

الحمد لله
الرحمن الرحيم



دانشگاه صنعتی شاهرود

دانشکده علوم ریاضی

رشته آمار، گرایش احتمال

رساله دکتری

برآورد پارامتر تنش - مقاومت براساس نمونه مجموعه‌ی رتبه دار رکوردی

نگارنده: امینه صادق‌پور

استاد راهنما

دکتر احمد نزاکتی

استاد مشاور

دکتر سیدمهدی صالحی

شهریور ۱۳۹۹

تقدیم به:

پدرم و مادرم، که از نگاهشان صلابت، از رفتارشان محبت و از صبرشان ایستادگی را آموختم.
همسرم، اسطوره زندگیم، پناه خستگیم و امید بودنم.
فرزند دلبندم، امید بخش جانم که خستگی های این راه را به امید و روشنی تبدیل کرده.
خانواده عزیزم، مهربان فرشتگانی که محطات زیبا و ناب زندگانیم می یون حضور سبز آن هاست.

سپاس‌گزاری

پیش از همه و بیش از همه شکر و سپاس خدایی را که توفیق را رفیق راهم ساخت تا انجام این تحقیق را به پایان رسانم. بر خود لازم می‌دانم از تمام عزیزانی که پیمودن این مسیر را برایم آسان‌تر کردند تشکر و قدردانی کنم. از اساتید بزرگوارم، جناب آقای دکتر احمد نزاکتی و آقای دکتر سید مهدی صالحی که در این مسیر با نهایت بردباری و مهربانی مرا از راهنمایی‌های ارزنده خویش بهره‌مند فرمودند بی‌نهایت سپاس‌گزارم. از اساتید محترم گروه آمار، که در طی این دوره از تحصیلم از علم و راهنمایی‌هایشان بهره بردم، متواضعانه سپاس‌گزارم.

از اساتید محترم، آقای دکتر مهدی مهدی زاده (دانشگاه حکیم سبزواری)، جناب آقای دکتر محمدرضا ربیعی و جناب آقای دکتر حسین باغیشنی (دانشگاه صنعتی شاهرود)، که زحمت تصحیح و داوری این رساله را متقبل شدند، سپاس‌گزارم.

از ریاست محترم دانشکده ریاضی جناب آقای دکتر مهدی قوتمند و مدیر محترم گروه آمار دانشگاه صنعتی شاهرود، جناب آقای دکتر مهرداد غزنوی به دلیل خدمات، راهنمایی‌ها و زحمات‌شان سپاس‌گزارم. از دوستان و همکلاسی‌های خوبم که مرا در رسیدن به اهدافم کمک کرده‌اند، تشکر کرده و برایشان آرزوی موفقیت و سربلندی می‌کنم.

در نهایت، از همه عزیزانی که تا بدین‌جا مرا یاری کرده‌اند، به‌ویژه خانواده مهربانم که همواره پشتیبان من بوده‌اند، سپاس‌گزارم. امیدوارم بتوانم از عهده ادای حق این عزیزان برآیم.

یقیناً این مجموعه بدون عیب و نقص نیست. لذا از خوانندگان عزیز تقاضا دارم عیب‌ها و کاستی‌های این رساله را از طریق پست الکترونیکی^۱ به اطلاع اینجانب برسانند.

امینه صادق‌پور
شهریور ۱۳۹۹

^۱ sadeghpour.amineh@yahoo.com

تعهد نامه

اینجانب امینه صادق پور دانشجوی دکتری رشته آمار علوم ریاضی دانشگاه صنعتی شاهرود، نویسنده پایان نامه با عنوان برآورد پارامتر تنش-مقاومت براساس نمونه مجموعه‌ی رتبه دار رکوردی، تحت راهنمایی دکتر احمد نزاکتی 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 هیچ جا ارایه نشده است.
- حقوق معنوی این اثر، به دانشگاه صنعتی شاهرود تعلق دارد، و مقالات مستخرج با نام “دانشگاه صنعتی شاهرود” یا “Shahrood University of Technology” به چاپ خواهد رسید.
- حقوق معنوی تمام افرادی که در به دست آوردن نتایج اصلی پایان نامه تاثیرگذار بوده‌اند، در مقالات مستخرج از پایان نامه رعایت می‌گردد.
- در تمام مراحل انجام این پایان نامه، در مواردی که از موجود زنده (یا بافت‌های آنها) استفاده شده است، ضوابط و اصول اخلاقی رعایت شده است.
- در تمام مراحل انجام این پایان نامه، در مواردی که به حوزه اطلاعات شخصی افراد دسترسی یافته (یا استفاده شده است)، اصل رازداری و اصول اخلاق انسانی رعایت شده است.

امینه صادق پور

شهریور ۱۳۹۹

مالکیت نتایج و حق نشر

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

چکیده

در برخی از آزمایش‌ها به دنبال ارزیابی میزان مقاومت مؤلفه‌های مورد آزمایش در برابر فشار سایر متغیرها می‌باشیم. چنین حالتی را در مباحث قابلیت اعتماد، یک مدل تنش- مقاومت می‌گویند. با توجه به کاربردهای قابلیت اعتماد تنش-مقاومت در صنعت، علوم مهندسی و پزشکی، در این رساله برآوردهای نقطه‌ای و فاصله‌ای پارامتر تنش-مقاومت در توزیع نمایی تعمیم یافته براساس طرح نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی مورد مطالعه قرار گرفته است. همچنین میزان کارایی برآوردهای حاصل از طرح نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی با رکوردهای معمولی حاصل از روش نمونه‌گیری معکوس تحت الگوی نرخ خطر معکوس متناسب محاسبه و مقایسه شده‌اند.

کلمات کلیدی: پارامتر تنش-مقاومت، مدل نرخ خطر معکوس متناسب، نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی، داده‌های رکوردی، توزیع نمایی تعمیم یافته.

فهرست مقالات مستخرج از پایان نامه

1. Sadeghpour, A., Salehi, M. and Nezakati, A. (2020). *Estimation of stress-strength reliability using record ranked set sampling scheme under the generalized exponential distribution*, Journal of Statistical Computation and Simulation, 90, 51–74.
2. Sadeghpour, A., Nezakati, A. and Salehi, M. (2020). *Comparison of two sampling schemes in estimating the stress-strength reliability under the proportional reversed hazard rate model*, Statistic, Optimization Information Computing. DOI: 10.19139/soic-2310-5070-781.
3. Sadeghpour, A., Nezakati, A. (1398). *Estimation of stress-strength reliability using record ranked set sampling scheme from the inverse Rayleigh distribution*. 14th Iranian Statistics Conference, p 853, Shahrood, Iran.

۴. صادق پور، ا.، نزاکی، ا.، صالحی، م. (۱۳۹۶)، ”برآورد نقطه‌ای پارامتر تنش-مقاومت در توزیع نمایی دو پارامتری براساس نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی“، کنفرانس روش‌های مدرن در قیمت گذاری بیمه‌ای و آمارهای صنعتی، ص ۲۵، همدان.

۵. صادق پور، ا.، نزاکی، ا.، صالحی، م. (۱۳۹۶)، ”برآورد نقطه‌ای پارامتر تنش-مقاومت در توزیع نمایی تعمیم یافته براساس نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی“، یازدهمین سمینار احتمال و فرآیندهای تصادفی، ص ۳۵۹، قزوین.

پیشگفتار

قابلیت اعتماد تنش-مقاومت، احتمال سالم ماندن مؤلفه یا به عبارت دیگر احتمال عدم شکست می‌باشد و به صورت $R = Pr(X < Y)$ تعریف می‌شود. در حقیقت اگر متغیر تصادفی X بیان‌گر تنش (فشار) وارد شده بر یک مؤلفه و متغیر تصادفی Y نشان‌دهنده‌ی مقاومت آن در برابر فشار X باشد، در این صورت اگر $X > Y$ باشد، آن‌گاه مؤلفه از کار افتاده و در نتیجه شکست رخ می‌دهد. برعکس اگر $X < Y$ باشد، آن مؤلفه به عملکرد خود ادامه می‌دهد.

در مبحث قابلیت اعتماد یکی از موضوعات مورد علاقه استنباط در مورد مدل تنش-مقاومت در حالتی است که متغیرهای تصادفی X و Y توزیع‌های مستقل از یکدیگر دارند. از آن‌جا که در بسیاری از آزمایش‌های دنباله‌ای فقط مقادیر رکوردها گزارش می‌شوند و معمولاً تعداد رکوردهای مشاهده شده در یک آزمایش دنباله‌ای کوچک است، محقق ناچار است برای استنباط در مورد توزیع جامعه‌ی آماری از تعداد محدودی آماری رکوردی استفاده کند.

با توجه به این که طرح نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی در بعضی موارد می‌تواند جایگزین مناسبی برای طرح رکوردهای معمولی باشد، در این رساله برآوردهای نقطه‌ای و فاصله‌ای پارامتر تنش-مقاومت با استفاده از طرح نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی مورد مطالعه قرار گرفته است. در این رساله مشاهده خواهیم کرد چنانچه توزیع جامعه‌ی آماری از الگوی نرخ خطر معکوس متناسب پیروی کند، برآوردهای نقطه‌ای و فاصله‌ای به‌دست آمده از طرح نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی، با تعداد آزمایش‌های بسیار کم‌تر نسبت به طرح رکوردهای معمولی، کاراتر می‌باشند. این مجموعه شامل پنج فصل می‌باشد که خلاصه‌ای از آن در زیر آمده است.

- در فصل اول، به معرفی اجمالی برخی مفاهیم و تعاریف اولیه مورد استفاده در رساله می‌پردازیم. نظریه قابلیت اعتماد و یکی از مباحث مطرح در این نظریه، یعنی مدل تنش-مقاومت که در سرتاسر این رساله نقش کلیدی ایفا می‌کند معرفی می‌کنیم. همچنین به بیان کاربردهایی از قابلیت اعتماد تنش-مقاومت و تاریخچه‌ی آن می‌پردازیم. با توجه به این که طرح استفاده شده در این رساله نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی می‌باشد که تلفیقی از مجموعه‌ی رتبه‌دار و آماره‌های رکوردی است، در ادامه مطالبی در مورد رکوردها و ویژگی‌های آن و نمونه‌ی مجموعه‌ی رتبه‌دار و در نهایت طرح نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی بیان شده است. با توجه به این که در فصل‌های بعدی از برخی فواصل اطمینان

بوت استرپ برای برآورد پارامتر تنش-مقاومت استفاده کرده‌ایم، به معرفی این پارامتر، تعاریف و الگوریتم‌های مربوط پرداخته‌ایم. تعریفی از خانواده‌های نرخ خطر ارائه شده است و در پایان برخی روش‌های شبیه‌سازی استفاده شده در رساله معرفی شده است.

- در فصل دوم، با در نظر گرفتن توزیع نمایی تعمیم‌یافته به عنوان توزیع پایه جوامع X و Y و با استفاده از طرح نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی، برآوردگر درست‌نمایی ماکسیمم و برآوردگر نااریب با کمترین واریانس پارامتر تنش-مقاومت را محاسبه نمودیم. فاصله اطمینان دقیق براساس کمیت محور و برخی فواصل اطمینان بوت استرپ را به دست آوردیم و به مقایسه فواصل فوق پرداختیم. سپس موارد مطرح شده در این فصل را در یک مثال واقعی بررسی کردیم.

- در فصل سوم با در نظر گرفتن مفروضات فصل دوم به استنباط‌های بیزی پارامتر تنش-مقاومت پرداخته‌ایم. برآوردگر بیزی پارامتر تنش-مقاومت را به دست آوردیم و برای توزیع‌های پیشین سره و ناسره، برآوردگرهای فاصله‌ای بیزی و چگال‌ترین پسین (HPD) را به دست آوردیم و با استفاده از مطالعات شبیه‌سازی و مثال واقعی به بررسی موارد فوق پرداخته‌ایم.

- در فصل چهارم، با فرض مدل نرخ خطر معکوس متناسب به عنوان توزیع جوامع X و Y ، پارامتر تنش-مقاومت را برآورد نمودیم. با توجه به نتایج به دست آمده، مشاهده شد، در مدل نرخ خطر معکوس متناسب، تمام محاسبات به توزیع پایه بستگی ندارد. لذا با در نظر گرفتن توزیع نمایی تعمیم‌یافته به عنوان توزیعی از خانواده نرخ خطر معکوس متناسب، کارایی نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی مورد بررسی واقع شد و برآوردهای نقطه‌ای و فاصله‌ای حاصل از دو طرح نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی و رکوردهای معمولی مقایسه شدند و در نهایت این فصل را با ذکر یک مثال واقعی به پایان رساندیم.

- در فصل پنجم نیز با در نظر گرفتن توزیع نمایی دو پارمتری برآوردگر درست‌نمایی ماکسیمم و برآوردگر نااریب با کمترین واریانس را برای پارامتر تنش-مقاومت در حالت‌های مختلف محاسبه کردیم.

به علت حجم زیاد برنامه‌های استفاده شده و جلوگیری از افزایش کاذب صفحات رساله، گزیده‌ای کوچک از کدهای برنامه‌نویسی R در ضمیمه (الف) آورده شده است.

فهرست مطالب

ع	لیست تصاویر
ف	لیست جداول
ک	فهرست نمادها
ل	فهرست کلمات اختصاری
۱	۱ مقدمات و مفاهیم اولیه
۱	۱.۱ نظریه قابلیت اعتماد
۲	۲.۱ پارامتر تنش-مقاومت
۴	۱.۲.۱ کاربردهای قابلیت اعتماد تنش-مقاومت
۴	۳.۱ آماره‌های مرتب و کاربرد آن‌ها
۵	۴.۱ رکوردها و ویژگی‌های آن
۸	۵.۱ نمونه‌ی مجموعه‌ی رتبه‌دار
۱۱	۶.۱ معرفی نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی
۱۵	۱.۶.۱ ویژگی‌های حاصل از روش $RRSS$
۱۶	۷.۱ مقدمه‌ای بر روش بوت استرپ
۱۷	۱.۷.۱ فواصل اطمینان بوت استرپ
۱۷	۸.۱ خانواده‌های نرخ خطر
۲۰	۹.۱ روش‌های شبیه‌سازی
۲۰	۱.۹.۱ روش مونت کارلو
۲۱	۱۰.۱ روش مونت کارلو با نمونه‌گیری از نقاط مهم
۲۳	۲ برآورد پارامتر R در توزیع نمایی تعمیم یافته با استفاده از $RRSS$
۲۳	۱.۲ مقدمه
۲۴	۲.۲ برآورد نقطه‌ای پارامتر تنش-مقاومت

۲۵	برآوردگر ML	۱.۲.۲
۳۰	برآوردگر $UMVU$	۲.۲.۲
۳۲	برآورد فاصله‌ای پارامتر تنش-مقاومت	۳.۲
۳۳	فاصله اطمینان براساس کمیت محوری	۱.۳.۲
۳۴	فواصل اطمینان بوت‌استرپ پارامتری	۲.۳.۲
۳۶	مقایسه‌ی برآوردهای فاصله‌ای براساس شبیه‌سازی	۴.۲
۳۷	مثال واقعی	۵.۲
۴۱	استنباط بیزی پارامتر R در توزیع نمایی تعمیم یافته با استفاده از $RRSS$	۳
۴۱	مقدمه	۱.۳
۴۲	برآوردگر بیزی پارامتر تنش-مقاومت	۲.۳
۴۶	برآوردهای فاصله‌ای بیزی پارامتر تنش-مقاومت	۳.۳
۴۹	مقایسه برآوردهای فاصله‌ای بیزی براساس شبیه‌سازی	۴.۳
۵۰	مثال واقعی	۵.۳
۵۳	مقایسه R براساس دو طرح نمونه‌گیری در خانواده نرخ خطر معکوس متناسب	۴
۵۳	مقدمه	۱.۴
۵۴	معرفی دو طرح در داده‌های رکوردی	۲.۴
۵۴	برآوردهای نقطه‌ای پارامتر R در خانواده $PRHR$	۳.۴
۵۵	برآوردگر درست‌نمایی ماکسیمم	۱.۳.۴
۶۰	برآوردگر $UMVU$	۲.۳.۴
۶۰	برآوردگر بیزی	۳.۳.۴
۶۲	برآوردهای فاصله‌ای پارامتر R در خانواده $PRHP$	۴.۴
۶۳	فاصله اطمینان براساس کمیت محور	۱.۴.۴
۶۳	فاصله اطمینان مجانبی	۲.۴.۴
۶۴	فواصل معتبر بیزی	۳.۴.۴
۶۵	فواصل اطمینان بوت‌استرپ	۴.۴.۴
۶۵	مقایسه‌ی برآوردها به کمک شبیه‌سازی	۵.۴
۶۶	مقایسه برآوردهای فاصله‌ای حاصل از دو طرح	۱.۵.۴
۶۹	مثال واقعی	۶.۴
۷۵	برآورد پارامتر تنش-مقاومت براساس $RRSS$ در توزیع نمایی دو پارامتری	۵
۷۶	برآوردگر درست‌نمایی ماکسیمم	۱.۵
۷۸	برآورد نااریب با کمترین واریانس	۲.۵
۸۰	شبیه‌سازی	۱.۲.۵

۸۵	مراجع
۹۵	آ کدهای برنامه‌نویسی
۱۰۵	ب اثبات روابط
۱۰۷	واژه‌نامه فارسی به انگلیسی
۱۰۹	واژه‌نامه انگلیسی به فارسی
۱۱۲	نمایه

لیست تصاویر

۳۰		نمودار $MSE(\hat{R}_{ML}, \mathcal{R})$ در مقابل $\mathcal{R}$	۱.۲
۳۳		نمودار $Var(\hat{R}_{UMVU})$ در مقابل $\mathcal{R}$	۲.۲
۳۳		نمودار $e(\hat{R}_{UMVU}, \hat{R}_{ML})$ در مقابل $\mathcal{R}$	۳.۲
۳۹		مقایسه $\hat{R}_{UMVU}$ و $\hat{R}_{ML}$	۴.۲
		نمودار $MSE(\hat{R}_B, \mathcal{R})$ در مقابل $\mathcal{R}$ برای حالت ۳ و با ابر-پارامترهای $a_i =$	۱.۳
۴۷		$b_i = ۲$ و $c_i = d_i = ۱$	
۵۷		نمودار $MSE(\hat{R}_{ML}, \mathcal{R})$ و $Bias(\hat{R}_{ML}, \mathcal{R})$ در مقابل $\mathcal{R}$	۱.۴
۵۹		نمودار $e(\hat{R}_{ML}^{(R)}; \hat{R}_{ML}^{(RRSS)})$ در مقابل $\mathcal{R}$	۲.۴
۵۹		نمودار کارایی نسبی $\hat{R}_{ML}$ براساس رابطه (۱۲.۴)	۳.۴
۶۱		نمودار $e(\hat{R}_{UMVU}^{(R)}; \hat{R}_{UMVU}^{(RRSS)})$ در مقابل $\mathcal{R}$	۴.۴
۶۱		کارایی نسبی $\hat{R}_{UMVU}$ براساس (؟؟)	۵.۴
۶۳		نمودار کارایی نسبی $\hat{R}_B$ براساس دو طرح	۶.۴
۷۳		مقایسه $\hat{R}_{UMVU}$ و $\hat{R}_{ML}$	۷.۴

لیست جداول

۳۷	مقادیر EL و CP فواصل اطمینان 95% برای R	۱.۲
۳۸	مقادیر EL و CP فواصل اطمینان 90% برای R	۲.۲
۳۹	مقادیر EL براساس داده‌های واقعی	۳.۲
۵۰	مقادیر EL و CP فواصل بیزی 95% برای R	۱.۳
۵۱	مقادیر EL و CP فواصل بیزی 90% برای R	۲.۳
۵۱	مقادیر EL براساس داده‌های واقعی	۳.۳
۶۷	EL و CP برای R به‌ازای $\gamma = 0.05$	۱.۴
۶۸	ادامه	۲.۴
۶۹	EL و CP برای R به‌ازای $\gamma = 0.1$	۳.۴
۷۰	ادامه	۴.۴
۷۱	مقادیر δ برای ترکیب‌هایی از m و n زمانی که $\gamma = 0.05$	۵.۴
۷۱	مقادیر δ برای ترکیب‌هایی از m و n زمانی که $\gamma = 0.1$	۶.۴
۷۲	مقادیر $RRSS$ پایین استخراج شده از داده‌های هواپیما	۷.۴
۷۲	مقادیر متوسط طول R براساس مجموعه داده‌ی واقعی	۸.۴
۸۱	نتایج مربوط به MLE و $UMVUE$	۱.۵

فهرست نمادها

$\mathbf{X}$	متغیر تنش
$\mathbf{Y}$	متغیر مقاومت
$\mathcal{R}$	پارامتر تنش مقاومت
$I(A)$	تابع نشانگر مجموعه A
$\mathbf{A}^\top$	ترانواده ماتریس $\mathbf{A}$
$Pr(A)$	احتمال پیشامد A
$E(\cdot)$	امید ریاضی
$Var(\cdot)$	واریانس
$Bias(\cdot)$	اریبی
$MSE(\cdot)$	میانگین مربعات خطا
$\mathbf{X}_i$	متغیر تصادفی i ام
$X = (X_1, \dots, X_m)$	نمونه‌ی تصادفی به اندازه m
$\mathbf{U}_i$	i امین رکورد بالا
$\mathbf{L}_i$	i امین رکورد پایین
$U_{i,i}$	i امین واحد یک $RRSS$ بالا
$L_{i,i}$	i امین واحد یک $RRSS$ پایین
$U_{(m)}^{RRSS} = (U_{1,1}, \dots, U_{m,m})$	یک $RRSS$ بالا به اندازه m
$L_{(m)}^{RRSS} = (L_{1,1}, \dots, L_{m,m})$	یک $RRSS$ پایین به اندازه m
$f(x)$	تابع چگالی احتمال متغیر تصادفی X
$\bar{F}(x)$	تابع بقای متغیر تصادفی X
$f_{i,i}(\cdot)$	تابع چگالی احتمال متغیر تصادفی $U_{i,i}$
$F_{i,i}(\cdot)$	تابع توزیع احتمال متغیر تصادفی $U_{i,i}$
$\text{Gamma}(\alpha, \beta)$	توزیع گاما با پارامتر α, β
$\text{Beta}(\alpha, \beta)$	توزیع بتا با پارامتر α, β
$\mathcal{F}(m, n)$	توزیع فیشر با درجات آزادی m و n
$\xrightarrow{d}$	همگرایی در توزیع

$\xrightarrow{p}$	همگرایی در احتمال
$\log x$	تابع لگاریتم طبیعی
$\gamma(\cdot)$	تابع گامای کامل
$B(\cdot, \cdot)$	تابع بتای کامل
$[a]$	کوچک‌ترین عدد صحیح بزرگ‌تر از a

فهرست کلمات اختصاری

CDF (Cumulative Distribution Function)

CP (Covering Probability)

EL (Expected Length)

HPD(Highest Posterior Density)

i.i.d (independent identically distribution)

MLE (Maximum likelihood Estimator)

MSE (Mean Square Error)

RSS (Ranked Set Sampling)

RRSS (Record Ranked Set Sampling)

RE (Relational Efficiency)

SEL(Square Error Loss Function)

PHR(Proportional Hazard Rate)

PRHR(Proportional Reversed Hazard Rate)

UMVUE (Uniformly Minimum Unbiased Estimator)

فصل ۱

مقدمات و مفاهیم اولیه

در این فصل به معرفی اجمالی برخی مفاهیم و تعاریف اولیه مورد استفاده در رساله می‌پردازیم. از آنجا که قابلیت اعتماد تنش-مقاومت در این رساله مورد توجه قرار گرفته‌است، ابتدا نظریه قابلیت اعتماد و در ادامه پارامتر تنش-مقاومت و کاربردهای آن مطرح خواهد شد. طرحی که در این رساله استفاده شده‌است تلفیقی از نمونه‌ی مجموعه‌ی رتبه‌دار و آماره‌های رکوردی می‌باشد، بنابراین ابتدا تعریفی از مفاهیم فوق و در نهایت نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی ارائه خواهد شد. با توجه به این که برخی فواصل اطمینان بوت‌استری را برای برآورد پارامتر تنش-مقاومت استفاده کرده‌ایم، تعاریف و الگوریتم‌های مربوط به فواصل فوق را بیان خواهیم کرد. در نهایت تعریفی از الگوریتم‌های تکراری که در این رساله استفاده شده‌است را ارائه خواهیم کرد.

۱.۱ نظریه قابلیت اعتماد

یکی از کاربردهای آمار به منظور نمایش رفتار پدیده‌های تصادفی، در نظریه قابلیت اعتماد نهفته است. از آنجایی که این نظریه بیشتر در مواردی مانند مهندسی و تعیین طول عمر دستگاه‌ها به کار گرفته می‌شود گاهی به آن مهندسی قابلیت اعتماد نیز می‌گویند. فرض کنید یک سیستم یا یک قطعه الکتریکی دارای طول عمر T باشد، در حقیقت T یک متغیر تصادفی است که می‌تواند در فاصله $(0, \infty)$ مقدار بگیرد. متغیر T می‌تواند پیوسته یا گسسته باشد. برای مثال اگر سیستم را لامپ فلورسنت در نظر بگیریم طول عمر آن یک متغیر تصادفی T است که هر مقدار را در

فاصله $(0, \infty)$ می‌گیرد. حال اگر سیستم را یک دستگاه فتوکپی در نظر بگیریم طول عمر آن، تعداد کپی‌هایی است که قبل از خرابی گرفته می‌شود. تابع محوری در مبحث قابلیت اعتماد جهت تعیین خواص متغیر طول عمر T ، تابع قابلیت اعتماد است. تابع قابلیت اعتماد یا تابع بقای T در لحظه $t > 0$ را با $S(t)$ نمایش می‌دهند که به صورت زیر تعریف می‌شود:

$$S(t) = Pr(T > t) = \int_t^{\infty} f(x)dx = 1 - F(t)$$

که در آن f تابع چگالی و F تابع توزیع احتمال است. واضح است که $S(0) = 1$ ، یعنی در زمان $t = 0$ قابلیت اعتماد سیستم یک است و $S(\infty) = 0$ ، یعنی در دراز مدت سیستم از کار می‌افتد. همچنین $S(t)$ یک تابع غیرصعودی است، به این معنی که اگر $0 < t_1 < t_2$ آنگاه $S(t_1) \geq S(t_2)$ ، یعنی در زمان t_1 قابلیت اعتماد سیستم کمتر از زمان t_2 نیست. لازم به ذکر است مفهوم طول عمر را برای افراد بخصوص بیماران نیز می‌توان بکار برد که این موضوع اهمیت مبحث قابلیت اعتماد در پزشکی و داروسازی را نشان می‌دهد.

تابع نرخ خطر، که در مباحث قابلیت اعتماد و طول عمر به تابع نرخ شکست نیز معروف است، نقش اساسی در اغلب تحلیل‌های بقا ایفا می‌کند. تابع نرخ خطر متغیر طول عمر T را با $h(t)$ نمایش می‌دهند و به صورت زیر تعریف می‌شود:

$$h(t) = \frac{f(t)}{S(t)}, \quad t > 0.$$

باید توجه داشت تابع نرخ خطر، احتمال از کار افتادن سیستم نیست، در حقیقت این کمیت نرخ از کارافتادگی سیستم را در زمان t نشان می‌دهد و معیاری برای اندازه‌گیری سالخوردگی است.

۲.۱ پارامتر تنش-مقاومت

اگر متغیر تصادفی X بیان‌گر فشار (تنش) وارد شده بر یک مؤلفه و متغیر تصادفی Y نشان دهنده‌ی مقاومت آن در برابر فشار X باشد، در این صورت اگر $X > Y$ ، آنگاه مؤلفه از کار افتاده و در نتیجه شکست رخ می‌دهد. برعکس، اگر $X < Y$ ، آن مؤلفه به عملکرد خود ادامه می‌دهد. در نتیجه، قابلیت اعتماد تنش-مقاومت همان احتمال سالم ماندن مؤلفه یا به عبارت دیگر احتمال عدم شکست است و به صورت $R = Pr(X < Y)$ تعریف می‌شود.

به لحاظ تاریخی مدل تنش-مقاومت بیشتر ریشه در وضعیت ناپارامتری دارد. نقطه‌ی شروع آن را می‌توان در کارهای ویلکاکسون^۱ (۱۹۴۵) و من و ویتنی^۲ (۱۹۴۷) ملاحظه کرد. هدف اصلی تحقیق آن‌ها مقایسه متغیرهای تصادفی X و Y و شرح نتایج مربوط به رفتار آن‌ها بود. پس از آن تا سال ۱۹۸۰ مطالعات بسیاری در این زمینه انجام شد که فعالیت‌های انجام شده در این دهه‌ها

^۱ Wilcoxon

^۲ Whitney

بیشتر تحت پذیره استقلال متغیرهای تصادفی تنش-مقاومت بوده است. از بین آن‌ها می‌توان به باوم^۳ (۱۹۵۶)، بیرن باوم و مک کارتی^۴ (۱۹۵۸)، اون^۵ (۱۹۶۴) و تانگ^۶ (۱۹۷۴) اشاره کرد. باتاچاریا و جانسون^۷ (۱۹۷۴) نیز مطالعاتی در خصوص سیستم‌های قابلیت اعتماد انجام داده‌اند. پس از آن، تعاریف مختلفی از توزیع نمایی دو متغیره ارائه گردید و مطالعاتی روی متغیرهای تصادفی نمایی وابسته با انواع مختلف وابستگی مطرح شدند. از جمله می‌توان به کارهای اواد و همکاران^۸ (۱۹۸۱) و کلین و باسو^۹ (۱۹۸۵) اشاره کرد. به طور کلی تا کنون پژوهش‌های بسیاری در خصوص برآورد پارامتر R تحت پذیره‌های توزیعی مختلف روی X و Y برای انواع گوناگونی از داده‌ها از جمله داده‌های کامل و داده‌های سانسور شده، مورد مطالعه قرار گرفته‌اند که کتاب کاتز^{۱۰} و همکاران (۲۰۰۳) مرجع مفیدی در این زمینه است. اریلماز^{۱۱} (۲۰۰۸) قابلیت اعتماد تنش-مقاومت را در سیستمی شامل n مؤلفه‌ی مستقل که هر مؤلفه شامل m زیر مؤلفه بود را به دست آورد. وی همچنین در سال ۲۰۱۰ به مطالعه‌ی قابلیت اعتماد تنش-مقاومت در سیستم‌های منسجم پرداخت. از جمله پژوهش‌های اخیر در ارتباط با پارامتر تنش-مقاومت می‌توان به سارا کوگلو^{۱۲} و کایا^{۱۳} (۲۰۰۷)، رضائی و همکاران (۲۰۱۰)، بکلیزی^{۱۴} (۲۰۱۴)، ندار و کیزیل آسلان^{۱۵} (۲۰۱۴)، اصغرزاده و همکاران (۲۰۱۳)، بصیرت و همکاران (۲۰۱۶) اشاره کرد. پاکدامن و احمدی (۲۰۱۸) برای افزایش قابلیت اعتماد تنش-مقاومت به طراحی سیستم مقاومتی پرداخته‌اند. تحقیقاتی نیز در این راستا انجام شده است که می‌توان رساله دکتری بصیرت (۱۳۹۳) و صالحی (۱۳۹۴) را نام برد.

^۳Birnbaum

^۴McCarty

^۵Own

^۶Tong

^۷Batacharya and Johnson

^۸Awad

^۹Klein and Basu

^{۱۰}Kotz

^{۱۱}Eryilmaz

^{۱۲}Saracoglu

^{۱۳}Kaya

^{۱۴}Baklizi

^{۱۵}Nadar and Kizilaslan

۱.۲.۱ کاربردهای قابلیت اعتماد تنش-مقاومت

کاربردهای عمده قابلیت اعتماد تنش-مقاومت در صنعت، علوم مهندسی و پزشکی است. مثال‌هایی از کاربردهای تنش-مقاومت در جانسون^{۱۶} (۱۹۸۸) ارائه شده‌اند که به برخی از آن‌ها اشاره می‌کنیم.

- فرض کنید X متغیر تصادفی بیشترین فشار وارده بر پیستون موتور یک موشک و Y متغیر تصادفی مقاومت آن پیستون در برابر فشار وارد شده باشد، در این صورت R احتمال روشن شدن موفقیت آمیز موتور است.
- فرض کنید X و Y متغیرهای تصادفی طول عمر قطعه‌های الکترونیکی به ترتیب A و B باشند، در این صورت R احتمال این است که A زودتر از B خراب شود.
- فرض کنید X و Y نشان‌دهنده‌ی تعداد دفعات بهبود بیماران سرطانی باشد که از دو روش شیمی درمانی متفاوت استفاده می‌کنند. در این صورت R مقایسه‌ای بین دو روش فوق برای یافتن روش مؤثر درمان است.
- اگر X و Y نشان‌دهنده‌ی استحکام دو ماده در یک طراحی مهندسی باشند، R می‌تواند بیانگر این باشد که استحکام X کمتر از استحکام Y است.
- یکی از مشخصات ساختمان ایده‌آل، داشتن مقاومت کافی در برابر زلزله است. اگر فرض کنید متغیر تصادفی Y نشان‌دهنده مقاومت یک ساختمان در هنگام زمین لرزه و متغیر تصادفی X فشار ناشی از زمین لرزه باشد، آن‌گاه R می‌تواند احتمال طراحی یک ساختمان ایمن باشد.
- در طراحی پل‌ها نیز اگر متغیر تصادفی Y نشان‌دهنده مقاومت پایه‌های پل و متغیر تصادفی X نشان‌دهنده فشار ناشی از وزن پل باشد که بر پایه‌ها وارد می‌شود، آن‌گاه R احتمال طراحی موفقیت آمیز پل می‌باشد.

۳.۱ آماره‌های مرتب و کاربرد آن‌ها

فرض کنید $X_1, \dots, X_n$ متغیر تصادفی باشند. اگر آن‌ها را به ترتیب صعودی مرتب کنیم و مقادیر مرتب شده را به صورت $X_{(1)} \leq X_{(2)} \leq \dots \leq X_{(n)}$ نشان دهیم، در این صورت $X_{(j)}$ ، $j = 1, \dots, n$ ، را j امین آماره مرتب در نمونه‌ای به حجم n گویند. کوچک‌ترین و بزرگ‌ترین آماره مرتب، یعنی $X_{(1)}$ و $X_{(n)}$ به مقادیر فرین معروفند و در مطالعات آماری بیشتر مورد توجه هستند.

^{۱۶}Johnson

آماره‌های مرتب در قسمت‌های مختلف توصیفی و استنباطی علم آمار، دارای کاربرد می‌باشند که می‌توان چند مورد را به شرح ذیل ذکر کرد.

- در محاسبه تابع توزیع تجربی نیازمند استفاده از آماره‌های مرتب هستیم و از آن‌جا که در آزمون‌های نیکویی برازش، اغلب تمرکز روی تغییرات بین تابع توزیع تجربی و تابع توزیع فرضی است، آماره‌های مرتب در این آزمون‌ها نقش اساسی دارند.
 - در موضوعات کنترل کیفیت برای بررسی در کنترل بودن تولید، اغلب از نمودار میانگین و دامنه تغییرات یا از نمودار میانه و دامنه تغییرات استفاده می‌شود که برای محاسبه میانه و دامنه تغییرات ناچار به استفاده از آماره‌های مرتب می‌باشیم.
 - آماره‌های مرتب، در بسیاری از موارد به عنوان آماره‌های بسنده معرفی می‌شوند و برآوردهای ناریب با کمترین واریانس، فواصل اطمینان و پرتوان‌ترین آزمون برای پارامترهای مجهول را فراهم می‌کنند.
 - در آزمایشات طول عمر، اغلب به دلیل عواملی از قبیل زمان یا اعتبار مالی، امکان شرکت دادن تمام واحدها در آزمایش وجود ندارد. بنابراین مشاهده طول عمر تمام واحدها امکان‌پذیر نیست، به همین دلیل از روش‌های سانسور استفاده می‌شوند و آماره‌های مرتب نقش مهمی در سانسورها دارند.
- برای جزئیات بیشتر در مورد خواص آماره‌های مرتب و کاربرد آن‌ها، می‌توان به کتاب‌های دیوید و ناگاراچا^{۱۷} (۲۰۰۳) و آرنولد^{۱۸} و همکاران (۲۰۰۸) مراجعه کرد.

۴.۱ رکوردها و ویژگی‌های آن

وقتی کلمه رکورد را می‌شنویم به یاد مسابقات رشته‌های ورزشی می‌افتیم و تصور عموم برای این است که این واژه مختص رشته‌های ورزشی می‌باشد. اما حقیقت این است که هرگاه با مشاهدات متوالی سروکار داریم، مشاهداتی که از مقادیر قبلی خود کمتر یا بیشتر می‌باشند مورد توجه قرار می‌گیرند و در نتیجه موضوع رکورد می‌تواند مطرح شود. این مشاهدات دنباله‌ای می‌توانند مثلاً متوسط دما در ماه‌های متوالی و متوسط میزان بارش در ماه‌های متوالی، تعداد تصادفات در هفته‌های متوالی، بازدیدکنندگان از مراکز در هفته‌های متوالی باشند. کاربرد اصلی رکوردها بیشتر به این خاطر است که در بسیاری از موارد داده‌های اولیه در دسترس نیستند و فقط مقادیر رکوردی ثبت می‌شوند. لذا محقق ناچار است برای استنباط در مورد توزیع مشاهدات از مقادیر رکوردی

^{۱۷}David and Nagaraja

^{۱۸}Arnold

استفاده کند. مانند داده‌های هواشناسی، ژئوفیزیک، زلزله‌نگاری و سیلاب‌ها. کشور ما نیز عاری از این حوادث نیست. از جمله در سال‌های اخیر شاهد سیل‌های فراوانی در مناطق مختلف کشور، زلزله سال ۱۳۹۷ سرپل ذهاب و زلزله سال ۱۳۸۲ در بم بوده ایم که جان هزاران انسان به خطر افتاد و خسارت مالی زیادی در پی داشت. لذا بررسی، مطالعه و استفاده از رکوردها و پیش‌بینی و کنترل آن‌ها از طریق روش‌های آماری حائز اهمیت می‌باشد.

مطالعات آماری در خصوص داده‌های رکوردی ابتدا توسط چاندلر^{۱۹} (۱۹۵۲) انجام شد. گلیک^{۲۰} (۱۹۷۸) به طور جامع رکوردها را مورد بررسی قرار داد. محققین بسیاری به موضوع برآورد پارامتری و ناپارامتری به کمک مشاهدات رکوردی پرداخته‌اند، از جمله احسن اله و بوج^{۲۱} (۱۹۹۶) و کارلین و گلفند^{۲۲} (۱۹۹۳). تعدادی از پژوهشگران مانند کوچار^{۲۳} (۱۹۹۰، ۱۹۹۶) و گوپتا و کرمانی^{۲۴} (۱۹۸۸) نیز به مطالعه خواص یکتائی مقادیر رکوردی پرداخته‌اند. از دیگر پژوهش‌های استنباطی بر مبنای رکوردها می‌توان به احمدی و ارقامی (۲۰۰۱) و رزمخواه و همکاران (۱۳۸۶ ش) اشاره کرد. کتاب رکوردها اثر آرنولد و همکاران^{۲۵} (۱۹۹۸) و کتاب استنباط پارامتری و ناپارامتری براساس داده‌های شکست رکورد از گولاتی و پاجت^{۲۶} (۲۰۰۳) در حال حاضر به‌عنوان دو منبع اصلی در زمینه نظریه‌ی رکوردها هستند. احمدی و همکاران (۲۰۰۹) K -رکوردها را در رده کلی از توابع تحت تابع نرخ خطر متناسب در نظر گرفتند و در سال (۲۰۱۰) معیار نزدیکی پیتمن را براساس رکوردها در خانواده توزیع‌های مکان-مقیاس بررسی کردند.

مثال کاربردی

موارد متعددی وجود دارند که آزمایش‌ها به صورت دنباله‌ای انجام می‌شوند و تنها آماره‌های رکوردی ثبت می‌شوند و از تمامی داده‌ها اطلاعی نداریم. بنابراین محقق ناچار است در استنباط‌های خود از تنها اطلاعات موجود، یعنی مشاهدات رکوردی استفاده کند. به دو مثال کاربردی زیر که در آن‌ها آماره‌های رکوردی تنها مشاهدات ثبت شده هستند توجه کنید (برگرفته از صالحی، ۱۳۹۴).

مثال ۱.۴.۱. روند اندازه‌گیری ضخامت یک کالای تولید شده (X) با استفاده از یک میکرومتر

^{۱۹}Chandler

^{۲۰}Glick

^{۲۱}Ahsanullah and Bhoj

^{۲۲}Carlin and Gelfand

^{۲۳}Kochar

^{۲۴}Gupta and Kirmani

^{۲۵}Arnold

^{۲۶}Gulati and Padget

را در نظر بگیرید. برای اندازه‌گیری نازک‌ترین کالا در میان اولین n کالای تولید شده، یک کالا را به طور تصادفی انتخاب کرده و ضخامت آن را با استفاده از میکرومتر اندازه‌گیری می‌کنند. ضخامت آن برابر است با $X_1 = L_1$. ضخامت این کالا (L_1) یا به عبارت دیگر، فاصله‌ی ایجاد شده بین دو دهنه‌ی میکرومتر را به عنوان شاخص برای اندازه‌ی ضخامت کالاهای بعدی در نظر می‌گیرند. اندازه‌گیری دقیق بعدی زمانی صورت می‌گیرد که ضخامت کالایی از شاخص مذکور کمتر باشد (قطر کالا بین دو دهنه‌ی میکرومتر قرار گیرد). به عنوان مثال فرض کنید اندازه‌گیری دقیق دوم برای دهمین کالا صورت پذیرد، یعنی داریم $L_2 = X_1$. از این به بعد، شاخص اندازه‌گیری دقیق بعدی L_2 خواهد بود. اگر این روند را تا کالای n ام انجام دهیم، تعداد اندازه‌گیری دقیق با احتمال بسیار زیاد کمتر از n خواهد بود. به طور مثال فرض کنید $L_7 = X_{n-5}$ آخرین اندازه‌گیری دقیق انجام شده باشد. بنابراین، تنها اطلاعات در دسترس از ضخامت کالاهای یافته‌های $L_i, i = 1, \dots, 7$ ، است که همان رکوردهای پایین از جامعه‌ی X می‌باشند. حال اگر هدف اندازه‌گیری نازک‌ترین کالا با روش فوق باشد، فقط تعداد محدودی از رکوردهای پایین را در اختیار خواهیم داشت.

مثال ۲.۴.۱. آزمون شکستن الوارهای چوبی را در نظر بگیرید و فرض کنید Y متغیر تصادفی میزان مقاومت الوارها باشد. آزمایش یاد شده بدین ترتیب است که اولین الوار آنقدر تحت فشار قرار می‌گیرد تا بشکند. فرض کنید L_1 میزان مقاومت اولین الوار تا زمان شکستن آن باشد، در این صورت داریم $L_1 = Y_1$. بعد از آن بقیه الوارها یکی پس از دیگری در معرض فشار قرار می‌گیرند، تا زمانی که بشکنند. یا تحت فشاری به میزان L_1 قرار بگیرند. چنانچه الواری در برابر فشار وارده‌ی L_1 نشکند، از آزمایش خارج شده و به عنوان یک نمونه‌ی سالم یا دست دوم مورد استفاده قرار می‌گیرد، اما اگر قبل از این که تحت فشار معادل L_1 قرار بگیرد بشکند، یعنی یک رکورد پایین رخ داده‌است و به طور مثال داریم $L_2 = Y_1$. این عمل تا زمانی انجام می‌پذیرد که $(L_1, \dots, L_m)$ ثبت شوند. در واقع در این آزمایش، مقاومت الوارها که همان m رکورد پایین از جامعه‌ی Y می‌باشند، تنها اطلاعاتی از جامعه‌ی آماری هستند که جمع‌آوری می‌شوند.

تعریف ۱.۴.۱. (رکوردهای بالا و پایین). فرض کنید $\{X_i, i \geq 1\}$ دنباله‌ای از متغیرهای تصادفی مستقل، هم توزیع (iid) و پیوسته با تابع توزیع تجمعی $F(x)$ و تابع چگالی $f(x)$ باشد. X_i را یک رکورد بالا گویند، اگر

$$X_i > \max\{X_1, X_2, \dots, X_{i-1}\}, \quad i = 2, 3, \dots$$

طبق قرارداد X_1 را اولین رکورد بالا تعریف می‌کنند. زمانی که در آن رکورد بالا رخ می‌دهد، خود یک متغیر تصادفی است و به صورت زیر تعریف می‌شود:

$$T_1 = 1, \quad T_n = \min\{i : X_i > X_{T_{n-1}}\}$$

در این صورت دنباله رکوردهای مستخرج از X_i عبارتست از

$$X_{T_1}, X_{T_2}, \dots$$

اگر n امین رکورد بالا را با U_n نشان دهیم، آن گاه

$$U_n = X_{T_n}, \quad n = 1, 2, \dots$$

به طور مشابه X_i را یک رکورد پایین گویند، اگر

$$X_i < \min\{X_1, X_2, \dots, X_{i-1}\}, \quad i = 2, 3, \dots$$

طبق قرارداد X_1 را اولین رکورد پایین تعریف می کنند. اگر n امین رکورد پایین را با L_n نشان دهیم، آن گاه

$$L_n = X_{S_n}, \quad n = 1, 2, \dots$$

که در آن

$$S_1 = 1, \quad S_n = \min\{i : X_i < X_{S_{n-1}}\}$$

توزیع احتمالی رکوردها

بسیاری از خواص رکوردهای بالا را می توان برحسب تابع نرخ از کارافتادگی تجمعی $R(u)$ بیان کرد که در آن $\bar{F}(u) = 1 - F(u)$ و در صورت نامنفی بودن متغیر تصادفی، $\bar{F}(u)$ همان تابع بقا در نظریه اعتماد است. تابع چگالی کناری U_i ها عبارت است از

$$f_{U_i}(u) = \frac{R^{i-1}(u)}{(i-1)!} f(u), \quad -\infty < u < \infty, \quad i = 1, \dots, n,$$

و تابع چگالی احتمال توأم $U_1, \dots, U_n$ به صورت زیر می باشد

$$f_{U_1, \dots, U_n}(u_1, \dots, u_n) = f(u_n) \prod_{i=1}^{n-1} r(u_i), \quad -\infty < u_1 < \dots < u_n < \infty,$$

که در آن

$$r(u) = \frac{d}{du} R(u) = \frac{f(u)}{\bar{F}(u)}$$

تابع نرخ خطر است.

اگر در روابط فوق F را به جای $\bar{F}$ قرار دهیم، به ترتیب تابع چگالی احتمال کناری و تابع چگالی احتمال توأم رکوردهای پایین به دست می آیند.

۵.۱ نمونه‌ی مجموعه‌ی رتبه‌دار

در مبحث نمونه‌گیری، پژوهش‌گر همواره با مسأله‌ی انتخاب روش مناسب برای جمع‌آوری داده‌ها مواجه است، به‌طوری‌که به دنبال روش‌هایی می‌باشد که دقت را برای برآورد ویژگی‌های یک جامعه‌ی آماری افزایش داده و هزینه را کاهش دهد. در شرایطی که اندازه‌گیری دقیق واحدهای جامعه‌ی آماری سخت، پرهزینه یا زمان‌بر باشد، اما بتوان واحدهای جامعه‌ی آماری را به سادگی و با کمترین هزینه به صورت چشمی یا قضاوتی مرتب کرد، از نمونه‌گیری مجموعه‌ی رتبه‌دار

(RSS) استفاده می‌شود. در این روش نمونه‌گیری، تعداد زیادی نمونه به زیرنمونه‌هایی کوچک‌تر با اندازه‌های مساوی تقسیم می‌شوند، پس از آن هر زیرمجموعه به طور مجزا مرتب می‌شود و از هر زیر نمونه یک واحد انتخاب می‌شود. به طور مثال فرض کنید مایل به تخمین میانگین طول عمر درختان یک جنگل هستیم. می‌توان درختان را براساس ارتفاع آن‌ها (که به طول عمر آن‌ها وابستگی دارد) به صورت چشمی مرتب کرد. برای دیدن مثال‌های بیشتر به چن^{۲۷} و همکاران (۲۰۰۳) مراجعه کنید. در اغلب موارد، روش نمونه‌گیری مجموعه‌ی رتبه‌دار بر نمونه‌گیری تصادفی برتری دارد و برآوردگرهایی که از این روش به‌دست می‌آیند، عملکرد بهتری نسبت به برآوردگرهای حاصل از روش نمونه‌گیری تصادفی دارند.

تعریف ۱.۵.۱. فرض کنید $X_1, X_2, \dots, X_m$ یک نمونه تصادفی از توزیعی با تابع توزیع احتمال $F(\cdot)$ و تابع چگالی احتمال $f(\cdot)$ باشد. این نمونه را به صورت چشمی یا قضاوتی مرتب کرده، سپس کوچک‌ترین آن‌ها را انتخاب می‌کنیم و با $X_{1,1}$ نشان می‌دهیم و بقیه را کنار می‌گذاریم. نمونه تصادفی دیگری به اندازه m گرفته و پس از مرتب کردن آن به صورت چشمی یا قضاوتی، دومین کوچک‌ترین واحد آن را انتخاب می‌کنیم و با $X_{2,2}$ نشان می‌دهیم و سپس بقیه را کنار می‌گذاریم. این کار را m بار انجام می‌دهیم و $X_{m,m}$ را از نمونه m ام انتخاب می‌کنیم و با $X_{m,m}$ نشان می‌دهیم. در نهایت به $X_{(m)}^{RSS} = (X_{1,1}, X_{2,2}, \dots, X_{m,m})$ یک نمونه‌ی مجموعه‌ی رتبه‌دار^{۲۸} (RSS) به اندازه‌ی m گفته می‌شود. به عبارت دیگر، چنانچه $X_{(i:m)k}$ نشان‌دهنده i امین آماره‌ی ترتیبی از نمونه تصادفی k ام باشد، آن‌گاه روش نمونه‌گیری مجموعه‌ی رتبه‌دار در شکل زیر خلاصه می‌شود

$$\begin{array}{ccccccc} \text{نمونه اول} & : & X_{(1:m)1} & X_{(2:m)1} & \dots & X_{(m:m)1} & \rightarrow X_{1,1} = X_{(1:m)1} \\ \text{نمونه دوم} & : & X_{(1:m)2} & X_{(2:m)2} & \dots & X_{(m:m)2} & \rightarrow X_{2,2} = X_{(2:m)2} \\ & \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ \text{نمونه } m\text{ام} & : & X_{(1:m)m} & X_{(2:m)m} & \dots & X_{(m:m)m} & \rightarrow X_{m,m} = X_{(m:m)m} \end{array} \quad (1.1)$$

در واقع ۱.۱ روش نمونه‌گیری مجموعه‌ی رتبه‌دار تک چرخه‌ای را نشان می‌دهد. از تعریف نمونه‌ی مجموعه‌ی رتبه‌دار مشخص می‌شود که واحدهای آن مستقل هستند، اما دلیلی ندارد که مرتب باشند. اگر شکل فوق را r بار تکرار کنیم، یک نمونه‌ی مجموعه‌ی رتبه‌دار r -حلقه‌ای به اندازه‌ی نمونه‌ی rm حاصل خواهد شد. به دلیل این که مرتب کردن واحدهای نمونه به صورت چشمی یا قضاوتی صورت می‌گیرد، ممکن است به ازای اندازه‌های نمونه‌ی بزرگ، خطای رتبه‌بندی افزایش یابد. بنابراین در عمل m را کم‌تر از ۱۰ در نظر می‌گیرند و با افزایش تعداد حلقه‌ها به اندازه‌ی نمونه دلخواه می‌رسند.

