواریانس بین آنها منفی است. بنابراین از تعریف پیوندی منفی نتیجه می‌شود که ویژگی NA برای Z بدیهی است. با توجه به ویژگی (۶) به NA بودن Z پی می‌بریم.

(۲) جعبه‌ای شامل N توپ با رنگ‌های متمایز را در نظر بگیرید. نمونه‌ای شامل n توپ (بدون جای‌گذاری) از این مجموعه انتخاب کرده و متغیرهای تصادفی $X_i, i = 1, 2, \dots, N$ را تابع نشان‌گر وجود یک توپ از رنگ i ام در نمونه تعریف کنید. به عبارت دیگر $X_i = 1$ است اگر توپی با رنگ i ام در نمونه موجود باشد و در غیر این صورت $X_i = 0$ است. با توجه به این که از هر رنگ فقط یک توپ موجود است پس بردار تصادفی $(X_1, X_2, \dots, X_N)$ برای هر $n < N$ ، زیرمجموعه $(X_1, X_2, \dots, X_N)$ بوده که دارای توزیع جایگشتی است. بنابراین با توجه به ویژگی‌های (۵) و (۶) پس NA می‌باشد. به طور کلی فرض کنید N_i توپ از رنگ i ام در جعبه موجود باشد. پس $\sum_{i=1}^k Z_i = N$ است. اگر Y_i تعداد توپ‌هایی از رنگ i ام در نمونه n تایی باشد آن‌گاه Y_i را می‌توان به عنوان مجموع N_i متغیر تصادفی نشان‌گری که در قسمت قبل تعریف کردیم، در نظر گرفت. با توجه به ویژگی (۶) Y_i ها هم NA هستند.

۵.۲.۱ وابستگی زیرجمعی منفی

وابسته زیرجمعی منفی^۱ (NSD) نوعی وابستگی است که توسط هو (۲۰۰۰) مطرح شد که در پایه آن از توابع زیرجمعی^۲ استفاده شده است.

تعریف ۳.۲.۱. (کمپرن، ۱۹۷۷) تابع $\phi : R^n \rightarrow R$ زیرجمعی نامیده می‌شود اگر برای هر $\mathbf{x}, \mathbf{y} \in R^n$

$$\phi(\mathbf{x} \vee \mathbf{y}) + \phi(\mathbf{x} \wedge \mathbf{y}) \geq \phi(\mathbf{x}) + \phi(\mathbf{y}).$$

در رابطه بالا $\vee$ نشان دهنده‌ی ماکسیمم مولفه‌ای^۳ و $\wedge$ نشان دهنده‌ی مینیمم مولفه‌ای^۴ است، که به صورت زیر تعریف می‌شود.

$$\mathbf{x} \vee \mathbf{y} = (x_1 \vee y_1, x_2 \vee y_2, \dots, x_n \vee y_n)$$

و

$$\mathbf{x} \wedge \mathbf{y} = (x_1 \wedge y_1, x_2 \wedge y_2, \dots, x_n \wedge y_n)$$

کمپرن (۱۹۷۷) روشی را بیان کرد که می‌توان با استفاده از آن زیرجمعی بودن یا نبودن یک تابع را مشخص کرد.

^۱Negatively Superadditive Dependent

^۲ Superadditive Function

^۳Maximum Componentwise

^۴Minimum Componentwise

لم ۱.۲.۱. فرض کنید تابع ϕ دارای مشتق جزئی مرتبه دوم پیوسته باشد، آن گاه ϕ زبرجمعی است هرگاه برای $1 \leq i \neq j \leq n$ داریم:

$$\frac{\partial^2 \phi(\mathbf{x})}{\partial x_i \partial x_j} \geq 0 \quad (1.1)$$

تعریف ۴.۲.۱. بردار تصادفی $\mathbf{X} = (X_1, X_2, \dots, X_n)$ را NSD گویند، هر گاه برای هر تابع زبرجمعی ϕ

$$E\phi(X_1, X_2, \dots, X_n) \leq E\phi(Y_1, Y_2, \dots, Y_n)$$

که Y_i ها متغیرهای تصادفی مستقل هستند، هم چنین به ازای هر $i = 1, 2, \dots, n$ ، X_i و Y_i هم توزیع هستند.

تعریف ۵.۲.۱. $\{X_n, n \geq 1\}$ را دنباله‌ای از متغیرهای تصادفی NSD گویند هرگاه برای هر $n \geq 1$ ، $(X_1, X_2, \dots, X_n)$ ، NSD باشد.

لم ۲.۲.۱. (هو، ۲۰۰۰) فرض کنید بردار $(X_1, X_2, \dots, X_n)$ ، NSD باشد. آن گاه $1 - (-X_1, -X_2, \dots, -X_n)$ نیز NSD هستند.

۲- برای توابع صعودی دلخواه $f_1, f_2, \dots, f_n$ ، $(f_1(X_1), f_2(X_2), \dots, f_n(X_n))$ ، نیز NSD هستند.

۳- برای هر $x_1, x_2, \dots, x_n \in R$ ، می‌توان نتیجه گرفت

$$P(X_1 \leq x_1, X_2 \leq x_2, \dots, X_n \leq x_n) \leq P(X_1 \leq x_1)P(X_2 \leq x_2) \dots P(X_n \leq x_n)$$

برهان. با توجه به فرض لم، داریم

۱- از آنجا که بردار $(X_1, X_2, \dots, X_n)$ ، NSD است و با توجه به لم ۱.۲.۱ کافی است نشان دهیم که برای هر تابع زبرجمعی ϕ و به ازای هر $i \neq j$ ، تابع $\phi(-X_1, -X_2, \dots, -X_n)$ نیز زبرجمعی است. در نتیجه NSD بودن آن روشن است. زیرا

$$\begin{aligned} \frac{\partial^2 \phi(-x_1, -x_2, \dots, -x_n)}{\partial x_i \partial x_j} &= \frac{\partial^2 \phi(-x_1, -x_2, \dots, -x_n)}{\partial (-x_i) \partial (-x_j)} \times \frac{\partial(-x_i) \partial(-x_j)}{\partial x_i \partial x_j} \\ &= \frac{\partial^2 \phi(x_1, x_2, \dots, x_n)}{\partial x_i \partial x_j} \geq 0 \end{aligned}$$

۲- چون f تابعی صعودی است، با توجه به لم ۱.۲.۱ کافی است نشان دهیم برای هر تابع زبرجمعی ϕ و هر $j \neq i$ ، تابع $\phi(f_1(x_1), f_2(x_2), \dots, f_n(x_n))$ نیز زبرجمعی است. در نتیجه برهان کامل می‌شود، زیرا

$$\begin{aligned} & \frac{\partial^2 \phi(f_1(x_1), f_2(x_2), \dots, f_n(x_n))}{\partial x_i \partial x_j} \\ &= \frac{\partial^2 \phi(f_1(x_1), f_2(x_2), \dots, f_n(x_n))}{\partial f_i(x_i) \partial f_j(x_j)} \cdot \frac{\partial f_i(x_i)}{\partial x_i} \cdot \frac{\partial f_j(x_j)}{\partial x_j} \geq 0. \end{aligned}$$

۳- با توجه به تعریف وابستگی زبرجمعی منفی، همچنین بر اساس قسمت دوم این لم با قرار دادن تابع صعودی نشانگر به جای تابع صعودی f ، می توان نشان داد

$$\begin{aligned} P(X_1 \leq x_1, X_2 \leq x_2, \dots, X_n \leq x_n) &= E(I\{X_1 \leq x_1, X_2 \leq x_2, \dots, X_n \leq x_n\}) \\ &= E\left(\prod_{i=1}^n I\{X_i \leq x_i\}\right) \\ &\leq E\left(\prod_{i=1}^n I\{Y_i \leq x_i\}\right) \\ &= \prod_{i=1}^n E(I\{Y_i \leq x_i\}) \\ &= \prod_{i=1}^n P(Y_i \leq x_i) \\ &= \prod_{i=1}^n P(X_i \leq x_i). \end{aligned}$$

□

ملاحظه ۱.۲.۱. قسمت ۲ لم ۲.۲.۱ اگر توابع $f_1, f_2, \dots, f_n$ ناصعودی باشند نیز برقرار است.

ملاحظه ۲.۲.۱. (کریستوفیدز و واگلاتو، ۲۰۰۴) اگر $X_1, X_2, \dots, X_n$ پیوندی منفی باشند، آن گاه وابسته زبرجمعی منفی هستند ولی عکس آن همیشه برقرار نیست. در زیر برخی ویژگی های متغیرهای تصادفی وابسته زبرجمعی منفی را از هو (۲۰۰۰)، بدون برهان بیان می کنیم.

ملاحظه ۳.۲.۱. وابستگی پیوندی منفی یک زوج متغیر تصادفی با وابستگی زبرجمعی منفی آن معادل است.

ملاحظه ۴.۲.۱. برای یک زوج متغیر تصادفی NQD با NA معادل است، لذا روابطی که میان NSD و NA برقرار هستند، میان NSD و NQD نیز برقرار خواهند بود.

۶.۲.۱ مثال‌ها و کاربردهایی از وابسته زبرجمعی منفی

خانواده‌ای از توزیع‌های چندمتغیره با ویژگی NSD

● خانواده توزیع‌های بیضی‌گون: برای این منظور، در ابتدا به تعریفی از خانواده توزیع‌های بیضی‌گون می‌پردازیم.

تعریف ۶.۲.۱. (بولک و سامپسون، ۱۹۸۸) خانواده $F(x, \Sigma)$ خانواده‌ای از توزیع‌های بیضی‌گون^۱ پایه‌گذاری شده بر اساس تابع g است، هرگاه تابع چگالی آن به صورت زیر تعریف شود:

$$f(x, \Sigma) = \left| \Sigma \right|^{-\frac{1}{2}} g(x^T \Sigma x) \quad : x \in \mathbb{R}^m$$

که در آن $\int_0^\infty g(t) t^{\frac{m}{2}-1} dt < \infty$.

بولک و سامپسون (۱۹۸۸) نشان دادند که $F(x, \Sigma)$ دارای وابستگی از نوع NSD خواهد بود اگر برای هر $i \neq j$ ، $\delta_{ij} \leq 0$ باشد به طوری که δ_{ij} ها مولفه‌های ماتریس Σ هستند. در این میان، توزیع نرمال چندمتغیره نیز به عنوان حالت خاصی از توزیع‌های بیضی‌گون محسوب می‌شود. بنابراین توزیع نرمال چندمتغیره NSD خواهد بود اگر عناصر خارج از قطر اصلی ماتریس کواریانس آن نامثبت باشند.

● خانواده توابع فراوانی پولیای نوع ۲^۲: به جهت آشنایی با این دست از توزیع‌ها تعریفی از آن را ارائه می‌کنیم.

تعریف ۷.۲.۱. (افرون، ۱۹۶۵) تابع چگالی احتمال روی خط حقیقی $r(x)$ را تابع فراوانی پولیای نوع ۲ (PF_2) نامند، هرگاه به ازای $x_2 \geq x_1$ و $z_2 \geq z_1$ داشته باشیم

$$\begin{bmatrix} r(x_1 - z_1) & r(x_1 - z_2) \\ r(x_2 - z_1) & r(x_2 - z_2) \end{bmatrix} \geq 0$$

این خانواده از توزیع‌ها دارای ویژگی‌هایی است که در زیر به طور مختصر به آن‌ها اشاره می‌نماییم.

آ. تعداد زیادی از توزیع‌های رایج در نظریه احتمال، PF_2 هستند که به عنوان حالت خاص می‌توان به توزیع‌های پواسن، دوجمله‌ای، گاما دوجمله‌ای منفی و همچنین توزیع نرمال با تابع چگالی

$$r(x) = (\sqrt{2\pi}\sigma)^{-1} \exp\left\{-\frac{1}{2}\left[\frac{x-m}{\sigma}\right]^2\right\}$$

با $\sigma > 0$ ، اشاره کرد.

^۱Elliptically Contoured Distribution

^۲Polya Frequency Function of order 2

ب. توابع چگالی PF_F به طور لگاریتمی مقعرند^۱. بنابراین تقریباً همه جا دارای مشتق اول و دوم هستند.

پ. توابع چگالی PF_F کران دارند و همچنین بر روی یک فاصله، غیرصفر و در خارج از آن، صفر هستند.

ت. فرض کنید $r(x)$ و $q(x)$ ، PF_F باشند آن گاه پیچش آن‌ها به شکل

$$P(x) = \int_{-\infty}^{\infty} r(t)q(x-t)dt$$

نیز PF_F است.

بنابراین توزیع‌هایی مانند پواسن، دوجمله‌ای، گاما و دوجمله‌ای منفی دارای ویژگی PF_F هستند.

بر اساس آن چه هو (۲۰۰۰) ارائه کرده است، اگر $X_1, X_2, \dots, X_n$ متغیرهای تصادفی مستقل با تابع چگالی احتمال PF_F باشند، $S_n = \sum_{i=1}^n X_i$ به طور تقریباً حتمی NSD است.

۳.۱ ترتیب‌های غیرتصادفی

ترتیب‌های غیرتصادفی را با o و O را نشان می‌دهند که نمادهایی برای توصیف مرتبه‌ی مجانبی کمیت‌های غیرتصادفی می‌باشند. ون و وارت (۱۹۹۸) تعاریفی از ترتیب‌های غیرتصادفی با توجه به ترتیب همگرایی در احتمال به همراه نمادهای مناسب ارائه داد که به شرح زیر می‌باشد.

تعریف ۱.۳.۱. دو دنباله از اعداد ثابت $\{a_n; n \geq 1\}$ و $\{b_n; n \geq 1\}$ را در نظر بگیرید. گوییم $b_n = O(a_n)$ ، هرگاه برای هر $\varepsilon > 0$ ، $K(\varepsilon) > 0$ و عدد صحیح $N(\varepsilon)$ وجود داشته باشند به طوری که برای هر $n \geq N$ ، آن گاه

$$|b_n| < K(\varepsilon)|a_n|$$

تعریف ۲.۳.۱. دو دنباله از اعداد ثابت $\{a_n; n \geq 1\}$ و $\{b_n; n \geq 1\}$ را در نظر بگیرید. گوییم $b_n = o(a_n)$ ، هرگاه

$$\lim_{n \rightarrow \infty} \left| \frac{b_n}{a_n} \right| = 0$$

با استفاده از تعاریفی که در بالا به آن اشاره شد، برای هر ثابت حقیقی مثل C ، ترتیب‌های $O(a_n)$ و $o(a_n)$ ، که به ترتیب معادل با $Ca_n O(1)$ ، $Ca_n o(1)$ هم‌چنین اگر $X_n = O(n^C)$ ، آن گاه $X_n = O(n^{C+1})$ قرار می‌گیرد اما عکس رابطه برقرار نیست.

^۱Log-Concave

فصل ۲

روش های آماری استوار

در این فصل به شرح روش های آماری استوار از جمله برآورد پارامترهای مکانی استوار، رگرسیون استوار و تاریخچه آن و مقایسه آن با کمترین توان های دوم در مورد نقاط پرت می پردازیم. مطالب این فصل از هوگ (۱۷۹۹)، پایان نامه کارشناسی ارشد اقبال (۱۳۸۴) و حیدری نژاد (۱۳۹۷) می پردازیم.

۱.۲ استواری

در متون آماری استوار بودن یک روش به معنی معتبر بودن و مناسب بودن آن است. به عبارت دیگر هرگاه ویژگی یک روش آماری تحت تأثیر انحراف کم از پذیره های اساسی قرار نگیرند، آن روش آماری را استوار گوییم. برای مثال $\bar{x}$ ، میانگین نمونه، برآوردگری برای میانگین جامعه است ولی استوار نیست. یعنی اگر حداقل یکی از اعضای نمونه مقدار خیلی بزرگی داشته باشد و از جامعه ی مورد نظر نیامده باشد، داده پرت باشد، بسیار متفاوت از میانگین واقعی جامعه است. در مقابل میانه ی نمونه برآوردگری استوار برای پارامتر مکان، در صورت تقارن توزیع جامعه، برای میانگین جامعه است. زیرا حتی اگر نقاط پرت در نمونه باشد، مقدار میانه تغییر نمی کند.

این سوال مطرح می شود که نرمال نبودن مشاهدات بر دقت و کارایی بازه اطمینان چقدر تأثیر می گذارد؟ یعنی انحراف از پذیره ی نرمال بودن جامعه به میزان کم به چه اندازه بر ضریب

اطمینان واقعی بازه و به چه میزان بر متوسط طول بازه تأثیر می گذارد؟

جوابی که می توان به این سوال داد این است که تأثیر نرمال نبودن، در حالی که جامعه متقارن باقی بماند بر ضریب اطمینان واقعی بسیار اندک است ولی اگر جامعه متقارن نباشد، ضریب اطمینان به شدت کاهش می یابد. بنابراین می توان گفت برآورد فاصله ای t - استیودنت مادامی که جامعه متقارن باشد، استوار و اگر توزیع از تقارن فاصله بگیرد ناستوار است. بر این اساس روش های استواری برای اصلاح روش کمترین توان های دوم به گونه ای هستند که نقاط دور افتاده تأثیر کمتری بر روی برآوردگر نهایی دارد. یکی از اقسام رگرسیون، رگرسیون استوار می باشد (در مورد تجزیه و تحلیل داده های پرت) که توسط جورج باکس (۱۹۵۳) برای اولین بار معرفی گردید. مقالات و کتاب هایی در این زمینه به چاپ رسیده است که می توان از آنها استفاده کرد. از جمله جان توکی (۱۹۶۲) که در مورد مکان و موقعیت داده ها توضیح می دهد شاید بیشترین سهم را هوپر داشته باشد که مقاله های اساسی او در سال های ۱۹۶۴، ۱۹۷۲ و ۱۹۷۳ نقطه عطف در این زمینه هستند. همچنین فرانک همپل از تأثیر منحنی در برآورد استواری که یک مفهوم مرکزی است صحبت می کند از این رو اگر کسی علاقه به تأثیر منحنی داشته باشد می تواند به مقاله همپل (۱۹۷۴) مراجعه کند. پس از آنکه هوپر (۱۹۷۳) M - برآوردگرها را مطالعه کرد، بسیاری از آمارشناسان علاقه مند به مطالعه این موضوع هستند. مراجعه به کارهای اخیر M - برآوردگرها را می توان در چن و ژائو (۱۹۹۵) یافت. پایداری مجانبی M - برآوردگرهای مدل های خطی با خطاهای وابسته عملاً مهم است اما از نظر علم نظری چالش برانگیز است. هوپر (۱۹۷۳) اظهار داشت که محدودیت درمورد فرض استقلال مهم است و بیان کرد که فرض استقلال را در نظریه M - برآوردگرهای کلاسیک کاهش خواهیم داد تا خطاهای وابسته مناسب تر شود. به طور کلی بدست آوردن خصوصیات استواری (سازگاری) قوی برای خطاهای وابسته کار آسانی نیست.

مدل خطی رگرسیون $Y_i = X_i^T \beta + e_i \quad i = 1, 2, \dots, n$ را در نظر بگیرید. وقتی مشاهدات Y_i ها دارای توزیع نرمال هستند برآورد خوبی از β حاصل می شود. اما اگر مشاهدات دارای توزیع نرمال نباشند نمی تواند برآورد خوبی را به دست آورد. زیرا در انتهای منحنی دارای مشاهداتی هستیم که منحنی دم پهن تر یا به عبارتی ضخیم تر از نرمال می باشد که روش کمترین توان های دوم کارا نیست که این شاخه های ضخیم توزیع را معمولاً نقاط پرت ایجاد می کنند. این نقاط پرت بر برآوردگرهای کمترین توان های دوم تأثیر می گذارند. داده های پرت در برازش کمترین توان های دوم را در حد زیادی به سمت خود می کشد که در نتیجه تعیین و تشخیص این دور افتاده ها را مشکل می شود زیرا باقی مانده ی مربوط به آن ها به طور ساختگی و مصنوعی کوچک می باشند. بنابراین باید دنبال راه حلی باشیم که این مشکل را برطرف کند که کاربران آماری رگرسیون استوار^۱ را پیشنهاد می دهند که برای کاهش اثر مشاهداتی به کار می رود. اگر روش کمترین توان های دوم تأثیرگذاری بالایی خواهند داشت یعنی اینکه روش استواری می خواهد

^۱Robust Regression

باقیمانده‌های مرتبط با دور افتاده‌های بزرگ را کنار بگذارد که به این وسیله تشخیص نقاط پرت خیلی ساده‌تر می‌شود. رگرسیون استوار علاوه بر حساس نبودن نسبت به داده‌های پرت زمانی که مشاهدات دارای توزیع نرمال هستند دارای کارایی ۹۰ تا ۹۵ درصد نسبت به کمترین توان‌های دوم می‌باشد.

