

بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ



دانشکده علوم ریاضی
رشته آمار، گرایش آمار ریاضی

پایان نامه کارشناسی ارشد

برآوردگر انقباضی ماتریس کوواریانس در بعد بالا

زهرا حسین پور

استاد راهنما

دکتر محمد آرشی

تیر ۱۳۹۷

شماره:

تاریخ:

باسمه تعالی



دانشگاه صنعتی شاهرود

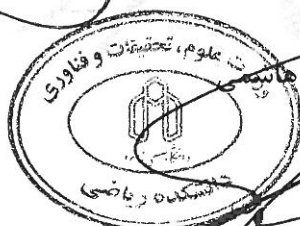
مدیریت تحصیلات تکمیلی

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

با نام و یاد خداوند متعال، ارزیابی جلسه دفاع از پایان نامه کارشناسی ارشد خانم زهرا حسین پور با شماره دانشجویی ۹۴۰۵۸۴۴ رشته آمار گرایش آمار ریاضی تحت عنوان: برآوردگر انقباضی ماتریس کوواریانس در بعد بالا که در تاریخ ۱۳۹۷/۰۴/۲۵ با حضور هیأت محترم داوران در دانشگاه صنعتی شاهرود برگزار گردید به شرح ذیل اعلام می گردد:

☐ مردود	☒ قبول (با درجه: خیلی خوب)
☐ عملی	☒ نظری
نوع تحقیق:	

عضو هیأت داوران	نام و نام خانوادگی	مرتبه علمی	امضاء
۱- استاد راهنمای اول	دکتر محمد آرشی	دانشیار	
۲- استاد راهنمای دوم	-----	-----	-----
۳- استاد مشاور	-----	-----	-----
۴- نماینده تحصیلات تکمیلی	دکتر محمدرضا ربیعی	استادیار	
۵- استاد ممتحن اول	دکتر داود شاهسونی	دانشیار	
۶- استاد ممتحن دوم	دکتر حسین باغیشنی	استادیار	



نام و نام خانوادگی رئیس دانشکده: دکتر ابراهیم هاشمی

تاریخ و امضاء و مهر دانشکده:

تبصره: در صورتی که کسی مردود شود حداکثر یکبار دیگر (در مدت مجاز تحصیل) می تواند از پایان نامه خود دفاع نماید (دفاع

مجدد نباید زودتر از ۴ ماه برگزار شود).

تقدیم به مقدس ترین واژه مادر لغت نامه دلم،
مادر مهربانم که زندگیم را دیون مهر و عطوفت آن می دانم.
پدر، مهربانی مشفق، بردبار و حامی.
همسرو فرزندم که نشانه لطف الهی در زندگی من هستند.
خواهر و برادرانم، همراهان، همبستگی و پشتوانه های زندگیم.

سپاس‌گزاری

خدایا! شکر که تو را حدی برایت متصور نیست. سپاس تو را که بی نهایت مهربانی. خدایا! حمد و سپاس از آن توست که محیطی بر تمام عالم و تو تکیه گاهی بر تمام موجودات. معبودا! تو آن قدر گویایی که من از وصف تو ناتوانم. بند بند وجودم از شعاع نور توست که گرم و روشن می شود.

اکنون که با عنایت خداوند متعال این دوره‌ی علمی به پایان رسیده است وظیفه خود را می‌دانم که تشکر قلبی خود را نسبت به کسانی که به طریقی در انجام این پژوهش مرا یاری نموده‌اند ابراز دارم.

در آغاز وظیفه‌ی خود می‌دانم از زحمات بی‌دریغ استاد راهنمای خود، جناب آقای دکتر محمد آرشی صمیمانه تشکر و قدردانی کنم که از راهنمایی‌های ارزنده ایشان در راستای پیشبرد پژوهش حاصل فراوان بردم و همواره شاگرد مکتب علم و انسانیت و منش والای ایشان هستم.

همچنین لازم می‌دانم از اساتید فرهیخته جناب آقای دکتر داود شاهسونی، جناب آقای دکتر حسین باغیثی که داوری این پایان‌نامه را به عهده گرفتند با تمام وجود تشکر و قدردانی نمایم.

در پایان از همه دوستان و عزیزانی که با اندیشه و قلبشان در این راه یاری‌ام نمودند، خالصانه تشکر و قدردانی می‌نمایم.

زهره حسین‌پور
تیر ۱۳۹۷

تعهد نامه

اینجانب زهرا حسین پور دانشجوی کارشناسی ارشد رشته آمار ریاضی دانشکده علوم ریاضی دانشگاه صنعتی شاهرود، نویسنده پایان نامه با عنوان برآوردگر انقباضی ماتریس کوواریانس در بعد بالا، تحت راهنمایی دکتر محمد آرشی متعهد می‌شوم:

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

زهرا حسین پور
تیر ۱۳۹۷

مالکیت نتایج و حق نشر

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

چکیده

در تحلیل‌های چندمتغیره یکی از موارد چالش برانگیز برآورد ماتریس کوواریانس یک متغیر تصادفی، زمانی که تعداد متغیرها (در این جا بعد) بزرگ باشد، است. همچنین در بعضی بحث‌های چندمتغیره لازم است که تابعی از ماتریس کوواریانس برآورد شود. فرض کنید p نشان‌دهنده بعد باشد. در این پایان‌نامه برآوردگر ماتریس کوواریانس و تابعی از آن را برای حالت بعد بالا ($p \rightarrow \infty$) مورد بررسی قرار می‌دهیم. در این راستا فرض می‌کنیم جامعه نمونه‌گیری دارای توزیع نرمال و غیرنرمال بوده و برآوردگرهای ناپارامتری انقباضی ماتریس کوواریانس در بعد بالا را معرفی کرده و برخی خواص آن‌ها را بررسی می‌کنیم. همچنین با استفاده از مثال‌های شبیه‌سازی، کارایی برآوردگرهای انقباضی معرفی شده برای ماتریس کوواریانس را در بعد بالا مورد ارزیابی قرار می‌دهیم.

کلمات کلیدی:

ماتریس کوواریانس، برآوردگر انقباضی، بعد بالا، سازگاری برآوردگر ناپارامتری.

پیش‌گفتار

مسئله‌ای را در نظر بگیرید که یک نمونه تصادفی به حجم n از یک بردار تصادفی p بعدی (p متغیره) در اختیار داریم و علاقه‌مندیم بر اساس این نمونه ماتریس کوواریانس Σ را که در اندازه $p \times p$ است برآورد کنیم. برآوردگر ماتریس کوواریانس نمونه S به راحتی قابل محاسبه است اما چنانچه بعد p بزرگ باشد خواص خوبی ندارد. در این حالت یافتن یک برآوردگر سازگار برای Σ در بعد بالا با توجه به کارایی آن از اهمیت بسزایی برخوردار است. یک راهکار مناسب و کارا استفاده از برآوردگر انقباضی Σ در بعد بالا است که تاکنون توسط محققین زیادی مورد مطالعه قرار گرفته است. در این راستا لدویت و ولف (۲۰۰۳)، سریواستاوا (۲۰۰۵) و فیشر و سان (۲۰۰۹) برآوردگرهای انقباضی ماتریس کوواریانس در بعد بالا را به خوبی بررسی کرده‌اند. در این پایان‌نامه، با توجه به اهمیت موضوع مورد بررسی، ساختار برآوردگرهای انقباضی خطی را برای ماتریس کوواریانس در بعد بالا بررسی کرده و تحقیقات جدید در خصوص برآوردگر انقباضی Σ و توابعی از آن را در جامعه نرمال و غیرنرمال ارائه می‌کنیم. با انجام مطالعات شبیه‌سازی برای مقادیر متفاوت حجم نمونه n و بعد متغیر p کارایی برآوردگرهای انقباضی را ارزیابی کرده و آن را از دیدگاه معیارهای متفاوت مقایسه می‌کنیم. این مجموعه شامل ۴ فصل و یک پیوست است که در ادامه به مطالب هر فصل به طور مختصر اشاره شده است.

در **فصل اول**، مقدمه‌ای بر برآوردگر ماتریس کوواریانس جامعه همراه با برآوردگرهای انقباضی نوع استاین Σ بیان می‌شود.

در **فصل دوم**، برآورد برخی توابع ماتریس کوواریانس را که در بحث آزمون فرضیه‌ها از آن‌ها استفاده می‌شود، در بعد بالا مورد بررسی قرار می‌دهیم.

فصل سوم به برآورد انقباضی ماتریس کوواریانس در بعد بالا در حالتی که توزیع جامعه نامعلوم است، اختصاص دارد.

فصل چهارم نیز در برگیرنده معرفی خانواده کلی برآوردگرهای انقباضی خطی Σ و مقایسه آن در بعد بالا است.

برنامه‌های کامپیوتری برای مثال‌های شبیه‌سازی به کار رفته، در پیوست آمده‌اند.

فهرست مطالب

م	فهرست جداول
۱	مقدمه
۴	۱.۱ تعاریف و توابع ریاضی
۱۱	۲.۱ برآورد
۱۲	۳.۱ برآوردگر انقباضی نوع استاین
۱۴	۴.۱ انواع پارامترهای انقباضی
۱۴	۱.۴.۱ پارامتر انقباضی LW
۱۵	۲.۴.۱ پارامتر انقباضی RBLW
۱۵	۳.۴.۱ پارامتر انقباضی SS
۱۵	۴.۴.۱ پارامتر انقباضی FS
۱۷	۲ برآورد توابع ماتریس کوواریانس در بعد بالا
۱۷	۱.۲ مقدمه
۱۸	۲.۲ نااریبی و سازگاری برآوردگرها
۲۵	۳.۲ شبیه‌سازی‌های عددی
۲۹	۳ برآوردگرهای ناپارامتری انقباضی نوع استاین ماتریس کوواریانس در بعد بالا
۲۹	۱.۳ مقدمه
۳۲	۲.۳ چارچوب بعد بالا
۳۲	۳.۳ نتایج اصلی
۳۴	۴.۳ شبیه‌سازی
۳۷	۴ مقایسه برآوردگر انقباضی خطی ماتریس کوواریانس در توزیع نرمال و غیر نرمال
۳۷	۱.۴ مقدمه
۳۸	۲.۴ برآوردگر انقباضی خطی
۳۹	۳.۴ برآورد وزن بهینه

۴۱		۴.۴	ماتریس هدف کروی
۴۷		۵.۴	مطالعه شبیه‌سازی
۵۰		۶.۴	خلاصه و پیشنهادات

۵۵		آ	برنامه‌های کامپیوتری
----	--	---	----------------------

۷۱		مراجع
----	--	-------

فهرست جداول

۲۶	۱.۲	مقدار واقعی a_2 و برآوردهای $\hat{a}_2$ و $\tilde{a}_2$ (و خطاهای استاندارد)
۲۷	۲.۲	مقدار واقعی a_1^* و برآوردهای $\hat{a}_1^*$ و $\tilde{a}_1^*$ (و خطاهای استاندارد)
۲۸	۳.۲	مقدار واقعی κ_{11}/p و برآوردهای $\hat{\kappa}_{11}/p$ (و خطاهای استاندارد)
۳۵	۱.۳	برآورد پارامتر انقباضی $\lambda = 1$ برای سناریوی ۱ به ازای $\Sigma = I_p$
۳۶	۲.۳	مقادیر $SPRIAL(\hat{S}^*)$ برای سناریوی ۲ به ازای $\Sigma = I_p$
	۳.۳	مقادیر $SPRIAL(\hat{S}^*)$ برای سناریوی ۳ به ازای Σ ماتریسی است که (a, b)
۳۶		امین عضو آن به صورت $\Sigma_{ab} = 1/I(a - b)$ است
	۱.۴	مقدار PRIAL برآوردهای انقباضی خطی برای مدل (الف) به ازای ρ
۵۱		$p = 100$ و $0.2, 0.9$ در حالت‌های ۱، ۲ و ۳
	۲.۴	مقدار PRIAL برآوردهای انقباضی خطی برای مدل (ب) به ازای ρ
۵۲		$p = 100$ و $0.2, 0.9$ در حالت‌های ۱، ۲ و ۳
	۳.۴	مقدار PRIAL برآوردهای انقباضی خطی برای مدل‌های (ج) و (د) به ازای
۵۳		$h = 0.7$ و $p = 100$ در حالت‌های ۲ و ۳
	۴.۴	مقدار PRIAL برآوردهای انقباضی خطی برای مدل‌های (الف) و (ب) به
۵۳		ازای $\rho = 0.2$ ، $p = 100$ و $n = 40$ در توزیع t_m برای $m = 10, 6, 5, 4, 3$

فصل ۱

مقدمه

گاهی اوقات در آمار، ماتریس کوواریانس یک متغیر تصادفی چندمتغیره مجهول است و باید برآورد شود. برآورد ماتریس‌های کوواریانس با این سوال مواجه می‌شود که چگونه می‌توان ماتریس واقعی کوواریانس را بر اساس نمونه مشاهده شده از توزیع چندمتغیره تقریب نمود. در مواردی که مشاهده‌ها کامل هستند، می‌توان از ماتریس کوواریانس نمونه استفاده کرد. معمولاً این ماتریس، یک برآوردگر نااریب و کارا برای ماتریس کوواریانس است. تحلیل آماری داده‌های چندمتغیره، اغلب شامل مطالعات اکتشافی^۱ از متغیرهایی که با یکدیگر ارتباط دارند، می‌باشند و ممکن است با مدل‌های آماری صریح با ماتریس کوواریانس متغیرها بیان شود. بنابراین، برآورد ماتریس کوواریانس به‌طور مستقیم از داده‌های مشاهده‌شده، دو نقش مهم دارد:

(۱) برای ارزیابی برآوردهای اولیه که می‌تواند برای مطالعه روابط درون متغیرها مورد استفاده قرار گیرد.

(۲) برای ارزیابی برآوردهای نمونه‌ای که می‌تواند برای بررسی مدل به‌کار رود.

برآوردهای ماتریس کوواریانس در مراحل اولیه تحلیل مولفه‌های اصلی^۲ و تحلیل عاملی^۳ مورد نیاز است. همچنین در مدل‌های رگرسیونی که برخی از متغیرهای پیشگو به یکدیگر وابستگی دارند، به‌کار برده می‌شود.

فرض کنید $x_1, x_2, \dots, x_n$ مشاهده‌های مستقل از بردار تصادفی p بعدی $\mathbf{X}$ باشند. ماتریس کوواریانس به‌صورت زیر تعریف می‌شود

$$\Sigma = E \left[(\mathbf{X} - E[\mathbf{X}])(\mathbf{X} - E[\mathbf{X}])^T \right].$$

^۱Exploratory

^۲Principle component

^۳Factor analysis

یک برآوردگر ناریب برای ماتریس Σ ، ماتریس کوواریانس نمونه است که به صورت زیر تعریف می شود

$$S = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})(x_i - \bar{x})^T, \quad (1.1)$$

که در آن

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i.$$

در مواردی که بعد داده ها بالاست مانند مطالعات مربوط به بیماری سرطان، برآوردگر ماتریس کوواریانس نمونه، اغلب ضعیف عمل می کند. برای مثال به میرهد^۴ (۱۹۸۲)؛ جان استون^۵ (۲۰۰۱)؛ بیکل و لوینا^۶ (۲۰۰۸)؛ و فن و ال وی^۷ (۲۰۰۸) مراجعه کنید. روش هایی اخیرا برای برآورد بهبودیافته ماتریس کوواریانس ارائه شده اند. در این راستا می توان به روش هایی در لم و فن^۸ (۲۰۰۹)، رتمن^۹ و همکاران (۲۰۰۹)، بیکل و لوینا (۲۰۰۸)، وو و پوراحمدی^{۱۰} (۲۰۰۳) و جان استون و لو^{۱۱} (۲۰۰۹) اشاره کرد.

در برآورد بردار میانگین توزیع نرمال چندمتغیره، استاین^{۱۲} (۱۹۵۶) نشان داد که می توان برآوردگر درستنمایی ماکسیمم^{۱۳} (MLE) را زمانی که بعد p غیر قابل اغماض است، بهبود بخشید. استاین (۱۹۶۴) به این موضوع پی برد که بررسی این عقیده بهتر از MLE می باشد. نظرات مختلفی در این زمینه ارائه شده اند که به چندین مورد اشاره می کنیم. هاف^{۱۴} (۱۹۸۱) و (۱۹۷۹) برآوردگر جدیدی برای ماتریس دقت^{۱۵} که معکوس ماتریس کوواریانس، یعنی Σ^{-1} ، است، مبتنی بر دو تابع زیان مختلف برای توزیع ویشارت ارائه داد. افرون و موریس^{۱۶} (۱۹۷۶) و یانگ و برگر^{۱۷} (۱۹۹۴) با استفاده از رویکرد بیزی به ترتیب ماتریس دقت و ماتریس کوواریانس را برآورد کردند. دی و سرینیواسان^{۱۸} (۱۹۸۵) به برآوردگری دست یافتند که تحت تابع زیان مشخص، مینیماکس است.

^۴Muirhead

^۵Johnstone

^۶Bickel and Levina

^۷Fan and Lv

^۸Lam and Fan

^۹Rothman

^{۱۰}Wu and Pourahmadi

^{۱۱}Lu

^{۱۲}Stein

^{۱۳}Maximum likelihood estimator

^{۱۴}Haff

^{۱۵}Precision matrix

^{۱۶}Efron and Morris

^{۱۷}Yang and Berger

^{۱۸}Dey and Srinivasan

اساس کار همه محققین کاهش متوسط توابع زیان $L(\hat{\Sigma}, \Sigma)$ یا $L(\hat{\Sigma}^{-1}, \Sigma^{-1})$ در مقایسه با $L(S, \Sigma)$ یا $L(S^{-1}, \Sigma^{-1})$ و دسترسی به برآوردگر مناسبی برای کاهش این معیار بوده است. این کار برای توابع زیان مختلفی با روش‌های متفاوتی انجام شده است که در این قسمت تعدادی از این روش‌ها را بیان خواهیم نمود. به‌طور خلاصه تعدادی از کارهایی که به برآورد انقباضی منتهی نمی‌شود را بیان خواهیم نمود. بیشتر تکنیک‌ها، شرایطی را برای ماتریس کوواریانس، برای به‌دست آوردن برآوردگری مناسب، در نظر گرفته‌اند. رورتو^{۱۹} (۲۰۰۰)، اسمیت و کهن^{۲۰} (۲۰۰۲)، و چن و دانسون^{۲۱} (۲۰۰۳)، از رویکرد تجزیه چولسکی^{۲۲} استفاده کردند.

همل و هارتلی^{۲۳} (۱۹۷۳) نشان دادند که چطور تابع هدف و مشتقات آن را برای برآورد درست‌نمایی ماکسیمم مولفه‌های واریانس می‌توان محاسبه کرد. این کار سپس توسط کربیل و سارل^{۲۴} (۱۹۷۶) دنبال شد. هاف (۱۹۹۱) نظر کلی برای برآوردگر بیزی بردار میانگین و ماتریس کوواریانس را با کمینه کردن زیان بیزی فراهم کرد. فرالی و برنز^{۲۵} (۱۹۹۵) نشان دادند که چگونه توابع احتمال و مشتقات آن‌ها توسط روش‌های ماتریس پراکنده (تنک^{۲۶}) محاسبه می‌شود و نتایج کلی را تعمیم دادند. پوراحمدی (۱۹۹۹) به بررسی مشترک مدل‌های میانگین-واریانس پرداخت. دنیلز و پوراحمدی^{۲۷} (۲۰۰۲) و دنیل (۲۰۰۵) از یک رویکرد بیزی توسط کاوش‌های قبلی برای مشاهدات وابسته استفاده کردند. وو و پوراحمدی (۲۰۰۳) یک برآوردگر ناپارامتری را برای ماتریس کوواریانس پیشنهاد نمودند.

دنیلز (۲۰۰۶) چندین برآوردگر برای ماتریس کوواریانس مبتنی بر چندین مدل رایج رگرسیون بیزی پیشنهاد کرد. چادهوری و همکاران^{۲۸} (۲۰۰۷) الگوریتمی را برای محاسبه برآورد درست‌نمایی ماکسیمم ماتریس کوواریانس تحت این محدودیت که کوواریانس صفر است، ارائه دادند. فن و ال وی (۲۰۰۸) در مدل کوواریانس به برآوردگر مناسبی با استفاده از تکنیک‌های تحلیل عاملی چندمتغیره دست یافتند.

کائو و بومن^{۲۹} (۲۰۰۹)، ماتریس را مبتنی بر فرضیه‌ای که مقادیر ویژه با استفاده از تبدیل ماتریس پراکنده ارائه شده است، برآورد نمودند. فان و وو (۲۰۰۸) روش نیمه پارامتری را که در آن واریانس با استفاده از یک تابع ناپارامتری تخمین زده می‌شود و همبستگی با یک تابع پارامتری برآورد می‌شود، در برآورد در نظر گرفتند.

^{۱۹}Roverato

^{۲۰}Smith and Kohn

^{۲۱}Chen and Dunson

^{۲۲}Cholesky decomposition

^{۲۳}Hemmerle and Hartley

^{۲۴}Corbeil and Searle

^{۲۵}Fraley and Burns

^{۲۶}Sparse matrix

^{۲۷}Daniels and Pourahmadi

^{۲۸}Chaudhuri et al.

^{۲۹}Cao and Buoman

مواردی که بیان شد تعدادی از پیشرفت‌های اخیر در برآورد ماتریس کوواریانس و ماتریس دقت می‌باشند. با این حال، در خصوص رده کلی از برآوردگرهای ماتریس کوواریانس و ماتریس دقت که به طور معمول به نام استاتین یا انقباضی یا برآوردگر انقباضی^{۳۰} نام گرفته است، مطالعات چندانی صورت نگرفته است. در این پایان‌نامه قصد داریم به مطالعه برآوردگرهای انقباضی ماتریس کوواریانس بپردازیم.

۱.۱ تعاریف و توابع ریاضی

برای اطلاع دقیق‌تر از تعاریف این بخش می‌توان به رودین^{۳۱} (۱۳۷۳)، میرهد (۱۹۸۲) و کرمی‌کبیر (۱۳۹۲) مراجعه کرد.

تعریف ۱.۱.۱ (تابع گاما^{۳۲}). تابع گاما را به ازای هر عدد حقیقی $a > 0$ با نماد $\Gamma(a)$ نشان داده و به صورت زیر تعریف می‌کنیم

$$\Gamma(a) = \int_0^{\infty} e^{-x} x^{a-1} dx.$$

تعریف ۲.۱.۱ (توزیع ویشارت^{۳۳}). فرض کنید $W = (w_{ij})$ ماتریس متقارن معین مثبت با احتمال یک باشد، و Σ نیز ماتریس $q \times q$ معین مثبت باشد. اگر $m \geq m$ (عدد صحیح)، آنگاه W دارای توزیع ویشارت با m درجه آزادی است اگر تابع چگالی پیوسته‌ی $\frac{1}{q}(q+1)$ عناصر مجزا W (بالا مثلی گفته می‌شود) به صورت زیر باشد

$$f(w_{11}, w_{12}, \dots, w_{qq}) = c^{-1} \det W^{\frac{m-q}{2} - \frac{1}{2}} \text{etr} \left(-\frac{1}{2} \Sigma^{-1} W \right)$$

که منظور از etr همان e^{trace} می‌باشد و

$$c = 2^{\frac{mq}{2}} \det \Sigma^{\frac{m}{2}} \Gamma_q \left(\frac{m}{2} \right)$$

که $\Gamma_q(a)$ تابع گامای چندمتغیره به صورت

$$\Gamma_q(a) = \pi^{\frac{q(q-1)}{2}} \prod_{j=1}^q \Gamma \left[a + \frac{1}{2}(1-j) \right]$$

می‌باشد. همچنین می‌توان گفت که $W \sim W_q(m, \Sigma)$ است.

تعریف ۳.۱.۱. فرض کنید X_n و Y_n دو دنباله تصادفی حقیقی مقدار باشند. آن‌گاه وقتی

$$n \rightarrow \infty$$

^{۳۰} Shrinkage estimator

^{۳۱} Rudin

^{۳۲} Gamma function

^{۳۳} Wishart distribution

(۱) $X_n = O(Y_n)$ ، O نماد اوی بزرگ) اگر و فقط اگر به ازای هر $\epsilon > 0$ ، وجود داشته باشند $\delta > 0$ و $N > 0$ به طوری که برای همه $n > N$ ،

$$P\left(\left|\frac{X_n}{Y_n}\right| > \delta\right) < \epsilon.$$

(۲) $X_n = o(Y_n)$ ، o نماد اوی کوچک) اگر و فقط اگر به ازای هر $\epsilon > 0$ ،

$$\lim_{n \rightarrow \infty} P\left(\left|\frac{X_n}{Y_n}\right| > \epsilon\right) = 0.$$

تعریف ۴.۱.۱ (مجموعه محدب^{۳۴}). مجموعه C محدب است، اگر پاره خط بین هر دو نقطه دلخواه در C ، درون C قرار گیرد. به عبارتی به ازای هر x_1 و x_2 عضو مجموعه C و به ازای هر θ ثابت، با شرط $0 \leq \theta \leq 1$ ، داشته باشیم

$$\theta x_1 + (1 - \theta)x_2 \in C.$$

تعریف ۵.۱.۱ (تابع محدب). تابع $f: \mathbb{R}^n \rightarrow \mathbb{R}$ محدب است، اگر دامنه f یک مجموعه محدب باشد و به ازای هر x و y عضو دامنه f و $0 \leq \theta \leq 1$ داشته باشیم

$$f(\theta x + (1 - \theta)y) \leq \theta f(x) + (1 - \theta)f(y).$$

تابع $f(\cdot)$ را مقعر گوییم اگر $-f(\cdot)$ محدب باشد.

تعریف ۶.۱.۱ (نرم^{۳۵}). تابع $f: \mathbb{R}^p \rightarrow \mathbb{R}$ را نرم گوییم، اگر

(۱) $f(\cdot)$ تابعی نامنفی باشد، یعنی به ازای هر $x \in \mathbb{R}^p$ داشته باشیم، $f(x) \geq 0$.

(۲) $f(\cdot)$ معین^{۳۶} باشد، یعنی $f(x) = 0$ اگر و تنها اگر $x = 0$.

(۳) $f(\cdot)$ همگن^{۳۷} باشد، یعنی به ازای هر $x \in \mathbb{R}^p$ و $t \in \mathbb{R}$ داشته باشیم، $f(tx) = |t|f(x)$.

(۴) $f(\cdot)$ در نامساوی مثلثی^{۳۸} صدق کند، یعنی به ازای هر $x, y \in \mathbb{R}^p$ داشته باشیم

$$f(x + y) \leq f(x) + f(y).$$

در این صورت از نماد $f(\cdot) = \|\cdot\|$ برای نمایش نرم استفاده می کنیم.

تعریف ۷.۱.۱ (نرم ماتریسی). فرض کنید A یک ماتریس $m \times n$ باشد، در این صورت نرم ماتریسی $\|A\|$ با خواص زیر تعریف می کنیم

^{۳۴}Convex

^{۳۵}Norm

^{۳۶}Definite

^{۳۷}Homogenous

^{۳۸}Triangle inequality

(۱) $\|A\| \geq 0$ ، $\|A\| = 0$ اگر و فقط اگر A یک ماتریس صفر باشد.

(۲) به ازای هر اسکالر α ، $\|\alpha A\| = |\alpha| \|A\|$.

(۳) $\|A + B\| \leq \|A\| + \|B\|$.

علاوه بر سه شرط فوق برای نرم ماتریسی، گاهی اوقات به شرایط زیر نیز نیاز دارند که لزوماً برای همه نرم‌های ماتریس برقرار نیست

$$\|Ax\| \leq \|A\| \cdot \|x\|,$$

$$\|AB\| \leq \|A\| \cdot \|B\|.$$

تعریف ۸.۱.۱ (نرم‌های مولفه‌ای). اگر با همی $m \times n$ مولفه‌ی ماتریس A مشابه مولفه‌های یک بردار mn بعدی رفتار شود، آن‌گاه p -نرم این بردار به صورت زیر تعریف می‌شود

$$\|A\|_p = \left\{ \sum_{i=1}^m \sum_{j=1}^n |\alpha_{ij}|^p \right\}^{\frac{1}{p}}.$$

به ویژه سه حالت $p = 1, 2, \infty$ را در نظر می‌گیریم.

(۱) $\|A\|_1$ ، $p = 1$ قدرمطلق مجموع همه عناصر ماتریس A است. به عبارتی

$$\|A\|_1 = \sum_{i=1}^m \sum_{j=1}^n |a_{ij}|.$$

(۲) $\|A\|_\infty$ ، $p = \infty$ ، نرم بیشینه، ماکزیمم مقدار مطلق در بین همه mn عنصر ماتریس A است. یعنی،

$$\|A\|_\infty = \max \{|a_{ij}| : 1 \leq i \leq m, 1 \leq j \leq n\}.$$

(۳) $\|A\|_2$ ، $p = 2$ ، نرم فروبنیوس^{۳۹} است که به صورت زیر می‌باشد

$$\|A\|_2 = \|A\|_F = \sqrt{\sum_{i=1}^m \sum_{j=1}^n |a_{ij}|^2} = \sqrt{\text{tr}(A^\top A)} = \sqrt{\sum_{i=1}^r \lambda_i} = \sqrt{\sum_{i=1}^r \sigma_i^2},$$

که در آن $r \leq \min(m, n)$ رتبه ماتریس A ، $\lambda_i = \sigma_i^2$ ، i امین مقدار ویژه نامنفی از $A^\top A$ است و $\sigma_i = \sqrt{\lambda_i}$ ، i امین مقدار تکین (ریشه مقدار ویژه) A ($i = 1, 2, \dots, r$) است.

نرم فروبنیوس ماتریس متعامد R که در شرط $R^\top = R^{-1}$ یا $R^\top R = R^{-1} R = I$ صدق کند، برابر با $\|R\|_F^2 = \text{tr}(R^\top R) = \text{tr}(I) = n$ می‌باشد.

^{۳۹}Frobenius norm

تعریف ۹.۱.۱ (عدد شرطی). نسبت بزرگترین مقدار ویژه ماتریس A به کوچکترین مقدار ویژه ماتریس A را عدد شرطی گوییم. به عبارتی

$$k(\lambda) = \left| \frac{\lambda_{max}}{\lambda_{min}} \right|.$$

تعریف ۱۰.۱.۱ (تابع زیان توان دوم). اگر بردار p مولفه‌ای $\delta = (\delta_1, \dots, \delta_p)^\top$ برآوردگری برای بردار پارامتر $\theta \in \mathbb{R}^p$ باشد، آن گاه زیان استفاده از δ در خصوص برآورد θ را با $L(\theta, \delta)$ نشان داده که برابر است با

$$L(\theta, \delta) = (\delta - \theta)^\top (\delta - \theta) = \|\delta - \theta\|^2 = \sum_{i=1}^p (\delta_i - \theta_i)^2.$$

چنانچه پارامتر مقیاس مجهول $\sigma^2 \geq 0$ در مدل موجود باشد، آن گاه از تابع زیان زیر استفاده می‌کنیم

$$L(\theta, \delta) = \frac{(\delta - \theta)^\top (\delta - \theta)}{\sigma^2} = \frac{\|\delta - \theta\|^2}{\sigma^2} = \frac{1}{\sigma^2} \sum_{i=1}^p (\delta_i - \theta_i)^2.$$

تعریف ۱۱.۱.۱ (تابع مخاطره). با استفاده از تعریف تابع زیان، تابع مخاطره δ را با $R(\theta, \delta)$ نشان داده که برابر است با

$$R(\theta, \delta) = E_\theta[L(\theta, \delta)] = \int_{\chi} L(\theta, \delta) f_\theta(x) dx,$$

که در آن $f_\theta(\cdot)$ تابع چگالی احتمال تعریف شده روی χ است.

تعریف ۱۲.۱.۱ (اثر یک ماتریس). در جبر خطی، اثر ماتریس $n \times n$ دلخواه A به صورت مجموع عناصر قطر اصلی A است. به عبارت دیگر

$$\text{tr}(A) = \sum_{i=1}^n a_{ii} = a_{11} + a_{22} + \dots + a_{nn},$$

که a_{ii} ، i امین مولفه روی قطر اصلی ماتریس A است.

ویژگی‌های اثر ماتریس

(۱) به ازای ماتریس‌های متقارن مربعی A و B و مقدار ثابت $c \in \mathbb{R}$

$$\text{tr}(A + B) = \text{tr}(A) + \text{tr}(B),$$

$$\text{tr}(cA) = c \text{tr}(A).$$

در حالت کلی، فرض کنید $k_1, \dots, k_m \in \mathbb{R}$ و $A_1, \dots, A_m$ ماتریس‌های مربعی باشند. در این صورت

$$\text{tr}\left(\sum_{i=1}^m k_i A_i\right) = \sum_{i=1}^m k_i \text{tr}(A_i).$$

(۲) به ازای هر ماتریس مربعی دلخواه A ، $\text{tr}(A) = \text{tr}(A^\top)$ ، زیرا ترانهاده کردن یک ماتریس، عناصر قطر اصلی را تغییر نمی‌دهد.

(۳) اثر حاصل ضرب ماتریس‌ها می‌تواند به صورت مجموع حاصل ضرب مولفه‌ها نوشته شود. یعنی

$$\text{tr}(A^\top B) = \text{tr}(AB^\top) = \text{tr}(B^\top A) = \text{tr}(BA^\top) = \sum_{i,j} a_{ij} b_{ij}.$$

بدین معنی است که حاصل ضرب توابع ماتریسی، مشابه حاصل ضرب داخلی بردارهاست. در حالت خاص اگر $A = (a_{ij})$ ماتریس $m \times n$ باشد، آن‌گاه

$$\text{tr}(A^\top A) = \sum_{i=1}^m \sum_{j=1}^n a_{ij}^2.$$

(۴) برای ماتریس‌های حقیقی، اثر ماتریس حاصل ضرب می‌تواند به صورت زیر نوشته شود

$$\text{tr}(A^\top B) = \sum_{i,j} (A \circ B)_{ij}.$$

که $\circ$ ضرب داخلی (ضرب مولفه به مولفه) را نشان می‌دهد.