نمونه‌گیری مجموعه‌ی رتبه‌دار برای نخستین بار توسط مک‌ایننتایر^{۲۹} (۱۹۵۲) پیشنهاد شد. او

^{۲۷}Chen

^{۲۸}Ranked Set Sampling

^{۲۹}McIntyre

نشان داد که میانگین نمونه‌ی مجموعه‌ی رتبه‌دار ($\bar{X}_{RSS}$)، برآوردگری کاراتر از میانگین نمونه‌ی تصادفی ($\bar{X}_{SRS}$) برای میانگین جامعه‌ی آماری است. پس از آن، نمونه‌گیری مجموعه‌ی رتبه‌دار تا مدتی به فراموشی سپرده شد، تا این که هالز و دل^{۳۰} (۱۹۶۶) برای برآورد وزن شاخ و برگ درخت کاج از روش نمونه‌گیری معرفی شده توسط مک اینتایر استفاده کردند و این روش را برای اولین بار نمونه‌ی مجموعه‌ی رتبه‌دار نامیدند و به طور تجربی و براساس یک تحقیق میدانی به برتری این روش نسبت به روش نمونه‌گیری تصادفی پی بردند. اولین پژوهش نظری در رابطه با نمونه‌گیری مجموعه‌ی رتبه‌دار توسط تاکاهاشی و واکیموتو^{۳۱} (۱۹۶۸) صورت گرفت. آن‌ها دوباره روش نمونه‌گیری مجموعه‌ی رتبه‌دار را زنده کردند و نشان دادند که اگر رتبه‌بندی در نمونه‌گیری مجموعه‌ی رتبه‌دار بدون خطا صورت گیرد- که در اصطلاح به آن رتبه‌بندی کامل گفته می‌شود- آن‌گاه واریانس برآوردگر $\bar{X}_{RSS}$ از واریانس برآوردگر $\bar{X}_{SRS}$ به طور اکید کمتر است. تاکاهاشی و واکیموتو نیز همانند مک اینتایر اسم خاصی برای روش نمونه‌گیری قرار ندادند و این روش را برآورد نااریب میانگین جامعه‌ی آماری براساس نمونه‌ی طبقه‌بندی شده توسط مرتب کردن نامیدند. آن‌ها نشان دادند در بعضی توزیع‌های تک‌مدی مثل توزیع نرمال، واریانس $\bar{X}_{RSS}$ به دست آمده از یک نمونه‌ی مجموعه‌ی رتبه‌دار به حجم n با واریانس $\bar{X}_{SRS}$ حاصل از یک نمونه‌ی تصادفی به حجم $2n$ برابر است. تاکاهاشی در سال ۱۹۷۰، در کنار نمونه‌گیری مجموعه‌ی رتبه‌دار ارائه شده توسط مک اینتایر، یک روش جدید برای تولید نمونه‌ی مجموعه‌ی رتبه‌دار ارائه کرد. چندی بعد، دل و کلاتر^{۳۲} (۱۹۷۲) بدون در نظر گرفتن شرط رتبه‌بندی کامل، به نتایج مشابه تاکاهاشی و واکیموتو (۱۹۶۸) رسیدند. استوکس^{۳۳} (۱۹۸۰) نشان داد که برآوردگر به دست آمده براساس نمونه‌گیری مجموعه‌ی رتبه‌دار برای واریانس جامعه‌ی آماری کاراتر از برآوردگر متناظر آن براساس نمونه‌ی تصادفی (S^2) است. تا قبل از ۱۹۸۰، اکثر پژوهش‌های انجام شده در رابطه با نمونه‌گیری مجموعه‌ی رتبه‌دار در حالت ناپارامتری انجام می‌شدند. بوهن^{۳۴} (۱۹۹۶) کارهای ناپارامتری انجام شده درباره‌ی مجموعه‌ی رتبه‌دار را مرور کرد. چند سال بعد از ۱۹۸۰، حالت پارامتری نمونه‌ی مجموعه‌ی رتبه‌دار مورد توجه قرار گرفت. از جمله‌ی آن‌ها می‌توان به استوکس (۱۹۹۵) اشاره کرد که برآوردگر درست‌نمایی ماکسیمم و بهترین برآوردگر خطی نااریب (BLU) را براساس نمونه‌ی مجموعه‌ی رتبه‌دار برای خانواده مکان-مقیاس به دست آورد. پاتیل و همکاران^{۳۵} (۱۹۹۹) تاریخچه‌ای از کارهای انجام شده در رابطه با نمونه‌گیری مجموعه‌ی رتبه‌دار فراهم کردند. چن در سال ۲۰۰۰ نتایج استوکس در رابطه با اطلاع فیشر را به حالت ناپارامتری تعمیم داد. وی به همراه همکارانش در سال ۲۰۰۳ کتابی در رابطه با نمونه‌گیری مجموعه‌ی

^{۳۰} Halls and Dell

^{۳۱} Takahashi and Wakimoto

^{۳۲} Clutter

^{۳۳} Stoks

^{۳۴} Bohn

^{۳۵} Patil

رتبه‌دار به چاپ رساند. ژنگ و الصالح^{۳۶} در سال ۲۰۰۴ به بررسی برآوردگر درست‌نمایی ماکسیمم اصلاح شده پرداختند. بهترین برآوردگر خطی نااریب براساس نمونه‌ی مجموعه‌ی رتبه‌دار توسط افراد زیادی از جمله بارت و مور^{۳۷} (۱۹۹۷) و بارتو^{۳۸} (۱۹۹۹) مورد توجه قرار گرفت. سینها (۲۰۰۵) یک جمع‌آوری از تحقیقات انجام‌شده‌ی پارامتری و ناپارامتری در خصوص نمونه‌گیری مجموعه‌ی رتبه‌دار تا آن زمان را ارائه کرد. بالا کریشان و لی^{۳۹} (۲۰۰۶) با مرتب کردن واحدهای نمونه‌ی مجموعه‌ی رتبه‌دار، آن را $ORSS$ نامیدند و پس از به‌دست آوردن ویژگی‌های توزیعی $ORSS$ یک فاصله‌ی اطمینان برای چندک‌های جامعه‌ی آماری براساس آن به‌دست آوردند. آن‌ها همچنین در سال ۲۰۰۸، مقادیر عددی گشتاورهای $ORSS$ را برای چند توزیع آماری به‌دست آوردند و با استفاده از آن‌ها برآوردگر BLU را براساس $ORSS$ محاسبه کردند. جعفری جوزانی و جانسون^{۴۰} (۲۰۱۱) برآوردیابی را براساس RSS در جوامع متناهی در نظر گرفتند (صالحی، ۱۳۹۰ ش، را ببینید).

در این فصل با بررسی معایب و مشکلاتی که در روش جمع‌آوری داده‌های رکوردی و نمونه‌گیری مجموعه‌ی رتبه‌دار وجود دارند به معرفی روش جدیدی تحت عنوان نمونه‌گیری مجموعه‌ی رتبه‌دار رکوردی^{۴۱} ($RRSS$) که تلفیقی از رکوردها و مجموعه‌ی رتبه‌دار است می‌پردازیم که نخستین بار توسط صالحی (۱۳۹۴ ش) در رساله دکتري مطرح شده است. در این رساله با در نظر گرفتن روش جدید مطرح شده سعی براین داریم با ایده گرفتن از طرح مطرح شده حالت‌های خاص دیگر را مورد بررسی قرار دهیم و با استفاده از روش نمونه‌گیری مجموعه‌ی رتبه‌دار رکوردی، خواص توزیع‌های مهم آماری را مورد ارزیابی قرار دهیم. در بخش ۶.۱ به معرفی مجموعه‌ی رتبه‌دار رکوردی می‌پردازیم و ویژگی‌های توزیعی روش فوق را بررسی می‌کنیم.

۶.۱ معرفی نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی

همان‌طور که در مثال‌های ۱.۴.۱ و ۲.۴.۱ ملاحظه شد، دو روش برای جمع‌آوری داده‌های رکوردی می‌توان در نظر گرفت. مواردی را همچون مثال ۱.۴.۱ در نظر بگیرید که در آن‌ها تعداد آزمایش‌ها ثابت است و تنها رکوردهای حاصل از آن ثبت می‌شوند. در نظریه رکوردها به این روش نمونه‌گیری، تصادفی یا نمونه‌ثابت می‌گویند که تعداد رکوردهای به‌دست آمده تصادفی است. در این روش که از لحاظ اقتصادی روشی مقرون به صرفه است، تنها رکوردهای به‌دست آمده تصادفی است و به طور متوسط تقریباً $\log(n)$ داده‌ی رکوردی از میان n متغیر تصادفی مستقل و هم توزیع

^{۳۶}Zheng and Al-saleh

^{۳۷}Bartnett and Moore

^{۳۸}Barreto

^{۳۹}Li

^{۴۰}Johnson

^{۴۱}Record Ranked Set Sampling

ثبت می‌شود.

در مثال ۲.۴.۱ مشاهده کردیم که به محض مشاهده شدن m امین رکورد، آزمایش متوقف شد که در مباحث مطالعات آماره‌های رکوردی به آن نمونه‌گیری معکوس گفته می‌شود. در این روش تعداد رکوردها ثابت است در حالی که تعداد آزمایش‌ها تصادفی می‌باشد.

برای این که به میزان دقت از پیش تعیین شده در تحلیل صفت‌های جامعه‌ی آماری مربوطه براساس آماره‌های رکوردی برسیم، لازم است تعداد زیادی آزمایش انجام گیرد تا به تعداد مورد نظر آماره‌های رکوردی برسیم (نمونه‌گیری معکوس رکوردها). برعکس در روش نمونه‌گیری تصادفی رکوردها، اگر تعداد آزمایش‌ها محدود باشد، به طور معمول تعداد کمی رکورد ثبت می‌شود و این امر باعث می‌شود که برآوردگرهای حاصل اغلب ناسازگار باشند (گولاتی و پاجت، ۱۹۹۴، و برگر و گولاتی، ۲۰۰۱ را ببینید). سامانیگو و وایتاگر^{۴۲} (۱۹۸۶) به این نتیجه رسیدند که اگر تعدادی رکورد تنها از یک دنباله از متغیرهای تصادفی در اختیار باشد، فقط در چارچوب آمار پارامتری می‌توان به برآوردگری کارا برای توابع توزیع احتمال جامعه براساس آن‌ها رسید. بدین منظور، آن‌ها در سال ۱۹۸۸ روشی در متون رکوردها پیشنهاد دادند که براساس آن، برآوردگر تابع توزیع احتمال جامعه‌ی آماری حتی براساس روش نمونه‌گیری تصادفی (نمونه ثابت) رکوردها نیز سازگار است. در واقع، آن‌ها روش نمونه‌گیری تصادفی و معکوس داده‌های رکوردی را که قبلاً به صورت یک نمونه‌ای انجام می‌شد به صورت چند نمونه‌ای تعمیم دادند. به طور دقیق‌تر، m دنباله‌ی تصادفی مستقل را در نظر بگیرید و فرض کنید نمونه‌گیری از دنباله‌ی i ام زمانی به اتمام می‌رسد که k_i امین رکورد مشاهده شود. بنابراین در مجموع $n = \sum_{i=1}^m k_i$ رکورد در اختیار داریم. شکل ۲.۱ چگونگی جمع‌آوری آماره‌های رکوردی را در این روش نمایش می‌دهد

$$\begin{aligned} 1: & X_{(1)1}, X_{(2)1}, \dots \rightarrow U_{(1)1} \dots \dots U_{(k_1)1} \\ 2: & X_{(1)2}, X_{(2)2}, \dots \rightarrow U_{(1)2} \dots U_{(k_2)2} \\ & \vdots \\ m: & X_{(1)m}, X_{(2)m}, \dots \rightarrow U_{(1)m} \dots \dots U_{(k_m)m} \end{aligned} \quad (2.1)$$

به طوری که $X_{(i)j}$ بیان‌گر i امین مشاهده از j امین دنباله و $U_{(i)j}$ بیان‌گر i امین رکورد معمولی بالا از آن دنباله می‌باشد. در ادامه محققانی روش چندنمونه‌ای سامانیگو و وایتاگر (۱۹۸۸) را مورد مطالعه و بررسی قرار دادند. به عنوان مثال، گولاتی و پاجت (۱۹۹۴) برآوردگری ناپارامتری برای چندک‌های جامعه‌ی آماری براساس روش چند نمونه‌ای بالا به‌دست آوردند. احمدی و ارقامی (۲۰۰۱) حالت دونمونه‌ای $m = 2$ را در نظر گرفتند در حالی که رزمخواه و همکاران (۱۳۸۶ش) حالت $k_i = k$ را مورد بررسی قرار دادند و آن‌را با روش یک نمونه‌ای با محاسبه کردن میزان اطلاع فیشر آن‌ها تحت توزیع‌های آماری مختلف مقایسه کردند. آن‌ها نتیجه گرفتند که میزان اطلاع فیشر نهفته در آماره‌های رکوردی حاصل از هر دو طرح بستگی به توزیع جامعه‌ی آماری دارد. بدیهی است که به دلیل مستقل بودن دنباله‌های متغیرهای تصادفی در شکل (۲.۱)،

دلیلی ندارد که ترتیبی بین رکوردهای حاصل از دنباله‌های متفاوت وجود داشته باشد، امینی و بالا کریشان (۲۰۱۳) مجموعه رکوردهای شکل ۲.۱ را مرتب کردند و پس از بررسی ویژگی‌های توزیعی رکوردهای مرتب شده، یک فاصله‌ی اطمینان ناپارامتری برای چندک‌های جامعه‌ی آماری و هم چنین فاصله پیش‌بینی برای آماره‌های رکوردی آینده براساس آن‌ها به دست آوردند. در بخش ۵.۱ مشاهده شد که در روش نمونه‌گیری مجموعه‌ی رتبه‌دار، نمونه‌های تصادفی جامعه‌ی مورد مطالعه یا متغیر فرعی آن به صورت چشمی یا قضاوتی مرتب می‌شوند و در نهایت واحدهای انتخاب شده به صورت دقیق اندازه‌گیری می‌شوند. از طرفی، همانطور که در ساختار نمونه‌گیری مجموعه‌ی رتبه‌دار در شکل ۲.۱ مشاهده می‌شود، عناصر روی قطر اصلی، یعنی $X_{1,1}, \dots, X_{m,m}$ ، به عنوان یک نمونه‌ی مجموعه‌ی رتبه‌دار انتخاب می‌شوند. در ادامه، یک الگو برای جمع‌آوری آماره‌های رکوردی برگرفته از روش چندنمونه‌ای و نمونه‌گیری مجموعه‌ی رتبه‌دار معرفی خواهیم کرد و به همین دلیل آن را «نمونه‌گیری مجموعه‌ی رتبه‌دار رکوردی» می‌نامیم. روش نمونه‌گیری مجموعه‌ی رتبه‌دار رکوردی برای اولین بار در رساله صالحی (۱۳۹۴) مطرح و مورد بررسی قرار گرفته است. قبل از تعریف طرح فوق در آماره‌های رکوردی، موارد زیر را در نظر بگیرید:

(۱) زمانی را در نظر بگیرید که تنها مشاهده‌ی ثبت شده، آخرین رکورد باشد و رکوردهای قبلی به هر دلیلی ثبت نشده باشند. به طور مثال در برخی داده‌های ورزشی و داده‌های هواشناسی این اتفاق رخ می‌دهد. مثلاً در طول هر ماه، رکورد دما ممکن است چندین بار شکسته شود، اما ممکن است فقط آخرین رکورد ثبت شود. همین رخداد می‌تواند برای میزان بارندگی ماهانه نیز اتفاق بیفتد. به عنوان مثال، شکل زیر نشان دهنده‌ی رکوردهای جمع‌آوری شده از میزان دمای ماه بهمن در m سال می‌باشد

$$\begin{array}{lcl}
 1: & NA & \dots \dots NA & U_{(k_1)}1 \\
 2: & NA & \dots NA & U_{(k_2)}2 \\
 & \vdots & & \\
 m: & NA & \dots \dots NA & U_{(k_m)}m
 \end{array} \quad (3.1)$$

همان‌طور که مشخص است، به عنوان مثال رکورد دمای ماه بهمن در سال دوم k_2 بار شکسته شده است و تنها آخرین رکورد در دسترس است. رکوردهای جمع‌آوری شده در شکل بالا مستقل از هم هستند. حال مثال دیگری در این زمینه و در مبحث سیستم‌ها ارائه می‌شود.

مثال ۱.۶.۱. یک سیستم قابل تعمیر موازی با تعمیرات مینیمال را در نظر بگیرید. پس از اعمال تعمیر مینیمال بر روی یک مؤلفه، آن مؤلفه دقیقاً به شرایط قبل از لحظه خرابی باز می‌گردد. فرض کنید این سیستم موازی از m مؤلفه‌ی همسان و مستقل تشکیل شده است که دارای طول عمری با تابع توزیع احتمال $F(\cdot)$ هستند. هم‌چنین فرض کنید که سیستم طوری طراحی شده است که i امین مؤلفه‌ی آن $k_i - 1$ بار قابل تعمیر است ($k_i \geq 1$)، یعنی مؤلفه i ام پس از k_i امین خرابی دیگر قابل تعمیر نیست. در نتیجه، $(k_1 + \dots + k_m)$ امین خرابی منجر به خرابی کل سیستم می‌شود. بنابراین طول عمر سیستم برابر خواهد بود با $\max\{T_1, \dots, T_m\}$ ، به‌طوری‌که T_i طول عمر i امین مؤلفه

است. از طرف دیگر، فرآیند تعمیر مینیمال و رکوردهای بالای مستخرج از مشاهدات مستقل و هم‌توزیع، دارای توزیع مشابه هستند. بنابراین T_i و $U_{(k_i)i}$ در طرح مزبور هم توزیع هستند. از آنجا که طول عمر کل سیستم برابر است با $\max\{U_{(k_1)1}, \dots, U_{(k_m)m}\}$ ، بنابراین نیاز است تا به تمامی $U_{(k_i)i}$ دسترسی داشته باشیم.

(۲) حالتی که اندازه‌گیری دقیق صفت موردنظر به هر دلیلی مثل گران بودن، از بین رفتن کامل واحدهای تحت آزمایش یا سخت بودن به راحتی امکان پذیر نباشد اما بتوان به صورت چشمی، قضاوتی یا توسط وسیله ای رکوردهای صفت موردنظر را مشخص کرد. به طور دقیق تر، m دنباله‌ی مستقل از متغیرهای تصادفی را در نظر بگیرید. از دنباله ی اول، اولین مشاهده را به عنوان اولین رکورد در نظر می‌گیریم. از دنباله ی دوم، به صورت چشمی دومین رکورد را مشخص می‌کنیم و این عمل را تا m امین دنباله انجام می‌دهیم و m امین رکورد را از آن دنباله به صورت چشمی اندازه‌گیری می‌کنیم، در این صورت m رکورد در اختیار خواهیم داشت. شکل (۴.۱) را ببینید.

$$\begin{aligned} 1: & X_{(1)1}, X_{(2)1}, \dots \rightarrow U_{(1)1} \\ 2: & X_{(1)2}, X_{(2)2}, \dots \rightarrow U_{(2)2} \\ & \vdots \\ m: & X_{(1)m}, X_{(2)m}, \dots \rightarrow U_{(m)m} \end{aligned} \quad (4.1)$$

(۳) حالتی را در نظر بگیرید که شرایط برقراری هیچ کدام از دو مورد بالا مهیا نباشد. فرض کنید m دنباله مستقل از متغیرهای تصادفی در اختیار است و نمونه‌گیری در هر دنباله به نحوی است که تنها رکوردها ثبت می‌شوند (مثال ۱.۴.۱ و ۲.۴.۱ را ببینید). هم چنین فرض کنید که i امین نمونه‌گیری دنباله‌ای، زمانی متوقف می‌شود که i امین رکورد ثبت شود. از دنباله‌ی i ام تنها i امین رکورد را برای استنباط در رابطه با صفت‌های جامعه‌ی آماری حفظ می‌کنیم. از دیدگاه دیگر، این روش مشابه روش چندنمونه‌ای سامانیکو و وایتاکر (۱۹۸۸) است به شرطی که $k_i = i$ باشد و فقط i امین رکورد از دنباله ی i ام را در نظر بگیریم. شکل زیر به طور دقیق‌تر این روش را نشان می‌دهد

$$\begin{aligned} 1: & X_{(1)1}, X_{(2)1}, \dots \rightarrow U_{(1)1} \\ 2: & X_{(1)2}, X_{(2)2}, \dots \rightarrow U_{(1)2} \quad U_{(2)2} \\ & \vdots \quad \quad \quad \vdots \quad \quad \quad \ddots \\ m: & X_{(1)m}, X_{(2)m}, \dots \rightarrow U_{(1)m} \quad U_{(2)m} \quad \dots \quad U_{(m)m} \end{aligned} \quad (5.1)$$

بنابراین در این حالت نیز m رکورد $(U_{(1)1}, \dots, U_{(m)m})$ را در اختیار خواهیم داشت. لازم به ذکر است که در عمل نیز ممکن است تنها آخرین رکورد در هر دنباله در دسترس باشد یا حتی تنها آخرین رکورد برای استنباط درباره‌ی پارامتر مجهول جامعه‌ی آماری کافی باشد.

تعریف ۱.۶.۱. فرض کنید $U_{(m)}^{RRSS} = (U_{(1)1}, \dots, U_{(m)m})$ آماره‌های رکوردی به‌دست آمده توسط یکی از حالات بالا باشد. مشاهده می‌شود که ساختار تمامی آن‌ها مشابه نمونه‌گیری مجموعه‌ی رتبه‌دار معرفی شده در شکل ۱.۱ است. اگر قرار دهیم $U_{i,i} = U_{(i)i}$ ، $i = 1, \dots, m$ ، آن‌گاه بردار $(U_{1,1}, U_{2,2}, \dots, U_{m,m})$ را یک نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی ($RRSS$) به اندازه‌ی m می‌گوییم.

صالحی و همکاران (۲۰۱۵)، برآوردهایی از قابلیت اعتماد تنش-مقاومت را براساس نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی در توزیع نمایی یک پارامتری بررسی نمودند. یکی از اهداف آن‌ها مقایسه‌ی $RRSS$ و آماره‌های رکوردی معمولی بر مبنای دو دیدگاه پارامتری و ناپارامتری بود. برای این منظور در مبحث ناپارامتری پیش‌بینی آماره‌های ترتیبی و رکوردهای آینده را مطالعه کردند و در مبحث پارامتری، با در نظر گرفتن خانواده توزیع‌های نرخ خطر متناسب میزان اطلاع فیشر، معیارهای کارایی نسبی و پیتمن-نزدیکی برآوردهای حاصل از دو طرح را مقایسه کردند. در نهایت نتیجه گرفتند برآوردهای حاصل از طرح $RRSS$ بسیار کاراتر از رکوردهای معمولی است و اطلاع فیشر نهفته در یک نمونه‌ی به اندازه n از $RRSS$ معادل یک نمونه‌ی به اندازه $n(n+1)/2$ براساس رکوردهای معمولی است.

۱.۶.۱ ویژگی‌های حاصل از روش $RRSS$

فرض کنید جامعه‌ی آماری تحت بررسی دارای تابع چگالی $f(\cdot)$ و تابع توزیع احتمال $F(\cdot)$ بوده و X_1 یک متغیر تصادفی مشاهده شده از آن باشد. هم چنین $f_{i,i}(x)$ و $F_{i,i}(x)$ را به ترتیب به عنوان تابع چگالی و تابع توزیع احتمال $L_{i,i}$ در نظر بگیرید. در این صورت براساس قرارداد داریم $L_{1,1} \stackrel{d}{=} X_1$. اما برای $i > 1$ ، با توجه به این که $L_{i,i}$ هم توزیع با i امین رکورد معمولی پایین است داریم

$$f_{i,i}(l) = \frac{\{-\log F(l)\}^{i-1}}{(i-1)!} f(l), \quad u \in S_X, \quad (6.1)$$

و

$$F_{i,i}(l) = 1 - F(l) \sum_{t=0}^{i-1} \frac{\{-\log F(l)\}^t}{t!}, \quad l \in \mathbb{R} \quad (7.1)$$

با توجه به استقلال $l_{i,i}$ ها و (۶.۱)، تابع چگالی احتمال توأم بردار $L_{(m)}^{RRSS}$ برابر است با

$$f_{L_{(m)}^{RRSS}}(l_{(m)}^{RRSS}) = \prod_{i=1}^m \frac{\{-\log F(l_{i,i})\}^{i-1}}{(i-1)!} f(l_{i,i}), \quad (8.1)$$

که در آن $I_{(m)}^{RRSS} = (l_{1,1}, \dots, l_{m,m})$ مقدار مشاهده شده‌ی $L_{(m)}^{RRSS}$ است. اگر $U_{(m)}^{RRSS} = (U_{1,1}, U_{2,2}, \dots, U_{m,m})$ را به عنوان یک $RRSS$ بالا به اندازه‌ی m در نظر بگیریم، آن‌گاه با جایگزین کردن $\bar{F}$ به جای F در روابط (۶.۱) و (۷.۱)، به ترتیب تابع چگالی احتمال

حاشیه‌ای و تابع توزیع احتمال حاشیه‌ای $U_{i,i}$ و هم چنین تابع چگالی احتمال توأم $U_{(m)}^{RRSS}$ حاصل می‌شوند. به طوری که $\bar{F}(\cdot) = 1 - F(\cdot)$.

۷.۱ مقدمه‌ای بر روش بوت استرپ

در استنباط آماری در حالت ایده‌آل با جامعه‌ای سروکار داریم که توزیع آن معلوم و تعداد داده‌ها به اندازه‌ی کافی زیاد می‌باشد. در صورتی که توزیع جامعه‌ی آماری نامعلوم ولی همچنان تعداد داده‌ها به اندازه‌ی کافی زیاد باشد آن‌گاه با استفاده از قضیه حد مرکزی به نتایج مطلوب دست می‌یابیم. بدترین حالت زمانی است که توزیع جامعه‌ی آماری نامعلوم و نمونه‌ی مشاهده شده کوچک باشد. حال این سؤال مطرح است که در این شرایط چگونه می‌توان استنباط خوب و معتبری ارائه داد.

در واقع تنها اطلاعاتی که از جامعه‌ی آماری در اختیار داریم، نمونه‌ی مشاهده شده‌ی $x = (x_1, \dots, x_n)$ به اندازه‌ی کوچک n از جامعه‌ای با تابع توزیع احتمال مجهول $F(\cdot)$ می‌باشد و هدف برآورد پارامتر $\theta = t(F)$ براساس x است. برای این منظور یک برآوردگر را به صورت $\hat{\theta} = T(\hat{F}_n)$ محاسبه می‌کنیم، به طوری که $\hat{F}_n$ همان تابع توزیع احتمال تجربی مشاهدات x است. اما سؤال این است که میزان دقت برآوردگر $\hat{\theta}$ چقدر است؟ افرون^{۴۳} (۱۹۷۹)، یک روش بازنمونه‌گیری برای برآورد انحراف معیار $\hat{\theta}$ تحت عنوان نمونه‌گیری بوت استرپ ارائه داد. در ادامه پژوهش‌های بیشتری در این رابطه ارائه گردید که کتاب جامع مقدمه‌ای بر بوت استرپ از افرون و تیبشیرانی^{۴۴} (۱۹۹۳)، تقریباً تمام پژوهش‌های انجام شده تا سال ۱۹۹۳ را پوشش می‌دهد. در عین حال که روش بوت استرپ ساختار ساده‌ای دارد ولی نیازمند کامپیوترهای پرقدرت برای شبیه‌سازی است. جهت محاسبه فواصل اطمینان بوت استرپی به کمک نرم افزار R می‌توانید به فصل پنجم کتاب ریزو^{۴۵} (۲۰۰۷)، مراجعه کنید. بوت استرپ در دو حالت پارامتری و ناپارامتری به کار گرفته می‌شود که در اینجا به شرح هر یک می‌پردازیم.

● **بوت استرپ ناپارامتری:** حالتی را در نظر بگیرید که یک نمونه‌ی مشاهده شده‌ی $x = (x_1, \dots, x_n)$ با اندازه کوچک n از جامعه‌ای با تابع احتمال مجهول $F(\cdot)$ در اختیار داریم. برای برآورد پارامتر $\theta = t(F)$ براساس x بایستی برآوردگر $\hat{\theta} = T(\hat{F}_n)$ را محاسبه نماییم به طوری که $\hat{F}_n$ همان تابع توزیع تجربی مشاهدات x می‌باشد.

● **بوت استرپ پارامتری:** چنانچه توزیع جامعه‌ی آماری مشخص و پارامترهای آن مجهول باشد، ابتدا پارامترهای مجهول را به کمک یکی از روش‌های برآوردیابی کلاسیک به دست می‌آوریم. سپس با جایگذاری مقدار برآورد شده به جای پارامتر مجهول، از خود توزیع

^{۴۳}Efron

^{۴۴}Tibshirani

^{۴۵}Rizzo

الگوریتم ۱ الگوریتم عمومی بوت‌استرپ ناپارامتری

۱. نمونه تصادفی $x^* = (x_1^*, \dots, x_n^*)$ را از تابع توزیع تجربی F_n تولید می‌کنیم، بدیهی است نمونه‌ی بوت‌استرپی x^* یک نمونه تصادفی با جایگذاری به اندازه n از x می‌باشد.
۲. متناظر با هر نمونه بوت‌استرپی مثل x^* ، می‌توان یک مشاهده‌ی بوت‌استرپی برای $\hat{\theta}$ به صورت $\hat{\theta}^* = T(x^*)$ به دست آورد.
۳. مرحله ۱ و ۲ را B مرتبه تکرار کنید تا B نمونه‌ی بوت‌استرپی $(x_1^*, \dots, x_B^*)$ که هر کدام شامل n داده‌ی تولید شده‌ی با جایگذاری از x هستند را انتخاب کنید (مقدار B در بازه‌ی ۲۰۰-۲۵۰ انتخاب می‌شود).
۴. مقدار مشاهده‌ی بوت‌استرپی برای $\hat{\theta}$ را به ازای هر $b = 1, \dots, B$ ، به صورت $\hat{\theta}_{(b)}^* = T(x_b^*)$ محاسبه کنید.
۵. با استفاده از B تکرار بوت‌استرپی، برآورد آماره مورد نظر را مانند میانگین و انحراف معیار به دست آورید.

جامعه‌ی آماری برای بازنمونه‌گیری استفاده می‌کنیم. در حقیقت، تفاوت بوت‌استرپ پارامتری و ناپارامتری تنها در انتخاب روش تولید نمونه‌های بوت‌استرپی می‌باشد. به‌طور دقیق‌تر، فرض کنید نمونه‌ی x را از جامعه‌ای با تابع توزیع احتمال $F(., \psi)$ مشاهده کرده‌ایم. ابتدا نمونه‌ی بوت‌استرپی x^* را از توزیع $F(., \hat{\psi})$ تولید می‌کنیم که $\hat{\psi}$ برآورد پارامتر مجهول ψ براساس یکی از روش‌های کلاسیک برآوردیابی می‌باشد. حال با انجام مراحل ۴ تا ۵ الگوریتم (۱) بوت‌استرپ آماره مورد نظر به دست می‌آید.

۱.۷.۱ فواصل اطمینان بوت‌استرپ

در این بخش، چند فاصله اطمینان بوت‌استرپی که در کتاب افرون و تیشیرانی (۱۹۹۴) استفاده شده‌است را معرفی می‌کنیم. بدین منظور فرض کنید x یک نمونه مشاهده شده از جامعه‌ای با پارامتر مجهول θ باشد. الگوریتم‌های زیر فواصل اطمینان بوت‌استرپی با ضریب تقریبی $1 - \alpha$ تولید می‌کنند.

۸.۱ خانواده‌های نرخ خطر

فرض کنید متغیر تصادفی مطلقاً پیوسته X دارای تابع توزیع تجمعی

$$G(x; \theta, \alpha) = [F(x; \theta)]^\alpha, \quad x \in \mathbb{R}, \quad \theta \in \Theta, \quad \alpha > 0,$$

الگوریتم ۲ فاصله اطمینان بوت استرپ پایه

۱. B نمونه بوت استرپی $x_1^*, \dots, x_B^*$ از بردار مشاهده شده x تولید کنید.
۲. برای هر یک از نمونه‌های بوت استرپی فوق، مشاهده بوت استرپی $d_{(b)}^* = \theta_{(b)ML}^* - \hat{\theta}_{ML}$ را محاسبه کنید. $b = 1, \dots, B$
۳. آن گاه بازه $(\hat{\theta}_{ML} - d_{1-\alpha/2}^*, \hat{\theta}_{ML} - d_{\alpha/2}^*)$ فاصله اطمینان بوت استرپ پایه با ضریب اطمینان تقریبی $1 - \alpha$ برای پارامتر θ می‌باشد، به‌طوری‌که d_α^* چندک α ام بردار $d^* = (d_{(1)}^*, \dots, d_{(B)}^*)^\top$ می‌باشد.

الگوریتم ۳ فاصله اطمینان بوت استرپ صدکی

۱. B نمونه بوت استرپی $x_1^*, \dots, x_B^*$ از بردار مشاهده شده x تولید کنید.
۲. برای هر یک از نمونه‌های بوت استرپی فوق، مشاهده بوت استرپی $\theta_{(b)}^*$ را محاسبه کنید.
۳. اگر $\theta_{(\alpha)}^*$ بیان گر چندک α ام بردار $(\theta_{(1)}^*, \dots, \theta_{(B)}^*)$ باشد، آن گاه بازه $(\theta_{\alpha/2}^*, \theta_{1-\alpha/2}^*)$ فاصله اطمینان بوت استرپ صدکی با ضریب اطمینان تقریبی $1 - \alpha$ برای پارامتر θ است.

الگوریتم ۴ فاصله اطمینان t -بوت استرپ

۱. B نمونه بوت استرپی $x_1^*, \dots, x_B^*$ از بردار مشاهده شده x تولید کنید.
۲. برای هر نمونه بوت استرپی، $t_{(b)}^* = \frac{\hat{\theta}_{(b)}^* - \hat{\theta}}{\hat{\sigma}(\hat{\theta}^*)}$ را محاسبه کنید. که در آن $\hat{\sigma}(\hat{\theta}^*)$ می‌تواند هر برآورگر استاندارد مانند برآوردگر درست‌نمایی ماکسیمم $\sigma(\hat{\theta}^*)$ یا برآورد بوت استرپی خطای معیار باشد. که برای به‌دست آوردن برآورد بوت استرپی انحراف معیار کافیست x^* را به‌جای x قرار داده و نمونه‌های بوت استرپی را از x^* تولید کنیم.
۳. آن گاه بازه $(\hat{\theta} - t_{1-\alpha/2}^*, \hat{\theta} - t_{\alpha/2}^*)$ فاصله اطمینان t -بوت استرپ با ضریب اطمینان تقریبی $1 - \alpha$ برای پارامتر θ می‌باشد، به‌طوری‌که t_α^* چندک α ام بردار $t^* = (t_{(1)}^*, \dots, t_{(B)}^*)^\top$ می‌باشد.

باشد. خانواده‌ای از توزیع‌های $\{G(x; \theta, \alpha), \theta \in \Theta, \alpha > 0\}$ ، مدل نرخ خطر معکوس متناسب^{۴۶} $(PRHR)$ با توزیع پایه $F(x; \theta)$ نامیده می‌شود، که در آن $F(x; \theta)$ یک تابع توزیع پیوسته دلخواه است. همچنین $G(x; \theta, \alpha)$ را تابع توزیع تعمیم‌یافته از تابع توزیع پایه $F(x; \theta)$ می‌نامند. متغیر تصادفی X از این خانواده دارای تابع چگالی احتمال

$$g(x; \theta, \alpha) = \alpha f(x; \theta) [F(x; \theta)]^{\alpha-1}, \quad x \in \mathbb{R}, \quad \theta \in \Theta, \quad \alpha > 0$$

است. خانواده $PRHR$ توسط گوپتا^{۴۷} و همکاران (۱۹۹۸) معرفی شد. همچنین گوپتا و همکاران (۲۰۰۱)، نتایجی را براساس خانواده فوق به‌دست آوردند. با توجه به کاربردهای وسیع خانواده توزیع‌های $PRHR$ در تحلیل‌های بقا و قابلیت اعتماد بخصوص در مطالعات مربوط به سیستم‌های موزی و داده‌های طول عمر سانسور شده از چپ، این مفهوم به طور گسترده‌ای مورد مطالعه واقع شده است. برای مثال می‌توان به سنگوپتا و ناندا^{۴۸} (۱۹۹۹) و چاندرا و روی^{۴۹} (۲۰۰۱)، اشاره کرد. خانواده $PRHR$ بسیاری از مدل‌های طول عمر، از جمله توزیع نمایی تعمیم‌یافته، وایبول، رایلی، پارتو، لوجیستیک، رایلی معکوس، بور نوع III، گامبل نمایی و توزیع‌های توانی را شامل می‌شود. برای آشنایی بیشتر با خانواده فوق به گوپتا و ناندا^{۵۰} (۲۰۰۱) و گوپتا و گوپتا^{۵۱} (۲۰۰۷) مراجعه نمایید. همچنین برای متغیر تصادفی مطلقاً پیوسته X با تابع توزیع تجمعی

$$G(x; \theta, \alpha) = 1 - [1 - F(x; \theta)]^\alpha, \quad x \in \mathbb{R}, \quad \theta \in \Theta, \quad \alpha > 0,$$

خانواده‌ای از توابع توزیع $\{G(x; \theta, \alpha), \theta \in \Theta, \alpha > 0\}$ مدل نرخ خطر متناسب با توزیع پایه $F(x; \theta)$ مدل نرخ خطر متناسب^{۵۲} (PHR) نامیده می‌شود. بنابراین متغیر تصادفی X دارای تابع چگالی احتمال به صورت

$$g(x; \theta, \alpha) = \alpha f(x; \theta) [1 - F(x; \theta)]^{\alpha-1}, \quad x \in \mathbb{R}, \quad \theta \in \Theta, \quad \alpha > 0$$

است. این خانواده شامل چند توزیع طول عمر مشهور از جمله بور نوع XII، لوماکس و وایبول است (احمدی و همکاران، ۲۰۰۸ و ۲۰۰۹؛ مارشال و الکین^{۵۳} ۲۰۰۷). در هر دو مدل، θ می‌تواند بردار پارامتر توزیع پایه باشد ولی α پارامتر عددی تابع توزیع $G(x; \theta, \alpha)$

^{۴۶}Proportional reversed hazard rate model

^{۴۷}Gupta

^{۴۸}Sengupta and Nanda

^{۴۹}Chandra and Roy

^{۵۰}Gupta and Nanda

^{۵۱}Gupta and Gupta

^{۵۲}Proportional hazard rate model

^{۵۳}Marshall and Olkin

است. با توجه به کاربرد وسیعی که خانواده‌های فوق در مدل‌بندی و تحلیل داده‌های طول عمر دارند، محققان کلاس‌های مختلفی از توزیع‌های این خانواده‌ها را با $F(x; \theta)$ های مختلف معرفی و به‌طور وسیعی مورد مطالعه قرار دادند.

۹.۱ روش‌های شبیه‌سازی

در این بخش، دو روش شبیه‌سازی را که در حل مسائل آماری این رساله مورد استفاده قرار می‌گیرند ارائه می‌دهیم. در حقیقت از دیدگاه کلاسیک علاقه‌مند به شبیه‌سازی یک توزیع نمونه‌ای که برخاسته از یک مدل پیچیده است و از دیدگاه بیزی علاقه‌مند به محاسبه برخی مشخصه‌های توزیع پسین هستیم. در ریاضی زنجیر مارکوف^{۵۴}، یک فرآیند تصادفی با خاصیت مارکوفی است. در خاصیت مارکوفی، به شرط داده شدن وضعیت حال، وضعیت آینده مستقل از وضعیت گذشته است.

عمده‌ترین محدودیتی که در انجام اکثر روش‌های بیزی روبرو می‌شویم، به‌دست آوردن توزیع پسین است که اغلب نیازمند انتگرال‌گیری از توابعی با ابعاد بالا است. چندین روش برای حل این‌گونه انتگرال‌ها توسط اشخاصی چون اسمیت^{۵۵} (۱۹۹۱)، تنر^{۵۶} (۱۹۹۶) ارائه شده‌است که یکی از این روش‌ها زنجیر مارکوف مونت کارلو^{۵۷} (MCMC) است که اساس آن شبیه‌سازی درستی از توزیع پیچیده مورد نظر می‌باشد. در واقع روش MCMC همان روش مونت کارلو^{۵۸} (MC) است با این تفاوت که در آن نمونه‌گیری براساس ویژگی زنجیر مارکوف صورت می‌گیرد. در اینجا دو روش MC و مونت کارلو با نمونه‌گیری از نقاط مهم^{۵۹} شرح داده می‌شود.

۱.۹.۱ روش مونت کارلو

فرض کنید علاقه‌مندیم انتگرال

$$I = \int_a^b g(\theta) d\theta$$

را محاسبه کنیم که $g(\theta)$ یک تابع هموار است. انتگرال فوق را می‌توان به صورت زیر نوشت:

$$I = \int_a^b [(b-a)g(\theta)] \frac{1}{b-a} d\theta$$

^{۵۴}Markov Chain

^{۵۵}Smith

^{۵۶}Tanner

^{۵۷}Markov Chain Monte Carlo

^{۵۸}Monte Carlo

^{۵۹}Monte carlo with important sampling

که مانند محاسبه امید ریاضی $[(b-a)g(\theta)]$ نسبت به توزیع یکنواخت روی (a, b) است. یعنی

$$I = E_{U(a,b)}[(b-a)g(\theta)]$$

که با استفاده از روش گشتاوری

$$\hat{I} = \frac{1}{n} \sum_{i=1}^n [(b-a)g(\theta_i)]$$

که $\theta_1, \dots, \theta_n$ یک نمونه تصادفی از توزیع $U(a, b)$ است.

۱۰.۱ روش مونت کارلو با نمونه‌گیری از نقاط مهم

در بسیاری از موارد می‌خواهیم $\mu = E(f(x))$ را محاسبه کنیم که $f(x)$ در خارج از ناحیه A که $p(X \in A)$ کوچک است، تقریباً صفر است. مجموعه A ممکن است حجم کمی داشته باشد یا ممکن است در انتهای توزیع X باشد. یک نمونه ساده مونت کارلویی از توزیع X ممکن است حتی یک نقطه در ناحیه نداشته‌باشد. چنین مشکلاتی در فیزیک انرژی بالا، استنباط بیزی و شبیه‌سازی رویدادهای نادر برای تأمین مالی و بیمه پیش می‌آید.

به‌طور شهودی روشن است که باید نمونه‌هایی از ناحیه‌های جالب و مهم را به‌دست آوریم. ما این کار را با نمونه‌گیری از توزیعی انجام می‌دهیم که منطقه‌ی مهم را سنگین کند، به همین دلیل این روش، نمونه‌گیری از نقاط مهم نامیده می‌شود. فرض کنید انتگرال مورد علاقه به‌صورت امید ریاضی یک تابع g نسبت به چگالی p است.

$$I = \int g(x)p(x)dx = \int g(x)\frac{p(x)}{h(x)}h(x)dx$$

که $h(x)$ یک تابع مثبت برای x هایی است که $p(x) > 0$ و $\int h(x)dx = 1$. بنابراین $h(x)$ یک تابع چگالی است. یک روش جایگزین با استفاده از برآورد گشتاوری، I را به‌صورت زیر برآورد می‌کند:

$$\bar{I} = \frac{1}{n} \sum_{i=1}^n g(X_i)w(X_i)$$

که در آن $w(X_i) = \frac{p(X_i)}{h(X_i)}$ که به ازای $i = 1, 2, \dots, n$ ، $X_i \sim h(X)$ ، چگالی نمونه‌گیری با اهمیت نامیده می‌شود.

در حقیقت این روش برای کاهش واریانس به‌کار می‌رود. به‌طوری‌که با انتخاب $g(x)w(x)$ تقریباً ثابت در رابطه زیر

$$Var(\bar{I}) = \frac{1}{n} \int (g(x)w(x) - I)^2 h(x)dx$$

می‌توانیم تا جایی که می‌خواهیم واریانس را کوچک کنیم.

فصل ۲

برآورد پارامتر $\mathcal{R}$ در توزیع نمایی تعمیم یافته با استفاده از $RRSS$

۱.۲ مقدمه

همان‌طور که در بخش ۲.۱ بیان شد، یکی از مفاهیم مهم در تحلیل‌های قابلیت اعتماد، مسأله‌ی برآورد پارامتر تنش-مقاومت می‌باشد. مقالات متعددی به مسائل احتمالاتی خاص، مربوط به ارزیابی $\mathcal{R} = Pr(X < Y)$ و ساختن برآوردهای کارآمد و قابل اطمینانی از پارامتر موردنظر بر مبنای مقادیر نمونه‌ای با فرض‌های مختلف از توزیع‌های X و Y پرداخته‌اند. به عنوان مثال برخی از نویسندگان با در نظر گرفتن داده‌های کامل به برآورد پارامتر تنش-مقاومت پرداخته‌اند. از جمله اون^۱ و همکاران (۱۹۶۴)، کران بالایی $\mathcal{R}$ را در حالت ناپارامتری محاسبه نمودند. کندو و گوپتا^۲ (۲۰۰۵) با در نظر گرفتن توزیع نمایی تعمیم یافته استنباط‌هایی در خصوص $\mathcal{R}$ ارائه دادند. رقب^۳ و همکاران (۲۰۰۸) با در نظر گرفتن توزیع نمایی تعمیم یافته سه پارامتری به

^۱Owen

^۲Kundu and Gupta

^۳Raqab

استنباط پارامتر تنش-مقاومت پرداخته‌اند. از سوی دیگر در سال‌های اخیر، استفاده از داده‌های ناقص در تحلیل پارامتر R مورد توجه پژوهش‌گران زیادی واقع شد. به عنوان مثال، اصغرزاده و همکاران (۲۰۱۱)، لئو و تسای^۴ (۲۰۱۲)، ساراک‌اوغلو^۵ و همکاران (۲۰۱۲) و بصیرت و همکاران (۲۰۱۳)، با در نظر گرفتن نمونه‌هایی از سانسور پیش‌رونده به برآورد پارامتر تنش-مقاومت پرداختند. همچنین بکلیزی^۶ (۲۰۰۸) پارامتر R را با در نظر گرفتن توزیع نمایی تعمیم‌یافته و براساس مقادیر رکوردی مورد بررسی قرار داد. موتلاک^۷ و همکاران (۲۰۱۰) استنباط‌هایی از پارامتر تنش-مقاومت با استفاده از مجموعه‌ی رتبه‌دار (RSS) در توزیع نمایی ارائه نمودند. همان‌طور که در بخش ۶.۱ بیان شد، صالحی و احمدی (۲۰۱۴)، طرحی تحت عنوان $RRSS$ در مبحث رکوردها پیشنهاد دادند که با استفاده از آن می‌توان به نتایج دقیق‌تری دست یافت. صالحی و احمدی در همان سال با استفاده از این طرح به پیش‌بینی آماره‌های ترتیبی و رکوردهای آینده پرداختند و در سال ۲۰۱۵ نیز پارامتر تنش-مقاومت را به صورت نقطه‌ای و فاصله‌ای براساس $RRSS$ مشاهده شده از توزیع نمایی برآورد نمودند. در سال‌های اخیر مقالات متعددی با در نظر گرفتن طرح $RRSS$ ارائه گردید. از جمله اصغرزاده و همکاران (۲۰۱۷)، اندازه اطلاع R را به‌دست آوردند. پل و توماس^۸ (۲۰۱۷) با در نظر گرفتن حالتی که اندازه‌گیری متغیرهای موردنظر پرهزینه یا عملاً غیرممکن می‌باشد، متغیرهای همراه $RRSS$ را ارائه نمودند. صفریان و همکاران (۲۰۱۹)، با در نظر گرفتن این طرح برآوردهای بهبود یافته پارامتر تنش-مقاومت را به‌دست آوردند. حال در این فصل قصد داریم پارامتر تنش-مقاومت R را به صورت نقطه‌ای و فاصله‌ای براساس $RRSS$ مشاهده شده از توزیع نمایی تعمیم یافته برآورد کنیم. بنابراین، در بخش ۲.۲، دو برآورد نقطه‌ای درست‌نمایی ماکسیمم (ML) و برآورد به طور یکنواخت با کمترین واریانس ($UMVU$) را برای R به‌دست می‌آوریم و برآوردهای فوق را براساس معیار کارایی نسبی مقایسه می‌کنیم. در ادامه و در بخش ۳.۲، فاصله‌ی اطمینان دقیق و چند فاصله‌ی اطمینان تقریبی بر پایه‌ی بوت‌استرپ پارامتری برای R به‌دست می‌آوریم. در نهایت، در بخش‌های ۴.۲ و ۵.۲، به ترتیب با استفاده از یک مطالعه‌ی شبیه‌سازی و همچنین یک مثال واقعی، کارایی برآوردهای به‌دست آمده در بخش‌های ۲.۲ و ۳.۲ را بررسی می‌کنیم.

۲.۲ برآورد نقطه‌ای پارامتر تنش-مقاومت

توزیع نمایی تعمیم‌یافته به ترتیب با تابع چگالی و تابع توزیع احتمال

$$f_{GE}(x; \theta, \lambda) = \theta \lambda e^{-\lambda x} (1 - e^{-\lambda x})^{\theta-1}, \quad (1.2)$$

^۴Lio and Tsai

^۵Saracoglu

^۶Baklizi

^۷Mutllak

^۸Paul and Thomas

و

$$F_{GE}(x; \theta, \lambda) = (1 - e^{-\lambda x})^\theta \quad x > 0, \quad (2.2)$$

در اختیار داریم، که در آن $\theta > 0$ و $\lambda > 0$ به ترتیب پارامترهای شکل و نرخ هستند.

