



دانشکده علوم ریاضی
گروه ریاضی کاربردی

پایان نامه کارشناسی ارشد

رهیافت بیزی قابلیت اطمینان سیستم‌های موازی در مدل‌های فشار-مقاومت

محسن مهدی زاده

استاد راهنما

دکتر احمد نزاکتی رضا زاده

استاد مشاور

دکتر حسین باغی‌شینی

شهریور ۱۳۹۴

تقدیم بہ پدر و مادر عزیزم۔

سپاس خدای را که سخنوران، در ستودن او بمانند و شمارندگان، شمردن نعمت‌های
او ندانند و کوشندگان، حق او را کزاردن نتوانند. و سلام و درود بر محمد و خاندان
پاک او، طاهران معصوم، هم آنان که جودمان و امدار و جودشان است؛ و نفرین پیوسته
بر دشمنان ایشان تا روز رستاخیز...

بدون شک جایگاه و منزلت معلم، اجل از آن است که در مقام قدردانی از زحمات بی‌شائبه‌ی او، با زبان
قاصر و دست ناتوان، چیزی بنگاریم.

اما از آنجایی که تجلیل از معلم، سپاس از انسانی است که هدف و غایت آفرینش را تامین می‌کند و
سلامتی امانت‌هایی را که به دستش سپرده‌اند، تضمین؛ بر حسب وظیفه و از باب ”من لم یشکر المنعم
من المخلوقین لم یشکر الله عزّو جلّ“: از پدر و مادر، این دو معلم بزرگوار که همواره بر کوتاهی و درشتی
من، قلم عفو کشیده و کریمانه از کنار غفلت‌هایم گذشته‌اند و در تمام عرصه‌های زندگی یار و یابوری بی
چشم داشت برای من بوده‌اند؛ از استاد با کمالات و شایسته؛ جناب آقای دکتر نزاکتی که در کمال سعه
صدر، با حسن خلق و فروتنی، از هیچ کمکی در این عرصه بر من دریغ ننمودند و زحمت راهنمایی این
رساله را بر عهده گرفتند؛ از استاد صبور و با تقوا، جناب آقای دکتر باغیثی، که زحمت مشاوره این
رساله را در حالی متقبل شدند که بدون مساعدت ایشان، این پروژه به نتیجه مطلوب نمی‌رسید؛ کمال
تشکر و قدردانی دارم.

باشد که این خردترین، بخشی از زحمات آنان را سپاس گوید.

محسن مهدی زاده
شهریور ۱۳۹۴

تعمدنامه

اینجانب محسن مهدی زاده دانشجوی کارشناسی ارشد رشته آمار دانشکده علوم ریاضی دانشگاه شاهرود، نویسنده پایان نامه با عنوان رهیافت بیزی قابلیت اطمینان سیستم های موازی در مدل های فشار-مقاومت، تحت راهنمایی دکتر احمد نراکتی رضا زاده متعهد می شوم:

- تحقیقات در این پایان نامه توسط اینجانب انجام شده است و از صحت و اصالت برخوردار است.
- در استفاده از نتایج پژوهش های دیگر پژوهشگران، به مرجع مورد استفاده استناد شده است.
- مطالب این پایان نامه، تاکنون توسط خود، یا فرد دیگری برای دریافت هیچ نوع مدرک یا امتیازی در هیچ جا ارایه نشده است.
- حقوق معنوی این اثر، به دانشگاه شاهرود متعلق دارد، و مقالات مستخرج با نام “ دانشگاه شاهرود “ یا “ University of Shahrood “ به چاپ خواهد رسید.
- حقوق معنوی تمام افرادی که در به دست آوردن نتایج اصلی پایان نامه تاثیرگذار بوده اند، در مقالات مستخرج از پایان نامه رعایت می گردد.
- در تمام مراحل انجام این پایان نامه، در مواردی که از موجود زنده (یا بافت های آنها) استفاده شده است، ضوابط و اصول اخلاقی رعایت شده است.
- در تمام مراحل انجام این پایان نامه، در مواردی که به حوزه اطلاعات شخصی افراد دسترسی یافته (یا استفاده) شده است، اصل رازداری و اصول اخلاق انسانی رعایت شده است.

محسن مهدی زاده
شهریور ۱۳۹۴

مالکیت نتایج و حق نشر

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

چکیده

امروزه قابلیت اطمینان و میزان کارایی سیستم‌ها مهم‌ترین اصل در میزان به‌کارگیری و استفاده از سیستم‌ها می‌باشد. هدف ما برآورد قابلیت اطمینان سیستم‌ها در دو نوع سری و موازی در مدل‌های فشار-مقاومت است. لذا در راستای دستیابی به برآوردی بهتر و حصول نتایجی مفید برای بالا بردن راندمان کاری سیستم‌ها، قابلیت اطمینان سیستم‌ها، با فرض مستقل بودن مولفه‌ها (مقاومت‌ها)، تحت توزیع‌های مختلف آماری (نمایی، وایبل، گاما، پاراتو...)، به دو روش ماکزیم درسنمایی و بیزی، برآورد، و در بعضی موارد باهم مقایسه شده‌اند. علاوه بر برآورد قابلیت اطمینان توزیع برآوردگر حاصله را نیز به دست می‌آوریم. در ادامه روند دستیابی به برآوردگر بهینه بخشی از برآوردگرها را شبیه سازی می‌کنیم. تحت بعضی از توزیع‌ها، برآورد قابلیت اطمینان دارای فرم بسته و پیچیده‌ای می‌باشد، در این‌گونه موارد، برای شبیه سازی، استفاده از روش‌های نمونه‌گیری زنجیر مارکوفی مونت کارکو (MCMC) پیشنهاد می‌شود.

کلمات کلیدی

قابلیت اطمینان، سیستم‌های سری و موازی، مدل‌های فشار-مقاومت، روش‌های نمونه‌گیری زنجیر مارکوفی مونت کارلویی.

پیشگفتار

امروزه در بخش صنعت و دیگر بخش‌ها با سیستم‌هایی سروکار داریم که متشکل از تعدادی مولفه (مقاومت) است، که برای انجام یک هدف به یکدیگر متصل شده‌اند. هر مولفه به‌طور مستقل یا وابسته به مولفه‌های دیگر وظیفه‌ای را انجام می‌دهد و بسته به نوع اتصال آن‌ها در سیستم، سیستم‌ها دارای دو نوع سری (متوالی) و موازی هستند، که در ادامه بطور کامل تعریف می‌شوند. در کنار تجهیز سیستم‌ها، قابلیت اطمینان آن‌ها به‌طور جدی مطرح بوده و جزء لاینفک عملکرد آن‌ها به‌شمار می‌رود. بنابراین می‌توان گفت برآورد قابلیت اطمینان سیستم اگر چه به‌ظاهر کم اهمیت است ولی چون در ابعاد وسیع اعمال می‌گردد، صرفه‌جویی زیادی در هزینه دارد. تاکنون برای برآورد قابلیت اطمینان سیستم‌ها روش‌های زیادی ارائه شده است. ما نیز در ادامه روند پیدا نمودن روشی مناسب برای برآورد قابلیت اطمینان سیستم‌ها، از چند روش آماری استفاده کرده، و کاربرد علم آمار را در علوم صنعتی بیان و روش‌های برآوردی که تلفیقی از علم آمار و علوم مهندسی است ارائه می‌دهیم. هدف یافتن برآورد بیزی مناسب قابلیت اطمینان سیستم‌های موازی، در مدل‌های فشار-مقاومت است. در این مدل‌ها فرض می‌کنیم مولفه‌ها از هم مستقل و تحت فشار یکسان هستند. لذا قابلیت اطمینان سیستم‌ها را به دو روش ماکزیمم درستنمایی و بیزی، به‌طوری که مولفه‌ها دارای طول عمری، تحت توزیع‌های از قبیل نمایی تک پارامتری، دو پارامتری، وایبل، گاما، پاراتو باشند را به دست می‌آوریم. مطالب پایان‌نامه، به اختصار، شامل موارد زیر می‌باشد:

- در فصل اول، مفاهیم و تعاریف اولیه را معرفی می‌کنیم که شامل سیستم‌ها، قابلیت اطمینان، الگوریتم‌های نمونه‌گیری مونت کارلو و روش‌های برآوردیابی ماکزیمم درستنمایی و بیزی می‌باشد.
- در فصل دوم، با فرض اینکه سیستم‌ها سری و مولفه‌های سیستم دارای طول عمری با توابع توزیع نمایی دو پارامتری، گاما، وایبل و پاراتو هستند، برآورد ماکزیمم درستنمایی و بیزی قابلیت اطمینان سیستم‌ها را به‌دست آورده و همچنین توزیع برآوردگر را نیز مشخص می‌کنیم.
- در فصل سوم، قابلیت اطمینان سیستم‌های موازی را با شرط مستقل بودن مولفه‌ها و اینکه تحت فشار یکسان قرار دارند، به دو روش ماکزیمم درستنمایی و بیزی، به‌گونه‌ای که مولفه‌ها دارای طول عمری با توزیع نمایی تک پارامتری و وایبل هستند، به‌دست آورده، همچنین با استفاده روش‌های نمونه‌گیری مونت کارلویی زنجیر مارکوفی MCMC نتایج را شبیه‌سازی می‌کنیم.
- در فصل چهارم، همانند فصل سوم برآورد قابلیت اطمینان سیستم‌های موازی در مدل‌های فشار-مقاومت را به دو روش ماکزیمم درستنمایی و بیزی محاسبه می‌کنیم، با این تفاوت که در این فصل فرض می‌کنیم مولفه‌ها دارای طول عمری با توزیع نمایی دو پارامتری هستند. همچنین با استفاده از روش‌های نمونه‌گیری بیان شده نتایج حاصله را شبیه‌سازی می‌کنیم.
- پیوست آ شامل تعاریف لازم و پیوست ب شامل کدهای نوشته‌شده در محیط نرم‌افزار R، است.

فهرست نشانه‌ها و نمادها

X	متغیر تصادفی
k	تعداد مولفه‌های سیستم
t	متغیر شکست
θ	پارامتر نمونه
Θ	فضای پارامتری
$\mathbf{X}$	بردار تصادفی
$F(t)$	تابع طول عمر
$R(t)$	تابع قابلیت اطمینان
$f(t)$	تابع چگالی طول عمر
$\frac{f(t)}{R(t)}$	تابع نرخ خطر
Δt	فاصله زمانی
$\pi(\theta y)$	توزیع پسین
$\pi(\theta)$	توزیع پیشین
$MCMC$	روش‌های نمونه‌گیری زنجیر مارکوفی مونت کارلویی
$K(. .)$	هسته زنجیر مارکوف
$L(\theta)$	تابع درستمایی
$\hat{R}$	برآورد ماکزیمم درستمایی
$\hat{R}_b$	برآوردگر بیزی
$R(\theta, \delta)$	تابع ریسک
$L(\theta, d)$	تابع زیان
$\delta(X)$	آماره آزمون
μ^*	ماکزیمم μ_i ها
$G(y)$	تابع بقای سیستم‌های موازی
$H(z)$	تابع بقای سیستم‌های سری
Λ	ماتریس قطری
I_{ij}	ماتریس اطلاع فیشر
$\xrightarrow{D}$	همگرایی در توزیع
$\xrightarrow{P}$	همگرایی در احتمال

$\approx$ تقریباً مساوی
 AN توزیع نرمال مجانبی

فهرست مطالب

۱	تعاریف و مفاهیم اولیه	۱
۱	تاریخچه	۱.۱
۲	مقدمه	۲.۱
۳	استنباط آماری	۳.۱
۴	دیدگاه بسامدی	۱.۳.۱
۵	دیدگاه بیزی در استنباط آماری	۲.۳.۱
۶	معرفی برخی از توزیع های پیشین	۳.۳.۱
۷	دیدگاه بیزی و روش های مونت کارلو	۴.۱
۸	روش های مونت کارلو	۱.۴.۱
۹	انتگرال گیری مونت کارلو	۲.۴.۱
۱۰	تحلیل بقا	۵.۱
۱۰	تابع قابلیت اطمینان	۶.۱
۱۱	خواص تابع قابلیت اطمینان	۱.۶.۱
۱۱	مقاسیه توابع قابلیت اطمینان	۲.۶.۱
۱۲	تابع نرخ خطر	۷.۱
۱۳	سیستم های سری و موازی	۸.۱
۱۴	برآورد قابلیت اطمینان	۹.۱
۱۴	برآورد پارامتری	۱.۹.۱
۱۵	مدل فشار-مقاومت	۱۰.۱
۱۵	قابلیت اطمینان در مدل های فشار-مقاومت	۱.۱۰.۱
۱۶	روش ماکزیمم درست نمایی (روش MLE)	۲.۱۰.۱
۱۹	برآورد بیزی	۱۱.۱
۱۹	برآورد نقطه ای	۱.۱۱.۱
۲۱	برآورد فاصله ای	۲.۱۱.۱
۲۱	قابلیت اطمینان و دیدگاه بیزی	۱۲.۱