امروزه آماردانان بیشتر از روش‌های استواری و تعمیم آن‌ها استفاده می‌کنند، که ما در این بخش به معرفی چند مورد از آن‌ها می‌پردازیم از جمله M – برآوردگر^۱، L – برآوردگر^۲. هدف اصلی ما در این پایان‌نامه بررسی برخی از روش‌های استواری اشاره شده است، روش‌هایی که تاکید می‌کنند، این است که به وسیله برخی از این سازگاری‌ها می‌توان هر زمان کمترین توان‌های دوم یا تعمیم آن‌ها استفاده شود مورد استفاده قرار گیرد.

۲.۲ برآورد استوار پارامتر مکان

در یک توزیع با تابع چگالی احتمال $f(x)$ و پارامتر مکان θ نامعلوم $-\infty < \theta < +\infty$ برای به‌دست آوردن تابع چگالی $f(x - \theta)$ فرض کنید که $x_1, x_2, \dots, x_n$ یک نمونه‌ی تصادفی از این توزیع را نشان می‌دهد یکی از روش‌های مشهور برای برآورد پارامتر مکان استفاده از روش درست‌نمایی ماکسیمم است. اگر تابع درست‌نمایی را با $L(\theta)$ نشان دهیم.

$$\ln L(\theta) = \sum_{i=1}^n \ln f(x_i - \theta) = - \sum_{i=1}^n \rho(x_i - \theta)$$

که $\rho(x) = -\ln f(x)$ با مشتق گرفتن نتیجه به صورت زیر حاصل می‌شود

$$\frac{\partial \ln L(\theta)}{\partial \theta} = - \sum_{i=1}^n \frac{f'(x_i - \theta)}{f(x_i - \theta)} = \sum_{i=1}^n \psi(x_i - \theta)$$

اگر $\rho' = \psi(x)$ آنگاه

$$\sum_{i=1}^n \psi(x_i - \theta) = 0$$

$L(\theta)$ را ماکسیمم می‌کند. به θ ای که در رابطه بالا صدق می‌کند برآوردگر درست‌نمایی ماکسیمم گوییم و آن را با $\hat{\theta}$ نمایش می‌دهیم به مثال‌های زیر توجه کنید:

۱- نرمال؛ $\rho(x) = \frac{x^2}{2} + c$ در این صورت $\psi(x) = x$ و

$$\sum_{i=1}^n (x_i - \theta) = 0 \implies \hat{\theta} = \bar{x}$$

۲- نمایی دوگانه

^۱M- estimator

^۲L- estimator

$$\rho(x) = |x| + c$$

با توجه به اینکه

$$\psi(x) = \begin{cases} -1 & x < 0 \\ 1 & x > 0 \end{cases}$$

است در این صورت $\hat{\theta}$ میانه نمونه می باشد.

۳- کوشی

$$\rho(x) = \ln(1 + x^2) + c$$

و با توجه به اینکه

$$\psi(x) = \frac{2x}{1 + x^2}$$

بنابراین $\hat{\theta}$ با روش تکرار حل می شود.

توجه داشته باشید که تابع ψ در مثال های مربوط نمایی دوگانه و کوشی کراندار است هم چنین در مثال مربوط به توزیع کوشی با افزایش x به طور مجانبی به صفر نزدیک می شود که این جواب ها خیلی تحت تاثیر نقاط پرت قرار نمی گیرند ولی از سوی دیگر برآوردگر کمترین توان دوم $\bar{x}$ در مثال نرمال به شدت تحت تاثیر نقاط پرت قرار دارد برآوردگر کمترین توان دوم $\bar{x}$ به اندازه حالتی که توزیع دمپهن مانند کوشی در نظر گرفته شود خوب نیست. بنابراین در برآورد استوار ما به دنبال برآوردگری هستیم که تقریباً کارا باشد یعنی کارایی (۹۰ تا ۹۵ درصدی) برای حالتی که توزیع نرمال است تحت فرضیه نرمال گاهی اوقات مقدار کمی از کارایی (تقریباً نزدیک به ۵ درصد) کاسته می شود.

تعریف ۱.۲.۲. اگر U_1 و U_2 دو برآوردگر نااریب برای θ باشند برآوردگری کاراتر است که واریانس کمتری داشته باشد. کارایی برآوردگر U_1 نسبت به U_2 به صورت زیر تعریف می کنیم

$$e(U_1, U_2) = \frac{Var(U_2)}{Var(U_1)}$$

اگر $e < 1$ باشد، U_2 کاراتر از U_1 است و اگر $e > 1$ باشد، U_1 کاراتر است.

برای یک دلیل نظری خاص (مینیمم کردن ماکسیمم واریانس مجانبی از برآوردگرهای مربوط به یک کلاس از توزیع ها)، هوبر (۱۹۶۴) پیشنهاد داد که برای برآوردگرهای استوار، برآوردکننده ماکسیمم درست نمایی از پارامتر مکان مربوط به چگالی مانند یک نرمال در مرکز و نمایی دوگانه در دم ها استفاده می شود. به طور خاص تابع ρ تابع هوبر است اگر

$$\rho(x) = \begin{cases} \frac{1}{4}x^2 & |x| \leq k \\ k|x| - \frac{k^2}{4} & k < |x| \end{cases} \quad (1.2)$$

همچنین تابع ψ مربوط به آن به صورت زیر می باشد

$$\psi(x) = \begin{cases} -k & x < -k \\ x & -k \leq x \leq k \\ k & k < x \end{cases} \quad (2.2)$$

که با روش تکراری حل می شود. یک برآوردگر از این نوع (لزوما این تابع ψ نیست)، برآوردگر برای این حالت را با $\hat{\theta}$ نمایش می دهیم که این را M - برآوردگر می نامیم، یکی دیگر از ویژگی های M - برآوردگرها به صورت زیر است.

فرض کنید $\hat{\theta}$ یک برآوردگر برای نمونه خاص $x_1, x_2, \dots, x_n$ باشد، سپس این موارد توسط بعضی از آنها جایگزین شده باشد به عنوان مثال انحراف از $\hat{\theta}$ سه برابر شده باشد. این برآوردگر جدید $\hat{\theta}$ با استفاده از نمونه تعدیل یافته لزوما یکسان نیست. برای به دست آوردن مقیاس ثابت از این برآوردگر ما می توانیم معادله زیر را حل کنیم

$$\sum_{i=1}^n \psi\left(\frac{x_i - \theta}{d}\right) = 0 \quad (3.2)$$

که d یک برآورد استوار برابر پارامتر مقیاس است. به طور خاص d یک آماره خوب استفاده شده در این راه حل است

$$d = \frac{\text{median}|x_i - \text{median}(x_i)|}{(0.6745)} \quad (4.2)$$

صورت کسر d ، میانه انحراف از میانه تقسیم بر ۰/۶۷۴۵ است. اگر n بزرگ باشد تقریباً $d \approx \sigma$ در واقع نمونه از یک توزیع نرمال حاصل می شود. معمولاً انحراف معیار استاندارد نمونه s به عنوان یک مقدار d استفاده نمی شود زیرا خیلی تحت تاثیر نقاط دور افتاده قرار می گیرد و در نتیجه استوار نیست.

این طرح ویژه از انتخاب d مقدار مناسب را پیشنهاد می دهد. اگر توزیع پایه ای نرمال باشد، ثابت هموارساز k باعث بالا بردن کارایی $\hat{\theta}$ می شود. در حالت نرمال، ما بیشترین محدودیت ها را برای برآورد نابرابری خواهیم داشت

$$\left| \frac{x_i - \theta}{d} \right| \leq k \quad (5.2)$$

زیرا

$$\psi\left(\frac{x_i - \theta}{d}\right) = \frac{x_i - \theta}{d} \quad (6.2)$$

در واقع اگر همه ی محدودیت ها از این نابرابری باشند سپس $\hat{\theta} = \bar{x}$ به عنوان برآوردگری از توزیع نرمال است.

اگر $\sigma \approx d$ ، آن گاه k معمولاً نزدیک به $1/5$ است. وقتی $k = 1/5$ است از روش هوبر استفاده می شود. اگر σ معلوم (d معلوم) باشد، تحت فرض نرمال دارای کارایی بیش از 95% درصدی می باشد هوبر (۱۹۶۴). اندروز و همکاران (۱۹۷۲) به این نتیجه رسیدند که در شرایط توزیع دم سنگین کارایی بهتر می شود، اما در بسیاری از موارد که توزیع ها غیرنرمال هستند و σ نامعلوم، پارامتر σ باید برآورد شود که این باعث می شود مقدار بسیار کمی از کارایی کاهش یابد زمانی که $n \geq 10$ است.

۳.۲ رگرسیون استوار

در حالتی که رگرسیون کلاسیک تحت تاثیر نقاط دور افتاده قرار می گیرند از این رو رگرسیون استوار با استفاده از روش های استوار تاثیر این نقاط دور افتاده را کم می کند. مدل رگرسیون زیر را در نظر بگیرید

$$Y = X^T \beta + \varepsilon \quad (1.2)$$

Y برداری تصادفی $n \times 1$ ، X ماتریس طرح است با بعد $n \times p$ طوری که $X'X$ یک ماتریس رتبه کامل است، β یک بردار $n \times 1$ از ضرایب مجهول است و ε بردار تصادفی $n \times 1$ است که عناصر آن نمونه ای از یک توزیع متقارن حول صفر (مثلاً نرمال) هستند. مشابه رهیافت پارامتر مکان ما می خواهیم مجموعه زیر را مینیمم کنیم

$$\sum_{i=1}^n \rho\left(\frac{Y_i - X_i^T \beta}{d}\right) \quad (2.2)$$

که در آن Y_i ، i امین عنصر از Y و X_i ، i مین سطر از ماتریس X است و d برآورد پارامتر مقیاس توزیع مربوط به ε است. معادله مشتقات جزئی مرتبه اول مربوط به عناصر β را که با β_j نمایش می دهیم برابر صفر قرار می دهیم. همانطور که مشاهده می کنیم این کار هم ارز با یافتن ماکسیمم جواب از رابطه زیر می باشد

$$\sum_{i=1}^n x_i \psi\left(\frac{Y_i - X_i^T \beta}{d}\right) = 0 \quad i = 0, 1, \dots, p \quad (3.2)$$

فرض کنید بخواهیم ابتدا برآورد β را به دست بیاوریم برای این کار به دو چیز نیاز داریم

- برآورد استوار d حاصل از پارامتر مقیاس
- نقطه شروع برای بدست آوردن $\hat{\beta}$

اگر برآوردگر L (مینیمم مجموع قدر مطلق خطاها) بتواند ضرایب رگرسیونی را به خوبی پیدا کند در این صورت کار ما را آسان تر می کند به عبارتی شبیه به میانه در یک نمونه واحد

است. برای برآورد استوار d حاصل از پارامتر مقیاس، میانه مقادیر قدر مطلق باقیمانده های غیر صفر را تقسیم بر 0.6745 مانند پارامتر مکان می باشد. اگر چه آماردانان تلاش زیادی در زمینه L - برآوردها کردند اما با این حال ناراضی بودند بنابراین برای برآورد پارامتر β و مقیاس به صورت همزمان سراغ الگوریتم معرفی شده توسط داتر (۱۹۷۷) می رویم که این الگوریتم فقط برای تابع ψ در هوبر (۲۰۰۰) توضیح داده شده است. همان طور که مشاهده می کنیم این روش شبیه کمترین توان های دوم می باشد. البته باید به این نکته هم توجه کنیم که این روش فقط یکی از روش هایی است که در داتر (۱۹۷۷) ذکر شده است. ابتدا با برخی از برآوردهای اولیه β_0 و d_0 شروع می کنیم، و این دو را برآوردهای معمولی برای پارامتر β و مقیاس در نظر می گیریم که به صورت زیر عمل می کنیم

۱. باقیمانده ها را به صورت زیر محاسبه می کنیم

$$z_i = Y_i - X_i^T \beta_0 \quad i = 1, 2, \dots, n$$

۲. یک برآوردها از پارامتر مقیاس به صورت زیر به دست می آوریم

$$d_1^2 = \frac{1}{(n-1)E(\psi^2)} \sum_{i=1}^n \left(\psi\left(\frac{z_i}{d_0}\right) \right)^2 d_0^2$$

که در آن $E(\psi^2)$ امید ریاضی مقادیر $\psi^2(w)$ هوبر است و w یک متغیر تصادفی است که از توزیع نرمال استاندارد پیروی می کند.

۳. باقیمانده های وینزوریده^۱ شده به صورت زیر تعیین می شوند

$$\Delta_i = \psi\left(\frac{z_i}{d_1}\right) d_1$$

توجه کنید که باقیمانده های وینزوریده شده $\psi\left(\frac{z_i}{d_1}\right)$ معادل z_i هستند اگر $\left|\frac{z_i}{d_1}\right| \leq k$ برابر است با $(-kd_1)kd_1 < z_i < kd_1$ اگر $z_i < -kd_1$ برای کسب اطلاعات بیشتر در مورد وینزوریده کردن به دیکسون و توکی (۱۹۶۸) مراجعه شود.

۴. برآورد کمترین توان های دوم^۲ ضرایب رگرسیونی به شرطی که باقیمانده های وینزوریده شده مشاهده شده باشند به صورت زیر می باشد

$$\hat{\xi}_0 = (X'X)^{-1} X' \Delta_0 \quad (4.2)$$

که Δ_0 یک بردار ستونی $1 \times p$ حاصل از $\Delta_1, \Delta_2, \dots, \Delta_n$ است.

۵. یک برآورد جدید از β محاسبه کنید یعنی

$$\beta = \beta_0 + q\hat{\xi}_0$$

¹Winsorize

²Least Squares Estimates

داتر (۱۹۷۷) دریافت یک انتخاب مناسب از فاکتور q به شکل زیر می باشد

$$q = \min\left[\frac{1}{\Phi(k) - \Phi(-k)}, 1/9\right]$$

که Φ تابع توزیع نرمال استاندارد می باشد.

با برآوردهای جدید β_1 و d_1 به عنوان مقادیر اولیه برای مراحل یک تا پنج تکرار می شوند. این روند تکراری را $i = 1, 2, \dots, p$ ادامه می دهیم در این صورت

$$|q\hat{\xi}_m^i| < \varepsilon d_{m+1} \quad |d_{m+1} - d_m| < \varepsilon d_{m+1}$$

که در آن $\varepsilon > 0$ یک سطح تحمل مناسب $\hat{\xi}_m^i$ ، i - امین مولفه از $\hat{\xi}$ و x_{ii} i امین عنصر قطری از $(X'X)^{-1}$ است در این زمان تکرار را متوقف کنید و β و پارامتر مقیاس را با به بکار بردن β_{m+1} و d_{m+1} برآورد کنید. خواننده باید دقیقاً مشخص کند که این روش برای کمترین توان های دوم است یا خیر.

البته مرحله چهارم کمترین توان های دوم برای باقیمانده های وینزوریده شده است و مرحله پنج کمترین توان دوم را با β شامل $q = 1$ و Δ هایی برابر حاصل از باقیمانده های واقعی (غیر وینزوریده شده) برای هر برآورد اولیه β_0 ارایه می دهد. هم چنین جالب است که برای محاسبه ی $(X'X)$ نیاز به محاسبه ی یک بار روند تکراری است که این خود باعث صرفه جویی در محاسبات می شود سپس یک برآورد استوار خوب از β_0 پارامتر مقیاس با استفاده از الگوریتم معرفی شده توسط داتر با روش های کمترین مقدار مطلق یا دیگر روش ها (برخی روش های ناپارامتری که ممکن است بسیار مناسب تر باشد) به دست آوریم. ما می توانیم مشکل نقاط پرت را با به کار بردن برآوردهای استوار d_0 و β_0 یک تابع صعودی ψ بر طرف کنیم مانند موجی یا وزنی در این موارد کمترین توان های دوم موزون می تواند جایگزین مناسبی برای استفاده باشد.

۱.۳.۲ M - برآوردگر

معادله (۲.۲) را نسبت به پارامترهای β_j ، $j = 1, 2, \dots, k$ مینیمم کنیم با مشتق گیری از معادله (۲.۲) نسبت به تک تک پارامترها تعداد $p = k + 1$ معادله به صورت زیر حاصل می شود

$$\sum_{i=1}^n x_{ij} \psi \left(\frac{Y_i - X_i^T \beta}{d} \right) = 0 \quad (5.2)$$

که $X'_i = (X_{i1}, X_{i2}, \dots, X_{ik})$ می باشد. این معادله ها در کل جواب واضحی ندارند و جواب عددی مکرر نیاز است. همانند بیتون و توکی (۱۹۷۴) وزن ها را به صورت زیر تعریف می کنیم

$$W_{i\beta} = \frac{\psi \left\{ \frac{(Y_i - X_i^T \beta)}{d} \right\}}{\frac{(Y_i - X_i^T \beta)}{d}} \quad (6.2)$$

در صورتی که دقیقا $Y_i - X_i^T \beta$ اتفاق بیافتد معادله (۵.۲) را می‌توان به صورت زیر بازنویسی کرد

$$\sum_{i=1}^n x_{ij} W_{i\beta} (Y_i - X_i^T \beta) = 0 \quad j = 0, 1, 2, \dots, n \quad (7.2)$$

آن‌گاه

$$\sum_{i=1}^n x_{ij} W_{i\beta} X_i'^T \beta = \sum_{i=1}^n x_{ij} W_{j\beta} Y_i \quad j = 0, 1, 2, \dots, n \quad (8.2)$$

که می‌توان به صورت ماتریسی زیر بازنویسی کرد

$$X' W_{\beta} X \beta = X' W_{\beta} Y \quad (9.2)$$

که در روابط بالا $W_{\beta} = (W_{1\beta}, W_{2\beta}, \dots, W_{n\beta})$ این معادلات به وضوح به صورت معادلات نرمال کمترین توان‌های دوم تعمیم یافته هستند. (به طور دقیق‌تر، معادلات کمترین توان‌های دوم وزنی هستند زیرا W_{β} یک ماتریس قطری از وزن‌ها می‌باشد). مشکل حل کردن آن این است که W_{β} به β بستگی دارد. با استفاده از یک روش ساده قدیمی این مشکل را حل می‌نماییم به این صورت که یک برآورد اولیه $\hat{\beta}_0$ از β را در نظر می‌گیریم (با استفاده از کمترین توان‌های دوم) مقدار W_0 را محاسبه می‌کنیم، یعنی W_{β} برای همان $\hat{\beta}_0$ و معادله (۹.۲) را برای رسیدن به جواب $\hat{\beta}_1$ حل می‌کنیم. با استفاده از $\hat{\beta}_1$ به W_1 می‌رسیم. سپس با استفاده از W_1 به $\hat{\beta}_1$ می‌رسیم و به همین ترتیب ادامه می‌دهیم. معمولا همگرایی به سرعت رخ می‌دهد. (این روند می‌تواند بعد از یک تعداد منتخبی از گام‌ها متوقف شود). تنها یک برنامه برای کمترین توان‌های دوم احتیاج است، مثلا در MINITAB می‌تواند مکرر به صورت زیر نوشته شود

$$\hat{\beta}_{q+1} = (X' W_q X)^{-1} X' W_q Y \quad (10.2)$$

و این روند ممکن است زمانی که همه‌ی برآوردها کمتر از یک مقدار از پیش تعیین شده مثلا ۱. درصد یا ۱۰. درصد، تغییر کنند یا بعد از یک تعداد گام تعیین شده، متوقف می‌شود. از معادله (۶.۲) مشاهده می‌شود که وزن‌ها به مقدار $\frac{(\hat{\rho}(x))}{x}$ وابسته‌اند که در آن $x = \frac{(Y_i - X_i^T \beta)}{d}$ می‌باشد. بنابراین می‌توانیم برای همه‌ی انتخاب‌های $\rho(x)$ که در جدول زیر آمده به تابع $W(x) = \frac{(\hat{\rho}(x))}{x} x$ نگاه کنیم. به وضوح انتخاب یک $\rho(x)$ مشخص مبنای انتخاب یک تابع وزنی است که کاربر برای اختصاص دادن به باقیمانده‌های مقیاس شده در انواع مقادیر مختلف نیاز دارد یادآوری می‌کنیم که وزن‌ها در (۶.۲) بر حسب $x = \frac{(Y_i - X_i^T \beta)}{d}$ تعریف شده‌اند و نیاز به برآورد مقیاس d دارد. یک انتخاب استوار که تا حدی محبوبیت پیدا کرده به صورت زیر است

$$d = \frac{\text{median}|e_i - \text{median}(e_i)|}{0.6745} \quad (11.2)$$

اگر n به قدر کافی بزرگ و خطاها به صورت نرمال توزیع شده باشند برآوردگری تقریباً نااریب که از $Y_i = \sigma$ به دست خواهد آورد. چندین عامل مقیاسی جایگزین نیز پیشنهاد و بررسی شده‌اند، برای مثال هیل و هولاند (۱۹۷۷) بررسی کردند انتخاب واقعی d چه اثری روی باقیمانده‌ها خواهد گذاشت، به این معنا که برآوردهای پارامتر رگرسیونی استوار بر اساس انتخاب d ثابت نیستند.