قضیه ۱.۲.۲. فرض کنید $X \sim GE(\theta_1, \lambda_1)$ و $Y \sim GE(\theta_2, \lambda_2)$ دو متغیر تصادفی مستقل باشند و $\mathcal{R} = Pr(X < Y)$ همان پارامتر تنش-مقاومت باشد، در این صورت به راحتی ثابت می‌شود که

$$\mathcal{R} = 1 - \int_0^1 \theta_1 (1-t)^{\theta_1-1} (1-t^{\frac{\lambda_2}{\lambda_1}})^{\theta_2} dt. \quad (3.2)$$

برهان. با استفاده از تبدیل $U = Y - X$ و $V = X$ داریم

$$\begin{aligned} \mathcal{R} = Pr(X < Y) &= \int_0^\infty f_U(u) du = \int_0^\infty \int_0^\infty f_{U,V}(u, v) dv du \\ &= \int_0^\infty \int_0^\infty \theta_1 \lambda_1 e^{-\lambda_1 v} (1 - e^{-\lambda_1 v})^{\theta_1-1} \theta_2 \lambda_2 e^{-\lambda_2(u+v)} (1 - e^{-\lambda_2(u+v)})^{\theta_2-1} dudv, \end{aligned} \quad (4.2)$$

حال با جایگذاری $z = 1 - e^{-\lambda_2(u+v)}$ در (۴.۲)، داریم

$$\begin{aligned} \mathcal{R} &= \int_0^\infty \left(\int_{1-e^{-\lambda_2 v}}^1 \theta_2 z^{\theta_2-1} dz \right) \theta_1 \lambda_1 e^{-\lambda_1 v} (1 - e^{-\lambda_1 v})^{\theta_1-1} dv \\ &= \int_0^\infty \theta_1 \lambda_1 e^{-\lambda_1 v} (1 - e^{-\lambda_1 v})^{\theta_1-1} [1 - (1 - e^{-\lambda_2 v})]^{\theta_2} dv. \end{aligned}$$

□

با استفاده از تبدیل $t = e^{-\lambda_1 v}$ اثبات کامل می‌شود.

اگر $\mathbf{r} = (r_{1,1}, r_{2,2}, \dots, r_{n,n})^\top$ مقدار مشاهده شده از بردار $\mathbf{R} = (R_{1,1}, R_{2,2}, \dots, R_{n,n})^\top$ ، یک $RRSS$ پایین به اندازه n از $GE(\theta_1, \lambda_1)$ و $\mathbf{s} = (s_{1,1}, s_{2,2}, \dots, s_{m,m})^\top$ مقدار مشاهده شده از بردار $\mathbf{S} = (S_{1,1}, S_{2,2}, \dots, S_{m,m})^\top$ یک $RRSS$ پایین به اندازه m از $GE(\theta_2, \lambda_2)$ ، تنها اطلاعاتی باشند که از جوامع X و Y در اختیار داریم، می‌خواهیم برآوردهای نقطه‌ای ML و $UMVU$ را برای $\mathcal{R} = Pr(X < Y)$ محاسبه کنیم.

۱.۲.۲ برآوردگر ML

برای محاسبه برآوردگر $\mathcal{R}$ از روش درستنمایی ماکسیمم ۳ حالت زیر را در نظر می‌گیریم:

حالت اول: $\lambda_1 \neq \lambda_2$ و $\theta_1 \neq \theta_2$

همان‌طور که در بخش ۶.۱ مشاهده کردیم، واحدهای $RRSS$ مستقل از یکدیگر می‌باشند، لذا با جایگذاری (۱.۲) و (۲.۲) در (۸.۱) تابع درستنمایی توأم به صورت زیر به دست می‌آید.

$$\begin{aligned} L(\theta_1, \theta_2, \lambda_1, \lambda_2) &= \prod_{i=1}^n \left[\frac{\{-\log(1 - e^{-\lambda_1 r_{i,i}})^{\theta_1}\}^{i-1} \theta_1 \lambda_1 e^{-\lambda_1 r_{i,i}} (1 - e^{-\lambda_1 r_{i,i}})^{\theta_1-1}}{(i-1)!} \right] \\ &\times \prod_{j=1}^m \left[\frac{\{-\log(1 - e^{-\lambda_2 s_{j,j}})^{\theta_2}\}^{j-1} \theta_2 \lambda_2 e^{-\lambda_2 s_{j,j}} (1 - e^{-\lambda_2 s_{j,j}})^{\theta_2-1}}{(j-1)!} \right]. \end{aligned} \quad (5.2)$$

و به همین ترتیب تساوی‌های مربوط به مشتق لگاریتم طبیعی درست‌نمایی نسبت به پارامترهای $(\theta_1, \theta_2, \lambda_1, \lambda_2)$ به ازای $\mathbf{r}$ و $\mathbf{s}$ داده‌شده به صورت زیر می‌باشد:

$$\frac{\partial \ell(\theta_1, \theta_2, \lambda_1, \lambda_2; \mathbf{r}, \mathbf{s})}{\partial \theta_1} = \sum_{i=1}^n \frac{(i-1)}{\theta_1} + \frac{n}{\theta_1} + \sum_{i=1}^n \log(1 - e^{-\lambda_1 r_{i,i}}) = 0, \quad (6.2)$$

$$\frac{\partial \ell(\theta_1, \theta_2, \lambda_1, \lambda_2; \mathbf{r}, \mathbf{s})}{\partial \theta_2} = \sum_{j=1}^m \frac{(j-1)}{\theta_2} + \frac{m}{\theta_2} + \sum_{j=1}^m \log(1 - e^{-\lambda_2 s_{j,j}}) = 0, \quad (7.2)$$

$$\begin{aligned} \frac{\partial \ell(\theta_1, \theta_2, \lambda_1, \lambda_2; \mathbf{r}, \mathbf{s})}{\partial \lambda_1} &= \sum_{i=1}^n (i-1) \frac{r_{i,i} e^{-\lambda_1 r_{i,i}}}{(1 - e^{-\lambda_1 r_{i,i}}) \log(1 - e^{-\lambda_1 r_{i,i}})} + \frac{n}{\lambda_1} \\ &\quad - \sum_{i=1}^n r_{i,i} + (\theta_1 - 1) \sum_{i=1}^n \frac{r_{i,i} e^{-\lambda_1 r_{i,i}}}{1 - e^{-\lambda_1 r_{i,i}}} = 0, \end{aligned} \quad (8.2)$$

$$\begin{aligned} \frac{\partial \ell(\theta_1, \theta_2, \lambda_1, \lambda_2; \mathbf{r}, \mathbf{s})}{\partial \lambda_2} &= \sum_{j=1}^m (j-1) \frac{s_{j,j} e^{-\lambda_2 s_{j,j}}}{(1 - e^{-\lambda_2 s_{j,j}}) \log(1 - e^{-\lambda_2 s_{j,j}})} + \frac{m}{\lambda_2} \\ &\quad - \sum_{j=1}^m s_{j,j} + (\theta_2 - 1) \sum_{j=1}^m \frac{s_{j,j} e^{-\lambda_2 s_{j,j}}}{1 - e^{-\lambda_2 s_{j,j}}} = 0, \end{aligned} \quad (9.2)$$

با حل معادلات (۶.۲) و (۷.۲) خواهیم داشت

$$\hat{\theta}_1(\lambda_1) = -\frac{N}{\sum_{i=1}^n \log(1 - e^{-\lambda_1 r_{i,i}})}, \quad (10.2)$$

$$\hat{\theta}_2(\lambda_2) = -\frac{M}{\sum_{j=1}^m \log(1 - e^{-\lambda_2 s_{j,j}})}, \quad (11.2)$$

که $N := n(n+1)/2$ و $M := m(m+1)/2$ می‌باشد. با حل معادلات غیرخطی حاصل شده از جایگذاری $\hat{\theta}_1(\lambda_1)$ و $\hat{\theta}_2(\lambda_2)$ در رابطه‌ی (۸.۲) و (۹.۲) با استفاده از روش‌های عددی، MLE پارامترهای مجهول را به دست می‌آوریم. در نهایت با استفاده از رابطه (۳.۲) و ویژگی پایایی برآوردگرهای درست‌نمایی ماکسیمم، برآوردگر $\hat{\mathcal{R}}_{ML}$ را محاسبه می‌کنیم.

حالت دوم: $\lambda_1 \neq \lambda_2$ و $\theta_1 = \theta_2 = \theta$

در حالتی که پارامترهای شکل برابر و پارامترهای مقیاس متفاوت باشد، با انجام برخی محاسبات جبری، برآورد درست‌نمایی ماکسیمم پارامترها با حل معادلات زیر به دست می‌آید.

$$\sum_{i=1}^n \frac{(i-1) r_{i,i} e^{-\lambda_1 r_{i,i}}}{(1 - e^{-\lambda_1 r_{i,i}}) \log(1 - e^{-\lambda_1 r_{i,i}})} + \frac{n}{\lambda_1} - \sum_{i=1}^n r_{i,i} + (\hat{\theta}(\lambda_1, \lambda_2) - 1) \frac{r_{i,i} e^{-\lambda_1 r_{i,i}}}{(1 - e^{-\lambda_1 r_{i,i}})} = 0, \quad (12.2)$$

$$\sum_{j=1}^m \frac{(j-1)s_{j,j}e^{-\lambda_2 s_{j,j}}}{(1-e^{-\lambda_2 s_{j,j}})\log(1-e^{-\lambda_2 s_{j,j}})} + \frac{m}{\lambda_2} - \sum_{i=1}^n s_{j,j} + (\hat{\theta}(\lambda_1, \lambda_2) - 1) \frac{s_{j,j}e^{-\lambda_2 s_{j,j}}}{(1-e^{-\lambda_2 s_{j,j}})} = 0, \quad (13.2)$$

که

$$\hat{\theta}(\lambda_1, \lambda_2) = -\frac{M+N}{\sum_{i=1}^n \log(1-e^{-\lambda_1 r_{i,i}}) + \sum_{j=1}^m \log(1-e^{-\lambda_2 s_{j,j}})}. \quad (14.2)$$

حالت سوم: $\lambda_1 = \lambda_2 = \lambda$ و $\theta_1 \neq \theta_2$

با در نظر گرفتن فرض‌های اولیه در مورد R و S، تابع درستنمایی ماکسیمم به صورت زیر به دست می‌آید.

$$L(\theta_1, \theta_2, \lambda) = \prod_{i=1}^n \left[\frac{\{-\log(1-e^{-\lambda r_{i,i}})\}^{\theta_1-1} \theta_1 \lambda e^{-\lambda r_{i,i}} (1-e^{-\lambda r_{i,i}})^{\theta_1-1}}{(i-1)!} \right] \times \prod_{j=1}^m \left[\frac{\{-\log(1-e^{-\lambda s_{j,j}})\}^{\theta_2-1} \theta_2 \lambda e^{-\lambda s_{j,j}} (1-e^{-\lambda s_{j,j}})^{\theta_2-1}}{(j-1)!} \right], \quad (15.2)$$

و با مشتق‌گیری از تابع درستنمایی نسبت به پارامترهای مجهول داریم:

$$\frac{\partial \ell(\theta_1, \theta_2, \lambda)}{\partial \theta_1} = \sum_{i=1}^n \frac{(i-1)}{\theta_1} + \frac{n}{\theta_1} + \sum_{i=1}^n \log(1-e^{-\lambda r_{i,i}}) = 0, \quad (16.2)$$

$$\frac{\partial \ell(\theta_1, \theta_2, \lambda)}{\partial \theta_2} = \sum_{j=1}^m \frac{(j-1)}{\theta_2} + \frac{m}{\theta_2} + \sum_{j=1}^m \log(1-e^{-\lambda s_{j,j}}) = 0, \quad (17.2)$$

$$\begin{aligned} \frac{\partial \ell(\theta_1, \theta_2, \lambda)}{\partial \lambda} &= \sum_{i=1}^n (i-1) \frac{r_{i,i} e^{-\lambda r_{i,i}}}{(1-e^{-\lambda r_{i,i}})\log(1-e^{-\lambda r_{i,i}})} + \frac{n}{\lambda} - \sum_{i=1}^n r_{i,i} + (\theta_1 - 1) \\ &\times \sum_{i=1}^n \frac{r_{i,i} e^{-\lambda r_{i,i}}}{1-e^{-\lambda r_{i,i}}} + \sum_{j=1}^m (j-1) \frac{s_{j,j} e^{-\lambda s_{j,j}}}{(1-e^{-\lambda s_{j,j}})\log(1-e^{-\lambda s_{j,j}})} + \frac{m}{\lambda} \\ &- \sum_{j=1}^m s_{j,j} + (\theta_2 - 1) \sum_{j=1}^m \frac{s_{j,j} e^{-\lambda s_{j,j}}}{1-e^{-\lambda s_{j,j}}} = 0. \end{aligned} \quad (18.2)$$

بنابراین MLE پارامترهای (۱۶.۲) و (۱۷.۲) به صورت تابعی از λ محاسبه می‌گردد.

$$\hat{\theta}_1(\lambda) = -\frac{N}{\sum_{i=1}^n \log(1-e^{-\lambda r_{i,i}})}, \quad (19.2)$$

$$\hat{\theta}_2(\lambda) = -\frac{M}{\sum_{j=1}^m \log(1-e^{-\lambda s_{j,j}})}. \quad (20.2)$$

با جایگذاری $\hat{\theta}_1(\lambda)$ و $\hat{\theta}_2(\lambda)$ در لگاریتم رابطه (۱۸.۲) و ماکسیمم نمودن آن بر حسب λ ، برآورد درستنمایی ماکسیمم λ را با استفاده از تساوی زیر محاسبه می‌کنیم

$$h(\hat{\lambda}) = \hat{\lambda}. \quad (21.2)$$

که

$$h(\lambda) = (m+n) \left[\sum_{i=1}^n r_{i,i} + \sum_{j=1}^m s_{j,j} + \left(\frac{N}{\sum_{i=1}^n \log(1 - e^{-\lambda r_{i,i}})} + 1 \right) \sum_{i=1}^n \frac{r_{i,i} e^{-\lambda r_{i,i}}}{(1 - e^{-\lambda r_{i,i}})} \right. \\ \left. - \sum_{i=1}^n \frac{(i-1) r_{i,i} e^{-\lambda r_{i,i}}}{(1 - e^{-\lambda r_{i,i}}) \log(1 - e^{-\lambda r_{i,i}})} + \left(\frac{M}{\sum_{j=1}^m \log(1 - e^{-\lambda s_{j,j}})} + 1 \right) \sum_{j=1}^m \frac{s_{j,j} e^{-\lambda s_{j,j}}}{(1 - e^{-\lambda s_{j,j}})} \right. \\ \left. - \sum_{j=1}^m \frac{(j-1) s_{j,j} e^{-\lambda s_{j,j}}}{(1 - e^{-\lambda s_{j,j}}) \log(1 - e^{-\lambda s_{j,j}})} \right]^{-1}.$$

با توجه به این که تابع $h(\lambda)$ در بازه $(0, \infty)$ پیوسته و مشتق پذیر است و به ازای هر $\lambda \in (0, \infty)$ ، $h(\lambda) \in (0, \infty)$ و همچنین عددی مانند $0 < k < 1$ وجود دارد به طوری که $|h'(\lambda)| \leq k < 1$ ، بنابراین $\hat{\lambda}$ یک نقطه‌ی ثابت از حل معادله غیرخطی (۲۱.۲) است که با استفاده از یک فرآیند تکراری ساده به صورت زیر محاسبه می گردد.

$$h(\lambda_{(j)}) = \lambda_{(j+1)},$$

که $\lambda_{(j)}$ ، j امین تکرار برای $\hat{\lambda}$ می باشد. لازم به توضیح است که فرآیند تکرار زمانی متوقف می گردد که $|\lambda_{(j)} - \lambda_{(j+1)}| < \epsilon$ به طوری که ϵ مقداری به اندازه کافی کوچک است. با توجه به این که در این حالت پارامتر $\mathcal{R}$ ارائه شده در (۳.۲) به صورت رابطه زیر به دست می آید

$$\mathcal{R} = \frac{\theta_2}{\theta_1 + \theta_2} \quad (22.2)$$

بنابراین برطبق رابطه (۲۲.۲) ملاحظه می کنیم در این حالت، پارامتر $\mathcal{R}$ مستقل از پارامتر λ می باشد، لذا بدون از دست دادن کلیت مسئله فرض می کنیم $\lambda = 1$ باشد. بنابراین برآورد درستنمایی ماکسیمم پارامترهای θ_1 و θ_2 به صورت زیر به دست می آید.

$$\hat{\theta}_1 = -\frac{N}{\sum_{i=1}^n \log(1 - e^{-r_{i,i}})}, \quad (23.2)$$

و

$$\hat{\theta}_2 = -\frac{M}{\sum_{j=1}^m \log(1 - e^{-s_{j,j}})}. \quad (24.2)$$

و $\hat{\mathcal{R}}_{ML}$ نیز به صورت زیر نتیجه می دهد.

$$\hat{\mathcal{R}}_{ML} = \frac{\hat{\theta}_2}{\hat{\theta}_1 + \hat{\theta}_2} = \left[1 + \frac{N}{M} \frac{T_2}{T_1} \right]^{-1}, \quad (25.2)$$

که در آن $T_1 = -\sum_{i=1}^n \log(1 - e^{-R_{i,i}})$ و $T_2 = -\sum_{j=1}^m \log(1 - e^{-S_{j,j}})$ می باشد. همان گونه که قبلاً نیز بیان شده است $R_{i,i}$ ها و $S_{j,j}$ ها، متغیرهای تصادفی مستقل می باشند، با در نظر گرفتن $Z_i = -\log(1 - e^{-R_{i,i}})$ می توان نشان داد که Z_i دارای توزیع گاما با پارامتر شکل i و پارامتر مقیاس θ_1 می باشد و یا به طور معادل $Z_i \sim \text{Gamma}(i, \theta_1)$. بنابراین

$$T_1 = -\sum_{i=1}^n \log(1 - e^{-R_{i,i}}) \sim \text{Gamma}(N, \theta_1), \quad (26.2)$$

و به‌طور مشابه

$$T_{\gamma} = - \sum_{j=1}^m \log(1 - e^{-S_{j,j}}) \sim \text{Gamma}(M, \theta_{\gamma}). \quad (27.2)$$

حال گشتاورهای $\hat{R}_{ML}$ را طبق قضیه زیر محاسبه می‌کنیم

قضیه ۲.۲.۲. گشتاور مرتبه k ام $\hat{R}_{ML}$ به‌صورت زیر می‌باشد

$$E(\hat{R}_{ML}^k) = \frac{\Gamma(M+N)\Gamma(N+k)}{\Gamma(N)\Gamma(N+M+k)} \left(\frac{N}{M} \frac{1-\mathcal{R}}{\mathcal{R}} \right)^N \times {}_2F_1 \left(M+N, N+k, M+N+k, 1 - \frac{N}{M} \frac{1-\mathcal{R}}{\mathcal{R}} \right), \quad (28.2)$$

که در آن ${}_2F_1$ نشان‌دهنده‌ی تابع فوق هندسی می‌باشد.

برهان. تابع چگالی $\hat{R}_{ML}$ ، با در نظر گرفتن روابط (۲۶.۲) و (۲۷.۲) و استفاده از روش تبدیل به‌صورت زیر محاسبه می‌گردد.

$$f_{\hat{R}_{ML}}(x) = \frac{\Gamma(M+N)}{\Gamma(M)\Gamma(N)} \left(\frac{M\theta_{\gamma}}{N\theta_1} \right)^M (1-x)^{M-1} x^{-(M+1)} \left(1 + \frac{M\theta_{\gamma}}{N\theta_1} \left(\frac{1-x}{x} \right) \right)^{-(M+N)}, \quad (29.2)$$

که $\Gamma(\cdot)$ نشان‌دهنده‌ی تابع گامای کامل می‌باشد. با در نظر گرفتن تابع فوق هندسی که به فرم زیر تعریف می‌شود (برای جزئیات بیشتر به بیلی^۹ (۱۹۳۵) مراجعه نمایید).

$${}_2F_1 = F(a, b, c; z) = \frac{\Gamma(c)}{\Gamma(b)\Gamma(c-b)} \int_0^1 t^{b-1} (1-t)^{c-b-1} (1-tz)^{-a} dt$$

جهت استفاده از ویژگی‌های تابع فوق هندسی در محاسبه گشتاورها، تابع چگالی به‌دست آمده در رابطه (۲۹.۲) را پس از انجام محاسباتی ساده به فرم زیر نمایش می‌دهیم

$$f_{\hat{R}_{ML}}(x) = \frac{\Gamma(M+N)}{\Gamma(M)\Gamma(N)} \left(\frac{N}{M} \frac{1-\mathcal{R}}{\mathcal{R}} \right)^N x^{N-1} (1-x)^{M-1} \times \left(1 - x \left(1 - \frac{N}{M} \frac{1-\mathcal{R}}{\mathcal{R}} \right) \right)^{-(M+N)}, \quad 0 < x < 1. \quad (30.2)$$

با استفاده از رابطه (۳۰.۲) و تعریف تابع فوق هندسی نتیجه مورد نظر به‌دست می‌آید. $\square$

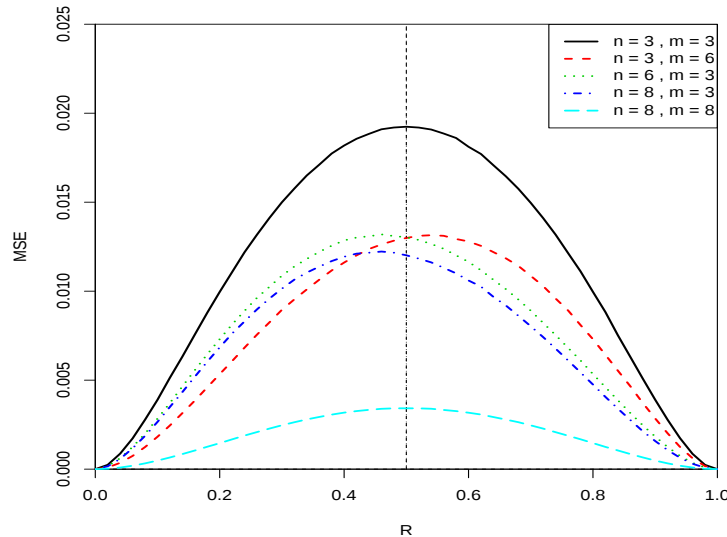
ملاحظه ۱.۲.۲. با استفاده از رابطه (۲۸.۲) امید ریاضی $\hat{R}_{ML}$ به صورت زیر می‌باشد

$$E(\hat{R}_{ML}) = \left(\frac{N}{M} \frac{1-\mathcal{R}}{\mathcal{R}} \right)^N \frac{N}{M+N} {}_2F_1 \left(M+N, N+1, M+N+1; 1 - \frac{N}{M} \frac{1-\mathcal{R}}{\mathcal{R}} \right),$$

که نشان می‌دهد $\hat{R}_{ML}$ یک برآوردگر اریب برای $\mathcal{R}$ می‌باشد و می‌توان جهت محاسبه $MSE(\hat{R}_{ML}, \mathcal{R}) = Var(\hat{R}_{ML}) + (E(\hat{R}_{ML}) - \mathcal{R})^2$ آن را به کار برد.

^۹Bailey

در شکل ۱.۲ نمودار $MSE(\hat{\mathcal{R}}_{ML}, \mathcal{R})$ در مقابل $\mathcal{R}$ برای مقادیر مختلف n و m رسم شده است. این نمودار نشان می‌دهد که با افزایش اندازه نمونه، MSE کاهش می‌یابد. همچنین زمانی که $n = m$ می‌باشد، تابع MSE حول $\mathcal{R} = 0.5$ متقارن است درحالی که به ازای $n \neq m$ از نامتقارن است. درضمن MSE بیشترین مقدار خود را در $\mathcal{R}_{max} \propto \frac{n}{n+m}$ اختیار می‌کند که $\mathcal{R}_{max} = \operatorname{argmax}_{\mathcal{R}} MSE(\hat{\mathcal{R}}_{ML}, \mathcal{R})$ می‌باشد.



شکل ۱.۲: نمودار $MSE(\hat{\mathcal{R}}_{ML}, \mathcal{R})$ در مقابل $\mathcal{R}$.

۲.۲.۲ برآوردگر $UMVU$

در این بخش با در نظر گرفتن توزیع نمایی تعمیم یافته به عنوان توزیع جوامع X و Y و براساس طرح $RRSS$ برآوردگر $UMVU$ پارامتر تنش-مقاومت را محاسبه می‌کنیم.

قضیه ۳.۲.۲. فرض کنید $\mathbf{r} = (r_{1,1}, r_{1,2}, \dots, r_{n,n})^T$ مشاهدات مربوط به بردار تصادفی $\mathbf{R}$ یک $RRSS$ پایین به اندازه n از توزیع $GE(\theta_1, 1)$ است. همچنین $\mathbf{s} = (s_{1,1}, s_{1,2}, \dots, s_{n,n})^T$ مشاهدات بردار تصادفی $\mathbf{S} = (S_{1,1}, S_{1,2}, \dots, S_{n,n})^T$ که یک $RRSS$ پایین به اندازه m از $GE(\theta_2, 1)$ می‌باشد. بنابراین $UMVUE$ پارامتر $\mathcal{R}$ را می‌توان به صورت عبارتی از برآوردگر درست‌نمایی ماکسیمم به صورت زیر به دست آورد.

$$\hat{\mathcal{R}}_{UMVU} = \begin{cases} \sum_{s=0}^{M-1} \frac{\binom{M-1}{s}}{\binom{N+s-1}{s}} \left(-\frac{\hat{\mathcal{R}}_{ML}}{1 - \hat{\mathcal{R}}_{ML}} \frac{N}{M} \right)^s & \hat{\mathcal{R}}_{ML} < \frac{M}{N+M}, \\ 1 - \sum_{s=0}^{N-1} \frac{\binom{N-1}{s}}{\binom{M+s-1}{s}} \left(-\frac{1 - \hat{\mathcal{R}}_{ML}}{\hat{\mathcal{R}}_{ML}} \frac{M}{N} \right)^s & \hat{\mathcal{R}}_{ML} \geq \frac{M}{N+M}. \end{cases} \quad (31.2)$$

برهان. به راحتی می‌توان نشان داد که $(T_1, T_2)^\top$ ارائه شده در روابط (۲۶.۲) و (۲۷.۲) یک آماره بسنده کامل برای (θ_1, θ_2) می‌باشد. اگر فرض کنیم $U = -\log(1 - e^{-R_{1,1}})$ و $V = -\log(1 - e^{-S_{1,1}})$ باشد در این صورت آماره

$$I(U < V) = \begin{cases} 1 & U < V, \\ 0 & o.w, \end{cases}$$

یک برآورگر نااریب برای $\mathcal{R}$ می‌باشد. بنابراین با استفاده از روش راثو- بلکول^{۱۰} و لهن و شفه^{۱۱} (لهن و کسلا^{۱۲}، ۱۹۹۸)، برآورگر $UMVU$ برای $\mathcal{R}$ با جایگذاری $(t_1, t_2)^\top$ به جای $(T_1, T_2)^\top$ در احتمال شرطی زیر به دست می‌آید.

$$P(U < V | T_1 = t_1, T_2 = t_2) = \int_C \int f_{U|T_1=t_1}(u) f_{V|T_2=t_2}(v) du dv, \quad (32.2)$$

که انتگرال اخیر روی ناحیه $\mathcal{C} = \{(u, v) : 0 < u < t_1, 0 < v < t_2, u < v\}$ می‌باشد. قبل از محاسبه این انتگرال، لازم است توزیع شرطی $U|(T_1 = t_1)$ و $V|(T_2 = t_2)$ را بدانیم. با اندکی محاسبات جبری خواهیم داشت

$$f_{U|(T_1=t_1)}(u) = \frac{(N-1)(t_1-u)^{N-2}}{t_1^{N-1}} \quad 0 < u < t_1, \quad (33.2)$$

و

$$f_{V|(T_2=t_2)}(v) = \frac{(M-1)(t_2-v)^{M-2}}{t_2^{M-1}} \quad 0 < v < t_2. \quad (34.2)$$

اگر $t_1 < t_2$ ، سمت راست رابطه (۳۲.۲) (RHS) برابر می‌شود با

$$\begin{aligned} RHS &= \frac{(N-1)(M-1)}{t_1^{N-1} t_2^{M-1}} \int_0^{t_1} \int_u^{t_2} (t_1-u)^{N-2} (t_2-v)^{M-2} dv du \\ &= \frac{(N-1)}{t_1^{N-1} t_2^{M-1}} \int_0^{t_1} (t_1-u)^{N-2} (t_2-u)^{M-1} du. \end{aligned}$$

با در نظر گرفتن تبدیل $z = \frac{u}{t_1}$ و استفاده از بسط دوجمله‌ای برای $(t_2 - z t_1)^{M-1}$ نتیجه می‌شود

$$\begin{aligned} RHS &= (N-1) \sum_{s=0}^{M-1} \left(\frac{t_1}{t_2}\right)^s (-1)^s \binom{M-1}{s} \int_0^1 z^s (1-z)^{N-2} dz \\ &= \sum_{s=0}^{M-1} (-1)^s \frac{\binom{M-1}{s}}{\binom{N+s-1}{s}} \left(\frac{t_1}{t_2}\right)^s. \end{aligned}$$

^{۱۰} Rao-Blackwell

^{۱۱} Lehmann-Scheffe

^{۱۲} Cassela

به طور مشابه، اگر $t_1 > t_2$ ، آن گاه

$$RHS = 1 - \sum_{s=0}^{N-1} (-1)^s \frac{\binom{N-1}{s}}{\binom{M+s-1}{s}} \left(\frac{t_2}{t_1}\right)^s.$$

از این رو داریم

$$\hat{R}_{UMVU} = \begin{cases} \sum_{s=0}^{M-1} (-1)^s \frac{\binom{M-1}{s}}{\binom{N+s-1}{s}} \left(\frac{T_1}{T_2}\right)^s & T_1 < T_2, \\ 1 - \sum_{s=0}^{N-1} (-1)^s \frac{\binom{N-1}{s}}{\binom{M+s-1}{s}} \left(\frac{T_2}{T_1}\right)^s & T_1 > T_2, \end{cases} \quad (35.2)$$

که T_1 و T_2 به ترتیب در رابطه های (۲۶.۲) و (۲۷.۲) بیان شده است. در نهایت با استفاده از (۲۵.۲) نتیجه مورد نظر محاسبه می گردد. □

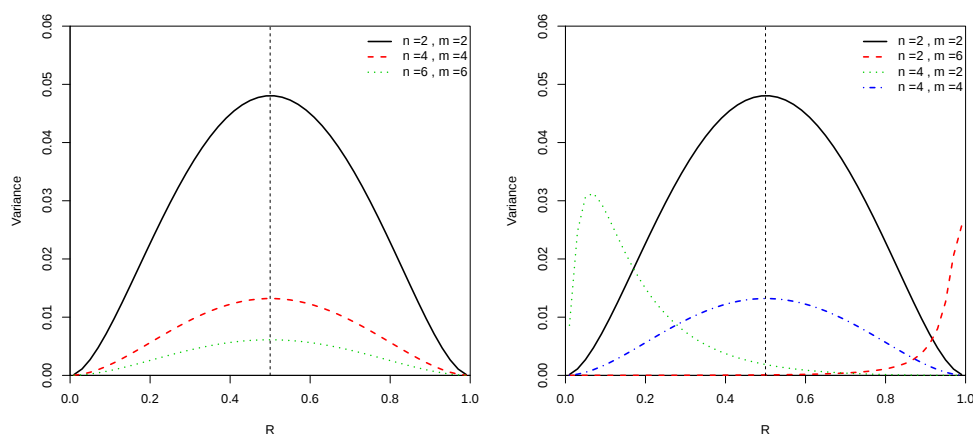
با استفاده از رابطه (۳۱.۲)، نمودار مقادیر عددی $Var(\hat{R}_{UMVU})$ در مقابل R را برای ترکیب های مختلفی از n و m در شکل ۲.۲ رسم نمودیم. مشاهده می کنیم رفتار $Var(\hat{R}_{UMVU})$ زمانی که $n = m$ ، تقریباً مشابه رفتار $MSE(\hat{R}_{ML}, R)$ می باشد اما زمانی که $n \neq m$ ، نمودار فوق رفتار متفاوتی دارد و نامتقارن می باشد. بنابراین برای مقایسه بهتر عملکرد این دو برآوردگر نقطه ای، از معیار کارایی نسبی استفاده می کنیم. با در نظر گرفتن

$$e(\hat{R}_{UMVU}, \hat{R}_{ML}) = \frac{Var(\hat{R}_{UMVU})}{MSE(\hat{R}_{ML}, R)},$$

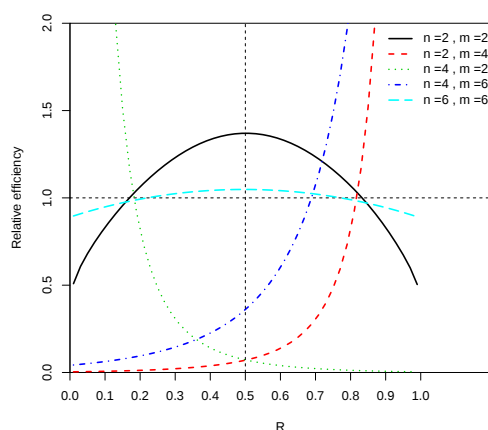
نمودار کارایی نسبی برآورگرهای نقطه ای به دست آمده را برای نمونه هایی با اندازه های مختلف در شکل (۳.۲) رسم نمودیم. با توجه به شکل، ملاحظه می کنیم نمودار فوق در حالتی که $n = m$ حول $R = 0.5$ متقارن می باشد و به ازای $n < m$ و یا $n > m$ از حالت تقارن فاصله می گیرند. در حالتی که $n = m$ ، برای مقادیر نزدیک به مقادیر میانی R ، کارایی نسبی برآوردگر $\hat{R}_{ML}$ بهتر از $\hat{R}_{UMVU}$ می باشد. با این حال، ملاحظه می شود برای مقادیر ابتدایی و انتهایی R ، عملکرد $\hat{R}_{UMVU}$ بهتر از $\hat{R}_{ML}$ می باشد. این ویژگی در داده های واقعی که در بخش ۵.۲ مورد بررسی قرار می گیرند نیز مشاهده می گردد.

۳.۲ برآورد فاصله ای پارامتر تنش-مقاومت

برآوردهای نقطه ای و فاصله ای به طور نزدیکی به هم مرتبط می باشند با این حال فواصل اطمینان شامل اطلاعات بیشتری می باشد و استنباط های بهتری فراهم می کند. در ادامه یک فاصله ای اطمینان دقیق برای حالت سوم بخش قبل یعنی حالتی که $\lambda_1 = \lambda_2 = 1$ و $\theta_1 \neq \theta_2$ باشد را با استفاده از کمیت محور و همچنین سه فاصله اطمینان تقریبی براساس بوت استرپ پارامتری برای R به دست می آوریم.



شکل ۲.۲: نمودار $Var(\hat{R}_{UMVU})$ در مقابل R .



شکل ۳.۲: نمودار $e(\hat{R}_{UMVU}, \hat{R}_{ML})$ در مقابل R .

۱.۳.۲ فاصله اطمینان براساس کمیت محوری

با استفاده از رابطه‌ی (۲۵.۲)–(۲۷.۲) داریم

$$\hat{R}_{ML} \stackrel{d}{=} \left(1 + \frac{1-R}{R} W \right)^{-1}, \quad (36.2)$$

که در آن $\stackrel{d}{=}$ نشان‌دهنده‌ی هم‌توزیع بودن و $W \sim F_{2M, 2N}$ یک متغیر تصادفی با توزیع فیشر به ترتیب با $2M$ و $2N$ درجه‌ی آزادی است. بنابراین با استفاده از (۳۶.۲) یک فاصله اطمینان دقیق

$100(1 - \alpha)\%$ برای $\mathcal{R}$ به صورت زیر به دست می آید.

$$\left\{ \left(1 + \frac{1 - \hat{R}_{ML}}{\hat{R}_{ML}} F_{1-\frac{\alpha}{2}}(2N, 2M) \right)^{-1}, \left(1 + \frac{1 - \hat{R}_{ML}}{\hat{R}_{ML}} F_{\frac{\alpha}{2}}(2N, 2M) \right)^{-1} \right\}, \quad (37.2)$$

که در آن $F_{\alpha}(2N, 2M)$ چندک α ام توزیع F با $2N$ و $2M$ درجه آزادی است. مقادیر عددی میانگین طول فاصله اطمینان (۳۷.۲) را براساس (۳۶.۲) در جدول های ۱.۲ و ۲.۲ برحسب سطح معنی داری $\alpha = 0.05, 0.1$ محاسبه کرده ایم. برطبق نتایج به دست آمده، مشاهده می شود با افزایش اندازه نمونه ها، طول فاصله اطمینان کاهش می یابد. همچنین زمانی که $\mathcal{R} = 0.5$ ، میانگین طول فاصله اطمینان بیشترین مقدار را دارد که شکل های ۱.۲ و ۲.۲ مربوط به میانگین توان دوم خطای برآوردگر $\hat{R}_{ML}$ و واریانس $\hat{R}_{UMVU}$ نیز براین مطلب صحنه می گذارند. درحقیقت بیشتر بودن میانگین طول فاصله اطمینان و همچنین میانگین مربع خطا در این نقطه را می توان به این نکته دلالت داد که تشخیص قدرت متغیرهای تنش و مقاومت در حوالی $\mathcal{R} = 0.5$ سخت تر می باشد.

۲.۳.۲ فواصل اطمینان بوت استرپ پارامتری

در بخش ۷.۱ درخصوص انواع بوت استرپ و فواصل اطمینان تقریبی بر پایه ی بوت استرپ بحث شد. همان طور که بیان شد، در صورتی که توزیع جامعه ی آماری مشخص باشد و فقط پارامترهای آن مجهول باشند، می توان از روش بوت استرپ پارامتری برای تولید نمونه های بوت استرپی استفاده نمود. در این بخش، با فرض این که پارامتر مقیاس λ برابر با یک باشد چند فاصله بوت استرپی را به دست می آوریم. فواصل اطمینان بوت استرپ پایه ای، بوت استرپ صدکی و همچنین t -بوت استرپ را با استفاده از روش های ارائه شده در افرون و تیشیرانی (۱۹۹۳) و صالحی و احمدی (۲۰۱۵) به دست می آوریم. با توجه به این که تمامی نتایج به دست آمده در بخش های قبلی از این فصل، براساس $T_1 = -\sum_{i=1}^n \log(1 - e^{-R_{i,i}})$ و $T_2 = -\sum_{j=1}^m \log(1 - e^{-S_{j,j}})$ می باشند، لذا می توان از بوت استرپ پارامتری به جای بوت استرپ ناپارامتری برای تولید نمونه های بوت استرپی استفاده نمود. سه گام معمول از روش بوت استرپ برای یافتن فواصل اطمینان به شرح الگوریتم زیر می باشد.

فاصله اطمینان بوت استرپ پایه

با انجام تمام مراحل الگوریتم ۵، مشاهدات بوت استرپی $d^* = (d_{(1)}^*, \dots, d_{(B)}^*)^T$ را به دست آورید که در آن $b = 1, \dots, B$ ، $d_{(b)}^* = \mathcal{R}_{(b)ML}^* - \hat{R}_{ML}$ است. درنهایت بازه اطمینان بوت استرپ پایه $100(1 - \alpha)\%$ برای $\mathcal{R}$ به صورت زیر به دست می آید.

$$(\hat{R}_{ML} - d_{1-\alpha/2}^*, \hat{R}_{ML} - d_{\alpha/2}^*) \quad (38.2)$$

که در آن d_{γ}^* چندک γ ام d^* می باشد.

الگوریتم ۵ الگوریتم عمومی بوت‌استرپ پارامتری

۱. براساس نمونه‌های مستقل مشاهده شده $t_{\lambda i}$ و $t_{\lambda j}$ از توزیع $t_{\lambda i} \sim \text{Gamma}(i, \theta_1)$ ، $t_{\lambda j} \sim \text{Gamma}(j, \theta_2)$ ، $j = 1, \dots, m$ و $i = 1, \dots, n$ برآوردهای $\hat{\theta}_1$ ، $\hat{\theta}_2$ و $\hat{R}_{ML}$ را به ترتیب از روابط (۲۳.۲) ، (۲۴.۲) و (۲۵.۲) به دست آورید.
۲. نمونه‌های بوت‌استرپی $t_{\lambda i}^* \sim \text{Gamma}(i, \hat{\theta}_1)$ که $i = 1, \dots, n$ و $t_{\lambda j}^* \sim \text{Gamma}(j, \hat{\theta}_2)$ که $j = 1, \dots, m$ را تولید کنید و با استفاده از آن‌ها برآوردهای بوت‌استرپی $\hat{\theta}_1^*$ ، $\hat{\theta}_2^*$ و $\hat{R}_{ML}^*$ را به دست آورید.
۳. گام دو را برای $b = 1, \dots, B$ تکرار کنید تا B مشاهده بوت‌استرپی $\hat{R}_{(b)ML}^*$ به دست آورید.

فاصله اطمینان بوت‌استرپ صدکی

ابتدا مراحل ۱-۳ الگوریتم ۵ را انجام دهید. حال اگر $\hat{F}$ تابع توزیع تجربی مشاهدات $\hat{R}_{(b)ML}^*$ ، $b = 1, \dots, B$ باشد، آن گاه بازه‌ی

$$\left(\hat{F}^{-1}\left(\frac{\alpha}{2}\right), \hat{F}^{-1}\left(1 - \frac{\alpha}{2}\right) \right) \quad (39.2)$$

یک فاصله اطمینان بوت‌استرپ صدکی برای $\mathcal{R}$ با ضریب تقریبی $100(1 - \alpha)\%$ است.

فاصله اطمینان t -بوت‌استرپ

برای به دست آوردن فاصله اطمینان t -بوت‌استرپ با انحراف معیار مجانبی، علاوه بر انجام مراحل بیان شده در الگوریتم ۵، به ازای هر تکرار گام دوم، بوت‌استرپ پارامتری دیگری برای به دست آوردن انحراف معیار مجانبی انجام دهید. بدین صورت که، مرحله دوم الگوریتم فوق را برای $t_{\lambda i}^{**} \sim \text{Gamma}(i, \hat{\theta}_1^*)$ و $t_{\lambda j}^{**} \sim \text{Gamma}(j, \hat{\theta}_2^*)$ که $i = 1, \dots, n$ ، $j = 1, \dots, m$ را تولید کنید. سپس با استفاده از مشاهدات $\hat{R}_{(b')ML}^{**}$ ، برآورد بوت‌استرپ انحراف استاندارد $\hat{R}_{(b)ML}^*$ را به صورت زیر به دست آورید

$$\hat{\sigma}(\hat{R}_{(b)ML}^*) = \sqrt{\frac{1}{B' - 1} \sum_{b'=1}^{B'} \left(\hat{R}_{(b')ML}^{**} - \bar{R}^{**} \right)^2}, \quad (40.2)$$

که در این رابطه $\bar{R}^{**}$ میانگین نمونه $\hat{R}_{(b')ML}^{**}$ ها یعنی $\bar{R}^{**} = \frac{1}{B'} \sum_{b'=1}^{B'} \hat{R}_{(b')ML}^{**}$ می‌باشد. حال بردار $t^* = (t_{(1)}^*, \dots, t_{(B)}^*)$ را در نظر بگیرید، به طوری که

$$t_{(b)}^* = \frac{\hat{R}_{(b)ML}^* - \hat{R}_{ML}}{\hat{\sigma}(\hat{R}_{(b)ML}^*)}, \quad b = 1, \dots, B.$$

در این صورت فاصله اطمینان t -بوت استرپ $\% (1 - \alpha) 100$ برای R به صورت زیر محاسبه می شود

$$(\hat{R}_{ML} - t_{1-\alpha/2}^* \hat{\sigma}(\hat{R}_{ML}), \hat{R}_{ML} - t_{\alpha/2}^* \hat{\sigma}(\hat{R}_{ML})),$$

که t_{γ}^* چندک γ ام t^* و $\hat{\sigma}(\hat{R}_{ML})$ نیز برآورد بوت استرپی انحراف معیار برآوردگر $\hat{R}_{ML}$ می باشد.

۴.۲ مقایسه‌ی برآوردگرهای فاصله‌ای براساس شبیه‌سازی

در این بخش، براساس نتایج شبیه‌سازی مونت کارلو و بوت استرپ، فواصل اطمینان مطرح شده در بخش ۳.۲، مقایسه شده‌اند. در فرآیند شبیه‌سازی تمام ترکیب‌های متفاوت $n = 2, 4, 6$ و $m = 2, 4$ ، در نظر گرفته شده‌است. همچنین $R = 0.1, 0.5, 0.95$ و $\theta_2 = 0.4$ فرض شده‌است و به تبع مفروضات فوق مقادیر مربوط به θ_1 از رابطه‌ی (۲۲.۲) به دست آمده‌است. مقدار $B = 1000$ در نظر گرفته شده‌است و به ازای هر کدام از ترکیب‌ها، ۱۰۰۰۰ نمونه‌ی r و s به ترتیب از توزیع $GE(\theta_1, 1)$ و $GE(\theta_2, 1)$ تولید شده‌است. همان طور که می‌دانیم برای محاسبه‌ی انحراف معیار بوت استرپی نیاز به تکرار زیاد نیست، لذا انتخاب $B' = 25$ برآوردگر منطقی از انحراف معیار بوت استرپی را نتیجه می‌دهد (افرون و تیشیرانی، ۱۹۹۳). نتایج مربوط به شبیه‌سازی فواصل اطمینان را که شامل میانگین طول فاصله اطمینان (EL) و احتمال پوشش (CP) است در جدول‌های ۱.۲ و ۲.۲ به ازای $\alpha = 0.05, 0.1$ به طور جداگانه‌ای محاسبه و نمایش داده شده‌است. اگر (L, U) یک فاصله اطمینان برای R باشد و (L_i, U_i) ، $i = 1, \dots, 10000$ ، مشاهدات آن باشد، آن‌گاه میانگین طول و احتمال پوشش آن به ترتیب برابر است با:

$$EL = \frac{1}{10000} \sum_{i=1}^{10000} (U_i - L_i), \quad CP = \frac{1}{10000} \sum_{i=1}^{10000} I(U_i \leq R \leq L_i) \quad (41.2)$$

که $I(A)$ تابع نشان‌گر مجموعه‌ی A می‌باشد. همچنین در این جدول‌ها نمادهای زیر به کار رفته‌است.

- ML : فاصله اطمینان دقیق بر پایه‌ی کمیت محور براساس رابطه (۳۷.۲).
- $Basic$: فاصله اطمینان بوت استرپ پایه برطبق رابطه (۳۸.۲).
- $Perc$: فاصله اطمینان بوت استرپ صدکی به دست آمده از رابطه (۳۹.۲).
- $Boot - t$: فاصله اطمینان t -بوت استرپ با انحراف معیار بوت استرپی براساس رابطه (۴۰.۲).

نتایج زیر از جدول‌های ۱.۲ و ۲.۲ کسب می‌شوند.

- میانگین طول تمامی فواصل اطمینان با افزایش حجم نمونه کاهش می‌یابد.
- بیشترین مقدار میانگین طول در نقطه وسط یعنی $R = 0.5$ رخ می‌دهد که این رفتار را قبلاً در نمودارهای مربوط به میانگین توان دوم برآورگرهای نقطه‌ای نیز شاهد بودیم.

جدول ۱.۲: مقادیر EL و CP فواصل اطمینان ۹۵٪ برای R .

EL.Basic	CP.Basic	EL.ML	CP.ML	EL.Boot-t	CP.Boot-t	EL.Perc	CP.Perc	m	n	$\mathcal{R}$	Row
۰.۳۸۰	۰.۷۶۴	۰.۳۸۲	۰.۹۴۸	۰.۴۲۷	۰.۹۵۴	۰.۳۸۲	۰.۹۴۸	۲	۲	۰.۱	۱
۰.۲۳۳	۰.۷۸۲	۰.۳۱۷	۰.۹۵۲	۰.۵۳۳	۰.۹۸۶	۰.۲۳۴	۰.۹۱۸	۴	۲	۰.۱	۲
۰.۲۰۳	۰.۷۷۷	۰.۳۰۰	۰.۹۵۱	۰.۶۲۴	۰.۹۹۴	۰.۲۰۴	۰.۸۹۴	۶	۲	۰.۱	۳
۰.۳۵۳	۰.۸۲۵	۰.۲۷۰	۰.۹۵۰	۰.۳۲۳	۰.۹۶۶	۰.۳۵۵	۰.۹۱۷	۲	۴	۰.۱	۴
۰.۱۷۷	۰.۸۵۹	۰.۱۷۸	۰.۹۵۲	۰.۱۹۰	۰.۹۶۸	۰.۱۷۸	۰.۹۵۱	۴	۴	۰.۱	۵
۰.۱۳۹	۰.۸۶۷	۰.۱۵۱	۰.۹۵۲	۰.۱۹۰	۰.۹۸۵	۰.۱۳۹	۰.۹۴۵	۶	۴	۰.۱	۶
۰.۳۴۱	۰.۸۳۹	۰.۲۴۵	۰.۹۵۲	۰.۳۲۶	۰.۹۸۲	۰.۳۴۳	۰.۸۹۵	۲	۶	۰.۱	۷
۰.۱۵۶	۰.۸۷۹	۰.۱۴۶	۰.۹۵۰	۰.۱۷۳	۰.۹۸۰	۰.۱۵۷	۰.۹۴۴	۴	۶	۰.۱	۸
۰.۱۱۵	۰.۸۹۹	۰.۱۱۵	۰.۹۵۱	۰.۱۲۱	۰.۹۶۳	۰.۱۱۵	۰.۹۵۰	۶	۶	۰.۱	۹
۰.۶۴۲	۰.۷۱۲	۰.۶۴۴	۰.۹۴۸	۰.۸۸۶	۰.۹۵۳	۰.۶۴۴	۰.۹۴۸	۲	۲	۰.۵	۱۰
۰.۵۴۸	۰.۷۸۲	۰.۵۵۸	۰.۹۵۲	۰.۸۸۲	۰.۹۸۴	۰.۵۵۰	۰.۹۱۸	۴	۲	۰.۵	۱۱
۰.۵۲۰	۰.۷۹۴	۰.۵۳۴	۰.۹۵۱	۰.۹۱۱	۰.۹۹۴	۰.۵۲۲	۰.۸۹۴	۶	۲	۰.۵	۱۲
۰.۵۴۷	۰.۷۷۸	۰.۵۵۸	۰.۹۵۰	۰.۸۸۶	۰.۹۸۴	۰.۵۴۹	۰.۹۱۷	۲	۴	۰.۵	۱۳
۰.۴۰۴	۰.۸۶۶	۰.۴۰۶	۰.۹۵۲	۰.۵۰۵	۰.۹۷۲	۰.۴۰۶	۰.۹۵۱	۴	۴	۰.۵	۱۴
۰.۳۵۴	۰.۸۸۹	۰.۳۵۶	۰.۹۵۲	۰.۴۹۴	۰.۹۸۷	۰.۳۵۵	۰.۹۴۵	۶	۴	۰.۵	۱۵
۰.۵۲۰	۰.۷۹۳	۰.۵۳۳	۰.۹۵۲	۰.۹۱۴	۰.۹۹۴	۰.۵۲۲	۰.۸۹۵	۲	۶	۰.۵	۱۶
۰.۳۵۴	۰.۸۸۴	۰.۳۵۶	۰.۹۵۰	۰.۴۹۳	۰.۹۸۶	۰.۳۵۵	۰.۹۴۴	۴	۶	۰.۵	۱۷
۰.۲۹۰	۰.۹۰۸	۰.۲۹۱	۰.۹۵۱	۰.۳۳۳	۰.۹۶۹	۰.۲۹۱	۰.۹۴۹	۶	۶	۰.۵	۱۸
۰.۲۴۹	۰.۷۶۰	۰.۲۵۰	۰.۹۴۸	۰.۲۴۴	۰.۹۴۵	۰.۲۵۱	۰.۹۴۸	۲	۲	۰.۹۵	۱۹
۰.۲۲۹	۰.۸۱۸	۰.۱۶۲	۰.۹۵۲	۰.۱۷۷	۰.۹۶۳	۰.۲۳۰	۰.۹۱۸	۴	۲	۰.۹۵	۲۰
۰.۲۲۳	۰.۸۳۷	۰.۱۴۶	۰.۹۵۱	۰.۱۸۰	۰.۹۷۹	۰.۲۲۴	۰.۸۹۴	۶	۲	۰.۹۵	۲۱
۰.۱۳۴	۰.۷۶۳	۰.۱۹۶	۰.۹۵۰	۰.۲۹۸	۰.۹۸۵	۰.۱۳۵	۰.۹۱۷	۲	۴	۰.۹۵	۲۲
۰.۰۹۹	۰.۸۵۰	۰.۰۹۹	۰.۹۵۲	۰.۰۹۹	۰.۹۶۱	۰.۰۹۹	۰.۹۵۱	۴	۴	۰.۹۵	۲۳
۰.۰۸۸	۰.۸۸۰	۰.۰۸۱	۰.۹۵۲	۰.۰۹۱	۰.۹۷۸	۰.۰۸۸	۰.۹۴۵	۶	۴	۰.۹۵	۲۴
۰.۱۱۵	۰.۷۷۵	۰.۱۸۵	۰.۹۵۲	۰.۳۵۵	۰.۹۹۲	۰.۱۱۵	۰.۸۹۵	۲	۶	۰.۹۵	۲۵
۰.۰۷۷	۰.۸۵۹	۰.۰۸۴	۰.۹۵۰	۰.۱۰۱	۰.۹۸۳	۰.۰۷۷	۰.۹۴۴	۴	۶	۰.۹۵	۲۶
۰.۰۶۳	۰.۸۹۴	۰.۰۶۳	۰.۹۵۱	۰.۰۶۴	۰.۹۵۷	۰.۰۶۳	۰.۹۵۰	۶	۶	۰.۹۵	۲۷

- همان‌طور که انتظار داشتیم، سطح پوشش فاصله اطمینان دقیق بر اساس کمیت محور (۳۷.۲) ، به سطح معنی‌داری مد نظر α نزدیک‌تر می‌باشد. با این حال، براساس معیار میانگین طول، فاصله اطمینان بوت‌استرپ پایه بهتر از سایر فواصل اطمینان عمل می‌کند. همچنین فاصله اطمینان t - بوت‌استرپ به سطح اطمینان مورد نظر نزدیک‌تر است.