۲۳	۲	برآورد قابلیت اطمینان سیستم‌های سری در مدل‌های فشار-مقاومت
۲۳	۱.۲	مقدمه
۲۳	۲.۲	مدل فشار-مقاومت نمایی دو پارامتری
۲۳	۱.۲.۲	قابلیت اطمینان سیستم‌های سری
۲۵	۲.۲.۲	برآورد ماکزیمم درست‌نمایی
۲۵	۳.۲.۲	برآورد بیزی
۲۶	۴.۲.۲	توزیع برآوردگر
۲۷	۳.۲	مدل فشار-مقاومت توزیع گاما
۲۸	۱.۳.۲	قابلیت اطمینان سیستم‌های سری
۲۹	۲.۳.۲	برآورد ماکزیمم درست‌نمایی
۲۹	۳.۳.۲	برآورد بیزی
۳۰	۴.۳.۲	توزیع برآوردگر
۳۱	۴.۲	مدل فشار-مقاومت توزیع وایبل
۳۱	۱.۴.۲	قابلیت اطمینان سیستم‌های سری
۳۱	۲.۴.۲	برآورد ماکزیمم درست‌نمایی
۳۲	۳.۴.۲	توزیع برآوردگر
۳۳	۴.۴.۲	برآورد بیزی
۳۳	۵.۴.۲	مدل نمایی
۳۴	۵.۲	مدل فشار-مقاومت توزیع پاراتو
۳۴	۱.۵.۲	قابلیت اطمینان سیستم‌های سری
۳۵	۲.۵.۲	برآورد ماکزیمم درست‌نمایی
۳۵	۳.۵.۲	برآورد بیزی
۳۶	۴.۵.۲	توزیع برآوردگر
۴۱	۳	برآورد قابلیت اطمینان سیستم‌های موازی
۴۱	۱.۳	مدل فشار-مقاومت توزیع نمایی
۴۲	۱.۱.۳	برآورد ماکزیمم درست‌نمایی
۴۳	۲.۱.۳	برآورد بیزی
۴۴	۲.۳	مدل فشار-مقاومت توزیع وایبل
۴۶	۱.۲.۳	برآورد ماکزیمم درست‌نمایی
۴۷	۲.۲.۳	برآورد بیزی
۴۸	۳.۳	مطالعه شبیه‌سازی
۴۸	۱.۳.۳	توزیع نمایی
۵۱	۲.۳.۳	توزیع وایبل

۵۵	۴	برآورد قابلیت اطمینان سیستم تحت توزیع نمایی دو پارامتری
۵۵	۱۰۴	مدل فشار-مقاومت نمایی دو پارامتری
۵۵	۱۰۱۰۴	قابلیت اطمینان سیستم‌های موازی
۵۷	۲۰۱۰۴	برآورد ماکزیمم درست‌نمایی
۵۸	۳۰۱۰۴	برآورد بیزی
۵۹	۴۰۱۰۴	مطالعه شبیه‌سازی
۶۷	آ	نمادها و تعاریف‌ها
۶۷	۱.آ	توزیع‌های پارامتری مهم در قابلیت اطمینان
۷۰	۲.آ	توزیع مجانبی
۷۵	ب	کد برنامه R برای اجرای فصل سوم و چهارم
۸۹		مراجع
۹۳		واژه‌نامه فارسی به انگلیسی
۹۵		واژه‌نامه انگلیسی به فارسی
۹۶		نمایه

فصل ۱

تعاریف و مفاهیم اولیه

۱.۱ تاریخچه

در قرون وسطی به دلیل سادگی فرآیند تولید، هر کارگر می‌توانست تمام قسمت‌های یک کالا را به تنهایی بسازد. از این رو لذت حاصل از تولید کل کالا به جایی جزئی از آن کافی بود تا کارگر وقت بیشتری را برای رسیدن به کیفیت بالای کالا صرف نماید. انقلاب صنعتی باعث کاهش این انگیزه شد. چرا که دیگر کارگر برخلاف گذشته، سازنده یک کالا نبود بلکه تنها جزء کوچکی از فرآیند ساخت آن را بر عهده داشت. در انقلاب صنعتی، روسای کارخانجات بزرگ نمی‌توانستند شخصا بر تمام وقایع نظارت داشته باشند. بنابراین ناچار بودند به طریق دیگری مشکلات را حل نمایند. این امر به منظور حفظ منافع اقتصادی و ایمنی مصرف‌کننده و نیز افزایش میزان تولید و به وجود آمدن رقابت مورد توجه جدی قرار گرفت و به این منظور به کارگیری روش‌های بازرسی برای جلوگیری از عرضه محصولات نامرغوب یا معیوب به بازار به سرعت گسترش یافت. حتی بسیاری از واحدهای تولیدی به منظور اطمینان خاطر مصرف‌کنندگان و گاه به عنوان ابزارهای تبلیغاتی اعلام می‌کردند که در تولید خود از روش‌های بازرسی صد در صد بهره می‌برند. بنابراین اولین مرحله کنترل کیفیت پدیدار شد که هدف از آن فقط جداسازی محصولات معیوب از سالم بود و به منظور کاهش تعداد محصولات معیوب، ابداع روش‌های علمی جدیدتر ضرورت یافت. در سال ۱۹۹۴ دکتر والتر شوهارت^۱ آمریکایی اولین نمودارهای آماری را به منظور کنترل فرآیند تولید ابداع و معرفی نمود. بنابراین وی را پایه‌گذار کنترل کیفیت آماری می‌شناسند. ولی استفاده از علم آمار در صنعت از این زمان آغاز نشد. علت این امر اعتقاد نداشتن مدیران تولید به روش‌های آماری و همچنین کمبود متخصص علم آمار در مراکز تولیدی بود. در سال ۱۹۹۴ در طی جنگ جهانی دوم خرید میلیون‌ها تن مواد غذایی، مهمات، پوشاک، دارو و ... توسط ارتش آمریکا بدون آنکه روش علمی برای کنترل بازرسی آن وجود داشته باشد سران ارتش آمریکا را وادار نمود که به سراغ کنترل کیفیت آماری بروند. این اقدام ارتش آمریکا از یک طرف منجر به پیروزی آن در جنگ جهانی دوم و از طرف دیگر سبب پایه‌گذاری علم کنترل کیفیت آماری در دنیا گردید. در کنترل کیفیت آماری که امروزه به طور گسترده‌ای

^۱Walter.Shewhar

در صنایع پیشرفته دنیا به کار می رود. سعی بر این است که ضایعات تولید تا حد امکان کاهش یابد. چنانچه مدیری بخواهد کنترل کیفیت را اجرا کند باید آمار بگیرد. به عنوان مثال تعیین نماید که آیا ضایعات کارخانه یا کارگاه نسبت به دیروز افزایش یافته است یا کاهش؟ آیا ضایعات این کامیون مواد خام نسبت به کامیون قبلی فرق دارد یا خیر؟ سپس لازم است این آمارها را به صورت نمودار در آورده و بعد از تجزیه و تحلیل نمودارها ریشه نقایص را بیابد. به این ترتیب می توان با استفاده از کنترل کیفیت آماری کنترل موثری بر تولید داشت. در سال ۱۹۹۴ اولین حلقه های کنترل کیفیت برای بهبود روش های کنترل کیفیت در ژاپن تشکیل شد. این حلقه ها عبارتست از تقسیم مجموعه عوامل موثر بر کیفیت به حلقه های مختلف و تقسیم حلقه های بزرگ به حلقه های کوچک؛ هدف اصلی از تشکیل این حلقه ها آن بود که شرایط مناسب برای بهبود کیفیت محصول تولیدی فراهم شود. کار حلقه های کنترل کیفیت از همان آغاز با موفقیت های چشمگیری رو به رو شد و به همین دلیل واحدهای تولیدی در بسیاری از کشورها این حلقه ها را در صنعت خود تشکیل دادند. در سالهای اخیر در اکثر واحدهای تولیدی، سیستمهای مختلف کنترل کیفیت تضمین کننده سلامت، بهداشت، رفاه و ... افراد جامعه موجود است. اولین دیدگاههای تئوریک برای تحلیل بقا^۲ (طول عمر) که سریعاً توسعه یافت و ویژگیهای کاربردی خویش را نشان داد مبتنی بر روش های کلاسیک مرسوم به «جدول طول عمر» که توسط کاپلان و مایر^۳ در سال ۱۹۸۵ ارائه شد. آن ها در بررسی خویش با توجه به داده های ناتمام موجود در مطالعات بقا روش ناپارامتری را برای برآورد احتمالات بقا ارائه نمودند. این روش ها سریعاً مورد توجه سایر آمار شناسان قرار گرفت. از جمله این افراد اسوردراپ^۴ بود که در سال ۱۹۶۵، نه تنها از این روش ها استفاده نمود بلکه برای اولین بار ضعف های این روش را نیز مطرح نموده و استفاده از مدل های مبنی بر فرآیندهای تصادفی را پیشنهاد نمود. به دنبال اسوردراپ، هوم^۵ نیز در سال ۱۹۶۵ طی مقاله معروف خود ایده هایی را در مورد استفاده از فرآیندهای تصادفی در مدل های تحلیل بقا مطرح نمود، او برای اولین بار امکان استفاده از فرآیندهای زنجیر مارکف را اعلام نمود.