(۵) ماتریس‌ها در اثر حاصل ضرب می‌توانند بدون تغییر جابه‌جا شوند. یعنی اگر A ماتریس $m \times n$ و B ماتریس $n \times m$ باشند، آن‌گاه

$$\text{tr}(AB) = \text{tr}(BA).$$

در حالت کلی، اثر تحت جایگشت‌های دایره‌ای^{۴۰} پایا است یعنی

$$\text{tr}(ABCD) = \text{tr}(BCDA) = \text{tr}(CDAB) = \text{tr}(DABC).$$

این ویژگی به خاصیت دایره‌ای معروف است. دقت کنید که برای هر جایگشت دلخواهی رابطه برقرار نیست. در حالت کلی

$$\text{tr}(ABC) \neq \text{tr}(ACB).$$

با این وجود اگر حاصل ضرب ماتریس‌های متقارن در نظر گرفته شود، هر جایگشتی مجاز است زیرا

$$\begin{aligned} \text{tr}(ABC) &= \text{tr}(A^\top B^\top C^\top) = \text{tr}(A^\top (CB)^\top) = \text{tr}((CB)^\top A^\top) \\ &= \text{tr}((ABC)^\top) \end{aligned}$$

تساوی آخر از ویژگی (۲) نتیجه شده است.

^{۴۰}Cyclic Permutations

(۶) حاصلضرب ماتریس‌ها با حاصلضرب اثر ماتریس‌ها برابر نیست به عبارت دیگر

$$\text{tr}(AB) \neq \text{tr}(A) \text{tr}(B).$$

اما اثر حاصلضرب کرونکر^{۴۱} دو ماتریس، حاصلضرب اثر آنهاست یعنی

$$\text{tr}(A \otimes B) = \text{tr}(A) \text{tr}(B).$$

(۷) اگر A متقارن و B نامتقارن باشد، آن گاه $\text{tr}(AB) = 0$.

(۸) اثر ماتریس همانی، بعد فضای ماتریس است، این ویژگی منجر به تعمیم بعد با استفاده از اثر می‌شود.

(۹) اثر یک ماتریس خودتوان A (یعنی $A^2 = A$) برابر رتبه ماتریس A است. برای جزئیات بیشتر، لانگ^{۴۲} (۱۹۷۸) را ببینید.

تعریف ۱۳.۱.۱ (برآوردگر انقباضی). در آمار برآوردگر انقباضی، برآوردگری است که به طور واضح یا مجازی در اثر انقباض به دست می‌آید. به عبارت دیگر، برآوردگر انقباضی برآوردگری است که از ترکیب اطلاعات نمونه‌ای یعنی برآوردگر خام یا اولیه مثلاً MLE، برآوردگر کمترین توان‌های دوم، بیزی و غیره با اطلاعات غیر نمونه‌ای که معمولاً در قالب یک یا چند محدودیت به مدل اضافه می‌شوند، به صورت برآوردگر بهبود یافته جدید به دست می‌آید. این بهبود در جهت بهبود برآوردگر اولیه است. در حالت کلی بهبود برآوردگرها را می‌توان با معیار مخاطره یا توان دوم خطا^{۴۳} (MSE) اندازه‌گیری کرد. در این صورت انقباض می‌تواند به سمت صفر یا یک مقدار ثابت باشد. اثر این انقباض معمولاً به صورت تبدیل یک برآوردگر ناریب به اریب ولی با MSE کمتر می‌باشد.

در عین حال، می‌توان برآوردگرهایی انقباضی از نوع استاین و ناریب را نیز یافت. برای آگاهی بیشتر در این خصوص به نوروزی راد (۱۳۹۰) مراجعه کنید. شکل کلی برآوردگر انقباضی به صورت $X + g(X)$ است که در آن g یک تابع اندازه‌پذیر می‌باشد. پس از حل نامعادله

$$\forall \theta \in \Theta : R(\theta, X + g(X)) \leq R(\theta, X),$$

به کمک لم استاین، تابع $g(\cdot)$ را در هر یک از حالت‌های زیر به دست می‌آوریم

$$(۱) \text{ برآوردگر انقباضی به سمت صفر، } (1 - c)X = X - cX,$$

^{۴۱} Kronecher product

^{۴۲} Lang

^{۴۳} Mean squared error

(۲) برآوردگر انقباضی به سمت بردار ثابت m

$$m + c(X - m) = (1 - c)m + cX,$$

که در آن c ثابت انقباضی است و بسته به نوع توزیع و تابع زیان، مقادیر متفاوتی می پذیرد.

تعریف ۱۴.۱.۱ (سازگاری). فرض کنید $X_1, \dots, X_n$ و دنباله‌ای از متغیرهای تصادفی مستقل با توزیعی از خانواده توزیع‌های $\{F_\theta : \theta \in \Theta\}$ باشد. همچنین فرض کنید T_n یک برآوردگر $\gamma(\theta)$ بر اساس n مشاهده اول باشد. هدف از سازگاری T_n بررسی رفتار حدی دنباله $\{T_n\}_{n=1}^\infty$ از برآوردگرهاست. گوییم دنباله برآوردگرهای T_n برای $\gamma(\theta)$ سازگار است اگر چنانچه n به سمت ∞ میل کند، آن گاه

$$T_n \xrightarrow{P} \gamma(\theta),$$

یعنی برای هر $\epsilon > 0$ و $\gamma > 0$ عدد صحیح مثبتی مانند $n_0 = n_0(\epsilon, \delta)$ وجود دارد به طوری که برای هر $n > n_0$,

$$P_\theta(|T_n - \gamma(\theta)| > \epsilon) < \delta,$$

یا برای هر $\epsilon > 0$,

$$\lim_{n \rightarrow \infty} P_\theta(|T_n - \gamma(\theta)| > \epsilon) = 0.$$

هر دنباله از برآوردگرها که در شرط فوق صدق نکند ناسازگار می نامیم.

اندازه‌گیری‌های حاصل روی یک مشاهده، به صورت برداری به شکل زیر است

$$\mathbf{x}_i = \begin{pmatrix} x_{i1} \\ x_{i2} \\ \vdots \\ x_{ip} \end{pmatrix},$$

که x_i ($i = 1, \dots, n$)، i امین مشاهده از نمونه تصادفی را ارائه می دهد. مجموعه تمام اندازه‌گیری‌های حاصل از مشاهدات، ماتریس زیر را تشکیل می دهد

$$\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_n),$$

که $\mathbf{X}$ ماتریسی $p \times n$ است.

لم ۱۵.۱.۱ (شفر و استریمر، ۲۰۰۵). فرض کنید $x_{n1}, \dots, x_{ni}$ نمونه‌ای تصادفی از متغیر تصادفی X_i باشد. همچنین فرض کنید k ، k امین مشاهده X_i و x_{ki} میانگین نمونه باشد. همچنین فرض کنید $S_{ii}, S_{ij} = \frac{1}{n-1} \sum_{k=1}^n (x_{ki} - \bar{x}_i)(x_{kj} - \bar{x}_j)$ عناصر قطری ماتریس کوواریانس نمونه، $\bar{w}_{ij} = \frac{1}{n} \sum_{k=1}^n w_{kij}$ و $w_{kij} = (x_{ki} - \bar{x}_i)(x_{kj} - \bar{x}_j)$ باشد. آن گاه کوواریانس تجربی به صورت زیر است

$$\widehat{Cov}(X_i, X_j) = S_{ij} = \frac{n}{n-1} \bar{w}_{ij},$$

و واریانس نیز به صورت

$$\widehat{Var}(X_i) = S_{ii} = \frac{n}{n-1} \bar{w}_{ii},$$

می باشد. همچنین به صورت مشابه

$$\begin{aligned} \widehat{Var}(S_{ij}) &= \frac{n^2}{(n-1)^2} \widehat{Var}(\bar{w}_{ij}) \\ &= \frac{n}{(n-1)^2} \sum_{k=1}^n (w_{kij} - \bar{w}_{ij})^2. \end{aligned}$$

۲.۱ برآورد

تحلیل چندمتغیره بخشی از آمار است که مربوط به تحلیل داده ها شامل یک یا بیشتر از یک اندازه گیری (متغیر، صفت) در افراد یا اشیا می باشد. نتایج حاصل از اندازه گیری را بعد نامیده و با p نشان می دهیم. نمونه n تایی را از جمعیتی با متغیرهای p بعدی در نظر بگیرید. اگر فرض کنیم توزیع بردار x_i متعلق به خانواده توزیع های احتمال p بعدی زیر باشد

$$\mathcal{F} = \{f(x; \mu, \Sigma) \mid \mu \in R^p, \Sigma > 0\}, \quad (2.1)$$

آن گاه بردار میانگین μ و ماتریس کوواریانس Σ به صورت زیر می باشند

$$\mu = \begin{pmatrix} EX_1 \\ EX_2 \\ \vdots \\ EX_p \end{pmatrix} = \begin{pmatrix} \mu_1 \\ \mu_2 \\ \vdots \\ \mu_p \end{pmatrix},$$

$$\Sigma = \begin{pmatrix} cov(X_1, X_1) & cov(X_1, X_2) & \dots & cov(X_1, X_p) \\ cov(X_2, X_1) & cov(X_2, X_2) & \dots & cov(X_2, X_p) \\ \vdots & \vdots & \ddots & \vdots \\ cov(X_p, X_1) & cov(X_p, X_2) & \dots & cov(X_p, X_p) \end{pmatrix} = \begin{pmatrix} \sigma_{11} & \sigma_{12} & \dots & \sigma_{1p} \\ \sigma_{21} & \sigma_{22} & \dots & \sigma_{2p} \\ \vdots & \vdots & \ddots & \vdots \\ \sigma_{p1} & \sigma_{p2} & \dots & \sigma_{pp} \end{pmatrix},$$

به طوری که $\mu_i = E[X_i]$ و $\Sigma_{ij} = E[(X_i - \mu_i)(X_j - \mu_j)]$ برای $i \neq j$. به طور کلی در خانواده توزیع های (۲.۱) علاقمند به برآورد دو پارامتر بردار میانگین μ و ماتریس کوواریانس Σ می باشیم. برآورد معمول برای μ ، میانگین نمونه $\bar{x}$ و برای Σ ، ماتریس کوواریانس نمونه S می باشد که بر اساس یک اندازه نمونه n تایی به صورت زیر می باشند

$$\bar{x} = \begin{pmatrix} \bar{x}_1 \\ \vdots \\ \bar{x}_p \end{pmatrix},$$

که در آن

$$\bar{x}_i = \frac{1}{n} \sum_{j=1}^n x_{ij},$$

میانگین i امین نمونه است. ماتریس کوواریانس نمونه نیز به صورت زیر قابل محاسبه است

$$S = \frac{1}{n-1} \sum_{j=1}^n (x_j - \bar{x})(x_j - \bar{x})^\top, \quad (3.1)$$

که S را می توان به صورت زیر نوشت

$$S = \frac{1}{n-1} (X - \bar{X})(X - \bar{X})^\top,$$

که $\bar{X}$ ماتریس $p \times n$ با هر ستون متشکل از $\bar{x}$ است.

براساس یافته های فیشر و سان^{۴۴} (۲۰۱۱)، $\bar{x}$ و S به دلیل ناریبی و سازگاری برآوردگرهای مناسبی برای μ و Σ هستند.

قضیه ۱.۲.۱ (میرهد، ۱۹۸۲). اگر $x_1, \dots, x_n$ بردارهای تصادفی مستقل دارای توزیع $\mathcal{N}_p(\mu, \Sigma)$ باشند، به طوری که $n > p$ ، آن گاه برآوردگر درست نمایی ماکزیمم μ و Σ به صورت

$$\bar{x} = \frac{1}{N} \sum_{i=1}^n x_i,$$

و $\hat{\Sigma} = \frac{n}{n+1} S$ می باشند که S به صورت (۳.۱) تعریف می شود.

۳.۱ برآوردگر انقباضی نوع استاین

در این بخش به دنبال برآوردگری مناسب برای ماتریس کوواریانس زمانی که حجم نمونه کمتر از تعداد متغیرها ($p > n$) است، می باشیم. رویکردی رایج در بهبود برآورد ماتریس کوواریانس در این حالت، «انقباضی» یا «برآورد اریب» است که به طور گسترده توسط افرون و موریس (۱۹۷۵) بررسی شده است.

فرض کنید S برآوردگر ارائه شده در (۳.۱) است که میانگین توان های دوم خطای آن به صورت زیر محاسبه می شود

$$MSE(S) = (Bias(S))^\top (Bias(S)) + Cov(S),$$

که در آن $Bias(S) = S - E(S)$. برای ماتریس کوواریانس نمونه S اریبی صفر است (فیشر و سان، ۲۰۰۹)، از این رو بهبود دقت آن در گرو کاهش واریانس است. با تعریف برآوردگری اریب، می توان واریانس را کاهش داد، که ایده اصلی برآوردگر جیمز و استاین (۱۹۶۱) است.

^{۴۴} Sun and Fisher

در واقع در این روش ترکیبی محدب از ماتریس کوواریانس نمونه با ماتریسی هدف به صورت زیر در نظر گرفته می شود

$$S^* = (1 - \lambda)S + \lambda T, \quad (4.1)$$

که $\lambda \in (0, 1)$ پارامتر (ضریب) انقباضی^{۴۵} نام دارد و T ماتریس هدف^{۴۶} است که باید معین مثبت باشد. این ترکیب را طوری در نظر می گیریم که S^* واریانس کمتری نسبت به S داشته باشد. با توجه به این که ماتریس هدف T از قبل تعیین شده است، مساله اصلی در ارتباط با برآوردگر انقباضی، تعیین پارامتر انقباضی λ است. پارامتر انقباضی λ طوری انتخاب می شود که برآوردگر S را بهبود بخشد. اگر n در مقایسه با p بزرگ باشد، S باید واریانس کمتری داشته باشد و برآوردگر خوبی باشد، از این رو $\lambda \rightarrow 0$. اگر p بزرگ باشد، S واریانس بزرگتری دارد و باید وزنی را برای ماتریس هدف فراهم کرد به طوری که $\lambda \rightarrow 1$. انتخاب پارامتر انقباضی (λ) را می توان با روش های بوت استرپ^{۴۷}، اعتبارسنجی متقابل^{۴۸} و سایر روش ها تعیین کرد. زمانی که بعد بالاست، مقادیر ویژه ماتریس کوواریانس برآوردگرهای ضعیفی برای مقادیر ویژه واقعی هستند. طبق آن، کوچکترین مقادیر ویژه به سمت صفر میل می کند، همچنان که بزرگترین آن به سمت بی نهایت میل می کند. ایده نظری، انقباض مقادیر ویژه S به مقادیر ویژه T است.

ماتریس هدف انتخاب شده معین مثبت (معکوس پذیر) بوده و باید ماتریس کوواریانس واقعی را برای کاربرد ما ارائه دهد. لازم به ذکر است که این شرط دشوار بوده و باید داده ها در دسترس باشند. برآوردگر انقباضی چندین امتیاز نسبت به ماتریس کوواریانس نمونه دارد. برآوردگر انقباضی S^* برای هر بعدی، ماتریس معین مثبت و نامنفرد (معکوس پذیر) است. لدویت و ولف^{۴۹} (۲۰۰۳) قضیه ای را ارائه دادند که بیانگر پارامتر انقباضی مطلوب از نظر واریانس-کوواریانس در ماتریس کوواریانس نمونه و ماتریس هدف است. تحت میانگین توان های دوم خطا و با در نظر گرفتن نرم فربنیوس، یک مقدار بهینه برای پارامتر انقباضی وجود دارد. لدویت و ولف (۲۰۰۴) جزئیات پارامتر بهینه را برای ماتریس هدف $T = \mu_\lambda I$ ، در نظر گرفتند که μ_λ میانگین مقادیر ویژه Σ و I ماتریس همانی است. پارامتر انقباضی بهینه در نظر گرفته شده به صورت زیر است

$$\lambda = \frac{\beta^2}{\delta^2}$$

که در آن

$$\beta^2 = E[\|S - \Sigma\|^2],$$

^{۴۵}Shrinkage intensity

^{۴۶}Target matrix

^{۴۷}Bootstrap

^{۴۸}Cross-validation

^{۴۹}Ledoit and Wolf

$$\delta^2 = E[\|S - \mu_\lambda I\|^2].$$

متاسفانه μ_λ ، β^2 و δ^2 همه به اطلاعاتی در مورد Σ نیاز دارند و از آن جایی که مجهول هستند، باید برآورد شوند.

لدویت و ولف (۲۰۰۴) برآوردهای سازگاری برای μ_λ و λ فراهم کردند. چن و همکاران^{۵۰} (۲۰۱۰) با در نظر گرفتن قضیه راثو-بلکول، برآوردگر لدویت و ولف (۲۰۰۴) را بهبود بخشیدند. شفر و استریمر (۲۰۰۵) شاخص ناریبی را نسبت به سازگاری در برآوردهای μ_λ و λ در نظر گرفتند.

۴.۱ انواع پارامترهای انقباضی

فرض کنید $\{x_1, \dots, x_n\}$ نمونه‌ای تصادفی از توزیع p متغیره نرمال با بردارهای میانگین صفر و کوواریانس Σ باشد. نکته‌ای که باید بدان توجه نمود این است که فرض $n \geq p$ را نداریم. به دنبال برآوردگری مانند $\hat{\Sigma}$ هستیم که تابع زیان زیر را کمینه کند

$$E \left[\left\| \hat{\Sigma}(\{x_i\}_{i=1}^n) - \Sigma \right\|_F^2 \right], \quad (5.1)$$

که در آن $\hat{\Sigma}$ برآوردگری برای Σ می‌باشد. در واقع باید پارامتر انقباضی λ را (از روی مشاهدات) طوری یافت که تابع زیان (۵.۱) را کمینه کند. حال به معرفی پارامترهای انقباضی می‌پردازیم.

۱.۴.۱ پارامتر انقباضی LW

لدویت و ولف (۲۰۰۴) تحت فرض این که x_i ها دارای توزیع نرمال p متغیره با میانگین صفر و ماتریس کوواریانس Σ هستند، پارامتر انقباضی λ را به صورت زیر معرفی کردند

$$\lambda_{LW} = \min \left\{ \frac{\sum_{i=1}^n \|x_i x_i^\top - \hat{S}\|_F^2}{n^2 \text{tr}(\hat{S}^2) - \frac{1}{p} \text{tr}^2(\hat{S})}, 1 \right\}, \quad (6.1)$$

که در آن

$$\hat{S} = \frac{1}{n} \sum_{i=1}^n x_i x_i^\top.$$

^{۵۰}Chen et al.

۲.۴.۱ پارامتر انقباضی RBLW

چن و همکاران (۲۰۱۰) با در نظر گرفتن قضیه راثو بلکول، برآوردگر لدویت و ولف (۲۰۰۴) را بهبود دادند. آن‌ها تحت فرضیه‌ی نرمال بودن داده‌ها و همچنین با در نظر گرفتن S به عنوان آماره بسنده، پارامتر انقباضی را به صورت زیر تعریف کردند

$$\lambda_{\text{RBLW}} = \min \left\{ \frac{\frac{n-2}{n} \text{tr}(\hat{S}^2) + \text{tr}^2(\hat{S})}{(n+2)[\text{tr}(\hat{S}^2) - \frac{1}{p} \text{tr}^2(\hat{S})]}, 1 \right\}. \quad (7.1)$$

برآوردگر انقباضی با استفاده از این برآوردگر برای پارامتر بهینه بر برآوردگر لدویت و ولف ارجحیت دارد. آن‌ها برای داده‌های سری زمانی و شبیه‌سازی‌های خود از این رویکرد استفاده کردند و در مثال‌های خود تاثیر برآوردگرها را مشاهده کردند.

۳.۴.۱ پارامتر انقباضی SS

شفر و استریم (۲۰۰۵) نیز تحت فرض نرمال بودن داده‌ها، پارامتر انقباضی را به صورت زیر ارائه دادند

$$\lambda_{\text{SS}} = \min \left\{ \frac{\sum_{i=1}^n \sum_{j=1, j \neq i}^n \widehat{\text{Var}}(S_{ij}) + \sum_{i=1}^n \widehat{\text{Var}}(S_{ii})}{\sum_{i=1}^n \sum_{j=1, j \neq i}^n S_{ij}^2 + \sum_{i=1}^n (S_{ii} - \mu_\lambda)^2}, 1 \right\}, \quad (8.1)$$

که در آن S_{ii} عناصر قطری ماتریس کوواریانس نمونه S و μ_λ میانگین عناصر قطری S_{ii} است. برای جزئیات بیشتر به شفر و استریم (۲۰۰۵) مراجعه کنید.

۴.۴.۱ پارامتر انقباضی FS

فیشر و سان (۲۰۱۱) پارامتر انقباضی زیر را پیشنهاد دادند

$$\lambda_{\text{FS}} = \min \left\{ \frac{\frac{1}{n} \hat{a}_1 + \frac{p}{n} \hat{a}_1^2}{\frac{n+1}{n} \hat{a}_2 + \frac{p-a}{n} \hat{a}_1}, 1 \right\},$$

که در آن $a_i = \frac{1}{p} \text{tr}(\Sigma^i)$. در این صورت

$$\begin{aligned} \hat{a}_1 &= \frac{\text{tr}(S)}{p}, \\ \hat{a}_2 &= \frac{n^2}{(n+1)(n+2)p} \left[\text{tr}(S^2) - \frac{\text{tr}^2(S)}{n} \right], \end{aligned}$$

که برآوردگرهایی سازگار برای a_2 و a_1 هستند.

فصل ۲

برآورد توابع ماتریس کوواریانس در بعد بالا

۱.۲ مقدمه

برآورد ماتریس کوواریانس بزرگ موضوعی مهم در بطن تحلیل داده‌های موضوعات مدرن و پیچیده مانند ژنومیک، تحقیقات سرطان، مطالعات بالینی، شناسایی الگو و هندسه محدب محاسباتی است. معمولاً در این مسائل هدف برآورد ماتریس کوواریانس Σ بر پایه یک نمونه n تایی از بردارهای تصادفی p متغیره مستقل و هم‌توزیع $x_1, \dots, x_n$ با بردار میانگین μ در حالت n کوچک و p بزرگ است. در این موارد، معمولاً تعداد مشاهدات n در مقایسه با تعداد پارامترها بسیار کوچکتر است.

فرض کنید $x_1, \dots, x_n$ بردارهای تصادفی p متغیره مستقل باشند که هر یک از x_i ها می‌تواند به صورت زیر بیان شود

$$x_i = \mu + \Sigma^{\frac{1}{2}} z_i, \quad (1.2)$$

که در آن z_i دارای توزیع دلخواه $\mathcal{F}$ با بردار میانگین صفر و ماتریس کوواریانس Σ می‌باشد، که Σ ماتریس معین مثبت $p \times p$ است.

در آزمون‌های آماری مربوط به تحلیل‌های چندمتغیره برای داده‌های بعد بالا، اغلب به برآورد $\text{tr}(\Sigma^2)$ نیاز داریم (اسکات^۱ (۲۰۰۷)؛ سریواستاوا و فوجیکوشی^۲ (۲۰۰۶)؛ چن و کین (۲۰۱۰)؛ فوجیکوشی و همکاران (۲۰۰۴)).

سریواستاوا (۲۰۰۵) برآوردگر ناریب و سازگار $a_2 = \text{tr}(\Sigma^2)/p$ را با در نظر گرفتن فرض نرمال در قالب قضیه زیر پیشنهاد کرد.

^۱Schott

^۲Srivastava and Fujikoshi

قضیه ۱.۱.۲ (سریواستوا، ۲۰۰۵). برآوردگر ناریب و سازگار $a_2 = \text{tr}(\Sigma^2)/p$ ، تحت فرض نرمال بودن $\mathcal{F}$ ، به صورت

$$\tilde{a}_2 = \frac{n^2}{(n-1)(n+2)} \frac{1}{p} \left[\text{tr}(\mathcal{S}^2) - \frac{1}{n} (\text{tr}(\mathcal{S}))^2 \right],$$

است.

پس از آن، سریواستوا و همکاران (۲۰۱۱) سازگاری برآوردگر $\tilde{a}_2$ را حتی زمانی که $\mathcal{F}$ نرمال نباشد بررسی کردند. نکته‌ای که باید به آن توجه کرد این است که برآوردگر آنها ناریب نیست. چن و کین (۲۰۱۰) و چن و همکاران (۲۰۱۰) نیز برآوردگر سازگار دیگری برای a_2 پیشنهاد کردند. برآوردگر پیشنهادی چن و کین (۲۰۱۰)، ناریب نیست و همین‌طور برآوردگر پیشنهادی آنها نیز صورت پیچیده‌ای دارد. علاوه بر این، سازگاری این برآوردگرها تحت شرایط گشتاوری قویتری نیز به دست آمده است. بنابراین برآوردگرهای سازگار ناریب برای $(\text{tr}(\Sigma))^2$ و κ_{11} علاوه بر $\text{tr}(\Sigma^2)$ نیز پیشنهاد شده است که در آن κ_{ij} به صورت

$$\kappa_{ij} = E[\mathbf{z}^\top \Sigma^i \mathbf{z} \mathbf{z}^\top \Sigma^j \mathbf{z}] - 2 \text{tr}(\Sigma^{i+j}) - \text{tr}(\Sigma^i) \text{tr}(\Sigma^j), \quad (2.2)$$

تعریف می‌شود که $\mathbf{z} \sim \mathcal{F}$. برای توزیع نرمال، $\kappa_{ij} = 0$ (مگنوس و نیودکر، ۱۹۷۹). معیار κ_{11} یکی از مهمترین پارامترهایی است که تفاوت حالت غیرنرمال را نشان می‌دهد. مطالب این فصل برگرفته از مرجع هیمنو و یامادا^۳ (۲۰۱۴) است.

۲.۲ ناریبی و سازگاری برآوردگرها

ابتدا شرایط نظم زیر را در نظر بگیرید. وقتی $p \rightarrow \infty$ ، شرایط زیر برقرار هستند:

$$A_1 : \frac{n}{p} \rightarrow c_0 \in (0, \infty),$$

$$A_2 : a_i = \frac{\text{tr} \Sigma^i}{p} \rightarrow a_{i0} \in (0, \infty) \quad i = 1, \dots, 4.$$

در برآورد $\text{tr}(\Sigma^2)$ اغلب از $\text{tr}(\mathcal{S}^2)$ و $(\text{tr}(\mathcal{S}))^2$ استفاده می‌کنیم، که امید ریاضی این آماره‌ها در قضیه زیر بیان شده است.

قضیه ۱.۲.۲. تحت مفروضات بخش ۱.۲، بر اساس نمونه تصادفی $x_1, \dots, x_n$ ، امید ریاضی $\text{tr}(\mathcal{S}^2)$ و $(\text{tr}(\mathcal{S}))^2$ به صورت زیر به دست می‌آیند

$$E[\text{tr}(\mathcal{S}^2)] = \frac{1}{n} \kappa_{11} + \frac{n}{n-1} \text{tr}(\Sigma^2) + \frac{1}{n-1} (\text{tr}(\Sigma))^2, \quad (3.2)$$

$$E[(\text{tr}(\mathcal{S}))^2] = \frac{1}{n} \kappa_{11} + \frac{2}{n-1} \text{tr}(\Sigma^2) + (\text{tr}(\Sigma))^2. \quad (4.2)$$

برهان. می‌توان ماتریس S را به صورت زیر بازنویسی کرد

$$\begin{aligned} S &= \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})(x_i - \bar{x})^\top \\ &= \frac{1}{n-1} \sum_{i=1}^n (x_i - \mu + \mu - \bar{x})(x_i - \mu + \mu - \bar{x})^\top. \end{aligned} \quad (5.2)$$

فرض کنید به ازای هر $i = 1, \dots, n$ ، $y_i = x_i - \mu$ در این صورت

$$\begin{aligned} x_i - \mu + \mu - \bar{x} &= y_i + \left(\mu - \frac{1}{n} \sum_{i=1}^n x_i\right) \\ &= y_i + \left(\frac{1}{n} \sum_{i=1}^n \mu - \frac{1}{n} \sum_{i=1}^n x_i\right) \\ &= y_i + \frac{1}{n} \sum_{i=1}^n (\mu - x_i) \\ &= y_i - \frac{1}{n} \sum_{j=1}^n y_j \\ &= y_i - \frac{1}{n} y_i - \frac{1}{n} \sum_{\substack{j=1 \\ j \neq i}}^n y_j \\ &= \frac{n-1}{n} y_i - \frac{1}{n} \sum_{\substack{j=1 \\ j \neq i}}^n y_j \end{aligned} \quad (6.2)$$

بنابراین با جای گذاری (6.2) در (5.2) نتیجه می‌شود

$$\begin{aligned} S &= \frac{1}{n-1} \sum_{i=1}^n \left(\frac{n-1}{n} y_i - \frac{1}{n} \sum_{\substack{j=1 \\ j \neq i}}^n y_j\right) \left(\frac{n-1}{n} y_i - \frac{1}{n} \sum_{\substack{j=1 \\ j \neq i}}^n y_j\right)^\top \\ &= \frac{1}{n-1} \sum_{i=1}^n \frac{(n-1)^2}{n^2} y_i y_i^\top - \frac{1}{n-1} \sum_{i=1}^n \frac{(n-1)}{n^2} y_i \left(\sum_{\substack{j=1 \\ j \neq i}}^n y_j\right)^\top \\ &\quad - \frac{1}{n-1} \sum_{i=1}^n \frac{(n-1)}{n^2} \left(\sum_{\substack{j=1 \\ j \neq i}}^n y_j\right) y_i^\top + \frac{1}{n-1} \frac{1}{n^2} \sum_{i=1}^n \left(\sum_{\substack{j=1 \\ j \neq i}}^n y_j\right) \left(\sum_{\substack{j=1 \\ j \neq i}}^n y_j\right)^\top \\ &= \frac{n-1}{n^2} \sum_{i=1}^n y_i y_i^\top - \frac{1}{n^2} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n y_i y_j^\top - \frac{1}{n^2} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n y_j y_i^\top \\ &\quad + \frac{1}{n^2(n-1)} \sum_{i=1}^n \left(\sum_{\substack{j=1 \\ j \neq i}}^n y_j\right) \left(\sum_{\substack{j=1 \\ j \neq i}}^n y_j\right)^\top \\ &= \frac{n-1}{n^2} \sum_{i=1}^n y_i y_i^\top - \frac{2}{n^2} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n y_i y_j^\top \\ &\quad + \frac{1}{n^2(n-1)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n y_j y_j^\top + \frac{1}{n^2(n-1)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq i \\ k \neq j}}^n y_j y_k^\top \end{aligned}$$

$$= \frac{1}{n} \sum_{i=1}^n \mathbf{y}_i \mathbf{y}_i^\top - \frac{1}{n^2} \sum_{i=1}^n \mathbf{y}_i \mathbf{y}_i^\top - \frac{2}{n^2} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \mathbf{y}_i \mathbf{y}_j^\top \\ + \frac{1}{n^2(n-1)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \mathbf{y}_j \mathbf{y}_j^\top + \frac{1}{n^2(n-1)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq i \\ k \neq j}}^n \mathbf{y}_j \mathbf{y}_k^\top.$$

می‌دانیم

$$\sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \mathbf{y}_j \mathbf{y}_j^\top = \sum_{i=1}^n (n-1) \mathbf{y}_i \mathbf{y}_i^\top = (n-1) \sum_{i=1}^n \mathbf{y}_i \mathbf{y}_i^\top$$

و

$$\sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq i \\ k \neq j}}^n \mathbf{y}_j \mathbf{y}_k^\top = \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n (n-2) \mathbf{y}_i \mathbf{y}_j^\top = (n-2) \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \mathbf{y}_i \mathbf{y}_j^\top.$$

در این صورت

$$\begin{aligned} S &= \frac{1}{n} \sum_{i=1}^n \mathbf{y}_i \mathbf{y}_i^\top - \frac{2}{n^2} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \mathbf{y}_i \mathbf{y}_j^\top \\ &\quad + \frac{(n-2)}{n^2(n-1)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \mathbf{y}_i \mathbf{y}_j^\top \\ &= \frac{1}{n} \sum_{i=1}^n \mathbf{y}_i \mathbf{y}_i^\top + \left(\frac{n-2}{n^2(n-1)} - \frac{2}{n^2} \right) \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \mathbf{y}_i \mathbf{y}_j^\top \\ &= \frac{1}{n} \sum_{i=1}^n \mathbf{y}_i \mathbf{y}_i^\top - \frac{1}{n(n-1)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \mathbf{y}_i \mathbf{y}_j^\top. \end{aligned}$$

از طرفی با توجه به مدل (۱.۲)

$$\mathbf{y}_i = \Sigma^{\frac{1}{2}} \mathbf{z}_i$$

و می‌توان نوشت

$$S = \frac{1}{n} \sum_{i=1}^n \Sigma^{\frac{1}{2}} \mathbf{z}_i \mathbf{z}_i^\top \Sigma^{\frac{1}{2}} - \frac{1}{n(n-1)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \Sigma^{\frac{1}{2}} \mathbf{z}_i \mathbf{z}_j^\top \Sigma^{\frac{1}{2}}.$$