هلند و والش (۱۹۷۷) از مقایسه برآوردهای استوار حاصل از ۸ تابع ψ به دست آمده که از جمله آن‌ها بسته‌های آماری استوار که در مرکز کامپیوتری دفتر ملی تحقیقات اقتصادی توسعه یافته است. یکی از ویژگی‌های مطلوب تابع ψ این است که تابع موج اندروز یا تابع وزن دوتایی توکی برای شناسایی نقاط دورافتاده زمانی که ما از روش هویر استفاده می‌کنیم بهتر می‌شوند و در بسیاری از موارد به نقاط پرت یا نقطه داده‌های بد وزن صفر داده می‌شود. البته هر زمان که روش‌های کمترین توان‌های دوم (تعمیم آن‌ها) به کار می‌روند به ظاهر امکان استفاده از روش‌های استواری نیز وجود دارد. که شامل $ANOVA$ ، رگرسیون، سری زمانی، اسپلین، تجزیه و تحلیل چندمتغیره و تشخیص، برای اطلاعات بیشتر درمورد توابع ρ و توابع وزن به جدول صفحه‌ی بعد رجوع کنید.

جدول ۱.۲: بعضی از توابعی که برای توزیع وزن مانده‌ها پیشنهاد شده‌اند

معیار	$\rho(x)$	دامنه $\rho(x)$	$\psi(x) = \frac{\partial \rho(x)}{\partial x}$	$w(x) = \frac{\psi'(x)}{x}$	مقدار ثابت‌های تنظیم ثابت‌های تنظیم کننده و
۱ حداقل مربعات (توزیع مانده‌های نرمال)	$\frac{1}{2}x^2$	$-\infty \leq x \leq \infty$	x	۱	
۲ Huber	$\begin{cases} \frac{1}{2}x^2 \\ a u - \frac{1}{2}a^2 \end{cases}$	$-a \leq x \leq +a$ $u < -a, x > a$	x $a * \text{sign}(x)$	۱ $\frac{a}{ x }$	$a = 2$
۳ Ramsay	$\frac{1}{a} \sqrt{1 - e^{a x } (1 + a x)}$	$-\infty \leq x \leq \infty$	$ue^{a x }$	$e^{a x }$	$a = 0.3$
۴ Andrew	$\frac{a}{\sqrt{a}} \left\{ 1 - \cos\left(\frac{x}{a}\right) \right\}$	$-a \leq x \leq +a\pi$ $x < -a\pi, x > a\pi$	$\sin\left(\frac{x}{a}\right)$	$\sin\left(\frac{x}{a}\right) \frac{x}{a}$	$a = 1.339$
۵ Tukey	$\begin{cases} \frac{1}{2}x^2 - \frac{x^3}{6a^2} \\ \frac{1}{2}a^2 \end{cases}$	$-a \leq x \leq +a$ $x < -a, x > a$	$x(1 - \frac{x^2}{a^2})$	$1 - \frac{x^2}{a^2}$	$a = 5 \leq a \leq 6$
۶ Hampel $a, b, c > 0$	$\begin{cases} \frac{1}{2}x^2 \\ a x - \frac{1}{2}a^2 \\ a \frac{(c x - \frac{1}{2}x^2)}{c-b} - \frac{Va^2}{2} \\ a(b+c-a) \end{cases}$	$-a \leq x \leq +a$ $-b \leq x \leq -a, a \leq x \leq b$ $-c \leq x \leq -b, b \leq x \leq c$ $x < -c, ux > c$	x $a \times \text{sign}(x)$ $c \text{sign}(x) - (x)a/c$	۱ $\frac{a}{ x }$ $(c/ x - 1)(a/c - b)$	$a = 1.7, b = 3.4, c = 1.5$

۲.۳.۲ L-برآوردگر

یک نمونه تصادفی از یک توزیع پیوسته با پارامتر مکان θ در نظر بگیرید آماره‌های ترتیبی نمونه، $X_{(1)}, X_{(2)}, \dots, X_{(n)}$ یک L - برآوردگر ترکیب خطی از آماره‌های ترتیبی می‌باشد. که موارد زیر نمونه‌هایی از L - برآوردگرها می‌باشند

۱- میانه نمونه

$$\bar{X}_\alpha = \frac{\sum_{i=[n\alpha]+1}^{n-[n\alpha]} X_{(i)}}{n - 2[n\alpha]} \quad i = [n\alpha] + 1, \dots, n - [n\alpha] \quad (12.2)$$

۲- برآوردگر گست‌ویث^۱ که یک میانگین وزنی با صدکهای $\frac{1}{4}, \frac{3}{4}, \frac{5}{8}, \frac{3}{8}$ با وزن‌های $\frac{1}{4}, \frac{3}{8}, \frac{3}{8}, \frac{1}{4}$ می‌باشد.

۳- توکی سه میانگین، که میانگین موزون اولین، دومین و سومین، چارک‌های مربوط به آن‌ها دارای وزن $\frac{1}{4}, \frac{1}{4}, \frac{1}{4}$ و $\frac{1}{4}$ می‌باشند.

این سه و سایر L - برآوردگرها توسط اندروز و همکاران (۱۹۷۲)، توضیح داده شده است که به طور کلی L - برآوردگرها دارای وضعیت رگرسیون روشنی درمورد M - برآوردگرها نیست زیرا با استفاده از ρ به صورت زیر می‌باشد

$$\rho(x) = |x| \quad (13.2)$$

که میانه را به عنوان یک برآوردگر (سطح میانه و سطح وضعیت رگرسیون) در نظر می‌گیرد. و این می‌تواند به آسانی تعدیل بیابد برای به دست آوردن سایر صدک‌ها تابع ρ ، به صورت زیر می‌باشد

$$\psi(x) = \begin{cases} -(1-p)x & x < 0 \\ p(x) & x \geq 0 \end{cases} \quad (14.2)$$

صدک $(100p)$ ام در یک تک نمونه و برآوردی از سطح تراز صدک $(100p)$ ام است، یا اینکه سطح وضعیت رگرسیون که برای اطلاعات بیشتر به کوکتر و باست (۱۹۷۸) مراجعه کنید.

¹Gastwirth's Estimator

فصل ۳

آزمون‌های استوارمدل خطی با خطاهای وابسته زبرجمعی منفی

۱.۳ مقدمه

در این فصل، به آزمون فرضیه‌های پارامترهای رگرسیون در مدل‌های خطی با خطاهای وابسته زبرجمعی منفی NSD می‌پردازیم. توزیع مجانبی آماره آزمون بدست آمده و برآوردهایی از پارامترهای زیادی در توزیع بدون علامت ایجاد می‌شود. اطلاعات این فصل از یو و همکاران (۲۰۱۷) و هوپر (۱۹۸۱) گرفته شده است. مدل خطی رگرسیون زیر را در نظر بگیرید

$$Y_i = X_i' \beta + e_i \quad i = 1, 2, \dots, n \quad (1.3)$$

که در (۱.۳) X_i بردار $1 \times p$ معلوم می‌باشد، $\beta \in \mathbb{R}^p$ پارامتر مجهول رگرسیونی، e_i جمله خطاهای تصادفی هستند. $\hat{\beta}_n$ یک M - برآورد برای β به صورت زیر تعریف می‌شود

$$\min \sum_{i=1}^n \rho(e)_i = \sum_{i=1}^n \rho(Y_i - X_i' \beta) \quad (2.3)$$

در رابطه (۲.۳) ρ تابعی محدب و غیریکنواخت می‌باشد، برای نمونه می‌توان از ρ هوپر استفاده کرد.

$$\rho(x) = \left(x^2 I(|x| \leq k) \right) / 2 + \left(k|x| - k^2/2 \right) I(|x| > k), \quad k > 0$$

رفتار مجانبی $\hat{\beta}_n$ برای انتخاب‌های مختلف تابع ρ توسط بای و همکاران (۱۹۹۲)، الش (۱۹۸۶)، یوهای (۱۹۷۴) و یوهای و مارونا (۱۹۷۹) بررسی شده است. غالب نتایج متکی به فرض استقلال باقیمانده‌هاست. برای مثال هوبر (۱۹۸۱) نشان داد شرط ضروری برای درستی نتایج، استقلال باقیمانده‌هاست. برلینت و همکاران (۲۰۰۰) با فرض باقیمانده‌های وابسته شرایط سازگاری M - برآوردها را در مدل‌های خطی با خطاهای آمیخته قوی را ارائه کردند. کوی و همکاران (۲۰۰۴) توزیع مجانبی M - برآوردها را در مدل‌های خطی با خطاهای همبسته فضایی مشخص کردند. وو (۲۰۰۷) در مورد نمایش‌های بهادر قوی و ضعیف برای M - برآوردها در مدل‌های خطی با جملات خطای وابسته تحقیق کرد. وو (۲۰۱۱) شرایط سازگاری قوی M - برآوردها را در مدل‌های خطی با خطاهای پیوندی منفی بررسی کرد.

هو (۲۰۰۰) با ارائه مثالی نشان داد که NSD به معنای پیوند منفی نیست، همچنین برخی از خواص پایه‌ای و دو قضیه اساسی برای NSD ارائه داد و نشان داد که تعدادی از توزیع‌های چند متغیره شناخته شده دارای خاصیت NSD هستند. NSD توسط آلام و ساکسنا (۱۹۹۸) مطرح و این که NSD اشاره به NSD دارد، نلسن (۲۰۰۶) خاصیت وابستگی را می‌توان برای کاربرد نظریه مفصل اعمال کرد، جاورکی و همکاران (۲۰۱۳) با دید اقتصادی به متغیرهای تصادفی NSD نگاه کردند و بیشتر تمرکزشون روی کاربرد اقتصادی قرار دادن، کریسدوندز و واگلاتو (۲۰۰۴) متغیرهای وابسته‌ای که NSD باشند دارای خاصیت NSD نیز هستند مورد بررسی قرار دادند، اقبال و همکاران (۲۰۱۰) قانون قوی اعدادبزرگ برای فرم‌های درجه دوم در متغیرهای تصادفی NSD تحت این فرض که این متغیرها نامنفی هستند، و یک دنباله $\{X_n, n \geq 1\}$

با $EX_n^r \leq \infty$ برای هر $n \geq 1$ و $r \geq 1$ مطرح کردند، همچنین اقبال و همکاران (۲۰۱۱) نامساوی کلموگروف را برای فرم‌های درجه دوم در دنباله نامنفی از متغیرهای تصادفی کراندار یکنواخت NSD ارائه دادند. شن و همکاران (۲۰۱۳) قضیه ثابت همگرایی قریب به یقین و همچنین پایداری قوی برای مقادیر وزنی از متغیرهای تصادفی NSD را مورد مطالعه قرار دادند.

هدف اصلی این بخش بررسی مساله آزمون استوار M برای پارامترهای رگرسیون در مدل $Y_i = X_i' \beta + e_i$ با خطاهای تصادفی NSD است. فرضیه‌های زیر را در نظر بگیرید:

$$H_0 : H^T (\beta - b) = 0 \quad H_1 : H^T (\beta - b) \neq 0 \quad (3.3)$$

که H یک ماتریس معلوم $p \times q$ با رتبه q درجه، $(0 < q \leq p)$ ، b یک بردار p تایی معلوم است. فرضیه مقابل را می‌توان با فرضیه زیر جایگزین کرد:

$$H_{\gamma,n} : H^T (\beta - b) = H^T \omega_n \quad (4.3)$$

که ω_n بردار p معلوم است. به طوریکه

$$\|S_n^{\frac{1}{2}}\omega_n\| = O(1) \quad (5.3)$$

که $S_n = \sum_{t=1}^n x_t x_t^\top$ و $\|\cdot\|$ یک نرم اقلیدسی است. روابط زیر را در نظر بگیرید:

$$\min_{H^\top(\beta-b)=0} \sum_{t=1}^n \rho(y_t - x_t^\top \beta) = \sum_{t=1}^n \rho(y_t - x_t^\top \tilde{\beta})$$

$$\min_{\beta \in \mathbb{R}^p} \sum_{t=1}^n \rho(y_t - x_t^\top \beta) = \sum_{t=1}^n \rho(y_t - x_t^\top \hat{\beta})$$

$$M_n = \sum_{t=1}^n \rho(y_t - x_t^\top \tilde{\beta}) - \sum_{t=1}^n \rho(y_t - x_t^\top \hat{\beta})$$

. که $\hat{\beta}$ و $\tilde{\beta}$ - براوردگرهای پارامترهای مدل $Y_i = X_i' \beta + e_i$ در مدل محدود و بدون محدودیت است. برای آزمون فرضیه (۳.۳) ما معیار M را در نظر می‌گیریم که M_n را معیار اندازه‌گیری سطح انحراف از فرضیه صفر می‌داند. چندین نتیجه‌گیری کلاسیک زمانی که خطاها به نظر می‌رسد مستقل هستند، در مقاله‌های (۸)، (۵۴) و (۵۵) ارائه شده‌اند، این مورد را به خطاهای تصادفی NSD تعمیم می‌دهد. در این بخش C یک مقدار ثابت مثبت می‌باشد. $|\tau| = \max_{1 \leq t \leq p} \{|\tau_1|, |\tau_2|, \dots, |\tau_p|\}$ به طوریکه τ یک بردار p تایی است. و همچنین داریم $x^+ = xI(x \geq 0)$ و $x^- = -xI(x \leq 0)$. یک دنباله تصادفی X_n گفته می‌شود روی نرم L_q ، و $q > 0$ است اگر $E|X_n|^q < \infty$.

ρ را یک تابع محدب و غیریکنوا فرض می‌کنیم، مشتقات چپ و راست تابع ψ را با ψ_- و ψ_+ نمایش می‌دهیم، $\psi_-(u) \leq \psi(u) \leq \psi_+(u)$. چندین فرضیه به شرح زیر در نظر می‌گیریم:

• A۱: تابع $G(u) = E\psi(e_t + u)$ به طوریکه $G(\circ) = E\psi(e_t) = \circ$ و تابع λ دارای مشتق مثبت $u = \circ$ است.

$$\bullet A2: \lim_{u \rightarrow \circ} E|\psi(e_1 + u) - \psi(e_1)|^2 = \circ, \text{ و } \circ < E\psi^2(e_1) = \sigma^2 < \infty$$

$$\bullet A3: \text{مقدار ثابت مثبت } \Delta \text{ وجود دارد طوری که برای هر } h \in (\circ, \Delta) \text{ و تابع}$$

$$\psi(u+h) - \psi(u)$$

نسبت به u یکنوا است.

$$\bullet A4:$$

$$\sum_{t=2}^{\infty} |Cov(\psi(e_1), \psi(e_t))| < \infty$$

$$\bullet A5: \text{از آنجایی که } S_n = \sum_{t=1}^n x_t x_t^\top \text{ فرض می‌کنیم } S_n > \circ, \text{ برای } n \text{ به قدر کافی بزرگ}$$

$$d_n = \max_{1 \leq t \leq n} x_t^\top S_n^{-1} x_t = O(n^{-1}).$$

ملاحظه ۱.۱.۳. شرایط A_1 تا A_4 استاندارد و اغلب در نظریه مجانبی M - برآوردها در مدل‌های رگرسیون خطی کاربرد دارند. برای مثال به ژائو (۲۰۰۲) مراجعه کنید.

ملاحظه ۲.۱.۳. شرط A_5 منطقی است زیرا معادل با کران $\max_{1 \leq t \leq n} |x_t x_t^\top|$ ، است، در این جا یک حالت شرط $d_n = O(n^{-\delta})$ برای $0 < \delta \leq 1$ ، که بوسط اوانگ و همکاران (۲۰۱۵) مورد استفاده قرار گرفت. مثال‌هایی که در زیر آورده شده شرایط A_5 برای آن‌ها برقرار است و توابع $\psi(x)$ که در مثال‌های زیر آورده شده شرایط A_1 و A_5 برقرار است.

مثال ۱.۱.۳. برآوردگر کمترین توان‌های دوم^۱ به صورت زیر می باشد

$$\rho(x) = x^2$$

که تابع $\psi(x)$ آن بصورت زیر می باشد:

$$\psi(x) = 2x.$$

مثال ۲.۱.۳. برآورد کمترین انحراف مطلق^۲ به صورت زیر می باشد

$$\rho(x) = |x|$$

که تابع $\psi(x)$ آن به صورت زیر می باشد:

$$\psi(x) = \text{sign}(x) = \begin{cases} 1 & x > 0 \\ 0 & x = 0 \\ -1 & x < 0 \end{cases}$$

لم ۱.۱.۳. فرض کنید $\{X_n, n \geq 0\}$ دنباله ای از متغیرهای تصادفی NSD باشد که $EX_n = 0$ و برای $p \geq 2$ و $E|X_n|^p < \infty$ آنگاه ثابت مثبت C_p وابسته به p وجود دارد به طوری که برای هر $n \geq 1$

برای اثبات این قضیه می توانید به (۴۰) مراجعه کنید.

لم ۲.۱.۳. وجود دارد یک D زیرمجموعه محدب از $\mathbb{R}^n$ و $\{f_n\}$ تابع محدب روی D برای هر $x \in D$

$$f_n(x) \xrightarrow{\mathbb{P}} f(x)$$

اگر f یک تابع واقعی روی D ، و سپس f محدب باشد. و برای همه زیرمجموعه‌های فشرده $D_0 \subset D$

¹Least Squares Estimator

²Absolute Deviation Estimator

$$\sup_{x \in D_0} |f_n(x) - f(x)| \xrightarrow{\mathbb{P}} 0 \quad (6.3)$$

علاوه بر این اگر f یک تابع متمایز روی D ، $g(x)$ و $g_n(x)$ به ترتیب گرادیان (شیب) و گرادیان فرعی از f را نشان می‌دهند. سپس رابطه (۶.۳) دلالت بر این دارد که برای همه D_0 داریم

$$\sup_{x \in D_0} |g_n(x) - g(x)| \xrightarrow{\mathbb{P}} 0$$

برای اثبات این قضیه می‌توانید به (۷) مراجعه کنید.