۲.۱ مقدمه

در راستای برآورد قابلیت اطمینان سیستم ها و به دست آوردن توزیع مناسب برآوردگرها، مقالات و کتب بسیاری توسط محققان و نویسندگان علوم مختلف به چاپ رسیده، و چون قابلیت اطمینان سیستم ها جزء لاینفک و وابسته به کارایی و میزان رضایت بخشی یک سیستم است، لذا تحلیل و بررسی این امر و دستیابی به قابلیت اطمینان بالاتر همچنان یکی از موضوعات مهم روز است. بنابر مستندات در راستای برآورد یابی، دریک تحقیق، چرچ و هریس^۶ (۱۹۷۰) یک راه حل پارامتری برای مسئله برآورد

^۲Survival

^۳Kaplan and Meier

^۴Sverdrup

^۵Hoem

^۶Chur and Harris

قابلیت اطمینان تحت فرضیه‌ای که توزیع Y مشخص باشد و $F_X(x)$ و $G_Y(y)$ هر دو نرمال باشند ارائه دادند. همچنین ثابت کردند که قابلیت اطمینان دارای توزیع مجانبی است و توزیع مجانبی $\hat{R}$ را می‌توان از جداول توزیع مجانبی استاندارد بدست آورد. باتاچاریا و جانسون (۱۹۷۴)^۷ برآورد نااریب با کمترین واریانس^۸ ($MVUE$) را برای قابلیت اطمینان سیستم، هنگامیکه مولفه‌های فشار و مقاومت هر دو تحت توزیع نمایی با پارامترهای نامعلوم باشند، را بررسی کردند. نظریه نمونه بزرگ در بخش سوم این پایان‌نامه نشان می‌دهد که برآوردهای بیزی و درست‌نمایی قابلیت اطمینان برای بطور مجانبی هم ارز هستند. آنها رابطه بین آن دو را به ازای n های مختلف و با استفاده از روش‌های انتگرال‌گیری عددی مونت کارلو، شبیه‌سازی کردند. در ادامه این روند دراپر و گاتمن (۱۹۷۸)^۹ یک رهیافت بیزی برای استنباط درباره قابلیت اطمینان سیستم‌های چند عضوی s عضو خارج از k عضو ارائه دادند. آنها فرضیات خود را اینگونه بیان کردند که مقاومت‌ها و فشار وارده مستقل بوده و دارای توزیع نمایی باشند. نهایتاً نتایج بدست آمده را با یک مثال عددی توضیح دادند. همچنین ویر (۱۹۸۱)^{۱۰} برآورد ماکزیمم درست‌نمایی و برآورد بیزی قابلیت اطمینان را در حالتی که سیستم تنها دارای دو متغیر باشد، برای هر دو توزیع پیشین دقیق و توأم مورد بحث و بررسی قرار داد.

مسئله برآورد قابلیت اطمینان سیستم دو عضوی، هنگامیکه (X_1, X_2) مولفه‌های مقاومت بوده و تحت فشار X_3 دارای توزیع نمایی باشد توسط هنگال (۱۹۹۶)^{۱۱} بررسی شد.

یک مرور کلی و جامع از مدل‌های فشار-مقاومت توسط کاتز و همکاران (۲۰۰۳)^{۱۲} جمع‌آوری شد.

۳.۱ استنباط آماری

استنباط آماری، یک فرآیند استنتاجی در مورد ویژگی‌های یک جامعه، شامل آزمون فرضیه و برآوردیابی است. جامعه شامل همه داده‌های مشاهده شده است. از آنجایی که دسترسی یا تحلیل همه مشاهدات جامعه، به دلیل حجم بالای آن، مشکل است، استنباط آماری براساس یک نمونه معقول از جامعه انجام و نتایج حاصل از آن به جامعه مورد نظر تعمیم داده می‌شود. بنابراین، هدف استنباط آماری تعمیم نتایج حاصل از مشاهدات نمونه‌ای به جامعه مورد نظر است، برای این‌که بتوانیم نتیجه‌ای درست از استنباط آماری داشته باشیم، رعایت یک نکته ضروری است، و آن انتخاب مدل واقعی یا مدل نزدیک به فرآیند تولید داده‌ها است. چرا که عمده مشکلات در استنباط آماری به انتخاب مدل نادرست مربوط می‌شود در استنباط آماری مکاتب فلسفی متفاوتی وجود دارند. این مکاتب یا الگوها، منحصر به فرد نیستند و روشی که تحت یک الگو روی مشاهدات به خوبی اجرا می‌گردد غالباً تفسیر متفاوتی نسبت به سایر

^۷Bhattacharya and Johnson

^۸Minimum-Variance unbiased Estimator

^۹Draper and Guttman

^{۱۰}Weier

^{۱۱}Hanagel

^{۱۲}Kotz, Lumelskii and Pensky

الگوها روی نمونه مورد نظر دارد. بندی اپاتای و فورستر^{۱۳} (۲۰۱۱) چهار الگو را برای استنباط آماری معرفی می‌کنند:

• آمار کلاسیک یا آمار بسامدی

• آمار بیزی

• آمار مبتنی بر درست‌نمایی

• آمار مبتنی بر معیار آکاییک

ما در تحلیل و نگارش این پایان‌نامه بیشتر با دو روش بیزی و درست‌نمایی سروکار داریم. در استنباط آماری براساس پارامتر مورد نظر دو دیدگاه وجود دارد: دیدگاه بسامدی، دیدگاه بیزی. در زیر مختصر شرحی از آن‌ها بیان می‌کنیم.