از عبارت بالا می‌توانیم $\text{tr}(S^2)$ و $(\text{tr}(S))^2$ را به صورت زیر بیان کنیم

$$\text{tr}(S^2) = \frac{1}{n^2} \sum_{i=1}^n (\mathbf{z}_i^\top \Sigma \mathbf{z}_i)^2 + \frac{1}{n^2} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n (\mathbf{z}_i^\top \Sigma \mathbf{z}_j)^2$$

$$\begin{aligned}
 & + \frac{1}{n^{\gamma}(n-1)^{\gamma}} \sum_{i=1}^n \sum_{j=1}^n \sum_{\substack{k=1 \\ k \neq j}}^n \sum_{\substack{l=1 \\ l \neq i}}^n z_i^{\top} \Sigma z_j z_k^{\top} \Sigma z_l \\
 & - \frac{2}{n^{\gamma}(n-1)^{\gamma}} \sum_{i=1}^n \sum_{j=1}^n \sum_{\substack{k=1 \\ k \neq i}}^n z_i^{\top} \Sigma z_j z_j^{\top} \Sigma z_k \\
 = & \frac{1}{n^{\gamma}} \sum_{i=1}^n (z_i^{\top} \Sigma z_i)^{\gamma} + \frac{n^{\gamma} - 2n + 2}{n^{\gamma}(n-1)^{\gamma}} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n (z_i^{\top} \Sigma z_j)^{\gamma} \\
 & + \frac{1}{n^{\gamma}(n-1)^{\gamma}} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n z_i^{\top} \Sigma z_i z_j^{\top} \Sigma z_j \\
 & - \frac{4}{n^{\gamma}(n-1)^{\gamma}} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n z_i^{\top} \Sigma z_i z_i^{\top} \Sigma z_j \\
 & + \frac{2}{n^{\gamma}(n-1)^{\gamma}} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq i \\ k \neq j}}^n z_i^{\top} \Sigma z_i z_j^{\top} \Sigma z_k \\
 & - \frac{2n - 4}{n^{\gamma}(n-1)^{\gamma}} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq i \\ k \neq j}}^n z_i^{\top} \Sigma z_j z_i^{\top} \Sigma z_k \\
 & + \frac{1}{n^{\gamma}(n-1)^{\gamma}} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq i \\ k \neq j}}^n \sum_{\substack{l=1 \\ l \neq i \\ l \neq j \\ l \neq k}}^n z_i^{\top} \Sigma z_j z_k^{\top} \Sigma z_l
 \end{aligned}$$

$$\begin{aligned}
 (\text{tr}(S))^{\gamma} &= \frac{1}{n^{\gamma}} \sum_{i=1}^n (z_i^{\top} \Sigma z_i)^{\gamma} + \frac{1}{n^{\gamma}} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n z_i^{\top} \Sigma z_i z_j^{\top} \Sigma z_j \\
 & + \frac{1}{n^{\gamma}(n-1)^{\gamma}} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq i \\ k \neq j}}^n \sum_{\substack{l=1 \\ l \neq i \\ l \neq j \\ l \neq k}}^n z_i^{\top} \Sigma z_j z_k^{\top} \Sigma z_l \\
 & - \frac{2}{n^{\gamma}(n-1)^{\gamma}} \sum_{i=1}^n \sum_{j=1}^n \sum_{\substack{k=1 \\ k \neq j}}^n z_i^{\top} \Sigma z_i z_j^{\top} \Sigma z_k \\
 = & \frac{1}{n^{\gamma}} \sum_{i=1}^n (z_i^{\top} \Sigma z_i)^{\gamma} + \frac{2}{n^{\gamma}(n-1)^{\gamma}} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n (z_i^{\top} \Sigma z_j)^{\gamma} \\
 & + \frac{1}{n^{\gamma}} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n z_i^{\top} \Sigma z_i z_j^{\top} \Sigma z_j \\
 & - \frac{4}{n^{\gamma}(n-1)^{\gamma}} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n z_i^{\top} \Sigma z_i z_i^{\top} \Sigma z_j
 \end{aligned}$$

$$\begin{aligned}
 & -\frac{2}{n^2(n-1)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq j \\ k \neq i}}^n z_i^\top \Sigma z_i z_j^\top \Sigma z_k \\
 & + \frac{4}{n^2(n-1)^2} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq j \\ k \neq i}}^n z_i^\top \Sigma z_i z_i^\top \Sigma z_k \\
 & + \frac{1}{n^2(n-1)^2} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq j \\ k \neq i}}^n \sum_{\substack{l=1 \\ l \neq k \\ l \neq j \\ l \neq i}}^n z_i^\top \Sigma z_j z_k^\top \Sigma z_l
 \end{aligned}$$

□

از آنجایی که سه پارامتر مجهول وجود دارند، برآورد $\text{tr}(\Sigma^2)$ تنها با استفاده از $\text{tr}(S^2)$ و $(\text{tr}(S))^2$ غیرممکن است. بنابراین از آماره زیر استفاده می‌کنیم

$$Q = \frac{1}{n-1} \sum_{i=1}^n \left((x_i - \bar{x})^\top (x_i - \bar{x}) \right)^2 \quad (7.2)$$

قضیه ۲.۲.۲. برای مدل (۱.۲) برآوردهای نااریب $\text{tr}(\Sigma^2)$ و κ_{11} به صورت زیر به دست می‌آیند

$$\begin{aligned}
 \widehat{\text{tr}(\Sigma^2)} &= \frac{n-1}{n(n-2)(n-3)} \{ (n-1)(n-2) \text{tr}(S^2) + (\text{tr}(S))^2 - nQ \} \\
 (\widehat{\text{tr}(\Sigma)})^2 &= \frac{n-1}{n(n-2)(n-3)} \{ 2 \text{tr}(S^2) + (n^2 - 3n + 1)(\text{tr}(S))^2 - nQ \} \\
 \hat{\kappa}_{11} &= \frac{-1}{(n-2)(n-3)} \{ 2(n-1)^2 \text{tr}(S^2) + (n-1)^2 (\text{tr}(S))^2 - n(n+1)Q \}
 \end{aligned}$$

برهان. می‌توان Q را به صورت زیر نوشت

$$\begin{aligned}
 Q &= \frac{1}{n-1} \sum_{i=1}^n \{ (y_i^\top y_i)^2 + 4(y_i^\top \bar{y})^2 + (\bar{y}^\top \bar{y})^2 \\
 & \quad + 2y_i^\top y_i \bar{y}^\top \bar{y} - 4y_i^\top y_i y_i^\top \bar{y} - 4y_i^\top \bar{y} \bar{y}^\top \bar{y} \} \\
 &= \frac{1}{n-1} \sum_{i=1}^n (y_i^\top y_i)^2 + \frac{4}{n^2(n-1)} \sum_{i=1}^n \left(y_i^\top \sum_{j=1}^n y_j \right)^2 \\
 & \quad + \frac{2}{n^2(n-1)} \sum_{i=1}^n y_i^\top y_i \left(\sum_{j=1}^n \sum_{k=1}^n y_j^\top y_k \right) \\
 & \quad - \frac{4}{n(n-1)} \sum_{i=1}^n y_i^\top y_i y_i^\top \left(\sum_{j=1}^n y_j \right) - \frac{3n}{n-1} (\bar{y}^\top \bar{y})^2 \\
 &= \frac{n^2 - 3n + 3}{n^3} \sum_{i=1}^n (z_i^\top \Sigma z_i)^2 + \frac{4n-6}{n^2(n-1)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n (z_i^\top \Sigma z_j)^2 \\
 & \quad + \frac{2n-3}{n^2(n-1)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n z_i^\top \Sigma z_i z_j^\top \Sigma z_j
 \end{aligned}$$

$$\begin{aligned}
 & - \frac{4(n^2 - 3n + 3)}{n^3(n-1)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n z_i^\top \Sigma z_i z_j^\top \Sigma z_j \\
 & + \frac{2n-6}{n^3(n-1)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq j \\ k \neq i}}^n z_i^\top \Sigma z_i z_j^\top \Sigma z_k \\
 & + \frac{4n-12}{n^3(n-1)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq j \\ k \neq i}}^n z_i^\top \Sigma z_j z_i^\top \Sigma z_k \\
 & - \frac{3}{n^3(n-1)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq j \\ k \neq i}}^n \sum_{\substack{l=1 \\ l \neq k \\ l \neq j \\ l \neq i}}^n z_i^\top \Sigma z_j z_k^\top \Sigma z_l.
 \end{aligned} \tag{۸.۲}$$

مقدار امید ریاضی Q به صورت زیر بیان می شود

$$E[Q] = \frac{n^2 - 3n + 3}{n^2} \kappa_{11} + \frac{2(n-1)}{n} \text{tr}(\Sigma^2) + \frac{n-1}{n} (\text{tr}(\Sigma))^2. \tag{۹.۲}$$

از حل کردن همزمان معادله های (۳.۲)، (۴.۲) و (۹.۲) قضیه اثبات می شود. $\square$

اکنون سازگاری $\hat{a}_2 = \widehat{\text{tr}(\Sigma^2)}/p$ و $\hat{a}_1 = (\widehat{\text{tr}(\Sigma)})^2/p^2$ را بررسی می کنیم. با استفاده از تعاریف $\widehat{\text{tr}(\Sigma^2)}$ و $(\widehat{\text{tr}(\Sigma)})^2$ ، این آماره ها به صورت زیر بیان می شوند

$$\begin{aligned}
 \widehat{\text{tr}(\Sigma^2)} &= \frac{1}{n(n-1)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n (z_i^\top \Sigma z_j)^2 \\
 & - \frac{2}{n(n-1)(n-2)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq j \\ k \neq i}}^n z_i^\top \Sigma z_j z_i^\top \Sigma z_k \\
 & + \frac{1}{n(n-1)(n-2)(n-3)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq j \\ k \neq i}}^n \sum_{\substack{l=1 \\ l \neq k \\ l \neq j \\ l \neq i}}^n z_i^\top \Sigma z_j z_k^\top \Sigma z_l \\
 (\widehat{\text{tr}(\Sigma)})^2 &= \frac{1}{n(n-1)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n z_i^\top \Sigma z_i z_j^\top \Sigma z_j \\
 & - \frac{2}{n(n-1)(n-2)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq j \\ k \neq i}}^n z_i^\top \Sigma z_i z_j^\top \Sigma z_k \\
 & + \frac{1}{n(n-1)(n-2)(n-3)} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq j \\ k \neq i}}^n \sum_{\substack{l=1 \\ l \neq k \\ l \neq j \\ l \neq i}}^n z_i^\top \Sigma z_j z_k^\top \Sigma z_l.
 \end{aligned}$$

لم ۳.۲.۲ (هیمنو و یامادا، ۲۰۰۹). برای مدل (۱.۲) واریانس‌های $\widehat{\text{tr}(\Sigma^2)}$ و $\widehat{(\text{tr}(\Sigma))^2}$ که در قضیه ۲.۲.۲ تعریف شد، به صورت زیر به دست می‌آیند

$$\begin{aligned} \text{Var}(\widehat{\text{tr}(\Sigma^2)}) &= \frac{2}{n(n-1)} E[(z_1^\top \Sigma z_2)^4] + \frac{4(n-3)}{n(n-2)} \kappa_{22} \\ &\quad - \frac{\lambda}{n(n-1)(n-2)} E[(z_1^\top \Sigma z_2)^2 z_1^\top \Sigma^2 z_2] \\ &\quad - \frac{2(n+1)(n-4)}{n(n-1)(n-2)(n-3)} (\text{tr}(\Sigma^2))^2 \\ &\quad + \frac{4(2n^2 - 13n + 23)}{n(n-2)(n-3)} \text{tr}(\Sigma^4) \\ \text{Var}((\widehat{\text{tr}(\Sigma)})^2) &= \frac{2}{n(n-1)} \kappa_{11}^2 + \frac{4(2n-5)}{n(n-1)(n-2)} \kappa_{11} \text{tr}(\Sigma^2) + \frac{4}{n} \kappa_{11} (\text{tr}(\Sigma))^2 \\ &\quad - \frac{\lambda}{n(n-1)(n-2)} E[z_1^\top \Sigma z_1 z_1^\top \Sigma^2 z_2 z_2^\top \Sigma z_2] \\ &\quad + \frac{\lambda(n^2 - 6n + 10)}{n(n-1)(n-2)(n-3)} (\text{tr}(\Sigma^2))^2 \\ &\quad + \frac{4(2n-3)}{n(n-1)} \text{tr}(\Sigma^2) (\text{tr}(\Sigma))^2 \\ &\quad + \frac{16}{n(n-1)(n-2)(n-3)} \text{tr}(\Sigma^4). \end{aligned}$$

از لم ۳.۲.۲ قضیه زیر نتیجه می‌شود. ابتدا شرایط زیر را در نظر بگیرید

$$\begin{aligned} C_1 : E[(z_1^\top \Sigma z_2)^4] &= o(p^4), \quad \kappa_{22} = o(p^3) \quad z_1, z_2 \sim \mathcal{F}, \\ C_2 : \kappa_{11} &= o(p^3). \end{aligned}$$

قضیه ۴.۲.۲ (هیمنو و یامادا، ۲۰۰۹). مدل (۱.۲) را در نظر بگیرید. با توجه به شرایط نظم A_1 و A_2 و همچنین شرط C_1 ، $\hat{a}_2$ برآوردگر سازگار a_2 می‌باشد. علاوه بر این تحت شرایط نظم A_1 ، A_2 و همچنین شرط C_2 ، $\hat{a}_1$ برآوردگر سازگار a_1 می‌باشد.

نتایج زیر با استفاده از نامساوی کوشی – شوارتز به دست آمده‌اند

تحت شرط C_1

$$\begin{aligned} E[(z_1^\top \Sigma z_2)^2 z_1^\top \Sigma^2 z_2] &\leq \sqrt{E[(z_1^\top \Sigma z_2)^4] E[(z_1^\top \Sigma^2 z_2)^2]} \\ &= \sqrt{E[(z_1^\top \Sigma z_2)^4] \text{tr}(\Sigma^4)} = \sqrt{o(p^4) o(p)} \quad (10.2) \end{aligned}$$

در نتیجه، $E[(z_1^\top \Sigma z_2)^2 z_1^\top \Sigma^2 z_2] = o(p^{\frac{5}{2}})$ تحت شرط C_2

$$\begin{aligned} E[(z_1^\top \Sigma z_1)^2 z_1^\top \Sigma^2 z_2 z_2^\top \Sigma z_2] &\leq \sqrt{E[(z_1^\top \Sigma z_1)^2 (z_2^\top \Sigma z_2)^2] E[(z_1^\top \Sigma^2 z_2)^2]} \\ &= \sqrt{\{\kappa_{11} + 2 \text{tr}(\Sigma^2) + (\text{tr}(\Sigma))^2\}^2 \text{tr}(\Sigma^4)} \\ &= o(p^{\frac{5}{2}}). \quad (11.2) \end{aligned}$$

۳.۲ شبه‌سازی‌های عددی

در این بخش، کارایی برآوردهای پیشنهادی را با استفاده از شبه‌سازی مونت کارلو ارزیابی می‌کنیم.

با توجه به این که $\hat{a}_2$ ، $\hat{a}_1$ و $\hat{\kappa}_{11}$ برای μ پایا هستند، فرض کنیم $x_i = \Sigma_i^{\frac{1}{2}} z_i$ (به ازای $i = 1, \dots, n$) که z_i ها نمونه تصادفی و مستقل از رده $\mathcal{F}$ هستند. ماتریس Σ را به صورت $\Sigma_{ij} = (\sigma/2^{|i-j|})$ و حجم نمونه n را برابر با 40 ، 80 یا 120 و نیز بعد پارامتر p را با 40 ، 80 یا 120 در نظر می‌گیریم. چهار حالت زیر را برای رده توزیع‌های $\mathcal{F}$ فرض می‌کنیم

D_1 توزیع نرمال استاندارد، $N(0, 1)$.

D_2 توزیع t با p درجه آزادی t_p .

D_3 توزیع χ^2 با 1 درجه آزادی، χ_1^2 .

D_4 توزیع p متغیره t با 10 درجه آزادی، Mt_{10} .

دقت کنید که توزیع $\mathcal{F}$ استاندارد شده است، بنابراین $\text{Var}(z) = I_p$. تحت این شرط، برآوردهای $\hat{a}_2$ ، $\hat{a}_1$ و $\hat{\kappa}_{11} = (\frac{\text{tr}(\mathcal{S})}{p})^2$ و نیز $\frac{\hat{\kappa}_{11}}{p^2}$ (واریانس نمونه ناریب از این برآوردها) و مقادیر واقعی پارامترها را وقتی $i = 1$ در حالت‌های D_1 ، D_2 ، D_3 و به ازای $i = 2$ در حالت D_4 با تکرار 10000 به دست می‌آوریم.

جدول ۱.۲ مقادیر واقعی a_2 و برآوردهای $\hat{a}_2$ و $\tilde{a}_2$ را نشان می‌دهد. مقادیر داخل پرانتز خطاهای استاندارد را بیان می‌کند. کارایی $\tilde{a}_2$ برای همه توزیع‌ها خوب است به ویژه اگر توزیع نرمال باشد خطای استاندارد $\hat{a}_2$ نزدیک خطای استاندارد $\tilde{a}_2$ است. اگر توزیع نرمال نباشد، کارایی $\tilde{a}_2$ چندان خوب نیست.

جدول ۲.۲ مقادیر واقعی a_1 و برآوردهای $\hat{a}_1$ و $\tilde{a}_1$ را نشان می‌دهد. بر اساس این جدول به نظر می‌رسد $\tilde{a}_1$ بیش‌برازش شده است. اما بر اساس این جدول تفاوت بین $\hat{a}_1$ و $\tilde{a}_1$ به سختی تایید می‌شود.

جدول ۳.۲ مقادیر واقعی $\frac{\kappa_{11}}{p^2}$ و برآوردهای $\frac{\hat{\kappa}_{11}}{p^2}$ را به ازای $i = 1$ برای حالت‌های D_1 ، D_2 ، D_3 و $i = 2$ در حالت D_4 نشان می‌دهد. وقتی $\mathcal{F}$ متقارن باشد، به نظر می‌رسد تقریب مناسب باشد. با این وجود اگر $\mathcal{F}$ کی‌دو باشد، تقریب‌ها بد می‌شود. در این حالت به p و n بزرگتری نیاز داریم. همچنین به ازای n معین، با افزایش p ، مقدار برآوردها افزایش می‌یابد.

جدول ۱.۲: مقدار واقعی a_2 و برآوردهای $\hat{a}_2$ و $\tilde{a}_2$ (و خطاهای استاندارد)

توزیع	n	p	a_2	$\hat{a}_2$	$\tilde{a}_2$
D_1	۴۰	۴۰	۱/۰۸۱	۱/۰۸۱(۰/۱۱۰)	۱/۰۸۱(۰/۱۰۸)
		۸۰	۱/۰۸۲	۱/۰۸۱(۰/۰۹۷)	۱/۰۸۱(۰/۰۹۵)
		۱۲۰	۱/۰۸۳	۱/۰۸۲(۰/۰۸۴)	۱/۰۸۲(۰/۰۸۲)
	۸۰	۴۰	۱/۰۸۱	۱/۰۸۱(۰/۰۷۶)	۱/۰۸۱(۰/۰۷۴)
		۸۰	۱/۰۸۲	۱/۰۸۲(۰/۰۶۳)	۱/۰۸۲(۰/۰۶۲)
		۱۲۰	۱/۰۸۳	۱/۰۸۳(۰/۰۵۸)	۱/۰۸۳(۰/۰۵۵)
	۱۲۰	۴۰	۱/۰۸۱	۱/۰۸۱(۰/۰۶۵)	۱/۰۸۱(۰/۰۶۲)
		۸۰	۱/۰۸۲	۱/۰۸۳(۰/۰۵۰)	۱/۰۸۳(۰/۰۴۸)
		۱۲۰	۱/۰۸۳	۱/۰۸۳(۰/۰۴۰)	۱/۰۸۳(۰/۰۴۹)
D_2	۴۰	۴۰	۱/۰۸۱	۱/۰۸۰(۰/۱۱۲)	۱/۱۰۴(۰/۱۵۶)
		۸۰	۱/۰۸۲	۱/۰۸۴(۰/۰۹۵)	۱/۰۹۲(۰/۰۹۰)
		۱۲۰	۱/۰۸۳	۱/۰۸۲(۰/۰۸۷)	۱/۰۸۸(۰/۰۸۱)
	۸۰	۴۰	۱/۰۸۱	۱/۰۸۲(۰/۰۷۴)	۱/۰۹۴(۰/۰۷۴)
		۸۰	۱/۰۸۲	۱/۰۸۲(۰/۰۵۸)	۱/۰۸۷(۰/۰۵۷)
		۱۲۰	۱/۰۸۳	۱/۰۸۳(۰/۰۴۳)	۱/۰۸۶(۰/۰۴۲)
	۱۲۰	۴۰	۱/۰۸۱	۱/۰۸۲(۰/۰۶۹)	۱/۰۸۹(۰/۰۶۷)
		۸۰	۱/۰۸۲	۱/۰۸۳(۰/۰۵۹)	۱/۰۸۵(۰/۰۵۷)
		۱۲۰	۱/۰۸۳	۱/۰۸۲(۰/۰۳۷)	۱/۰۸۴(۰/۰۳۸)
D_3	۴۰	۴۰	۱/۰۸۱	۱/۰۸۱(۰/۲۱۳)	۱/۳۶۷(۰/۳۴۲)
		۸۰	۱/۰۸۲	۱/۰۸۱(۰/۱۷۱)	۱/۰۸۱(۰/۰۹۵)
		۱۲۰	۱/۰۸۳	۱/۰۸۴(۰/۱۴۹)	۱/۳۷۰(۰/۲۱۵)
	۸۰	۴۰	۱/۰۸۱	۱/۰۸۲(۰/۱۵۶)	۱/۲۲۴(۰/۱۹۲)
		۸۰	۱/۰۸۲	۱/۰۸۴(۰/۱۱۲)	۱/۲۳۱(۰/۱۵۲)
		۱۲۰	۱/۰۸۳	۱/۰۸۳(۰/۰۹۱)	۱/۲۲۸(۰/۱۱۲)
	۱۲۰	۴۰	۱/۰۸۱	۱/۰۸۱(۰/۱۲۴)	۱/۰۸۱(۰/۱۵۴)
		۸۰	۱/۰۸۲	۱/۰۸۲(۰/۰۹۳)	۱/۰۸۱(۰/۰۹۹)
		۱۲۰	۱/۰۸۳	۱/۰۸۲(۰/۰۷۴)	۱/۰۸۰(۰/۰۶۹)
D_4	۴۰	۴۰	۱/۰۸۱	۱/۰۸۲(۰/۲۴۰)	۱/۴۲۱(۰/۵۹۸)
		۸۰	۱/۰۸۲	۱/۰۸۰(۰/۲۳۱)	۱/۷۱۹(۰/۶۹۸)
		۱۲۰	۱/۰۸۳	۱/۰۸۴(۰/۲۲۷)	۲/۰۵۹(۰/۹۸۴)
	۸۰	۴۰	۱/۰۸۱	۱/۰۷۸(۰/۱۶۱)	۱/۲۴۰(۰/۲۴۸)
		۸۰	۱/۰۸۲	۱/۰۸۲(۰/۱۵۲)	۱/۴۱۵(۰/۳۴۹)
		۱۲۰	۱/۰۸۳	۱/۰۸۵(۰/۱۴۸)	۱/۵۴۹(۰/۴۱۷)
	۱۲۰	۴۰	۱/۰۸۱	۱/۰۸۱(۰/۱۲۷)	۱/۱۸۴(۰/۱۷۹)
		۸۰	۱/۰۸۲	۱/۰۸۱(۰/۱۲۵)	۱/۳۰۷(۰/۲۱۳)
		۱۲۰	۱/۰۸۳	۱/۰۸۱(۰/۱۲۲)	۱/۴۱۸(۰/۲۷۹)

جدول ۲.۲: مقدار واقعی a_1^2 و برآوردهای $\hat{a}_1^2$ و $\tilde{a}_1^2$ (و خطاهای استاندارد)

$\tilde{a}_1^2$	$\hat{a}_1^2$	a_1^2	p	N	توزیع
۱/۰۰۱(۰/۰۷۵)	۱/۰۰۰(۰/۰۷۵)	۱/۰۰۰	۴۰	۴۰	D_1
۱/۰۰۱(۰/۰۵۳)	۱/۰۰۱(۰/۰۵۳)	۱/۰۰۰	۸۰		
۱/۰۰۰(۰/۰۴۲)	۰/۹۹۹(۰/۰۴۲)	۱/۰۰۰	۱۲۰		
۱/۰۰۰(۰/۰۵۲)	۱/۰۰۰(۰/۰۵۲)	۱/۰۰۰	۴۰	۸۰	
۱/۰۰۱(۰/۰۳۷)	۱/۰۰۰(۰/۰۳۷)	۱/۰۰۰	۸۰		
۱/۰۰۱(۰/۰۳۰)	۱/۰۰۱(۰/۰۳۰)	۱/۰۰۰	۱۲۰		
۱/۰۰۰(۰/۰۴۳)	۱/۰۰۰(۰/۰۴۳)	۱/۰۰۰	۴۰	۱۲۰	
۱/۰۰۰(۰/۰۳۰)	۱/۰۰۰(۰/۰۳۰)	۱/۰۰۰	۸۰		
۱/۰۰۰(۰/۰۲۵)	۱/۰۰۰(۰/۰۲۵)	۱/۰۰۰	۱۲۰		
۱/۰۰۱(۰/۰۹۰)	۰/۹۹۹(۰/۰۹۰)	۱/۰۰۰	۴۰	۴۰	D_2
۱/۰۰۲(۰/۰۵۷)	۱/۰۰۱(۰/۰۵۷)	۱/۰۰۰	۸۰		
۱/۰۰۱(۰/۰۴۶)	۱/۰۰۰(۰/۰۴۶)	۱/۰۰۰	۱۲۰		
۱/۰۰۲(۰/۰۶۴)	۱/۰۰۱(۰/۰۶۴)	۱/۰۰۰	۴۰	۸۰	
۱/۰۰۲(۰/۰۴۰)	۱/۰۰۰(۰/۰۴۰)	۱/۰۰۰	۸۰		
۱/۰۰۱(۰/۰۳۲)	۱/۰۰۰(۰/۰۳۲)	۱/۰۰۰	۱۲۰		
۱/۰۰۲(۰/۰۵۲)	۱/۰۰۰(۰/۰۵۲)	۱/۰۰۰	۴۰	۱۲۰	
۱/۰۰۱(۰/۰۳۲)	۱/۰۰۰(۰/۰۳۲)	۱/۰۰۰	۸۰		
۱/۰۰۰(۰/۰۲۶)	۱/۰۰۰(۰/۰۲۶)	۱/۰۰۰	۱۲۰		
۱/۰۱۱(۰/۱۹۳)	۱/۰۰۲(۰/۱۹۰)	۱/۰۰۰	۴۰	۴۰	D_3
۱/۰۰۳(۰/۱۳۴)	۰/۹۹۹(۰/۱۳۳)	۱/۰۰۰	۸۰		
۱/۰۰۴(۰/۱۱۰)	۱/۰۰۱(۰/۱۰۹)	۱/۰۰۰	۱۲۰		
۱/۰۰۶(۰/۱۳۲)	۱/۰۰۱(۰/۱۳۱)	۱/۰۰۰	۴۰	۸۰	
۱/۰۰۴(۰/۰۹۵)	۱/۰۰۲(۰/۰۹۵)	۱/۰۰۰	۸۰		
۱/۰۰۲(۰/۰۷۷)	۱/۰۰۱(۰/۰۷۷)	۱/۰۰۰	۱۲۰		
۱/۰۰۳(۰/۱۱۰)	۱/۰۰۰(۰/۱۱۰)	۱/۰۰۰	۴۰	۱۲۰	
۱/۰۰۲(۰/۰۷۷)	۱/۰۰۰(۰/۰۷۷)	۱/۰۰۰	۸۰		
۱/۰۰۱(۰/۰۶۳)	۰/۹۹۹(۰/۰۶۳)	۱/۰۰۰	۱۲۰		
۱/۰۱۰(۰/۲۱۲)	۱/۰۰۰(۰/۲۰۵)	۱/۰۰۰	۴۰	۴۰	D_4
۱/۰۰۵(۰/۱۹۵)	۰/۹۹۶(۰/۱۹۱)	۱/۰۰۰	۸۰		
۱/۰۱۱(۰/۱۹۵)	۱/۰۰۲(۰/۱۹۰)	۱/۰۰۰	۱۲۰		
۱/۰۰۴(۰/۱۴۵)	۰/۹۹۹(۰/۱۴۳)	۱/۰۰۰	۴۰	۸۰	
۱/۰۰۴(۰/۱۳۷)	۰/۹۹۹(۰/۱۳۵)	۱/۰۰۰	۸۰		
۱/۰۰۳(۰/۱۳۶)	۰/۹۹۹(۰/۱۳۴)	۱/۰۰۰	۱۲۰		
۱/۰۰۳(۰/۱۱۷)	۱/۰۰۰(۰/۱۱۶)	۱/۰۰۰	۴۰	۱۲۰	
۱/۰۰۲(۰/۱۱۲)	۰/۹۹۹(۰/۱۱۱)	۱/۰۰۰	۸۰		
۱/۰۰۲(۰/۱۱۰)	۰/۹۹۹(۰/۱۰۹)	۱/۰۰۰	۱۲۰		

جدول ۳.۲: مقدار واقعی κ_{11}/p و برآوردهای $\hat{\kappa}_{11}/p$ (و خطاهای استاندارد)

توزیع	n	p	κ_{11}/p	$\hat{\kappa}_{11}/p$
D_1	۴۰	۴۰	۰/۰۰۰	-۰/۰۰۴(۰/۵۶۴)
		۸۰	۰/۰۰۰	۰/۰۰۵(۰/۵۱۵)
		۱۲۰	۰/۰۰۰	۰/۰۰۱(۰/۵۲۷)
	۸۰	۴۰	۰/۰۰۰	-۰/۰۰۷(۰/۳۴۹)
		۸۰	۰/۰۰۰	۰/۰۰۰(۰/۳۱۷)
		۱۲۰	۰/۰۰۰	۰/۰۰۴(۰/۳۵۹)
	۱۲۰	۳۸	۰/۰۰۰	۰/۰۰۲(۰/۲۹۴)
		۸۰	۰/۰۰۰	۰/۰۰۱(۰/۲۷۴)
		۱۲۰	۰/۰۰۰	۰/۰۰۴(۰/۲۸۷)
D_2	۴۰	۴۰	۱/۰۰۰	۰/۹۹۷(۰/۹۴۴)
		۸۰	۰/۳۷۵	۰/۳۶۹(۰/۶۴۱)
		۱۲۰	۰/۲۳۱	۰/۲۳۶(۰/۵۴۸)
	۸۰	۴۰	۱/۰۰۰	۱/۰۱۱(۰/۶۸۸)
		۸۰	۰/۳۷۵	۰/۳۶۴(۰/۴۲۲)
		۱۲۰	۰/۲۳۱	۰/۲۲۸(۰/۳۹۴)
	۱۲۰	۴۰	۱/۰۰۰	۰/۹۹۶(۰/۵۷۸)
		۸۰	۰/۳۷۵	۰/۳۷۵(۰/۳۴۹)
		۱۲۰	۰/۲۳۱	۰/۲۲۴(۰/۳۱۶)
D_3	۴۰	۴۰	۱۲/۰۰۰	۱۲/۰۵۸(۶/۸۷۹)
		۸۰	۱۲/۰۰۰	۱۱/۹۶۰(۵/۵۹۸)
		۱۲۰	۱۲/۰۰۰	۱۱/۹۷۵(۴/۹۲۶)
	۸۰	۴۰	۱۲/۰۰۰	۱۱/۹۷۰(۴/۹۸۵)
		۸۰	۱۲/۰۰۰	۱۲/۰۴۷(۳/۹۳۵)
		۱۲۰	۱۲/۰۰۰	۱۲/۰۶۰(۳/۴۵۷)
	۱۲۰	۴۰	۱۲/۰۰۰	۱۲/۰۲۰(۴/۲۲۴)
		۸۰	۱۲/۰۰۰	۱۱/۹۹۰(۳/۲۲۷)
		۱۲۰	۱۲/۰۰۰	۱۱/۹۵۷(۲/۷۵۶)
D_4	۴۰	۴۰	۰/۳۵۱	۰/۳۵۶(۰/۵۸۷)
		۸۰	۰/۳۴۲	۰/۳۴۰(۰/۳۲۷)
		۱۲۰	۰/۳۳۹	۰/۳۴۱(۰/۳۲۹)
	۸۰	۴۰	۰/۳۵۱	۰/۳۵۱(۰/۲۳۴)
		۸۰	۰/۳۴۲	۰/۳۴۲(۰/۲۳۸)
		۱۲۰	۰/۳۳۹	۰/۳۳۹(۰/۲۴۸)
	۱۲۰	۴۰	۰/۳۵۱	۰/۳۵۴(۰/۲۱۷)
		۸۰	۰/۳۴۲	۰/۳۴۰(۰/۱۸۴)
		۱۲۰	۰/۳۳۹	۰/۳۴۱(۰/۱۹۱)

فصل ۳

برآوردگرهای ناپارامتری انقباضی نوع استاین ماتریس کوواریانس در بعد بالا

۱.۳ مقدمه

برآورد ماتریس کوواریانس، نقش مهمی در مسائل بعد بالا، وقتی تعداد متغیرها بیشتر از تعداد مشاهدات باشد، دارد. در مطالعات انجام شده، رهیافت‌های انقباضی برای برآورد ماتریس کوواریانس به کار گرفته شده است تا محدودیت‌های ماتریس کوواریانس نمونه (از جمله معکوس پذیر نبودن) را حل کند. یک خانواده جدید از برآوردگرهای ناپارامتری انقباضی نوع استاین ماتریس کوواریانس که مولفه‌های آن به صورت ترکیب خطی محدب از ماتریس کوواریانس نمونه و ماتریس هدف معکوس پذیر از قبل تعیین شده، پیشنهاد شده است. این ترکیب خطی محدب شامل یک پارامتر انقباضی یا همان ضریب موازنه است که باید به درستی تعیین شود. در این فصل، یک فرآیند ساده اما تاثیرگذار برای یافتن برآوردگرهای ناپارامتری سازگار برای پارامتر انقباضی بهینه برای یک ماتریس هدف رایج معرفی شده است.