۵.۲ مثال واقعی

در این بخش قصد داریم برآوردگرهای نقطه‌ای و فاصله‌ای به‌دست آمده در بخش‌های ۲.۲ و ۳.۲ را با استفاده از یک مجموعه داده‌ی واقعی، آزمایش کنیم. بدین منظور از داده‌های مربوط به آزمایشات موتور راکت که قبلاً توسط گاتمن^{۱۳} و همکاران (۱۹۸۸) و کوتز^{۱۴} و همکاران (۲۰۰۳)، صفحه ۲۰۵، گزارش شده‌است، استفاده می‌کنیم. در اینجا متغیر تصادفی Y ، مقاومت موتور راکت در برابر فشار عملیاتی X است. برطبق داده‌های گزارش‌شده، فشار X به دمای Z بستگی دارد که در بررسی انجام شده، داده‌های X مربوط به دمای $Z = ۵۹^\circ$ در نظر گرفته شده‌است. این داده‌ها شامل ۵۱ مقدار برای (X, Z) و ۱۷ مقدار برای متغیر

^{۱۳}Guttman

^{۱۴}Kotz

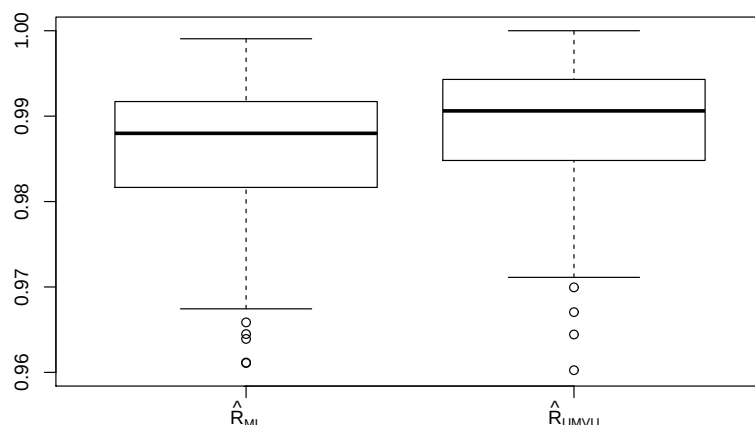
جدول ۲.۲: مقادیر EL و CP فواصل اطمینان ۹۰٪ برای R .

EL.Basic	CP.Basic	EL.ML	CP.ML	EL.Boot-t	CP.Boot-t	EL.Perc	CP.Perc	m	n	R	Row
۰.۳۱۰	۰.۷۴۴	۰.۳۱۱	۰.۸۹۷	۰.۳۳۲	۰.۹۰۶	۰.۳۱۱	۰.۸۹۷	۲	۲	۰.۱	۱
۰.۱۹۴	۰.۷۶۳	۰.۲۵۲	۰.۹۰۰	۰.۳۹۸	۰.۹۷۰	۰.۱۹۵	۰.۸۷۲	۴	۲	۰.۱	۲
۰.۱۷۱	۰.۷۵۸	۰.۲۳۷	۰.۹۰۱	۰.۴۷۵	۰.۹۸۴	۰.۱۷۲	۰.۸۴۸	۶	۲	۰.۱	۳
۰.۲۸۵	۰.۸۰۴	۰.۲۲۷	۰.۸۹۸	۰.۲۷۵	۰.۹۳۰	۰.۲۸۶	۰.۸۷۰	۲	۴	۰.۱	۴
۰.۱۴۶	۰.۸۳۲	۰.۱۴۷	۰.۹۰۱	۰.۱۵۱	۰.۹۱۸	۰.۱۴۷	۰.۹۰۱	۴	۴	۰.۱	۵
۰.۱۱۶	۰.۸۳۸	۰.۱۲۴	۰.۹۰۲	۰.۱۵۱	۰.۹۵۵	۰.۱۱۶	۰.۸۹۶	۶	۴	۰.۱	۶
۰.۲۷۴	۰.۸۱۸	۰.۲۰۷	۰.۸۹۷	۰.۲۸۰	۰.۹۵۶	۰.۲۷۵	۰.۸۵۰	۲	۶	۰.۱	۷
۰.۱۲۹	۰.۸۵۱	۰.۱۲۱	۰.۹۰۰	۰.۱۴۰	۰.۹۴۶	۰.۱۲۹	۰.۸۹۵	۴	۶	۰.۱	۸
۰.۰۹۶	۰.۸۶۸	۰.۰۹۶	۰.۹۰۱	۰.۰۹۸	۰.۹۱۱	۰.۰۹۶	۰.۹۰۲	۶	۶	۰.۱	۹
۰.۵۵۶	۰.۶۸۰	۰.۵۵۸	۰.۸۹۷	۰.۷۶۰	۰.۹۰۲	۰.۵۵۷	۰.۸۹۶	۲	۲	۰.۵	۱۰
۰.۴۷۲	۰.۷۳۹	۰.۴۷۹	۰.۹۰۰	۰.۷۹۰	۰.۹۶۰	۰.۴۷۳	۰.۸۷۲	۴	۲	۰.۵	۱۱
۰.۴۴۸	۰.۷۵۳	۰.۴۵۸	۰.۹۰۱	۰.۸۴۱	۰.۹۸۲	۰.۴۴۹	۰.۸۴۸	۶	۲	۰.۵	۱۲
۰.۴۷۲	۰.۷۴۰	۰.۴۷۹	۰.۸۹۸	۰.۷۹۲	۰.۹۵۸	۰.۴۷۳	۰.۸۶۹	۲	۴	۰.۵	۱۳
۰.۳۴۴	۰.۸۱۶	۰.۳۴۵	۰.۹۰۱	۰.۴۰۵	۰.۹۳۲	۰.۳۴۵	۰.۹۰۰	۴	۴	۰.۵	۱۴
۰.۳۰۰	۰.۸۴۰	۰.۳۰۱	۰.۹۰۲	۰.۳۹۹	۰.۹۶۱	۰.۳۰۱	۰.۸۹۶	۶	۴	۰.۵	۱۵
۰.۴۴۸	۰.۷۵۵	۰.۴۵۸	۰.۸۹۷	۰.۸۴۲	۰.۹۸۱	۰.۴۴۹	۰.۸۴۹	۲	۶	۰.۵	۱۶
۰.۳۰۰	۰.۸۳۵	۰.۳۰۱	۰.۹۰۰	۰.۳۹۹	۰.۹۵۹	۰.۳۰۱	۰.۸۹۵	۴	۶	۰.۵	۱۷
۰.۲۴۵	۰.۸۶۲	۰.۲۴۶	۰.۹۰۱	۰.۲۷۱	۰.۹۲۵	۰.۲۴۶	۰.۹۰۱	۶	۶	۰.۵	۱۸
۰.۱۹۶	۰.۷۴۶	۰.۱۹۷	۰.۸۹۷	۰.۱۸۳	۰.۹۰۱	۰.۱۹۷	۰.۸۹۶	۲	۲	۰.۹۵	۱۹
۰.۱۷۸	۰.۷۹۹	۰.۱۳۴	۰.۹۰۰	۰.۱۴۷	۰.۹۲۸	۰.۱۷۹	۰.۸۷۲	۴	۲	۰.۹۵	۲۰
۰.۱۷۳	۰.۸۲۰	۰.۱۲۲	۰.۹۰۱	۰.۱۵۳	۰.۹۵۱	۰.۱۷۴	۰.۸۴۸	۶	۲	۰.۹۵	۲۱
۰.۱۱۰	۰.۷۴۹	۰.۱۵۱	۰.۸۹۸	۰.۲۱۷	۰.۹۶۸	۰.۱۱۰	۰.۸۷۰	۲	۴	۰.۹۵	۲۲
۰.۰۸۱	۰.۸۲۶	۰.۰۸۱	۰.۹۰۱	۰.۰۷۹	۰.۹۰۴	۰.۰۸۱	۰.۹۰۱	۴	۴	۰.۹۵	۲۳
۰.۰۷۲	۰.۸۵۵	۰.۰۶۷	۰.۹۰۲	۰.۰۷۳	۰.۹۴۳	۰.۰۷۲	۰.۸۹۶	۶	۴	۰.۹۵	۲۴
۰.۰۹۶	۰.۷۵۸	۰.۱۴۲	۰.۸۹۷	۰.۲۶۲	۰.۹۸۳	۰.۰۹۶	۰.۸۵۰	۲	۶	۰.۹۵	۲۵
۰.۰۶۴	۰.۸۳۳	۰.۰۶۹	۰.۹۰۰	۰.۰۸۱	۰.۹۴۹	۰.۰۶۴	۰.۸۹۵	۴	۶	۰.۹۵	۲۶
۰.۰۵۲	۰.۸۶۲	۰.۰۵۲	۰.۹۰۱	۰.۰۵۲	۰.۹۰۴	۰.۰۵۲	۰.۹۰۱	۶	۶	۰.۹۵	۲۷

Y می‌باشند. در ابتدا برای بررسی مطابقت توزیع نمایی تعمیم یافته بر داده‌های فوق از معیار کولوموگروف-اسمیرنوف^{۱۵} استفاده نمودیم. نتایج مربوط به P -value آزمون کولوموگروف-اسمیرنوف به ترتیب ۰.۱۵۴ و ۰.۶۸۸ به دست آمد که نشان می‌دهد توزیع نمایی تعمیم یافته بر داده‌های فوق مطابقت دارد. مقادیر $RRSS$ پایین استخراج شده براساس متغیر X در دمای $Z = 59^\circ$ و متغیر Y براساس دیاگرام (۴.۱)، به ترتیب $r = (7/74010, 7/7227, 7/29040)$ و $s = (15/30, 16/30, 16)$ می‌باشد که $n = m = 3$ است. همان‌طور که قبلاً در بخش ۶.۱ بیان شد، داده‌های مربوط به $RRSS$ ها لزوماً مرتب نمی‌باشند که داده‌های فوق نیز این موضوع را تأیید می‌کند. برآوردگرهای نقطه‌ای مربوط به قابلیت اعتماد تنش-مقاومت براساس $RRSS$ پایین به دست آمده برابر است با $\hat{R}_{ML} = 0/9879$ و $\hat{R}_{UMVU} = 0/9899$. برای بررسی عملکرد برآوردهای نقطه‌ای و فاصله‌ای نسبت به یکدیگر از روش بوت‌استرپ استفاده نمودیم تا با باز نمونه‌گیری داده‌ها، مشاهدات بیشتری برای بررسی برآوردگرها فراهم نماییم. بدین منظور ۲۰۰ نمونه‌ی بوت‌استرپی از داده‌های اصلی تولید و سپس $RRSS$ های پایین را استخراج نمودیم. نمودار جعبه‌ای مشاهدات برآوردگرهای نقطه‌ای در شکل ۴.۲ نمایش داده شده است. مشاهده می‌شود دامنه‌ی میان چارکی $\hat{R}_{UMVU}$ کمتر از $\hat{R}_{ML}$ است، بنابراین با توجه به مقادیر برآورد شده، همان‌طور که قبلاً در شکل ۳.۲ نیز مشاهده کردیم، برای مقادیر انتهایی R ، برآوردگر $\hat{R}_{UMVU}$

^{۱۵}Kolmogorov- Smirnov

بهرتر از $\hat{R}_{ML}$ عمل می‌کند. برای بررسی فواصل اطمینان معرفی شده در بخش ۳.۲، فرآیند



شکل ۴.۲: مقایسه $\hat{R}_{UMVU}$ و $\hat{R}_{ML}$.

بوت‌استرپ ذکر شده در بالا عیناً ۱۰۰۰ بار دیگر تکرار شد و مقادیر مربوط به میانگین طول فواصل اطمینان بر طبق جدول ۳.۲ گزارش گردید. بر طبق نتایج جدول فوق، مشاهده می‌شود

جدول ۳.۲: مقادیر EL براساس داده‌های واقعی

EL.Basic	EL.ML	EL.Boot-t	EL.Perc	α
۰.۰۴۵	۰.۰۳۴	۰.۰۳۴	۰.۰۴۷	۰.۰۵
۰.۰۳۳	۰.۰۲۷	۰.۰۲۴	۰.۰۳۴	۰.۱

میانگین طول فاصله اطمینان همانند نتایج مربوط به شبیه‌سازی می‌باشد. به‌طور دقیق‌تر، فواصل t -بوت‌استرپ و ML بهتر از فاصله اطمینان بوت‌استرپ پایه و صدکی عمل کرده‌است. در حقیقت نتایج به‌دست آمده از طریق داده‌های واقعی، کاملاً منطبق بر نتایج بخش قبل می‌باشند.

نتیجه‌گیری

در این فصل، برآوردهایی از قابلیت اعتماد تنش-مقاومت را تحت شرایطی در نظر گرفتیم که تنها اطلاعات موجود از جامعه، $RRSS$ ‌های پایین به‌دست آمده از توزیع نمایی تعمیم یافته می‌باشند. برآوردهای ML و $UMVU$ را به‌دست آوردیم و خواص توزیعی برآوردهای فوق را بررسی نمودیم. در نهایت رفتار برآوردهای فوق نسبت به R و همچنین در مقایسه با یکدیگر، براساس معیار کارایی نسبی بررسی گردید. مشاهده شد، زمانی که R نزدیک به ۰.۵ است برآوردهای ML بهتر

از $UMVU$ عمل می‌کند در حالی که برای مقادیر فرین R ، عملکرد برآوردگر $UMVU$ بهتر بود. فواصل اطمینان متعددی را به دست آوردیم و به کمک شبیه سازی مونت کارلو و با در نظر گرفتن دو معیار میانگین طول فاصله اطمینان و سطح پوشش برآوردگرهای فوق را مقایسه نمودیم. با بررسی نتایج به دست آمده، دریافتیم فاصله اطمینان دقیق، براساس کمیت محور بهتر از سایر فواصل اطمینان عمل می‌کند. در نهایت برآوردگرهای نقطه‌ای و فاصله‌ای به دست آمده را در یک مثال واقعی بررسی نمودیم که نتایج حاصل کاملاً منطبق بر نتایج نظری و شبیه سازی‌های صورت گرفته در این فصل می‌باشند.

فصل ۳

استنباط بیزی پارامتر R در توزیع نمایی تعمیم یافته با استفاده از $RRSS$

۱.۳ مقدمه

هدف اصلی ما در این فصل، استنباط در خصوص پارامتر تنش-مقاومت در توزیع نمایی تعمیم یافته براساس طرح $RRSS$ از نقطه نظر بیزی می‌باشد. در سال‌های اخیر رویکرد و روش بیزی به‌طور گسترده‌ای در علوم مختلف به کار برده شده‌است. روش بیزی را می‌توان به عنوان روش دیگری که در مقابل روش کلاسیک می‌باشد، در نظر گرفت. در حقیقت به کار بردن توزیع‌های پیشین دلیلی است بر اینکه روش بیزی نسبت به روش کلاسیک دارای مزیت و برتری می‌باشد. در روش بیزی پارامترها به عنوان یک متغیر تصادفی در نظر گرفته می‌شوند که دارای یک توزیع احتمالاتی می‌باشند که همان توزیع‌های پیشین هستند. در روش‌های بیزی توزیع احتمالات ذهنی و بر پایه درجه و میزان اعتقاد محقق از تجربیات گذشته می‌باشد و نتایج اصلی و تمام استنباط‌ها براساس توزیع‌های پسین می‌باشد یعنی اطلاعات پیشین با استفاده از داده‌ها و مشاهداتی که از نمونه به دست می‌آیند به‌روز شده و براساس این اطلاعات استنباط و برآورد انجام می‌شود. همچنین در این روش فواصل معتبر فواصلی هستند که تحلیل‌گر با ترکیب دانش و اطلاعات

پیشین با داده‌ها و مشاهدات نمونه، با اطمینان $100(1 - \alpha)\%$ مطمئن است که پارامتر درست در این بازه قرار می‌گیرد. در این دیدگاه بکلیزی^۱ (۲۰۰۸)، با در نظر گرفتن توزیع نمایی تعمیم یافته و با استفاده از مقادیر رکوردی به استنباط در خصوص پارامتر تنش-مقاومت پرداخت. از جمله افراد دیگر که از دیدگاه بیزی به مسئله تنش-مقاومت پرداخته‌اند می‌توان به کندو و گوپتا^۲ (۲۰۰۵)، اصغرزاده و همکاران (۲۰۱۶)، اشاره کرد. در ادامه، این فصل به شرح زیر سازماندهی شده است:

در بخش ۲.۳، برآورگر بیزی پارامتر تنش-مقاومت را با در نظر گرفتن سه حالت مختلف براساس $RRSS$ های پایین از توزیع نمایی تعمیم یافته به دست می‌آوریم و سپس به کمک نمودار، رفتار برآورگر فوق را از منظر میانگین توان دوم خطا در مقابل R ، بررسی می‌نماییم. در بخش ۳.۳، فواصل معتبر بیزی و فواصل HPD را بر طبق توزیع های پیشین مزدوج و همچنین ناآگاهی بخش به دست می‌آوریم. در نهایت، به ترتیب در بخش های ۴.۳ و ۵.۳، به کمک شبیه سازی فواصل بیزی به دست آمده را مقایسه نموده و با یک مثال واقعی نتایج حاصل شده را بررسی می‌نماییم.

۲.۳ برآورگر بیزی پارامتر تنش-مقاومت

در این بخش با در نظر گرفتن توزیع نمایی تعمیم یافته به عنوان توزیع جوامع X و Y و براساس طرح $RRSS$ ، به محاسبه ی برآورگر بیزی برای $R = Pr(X < Y)$ می‌پردازیم. در حقیقت فرض براین است که پارامترهای $\theta_1, \theta_2, \lambda_1, \lambda_2$ متغیرهای تصادفی مستقل از هم باشند و تحت حالت های زیر برآورد بیزی R را به دست می‌آوریم.

حالت اول: $\lambda_1 \neq \lambda_2$ و $\theta_1 \neq \theta_2$

یک انتخاب معمول برای توزیع های پیشین پارامترهای $\theta_1, \theta_2, \lambda_1, \lambda_2$ به صورت زیر می‌باشد

$$\theta_i \sim \text{Gamma}(a_i, b_i), \quad \lambda_i \sim \text{Gamma}(c_i, d_i), \quad (i = 1, 2) \quad (1.3)$$

که در آن $a_i, b_i, c_i, d_i (> 0)$ ابر-پارامترهای مربوط به توزیع های پیشین مستقل از هم هستند. تابع

^۱ Baklizi

^۲ Kundu and Gupta

چگالی پسین توأم پارامترهای $\theta_1, \theta_2, \lambda_1, \lambda_2$ متناسب است با

$$\begin{aligned} \pi(\theta_1, \theta_2, \lambda_1, \lambda_2 | \mathbf{r}, \mathbf{s}) &\propto \theta_1^{a_1+N-1} e^{-\theta_1(b_1+z_{\lambda_1}(\mathbf{r}))} \times \theta_2^{a_2+M-1} e^{-\theta_2(b_2+z_{\lambda_2}(\mathbf{s}))} \\ &\times \lambda_1^{n+c_1-1} \prod_{i=1}^n \{\log u_i\}^{i-1} e^{\left(-\lambda_1 \left[\sum_{i=1}^n r_{i,i} + d_1\right] + z_{\lambda_1}(\mathbf{r})\right)} \\ &\times \lambda_2^{m+c_2-1} \prod_{j=1}^m \{\log v_j\}^{j-1} e^{\left(-\lambda_2 \left[\sum_{j=1}^m s_{j,j} + d_2\right] + z_{\lambda_2}(\mathbf{s})\right)} \end{aligned} \quad (2.3)$$

که از حاصل ضرب تابع درستنمایی (۵.۲) در توزیع‌های پیشین (۱.۳) به دست می‌آید. در این توزیع‌های پسین داریم

$$u_i = 1 - e^{-\lambda_1 r_{i,i}}, v_j = 1 - e^{-\lambda_2 s_{j,j}}, z_{\lambda}(\mathbf{x}) = -\sum_{i=1}^k \log(1 - e^{-\lambda x_i}) \quad (3.3)$$

به طوری که $\mathbf{x} = (x_1, x_2, \dots, x_k)^T$ و λ می‌تواند هر عدد حقیقی مثبت باشد. تحت تابع زیان مربع خطا (SEL)، برآوردگر بیزی $\mathcal{R}$ با توجه به رابطه‌ی (۳.۲) برابر است با $E(\mathcal{R} | \mathbf{R}, \mathbf{S})$ یا به طور معادل

$$\hat{\mathcal{R}}_B = \int_0^\infty \int_0^\infty \int_0^\infty \int_0^\infty \mathcal{R} \pi(\theta_1, \theta_2, \lambda_1, \lambda_2 | \mathbf{r}, \mathbf{s}) d\theta_1 d\theta_2 d\lambda_1 d\lambda_2. \quad (4.3)$$

بی‌گمان انتگرال (۴.۳) راه حل صریحی ندارد. از این رو به ناچار از بعضی الگوریتم‌های عددی یا الگوریتم‌های مبتنی بر شبیه‌سازی مونت کارلویی استفاده می‌کنیم. در حقیقت برای حل انتگرال فوق دو راه حل پیشنهاد می‌گردد.

۱- استفاده از روش نمونه‌گیری گیبس^۳ برای تولید متغیرهای تصادفی از رابطه (۲.۳) و سپس به کار بردن انتگرال‌گیری مونت کارلویی

۲- به کارگیری نمونه‌گیری از نقاط مهم^۴

که در اینجا روش دوم به طور دقیق توصیف می‌گردد. تابع چگالی رابطه (۲.۳) را به صورت زیر بازنویسی می‌کنیم

$$\pi(\theta_1, \theta_2, \lambda_1, \lambda_2 | \mathbf{r}, \mathbf{s}) \propto f_1(\theta_1 | \lambda_1, \mathbf{r}) f_2(\lambda_1 | \mathbf{r}) h_1(\theta_1, \lambda_1 | \mathbf{r}) \times g_1(\theta_2 | \lambda_2, \mathbf{s}) g_2(\lambda_2 | \mathbf{s}) h_2(\theta_2, \lambda_2 | \mathbf{s}), \quad (5.3)$$

که

$$f_1(\theta_1 | \lambda_1, \mathbf{r}) \sim \text{Gamma}(a_1 + N, b_1 + z_{\lambda_1}(\mathbf{r}))$$

^۳Gibbs

^۴Importance sampling

و

$$g_1(\theta_1 | \lambda_1, \mathbf{s}) \sim \text{Gamma}(a_1 + M, b_1 + z_{\lambda_1}(\mathbf{s}))$$

همچنین

$$f_1(\lambda_1 | \mathbf{r}) \propto \lambda_1^{n+c_1-1} e^{\left(-\lambda_1 \left[\sum_{i=1}^n r_{i,i} + d_1 \right] + z_{\lambda_1}(\mathbf{r})\right)}$$

و

$$g_2(\lambda_2 | \mathbf{s}) \propto \lambda_2^{m+c_2-1} e^{\left(-\lambda_2 \left[\sum_{j=1}^m s_{j,j} + d_2 \right] + z_{\lambda_2}(\mathbf{s})\right)}$$

دو تابع چگالی سره می باشند. به علاوه

$$h(\theta_1, \theta_2, \lambda_1, \lambda_2 | \mathbf{r}, \mathbf{s}) = \prod_{i=1}^n \{\log(1 - e^{-\lambda_1 r_{i,i}})\}^{i-1} \times \prod_{j=1}^m \{\log(1 - e^{-\lambda_2 s_{j,j}})\}^{j-1}.$$

در نهایت با استفاده از تکنیک نمونه گیری از نقاط مهم و با استفاده از الگوریتم زیر می توان به یک تقریب خوب از رابطه (۴.۳)، دست یافت.

الگوریتم ۶ الگوریتم نمونه گیری از نقاط مهم

۱. تولید $\theta_1 \sim \text{Gamma}(a_1 + N, b_1 + z_{\lambda_1}(\mathbf{r}))$ و $\theta_2 \sim \text{Gamma}(a_2 + M, b_2 + z_{\lambda_2}(\mathbf{s}))$.
۲. تولید λ_1 و λ_2 به ترتیب از $f_1(\lambda_1 | \mathbf{r})$ و $g_2(\lambda_2 | \mathbf{s})$ ، از طریق روش های مونت کارلو ($MCMC$) همانند قدم زدن تصادفی.
۳. تکرار مراحل یک و دو B بار برای به دست آوردن مشاهدات $(\theta_1^{(t)}, \theta_2^{(t)}, \lambda_1^{(t)}, \lambda_2^{(t)})$ ، $t = 1, \dots, B$ و سپس محاسبه $\mathcal{R}^{(t)}$ از رابطه (۳.۲).
۴. سپس برآورد رابطه (۴.۳)، که با $\hat{\mathcal{R}}_B$ نشان می دهیم با استفاده از میانگین وزنی

$$\hat{\mathcal{R}}_B = \sum_{t=1}^B w_t \mathcal{R}^{(t)}, \quad (۶.۳)$$

$$w_t = \frac{h(\theta_1^{(t)}, \theta_2^{(t)}, \lambda_1^{(t)}, \lambda_2^{(t)} | \mathbf{r}, \mathbf{s})}{\sum_{t=1}^B h(\theta_1^{(t)}, \theta_2^{(t)}, \lambda_1^{(t)}, \lambda_2^{(t)} | \mathbf{r}, \mathbf{s})} \text{ که } w_t \text{ می باشد.}$$

حالت دوم: $\lambda_1 \neq \lambda_2$ و $\theta_1 = \theta_2 = \theta$

در این حالت فرض کنید پارامترهای $(\theta, \lambda_1, \lambda_2)$ متغیرهای تصادفی مستقل با تابع چگالی پیشین

گاما به صورت زیر باشند؛

$$\theta \sim \text{Gamma}(a, b), \quad \lambda_i \sim \text{Gamma}(c_i, d_i), \quad (i = 1, 2) \quad (۷.۳)$$

که $a, b, c_i, d_i (> 0), i = 1, 2$ ابر- پارامترها می باشند. با روشی مشابه حالت اول، تابع چگالی توأم پسین $(\theta, \lambda_1, \lambda_2)$ به صورت زیر به دست می آید.

$$\begin{aligned} \pi(\theta, \lambda_1, \lambda_2 | \mathbf{r}, \mathbf{s}) &\propto \theta^{N+M+a-1} e^{-\theta(b+z_{\lambda_1}(\mathbf{r})+z_{\lambda_2}(\mathbf{s}))} \\ &\times \lambda_1^{n+c_1-1} \prod_{i=1}^n \{\log u_i\}^{i-1} e^{\left(-\lambda_1 \left[\sum_{i=1}^n r_{i,i} + d_1\right] + z_{\lambda_1}(\mathbf{r})\right)} \\ &\times \lambda_2^{m+c_2-1} \prod_{j=1}^m \{\log v_j\}^{j-1} e^{\left(-\lambda_2 \left[\sum_{j=1}^m s_{j,j} + d_2\right] + z_{\lambda_2}(\mathbf{s})\right)} \end{aligned} \quad (۸.۳)$$

که $z_{\lambda}(\mathbf{x})$ و u_i و v_j در رابطه (۳.۳) بیان شده است. همانند حالت اول، به دست آوردن $\hat{R}_B$ از طریق رابطه بسته امکان پذیر نمی باشد. بنابراین مشابه با روشی که در حالت اول به طور کامل توضیح داده شد، مقدار تقریبی $\hat{R}_B$ را محاسبه می کنیم.

حالت سوم: $\lambda_1 = \lambda_2 = \lambda$ و $\theta_1 \neq \theta_2$ با استفاده از (۲۲.۲)، مشاهده می شود در این حالت، پارامتر $\mathcal{R}$ به λ بستگی ندارد. در نتیجه بدون از دست دادن کلیت مسئله، فرض کنید $\lambda = 1$. همچنین فرض کنید θ_1 و θ_2 متغیرهای تصادفی به ترتیب با توزیع گاما با پارامترهای (a_1, b_1) و (a_2, b_2) باشند. بنابراین

$$\theta_1 | \mathbf{r} \sim \text{Gamma}(a_1 + N, b_1 + t_1) \quad \text{و} \quad \theta_2 | \mathbf{s} \sim \text{Gamma}(a_2 + M, b_2 + t_2)$$

که در آن $t_i, i = 1, 2$ مقادیر مشاهده شده متغیرهای تصادفی $T_i, i = 1, 2$ می باشند. در نتیجه

$$2(b_1 + T_1) \theta_1 | \mathbf{r} \sim \chi_{2(a_1+N)}^2 \quad \text{و} \quad 2(b_2 + T_2) \theta_2 | \mathbf{s} \sim \chi_{2(a_2+M)}^2. \quad (۹.۳)$$

بنابراین تابع چگالی پسین $\mathcal{R}$ به شرط $(\mathbf{R}, \mathbf{S})$ به فرم زیر خواهد بود:

$$\pi_{\mathcal{R} | (\mathbf{R}, \mathbf{S})}(r) = C \frac{r^{a_2+M-1} (1-r)^{a_1+N-1}}{[(b_1+t_1)(1-r) + (b_2+t_2)r]^{M+N+a_1+a_2}}, \quad 0 < r < 1, \quad (۱۰.۳)$$

که

$$C = \frac{1}{B(a_1 + N, a_2 + M)} (b_1 + t_1)^{a_1+N} (b_2 + t_2)^{a_2+M}.$$

و $B(.,.)$ تابع بتا می باشد.

در نهایت برآورد بیزی $\mathcal{R}$ تحت تابع زیان درجه دوم خطا به صورت زیر به دست می آید.

$$\begin{aligned}\hat{\mathcal{R}}_B &= \int_0^1 r \pi_{\mathcal{R}|\mathbf{(R,S)}}(r) dr = \int_0^1 C \frac{r^{a_2+M}(1-r)^{a_1+N-1}}{[(b_1+T_1)(1-r) + (b_2+T_2)r]^{M+N+a_1+a_2}} dr \\ &= \frac{a_2+M}{M+N+a_1+a_2} \left(\frac{b_2+T_2}{b_1+T_1} \right)^{a_2+M} \\ &\times {}_2F_1 \left(M+N+a_1+a_2, M+a_2+1, M+N+a_1+a_2+1, 1 - \frac{b_2+T_2}{b_1+T_1} \right).\end{aligned}\quad (11.3)$$

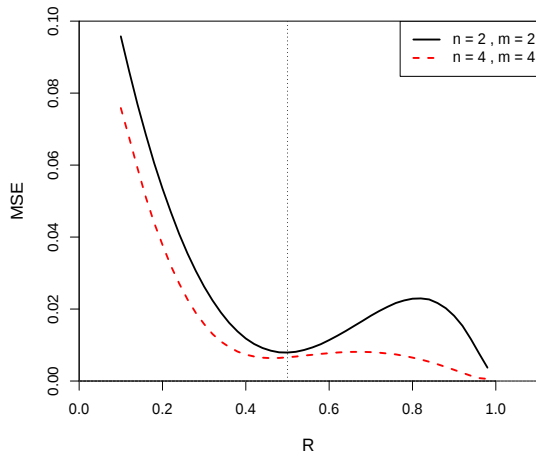
از این رو فقط در این حالت برآوردگر بیزی $\mathcal{R}$ به صورت بسته به دست می آید. در این بخش عملکرد برآورگرهای بیزی در حالت های اول و سوم، با کمک شبیه سازی مورد بررسی قرار گرفت. بدین منظور، دو مجموعه از ابر-پارامترها را در نظر گرفتیم، اولی یک پیشین جفریز را تولید می کند و دیگری اشاره به پیشین مناسبی دارد به طوری که، $a_i = b_i = 2$ و $c_i = d_i = 1$ ، $i = 1, 2$. سپس به ازای θ_i و λ_i ، $i = 1, 2$ داده شده که $\mathcal{R}$ را نتیجه می دهد، $RRSS$ های $\mathbf{r}$ و $\mathbf{s}$ را از توزیع نمایی تعمیم یافته تولید می کنیم. با استفاده از الگوریتم (۶) برای حالت اول و رابطه (۱۱.۳) برای حالت سوم، مشاهدات $\hat{\mathcal{R}}_B$ را به دست می آوریم. به وضوح با تکرار ۱۰۰۰ بار این فرآیند، یک مجموعه از مشاهدات برای $\hat{\mathcal{R}}_B$ به دست می آید. سپس، مقادیر تابع مخاطره تحت زیان درجه دوم خطا (که در این جا با MSE نشان داده شده است) از طریق میانگین نمونه $(\hat{\mathcal{R}}_B - \mathcal{R})^2$ ها برآورد می شود. لازم به ذکر است مقدار θ_2 را ثابت و برابر یک می باشد و θ_1 را از یک مقدار کم به مقدارهای بزرگتر طوری در نظر گرفتیم که $\mathcal{R}$ مقادیر بین دامنه $(0, 1)$ را اختیار کند. شکل ۱.۳ (ب)، نتایج متناظر با پیشین مناسب برای حالت سوم می باشد. (که بسیار شبیه به حالت اول می باشد، بنابراین فقط یکی از آن ها در این قسمت نمایش داده شده است.)

با استفاده از شکل ۱.۳ (ب)، مشاهده می شود که برخلاف نمودارهای مربوط به برآوردهای نقطه ای روش کلاسیک که در فصل قبل در نمودارهای ۱.۲ و ۲.۲ نمایش داده شد، بیشترین مخاطره در اطراف مقادیر میانی $\mathcal{R}$ اتفاق نمی افتد و کاملاً بستگی به انتخاب ابر-پارامترها دارد. همچنین قابل توجه است که، همان طور که انتظار داشتیم، با توجه به شکل ۱.۳ (الف)، پیشین جفریز، شکلی تقریباً مشابه با شکل های مربوط به برآورگرهای کلاسیک رسم شده در نمودارهای ۱.۲ و ۲.۲ از فصل قبل دارد.

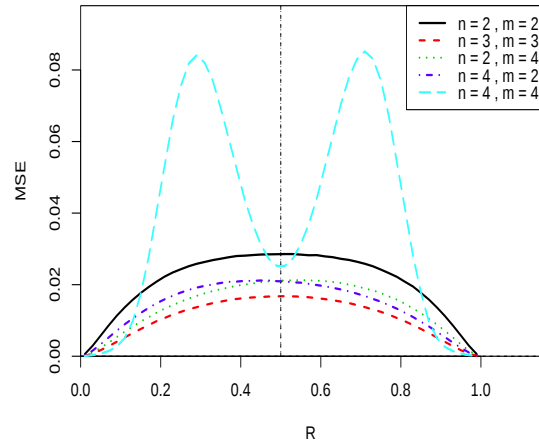
۳.۳ برآورگرهای فاصله ای بیزی پارامتر تنش-مقاومت

با استفاده از رابطه (۹.۳)، به طور واضح نتیجه می شود که

$$\frac{\theta_1(b_1+T_1)(a_2+M)}{\theta_2(b_2+T_2)(a_1+N)} \mid (\mathbf{r}, \mathbf{s}) \sim F(2(a_1+N), 2(a_2+M))$$



(ب)



(الف)

شکل ۱۳: نمودار $MSE(\hat{R}_B, R)$ در مقابل R برای حالت ۳ و با ابر-پارامترهای $a_i = b_i = 2$ و $c_i = d_i = 1$.

بنابراین فاصله معتبر بیزی $100(1 - \alpha)\%$ برای R به صورت زیر به دست می‌آید.

$$\left[\left(1 + AF_{1-\frac{\alpha}{2}}(Y(a_1 + N), Y(a_2 + M)) \right)^{-1}, \left(1 + AF_{\frac{\alpha}{2}}(Y(a_1 + N), Y(a_2 + M)) \right)^{-1} \right], \quad (12.3)$$

$$A = \frac{(a_1 + N)(b_2 + T_2)}{(a_2 + M)(b_1 + T_1)}$$

که با استفاده از توزیع پیشین ناآگاهی بخش نیز می‌توان به صورت مشابه، فاصله‌ی معتبر بیزی را به دست آورد. همان‌طور که می‌دانید تابع چگالی پیشین جفریز به صورت $\pi(\theta) \propto \sqrt{|I(\theta)|}$ تعریف می‌شود که در آن $I(\theta)$ همان ماتریس اطلاع فیشر می‌باشد. با به کار بردن رابطه (۱۶.۲) و عملیات ساده جبری داریم:

$$I(\theta_1) = \left[-E \frac{d^2}{d\theta_1^2} \log(\theta_1 | r_{i,i}) \right] = \frac{N}{\theta_1^2}$$

بنابراین توزیع پیشین پارامترهای θ_1 و θ_2 به ترتیب به صورت $1/\theta_1$ و $1/\theta_2$ به دست می‌آید لذا با استدلال‌های مستقیم، می‌توان نشان داد که

$$\mathcal{R}(\mathbf{r}, \mathbf{s}) \stackrel{d}{=} \left(1 + \frac{NT_2}{MT_1} W \right)^{-1} \quad (13.3)$$

که در آن $W \sim F_{\Upsilon N, \Upsilon M}$.

بنابراین براساس (۱۳.۳) فاصله معتبر بیزی دیگری برای $\mathcal{R}$ به فرم زیر به دست می آید.

$$\left[\left(1 + \frac{NT_{\Upsilon}}{MT_{\Upsilon}} F_{1-\frac{\alpha}{\Upsilon}}(\Upsilon N, \Upsilon M) \right)^{-1}, \left(1 + \frac{NT_{\Upsilon}}{MT_{\Upsilon}} F_{\frac{\alpha}{\Upsilon}}(\Upsilon N, \Upsilon M) \right)^{-1} \right]. \quad (14.3)$$

در ادامه می خواهیم چگال ترین فاصله پسین (HPD)، را برای $\mathcal{R}$ به دست آوریم. بنابراین در ابتدا تعریف رسمی از HPD را ارائه می کنیم.

تعریف ۱.۳.۳. فرض کنید $\Pi(\theta|D)$ ، تابع توزیع پسین پیوسته و تک مدی باشد، در این صورت بازه HPD ، $100(1-\alpha)\%$ برای θ به صورت (θ_L, θ_U) خواهد بود که θ_U و θ_L از بهینه سازی معادله ی زیر به دست می آیند. (برای مثال برگر^۵ (۱۹۸۵)، را ببینید).

$$\min_{\theta_L < \theta_U} (|\pi(\theta_U|D) - \pi(\theta_L|D)| + |\Pi(\theta_U|D) - \Pi(\theta_L|D) - (1-\alpha)|).$$

بنابراین، یک فاصله HPD ، $100(1-\alpha)\%$ به صورت

$$C(\pi_{\alpha}) = \{\theta : \pi(\theta|D) \geq \pi_{\alpha}\}, \quad (15.3)$$

می باشد که D نشان دهنده داده ها و π_{α} بزرگ ترین ثابت می باشد که $\Pr(\theta \in C(\pi_{\alpha})) \geq 1-\alpha$ و $\Pi(\theta|D)$ و $\pi(\theta|D)$ به ترتیب تابع چگالی پسین و تابع توزیع پسین θ می باشند.

به نظر می رسد که فاصله HPD برای $\mathcal{R}$ فرم بسته ای نداشته باشد، اما با استفاده از رابطه (۱۰.۳) داریم:

$$\begin{aligned} \log \pi_{\mathcal{R}|\mathbf{(R,S)}}(r) &= (a_{\Upsilon} + M - 1) \log r + (a_1 + N - 1) \log(1-r) \\ &\quad - (M + N + a_1 + a_{\Upsilon}) \log[(b_1 + t_1)(1-r) + (b_{\Upsilon} + t_{\Upsilon})r] \end{aligned}$$

می توان نشان داد که $\frac{d}{dr} \log \pi_{\mathcal{R}|\mathbf{(R,S)}}(r) = 0$ دارای دو ریشه ۱، $r = 0$ می باشد، به طوری که

$$\lim_{r \rightarrow 0^+} \frac{d}{dr} \log \pi_{\mathcal{R}|\mathbf{(R,S)}}(r) > 0$$

و

$$\lim_{r \rightarrow 1^-} \frac{d}{dr} \log \pi_{\mathcal{R}|\mathbf{(R,S)}}(r) < 0$$

بنابراین به راحتی می توان نشان داد که $\pi_{\mathcal{R}|\mathbf{(R,S)}}(r)$ تابع چگالی تک مدی می باشد. چن و شائو^۶ (۱۹۹۹)، با یک تکنیک مونت کارلویی ساده تقریبی برای بازه HPD به دست آوردند. در لم زیر تکنیک فوق بیان می گردد.

^۵Berger

^۶Chen and Shao

لم ۱.۳.۳. فرض کنید $R_i, i = 1, 2, \dots, n$ نمونه تصادفی از (۱۰.۳) باشد و همچنین $R_{(i)}$ متناظر با مقادیر مرتب‌شده R_i می‌باشند. در نتیجه فاصله HPD ، $100(1-\alpha)\%$ برای R را می‌توان به صورت زیر تقریب زد

$$C_{j^*}(n) = \left(R_{(j^*)}, R_{(j^* + [(1-\alpha)n])} \right), \quad (16.3)$$

که در این رابطه j^* طوری انتخاب می‌شود که

$$R_{(j^* + [(1-\alpha)n])} - R_{(j^*)} = \min_{1 \leq j \leq n - [(1-\alpha)n]} (R_{(j + [(1-\alpha)n])} - R_{(j)}), \quad (17.3)$$

همچنین [.] نشان‌دهنده‌ی تابع جزء صحیح می‌باشد.

۴.۳ مقایسه برآوردهای فاصله‌ای بیزی براساس شبیه‌سازی

در این بخش برای مقایسه‌ی فواصل بیزی مطرح شده در بخش ۳.۳، از شبیه‌سازی مونت کارلویی استفاده کرده‌ایم و در آن مفروضات مطرح شده در بخش ۳.۲ برای تمام ترکیب‌های متفاوت $n = 2, 4, 6$ و $m = 2, 4, 6$ در دو جدول مجزا برای سطح معنی‌داری $\alpha = 0.05, 0.1$ در نظر گرفته شده‌است. به علاوه، توزیع پیشین پارامترهای θ_1 و θ_2 هر دو توزیع $Gamma(1, 2)$ می‌باشد. در ضمن برای حالتی که توزیع پیشین ناآگاهی بخش می‌باشد، توزیع پیشین جفریز در نظر گرفته شده‌است. جدول‌های ۱.۳ و ۲.۳ نتایج شبیه‌سازی را نشان می‌دهند که در آن نمادهای زیر به کار رفته‌اند:

- $BCI1$: فاصله بیزی به دست آمده براساس توزیع پیشین مزدوج.
 - $HPD1$: فاصله HPD به دست آمده براساس توزیع پیشین مزدوج.
 - $BCI2$: فاصله بیزی به دست آمده براساس توزیع پیشین جفریز.
 - $HPD2$: فاصله HPD به دست آمده براساس توزیع پیشین جفریز.
- به علاوه منظور از EL و CP در جدول‌های ۱.۳ و ۲.۳، به ترتیب میانگین طول و احتمال پوشش فواصل معتبر بیزی می‌باشد که در رابطه (۴۱.۲) معرفی شده‌اند. از جدول‌های ۱.۳ و ۲.۳ نتایج زیر کسب می‌شوند:
- با افزایش اندازه‌ی نمونه، میانگین طول فواصل بیزی کاهش می‌یابد.
 - برای اندازه‌ی نمونه‌های مختلف مشاهده می‌شود، میانگین طول فواصل بیزی، بیشترین مقدار خود را در $R = 0.5$ اختیار می‌کند و کمترین میانگین طول فواصل بیزی در $R = 0.95$ می‌باشد.
 - میانگین طول فواصل HPD کمتر از فاصله معتبر بیزی متناظرش می‌باشد.

جدول ۱.۳: مقادیر EL و CP فواصل بیزی ۹۵٪ برای R

EL.HPD۲	CP.HPD۲	EL.BCI۲	CP.BCI۲	EL.HPD۱	CP.HPD۱	EL.BCI۱	CP.BCI۱	m	n	R	Row
۰.۳۳۳	۰.۹۳۰	۰.۳۸۲	۰.۹۴۸	۰.۴۸۰	۰.۸۸۱	۰.۵۱۲	۰.۸۰۴	۲	۲	۰.۱	۱
۰.۲۷۱	۰.۹۴۲	۰.۳۱۷	۰.۹۵۲	۰.۴۱۲	۰.۷۸۲	۰.۴۴۲	۰.۵۸۰	۴	۲	۰.۱	۲
۰.۲۵۶	۰.۹۴۰	۰.۳۰۰	۰.۹۵۱	۰.۳۸۷	۰.۷۲۴	۰.۴۱۸	۰.۴۵۲	۶	۲	۰.۱	۳
۰.۲۵۰	۰.۹۳۰	۰.۲۷۰	۰.۹۵۰	۰.۲۹۹	۰.۹۶۳	۰.۳۱۵	۰.۹۴۸	۲	۴	۰.۱	۴
۰.۱۶۶	۰.۹۳۹	۰.۱۷۸	۰.۹۵۲	۰.۲۱۹	۰.۹۰۴	۰.۲۳۰	۰.۸۴۹	۴	۴	۰.۱	۵
۰.۱۴۱	۰.۹۴۳	۰.۱۵۱	۰.۹۵۲	۰.۱۸۸	۰.۸۸۰	۰.۱۹۸	۰.۸۰۲	۶	۴	۰.۱	۶
۰.۲۳۱	۰.۹۳۲	۰.۲۴۵	۰.۹۵۲	۰.۲۴۳	۰.۹۷۷	۰.۲۵۳	۰.۹۷۵	۲	۶	۰.۱	۷
۰.۱۴۰	۰.۹۴۰	۰.۱۴۶	۰.۹۵۰	۰.۱۶۲	۰.۹۴۳	۰.۱۶۸	۰.۹۲۳	۴	۶	۰.۱	۸
۰.۱۱۱	۰.۹۴۳	۰.۱۱۵	۰.۹۵۱	۰.۱۳۲	۰.۹۱۸	۰.۱۳۶	۰.۸۹۳	۶	۶	۰.۱	۹
۰.۶۲۱	۰.۹۰۴	۰.۶۴۴	۰.۹۴۸	۰.۵۸۰	۰.۹۵۴	۰.۵۹۶	۰.۹۸۳	۲	۲	۰.۵	۱۰
۰.۵۴۲	۰.۹۱۵	۰.۵۵۸	۰.۹۵۲	۰.۵۰۴	۰.۹۵۲	۰.۵۱۵	۰.۹۷۹	۴	۲	۰.۵	۱۱
۰.۵۱۹	۰.۹۱۷	۰.۵۳۴	۰.۹۵۱	۰.۴۷۷	۰.۹۵۵	۰.۴۸۸	۰.۹۸۰	۶	۲	۰.۵	۱۲
۰.۵۴۱	۰.۹۱۰	۰.۵۵۸	۰.۹۵۰	۰.۵۰۳	۰.۹۵۰	۰.۵۱۵	۰.۹۷۷	۲	۴	۰.۵	۱۳
۰.۳۹۹	۰.۹۲۴	۰.۴۰۶	۰.۹۵۲	۰.۳۸۴	۰.۹۳۶	۰.۳۹۰	۰.۹۶۱	۴	۴	۰.۵	۱۴
۰.۳۵۰	۰.۹۳۰	۰.۳۵۶	۰.۹۵۲	۰.۳۳۹	۰.۹۴۰	۰.۳۴۴	۰.۹۶۰	۶	۴	۰.۵	۱۵
۰.۵۱۹	۰.۹۱۶	۰.۵۳۳	۰.۹۵۲	۰.۴۷۷	۰.۹۵۴	۰.۴۸۸	۰.۹۷۹	۲	۶	۰.۵	۱۶
۰.۳۵۰	۰.۹۳۱	۰.۳۵۶	۰.۹۵۰	۰.۳۳۹	۰.۹۴۰	۰.۳۴۴	۰.۹۵۷	۴	۶	۰.۵	۱۷
۰.۲۸۷	۰.۹۳۶	۰.۲۹۱	۰.۹۵۱	۰.۲۸۱	۰.۹۳۷	۰.۲۸۵	۰.۹۵۶	۶	۶	۰.۵	۱۸
۰.۲۱۳	۰.۹۴۰	۰.۲۵۰	۰.۹۴۸	۰.۲۱۳	۰.۹۵۶	۰.۲۳۹	۰.۹۳۱	۲	۲	۰.۹۵	۱۹
۰.۱۵۰	۰.۹۴۲	۰.۱۶۲	۰.۹۵۲	۰.۱۶۲	۰.۹۳۷	۰.۱۷۴	۰.۹۱۵	۴	۲	۰.۹۵	۲۰
۰.۱۱۶	۰.۹۴۳	۰.۱۴۶	۰.۹۵۱	۰.۱۵۱	۰.۹۲۴	۰.۱۵۸	۰.۹۱۰	۶	۲	۰.۹۵	۲۱
۰.۱۶۴	۰.۹۴۳	۰.۱۹۶	۰.۹۵۰	۰.۱۴۱	۰.۹۷۰	۰.۱۶۴	۰.۹۶۵	۲	۴	۰.۹۵	۲۲
۰.۰۹۲	۰.۹۴۲	۰.۰۹۹	۰.۹۵۲	۰.۰۹۳	۰.۹۴۹	۰.۱۰۰	۰.۹۴۷	۴	۴	۰.۹۵	۲۳
۰.۰۷۷	۰.۹۴۵	۰.۰۸۱	۰.۹۵۲	۰.۰۷۹	۰.۹۴۶	۰.۰۸۳	۰.۹۴۰	۶	۴	۰.۹۵	۲۴
۰.۱۵۳	۰.۹۴۶	۰.۱۸۵	۰.۹۵۲	۰.۱۲۴	۰.۹۷۴	۰.۱۴۷	۰.۹۷۵	۲	۶	۰.۹۵	۲۵
۰.۰۷۷	۰.۹۴۴	۰.۰۸۴	۰.۹۵۰	۰.۰۷۶	۰.۹۵۱	۰.۰۸۲	۰.۹۵۳	۴	۶	۰.۹۵	۲۶
۰.۰۶۰	۰.۹۴۵	۰.۰۶۳	۰.۹۵۱	۰.۰۶۰	۰.۹۴۵	۰.۰۶۳	۰.۹۵۰	۶	۶	۰.۹۵	۲۷

- زمانی که $R = ۱/۰$ ، با در نظر گرفتن دو معیار EL و CP ، فاصله $HPD۲$ بهتر از $HPD۱$ و $BCI۲$ بهتر از $BCI۱$ عمل می کند ولی در $R = ۵/۰$ کاملاً برعکس عمل می کنند. به عبارت دیگر $HPD۱$ و $BCI۱$ بهتر از موارد متناظرشان عمل می کنند.
- در $R = ۱/۰, ۹۵/۰$ ، احتمال پوشش $HPD۱$ بیشتر از $HPD۲$ می باشد.
- زمانی که $R = ۱/۰, ۹۵/۰$ تقریباً همه جا $HPD۱$ بهتر از $BCI۱$ می باشد. با این حال، زمانی که $R = ۵/۰$ ، احتمال پوشش $BCI۱$ بیشتر از $HPD۱$ و همچنین $BCI۲$ بیشتر از $HPD۲$ می باشد.
- با مقایسه نتایج مربوط به دو جدول، مشاهده می شود میانگین طول فواصل بیزی با در نظر گرفتن $\alpha = ۱/۰$ در تمام موارد نسبت به مورد مشابهش با سطح معنی داری $\alpha = ۵/۰$ کاهش یافته است و لیکن احتمال پوشش در جدول ۱.۳ به ضریب اطمینان مدنظر α نزدیک تر می باشد.