۱.۳.۱ دیدگاه بسامدی

در مبحث بسامدی استنباط آماری، پارامتر مطلوب و مورد بررسی ثابت اما نامعلوم است و احتمال به عنوان فراوانی نسبی رخداد پیشامد مورد نظر به ازای تکرار زیاد آزمایش تصادفی، روی مقادیر فضای نمونه و با توجه به پارامتر نامعلوم، تفسیر می‌شود. در این مکتب، معمولاً عملکرد هر فرآیند آماری با میانگین‌گیری روی کل فضای نمونه تعیین می‌شود. این میانگین‌گیری می‌تواند قبل از اجرای آزمایش انجام شود و به داده‌ها بستگی ندارد که باعث شده است به این دیدگاه، از این منظر، نقدهای وارد شود. به طور کلی، دو دیدگاه در استنباط بسامدی وجود دارد:

* فیشر (۱۹۲۶)، استنباط آماری را مبتنی درست‌نمایی می‌دانست. تابع درست‌نمایی ضابطه مشابهی همچون چگالی توام نمونه دارد، با این تفاوت که، مشاهدات ثابت در نظر گرفته می‌شوند و پارامترها مجاز به تغییر روی همه مقادیر ممکن هستند. او پیچیدگی داده‌ها را با کمک آماره‌های بسنده، کاهش داد به طوری که این آماره‌ها همه اطلاعات موجود در داده‌ها، در مورد پارامتر مورد بررسی را دارند (کسلو برگر ۲۰۰۲). فیشر همچنین نظریه برآوردیابی ماکزیم درست‌نمایی^{۱۴} (MLE) را ایجاد و توزیع جانبی برآوردگرهای حاصل از این روش برآوردیابی را معرفی کرد (فیشر ۱۹۲۶). او کارایی یک برآوردگر را با استفاده از معیاری به نام اطلاع فیشر که در ادامه به طور کامل بیان شده، اندازه‌گیری کرد.

* نیمن و پیرسون (۱۹۳۳). پایه‌های نظریه تصمیم را ایجاد کردند و سپس استنباط آماری را در آن نهادند. نظریه آن‌ها بر خلاف فیشر اساساً قیاسی بود. آن‌ها نواحی اطمینان را معرفی و سعی کردند تا راه‌حل‌های بهینه‌ای در رده‌های مجاز بیابند (موری و همکاران، ۲۰۰۶)^{۱۵} و اگر نیاز بود رده‌ها را محدود می‌کردند تا بتوانند راه‌حلی بیابند.

^{۱۳}Bandyopadhyay, S.P. and Forster

^{۱۴}Maximum Likelihood Estimation

^{۱۵}Murray and Associates

۲.۳.۱ دیدگاه بیزی در استنباط آماری

دیدگاه بیزی استنباط آماری، برای اولین بار توسط یک کشیش به نام توماس بیز^{۱۶} در سال ۱۷۶۳ مطرح و پس از مرگ او توسط دوستش، ریچارد پرایس، در مجله‌ای علمی در انجمن سلطنتی انگلستان به چاپ رسید (کسلو برگر، ۲۰۰۴). استنباط بیزی یکی از روش‌های حل مسائل آنالیز بقا است، و بسیاری از مسئله‌های واقعی به‌طور طبیعی به‌وسیله انتگرال‌ها با نسبت توزیع پسین فرمول بندی می‌شود. قضیه بیز که بدون شک یکی از کاربردی‌ترین مسائل در دیدگاه بیزی است، بنابر احتمالی شرطی بیان می‌شود و روشی را برای به‌روز کردن احتمال رخ داده‌های مشاهده نشده که در ارتباط با رویداد دیگری رخ داده‌اند، ارائه می‌دهد. یعنی با توجه به وقوع رخداد مرتبط، یک احتمال پیشین برای رخداد مشاهده نشده داریم، سپس با جمع‌آوری اطلاعات جدید احتمال پسین را به روز می‌کنیم. درآمار بیزی، قضیه بیز به عنوان اساس استنباط بیزی درباره پارامتر نامعلوم توزیع آماری استفاده می‌شود. لزومات قضیه بیز را به‌صورت زیر بیان می‌کنیم:

- از آنجایی که مقدار واقعی پارامتر نامعلوم است، در آمار بیزی فرض می‌شود این مقدار واقعی تحقیقی از یک توزیع احتمالی است. به عبارتی، پارامتر، یک متغیر تصادفی محسوب می‌شود. و این همان تضادی است که دیدگاه بیزی با دیدگاه بسامدی دارد، چرا که در دیدگاه بسامدی پارامتر نامعلوم ثابت فرض می‌شود. استفاده از استنباط بیزی مستلزم سه نوع توزیع است. این توزیع‌ها عبارتند از توزیع پیشین، توزیع پسین و توزیع توام متغیرها. قضیه بیز وظیفه به روز کردن استنباط درباره پارامتر واقعی، با توجه به اطلاعات مشاهده‌شده را بر عهده دارد. بنابراین توزیع احتمال پیشینی داریم که توسط آن شانس مشاهده هر مقدار ممکن پارامتر قبل از مشاهده داده‌ها، به گونه‌ای که قابل توجیه باشد، محاسبه می‌شود. توزیع احتمال پیشین اختیاری است، زیرا هر کس دیگری می‌تواند اعتقاد پیشینی در مورد پارامتر نامعلوم داشته باشد.

- هرگزازه احتمالی درباره پارامتر باید به عنوان «درجه باور» فرد درباره آن پارامتر تفسیر گردد. از قاعده احتمالی به طور مستقیم استفاده می‌کنیم تا استنباط درباره پارامتر را با توجه به داده‌های مشاهده شده انجام دهیم.

- قضیه بیز دو منبع اطلاعاتی درباره مقدار پارامتر نامعلوم را ترکیب می‌کند: اطلاعات پیشین و داده‌های مشاهده شده. اطلاعات پیشین، در قالب یک توزیع پیشین، به عنوان وزن و باور نسبی از مقدار واقعی پارامتر قبل از مشاهده داده‌ها تعیین می‌شود. اطلاعات موجود در داده‌های مشاهده شده نیز توسط تابع درستنمایی (تابع توزیع توام) خلاصه می‌شود. قضیه بیز این دو منبع را ترکیب می‌کند و توزیع پسین را که وزن نسبی باورمان از مقدار پارامتر بعد از مشاهده داده‌ها می‌باشد را ارائه می‌دهد.

با توجه به داده‌هایی که داریم، قضیه بیز تنها روش سازگار با تغییر باورمان در مورد پارامتر است. استنباط بیزی فقط وابسته به داده‌های است که مشاهده شده‌اند، نه آن‌هایی که می‌توانند رخ بدهند اما رخ نداده‌اند. بنابراین استنباط بیزی با اصل درستنمایی سازگار است (برگر و ولپرت ۱۹۹۸). خاصیت اصلی استنباط بیزی این است که پسین همیشه بایک روش واحد و خودکار مشخص می‌شود: اطلاعات اولیه درباره پارامتر را از چگالی پیشین و اطلاعات موجود در داده‌های مشاهده شده را توسط تابع درستنمایی باهم

^{۱۶}Bayes, Thomas

ترکیب می‌کند و چگالی پسین را می‌سازد.

توزیع پیشین: توزیع پیشین تبلور کاربر آماراز خلاصه اطلاعات و دانسته‌های او در این باره است که احتمال قرار داشتن پارامتر در چه بخش‌های از فضای پارامتری بیش از همه است. به بیان دیگر قبل از مشاهده و جمع‌آوری هرگونه داده‌ای، اطلاعات قبلی کاربر وی را متقاعد می‌کند که شانس قرار گرفتن پارامتر مورد نظر در فضای پارامتری چگونه است و کاربر می‌تواند اطلاعات خود را در قالب یک تابع توزیع بیان کند.