لم ۳.۱.۳. فرض کنید $\{X_n, n \geq 1\}$ یک دنباله تصادفی از توزیع NSD با واریانس محدود و آرایه‌ای از مقدار حقیقی $\{a_{nj}, 1 \leq j \leq n\}$ متقاعد است. $\sum_{j=1}^n a_{nj}^2 = O(1)$ ، $\max_{1 \leq j \leq n} |a_{nj}| \rightarrow 0$ برای هر $j = 1, 2, \dots, n$.

$$\left| E \exp \left(i \sum_{j=1}^n r_j Z_{nj} \right) - \prod_{j=1}^n E \exp (i r_j Z_{nj}) \right| \leq \frac{1}{4} \sum_{j \neq \ell, j, \ell=1}^n |r_j r_\ell \text{Cov}(Z_{nj}, Z_{n\ell})|$$

وجود دارد $\{Z_{nj} = \sum_{\ell \in \gamma_j} a_{n\ell} X_\ell\}$ زیرمجموعه‌های ناپیوسته از $\{1, 2, \dots, n\}$

برهان: برای دوتایی متغیرهای تصادفی NSD ، X, Y با استفاده از ویژگی ۳ در لم ۲.۲.۱ داریم

$$H(x, y) = P(X \leq x, Y \leq y) - P(X \leq x)P(Y \leq y) \leq 0. \quad (7.3)$$

اگر $F(x, y)$ تابع توزیع توام از (X, Y) و $F_X(x), F_Y(y)$ تابع توزیع حاشیه‌ای از X, Y داریم:

$$\begin{aligned} \text{Cov}(X, Y) &= E(XY) - E(X)E(Y) = \int \int [F(x, y) - F_X(x)F_Y(y)] dx dy \\ &= \int \int H(x, y) dx dy \end{aligned} \quad (8.3)$$

از رابطه فوق بدست می‌آوریم:

$$\text{Cov}(f(X), g(Y)) = \int \int f'(X)g'(Y)H(x, y) dx dy$$

وجود دارد f, g توابع برآورد شده روی R به طوریکه $f', g' < \infty$ از ترکیب رابطه (۷.۳) و (۸.۳) بدست می‌آوریم:

$$|\text{Cov}(f(X), g(Y))| \leq \int \int |f'(X)||g'(Y)||H(x, y)| dx dy \leq \|f'\|_\infty \|g'\|_\infty |\text{Cov}(X, Y)|. \quad (9.3)$$

به طوریکه $f(X) = \exp(irX), g(Y) = \exp(iuY)$ به راحتی دیده می‌شود

$$\|f'(X)\|_{\infty} \leq r < \infty, \quad \|g'(Y)\|_{\infty} \leq u < \infty$$

در نتیجه برای هر مقادیر r, u داریم

$$|E\exp(irX + iuY) - E\exp(irX)E\exp(iuY)| \leq |ru\text{Cov}(X, Y)| \quad (۱۰.۳)$$

این لم برای $n = ۱$ بدیهی است و $n = ۲$ برای رابطه (۱۰.۳) صدق می‌کند. فرض کنید نتیجه برای همه $n \leq M$ ($n \geq ۳$) برای $n = M+۱$ وجود دارد $\epsilon^۲ = ۱, \delta^۲ = ۱, k \in ۱, ۲, \dots, M$ به طوریکه

$$\begin{cases} \epsilon r_j \geq ۰, & ۱ \leq j \leq k \\ \delta r_j \geq ۰, & ۱+k \leq j \leq M+۱ \end{cases}$$

اگر $X = \sum_{j=۱}^k \epsilon r_j Z_{nj}, Y = \sum_{j=k+۱}^{M+۱} \delta r_j Z_{nj}$ آن‌گاه

$$\sum_{j=۱}^n r_j Z_{nj} = \epsilon X + \delta Y.$$

توجه داشته باشید X, Y توابع ناصعودی هستند. با توجه به فرضیات داریم:

$$\begin{aligned} & \left| E\exp\left(i \sum_{j=۱}^n r_j Z_{nj}\right) - \prod_{j=۱}^n E\exp(ir_j Z_{nj}) \right| \\ & \leq |E(T_۱ T_۲) - E(T_۱)E(T_۲)| + \left| E(T_۱)E(T_۲) - E(T_۱) \prod_{j=۱+k}^{M+۱} E(R) \right| \\ & + \left| E(T_۱) \prod_{j=۱+k}^{M+۱} E(R) - \prod_{j=۱}^k E(R) \right| \\ & \leq |\epsilon \delta| |\text{Cov}(X, Y)| + \left| E(T_۲) - \prod_{j=۱+k}^{M+۱} E(R) \right| + \left| E(T_۱) - \prod_{j=۱+k}^{M+۱} E(R) \right| \\ & \leq \left| \text{Cov}\left(\sum_{j=۱}^k \epsilon r_j Z_{nj}, \sum_{j=k+۱}^{M+۱} \delta r_j Z_{nj}\right) \right| + \frac{۱}{r} \sum_{j \neq \ell, j, \ell = k+۱}^{M+۱} |r_j r_{\ell} \text{Cov}(Z_{nj}, Z_{n\ell})| \\ & + \frac{۱}{r} \sum_{j \neq \ell, j, \ell = ۱}^{M+۱} |r_j r_{\ell} \text{Cov}(Z_{nj}, Z_{n\ell})| \\ & \leq \frac{۱}{r} \sum_{j \neq \ell, j, \ell = ۱}^n |r_j r_{\ell} \text{Cov}(Z_{nj}, Z_{n\ell})| \end{aligned}$$

به طوریکه $T_۱ = \exp(i\epsilon X), T_۲ = \exp(i\delta Y), R = \exp(ir_j Z_{nj})$. در نتیجه اثبات لم کامل می‌شود.

لم ۴.۱.۳. اگر $X_j \xrightarrow{L} X$ ، $X_{nj} \xrightarrow{L} X_j$ برای هر j و اندازه یکنواخت ϱ و برای هر $\varepsilon > 0$ ،

$$\lim_{j \rightarrow \infty} \lim_{n \rightarrow \infty} \sup \{ \varrho(X_{nj}, Y_n) \geq \varepsilon \} = 0$$

سپس

$$Y_n \xrightarrow{L} X.$$

برای اثبات این قضیه می‌توانید به (۱۰) مراجعه کنید.

لم ۵.۱.۳. فرض کنید $\{X_n, n \geq 1\}$ ، $\{a_{nj} | 1 \leq j \leq n\}$ همان شرایط لم از ۳.۱.۳ را داشته باشند فرض کنید $\{X_n, n \geq 1\}$ از خانواده انتگرال از $L_2 - norm$ سپس

$$\sigma_n^{-1} \sum_{j=1}^n a_{nj} X_j \xrightarrow{D} N(0, 1)$$

که $\sigma_n^2 = \text{var}(\sum_{j=1}^n a_{nj} X_j)$

برهان: بدون از دست دادن کلیت، فرض می‌کنیم $a_{nj} = 0$ برای همه $j > n$. با استفاده از رابطه (۸.۳)، $\text{Cov}(X, Y) \leq 0$ ، سپس برای $1 \leq m \leq n-1$ داریم:

$$\begin{aligned} & \sum_{l,j=1, |l-j| \geq m}^n |a_{nl} a_{nj} \text{cov}(X_l, X_j)| = 2 \sum_{l,j=1, l-j \geq m}^n |a_{nl} a_{nj} \text{cov}(X_l, X_j)| \\ & \leq \sum_{l,j=1, l-j \geq m}^n (a_{nl}^2 + a_{nj}^2) |\text{cov}(X_j, X_l)| \\ & \leq \sum_{l,j=1, |l-j| \geq m}^n 2 a_{nj}^2 |\text{cov}(X_j, X_l)| \\ & = \sum_{l,j=1, |l-j| \geq m}^n \sum_{|l-j| \geq m}^n a_{nj}^2 |\text{cov}(X_j, X_l)| \\ & = \sum_{j=1}^n a_{nj}^2 \sum_{|l-j| \geq m}^n |\text{cov}(X_j, X_l)| \\ & \leq \sum_{j=1}^n a_{nj}^2 \left(\sup_j \sum_{|l-j| \geq m, l=1}^n |\text{cov}(X_j, X_l)| \right) \\ & = \sum_{j=1}^n a_{nj}^2 \left(\sup_j \left| \sum_{|l-j| \geq m}^n \text{cov}(X_j, X_l) \right| \right) \\ & = \sup_j \left| \sum_{l=1, |l-j| \geq m}^n \text{cov}(X_l, X_j) \right| \left(\sum_{j=1}^n a_{nj}^2 \right) \end{aligned}$$

اگر در شرط A۴ و $\psi(x) = x$ در نظر بگیریم برای همه $\ell \geq 1$ و j های به قدر کافی بزرگ داریم

$$\sum_{j:|\ell-j|\geq m}^{\infty} |Cov(X_{\ell}, X_j)| \rightarrow 0.$$

بنابراین، برای ε کوچک ثابت، وجود دارد عدد صحیح مثبت $m = m_{\varepsilon}$ چنان که داریم

$$\sum_{j,\ell=1,|\ell-j|\geq m} |a_{n\ell}a_{nj}Cov(X_{\ell}, X_j)| \leq \varepsilon \quad (11.3)$$

$Y_{nj} = \sum_{k=m_{j+1}}^{m(j+1)} a_{nk}X_k$, $j = 0, 1, 2, \dots, n$ ، در نظر می‌گیریم که $[x]$ قسمت صحیح x است، $N_0 = [1/\varepsilon]$

$$\gamma_j = \left\{ j : 2N_0\ell \leq j \leq 2N_0\ell + N_0, |Cov(Y_{nj}, Y_{n(j+1)})| \leq \frac{2}{N_0} \sum_{j=2N_0\ell}^{2N_0\ell+N_0} Var(Y_{n\ell}) \right\}$$

تعریف می‌کنیم

$$s_0 = 0, s_{j+1} = \min\{s : s > s_j, s \in \gamma_j\}$$

با بکار بردن روابط زیر داریم

$$Z_{nj} = \sum_{\ell=s_j+1}^{s_{j+1}} Y_{n\ell}, j = 0, 1, 2, \dots, n$$

$$\Lambda_j = \{m(s_j + 1) + 1, \dots, m(s_{j+1} + 1)\}$$

توجه کنید که

$$Z_{nj} = \sum_{\ell \in \Lambda_j} a_{n\ell}X_{\ell}, j = 0, 1, 2, \dots, n$$

به راحتی می‌توان دید، $\#\Lambda_j \leq 3N_0m$ وجود دارد که در آن $\#$ برابر تعداد عناصر یک مجموعه است. اثبات خواهیم کرد که $\{Z_{nj}, 1 \leq j \leq n\}$ ، شرایط لیندبرگ را دارد. از آنجا که $B_n^2 =$

$$\sum_{j=1}^n EZ_{nj}^2$$

است با استفاده از لم ۱.۱.۳ پس داریم

$$\begin{aligned}
 B_n^2 &= \sum_{j=1}^n E \left(\sum_{\ell \in \Lambda_j} a_{n\ell} X_\ell \right)^2 \\
 &\leq \sum_{j=1}^n a_{n\ell}^2 E \left(\sum_{\ell \in \Lambda_j} X_\ell \right)^2 \\
 &\leq \sum_{j=1}^n \sum_{\ell \in \Lambda_j} a_{n\ell}^2 E(X_\ell)^2 \\
 &\leq \sum_{j=1}^n E(a_{n\ell}^2 X_j^2) = \sigma_n^2 < \infty
 \end{aligned}$$

طبق لم ۳.۱.۳ واریانس $E(X_n)^2 < \infty$ و $\sum_{j=1}^n = O(1)$ پس طبق لم ۳.۱.۳ $E(X_j)^2$ واریانس متناهی است و وقتی واریانس متناهی باشد E هم متناهی است. گوییم $\{Z_{nj}^2\}$ یکنواخت انتگرال پذیر است.

از آنجایی که $\{X_j^2, j \geq 1\}$ به طور یکنواخت انتگرال پذیر است. از این رو برای هر ε مثبت، داریم $\{Z_{nj}, 1 \leq j \leq n\}$ بنابراین

$$\begin{aligned}
 \frac{1}{B_n^2} \sum_{j=1}^n E Z_{nj}^2 I\{|z_{nj}| \geq \varepsilon B_n\} &\leq \sum_{j=1}^n \left(\sum_{\ell \in \Lambda_j} a_{n\ell}^2 \right) \max_{\ell \in \Lambda_j} E X_\ell^2 I\left\{ \sum_{\ell \in \Lambda_j} |X_\ell| \geq \varepsilon / \max_{\ell \in \Lambda_j} |a_{n\ell}| \right\} \\
 &\leq \left(\sum_{j=1}^n a_{n\ell}^2 \right) \max_{1 \leq j \leq n} E X_\ell^2 I\left\{ \sum_{\ell \in \Lambda_j} X_\ell^2 \geq \varepsilon^2 / \left(\max_{\ell \in \Lambda_j} |a_{n\ell}| \right)^2 \right\} \\
 &\leq \left(\sum_{j=1}^n a_{n\ell}^2 \right) \max_{1 \leq j \leq n} E \sum_{\ell \in \Lambda_j} X_\ell^2 I\left\{ \sum_{\ell \in \Lambda_j} X_\ell^2 \geq \varepsilon^2 / \left(\max_{\ell \in \Lambda_j} |a_{n\ell}| \right)^2 \right\}
 \end{aligned}$$

حال $\{Z_{nj}^*, j = 1, 2, \dots, n\}$ دنباله مستقل تصادفی که دارای توزیع حاشیه‌ای یکسانی با Z_{nj} برای هر j و $F(Z_{n1}, Z_{n2}, \dots, Z_{nn})$ و $G(Z_{n1}^*, Z_{n2}^*, \dots, Z_{nn}^*)$ توابع ویژه به ترتیب برای $\sum_{j=1}^n Z_{nj}^*$ و $\sum_{j=1}^n Z_{nj}$ با انتخاب $r = \max(r_\ell, r_j)$ ، با استفاده از لم ۳.۱.۳ و رابطه (۱۱.۳) داریم

$$\begin{aligned}
 |F - G| &= \left| \text{Exp}\left(i \sum_{j=1}^n r_j Z_{nj}\right) - \text{Exp}\left(i \sum_{j=1}^n r_j Z_{nj}^*\right) \right| \\
 &= \left| \text{Exp}\left(i \sum_{j=1}^n r_j Z_{nj}\right) - \prod_{j=1}^n \text{Exp}(ir_j Z_{nj}) \right| \\
 &\leq \frac{1}{r} \sum_{j \neq \ell, j, \ell=1}^n |r_j r_\ell \text{Cov}(Z_{nj}, Z_{n\ell})| \\
 &\leq r^2 \left(\sum_{1 \leq \ell < j \leq n, |\ell-j|=1}^n |\text{Cov}(Z_{nj}, Z_{n\ell})| + \sum_{1 \leq \ell < j \leq n, |\ell-j|>1}^n |\text{Cov}(Z_{nj}, Z_{n\ell})| \right) \\
 &\leq r^2 \left(\sum_{j=1}^n |\text{Cov}(Y_{n\ell_s}, Y_{n\ell_{s+1}})| + \sum_{1 \leq \ell < j \leq n, |\ell-j|>m}^n |a_{n\ell} a_{nj}| |\text{Cov}(X_{nj}, X_{n\ell})| \right) \\
 &\leq r^2 \left(\frac{C}{N_o} \sum_{j=1}^n \text{Var}(Y_{nj}) + \varepsilon \right) \leq \varepsilon (r^2 + C\sigma_n^2)
 \end{aligned}$$

با جمع‌بندی قضیه‌ها، Z_{nj}^* حالت نرمال مجانبی را کسب می‌کند. با استفاده از لم ۴.۱.۳، توزیع مشخصه به طور یکسان از $\{X_j\}$ نتیجه می‌گیریم

$$B_n^{-1} \sum_{j=1}^n Z_{nj} = \sigma_n^{-1} \sum_{j=1}^n a_{nj} X_j \xrightarrow{D} N(0, 1)$$

بنابراین اثبات لم کامل می‌شود.

لم ۶.۱.۳. در مدل $Y_i = X_i' \beta + e_i \quad i = 1, 2, \dots, n$ فرض می‌کنیم دنباله $\{e_i, 1 \leq i \leq n\}$ متغیر تصادفی NSD به طور یکسان شرایط $A1 - A4$ ایجاد کردند. برای هر δ مثبت، n بزرگ ایجاد کردند، سپس

$$\begin{aligned}
 \sup_{|\zeta| \leq \delta} \left| \sum_{t=1}^n \rho(e_t - x_n t^\top \zeta) - \rho(e_t) + \psi(e_t) x_n t^\top \zeta - \frac{1}{r} \lambda \zeta^\top \right| &\xrightarrow{p} 0, \\
 \sup_{|\zeta| \leq \delta} \left| \sum_{t=1}^n \psi(e_t - x_n t^\top \zeta) - \psi(e_t) x_n t^\top + \lambda \zeta \right| &\xrightarrow{p} 0.
 \end{aligned}$$

و ζ یک بردار p تایی است.

برهان:

$$\begin{aligned}
 f_n(\zeta) &= \sum_{t=1}^n \rho(e_t - x_n t^\top \zeta) - \rho(e_t) + \psi(e_t) x_n t^\top \zeta \\
 &= \sum_{t=1}^n \int_0^\zeta x_n t^\top \zeta \{\psi(e_t + u) - \psi(e_t)\} du
 \end{aligned}$$

برای ζ ثابت، از شرط $A5$ پیروی می‌کند چنانچه

$$\max_{1 \leq t \leq n} |x_n t^\top \zeta| \rightarrow O(n^{-1/2}) \quad (12.3)$$

از شرط $A1$ و رابطه‌ی فوق دنباله‌ای از مقادیر مثبت $\varepsilon_n \rightarrow 0$ و $\theta_{nt} \in (-1, 1)$ وجود دارد، چنانچه، برای n های به قدر کافی بزرگ داریم

$$\begin{aligned} Ef_n(\zeta) &= \sum_{t=1}^n \int_0^{-x_n t^\top \zeta} E(\psi(e_t + u) - \psi(e_t)) du \\ &= \sum_{t=1}^n \int_0^{-x_n t^\top \zeta} \{\lambda u + o(|u|)\} du \\ &= \frac{1}{2} \lambda \sum_{t=1}^n (x_n t^\top \zeta)^2 (1 + \varepsilon_n \theta_{nt}) \rightarrow \frac{1}{2} \lambda \zeta^\top \zeta. \end{aligned}$$

از دیدگاه همگرایی $\psi(e_t + u) - \psi(e_t)$ ، مجموعه‌ای از $f_n(\zeta)$ همچنین همگرایی در رابطه با e_t از ویژگی ۲ در لم ۲.۲.۱ با جدا کردن مجموعه از $f_n(\zeta)$ در دو قسمت منفی و مثبت، با استفاده از ویژگی ۱ در لم ۲.۲.۱ که NSD هستند. پس با استفاده از نامساوی شوارتز و رابطه (۱۲.۳) بدست می‌آوریم

$$\begin{aligned} var[f_n(\zeta)] &= E \left\{ \sum_{t=1}^n \left[\int_0^{-x_n t^\top \zeta} (\psi(e_t + u) - \psi(e_t)) du \right]^+ \right. \\ &\quad \left. - \sum_{t=1}^n \left[\int_0^{-x_n t^\top \zeta} (\psi(e_t) + u) - \psi(e_t) \right]^- du \right\}^2 \\ &\leq E \left\{ \sum_{t=1}^n \left[\int_0^{-x_n t^\top \zeta} (\psi(e_t + u) - \psi(e_t)) du \right]^+ \right\}^2 \\ &\quad + E \left\{ \sum_{t=1}^n \left[\int_0^{-x_n t^\top \zeta} (\psi(e_t + u) - \psi(e_t)) du \right]^- \right\}^2 \\ &\leq \sum_{t=1}^n E \left[\int_0^{-x_n t^\top \zeta} (\psi(e_t + u) - \psi(e_t)) du \right]^2 \\ &\leq \sum_{t=1}^n |x_{nt}^\top \zeta| \left| \int_0^{-x_n t^\top \zeta} E[\psi(e_t + u) - \psi(e_t)]^2 du \right| \\ &= o(1) \sum_{t=1}^n (x_{nt}^\top \zeta)^2 \rightarrow 0. \end{aligned}$$

بنابراین برای n های به قدر کافی بزرگ داریم

$$f_n(\zeta) \xrightarrow{p} \frac{1}{2} \lambda \zeta^\top \zeta. \quad (13.3)$$

با استفاده از رابطه ۱۳.۳ و لم ۲.۱.۳ اثبات شد.