۵.۳ مثال واقعی

در این بخش، فواصل معتبر بیزی و فواصل HPD به دست آمده در بخش ۳.۳ را با استفاده از یک مجموعه داده واقعی آزمایش می کنیم. بدین منظور از داده های مربوط به آزمایشات موتور

جدول ۲.۳: مقادیر EL و CP فواصل بیزی ۹۰٪ برای R

EL.HPD۲	CP.HPD۲	EL.BCI۲	CP.BCI۲	EL.HPD۱	CP.HPD۱	EL.BCI۱	CP.BCI۱	m	n	R	Row
۰.۲۶۷	۰.۸۶۸	۰.۳۱۱	۰.۸۹۷	۰.۴۰۴	۰.۷۹۷	۰.۴۳۳	۰.۶۳۵	۲	۲	۰.۱	۱
۰.۲۱۲	۰.۸۸۱	۰.۲۵۲	۰.۹۰۰	۰.۳۴۱	۰.۶۵۸	۰.۳۶۸	۰.۳۸۵	۴	۲	۰.۱	۲
۰.۱۹۹	۰.۸۸۷	۰.۲۳۷	۰.۹۰۱	۰.۳۱۹	۰.۵۸۰	۰.۳۴۷	۰.۲۴۵	۶	۲	۰.۱	۳
۰.۲۰۹	۰.۸۶۶	۰.۲۲۷	۰.۸۹۸	۰.۲۵۲	۰.۹۲۰	۰.۲۶۶	۰.۸۶۹	۲	۴	۰.۱	۴
۰.۱۳۷	۰.۸۸۳	۰.۱۴۷	۰.۹۰۱	۰.۱۸۲	۰.۸۴۰	۰.۱۹۱	۰.۷۴۶	۴	۴	۰.۱	۵
۰.۱۱۶	۰.۸۸۶	۰.۱۲۴	۰.۹۰۲	۰.۱۵۶	۰.۷۹۸	۰.۱۶۴	۰.۶۷۶	۶	۴	۰.۱	۶
۰.۱۹۵	۰.۸۷۲	۰.۲۰۷	۰.۸۹۷	۰.۲۰۵	۰.۹۴۲	۰.۲۱۳	۰.۹۲۵	۲	۶	۰.۱	۷
۰.۱۱۷	۰.۸۸۵	۰.۱۲۱	۰.۹۰۰	۰.۱۳۶	۰.۸۹۳	۰.۱۴۱	۰.۸۴۵	۴	۶	۰.۱	۸
۰.۰۹۲	۰.۸۹۱	۰.۰۹۶	۰.۹۰۱	۰.۱۱۰	۰.۸۵۶	۰.۱۱۳	۰.۸۱۰	۶	۶	۰.۱	۹
۰.۵۳۴	۰.۸۱۵	۰.۵۵۸	۰.۸۹۷	۰.۴۹۹	۰.۸۸۶	۰.۵۱۴	۰.۹۴۷	۲	۲	۰.۵	۱۰
۰.۴۶۴	۰.۸۳۷	۰.۴۷۹	۰.۹۰۰	۰.۴۳۱	۰.۸۹۶	۰.۴۴۱	۰.۹۴۰	۴	۲	۰.۵	۱۱
۰.۴۴۴	۰.۸۴۷	۰.۴۵۸	۰.۹۰۱	۰.۴۰۷	۰.۸۹۳	۰.۴۱۷	۰.۹۴۳	۶	۲	۰.۵	۱۲
۰.۴۶۴	۰.۸۳۴	۰.۴۷۹	۰.۸۹۸	۰.۴۳۰	۰.۸۹۰	۰.۴۴۱	۰.۹۳۹	۲	۴	۰.۵	۱۳
۰.۳۳۹	۰.۸۶۲	۰.۳۴۵	۰.۹۰۱	۰.۳۲۶	۰.۸۷۶	۰.۳۳۱	۰.۹۱۳	۴	۴	۰.۵	۱۴
۰.۲۹۷	۰.۸۷۳	۰.۳۰۱	۰.۹۰۲	۰.۲۸۷	۰.۸۸۰	۰.۲۹۱	۰.۹۱۲	۶	۴	۰.۵	۱۵
۰.۴۴۴	۰.۸۴۵	۰.۴۵۸	۰.۸۹۷	۰.۴۰۷	۰.۸۹۷	۰.۴۱۷	۰.۹۴۲	۲	۶	۰.۵	۱۶
۰.۲۹۷	۰.۸۷۰	۰.۳۰۱	۰.۹۰۰	۰.۲۸۷	۰.۸۸۳	۰.۲۹۱	۰.۹۱۰	۴	۶	۰.۵	۱۷
۰.۲۴۳	۰.۸۸۲	۰.۲۴۶	۰.۹۰۱	۰.۲۳۸	۰.۸۸۲	۰.۲۴۱	۰.۹۰۵	۶	۶	۰.۵	۱۸
۰.۱۶۵	۰.۸۸۷	۰.۱۹۷	۰.۸۹۷	۰.۱۶۹	۰.۹۱۴	۰.۱۹۲	۰.۸۷۳	۲	۲	۰.۹۵	۱۹
۰.۱۲۳	۰.۸۸۶	۰.۱۳۴	۰.۹۰۰	۰.۱۳۴	۰.۸۹۱	۰.۱۴۵	۰.۸۴۹	۴	۲	۰.۹۵	۲۰
۰.۱۱۳	۰.۸۸۵	۰.۱۲۲	۰.۹۰۱	۰.۱۲۶	۰.۸۷۰	۰.۱۳۲	۰.۸۳۹	۶	۲	۰.۹۵	۲۱
۰.۱۲۴	۰.۸۹۱	۰.۱۵۱	۰.۸۹۸	۰.۱۰۹	۰.۹۲۹	۰.۱۲۸	۰.۹۲۵	۲	۴	۰.۹۵	۲۲
۰.۰۷۵	۰.۸۸۸	۰.۰۸۱	۰.۹۰۱	۰.۰۷۶	۰.۸۹۸	۰.۰۸۲	۰.۸۹۴	۴	۴	۰.۹۵	۲۳
۰.۰۶۳	۰.۸۹۴	۰.۰۶۷	۰.۹۰۲	۰.۰۶۶	۰.۸۹۴	۰.۰۶۹	۰.۸۸۸	۶	۴	۰.۹۵	۲۴
۰.۱۱۵	۰.۸۹۷	۰.۱۴۲	۰.۸۹۷	۰.۰۹۶	۰.۹۲۹	۰.۱۱۴	۰.۹۳۴	۲	۶	۰.۹۵	۲۵
۰.۰۶۳	۰.۸۹۳	۰.۰۶۹	۰.۹۰۰	۰.۰۶۲	۰.۹۰۲	۰.۰۶۷	۰.۹۰۲	۴	۶	۰.۹۵	۲۶
۰.۰۵۰	۰.۸۹۵	۰.۰۵۲	۰.۹۰۱	۰.۰۵۰	۰.۸۹۶	۰.۰۵۲	۰.۹۰۰	۶	۶	۰.۹۵	۲۷

راکت که در بخش ۵.۲ از فصل قبل بیان شده است استفاده می کنیم. با در نظر گرفتن شرایط مطرح شده در مثال عددی ارائه شده در بخش ۵.۲ و $RRSS$ های پایین به دست آمده در مثال فوق، فواصل بیزی را مقایسه می کنیم. همچنین توزیع $Gamma(2, 1)$ ، به عنوان توزیع پیشین هردو پارامتر θ_1 و θ_2 به کار برده می شود.

مقادیر مربوط به میانگین طول فواصل بیزی بر طبق جدول ۳.۳ گزارش شده است. که در این جدول نمادهای زیر به کار رفته است.

$BCI1$: فاصله معتبر بیزی براساس پیشین های مزدوج

$BCI2$: فاصله معتبر بیزی براساس پیشین های جفریز

$HPD1$: چگال ترین فاصله بیزی براساس پیشین های مزدوج

$HPD2$: چگال ترین فاصله بیزی براساس پیشین های جفریز

بر طبق نتایج جدول ۳.۳، مشاهده می شود فاصله $HPD2$ بهتر از سایر فواصل رفتار می کند.

جدول ۳.۳: مقادیر EL براساس داده های واقعی

HPD۲	BCI۲	HPD۱	BCI۱	α
۰.۰۳۵	۰.۰۴۷	۰.۰۴۹	۰.۰۴۹	۰.۰۵
۰.۰۲۸	۰.۰۳۷	۰.۰۳۸	۰.۰۳۸	۰.۱

خلاصه فصل

در این فصل برآوردهای بیزی در حالت‌های مختلف را برای پارامتر تنش-مقاومت، براساس $RRSS$ ‌های پایین از توزیع نمایی تعمیم یافته به‌دست آوردیم. سپس با کمک شبیه‌سازی نمودار مربوط به $MSE(\hat{\mathcal{R}}_B, \mathcal{R})$ در مقابل $\mathcal{R}$ رسم گردید. نکته جالب در خصوص برآوردهای بیزی این بود که این برآوردها با پیشین جفریز، رفتاری مشابه با برآوردهای کلاسیک داشت در حالی که با پیشین مناسب رفتاری متفاوت داشت به‌طوری که بیشترین مقدار میانگین توان دوم خطا، در حوالی نقاط میانی $\mathcal{R}$ رخ نداد. برای حالت سوم، فواصل معتبر بیزی و HPD را با در نظر گرفتن پیشین‌های مزدوج و همچنین ناآگاهی بخش، به‌دست آوردیم. سپس فواصل فوق را به کمک شبیه‌سازی و برطبق معیارهای میانگین طول و احتمال پوشش، مقایسه نمودیم. مشاهده شد فواصل HPD مناسب‌تر از سایر فواصل بیزی می‌باشند. در نهایت فواصل به‌دست آمده از طریق داده‌های واقعی نیز مورد بررسی قرار گرفتند که نتایجی کاملاً منطبق بر شبیه‌سازی داشتند.

فصل ۴

مقایسه R براساس دو طرح نمونه‌گیری در خانواده نرخ خطر معکوس متناسب

۱.۴ مقدمه

در این فصل استنباط‌های درست‌نمایی و بیزی قابلیت اعتماد تنش-مقاومت را در خانواده توزیع‌های با نرخ خطر معکوس متناسب ($PRHR$)، براساس طرح $RRSS$ مورد مطالعه قرار می‌دهیم و برآوردهای حاصل را با برآوردهای متناظر خود که براساس رکوردهای معمولی به‌دست آمده‌اند مقایسه می‌کنیم. مدل $PRHR$ معادل مدل $G(t) = [F_0(t)]^\theta$ می‌باشد که $F_0(t)$ تابع توزیع پایه است که مستقل از هر پارامتر نامعلومی می‌باشد. بدلیل توجه فراوانی که به قابلیت اعتماد تنش-مقاومت در علوم مختلف از جمله در صنعت، داروسازی و مهندسی می‌باشد، به دنبال روش‌هایی مناسب برای انتخاب مشاهدات می‌باشیم به‌طوری‌که در حداقل زمان و با کمترین هزینه به برآوردهایی با قابلیت اعتماد بالا دست یابیم. با توجه به‌اینکه خانواده $PRHR$ بسیاری از مدل‌های طول عمر را شامل می‌شود، بنابراین در این فصل قصد داریم استنباط‌های از R را براساس طرح $RRSS$ به‌دست آوریم و نتایج به‌دست آمده را با نتایج حاصل از رکوردهای معمولی مقایسه کنیم. در بخش ۲.۴ دو طرح موردنظر برای تولید داده‌های رکوردی را معرفی می‌کنیم. در ادامه و در بخش ۳.۴ برآوردهای نقطه‌ای کلاسیک و برآوردهای بیزی را به‌دست می‌آوریم و

برآوردگرهای فوق را از طریق معیار توان دوم خطا و اریبی بررسی می‌کنیم. سپس با در نظر گرفتن اندازه نمونه مختلف برای جوامع X و Y کارایی برآوردگرها را نسبت به برآوردگرهای متناظر حاصل از رکوردهای معمولی مقایسه می‌کنیم. در بخش ۴.۴ فواصل اطمینان مجانبی و دقیق بر پایه کمیت محوری و چند فاصله اطمینان بوت‌استری می‌محاسبه می‌گردد و در بخش ۵.۴ به کمک شبیه‌سازی فواصل به‌دست آمده مقایسه می‌شوند و براساس مقاله بکلیزی (۲۰۰۸)، مقایسه‌ای بین فواصل اطمینان به‌دست آمده از دو طرح RSS و رکوردهای معمولی انجام می‌دهیم. در نهایت در بخش ۶.۴ با یک مثال واقعی نتایج حاصل را بررسی می‌کنیم.

۲.۴ معرفی دو طرح در داده‌های رکوردی

با توجه به اینکه هدف اصلی در این بخش استنباط و مقایسه‌ی برآوردگرهای پارامتر تنش-مقاومت در مدل $PRHR$ است، لذا دو طرح زیر را در مبحث داده‌های رکوردی در نظر بگیرید:

طرح (R) : X و Y متغیرهای تصادفی مستقل از خانواده توزیع‌های $PRHR$ ، به ترتیب با پارامترهای α و β هستند. فرض کنید آزمایش‌ها به صورت دنباله‌ای انجام می‌شوند و داده‌های رکوردی تنها مشاهدات در دسترس برای انجام تحلیل جامعه‌ی آماری می‌باشند. نمونه‌گیری از جامعه‌ی X تا زمان مشاهده شدن رکورد n ام ادامه می‌یابد و به محض مشاهده آن آزمایش متوقف می‌شود. به‌طوری‌که U_i نشان‌دهنده‌ی i امین رکورد از جامعه X است. در نتیجه بردار $U^{(R)} = (U_1, \dots, U_n)^\top$ را در اختیار داریم. همچنین نمونه‌گیری از جامعه Y تا زمان مشاهده رکورد m ام ادامه می‌یابد. در نتیجه، اگر S_j بیان‌گر j امین رکورد مشاهده‌شده باشد، آن‌گاه $S^{(R)} = (S_1, \dots, S_m)^\top$ را به عنوان بردار تصادفی از مقادیر رکوردی جامعه Y ، در اختیار داریم. که نماد (R) نشان‌دهنده‌ی روش رکوردی در جمع‌آوری مشاهدات است.

طرح $(RRSS)$: n دنباله مستقل از متغیرهای تصادفی را در نظر بگیرید و آزمایش در دنباله‌ی i ام تا زمان مشاهده‌ی i امین رکورد ادامه می‌یابد. اگر $U_{i,i}$ نشان‌دهنده‌ی i امین رکورد از i امین دنباله باشد، آن‌گاه $U^{(RRSS)} = (U_{1,1}, \dots, U_{n,n})^\top$ بردار مشاهدات از جامعه‌ی X می‌باشد. به همین ترتیب اگر m دنباله مستقل را در نظر بگیرید که در هر کدام از آن‌ها آزمایش‌ها به صورت دنباله‌ای انجام شود و به محض مشاهده j امین رکورد از j امین دنباله آزمایش متوقف گردد، در این صورت $S^{(RRSS)} = (S_{1,1}, \dots, S_{m,m})^\top$ بردار مشاهدات در دسترس از جامعه‌ی Y می‌باشد.

۳.۴ برآوردگرهای نقطه‌ای پارامتر R در خانواده $PRHR$

فرض کنید متغیر تصادفی مطلقاً پیوسته T به ترتیب دارای تابع توزیع و تابع چگالی به صورت زیر است.

$$G(t) = [F_\circ(t)]^\theta \quad \text{و} \quad g(t) = \theta f_\circ(t) [F_\circ(t)]^{\theta-1},$$

که $F_{\circ}(t)$ تابع توزیع پایه و مستقل از پارامتر شکل θ است و $\theta > 0$. این خانواده از توزیع‌ها به مدل $PRHR$ معروف هستند. فرض کنید X و Y دو متغیر تصادفی مستقل از مدل $PRHR$ به ترتیب با پارامترهای α و β باشند. به راحتی نتیجه می‌شود در این حالت $\mathcal{R} = Pr(X < Y) = \frac{\beta}{\alpha + \beta}$. فرض کنید $\mathbf{u} = (u_{1,1}, u_{2,2}, \dots, u_{n,n})^T$ بردار مشاهدات تصادفی $\mathbf{U} = (U_{1,1}, U_{2,2}, \dots, U_{n,n})^T$ یک $RRSS$ پایین به اندازه n از مدل $PRHR$ با تابع توزیع تجمعی $G(u) = (F_{\circ}(u))^{\alpha}$ باشد و $\mathbf{s} = (s_{1,1}, s_{2,2}, \dots, s_{m,m})^T$ نیز بردار مشاهدات تصادفی $\mathbf{S} = (S_{1,1}, S_{2,2}, \dots, S_{m,m})^T$ یک $RRSS$ پایین به اندازه m از مدل $PRHR$ با تابع توزیع تجمعی $G(s) = (F_{\circ}(s))^{\beta}$ باشد. در این بخش قصد داریم استنباط‌هایی از پارمتر تنش-مقاومت $\mathcal{R}$ ارائه دهیم.

۱.۳.۴ برآوردهای درست‌نمایی ماکسیم

برای برآورد $\mathcal{R}$ از طریق روش درست‌نمایی ماکسیم، لازم است MLE پارمترهای α و β را به دست آوریم. با استفاده از رابطه (۶.۱)، تابع درست‌نمایی به صورت زیر محاسبه می‌گردد.

$$L(\alpha, \beta) = \prod_{i=1}^n \left\{ \frac{\{-\log(F(u_{i,i})^{\alpha})\}^{i-1}}{(i-1)!} \alpha f(u_{i,i}) [F(u_{i,i})^{\alpha-1}] \right\} \times \prod_{j=1}^m \left\{ \frac{\{-\log(F(s_{j,j})^{\beta})\}^{j-1}}{(j-1)!} \beta f(s_{j,j}) [F(s_{j,j})^{\beta-1}] \right\}, \quad (1.4)$$

و لگاریتم تابع درست‌نمایی به صورت

$$\begin{aligned} \ell(\alpha, \beta) = & \sum_{i=1}^n (i-1) \log(-\alpha \log(F(u_{i,i}))) - \log\left(\prod_{i=1}^n (i-1)!\right) + n \log \alpha + \sum_{i=1}^n \log(f(u_{i,i})) \\ & + (\alpha-1) \sum_{i=1}^n \log(F(u_{i,i})) + \sum_{j=1}^m (j-1) \log(-\beta \log(F(s_{j,j}))) - \log\left(\prod_{j=1}^m (j-1)!\right) \\ & + m \log \beta + \sum_{j=1}^m \log(f(s_{j,j})) + (\beta-1) \sum_{j=1}^m \log(F(s_{j,j})). \end{aligned}$$

می‌باشد، بنابراین MLE پارمترهای α و β به صورت زیر نتیجه می‌شوند:

$$\hat{\alpha}_{ML} = -\frac{N}{\sum_{i=1}^n \log(F(u_{i,i}))}, \quad (2.4)$$

و

$$\hat{\beta}_{ML} = -\frac{M}{\sum_{j=1}^m \log(F(s_{j,j}))}, \quad (3.4)$$

که در آن‌ها $N = \frac{n(n+1)}{2}$ و $M = \frac{m(m+1)}{2}$. در نتیجه با استفاده از ویژگی پایایی برآوردگرهای درست‌نمایی ماکسیمم

$$\hat{\mathcal{R}}_{ML} = \frac{\hat{\beta}_{ML}}{\hat{\alpha}_{ML} + \hat{\beta}_{ML}} = \left(1 + \frac{NH'_m}{MT'_n}\right)^{-1}, \quad (4.4)$$

به‌طوری‌که $T'_n = -\sum_{i=1}^n \log(F(U_{i,i}))$ و $H'_m = -\sum_{j=1}^m \log(F(S_{j,j}))$. همان‌طور که قبلاً بحث شده است $U_{i,i}$ ها و $S_{j,j}$ ها متغیرهای تصادفی مستقل از هم می‌باشند. به علاوه می‌توان نشان داد که $T_i = -\log(F(U_{i,i}))$ که $i = 1, \dots, n$ دارای توزیع گاما است؛ به عبارتی $T_i \sim \text{Gamma}(i, \alpha)$. بنابراین

$$T'_n = -\sum_{i=1}^n \log(F(U_{i,i})) \sim \text{Gamma}(N, \alpha), \quad (5.4)$$

و به طور مشابه

$$H'_m = -\sum_{j=1}^m \log(F(S_{j,j})) \sim \text{Gamma}(M, \beta). \quad (6.4)$$

با استفاده از رابطه‌های (5.4) و (6.4) به‌وضوح نتیجه می‌شود

$$\hat{\mathcal{R}}_{ML} \stackrel{d}{=} \left(1 + \frac{\alpha}{\beta} W\right)^{-1}, \quad (7.4)$$

W دارای توزیع فیشر به‌ترتیب با $2M$ و $2N$ درجه آزادی است. برطبق توزیع W و با بکارگیری روش تبدیل، تابع چگالی $\hat{\mathcal{R}}_{ML}$ به صورت زیر به‌دست می‌آید.

$$f_{\hat{\mathcal{R}}_{ML}}(r) = \frac{\Gamma(M+N)}{\Gamma(M)\Gamma(N)} \left(\frac{N\alpha}{M\beta}\right)^N r^{N-1} (1-r)^{M-1} \left(1 - r \left(1 - \frac{N\alpha}{M\beta}\right)\right)^{-(M+N)}.$$

حال برای محاسبه $Var(\hat{\mathcal{R}}_{ML})$ ، $E(\hat{\mathcal{R}}_{ML})$ و $E(\hat{\mathcal{R}}_{ML}^2)$ را به‌صورت زیر محاسبه می‌کنیم (بیلی، ۱۹۳۵):

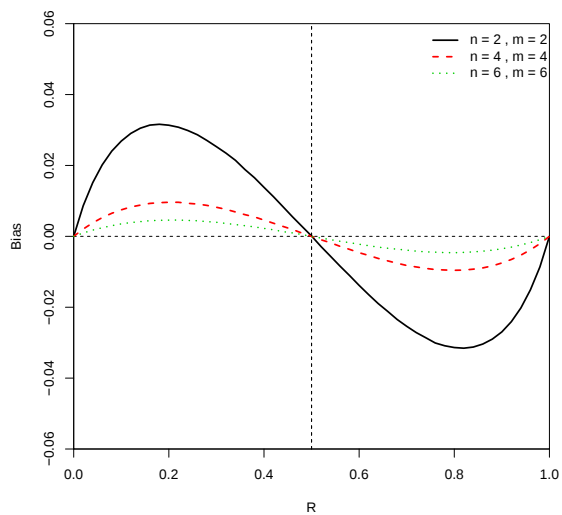
$$E(\hat{\mathcal{R}}_{ML}) = \left(\frac{N\alpha}{M\beta}\right)^N \left(\frac{N}{N+M}\right) {}_2F_1\left(M+N, N+1, M+N+1, 1 - \frac{N\alpha}{M\beta}\right), \quad (8.4)$$

و

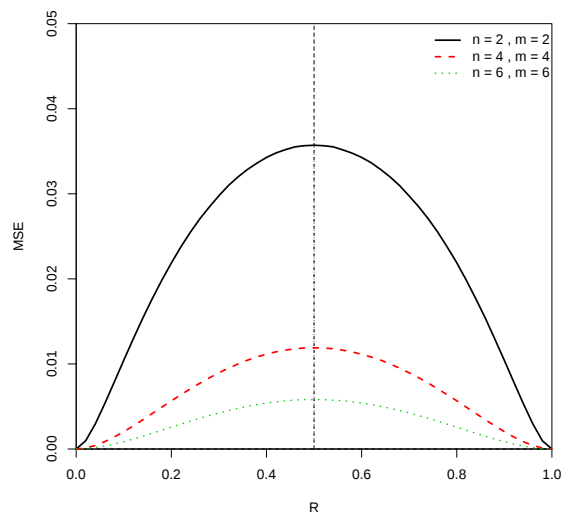
$$E(\hat{\mathcal{R}}_{ML}^2) = \frac{N(N+1)}{(M+N+1)(M+N)} \left(\frac{N\alpha}{M\beta}\right)^N {}_2F_1\left(M+N, N+2, M+N+2, 1 - \frac{N\alpha}{M\beta}\right), \quad (9.4)$$

با جایگذاری روابط (8.4) و (9.4) در $Var(\hat{\mathcal{R}}_{ML})$ ، می‌توان $MSE(\hat{\mathcal{R}}_{ML})$ را به‌دست آورد. در شکل ۱.۴، نمودار مقادیر MSE و اریبی $\hat{\mathcal{R}}_{ML}$ برای ترکیب‌های مختلف n و m در مقابل $\mathcal{R}$ رسم شده است.

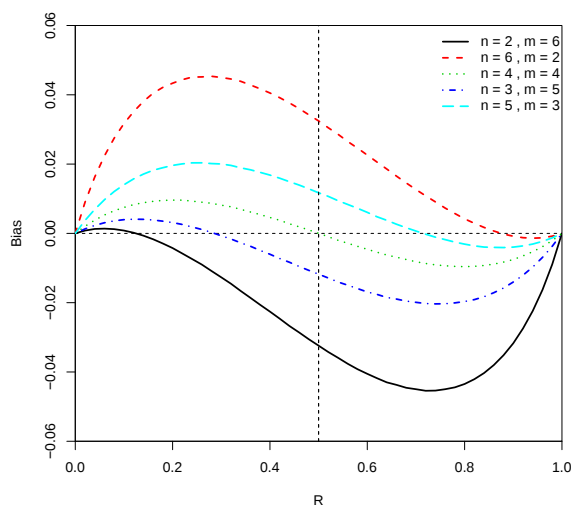
برطبق شکل ۱.۴، نکات زیر مشاهده می‌شود:



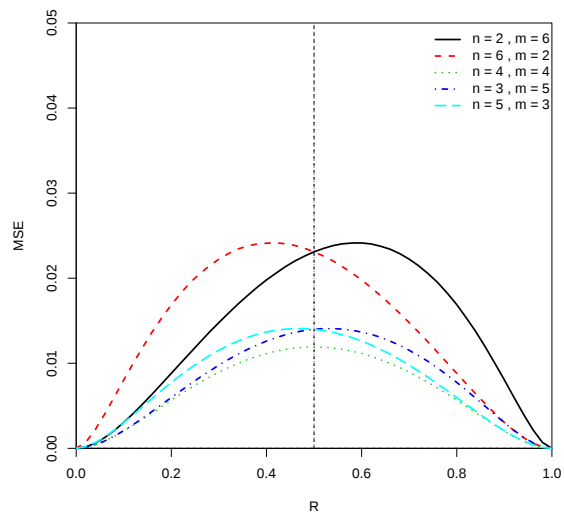
(ب)



(الف)



(ت)



(پ)

شکل ۱.۴: نمودار $MSE(\hat{\mathcal{R}}_{ML}, \mathcal{R})$ و $Bias(\hat{\mathcal{R}}_{ML}, \mathcal{R})$ در مقابل $\mathcal{R}$

• زمانی که $n = m$ ، نمودارهای $MSE(\hat{R}_{ML}, R)$ و $Bias(\hat{R}_{ML}, R)$ به ترتیب حول $R = 0.5$ و در نقطه $(0.5, 0)$ متقارن هستند. همچنین اریبی برای مقادیر $R = 0, 0.5, 1$ برابر صفر است. (شکل‌های (الف) و (ب) را ببینید).

• بر طبق نمودار قسمت (پ) زمانی که $n \neq m$ ، منحنی $MSE(\hat{R}_{ML}, R)$ نامتقارن می‌باشد. زمانی که $n < m$ ، منحنی چوله به راست و زمانی که $n > m$ چوله به چپ می‌باشد. به علاوه، در حالتی که $n = m$ ، $MSE(\hat{R}_{ML}, R)$ به طور معنی‌داری کم‌تر از حالتی می‌باشد که $n \neq m$.

• زمانی که اندازه نمونه افزایش می‌یابد، $MSE(\hat{R}_{ML}, R)$ کاهش می‌یابد و مشاهده می‌کنیم که $R_{max} \propto \frac{n}{n+m}$ ، که در آن R_{max} ماکسیمم کننده $MSE(\hat{R}_{ML}, R)$ می‌باشد.

مقایسه‌ی برآوردگر ML حاصل از دو طرح نمونه‌گیری

با توجه به رابطه (۸.۴) و (۹.۴)، توزیع $MSE(\hat{R}_{ML}, R)$ به توزیع پایه F_0 بستگی ندارد. لذا با در نظر گرفتن توزیع نمایی تعمیم‌یافته به ترتیب با تابع چگالی و تابع توزیع

$$f_0(t; \theta) = \theta e^{-t}(1 - e^{-t})^{\theta-1}, \quad (10.4)$$

و

$$F_0(t; \theta) = (1 - e^{-t})^\theta \quad t > 0, \quad (11.4)$$

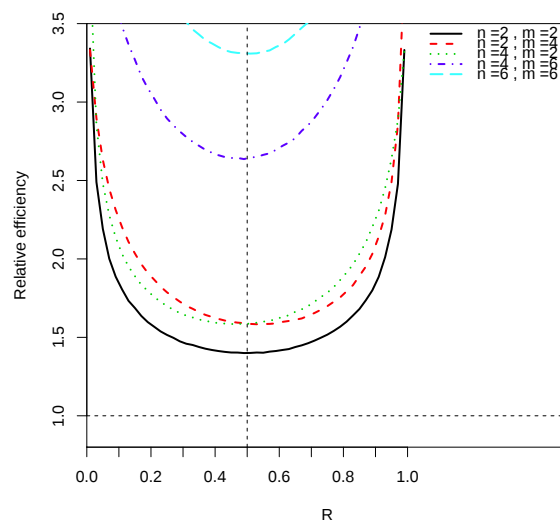
به عنوان توزیعی از خانواده $PRHR$ ، برآوردگرهای ML به دست آمده از دو طرح معرفی شده در بخش ۲.۴ را با استفاده از معیار کارایی نسبی مقایسه می‌کنیم.

$$e(\hat{R}_{ML}^{(RRSS)}, \hat{R}_{ML}^{(R)}) = \frac{MSE(\hat{R}_{ML}^{(R)}, R)}{MSE(\hat{R}_{ML}^{(RRSS)}, R)}. \quad (12.4)$$

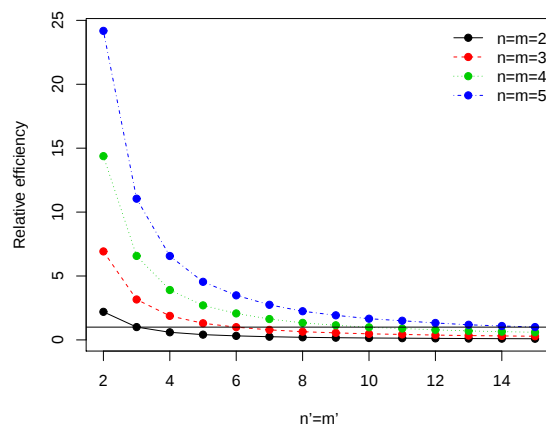
نمودار کارایی نسبی برآوردگرهای درست‌نمایی ماکسیمم با در نظر گرفتن نمونه‌هایی با اندازه برابر برای دو طرح، با توجه به رابطه (۱۲.۴) و برای ترکیب‌های مختلف (n, m) در شکل ۲.۴ نمایش داده شده است. با توجه به شکل ۲.۴، ملاحظه می‌شود، برای ترکیب‌های مختلف (n, m) ، کارایی نسبی برآوردگر ML به دست آمده از طرح $RRSS$ بهتر از برآورگر متناظر به دست آمده از رکوردهای معمولی است. با توجه به نمودار فوق، مشاهده می‌شود با افزایش حجم نمونه کارایی $\hat{R}_{ML}^{(RRSS)}$ به طور معنی‌داری بهتر از $\hat{R}_{ML}^{(R)}$ است.

همچنین برای تمرکز بیشتر در مورد مقایسه این برآوردگرها، نمودار کارایی نسبی (۱۲.۴) را با در نظر گرفتن ۱۵، $n' = m' = 2, \dots, 5$ برای $n = m = 2, \dots, 5$ داده شده و به ازای $R = 0.5, 0.95$ در شکل ۳.۴ رسم نمودیم.

با توجه به شکل ۳.۴، به آسانی مشاهده می‌شود، زمانی که $n' = m' \leq N = M$ ، کارایی نسبی بزرگتر از یک است. به طور واضح، زمانی که $n' \leq N$ و $m' \leq M$ برآوردگر درست‌نمایی ماکسیمم



شکل ۲.۴: نمودار $e(\hat{R}_{ML}^{(R)}; \hat{R}_{ML}^{(RRSS)})$ در مقابل R



شکل ۳.۴: نمودار کارایی نسبی $\hat{R}_{ML}$ براساس رابطه (۱۲.۴)

براساس طرح $RRSS$ بهتر از رکوردهای معمولی است. به عنوان مثال، به ازای $n = m = 4$ (یا به طور معادل $N = M = 10$)، برآوردگر $\hat{R}_{ML}^{(RRSS)}$ کاراتر از برآوردگر $\hat{R}_{ML}^{(R)}$ به ازای $n' = m' = 8$ می باشد.

۲.۳.۴ برآودگر $UMVU$

در این بخش با در نظر گرفتن مدل $PRHR$ برآودگر $UMVU$ پارامتر R را براساس طرح $RRSS$ به‌دست می‌آوریم. با مراجعه به رابطه (۷.۴)، مشاهده می‌کنیم توزیع $\hat{R}_{ML}$ به توزیع پایه بستگی ندارد، بنابراین با توجه به قضیه ۳.۲.۲ از فصل دوم نتیجه می‌شود که با در نظر گرفتن مدل $PRHR$ نتایجی مشابه برای $\hat{R}_{UMVU}$ به‌صورت زیر به‌دست می‌آید.

$$\hat{R}_{UMVU} = \begin{cases} \sum_{s=0}^{M-1} \frac{\binom{M-1}{s}}{\binom{N+s-1}{s}} \left(-\frac{1 - \hat{R}_{ML} M}{\hat{R}_{ML} N} \right)^s, & \hat{R}_{ML} < \frac{M}{M+N}, \\ 1 - \sum_{s=0}^{N-1} \frac{\binom{N-1}{s}}{\binom{M+s-1}{s}} \left(-\frac{\hat{R}_{ML} N}{1 - \hat{R}_{ML} M} \right)^s, & \hat{R}_{ML} \geq \frac{M}{M+N}. \end{cases} \quad (۱۳.۴)$$

مقایسه‌ی برآودگر $UMVU$ حاصل از دو طرح نمونه‌گیری

برآودگر $UMVU$ پارامتر R به صورت تابعی از $\hat{R}_{ML}$ در رابطه (۱۳.۴) بیان شده‌است. لذا با توجه به رابطه (۷.۴)، $Var(\hat{R}_{UMVU})$ به توزیع پایه X و Y بستگی ندارد. بنابراین برای مقایسه برآودگرهای $UMVU$ به‌دست آمده از دو طرح با توجه به

$$e(\hat{R}_{UMVU}^{(RRSS)}, \hat{R}_{UMVU}^{(R)}) = \frac{Var(\hat{R}_{UMVU}^{(R)})}{Var(\hat{R}_{UMVU}^{(RRSS)})}.$$

و با در نظر گرفتن توزیع نمایی تعمیم‌یافته بر طبق روابط (۱۰.۴) و (۱۱.۴) به‌عنوان توزیعی از خانواده‌ی $PRHR$ ، نمودار کارایی نسبی برآورگرهای فوق را در مقابل R در شکل ۴.۴، رسم نمودیم.

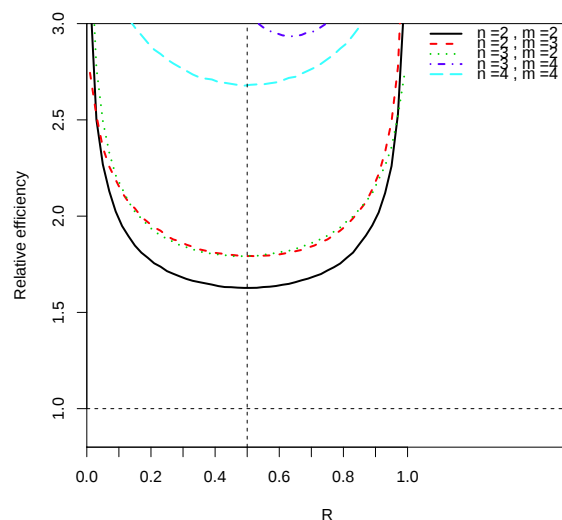
با توجه به شکل ۴.۴، مشاهده می‌شود برای تمام ترکیب‌های مختلف (n, m) در نظر گرفته شده، کارایی نسبی برآودگر $UMVU$ به‌دست آمده از طرح $RRSS$ بهتر از برآودگر متناظر به‌دست آمده از رکوردهای معمولی است.

همانند برآودگر ML ، مقادیر کارایی نسبی براساس رابطه (؟؟)، برای $n = m = 2, \dots, 5$ و $n' = m' = 2, \dots, 15$ در شکل ۵.۴ رسم شده‌است.

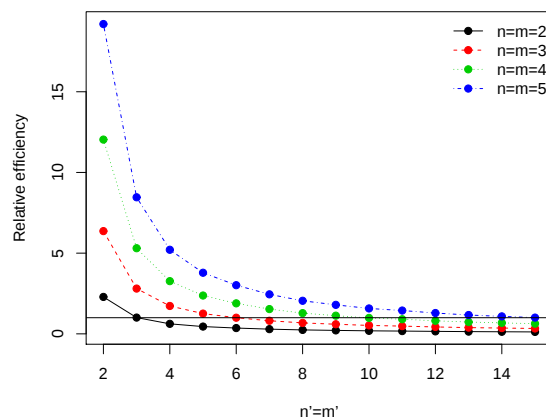
با توجه به شکل ۵.۴، مشاهده می‌شود عملکرد $\hat{R}_{UMVU}^{(R)}$ زمانی که تعداد مشاهدات نمونه افزایش می‌یابد ($n' > N$ و $m' > M$)، بهتر از $\hat{R}_{UMVU}^{(RRSS)}$ است که این موضوع از نظر اقتصادی مقرون به‌صرفه نمی‌باشد. بنابراین برآودگر $\hat{R}_{UMVU}^{(RRSS)}$ برآودگر کاراتری می‌باشد.

۳.۳.۴ برآودگر بیزی

در این بخش، استنباط‌های بیزی پارامتر تنش-مقاومت را با فرض این که α و β دارای توزیع پیشین گاما باشند مورد بررسی قرار می‌دهیم، به‌طوری که $\alpha \sim \text{Gamma}(\lambda_1, \mu_1)$ و $\beta \sim \text{Gamma}(\lambda_2, \mu_2)$.



شکل ۴.۴: نمودار $e(\hat{\mathcal{R}}_{UMVU}^{(R)}; \hat{\mathcal{R}}_{UMVU}^{(RRSS)})$ در مقابل $\mathcal{R}$



شکل ۵.۴: کارایی نسبی $\hat{\mathcal{R}}_{UMVU}$ براساس (؟؟)

بنابراین توزیع پسین α و β به صورت زیر نتیجه می‌شود

$$\alpha | \mathbf{r} \sim \text{Gamma} \left(\lambda_1 + N, \mu_1 + t'_n \right), \quad (14.4)$$

و

$$\beta | \mathbf{s} \sim \text{Gamma} \left(\lambda_2 + M, \mu_2 + h'_m \right), \quad (15.4)$$

که h'_m و t'_n مقادیر مشاهده‌شده متغیرهای تصادفی H'_m و T'_n به ترتیب در روابط (۵.۴) و (۶.۴) می‌باشند. با فرض استقلال α و β ، توزیع پسین $\mathcal{R}$ به صورت زیر به‌دست می‌آید

$$f_{\mathcal{R}|\mathbf{R},\mathbf{S}}(r) = \frac{\Gamma(\lambda_1 + \lambda_2 + N + M)}{\Gamma(\lambda_1 + N)\Gamma(\lambda_2 + M)} \left(\frac{\mu_2 + H'_m}{\mu_1 + T'_n} \right)^{\lambda_2 + M} \frac{r^{\lambda_2 + M - 1} (1 - r)^{\lambda_1 + N - 1}}{\left(1 - r \left(1 - \frac{\mu_2 + H'_m}{\mu_1 + T'_n} \right) \right)^{\lambda_1 + \lambda_2 + N + M}}. \quad (۱۶.۴)$$

و همچنین برآوردگر بیزی $\mathcal{R}$ تحت زیان درجه دوم خطا به صورت زیر است. (کسلا و برگر^۱ ۱۹۹۰).

$$\hat{\mathcal{R}}_B = \int_0^1 r f_{\mathcal{R}|\mathbf{R},\mathbf{S}}(r) dr = \frac{\lambda_2 + M}{\lambda_1 + \lambda_2 + N + M} \left(\frac{\mu_2 + H'_m}{\mu_1 + T'_n} \right)^{\lambda_2 + M} \times {}_2F_1 \left(\lambda_1 + \lambda_2 + N + M, \lambda_2 + M + 1, \lambda_1 + \lambda_2 + N + M + 1, \left(1 - \frac{\mu_2 + H'_m}{\mu_1 + T'_n} \right) \right).$$

مقایسه‌ی برآوردگر بیزی حاصل از دو طرح نمونه‌گیری

با استفاده از رابطه (۱۶.۴)، به آسانی می‌توان نشان داد که $MSE(\hat{\mathcal{R}}_B, \mathcal{R})$ بستگی به انتخاب توزیع پایه برای توابع X و Y ندارد. بار دیگر با در نظر گرفتن توزیع نمایی تعمیم یافته برآوردگرهای بیز به‌دست آمده از دو طرح را مقایسه می‌کنیم.

$$e(\hat{\mathcal{R}}_B^{(RRSS)}, \hat{\mathcal{R}}_B^{(R)}) = \frac{MSE(\hat{\mathcal{R}}_B^{(R)}, \mathcal{R})}{MSE(\hat{\mathcal{R}}_B^{(RRSS)}, \mathcal{R})}.$$

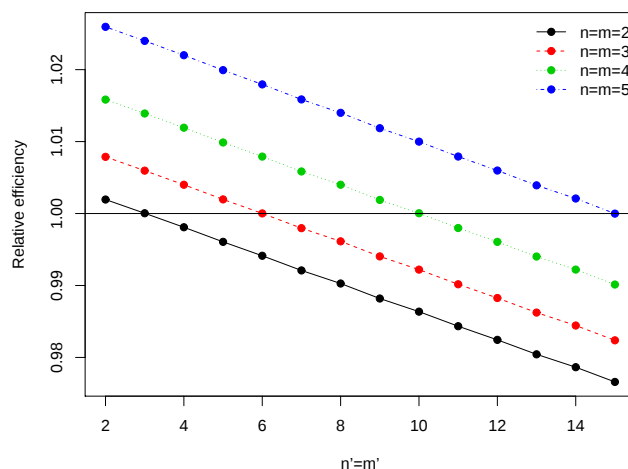
مقادیر کارایی نسبی برآوردگرهای بیز به‌دست آمده براساس دو طرح، در شکل ۶.۴ رسم شده‌است.

با توجه به شکل ۶.۴، مشاهده می‌شود زمانی که $n' = m' \leq N = M$ ، $e(\hat{\mathcal{R}}_B^{(RRSS)}, \hat{\mathcal{R}}_B^{(R)}) > 1$ ، بنابراین، کارایی نسبی برآوردگر $\hat{\mathcal{R}}_B^{(RRSS)}$ بهتر از برآوردگر $\hat{\mathcal{R}}_B^{(R)}$ می‌باشد.

۴.۴ برآوردگرهای فاصله‌ای پارامتر $\mathcal{R}$ در خانواده $PRHP$

در این بخش، چند فاصله اطمینان برای $\mathcal{R}$ ارائه می‌دهیم.

^۱ Casella and Berger



شکل ۶.۴: نمودار کارایی نسبی $\hat{\mathcal{R}}_B$ براساس دو طرح

۱.۴.۴ فاصله اطمینان براساس کمیت محور

با استفاده از رابطه (۷.۴) و با فرض $W = \frac{\beta N H'_m}{\alpha M T'_n}$ فاصله اطمینان $100(1-\gamma)\%$ برای $\mathcal{R}$ به صورت زیر به دست می‌آید:

$$\left[\left(1 + \frac{1 - \hat{\mathcal{R}}_{ML}}{\hat{\mathcal{R}}_{ML}} F_{1-\frac{\gamma}{2}}(\gamma, \gamma) \right)^{-1}, \left(1 + \frac{1 - \hat{\mathcal{R}}_{ML}}{\hat{\mathcal{R}}_{ML}} F_{\frac{\gamma}{2}}(\gamma, \gamma) \right)^{-1} \right],$$

که در آن $F_{\gamma}(\gamma, \gamma)$ ، چندک γ ام توزیع F به ترتیب با $2M$ و $2N$ درجه آزادی است.

۲.۴.۴ فاصله اطمینان مجانبی

در این بخش می‌خواهیم یک فاصله اطمینان مجانبی برای $\mathcal{R}$ براساس $\hat{\mathcal{R}}_{ML}$ به دست آوریم. اگر فرض کنید تعداد مشاهدات به اندازه کافی بزرگ باشد به طوری که $n, m \rightarrow \infty$ ، در این صورت $(\hat{\alpha}_{ML} - \alpha) \xrightarrow{d} N(0, \sigma_{\alpha}^2)$ و $(\hat{\beta}_{ML} - \beta) \xrightarrow{d} N(0, \sigma_{\beta}^2)$ که در آن $\xrightarrow{d}$ نماد همگرایی در توزیع است و σ_{α}^2 و σ_{β}^2 واریانس مجانبی است و به صورت زیر محاسبه می‌گردد.

$$\sigma_{\alpha}^2 = -1/E \left(\frac{\partial^2}{\partial \alpha^2} \log L(\alpha, \mathbf{U}) \right)$$

$$\sigma_{\beta}^2 = -1/E \left(\frac{\partial^2}{\partial \beta^2} \log L(\beta, \mathbf{S}) \right)$$

که L در رابطه (۱.۴) ارائه شده است. به آسانی نشان می‌دهیم که $\sigma_1^2 = \frac{\alpha^2}{N}$ و $\sigma_2^2 = \frac{\beta^2}{M}$. فرض کنید n و m به اندازه کافی بزرگ باشد، به طوری که $p \rightarrow \frac{M}{N}$ ، که در آن $p \in (0, 1)$ بنابراین داریم

$$\sqrt{N} \begin{pmatrix} \hat{\alpha} - \alpha \\ \hat{\beta} - \beta \end{pmatrix} \xrightarrow{d} N_2 \left(\begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} \alpha^2 & 0 \\ 0 & \frac{\beta^2}{p} \end{pmatrix} \right).$$

با فرض $h(t_1, t_2) = \frac{t_2}{t_1 + t_2}$ و با استفاده از روش دلتا (واسرمن^۲ (۲۰۰۶) صفحه ۵-۶ را ببینید). داریم

$$(\hat{\mathcal{R}}_{ML} - \mathcal{R}) \xrightarrow{d} N(0, \sigma^2),$$

که

$$\sigma = \frac{\alpha\beta}{(\alpha + \beta)^2} \sqrt{\frac{1}{N} \left(1 + \frac{1}{p}\right)}, \quad (17.4)$$

براساس این نتایج مجانبی و با به کار بردن قضیه اسلاتسکی^۳، فاصله اطمینان مجانبی $100(1 - \gamma)\%$ برای $\mathcal{R}$ به صورت زیر نتیجه می‌شود

$$(\hat{\mathcal{R}}_{ML} - Z_{1-\frac{\gamma}{2}} \hat{\sigma}, \hat{\mathcal{R}}_{ML} + Z_{1-\frac{\gamma}{2}} \hat{\sigma}).$$

که Z_γ ، چندک γ توزیع نرمال استاندارد است و $\hat{\sigma}$ با جایگذاری $\hat{\alpha}_{ML}$ و $\hat{\beta}_{ML}$ به ترتیب به جای α و β در (۱۷.۴) به دست می‌آید.