در این فصل به روش انقباضی نوع استاین (استاین، ۱۹۵۶) که توسط لدویت و ولف (۲۰۰۴) برای برآورد ماتریس کوواریانس به کار برده می‌شود، می‌پردازیم. این روش برآوردگرهایی را نتیجه می‌دهد که ویژگی‌های زیر را دارند:

- (۱) معکوس پذیر هستند.
- (۲) به جایگشت‌های مختلف از p متغیر پایا است.
- (۳) نسبت به فاصله از توزیع نرمال چندمتغیره استوار است.
- (۴) به صورت یک فرم ساده بیان می‌شود.

(۵) صرف نظر از تعداد متغیرهای p ، هزینه محاسباتی کمی دارد.

لدویت و ولف (۲۰۰۴) برآوردهای نوع استاین ماتریس کوواریانس Σ را به صورت

$$S^* = (1 - \lambda)S + \lambda\nu I_p, \quad (1.3)$$

پیشنهاد دادند که انتخاب بهینه λ تابع مخاطره $E[\|S^* - \Sigma\|_F^2]$ را کمینه می کند و در قضیه زیر بیان شده است.

قضیه ۱.۱.۳. مقدار بهینه λ در برآوردهای S^* در رابطه (۱.۳) برابر است با

$$\lambda = \frac{\frac{1}{p}E[\text{tr}((S - \Sigma)^\top (S - \nu I_p))]}{E[\|S - \nu I_p\|_F^2]}$$

$$\nu = \frac{\text{tr}(\Sigma)}{p}$$

برهان. ابتدا مخاطره برآوردهای S^* را به دست می آوریم. داریم

$$\begin{aligned} R(\Sigma, S^*) &= E[\|S^* - \Sigma\|_F^2] \\ &= \frac{1}{p}E[\text{tr}(S^* - \Sigma)^\top (S^* - \Sigma)] \\ &= \frac{1}{p}E[\text{tr}((1 - \lambda)S + \lambda\nu I_p - \Sigma)^\top ((1 - \lambda)S + \lambda\nu I_p - \Sigma)] \\ &= \frac{1}{p}E[\text{tr}(S - \lambda(S - \nu I_p) - \Sigma)^\top (S - \lambda(S - \nu I_p) - \Sigma)] \\ &= \frac{1}{p}E[\text{tr}((S - \Sigma) - \lambda(S - \nu I_p))^\top ((S - \Sigma) - \lambda(S - \nu I_p))] \\ &= \frac{1}{p}E[\text{tr}((S - \Sigma)^\top (S - \Sigma))] - 2\lambda \frac{1}{p}E[\text{tr}((S - \Sigma)^\top (S - \nu I_p))] \\ &\quad + \lambda^2 \frac{1}{p}E[\text{tr}((S - \nu I_p)^\top (S - \nu I_p))] \\ &= E[\|S - \Sigma\|_F^2] - 2\lambda \frac{1}{p}E[\text{tr}((S - \Sigma)^\top (S - \nu I_p))] + \lambda^2 E[\|S - \nu I_p\|_F^2]. \\ \frac{\partial R(\Sigma, S^*)}{\partial \lambda} &= -2 \frac{1}{p}E[\text{tr}((S - \Sigma)^\top (S - \nu I_p))] + 2\lambda E[\|S - \nu I_p\|_F^2] = 0 \end{aligned} \quad (2.3)$$

اکنون با حل معادله نتیجه می شود

$$\lambda = \frac{\frac{1}{p}E[\text{tr}((S - \Sigma)^\top (S - \nu I_p))]}{E[\|S - \nu I_p\|_F^2]}.$$

برای به دست آوردن مقدار بهینه ν ، مخاطره $R(S^*, \Sigma)$ در رابطه (۲.۳) را بر حسب ν به صورت زیر بازنویسی می کنیم

$$\begin{aligned} R(S^*, \Sigma) &= E[\|S - \Sigma\|_F^2] - 2\lambda \frac{1}{p}E[\text{tr}(S^\top S) - \text{tr}(\Sigma^\top S) - \nu \text{tr}(S^\top) \\ &\quad + \nu \text{tr}(\Sigma^\top)] + \lambda^2 E[\text{tr}(S^\top S) - 2\nu \text{tr}(S) + \nu^2 \text{tr}(I_p)] \\ &= E[\|S - \Sigma\|_F^2] - 2\lambda \frac{1}{p}E[S^\top - \text{tr}(\Sigma^\top \Sigma) - \nu \text{tr}(\Sigma^\top) + \nu \text{tr}(\Sigma^\top)] \end{aligned}$$

$$\begin{aligned}
 & +\lambda^2 E[\text{tr}(\mathbf{S}^2) - 2\nu \text{tr}(\mathbf{S}) + p\nu^2] \\
 = & E[\|\mathbf{S} - \Sigma\|_F^2] - 2\lambda \frac{1}{p} [E(\text{tr}(\mathbf{S}^2)) - \text{tr}(\Sigma^2)] \\
 & +\lambda^2 E[(\text{tr}(\mathbf{S}^2))] - 2\nu \text{tr}(\Sigma) + p\nu^2
 \end{aligned} \tag{۳.۳}$$

بنابراین

$$\frac{\partial R(\Sigma, \mathbf{S}^*)}{\partial \nu} = \lambda^2 (-2 \text{tr}(\Sigma) + 2p\nu) = 0$$

و $\nu = \text{tr}(\Sigma)/p$ نتیجه می‌شود و برهان کامل می‌شود. $\square$

با توجه به تعریف $\mathbf{S}^*$ ، پارامتر انقباضی بهینه λ پیشنهاد می‌کند که چه اندازه باید مقادیر ویژه ماتریس کوواریانس نمونه به سمت مقدار ویژه ماتریس هدف νI_p منقبض شود. برای مثال $\lambda = 0$ نتیجه می‌دهد که ماتریس νI_p هیچ سهمی در برآوردگر $\mathbf{S}^*$ ندارد در حالی که $\lambda = 1$ ، بدین معنی است که ماتریس $\mathbf{S}$ سهمی در $\mathbf{S}^*$ ندارد. مقادیر بین ۰ و ۱ برای λ ، به‌طور همزمان برای $\mathbf{S}$ و νI_p در $\mathbf{S}^*$ سهمی قائل می‌شود. با وجود این تفسیر جالب، ماتریس $\mathbf{S}^*$ برآوردگر ماتریس کوواریانس نیست، زیرا ν و λ به ماتریس مشاهده نشده Σ بستگی دارند. بدین منظور، لدویت و ولف (۲۰۰۴) برآوردهای ناپارامتری n سازگار را برای ν و λ در (۱.۳) پیشنهاد داده‌اند و برآوردگر نتیجه را به عنوان برآوردگر انقباضی ماتریس Σ معرفی کردند. اگرچه به نظر می‌رسد ν با $\hat{\nu} = \frac{1}{p} \text{tr}(\mathbf{S})$ به‌صورت مناسبی برآورد شده است. لدویت و ولف (۲۰۰۴) در قالب یک مطالعه شبیه‌سازی نشان دادند برآوردگر پیشنهادی λ ی آن‌ها، در حالت بعد بالا و نیز به ازای $\Sigma = I_p$ ، برآوردگر اریبی است. به عنوان یک اثبات ساده از این موضوع، فرض کنید $\lambda = 1$. در این حالت انتظار داریم برآوردگر جای‌گذاری $\mathbf{S}^*$ تا حد امکان به ماتریس هدف νI_p نزدیک باشد که به سادگی اریب بودن قابل مشاهده است. در این راستا فیشر و سان (۲۰۱۱) برآوردهای ماتریس کوواریانس انقباضی نوع استاین را برای سایر ماتریس‌های هدف پیشنهاد کردند. اما با توجه به اینکه ساختار فرضی آن‌ها بر اساس مدل نرمال چندمتغیره است، برآوردگر پیشنهادی آن‌ها به‌صورت پارامتری معرفی شده است.

با توجه به مطالب بیان‌شده، برآوردگر پارامتر انقباضی بهینه با برآوردگر سازگار λ در حالت بعد بالا بهبود بخشیده می‌شود. برای ساختن برآوردگر λ ، از سه مرحله ساده استفاده می‌شود. (۱) با در نظر گرفتن فرضیه مدل نرمال چندمتغیره، عبارت‌های صورت و مخرج λ را بسط می‌دهیم.

(۲) اثبات می‌کنیم که این نسبت، λ^* ، به‌طور مجانبی معادل λ است.

(۳) هر پارامتر مجهول در λ^* را با برآوردهای سازگار و ناریب جای‌گذاری می‌کنیم.

آخرین مرحله، برآوردگر ناپارامتری و سازگار λ را تعیین می‌کند. علاوه بر این، فرضیه نرمال بودن در مقاله فیشر و سان (۲۰۱۱) برای ماتریس هدف νI_p در (۱.۳) را نادیده می‌گیریم و نشان می‌دهیم که چگونه پارامترهای انقباضی بهینه در حالت بعد بالا به‌طور سازگار برآورد می‌شوند. به عبارت دیگر، خانواده جدیدی از برآوردهای ناپارامتری انقباضی نوع استاین

برای حالت بعد بالا مناسب است که ویژگی‌های جذاب بیان شده را تعیین می‌کند و می‌تواند برای ماتریس‌های هدف دلخواه دیگری نیز به کاربرده شوند، پیشنهاد می‌شود. نتایج این فصل عمدتاً برگرفته از مرجع تولومیس (۲۰۱۵) است.

۲.۳ چارچوب بعد بالا

فرض کنید $x_1, \dots, x_n$ یک نمونه تصادفی از بردارهای تصادفی p متغیره از مدل ناپارامتری

$$x_i = \Sigma^{\frac{1}{2}} z_i + \mu, \quad (4.3)$$

باشد، که در آن $\mu = E[x_i]$ بردار میانگین p متغیره، $\Sigma = \text{Cov}[x_i] = \Sigma^{\frac{1}{2}} \Sigma^{\frac{1}{2}}$ ماتریس کوواریانس $p \times p$ و $z_1, \dots, z_n$ یک نمونه هم‌توزیع و مستقل از بردارهای تصادفی p متغیره هستند. به جای مفروضات توزیعی، محدودیت‌های لحظه‌ای روی متغیرهای تصادفی در z_i اعمال می‌شود.

به‌طور خاص، فرض کنید z_{ia} ، a امین متغیر تصادفی در z_i باشد و فرض کنید $E[z_{ia}] = 0$ ، $E[z_{ia}^2] = 1$ ، $E[z_{ia}^3] = 3 + B$ ، که در آن $-2 \leq B < \infty$ ، و برای اعداد صحیح نامنفی و دلخواه $l_1, \dots, l_4$ به‌طوری که $\sum_{\nu=1}^4 l_{\nu} \leq 4$ ،

$$E[z_{ia_1}^{l_1} z_{ia_2}^{l_2} z_{ia_3}^{l_3} z_{ia_4}^{l_4}] = E[z_{ia_1}^{l_1}] E[z_{ia_2}^{l_2}] E[z_{ia_3}^{l_3}] E[z_{ia_4}^{l_4}]$$

که در آن $a_1, \dots, a_4$ از یکدیگر متمایز هستند.

کمیت λ را نیز به‌صورت زیر در نظر بگیرید

$$\lambda = \frac{E[\|S - \Sigma\|_F^2]}{E[\|S - \nu I_p\|_F^2]}$$

۳.۳ نتایج اصلی

تحت مدل ناپارامتری (۴.۳)، مقدار مورد انتظار در صورت و مخرج λ برابر با

$$\lambda = \frac{E[\|S - \Sigma\|_F^2]}{E[\|S - \nu I_p\|_F^2]} = \frac{\text{tr}(\Sigma^2) + \text{tr}^2(\Sigma) + B \text{tr}(D_{\Sigma}^2)}{n \text{tr}(\Sigma^2) + \frac{p-n+1}{p} \text{tr}^2(\Sigma) + B \text{tr}(D_{\Sigma}^2)}, \quad (5.3)$$

که در آن برای هر ماتریس معین مثبت A ،

$$\text{tr}(D_A^2) = \text{tr}(A \circ A) \leq \text{tr}(A^2) \leq \text{tr}^2(A),$$

اگر از سهم $B \text{tr}(D_{\Sigma}^2)$ برای λ چشم‌پوشی کنیم، خواهیم داشت

$$\lambda^* = \frac{E[\|S - \Sigma\|_F^2]}{E[\|S - \nu I_p\|_F^2]} = \frac{\text{tr}(\Sigma^2) + \text{tr}^2(\Sigma)}{n \text{tr}(\Sigma^2) + \frac{p-n+1}{p} \text{tr}^2(\Sigma)},$$

که در آن λ^* میزان انقباض بهینه مدل نرمال چندمتغیره است. یا به طور کلی، اگر در مدل (۴.۳) $B = 0$ ، آن گاه $\lambda = \lambda^*$ در حالت کلی

$$\begin{aligned} |\lambda - \lambda^*| &= \frac{(n-1) \left(\text{tr}(\Sigma^\vee) - \frac{1}{p} \text{tr}^\vee(\Sigma) \right)}{n \left(\text{tr}(\Sigma^\vee) - \frac{1}{p} \text{tr}^\vee(\Sigma) \right) + \frac{p+1}{p} \text{tr}^\vee(\Sigma)} \\ &\quad \times \frac{|B| \text{tr}(D_\Sigma^\vee)}{n \left(\text{tr}(\Sigma^\vee) - \frac{1}{p} \text{tr}^\vee(\Sigma) \right) + \frac{p+1}{p} \text{tr}^\vee(\Sigma) + B \text{tr}(D_\Sigma^\vee)} \\ &\leq \frac{|B| \frac{\text{tr}(D_\Sigma^\vee)}{\text{tr}^\vee(\Sigma)}}{n \left(\frac{\text{tr}(\Sigma^\vee)}{\text{tr}^\vee(\Sigma)} - \frac{1}{p} \right) + \frac{p+1}{p} + B \frac{\text{tr}(D_\Sigma^\vee)}{\text{tr}^\vee(\Sigma)}} \rightarrow 0. \end{aligned}$$

بنابراین میزان انقباض بهینه بدست آمده تحت حالت نرمال به طور مجانبی معادل میزان انقباض بهینه با توجه به مدل (۴.۳) است.

این بدین معناست که برای برای برآورد λ از برآوردگر ناپارامتری و سازگار λ^* استفاده شود. برای انجام این کار پارامترهای مجهول $\text{tr}(\Sigma)$ و $\text{tr}(\Sigma^\vee)$ در λ^* را با برآوردگرهای نارایب زیر جایگزین می کنیم

$$Y_{1N} = U_{1N} - U_{\epsilon N} = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i^\top \mathbf{x}_i - \frac{1}{P_N} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \mathbf{x}_j^\top \mathbf{x}_i,$$

9

$$\begin{aligned} Y_{\vee N} &= U_{\vee N} - \vee U_{\Delta N} + U_{\epsilon N} \\ &= \frac{1}{P_N} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n (\mathbf{x}_i^\top \mathbf{x}_j)^\vee - \vee \frac{1}{P_N} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq i \\ k \neq j}}^n \mathbf{x}_i^\top \mathbf{x}_j \mathbf{x}_i^\top \mathbf{x}_k \\ &\quad + \frac{1}{P_N} \sum_{i=1}^n \sum_{\substack{j=1 \\ j \neq i}}^n \sum_{\substack{k=1 \\ k \neq i \\ k \neq j}}^n \sum_{\substack{l=1 \\ l \neq i \\ l \neq j \\ l \neq k}}^n \mathbf{x}_i \mathbf{x}_j^\top \mathbf{x}_k \mathbf{x}_l^\top. \end{aligned}$$

که در آن $P_t^s = \frac{s!}{(s-t)!}$ و U_{1N} و $U_{\vee N}$ به ترتیب برآوردگرهای نارایب $\text{tr}(\Sigma)$ و $\text{tr}(\Sigma^\vee)$ وقتی که داده ها حول بردار میانگین هستند (برای مثال $\mu = 0$)، می باشند.

$U_{\epsilon N}$ ، U_{Δ} و U_{ϵ} به ازای $\mu \neq 0$ باعث تضمین نارایبی Y_{1N} و $Y_{\vee N}$ می شوند. بنابراین برآوردگر سازگار λ به صورت زیر است:

$$\hat{\lambda} = \frac{Y_{\vee N} + Y_{1N}^\vee}{nY_{\vee N} + \frac{p-n+1}{p}Y_{1N}^\vee}.$$

برآوردگر انقباضی نوع استاین پیشنهاد شده برای Σ به صورت زیر می باشد

$$\hat{S}^* = (1 - \hat{\lambda})S + \hat{\lambda}\hat{\nu}I_p,$$

که با جایگزینی $\hat{\nu} = \frac{Y_{1N}}{p}$ و $\hat{\lambda}$ در (۱.۳) بدست می آید.

۴.۳ شبیه‌سازی

در این بخش از یک شبیه‌سازی برای بررسی کارایی برآوردگرهای پیشنهادی نوع استاین ماتریس کوواریانس استفاده کرده‌ایم. فرض کنید مشاهدات p بعدی $x_1, x_2, \dots, x_n$ از مدل $x_i = \mu + \Sigma^{1/2} z_i$ تولید شده باشند که سه توزیع زیر را برای متغیرهای $z_1, z_2, \dots, z_n$ و z_n در نظر گرفته‌ایم.

سناریوی ۱ یک سناریوی نرمال که $z_{ia} \sim \mathcal{N}(0, 1)$.

سناریوی ۲ یک سناریوی گاما با $z_{ia} = (z_{ia}^* - 8)/4$ که در آن $z_{ia}^* \sim \Gamma(4, 0.5)$ و بنابراین $B = 12$.

سناریوی ۳ ترکیبی از سناریوهای ۱ و ۲ که در این سناریو، $p/2$ مولفه اول بردار z_i از سناریوی اول و سایر مولفه‌ها از سناریوی دوم تولید شده‌اند.

سناریوی ۲ و ۳ به منظور بررسی ماهیت ناپارامتری روش پیشنهادی در این فصل در نظر گرفته شده‌اند. برای بررسی ساختار بعد بالا، فرض کنید n از ۱۰ تا ۱۰۰ با گام ۱۰ تغییر کند و همچنین $p = 100, 500, 1000, 1500, 2500$. خانواده پیشنهادی از برآوردگرهای کوواریانس با استفاده از $\hat{S}^*$ ، که همان برآوردگر ماتریس کوواریانس پیشنهادی به ازای $T = \nu I_p$ است، محاسبه و سپس با برآوردگرهای $\hat{S}_{FS}^*$ و $\hat{S}_{LW}^*$ ، که به ترتیب ماتریس‌های متناظر با روش‌های لدویت و ولف (۲۰۰۴) و فیشر و سان (۲۰۱۱) هستند، مقایسه می‌شوند. مشابه لدویت و ولف (۲۰۰۴)، فرض می‌کنیم $\mu = 0$ و تغییرات لازم را روی $\hat{S}^*$ و $\hat{S}_{FS}^*$ اعمال می‌کنیم. با توجه به اینکه $\hat{S}_{FS}^*$ بر اساس فرض نرمال بودن ساخته شده است، کارایی آن فقط در سناریوی ۱ بررسی شده است. لازم به توضیح است که برآوردگر شفر و استریمر (۲۰۰۵) در مطالعات شبیه‌سازی در نظر گرفته نشده است زیرا روش آنها، برآوردگر سازگار λ را در مسایل بعد بالا تضمین نمی‌کند. به ازای Σ ، n و p از قبل تعیین شده در هر یک از سناریوهای تعیین شده، معیار درصد بهبودی نسبی در میانگین خطا^۱ PRIAL شبیه‌سازی شده (SPRIAL) ماتریس $\hat{S}^*$ و $\hat{S}_{FS}^*$ را برای برآورد Σ محاسبه کرده‌ایم. معیار SPRIAL برای ماتریس $\hat{\Sigma}$ به صورت زیر تعریف می‌شود:

$$\text{SPRIAL}(\hat{\Sigma}) = \frac{\sum_{b=1}^{1000} \|\hat{S}_{LW,b} - \Sigma\|_F^2 - \sum_{b=1}^{1000} \|\hat{\Sigma}_b - \Sigma\|_F^2}{\sum_{b=1}^{1000} \|\hat{S}_{LW,b} - \Sigma\|_F^2} \times 100\%.$$

که در آن $\hat{S}_{LW,b}^*$ برآوردگر لدویت و ولف (۲۰۰۴) در تکرار b ام ($b = 1, \dots, 1000$) و $\hat{\Sigma}_b$ برآوردگر متناظر از فرآیند برآورد رقیب در برآورد ماتریس کوواریانس را نشان می‌دهد. بر اساس تعریف، $\text{SPRIAL}(\hat{\Sigma}) = 100\%$ و $\text{SPRIAL}(\hat{S}_{LW}^*) = 0\%$. بنابراین مقادیر مثبت (منفی) از $\text{SPRIAL}(\hat{\Sigma})$ نتیجه می‌دهد که برآوردگر $\hat{\Sigma}$ در مقایسه با برآوردگر $\hat{S}_{LW}^*$ کارایی بیشتری (کمتری) دارد در

^۱Percentage relative improvement in average loss

حالی که مقادیر نزدیک صفر، نتیجه می‌دهد که دو برآوردگر تقریباً کارا هستند. دقت کنید از آن جایی که لدویت و ولف (۲۰۰۴) کارایی برآوردگر پیشنهادی‌شان را برای ماتریس کوواریانس نمونه نشان دادند، برآوردگر $\hat{\Sigma}_{LW}^*$ را به عنوان برآوردگر پایه‌ای در نظر گرفتیم.

برای بررسی حالتی که ماتریس هدف با ماتریس کوواریانس واقعی برابر باشد، $\Sigma = I_p$ را در نظر می‌گیریم. در این حالت، برآورد دقیقی از λ حیاتی است زیرا صرف‌نظر از مقادیر n و p ، مقدار λ برابر ۱ است.

جدول ۱.۳ نتایج شبه‌سازی را برای سناریوی ۱ برای میانگین پارامتر انقباضی بهینه بر اساس روش پیشنهادی ($\hat{\lambda}$)، رهیافت لدویت و ولف (۲۰۰۴) ($\hat{\lambda}_{LW}$) و رهیافت فیشر و سان (۲۰۱۱) ($\hat{\lambda}_{FS}$) برای $n = 10, 50, 100$ و $p = 100, 1000, 2500$ نشان می‌دهد. برآوردگرهای $\hat{\lambda}$ ، $\hat{\lambda}_{LW}$ و $\hat{\lambda}_{FS}$ در بازه (۰، ۱) محدود شده‌اند. در مقایسه با $\hat{\lambda}_{LW}$ ، برآوردگر $\hat{\lambda}$ در برآورد λ دقت بیشتری دارند. این برآوردگر ذاتاً اریبی کمتری دارد و به ازای همه مقادیر n و p ، خطای استاندارد کوچکتری دارد. اگرچه اریبی برآوردگر $\hat{\lambda}_{LW}$ با افزایش n کاهش می‌یابد، به نظر می‌رسد که برآوردگر $\hat{\lambda}_{LW}$ ، یک برآوردگر اریب حتی در حالت $n = 100$ است. همان‌طور که انتظار می‌رفت، تفاوت معنی‌داری بین برآوردگرهای $\hat{\lambda}$ و $\hat{\lambda}_{LW}$ در سناریوی ۱ وجود ندارد. جدول ۲.۳، مقادیر $SPRIAL(\hat{S}^*)$ را در سناریوی ۲ و جدول ۳.۳، این مقادیر را برای سناریوی ۳ نشان می‌دهد.

جدول ۱.۳: برآورد پارامتر انقباضی $\lambda = 1$ برای سناریوی ۱ به ازای $\Sigma = I_p$

$\hat{\lambda}_{FS}$	$\hat{\lambda}_{LW}$	$\hat{\lambda}$	p	n
۰/۹۹۱۵(۰/۰۱۰)	۰/۸۹۹۵(۰/۰۱۸)	۰/۹۹۱۵(۰/۰۱۲)	۱۰۰	
۰/۹۹۹۱(۰/۰۰۱)	۰/۸۹۹۹(۰/۰۰۱)	۰/۹۹۸۸(۰/۰۰۱)	۱۰۰۰	۱۰
۰/۹۹۹۶(۰/۰۰۰)	۰/۹۰۰۰(۰/۰۰۰)	۰/۹۹۹۵(۰/۰۰۰)	۲۵۰۰	
۰/۹۹۲۰(۰/۰۰۹)	۰/۹۸۷۸(۰/۰۱۴)	۰/۹۹۱۸(۰/۰۱۰)	۱۰۰	
۰/۹۹۹۳(۰/۰۰۱)	۰/۹۷۹۹(۰/۰۰۲)	۰/۹۹۹۰(۰/۰۰۱)	۱۰۰۰	۵۰
۰/۹۹۹۶(۰/۰۰۵)	۰/۹۷۹۹(۰/۰۰۰)	۰/۹۹۹۵(۰/۰۰۰)	۲۵۰۰	
۰/۹۹۲۰(۰/۰۱۱)	۰/۹۸۵۵(۰/۰۱۴)	۰/۹۹۱۹(۰/۰۱۱)	۱۰۰	
۰/۹۹۹۲(۰/۰۰۱)	۰/۹۸۹۸(۰/۰۰۲)	۰/۹۹۹۱(۰/۰۰۱)	۱۰۰۰	۱۰۰
۰/۹۹۹۶(۰/۰۰۰)	۰/۹۸۹۹(۰/۰۰۰)	۰/۹۹۹۸(۰/۰۰۰)	۲۵۰۰	

جدول ۲.۳: مقادیر $SPRIAL(\hat{S}^*)$ برای سناریوی ۲ به ازای $\Sigma = I_p$

n	p				
	۱۰۰	۵۰۰	۱۰۰۰	۱۵۰۰	۲۵۰۰
۱۰	۶۱/۱۵	۸۹/۲۴	۹۲/۰۴	۹۴/۶۰	۹۸/۶۵
۲۰	۴۲/۹۹	۹۰/۰۲	۹۳/۴۲	۹۴/۶۳	۹۸/۷۵
۳۰	۲۹/۸۰	۸۸/۸۷	۹۳/۹۴	۹۵/۰۶	۹۸/۹۰
۴۰	۲۲/۰۲	۸۳/۲۵	۹۴/۹۵	۹۵/۷۵	۹۸/۹۲
۵۰	۱۸/۹۲	۸۲/۶۳	۹۳/۲۱	۹۶/۰۲	۹۸/۹۷
۶۰	۱۷/۲۴	۷۹/۸۲	۹۲/۵۴	۹۵/۴۸	۹۸/۲۴
۷۰	۱۶/۷۶	۷۸/۲۴	۹۲/۴۳	۹۴/۸۷	۹۸/۵۴
۸۰	۱۵/۴۸	۷۳/۵۹	۹۰/۵۷	۹۳/۰۶	۹۸/۲۳
۹۰	۱۴/۲۶	۶۸/۳۶	۸۵/۶۹	۹۲/۸۵	۹۷/۸۶
۱۰۰	۱۳/۶۹	۶۱/۰۳	۸۴/۰۶	۹۱/۶۹	۹۸/۲۶

جدول ۳.۳: مقادیر $SPRIAL(\hat{S}^*)$ برای سناریوی ۳ به ازای Σ ماتریسی است که (a, b) امین عضو آن به صورت $\Sigma_{ab} = 0/1I(|a - b|)$ است

n	p				
	۱۰۰	۵۰۰	۱۰۰۰	۱۵۰۰	۲۵۰۰
۱۰	۷۸/۵۶	۹۸/۳۶	۹۸/۵۶	۹۸/۹۷	۹۹/۵۴
۲۰	۳۸/۷۵	۷۸/۶۹	۸۲/۳۶	۸۵/۶۹	۹۰/۶۳
۳۰	۱۸/۳۶	۴۲/۶۹	۶۳/۲۶	۷۵/۶۹	۸۳/۲۶
۴۰	۱۵/۲۴	۲۲/۳۹	۴۲/۲۸	۵۹/۳۰	۸۳/۲۷
۵۰	۱۴/۶۷	۱۹/۲۴	۳۵/۶۹	۴۲/۱۶	۴۷/۶۹
۶۰	۱۴/۳۱	۱۵/۳۶	۲۶/۷۸	۳۶/۴۷	۳۹/۸۴
۷۰	۱۴/۰۳	۱۳/۶۸	۱۵/۷۴	۲۰/۰۶	۲۴/۰۳
۸۰	۱۳/۷۵	۱۳/۰۶	۱۴/۶۹	۱۵/۰۳	۲۰/۱۶
۹۰	۱۳/۴۷	۱۲/۹۸	۱۴/۰۵	۱۴/۳۶	۱۹/۲۴
۱۰۰	۱۳/۰۲	۱۲/۲۳	۱۳/۵۷	۱۴/۱۰	۱۸/۳۲

فصل ۴

مقایسه برآوردگر انقباضی خطی ماتریس کوواریانس در توزیع نرمال و غیر نرمال

۱.۴ مقدمه

هنگامی که تعداد متغیرها یعنی p بزرگتر از اندازه نمونه n است، ماتریس کوواریانس نمونه معکوس پذیر نیست. وقتی p بزرگ و نزدیک به n باشد، معکوس ماتریس کوواریانس نمونه حتی اگر $n > p$ ، ممکن است سازگار نباشد. بنابراین به برآوردگری برای ماتریس کوواریانس نیاز داریم که هم معکوس پذیر و هم سازگار باشد.

تاکنون رهیافت‌های بسیاری برای رسیدن به این هدف بررسی شده‌اند. یک رهیافت ساده روش انقباضی خطی است که ترکیب محدب از ماتریس کوواریانس نمونه و ماتریس هدف معین مثبت است. این رهیافت توسط دنیلز و کاس (۲۰۰۱)، لدویت و ولف (۲۰۰۳ و ۲۰۰۴)، شافر و استریمر (۲۰۰۵)، سریواستارا و کوبوکاوا^۱ (۲۰۰۷)، کنو^۲ (۲۰۰۹)، چن و همکاران (۲۰۱۰)، فیشر و سان (۲۰۱۱)، بای و شی (۲۰۱۱) و تولومیس (۲۰۱۵) مورد مطالعه قرار گرفته است. لدویت و ولف (۲۰۱۲ و ۲۰۱۵) روش‌های انقباضی غیرخطی را بسط داده‌اند. از رهیافت‌های دیگر می‌توان به تکنیک‌های منظم‌سازی^۳ و آستانه‌سازی^۴ مانند بیکل و لوینا (۲۰۰۸)، راتمن^۵

^۱Srivastara and Kubokawa

^۲Konno

^۳Regularization

^۴Thresholding

^۵Rothm

و همکاران (۲۰۰۹)، کای و ژائو^۶ (۲۰۱۰)، کای و لیو^۷ (۲۰۱۱) و فن^۸ و همکاران (۲۰۱۱) و (۲۰۱۳) اشاره کرد.

در این فصل، برآوردگرهای انقباضی خطی ساده را بررسی می‌کنیم و یک روش منطقی برای برآورد وزن‌های بهینه در حالت‌های ناپارامتری ارائه می‌شود. لازم به ذکر است که برآوردگرهای انقباضی خطی، حتی در حالت‌های بعد پایین مینیماکس یا مجاز نیستند. مطالب این فصل برگرفته از مرجع ایکدا و همکاران^۹ (۲۰۱۶) است.

۲.۴ برآوردگر انقباضی خطی

فرض کنید S یک ماتریس کوواریانس نمونه با $E(S) = \Sigma$ باشد و همچنین $\hat{\Lambda}$ ماتریس هدف معین مثبت بر پایه S باشد. برآوردگرهای انقباضی خطی، ترکیب‌های محدب S و $\hat{\Lambda}$ هستند که به صورت

$$\begin{aligned}\hat{\Sigma}_w(\hat{\Lambda}) &= (1-w)S + w\hat{\Lambda} \\ &= S - w(S - \hat{\Lambda})\end{aligned}$$

تعریف می‌شوند که در آن w یک ثابت دلخواه است که $0 \leq w \leq 1$. این نوع از برآوردگرها را، برآوردگرهای انقباضی تکی^{۱۰} می‌نامند. می‌خواهیم وزن w را به صورت ناپارامتری بر اساس S با یک وزن بهینه برآورد کنیم. برآورد خطی ناپارامتری توسط هارتیگان^{۱۱} (۱۹۶۹)، رابینز^{۱۲} (۱۹۸۳) و گاش و لاهیری^{۱۳} (۱۹۸۷) مطالعه شده است. مسئله برآورد Σ با برآوردگر $\hat{\Sigma}$ را نسبت به تابع زیان $L(\Sigma, \hat{\Sigma}) = \text{tr}(\hat{\Sigma} - \Sigma)^2$ و در نتیجه تابع مخاطره $R(\Sigma, \hat{\Sigma}) = E[\text{tr}(\hat{\Sigma} - \Sigma)^2]$ در نظر بگیرید. تابع مخاطره بر اساس $\hat{\Sigma}_w(\Lambda)$ به صورت زیر خواهد بود

$$\begin{aligned}R(\Sigma, \hat{\Sigma}_w) &= E[\text{tr}(\hat{\Sigma}_w(\hat{\Lambda}) - \Sigma)^2] \\ &= E[\text{tr}(S - w(S - \hat{\Lambda}) - \Sigma)^2] \\ &= E[\text{tr}((S - \Sigma) - w(S - \hat{\Lambda}))^2] \\ &= E[\text{tr}((S - \Sigma)^2 - 2w(S - \Sigma)(S - \hat{\Lambda}) + w^2(S - \hat{\Lambda})^2)] \\ &= E[\text{tr}(S - \Sigma)^2] - 2wE[\text{tr}(S - \Sigma)(S - \hat{\Lambda})] + w^2E[\text{tr}(S - \hat{\Lambda})^2].\end{aligned}$$

^۶Cai and Zhou

^۷Cai and Liu

^۸Fan

^۹Ikeda et al.