لم ۷.۱.۳. تحت شرایط از لم ۶.۱.۳ و با جایگزین کردن رابطه $H_{\Upsilon,n} : H^\top(\beta - b) = H^\top \omega_n$ برای هر δ ثابت مثبت و n های به قدر کافی بزرگ. $\|S_n^{1/2} \omega_n\| = O(1)$

$$\sup_{|\xi_1| \leq \delta} \left| \sum_{t=1}^n \{ \rho(y_{nt} - x_{nt}^\top \eta) - \rho(e_t) + \psi(e_t) x_{nt}^\top \xi_1 \} - \frac{1}{\Upsilon} \lambda \|\xi_1\|^2 \right| \xrightarrow{p} 0 \quad (14.3)$$

$$\sup_{|\xi_1| \leq \delta} \left| \sum_{t=1}^n [\psi(y_{nt} - x_{nt}^\top \eta) - \psi(e_t)] x_{nt}^\top + \lambda \xi_1 \right| \xrightarrow{p} 0 \quad (15.3)$$

$$\sup_{|\xi_2| \leq \delta} \left| \sum_{t=1}^n \{ \rho(y_{nt} - x_{nt}^\top K_n \zeta) - \rho(e_t) \} + \sum_{t=1}^n \psi(e_t) x_{nt}^\top (K_n \zeta - \beta(n)) - \frac{1}{\Upsilon} \lambda \|\xi_2\|^2 + \|\omega(n)\|^2 \right| \xrightarrow{p} 0 \quad (16.3)$$

$$\sup_{|\xi_2| \leq \delta} \left| \sum_{t=1}^n [\psi(y_{nt} - x_{nt}^\top K_n \zeta) - \psi(e_t)] x_{nt}^\top + \lambda (K_n \zeta - \beta(n)) \right| \xrightarrow{p} 0 \quad (17.3)$$

از آنجایی که $\xi_1 = \eta - \beta(n)$, $\xi_2 = \zeta - \tau(n)$

اثبات رابطه (۱.۳):

$$\begin{aligned} & \sup_{|\xi_1| \leq \delta} \left| \sum_{t=1}^n \{ \rho(y_{nt} - X_{nt}^\top (\xi_1 + \beta(n))) - \rho(e_t) + \psi(e_t) X_{nt}^\top \xi_1 \} - \frac{1}{\Upsilon} \lambda \|\xi_1\|^2 \right| \\ &= \sup_{|\xi_1| \leq \delta} \left| \sum_{t=1}^n \{ \rho(y_{nt} - X_{nt}^\top \beta(n) - X_{nt}^\top \xi_1) - \rho(e_t) + \psi(e_t) X_{nt}^\top \xi_1 \} - \frac{1}{\Upsilon} \lambda \|\xi_1\|^2 \right| \\ &= \sup_{|\xi_1| \leq \delta} \left| \sum_{t=1}^n \{ (e_t - X_{nt}^\top \xi_1) - \rho(e_t) + \psi(e_t) X_{nt}^\top \xi_1 \} - \frac{1}{\Upsilon} \lambda \|\xi_1\|^2 \right| \xrightarrow{p} 0 \end{aligned}$$

اثبات رابطه (۲.۳):

$$\begin{aligned} & \sup_{|\xi_1| \leq \delta} \left| \sum_{t=1}^n [\psi(y_{nt} - x_{nt}^\top (\xi_1 + \beta(n))) - \psi(e_t)] X_{nt}^\top + \lambda \xi_1 \right| \\ &= \sup_{|\xi_1| \leq \delta} \left| \sum_{t=1}^n [\psi(y_{nt} - X_{nt}^\top \beta(n) - X_{nt}^\top \xi_1) - \psi(e_t)] X_{nt}^\top + \lambda \xi_1 \right| \\ &= \sup_{|\xi_1| \leq \delta} \left| \sum_{t=1}^n [\psi(e_t - X_{nt}^\top \xi_1) - \psi(e_t)] X_{nt}^\top + \lambda \xi_1 \right| \xrightarrow{p} 0 \end{aligned}$$

اثبات رابطه‌ی (۱۴.۳) و (۱۵.۳) مشابه اثبات رابطه‌ی (۱۶.۳) و (۱۷.۳) هستند.
 • اثباتی از قضیه‌ها
 با در نظر گرفتن مدل مجدد

$$y_{nt} = \mathbf{x}_{nt}^\top \beta(n) + e_t \quad t = 1, 2, \dots, n \quad (18.3)$$

از آنجایی که $\mathbf{x}_{nt} = \mathbf{S}_n^{-1/2} \mathbf{x}_t$ ، $\beta(n) = \mathbf{S}_n^{1/2}(\beta - \mathbf{b})$ ، $y_{nt} = y_t - \mathbf{x}_t^\top \mathbf{b}$ به راحتی می‌توان آن را بررسی کرد

$$\sum_{t=1}^n \|\mathbf{x}_{nt} \mathbf{x}_{nt}^\top\| = p \quad (19.3)$$

فرض کنید $q < p$ باشد. ماتریس $K_p \times (p - q)$ با $(p - q)$ درجه آزادی چنانچه که $\mathbf{H}^\top \mathbf{K} = \mathbf{0}$ و $\mathbf{K}^\top \omega_n = \mathbf{0}$ سپس برای برخی از $\gamma \in R^{p-q}$ ، که H_0 و $H_{\gamma,n}$ را می‌توان به صورت زیر نوشت

$$H_0 : \beta - \mathbf{b} = \mathbf{K}\gamma \quad H_{\gamma,n} : \beta - \mathbf{b} = \mathbf{K}\gamma + \omega_n \quad (20.3)$$

دلالت بر اینکه $\mathbf{H}_n = \mathbf{S}_n^{-1/2} \mathbf{H}(\mathbf{H}^\top \mathbf{S}_n^{-1} \mathbf{H})^{-1/2}$ ، $\mathbf{K}_n = \mathbf{S}_n^{1/2} \mathbf{K}(\mathbf{K}^\top \mathbf{S}_n \mathbf{K})^{-1/2}$ سپس

$$\mathbf{H}_n^\top \mathbf{H}_n = \mathbf{I}_q, \quad \mathbf{K}_n^\top \mathbf{K}_n = \mathbf{I}_{p-q}, \quad \mathbf{H}_n^\top \mathbf{K}_n = \mathbf{0}, \quad \mathbf{H}_n \mathbf{H}_n^\top + \mathbf{K}_n \mathbf{K}_n^\top = \mathbf{I}_p \quad (21.3)$$

با توجه به اینکه $\gamma_0(n) = (\mathbf{K}^\top \mathbf{S}_n \mathbf{K})^{1/2} \gamma$ تحت فرضیه صفر مدل (۱۸.۳) می‌توان به صورت زیر بازنویسی کرد

$$y_{nt} = \mathbf{x}_{nt}^\top \mathbf{K}_n \gamma_0(n) + e_t \quad t = 1, 2, \dots, n$$

مجموعه $\omega(n) = \mathbf{H}_n^\top \mathbf{S}_n^{1/2} \omega_n$ ، $\gamma(n) = \gamma_0(n) + \mathbf{K}_n^\top \mathbf{S}_n^{1/2} \omega_n$ تحت جایگزین‌های مکانی (۲۰.۳).

$$\beta(n) = \mathbf{K}_n \gamma(n) + \mathbf{H}_n \omega(n) \quad (22.3)$$

تعریف می‌کنیم $\hat{\beta}(n) = \mathbf{S}_n^{1/2}(\hat{\beta} - \mathbf{b})$ و $\hat{\gamma}(n)$ حاصل می‌شود

$$\min_{\zeta \in R^{p-q}} \sum_{t=1}^n \rho(y_{nt} - \mathbf{x}_{nt}^\top \mathbf{K}_n \zeta) = \sum_{t=1}^n \rho(y_{nt} - \mathbf{x}_{nt}^\top \mathbf{K}_n \hat{\gamma}(n))$$

$\hat{\beta}(n)$ ، $\hat{\gamma}(n)$ به ترتیب M - برآوردهایی از $\beta(n)$ و $\gamma(n)$ هستند. بنابراین

$$\tilde{\beta} = \mathbf{b} + \mathbf{S}_n^{-1/2} \mathbf{K}_n \hat{\gamma}(n)$$

قضیه ۱.۱.۳. در مدل $Y_{it} = X_i' \beta + e_i, i = 1, 2, \dots, n$ فرض کنید $\{e_t, 1 \leq t \leq n\}$ که یک دنباله از توزیع NSD متغیر تصادفی یک خانواده کامل است. از $L^2 - norm$ و شرایط $A1 - A5$ را شامل می‌شوند. سپس $2\lambda\sigma^{-2}M_n$ مجانب مرکزی نیست توزیع χ^2 با p درجه آزادی است به عبارت دیگر پارامتر $v(n)$ مرکزی نیست داریم

$$2\lambda\sigma^{-2}M_n \xrightarrow{D} \chi_{p, v(n)}^2$$

وجود دارد $H_n = S_n^{-1/2} H (H^T S_n^{-1/2} H)^{-1/2}$ ، $w(n) = H_n^T S_n^{1/2} w(n)$ ، $v(n) = \lambda^2 \sigma^{-2} \|w(n)\|^2$ به طور خاص زمانی که $\|S_n^{1/2} w(n)\| \rightarrow 0$ به این معنی است که پارامترهای واقعی از صفر منحرف شده‌اند. پس $2\lambda\sigma^{-2}M_n$ توزیع مربع χ^2 با p درجه آزادی

$$2\lambda\sigma^{-2}M_n \xrightarrow{D} \chi_p^2$$

برای یک سطح معنی‌دار مشخص، ما می‌توانیم ناحیه رد را به صورت زیر تعیین کنیم:

$$W = (\circ, \chi_p^2(1 - \alpha/2)) \cup (\chi_p^2(\alpha/2), +\infty) \quad (23.3)$$

به طوریکه $\chi_p^2(1 - \alpha/2)$ ، $\chi_p^2(\alpha/2)$ به ترتیب چنک $\alpha/2$ از χ^2 مرکزی توزیعی با p درجه آزادی.

اثبات: طبق رابطه $|\hat{\beta}_n| = O_p(1)$ و رابطه زیر در لم ۷.۱.۳ داریم

$$\sup_{|\xi_2| \leq \delta} \left| \sum_{t=1}^n \{\rho(y_{nt} - x_{nt}^T K_n \zeta) - \rho(e_t)\} + \sum_{t=1}^n \psi(e_t) x_{nt}^T (K_n \zeta - \beta(n)) - \frac{1}{\sqrt{\lambda}} \lambda \|\xi_2\|^2 + \|w(n)\|^2 \right| \xrightarrow{p} 0$$

$$\begin{aligned} & \sum_{t=1}^n [\rho(y_{nt} - x_{nt}^T \hat{\beta}(n)) - \rho(e_t)] + \sum_{t=1}^n \psi(e_t) x_{nt}^T (\hat{\beta} - \beta(n)) - \frac{\lambda}{\sqrt{\lambda}} \|\hat{\beta} - \beta(n)\|^2 \xrightarrow{p} 0 \\ & , \sum_{t=1}^n [\rho(y_{nt} - x_{nt}^T \hat{\beta}(n)) - \rho(e_t)] + \frac{1}{\sqrt{\lambda}} \left\| \sum_{t=1}^n \psi(e_t) x_{nt} \right\|^2 \xrightarrow{p} 0. \end{aligned} \quad (24.3)$$

با استفاده از رابطه بالا داریم

$$\begin{aligned} & \sum_{t=1}^n [\rho(y_{nt} - x_{nt}^T K_n \hat{\gamma}(n)) - \rho(e_t)] - \sum_{t=1}^n \psi(e_t) x_{nt}^T (K_n \hat{\gamma}(n) - K_n \gamma(n) - H_n w(n)) \\ & + \frac{1}{\sqrt{\lambda}} \left\| \sum_{t=1}^n \psi(e_t) K_n^T x_{nt} \right\|^2 - \frac{\lambda}{\sqrt{\lambda}} \|w(n)\|^2 \xrightarrow{p} 0 \\ & = \sum_{t=1}^n [(y_{nt} - x_{nt}^T K_n \hat{\gamma}(n)) - \rho(e_t)] - \sum_{t=1}^n \psi(e_t) x_{nt}^T H_n w(n) \\ & + \frac{1}{\sqrt{\lambda}} \left\| \sum_{t=1}^n \psi(e_t) K_n^T X_{nt} \right\|^2 - \sum_{t=1}^n \psi(e_t) x_{nt}^T H_n w(n) - \frac{\lambda}{\sqrt{\lambda}} \|w(n)\|^2 \xrightarrow{p} 0 \end{aligned} \quad (25.3)$$

با استفاده از رابطه‌های زیر

$$H_n^\top H_n = I_q, K_n^\top K_n = I_{p-q}, H_n^\top K_n = 0, H_n H_n^\top + K_n K_n^\top = I_p$$

داریم

$$\begin{aligned} \lambda \sigma_n^\top M_n &= \left\| \sum_{t=1}^n H_n^\top x_{nt} \psi(et) \right\|^\top + \lambda^\top \|\omega(n)\|^\top + \lambda \omega^\top(n) \sum_{t=1}^n H_n^\top x_{nt} \psi(et) + o_p(1) \\ &= \left\| \sum_{t=1}^n H_n^\top x_{nt} \psi(et) + \lambda \omega(n) \right\|^\top + o_p(1) \end{aligned} \quad (26.3)$$

از آنجایی که $E\psi(et) = 0$, $E\psi^\top(et) = \sigma^\top < \infty$, $\max_{1 \leq t \leq n} \|x_{nt}\| = O(n^{-1/2})$ و کراندار بودن $\sigma_n^\top$ داریم

$$\begin{aligned} \|Var\left(\sum_{t=1}^n x_{nt} \psi(et)\right)\| &= \left\| \sum_{t=1}^n (x_{nt})^\top E\psi^\top(et) \right\| + \left\| \sum_{t=1}^{n-1} \sum_{j=t+1}^n x_{nt} x_{nj}^\top E(\psi(et)\psi(e_j)) \right\| \\ &= \rho \sigma^\top + \left\| S_n^{-1} \|x_t x_j^\top\| \sum_{t=1}^{n-1} \sum_{j=t+1}^n |E(\psi(et)\psi(e_j))| \right\| \\ &= \rho \sigma^\top + \left\| C \|x_t x_j^\top\| \right\| \\ &= \rho \sigma^\top + \rho C = O(1) \end{aligned}$$

از دید $H_n^\top H_n = I_q$ و لم ۵.۱.۳ داریم:

$$\sigma^{-1} \sum_{t=1}^n H_n^\top x_{nt} \psi(et) \xrightarrow{D} N(0, 1) \quad (27.3)$$

بنابراین قضیه فوق از رابطه‌های (۲۶.۳) و (۲۷.۳) پیروی می‌کند.

لم ۸.۱.۳. تحت شرایطی از قضیه ۱.۱.۳، زمانی که $n \rightarrow \infty$ داریم

$$\hat{\beta}(n) - \beta(n) = \lambda^{-1} \sum_{t=1}^n x_{nt} \psi(et) + o_p(1) \quad (28.3)$$

$$\hat{\gamma}(n) - \gamma(n) = \lambda^{-1} \sum_{t=1}^n K_n^\top x_{nt} \psi(et) + o_p(1) \quad (29.3)$$

برهان: از برآورد از رابطه $\sum_{t=1}^n \rho(y_t - x_t^\top \beta)$ تعریف کنیم

$$\left\| S_n^{-1/2} \sum_{t=1}^n \psi(y_t - x_t^\top \hat{\beta}) x_t \right\| = o_p(1). \quad (30.3)$$

به طوریکه $\hat{\beta}(n) = S_n^{1/2} \hat{\beta}_n$ رابطه (۳۰.۳) به صورت زیر بازنویسی می‌کنیم

$$\left\| \sum_{t=1}^n \psi(e_t - x_{nt}^\top \hat{\beta}(n)) x_{nt} \right\| = o_p(1) \quad (31.3)$$

با یک استدلال معمولی، باید ثابت کنیم

$$\hat{\beta}(n) = O_p(1) \quad (32.3)$$

u یک واحد شمارای متراکم در حدود $\mathbb{R}^p$ است. می‌نویسیم:

$$D(\zeta, L) = \sum_{t=1}^n \psi(e_t - L x_{nt}^\top \zeta) x_{nt}^\top \zeta$$

به طوریکه $\zeta \in \mathbb{R}^p$ ، $L \geq 0$ بدیهی است که، $D(\cdot, L)$ ، ζ داده شده صعودی روی L زمانی که ψ صعودی باشد. برای هر $\varepsilon > 0$ ،

$$L_\circ = \sqrt{p} \sigma / (\lambda \sqrt{\varepsilon}).$$

پس با استفاده از رابطه (۳۱.۳) وجود دارد n_1 که $n \geq n_1$ ،

$$Pr \left\{ \left| D \left(\frac{\hat{\beta}(n)}{\|\hat{\beta}(n)\|}, \|\hat{\beta}(n)\| \right) \right| \geq L_\circ \lambda \right\} < \varepsilon.$$

توجه کنید که $D(\zeta, \cdot)$ یک تابع صعودی روی ζ برای L سپس داریم:

$$Pr \{ \|\hat{\beta}(n)\| L_\circ \} < Pr \{ \sup_{\zeta \in u} D(\zeta, L_\circ) \leq -L_\circ \lambda \} + \varepsilon.$$

با توجه در لم ۷.۱.۳ و $\max_{1 \leq t \leq n} |x_{nt}^\top \zeta| = O(n^{-1/2})$ داریم

$$\sup_{\zeta \in u} \left| \sum_{t=1}^n [\psi(e_t - L_\circ x_{nt}^\top \zeta) - \psi(e_t)] x_{nt}^\top \zeta + L_\circ \lambda \right| \rightarrow 0 \quad (33.3)$$

از طرف دیگر با استفاده از نامساوی شوارتز داریم:

$$\sup_{\zeta \in u} \left| \sum_{t=1}^n \psi(e_t) x_{nt}^\top \zeta \right| \leq \left\| \sum_{t=1}^n \psi(e_t) x_{nt} \right\| \quad (34.3)$$

از ترکیب رابطه (۳۳.۳) و (۳۴.۳) وجود دارد n_2 ($n_1 \leq n_2 \leq n$) چنانچه

$$Pr \left\{ \sup_{\zeta \in u} D(\zeta, L_\circ) < -L_\circ \lambda + \left\| \sum_{t=1}^n \psi(e_t) x_{nt} \right\| \right\} > 1 - \varepsilon \quad (35.3)$$

با استفاده از نامساوی چبیشف و نابرابری C_r و لم ۷.۱.۳ داریم:

$$\begin{aligned} Pr \left\{ \left\| \sum_{t=1}^n \psi(e_t) x_{nt} \right\| \geq L_\circ \lambda \right\} &\leq \left(\frac{1}{L_\circ \lambda} \right)^2 E \left\| \sum_{t=1}^n \psi(e_t) x_{nt} \right\|^2 \\ &\leq \left(\frac{1}{L_\circ \lambda} \right)^2 E \sum_{t=1}^n \|\psi(e_t) x_{nt}\|^2 \\ &\leq \left(\frac{1}{L_\circ \lambda} \right)^2 E \sum_{t=1}^n \sigma^2 \|x_{nt}\|^2 \leq \varepsilon \end{aligned} \quad (36.3)$$

از رابطه (۳۵.۳) و (۳۶.۳) می‌توان نوشت

$$Pr\{\sup_{\zeta \in u} D(\zeta, L_o) < -L_o \lambda\} > 1 - 2\varepsilon$$

همچنین، زمانی که $n \geq n_2$ بدست می‌آوریم

$$Pr\{\|\hat{\beta}(n)\| \geq L_o\} < \varepsilon - 1 - (1 - 2\varepsilon) = 3\varepsilon \quad (37.3)$$

بنابراین نتیجه رابطه (۳۲.۳) از رابطه (۳۷.۳) پیروی می‌کند و استبدادی از ε شکل می‌گیرد. با استفاده از لم ۷.۱.۳ و $\max_{1 \leq t \leq n} |x_{nt}^\top \zeta| = O(n^{-1/2})$ از این قرار است که

$$I(\|\hat{\beta}(n)\| \geq L_o) \left| \sum_{t=1}^n [\psi(e_t - x - nt^\top \hat{\beta}(n)) - \psi(e_t)] x_{nt} + \lambda \hat{\beta}(n) \right| \xrightarrow{p} 0$$

دلالت بر این دارد

$$\hat{\beta}(n) = \lambda^{-1} \sum_{t=1}^n x_{nt} \psi(e_t) + o_p(1)$$

در نتیجه، رابطه (۲۸.۳) ثابت شده است. همان‌طور که تعریف شد، رابطه $\beta(n) = K_n \gamma(n) + H_n \omega(n)$ و $\hat{\beta}(n)$ می‌توان به طور مشابه به صورت $\hat{\beta}(n) = K_n \hat{\gamma}(n) + H_n \omega(n)$ نوشته شود. پس به ترتیب با تعویض $\beta(n)$ ، $\hat{\beta}(n)$ و $\gamma(n)$ ، $\hat{\gamma}(n)$ ، رابطه (۲۹.۳) ثابت شده است با $K_n^\top K_n = I_{p-q}$.