۳.۴.۴ فواصل معتبر بیزی

با استفاده از این واقعیت که $(\mathbf{r}, \mathbf{s}) \sim F(\gamma(\lambda_1 + N), \gamma(\lambda_2 + M))$ فاصله معتبر^۴ بیزی $100(1 - \gamma)\%$ برای $\mathcal{R}$ به صورت زیر به دست می‌آید.

$$\left[(1 + AF_{1-\frac{\gamma}{2}})^{-1}, (1 + AF_{\frac{\gamma}{2}})^{-1} \right]$$

$$.A = \left(\frac{\mu_2 + H'_m}{\mu_1 + T'_n} \right) \left(\frac{\lambda_1 + N}{\lambda_2 + M} \right) \text{ که}$$

از طریق روش زنجیر مارکوف مونت کارلو^۵ (MCMC) که توسط چن و شائو^۶ (۱۹۹۹) ارائه شده است و با استفاده از لم ۱.۳.۳، فاصله HPD برای $\mathcal{R}$ را به دست می‌آوریم.

^۲Wasserman

^۳Slutsky

^۴Credible interval

^۵Monte Carlo Markov Chain

^۶Chen and Shao

۴.۴.۴ فواصل اطمینان بوت‌استرپ

در این بخش چند فاصله بوت‌استرپی را ارائه می‌دهیم. لازم است بدانید تمام فرایندهای این بخش بستگی به T'_n و H'_m می‌باشند که به ترتیب در روابط (۵.۴) و (۶.۴) داده شده‌است. بنابراین در این بخش از روش بوت‌استرپ پارامتری استفاده می‌نماییم و فواصل بوت‌استرپ صدکی و t -بوت‌استرپ و بوت‌استرپ پایه را به دست می‌آوریم. فرآیند بوت‌استرپ تحت گام‌های زیر توضیح داده شده‌است.

گام ۱. $\hat{\alpha}$ ، $\hat{\beta}$ و $\hat{R}_{ML}$ را به ترتیب با استفاده از روابط (۵.۴)، (۶.۴) و (۷.۴)، با در نظر گرفتن نمونه‌های مشاهده‌شده مستقل u_i ، s_j به ترتیب از توزیع‌های $Gamma(i, \alpha)$ ، $i = 1, \dots, n$ و $Gamma(j, \beta)$ ، $j = 1, \dots, m$ محاسبه می‌کنیم.

گام ۲. $u_i^* \sim Gamma(i, \hat{\alpha})$ و $s_j^* \sim Gamma(j, \hat{\beta})$ را به منظور محاسبه $\hat{\alpha}^*$ ، $\hat{\beta}^*$ و همچنین $\hat{R}_{ML}^*$ تولید می‌کنیم.

گام ۳. مرحله ۲ را B بار ($b = 1, \dots, B$) برای محاسبه $\hat{R}_{(B)ML}^*$ تکرار می‌کنیم.

t -بوت‌استرپ براساس انحراف معیار مجانبی σ

گام ۲ را $b = 1, \dots, B$ بار برای به دست آوردن $\hat{\sigma}^*$ در رابطه (۱۷.۴) به ترتیب با جایگذاری $\hat{\alpha}^*$ و $\hat{\beta}^*$ به جای $\hat{\alpha}$ و $\hat{\beta}$ فرض کنید $Z^* = (Z_{(1)}^*, \dots, Z_{(B)}^*)^\top$ که

$$z_{(b)}^* = \frac{\hat{R}_{(b)ML}^* - \hat{R}_{ML}}{\hat{\sigma}_{(b)}^*}, \quad b = 1, \dots, B$$

در این صورت بازه اطمینان t -بوت‌استرپ $\% (1 - \gamma)(100)$ برای $\mathcal{R}$ به صورت زیر است:

$$\left(\hat{R}_{ML} - z_{1-\frac{\gamma}{2}}^* \hat{\sigma}, \hat{R}_{ML} - z_{\frac{\gamma}{2}}^* \hat{\sigma} \right).$$

که Z_γ^* چندک γ ام Z^* می‌باشد.

برای محاسبه فواصل اطمینان بوت‌استرپ صدکی و پایه و همچنین فاصله اطمینان t -بوت‌استرپ براساس برآورد بوت‌استرپی واریانس به بخش ۲.۳.۲ از فصل دوم مراجعه نمایید.

۵.۴ مقایسه‌ی برآوردگرها به کمک شبیه‌سازی

در این بخش عملکرد فواصل اطمینان به دست آمده از طریق شبیه‌سازی مقایسه می‌گردد. بدین منظور تمام ترکیب‌های مختلف $n = 3, 5, 7$ و $m = 3, 5, 7$ با فرض $\beta = 0.3$ و $\alpha = 0.1, 0.3, 0.5, 0.7, 0.9$ در $\mathcal{R}$ در نظر گرفته شده‌است. برای هر ترکیب ۱۰۰۰۰ نمونه شبیه‌سازی شده‌است. برای برآورد واریانس بوت‌استرپی، $B = 1000$ و $B' = 25$ در نظر گرفته شده‌است. همچنین توزیع پیشین $Gamma(2, 4)$

برای α و β انتخاب شده‌است. فواصل اطمینان زیر را به‌دست‌آورده‌ایم که نتایج شبیه‌سازی در جدول‌های ۱.۴ تا ۴.۴ ارائه شده‌است.

$Perc$: فاصله اطمینان بوت‌استرپ صدکی

$Boot - t1$: فاصله اطمینان t -بوت‌استرپ براساس انحراف معیار مجانبی σ

$Boot - t2$: فاصله اطمینان t -بوت‌استرپ براساس برآورد واریانس بوت‌استرپی

$Basic$: فاصله اطمینان بوت‌استرپ پایه

MLE : فاصله اطمینان براساس کمیت‌محور

$AMLE$: فاصله اطمینان مجانبی

$Bayes$: فاصله اطمینان بیزی

HPD : فاصله تقریبی

بر طبق نتایج به‌دست آمده از جدول‌های ۱.۴ تا ۴.۴، نکات زیر مشاهده می‌شود

- همان‌طور که انتظار داشتیم سطح پوشش و میانگین طول فاصله اطمینان با افزایش اندازه نمونه کاهش می‌یابد.

- واضح است که ماکسیم مقدار میانگین طول فاصله اطمینان در $R = 0.5$ می‌باشد و برای مقادیر فرین R کاهش می‌یابد که نتایج شبیه‌سازی با نتایج به‌دست آمده در بخش ۳.۴ مطابقت دارد.

- فاصله اطمینان $Boot - t1$ بخصوص برای اندازه نمونه‌های بزرگ بهتر از فاصله اطمینان $Boot - t2$ عمل می‌کند.

- سطح پوشش فاصله اطمینان به‌دست آمده براساس کمیت‌محور و فاصله اطمینان بوت‌استرپ صدکی بسیار نزدیک به سطح اطمینان γ می‌باشد. با این وجود، براساس معیار EL ، فاصله اطمینان مجانبی بهتر از سایر فواصل عمل می‌کند. در نهایت با در نظر گرفتن دو معیار یعنی سطح پوشش و میانگین طول، مشاهده می‌شود فاصله اطمینان بوت‌استرپ صدکی و فاصله اطمینان دقیق براساس کمیت‌محور بهتر می‌باشند.

- همان‌طور که انتظار داشتیم فاصله HPD بهتر از فاصله معتبر بیزی رفتار می‌کند.

- با در نظر گرفتن سطح اطمینان ۰.۰۵ و ۰.۱، مشاهده می‌شود زمانی که $\gamma = 0.1$ سطح پوشش و میانگین طول فاصله اطمینان کمتر از زمانی می‌باشد که $\gamma = 0.05$.

۱.۵.۴ مقایسه برآوردگرهای فاصله‌ای حاصل از دو طرح

صالحی و همکاران (۲۰۱۶)، براساس روش نمونه‌گیری معکوس رکوردها، رکوردهای معمولی و $RRSS$ را زمانی که توزیع جامعه‌ی آماری از الگوی نرخ خطر متناسب پیروی می‌کرد مقایسه نمودند. آن‌ها چند برآوردگر پارامتر مقیاس توزیع نمایی یک پارامتری را براساس معیار میانگین

جدول ۱.۴: CP و EL برای R به‌ازای $\gamma = 5\%$.

ردیف	R	m	n	<i>perc</i>		<i>Boot - t\</i>		<i>Boot - t\</i>		<i>Basic</i>	
				EL	CP	EL	CP	EL	CP	EL	CP
۱	۰.۱	۳	۳	۰.۹۴۷	۰.۲۴۴	۰.۸۴۳	۰.۲۰۱	۰.۸۸۱	۰.۲۱۸	۰.۸۲۶	۰.۲۴۳
۲	۰.۱	۵	۳	۰.۹۳۷	۰.۱۷۶	۰.۸۲۳	۰.۱۶۷	۰.۸۴۶	۰.۱۷۳	۰.۸۴۰	۰.۱۷۶
۳	۰.۱	۷	۳	۰.۹۲۲	۰.۱۵۶	۰.۸۱۱	۰.۱۵۶	۰.۸۲۹	۰.۱۶۱	۰.۸۴۱	۰.۱۵۵
۴	۰.۱	۳	۵	۰.۹۳۹	۰.۲۱۷	۰.۹۰۱	۰.۱۹۴	۰.۹۳۰	۰.۲۱۰	۰.۸۶۶	۰.۲۱۶
۵	۰.۱	۵	۵	۰.۹۴۹	۰.۱۴۰	۰.۸۹۲	۰.۱۴۷	۰.۹۰۸	۰.۱۵۵	۰.۸۸۶	۰.۱۳۹
۶	۰.۱	۷	۵	۰.۹۴۵	۰.۱۱۶	۰.۸۸۵	۰.۱۲۷	۰.۸۹۶	۰.۱۳۴	۰.۸۹۱	۰.۱۱۶
۷	۰.۱	۳	۷	۰.۹۲۹	۰.۲۰۶	۰.۹۱۷	۰.۱۸۶	۰.۹۴۴	۰.۲۰۳	۰.۸۷۷	۰.۲۰۵
۸	۰.۱	۵	۷	۰.۹۴۹	۰.۱۲۴	۰.۹۱۰	۰.۱۲۳	۰.۹۲۸	۰.۱۳۲	۰.۸۹۳	۰.۱۲۳
۹	۰.۱	۷	۷	۰.۹۴۹	۰.۰۹۸	۰.۹۰۷	۰.۱۰۱	۰.۹۲۳	۰.۱۰۷	۰.۹۰۴	۰.۰۹۸
۱۰	۰.۳	۳	۳	۰.۹۴۷	۰.۴۴۹	۰.۸۸۳	۰.۵۳۲	۰.۹۰۷	۰.۵۳۹	۰.۸۲۱	۰.۴۴۷
۱۱	۰.۳	۵	۳	۰.۹۳۷	۰.۳۶۶	۰.۸۶۹	۰.۴۴۵	۰.۸۸۲	۰.۴۵۱	۰.۸۴۵	۰.۳۶۴
۱۲	۰.۳	۷	۳	۰.۹۲۳	۰.۳۳۷	۰.۸۵۴	۰.۴۱۸	۰.۸۶۴	۰.۴۲۴	۰.۸۵۴	۰.۳۳۶
۱۳	۰.۳	۳	۵	۰.۹۳۹	۰.۳۹۹	۰.۹۱۷	۰.۴۶۰	۰.۹۳۱	۰.۴۷۵	۰.۸۶۶	۰.۳۹۸
۱۴	۰.۳	۵	۵	۰.۹۴۹	۰.۲۹۴	۰.۹۱۸	۰.۳۲۹	۰.۹۳۰	۰.۳۴۰	۰.۸۹۴	۰.۲۹۳
۱۵	۰.۳	۷	۵	۰.۹۴۵	۰.۲۵۵	۰.۹۱۱	۰.۲۸۵	۰.۹۲۱	۰.۲۹۵	۰.۹۰۲	۰.۲۵۴
۱۶	۰.۳	۳	۷	۰.۹۲۸	۰.۳۸۰	۰.۹۱۸	۰.۴۲۱	۰.۹۳۰	۰.۴۴۱	۰.۸۷۴	۰.۳۷۸
۱۷	۰.۳	۵	۷	۰.۹۴۹	۰.۲۶۳	۰.۹۲۹	۰.۲۸۰	۰.۹۴۱	۰.۲۹۲	۰.۹۰۲	۰.۲۶۱
۱۸	۰.۳	۷	۷	۰.۹۴۹	۰.۲۱۷	۰.۹۲۷	۰.۲۳۰	۰.۹۳۷	۰.۲۴۰	۰.۹۱۴	۰.۲۱۶
۱۹	۰.۵	۳	۳	۰.۹۴۷	۰.۵۰۱	۰.۸۹۳	۰.۶۵۷	۰.۹۱۰	۰.۶۵۰	۰.۸۱۹	۰.۴۹۹
۲۰	۰.۵	۵	۳	۰.۹۳۷	۰.۴۳۳	۰.۸۹۷	۰.۵۴۲	۰.۹۰۹	۰.۵۴۹	۰.۸۵۰	۰.۴۳۱
۲۱	۰.۵	۷	۳	۰.۹۲۲	۰.۴۰۷	۰.۸۹۱	۰.۵۰۲	۰.۸۹۸	۰.۵۱۴	۰.۸۵۷	۰.۴۰۶
۲۲	۰.۵	۳	۵	۰.۹۳۹	۰.۴۳۳	۰.۹۰۳	۰.۵۴۳	۰.۹۱۵	۰.۵۴۹	۰.۸۵۹	۰.۴۳۲
۲۳	۰.۵	۵	۵	۰.۹۴۹	۰.۳۳۹	۰.۹۲۳	۰.۳۸۵	۰.۹۳۳	۰.۳۹۶	۰.۸۹۲	۰.۳۳۸
۲۴	۰.۵	۷	۵	۰.۹۴۶	۰.۳۰۱	۰.۹۲۹	۰.۳۳۳	۰.۹۳۸	۰.۳۴۴	۰.۹۰۵	۰.۳۰۰
۲۵	۰.۵	۳	۷	۰.۹۲۸	۰.۴۰۸	۰.۸۹۶	۰.۵۰۲	۰.۹۰۴	۰.۵۱۳	۰.۸۶۴	۰.۴۰۶
۲۶	۰.۵	۵	۷	۰.۹۴۹	۰.۳۰۱	۰.۹۲۸	۰.۳۳۲	۰.۹۳۶	۰.۳۴۳	۰.۹۰۶	۰.۳۰۰
۲۷	۰.۵	۷	۷	۰.۹۴۹	۰.۲۵۴	۰.۹۳۴	۰.۲۷۲	۰.۹۴۳	۰.۲۸۲	۰.۹۱۷	۰.۲۵۳
۲۸	۰.۷	۳	۳	۰.۹۴۷	۰.۴۴۷	۰.۸۸۴	۰.۵۳۰	۰.۹۰۷	۰.۵۳۸	۰.۸۲۲	۰.۴۴۵
۲۹	۰.۷	۵	۳	۰.۹۳۷	۰.۳۹۹	۰.۹۰۷	۰.۴۵۹	۰.۹۲۵	۰.۴۷۴	۰.۸۵۷	۰.۳۹۷
۳۰	۰.۷	۷	۳	۰.۹۲۳	۰.۳۸۰	۰.۹۱۲	۰.۴۲۲	۰.۹۲۶	۰.۴۴۲	۰.۸۷۱	۰.۳۷۸
۳۱	۰.۷	۳	۵	۰.۹۳۹	۰.۳۶۵	۰.۸۷۰	۰.۴۴۴	۰.۸۸۵	۰.۴۵۰	۰.۸۵۲	۰.۳۶۴
۳۲	۰.۷	۵	۵	۰.۹۴۹	۰.۲۹۴	۰.۹۱۶	۰.۳۲۸	۰.۹۲۸	۰.۳۳۹	۰.۸۸۹	۰.۲۹۳
۳۳	۰.۷	۷	۵	۰.۹۴۵	۰.۲۶۳	۰.۹۳۱	۰.۲۸۰	۰.۹۴۲	۰.۲۹۳	۰.۹۱۰	۰.۲۶۲
۳۴	۰.۷	۳	۷	۰.۹۲۸	۰.۳۳۸	۰.۸۶۶	۰.۴۱۸	۰.۸۷۵	۰.۴۲۴	۰.۸۶۵	۰.۳۳۶
۳۵	۰.۷	۵	۷	۰.۹۴۹	۰.۲۵۶	۰.۹۱۳	۰.۲۸۴	۰.۹۲۴	۰.۲۹۵	۰.۹۰۲	۰.۲۵۵
۳۶	۰.۷	۷	۷	۰.۹۴۹	۰.۲۱۷	۰.۹۳۰	۰.۲۳۰	۰.۹۴۱	۰.۲۴۰	۰.۹۱۵	۰.۲۱۶
۳۷	۰.۹	۳	۳	۰.۹۴۷	۰.۲۴۳	۰.۸۴۰	۰.۲۰۱	۰.۸۸۰	۰.۲۱۷	۰.۸۲۲	۰.۲۴۱
۳۸	۰.۹	۵	۳	۰.۹۳۷	۰.۲۱۶	۰.۸۹۲	۰.۱۹۴	۰.۹۲۵	۰.۲۱۰	۰.۸۵۸	۰.۲۱۵
۳۹	۰.۹	۷	۳	۰.۹۲۲	۰.۲۰۶	۰.۹۱۵	۰.۱۸۷	۰.۹۴۳	۰.۲۰۴	۰.۸۷۴	۰.۲۰۵
۴۰	۰.۹	۳	۵	۰.۹۳۹	۰.۱۷۶	۰.۸۲۷	۰.۱۶۶	۰.۸۴۸	۰.۱۷۳	۰.۸۳۹	۰.۱۷۵
۴۱	۰.۹	۵	۵	۰.۹۴۹	۰.۱۴۰	۰.۸۸۸	۰.۱۴۷	۰.۹۰۸	۰.۱۵۵	۰.۸۷۸	۰.۱۳۹
۴۲	۰.۹	۷	۵	۰.۹۴۵	۰.۱۲۵	۰.۹۲۰	۰.۱۲۴	۰.۹۳۵	۰.۱۳۲	۰.۹۰۱	۰.۱۲۴
۴۳	۰.۹	۳	۷	۰.۹۲۸	۰.۱۵۶	۰.۸۲۱	۰.۱۵۷	۰.۸۳۹	۰.۱۶۱	۰.۸۵۰	۰.۱۵۶
۴۴	۰.۹	۵	۷	۰.۹۴۹	۰.۱۱۷	۰.۸۸۹	۰.۱۲۷	۰.۹۰۲	۰.۱۳۴	۰.۸۹۶	۰.۱۱۶
۴۵	۰.۹	۷	۷	۰.۹۴۹	۰.۰۹۹	۰.۹۱۳	۰.۱۰۱	۰.۹۲۵	۰.۱۰۷	۰.۹۰۷	۰.۰۹۸

توان دوم خطا و هم‌چنین شاخص پیتمن- نزدیکی مقایسه نمودند و از هردو معیار به این نتیجه رسیدند که طرح $RRSS$ برآوردگر کاراتری نسبت به رکوردهای معمولی دارد. بکلیزی (۲۰۰۸)، برآوردهای درست‌نمایی و بیزی قابلیت اعتماد تنش-مقاومت را براساس مقادیر رکورد پایین از توزیع نمایی تعمیم یافته بررسی نمود. او فواصل اطمینان دقیق و تقریبی را به‌دست آورد و عملکرد فواصل فوق را از طریق شبیه‌سازی مطالعه نمود. با توجه به این‌که توزیع نمایی تعمیم یافته جزء خانواده $PRHR$ می‌باشد، بنابراین، در این بخش قصد داریم فواصل اطمینان به دست‌آمده در جدول‌های ۱ و ۲ بکلیزی (۲۰۰۸) را با فواصل متناظر به‌دست‌آمده براساس طرح

جدول ۲.۴: ادامه

ردیف	R	n	m	MLE		AMLE		Bayes		HPD	
				EL	CP	EL	CP	EL	CP	EL	CP
۱	۰.۱	۳	۳	۰.۲۴۴	۰.۹۴۹	۰.۷۸۷	۰.۱۴۴	۰.۷۷۳	۰.۳۰۸	۰.۸۵۳	۰.۲۹۳
۲	۰.۱	۳	۵	۰.۱۷۷	۰.۹۳۸	۰.۶۲۴	۰.۰۸۰	۰.۷۵۸	۰.۲۵۵	۰.۸۶۰	۰.۲۴۰
۳	۰.۱	۳	۷	۰.۱۵۶	۰.۹۲۲	۰.۵۳۳	۰.۰۶۰	۰.۷۴۹	۰.۲۳۳	۰.۸۵۷	۰.۲۱۹
۴	۰.۱	۵	۳	۰.۲۱۷	۰.۹۴۰	۰.۷۱۰	۰.۰۹۹	۰.۸۶۷	۰.۲۲۱	۰.۹۰۶	۰.۲۱۲
۵	۰.۱	۵	۵	۰.۱۴۰	۰.۹۵۱	۰.۶۴۸	۰.۰۶۳	۰.۸۵۶	۰.۱۶۶	۰.۹۰۰	۰.۱۶۰
۶	۰.۱	۷	۵	۰.۱۱۶	۰.۹۴۷	۰.۵۵۳	۰.۰۴۴	۰.۸۵۶	۰.۱۴۴	۰.۹۰۳	۰.۱۳۸
۷	۰.۱	۷	۳	۰.۲۰۶	۰.۹۲۸	۰.۶۴۵	۰.۰۷۹	۰.۹۰۴	۰.۱۸۸	۰.۹۲۴	۰.۱۸۳
۸	۰.۱	۷	۵	۰.۱۲۴	۰.۹۴۹	۰.۵۷۲	۰.۰۴۷	۰.۹۰۵	۰.۱۳۳	۰.۹۱۷	۰.۱۳۰
۹	۰.۱	۷	۷	۰.۰۹۹	۰.۹۵۰	۰.۵۵۳	۰.۰۳۷	۰.۸۹۴	۰.۱۱۱	۰.۹۲۷	۰.۱۰۷
۱۰	۰.۳	۳	۳	۰.۴۴۹	۰.۹۴۹	۰.۷۰۸	۰.۲۶۴	۰.۹۵۴	۰.۴۱۴	۰.۹۳۹	۰.۴۰۳
۱۱	۰.۳	۳	۵	۰.۳۶۶	۰.۹۳۸	۰.۵۸۶	۰.۱۶۶	۰.۹۶۱	۰.۳۵۶	۰.۹۵۷	۰.۳۴۵
۱۲	۰.۳	۳	۷	۰.۳۳۷	۰.۹۲۲	۰.۵۱۱	۰.۱۲۹	۰.۹۷۱	۰.۳۳۱	۰.۹۶۳	۰.۳۲۰
۱۳	۰.۳	۵	۳	۰.۴۰۰	۰.۹۴۰	۰.۶۲۶	۰.۱۸۱	۰.۹۴۵	۰.۳۵۶	۰.۹۲۷	۰.۳۴۹
۱۴	۰.۳	۵	۵	۰.۲۹۴	۰.۹۵۱	۰.۶۰۹	۰.۱۳۲	۰.۹۵۱	۰.۲۸۴	۰.۹۴۲	۰.۲۸۸
۱۵	۰.۳	۵	۷	۰.۲۵۵	۰.۹۴۷	۰.۵۳۴	۰.۰۹۷	۰.۹۵۷	۰.۲۵۱	۰.۹۴۶	۰.۲۴۵
۱۶	۰.۳	۷	۳	۰.۳۸۰	۰.۹۲۸	۰.۵۶۱	۰.۱۴۵	۰.۹۴۳	۰.۳۳۲	۰.۹۲۷	۰.۳۲۷
۱۷	۰.۳	۷	۵	۰.۲۶۳	۰.۹۴۹	۰.۵۴۴	۰.۱۰۰	۰.۹۴۸	۰.۲۵۱	۰.۹۳۷	۰.۲۴۷
۱۸	۰.۳	۷	۷	۰.۲۱۷	۰.۹۵۰	۰.۵۳۴	۰.۰۸۲	۰.۹۵۲	۰.۲۱۳	۰.۹۴۵	۰.۲۰۹
۱۹	۰.۵	۳	۳	۰.۵۰۱	۰.۹۴۹	۰.۶۸۹	۰.۲۹۵	۰.۹۴۹	۰.۴۴۶	۰.۹۲۰	۰.۴۳۷
۲۰	۰.۵	۳	۵	۰.۴۳۳	۰.۹۳۸	۰.۵۸۷	۰.۱۹۶	۰.۹۴۶	۰.۳۹۲	۰.۹۲۳	۰.۳۸۵
۲۱	۰.۵	۳	۷	۰.۴۰۸	۰.۹۲۲	۰.۵۱۷	۰.۱۵۶	۰.۹۴۴	۰.۳۶۸	۰.۹۱۸	۰.۳۶۱
۲۲	۰.۵	۵	۳	۰.۴۳۳	۰.۹۴۰	۰.۵۹۰	۰.۱۹۷	۰.۹۵۱	۰.۳۹۲	۰.۹۲۲	۰.۳۸۵
۲۳	۰.۵	۵	۵	۰.۳۳۹	۰.۹۵۱	۰.۵۹۸	۰.۱۵۳	۰.۹۵۰	۰.۳۲۱	۰.۹۳۲	۰.۳۱۶
۲۴	۰.۵	۵	۷	۰.۳۰۱	۰.۹۴۷	۰.۵۳۳	۰.۱۱۴	۰.۹۴۸	۰.۲۸۷	۰.۹۳۲	۰.۲۸۳
۲۵	۰.۵	۷	۳	۰.۴۰۸	۰.۹۲۸	۰.۵۲۷	۰.۱۵۶	۰.۹۴۷	۰.۳۶۸	۰.۹۲۱	۰.۳۶۱
۲۶	۰.۵	۷	۵	۰.۳۰۱	۰.۹۴۹	۰.۵۲۸	۰.۱۱۴	۰.۹۴۹	۰.۲۸۷	۰.۹۳۱	۰.۲۸۲
۲۷	۰.۵	۷	۷	۰.۲۵۴	۰.۹۵۰	۰.۵۲۷	۰.۰۹۶	۰.۹۴۸	۰.۲۴۶	۰.۹۳۷	۰.۲۴۳
۲۸	۰.۷	۳	۳	۰.۴۴۸	۰.۹۴۹	۰.۷۰۶	۰.۲۶۳	۰.۹۳۶	۰.۴۰۱	۰.۹۱۱	۰.۳۹۰
۲۹	۰.۷	۳	۵	۰.۳۹۹	۰.۹۳۸	۰.۶۲۲	۰.۱۸۱	۰.۹۲۲	۰.۳۵۵	۰.۹۰۵	۰.۳۴۹
۳۰	۰.۷	۳	۷	۰.۳۸۰	۰.۹۲۲	۰.۵۴۸	۰.۱۴۵	۰.۹۱۲	۰.۳۳۴	۰.۸۹۶	۰.۳۲۹
۳۱	۰.۷	۵	۳	۰.۳۶۶	۰.۹۴۰	۰.۵۹۲	۰.۱۶۶	۰.۹۵۰	۰.۳۴۰	۰.۹۲۴	۰.۳۲۹
۳۲	۰.۷	۵	۵	۰.۲۹۴	۰.۹۵۱	۰.۶۰۸	۰.۱۳۲	۰.۹۴۴	۰.۲۸۰	۰.۹۳۰	۰.۲۷۳
۳۳	۰.۷	۵	۷	۰.۲۶۳	۰.۹۴۷	۰.۵۴۴	۰.۱۰۰	۰.۹۴۰	۰.۲۵۱	۰.۹۳۱	۰.۲۴۶
۳۴	۰.۷	۷	۳	۰.۳۳۸	۰.۹۲۸	۰.۵۲۲	۰.۱۲۹	۰.۹۴۷	۰.۳۱۵	۰.۹۲۱	۰.۳۰۲
۳۵	۰.۷	۷	۵	۰.۲۵۶	۰.۹۴۹	۰.۵۳۱	۰.۰۹۷	۰.۹۴۸	۰.۲۴۷	۰.۹۳۲	۰.۲۴۰
۳۶	۰.۷	۷	۷	۰.۲۱۸	۰.۹۵۰	۰.۵۳۳	۰.۰۸۳	۰.۹۴۷	۰.۲۱۲	۰.۹۳۷	۰.۲۰۸
۳۷	۰.۹	۳	۳	۰.۲۴۲	۰.۹۴۹	۰.۷۸۴	۰.۱۴۳	۰.۹۲۷	۰.۲۲۱	۰.۹۲۷	۰.۲۰۷
۳۸	۰.۹	۳	۵	۰.۲۱۶	۰.۹۳۸	۰.۷۰۲	۰.۰۹۸	۰.۹۰۵	۰.۱۹۲	۰.۹۱۹	۰.۱۸۴
۳۹	۰.۹	۳	۷	۰.۲۰۶	۰.۹۲۲	۰.۶۳۰	۰.۰۸۰	۰.۸۹۰	۰.۱۸۰	۰.۹۰۵	۰.۱۷۳
۴۰	۰.۹	۵	۳	۰.۱۷۶	۰.۹۴۰	۰.۶۳۲	۰.۰۸۰	۰.۹۴۸	۰.۱۶۸	۰.۹۳۲	۰.۱۵۶
۴۱	۰.۹	۵	۵	۰.۱۴۰	۰.۹۵۱	۰.۶۴۶	۰.۰۶۳	۰.۹۴۲	۰.۱۳۵	۰.۹۳۹	۰.۱۲۹
۴۲	۰.۹	۵	۷	۰.۱۲۵	۰.۹۴۷	۰.۵۷۸	۰.۰۴۸	۰.۹۳۵	۰.۱۲۰	۰.۹۳۷	۰.۱۱۶
۴۳	۰.۹	۷	۳	۰.۱۵۷	۰.۹۲۸	۰.۵۴۸	۰.۰۶۰	۰.۹۴۷	۰.۱۵۰	۰.۹۲۴	۰.۱۳۸
۴۴	۰.۹	۷	۵	۰.۱۱۷	۰.۹۴۹	۰.۵۵۴	۰.۰۴۵	۰.۹۴۷	۰.۱۱۵	۰.۹۳۸	۰.۱۰۹
۴۵	۰.۹	۷	۷	۰.۰۹۹	۰.۹۵۰	۰.۵۵۱	۰.۰۳۸	۰.۹۴۵	۰.۰۹۷	۰.۹۴۲	۰.۰۹۴

$RRSS$ مقایسه کنیم. معیار کارایی نسبی را به صورت $\delta = \frac{EL(R)}{EL(RRSS)}$ در نظر می گیریم. که در آن $EL^{(R)}$ و $EL^{(RRSS)}$ به ترتیب متوسط طول فواصل اطمینان براساس رکوردهای معمولی و طرح $RRSS$ می باشند. در حقیقت زمانی که $\delta > 1$ ، فاصله اطمینان براساس $RRSS$ عملکرد بهتری در مقایسه با فاصله اطمینان براساس رکوردهای معمولی دارد. حال با در نظر گرفتن توزیع نمایی تعمیم یافته، مقادیر δ را برای ترکیب های مختلف n و m در نظر گرفته شده در بکلیزی (۲۰۰۸)، محاسبه نمودیم و نتایج را در جدول های ۵.۴ و ۶.۴ ارائه دادیم.

ملاحظه می شود در تمامی حالت ها δ مقداری مثبت اختیار می کند، بنابراین همان طور که

جدول ۳.۴: CP و EL برای R به ازای $\gamma = 1\%$

Basic		Boot - t^*		Boot - t		perc		m	n	R	ردیف
EL	CP	EL	CP	EL	CP	EL	CP				
۰.۱۹۹	۰.۸۰۲	۰.۱۹۹	۰.۸۵۰	۰.۱۸۷	۰.۸۱۶	۰.۲۰۰	۰.۸۹۷	۳	۳	۰.۱	۱
۰.۱۴۷	۰.۸۱۱	۰.۱۶۲	۰.۸۱۵	۰.۱۵۸	۰.۷۹۹	۰.۱۴۷	۰.۸۸۸	۵	۳	۰.۱	۲
۰.۱۳۰	۰.۸۱۲	۰.۱۵۱	۰.۷۹۸	۰.۱۴۸	۰.۷۸۴	۰.۱۳۱	۰.۸۷۷	۷	۳	۰.۱	۳
۰.۱۷۶	۰.۸۴۰	۰.۱۸۲	۰.۸۹۷	۰.۱۶۹	۰.۸۷۰	۰.۱۷۶	۰.۸۹۲	۳	۵	۰.۱	۴
۰.۱۱۵	۰.۸۵۵	۰.۱۲۶	۰.۸۷۷	۰.۱۲۰	۰.۸۶۲	۰.۱۱۶	۰.۸۹۸	۵	۵	۰.۱	۵
۰.۰۹۷	۰.۸۶۲	۰.۱۰۸	۰.۸۶۷	۰.۱۰۴	۰.۸۵۵	۰.۰۹۷	۰.۸۹۷	۷	۵	۰.۱	۶
۰.۱۶۷	۰.۸۵۳	۰.۱۶۷	۰.۹۰۶	۰.۱۵۴	۰.۸۸۲	۰.۱۶۷	۰.۸۸۲	۳	۷	۰.۱	۷
۰.۱۰۲	۰.۸۶۰	۰.۱۰۷	۰.۸۹۱	۰.۱۰۲	۰.۸۷۵	۰.۱۰۳	۰.۸۹۸	۵	۷	۰.۱	۸
۰.۰۸۲	۰.۸۷۰	۰.۰۸۷	۰.۸۸۴	۰.۰۸۴	۰.۸۷۵	۰.۰۸۲	۰.۸۹۸	۷	۷	۰.۱	۹
۰.۳۷۹	۰.۷۸۴	۰.۴۵۳	۰.۸۶۰	۰.۴۵۱	۰.۸۳۴	۰.۳۸۰	۰.۸۹۸	۳	۳	۰.۳	۱۰
۰.۳۱۰	۰.۸۰۴	۰.۳۸۰	۰.۸۴۳	۰.۳۷۶	۰.۸۲۷	۰.۳۱۱	۰.۸۸۹	۵	۳	۰.۳	۱۱
۰.۲۸۶	۰.۸۱۱	۰.۳۵۵	۰.۸۳۱	۰.۳۵۰	۰.۸۲۰	۰.۲۸۶	۰.۸۷۸	۷	۳	۰.۳	۱۲
۰.۳۳۵	۰.۸۲۵	۰.۳۸۰	۰.۸۸۴	۰.۳۷۰	۰.۸۶۶	۰.۳۳۶	۰.۸۹۳	۳	۵	۰.۳	۱۳
۰.۲۴۷	۰.۸۵۱	۰.۲۷۵	۰.۸۸۵	۰.۲۶۸	۰.۸۷۴	۰.۲۴۸	۰.۸۹۹	۵	۵	۰.۳	۱۴
۰.۲۱۵	۰.۸۵۷	۰.۲۴۰	۰.۸۷۹	۰.۲۳۳	۰.۸۷۱	۰.۲۱۵	۰.۸۹۸	۷	۵	۰.۳	۱۵
۰.۳۱۸	۰.۸۳۳	۰.۳۵۳	۰.۸۸۲	۰.۳۴۰	۰.۸۶۷	۰.۳۱۹	۰.۸۸۳	۳	۷	۰.۳	۱۶
۰.۲۲۰	۰.۸۵۶	۰.۲۳۸	۰.۸۹۲	۰.۲۳۰	۰.۸۷۸	۰.۲۲۱	۰.۸۹۹	۵	۷	۰.۳	۱۷
۰.۱۸۲	۰.۸۶۷	۰.۱۹۶	۰.۸۹۴	۰.۱۹۰	۰.۸۸۰	۰.۱۸۳	۰.۹۰۰	۷	۷	۰.۳	۱۸
۰.۴۲۷	۰.۷۷۵	۰.۵۲۴	۰.۸۶۱	۰.۵۳۱	۰.۸۴۲	۰.۴۲۸	۰.۸۹۸	۳	۳	۰.۵	۱۹
۰.۳۶۸	۰.۸۰۲	۰.۴۳۹	۰.۸۵۹	۰.۴۳۴	۰.۸۴۵	۰.۳۶۹	۰.۸۸۸	۵	۳	۰.۵	۲۰
۰.۳۴۶	۰.۸۱۴	۰.۴۱۱	۰.۸۵۳	۰.۴۰۳	۰.۸۴۱	۰.۳۴۷	۰.۸۷۸	۷	۳	۰.۵	۲۱
۰.۳۶۸	۰.۸۱۱	۰.۴۳۹	۰.۸۶۴	۰.۴۳۴	۰.۸۵۱	۰.۳۶۹	۰.۸۹۳	۳	۵	۰.۵	۲۲
۰.۲۸۶	۰.۸۴۳	۰.۳۲۱	۰.۸۸۴	۰.۳۱۴	۰.۸۷۴	۰.۲۸۷	۰.۸۹۹	۵	۵	۰.۵	۲۳
۰.۲۵۴	۰.۸۵۹	۰.۲۸۰	۰.۸۸۷	۰.۲۷۳	۰.۸۷۷	۰.۲۵۴	۰.۸۹۷	۷	۵	۰.۵	۲۴
۰.۳۴۶	۰.۸۲۱	۰.۴۱۱	۰.۸۶۰	۰.۴۰۳	۰.۸۵۲	۰.۳۴۷	۰.۸۸۳	۳	۷	۰.۵	۲۵
۰.۲۵۴	۰.۸۵۳	۰.۲۸۰	۰.۸۸۸	۰.۲۷۲	۰.۸۷۸	۰.۲۵۴	۰.۹۰۰	۵	۷	۰.۵	۲۶
۰.۲۱۴	۰.۸۶۶	۰.۲۳۱	۰.۸۹۵	۰.۲۲۵	۰.۸۸۳	۰.۲۱۴	۰.۹۰۰	۷	۷	۰.۵	۲۷
۰.۳۷۸	۰.۷۸۳	۰.۴۵۱	۰.۸۶۰	۰.۴۴۹	۰.۸۳۶	۰.۳۷۹	۰.۸۹۸	۳	۳	۰.۷	۲۸
۰.۳۳۵	۰.۸۱۸	۰.۳۷۹	۰.۸۷۵	۰.۳۶۹	۰.۸۵۷	۰.۳۳۶	۰.۸۸۹	۵	۳	۰.۷	۲۹
۰.۳۱۸	۰.۸۲۹	۰.۳۵۵	۰.۸۷۵	۰.۳۴۰	۰.۸۵۸	۰.۳۱۹	۰.۸۷۸	۷	۳	۰.۷	۳۰
۰.۳۰۹	۰.۸۱۰	۰.۳۷۹	۰.۸۴۶	۰.۳۷۵	۰.۸۳۳	۰.۳۱۰	۰.۸۹۳	۳	۵	۰.۷	۳۱
۰.۲۴۷	۰.۸۴۳	۰.۲۷۵	۰.۸۸۲	۰.۲۶۷	۰.۸۶۹	۰.۲۴۸	۰.۸۹۹	۵	۵	۰.۷	۳۲
۰.۲۲۱	۰.۸۶۱	۰.۲۳۹	۰.۸۹۶	۰.۲۳۱	۰.۸۸۵	۰.۲۲۱	۰.۸۹۷	۷	۵	۰.۷	۳۳
۰.۲۸۶	۰.۸۲۲	۰.۳۵۵	۰.۸۳۸	۰.۳۵۰	۰.۸۲۹	۰.۲۸۷	۰.۸۸۳	۳	۷	۰.۷	۳۴
۰.۲۱۵	۰.۸۶۱	۰.۲۴۰	۰.۸۸۳	۰.۲۳۳	۰.۸۷۳	۰.۲۱۶	۰.۸۹۹	۵	۷	۰.۷	۳۵
۰.۱۸۲	۰.۸۶۶	۰.۱۹۶	۰.۸۹۲	۰.۱۹۰	۰.۸۸۱	۰.۱۸۳	۰.۹۰۰	۷	۷	۰.۷	۳۶
۰.۱۹۸	۰.۸۰۰	۰.۱۹۸	۰.۸۴۶	۰.۱۸۶	۰.۸۱۳	۰.۱۹۹	۰.۸۹۸	۳	۳	۰.۹	۳۷
۰.۱۷۵	۰.۸۳۵	۰.۱۸۱	۰.۸۸۹	۰.۱۶۹	۰.۸۶۴	۰.۱۷۶	۰.۸۸۹	۵	۳	۰.۹	۳۸
۰.۱۶۷	۰.۸۴۵	۰.۱۶۸	۰.۹۰۱	۰.۱۵۴	۰.۸۷۷	۰.۱۶۸	۰.۸۷۸	۷	۳	۰.۹	۳۹
۰.۱۴۶	۰.۸۱۴	۰.۱۶۱	۰.۸۱۶	۰.۱۵۷	۰.۸۰۰	۰.۱۴۶	۰.۸۹۳	۳	۵	۰.۹	۴۰
۰.۱۱۵	۰.۸۴۸	۰.۱۲۶	۰.۸۷۳	۰.۱۲۰	۰.۸۵۶	۰.۱۱۶	۰.۸۹۹	۵	۵	۰.۹	۴۱
۰.۱۰۳	۰.۸۶۵	۰.۱۰۸	۰.۸۹۸	۰.۱۰۲	۰.۸۸۳	۰.۱۰۳	۰.۸۹۷	۷	۵	۰.۹	۴۲
۰.۱۳۱	۰.۸۲۲	۰.۱۵۲	۰.۸۰۷	۰.۱۴۹	۰.۷۹۶	۰.۱۳۱	۰.۸۸۳	۳	۷	۰.۹	۴۳
۰.۰۹۷	۰.۸۶۴	۰.۱۰۸	۰.۸۷۱	۰.۱۰۴	۰.۸۶۲	۰.۰۹۸	۰.۸۹۹	۵	۷	۰.۹	۴۴
۰.۰۸۲	۰.۸۷۰	۰.۰۸۷	۰.۸۸۷	۰.۰۸۴	۰.۸۷۶	۰.۰۸۲	۰.۹۰۰	۷	۷	۰.۹	۴۵

انتظار داشتیم فواصل اطمینان براساس طرح نمونه گیری $RRSS$ در خانواده $PRHR$ بهتر از رکورهای معمولی رفتار می کنند.

۶.۴ مثال واقعی

برای شرح بهتر فرایندهای ارائه شده در بخش های قبل، یک مجموعه داده واقعی در نظر می گیریم. در این بررسی، داده های مربوط به زمان شکست های متوالی سیستم تهویه هوا در هر یک از

جدول ۴.۴: ادامه

ردیف	R	n	m	MLE		AMLE		Bayes		HPD	
				EL	CP	EL	CP	EL	CP	EL	CP
۱	۰.۱	۳	۳	۰.۸۹۹	۰.۲۰۰	۰.۷۱۴	۰.۱۲۱	۰.۶۳۹	۰.۲۵۷	۰.۷۶۷	۰.۲۴۴
۲	۰.۱	۳	۵	۰.۸۸۹	۰.۱۴۷	۰.۵۴۸	۰.۰۶۷	۰.۶۰۸	۰.۲۱۱	۰.۷۷۲	۰.۱۹۸
۳	۰.۱	۳	۷	۰.۸۷۷	۰.۱۳۱	۰.۴۶۵	۰.۰۵۰	۰.۵۹۶	۰.۱۹۲	۰.۷۶۴	۰.۱۸۰
۴	۰.۱	۵	۳	۰.۸۹۳	۰.۱۷۶	۰.۶۲۷	۰.۰۸۳	۰.۷۷۲	۰.۱۸۵	۰.۸۳۶	۰.۱۷۸
۵	۰.۱	۵	۵	۰.۹۰۱	۰.۱۱۶	۰.۵۶۴	۰.۰۵۳	۰.۷۶۶	۰.۱۳۹	۰.۸۳۴	۰.۱۳۴
۶	۰.۱	۷	۵	۰.۸۹۸	۰.۰۹۷	۰.۴۸۰	۰.۰۳۷	۰.۷۵۸	۰.۱۲۰	۰.۸۳۱	۰.۱۱۵
۷	۰.۱	۷	۳	۰.۸۸۵	۰.۱۶۷	۰.۵۶۲	۰.۰۶۷	۰.۸۲۸	۰.۱۵۸	۰.۸۵۹	۰.۱۵۳
۸	۰.۱	۷	۵	۰.۹۰۰	۰.۱۰۳	۰.۴۹۶	۰.۰۴۰	۰.۸۳۲	۰.۱۱۱	۰.۸۵۲	۰.۱۰۹
۹	۰.۱	۷	۷	۰.۸۹۹	۰.۰۸۲	۰.۴۷۳	۰.۰۳۱	۰.۸۱۸	۰.۰۹۲	۰.۸۶۵	۰.۰۸۹
۱۰	۰.۳	۳	۳	۰.۸۹۹	۰.۳۸۰	۰.۶۲۹	۰.۲۲۲	۰.۹۰۸	۰.۳۵۱	۰.۸۸۳	۰.۳۴۱
۱۱	۰.۳	۳	۵	۰.۸۸۹	۰.۳۱۱	۰.۵۰۸	۰.۱۳۹	۰.۹۱۹	۰.۲۹۹	۰.۹۱۰	۰.۲۹۰
۱۲	۰.۳	۳	۷	۰.۸۷۷	۰.۲۸۷	۰.۴۴۴	۰.۱۰۸	۰.۹۳۳	۰.۲۷۷	۰.۹۱۲	۰.۲۶۸
۱۳	۰.۳	۵	۳	۰.۸۹۳	۰.۳۳۶	۰.۵۴۲	۰.۱۵۲	۰.۸۹۳	۰.۳۰۲	۰.۸۶۵	۰.۲۹۶
۱۴	۰.۳	۵	۵	۰.۹۰۱	۰.۲۴۸	۰.۵۳۴	۰.۱۱۱	۰.۹۰۳	۰.۲۳۹	۰.۸۸۸	۰.۲۳۴
۱۵	۰.۳	۵	۷	۰.۸۹۸	۰.۲۱۵	۰.۴۶۴	۰.۰۸۱	۰.۹۱۱	۰.۲۱۱	۰.۸۹۵	۰.۲۰۶
۱۶	۰.۳	۷	۳	۰.۸۸۵	۰.۳۱۹	۰.۴۸۱	۰.۱۲۲	۰.۸۹۰	۰.۲۸۱	۰.۸۶۳	۰.۲۷۷
۱۷	۰.۳	۷	۵	۰.۹۰۰	۰.۲۲۱	۰.۴۶۸	۰.۰۸۴	۰.۸۹۸	۰.۲۱۱	۰.۸۷۸	۰.۲۰۹
۱۸	۰.۳	۷	۷	۰.۸۹۹	۰.۱۸۳	۰.۴۵۷	۰.۰۶۹	۰.۹۰۰	۰.۱۷۹	۰.۸۹۴	۰.۱۷۶
۱۹	۰.۵	۳	۳	۰.۸۹۹	۰.۴۲۸	۰.۶۱۰	۰.۲۴۸	۰.۸۹۹	۰.۳۸۰	۰.۸۵۰	۰.۳۷۲
۲۰	۰.۵	۳	۵	۰.۸۸۹	۰.۳۶۹	۰.۵۰۸	۰.۱۶۵	۰.۸۹۴	۰.۳۳۲	۰.۸۵۷	۰.۳۲۷
۲۱	۰.۵	۳	۷	۰.۸۷۷	۰.۳۴۷	۰.۴۴۷	۰.۱۳۱	۰.۸۸۶	۰.۳۱۱	۰.۸۵۰	۰.۳۰۵
۲۲	۰.۵	۵	۳	۰.۸۹۳	۰.۳۶۹	۰.۵۱۲	۰.۱۶۵	۰.۸۹۹	۰.۳۳۲	۰.۸۵۴	۰.۳۲۷
۲۳	۰.۵	۵	۵	۰.۹۰۱	۰.۲۸۷	۰.۵۲۶	۰.۱۲۸	۰.۸۹۸	۰.۲۷۱	۰.۸۷۲	۰.۲۶۷
۲۴	۰.۵	۵	۷	۰.۸۹۸	۰.۲۵۴	۰.۴۶۲	۰.۰۹۶	۰.۸۹۸	۰.۲۴۲	۰.۸۷۴	۰.۲۳۹
۲۵	۰.۵	۷	۳	۰.۸۸۵	۰.۳۴۷	۰.۴۵۲	۰.۱۳۱	۰.۸۹۵	۰.۳۱۲	۰.۸۵۰	۰.۳۰۶
۲۶	۰.۵	۷	۵	۰.۹۰۰	۰.۲۵۴	۰.۴۵۷	۰.۰۹۶	۰.۹۰۰	۰.۲۴۲	۰.۸۶۹	۰.۲۳۹
۲۷	۰.۵	۷	۷	۰.۸۹۹	۰.۲۱۴	۰.۴۵۱	۰.۰۸۱	۰.۸۹۷	۰.۲۰۷	۰.۸۸۳	۰.۲۰۵
۲۸	۰.۷	۳	۳	۰.۸۹۹	۰.۳۷۹	۰.۶۳۱	۰.۲۲۱	۰.۸۸۱	۰.۳۳۹	۰.۸۴۳	۰.۳۳۰
۲۹	۰.۷	۳	۵	۰.۸۸۹	۰.۳۳۶	۰.۵۳۷	۰.۱۵۲	۰.۸۵۸	۰.۳۰۱	۰.۸۳۷	۰.۲۹۶
۳۰	۰.۷	۳	۷	۰.۸۷۷	۰.۳۱۹	۰.۴۷۵	۰.۱۲۲	۰.۸۴۶	۰.۲۸۳	۰.۸۲۴	۰.۲۷۹
۳۱	۰.۷	۵	۳	۰.۸۹۳	۰.۳۱۰	۰.۵۱۴	۰.۱۳۹	۰.۸۹۸	۰.۲۸۶	۰.۸۵۴	۰.۲۷۶
۳۲	۰.۷	۵	۵	۰.۹۰۱	۰.۲۴۸	۰.۵۳۵	۰.۱۱۱	۰.۸۹۲	۰.۲۳۵	۰.۸۷۱	۰.۲۳۰
۳۳	۰.۷	۵	۷	۰.۸۹۸	۰.۲۲۱	۰.۴۷۴	۰.۰۸۴	۰.۸۸۷	۰.۲۱۱	۰.۸۷۰	۰.۲۰۸
۳۴	۰.۷	۷	۳	۰.۸۸۵	۰.۲۸۷	۰.۴۴۹	۰.۱۰۹	۰.۸۹۶	۰.۲۶۳	۰.۸۴۹	۰.۲۵۳
۳۵	۰.۷	۷	۵	۰.۹۰۰	۰.۲۱۶	۰.۴۶۰	۰.۰۸۲	۰.۸۹۹	۰.۲۰۷	۰.۸۷۰	۰.۲۰۲
۳۶	۰.۷	۷	۷	۰.۸۹۹	۰.۱۸۳	۰.۴۵۶	۰.۰۶۹	۰.۸۹۴	۰.۱۷۸	۰.۸۸۳	۰.۱۷۵
۳۷	۰.۹	۳	۳	۰.۸۹۹	۰.۱۹۹	۰.۷۱۱	۰.۱۲۰	۰.۸۶۶	۰.۱۸۲	۰.۸۷۰	۰.۱۷۰
۳۸	۰.۹	۳	۵	۰.۸۸۹	۰.۱۷۶	۰.۶۱۹	۰.۰۸۳	۰.۸۳۴	۰.۱۶۰	۰.۸۶۰	۰.۱۵۴
۳۹	۰.۹	۳	۷	۰.۸۷۷	۰.۱۶۸	۰.۵۵۰	۰.۰۶۷	۰.۸۲۰	۰.۱۵۱	۰.۸۴۲	۰.۱۴۶
۴۰	۰.۹	۵	۳	۰.۸۹۳	۰.۱۴۶	۰.۵۵۰	۰.۰۶۷	۰.۸۹۹	۰.۱۳۸	۰.۸۶۹	۰.۱۲۷
۴۱	۰.۹	۵	۵	۰.۹۰۱	۰.۱۱۶	۰.۵۶۸	۰.۰۵۳	۰.۸۸۵	۰.۱۱۲	۰.۸۸۴	۰.۱۰۷
۴۲	۰.۹	۷	۵	۰.۸۹۸	۰.۱۰۳	۰.۵۰۳	۰.۰۴۰	۰.۸۷۸	۰.۱۰۰	۰.۸۸۳	۰.۰۹۷
۴۳	۰.۹	۷	۳	۰.۸۸۵	۰.۱۳۱	۰.۴۷۲	۰.۰۵۰	۰.۸۹۶	۰.۱۲۲	۰.۸۵۵	۰.۱۱۲
۴۴	۰.۹	۷	۵	۰.۹۰۰	۰.۰۹۸	۰.۴۷۸	۰.۰۳۷	۰.۸۹۶	۰.۰۹۵	۰.۸۷۷	۰.۰۹۰
۴۵	۰.۹	۷	۷	۰.۸۹۹	۰.۰۸۲	۰.۴۷۲	۰.۰۳۱	۰.۸۹۲	۰.۰۸۱	۰.۸۹۰	۰.۰۷۸

هوایماهای بوئینگ ۷۲۰ در نظر گرفته شده است. داده های ذکر شده در پروسیچان^۷ (۱۹۶۳)، در دسترس می باشند. به طور تصادفی داده های مربوط به دو هوایمای ۷۹۰۸ و ۸۰۴۵ را برای مطالعه قابلیت اعتماد تنش-مقاومت انتخاب نمودیم. سپس برای بررسی این که آیا توزیع نمایی تعمیم یافته بر داده های فوق مطابقت دارد یا خیر از آزمون کولموگروف-اسمیرنوف^۸ $(K - S)$ ، استفاده نمودیم. مشاهده شد P -مقدار آزمون $K - S$ برای مشاهدات X و Y به ترتیب ۰.۳۷۲ و ۰.۳۷۲.