^{۱۰}Single Shrinkage Estimate

^{۱۱}Hartigan

^{۱۲}Robins

^{۱۳}Gosh and Lahiri

برای به دست آوردن وزن بهینه کافیتست مخاطره کمینه شود. برای این منظور با مشتق گیری نسبت به w داریم

$$\frac{d}{dw} R(\Sigma, \hat{\Sigma}_w(\Lambda)) = \text{tr} E[\text{tr}(S - \hat{\Lambda})^2] - \text{tr} E[\text{tr}(S - \Sigma)(S - \hat{\Lambda})] = 0$$

و در نتیجه

$$w^* = \frac{E[\text{tr}(S - \Sigma)(S - \hat{\Lambda})]}{E[\text{tr}(S - \hat{\Lambda})^2]} \quad (1.4)$$

در این فصل، یک روش نااریب برای برآورد صورت و مخرج w^* در حالت های غیر نرمال پیشنهاد می شود.

مثال های معمولی از ماتریس هدف $\hat{\Lambda}$ ، ماتریس کروی $\hat{\Lambda}_1 I_p$ به ازای $\hat{\Lambda}_1 = \text{tr}(S)/p$ و ماتریس قطری $D_s = \text{diag}(s_{11}, \dots, s_{pp})$ برای $S = (s_{ij})$ هستند. برای این ماتریس های هدف، برآوردگرهای ناپارامتری $\hat{w}$ به دست آمده اند. یک برآوردگر انقباضی خطی دیگر، برآوردگر انقباضی دوگانه^{۱۴} خواهد بود. وقتی چندین ماتریس هدف در نظر گرفته شوند، منطقی است که S را به سمت آن ماتریس های هدف منقبض کنیم. برای مثال، فرض کنید این ماتریس های هدف، ماتریس کروی $\hat{\Lambda}_1 I_p$ و ماتریس قطری D_s به عنوان دو نامزد از ماتریس های هدف باشند. می توان ماتریس هدف را به صورت ترکیب محدب زیر در نظر گرفت

$$\Sigma_D = (1 - \alpha)\hat{\Lambda}_1 I_p + \alpha D_s.$$

در این صورت رهیافت انقباضی خطی برآوردگر زیر پیشنهاد می شود

$$\begin{aligned} \hat{\Sigma}(\hat{\Lambda}_1 I_p, D_s) &= (1 - \beta)S + \beta \Sigma_D \\ &= (1 - \beta)S + \beta \left\{ (1 - \alpha)\hat{\Lambda}_1 I_p + \alpha D_s \right\} \end{aligned}$$

که ثابت های α و β در بازه $[0, 1]$ قرار دارند. اگر $w_1 = (1 - \alpha)\beta$ و $w_2 = \alpha\beta$ ، آن گاه برآوردگر $\hat{\Sigma}(\hat{\Lambda}_1 I_p, D_s)$ به صورت زیر بیان می شود

$$\begin{aligned} \hat{\Sigma}_{w_1, w_2}(\hat{\Lambda}_1 I_p, D_s) &= (1 - w_1 - w_2)S + w_1 \hat{\Lambda}_1 I_p + w_2 D_s \\ &= S - w_1(S - \hat{\Lambda}_1 I_p) - w_2(S - D_s) \end{aligned}$$

که در این صورت این برآوردگر را برآوردگر انقباضی دوگانه می نامند.

۳.۴ برآورد وزن بهینه

بردارهای تصادفی p متغیره مستقل و هم توزیع $x_1, \dots, x_n$ با بردار میانگین μ و ماتریس کوواریانس $\Sigma = \Sigma_1^T \Sigma_1^T$ را در نظر بگیرید که Σ_1^T تجزیه چولسکی^{۱۵} با عناصر قطری مثبت

^{۱۴} Double Shrinkage

^{۱۵} Cholesky Decomposition

است. فرض کنید بردارهای مشاهدات x_i به صورت زیر تولید شده‌اند

$$x_i = \mu + \sum_{k=1}^p u_{ki} \quad i = 1, \dots, n,$$

به طوری که

$$E(u_i) = 0, \quad \text{Cov}(u_i) = I_p,$$

و به ازای $\gamma_1, \dots, \gamma_k$ که $0 \leq \sum_{k=1}^p \gamma_k \leq 4$

$$E \left[\prod_{k=1}^p u_{ki}^{\gamma_k} \right] = \prod_{k=1}^p E(u_{ki}^{\gamma_k}) \quad i = 1, \dots, n,$$

که در آن u_{ki} ، k امین مولفه از بردار $u_i = (u_{i1}, \dots, u_{ip})$ است. این فرضیه برای وجود تابع مخاطره ضروری است. سومین و چهارمین گشتاور u_{ki} به صورت $E(u_{ki}^3) = \kappa_3$ و $E(u_{ki}^4) = \kappa_4 + 3\kappa_3^2$ نوشته می‌شوند. در حالت توزیع نرمال، $\kappa_3 = \kappa_4 = 0$. فرض کنید $\Sigma = (\sigma_{ij})$. آن گاه از نمادهای زیر استفاده می‌شود:

$$\begin{aligned} a_1 &= \frac{1}{p} \text{tr}(\Sigma) \\ a_2 &= \frac{1}{p} \text{tr}(\Sigma^2) \\ a_{20} &= \frac{1}{p} \sum_{i=1}^p \sigma_{ii}^2. \end{aligned}$$

برآوردگر نااریب Σ ، ماتریس کوواریانس S است اما در حالت $p > n$ و حتی زمانی که p نزدیک به n است، خوش-وضع^{۱۶} نیست. بنابراین منطقی است که ترکیب‌های محدب S و ماتریس معین مثبت بر پایه S در نظر بگیریم. برای مثال، محدودیت کروی بودن $\Sigma = \sigma^2 I_p$ را در نظر بگیرید. آن گاه σ^2 با $\hat{a}_1 = \text{tr}(S)/p$ برآورد می‌شود و ترکیب‌های محدب $(1-w)S + w\hat{a}_1 I_p$ به ازای $0 \leq w \leq 1$ پیشنهاد شده است. این برآوردگرهای انقباضی خطی، ماتریس S را به سمت ماتریس هدف معین مثبت $\hat{a}_1 I_p$ منقبض می‌کنند. ماتریس هدف دیگر، ماتریس قطری $D_s = \text{diag}(s_{11}, \dots, s_{pp})$ برای $S = (s_{ij})$ است. چنین برآوردگرهای انقباضی خطی توسط لدویت و ولف (۲۰۰۴)، فیشر و سان (۲۰۱۱)، تولومیس (۲۰۱۵) و سایرین مورد بررسی قرار گرفته‌اند.

در حالت کلی، ممکن است محدودیت $\Sigma = \Lambda(\theta)$ در نظر گرفته شود. دقت کنید که برای محدودیت کروی بودن، $\Lambda(\theta) = \sigma^2 I_p$ و برای محدودیت قطری بودن، $\Lambda(\theta) = \text{diag}(\sigma_{11}, \dots, \sigma_{pp})$ است. فرض کنید $\hat{\theta}$ برآوردگر θ بر حسب S تحت محدودیت $\Sigma = \Lambda(\theta)$ باشد. در این صورت، برآوردگرهای انقباضی خطی کلی به صورت

$$\hat{\Sigma}_w(\Lambda) = (1-w)S + w\hat{\Lambda}$$

^{۱۶}Well-conditioned

است که $\hat{\Lambda} = \Lambda(\hat{\theta})$ و $0 \leq w \leq 1$.

وزن بهینه رابطه (۱.۴) را می‌توان به صورت زیر نوشت

$$w^* = \frac{E[\text{tr}(\mathbf{S} - \hat{\Lambda})(\mathbf{S} - \hat{\Lambda})]}{E[\text{tr}(\mathbf{S} - \hat{\Lambda})^2]} = \frac{E[\text{tr}(\mathbf{S}(\mathbf{S} - \hat{\Lambda}))] - E[\text{tr}(\Sigma(\mathbf{S} - \hat{\Lambda}))]}{E[\text{tr}(\mathbf{S} - \hat{\Lambda})^2]}.$$

همچنین

$$\begin{aligned} E[\text{tr}(\Sigma(\mathbf{S} - \hat{\Lambda}))] &= \text{tr}(E(\Sigma(\mathbf{S} - \hat{\Lambda}))) \\ &= \text{tr}(\Sigma E(\mathbf{S}) - E(\Sigma \hat{\Lambda})) \\ &= \text{tr}(\Sigma^2) - E[\text{tr}(\Sigma \hat{\Lambda})]. \end{aligned}$$

دقت کنید که $\text{tr}[\mathbf{S}(\mathbf{S} - \hat{\Lambda})]$ و $\text{tr}[(\mathbf{S} - \hat{\Lambda})^2]$ برآوردگرهای نااریب دقیق $E[\text{tr}(\mathbf{S}(\mathbf{S} - \hat{\Lambda}))]$ و $E[\text{tr}(\mathbf{S} - \hat{\Lambda})^2]$ هستند. فرض کنید

$$G(\Lambda) = E \text{tr}[\Sigma(\mathbf{S} - \hat{\Lambda})] = \text{tr}(\Sigma^2) - E[\text{tr}(\Sigma \hat{\Lambda})].$$

یک برآوردگر به طور مجانبی نااریب از $G(\Lambda)$ را با $\hat{G}$ نشان می‌دهیم، در این صورت

$$E[\hat{G}] \approx G(\Lambda) = \text{tr}[\Sigma E(\mathbf{S} - \hat{\Lambda})].$$

لذا برآوردگر ناپارامتری w^* به صورت

$$\hat{w}^* = \frac{\text{tr}[\mathbf{S}(\mathbf{S} - \hat{\Lambda})] - \hat{G}}{\text{tr}[(\mathbf{S} - \hat{\Lambda})^2]}$$

خواهد بود. از آن جایی که w در بازه $(0, 1)$ محدود شده است، از برآوردگر

$$\hat{w} = \max\{0, \min\{1, \hat{w}^*\}\}$$

استفاده می‌کنیم.

۴.۴ ماتریس هدف کروی

حالتی را در نظر بگیرید که $\Lambda_1 = \sigma^2 \mathbf{I}_p$ و به ازای $\hat{a}_1 = \text{tr}(\mathbf{S})/p$ ،

$$\hat{\Lambda}_1 = \hat{a}_1 \mathbf{I}_p.$$

در این حالت، برآوردگر انقباضی خطی به صورت زیر نوشته می‌شود

$$\hat{\Sigma}_w(\hat{a}_1 \mathbf{I}_p) = (1 - w)\mathbf{S} + w\hat{a}_1 \mathbf{I}_p.$$

می‌دانیم

$$E[\hat{a}_1] = E\left[\frac{1}{p} \text{tr}(\mathbf{S})\right] = \frac{1}{p} \text{tr}(E(\mathbf{S})) = \frac{1}{p} \text{tr}(\Sigma) = a_1$$

از آنجایی که $E\hat{a}_1 = a_1$ داریم

$$\begin{aligned} G(\Lambda_1) &= E[\text{tr}(\Sigma(S - \hat{a}_1 I_p))] \\ &= \text{tr}(\Sigma E(S - \hat{a}_1 I_p)) \\ &= \text{tr}(\Sigma(\Sigma - a_1 I_p)) \\ &= \text{tr}(\Sigma^2) - a_1 \text{tr}(\Sigma) \\ &= p \left\{ \frac{1}{p} \text{tr}(\Sigma^2) - a_1 \frac{1}{p} \text{tr}(\Sigma) \right\} \\ &= p \{ a_2 - a_1^2 \}. \end{aligned}$$

بنابراین باید بدون فرض نرمال بودن، a_2 و a_1^2 را به طور نااریب برآورد کنیم. یک برآوردگر برای a_2 را در حالت توزیع نرمال به صورت زیر در نظر بگیرید

$$\hat{a}_2 = \frac{(n-1)^2}{(n-2)(n+1)} \left[\frac{1}{p} \text{tr}(S^2) - \frac{p}{n-1} \hat{a}_1^2 \right]. \quad (2.4)$$

برای اثبات اینکه برآوردگر $\hat{a}_2$ در رابطه (۲.۴) برای a_2 نااریب است، ابتدا لم زیر را در نظر بگیرید.

لم ۱.۴.۴. فرض کنید ν متغیر تصادفی از توزیع χ^2 با n درجه آزادی باشد. آنگاه

$$\begin{aligned} E(\nu^r) &= n(n+2), \dots, n(n+2r-2) \quad r = 1, 2, \dots \\ \text{Var}(\nu) &= 2n \\ \text{Var}(\nu^2) &= 8n(n+2)(n+4) \\ E(\nu - n)^3 &= 8n \\ E(\nu - n)^4 &= 12n(n+4) \\ E(\nu^2 - n(n+2))^4 &= 24n(n+2)(24n^2 + O(n^3)) \end{aligned}$$

قضیه ۲.۴.۴. در حالت نرمال، برآوردگر $\hat{a}_2$ در رابطه (۲.۴) برای a_2 نااریب است،

برهان. عبارت‌های $\text{tr}(S)$ ، $\text{tr}(S^2)$ و $(\text{tr}(S))^2$ را بر حسب متغیرهای تصادفی کی دو می‌نویسیم. فرض کنید

$$V = (n-1)S = YY^T \sim \mathcal{W}_p(\Sigma, n)$$

که در آن $Y = (y_1, y_2, \dots, y_n)^T$ و $y_i \stackrel{iid}{\sim} \mathcal{N}_p(0, \Sigma)$. همچنین یک ماتریس متعامد $p \times n$ باشد به طوری که $\Gamma \Sigma \Gamma^T = \Lambda$ که در آن $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_p)$ و λ_i ها مقادیر ویژه ماتریس Σ هستند. فرض کنید $U = (u_1, \dots, u_n)^T$ به طوری که $u_i \stackrel{iid}{\sim} \mathcal{N}_p(0, I_p)$ و $Y = \Sigma^{\frac{1}{2}} U$ و $\Sigma^{\frac{1}{2}} \Sigma^{\frac{1}{2}} = \Sigma$ بنابراین

$$\text{tr}(S) = \text{tr} \left(\frac{1}{n-1} YY^T \right)$$

$$\begin{aligned}
 &= \text{tr} \left(\frac{1}{n-1} \mathbf{Y}^\top \mathbf{Y} \right) \\
 &= \text{tr} \left(\frac{1}{n-1} \mathbf{U}^\top \boldsymbol{\Sigma}^\dagger \boldsymbol{\Sigma}^\dagger \mathbf{U} \right) \\
 &= \text{tr} \left(\frac{1}{n-1} \mathbf{U}^\top \boldsymbol{\Sigma} \mathbf{U} \right) \\
 &= \text{tr} \left(\frac{1}{n-1} \mathbf{U} (\boldsymbol{\Gamma}^\top \boldsymbol{\Lambda} \boldsymbol{\Gamma}) \mathbf{U} \right) \\
 &= \text{tr} \left(\frac{1}{n-1} (\boldsymbol{\Gamma}^\top \mathbf{U}) \boldsymbol{\Lambda} (\boldsymbol{\Gamma} \mathbf{U}) \right) \\
 &= \text{tr} \left(\frac{1}{n-1} \mathbf{W}^\top \boldsymbol{\Lambda} \mathbf{W} \right) \quad \mathbf{W} = \boldsymbol{\Gamma} \mathbf{U} \\
 &= \frac{1}{n-1} \sum_{i=1}^p \lambda_i \mathbf{W}_i^\top \mathbf{W}_i \tag{۳.۴}
 \end{aligned}$$

به طوری که $\mathbf{W} = \boldsymbol{\Gamma}_{p \times n} \mathbf{U}_{n \times 1} = (\mathbf{w}_1, \dots, \mathbf{w}_p)$ و $\mathbf{w}_i \stackrel{iid}{\sim} \mathcal{N}_n(\mathbf{0}, \mathbf{I}_n)$ بنابراین، اگر $v_{ii} = \mathbf{w}_i^\top \mathbf{w}_i$ آن گاه $\nu_{ii} \sim \chi_n^2$ پس

$$\begin{aligned}
 (n-1)p\hat{a}_1 &= (n-1)(\text{tr}(\mathbf{S})) = \sum_{i=1}^p \lambda_i \nu_{ii} \\
 (n-1)^2 (\text{tr}(\mathbf{S}))^2 &= \left[\sum_{i=1}^p \lambda_i^2 \nu_{ii}^2 + 2 \sum_{i=1}^p \sum_{\substack{j=1 \\ j \neq i}}^p \lambda_i \lambda_j \nu_{ii} \nu_{jj} \right]
 \end{aligned}$$

9

$$\begin{aligned}
 (n-1)^2 \text{tr}(\mathbf{S}^2) &= \text{tr} \left((\mathbf{W}^\top \boldsymbol{\Lambda} \mathbf{W})(\mathbf{W}^\top \boldsymbol{\Lambda} \mathbf{W}) \right) \\
 &= \text{tr} \left[\left(\sum_{i=1}^p \lambda_i \mathbf{W}_i \mathbf{W}_i^\top \right) \left(\sum_{i=1}^p \lambda_i \mathbf{W}_i \mathbf{W}_i^\top \right) \right] \\
 &= \text{tr} \left[\sum_{i=1}^p \lambda_i^2 \mathbf{W}_i \mathbf{W}_i^\top \mathbf{W}_i \mathbf{W}_i^\top + 2 \sum_{i=1}^p \sum_{\substack{j=1 \\ j \neq i}}^p \lambda_j \lambda_i \mathbf{W}_i \mathbf{W}_i^\top \mathbf{W}_j \mathbf{W}_j^\top \right] \\
 &= \text{tr} \left[\sum_{i=1}^p \lambda_i^2 (\mathbf{W}_i \mathbf{W}_i^\top)^2 + 2 \sum_{i=1}^p \sum_{\substack{j=1 \\ j \neq i}}^p \lambda_i \lambda_j (\mathbf{W}_i \mathbf{W}_j)^\top \mathbf{W}_i \mathbf{W}_j \right] \\
 &= \sum_{i=1}^p \lambda_i^2 (\mathbf{W}_i \mathbf{W}_i^\top)^2 + 2 \sum_{i=1}^p \sum_{\substack{j=1 \\ j \neq i}}^p \lambda_i \lambda_j (\mathbf{W}_i \mathbf{W}_j)^\top \mathbf{W}_i \mathbf{W}_j \\
 &= \sum_{i=1}^p \lambda_i^2 \nu_{ii}^2 + 2 \sum_{i=1}^p \sum_{\substack{j=1 \\ j \neq i}}^p \lambda_i \lambda_j \nu_{ij}
 \end{aligned}$$

که در آن

$$\nu_{ij} = \mathbf{W}_i^\top \mathbf{W}_j \quad i \neq j$$

9

$$\nu_{ii} = \mathbf{W}_i^\top \mathbf{W}_i.$$

با توجه به اینکه $\frac{n^2}{(n-1)(n+2)} \simeq 1$ ، بنابراین

$$\begin{aligned}
 a_2 &\simeq \frac{1}{p} \left[\text{tr}(\mathbf{S}^2) - \frac{1}{(n-1)} (\text{tr}(\mathbf{S}))^2 \right] \\
 &= \frac{1}{(n-1)^2 p} \left[\sum_{i=1}^p \lambda_i^2 \nu_{ii}^2 + 2 \sum_{i=1}^p \sum_{\substack{j=1 \\ j \neq i}}^p \lambda_i \lambda_j \nu_{ij}^2 - \frac{1}{(n-1)} \left(\sum_{i=1}^p \lambda_i^2 \nu_{ii}^2 + 2 \sum_{i=1}^p \sum_{\substack{j=1 \\ j \neq i}}^p \lambda_i \lambda_j \nu_{ij}^2 \right) \right] \\
 &= \frac{1}{(n-1)^2 p} \left[\left(1 - \frac{1}{(n-1)} \right) \sum_{i=1}^p \lambda_i^2 \nu_{ii}^2 + 2 \sum_{i=1}^p \sum_{\substack{j=1 \\ j \neq i}}^p \lambda_i \lambda_j \left(\nu_{ij}^2 - \frac{1}{(n-1)} \nu_{ii} \nu_{jj} \right) \right] \\
 &= \frac{1}{(n-1)^2 p} \left[\left(\frac{n-2}{n-1} \right) \sum_{i=1}^p \lambda_i^2 \nu_{ii}^2 + 2 \sum_{i=1}^p \sum_{\substack{j=1 \\ j \neq i}}^p \lambda_i \lambda_j \left(\nu_{ij}^2 - \frac{1}{(n-1)} \nu_{ii} \nu_{jj} \right) \right] \\
 &= \left[\frac{n-2}{(n-1)^2 p} \sum_{i=1}^p \lambda_i^2 \nu_{ii}^2 \right] + \left[\frac{2}{(n-1)^2 p} \sum_{i=1}^p \sum_{\substack{j=1 \\ j \neq i}}^p \lambda_i \lambda_j z_{ij}^2 \right] \\
 &= q_1 + q_2
 \end{aligned}$$

که در آن

$$\begin{aligned}
 z_{ij} &= \nu_{ij}^2 - \frac{1}{(n-1)} \nu_{ii} \nu_{jj} \\
 &= (\mathbf{W}_i^\top \mathbf{W}_j) - \frac{1}{(n-1)} (\mathbf{W}_i^\top \mathbf{W}_i) (\mathbf{W}_j^\top \mathbf{W}_j).
 \end{aligned}$$

با توجه به رابطه (۴.۴)،

$$\begin{aligned}
 E(a_2) &= \frac{1}{p} E \left[\text{tr}(\mathbf{S}^2) - \frac{1}{n} (\text{tr}(\mathbf{S}))^2 \right] \\
 &= \frac{n-2}{(n-1)^2 p} E \left[\sum_{i=1}^p \lambda_i^2 \nu_{ii}^2 \right] + \frac{2}{(n-1)^2 p} E \left[\sum_{i=1}^p \sum_{\substack{j=1 \\ j > i}}^p \lambda_i \lambda_j z_{ij}^2 \right] \\
 &= \frac{n-2}{(n-1)^2 p} (n-1)(n+1) \sum_{i=1}^p \lambda_i^2 + \frac{2}{(n-1)^2 p} \sum_{i=1}^p \sum_{\substack{j=1 \\ j > i}}^p \lambda_i \lambda_j E(z_{ij}^2) \\
 &= \frac{(n-2)(n+1)}{(n-1)^2} \left(\frac{\text{tr}(\Sigma^2)}{p} \right)
 \end{aligned}$$

بنابراین

$$\frac{(n-1)^2}{(n-2)(n+1)} \frac{1}{p} E \left[\text{tr}(\mathbf{S}^2) - \frac{1}{(n-1)} (\text{tr}(\mathbf{S}))^2 \right]$$

□

برآوردگر نااریب $\text{tr}(\Sigma^2)/p$ است.

با این وجود تحت فرض غیرنرمال بودن، سریواستوا و همکاران (۲۰۱۴) نشان دادند که

$$E[\hat{a}_2] = a_2 + \frac{n-1}{n(n+1)} \kappa_4 a_4.$$

که نتیجه می‌دهد وقتی $a_{\gamma_0} = O(1)$ ، برآوردگر $\hat{a}_{\gamma}$ اریبی مرتبه دوم دارد. سریواستوا و همکاران (۲۰۱۴) و هیمنو و یامادا (۲۰۱۴) به‌طور مستقل از یکدیگر، برآوردگر دقیق ناریب $\hat{a}_{\gamma C}$ را به‌صورت زیر پیشنهاد دادند

$$\hat{a}_{\gamma C} = \frac{n-1}{n(n-2)(n-3)} \left\{ (n-1)(n-2) \frac{1}{p} \text{tr}(\mathbf{S}^{\gamma}) + p\hat{a}_{\gamma}^2 - \frac{1}{p} nQ \right\}$$

که در آن

$$Q = \frac{1}{n-1} \sum_{i=1}^n \left\{ (\mathbf{x}_i - \bar{\mathbf{x}})^{\top} (\mathbf{x}_i - \bar{\mathbf{x}}) \right\}^2 \quad (4.4)$$

همچنین نشان دادند که برآوردگر $\hat{a}_{\gamma C}$ با برآوردگر پیشنهادی چن و همکاران (۲۰۱۰) یکسان است. در مقایسه با شکل ساده چن و همکاران (۲۰۱۰) عبارت $\hat{a}_{\gamma C}$ ساده است و به آسانی به‌کار برده می‌شود.

درباره برآورد a_{γ}^2 ، سریواستوا و همکاران (۲۰۱۱) نشان دادند که در حالت غیرنرمال،

$$\text{Var}(\hat{a}_{\gamma}) = \frac{2}{(n-1)p} a_{\gamma} + \frac{1}{np} \kappa_4 a_{\gamma_0}.$$

بنابراین

$$E[\hat{a}_{\gamma}^2] = a_{\gamma}^2 + \frac{2}{(n-1)p} a_{\gamma} + \frac{1}{np} \kappa_4 a_{\gamma_0}.$$

که بدین معنی است که وقتی $a_{\gamma} = O(1)$ و $a_{\gamma_0} = O(1)$ ، در این صورت

$$E[\hat{a}_{\gamma}^2] = a_{\gamma}^2 + O\left(\frac{1}{np}\right).$$

بنابراین از برآوردگر $\hat{a}_{\gamma}^2$ برای a_{γ}^2 استفاده می‌شود.

با استدلال‌های بالا، می‌توان $G(\Lambda_1)$ را با $\hat{G}_1 = p(\hat{a}_{\gamma C} - \hat{a}_{\gamma}^2)$ برآورد کرد و برآوردگر نتیجه‌شده

با وزن بهینه

$$\hat{w}_{U1} = \max \left\{ 0, \min \left\{ \frac{\frac{1}{p} \text{tr}(\mathbf{S}^{\gamma}) - \hat{a}_{\gamma C}}{\frac{1}{p} \text{tr}(\mathbf{S}^{\gamma}) - \hat{a}_{\gamma}^2}, 1 \right\} \right\}$$

به‌صورت

$$\hat{\Sigma}_U = \hat{w}_U \mathbf{S} + (1 - \hat{w}_U) \hat{a}_1 \mathbf{I}_p$$

می‌باشد. مشابه لدویت و ولف (۲۰۰۴) معیار PRIAL روی ماتریس کوواریانس نمونه به‌صورت زیر تعریف می‌شود

$$\text{PRIAL} = \frac{E[\text{tr}(\mathbf{S} - \Sigma)^2] - E[\text{tr}(\hat{\Sigma}_{w^*}(\hat{a}_1 \mathbf{I}_p) - \Sigma)^2]}{E[\text{tr}(\mathbf{S} - \Sigma)^2]}$$

قضیه ۳.۴.۴. به ازای مقدار بزرگ p ، فرض کنید $a_1 = O(1)$ ، $a_{\gamma_0} = O(1)$ و $a_{\gamma} = O(p^{\delta})$. به ازای

$\delta \geq 0$

(حالت ۱) اگر $\frac{na_2}{p} \rightarrow \infty$ ، $\text{PRIAL} \rightarrow \circ$ آن گاه $\circ$.

(حالت ۲) اگر $\frac{na_2}{p} \rightarrow \circ$ ، $\text{PRIAL} \rightarrow ۱$ آن گاه ۱ .

(حالت ۳) اگر برای C ثابت مثبت، $\frac{na_2}{p} \rightarrow C$ ، در این صورت PRIAL به مقدار

$$\frac{a_1^2 + \frac{1}{n}C}{C - \frac{n}{p}a_1^2 + a_1^2 + \frac{1}{n}C}$$

میل می کند.

برهان. می دانیم

$$\begin{aligned} E \left[\text{tr}(\hat{\Sigma}_{w^*} - \Sigma)^2 \right] &= E \left[\text{tr}(\mathbf{S} - \Sigma) \right] - 2 E \left[\text{tr}(\mathbf{S} - \Sigma)(\mathbf{S} - \hat{a}_1 \mathbf{I}_p) \right] \\ &\quad + E \left[\text{tr}(\mathbf{S} - \hat{a}_1 \mathbf{I}_p)^2 \right] \end{aligned}$$

در این صورت PRIAL به صورت زیر به دست می آید.

$$\text{PRIAL} = w^* \frac{2 E \left[\text{tr}(\mathbf{S} - \Sigma)(\mathbf{S} - \hat{a}_1 \mathbf{I}_p) \right] - w^* E \left[\text{tr}(\mathbf{S} - \hat{a}_1 \mathbf{I}_p)^2 \right]}{E \left[\text{tr}(\mathbf{S} - \Sigma)^2 \right]}$$

با توجه به اینکه

$$w^* = \frac{E \left[\text{tr}(\mathbf{S} - \Sigma)(\mathbf{S} - \hat{a}_1 \mathbf{I}_p) \right]}{E \left[\text{tr}(\mathbf{S} - \hat{a}_1 \mathbf{I}_p)^2 \right]}$$

عبارت PRIAL برابر است با

$$\text{PRIAL} = \frac{E \left[\text{tr}(\mathbf{S} - \Sigma)(\mathbf{S} - \hat{a}_1 \mathbf{I}_p) \right]}{E \left[\text{tr}(\mathbf{S} - \hat{a}_1 \mathbf{I}_p)^2 \right]} \frac{A}{E \left[\text{tr}(\mathbf{S} - \Sigma)^2 \right]}$$

که

$$A = 2 E \left[\text{tr}(\mathbf{S} - \Sigma)(\mathbf{S} - \hat{a}_1 \mathbf{I}_p) \right] - \frac{E \left[\text{tr}(\mathbf{S} - \Sigma)(\mathbf{S} - \hat{a}_1 \mathbf{I}_p) \right]}{E \left[\text{tr}(\mathbf{S} - \hat{a}_1 \mathbf{I}_p)^2 \right]} E \left[\text{tr}(\mathbf{S} - \hat{a}_1 \mathbf{I}_p)^2 \right]$$

صورت و مخرج کسر فوق را می توان به ترتیب به صورت زیر نوشت:

$$\begin{aligned} E \left[\text{tr}(\mathbf{S} - \Sigma)(\mathbf{S} - \hat{a}_1 \mathbf{I}_p) \right] &\left\{ 2 E \left[\text{tr}(\mathbf{S} - \Sigma)(\mathbf{S} - \hat{a}_1 \mathbf{I}_p) \right] E \left[\text{tr}(\mathbf{S} - \hat{a}_1 \mathbf{I}_p)^2 \right] \right. \\ &\quad \left. - E \left[\text{tr}(\mathbf{S} - \Sigma)(\mathbf{S} - \hat{a}_1 \mathbf{I}_p) \right] E \left[\text{tr}(\mathbf{S} - \hat{a}_1 \mathbf{I}_p)^2 \right] \right\} \end{aligned}$$

و

$$E \left[\text{tr}(\mathbf{S} - \hat{a}_1 \mathbf{I}_p)^2 \right] E \left[\text{tr}(\mathbf{S} - \Sigma)^2 \right]$$

بنابراین PRIAL برابر است با

$$\text{PRIAL} = \frac{E \left[\text{tr}(\mathbf{S} - \Sigma)(\mathbf{S} - \hat{a}_1 \mathbf{I}_p) \right] E \left[\text{tr}(\mathbf{S} - \Sigma)(\mathbf{S} - \hat{a}_1 \mathbf{I}_p) \right] E \left[\text{tr}(\mathbf{S} - \hat{a}_1 \mathbf{I}_p)^2 \right]}{E \left[\text{tr}(\mathbf{S} - \hat{a}_1 \mathbf{I}_p)^2 \right] E \left[\text{tr}(\mathbf{S} - \Sigma)^2 \right]}$$

$$= \frac{\{E[\text{tr}(\mathbf{S} - \Sigma)(\mathbf{S} - \hat{\mathbf{a}}\mathbf{I}_p)]\}^2}{E[\text{tr}(\mathbf{S} - \hat{\mathbf{a}}\mathbf{I}_p)]^2 E[\text{tr}(\mathbf{S} - \Sigma)]^2}$$

با به کار بردن لم ۱.۴.۴، عبارت PRIAL به صورت زیر بیان می‌شود

$$\begin{aligned} \text{PRIAL} &= \frac{\left\{ \frac{1}{n-1}a_2 + \frac{p}{n-1}a_1^2 + \frac{1}{n}\kappa_4 a_{20} \right\} \left\{ a_2 - a_1^2 + \frac{p}{n-1}a_1^2 + \frac{p-2}{(n-1)p}a_2 + \frac{p-1}{np}\kappa_4 a_{20} \right\}}{\left\{ \frac{p}{n-1}a_1^2 + \frac{p-2}{(n-1)p}a_2 + \frac{p-1}{np}\kappa_4 a_{20} \right\}^2} \\ &= \frac{\left\{ a_1^2 + \frac{p-2}{p}a_2 + \frac{(p-1)(n-1)}{np}\kappa_4 a_{20} \right\}^2}{\left\{ a_1^2 + \frac{1}{p}a_2 + \frac{n-1}{n}\frac{1}{p}\kappa_4 a_{20} \right\} \left\{ a_1^2 + \frac{n-1}{p}a_2 - \frac{n-1}{p}a_1^2 + \frac{p-2}{p^2}a_2 + \frac{(n-1)(p-1)}{np}\frac{1}{p}\kappa_4 a_{20} \right\}} \end{aligned}$$

در این عبارت، می‌توان مقادیر PRIAL را در سه حالت بررسی کرد

(حالت اول) یعنی $\frac{n}{p}a_2 \rightarrow \infty$. با توجه به فرض‌های قضیه به راحتی نتیجه می‌شود $\text{PRIAL} \rightarrow 0$

(حالت دوم) یعنی $\frac{n}{p}a_2 \rightarrow 0$ ، عبارت PRIAL به ازای $p \rightarrow \infty$ برابر است با

$$\frac{\{a_1^2 + 0 + 0\}^2}{\{a_1^2 + 0 + 0\} \{0 + 0 + a_1^2 + 0\}}$$

(حالت سوم) یعنی $\frac{n}{p}a_2 \rightarrow 0$ ، عبارت PRIAL برابر است با

$$\begin{aligned} &\lim_{p \rightarrow \infty} \frac{\left\{ a_1^2 + \left(1 - \frac{p-2}{p}\right)\frac{1}{n}C + \frac{(p-1)(n-1)}{np}\frac{1}{p}\kappa_4 a_{20} \right\}^2}{\left\{ a_1^2 + \frac{1}{n}C + \frac{n-1}{n}\frac{1}{p}\kappa_4 a_{20} \right\} \left\{ a_1^2 + C - \frac{n-1}{p}a_1^2 + \left(1 - \frac{2}{p}\right)\frac{1}{n}C + \frac{(n-1)(p-1)}{np}\frac{1}{p}\kappa_4 a_{20} \right\}} \\ &= \frac{\left\{ a_1^2 + \frac{1}{n}C \right\}^2}{\left\{ a_1^2 + \frac{1}{n}C \right\} \left\{ C - \frac{n}{p}a_1^2 + a_1^2 + \frac{1}{n}C \right\}} \end{aligned}$$

□

و اثبات کامل می‌شود.