قضیه ۲.۱.۳. از آنجایی که

$$\hat{\sigma}_n^2 = n^{-1} \sum_{t=1}^n \psi^2(y_t - x_t^\top \hat{\beta}(n))$$

$$\hat{\lambda}_n = (2nh)^{-1} \sum_{t=1}^n \psi(y_t - x_t^\top \hat{\beta}_n + h) - \psi(y_t - x_t^\top \hat{\beta}_n - h)$$

وجود دارد $h = h_n > 0$ و h_n ترتیبی است برای پیروی کردن

$$h_n/d_n^{1/2} \rightarrow \infty \quad h_n \rightarrow 0 \quad \lim_{n \rightarrow \infty} nh_n^2 > 0$$

با توجه به قضیه یک داریم

$$\hat{\sigma}_n^2 \xrightarrow{P} \sigma^2$$

$$\hat{\lambda}_n^2 \rightarrow \lambda$$

تحت این فرض $\|S_n^{1/2} \omega_n\| \rightarrow 0$ به ترتیب با تعویض λ, σ^2 برآورد ثابت آن‌ها $\hat{\lambda}_n$ و $\hat{\sigma}_n$ سپس داریم:

$$2\hat{\lambda}_n \hat{\sigma}_n^{-2} M_n \xrightarrow{D} \chi_p^2$$

اثبات: مدل زیر را در نظر بگیرید $y_{nt} = x_{nt}^\top \beta(n) + e_t \quad t = 1, 2, \dots, n$ بدون از دست دادن کلیت، فرض کنید پارامتر $\beta(n)$ صحیح، و برابر صفر باشد. برای هر $\delta > 0$ می‌نویسیم:

$$V_n = E|\psi(e_1 + \delta d_n^{1/2}) - \psi(e_1 - \delta d_n^{1/2})|$$

با استفاده از همگرایی ψ ، نامساوی شوارتز و شرط A_2 برای n های به اندازه کافی بزرگ داریم:

$$\begin{aligned} EI(\|\hat{\beta}(n)\| \leq \delta) & \left| \hat{\sigma}_n^2 - n^{-1} \sum_{t=1}^n \psi^2(e_t) \right| \\ & \leq V_n + 2E|\psi(e_1)[\psi(e_1 + \delta d_n^{1/2}) - \psi(e_1 - \delta d_n^{1/2})]| \\ & \leq V_n + 2\delta V_n^{1/2} \rightarrow 0 \end{aligned}$$

با استفاده از لم ۸.۱.۳ داریم:

$$\hat{\sigma}_n^2 = n^{-1} \sum_{t=1}^n \psi(e_t + h) - \psi(e_t - h) \xrightarrow{P} E\psi^2(e_1) = \sigma^2.$$

در نتیجه رابطه $\hat{\sigma}_n^2 \xrightarrow{P} \sigma^2$ ثابت شده است. چن و همکاران به منظور اثبات رابطه $\hat{\lambda}^2 \xrightarrow{P} \lambda$ می‌توان اثبات کرد که

$$(\Psi nh)^{-1} \sum_{t=1}^n \psi(e_t + h) - \psi(e_t - h) \xrightarrow{P} \lambda.$$

در واقع همگرایی از $\psi(e_t + h) - \psi(e_t - h)$ و فرض رابطه

$$h_n/d_n^{1/2} \rightarrow \infty, \quad h_n \rightarrow 0, \quad \lim_{n \rightarrow \infty} nh_n^2 > 0$$

با استفاده از لم ۶.۱.۳ و لم ۱.۱.۳ داریم:

$$\begin{aligned} \text{Var}\left\{(\Psi nh)^{-1} \sum_{t=1}^n [\psi(e_t + h) - \psi(e_t - h)]\right\} & \leq (\Psi n^2 h^2)^{-1} E\left[\sum_{t=1}^n \psi(e_t + h) - \psi(e_t - h)\right]^2 \\ & \leq (\Psi nh^2)^{-1} E[\psi(e_t + h) - \psi(e_t - h)]^2 \rightarrow 0. \end{aligned}$$

از طرف دیگر از آنجایی که $\lim_{n \rightarrow \infty} [G(h) - G(-h)]/(\Psi h) = \lambda$

$$(\Psi nh)^{-1} \sum_{t=1}^n [\psi(e_t + h) - \psi(e_t - h)] = [G(h) - G(-h)]/(\Psi h) + o_p(1) \rightarrow \lambda.$$

اثبات قضیه کامل شد.

فصل ۴

ارزیابی برآوردها و آزمون‌های M با مطالعه شبیه‌سازی

در این فصل، برای ارزیابی روش برآوردیابی استوار و توان آزمون‌های M مطرح‌شده در این پایان‌نامه، از مطالعه شبیه‌سازی مونت کارلویی و نتایج گزارش‌شده در مقاله یو و همکاران (۲۰۱۷) استفاده کرده‌ایم. برای اجرای شبیه‌سازی از نرم‌افزار R استفاده شده و کدهای مربوط برای این هدف در پیوست آمده است.

۱.۴ مثال شبیه‌سازی

در این مطالعه، تحت فرضیه صفر، برآوردگرهای ضرایب رگرسیونی و پارامترهای مازاد توسط روش‌های M مطرح‌شده شامل روش‌های LS، LAD و هوبر محاسبه شدند. همچنین، تحت فرضیه جانشین موضعی، توابع توان آزمون‌های M پیشنهادی با محاسبه ناحیه رد داده‌شده در قضیه ۱.۱.۳ به دست آمدند.

برای تولید یک دنباله از متغیرهای NSD، از رابطه زیر استفاده شد:

$$X_t = a_n Y_t + b_n Z_t \quad t = 1, 2, \dots, n \quad (1.4)$$

که در آن a_n و b_n دنباله‌هایی از اعداد حقیقی مثبت هستند و Y_t و Z_t متغیرهای تصادفی وابسته

منفی (متناظر با $\rho_0 < 0$) هستند که از توزیع توام

$$(Y, Z) \sim N(\mu_1, \mu_2, \sigma_1^2, \sigma_2^2, \rho_0)$$

پیروی می‌کنند. برای اطمینان از درستی فرآیند تولید مشاهدات از یک دنباله NSD، ابتدا نشان خواهیم داد $(e_1, e_2, \dots, e_n)$ که به کمک رابطه (۱.۴) ساخته می‌شوند، یک دنباله متناهی NSD است.

با توجه به توزیع (Y_t, Z_t) ، می‌توان نشان داد

$$\text{Cov}(X_{tt}, X_j) < 0 \quad 1 \leq t < j \leq n.$$

از طرفی، هو (2000) نشان داده است اگر تابع ϕ دارای مشتق‌های جزئی مرتبه دوم پیوسته باشد، آن‌گاه زبرجمعی بودن ϕ معادل است با

$$\frac{\partial^2 \phi}{\partial x_t \partial x_j} \geq 0 \quad 1 \leq t \neq j \leq n.$$

پس تابع $\phi(x_1, x_2, \dots, x_n) = \exp(\sum_{t=1}^n x_t^2)$ یک تابع زبرجمعی است. اکنون با توجه به این‌که مولفه‌های $\{X_t^*, 1 \leq t \leq n\}$ توزیع‌های کناری مشابهی با مولفه‌های $\{X_t, 1 \leq t \leq n\}$ دارند، با استفاده از نامساوی ینسن، می‌توان نوشت

$$\begin{aligned} \frac{E\phi(X_1^*, X_2^*, \dots, X_n^*)}{E\phi(X_1, X_2, \dots, X_n)} &= E \exp \left\{ \left(\sum_{t=1}^n X_t^* \right)^2 - \left(\sum_{t=1}^n X_t \right)^2 \right\} \\ &\geq \exp E \left\{ \left(\sum_{t=1}^n X_t^* \right)^2 - \left(\sum_{t=1}^n X_t \right)^2 \right\} \geq 1. \end{aligned}$$

بنابراین $(X_1, \dots, X_n)$ یک دنباله NSD است.

برای برآورد پارامترها در مثال شبیه‌سازی، تابع هوبر به صورت

$$\rho(x) = (x^2 I(|x| \leq k)) / 2 + (k|x| - k^2/2) I(|x| > k)$$

در نظر گرفته شد، که در آن $k = 1/345\sigma_0$. برای تولید متغیرهای توضیحی رگرسیونی x_t دو سناریو مختلف در نظر گرفته شدند:

(I) با تولید u_t ، برای $t = 1, \dots, n$ ، از توزیع یکنواخت استاندارد، متغیر توضیحی به صورت $x_t = 5u_t$ تعریف شد.

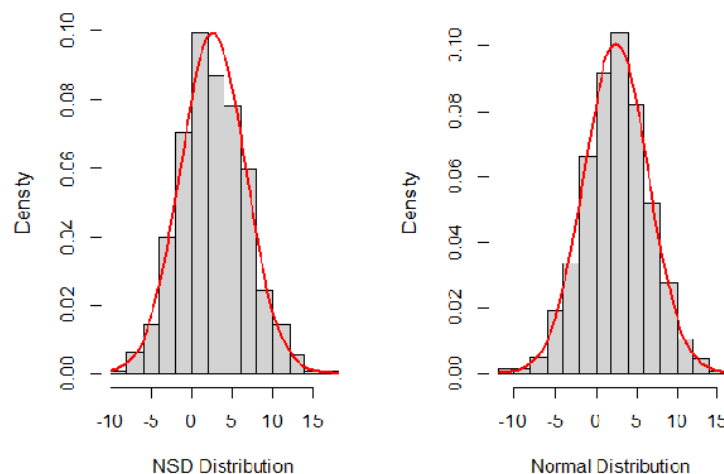
(II) با همان تحقق‌های یکنواخت استاندارد، متغیر توضیحی به صورت $\sin(2t) + 1/5u_t$ در نظر گرفته شد.

سپس مشاهدات از مدل خطی با جمله خطای NSD به صورت زیر تولید شدند:

$$y_t = \beta_0 + \beta_1 x_t + e_t, \quad t = 1, 2, \dots, n$$

که در آن خطاهای NSD $\{e_t = Y_t + Z_t, t = 1, 2, \dots, n\}$ از یک آمیخته چندمتغیره از توزیع نرمال $(Y, Z) \sim N(\mu_1, \mu_2, \sigma_1^2, \sigma_2^2, \rho_0)$ با $\rho_0 < 0$ پیروی می‌کنند. به‌طور دقیق توزیع نرمال دومتغیره به صورت $(Y, Z) \sim N(0, 0, 1, 16, -0.5)$ در نظر گرفته شد. مقادیر واقعی ضرایب رگرسیونی نیز برابر $(\beta_0, \beta_1)^T = (1, 2)^T$ در نظر گرفته شدند. این مقادیر همان وضعیت فرضیه صفر را تعیین می‌کنند؛ یعنی $H_0: (\beta_0, \beta_1)^T = (1, 2)^T$. افزون بر این برای بررسی اندازه نمونه، سه حجم نمونه $n = 100, 500, 1000$ انتخاب شدند. برای کاهش تاثیر تولید تصادفی مشاهدات در برآورد پارامترها، شبیه‌سازی مدل خطی ۱۰۰۰ بار تکرار شد و میانگین برآوردهای پارامترها محاسبه شدند.

شکل ۱.۴ منحنی تابع چگالی توزیع‌های برازش‌شده NSD و نرمال را روی هیستوگرام باقی‌مانده‌های حاصل از مدل رگرسیونی برازش‌شده با روش برآورد M، برای حجم نمونه ۱۰۰۰، نشان می‌دهد. از روی نتیجه این شکل می‌توان نزدیکی دو توزیع برازش‌شده روی باقی‌مانده‌ها را نتیجه گرفت. البته، به‌طور نسبی، می‌توان گفت توزیع NSD برازش‌شده رفتاری نزدیک به یک توزیع بریده‌شده دارد.



شکل ۱.۴: نمودارهای هیستوگرام و منحنی‌های چگالی برازش‌شده باقی‌مانده‌های مدل رگرسیونی با خطای NSD و مستقل برای حجم نمونه $n = 1000$ به ازای یکی از روش‌های M-برآورد

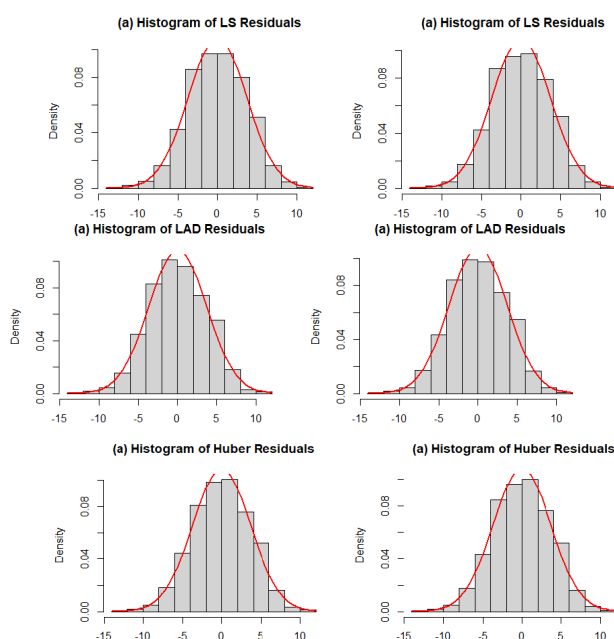
جدول ۱.۴ حاوی نتایج ارزیابی برآوردگرهای ضرایب رگرسیونی و پارامترهای مازاد تحت فرضیه صفر است. نتایج جدول نشان می‌دهند که روش‌های M-برآورد معتبر هستند؛ زیرا برآوردهای متناظر پارامترهای رگرسیونی به مقادیر واقعی $\beta_0 = 1$ و $\beta_1 = 2$ نزدیک هستند. از طرفی برآوردگرهای پارامترهای مازاد در روش‌های مختلف مناسب و موثر هستند. مثلاً برای تابع محدب روش LS، مقادیر واقعی پارامترهای مازاد $\sigma^2 = 13$ و $\lambda = 1$ هستند. برای دو روش دیگر LAD و هوبر نیز وضعیت مشابهی برقرار است. البته مقادیر واقعی پارامترهای مازاد در این روش‌ها متفاوت هستند ولی از روی علامت یکسان برآوردها و معنی‌دار بودن آنها، می‌توان به‌طور ضمنی تشخیص داد که عملکرد روش‌ها قابل قبول و موثر هستند. همچنین می‌توان

گفت که برآوردها حتی با حجم نمونه نه چندان بزرگ مثل $n = 100$ دقت و رفتار خوبی دارد، اما از روی نتایج جدول واضح است که با افزایش اندازه نمونه، دقت برآوردها افزایش می‌یابند.

جدول ۱.۴: ارزیابی ضرایب رگرسیون و پارامترهای زیاد

پارامتر	n	ls.I	ls.II	lad.I	lad.II	هو.بر.I	هو.بر.II
β_0	۱۰۰	۱/۱۶۳	۱/۵۰۴	۰/۹۹۰	۱/۴۹۳	۱/۰۸۹	۱/۵۶۸
	۵۰۰	۱/۱۷۷	۱/۰۱۳	۱/۴۸۶	۱/۱۶۰	۱/۲۱۷	۱/۰۲۵
	۱۰۰۰	۱/۴۲۳	۱/۱۹۳	۱/۲۴۷	۰/۸۹۲	۱/۳۶۶	۱/۱۶۹
β_1	۱۰۰	۱/۹۲۷	۱/۸۱۴	۱/۸۹۲	۱/۶۹۰	۱/۹۷۷	۱/۷۵۳
	۵۰۰	۱/۹۲۷	۲/۱۰۳	۱/۸۴۸	۱/۸۷۹	۱/۹۲۶	۲/۰۶۸
	۱۰۰۰	۱/۸۶۵	۱/۸۷۴	۱/۸۹۱	۲/۰۰۲	۱/۸۷۶	۱/۹۱۰
σ_n^2	۱۰۰	۱۳/۹۵۰	۱۰/۹۵۳	۱/۰۰۰	۱/۰۰۰	۱۰/۴۵۱	۸/۴۹۹
	۵۰۰	۱۳/۲۵۸	۱۲/۱۱۱	۱/۰۰۰	۱/۰۰۰	۹/۶۸۷	۹/۲۰۷
	۱۰۰۰	۱۲/۵۶۵	۱۲/۷۱۴	۱/۰۰۰	۱/۰۰۰	۹/۳۵۴	۹/۳۹۹
λ_n	۱۰۰	۱/۰۰۰	۱/۰۰۰	۰/۲۰۱	۰/۲۳۳	۰/۴۹۹	۰/۳۷۱
	۵۰۰	۱/۰۰۰	۱/۰۰۰	۰/۲۰۰	۰/۲۰۴	۰/۶۷۸	۰/۵۸۵
	۱۰۰۰	۱/۰۰۰	۱/۰۰۰	۰/۲۴۶	۰/۲۴۵	۰/۷۹۰	۰/۷۵۷

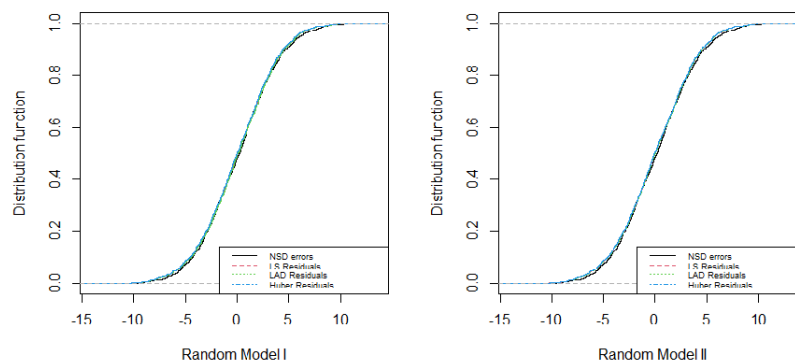
همان‌طور که انتظار داشتیم، بر اساس شکل ۲.۴، منحنی‌های توابع چگالی برازش‌شده باقی‌مانده‌ها به خطاهای NSD فرض‌شده نزدیک هستند و همه آن‌ها باز هم ویژگی یک توزیع بریده‌شده را نشان می‌دهند.



شکل ۲.۴: نمودارهای هیستوگرام و منحنی چگالی برازش‌شده باقی‌مانده‌های M – برآورد با متغیرهای توضیحی مختلف در دو سناریو و روش‌های مختلف M – برآورد به ازای حجم نمونه $n = 1000$

شکل ۳.۴ نموداری است که NSD بودن باقی‌مانده‌ها را ارزیابی می‌کند. برای این کار توزیع تجربی باقی‌مانده‌ها به‌عنوان برآوردی از توزیع واقعی آن‌ها با توابع توزیع برازش‌شده NSD، برای دو سناریو تولید متغیر توضیحی، مقایسه شده‌اند. نتیجه از NSD بودن باقی‌مانده‌ها به خوبی پشتیبانی می‌کند.