^۷Proschan

^۸Kolmogorov-Smirnov

جدول ۵.۴: مقادیر δ برای ترکیب‌هایی از m و n زمانی که $\gamma = 0.5$

ردیف	n	m	$\mathcal{R}$	MLE			Perc			Boot - t\			Boot - tY			HPD		
				CP^{RRSS}	CP^R	δ	CP^{RRSS}	CP^R	δ	CP^{RRSS}	CP^R	δ	CP^{RRSS}	CP^R	δ	CP^{RRSS}	CP^R	δ
۱	۵	۵	۰.۱	۱.۵۷۶	۰.۸۳۲	۰.۹۲۰	۲.۲۲۳	۰.۹۱۲	۰.۹۰۴	۲.۹۵۹	۰.۸۸۱	۰.۹۲۳	۳.۱۰۱	۰.۸۸۳	۰.۹۱۹	۱.۸۷۹	۰.۸۷۷	۰.۹۰۲
۲	۵	۵	۰.۳	۱.۴۷۶	۰.۸۰۶	۰.۹۲۰	۱.۷۰۱	۰.۹۱۹	۰.۹۰۲	۲.۵۹۹	۰.۸۹۱	۰.۹۲۹	۲.۳۱۵	۰.۸۸۸	۰.۹۲۲	۱.۶۷۹	۰.۸۶۵	۰.۸۹۵
۳	۵	۵	۰.۵	۱.۵۰۹	۰.۸۱۹	۰.۹۲۰	۱.۵۹۳	۰.۹۱۸	۰.۸۹۸	۲.۶۲۱	۰.۹۰۲	۰.۹۳۰	۲.۱۷۵	۰.۸۹۷	۰.۹۳۸	۱.۷۰۰	۰.۸۸۱	۰.۹۰۸
۴	۵	۱۰	۰.۱	۱.۸۴۷	۰.۸۸۶	۰.۸۸۸	۳.۰۴۹	۰.۸۴۵	۰.۸۷۰	۲.۵۷۸	۰.۹۳۲	۰.۸۸۴	۲.۱۴۳	۰.۹۴۹	۰.۹۶۹	۲.۱۷۶	۰.۹۰۰	۰.۹۲۵
۵	۵	۱۰	۰.۳	۱.۶۶۹	۰.۸۴۶	۰.۸۸۸	۲.۰۷۸	۰.۸۳۳	۰.۸۷۳	۲.۵۳۹	۰.۹۱۲	۰.۸۹۲	۱.۸۸۱	۰.۹۰۹	۰.۹۶۸	۱.۹۹۱	۰.۸۸۷	۰.۹۰۹
۶	۵	۱۰	۰.۵	۱.۵۸۳	۰.۸۱۴	۰.۸۸۸	۱.۶۹۱	۰.۸۳۵	۰.۸۶۹	۲.۵۹۳	۰.۸۳۷	۰.۸۹۷	۱.۸۴۲	۰.۸۵۰	۰.۹۶۶	۱.۹۲۵	۰.۸۸۹	۰.۹۰۷
۷	۵	۱۵	۰.۱	۱.۹۶۲	۰.۹۱۰	۰.۹۰۷	۳.۲۹۲	۰.۷۷۹	۰.۸۸۱	۲.۵۱۴	۰.۹۳۸	۰.۹۰۵	۱.۷۲۱	۰.۹۴۳	۰.۹۹۵	۲.۲۴۸	۰.۸۹۹	۰.۹۲۱
۸	۵	۱۵	۰.۳	۱.۷۲۶	۰.۸۳۳	۰.۹۰۷	۲.۱۲۴	۰.۷۷۴	۰.۸۸۰	۲.۵۰۴	۰.۸۶۲	۰.۹۱۶	۱.۵۵۳	۰.۸۵۴	۰.۹۹۵	۲.۰۲۷	۰.۸۶۵	۰.۹۱۲
۹	۵	۱۵	۰.۵	۱.۶۰۵	۰.۸۱۰	۰.۹۰۷	۱.۶۶۷	۰.۷۸۳	۰.۸۷۸	۲.۵۲۲	۰.۸۱۱	۰.۹۱۷	۱.۵۳۹	۰.۸۰۸	۰.۹۹۱	۱.۹۵۲	۰.۸۸۶	۰.۹۰۲
۱۰	۱۰	۱۰	۰.۱	۱.۸۹۷	۰.۸۶۸	۰.۸۹۸	۱.۵۶۰	۰.۸۸۷	۰.۸۹۲	۳.۱۵۲	۰.۸۲۰	۰.۹۰۷	۲.۴۷۶	۰.۷۶۳	۰.۹۷۰	۱.۸۷۷	۰.۸۹۹	۰.۹۰۲
۱۱	۱۰	۱۰	۰.۳	۲.۰۰۴	۰.۸۸۶	۰.۸۹۸	۱.۵۷۰	۰.۸۹۳	۰.۸۹۳	۲.۹۵۳	۰.۸۸۰	۰.۹۱۵	۲.۰۵۸	۰.۸۱۶	۰.۹۷۷	۱.۸۶۴	۰.۸۹۵	۰.۹۰۴
۱۲	۱۰	۱۰	۰.۵	۲.۰۹۸	۰.۸۸۷	۰.۸۹۸	۱.۶۷۹	۰.۸۷۷	۰.۸۸۷	۲.۸۶۴	۰.۹۱۶	۰.۹۱۶	۱.۸۹۳	۰.۸۳۳	۰.۹۶۶	۱.۹۱۰	۰.۸۹۳	۰.۹۲۱
۱۳	۱۰	۱۰	۰.۱	۲.۲۱۷	۰.۹۰۳	۰.۸۹۷	۲.۵۸۰	۰.۹۳۳	۰.۹۰۸	۳	۰.۹۱۵	۰.۹۰۶	۳.۱۶۷	۰.۸۹۶	۰.۹۷۵	۲.۵۵۹	۰.۸۸۶	۰.۹۰۳
۱۴	۱۰	۱۰	۰.۳	۲.۱۶۷	۰.۸۸۹	۰.۸۹۷	۲.۲۳۷	۰.۹۱۸	۰.۹۰۷	۲.۸۹۴	۰.۹۳۶	۰.۹۰۸	۲.۷۹۶	۰.۹۱۶	۰.۹۱۵	۲.۴۷۷	۰.۹۰۱	۰.۹۰۱
۱۵	۱۰	۱۰	۰.۵	۲.۱۵۸	۰.۸۵۸	۰.۸۹۷	۲.۱۰۹	۰.۹۱۰	۰.۹۰۶	۲.۸۸۹	۰.۹۱۴	۰.۹۱۰	۲.۶۶۳	۰.۹۱۱	۰.۹۱۹	۲.۴۷۸	۰.۹۰۷	۰.۹۱۰
۱۶	۱۰	۱۵	۰.۱	۲.۴۴۸	۰.۹۱۰	۰.۸۹۲	۳.۱۰۵	۰.۸۹۱	۰.۸۸۶	۳.۰۸۵	۰.۹۲۲	۰.۸۸۸	۲.۸۳۱	۰.۹۳۵	۰.۹۴۴	۲.۸۷۹	۰.۹۲۱	۰.۹۰۰
۱۷	۱۰	۱۵	۰.۳	۲.۲۷۱	۰.۸۸۵	۰.۸۹۲	۲.۵۱۵	۰.۹۰۶	۰.۸۸۶	۲.۹۴۹	۰.۹۴۲	۰.۸۹۲	۲.۵۲۷	۰.۹۳۵	۰.۹۳۷	۲.۷۲۷	۰.۹۱۵	۰.۸۹۷
۱۸	۱۰	۱۵	۰.۵	۲.۲۱۵	۰.۸۷۶	۰.۸۹۲	۲.۲۶۶	۰.۹۰۰	۰.۸۸۳	۲.۹۵۱	۰.۹۱۰	۰.۸۹۵	۲.۴۵۴	۰.۹۲۱	۰.۹۳۸	۲.۷۱۶	۰.۹۲۳	۰.۸۹۸
۱۹	۱۵	۱۵	۰.۱	۲.۳۳۰	۰.۸۹۱	۰.۸۹۷	۱.۴۱۳	۰.۸۴۱	۰.۸۸۳	۲.۳۹۱	۰.۷۷۲	۰.۹۰۱	۲.۵۰۹	۰.۷۱۰	۰.۹۹۲	۱.۸۶۹	۰.۸۸۸	۰.۹۰۳
۲۰	۱۵	۲۰	۰.۳	۲.۴۳۷	۰.۹۱۸	۰.۸۹۷	۱.۴۸۷	۰.۸۲۴	۰.۸۸۴	۳	۰.۸۱۴	۰.۹۰۴	۱.۷۳۱	۰.۷۲۹	۰.۹۶۶	۱.۸۵۵	۰.۸۸۵	۰.۹۰۴
۲۱	۱۵	۲۰	۰.۵	۲.۶۳۶	۰.۹۳۰	۰.۸۹۷	۱.۶۸۱	۰.۸۳۷	۰.۸۸۱	۲.۹۰۷	۰.۸۷۲	۰.۹۰۵	۱.۶۱۳	۰.۷۹۵	۰.۹۹۴	۱.۹۴۴	۰.۸۹۳	۰.۹۱۰
۲۲	۱۵	۲۲	۰.۱	۲.۶۲۱	۰.۹۱۵	۰.۸۹۰	۲.۴۵۸	۰.۹۱۰	۰.۸۸۸	۳.۲۳۷	۰.۸۷۴	۰.۸۸۶	۲.۸۸۶	۰.۸۵۱	۰.۹۴۰	۲.۶۵۵	۰.۹۱۵	۰.۸۹۴
۲۳	۱۵	۲۳	۰.۳	۲.۵۹۴	۰.۹۲۲	۰.۸۹۰	۲.۲۶۹	۰.۹۲۳	۰.۸۸۶	۳.۱۰۳	۰.۹۲۳	۰.۸۷۶	۲.۵۶۱	۰.۸۹۱	۰.۹۶۷	۲.۶۳۴	۰.۹۲۰	۰.۸۹۶
۲۴	۱۵	۲۴	۰.۵	۲.۶۵۲	۰.۹۲۰	۰.۸۹۰	۲.۲۵۹	۰.۹۲۷	۰.۸۸۸	۳.۱۱۸	۰.۹۲۴	۰.۸۷۸	۲.۴۸۲	۰.۹۱۹	۰.۹۴۷	۲.۷۰۳	۰.۹۳۱	۰.۸۹۳
۲۵	۱۵	۲۵	۰.۱	۲.۷۳۹	۰.۹۰۹	۰.۹۱۵	۲.۹۵۷	۰.۹۱۱	۰.۹۰۶	۳.۳۴۸	۰.۹۲۱	۰.۹۰۴	۳.۴۷۹	۰.۹۰۰	۰.۹۱۰	۳.۰۶۵	۰.۹۱۶	۰.۹۰۱
۲۶	۱۵	۲۶	۰.۳	۲.۶۷۴	۰.۹۱۳	۰.۹۱۵	۲.۶۶۰	۰.۹۲۷	۰.۹۰۵	۳.۲۷۱	۰.۹۴۵	۰.۹۰۴	۳.۲۰۵	۰.۹۲۶	۰.۹۱۲	۳.۰۴۸	۰.۹۳۲	۰.۹۰۱
۲۷	۱۵	۲۷	۰.۵	۲.۶۹۰	۰.۹۱۳	۰.۹۱۵	۲.۵۶۸	۰.۹۲۰	۰.۹۰۶	۳.۲۹۹	۰.۹۴۰	۰.۹۰۴	۳.۱۲۸	۰.۹۳۸	۰.۹۱۳	۳.۱۳۰	۰.۹۳۴	۰.۸۹۶

جدول ۶.۴: مقادیر δ برای ترکیب‌هایی از m و n زمانی که $\gamma = 0.1$

HPD			Boot - tY			Boot - t\			Perc			MLE						
CP ^{RRSS}	CP ^R	δ	CP ^{RRSS}	CP ^R	δ	CP ^{RRSS}	CP ^R	δ	CP ^{RRSS}	CP ^R	δ	CP ^{RRSS}	CP ^R	δ	R	m	n	ردیف
۰.۹۰۲	۰.۸۰۲	۱.۸۵۳	۰.۹۱۹	۰.۸۳۶	۲.۷۵۶	۰.۹۲۳	۰.۸۴۱	۲.۶۶۰	۰.۹۰۴	۰.۸۴۶	۲.۰۸۶	۰.۹۲۰	۰.۷۷۴	۱.۵۶۵	۰.۱	۵	۵	۱
۰.۸۹۵	۰.۷۸۵	۱.۷۶۴	۰.۹۳۲	۰.۸۵۲	۲.۱۳۵	۰.۹۲۹	۰.۸۶۴	۲.۳۶۱	۰.۹۰۲	۰.۸۵۶	۱.۶۶۹	۰.۹۲۰	۰.۷۶۰	۱.۴۷۶	۰.۳	۵	۵	۲
۰.۹۰۸	۰.۸۱۱	۱.۸۰۹	۰.۹۳۸	۰.۸۵۵	۲.۰۲۱	۰.۹۳۰	۰.۸۵۷	۲.۳۸۱	۰.۸۹۸	۰.۸۵۲	۱.۵۵۹	۰.۹۲۰	۰.۷۵۲	۱.۴۹۷	۰.۵	۵	۵	۳
۰.۹۲۵	۰.۸۱۷	۲.۳۰۳	۰.۹۶۹	۰.۹۰۰	۲.۰۶۲	۰.۸۸۴	۰.۸۹۵	۲.۵۱۱	۰.۸۷۰	۰.۷۳۶	۲.۹۴۲	۰.۸۸۸	۰.۸۱۵	۱.۸۹۱	۰.۱	۱۰	۵	۴
۰.۹۰۹	۰.۷۷۷	۲.۱۱۸	۰.۹۶۸	۰.۸۱۵	۱.۷۷۱	۰.۸۹۲	۰.۸۲۳	۲.۳۹۰	۰.۸۷۳	۰.۷۴۲	۲.۰۲۱	۰.۸۸۸	۰.۷۵۰	۱.۶۶۵	۰.۳	۱۰	۵	۵
۰.۹۰۷	۰.۸۱۴	۲.۰۴۹	۰.۹۶۶	۰.۸۰۰	۱.۷۳۰	۰.۸۹۷	۰.۷۹۲	۲.۴۱۹	۰.۸۶۹	۰.۷۵۹	۱.۶۶۴	۰.۸۸۸	۰.۷۵۸	۱.۵۷۵	۰.۵	۱۰	۵	۶
۰.۹۲۱	۰.۸۲۱	۲.۳۶۹	۰.۹۹۵	۰.۸۸۳	۱.۶۶۷	۰.۹۰۵	۰.۸۹۱	۲.۴۷۷	۰.۸۸۱	۰.۶۷۸	۳.۱۷۵	۰.۹۰۷	۰.۸۴۱	۲	۰.۱	۱۵	۵	۷
۰.۹۱۲	۰.۷۷۳	۲.۱۹۶	۰.۹۹۵	۰.۷۷۲	۱.۴۸۳	۰.۹۱۶	۰.۷۸۵	۲.۳۸۷	۰.۸۸۰	۰.۶۶۵	۲.۰۷۷	۰.۹۰۷	۰.۷۵۵	۱.۷۲۱	۰.۳	۱۵	۵	۸
۰.۹۰۲	۰.۷۹۴	۲.۰۶۶	۰.۹۹۱	۰.۷۳۷	۱.۴۴۱	۰.۹۱۷	۰.۷۳۵	۲.۴۰۱	۰.۸۷۸	۰.۶۷۱	۱.۶۳۲	۰.۹۰۷	۰.۷۳۶	۱.۵۹۱	۰.۵	۱۵	۵	۹
۰.۹۰۲	۰.۸۰۹	۱.۸۲۰	۰.۹۷۰	۰.۶۹۱	۲.۲۰۲	۰.۹۰۷	۰.۷۵۲	۲.۸۰۴	۰.۸۹۲	۰.۷۹۳	۱.۵۱۶	۰.۸۹۸	۰.۸۰۳	۱.۸۸۸	۰.۱	۵	۱۰	۱۰
۰.۹۰۴	۰.۸۰۸	۱.۹۲۷	۰.۹۷۷	۰.۷۴۴	۱.۸۵۹	۰.۹۱۵	۰.۷۹۹	۲.۶۴۱	۰.۸۹۳	۰.۷۹۳	۱.۵۴۴	۰.۸۹۸	۰.۸۲۳	۱.۹۷۵	۰.۳	۵	۱۰	۱۱
۰.۹۲۱	۰.۸۴۴	۲.۰۴۹	۰.۹۷۶	۰.۷۹۷	۱.۷۵۱	۰.۹۱۶	۰.۸۵۷	۲.۶۰۳	۰.۸۸۷	۰.۸۱۸	۱.۶۶۱	۰.۸۹۸	۰.۸۴۶	۲.۰۹۰	۰.۵	۵	۱۰	۱۲
۰.۹۰۳	۰.۸۵۴	۲.۵۹۶	۰.۹۰۵	۰.۸۶۷	۳	۰.۹۰۶	۰.۸۸۴	۲.۸۷۹	۰.۹۰۸	۰.۸۳۳	۲.۵۰۹	۰.۸۹۷	۰.۸۳۸	۲.۲۶۳	۰.۱	۱۰	۱۰	۱۳
۰.۹۰۱	۰.۸۴۶	۲.۶۱۲	۰.۹۱۵	۰.۸۵۸	۲.۶۴۵	۰.۹۰۸	۰.۸۸۱	۲.۷۳۹	۰.۹۰۷	۰.۸۴۶	۲.۱۷۶	۰.۸۹۷	۰.۸۲۸	۲.۱۵۳	۰.۳	۱۰	۱۰	۱۴
۰.۹۱۰	۰.۸۶۲	۲.۶۹۱	۰.۹۱۹	۰.۸۷۲	۲.۵۷۳	۰.۹۱۰	۰.۸۸۲	۲.۷۶۶	۰.۹۰۶	۰.۸۴۷	۲.۰۷۱	۰.۸۹۷	۰.۸۲۸	۲.۱۷۴	۰.۵	۱۰	۱۰	۱۵
۰.۹۰۰	۰.۸۷۲	۲.۸۹۸	۰.۹۳۴	۰.۹۰۲	۲.۶۴۴	۰.۸۸۸	۰.۹۱۰	۲.۹۵۹	۰.۸۸۶	۰.۸۲۹	۲.۹۱۷	۰.۸۹۲	۰.۸۴۷	۲.۳۸۸	۰.۱	۱۵	۱۰	۱۶
۰.۸۹۷	۰.۸۵۱	۲.۹۲۸	۰.۹۳۷	۰.۸۹۰	۲.۴۳۴	۰.۸۹۲	۰.۸۹۰	۲.۸۳۳	۰.۸۸۶	۰.۸۱۸	۲.۴۴۱	۰.۸۹۲	۰.۸۲۵	۲.۲۵۹	۰.۳	۱۵	۱۰	۱۷
۰.۸۹۸	۰.۸۵۲	۲.۹۴۶	۰.۹۳۸	۰.۸۶۱	۲.۳۵۸	۰.۸۹۵	۰.۸۵۵	۲.۸۳۰	۰.۸۸۳	۰.۸۱۷	۲.۲۰۳	۰.۸۹۲	۰.۸۱۵	۲.۱۹۵	۰.۵	۱۵	۱۰	۱۸
۰.۹۰۳	۰.۸۱۸	۱.۸۰۷	۰.۹۹۲	۰.۶۶۰	۱.۸۱۲	۰.۹۰۱	۰.۷۱۷	۲.۸۴۹	۰.۸۸۳	۰.۷۵۹	۱.۳۶۷	۰.۸۹۷	۰.۸۶۰	۲.۲۷۴	۰.۱	۵	۱۵	۱۹
۰.۹۰۴	۰.۸۰۷	۱.۹۰۶	۰.۹۶۶	۰.۶۸۲	۱.۵۵۲	۰.۹۰۴	۰.۷۵۱	۲.۸۸۲	۰.۸۸۴	۰.۷۵۱	۱.۴۶۱	۰.۸۹۷	۰.۸۶۹	۲.۴۳۳	۰.۳	۵	۱۵	۲۰
۰.۹۱۰	۰.۸۴۲	۲.۰۷۵	۰.۹۹۴	۰.۷۴۰	۱.۴۸۷	۰.۹۰۵	۰.۸۲۰	۲.۶۲۸	۰.۸۸۱	۰.۷۶۹	۱.۶۵۳	۰.۸۹۷	۰.۸۹۲	۲.۶۱۸	۰.۵	۵	۱۵	۲۱
۰.۸۹۴	۰.۸۵۸	۲.۶۷۳	۰.۹۴۰	۰.۸۰۳	۲.۶۷۲	۰.۸۷۴	۰.۸۳۷	۳.۰۶۱	۰.۸۸۸	۰.۸۳۹	۲.۳۲۰	۰.۸۹۰	۰.۸۶۱	۲.۵۹۲	۰.۱	۱۰	۱۵	۲۲
۰.۸۹۶	۰.۸۵۹	۲.۷۹۱	۰.۹۴۷	۰.۸۴۳	۲.۳۴۴	۰.۸۷۶	۰.۸۷۳	۲.۹۳۹	۰.۸۸۶	۰.۸۴۵	۲.۲۰۴	۰.۸۹۰	۰.۸۶۸	۲.۵۸۹	۰.۳	۱۰	۱۵	۲۳
۰.۸۹۳	۰.۸۷۷	۲.۹۳۸	۰.۹۴۷	۰.۸۵۷	۲.۳۶۴	۰.۸۷۸	۰.۸۸۸	۲.۹۳۳	۰.۸۸۸	۰.۸۴۷	۲.۱۹۵	۰.۸۹۰	۰.۸۷۱	۲.۶۳۹	۰.۵	۱۰	۱۵	۲۴
۰.۹۰۱	۰.۸۷۶	۳.۲۸۹	۰.۹۱۰	۰.۸۹۵	۳.۳۷۵	۰.۹۰۴	۰.۹۰۸	۳.۲۸۲	۰.۹۰۶	۰.۸۵۲	۲.۸۷۲	۰.۹۱۵	۰.۸۷۴	۲.۷۴۴	۰.۱	۱۵	۱۵	۲۵
۰.۹۰۱	۰.۸۷۱	۳.۲۶۱	۰.۹۱۲	۰.۸۷۹	۳.۰۶۵	۰.۹۰۴	۰.۸۹۶	۳.۱۳۳	۰.۹۰۵	۰.۸۵۵	۲.۵۸۴	۰.۹۱۵	۰.۸۵۷	۲.۷۷۴	۰.۳	۱۵	۱۵	۲۶
۰.۸۹۶	۰.۸۶۹	۳.۳۸۵	۰.۹۱۳	۰.۸۸۲	۳.۰۱۸	۰.۹۰۴	۰.۸۸۸	۳.۱۴۰	۰.۹۰۶	۰.۸۴۲	۲.۴۷۲	۰.۹۱۵	۰.۸۵۰	۲.۶۶۰	۰.۵	۱۵	۱۵	۲۷

۷۲ مقایسه R براساس دو طرح نمونه‌گیری در خانواده نرخ خطر معکوس متناسب

و ۰.۶۶۷ می‌باشد، که مطابقت توزیع نمایی تعمیم یافته را نشان می‌دهد. لازم به ذکر است هر هواپیمای دیگری که بر مدل در نظر گرفته شده مطابقت داشته باشد را می‌توان برای مطالعه انتخاب نمود. مقادیر $RRSS$ پایین این مجموعه داده‌ها براساس دیاگرام ۲.۱ در جدول ۷.۴ ارائه شده‌اند.

جدول ۷.۴: مقادیر $RRSS$ پایین استخراج شده از داده‌های هواپیما

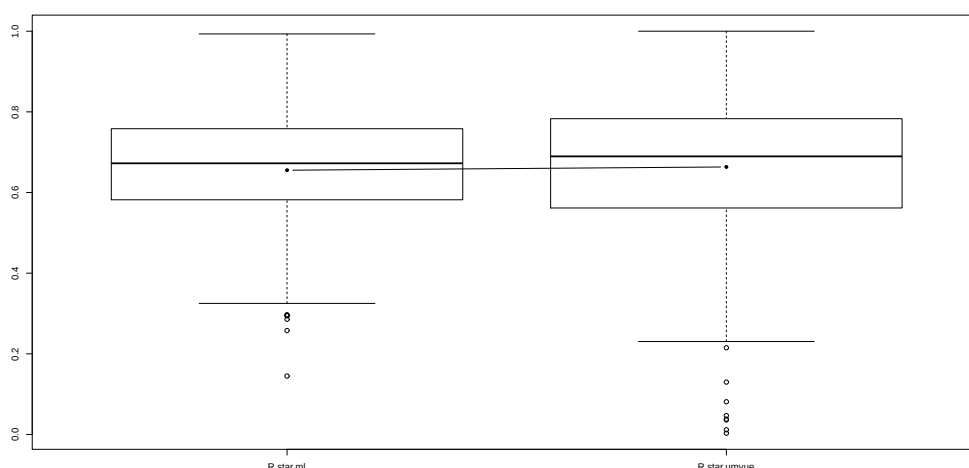
هواپیما												
۸۰۴۵	۸۰۴۴	۷۹۱۷	۷۹۱۶	۷۹۱۵	۷۹۱۴	۷۹۱۳	۷۹۱۲	۷۹۱۱	۷۹۱۰	۷۹۰۹	۷۹۰۸	۷۹۰۷
۱۰۲	۴۸۷	۱۳۰	۵۰	۳۵۹	۵۰	۹۷	۲۳	۵۵	۷۴	۹۰	۴۱۳	۱۹۴
۱۴	۷		۵	۳	۲۲	۱۱	۸۷	۵۶	۴۸		۹	
۳۲	۳		۱۲				۳	۳۳			۳۴	
۱۴							۱				۷	

جدول ۷.۴ نشان می‌دهد که تعداد $RRSS$ پایین برای دو هواپیمای بوئینگ ۷۹۰۸ و ۸۰۴۵ برابر ۴ می‌باشد. برای مطالعه $R = P(X < Y)$ ، که در این مثال هدف ما مقایسه سیستم تهویه دو هواپیمای بوئینگ ۷۹۰۸ و ۸۰۴۵ است، فرض کنید X ، متغیر فشار باشد که $u = (413, 9, 34, 7)$ مقادیر $RRSS$ پایین هواپیمای شماره ۷۹۰۸ می‌باشد و Y متغیر مقاومت باشد که $s = (102, 14, 32, 14)$ ($n = m = 4$) مقادیر $RRSS$ پایین از هواپیمای شماره ۸۰۴۵ می‌باشد. برای به دست آوردن مشاهدات بیشتر به منظور بررسی عملکرد برآوردهای نقطه‌ای و فاصله‌ای، با استفاده از روش بوت‌استرپ، باز نمونه‌گیری از داده‌ها صورت می‌گیرد. بدین منظور ۲۰۰ نمونه‌ی بوت‌استرپی از داده‌های اصلی تولید می‌شود و سپس مقادیر $RRSS$ پایین استخراج می‌گردد. براساس $RRSS$ های به دست آمده، برآوردهای نقطه‌ای پارامتر تنش-مقاومت به صورت $\hat{R}_{ML} = 0.655$ و $\hat{R}_{UMVU} = 0.663$ به دست می‌آید. نمودار جعبه‌ای متناظر با مشاهدات مربوط به برآوردهای نقطه‌ای در شکل ۷.۴ رسم شده‌است. برطبق مقادیر برآوردهای نقطه‌ای که حدوداً برابر ۰.۶ می‌باشد، انتظار داریم $\hat{R}_{ML}$ کارایی بهتری نسبت به $\hat{R}_{UMVU}$ داشته‌باشد. با توجه به شکل ۷.۴ می‌بینیم دامنه میان چارکی در $\hat{R}_{ML}$ کمتر از $\hat{R}_{UMVU}$ می‌باشد، که نتایج قبلی را تأیید می‌کند. برای بررسی فواصل اطمینان، فرآیند ذکر شده در بالا را ۱۰۰۰ بار دیگر تکرار می‌کنیم. توزیع‌های $Gamma(2, 1)$ و $Gamma(2, 2)$ به ترتیب به عنوان توزیع پیشین α و β در نظر گرفته شده‌است و میانگین طول فواصل اطمینان در جدول ۸.۴ نمایش داده شده‌است.

جدول ۸.۴: مقادیر متوسط طول R براساس مجموعه داده‌ی واقعی

HPD	Bayes	AMLE	MLE	Basic	Boot - t2	Boot - t1	Perc	γ
۰.۳۵۹	۰.۳۸۵	۰.۳۹۶	۰.۳۸۹	۰.۵۷۸	۰.۵۹۸	۰.۸۹۴	۰.۵۹۳	۰.۰۵
۰.۳۰۶	۰.۳۲۷	۰.۳۳۲	۰.۳۲۹	۰.۴۸۷	۰.۴۳۱	۰.۶۸۸	۰.۴۹۴	۰.۱

برطبق نتایج جدول ۸.۴، مشاهده می‌شود همانند بخش قبل متوسط طول $Boot - t1$ کم‌تر از $Boot - t2$ می‌باشد. همچنین، میانگین طول MLE و $AMLE$ بهتر از سایر فواصل اطمینان



شکل ۷.۴: مقایسه $\hat{R}_{ML}$ و $\hat{R}_{UMVU}$.

می باشد. به طور دقیق، فاصله اطمینان براساس کمیت محور دارای کمترین متوسط طول می باشد. همچنین فاصله HPD نیز بهتر از فاصله معتبر بیزی عمل می کند. در نتیجه نتایج این بخش کاملاً نتایج بخش قبل را تأیید می کند.

نتیجه گیری

در این فصل، برآوردگرهای ML ، $UMVU$ و همچنین برآوردگر بیزی پارامتر تنش-مقاومت R را براساس $RRSS$ پایین از مدل $PRHR$ به دست آوردیم. برآوردگرهای ML و $UMVU$ براساس معیار کارایی نسبی مقایسه شدند. مشاهده شد زمانی که R نزدیک به مقدار ۰.۵ می باشد، برآوردگر ML بهتر از برآوردگر $UMVU$ عمل می کند و برای مقادیر فرین R عملکرد برآوردگر $UMVU$ بهتر می باشد. همچنین در این فصل برآوردگرهای نقطه ای حاصل از دو طرح $RRSS$ و رکوردهای معمولی را براساس معیار کارایی نسبی مقایسه نمودیم. مشاهده شد برآوردگر R براساس طرح $RRSS$ بهتر از رکوردهای معمولی عمل می کند.

فواصل اطمینان مختلفی محاسبه گردید و براساس شبیه سازی مقایسه شدند. براساس مطالعات شبیه سازی و همچنین مثال عددی، مشاهده شد فاصله اطمینان صدکی و دقیق براساس ML بهتر از سایر فواصل اطمینان می باشد. همچنین فاصله اطمینان HPD عملکرد بهتری نسبت به فاصله اطمینان بیزی دارد. در نهایت برآوردگرهای فاصله ای را براساس دو طرح رکوردهای معمولی و $RRSS$ با در نظر گرفتن توزیع نمایی تعمیم یافته به عنوان توزیعی از خانواده $PRHR$ ، مقایسه نمودیم. مشاهده کردیم فواصل اطمینان براساس طرح $RRSS$ بهتر از فواصل مشابه براساس رکوردهای معمولی عمل می کنند. همچنین ثابت کردیم که تمام محاسبات به توزیع پایه در مدل $PRHR$ بستگی ندارد.

فصل ۵

برآورد پارامتر تنش-مقاومت براساس $RRSS$ در توزیع نمایی دو پارامتری

در این فصل، با در نظر گرفتن طرح $RRSS$ و با فرض این که متغیرهای تصادفی X و Y مستقل و هم توزیع با توزیع نمایی دو پارامتری باشند، به برآورد نقطه‌ای پارامتر تنش-مقاومت می پردازیم. مانیشا^۱ و همکاران (۲۰۰۵)، با فرض توزیع نمایی دو پارامتری برآوردگر نااریب با کمترین واریانس پارامتر تنش-مقاومت را محاسبه نمودند. همچنین بکلیزی در سال‌های (۲۰۰۸a) و (۲۰۱۴)، با استفاده از مشاهدات رکوردی، برآوردهای نقطه‌ای و فاصله‌ای پارامتر تنش-مقاومت را با در نظر گرفتن توزیع نمایی دو پارامتری به دست آورد. در ادامه و در بخش ۱.۵ با در نظر گرفتن حالت‌های مختلف و رکوردهای بالای نمونه‌ی مجموعه‌ی رتبه‌دار، برآوردگر درست‌نمایی ماکسیمم پارامتر تنش-مقاومت را به دست می آوریم و در بخش ۲.۵ نیز برآوردگر نااریب با کمترین واریانس پارامتر فوق را محاسبه می کنیم.

^۱Manisha

۱.۵ برآوردگر درستنمایی ماکسیمم

در این بخش، به محاسبه برآوردگر درستنمایی ماکسیمم پارامتر تنش-مقاومت با در نظر گرفتن توزیع نمایی دو پارامتری و طرح $RRSS$ می پردازیم. تابع چگالی توزیع نمایی دوپارامتری به صورت زیر می باشد

$$f(x|\mu, \theta) = \frac{1}{\theta} e^{-\frac{1}{\theta}(x-\mu)}, \quad x > \mu, \quad -\infty < \mu < \infty, \quad \theta > 0 \quad (1.5)$$

که در آن μ و θ به ترتیب پارامتر مکان و مقیاس می باشند. فرض کنید X و Y به ترتیب متغیرهای تصادفی مستقل از توزیع نمایی دو پارامتری باشند که به اختصار با نماد $E(\mu_1, \theta_1)$ و $E(\mu_2, \theta_2)$ نشان می دهیم. همچنین فرض می کنیم $\mathbf{r} = (r_{1,1}, r_{2,2}, \dots, r_{n,n})^T$ بردار مشاهدات $\mathbf{R} = (R_{1,1}, R_{2,2}, \dots, R_{n,n})^T$ نمونه ای به حجم n از رکوردهای بالای نمونه ای مجموعه ی رتبه دار با توزیع $E(\mu_1, \theta_1)$ می باشد و $\mathbf{s} = (s_{1,1}, s_{2,2}, \dots, s_{m,m})^T$ نیز بردار مشاهدات $\mathbf{S} = (S_{1,1}, S_{2,2}, \dots, S_{m,m})^T$ نمونه ای به حجم m از رکوردهای بالای مجموعه ی رتبه دار با توزیع $E(\mu_2, \theta_2)$ تنها اطلاعات در خصوص جامعه موردنظر می باشد.

لم ۱.۱.۵. فرض کنید $X \sim E(\mu_1, \theta_1)$ و $Y \sim E(\mu_2, \theta_2)$ دو متغیر تصادفی مستقل باشند و $\mathcal{R} = Pr(X < Y)$ همان پارامتر تنش-مقاومت باشد، در این صورت به راحتی ثابت می شود که

$$\hat{\mathcal{R}}_{ML} = \begin{cases} 1 - \frac{\theta_1}{\theta_1 + \theta_2} e^{-\frac{1}{\theta_1}(\mu_2 - \mu_1)} & \mu_2 - \mu_1 \geq 0 \\ \frac{\theta_2}{\theta_1 + \theta_2} e^{-\frac{1}{\theta_2}(\mu_1 - \mu_2)} & \mu_2 - \mu_1 < 0 \end{cases} \quad (2.5)$$

برهان. در حالت اول با در نظر گرفتن اشتراک نواحی

$$۱) \mu_2 - \mu_1 > 0, \quad ۲) X < Y, \quad ۳) X > \mu_1, \quad ۴) Y > \mu_2$$

احتمال مورد نظر برابر است با

$$Pr(X < Y) = \int_{\mu_2}^{\infty} \int_{\mu_1}^y \frac{1}{\theta_1} e^{-\frac{1}{\theta_1}(x-\mu_1)} \frac{1}{\theta_2} e^{-\frac{1}{\theta_2}(y-\mu_2)} dx dy$$

که با محاسبه انتگرال فوق نتیجه مورد نظر حاصل می شود. در حالت دوم به طریقی مشابه و با اشتراک ناحیه های

$$۱) \mu_2 - \mu_1 < 0, \quad ۲) X < Y, \quad ۳) X > \mu_1, \quad ۴) Y > \mu_2$$

□

نتیجه به دست می آید.

همان‌طور که در بخش ۶.۱ مشاهده کردیم، واحدهای $RRSS$ مستقل از یکدیگر می‌باشند، لذا با توجه به اینکه $X \sim E(\mu_1, \theta_1)$ و $Y \sim E(\mu_2, \theta_2)$ با جایگذاری تابع چگالی و تابع توزیع X و Y در (۸.۱) تابع درستنمایی توأم به صورت زیر به دست می‌آید:

$$L(\theta_1, \mu_1, \theta_2, \mu_2) = \frac{(\frac{1}{\theta_1})^{n(n+1)/2} \prod_{i=1}^n (r_{i,i} - \mu_1)^{i-1} e^{-\frac{1}{\theta_1} \sum_{i=1}^n (r_{i,i} - \mu_1)}}{\prod_{i=1}^n (i-1)!} \times \frac{(\frac{1}{\theta_2})^{m(m+1)/2} \prod_{j=1}^m (s_{j,j} - \mu_2)^{j-1} e^{-\frac{1}{\theta_2} \sum_{j=1}^m (s_{j,j} - \mu_2)}}{\prod_{j=1}^m (j-1)!} \quad (3.5)$$

و

$$\mu_1 = \min\{r_{i,i} \mid 1 < i < n\} \quad \text{و} \quad \mu_2 = \min\{s_{j,j} \mid 1 < j < m\} \quad (4.5)$$

حال با مشتق‌گیری از لگاریتم تابع درستنمایی توأم نسبت به پارامترهای $(\mu_1, \theta_1, \mu_2, \theta_2)$ به ازای r و s داده شده برآوردگرهای درستنمایی ماکسیمم پارامترها به صورت زیر به دست می‌آید.

$$\hat{\mu}_1 = \min\{r_{i,i} \mid 1 < i < n\} \quad \text{و} \quad \hat{\theta}_1 = \frac{\sum_{i=1}^n (r_{i,i} - \hat{\mu}_1)}{N} \quad (5.5)$$

و

$$\hat{\mu}_2 = \min\{s_{j,j} \mid 1 < j < m\} \quad \text{و} \quad \hat{\theta}_2 = \frac{\sum_{j=1}^m (s_{j,j} - \hat{\mu}_2)}{M} \quad (6.5)$$

که در آن $N = \frac{n(n+1)}{2}$ و $M = \frac{m(m+1)}{2}$.

در نتیجه با استفاده از (۲.۵) و ویژگی پایایی برآوردگرهای درستنمایی ماکسیمم، برآوردگر $\hat{R}_{ML}$ به فرم زیر به دست می‌آید.

$$\hat{R}_{ML} = \begin{cases} 1 - \frac{\hat{\theta}_1}{\hat{\theta}_1 + \hat{\theta}_2} e^{-\frac{1}{\hat{\theta}_1}(\hat{\mu}_2 - \hat{\mu}_1)} & \hat{\mu}_2 - \hat{\mu}_1 \geq 0 \\ \frac{\hat{\theta}_2}{\hat{\theta}_1 + \hat{\theta}_2} e^{-\frac{1}{\hat{\theta}_2}(\hat{\mu}_1 - \hat{\mu}_2)} & \hat{\mu}_2 - \hat{\mu}_1 < 0 \end{cases} \quad (7.5)$$

برای بررسی بیشتر دو حالت زیر را در نظر می‌گیریم:

پارامتر مکان در دو جامعه برابر باشند:

در حالتی که $\mu_1 = \mu_2 = \mu$ با استفاده از روابط (۳.۵) و (۴.۵)، برآورد درستنمایی ماکسیمم پارامترهای $(\mu, \theta_1, \theta_2)$ براساس $RRSS$ بالا به صورت زیر به دست می‌آید

$$\hat{\mu} = \min\{r_{i,i}, s_{j,j}, \quad 1 \leq i \leq n \quad 1 \leq j \leq m\} \quad (8.5)$$

$$\hat{\theta}_1 = \frac{1}{N} \sum_{i=1}^n (r_{i,i} - \hat{\mu}) \quad (9.5)$$

$$\hat{\theta}_2 = \frac{1}{M} \sum_{j=1}^m (s_{j,j} - \hat{\mu}) \quad (10.5)$$

بنابراین با توجه به اینکه در این حالت $\mathcal{R} = Pr(X < Y) = \frac{\theta_2}{\theta_1 + \theta_2}$ در نتیجه

$$\hat{\mathcal{R}}_{ML} = \frac{\hat{\theta}_2}{\hat{\theta}_1 + \hat{\theta}_2} \quad (11.5)$$

پارامترهای مقیاس دو جامعه برابر باشند:

در صورتی که فرض کنیم $\theta_1 = \theta_2 = \theta$ باشد و $\mathbf{r} = (r_{1,1}, r_{2,2}, \dots, r_{n,n})^T$ بردار مشاهدات $RRSS$ بالا به حجم n از توزیع $E(\mu_1, \theta)$ و $\mathbf{s} = (s_{1,1}, s_{2,2}, \dots, s_{m,m})^T$ بردار مشاهدات $RRSS$ بالا به حجم m از توزیع $E(\mu_2, \theta)$ باشد، برآورد درستنمایی ماکسیمم پارامترهای (μ_1, μ_2, θ) براساس $RRSS$ بالا به صورت زیر به دست می آید

$$\hat{\mu}_1 = \min\{r_{i,i}, 1 \leq i \leq n\} \quad (12.5)$$

$$\hat{\mu}_2 = \min\{s_{j,j}, 1 \leq j \leq m\} \quad (13.5)$$

$$\hat{\theta} = \frac{\sum_{i=1}^n (r_{i,i} - \hat{\mu}_1) + \sum_{j=1}^m (s_{j,j} - \hat{\mu}_2)}{N + M} \quad (14.5)$$

و به سادگی نتیجه می شود

$$\hat{\mathcal{R}}_{ML} = \begin{cases} 1 - \frac{1}{2} e^{-\frac{1}{\hat{\theta}}(\hat{\mu}_2 - \hat{\mu}_1)} & \hat{\mu}_2 - \hat{\mu}_1 \geq 0 \\ \frac{1}{2} e^{-\frac{1}{\hat{\theta}}(\hat{\mu}_2 - \hat{\mu}_1)} & \hat{\mu}_2 - \hat{\mu}_1 < 0 \end{cases} \quad (15.5)$$

۲.۵ برآورد نااریب با کمترین واریانس

در این بخش برآوردگر نااریب با کمترین واریانس، برای پارامتر تنش-مقاومت $\mathcal{R}$ را در حالتی که پارامترهای مکان در دو جامعه برابر باشند را بررسی می کنیم. با فرض این که $\mu_1 = \mu_2 = \mu$ که μ مقداری معلوم می باشد، می توان نشان داد که $\left(\sum_{i=1}^n R_{i,i}, \sum_{j=1}^m S_{j,j} \right)$ آماره بسنده کامل برای $(\theta_1, \theta_2)^T$ می باشد. می دانیم $R_{1,1} \stackrel{d}{=} X$ و $S_{1,1} \stackrel{d}{=} Y$ ، در این صورت $I(R_{1,1} < S_{1,1})$ یک برآوردگر نااریب برای پارامتر $\mathcal{R}$ می باشد. در نتیجه با استفاده از قضیه راثو-بلکول و لهن-شفه برآوردگر $UMVU$ پارامتر $\mathcal{R}$ را با در نظر گرفتن سه حالت زیر محاسبه می کنیم.

حالت اول: $\min\{m, n\} \geq 2$

$$\begin{aligned}\hat{\mathcal{R}}_{UMVU} &= E \left\{ I(R_{1,1} < S_{1,1}) \left| \sum_{i=1}^n (R_{i,i} - \mu), \sum_{j=1}^m (S_{j,j} - \mu) \right. \right\} \\ &= Pr \left\{ W < \frac{\sum_{j=1}^m (S_{j,j} - \mu)}{\sum_{i=1}^n (R_{i,i} - \mu)} \left| \sum_{i=1}^n (R_{i,i} - \mu), \sum_{j=1}^m (S_{j,j} - \mu) \right. \right\}\end{aligned}$$

که در آن $W = \frac{\sum_{j=1}^m (S_{j,j} - \mu)}{\sum_{i=1}^n (R_{i,i} - \mu)} \stackrel{d}{=} \frac{B(1, N-1)}{B(1, M-1)}$ آماره‌ی فرعی است و بنا به قضیه باسو،

مستقل از آماره‌ی بسنده‌ی کامل می‌باشد. بنابراین اگر $F_W(\cdot)$ بیان‌گر تابع توزیع احتمال متغیر تصادفی W باشد، آن‌گاه

$$\hat{\mathcal{R}}_{UMVU} = F_W \left(\frac{\sum_{j=1}^m (S_{j,j} - \mu)}{\sum_{i=1}^n (R_{i,i} - \mu)} \right) \quad (16.5)$$

حالت دوم: ($n=1$ و $m \geq 2$)

واضح است در این حالت $\left(R_{1,1}, \sum_{j=1}^m S_{j,j} \right)$ آماره‌ی بسنده‌ی کامل برای $(\theta_1, \theta_2)^T$ می‌باشد بنابراین همانند حالت قبل، برآوردگر $UMVU$ برای $\mathcal{R}$ عبارت‌است از

$$\hat{\mathcal{R}}_{UMVU} = Pr \left\{ W < \frac{\sum_{j=1}^m (S_{j,j} - \mu)}{(R_{1,1} - \mu)} \left| (R_{1,1} - \mu), \sum_{j=1}^m (S_{j,j} - \mu) \right. \right\}$$

که در آن $W = \frac{\sum_{j=1}^m (S_{j,j} - \mu)}{R_{1,1} - \mu} \stackrel{d}{=} \frac{1}{B(1, M-1)}$ آماره فرعی می‌باشد و بنا به قضیه باسو مستقل از آماره‌ی بسنده‌ی کامل می‌باشد بنابراین

$$\hat{\mathcal{R}}_{UMVU} = \left(1 - \frac{R_{1,1} - \mu}{\sum_{j=1}^m (S_{j,j} - \mu)} \right)^{M-1} \quad (17.5)$$

حالت سوم: ($n \geq 2$ و $m = 1$)

در این حالت $\left(\sum_{i=1}^n R_{i,1}, S_{1,1}\right)$ آماره‌ی بسنده‌ی کامل و همانند حالت‌های قبل $UMVU$ برای $\mathcal{R}$ برابر است با

$$\hat{\mathcal{R}}_{UMVU} = Pr \left\{ W < \frac{(S_{1,1} - \mu)}{\sum_{i=1}^n (R_{i,i} - \mu)} \middle| \sum_{i=1}^n (R_{i,i} - \mu), (S_{1,1} - \mu) \right\}$$

که در آن $W = \frac{(R_{1,1} - \mu)}{\sum_{i=1}^n (R_{i,i} - \mu)} \stackrel{d}{=} \frac{1}{B(1, N-1)}$ آماره فرعی و مستقل از آماره‌ی بسنده‌ی کامل

می‌باشد بنابراین

$$\hat{\mathcal{R}}_{UMVUE} = 1 - \left(1 - \frac{S_{1,1} - \mu}{\sum_{i=1}^n (R_{i,i} - \mu)}\right)^{N-1} \quad (18.5)$$

که $B(.,.)$ تابع بتا می‌باشد.