در کاربرد قضیه ۳.۴.۴ می‌توان گفت چنانچه مقدار na_2 خیلی بزرگتر از p شود، مقادیر PRIAL نزدیک صفر بوده و چنانچه خیلی کوچکتر از P باشد به طوری که بتوان نتیجه گرفت $\frac{na_2}{P} \rightarrow 0$ ، آنگاه مقادیر PRIAL نزدیک ۱ است. در این حالت برآوردگر $\hat{\Sigma}_w^*$ برآوردگری کارا بوده و بهتر از S است.

۵.۴ مطالعه شبیه‌سازی

در این بخش کارایی‌های عددی توابع مخاطره برآوردگرهای انقباضی خطی با به کارگیری شبیه‌سازی مونت کارلو مورد بررسی قرار می‌گیرد. در این راستا، فرض کنید $\Delta = \text{diag}(d_1, \dots, d_p)$ که در آن

$$d_i = 0/1 + 1/0 \times U_i \quad (5.4)$$

و U_i عدد تصادفی است که از توزیع یکنواخت $(0, 1)$ تولید شده است. همچنین σ_{ij} ، (i, j) امین مولفه ماتریس کوواریانس Σ باشد که چهار ساختار زیر را برای آن در نظر می گیریم

(مدل الف)

$$\sigma_{ij} = d_i d_j \rho^{|i-j|}$$

(مدل ب)

$$\sigma_{ij} = \rho^{|i-j|}$$

(مدل ج)

$$\sigma_{ij} = d_i d_j \frac{|i-j+1|^{2h} - 2|i-j|^{2h} + |i-j-1|^{2h}}{2}$$

(مدل د)

$$\sigma_{ij} = \frac{|i-j+1|^{2h} - 2|i-j|^{2h} + |i-j-1|^{2h}}{2}.$$

مدل های (ب) و (د) به ترتیب متناظر با ساختار اتورگرسیو و حرکت براونی کسری 17 (FBM) هستند، که هردو آن ها در چن و همکاران (۲۰۱۰) مورد بررسی قرار گرفته اند.

دقت کنید هر چهار مدل در نظر گرفته شده برای ماتریس کوواریانس در شرایط در نظر گرفته شده $a_i = O(1)$ به ازای $i = 1, 2$ و $a_{\infty} = O(1)$ صدق می کنند.

مشاهده های تصادفی $x_1, \dots, x_n$ و به صورت $x_i = \Sigma^{\frac{1}{2}} u_i$ تولید شده اند که $u_i = (u_{i1}, \dots, u_{pi})^T$ و $u_{pi}, \dots, u_{1i}$ به طور مستقل به یکی از سه حالت زیر توزیع شده اند

(حالت اول)

$$u_{ji} \sim \mathcal{N}(0, 1)$$

(حالت دوم)

$$z_{ji} \sim t_m \quad u_{ji} \sim \sqrt{\frac{(m-3)}{m}} z_{ji}$$

(حالت سوم)

$$z_{ji} \sim \chi_{\nu}^2 \quad u_{ji} \sim \frac{(z_{ji} - \nu)}{\sqrt{2\nu}}$$

به طوری که t_m توزیع t با m درجه آزادی و χ_{ν}^2 توزیع کی دو با ν درجه آزادی را نشان می دهند. حالت دوم، توزیع های دم پهن و حالت سوم توزیع های چوله را در بر می گیرند.

کارایی برآوردگرهای پیشنهادی با سایر برآوردگرها با نماد $\hat{\Sigma}_U$ ، $\hat{\Sigma}_{LW}$ ، $\hat{\Sigma}_{RBLW}$ ، $\hat{\Sigma}_{OAS}$ و $\hat{\Sigma}_{FS}$ نشان داده شده و برای راحتی در جدول ها از U ، LW ، $RBLW$ ، OAS و FS استفاده شده است.

¹⁷Fractional Brownian motion

مطالعه شبیه‌سازی برای $p = 100$ ، $n = 10, 40, 70, 100, 130$ ، $\rho = 0.2, 0.4$ ، $h = 0.7$ ، $m = 5$ و $\nu = 2$ انجام شده است. بر اساس ۱۰۰۰ تکرار، مخاطره تجربی این برآوردها محاسبه و معیار PRIAL برای هر برآوردها نسبت به ماتریس نمونه که به صورت زیر تعریف می‌شود، گزارش شده است

$$PRIAL = 100 \times \frac{E[tr(\mathbf{S} - \mathbf{\Sigma})^2] - E[tr(\hat{\mathbf{\Sigma}} - \mathbf{\Sigma})]}{E[tr(\mathbf{S} - \mathbf{\Sigma})]}$$

جدول ۱.۴ مقادیر PRIAL را برای مدل (الف) با $\rho = 0.2$ و $\rho = 0.9$ در توزیع‌های نرمال، $\sqrt{\frac{3}{5}}t_5$ و $\frac{1}{4}(\chi_2^2 - 2)$ که همان حالت‌های ۱، ۲ و ۳ هستند، نشان می‌دهد.

جدول ۲.۴، مقادیر PRIAL را برای مدل (ب) با در نظر گرفتن همان فرض‌های جدول ۱.۴ گزارش می‌کند. مقادیر PRIAL برای مدل‌های (ج) و (د) در توزیع‌های $\sqrt{\frac{3}{5}}t_5$ و $\frac{1}{4}(\chi_2^2 - 2)$ در جدول ۳.۴ آورده شده است.

از این جداول نتایج زیر برداشت می‌شود:

(۱) برآوردها پیشنهادی بهتر از سایر برآوردها رفتار می‌کند اگرچه این تفاوت اندک است. این حالت‌ها شبیه ماتریس‌های قطری هستند که عناصر قطری آن تغییرپذیری زیادی دارند.

(۲) کارایی برآوردهای $\hat{\Sigma}_{AOS}$ و $\hat{\Sigma}_{FS}$ بسیار نزدیک هم است.

(۳) از آنجایی که $\hat{\Sigma}_{RBLW}$ برآوردها راثو-بلکول شده‌ای از $\hat{\Sigma}_{LW}$ با در نظر گرفتن فرض نرمال است، این برآوردها در حالت اول (توزیع نرمال) بهتر از $\hat{\Sigma}_{LW}$ رفتار می‌کند، اما این ویژگی لزوماً در حالت‌های غیرنرمال (توزیع‌های t_5 و χ_2^2) برقرار نیست. دقت کنید $\hat{\Sigma}_U$ برآوردها انقباضی خطی ناپارامتری است درحالی که $\hat{\Sigma}_{OAS}$ و $\hat{\Sigma}_{FS}$ از برآوردهای وزن بهینه با فرض نرمال بودن استفاده می‌کند.

نتایج جدول‌های ۱.۴-۳.۴، این حقیقت را تایید می‌کند که برآوردهای $\hat{\Sigma}_{OAS}$ و $\hat{\Sigma}_{FS}$ کمی بهتر از $\hat{\Sigma}_U$ در بیشتر موارد در نظر گرفته شده برای حالت اول (توزیع نرمال) است، اما برآوردها پیشنهادی در بیشتر موارد گزارش شده برای توزیع‌های غیرنرمال بهتر است.

در جدول ۴.۴، کارایی برآوردهای انقباضی خطی برای مدل‌های (الف) و (ب) به ازای $\rho = 0.2$ ، $n = 40$ و $p = 100$ در حالت توزیع t با m درجه آزادی برای $m = 3, 4, 5, 6, 10$ بررسی شده است. حالت‌های ۳ و ۴ را به منظور بررسی کارایی برآوردها در حالت توزیع‌های دم پهن در نظر گرفتیم. مشاهده می‌شود که $\hat{\Sigma}_{LW}$ به $\hat{\Sigma}_{PBLW}$ ، $\hat{\Sigma}_{OAS}$ و $\hat{\Sigma}_{FS}$ برای مقادیر کوچک m برتری دارد و صرف نظر از مقدار m ، برآوردها پیشنهادی بهتر از $\hat{\Sigma}_{LW}$ است.

۶.۴ خلاصه و پیشنهادات

در این مجموعه برآورد ماتریس کوواریانس یک بردار p بعدی را وقتی p بزرگ است مورد مطالعه قرار دادیم. در این راستا برآوردگرهای انقباضی خطی ماتریس کوواریانس و برآوردگرهای سازگار توابعی از آن را بررسی کرده و با استفاده از مطالعات شبیه‌سازی مقایسه کردیم. از آنجایی که تعیین پارامتر انقباضی در برآوردگر انقباضی خطی ماتریس کوواریانس با بعد بالا از اهمیت بسزایی برخوردار است تحت سناریوهای مختلف شبیه‌سازی شده، مقدار بهینه آن را در هر حالت تعیین کردیم.

با توجه به مطالب ارائه‌شده می‌توان برای آینده این تحقیق موضوعات زیر را پیشنهاد داد:

(۱) با توجه به اینکه در تحلیل ممیزی خطی به معکوس برآوردگر ماتریس کوواریانس نیاز است، می‌توان از برآوردگر انقباضی مطرح‌شده در این مجموعه در مسائل با بعد بالا استفاده کرد.

(۲) در هیچ یک از حالت‌های مورد بحث دیدگاه بیزی برآوردگر کوواریانس بحث نشده که این موضوع می‌تواند در بعد بالا مورد توجه باشد.

(۳) صالح (۲۰۰۶) یک دسته جدیدی از برآوردگرهای انقباضی واریانس را بر پایه برآوردگر انقباضی نوع استاین پیشنهاد کرد که خواص جالب توجهی دارد. می‌توان با الگو گرفتن از آن نتایج را برای حالت ماتریسی تعمیم داد.

جدول ۱.۴: مقدار PRIAL برآوردهای انقباضی خطی برای مدل (الف) به ازای $\rho = 0/2, 0/9$ و $p = 100$ در حالت‌های ۱، ۲ و ۳

FS	OAS	$RBLW$	LW	U_2	U_1	n	
۸۳/۶۹	۸۳/۶۹	۸۲/۷۳	۰۰/۴۸	۹۳/۰۳	۸۳/۶۰	۱۰	نرمال مدل الف $\rho = 0/2$
۵۴/۵۶	۵۴/۵۶	۵۴/۵۴	۵۴/۴۴	۹۰/۰۹	۵۴/۵۵	۴۰	
۴۰/۵۵	۴۰/۵۵	۴۰/۵۵	۴۰/۵۲	۸۷/۳۸	۴۰/۵۴	۷۰	
۳۲/۱۸	۳۲/۱۸	۳۲/۱۸	۳۲/۱۷	۸۴/۸۲	۳۲/۱۸	۱۰۰	
۲۶/۸۷	۲۶/۸۷	۲۶/۸۷	۲۶/۸۷	۸۲/۴۶	۲۶/۸۷	۱۳۰	
۸۳/۲۷	۸۳/۳۱	۸۱/۰۹	۸۱/۳۸	۸۳/۴۰	۸۴/۹۷	۱۰	t_5 مدل الف $\rho = 0/2$
۵۵/۴۴	۵۵/۴۴	۵۵/۱۶	۵۷/۷۹	۷۹/۳۷	۵۷/۹۳	۴۰	
۴۱/۵۷	۴۱/۵۶	۴۱/۴۴	۴۴/۸۶	۷۵/۸۲	۴۴/۹۰	۷۰	
۳۳/۶۴	۳۳/۶۴	۳۳/۵۷	۳۶/۶۴	۷۳/۳۹	۳۶/۶۷	۱۰۰	
۲۷/۸۲	۲۷/۸۲	۲۷/۷۸	۳۱/۴۲	۷۰/۷۲	۳۱/۴۳	۱۳۰	
۸۳/۵۶	۸۳/۷۴	۸۱/۱۴	۸۱/۶۰	۸۰/۹۲	۸۵/۱۴	۱۰	χ^2_2 مدل الف $\rho = 0/2$
۵۶/۶۴	۵۶/۶۴	۵۶/۲۸	۵۸/۰۰	۷۷/۶۰	۵۸/۱۲	۴۰	
۴۲/۹۴	۴۲/۹۴	۴۲/۷۹	۴۴/۰۳	۷۵/۲۷	۴۴/۰۵	۷۰	
۳۴/۶۲	۳۴/۶۱	۳۴/۵۳	۳۵/۵۷	۷۲/۹۶	۳۵/۵۸	۱۰۰	
۲۹/۰۳	۲۹/۰۳	۲۸/۹۸	۲۹/۸۳	۷۰/۸۵	۲۹/۸۳	۱۳۰	
۴۸/۵۶	۴۸/۵۲	۴۸/۳۰	۴۶/۹۰	۴۹/۰۴	۴۷/۷۶	۱۰	χ^2_2 مدل الف $\rho = 0/9$
۱۷/۷۴	۱۷/۷۴	۱۷/۷۹	۱۷/۷۲	۱۸/۶۹	۱۷/۶۵	۴۰	
۱۰/۸۰	۱۰/۸۰	۱۰/۸۱	۱۰/۸۳	۱۱/۴۷	۱۰/۸۰	۷۰	
۷/۸۶	۷/۸۶	۷/۸۷	۷/۸۵	۸/۳۹	۷/۸۴	۱۰۰	
۶/۵۶	۶/۵۶	۶/۵۶	۶/۵۶	۶/۹۶	۶/۵۵	۱۳۰	
۴۹/۴۹	۴۹/۴۹	۴۸/۶۶	۴۸/۷۲	۴۸/۹۶	۴۹/۸۷	۱۰	χ^2_2 مدل الف $\rho = 0/9$
۱۸/۵۷	۱۸/۵۷	۱۸/۵۵	۱۹/۴۳	۱۹/۱۱	۱۹/۳۹	۴۰	
۱۱/۵۵	۱۱/۵۴	۱۱/۵۴	۱۲/۳۵	۱۱/۹۶	۱۲/۳۳	۷۰	
۸/۴۳	۸/۴۳	۸/۴۳	۹/۱۹	۸/۷۶	۹/۱۹	۱۰۰	
۶/۱۸	۶/۱۸	۶/۱۸	۶/۶۸	۶/۴۶	۶/۶۷	۱۳۰	
۵۰/۰۹	۵۰/۱۰	۴۹/۱۸	۴۸/۸۹	۴۹/۴۸	۵۰/۰۸	۱۰	χ^2_2 مدل الف $\rho = 0/9$
۱۸/۸۹	۱۸/۸۸	۱۸/۸۵	۱۹/۲۹	۱۹/۳۸	۱۹/۲۵	۴۰	
۱۱/۶۲	۱۱/۶۱	۱۱/۶۰	۱۱/۹۳	۱۲/۰۲	۱۱/۹۱	۷۰	
۸/۴۷	۸/۴۷	۸/۴۶	۸/۷۲	۸/۸۰	۸/۷۱	۱۰۰	
۶/۶۳	۶/۶۳	۶/۶۳	۶/۸۲	۶/۹۲	۶/۸۱	۱۳۰	

جدول ۲.۴: مقدار PRIAL برآوردگرهای انقباضی خطی برای مدل (ب) به ازای $\rho = ۰/۲, ۰/۹$ و $p = ۱۰۰$ در حالت‌های ۱، ۲ و ۳

FS	OAS	$RBLW$	LW	U_2	U_1	n	
۹۹/۲۱	۹۹/۲۳	۹۵/۳۳	۹۵/۳۳	۹۷/۲۷	۹۹/۲۱	۱۰	نرمال
۹۴/۶۰	۹۴/۶۰	۹۴/۵۳	۹۴/۵۳	۹۲/۶۵	۹۴/۶۰	۷۰	مدل ب
۹۰/۴۲	۹۰/۴۲	۹۰/۴۰	۹۰/۴۱	۸۸/۴۸	۹۰/۴۲	۱۳۰	$\rho = ۰/۲$
۹۸/۴۳	۹۸/۴۶	۹۵/۱۳	۹۶/۵۷	۹۲/۷۴	۹۹/۲۱	۱۰	t_5
۹۴/۰۴	۹۴/۰۴	۹۴/۷۶	۹۳/۹۰	۰۸۸/۰۳	۹۴/۸۳	۷۰	مدل ب
۹۰/۳۰	۹۰/۲۹	۹۰/۷۷	۹۵/۱۳	۸۴/۳۳	۹۰/۷۹	۱۳۰	$\rho = ۰/۲$
۹۸/۹۵	۹۸/۹۸	۹۵/۱۳	۹۴/۷۵	۹۲/۷۰	۹۹/۲۰	۱۰	χ^2_2
۹۴/۵۴	۹۴/۵۵	۹۴/۷۵	۹۰/۸۱	۸۷/۹۴	۹۴/۸۲	۷۰	مدل ب
۹۰/۵۵	۹۰/۵۵	۹۰/۸۱	۸۵/۲۹	۸۳/۹۶	۹۰/۸۳	۱۳۰	$\rho = ۰/۲$
۵۸/۱۵	۵۸/۱۳	۵۶/۰۴	۶۸/۲۵	۵۶/۷۲	۵۷/۶۴	۱۰	نرمال
۱۵/۲۸	۱۵/۲۸	۱۵/۲۹	۸۰/۲۳	۱۴/۸۰	۱۵۰/۲۷	۷۰	مدل ب
۹/۰۶	۹/۰۶	۹/۰۶	۹/۲۵	۸/۷۷	۹/۰۶	۱۳۰	$\rho = ۰/۹$
۵۸/۴۸	۵۸/۵۰	۵۷/۲۰	۵۶/۲۰	۵۶/۶۸	۵۹/۰۱	۱۰	t_5
۱۶/۳۲	۱۶/۳۱	۱۶/۸۴	۱۵/۲۴	۱۵/۶۷	۱۶/۸۴	۷۰	مدل ب
۹/۰۹	۹/۰۹	۹/۳۷	۹/۴۰	۸/۷۱	۹/۳۶	۱۳۰	$\rho = ۰/۹$
۵۹/۱۹	۵۹/۲۱	۵۷/۳۵	۵۸/۷۸	۵۷/۳۱	۵۹/۲۲	۱۰	χ^2_2
۱۶/۲۶	۱۶/۲۵	۱۶/۴۶	۱۶/۲۴	۱۵/۶۰	۱۶/۴۵	۷۰	مدل ب
۹/۴۱	۹/۴۰	۹/۵۳	۹/۴۰	۹/۰۱	۹/۵۳	۱۳۰	$\rho = ۰/۹$

جدول ۳.۴: مقدار PRIAL برآوردگرهای انقباضی خطی برای مدل‌های (ج) و (د) به ازای $h = 0.7$ و $p = 100$ در حالت‌های ۲ و ۳

FS	OAS	$RBLW$	LW	U_2	U_1	n	
۸۰/۳۵	۷۸/۴۸	۷۸/۴۷	۸۱/۷۱	۸۰/۸۸	۸۱/۸۷	۱۰	t_5
۹۴/۶۰	۹۴/۶۰	۹۴/۵۳	۹۴/۵۳	۹۲/۶۵	۹۴/۶۰	۷۰	مدل ج
۹۰/۴۲	۹۰/۴۲	۹۰/۴۰	۹۰/۴۱	۸۸/۴۸	۹۰/۴۲	۱۳۰	$h = 0.7$
۹۸/۴۳	۹۸/۴۶	۹۵/۱۳	۹۶/۵۷	۹۲/۷۴	۹۹/۲۱	۱۰	χ^2_2
۹۴/۰۴	۹۴/۰۴	۹۴/۷۶	۹۳/۹۰	۰۸۸/۰۳	۹۴/۸۳	۷۰	مدل ج
۹۰/۳۰	۹۰/۲۹	۹۰/۷۷	۹۵/۱۳	۸۴/۳۳	۹۰/۷۹	۱۳۰	$h = 0.7$
۹۸/۹۵	۹۸/۹۸	۹۵/۱۳	۹۴/۷۵	۹۲/۷۰	۹۹/۲۰	۱۰	t_5
۹۴.۵۴	۹۴.۵۵	۹۴/۷۵	۹۰۸۱	۸۷/۹۴	۹۴/۸۲	۷۰	مدل د
۹۰.۵۵	۹۰.۵۵	۹۰.۸۱	۸۵.۲۹	۸۳.۹۶	۹۰.۸۳	۱۳۰	$h = 0.7$
۵۸.۱۵	۵۸.۱۳	۵۶.۰۴	۶۸.۲۵	۵۶.۷۲	۵۷.۶۴	۱۰	χ^2_2
۱۵/۲۸	۱۵/۲۸	۱۵/۲۹	۸۰/۲۳	۱۴/۸۰	۱۵/۲۷	۷۰	مدل د
۹/۰۶	۹/۰۶	۹/۰۶	۹/۲۵	۸/۷۷	۹/۰۶	۱۳۰	$h = 0.7$

جدول ۴.۴: مقدار PRIAL برآوردگرهای انقباضی خطی برای مدل‌های (الف) و (ب) به ازای $\rho = 0.2$ ، $p = 100$ و $n = 40$ در توزیع t_m برای $m = 10, 6, 5, 4, 3$

FS	OAS	$RBLW$	LW	U_2	U_1	n	
۵۵/۲۱	۵۵/۱۱	۵۵/۱۷	۸۸/۶۰	۸۷/۷۶	۵۵/۲۹	۱۰	
۵۵/۵۲	۵۵/۳۳	۵۶/۲۴	۸۶/۱۲	۸۳/۷۰	۵۶/۳۶	۶	
۵۵/۰۶	۵۴/۷۸	۹۰/۴۰	۸۴/۵۷	۷۸/۵۰	۵۸/۴۵	۵	مدل الف
۵۰/۲۲	۵۰/۸۸	۵۰/۹۶	۵۰/۹۵	۵۰/۴۸	۶۴/۷۸	۴	
۲۳/۷۸	۲۳/۵۶	۲۳/۴۳	۲۳/۳۷	۲۲/۵۸	۸۸/۸۷	۳	
۹۰/۳۰	۹۰/۲۹	۹۰/۷۷	۹۵/۱۳	۸۴/۳۳	۹۰/۷۹	۱۰	
۹۸/۹۵	۹۸/۹۸	۹۵/۱۳	۹۴/۷۵	۹۲/۷۰	۹۹/۲۰	۶	
۹۴/۵۴	۹۴/۵۵	۹۴/۷۵	۹۰۸۱	۸۷/۹۴	۹۴/۸۲	۵	مدل ب
۹۰/۵۵	۹۰/۵۵	۹۰/۸۱	۸۵/۲۹	۸۳/۹۶	۹۰/۸۳	۴	
۵۸/۱۵	۵۸/۱۳	۵۶/۰۴	۶۸/۲۵	۵۶/۷۲	۵۷/۶۴	۳	

پیوست آ

برنامه‌های کامپیوتری

```
library(expm)
library(matrixcalc)
set.seed(8)
N.Sim = 100
N = 40
n = N - 1
p = 40

Sigma = matrix(0,p,p)
for(i in 1:p){
  for(j in 1:p){
    Sigma[i,j] = 0.2^(abs(i-j))
  }
}

a2 <- function(Sigma){
  matrix.trace(Sigma**Sigma)/p
}

a1.2 <- function(Sigma){
  (matrix.trace(Sigma)/p)^2
}

a2.tilde <- function(S){
  S2 <- S**S
  ((n^2)/((n-1)*(n+2)))*(1/p)*(matrix.trace(S2)-(1/n)*((matrix.trace(S))^2))
}

a2.hat <- function(S){
  S2 <- S**S
```

```
tr.Sigma2 = ((N-1)/(N*(N-2)*(N-3)))*((N-1)*(N-2)*matrix.trace(S2)
+ ((matrix.trace(S))^2) - N*Q)
tr.Sigma2 / p
}
```

```
a1.2.hat <- function(S){
S2 <- S%*%S
tr.Sigma.2 = ((N-1)/(N*(N-2)*(N-3)))*(2*matrix.trace(S2)
+ (N^2 - 3*N + 1)*((matrix.trace(S))^2) - N*Q)
tr.Sigma.2 / (p^2)
}
```

```
a1.2.tilde <- function(S){
((matrix.trace(S))/p)^2
}
```

```
A2.tilde <- numeric(N.Sim)
A2.hat <- numeric(N.Sim)
A1.2.tilde <- numeric(N.Sim)
A1.2.hat <- numeric(N.Sim)
```

```
for(i.sim in 1:N.Sim){
X = matrix(0,nrow = N, ncol = p)
for(i in 1:N){
z = rnorm(p)
X[i,] = sqrtm(Sigma)%*%z
}
S = var(X)
bar.X = colMeans(X)

sum = 0
for(i in 1:N){
sum = sum + (t(X[i,]-bar.X)%*%(X[i,]-bar.X))^2}
Q = sum/(N-1)
```

```
A2.tilde[i.sim] = a2.tilde(S)
A2.hat[i.sim] = a2.hat(S)
A1.2.tilde[i.sim] = a1.2.tilde(S)
A1.2.hat[i.sim] = a1.2.hat(S)
```

```
}
```

```
a2.tilde.result = mean(A2.tilde)
a2.hat.result = mean(A2.hat)
a1.2.tilde.result = mean(A1.2.tilde)
a1.2.hat.result = mean(A1.2.hat)
```

```
a2(Sigma)
a2.tilde.result
a2.hat.result

sd(A2.tilde)
sd(A2.hat)

a1.2(Sigma)
a1.2.tilde.result
a1.2.hat.result
sd(A1.2.tilde)
sd(A1.2.hat)

rm(list = ls(all = TRUE))
set.seed(1234231)
#####
library(mnormt)
library(Matrix)
library(expm)
library(matrixcalc)
library(MASS)
p = 100;
N = 10 #10 - 40 - 70 - 100 - 130
rho = 0.2 #0.9
h = 0.7
m = 5
nu = 2
n = 5000 #number of simulations
#####
mat.tr.S.Sigma1 = matrix(NA,nr = n,nc = 3)
mat.tr.S.Sigma2 = matrix(NA,nr = n,nc = 3)
mat.tr.S.Sigma3 = matrix(NA,nr = n,nc = 3)
mat.tr.S.Sigma4 = matrix(NA,nr = n,nc = 3)
#####
mat.tr.U1.Sigma1 = matrix(NA,nr = n,nc = 3)
mat.tr.U1.Sigma2 = matrix(NA,nr = n,nc = 3)
mat.tr.U1.Sigma3 = matrix(NA,nr = n,nc = 3)
mat.tr.U1.Sigma4 = matrix(NA,nr = n,nc = 3)
#####
mat.tr.U2.Sigma1 = matrix(NA,nr = n,nc = 3)
mat.tr.U2.Sigma2 = matrix(NA,nr = n,nc = 3)
mat.tr.U2.Sigma3 = matrix(NA,nr = n,nc = 3)
mat.tr.U2.Sigma4 = matrix(NA,nr = n,nc = 3)
#####
mat.tr.DS.Sigma1 = matrix(NA,nr = n,nc = 3)
```

```

mat.tr.DS.Sigma2 = matrix(NA,nr = n,nc = 3)
mat.tr.DS.Sigma3 = matrix(NA,nr = n,nc = 3)
mat.tr.DS.Sigma4 = matrix(NA,nr = n,nc = 3)
#####
mat.tr.LW.Sigma1 = matrix(NA,nr = n,nc = 3)
mat.tr.LW.Sigma2 = matrix(NA,nr = n,nc = 3)
mat.tr.LW.Sigma3 = matrix(NA,nr = n,nc = 3)
mat.tr.LW.Sigma4 = matrix(NA,nr = n,nc = 3)
#####
mat.tr.RBLW.Sigma1 = matrix(NA,nr = n,nc = 3)
mat.tr.RBLW.Sigma2 = matrix(NA,nr = n,nc = 3)
mat.tr.RBLW.Sigma3 = matrix(NA,nr = n,nc = 3)
mat.tr.RBLW.Sigma4 = matrix(NA,nr = n,nc = 3)
#####
mat.tr.OAS.Sigma1 = matrix(NA,nr = n,nc = 3)
mat.tr.OAS.Sigma2 = matrix(NA,nr = n,nc = 3)
mat.tr.OAS.Sigma3 = matrix(NA,nr = n,nc = 3)
mat.tr.OAS.Sigma4 = matrix(NA,nr = n,nc = 3)
#####
mat.tr.FS.Sigma1 = matrix(NA,nr = n,nc = 3)
mat.tr.FS.Sigma2 = matrix(NA,nr = n,nc = 3)
mat.tr.FS.Sigma3 = matrix(NA,nr = n,nc = 3)
mat.tr.FS.Sigma4 = matrix(NA,nr = n,nc = 3)
#####
#Sturctures of Sigma
U.rand <- runif(p,min = 0, max = 1)
D.vec <- 0.1 + 10*U.rand
D.mat <- diag(D.vec)

# Sigma matrix - structure 1 - Model A-D
# Sigma1 <- matrix(0,p,p)
# for(i in 1:p){
#   for(j in 1:p){
#     Sigma1[i,j] = D.vec[i]*D.vec[j]*rho^(abs(i-j))
#   }
# }

# write.matrix(Sigma1, file = "/home/mina_norouzirad/results/Sigma1.txt", sep = " ")
Sigma1 <- as.matrix(read.table("/home/mina_norouzirad/results/Sigma1.txt"))

# Sigma matrix - structure 2 - Model A-I
Sigma2 <- matrix(0,p,p)
for(i in 1:p){
  for(j in 1:p){
    Sigma2[i,j] = rho^(abs(i-j))
  }
}

# Sigma matrix - structure 3 - Model B-D
# Sigma3 <- matrix(0,p,p)
# for(i in 1:p){
#   for(j in 1:p){

```

```

#      Sigma3[i,j] = D.vec[i]*D.vec[j]*((abs(i - j +1)^(2*h)) - 2*(abs(i-j)^(2*h))
+ (abs(i - j - 1)^(2*h)))/2
#    }
#  }

# write.matrix(Sigma3, file =  "/home/mina_norouzirad/results/Sigma3.txt", sep = " ")
Sigma3 <-as.matrix(read.table("/home/mina_norouzirad/results/Sigma3.txt"))

# Sigma matrix - structure 4 - Model B-I
Sigma4 <- matrix(0,p,p)
for(i in 1:p){
  for(j in 1:p){
    Sigma4[i,j] = ((abs(i - j +1)^(2*h)) - 2*(abs(i-j)^(2*h)) + (abs(i - j - 1)^(2*h)))/2
  }
}

#Simulation loop
for(i.n in 2501:5000){
  #Case 1 - Normal distribution
  U1 <- matrix(rnorm(N*p,0,1),nr = N, nc = p)
  X1 <- matrix(0, nr = N, nc = p)
  for(i.1 in 1:N){
    X1[i.1,] <- sqrtm(Sigma1)%*%U1[i.1,]
  }

  #Case 2 - t distribution with m degrees of freedom
  Z2 <- matrix(rt(N*p,m),nr = N, nc = p)
  X2 <- matrix(0, nr = N, nc = p)
  U2 = sqrt((m-2)/m) * Z2
  for(i.2 in 1:N){
    X2[i.2,] <- sqrtm(Sigma1)%*%U2[i.2,]
  }

  #Case 3 - Chi square distribution with nu degrees of freedom
  Z3 <- matrix(rchisq(N*p,m,df = nu),nr = N, nc = p)
  X3 <- matrix(0, nr = N, nc = p)
  U3 = ((Z3 - nu)/(sqrt(2*nu)))
  for(i.3 in 1:N){
    X3[i.3,] <- sqrtm(Sigma1)%*%U3[i.3,]
  }

  #####
  S1 = cov(X1)
  S2 = cov(X2)
  S3 = cov(X3)

  #####
  #####
  S1.2 = S1%*%S1
  S2.2 = S2%*%S2
  S3.2 = S3%*%S3

  #####
  #####
  a1.hat.1 <- matrix.trace(S1)/p
  a1.hat.2 <- matrix.trace(S2)/p