توزیع پسین: به توزیع احتمال پارامترها بعد از احتساب اثر مشاهدات توزیع پسین گفته می‌شود. توزیع پسین با استفاده از قاعده بیز و توزیع توام متغیرها به صورت زیر تعیین می‌شود. فرض می‌کنیم θ مقدار مشاهد شده یک متغیر تصادفی، مثلاً w با تابع چگالی $g(\theta)$ است، که از آن با نام توزیع پیشین w یاد می‌کنیم در این حالت، $f_{\theta}(\cdot)$ را می‌توان به عنوان چگالی شرطی متغیر X به شرط $W = \theta$ تلقی کرده و چگالی توام آن‌ها را از رابطه زیر به دست آورد.

$$f_{X,W}(x, \theta) = f_{X|W=\theta}(x) \pi(\theta|x)$$

توزیع شرطی W به شرط $X=x$ قابل محاسبه است، توزیع پسین می‌نامیم. توزیع پسین W دارای چگالی پسین است که با استفاده از رابطه فوق و قاعده بیز به صورت زیر قابل محاسبه است.

$$\pi_{W|X=x}(\theta) = \frac{f_{X,W}(x, \theta)}{f_X(x)}$$

بنابراین تابع چگالی پسین را می‌توان بصورت زیر باز نویسی کرد.

$$\pi_{W|X=x}(\theta) = \pi(\theta|x) = \frac{f_{\theta}(x) \pi(\theta)}{\int_{\Theta} f_{\theta}(x) \partial G(\theta)}$$

بنابراین، یکی مهم‌ترین چالش‌ها در استنباط بیزی محاسبه تابع چگالی پسین است. به ویژه آن‌که معمولاً نمی‌توان فرض کرد که مشاهدات مصداقی از یک نمونه تصادفی از توزیع نرمال باشند. برای برخورد با این چالش، روش‌های عددی و نمونه‌گیری مختلفی معرفی شده‌اند (گلماند و اسمیت، ۱۹۹۰؛ رابرت و کسلا، ۲۰۰۲؛ تیرنی و همکاران، ۱۹۸۶). در ادامه چند نوع از توابع پیشین به کار برده شده در این پایان‌نامه را بیان می‌کنیم

۳.۳.۱ معرفی برخی از توزیع‌های پیشین

تعریف ۱.۳.۱. توزیع پیشین مزدوج: درمبحث توابع توزیع، خانواده‌های وجود دارند که توزیع پیشین و پسین پارامترها هر دو متعلق به یک خانواده مشترک می‌باشند. به این نوع توزیع‌ها، توزیع پیشین مزدوج می‌گویند. یک روش برای انتخاب توزیع پیشین مزدوج، بر پایه ساختار توزیع توام $X = (X_1, X_2, \dots, X_n)$ بر حسب تابعی از آماره بسنده است.

قضیه ۲.۳.۱: اگر بردار $x = (x_1, x_2, x_3, \dots, x_n)$ یک نمونه تصادفی از خانواده نمایی با تابع

درست‌نمایی به شکل زیر باشد.

$$L(x_i; \theta) = \prod_{i=1}^n h(x_i) c(\theta) \text{Exp}\left(\sum_{i=1}^k w_i(\theta) t_i(x)\right)$$

آنگاه خانواده توزیع های مزدوج برای فضای پارامتر Θ به صورت زیر است.

$$L(\theta; \tau) = \frac{1}{k(\tau)} c(\theta)^{\tau} \text{Exp}\left(\sum_{i=1}^k w_i(\theta) t_i(x)\right)$$

که در آن τ در رابطه زیر صدق می کند.

$$k(\tau) = \int_{\theta} c(\theta)^{\tau} \text{Exp}\left(\sum_{i=1}^k w_i(\theta) t_i(x)\right)$$

تعریف ۳.۳.۱. توزیع پیشین کم‌آگاهی بخش : در مسائلی که پژوهشگر اطلاعات قابل اتکایی در اختیار ندارد و یا حجم اطلاعات ناچیز باشد (مانند فضای پارامتر اعداد حقیقی مثبت است) می‌توان از توزیع پیشین کم‌آگاهی بخش استفاده نمود. ساده ترین و قدیمی ترین روش تعیین توزیع پیشین کم‌آگاهی بخش، اصل عدم تفاوت است. براساس این روش برای تمام برآمدهای ممکن، پارامتر احتمالات مساوی در نظر گرفته می‌شود. نتایج حاصل از استنباط بیزی برپایه اصل عدم تفاوت، تفاوت چندانی با روش های معمول (برای مثال روش ماکزیم درست‌نمایی) استنباط آماری ندارند.

تعریف ۴.۳.۱. توزیع پیشین آگاهی بخش : توزیع پیشین آگاهی بخش به طور معین و کاملی یک اطلاع خاص در ارتباط با یک پارامتر بیان می کند. اساس اطلاع اضافی توزیع پیشین تجربیات قبلی پژوهشگر است. لازم بذکر است که این مسئله با در نظر گرفتن یک فرض تفاوت اساسی دارد بنابراین می توان به اطلاع اضافی، شهود موجود برای مسئله گفت. بنابراین اگر تعیین توزیع پیشین بر اساس شهود موجود در پژوهش صورت پذیرد به طوری که توزیع پیشنهادی در برگیرنده اطلاعات قبلی باشد آنگاه از یک توزیع پیشین آگاهی بخش استفاده شده است.

۴.۱ دیدگاه بیزی و روش‌های مونت کارلو

اغلب در مسائل بیز سلسه مراتبی، فرم بسته‌ای برای توزیع پسین وجود ندارد. در این حالت برای استفاده از دیدگاه بیز در استنباط آماری لازم است تا از مدل بیز تجربی استفاده شود که بر مبنای نمونه ای از توزیع پسین شکل می‌گیرد. برای تولید نمونه از یک توزیع بدون فرم بسته الگوریتم هایی وجود دارد که مبنای اصلی آنها توزیع های شرطی است به این دسته از الگوریتم ها روش های مونت کارلو گفته می‌شود که به اختصار با نماد $(MCMC)$ نمایش داده می‌شود.