۱.۲.۵ شبیه‌سازی

در این بخش به مقایسه عددی برآوردگرهای نقطه‌ای معرفی شده در بخش‌های ۱.۵ و ۲.۵ می‌پردازیم و در آن تمام ترکیب‌های $n = (2, 4, 6)$ و $m = (2, 4, 6)$ و پارامتر μ را معلوم در نظر می‌گیریم. همچنین $\theta_1 = 1$ و θ_2 را طوری در نظر گرفتیم که براساس رابطه (۱۱.۵)، داشته باشیم $\mathcal{R} = (0/1, 0/3, 0/5, 0/95, 0/99)$. به ازای هر کدام از ترکیب‌های فوق، ۱۰۰۰۰ نمونه‌ی s و r از توزیع‌های $E(\mu_1, \theta_1)$ و $E(\mu_2, \theta_2)$ تولید کرده‌ایم. نتایج حاصل در جدول ۱.۵ آمده است. به توجه به نتایج ارائه شده در جدول ۱.۵، مشاهده می‌کنیم نقاطی از $\mathcal{R}$ که نزدیک ۰.۵ هستند برآوردگر $UMVU$ خطای بیشتری نسبت به برآوردگر ML دارد. همچنین همان‌گونه که انتظار داشتیم با افزایش حجم نمونه، خطا کاهش یافته است.

نتیجه‌گیری

در این فصل، با فرض این که دو نمونه مستقل $RRSS$ بالا، به اندازه n و m از توزیع نمایی دو پارامتری، تنها اطلاعات از جوامع X و Y می‌باشد، برآوردگرهای نقطه‌ای پارامتر تنش-مقاومت را با در نظر گرفتن حالت‌های مختلف بررسی نمودیم. در نهایت از طریق شبیه‌سازی، برآوردگرهای فوق را مقایسه نمودیم. مشاهده گردید در حالتی که $\mathcal{R} = 0/5$ می‌باشد برآوردگر ML بهتر از برآوردگر $UMVU$ عمل می‌کند.

جدول ۱.۵: نتایج مربوط به MLE و $UMVUE$

ردیف	$\mathcal{R}$	m	n	Bias MLE	MSE MLE	Var UMVUE
۱	۰.۱۰۰	۲	۲	۰.۰۱۸-	۰.۰۱۰	۰.۰۰۹
۲	۰.۱۰۰	۴	۲	۰.۰۷۱-	۰.۰۰۷	۰.۰۰۸
۳	۰.۱۰۰	۶	۲	۰.۰۸۵-	۰.۰۰۸	۰.۰۰۹
۴	۰.۱۰۰	۲	۴	۰.۰۵۰	۰.۰۰۸	۰.۰۰۸
۵	۰.۱۰۰	۴	۴	۰.۰۱۵-	۰.۰۰۲	۰.۰۰۲
۶	۰.۱۰۰	۶	۴	۰.۰۵۷-	۰.۰۰۴	۰.۰۰۴
۷	۰.۱۰۰	۲	۶	۰.۱۱۶	۰.۰۱۹	۰.۰۲۱
۸	۰.۱۰۰	۴	۶	۰.۰۱۰	۰.۰۰۲	۰.۰۰۲
۹	۰.۱۰۰	۶	۶	۰.۰۱۲-	۰.۰۰۱	۰.۰۰۱
۱۰	۰.۳۰۰	۲	۲	۰.۰۲۴-	۰.۰۴۷	۰.۰۵۶
۱۱	۰.۳۰۰	۴	۲	۰.۱۷۱-	۰.۰۵۱	۰.۰۵۸
۱۲	۰.۳۰۰	۶	۲	۰.۲۲۵-	۰.۰۶۲	۰.۰۶۸
۱۳	۰.۳۰۰	۲	۴	۰.۱۰۴	۰.۰۳۱	۰.۰۴۰
۱۴	۰.۳۰۰	۴	۴	۰.۰۲۱-	۰.۰۱۲	۰.۰۱۳
۱۵	۰.۳۰۰	۶	۴	۰.۱۳۹-	۰.۰۲۵	۰.۰۲۸
۱۶	۰.۳۰۰	۲	۶	۰.۲۱۲	۰.۰۵۷	۰.۰۷۷
۱۷	۰.۳۰۰	۴	۶	۰.۰۲۷	۰.۰۰۸	۰.۰۰۹
۱۸	۰.۳۰۰	۶	۶	۰.۰۱۶-	۰.۰۰۵	۰.۰۰۶
۱۹	۰.۵۰۰	۲	۲	۰.۰۰۳-	۰.۰۶۷	۰.۰۸۲
۲۰	۰.۵۰۰	۴	۲	۰.۲۱۱-	۰.۰۹۶	۰.۱۱۳
۲۱	۰.۵۰۰	۶	۲	۰.۳۱۲-	۰.۱۳۳	۰.۱۵۱
۲۲	۰.۵۰۰	۲	۴	۰.۱۲۴	۰.۰۳۵	۰.۰۵۰
۲۳	۰.۵۰۰	۴	۴	۰.۰۰۰	۰.۰۱۸	۰.۰۱۹
۲۴	۰.۵۰۰	۶	۴	۰.۱۶۶-	۰.۰۴۲	۰.۰۴۶
۲۵	۰.۵۰۰	۲	۶	۰.۲۱۴	۰.۰۵۵	۰.۰۷۸
۲۶	۰.۵۰۰	۴	۶	۰.۰۴۳	۰.۰۱۱	۰.۰۱۳
۲۷	۰.۵۰۰	۶	۶	۰.۰۰۰	۰.۰۰۸	۰.۰۰۸
۲۸	۰.۹۵۰	۲	۲	۰.۰۱۰	۰.۰۰۳	۰.۰۰۳
۲۹	۰.۹۵۰	۴	۲	۰.۰۵۸-	۰.۰۱۶	۰.۰۱۷
۳۰	۰.۹۵۰	۶	۲	۰.۱۳۶-	۰.۰۴۴	۰.۰۵۰
۳۱	۰.۹۵۰	۲	۴	۰.۰۲۷	۰.۰۰۱	۰.۰۰۱
۳۲	۰.۹۵۰	۴	۴	۰.۰۰۹	۰.۰۰۱	۰.۰۰۱
۳۳	۰.۹۵۰	۶	۴	۰.۰۳۲-	۰.۰۰۳	۰.۰۰۳
۳۴	۰.۹۵۰	۲	۶	۰.۰۳۴	۰.۰۰۱	۰.۰۰۲
۳۵	۰.۹۵۰	۴	۶	۰.۰۱۵	۰.۰۰۰	۰.۰۰۱
۳۶	۰.۹۵۰	۶	۶	۰.۰۰۷	۰.۰۰۰	۰.۰۰۰
۳۷	۰.۹۹۰	۲	۲	۰.۰۰۲	۰.۰۰۰	۰.۰۰۰
۳۸	۰.۹۹۰	۴	۲	۰.۰۱۴-	۰.۰۰۱	۰.۰۰۱
۳۹	۰.۹۹۰	۶	۲	۰.۰۳۸-	۰.۰۰۵	۰.۰۰۵
۴۰	۰.۹۹۰	۲	۴	۰.۰۰۶	۰.۰۰۰	۰.۰۰۰
۴۱	۰.۹۹۰	۴	۴	۰.۰۰۲	۰.۰۰۰	۰.۰۰۰
۴۲	۰.۹۹۰	۶	۴	۰.۰۰۷-	۰.۰۰۰	۰.۰۰۰
۴۳	۰.۹۹۰	۲	۶	۰.۰۰۷	۰.۰۰۰	۰.۰۰۰
۴۴	۰.۹۹۰	۴	۶	۰.۰۰۳	۰.۰۰۰	۰.۰۰۰
۴۵	۰.۹۹۰	۶	۶	۰.۰۰۲	۰.۰۰۰	۰.۰۰۰

آینده تحقیق

در این رساله، پارامتر تنش-مقاومت را براساس نمونه‌های به‌دست آمده از طرح $RRSS$ بررسی نمودیم. الگوی نرخ خطر معکوس متناسب که شامل تعداد زیادی از توزیع‌های مهم آماری می‌باشد به عنوان توزیع جامعه‌ی آماری در نظر گرفته شد و به بررسی پارامتر تنش-مقاومت براساس طرح $RRSS$ پرداختیم. نتایج به‌دست آمده از طرح فوق و آماره‌های رکوردی معمولی را مقایسه نمودیم. این تحقیق می‌تواند در موارد زیر نیز مورد بررسی و مطالعه قرار گیرد.

- در این رساله، طرح $RRSS$ ، در حقیقت براساس ساختار نمونه‌گیری مجموعه‌ی رتبه‌دار تک حلقه‌ای می‌باشد، به نظر می‌رسد می‌توان نتایج فوق را براساس روش نمونه‌گیری مجموعه‌ی رتبه‌دار رکوردی چند حلقه‌ای ادامه داد.
- طرح $RRSS$ را می‌توان از دیدگاه نظریه اطلاع مانند آنتروپی شانون و رنی و معیارهای فاصله از قبیل کولبک نیز بررسی نمود.
- برآورد چندک‌های جامعه براساس داده‌های ترتیبی حاصل از طرح $RRSS$ می‌تواند مورد بررسی محققان واقع شود.
- در این رساله پارامتر تنش-مقاومت با فرض استقلال متغیرهای X و Y می‌باشد، می‌توان با در نظر گرفتن فرض وابستگی متغیرهای X و Y نتایج فوق را بررسی نمود.

مراجع

- [۱] بصیرت، م. (۱۳۹۳ ش). استنباط آماری برای پارامتر تنش-مقاومت بر اساس داده‌های ترتیبی، رساله ی دکتری، دانشگاه فردوسی مشهد.
- [۲] رزمخواه، م.، احمدی، ج.، خطیب آستانه، ب. (۱۳۸۶ ش). مقایسه ی دوروش استخراج رکوردها از دیدگاه اطلاع فیشر، مجله ی علوم آماری، جلد ۱، شماره ی ۴۴، ۱.
- [۳] صالحی، م. (۱۳۹۰ ش). نمونه گیری مجموعه رتبه دار: برآورد و پیش بینی، پایان نامه کارشناسی ارشد، دانشگاه فردوسی مشهد.
- [۴] صالحی، م. (۱۳۹۴). نمونه ی مجموعه ی رتبه دار بر مبنای داده های رکوردی: پیش بینی و برآورد، رساله ی دکتری، دانشگاه فردوسی مشهد.
- [5] Ahmadi, J., Arghami, N.R. (2001), "Some univariate stochastic orders on record values", Communication in Statistics: Theory and Methods, 30, 69-74.
- [6] Arnold, B. C., Balakrishnan, N. and Nagaraja, H. N. (1998), "Records", John Wiley and Sons, New York.
- [7] Ahmadi, J., Jafari Jozani, M., Marchand, E., and Parsian, A. (2009), "Bayesian estimation based on k-record data from a general class of distributions under balanced type loss functions", Journal of Statistical Planning and Inference, 139, 1180-1189.
- [8] Ahmadi, J. and Balakrishnan, N. (2010), "Prediction of order statistics and record values from two independent sequences", Statistics, 44, 417-430.
- [9] Ahmadi, J., Jafari Jozani, M., Marchand, E. and Parsian, A. (2008), "Prediction of k-records from a general class of distributions under balanced type loss functions", Metrika, 70, 19-33
- [10] Ahmadi, J., Jafari Jozani, M., Marchand, E. and Parsian, A. (2009), "Bayesian estimation based on k-record data from a general class of distributions under balanced type loss functions", Journal of Statistical Planing and Inference, 139, 1180-1189.

-
- [11] Ahsanullah, M. and Dinesh S. B. (1996), "**Record values of extreme value distribution and a test for domain of attraction of type I extreme value distribution**", Sankhya V, 58, 151-158.
- [12] Amini, M. and Balakrishnan, N. (2013), "**Nonparametric meta-analysis of independent samples**", Computational Statistics and Data Analysis, 66, 70–81.
- [13] Arnold, B. C., Balakrishnan, N. and Nagaraja, H. N. (2008), "**A First Course in Order Statistics**", Wiley, New York.
- [14] Asgharzadeh, A., Valiollahi, R., Raqab, M. Z. (2011), "**Stress-strength reliability of Weibull distribution based on progressively censored samples**", Statistics and Operations Research Transactions, 35(2), 103-124.
- [15] Asgharzadeh, A., Valiollahi, R. and Raqab, M.Z. (2013), "**Estimation of the stress–strength reliability for the generalized logistic distribution**", Statistical Methodology, 15, 73–94.
- [16] Asgharzadeh, A., Valiollahi, R. and Raqab, M.Z. (2017), "**Estimation of $Pr(Y < X)$ for the two-parameter generalized exponential records**", Communications in Statistics: Simulation and Computation, 46, 379-394.
- [17] Awad, A. M., Azzam, M. M. and Hamadan, M. A. (1981), "**Some inference results in $Pr(Y < X)$ in the bivariate exponential model**", Communications in Statistics: Theory and Methods, 10, 2515-2524.
- [18] Baklizi, A. (2008a), "**Estimation of $Pr(X < Y)$ using record values in the one and two parameter exponential distributions**", Communications in Statistics: Theory and Methods, 37, 692-698.
- [19] Baklizi, A. (2008b), "**Likelihood and Bayesian estimation of $Pr(X < Y)$ using lower record values from the generalized exponential distribution**", Computational Statistics and Data Analysis, 52, 3468-3473.
- [20] Baklizi, A. (2014), "**Interval estimation of the stress-strength reliability in the two parameter exponential distribution based on records**", Journal of Statistical Computation and Simulation, 84, 2670-2679.
- [21] Balakrishnan, N., Kamps, U. and Kateri, M. (2009), "**Minimal repair under a step-stress test**", Statistics and Probability Letters, 79, 1548-1558.

- [22] Balakrishnan, N. and Li, T. (2008), "**Ordered ranked set samples and applications to inference**", Journal of Statistical Planning and Inference, 138, 3512–3524.
- [23] Barnett, V. and Moore, K. (1997), "**Best linear unbiased estimates in ranked-set sampling with particular reference to imperfect ordering**", Journal of Applied Statistics, 24, 697–710.
- [24] Barreto, M. C. C. and Barnett, V. (1999), "**Best linear unbiased estimators for the simple linear regression model using ranked set sampling**", Environmental and Ecological Statistics, 6, 119–133.
- [25] Basirat, M., Baratpour, S. and Ahmadi, J. (2013), "**Statistical inferences for stress-strength in the proportional hazard models based on progressive Type-II censored samples**", Journal of Statistical Computation and Simulation, 85, 431-449.
- [26] Basirat, M., Baratpour, S. and Ahmadi, J. (2016), "**On estimation of stress-strength parameter using record values from proportional hazard rate models**", Communications in Statistics-Theory and Methods, 45, 5787-5801.
- [27] Bailey, W. N. (1935), "**Generalized Hyper geometric Series**", University Press, Cambridge.
- [28] Berger, J. O. (1985), "**Statistical Decision Theory and Bayesian Analysis (2nd ed.)**", New York: Wiley.
- [29] Birnbaum, Z. W. (1956), "**On a use of Mann-Whitney statistics**", Proceedings of the third Berkeley Symposium in Mathematics, Statistics and Probability, 1, 13–17.
- [30] Birnbaum, Z. W. and McCarty, B.C. (1958), "**A distribution-free upper confidence bounds for $Pr(Y < X)$ based on independent samples of X and Y** ", The Annals of Mathematical Statistics, 29, 558–562.
- [31] Batacharya, G. K., Johnson, R. A. (1974), "**Estimation of reliability in a multicomponent stress-strength model**", Journal of the American Statistical Association, 69, 966-970.
- [32] Bohn, L. L. (1996), "**A review of nonparametric ranked set sampling methodology**", Communications in Statistics, Part A-Theory and Methods, 25, 2675–2685.
- [33] Carlin, B. P. and Gelfand, A. E. (1993), "**Parametric likelihood inference for record breaking problems**", Biometrika, 80, 507-515.

-
- [34] Casella, G. and Berger, R. L. (1990), "**Statistical Inference**", Pacific Grove, CA: Wadsworth/Brooks Cole.
- [35] Chandler, K. N. (1952), "**The distribution and frequency of record value**", Journal of the Royal Statistical Society, Series B, 14, 220–228.
- [36] Chandra, N.N., and Roy, D., (2001), "**Some results on the reversed hazard rate**", Probability in the Engineering and Informational Sciences 15, 95–102
- [37] Chen, M. and Shao, Q., (1999), "**Monte carlo estimation of Bayesian credible and HPD intervals**", Journal of Computational and Graphical Statistics, 8 (1), 69-92.
- [38] Chen, Z. (2000), "**The efficiency of ranked set sampling relative to simple random sampling under mutiple-parameter family**", Statistica Sinica, 10, 247–263.
- [39] Chen, Z., Bai, Z. and Sinha, A. K. (2003), "**Ranked Set Sampling**": Theory and Applications, Springer, New York.
- [40] David, H. A. and Nagaraja, H. N. (2003), "**Order Statistics**", Third edition, John Wiley Sons, Hoboken, New Jersey.
- [41] Dell, T. R., and Clutter, J. L. (1972), "**Ranked set sampling theory with order statistics background**", Biometrics, 28, 545–553.
- [42] Efron, B. (1979), "**Computers and the theory of the statistics: thinking the unthinkable**", Society for Industrial and Applied Mathematics, 21, 4, 430-443.
- [43] Efron, B. and Tibshirani, R. (1993), "**An Introduction to the Bootstrap**", Chapman and Hall, New York.
- [44] Eryilmaz, S. (2008), "**Consecutive k-out-of-n: G system in stress-strength setup**", Communications in Statistics - Simulation and Computation, 37, 579-589.
- [45] Eskandarzadeh M., Tahmasebi, S., Afshari, M. (2016), "**Information measures for record ranked set samples**", Ciencia e Natura, 38, 554-563.
- [46] Glich, N. (1978), "**Breaking record and breaking boards**", The American Mathematical Monthly, 85, 2-26.
- [47] Gulati, S. and Padgett, W. J. (1994), "**Nonparametric quantile estimation from record-breaking data**", Australian Journal of Statistics, 36, 211–223.

- [48] Gulati, S. and Padgett, W. J. (2003), "**Parametric and Nonparametric Inference from Record-breaking Data.**" Lecture Notes in Statistics, 172, Springer-Verlag, New York.
- [49] Gupta, R. C. and Kirmani, S. N. U. A. (1988), "**Closure and monotonicity properties of nonhomogeneous Poisson processes and record values**", Probability in the Engineering and Informational Sciences, 2, 475–484.
- [50] Gupta, R.D. and Nanda, A.K., (2001), "**Some results on (reversed) hazard rate ordering**", Communication in Statistics: Theory and Methods, 30, 2447–2458
- [51] Gupta, R.C., Gupta, P.L. and Gupta, R.D., (1998), "**Modeling failure time data by Lehman alternatives**", Communication in Statistics: Theory and Methods, 27, 887–904.
- [52] Gupta, R.C. and Gupta, R.D., (2007), "**Proportional reversed hazard rate model and its applications**", Journal of Statistical Planning and Inference 137, 3525–3536.
- [53] Guttman, I., Johnson, R.A., Battacharyya, G.K. and Reiser, B. (1998), "**Confidence limits for stress-strength models with explanatory variables**", Technometrics, 30(2), 161-168.
- [54] Halls and Dell (1966), "**Trial of ranked set sampling for forage yields**", Forest Science, 12, 22–26.
- [55] Jafari Jozani, M. and Johnson, B. C. (2011), "**Design based estimation for ranked set sampling in finite populations**", Environmental and Ecological Statistics, 18, 663–685.
- [56] Johnson, R. A. (1988), "**Stress-strength Models for Reliability**", Elsevier, 7, 27-54.
- [57] Hassan, S. A., Abd-Allah, M. and Nagy, F. H. (2018), "**Bayesian Analysis of Record Statistics Based on Generalized Inverted Exponential Model**", International Journal on Advanced Science Engineering Information Technology, 8(2), 323-335.
- [58] Klein, J. P., Basu, A. P. (1985), "**Estimating reliability for bivariate exponential distributions**", Sankhya, 47, 346-353.
- [59] Kotz, S., Lumelskii, Y. and Pensky, M. (2003), "**The Stress-Strength Model and its Generalizations: Theory and Applications**", World Scientific, Singapor.
- [60] Kochar, S. C. (1990), "**Some partial ordering results on record values**", Communications in Statistics: Theory and Methods, 19, 229-306.
- [61] Kochar, S. C. (1996), "**A note on dispersive ordering of record values**", Calcutta Statistical Association, 46, 63-67.

-
- [62] Kundu, D., Gupta, R. D. (2005), "**Estimation of $P(X < Y)$ for generalized exponential distribution**", *Metrika*, 61, 291-308.
- [63] Lehmann, E.L. and Casella, G. (1998), "**Theory of Point Estimation, Second Edition**", Springer Texts in Statistics, New York.
- [64] Lio, Y. L. and Tsai, T. R. (2012), "**Estimation of $Pr(X < Y)$ for Burr XII distribution based on the progressively first failure-censored samples**", *Journal of Applied Statistics*, 39(2), 309-322.
- [65] Manisha Pala, M. Masoom Alib, Jungsoo Wooc (2005), "**Estimation and testing of $P(Y > X)$ in two-parameter exponential distributions**", *Statistics*, 5, 415-428
- [66] Mann, H. B., Whitney, D. R. (1947), "**On a test whether one of two random variables is stochastically larger than the other**", *The Annals of Mathematical Statistics*, 18, 50-60.
- [67] Marshall, A. W. and Olkin, O. (2007), "**Life distributions**", Springer, New York
- [68] McIntyer, G.A. (1952), "**A method for unbiased selective sampling, using ranked sets**", *Australian Journal of Agricultural Research*, 3, 385-390.
- [69] Muttlak, H. A., Abu-Dayyeh, W. A., Saleh, M. F. and Al-Sawi, E. (2010), "**Estimating $Pr(Y < X)$ using ranked set sampling in case of the exponential distribution**", *Communication in Statistics: Theory and Methods*, 39, 1855–1868.
- [70] Nadar, M and Kizilaslan, F. (2014), "**Classical and Bayesian estimation of $Pr(X < Y)$ using upper record values from Kumaraswamy's distribution**", *Statistical Papers*, 55, 3, 751-783.
- [71] Owen, D. B., Craswell, K. J., Hanson, D. L. (1964), "**Non-parametric upper confidence bounds for $Pr(Y < X)$ and confidence limits for $Pr(Y < X)$ when X and Y are normal**", *Journal of the American Statistical Association*, 59, 906-924.
- [72] Pakdaman, Z., and Ahmadi, J., (2018), "**Point estimation of the stress- strength reliability parameter for parallel system with independent and non-identical components**", *Communications in Statistics: Simulation and Computation*, 47, 1193-1203.
- [73] Patil, G. P., Sinha, A. K. and Taillie, C. (1999), "**Ranked set sampling a bibliography**", *Environmental and Ecological Statistics*, 6, 91–98.

- [74] Paul, J. and Thomas, P.Y. (2017), "**Concomitant record ranked set sampling**", Communications in Statistics-Theory and methods, 46 (19), 9518-9540.
- [75] Proschan (1963), "**Theoretical explanation of observed decreasing failure rate**", Technometrics, 5, 375–383
- [76] Raqab, M. Z, Madi, T. and Kundu, D. (2008), "**Estimation of $Pr(Y < X)$ for the three-parameter generalized exponential distribution**", Communications in Statistics-Theory and Methods, 37, 2854-2865.
- [77] Rezaei, S., Tahmasbi, R. and Mahmoodi, M. (2010), "**Estimation of $Pr(Y < X)$ for generalized Pareto distribution**", Journal of Statistical Planning and Inference, 140, 480-494.
- [78] Rizzo, M. L. (2008), "**Statistical Computing with R**", Chapman Hall, Taylor Francis Group, LLC.
- [79] Safarian, A., Arashi, M. and Arabi Belaghi, R. (2019a), "**Improved point and interval estimation of the stress-strength reliability based on record ranked set sampling scheme**", Statistics, 53, 101–125.
- [80] Safarian, A., Arashi, M. and Arabi Belaghi, R. (2019b), "**Improved estimators for stress-strength reliability using record ranked set sampling scheme**", Communications in Statistics: Simulation and Computation, 48(9), 2708–2726
- [81] Salehi, M., Ahmadi, J. (2014), "**Record ranked set sampling scheme**", Metron, 72, 351-365
- [82] Salehi, M., Ahmadi, J. (2015a), "**Estimation of stress-strength using record ranked set sampling scheme from the exponential distribution**", Filomat, 29, 1149-1162.
- [83] Salehi, M. and Ahmadi, J. (2015b), "**Prediction of order statistics and record values based on ordered ranked set sampling**", Journal of Statistical Computation and Simulation, 85, 77–88.
- [84] Salehi, M., Ahmadi, J. and Dey, S. (2016), "**Comparison of two sampling schemes for generating record-breaking data from the proportional hazard rate models**", Communications in Statistics-Theory and Methods, 45, 3721-3733.

-
- [85] Samaniego, F. J., Whitaker, L. R. (1986), "**On estimating population characteristics from record-breaking observations**", I: Parametric Results, Naval research Logistics Quarterly, 33, 531–543.
- [86] Samaniego, F. J., Whitaker, L. R. (1988), "**On estimating population characteristics from record-breaking observations**", II: Nonparametric Results, Naval research Logistics Quarterly, 35, 221–236.
- [87] Saracoglu, B., Kaya, M. F. (2007), "**Maximum likelihood estimation and confidence intervals of system reliability for Gomperts distribution in stress-strength models**", Selcuk Journal of Applied Mathematics, 8, 25-36.
- [88] Saracoglu, B., Kinaci, I. and Kundu, D. (2012), "**On estimation of $\mathcal{R} = Pr(Y < X)$ for exponential distribution under progressive type-II censoring**", Journal of Statistical Computation and Simulation, 82, 729-744.
- [89] Sengupta, D. and Nanda, A.K., (1999), "**Log-concave and concave distributions in reliability**", Naval Res. Logistics Quart. 46, 419–433.
- [90] Sinha, A. K., R. W. (2005), "**On some recent developments in ranked set sampling**", Bulletin of Informatics and Cybernetics, 37.
- [91] Smith, A.F.M. (1991). "**Bayesian computational methods**", Philosophical Transactions of the Royal Society of London A, 369-386.
- [92] Stokes, S. L. (1980), "**Estimation of variance using judgment ordered ranked set samples**", Biometrics, 36, 35–42.
- [93] Stokes, S. L. (1995), "**Parametric ranked set sampling**", Annals of the Institute of Statistical Mathematics, 47, 465–482.
- [94] Takahasi, K. and Wakimoto, K. (1968), "**On unbiased estimates of the population mean based on the sample stratified by means of ordering**", Annals of the Institute of Statistical Mathematics, 20, 1–31.
- [95] Tanner, M.A. (1996). "**Tools for Statistical inference**", 3rd ed. Springer-verlag, New York.
- [96] Tong, H. (1974), "**A note on the estimation of $Pr(Y < X)$ in the exponential case**", Technometrics, 16, 625.
- [97] Wasserman, L. (2006), "**All of Non-parametric Statistics**", Springer, New York.

- [98] Wilcoxon, F. (1945), "**Individual comparisons by ranking methods**", Biometrics Bulletin, 1, 80 - 83.
- [99] Zheng, G. and Al-Saleh, M. F. (2000), "**Modified maximum likelihood estimators based on ranked set sampling**", Annals of the Institute of Statistical Mathematics, 54, 641–658.

پیوست آ

کدهای برنامه نویسی

به علت حجم زیاد برنامه های استفاده شده در این رساله و جلوگیری از ازدیاد کاذب صفحات، در زیر تنها نمونه ای از کدهای برنامه نویسی محاسبات عددی و شبیه سازی آورده شده اند و سایر کدها، نزد گردآورنده محفوظ می باشند. لذا، خوانندگان گرامی رساله می توانند در صورت تمایل با آدرس ایمیل sadeghpour.amineh@yahoo.com تماس حاصل نمایند. لازم به ذکر است که تمامی محاسبات عددی و شبیه سازی این رساله با استفاده از نرم افزار R انجام شده اند.

(۱) کد برنامه های زیر مربوط به نمودار ۱.۲ هستند.

```
k=5000000
n=c(2,3,4,5,6)
m=c(2,3,4,5,6)
M=m*(m+1)/2
N=n*(n+1)/2
R=seq(0,1,by=0.02)
mse1=numeric(length(R))
mse2=numeric(length(R))
mse3=numeric(length(R))
mse4=numeric(length(R))
```

```

mse5=numeric(length(R))
for(i in 1:length(R)){
#v=rf(k,2*M,2*N)
Rhat = (1+rf(k,2*M[5],2*N[1])*(1-R[i])/R[i])^(-1)
mse1[i] = mean((Rhat-R[i])^2)
Rhat = (1+rf(k,2*M[1],2*N[5])*(1-R[i])/R[i])^(-1)
mse2[i] = mean((Rhat-R[i])^2)
Rhat = (1+rf(k,2*M[3],2*N[3])*(1-R[i])/R[i])^(-1)
mse3[i] = mean((Rhat-R[i])^2)
Rhat = (1+rf(k,2*M[4],2*N[2])*(1-R[i])/R[i])^(-1)
mse4[i] = mean((Rhat-R[i])^2)
Rhat = (1+rf(k,2*M[2],2*N[4])*(1-R[i])/R[i])^(-1)
mse5[i] = mean((Rhat-R[i])^2)
}
plot(c(0,1), c(0,0.05), type = "n", xlab = "R", main = "",
      ylab="MSE", yaxt = "n", xaxs = "i", yaxs = "i")
abline(v=0.5,lty=2)
axis(2)
abline(h=0,lty=2)
lines(R,mse1,type='l',lwd=2)
lines(R,mse2, lty=2, col=2,lwd=2)
lines(R,mse3, lty=3, col=3,lwd=2)
lines(R,mse4, lty=4, col=4,lwd=2)
lines(R,mse5, lty=5, col=5,lwd=2)
abline(v=0.5 , lty=3)
legend("topright",legend=c(expression(paste("n = 2 , m = 6")),
      expression(paste("n = 6 , m = 2")),
      expression(paste("n = 4 , m = 4")),
      expression(paste("n = 3 , m = 5")),
      expression(paste("n = 5 , m = 3"))),
      lty=1:6,lwd=2, col=1:6, bty='n', cex=1)

```

۲) کد برنامه‌ی زیر مربوط به محاسبات شبیه‌سازی جدول‌های ۱.۲ و ۱.۳ می‌باشند.

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

```
library(xtable)

library(parallel) # for parallel computing
library(doParallel) # for parallel computing
library(foreach) # for parallel computing

# functions -----
Q.boot=function(t,RR){
  f=ecdf(RR)
  ff=function(x) f(x)-t
  uniroot(ff,int=c(0,10))$root
}
Q.boot=Vectorize(Q.boot,"t")
Q.boot(c(0.5,0.9),rexp(100))
f.hpd1=function(N, M, theta1, theta2, B.hpd, a, b, c, d, alpha) {
  B=B.hpd
  # T1=rgamma(1,N,1/theta1)
  # T2=rgamma(1,M,1/theta2)
  # Theta1i=rgamma(B,a+N,1/(b+T1))
  # Theta2i=rgamma(B,c+M,1/(d+T2))
  T1=rgamma(1,shape=N,scale=1/theta1)
  T2=rgamma(1,shape=M,scale=1/theta2)
  Theta1i=rgamma(B,shape=a+N,scale=1/(b+T1))
  Theta2i=rgamma(B,shape=c+M,scale=1/(d+T2))
  Ri=sort(Theta2i/(Theta2i+Theta1i))
  up=B-floor(B*(1-alpha))
  J=1:up
  jst=J[which.min(Ri[J+floor(B*(1-alpha))]-Ri[J])]
  # EL=Ri[jst+floor(B*(1-alpha))]-Ri[jst]
  # CP=ifelse(Ri[jst]<R & R<Ri[jst+floor(B*(1-alpha))],1,0)
  # return(c(EL = EL , CP=CP))
  # return(list(EL = EL , CP=CP))
  LCL=Ri[jst]
  UCL=Ri[jst+floor(B*(1-alpha))]
  return(c(LCL = LCL , UCL=UCL))
}
```

```
f.hpd2=function(N, M, theta1, theta2, B.hpd, alpha) {
  B=B.hpd
  T1=rgamma(1,shape=N,scale=1/theta1)
  T2=rgamma(1,shape=M,scale=1/theta2)
  Theta1i=rgamma(B,shape=N,scale=1/(T1))
  Theta2i=rgamma(B,shape=M,scale=1/(T2))
  Ri=sort(Theta2i/(Theta2i+Theta1i))
  up=B-floor(B*(1-alpha))
  J=1:up
  jst=J[which.min(Ri[J+floor(B*(1-alpha))]-Ri[J])]
  # EL=Ri[jst+floor(B*(1-alpha))]-Ri[jst]
  # CP=ifelse(Ri[jst]<R & R<Ri[jst+floor(B*(1-alpha))],1,0)
  # return(c(EL = EL , CP=CP))
  # return(list(EL = EL , CP=CP))
  LCL=Ri[jst]
  UCL=Ri[jst+floor(B*(1-alpha))]
  return(c(LCL = LCL , UCL=UCL))
}

f.boot_t1=function(n,m,theta2,R,alpha){
  p=min(m,n)/max(m,n)
  N=n*(n+1)/2
  M=m*(m+1)/2
  theta1=theta2*(1/R-1)
  r=c()
  for(i in 1:n) r=c(r,rgamma(1,shape=i,scale=1/theta1))
  s=c()
  for(j in 1:m) s=c(s,rgamma(1,shape=j,scale=1/theta2))
  theta1.hat=N/sum(r)
  theta2.hat=M/sum(s)
  R.hat=theta2.hat/(theta2.hat+theta1.hat)
  I=(n>=m)
  eta.hat=theta2.hat*theta1.hat/(theta1.hat+theta2.hat)^2
  *sqrt(1/(N*p^(1-I))+1/(p^I*M))
  R.star=eta.star=SD=rep(NA,B)
```

```
for(bbb in 1:B){
  r.star=c()
  for(i in 1:n) r.star=c(r.star,rgamma(1,shape=i,scale=1/theta1.hat))
  s.star=c()
  for(j in 1:m) s.star=c(s.star,rgamma(1,shape=j,scale=1/theta2.hat))
  theta1.star=N/sum(r.star)
  theta2.star=M/sum(s.star)
  R.star[bbb]=theta2.star/(theta2.star+theta1.star)
  R.tilde=c()
  for(bb in 1:30){
    r.tilde=c()
    for(i in 1:n) r.tilde=c(r.tilde,rgamma(1,shape=i,scale=1/theta1.star))
    s.tilde=c()
    for(j in 1:m) s.tilde=c(s.tilde,rgamma(1,shape=j,scale=1/theta2.star))
    theta1.tilde=N/sum(r.tilde)
    theta2.tilde=M/sum(s.tilde)
    R.tilde[bb]=theta2.tilde/(theta2.tilde+theta1.tilde)
  }
  SD[bbb]=sqrt(var(R.tilde))
  eta.star[bbb]=theta2.star*theta1.star/(theta1.star+theta2.star)^2
  *sqrt(1/(N*p^(1-I))+1/(p^I*M))
}
z.star=(R.star-R.hat)/eta.star
z1=quantile(z.star,alpha/2,type=1);z2=quantile(z.star,1-alpha/2,type=1)
z.tilde=(R.star-R.hat)/SD
t1=quantile(z.tilde,alpha/2,type=1);t2=quantile(z.tilde,1-alpha/2,type=1)
int.perc=c(Q.boot(alpha/2,R.star),Q.boot(1-alpha/2,R.star))
int.boot1=c(R.hat-z2*eta.hat,R.hat-z1*eta.hat)
int.boot2=c(R.hat-t2*eta.hat,R.hat-t1*eta.hat)
int.AMLE=c(R.hat-qnorm(alpha/2)*eta.hat,R.hat+qnorm(alpha/2)*eta.hat)
l.mle=1/(1+(1-R.hat)/R.hat*qf(1-alpha/2,2*N,2*M))
u.mle=1/(1+(1-R.hat)/R.hat*qf(alpha/2,2*N,2*M))
int.MLE=c(l.mle,u.mle)
deltastar=R.star-R.hat
```

```

# dd=quantile(deltastar,c(alpha,1-alpha))
dd=quantile(deltastar,c(alpha/2,1-alpha/2))
basic.boot=R.hat-c(dd[2],dd[1])
tt1=sum(r)
tt2=sum(s)
A=(a+N)*(d+tt2)/((c+M)*(b+tt1))
# l.bci1=1/(1+A*qf(alpha/2,2*(a+N),2*(c+M)))
# u.bci1=1/(1+A*qf(1-alpha/2,2*(a+N),2*(c+M)))
u.bci1=1/(1+A*qf(alpha/2,2*(a+N),2*(c+M)))
l.bci1=1/(1+A*qf(1-alpha/2,2*(a+N),2*(c+M)))
bci1=c(l.bci1,u.bci1)
# l.bci2=1/(1+N*tt2/(M*tt1)*qf(alpha/2,2*N,2*M))
# u.bci2=1/(1+N*tt2/(M*tt1)*qf(1-alpha/2,2*N,2*M))
u.bci2=1/(1+N*tt2/(M*tt1)*qf(alpha/2,2*N,2*M))
l.bci2=1/(1+N*tt2/(M*tt1)*qf(1-alpha/2,2*N,2*M))
bci2=c(l.bci2,u.bci2)
hpd1=f.hpd1(N, M, theta1, theta2, B.hpd, a, b, c, d, alpha)
hpd2=f.hpd2(N, M, theta1, theta2, B.hpd, alpha)
res=c(int.perc,
      # int.boot1,
      int.boot2,int.MLE,
      # int.AMLE,
      basic.boot,
      bci1,
      bci2,
      hpd1,
      hpd2,
      R.hat)
names(res)=c("L.Perc","U.Perc",
             # "L.Boot1","U.Boot1",
             "L.Boot2","U.Boot2",
             "L.MLE","U.MLE",
             # "L.AMLE","U.AMLE",
             "L.basic","U.basic",

```

\. \

```
      "L.bci1","U.bci1",
      "L.bci2","U.bci2",
      "L.HPD1","U.HPD1",
      "L.HPD2","U.HPD2",
      "R.hat")

  return(res)
}

# simulation -----
# B=10 # Only for test
# B=200
B=1000
# B.hpd=100 # Only for test
B.hpd=1000
#### Prior parameters ###
a=1; b=2
c=1; d=2
###
test=f.boot_t1(n=3,m=3,theta2=0.4,R=0.4,alpha=0.1)
test
#repl=100 # Only for test
# repl=1000
repl=10000
pt=proc.time()
ncore=detectCores()
cl=makeCluster(ncore)
registerDoParallel(cl)
table=foreach(R=c(0.1,0.5,0.95), .combine='rbind') %:%
  foreach(n=c(2,4,6), .combine='rbind') %:%
    foreach(m=c(2,4,6), .combine='rbind') %do% {
      set.seed(123)
      total=t(replicate(repl,expr=f.boot_t1(n,m,theta2=0.4,R,alpha=0.1)))
      L.Perc=total[,1];U.Perc=total[,2]
      # L.Boot1=pmax(total[,3],0);U.Boot1=pmin(total[,4],1)
      L.Boot2=pmax(total[,3],0);U.Boot2=pmin(total[,4],1)
```

```

L.MLE=total[,5];U.MLE=total[,6]
# L.AMLE=pmax(total[,9],0);U.AMLE=pmin(total[,10],1)
L.basic=total[,7];U.basic=total[,8]
L.bci1=total[,9];U.bci1=total[,10]
L.bci2=total[,11];U.bci2=total[,12]
L.hpd1=total[,13];U.hpd1=total[,14]
L.hpd2=total[,15];U.hpd2=total[,16]
#Bias=mean(total[,15])-R
#mse.R=mean((total[,15]-R)^2)
CV.Perc=mean((L.Perc<=R)&(R<=U.Perc))
L.Perc=mean(U.Perc-L.Perc)
# CV.Boot1=mean((L.Boot1<=R)&(R<=U.Boot1))
# L.Boot1=mean(U.Boot1-L.Boot1)
CV.Boot2=mean((L.Boot2<=R)&(R<=U.Boot2))
L.Boot2=mean(U.Boot2-L.Boot2)
CV.MLE=mean((L.MLE<=R)&(R<=U.MLE))
L.MLE=mean(U.MLE-L.MLE)
# CV.AMLE=mean((L.AMLE<=R)&(R<=U.AMLE))
# L.AMLE=mean(U.AMLE-L.AMLE)
CV.basic=mean((L.basic<=R)&(R<=U.basic))
L.basic=mean(U.basic-L.basic)
CV.bci1=mean((L.bci1<=R)&(R<=U.bci1))
L.bci1=mean(U.bci1-L.bci1)
CV.bci2=mean((L.bci2<=R)&(R<=U.bci2))
L.bci2=mean(U.bci2-L.bci2)
CV.hpd1=mean((L.hpd1<=R)&(R<=U.hpd1))
L.hpd1=mean(U.hpd1-L.hpd1)
CV.hpd2=mean((L.hpd2<=R)&(R<=U.hpd2))
L.hpd2=mean(U.hpd2-L.hpd2)
c(CV.Perc,L.Perc,
  # CV.Boot1,L.Boot1,
  CV.Boot2,
  L.Boot2,CV.MLE,L.MLE,
  # CV.AMLE,L.AMLE,

```

```
        CV.basic,L.basic,
        CV.bci1,L.bci1,CV.hpd1,L.hpd1,
        CV.bci2,L.bci2,CV.hpd2,L.hpd2)
    }
#   }
# }
stopCluster(cl)
pt=proc.time()-pt
pt
colnames(table)=c("CV.Perc","L.Perc",
                  # "CV.Boot1","L.Boot1",
                  "CV.Boot2","L.Boot2","CV.MLE","L.MLE",
                  # "CV.AMLE","L.AMLE",
                  "CV.basic","L.basic",
                  "CV.bci1","L.bci1","CV.hpd1","L.hpd1",
                  "CV.bci2","L.bci2","CV.hpd2","L.hpd2")
rownames(table)=NULL
round(table,4)
R=rep(c(0.1,0.5,0.95),each=9)
n=rep(c(2,4,6),each=3,times=3)
m=rep(c(2,4,6),times=9)
round(cbind(R,n,m,table),3)
xtable(round(cbind(R,n,m,table),3),digits=3)
id=1:8
round(cbind(R,n,m,table[,id]),3)
xtable(round(cbind(R,n,m,table[,id]),3),digits=3)
id=9:16
round(cbind(R,n,m,table[,id]),3)
xtable(round(cbind(R,n,m,table[,id]),3),digits=3)
write.table(table,"table1out.txt",row.names = F)
save(table,"table1out.RData")
load("table1out.RData")
```


پیوست ب

اثبات روابط

برهان. رابطه (۲۶.۲):

همان‌طور که قبلاً نیز بیان شده است $R_{i,i}$ ها متغیرهای تصادفی مستقل و $RRSS$ های پایین هستند. با توجه به رابطه (۶.۱) تابع چگالی $R_{i,i}$ به صورت زیر است:

$$f_{R_{i,i}}(r_{i,i}) = \frac{\{-\log(1 - e^{-r_{i,i}})\theta_1\}^{i-1}}{(i-1)!} \theta_1 e^{-r_{i,i}} (1 - e^{-r_{i,i}})^{\theta_1-1},$$

با در نظر گرفتن

$$Z_i = -\log(1 - e^{-R_{i,i}}),$$

و استفاده از روش تبدیل

$$f_{Z_i}(z_i) = \frac{\theta_1^i z_i^{i-1}}{(i-1)!} e^{-\theta_1 z_i}.$$

پس $(Z_i \sim \text{Gamma}(i, \theta_1))$ ، چون $T_1 = \sum_{i=1}^n Z_i$ بنابراین نتیجه مورد نظر به دست می‌آید.

□

رابطه (۲۷.۲) نیز مشابه روابط فوق محاسبه می‌شود.

برهان. رابطه (۳۳.۲):

$$\begin{aligned} f_{U|(T_1=t_1)}(u) &= \frac{f_u(-\log(1 - e^{-r_1,1}))f_{T_1}(\sum_{i=2}^n -\log(1 - e^{-r_{i,i}}) = t_1 - u)}{f_{T_1}(\sum_{i=1}^n -\log(1 - e^{-r_{i,i}}) = t_1)} \\ &= (N-1) \frac{(t_1 - u)^{N-2}}{t_1^{N-1}}. \end{aligned}$$

□

رابطه (۳۴.۲) نیز به همین ترتیب اثبات می‌شود.

برهان. رابطه (۳۶.۲):

با توجه به رابطه (۲۶.۲) و (۲۷.۲) و استفاده از روابط بین توزیع‌ها داریم:

$$\begin{aligned} 2\theta_1 T_1 &\sim \chi^2_{(2N)}, \\ 2\theta_2 T_2 &\sim \chi^2_{(2M)}. \end{aligned}$$

همچنین با استفاده از رابطه (۲۲.۲)، $\frac{\theta_1}{\theta_2} = \frac{1-R}{R}$ ، با جایگذاری روابط فوق در (۲۵.۲) و با در نظر گرفتن این نکته که حاصل تقسیم دو توزیع χ^2 توزیع فیشر می‌شود، رابطه هم توزیعی (۳۶.۲) به دست می‌آید.

□

واژه‌نامه فارسی به انگلیسی

Complete sufficient statistic	آماره بسنده کامل
Bias	اریبی
Coverage Probability (CP)	احتمال پوشش
Distribution-free	آزاد-توزیع
Order statistic	آماره ترتیبی
Proportional hazard rate model	الگوی نرخ خطر متناسب
Proportional reversed hazard rate model	الگوی نرخ خطر معکوس متناسب
Confidence interval (CI)	بازه‌ی اطمینان
Survival	بقا
Bootstrap	بوت استرپ
Estimation	برآورد
Estimator	برآوردگر
Maximum likelihood estimator (MLE)	برآوردگر درست‌نمایی ماکسیمم
Uniformly minimum variance unbiased estimator (UMVUE)	برآوردگر ناریب با واریانس به‌طور یکنواخت مینیمم
Point estimator	برآوردگر نقطه‌ای
Interval estimator	برآوردگر فاصله‌ای
Parameter	پارامتر
Density function	تابع چگالی
Probability distribution function	تابع توزیع احتمال
Probability density function	تابع چگالی احتمال
Joint probability density function	تابع چگالی احتمال توأم
Indicator Function	تابع نشان‌گر
Likelihood function	تابع درست‌نمایی
Floor function	تابع جزء صحیح
Quantile function	تابع چندک

Complete gamma function	تابع گامای کامل
Gamma function	تابع گاما
Risk function	تابع مخاطره
Stress-Strength	تنش-مقاومت
Empirical	تجربی
Cumulative Distribution	توزیع تجمعی
Asymptotic distribution	توزیع مجانبی
Exponential distribution	توزیع نمایی
Weibull distribution	توزیع وایبل
Location-Scale family	خانواده مکان-مقیاس
Record	رکورد
Characteristic	شاخص
Monte Carlo Simulation	شبیه‌سازی مونت کارلو
Expected Length (EI)	طول مورد انتظار
Life Time	طول عمر
Reliability	قابلیت اعتماد
Efficiency	کارایی
Independent	مستقل
Independent and identically distributed (iid)	مستقل و هم توزیع
Mean squared error (MSE)	میانگین توان دوم خطا
Unbiased	نااریب
Relational	نسبی
Ranked set sample (RSS)	نمونه‌ی مجموعه‌ی رتبه‌دار
Record ranked set sample (RRSS)	نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی

واژه‌نامه انگلیسی به فارسی

Asymptotic distribution	توزیع مجانبی
Bias	اریبی
Bootstrap	بوت استرپ
Characteristic	شاخص
Complete gamma function	تابع گامای کامل
Complete sufficient statistic	آماره بسنده کامل
Confidence interval (CI)	بازه‌ی اطمینان
Coverage Probability (CP)	احتمال پوشش
Cumulative Distribution	توزیع تجمعی
Density function	تابع چگالی
Distribution-free	آزاد-توزیع
Empirical	تجربی
Efficiency	کارایی
Estimation	برآورد
Estimator	برآوردگر
Expected Length (EI)	طول مورد انتظار
Exponential distribution	توزیع نمایی
Floor function	تابع جزء صحیح
Independent	مستقل
Independent and identically distributed (iid)	مستقل و هم توزیع
Indicator Function	تابع نشان‌گر
Interval estimator	برآوردگر فاصله‌ای
Joint probability density function	تابع چگالی احتمال توأم
Life Time	طول عمر
Likelihood function	تابع درستنمایی
Location-Scale family	خانواده مکان-مقیاس

Maximum likelihood estimator (MLE)	برآوردگر درست‌نمایی ماکسیمم
Mean squared error (MSE)	میانگین توان دوم خطا
Monte Carlo Simulation	شبیه‌سازی مونت کارلو
Order statistic	آماره ترتیبی
Parameter	پارامتر
Probability distribution function	تابع توزیع احتمال
Probability density function	تابع چگالی احتمال
Proportional hazard rate model	الگوی نرخ خطر متناسب
Proportional reversed hazard rate model	الگوی نرخ خطر معکوس متناسب
Point estimator	برآوردگر نقطه‌ای
Quantile function	تابع چندک
Ranked set sample (RSS)	نمونه‌ی مجموعه‌ی رتبه‌دار
Relational	نسبی
Reliability	قابلیت اعتماد
Record	رکورد
Record ranked set sample (RRSS)	نمونه‌ی مجموعه‌ی رتبه‌دار رکوردی
Risk function	تابع مخاطره
Stress-Strength	تنش-مقاومت
Survival	بقا
Unbiased	نااریب
Uniformly minimum variance unbiased estimator (UMVUE)	برآوردگر نااریب با واریانس به‌طور یکنواخت مینیمم
Weibull distribution	توزیع وایبل

Abstract

In some experiments, we would like to evaluate the strength of a component against the stress of other variables. Such case in reliability theory is called a stress-strength model. The stress-strength reliability has received considerable attention in branches of sciences such as industry, medicine and engineering. So, this thesis studies point, interval and Bayesian estimations of stress-strength reliability parameter based on record ranked set sampling scheme from the generalized exponential distribution. Also, the record ranked set sampling scheme is compared with ordinary records in case of point and interval estimations under the proportional reversed hazard rate model.

Keywords: Stress-Strength reliability, Proportional reversed hazard rate model, Record ranked set sampling, Record values, Generalized exponential distribution.



Shahrood University of Technology

Faculty of Mathematical Sciences

PhD Thesis in: Probability

Estimation of the Stress-Strength Reliability Based on Record Ranked Set Sample

By: Amineh Sadeghpour

Supervisor

Ahmad Nezakati

Advisor

Mahdi Salehi

April 2020