```

```

a1.hat.3 <- matrix.trace(S3)/p
#####
#####
X1.bar = apply(X1,2,mean)
X2.bar = apply(X2,2,mean)
X3.bar = apply(X3,2,mean)
#####
sum.1 = sum.2 = sum.3 = 0
for(i.N in 1:N){
  sum.1 = sum.1 + ((t(X1[i.N,]-X1.bar)%*(X1[i.N,]-X1.bar))^2)
  sum.2 = sum.2 + ((t(X2[i.N,]-X2.bar)%*(X2[i.N,]-X2.bar))^2)
  sum.3 = sum.3 + ((t(X3[i.N,]-X3.bar)%*(X3[i.N,]-X3.bar))^2)
}
#####
#####
Q.1 = sum.1/(N-1)
Q.2 = sum.2/(N-1)
Q.3 = sum.3/(N-1)
#####
#####
a2C.hat.1 <- ((N-1)/(N*(N-2)*(N-3)))*((N-1)*(N-2)*matrix.trace(S1.2)/p + p*a1.hat.1^2 - N*Q.1/p)
a2C.hat.2 <- ((N-1)/(N*(N-2)*(N-3)))*((N-1)*(N-2)*matrix.trace(S2.2)/p + p*a1.hat.2^2 - N*Q.2/p)
a2C.hat.3 <- ((N-1)/(N*(N-2)*(N-3)))*((N-1)*(N-2)*matrix.trace(S3.2)/p + p*a1.hat.3^2 - N*Q.3/p)
#####
#####
#U1 method
#####
#####
w.U1.1 = max(0,min(1,((matrix.trace(S1.2)/p - a2C.hat.1)/(matrix.trace(S1.2)/p - (a1.hat.1^2))))))
w.U1.2 = max(0,min(1,((matrix.trace(S2.2)/p - a2C.hat.2)/(matrix.trace(S2.2)/p - (a1.hat.2^2))))))
w.U1.3 = max(0,min(1,((matrix.trace(S3.2)/p - a2C.hat.3)/(matrix.trace(S3.2)/p - (a1.hat.3^2))))))
#####
#####
Sigma.U1.1 = (1-w.U1.1)*S1 + w.U1.1*a1.hat.1*diag(1,p,p)
Sigma.U1.2 = (1-w.U1.2)*S2 + w.U1.2*a1.hat.2*diag(1,p,p)
Sigma.U1.3 = (1-w.U1.3)*S3 + w.U1.3*a1.hat.3*diag(1,p,p)
#####
#####
a20.hat.K4.1 <- (-1/((N-2)*(N-3)*p))*(2*((N-1)^2)*matrix.trace(S1.2)
+ ((N-1)^2)*((matrix.trace(S1.2)^2)) - (N*(N+1)*Q.1))
a20.hat.K4.2 <- (-1/((N-2)*(N-3)*p))*(2*((N-1)^2)*matrix.trace(S2.2)
+ ((N-1)^2)*((matrix.trace(S2.2)^2)) - (N*(N+1)*Q.2))
a20.hat.K4.3 <- (-1/((N-2)*(N-3)*p))*(2*((N-1)^2)*matrix.trace(S3.2)
+ ((N-1)^2)*((matrix.trace(S3.2)^2)) - (N*(N+1)*Q.3))
#####
#####
D.S.1 = D.S.2 = D.S.3 = matrix(0, nr = p, nc = p)
diag(D.S.1) = diag(S1)
diag(D.S.2) = diag(S2)

```

```

diag(D.S.3) = diag(S3)
#####
#####
a20.hat.1 <- ((N-1)/((N+1)*p))*matrix.trace(D.S.1%^2) - a20.hat.K4.1/(N+1)
a20.hat.2 <- ((N-1)/((N+1)*p))*matrix.trace(D.S.2%^2) - a20.hat.K4.2/(N+1)
a20.hat.3 <- ((N-1)/((N+1)*p))*matrix.trace(D.S.3%^2) - a20.hat.K4.3/(N+1)
#####
#####
#U2 method
#####
w.U2.1 = max(0,min(1,((matrix.trace(S1.2)/p - matrix.trace(S1%*D.S.1)/p - a2C.hat.1 +
  a20.hat.1)/(matrix.trace(S1.2)/p - (a1.hat.1^2))))))
w.U2.2 = max(0,min(1,((matrix.trace(S2.2)/p - matrix.trace(S2%*D.S.2)/p - a2C.hat.2 +
  a20.hat.2)/(matrix.trace(S2.2)/p - (a1.hat.2^2))))))
w.U2.3 = max(0,min(1,((matrix.trace(S3.2)/p - matrix.trace(S3%*D.S.3)/p - a2C.hat.3 +
  a20.hat.3)/(matrix.trace(S3.2)/p - (a1.hat.3^2))))))
#####
#####
Sigma.U2.1 = (1-w.U2.1)*S1 + w.U2.1*D.S.1
Sigma.U2.2 = (1-w.U2.2)*S2 + w.U2.2*D.S.2
Sigma.U2.3 = (1-w.U2.3)*S3 + w.U2.3*D.S.3
#####
#####
mat.1 = matrix(c(matrix.trace((S1-a1.hat.1*diag(1,p,p))%^2),
  matrix.trace((S1-a1.hat.1*diag(1,p,p))%*(S1-D.S.1)),matrix.trace((S1-a1.hat.1*diag(1,p,p))
  %*(S1-D.S.1)),matrix.trace((S1-D.S.1)%^2)),nr = 2,nc = 2, byrow = TRUE)
mat.2 = matrix(c(matrix.trace((S2-a1.hat.2*diag(1,p,p))%^2),
  matrix.trace((S2-a1.hat.2*diag(1,p,p))%*(S2-D.S.2)),matrix.trace((S2-a1.hat.2*diag(1,p,p))
  %*(S2-D.S.2)),matrix.trace((S2-D.S.2)%^2)),nr = 2,nc = 2, byrow = TRUE)
mat.3 = matrix(c(matrix.trace((S3-a1.hat.3*diag(1,p,p))%^2),
  matrix.trace((S3-a1.hat.3*diag(1,p,p))%*(S3-D.S.3)),matrix.trace((S3-a1.hat.3*diag(1,p,p))
  %*(S3-D.S.3)),matrix.trace((S3-D.S.3)%^2)),nr = 2,nc = 2, byrow = TRUE)
#####
vec.1 = matrix(c(matrix.trace(S1.2) - p*a2C.hat.1, matrix.trace(S1.2)-matrix.trace(S1%*D.S.1)
  - p*(a2C.hat.1 - a20.hat.1)))
vec.2 = matrix(c(matrix.trace(S2.2) - p*a2C.hat.2, matrix.trace(S2.2)-matrix.trace(S2%*D.S.2)
  - p*(a2C.hat.2 - a20.hat.2)))
vec.3 = matrix(c(matrix.trace(S3.2) - p*a2C.hat.3, matrix.trace(S3.2)-matrix.trace(S3%*D.S.3)
  - p*(a2C.hat.3 - a20.hat.3)))
#####
w.star.1 = solve(mat.1)%*%vec.1
w.star.2 = solve(mat.2)%*%vec.2
w.star.3 = solve(mat.3)%*%vec.3
#####
beta.hat.star.1 = w.star.1[1] + w.star.1[2]
beta.hat.star.2 = w.star.2[1] + w.star.2[2]
beta.hat.star.3 = w.star.3[1] + w.star.3[2]
#####

```

```

alpha.hat.star.1 = w.star.1[2]/(w.star.1[1]+w.star.1[2])
alpha.hat.star.2 = w.star.2[2]/(w.star.2[1]+w.star.2[2])
alpha.hat.star.3 = w.star.3[2]/(w.star.3[1]+w.star.3[2])
#####
alpha.hat.1 = max(0,min(alpha.hat.star.1,1))
alpha.hat.2 = max(0,min(alpha.hat.star.2,1))
alpha.hat.3 = max(0,min(alpha.hat.star.3,1))
#####
beta.hat.1 = max(0,min(beta.hat.star.1,1))
beta.hat.2 = max(0,min(beta.hat.star.2,1))
beta.hat.3 = max(0,min(beta.hat.star.3,1))
#####
wD1.hat.1 = (1-alpha.hat.1)*beta.hat.1
wD1.hat.2 = (1-alpha.hat.2)*beta.hat.2
wD1.hat.3 = (1-alpha.hat.3)*beta.hat.3
#####
wD2.hat.1 = alpha.hat.1 * beta.hat.1
wD2.hat.2 = alpha.hat.2 * beta.hat.2
wD2.hat.3 = alpha.hat.3 * beta.hat.3
#####
#Double Shrinkage method
#####
Sigma.DS.1 = (1-wD1.hat.1 - wD2.hat.1)*S1 + wD1.hat.1*a1.hat.1*diag(1,p,p) + wD2.hat.1*D.S.1
Sigma.DS.2 = (1-wD1.hat.2 - wD2.hat.2)*S2 + wD1.hat.2*a1.hat.2*diag(1,p,p) + wD2.hat.2*D.S.2
Sigma.DS.3 = (1-wD1.hat.3 - wD2.hat.3)*S3 + wD1.hat.3*a1.hat.3*diag(1,p,p) + wD2.hat.3*D.S.3
#####
sum.LW.1 = sum.LW.2 = sum.LW.3 = 0
for(i.LW in 1:N){
  sum.LW.1 = sum.LW.1 + (matrix.trace(((X1[i.LW,]-X1.bar)%*t(X1[i.LW,]-X1.bar) - S1)^2)/p))
  sum.LW.2 = sum.LW.2 + (matrix.trace(((X2[i.LW,]-X2.bar)%*t(X2[i.LW,]-X2.bar) - S2)^2)/p))
  sum.LW.3 = sum.LW.3 + (matrix.trace(((X3[i.LW,]-X3.bar)%*t(X3[i.LW,]-X3.bar) - S3)^2)/p))
}
#####
W.LW.1 = max(0,min(1,sum.LW.1 / (((N-1)^2)*(matrix.trace(S1.2)-((matrix.trace(S1))^2)/p))))
W.LW.2 = max(0,min(1,sum.LW.2 / (((N-1)^2)*(matrix.trace(S2.2)-((matrix.trace(S2))^2)/p))))
W.LW.3 = max(0,min(1,sum.LW.3 / (((N-1)^2)*(matrix.trace(S3.2)-((matrix.trace(S3))^2)/p))))
#####
#Ledoit and Wolf method
#####
Sigma.LW.1 = (1-W.LW.1)*S1 + W.LW.1*a1.hat.1*diag(1,p,p)
Sigma.LW.2 = (1-W.LW.2)*S2 + W.LW.2*a1.hat.2*diag(1,p,p)
Sigma.LW.3 = (1-W.LW.3)*S3 + W.LW.3*a1.hat.3*diag(1,p,p)
#####
W.RBLW.1 = max(0,min(1,((N-3)*matrix.trace(S1.2)
+ (N-1)*(matrix.trace(S1))^2) / (((N-1)*(N+1))*(matrix.trace(S1.2)-((matrix.trace(S1))^2)/p))))
W.RBLW.2 = max(0,min(1,((N-3)*matrix.trace(S2.2)
+ (N-1)*(matrix.trace(S2))^2) / (((N-1)*(N+1))*(matrix.trace(S2.2)-((matrix.trace(S2))^2)/p))))
W.RBLW.3 = max(0,min(1,((N-3)*matrix.trace(S3.2)

```

```

+ (N-1)*(matrix.trace(S3))^2) / (((N-1)*(N+1))*(matrix.trace(S3.2)-((matrix.trace(S3))^2)/p)))
#####
#Rao-Blackwell-Ledoit-Wolf method
#####
Sigma.RBLW.1 = (1-W.RBLW.1)*S1 + W.RBLW.1*a1.hat.1*diag(1,p,p)
Sigma.RBLW.2 = (1-W.RBLW.2)*S2 + W.RBLW.2*a1.hat.2*diag(1,p,p)
Sigma.RBLW.3 = (1-W.RBLW.3)*S3 + W.RBLW.3*a1.hat.3*diag(1,p,p)
#####
W.OAS.1 = max(0,min(1,((p-2)*matrix.trace(S1.2) +
p*(matrix.trace(S1))^2) / (p*(N-2/p)*(matrix.trace(S1.2)-((matrix.trace(S1))^2)/p))))
W.OAS.2 = max(0,min(1,((p-2)*matrix.trace(S2.2)
+ p*(matrix.trace(S2))^2) / (p*(N-2/p)*(matrix.trace(S2.2)-((matrix.trace(S2))^2)/p))))
W.OAS.3 = max(0,min(1,((p-2)*matrix.trace(S3.2)
+ p*(matrix.trace(S3))^2) / (p*(N-2/p)*(matrix.trace(S3.2)-((matrix.trace(S3))^2)/p))))
#####
#Oracle-Approximating Shrinkage method
#####
Sigma.OAS.1 = (1-W.OAS.1)*S1 + W.OAS.1*a1.hat.1*diag(1,p,p)
Sigma.OAS.2 = (1-W.OAS.2)*S2 + W.OAS.2*a1.hat.2*diag(1,p,p)
Sigma.OAS.3 = (1-W.OAS.3)*S3 + W.OAS.3*a1.hat.3*diag(1,p,p)
#####
a2.hat.1 <- (((N-1)^2)/((N-1)*(N+1))) * (matrix.trace(S1.2)/p - (p/(N-1))*(a1.hat.1^2))
a2.hat.2 <- (((N-1)^2)/((N-1)*(N+1))) * (matrix.trace(S2.2)/p - (p/(N-1))*(a1.hat.2^2))
a2.hat.3 <- (((N-1)^2)/((N-1)*(N+1))) * (matrix.trace(S3.2)/p - (p/(N-1))*(a1.hat.3^2))
#####
W.FS.1 = max(0,min(1,(p*a1.hat.1^2 + a2.hat.1)/((N-1)*(a2.hat.1 - a1.hat.1^2)
+ p*a1.hat.1^2 + a2.hat.1)))
W.FS.2 = max(0,min(1,(p*a1.hat.2^2 + a2.hat.2)/((N-1)*(a2.hat.2 - a1.hat.2^2)
+ p*a1.hat.2^2 + a2.hat.2)))
W.FS.3 = max(0,min(1,(p*a1.hat.3^2 + a2.hat.3)/((N-1)*(a2.hat.3 - a1.hat.3^2)
+ p*a1.hat.3^2 + a2.hat.3)))
#####
#Fisher and Sun method
#####
Sigma.FS.1 = (1-W.FS.1)*S1 + W.FS.1*a1.hat.1*diag(1,p,p)
Sigma.FS.2 = (1-W.FS.2)*S2 + W.FS.2*a1.hat.2*diag(1,p,p)
Sigma.FS.3 = (1-W.FS.3)*S3 + W.FS.3*a1.hat.3*diag(1,p,p)
#####
mat.tr.S.Sigma1[i.n,1] <- matrix.trace((S1-Sigma1)%*%t((S1-Sigma1)))
mat.tr.S.Sigma2[i.n,1] <- matrix.trace((S1-Sigma2)%*%t((S1-Sigma2)))
mat.tr.S.Sigma3[i.n,1] <- matrix.trace((S1-Sigma3)%*%t((S1-Sigma3)))
mat.tr.S.Sigma4[i.n,1] <- matrix.trace((S1-Sigma4)%*%t((S1-Sigma4)))
#####
mat.tr.S.Sigma1[i.n,2] <- matrix.trace((S2-Sigma1)%*%t((S2-Sigma1)))
mat.tr.S.Sigma2[i.n,2] <- matrix.trace((S2-Sigma2)%*%t((S2-Sigma2)))
mat.tr.S.Sigma3[i.n,2] <- matrix.trace((S2-Sigma3)%*%t((S2-Sigma3)))
mat.tr.S.Sigma4[i.n,2] <- matrix.trace((S2-Sigma4)%*%t((S2-Sigma4)))
#####

```

```

mat.tr.S.Sigma1[i.n,3] <- matrix.trace((S3-Sigma1)%*%t((S3-Sigma1)))
mat.tr.S.Sigma2[i.n,3] <- matrix.trace((S3-Sigma2)%*%t((S3-Sigma2)))
mat.tr.S.Sigma3[i.n,3] <- matrix.trace((S3-Sigma3)%*%t((S3-Sigma3)))
mat.tr.S.Sigma4[i.n,3] <- matrix.trace((S3-Sigma4)%*%t((S3-Sigma4)))

#####

#####

mat.tr.U1.Sigma1[i.n,1] = matrix.trace((Sigma.U1.1-Sigma1)%*%t(Sigma.U1.1-Sigma1))
mat.tr.U1.Sigma2[i.n,1] = matrix.trace((Sigma.U1.1-Sigma2)%*%t(Sigma.U1.1-Sigma2))
mat.tr.U1.Sigma3[i.n,1] = matrix.trace((Sigma.U1.1-Sigma3)%*%t(Sigma.U1.1-Sigma3))
mat.tr.U1.Sigma4[i.n,1] = matrix.trace((Sigma.U1.1-Sigma4)%*%t(Sigma.U1.1-Sigma4))

#####

mat.tr.U1.Sigma1[i.n,2] = matrix.trace((Sigma.U1.2-Sigma1)%*%t(Sigma.U1.2-Sigma1))
mat.tr.U1.Sigma2[i.n,2] = matrix.trace((Sigma.U1.2-Sigma2)%*%t(Sigma.U1.2-Sigma2))
mat.tr.U1.Sigma3[i.n,2] = matrix.trace((Sigma.U1.2-Sigma3)%*%t(Sigma.U1.2-Sigma3))
mat.tr.U1.Sigma4[i.n,2] = matrix.trace((Sigma.U1.2-Sigma4)%*%t(Sigma.U1.2-Sigma4))

#####

mat.tr.U1.Sigma1[i.n,3] = matrix.trace((Sigma.U1.3-Sigma1)%*%t(Sigma.U1.3-Sigma1))
mat.tr.U1.Sigma2[i.n,3] = matrix.trace((Sigma.U1.3-Sigma2)%*%t(Sigma.U1.3-Sigma2))
mat.tr.U1.Sigma3[i.n,3] = matrix.trace((Sigma.U1.3-Sigma3)%*%t(Sigma.U1.3-Sigma3))
mat.tr.U1.Sigma4[i.n,3] = matrix.trace((Sigma.U1.3-Sigma4)%*%t(Sigma.U1.3-Sigma4))

#####

#####

mat.tr.U2.Sigma1[i.n,1] = matrix.trace((Sigma.U2.1-Sigma1)%*%t(Sigma.U2.1-Sigma1))
mat.tr.U2.Sigma2[i.n,1] = matrix.trace((Sigma.U2.1-Sigma2)%*%t(Sigma.U2.1-Sigma2))
mat.tr.U2.Sigma3[i.n,1] = matrix.trace((Sigma.U2.1-Sigma3)%*%t(Sigma.U2.1-Sigma3))
mat.tr.U2.Sigma4[i.n,1] = matrix.trace((Sigma.U2.1-Sigma4)%*%t(Sigma.U2.1-Sigma4))

#####

mat.tr.U2.Sigma1[i.n,2] = matrix.trace((Sigma.U2.2-Sigma1)%*%t(Sigma.U2.2-Sigma1))
mat.tr.U2.Sigma2[i.n,2] = matrix.trace((Sigma.U2.2-Sigma2)%*%t(Sigma.U2.2-Sigma2))
mat.tr.U2.Sigma3[i.n,2] = matrix.trace((Sigma.U2.2-Sigma3)%*%t(Sigma.U2.2-Sigma3))
mat.tr.U2.Sigma4[i.n,2] = matrix.trace((Sigma.U2.2-Sigma4)%*%t(Sigma.U2.2-Sigma4))

#####

mat.tr.U2.Sigma1[i.n,3] = matrix.trace((Sigma.U2.3-Sigma1)%*%t(Sigma.U2.3-Sigma1))
mat.tr.U2.Sigma2[i.n,3] = matrix.trace((Sigma.U2.3-Sigma2)%*%t(Sigma.U2.3-Sigma2))
mat.tr.U2.Sigma3[i.n,3] = matrix.trace((Sigma.U2.3-Sigma3)%*%t(Sigma.U2.3-Sigma3))
mat.tr.U2.Sigma4[i.n,3] = matrix.trace((Sigma.U2.3-Sigma4)%*%t(Sigma.U2.3-Sigma4))

#####

#####

mat.tr.DS.Sigma1[i.n,1] = matrix.trace((Sigma.DS.1-Sigma1)%*%t(Sigma.DS.1-Sigma1))
mat.tr.DS.Sigma2[i.n,1] = matrix.trace((Sigma.DS.1-Sigma2)%*%t(Sigma.DS.1-Sigma2))
mat.tr.DS.Sigma3[i.n,1] = matrix.trace((Sigma.DS.1-Sigma3)%*%t(Sigma.DS.1-Sigma3))
mat.tr.DS.Sigma4[i.n,1] = matrix.trace((Sigma.DS.1-Sigma4)%*%t(Sigma.DS.1-Sigma4))

#####

mat.tr.DS.Sigma1[i.n,2] = matrix.trace((Sigma.DS.2-Sigma1)%*%t(Sigma.DS.2-Sigma1))
mat.tr.DS.Sigma2[i.n,2] = matrix.trace((Sigma.DS.2-Sigma2)%*%t(Sigma.DS.2-Sigma2))
mat.tr.DS.Sigma3[i.n,2] = matrix.trace((Sigma.DS.2-Sigma3)%*%t(Sigma.DS.2-Sigma3))
mat.tr.DS.Sigma4[i.n,2] = matrix.trace((Sigma.DS.2-Sigma4)%*%t(Sigma.DS.2-Sigma4))

#####

```

```

mat.tr.DS.Sigma1[i.n,3] = matrix.trace((Sigma.DS.3-Sigma1)%*%t(Sigma.DS.3-Sigma1))
mat.tr.DS.Sigma2[i.n,3] = matrix.trace((Sigma.DS.3-Sigma2)%*%t(Sigma.DS.3-Sigma2))
mat.tr.DS.Sigma3[i.n,3] = matrix.trace((Sigma.DS.3-Sigma3)%*%t(Sigma.DS.3-Sigma3))
mat.tr.DS.Sigma4[i.n,3] = matrix.trace((Sigma.DS.3-Sigma4)%*%t(Sigma.DS.3-Sigma4))
#####
#####
mat.tr.LW.Sigma1[i.n,1] = matrix.trace((Sigma.LW.1-Sigma1)%*%t(Sigma.LW.1-Sigma1))
mat.tr.LW.Sigma2[i.n,1] = matrix.trace((Sigma.LW.1-Sigma2)%*%t(Sigma.LW.1-Sigma2))
mat.tr.LW.Sigma3[i.n,1] = matrix.trace((Sigma.LW.1-Sigma3)%*%t(Sigma.LW.1-Sigma3))
mat.tr.LW.Sigma4[i.n,1] = matrix.trace((Sigma.LW.1-Sigma4)%*%t(Sigma.LW.1-Sigma4))
#####
mat.tr.LW.Sigma1[i.n,2] = matrix.trace((Sigma.LW.2-Sigma1)%*%t(Sigma.LW.2-Sigma1))
mat.tr.LW.Sigma2[i.n,2] = matrix.trace((Sigma.LW.2-Sigma2)%*%t(Sigma.LW.2-Sigma2))
mat.tr.LW.Sigma3[i.n,2] = matrix.trace((Sigma.LW.2-Sigma3)%*%t(Sigma.LW.2-Sigma3))
mat.tr.LW.Sigma4[i.n,2] = matrix.trace((Sigma.LW.2-Sigma4)%*%t(Sigma.LW.2-Sigma4))
#####
mat.tr.LW.Sigma1[i.n,3] = matrix.trace((Sigma.LW.3-Sigma1)%*%t(Sigma.LW.3-Sigma1))
mat.tr.LW.Sigma2[i.n,3] = matrix.trace((Sigma.LW.3-Sigma2)%*%t(Sigma.LW.3-Sigma2))
mat.tr.LW.Sigma3[i.n,3] = matrix.trace((Sigma.LW.3-Sigma3)%*%t(Sigma.LW.3-Sigma3))
mat.tr.LW.Sigma4[i.n,3] = matrix.trace((Sigma.LW.3-Sigma4)%*%t(Sigma.LW.3-Sigma4))
#####
#####
mat.tr.RBLW.Sigma1[i.n,1] = matrix.trace((Sigma.RBLW.1-Sigma1)%*%t(Sigma.RBLW.1-Sigma1))
mat.tr.RBLW.Sigma2[i.n,1] = matrix.trace((Sigma.RBLW.1-Sigma2)%*%t(Sigma.RBLW.1-Sigma2))
mat.tr.RBLW.Sigma3[i.n,1] = matrix.trace((Sigma.RBLW.1-Sigma3)%*%t(Sigma.RBLW.1-Sigma3))
mat.tr.RBLW.Sigma4[i.n,1] = matrix.trace((Sigma.RBLW.1-Sigma4)%*%t(Sigma.RBLW.1-Sigma4))
#####
mat.tr.RBLW.Sigma1[i.n,2] = matrix.trace((Sigma.RBLW.2-Sigma1)%*%t(Sigma.RBLW.2-Sigma1))
mat.tr.RBLW.Sigma2[i.n,2] = matrix.trace((Sigma.RBLW.2-Sigma2)%*%t(Sigma.RBLW.2-Sigma2))
mat.tr.RBLW.Sigma3[i.n,2] = matrix.trace((Sigma.RBLW.2-Sigma3)%*%t(Sigma.RBLW.2-Sigma3))
mat.tr.RBLW.Sigma4[i.n,2] = matrix.trace((Sigma.RBLW.2-Sigma4)%*%t(Sigma.RBLW.2-Sigma4))
#####
mat.tr.RBLW.Sigma1[i.n,3] = matrix.trace((Sigma.RBLW.3-Sigma1)%*%t(Sigma.RBLW.3-Sigma1))
mat.tr.RBLW.Sigma2[i.n,3] = matrix.trace((Sigma.RBLW.3-Sigma2)%*%t(Sigma.RBLW.3-Sigma2))
mat.tr.RBLW.Sigma3[i.n,3] = matrix.trace((Sigma.RBLW.3-Sigma3)%*%t(Sigma.RBLW.3-Sigma3))
mat.tr.RBLW.Sigma4[i.n,3] = matrix.trace((Sigma.RBLW.3-Sigma4)%*%t(Sigma.RBLW.3-Sigma4))
#####
#####
mat.tr.OAS.Sigma1[i.n,1] = matrix.trace((Sigma.OAS.1-Sigma1)%*%t(Sigma.OAS.1-Sigma1))
mat.tr.OAS.Sigma2[i.n,1] = matrix.trace((Sigma.OAS.1-Sigma2)%*%t(Sigma.OAS.1-Sigma2))
mat.tr.OAS.Sigma3[i.n,1] = matrix.trace((Sigma.OAS.1-Sigma3)%*%t(Sigma.OAS.1-Sigma3))
mat.tr.OAS.Sigma4[i.n,1] = matrix.trace((Sigma.OAS.1-Sigma4)%*%t(Sigma.OAS.1-Sigma4))
#####
mat.tr.OAS.Sigma1[i.n,2] = matrix.trace((Sigma.OAS.2-Sigma1)%*%t(Sigma.OAS.2-Sigma1))
mat.tr.OAS.Sigma2[i.n,2] = matrix.trace((Sigma.OAS.2-Sigma2)%*%t(Sigma.OAS.2-Sigma2))
mat.tr.OAS.Sigma3[i.n,2] = matrix.trace((Sigma.OAS.2-Sigma3)%*%t(Sigma.OAS.2-Sigma3))
mat.tr.OAS.Sigma4[i.n,2] = matrix.trace((Sigma.OAS.2-Sigma4)%*%t(Sigma.OAS.2-Sigma4))
#####

```

```

mat.tr.OAS.Sigma1[i.n,3] = matrix.trace((Sigma.OAS.3-Sigma1)%*%t(Sigma.OAS.3-Sigma1))
mat.tr.OAS.Sigma2[i.n,3] = matrix.trace((Sigma.OAS.3-Sigma2)%*%t(Sigma.OAS.3-Sigma2))
mat.tr.OAS.Sigma3[i.n,3] = matrix.trace((Sigma.OAS.3-Sigma3)%*%t(Sigma.OAS.3-Sigma3))
mat.tr.OAS.Sigma4[i.n,3] = matrix.trace((Sigma.OAS.3-Sigma4)%*%t(Sigma.OAS.3-Sigma4))
#####
#####
mat.tr.FS.Sigma1[i.n,1] = matrix.trace((Sigma.FS.1-Sigma1)%*%t(Sigma.FS.1-Sigma1))
mat.tr.FS.Sigma2[i.n,1] = matrix.trace((Sigma.FS.1-Sigma2)%*%t(Sigma.FS.1-Sigma2))
mat.tr.FS.Sigma3[i.n,1] = matrix.trace((Sigma.FS.1-Sigma3)%*%t(Sigma.FS.1-Sigma3))
mat.tr.FS.Sigma4[i.n,1] = matrix.trace((Sigma.FS.1-Sigma4)%*%t(Sigma.FS.1-Sigma4))
#####
mat.tr.FS.Sigma1[i.n,2] = matrix.trace((Sigma.FS.2-Sigma1)%*%t(Sigma.FS.2-Sigma1))
mat.tr.FS.Sigma2[i.n,2] = matrix.trace((Sigma.FS.2-Sigma2)%*%t(Sigma.FS.2-Sigma2))
mat.tr.FS.Sigma3[i.n,2] = matrix.trace((Sigma.FS.2-Sigma3)%*%t(Sigma.FS.2-Sigma3))
mat.tr.FS.Sigma4[i.n,2] = matrix.trace((Sigma.FS.2-Sigma4)%*%t(Sigma.FS.2-Sigma4))
#####
mat.tr.FS.Sigma1[i.n,3] = matrix.trace((Sigma.FS.3-Sigma1)%*%t(Sigma.FS.3-Sigma1))
mat.tr.FS.Sigma2[i.n,3] = matrix.trace((Sigma.FS.3-Sigma2)%*%t(Sigma.FS.3-Sigma2))
mat.tr.FS.Sigma3[i.n,3] = matrix.trace((Sigma.FS.3-Sigma3)%*%t(Sigma.FS.3-Sigma3))
mat.tr.FS.Sigma4[i.n,3] = matrix.trace((Sigma.FS.3-Sigma4)%*%t(Sigma.FS.3-Sigma4))

}

#PRIAL
mean.tr.S.Sigma1 <- apply(mat.tr.S.Sigma1,2,mean)
mean.tr.S.Sigma2 <- apply(mat.tr.S.Sigma2,2,mean)
mean.tr.S.Sigma3 <- apply(mat.tr.S.Sigma3,2,mean)
mean.tr.S.Sigma4 <- apply(mat.tr.S.Sigma4,2,mean)
#####
mean.tr.U1.Sigma1 <- apply(mat.tr.U1.Sigma1,2,mean)
mean.tr.U1.Sigma2 <- apply(mat.tr.U1.Sigma2,2,mean)
mean.tr.U1.Sigma3 <- apply(mat.tr.U1.Sigma3,2,mean)
mean.tr.U1.Sigma4 <- apply(mat.tr.U1.Sigma4,2,mean)
#####
mean.tr.U2.Sigma1 <- apply(mat.tr.U2.Sigma1,2,mean)
mean.tr.U2.Sigma2 <- apply(mat.tr.U2.Sigma2,2,mean)
mean.tr.U2.Sigma3 <- apply(mat.tr.U2.Sigma3,2,mean)
mean.tr.U2.Sigma4 <- apply(mat.tr.U2.Sigma4,2,mean)
#####
mean.tr.DS.Sigma1 <- apply(mat.tr.DS.Sigma1,2,mean)
mean.tr.DS.Sigma2 <- apply(mat.tr.DS.Sigma2,2,mean)
mean.tr.DS.Sigma3 <- apply(mat.tr.DS.Sigma3,2,mean)
mean.tr.DS.Sigma4 <- apply(mat.tr.DS.Sigma4,2,mean)
#####
mean.tr.LW.Sigma1 <- apply(mat.tr.LW.Sigma1,2,mean)
mean.tr.LW.Sigma2 <- apply(mat.tr.LW.Sigma2,2,mean)
mean.tr.LW.Sigma3 <- apply(mat.tr.LW.Sigma3,2,mean)
mean.tr.LW.Sigma4 <- apply(mat.tr.LW.Sigma4,2,mean)