۱.۴.۱ روش‌های مونت‌کارلو

روش‌های نمونه‌گیری مونت‌کارلوی زنجیر مارکوفی^{۱۷} ($MCMC$) اولین بار توسط متروپولیس و همکاران (۱۹۵۳) به عنوان روشی برای شبیه‌سازی کارا از سطوح انرژی اتم‌ها در ساختارهای بلوری مطرح شد. به دنبال آن در استنباط آماری، به ویژه استنباط بیزی، هستینگز (۱۹۷۰) این ایده را پرورش داد و براساس نظریه زنجیرهای مارکوف، ویژگی روش‌های نمونه‌گیری $MCMC$ را به طور گسترده تحت مطالعه قرار داد. در بخش‌های قبلی اشاره شد که توزیع پسین در بیشتر مواقع دارای شکل بسته نیست. لذا برای انجام محاسبات نیاز به تقریب دارد. گاهی توزیع پسین آنقدر پیچیده است که تولید مستقیم نمونه از آن مشکل است. به همین دلیل از روش‌های تولید نمونه به‌طور غیر مستقیم استفاده می‌کنیم. یک روش غیر مستقیم برای نمونه‌گیری از $\pi(\theta|y)$ ساخت زنجیر مارکوف بازگشتی^{۱۸} و تحویل‌ناپذیری^{۱۹} با توزیع مانای $\pi(\theta|y)$ است. فرض کنید چنین زنجیری وجود دارد. بنابراین اگر، در اجرا، زنجیر به اندازه کافی بزرگ باشد، آنگاه مقادیر شبیه‌سازی شده از زنجیر را می‌توان به عنوان نمونه‌ای از توزیع پسین در نظر گرفت. در اجرای روش‌های نمونه‌گیری $MCMC$ ، مطالب مهمی از جمله مکانیسم تولید تحقق برای زنجیر، تعداد زنجیرها و طول زنجیرها و انتخاب نقطه شروع زنجیر مطرح هستند. برای کسب اطلاعات بیشتر در مورد این مطالب به اسمیت (۱۹۸۴)، گرین (۱۹۹۵)، آسموسن (۲۰۰۷) و ریکی (۲۰۱۰) مراجعه کنید. نظریه اصلی و زیربنایی روش‌های $MCMC$ این است که هر زنجیر بازگشتی و تحویل‌ناپذیر، هسته انتقال واحدی خواهد داشت، و هسته انتقال t مرحله‌ای، وقتی $t \rightarrow \infty$ ، به توزیع مانای خود همگرا می‌گردد. بنابراین برای تولید زنجیر با هسته انتقال K و توزیع مانای π باید شرط $\pi k = \pi$ برقرار باشد. فرض کنید $\theta \sim \pi$ ، اگر $\theta' \sim K(\theta'|\theta)$ آنگاه $\theta' \sim \pi$. چنین زنجیرهایی در شرط زیر، معروف به شرط تعادل دقیق^{۲۰}، صدق می‌کنند:

$$\pi(\theta|y)k(\theta'|\theta) = \pi(\theta'|\theta)k(\theta|\theta') \quad (1.1)$$

براساس هسته انتقال، الگوریتم‌های برای روش‌های نمونه‌گیری $MCMC$ ، وجود دارند که می‌توان به الگوریتم‌های نمونه‌گیری متروپولیس-هستینگز (متروپولیس و همکاران، ۱۹۵۳؛ هستینگز، ۱۹۷۰) و الگوریتم نمونه‌گیری گیبز (گمان، ۱۹۸۴؛ گلفاند و اسمیت، ۱۹۹۰) اشاره کرد. به‌طور کلی الگوریتم $MCMC$ را می‌توان به صورت جدول زیر ارائه کرد.

با توجه به جدول ۱.۱، احتمال پذیرش الگوریتم $MCMC$ را با P مشخص می‌کنیم. در این صورت با احتمال p ، θ به عنوان نمونه‌ای از توزیع $\pi(\theta|y)$ پذیرفته می‌شود. سپس می‌توانیم با توجه به خروجی الگوریتم که مجموعه‌ای از θ^i هارا تشکیل می‌دهد، انتگرال مونت‌کارلویی یا توزیع پسین را تقریب بزنیم. برای هر الگوریتم $MCMC$ مقدار P تغییر می‌کند. مثلاً برای الگوریتم متروپولیس - هستینگز مقدار احتمال پذیرش برابر است با

$$P(\theta, \theta^{(i-1)}) = \min \left[1, \frac{\pi(\theta|y)}{\pi(\theta^{(i-1)}|y)} \frac{K(\theta^{(i-1)}|\theta)}{K(\theta|\theta^{(i-1)})} \right] \quad (2.1)$$

^{۱۷}Markov Chain Monte Carlo^{۱۸}Recurrent^{۱۹}Irreducible^{۲۰}Detailed balance conditon

الگوریتم نمونه‌گیری MCMC
برای $i = 1$ تا T تکرار کن
۱- هسته مارکوف $k(\cdot \cdot)$ با توزیع مانای $\pi(\cdot)$ را ایجاد کن
۲- زنجیر مارکوف $\theta \sim \pi \rightarrow \theta^{(i)}$ را تولید کن
۳- θ را از هسته مارکوف $k(\theta \theta^{(i-1)})$ تولید کن
۴- قرار بده $\theta = \begin{cases} \theta^{(i)} & \text{با احتمال } p \\ \theta^{(i-1)} & \text{با احتمال } 1-p \end{cases}$

جدول ۱.۱: الگوریتم نمونه‌گیری MCMC

۲.۴.۱ انتگرال‌گیری مونت کارلو

روش مونت کارلو در اصل یک روش بسط داده شده توسط فیزیکدانان می‌باشد که از تولید اعداد تصادفی برای محاسبه انتگرال‌های پیچیده استفاده می‌کند. فرض کنید هدف محاسبه انتگرال $\int_a^b h(X)d(X)$ است. اگر بتوانیم $h(X)$ را به صورت حاصلرب یک تابع $P(X)$ و یک تابع چگالی احتمال $f(X)$ روی فاصله (a, b) در نظر بگیریم آنگاه:

$$\int_a^b h(x)d(x) = \int_a^b f(x)P(x)d(x) = E_{f(x)}[P(X)]$$

به عبارت دیگر این انتگرال می‌تواند به عنوان امید ریاضی تابع $P(X)$ روی چگالی $f(X)$ تعبیر می‌شود. اگر نمونه تصادفی $(x_1, x_2, \dots, x_n)$ از چگالی $f(X)$ استخراج شود آنگاه:

$$\int_a^b h(X)d(x) = E_{f(X)}[P(X)] \approx \frac{1}{n} \sum_{i=1}^n P(X)$$

روش اخیر انتگرال‌گیری مونت کارلو نام دارد و برای تقریب توزیع‌های پسین مورد نیاز در تحلیل بیز به کار گرفته می‌شود. مثلاً اگر هدف محاسبه امید توزیع پسین

$$E(\theta) = \int \theta p(\theta|y)d\theta$$

باشد و بتوان پارامترهای $\theta_1, \theta_2, \dots, \theta_n$ را از توزیع پسین شبیه‌سازی کرد. آنگاه می‌توان امید ریاضی توزیع پسین را با متوسط این نمونه تقریب زد یعنی

$$E(\theta) \approx \frac{1}{n} \sum \theta_i$$

که به آن برآورد مونت کارلو θ گویند. لازم به ذکر است طبق قانون ضعیف اعداد بزرگ اگر $n \rightarrow \infty$ آنگاه:

$$\frac{1}{n} \sum \theta_i \xrightarrow{p} E(\theta)$$

۵.۱ تحلیل بقا

آنالیز یا تحلیل بقا، یکی از مباحث مهم علم آمار است که در رشته‌های مختلفی از جمله علوم کامپوتری، کشاورزی، مهندسی و غیره کاربرد دارد. تحلیل بقا به مجموعه‌ای از روش‌های آماری تحلیل داده گفته می‌شود که در آن‌ها متغیر مطلوب، زمان وقوع یک پدیده است. این موضوع در علوم مهندسی نظریه قابلیت اطمینان نامیده می‌شود. تحلیل بقا به عنوان یک روش آماری که اساساً برای تحقیقات زیستی و مهندسی توسعه یافته، می‌تواند در آنالیز داده‌های طول عمر مورد استفاده قرار گیرد. آنالیز و تحلیل بقا براساس استفاده از توزیع و توابع خاص می‌باشد.