به‌عنوان آخرین بخش مطالعه شبیه‌سازی، سطح معنی‌داری تجربی آزمون‌های M پیشنهادی



شکل ۳.۴: مقایسه توابع توزیع NSD فرض شده برای باقی‌مانده‌های رگرسیونی با توزیع تجربی برآورد شده باقی‌مانده‌ها برای حجم نمونه $n = 1000$ و دو سناریو مختلف تولید متغیر توضیحی

ارزیابی شدند. بنا به قضیه ۱.۱.۳ و قضیه ۲.۱.۳، تحت فرضیه جانشین موضعی، $2\hat{\lambda}_n\hat{\sigma}_n^{-2}M_n$ دارای توزیع مجانبی کی دو مرکزی با ۲ درجه آزادی است. بنابراین فرضیه جانشین رد می‌شود اگر مقدار شبیه‌سازی شده $2\hat{\lambda}_n\hat{\sigma}_n^{-2}M_n$ در W قرار گیرد، به طوری که

$$W = (\circ, \chi_p^2(1 - \alpha/2)) \cup (\chi_p^2(\alpha/2), +\infty).$$

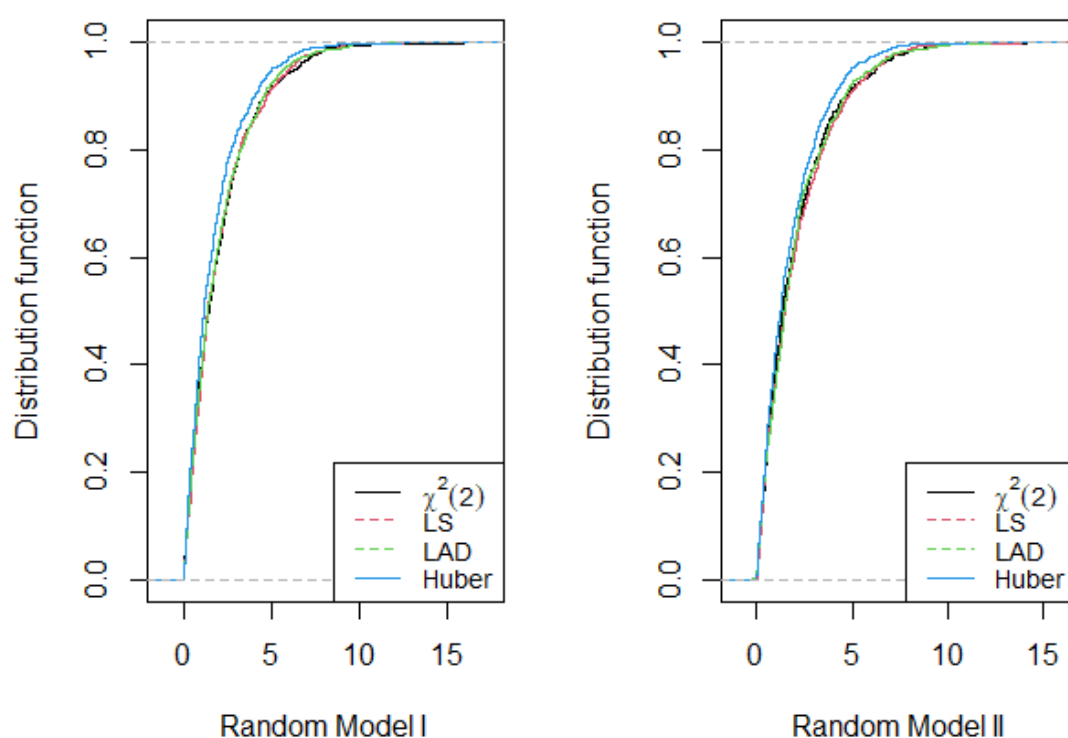
جدول ۲.۴ نتایج اندازه‌های تجربی آزمون‌ها را برای سطوح معنی داری $\alpha = 0.05$ و $\alpha = 0.01$ به ازای انتخاب‌های مختلف روش‌های M-برآورد، سناریوهای تولید متغیر توضیحی و اندازه‌های نمونه مختلف نشان می‌دهد.

جدول ۲.۴: اندازه‌های تجربی آزمون‌های M مختلف در مدل رگرسیون خطی با خطاهای NSD-مقادیری که با * مشخص شده‌اند، بیانگر اندازه‌های معنی داری اسمی هستند.

هوبر.II	هوبر.I	LAD.II	LAD.I	LS.II	LS.I	سطح معنی داری
۰۳۵.۰	۰۴۱.۰	۰۳۸.۰	۰۴۲.۰	۰۴۷.۰	۰۵۷.۰	۰/۰۵*
۰۰۷.۰	۰۰۹.۰	۰۰۷.۰	۰۱۵.۰	۰۱۳.۰	۰۱۶.۰	۰/۱۰*
۰۴۵.۰	۰۴۹.۰	۰۴۸.۰	۰۴۳.۰	۰۵۴.۰	۰۵۱.۰	۰/۰۵*
۰۱۰.۰	۰۰۸.۰	۰۱۲.۰	۰۱۳.۰	۰۰۹.۰	۰۰۹.۰	۰/۱۰*
۰۳۷.۰	۰۴۲.۰	۰۵۰.۰	۰۵۵.۰	۰۴۱.۰	۰۴۹.۰	۰/۰۵*
۰۱۱.۰	۰۰۹.۰	۰۱۴.۰	۰۱۸.۰	۰۰۵.۰	۰۱۰.۰	۰/۱۰*

نتایج جدول بیانگر آن است که اندازه‌های معنی داری تجربی به اندازه‌های اسمی خود نزدیک هستند. بنابراین آزمون‌های M پیشنهادی معتبر و قابل دفاع هستند.

شکل ۴.۴ نیز، با مقایسه تابع توزیع تجربی $2\hat{\lambda}_n\hat{\sigma}_n^{-2}M_n$ برای روش‌های M-برآورد مختلف و تابع توزیع کی دو با دو درجه آزادی، نشان می‌دهد که $2\hat{\lambda}_n\hat{\sigma}_n^{-2}M_n$ می‌تواند توزیع کی دو را به خوبی تقریب بزند. این نتیجه، تایید دیگری بر معتبر بودن آزمون‌های M پیشنهادی تحت فرضیه‌های جانشین موضعی است.



شکل ۴.۴: مقایسه توابع توزیع برازش شده $\hat{\lambda}_n \hat{\sigma}_n^2 M_n$ برای روش‌های M-برآورد مختلف و توزیع کی دو مرکزی با دو درجه آزادی به ازای اندازه نمونه $n = 1000$ و سناریوهای تولید متغیر توضیحی

نتیجه‌گیری مطالعه شبیه‌سازی

نتایج حاصل از شبیه‌سازی نشان می‌دهد که آزمون‌های M برای مدل خطی با خطاهای NSD به انتخاب روش‌های M-برآورد و همین‌طور متغیرهای توضیحی حساس نیستند و بنابراین یک تنومندی را نشان می‌دهند. این نتیجه موثر بودن آزمون M را در چارچوب وابستگی نوع NSD در مدل‌های خطی را تشریح می‌کند.

مراجع

[۱] اقبال، ن. (۱۳۸۴)، پایان نامه ارشد: ”روش های استواری” دانشکده علوم ریاضی، دانشگاه فردوسی مشهد.

[۲] حبیبی کمپی، س. (۱۳۹۵)، پایان نامه ارشد: ”همگرایی کامل برای مجموع موزون متغیرهای تصادفی وابسته زبرجمعی منفی و کاربرد آن در مدل رگرسیون خطی با متغیرهای تبیینی آلوده به خطا” دانشکده علوم ریاضی، دانشگاه صنعتی شاهرود.

[۳] حیدری نژاد، ح. (۱۳۹۷)، پایان نامه ارشد: ”سازگاری قوی M – برآوردها در مدل خطی با خطاهای وابسته زبرجمعی منفی” دانشکده علوم ریاضی، دانشگاه صنعتی شاهرود.

[۴] رحمانی، س. (۱۳۹۵)، پایان نامه ارشد: ”همگرایی کامل برآوردگر مدل رگرسیون ناپارامتری با جملات خطای وابسته زبرجمعی منفی” دانشکده علوم ریاضی، دانشگاه صنعتی شاهرود.

[۵] فریبا، ا. حسین، ف. الهام، خ. نشریه دانشجویی آمار(ندا) – سال سیزدهم – شماره اول، ”روش های شناسایی و تحلیل داده های پرت در تحلیل رگرسیون”، دانشگاه بیرجند، دانشکده ریاضی، گروه آمار.

[6] Alam, K. and Saxena, K.M.L. (1981). ”Positive dependence in multivariate distributions. Positive dependence in multivariate distributions” **Comm Statist Theor Meth.** A10, pp 1183-1196.

[7] Andersen P. K. and Gill R. D. (1982). ”Cox’s regression model for counting processes: a large sample study” **The annals of statistics** 1100-1120.

[8] Bai Z. D. Rao C. R. and Zhao L. C. (1993). ”MANOVA type tests under a convex discrepancy function for the standard multivariate linear model” **Journal of statistical planning and inference** 36(1), 77-90.

-
- [9] Berline, A., Liese, F. and Vajda, I. (2000). "Necessary and sufficient conditions for consistency of M estimates in regression models with general errors." **J. Statist. Plann. Infer.** 89, pp 243-267.
 - [10] Billingsley P. (1968). "Convergence of probability measures John Wiley" **New York**
 - [11] Block, H. W., and Sampson, A. R. (1988). Conditionally ordered distributions, **J Multi Anal**, 27,91-104.
 - [12] Box, G.E.P.(1953). "Non-Normality and Tests on Variances," **Biometrika**, 40,pp 318-335.
 - [13] Chen, X.R. and Wu, Y.H.(1988). "Strong consistency of M-estimates in linear models." **J. Multi. Anal.** 27, pp 116-130.
 - [14] Chen, X.R. and Zhao, L.C. (1995). "Strong consistency of M-estimates of multiple regression coefficients." **J.Syst. Sci. Complex.** 8, pp 82-87.
 - [15] Christofides, T.C. and Vaggelatou, E. (2004). "A connection between supermodular ordering and positive/ negative association." **J. Multi. Anal.** 88, pp 138-151.a
 - [16] Dixon,W.J. and Tukey, J.W. (1968). "Approximate Behavior of the Distribution of Winsorized (Trimming/Winsorization 2)," **Technometrics**, 10, pp 83-98.
 - [17] Dutter, R. (1977). "Numerical Solution of Robust Regression Problems: Computational Aspects, a Comparison, " **Journal of Statistical Computation and Simulation**, 5, pp 207-238.
 - [18] Efron, D. (1965). Increasing Properteies of polya frequency function, *Ann Math Statist*, **36**, 272-279.
 - [19] Eghbal N. Amini M. Bozorgnia A. (2010). "Some maximal inequalities for quadratic forms of negative superadditive dependence random variables" **Statistics and probability letters** 80(7-8), 587-591.
 - [20] Eghbal N. Amini M. and Bozorgnia A. (2011). "On the Kolmogorov inequalities for quadratic forms of dependent uniformly bounded random variables" **Statistics and probability letters** 81(8), 1112-1120.
 - [21] Gut, A. (1994). **Probability:A Graduate Course**, Spring T Stat, Sweden.

- [22] Hampel, F.R. (1974). "The Influence Curve and Its Robust Estimation", **Journal of the American Statistical Association**, 69, pp 383-393.
- [23] Hill, R.W. and Holland, P.W. (1977). "two Robust to least squares regression ," **Journal of the American Statistical Association**, 72, pp 828-833.
- [24] Hogg, R.E. (1799). "Statistical Robustness: One view of its use in applications today," **American Statistical**. 33:3, pp 108-115.
- [25] Holland, P,W. and Welsch R. E. (1977). "Robust Regression Using Iteratively Reweighted Least-Squares," **Communications in Statistics**, A6, pp 813-828.
- [26] Hu, T.Z. (2000). "Negatively supperadditive dependenc of random variables with applications." **Chinese J. Appl. Probab. Statist.** 16, pp 133-144.
- [27] Huber P. J. (1981). Robust Statistics, New York: John Wiley.
- [28] Huber, P.J. (1964),"Robust Estimation of a Location Parameter." **Annals of Mathematical Statisstics**, 35, pp 73-101.
- [29] Huber, P.J.(1972)."Robust Statistic: A Review," Annals of Mathematical Statisstics, 43, pp 1041-1067.
- [30] Huber, P.J.(1973)."Robust Regression: Asymptotics, Conjectures, and Monte Carlo," Anals of Statistics, 1, pp 799-821.
- [31] Jaworski, P., Durante, F. and Hardle, W.K. (2013)."Copulae in Mathematical and Quantitative Finance." **Heidelberg: Springer-Verlag**.
- [32] Kemperman, J.H.B. (1977)."On the FKG-inequalities for measures on a partially ordered space." **Proc. Akad. Wetenschappen Ser. A.** 80, pp 313-331.
- [33] Koul, H.L. and Surgailis, D. (2000)."Second-order behavior of M-estimators in linear regression with long-memory errors." **J.Statist. Plann. Infer.** 91, pp 399-412.
- [34] Moussa-Hamouda, E. and Leon, F. C.(1977). "The robustness of efficitency of abjusted trimmed estimation in linear regression." **Technometrics**. 19, pp 19-34.
- [35] Moussa-Hamouda, E. and Leon, F. C.(1974). "The 0-blue estimators for complete and censored samples in liner regression." **Technometrics**. 16, pp 441-446.

-
- [36] Nelsen, R.B.(2006). **An Introduction to Copulas**. New York: Springer Science + Business Media.
- [37] Rao, C.R. and Zhao, L.C. (1992). "On the consistency of M-estimate in a linear model obtained through an estimating equation." **Statist. Probab. Lett.** 14, pp 79-84.
- [38] Rao, C.R. and Zhao, L.C. (1992). "Linear representation of M-estimates in linear models. Canad." **J.Statist.** 20 pp 359-368.
- [39] Shen Y., Wang X.J., Yang W.Z. and Hu S.H. (2013). "Almost sure convergence theorem and strong stability for weighted sums of NSD random variables." **Acta. Math. Sinica English Ser.** 29, pp 743-756.
- [40] Shen A. Zhang Y. and Volodin, A. (2015). "Applications of the Rosenthal-type inequality for negatively superadditive dependent random variables" **Metrika** 78(3), 295-311.
- [41] Tukey, J.W. (1962). "The Future of Data Analysis, " **Annals of Mathematical Statistics, Statistics**, 33, pp 1-67.
- [42] Wang, X. and Hu, S. (2015). " On the strong consistency of M-estimates in linear models for negatively superadditive dependent errors," **Australian And New Zealand Journal of Statistics**, 57(2), pp 259-274.
- [43] Wang, X.J., Deng , X., Zheng L.L. and Hu S.H. (2014). "Complete convergence for arrays of rowwise negatively superadditive-dependent random variables and its applications." **Statistics** 48, pp 834-850.
- [44] Welsh A. H. (1986), "Bahadur representations robust scale estimators based on regression residuals" **The Annals of Statistics**, 1246-1251.
- [45] Wu, Q.Y. (2006). "Strong consistency of M-estimator in linear model for negatively associated samples." **J.Syst. Sci. Complex.** 19, pp 592-600.
- [46] Wu, Q.Y. and Jiang Y.Y. (2011). "The strong consistency of M-estimator in a linear model for negatively dependent random samples," **Communications in statistics theory and methods.** 40, pp 467-491.
- [47] Wu, Q.Y. (2007). "M-estimation of linear models with dependent errors." **Ann. Statist.** 35, pp 495-521.

- [48] Wu, Q.Y. (2011). The strong consistency of M-estimator in a linear model for negatively dependent random samples. *Commun. Statist. Theor. Meth.* 40, 467-491.
- [49] Yang, S.C. (2002). "Strong consistency of M-estimator in linear model." **Acta Math. Sinica** 45, pp 21-28.(in Chinese).
- [50] Yohai V. J., and Maronna R. A. (1979), "Asymptotic behavior of M-estimators for the linear model" **The Annals of Statistics** 258-268.
- [51] Yu, Y., Hu, H., Liu, L., and Huang, S. (2017). M-test in linear models with negatively superadditive dependent errors. **Journal of inequalities and applications**, 2017(1), .1-21
- [52] Zhao, L.C. (2000). "Some contributions to M-estimation in linear models." **J.Statist.Plann. Infer.** 88, pp 189-203.
- [53] Zhao, L.C. (2000). "Strong consistency of M-estimates in linear models." **China Ser. A** 45, pp 1420-1427.
- [54] Zhao L. C. and Chen X. R. (1991). "Asymptotic behavior of M-test statistics in linear models" **J. Comput. Inf. Syst** 16, 234-248.
- [55] Zhao L. (2008). "Bai's Contribution to M-estimation and Relevant tests in linear models" **In Advances In Statistics** (pp. 19-26).

پیوست آ

کدهای شبیه‌سازی جداول و نمودارهای فصل چهار

۱.آ کدهای مربوط به جدول دو و نمودار شکل چهار

```
##### Model I
```

```
rm(list=ls())
```

```
library(mvtnorm)
```

```
library(MASS)
```

```
set.seed(8888)
```

```
Beta_0 = 1 ; Beta_1 = 2
```

```
n =1000
```

```
ro = -.5
```

```
cov.0 = ro * sqrt(1*16)
```

```

R <- matrix(c(1, cov.0, cov.0, 16),
nrow = 2, ncol = 2)

mu <- c(X = 0, Y = 0)
X2_ls= X2_lad= X2_huber=c()

# Make loop to find the Statistics
rep=1000

for(rr in 1:rep){
YZ <- rmvnorm(n, mean = mu, sigma = R)

u = runif(n)

##### First model (I)
x = 5 * u
##### Second model (II)
#x = c(); for (t in 1:n) { x[t] = sin(2*t) + 1.5 * u[t]}
##### Generate Linear model with NSD errors
e = YZ[,1]+YZ[,2]
y = Beta_0 + Beta_1 * x + e

#plot(x,y)

#fit <- fitdistr(y, densfun="normal")
#fit
#hist(y, pch=20, breaks=13, prob=TRUE, main="" , xlab = "Normal Distribution")
#curve(dnorm(x, fit$estimate[1], fit$estimate[2]), col="red", lwd=2, add=T)
#par(mfrow=c(1,1))

##### LS
fit.ls = lm(y~x)
var.ls = var(predict(fit.ls)-y)