```

```
#####
mean.tr.RBLW.Sigma1 <- apply(mat.tr.RBLW.Sigma1,2,mean)
mean.tr.RBLW.Sigma2 <- apply(mat.tr.RBLW.Sigma2,2,mean)
mean.tr.RBLW.Sigma3 <- apply(mat.tr.RBLW.Sigma3,2,mean)
mean.tr.RBLW.Sigma4 <- apply(mat.tr.RBLW.Sigma4,2,mean)
#####
mean.tr.OAS.Sigma1 <- apply(mat.tr.OAS.Sigma1,2,mean)
mean.tr.OAS.Sigma2 <- apply(mat.tr.OAS.Sigma2,2,mean)
mean.tr.OAS.Sigma3 <- apply(mat.tr.OAS.Sigma3,2,mean)
mean.tr.OAS.Sigma4 <- apply(mat.tr.OAS.Sigma4,2,mean)
#####
mean.tr.FS.Sigma1 <- apply(mat.tr.FS.Sigma1,2,mean)
mean.tr.FS.Sigma2 <- apply(mat.tr.FS.Sigma2,2,mean)
mean.tr.FS.Sigma3 <- apply(mat.tr.FS.Sigma3,2,mean)
mean.tr.FS.Sigma4 <- apply(mat.tr.FS.Sigma4,2,mean)
#####
PRIAL.Sigma1.case1.U1 <- 100*(mean.tr.S.Sigma1[1]-mean.tr.U1.Sigma1[1])/mean.tr.S.Sigma1[1]
PRIAL.Sigma1.case2.U1 <- 100*(mean.tr.S.Sigma1[2]-mean.tr.U1.Sigma1[2])/mean.tr.S.Sigma1[2]
PRIAL.Sigma1.case3.U1 <- 100*(mean.tr.S.Sigma1[3]-mean.tr.U1.Sigma1[3])/mean.tr.S.Sigma1[3]
#####
PRIAL.Sigma2.case1.U1 <- 100*(mean.tr.S.Sigma2[1]-mean.tr.U1.Sigma2[1])/mean.tr.S.Sigma2[1]
PRIAL.Sigma2.case2.U1 <- 100*(mean.tr.S.Sigma2[2]-mean.tr.U1.Sigma2[2])/mean.tr.S.Sigma2[2]
PRIAL.Sigma2.case3.U1 <- 100*(mean.tr.S.Sigma2[3]-mean.tr.U1.Sigma2[3])/mean.tr.S.Sigma2[3]
#####
PRIAL.Sigma3.case1.U1 <- 100*(mean.tr.S.Sigma3[1]-mean.tr.U1.Sigma3[1])/mean.tr.S.Sigma3[1]
PRIAL.Sigma3.case2.U1 <- 100*(mean.tr.S.Sigma3[2]-mean.tr.U1.Sigma3[2])/mean.tr.S.Sigma3[2]
PRIAL.Sigma3.case3.U1 <- 100*(mean.tr.S.Sigma3[3]-mean.tr.U1.Sigma3[3])/mean.tr.S.Sigma3[3]
#####
PRIAL.Sigma4.case1.U1 <- 100*(mean.tr.S.Sigma4[1]-mean.tr.U1.Sigma4[1])/mean.tr.S.Sigma4[1]
PRIAL.Sigma4.case2.U1 <- 100*(mean.tr.S.Sigma4[2]-mean.tr.U1.Sigma4[2])/mean.tr.S.Sigma4[2]
PRIAL.Sigma4.case3.U1 <- 100*(mean.tr.S.Sigma4[3]-mean.tr.U1.Sigma4[3])/mean.tr.S.Sigma4[3]
#####
PRIAL.Sigma1.case1.U2 <- 100*(mean.tr.S.Sigma1[1]-mean.tr.U2.Sigma1[1])/mean.tr.S.Sigma1[1]
PRIAL.Sigma1.case2.U2 <- 100*(mean.tr.S.Sigma1[2]-mean.tr.U2.Sigma1[2])/mean.tr.S.Sigma1[2]
PRIAL.Sigma1.case3.U2 <- 100*(mean.tr.S.Sigma1[3]-mean.tr.U2.Sigma1[3])/mean.tr.S.Sigma1[3]
#####
PRIAL.Sigma2.case1.U2 <- 100*(mean.tr.S.Sigma2[1]-mean.tr.U2.Sigma2[1])/mean.tr.S.Sigma2[1]
PRIAL.Sigma2.case2.U2 <- 100*(mean.tr.S.Sigma2[2]-mean.tr.U2.Sigma2[2])/mean.tr.S.Sigma2[2]
PRIAL.Sigma2.case3.U2 <- 100*(mean.tr.S.Sigma2[3]-mean.tr.U2.Sigma2[3])/mean.tr.S.Sigma2[3]
#####
PRIAL.Sigma3.case1.U2 <- 100*(mean.tr.S.Sigma3[1]-mean.tr.U2.Sigma3[1])/mean.tr.S.Sigma3[1]
PRIAL.Sigma3.case2.U2 <- 100*(mean.tr.S.Sigma3[2]-mean.tr.U2.Sigma3[2])/mean.tr.S.Sigma3[2]
PRIAL.Sigma3.case3.U2 <- 100*(mean.tr.S.Sigma3[3]-mean.tr.U2.Sigma3[3])/mean.tr.S.Sigma3[3]
#####
PRIAL.Sigma4.case1.U2 <- 100*(mean.tr.S.Sigma4[1]-mean.tr.U2.Sigma4[1])/mean.tr.S.Sigma4[1]
PRIAL.Sigma4.case2.U2 <- 100*(mean.tr.S.Sigma4[2]-mean.tr.U2.Sigma4[2])/mean.tr.S.Sigma4[2]
PRIAL.Sigma4.case3.U2 <- 100*(mean.tr.S.Sigma4[3]-mean.tr.U2.Sigma4[3])/mean.tr.S.Sigma4[3]
#####
```

```

PRIAL.Sigma1.case1.DS <- 100*(mean.tr.S.Sigma1[1]-mean.tr.DS.Sigma1[1])/mean.tr.S.Sigma1[1]
PRIAL.Sigma1.case2.DS <- 100*(mean.tr.S.Sigma1[2]-mean.tr.DS.Sigma1[2])/mean.tr.S.Sigma1[2]
PRIAL.Sigma1.case3.DS <- 100*(mean.tr.S.Sigma1[3]-mean.tr.DS.Sigma1[3])/mean.tr.S.Sigma1[3]
#####
PRIAL.Sigma2.case1.DS <- 100*(mean.tr.S.Sigma2[1]-mean.tr.DS.Sigma2[1])/mean.tr.S.Sigma2[1]
PRIAL.Sigma2.case2.DS <- 100*(mean.tr.S.Sigma2[2]-mean.tr.DS.Sigma2[2])/mean.tr.S.Sigma2[2]
PRIAL.Sigma2.case3.DS <- 100*(mean.tr.S.Sigma2[3]-mean.tr.DS.Sigma2[3])/mean.tr.S.Sigma2[3]
#####
PRIAL.Sigma3.case1.DS <- 100*(mean.tr.S.Sigma3[1]-mean.tr.DS.Sigma3[1])/mean.tr.S.Sigma3[1]
PRIAL.Sigma3.case2.DS <- 100*(mean.tr.S.Sigma3[2]-mean.tr.DS.Sigma3[2])/mean.tr.S.Sigma3[2]
PRIAL.Sigma3.case3.DS <- 100*(mean.tr.S.Sigma3[3]-mean.tr.DS.Sigma3[3])/mean.tr.S.Sigma3[3]
#####
PRIAL.Sigma4.case1.DS <- 100*(mean.tr.S.Sigma4[1]-mean.tr.DS.Sigma4[1])/mean.tr.S.Sigma4[1]
PRIAL.Sigma4.case2.DS <- 100*(mean.tr.S.Sigma4[2]-mean.tr.DS.Sigma4[2])/mean.tr.S.Sigma4[2]
PRIAL.Sigma4.case3.DS <- 100*(mean.tr.S.Sigma4[3]-mean.tr.DS.Sigma4[3])/mean.tr.S.Sigma4[3]
#####
#####
PRIAL.Sigma1.case1.LW <- 100*(mean.tr.S.Sigma1[1]-mean.tr.LW.Sigma1[1])/mean.tr.S.Sigma1[1]
PRIAL.Sigma1.case2.LW <- 100*(mean.tr.S.Sigma1[2]-mean.tr.LW.Sigma1[2])/mean.tr.S.Sigma1[2]
PRIAL.Sigma1.case3.LW <- 100*(mean.tr.S.Sigma1[3]-mean.tr.LW.Sigma1[3])/mean.tr.S.Sigma1[3]
#####
PRIAL.Sigma2.case1.LW <- 100*(mean.tr.S.Sigma2[1]-mean.tr.LW.Sigma2[1])/mean.tr.S.Sigma2[1]
PRIAL.Sigma2.case2.LW <- 100*(mean.tr.S.Sigma2[2]-mean.tr.LW.Sigma2[2])/mean.tr.S.Sigma2[2]
PRIAL.Sigma2.case3.LW <- 100*(mean.tr.S.Sigma2[3]-mean.tr.LW.Sigma2[3])/mean.tr.S.Sigma2[3]
#####
PRIAL.Sigma3.case1.LW <- 100*(mean.tr.S.Sigma3[1]-mean.tr.LW.Sigma3[1])/mean.tr.S.Sigma3[1]
PRIAL.Sigma3.case2.LW <- 100*(mean.tr.S.Sigma3[2]-mean.tr.LW.Sigma3[2])/mean.tr.S.Sigma3[2]
PRIAL.Sigma3.case3.LW <- 100*(mean.tr.S.Sigma3[3]-mean.tr.LW.Sigma3[3])/mean.tr.S.Sigma3[3]
#####
PRIAL.Sigma4.case1.LW <- 100*(mean.tr.S.Sigma4[1]-mean.tr.LW.Sigma4[1])/mean.tr.S.Sigma4[1]
PRIAL.Sigma4.case2.LW <- 100*(mean.tr.S.Sigma4[2]-mean.tr.LW.Sigma4[2])/mean.tr.S.Sigma4[2]
PRIAL.Sigma4.case3.LW <- 100*(mean.tr.S.Sigma4[3]-mean.tr.LW.Sigma4[3])/mean.tr.S.Sigma4[3]
#####
#####
PRIAL.Sigma1.case1.RBLW <- 100*(mean.tr.S.Sigma1[1]-mean.tr.RBLW.Sigma1[1])/mean.tr.S.Sigma1[1]
PRIAL.Sigma1.case2.RBLW <- 100*(mean.tr.S.Sigma1[2]-mean.tr.RBLW.Sigma1[2])/mean.tr.S.Sigma1[2]
PRIAL.Sigma1.case3.RBLW <- 100*(mean.tr.S.Sigma1[3]-mean.tr.RBLW.Sigma1[3])/mean.tr.S.Sigma1[3]
#####
PRIAL.Sigma2.case1.RBLW <- 100*(mean.tr.S.Sigma2[1]-mean.tr.RBLW.Sigma2[1])/mean.tr.S.Sigma2[1]
PRIAL.Sigma2.case2.RBLW <- 100*(mean.tr.S.Sigma2[2]-mean.tr.RBLW.Sigma2[2])/mean.tr.S.Sigma2[2]
PRIAL.Sigma2.case3.RBLW <- 100*(mean.tr.S.Sigma2[3]-mean.tr.RBLW.Sigma2[3])/mean.tr.S.Sigma2[3]
#####
PRIAL.Sigma3.case1.RBLW <- 100*(mean.tr.S.Sigma3[1]-mean.tr.RBLW.Sigma3[1])/mean.tr.S.Sigma3[1]
PRIAL.Sigma3.case2.RBLW <- 100*(mean.tr.S.Sigma3[2]-mean.tr.RBLW.Sigma3[2])/mean.tr.S.Sigma3[2]
PRIAL.Sigma3.case3.RBLW <- 100*(mean.tr.S.Sigma3[3]-mean.tr.RBLW.Sigma3[3])/mean.tr.S.Sigma3[3]
#####
PRIAL.Sigma4.case1.RBLW <- 100*(mean.tr.S.Sigma4[1]-mean.tr.RBLW.Sigma4[1])/mean.tr.S.Sigma4[1]
PRIAL.Sigma4.case2.RBLW <- 100*(mean.tr.S.Sigma4[2]-mean.tr.RBLW.Sigma4[2])/mean.tr.S.Sigma4[2]

```

```

PRIAL.Sigma4.case3.RBLW <- 100*(mean.tr.S.Sigma4[3]-mean.tr.RBLW.Sigma4[3])/mean.tr.S.Sigma4[3]
#####
#####
PRIAL.Sigma1.case1.OAS <- 100*(mean.tr.S.Sigma1[1]-mean.tr.OAS.Sigma1[1])/mean.tr.S.Sigma1[1]
PRIAL.Sigma1.case2.OAS <- 100*(mean.tr.S.Sigma1[2]-mean.tr.OAS.Sigma1[2])/mean.tr.S.Sigma1[2]
PRIAL.Sigma1.case3.OAS <- 100*(mean.tr.S.Sigma1[3]-mean.tr.OAS.Sigma1[3])/mean.tr.S.Sigma1[3]
#####
PRIAL.Sigma2.case1.OAS <- 100*(mean.tr.S.Sigma2[1]-mean.tr.OAS.Sigma2[1])/mean.tr.S.Sigma2[1]
PRIAL.Sigma2.case2.OAS <- 100*(mean.tr.S.Sigma2[2]-mean.tr.OAS.Sigma2[2])/mean.tr.S.Sigma2[2]
PRIAL.Sigma2.case3.OAS <- 100*(mean.tr.S.Sigma2[3]-mean.tr.OAS.Sigma2[3])/mean.tr.S.Sigma2[3]
#####
PRIAL.Sigma3.case1.OAS <- 100*(mean.tr.S.Sigma3[1]-mean.tr.OAS.Sigma3[1])/mean.tr.S.Sigma3[1]
PRIAL.Sigma3.case2.OAS <- 100*(mean.tr.S.Sigma3[2]-mean.tr.OAS.Sigma3[2])/mean.tr.S.Sigma3[2]
PRIAL.Sigma3.case3.OAS <- 100*(mean.tr.S.Sigma3[3]-mean.tr.OAS.Sigma3[3])/mean.tr.S.Sigma3[3]
#####
PRIAL.Sigma4.case1.OAS <- 100*(mean.tr.S.Sigma4[1]-mean.tr.OAS.Sigma4[1])/mean.tr.S.Sigma4[1]
PRIAL.Sigma4.case2.OAS <- 100*(mean.tr.S.Sigma4[2]-mean.tr.OAS.Sigma4[2])/mean.tr.S.Sigma4[2]
PRIAL.Sigma4.case3.OAS <- 100*(mean.tr.S.Sigma4[3]-mean.tr.OAS.Sigma4[3])/mean.tr.S.Sigma4[3]
#####
#####
PRIAL.Sigma1.case1.FS <- 100*(mean.tr.S.Sigma1[1]-mean.tr.FS.Sigma1[1])/mean.tr.S.Sigma1[1]
PRIAL.Sigma1.case2.FS <- 100*(mean.tr.S.Sigma1[2]-mean.tr.FS.Sigma1[2])/mean.tr.S.Sigma1[2]
PRIAL.Sigma1.case3.FS <- 100*(mean.tr.S.Sigma1[3]-mean.tr.FS.Sigma1[3])/mean.tr.S.Sigma1[3]
#####
PRIAL.Sigma2.case1.FS <- 100*(mean.tr.S.Sigma2[1]-mean.tr.FS.Sigma2[1])/mean.tr.S.Sigma2[1]
PRIAL.Sigma2.case2.FS <- 100*(mean.tr.S.Sigma2[2]-mean.tr.FS.Sigma2[2])/mean.tr.S.Sigma2[2]
PRIAL.Sigma2.case3.FS <- 100*(mean.tr.S.Sigma2[3]-mean.tr.FS.Sigma2[3])/mean.tr.S.Sigma2[3]
#####
PRIAL.Sigma3.case1.FS <- 100*(mean.tr.S.Sigma3[1]-mean.tr.FS.Sigma3[1])/mean.tr.S.Sigma3[1]
PRIAL.Sigma3.case2.FS <- 100*(mean.tr.S.Sigma3[2]-mean.tr.FS.Sigma3[2])/mean.tr.S.Sigma3[2]
PRIAL.Sigma3.case3.FS <- 100*(mean.tr.S.Sigma3[3]-mean.tr.FS.Sigma3[3])/mean.tr.S.Sigma3[3]
#####
PRIAL.Sigma4.case1.FS <- 100*(mean.tr.S.Sigma4[1]-mean.tr.FS.Sigma4[1])/mean.tr.S.Sigma4[1]
PRIAL.Sigma4.case2.FS <- 100*(mean.tr.S.Sigma4[2]-mean.tr.FS.Sigma4[2])/mean.tr.S.Sigma4[2]
PRIAL.Sigma4.case3.FS <- 100*(mean.tr.S.Sigma4[3]-mean.tr.FS.Sigma4[3])/mean.tr.S.Sigma4[3]

result.Sigma1.case1 =
  c(PRIAL.Sigma1.case1.U1,PRIAL.Sigma1.case1.U2,PRIAL.Sigma1.case1.DS,PRIAL.Sigma1.case1.LW,
    ...PRIAL.Sigma1.case1.RBLW,PRIAL.Sigma1.case1.OAS,PRIAL.Sigma1.case1.FS)
result.Sigma1.case2 =
  c(PRIAL.Sigma1.case2.U1,PRIAL.Sigma1.case2.U2,PRIAL.Sigma1.case2.DS,PRIAL.Sigma1.case2.LW,
    ...PRIAL.Sigma1.case2.RBLW,PRIAL.Sigma1.case2.OAS,PRIAL.Sigma1.case2.FS)
result.Sigma1.case3 =
  c(PRIAL.Sigma1.case3.U1,PRIAL.Sigma1.case3.U2,PRIAL.Sigma1.case3.DS,PRIAL.Sigma1.case3.LW,
    ...PRIAL.Sigma1.case3.RBLW,PRIAL.Sigma1.case3.OAS,PRIAL.Sigma1.case3.FS)

result.Sigma2.case1 =
  c(PRIAL.Sigma2.case1.U1,PRIAL.Sigma2.case1.U2,PRIAL.Sigma2.case1.DS,PRIAL.Sigma2.case1.LW,

```

```
...PRIAL.Sigma2.case1.RBLW,PRIAL.Sigma2.case1.OAS,PRIAL.Sigma2.case1.FS)
result.Sigma2.case2 =
c(PRIAL.Sigma2.case2.U1,PRIAL.Sigma2.case2.U2,PRIAL.Sigma2.case2.DS,PRIAL.Sigma2.case2.LW,
...PRIAL.Sigma2.case2.RBLW,PRIAL.Sigma2.case2.OAS,PRIAL.Sigma2.case2.FS)
result.Sigma2.case3 =
c(PRIAL.Sigma2.case3.U1,PRIAL.Sigma2.case3.U2,PRIAL.Sigma2.case3.DS,PRIAL.Sigma2.case3.LW,
...PRIAL.Sigma2.case3.RBLW,PRIAL.Sigma2.case3.OAS,PRIAL.Sigma2.case3.FS)

result.Sigma3.case1 =
c(PRIAL.Sigma3.case1.U1,PRIAL.Sigma3.case1.U2,PRIAL.Sigma3.case1.DS,PRIAL.Sigma3.case1.LW,
...PRIAL.Sigma3.case1.RBLW,PRIAL.Sigma3.case1.OAS,PRIAL.Sigma3.case1.FS)
result.Sigma3.case2 =
c(PRIAL.Sigma3.case2.U1,PRIAL.Sigma3.case2.U2,PRIAL.Sigma3.case2.DS,PRIAL.Sigma3.case2.LW,
...PRIAL.Sigma3.case2.RBLW,PRIAL.Sigma3.case2.OAS,PRIAL.Sigma3.case2.FS)
result.Sigma3.case3 =
c(PRIAL.Sigma3.case3.U1,PRIAL.Sigma3.case3.U2,PRIAL.Sigma3.case3.DS,PRIAL.Sigma3.case3.LW,
...PRIAL.Sigma3.case3.RBLW,PRIAL.Sigma3.case3.OAS,PRIAL.Sigma3.case3.FS)

result.Sigma4.case1 =
c(PRIAL.Sigma4.case1.U1,PRIAL.Sigma4.case1.U2,PRIAL.Sigma4.case1.DS,PRIAL.Sigma4.case1.LW,
...PRIAL.Sigma4.case1.RBLW,PRIAL.Sigma4.case1.OAS,PRIAL.Sigma4.case1.FS)
result.Sigma4.case2 =
c(PRIAL.Sigma4.case2.U1,PRIAL.Sigma4.case2.U2,PRIAL.Sigma4.case2.DS,PRIAL.Sigma4.case2.LW,
...PRIAL.Sigma4.case2.RBLW,PRIAL.Sigma4.case2.OAS,PRIAL.Sigma4.case2.FS)
result.Sigma4.case3 =
c(PRIAL.Sigma4.case3.U1,PRIAL.Sigma4.case3.U2,PRIAL.Sigma4.case3.DS,PRIAL.Sigma4.case3.LW,
...PRIAL.Sigma4.case3.RBLW,PRIAL.Sigma4.case3.OAS,PRIAL.Sigma4.case3.FS)

Result = rbind(result.Sigma1.case1,result.Sigma1.case2,result.Sigma1.case3,
result.Sigma2.case1,result.Sigma2.case2,result.Sigma2.case3,
result.Sigma3.case1,result.Sigma3.case2,result.Sigma3.case3,
result.Sigma4.case1,result.Sigma4.case2,result.Sigma4.case3)
colnames(Result) <- c("U1","U2","DS","LW","RBLW","OAS","FS")
Result
```

مراجع

- [۱] آشوری، ا. (۱۳۹۳). برآوردگر انقباضی ماتریس کوورایانس توزیع نرمال چندمتغیره با بعد بالا، پایان نامه کارشناسی ارشد، دانشگاه صنعتی شاهرود.
- [۲] رودین، و. (۱۳۷۳) اصول آنالیز ریاضی، ترجمه عالمزاده، ع.ا.، علمی و فنی، تهران.
- [۳] کرمی کبیر، ح. (۱۳۹۲). برآورد انقباضی در مدل های محدودشده، پایان نامه کارشناسی ارشد، دانشگاه صنعتی شاهرود.
- [4] Abramsky, S., Jung, A., (1994). Domain theory, in: Handbook of Logic in Computer Science. Clarendon Press, Oxford. Vol. 3. pp. 1-68.
- [5] Bai, J., Shi, S., (2011). Estimating high dimensional covariance matrices and its applications. *Ann. Econ. Finance* **12**, 199-215.
- [6] Bickel, P., Levina, E., (2008). Covariance regularization by thresholding. *Ann. Statist.*, **36**, 2577-2604.
- [7] Birkhoff, G., (1967). Lattice Theory. *Amer. Math. Soc. 3rd Edition, AMS Colloq. Publication*, Vol. 25.
- [8] Cai, T., Liu, W., (2011). Adaptive threshold for sparse covariance matrix estimation. *J. Amer. Statist. Assoc.* **106**, 672-684.
- [9] Cai, T., Zhou, H., (2010). Optimal rates of convergence for sparse covariance matrix estimation. *Manuscript. University of Pennsylvania*.
- [10] Cao, G., and Bouman, C., (2009). Covariance estimation for high dimensional data vectors using the sparse matrix transform, in *Advances in Neural Information Processing Systems 21*, editors: Koller, D., Schuurmans, D., Bengio, L., and Bottou, L., 225-232.
- [11] Chaudhuri, S., Drton, M., Richardson, T. S. (2007). Estimation of covariance matrix with Zeros, *Biometrika*, **94**(1), 199-216.

-
- [12] Chen, Z. and Dunson, D., (2003). Random effects selection in linear mixed models. *Biometrics*, **59**, 762-769.
 - [13] Chen, S.X., Qin, Y.L., (2010). A two-sample test for high-dimensional data with applications to gene-set testing, *Ann. Statist.* **38**, 808-835.
 - [14] Chen, Y., Wiesel, A., Eldar, C.Y., Hero, A.O., (2010). Shrinkage algorithms for MMSE covariance estimation. *IEEE Trans. Signal Process.* **58**, 5016-5029.
 - [15] Chen, S.X., Zhang, L.X., Zhong, P.S., (2010). Tests for high-dimensional covariance matrices. *J. Amer. Statist. Assoc.* **105**, 810-819.
 - [16] Corbeil, R.R. and Searle, S.R., (1976). A comparison of variance component estimators, *Biometric*. **32**(4), 779-791.
 - [17] Daniels, M.J., (2005) Shrinkage priors for the dependence structure in longitudinal data. *J. Statist. Plan. Inference.* **127**, 119-130.
 - [18] Daniels, M.J., (2006). Bayesian modeling of several covariance matrices and some results on the propriety of the posterior for linear regression with correlated and/or heterogeneous errors. *J. Mult. Anal.* **97**, 1185-1207.
 - [19] Daniels, M.J. and Kass, R., (2001). Shrinkage estimators for a covariance matrix. *Biometrics*. **57**, 1173-1184.
 - [20] Daniels, M. J., and Pourahmadi, M., (2002). Bayesian analysis of covariance matrices and dynamic models of longitudinal data, *Biometrika*. **89**(3), 553-566.
 - [21] Dey, D.K. and Srinivasan, C., (1985). Estimation of a covariance matrix under Stein' loss, *Ann. Statist.* **4**, 1581-1591.
 - [22] Edwards, D.A., (1978). On the existence of probability measures with given marginals. *Ann. Inst. Fourier, Grenoble*. **28**, 53-78.
 - [23] Efron, B., and Morris, C., (1975). Data analysis using Stein's estimator and its generalizations, *J. Am. Statist. Assoc.* **70** (350), 311-319.
 - [24] Efron, B., and Morris, C., (1976). Multivariate empirical Bayes and estimation of covariance matrices, *Ann. Statist.* **4** (1), 22-32.
 - [25] Fan, J. and Lv, J. (2008). Sure independence screening for ultrahigh dimensional feature space, *J. Royal. Statist. Soc. Series B.* **70**(5), 849-911.

- [26] Fan, J., and Wu, Y., (2009). Semiparametric estimation of covariance matrices for longitudinal data, *J. Am. Statist. Assoc.*, **103**(484), 1520-1533.
- [27] Fraley, C. and Burns, P. J., (1995). Large-Scale Estimation of Variance and Covariance Components. *SIAM J. Sci. Comput.*, **16**(1), 192-209.
- [28] Fisher, T. j and Sun, X., (2011), Improved Stein-type shrinkage estimators for the high-dimensional multivariate normal covariance matrix. *Comp. Statist. Data Anal.*, **55**, 1909-1918.
- [29] Fujikoshi, Y., Himeno, T., Wakaki, H., (2004). Asymptotic results of a high dimensional MANOVA test and power comparison when the dimension is large compared to the sample size. *J. Japan Statist. Soc.* **34**, 19-26.
- [30] Ghosh, M., Lahiri, P., (1987). Robust empirical Bayes estimation of means from stratified samples. *J. Amer. Statist. Assoc.* **82**, 1153-1162.
- [31] Haff, L.R., (1979). Estimation of the inverse covariance matrix: random mixtures of the inverse Wishart matrix and the identity, *Ann. Statist.*, **7**(6), 1264-1276.
- [32] Haff, L.R., (1981). Corrections to Estimation of the inverse covariance matrix: random mixtures of the inverse Wishart matrix and the identity, *Ann. Statist.* **9**(5), 11-32.
- [33] Haff, L.R., (1991). The variational form of certain Bayes estimators. *Ann. Statist.*, **19** 1163-1190.
- [34] Hannart, A., Naveau, P., (2014). Estimating high dimensional covariance matrices: A new look at the Gaussian conjugate framework. *J. Mult. Anal.* **131**, 149-162.
- [35] Hartigan, J.A., (1969). Linear Bayes methods. *J. Royal. Statist. Soc. Ser. B* **31**, 446-454.
- [36] Hemmerle, W. J. and Hartley, H. O. (1973). Computing maximum likelihood estimates for the mixed A.O.V. model using the W transformation, *Technometric*, **15**(4), 819-831.
- [37] Himeno, T., Yamada, T., (2014). Estimation for some functions of covariance matrix in high dimension under non-normality and its applications. *J. Mult. Anal.* **130**, 27-44.
- [38] Horn, A., Tarski, A., (1948). Measures on Boolean algebras. *Trans. Amer. Math. Soc.*, **64**, 467-497.

-
- [39] Ikeda, Y., Kubokawa, T., Srivastava, M. S., (2016). Comparison of linear shrinkage estimators of a large covariance matrix in normal and non-normal distributions, *Comp. Statist. and Data Anal.*, **95**, 95-108.
- [40] James, W. and Stein, C., (1961). Estimation with quadratic loss. *In the Fourth Berkeley Symposium on Mathematical Statistics and Probability, CA: University of California Press.* pp. 361-379.
- [41] Johnstone, I.M., (2001). On the distribution of the largest eigenvalue on principle component analysis, *Ann. Statist.*, **29**(2), 295-357.
- [42] Johnstone, I. M., and Lu, A. Y., (2009). On consistency and sparsity for principle components analysis in high dimensions, *J. Amer. Statist. Assoc.*, **104**, 682-693.
- [43] Konno, Y., (2009). Shrinkage estimators for large covariance matrices in multivariate real and complex normal distributions under an invariant quadratic loss. *J. Mult. Anal.* **100**, 2237-2253.
- [44] Lam, C., Fan, J., (2009). Sparsistency and rates of convergence in large covariance matrix estimation. *Ann. Statist.*, **37**, 4254-4278.
- [45] Lang, S., (1978). *Linear algebra*, Springer.
- [46] Ledoit, O., Wolf, M., (2003), Improved estimation of the covariance matrix of stock returns with an application to portfolio selection. *J. Empir. Finance.*, **10**, 603-621.
- [47] Ledoit, O., Wolf, M., (2004). A well-conditioned estimator for large-dimensional covariance matrices. *J. Mult. Anal.* **88**, 365-411.
- [48] Ledoit, O., Wolf, M., (2012). Nonlinear shrinkage estimation of large-dimensional covariance matrices. *Ann. Statist.* **40**, 1024-1060.
- [49] Ledoit, O., Wolf, M., (2015). Spectrum estimation: A unified framework for covariance matrix estimation and PCA in large dimensions. *J. Mult. Anal.* **139**, 360-384.
- [50] Magnus, J. R., Neudecker, H., (1979). The commutation matrix: some properties and applications. *Ann. Statist.* **7**, 381-894.
- [51] Muirhead, R.J., (1982), *Aspects Stat. Theory*. John Wiley, New York.
- [52] Pourahmadi, M., (1999). Joint mean-covariance models with applications to longitudinal data: unconstrained parameterisation, *Biometrika*, **86**(3), 677-690.

- [53] Roeverto, A., (2000) Cholesky decomposition of hyper invers Wishart matri, *Biometrika*, **87**(1), 99-112.
- [54] Robbins, H., (1983) Some thoughts on empirical Bayes estimation. *Ann. Statist.* **11**, 713-723.
- [55] Rothman, A., Levina, E., Zhu, J., (2009). Generalized thresholding of large covariance matrices. *J. Amer. Statist. Assoc.*, **104**, 177-186.
- [56] Saleh, A. K. M. E., (2006). Theory of Preliminary Test and Stein-Type Estimation with Applications, *John Wiley and Sons*.
- [57] Schafer, J., Strimmer, K., (2005). A shrinkage approach to large-scale covariance matrix estimation and implications for functional genomics. *Statist. Appl. Genet. Molec. Biol.*
- [58] Schott, J. R., (2007). Some high-dimensional tests for a one-way MANOVA. *J. Mult. Anal.*, **98**, 1825-1839.
- [59] Smith, M., and Kohn, R., (2002), Parsimonious Covariance Matrix Estimation for Longitudinal Data, *J. Amer. Statist. Assoc.*, **97**, 1141-1153.
- [60] Srivastava, M. S., (2005). Some tests concerning the covariance matrix in high dimensional data. *J. Japan Statist. Soc.* **35**, 251-272.
- [61] Srivastava, M. S., Fujikoshi, Y., (2006). Multivariate analysis of variance with fewer observations than the dimension. *J. Mult. Anal.* **97**, 1927-1940.
- [62] Srivastava, M. S., Kollo, T., Rosen, D. V., (2011). Some tests for the covariance matrix with fewer observations than the dimension under non-normality. *J. Mult. Anal.* **102**, 1090-1103.
- [63] Srivastava, M.S., Kubokawa, T., (2007). Comparison of discrimination methods for high dimensional data. *J. Japan Statist. Soc.* **37**, 123-134.
- [64] Srivastava, M.S., Yanagihara, H., Kubokawa, T., (2014). Tests for covariance matrices in high dimension with less sample size. *J. Mult. Anal.* **130**, 289-309.
- [65] Stein, C., (1956), Inadmissibility of the usual estimator for the mean of a multivariate normal distribution. *In Proceedings of the Third Berkeley Symposium.*, **1**, 197-206.
- [66] Touloumis, A., (2015). Nonparametric Stein-type shrinkage covariance matrix estimators in high-dimensional setting. *Comp. Statist. Data Anal.* **83**, 251-261.

-
- [67] Wu, W. B., and Pourahmadi, M., (2003). Nonparametric estimation of large covariance matrices of longitudinal data, *Biometrika*, **90**(4), 831-844.
- [68] Yang, R. and Berger, J. O., (1994). Estimation of a covariance matrix using the reference prior, *Ann. Statist.*, **22** (3), 1195-1211.

Abstract

In multivariate analysis, one of the most challenging cases of estimating a covariance matrix of a random variable is when the number of variables (dimension) is large. Also, in some multivariate discussions, it is necessary to estimate the function of a covariance matrix. Suppose p represents the dimension. In this thesis, we examine the covariance matrix estimator and some functions of it for the high dimension ($p \rightarrow \infty$). In this regard, we assume that the population of sampling has normal and non-normal distribution and introduce non-parametric shrinkage estimators of the covariance matrix in the high dimension and examine some of their properties. Also, using simulation examples, we evaluate the efficiency of the proposed shrinkage estimators for the covariance matrix in the high dimension.

Keywords

Covariance Matrix, Shrinkage Estimator, High Dimension, Consistent of Nonparametric Estimator.



Shahrood University of Technology

Faculty of Mathematical Sciences
MSc Thesis in: Mathematical Statistics

Shrinkage Covariance Matrix Estimators in High-Dimensional Settings

By: Zahra Hosseinpour

Supervisor
Dr. Mohammad Arashi

July 2018