```

```
res1<-predict(fit.ls)-y
#fit <- fitdistr(res1, densfun="normal")
#fit
#hist(res1, pch=20, breaks=13, prob=TRUE, main="(a) Histogram of LS Residuals",

      xlab = "Model I")
#curve(dnorm(x, fit$estimate[1], fit$estimate[2]), col="red", lwd=2, add=T)

##### LAD

library(L1pack)
fit.lad = lad(y~x, method = "EM")
var.lad = var(predict(fit.lad)-y)

res2<-predict(fit.lad)-y
#fit <- fitdistr(res2, densfun="normal")
#fit
#hist(res2, pch=20, breaks=13, prob=TRUE, main="(a) Histogram of LAD Residuals",

      xlab = "Model I")
#curve(dnorm(x, fit$estimate[1], fit$estimate[2]), col="red", lwd=2, add=T)

##### Huber

fit.huber <- rlm(y~x, scale.est = "Huber", method = "M")
var.huber = var(predict(fit.huber)-y)

res3<-predict(fit.huber)-y
#fit <- fitdistr(res3, densfun="normal")
#fit
#hist(res3, pch=20, breaks=13, prob=TRUE, main="(a) Histogram of Huber Residuals",

      xlab = "Model I")
#curve(dnorm(x, fit$estimate[1], fit$estimate[2]), col="red", lwd=2, add=T)
```

```
##### another way of lad
#library("quantreg")
#fit.lad2 <- rq(y ~ x)
##### ecdf plots Figure 3

#par(mfrow=c(1,1))
#plot(ecdf(e), xlab="Random Model I", ylab = "Distribution function", main = "")
#lines(ecdf(res1), col = 2)
#lines(ecdf(res2), col = 3)
#lines(ecdf(res3), col = 4)
#legend("bottomright", c("NSD errors", "LS Residuals", "LAD Residuals",
  "Huber Residuals"), col=c(1,2,3,4),
#      lty= 1:4, cex = 0.55)

##### Mn Functions

Mn.ls <- sum((y-Beta_0-Beta_1*x)^2)/2-sum((y-fit.ls$coefficients[1]-fit.
ls$coefficients[2]*x)^2)/2

Mn.lad <- sum(abs(y-Beta_0-Beta_1*x))-sum(abs(y-fit.lad$coefficients[1]-fit.
lad$coefficients[2]*x))

w0=y-Beta_0-Beta_1*x
w1=y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x
#..... Find the Best k in Huber
MAR = median(abs(predict(fit.huber)-y))    #median absolute residual
k = 1.345 * (MAR/ 0.6745)
k
#.....

Mn.huber <- sum(0.5*((w0^2)*ifelse(abs(w0)<=k,1,0))+
(k*abs(w0)-0.5*(k^2))*ifelse(abs(w0)>k,1,0))-
```

```

sum(0.5*((w1^2)*ifelse(abs(w1)<=k,1,0))+
(k*abs(w1)-0.5*(k^2))*ifelse(abs(w1)>k,1,0))

Mn.ls; Mn.lad; Mn.huber

##### Sigma

sigma.ls = n^(-1)*sum((y-fit.ls$coefficients[1]-fit.ls$coefficients[2]*x)^2)

sigma.lad = n^(-1)*sum((sign(y-fit.lad$coefficients[1]-fit.lad$coefficients[2]*x))^2)

sigma.huber = n^(-1)*sum(((k*ifelse((y-fit.huber$coefficients[1]-fit.
huber$coefficients[2]*x)< -k,1,0))+
(y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x)*
ifelse(abs(y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x)<=k,1,0)
+k*ifelse((y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x)>k,1,0))^2)

##### Lambda

h=1

lamb.ls = (2*n*h)^(-1)*sum((y-fit.ls$coefficients[1]
-fit.ls$coefficients[2]*x+h)-(y-fit.ls$coefficients[1]-fit.ls$coefficients[2]*x-h))

lamb.lad = (2*n*h)^(-1)*sum(sign(y-fit.
lad$coefficients[1]-fit.lad$coefficients[2]*x+h)
-sign(y-fit.ls$coefficients[1]-fit.ls$coefficients[2]*x-h))

lamb.huber = (2*n*h)^(-1)*sum((-k*ifelse((y-fit.huber$coefficients[1]
-fit.huber$coefficients[2]*x+h)< -k,1,0))+
(y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x+h)
*ifelse(abs(y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x+h)<=k,1,0)
+k*ifelse((y-fit.huber$coefficients[1]
-fit.huber$coefficients[2]*x+h)>k,1,0)-(-k*ifelse((y-fit.huber$coefficients[1]
-fit.huber$coefficients[2]*x-h)< -k,1,0))+
(y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x-h)

```

```

*ifelse(abs(y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x-h)<=k,1,0)
+k*ifelse((y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x-h)>k,1,0))

#####

#                               Model I

##### Figure 4

X2_ls[rr] = 2*lamb.ls*(1/sigma.ls)*Mn.ls

X2_lad[rr] = 2*lamb.lad*(1/sigma.lad)*Mn.lad

X2_huber[rr] = 2*lamb.huber*(1/sigma.huber)*Mn.huber

}

par(mfrow=c(1,2))
plot(ecdf(rchisq(rep,2)), xlab="Random Model I", ylab = "Distribution function", main = "")
lines(ecdf(X2_ls), col = 2)
lines(ecdf(X2_lad), col = 3)
lines(ecdf(X2_huber), col = 4)
legend("bottomright", c(expression(chi^2*(2)), "LS","LAD", "Huber"), col=c(1,2,3,4),
lty= c(1,2,2,1), cex = 0.85)

##### density plot

plot(density(X2_huber), col=4, main= "densities plot Model I")
lines(density(rchisq(rep, 2)), col=1)
lines(density(X2_lad), col=3)
lines(density(X2_ls), col=2)
legend("topright", c(expression(chi^2*(2)),"LS","LAD","Huber"), col=1:4, lwd = 1)

##### Table 2

alpha=0.01
X2.L <- qchisq(alpha/2, df = 2, lower.tail = T)
X2.R <- qchisq(alpha/2, df = 2, lower.tail = F)

```

```

pvalue.ls = length(which(X2_ls<=X2.L))/rep+length(which(X2_ls>=X2.R))/rep
pvalue.lad = length(which(X2_lad<=X2.L))/rep+length(which(X2_lad>=X2.R))/rep
pvalue.huber = length(which(X2_huber<=X2.L))/rep+length(which(X2_huber>=X2.R))/rep

p100.05.II = c(pvalue.ls , pvalue.lad, pvalue.huber)
p100.01.II = c(pvalue.ls , pvalue.lad, pvalue.huber)

p500.05.II = c(pvalue.ls , pvalue.lad, pvalue.huber)
p500.01.II = c(pvalue.ls , pvalue.lad, pvalue.huber)

p1000.05.II = c(pvalue.ls , pvalue.lad, pvalue.huber)
p1000.01.II = c(pvalue.ls , pvalue.lad, pvalue.huber)
#####
p100.05.I = c(pvalue.ls , pvalue.lad, pvalue.huber)
p100.01.I = c(pvalue.ls , pvalue.lad, pvalue.huber)

p500.05.I = c(pvalue.ls , pvalue.lad, pvalue.huber)
p500.01.I = c(pvalue.ls , pvalue.lad, pvalue.huber)

p1000.05.I = c(pvalue.ls , pvalue.lad, pvalue.huber)
p1000.01.I = c(pvalue.ls , pvalue.lad, pvalue.huber)

LS.I=c(p100.05.I[1],p100.01.I[1],p500.05.I[1],p500.01.I[1],p1000.05.I[1]
,p1000.01.I[1])
LS.II=c(p100.05.II[1],p100.01.II[1],p500.05.II[1],p500.01.II[1]
,p1000.05.II[1],p1000.01.II[1])

LAD.I=c(p100.05.I[2],p100.01.I[2],p500.05.I[2],p500.01.I[2]
,p1000.05.I[2],p1000.01.I[2])
LAD.II=c(p100.05.II[2],p100.01.II[2],p500.05.II[2],p500.01.II[2]
,p1000.05.II[2],p1000.01.II[2])

Huber.I=c(p100.05.I[3],p100.01.I[3],p500.05.I[3],p500.01.I[3])

```

```
,p1000.05.I[3],p1000.01.I[3])
Huber.II=c(p100.05.II[3],p100.01.II[3],p500.05.II[3],p500.01.II[3],
,p1000.05.II[3],p1000.01.II[3])

table_2 = cbind(rep(c(0.05,0.1), times=3),LS.I,LS.II,LAD.I,LAD.II,Huber.I,Huber.II)
save(table_2, file = "table2.RData")

library(xtable)

table_2=data.frame(table_2)
xtable(table_2, digits = 3)

##### Table 1
##..... n= 100

out100_1 <- rbind(c(fit.ls$coefficients[1],
fit.lad$coefficients[1],fit.huber$coefficients[1]),
c(fit.ls$coefficients[2],fit.lad$coefficients[2],fit.huber$coefficients[2]),
c(sigma.ls, sigma.lad, sigma.huber),
c(lamb.ls, lamb.lad, lamb.huber))
colnames(out100_1)=c("LS","LAD","Huber")
rownames(out100_1)=c("B0","B1","Sigma2","Lambda")
#save(out100_1, file = "out100_1.Rdata")

##..... n=500

out500_1 <- rbind(c(fit.ls$coefficients[1],
fit.lad$coefficients[1],fit.huber$coefficients[1]),
c(fit.ls$coefficients[2],fit.lad$coefficients[2],fit.huber$coefficients[2]),
c(sigma.ls, sigma.lad, sigma.huber),
c(lamb.ls, lamb.lad, lamb.huber))
colnames(out500_1)=c("LS","LAD","Huber")
rownames(out500_1)=c("B0","B1","Sigma2","Lambda")
```

```

save(out500_1, file = "out500_1.Rdata")

##..... n=1000

out1000_1 <- rbind(c(fit.ls$coefficients[1],
fit.lad$coefficients[1],fit.huber$coefficients[1]),
c(fit.ls$coefficients[2],fit.lad$coefficients[2],
fit.huber$coefficients[2]),
c(sigma.ls, sigma.lad, sigma.huber),
c(lamb.ls, lamb.lad, lamb.huber))
colnames(out1000_1)=c("LS","LAD","Huber")
rownames(out1000_1)=c("B0","B1","Sigma2","Lambda")
save(out1000_1, file = "out1000_1.Rdata")

##### table of results

ls.I=c(out100_1[1,1],out500_1[1,1],out1000_1[1,1],
out100_1[2,1],out500_1[2,1],out1000_1[2,1],
out100_1[3,1],out500_1[3,1],out1000_1[3,1],
out100_1[4,1],out500_1[4,1],out1000_1[4,1])
ls.II=c(out100_2[1,1],out500_2[1,1],out1000_2[1,1],
out100_2[2,1],out500_2[2,1],out1000_2[2,1],
out100_2[3,1],out500_2[3,1],out1000_2[3,1],
out100_2[4,1],out500_2[4,1],out1000_2[4,1])

lad.I=c(out100_1[1,2],out500_1[1,2],out1000_1[1,2],
out100_1[2,2],out500_1[2,2],out1000_1[2,2],
out100_1[3,2],out500_1[3,2],out1000_1[3,2],
out100_1[4,2],out500_1[4,2],out1000_1[4,2])
lad.II=c(out100_2[1,2],out500_2[1,2],out1000_2[1,2],
out100_2[2,2],out500_2[2,2],out1000_2[2,2],
out100_2[3,2],out500_2[3,2],out1000_2[3,2],
out100_2[4,2],out500_2[4,2],out1000_2[4,2])

```

```

huber.I=c(out100_1[1,3],out500_1[1,3],out1000_1[1,3],
out100_1[2,3],out500_1[2,3],out1000_1[2,3],
out100_1[3,3],out500_1[3,3],out1000_1[3,3],
out100_1[4,3],out500_1[4,3],out1000_1[4,3])
huber.II=c(out100_2[1,3],out500_2[1,3],out1000_2[1,3],
out100_2[2,3],out500_2[2,3],out1000_2[2,3],
out100_2[3,3],out500_2[3,3],out1000_2[3,3],
out100_2[4,3],out500_2[4,3],out1000_2[4,3])

table=cbind(ls.I,ls.II,lad.I,lad.II,huber.I,huber.II)
colnames(table)=c("ls.I","ls.II","lad.I","lad.II","huber.I","huber.II")

nn = rep(c(100,500,1000),times=4)
model=rep(c("I","II"), times=3)
parameters=c("B0","B1","Sigma2","Lambda")

table=data.frame(table)
#save(table, file = "table.RData")
xtable(cbind(parameters,nn,table))

```

۲.آ کدهای مربوط به جدول یک نمودارهای یک الی سه

```
##### Model II
```

```
rm(list=ls())
```

```
library(mvtnorm)
```

```
library(MASS)
```

```
set.seed(8888)
```

```
Beta_0 = 1 ; Beta_1 = 2
```

```
n =1000
```

```
ro = -.5
cov.0 = ro * sqrt(1*16)
R <- matrix(c(1, -2, -2, 16),
nrow = 2, ncol = 2)

mu <- c(X = 0, Y = 0)
YZ <- rmvnorm(n, mean = mu, sigma = R)

u = runif(n)

##### First model (I)
#x = 5 * u
##### Second model (II)
x = c(); for (t in 1:n) { x[t] = sin(2*t) + 1.5 * u[t]}
##### Generate Linear model with NSD errors
e = YZ[,1]+YZ[,2]
y = Beta_0 + Beta_1 * x + e

#plot(x,y)

fit <- fitdistr(y, densfun="normal")
fit
hist(y, pch=20, breaks=13, prob=TRUE, main="" , xlab = "Normal Distribution")
curve(dnorm(x, fit$estimate[1], fit$estimate[2]), col="red", lwd=2, add=T)
par(mfrow=c(1,2))

##### LS
fit.ls = lm(y~x)
var.ls = var(predict(fit.ls)-y)

res1<-predict(fit.ls)-y
fit <- fitdistr(res1, densfun="normal")
fit
```

```
hist(res1, pch=20, breaks=13, prob=TRUE, main="(a) Histogram of LS Residuals",
     xlab = "Model II")
curve(dnorm(x, fit$estimate[1], fit$estimate[2]), col="red", lwd=2, add=T)
```

```
##### LAD
```

```
library(L1pack)
fit.lad = lad(y~x, method = "EM")
var.lad = var(predict(fit.lad)-y)
```

```
res2<-predict(fit.lad)-y
fit <- fitdistr(res2, densfun="normal")
fit
hist(res2, pch=20, breaks=13, prob=TRUE, main="(a) Histogram of LAD Residuals",
     xlab = "Model II")
curve(dnorm(x, fit$estimate[1], fit$estimate[2]), col="red", lwd=2, add=T)
```

```
##### Huber
```

```
library(MASS)

fit.huber <- rlm(y~x, scale.est = "Huber", method = "M")
var.huber = var(predict(fit.huber)-y)
```

```
res3<-predict(fit.huber)-y
fit <- fitdistr(res3, densfun="normal")
fit
hist(res3, pch=20, breaks=13, prob=TRUE, main="(a) Histogram of Huber Residuals",
     xlab = "Model II")
curve(dnorm(x, fit$estimate[1], fit$estimate[2]), col="red", lwd=2, add=T)
```

```
##### another way of lad
```

```
#library("quantreg")
#fit.lad2 <- rq(y ~ x)
```

```
##### ecdf plots
```

```

par(mfrow=c(1,2))
plot(ecdf(e), xlab="Random Model II", ylab = "Distribution function", main = "")
lines(ecdf(res1), col = 2)
lines(ecdf(res2), col = 3)
lines(ecdf(res3), col = 4)
legend("bottomright", c("NSD errors", "LS Residuals", "LAD Residuals",
"Huber Residuals"), col=c(1,2,3,4),
lty= 1:4, cex = 0.55)

##### Mn Functions

Mn.ls <- sum((y-Beta_0-Beta_1*x)^2)/2-sum((y-fit.ls$coefficients[1]
-fit.ls$coefficients[2]*x)^2)/2

Mn.lad <- sum(abs(y-Beta_0-Beta_1*x))-sum(abs(y-fit.lad$coefficients[1]
-fit.lad$coefficients[2]*x))

w0=y-Beta_0-Beta_1*x
w1=y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x

#..... Find the Best k in Huber

MAR = median(abs(predict(fit.huber)-y)) #median absolute residual
k = 1.345 * (MAR/ 0.6745)
k
#.....

Mn.huber <- sum(0.5*((w0^2)*ifelse(abs(w0)<=k,1,0))+
(k*abs(w0)-0.5*(k^2))*ifelse(abs(w0)>k,1,0))-
sum(0.5*((w1^2)*ifelse(abs(w1)<=k,1,0))+
(k*abs(w1)-0.5*(k^2))*ifelse(abs(w1)>k,1,0))

Mn.ls; Mn.lad; Mn.huber

```

```
##### Sigma
```

```
sigma.ls = n^(-1)*sum((y-fit.ls$coefficients[1]-fit.ls$coefficients[2]*x)^2)
```

```
sigma.lad = n^(-1)*sum((sign(y-fit.lad$coefficients[1]-fit.lad$coefficients[2]*x))^2)
```

```
sigma.huber = n^(-1)*sum(((k*ifelse((y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x)< -k,1,0))+
(y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x)
*ifelse(abs(y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x)<=k,1,0)
+k*ifelse((y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x)>k,1,0))^2)
```

```
##### Lambda
```

```
h=1
```

```
lamb.ls = (2*n*h)^(-1)*sum((y-fit.ls$coefficients[1]-fit.
ls$coefficients[2]*x+h)-(y-fit.ls$coefficients[1]-fit.ls$coefficients[2]*x-h))
```

```
lamb.lad = (2*n*h)^(-1)*sum(sign(y-fit.lad$coefficients[1]-fit.
lad$coefficients[2]*x+h)-sign(y-fit.ls$coefficients[1]-fit.ls$coefficients[2]*x-h))
```

```
lamb.huber = (2*n*h)^(-1)*sum(((k*ifelse((y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x+h)< -k,1,0))+
(y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x+h)*ifelse(abs(y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x+h)<=k,1,0)
+k*ifelse((y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x+h)>k,1,0)
-((k*ifelse((y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x-h)< -k,1,0))+
(y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x-h)
*ifelse(abs(y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x-h)<=k,1,0)
+k*ifelse((y-fit.huber$coefficients[1]-fit.huber$coefficients[2]*x-h)>k,1,0))
```

```
#####
```

```
#
# Model II
##### Table 2
X2_ls = 2*lamb.ls*(1/sigma.ls)*Mn.ls

X2_lad = 2*lamb.lad*(1/sigma.lad)*Mn.lad

X2_huber = 2*lamb.huber*(1/sigma.huber)*Mn.huber

alpha = 0.05

pchisq(X2_ls, df=2, lower.tail = F)
pchisq(X2_lad, df=2, lower.tail = F)
pchisq(X2_huber, df=2, lower.tail = F)

##### Table 1
##..... n=100

out100_2 <- rbind(c(fit.ls$coefficients[1],fit.
lad$coefficients[1],fit.huber$coefficients[1]),
c(fit.ls$coefficients[2],fit.lad$coefficients[2],fit.huber$coefficients[2]),
c(sigma.ls, sigma.lad, sigma.huber),
c(lamb.ls, lamb.lad, lamb.huber))
colnames(out100_2)=c("LS","LAD","Huber")
rownames(out100_2)=c("B0","B1","Sigma2","Lambda")
save(out100_2, file = "out100_2.Rdata")

##..... n=500

out500_2 <- rbind(c(fit.ls$coefficients[1],
fit.lad$coefficients[1],fit.huber$coefficients[1]),
c(fit.ls$coefficients[2],fit.lad$coefficients[2],fit.huber$coefficients[2]),
c(sigma.ls, sigma.lad, sigma.huber),
c(lamb.ls, lamb.lad, lamb.huber))
```

```
colnames(out500_2)=c("LS","LAD","Huber")
rownames(out500_2)=c("B0","B1","Sigma2","Lambda")
save(out500_2, file = "out500_2.Rdata")

##..... n=1000

out1000_2 <- rbind(c(fit.ls$coefficients[1],fit.lad$coefficients[1],
fit.huber$coefficients[1]),
c(fit.ls$coefficients[2],fit.lad$coefficients[2],fit.huber$coefficients[2]),
c(sigma.ls, sigma.lad, sigma.huber),
c(lamb.ls, lamb.lad, lamb.huber))
colnames(out1000_2)=c("LS","LAD","Huber")
rownames(out1000_2)=c("B0","B1","Sigma2","Lambda")
save(out1000_2, file = "out1000_2.Rdata")
```

Aabstract

In general, the robustness of a statistical method can be defined as its insensitivity to mild departures from the basic assumptions. In other words, if the characteristics of a statistical method are not affected by a small deviation from the basic assumptions, the method is described as robust. Because outliers strongly influence classical estimates of regression coefficients in linear models, robust estimators such as M-estimators are suggested. The level and power of classical tests are also significantly reduced in the presence of suspicious data. For solving the problem, robust tests are introduced. To perform a robust test about regression parameters, the asymptotic distribution of test statistic is determined. The methods are evaluated based on the reliability of the model estimators and power of tests. In addition to examining some basic features of negatively super-additive dependent random variables in this thesis, we study robust M-tests in linear models with negatively super-additive dependent errors under certain conditions. Using Monte Carlo simulation studies, we compare classic and robust methods based on estimators' reliability and power of tests as well.

Keywords: Robust tests, negatively super-additive dependent random variables, linear model, Monte Carlo simulation.



Faculty Of Mathematical Sciences

MSc Thesis in Mathematical Statistics

Robust M-tests in linear models with negatively superadditive dependent errors

By: Samane Darvishi Bohloli

Supervisor:

Negar Eghbal

Advisor:

Hossein Baghishani

February 2021