تابع بقا

تابع بقا یا طول عمر. بیانگر این است که یک مولفه حداقل تا زمان T ماندگاری داشته باشد. براین اساس تابع تجمعی طول عمر به صورت زیر است

$$F(t) = P(T \leq t)$$

۶.۱ تابع قابلیت اطمینان

احتمال موفقیت، یا احتمال اینکه سیستم بدون وقوع خرابی به وظایف تعیین شده در طراحی عمل می‌کند را قابلیت اطمینان^{۲۱} نامند. تعریف قابلیت اطمینان بر تعریف وقوع خرابی بنا شده است. برای اندازه‌گیری قابلیت اطمینان یک سیستم ابتدا سیستم به اجزایی شکسته می‌شود و قابلیت اطمینان سیستم بر حسب قابلیت اطمینان اجزای آن بیان می‌گردد. هر متغیر تصادفی نامنفی را یک متغیر تصادفی طول عمر^{۲۲} می‌نامند. فرض کنید X یک متغیر طول عمر (مقاومت) باشد. تابع قابلیت اطمینان متغیر X در نقطه t به صورت زیر تعریف می‌شود.

$$R(t) = P(X > t) = 1 - F(t) \quad \forall t > 0$$

در واقع تابع $R(t)$ احتمال این‌که فرد یا قطعه دارای طول عمر بیشتر از t باشد، را محاسبه می‌کند. به عنوان مثال هرگاه $R(1) = 0.7$ باشد، بدین معنی است که با احتمال ۰/۷ فری یا قطعه طول عمری بیشتر از یک واحد زمانی خواهد داشت.

هرگاه $F(x)$ تابع توزیع تجمعی احتمال X باشد داریم:

$$R(t) = P(X > t) = 1 - P(x < t) = 1 - F(t). \quad (3.1)$$

همچنین اگر $f(x)$ تابع چگالی احتمال X باشد داریم:

$$R(t) = P(X > t) = \int_t^{\infty} f(x) dx \quad (4.1)$$

^{۲۱} Reliability

^{۲۲} lifetime random variabl

با مشتق‌گیری نسبت به t از رابطه (۱.۴) داریم:

$$f(t) = -R'(t)$$

بنابراین با مشخص شدن یکی از توابع $R(t)$ ، $F(t)$ ، و $f(t)$ سایر توابع نیز به دست می‌آید.

۱.۶.۱ خواص تابع قابلیت اطمینان

هر تابع قابلیت اطمینان در سه شرط زیر صدق می‌کند:

$$R(0) = 1 \quad (۱)$$

$$R(\infty) = 0 \quad (۲)$$

(۳) $R(t)$ نسبت به t به طور یکنواخت نزولی است.
زیرا:

$$R(0) = P(X > 0) = \int_0^{+\infty} f(x)dx = \int_{-\infty}^0 d(x) + \int_0^{+\infty} f(x)dx = \int_{-\infty}^{+\infty} f(x)dx = 1 \quad (۵.۱)$$

از طرفی داریم

$$R(\infty) = \lim_{t \rightarrow \infty} R(t) = \lim_{x \rightarrow \infty} (1 - F(t)), \quad \lim_{x \rightarrow \infty} F(t) = 1, \quad R(\infty) = 1 - 1 = 0$$

فرض کنید $t_1 < t_2$ باشد در این صورت:

$$R(t_2) - R(t_1) = (1 - F(t_1)) - (1 - F(t_2)) = F(t_1) - F(t_2) \leq 0$$

لذا تابع $R(t)$ تابعی به طور یکنواخت نزولی است.

۲.۶.۱ مقایسه توابع قابلیت اطمینان

بدیهی است که در انتخاب بین دو کالا یا قطعه، قطعه‌ای انتخاب می‌شود که دارای قابلیت اطمینان بیشتری باشد. در نمودار فوق همواره تابع قابلیت اطمینان قطعه (۲) از قطعه (۱) بیشتر است به عبارت دیگر

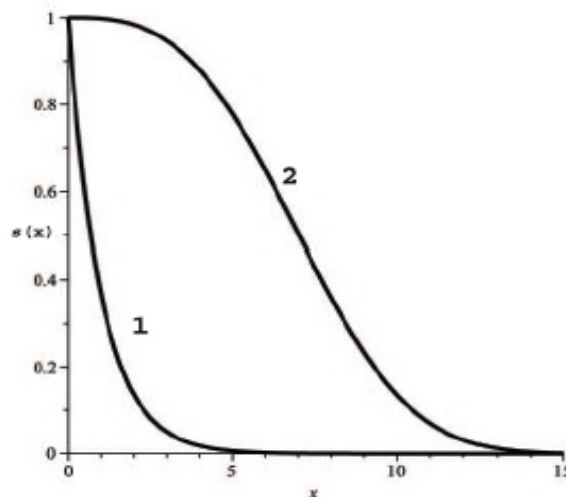
$$R_1(t) < R_2(t) \quad \forall t > 0$$

گاهی اوقات $R_1(t)$ و $R_2(t)$ یکدیگر را قطع می‌کنند. در این حالت نمی‌توان به صراحت مشخص کرد که کدام قطعه بهتر است. یکی از روش‌های معقول برای مقایسه، استفاده از امید ریاضی طول عمر قطعات است به عبارت دیگر هر مولفه‌ای (مقاومت) که دارای بیشترین مقدار میانگین طول عمر ^{۲۳} $MTTF$ باشد بهتر است. در مباحث قابلیت اطمینان، امید ریاضی طول عمر را $MTTF$ می‌نامند.

$$MTTF = E(X) = \int_0^{\infty} f(x)dx$$

به دلیل عدم نیاز به این مبحث در این پایان نامه از ذکر جزئیات چشم پوشی کرده و به تعریف آن بسنده می‌کنیم.

^{۲۳}Mean Time To Failure



شکل ۱.۱: نمودار مقایسه توابع قابلیت اطمینان

۷.۱ تابع نرخ خطر

تابع نرخ خطر یا همان تابع طول عمر اساس و پایه اصلی در برآورد قابلیت اطمینان یک سیستم است. فرض کنید X متغیر تصادفی طول عمر باشد، که تا زمان t عمر کرده باشد یعنی پیشامد $X > t$ رخ داده است. Δt را فاصله کوتاه در نظر بگیرید بنابراین $P(t < X < t + \Delta t | X > t)$ احتمال اینکه فرد یا قطعه در فاصله زمانی Δt بعد از t از کار بیفتد را محاسبه می‌کند. کمیت

$$\frac{P(t < X < t + \Delta t | X > t)}{\Delta t}$$

نرخ از کار افتادگی یا از بین رفتن را بر حسب واحد زمانی محاسبه می‌کند. فرض کنید $\Delta t \rightarrow 0$ در این صورت داریم:

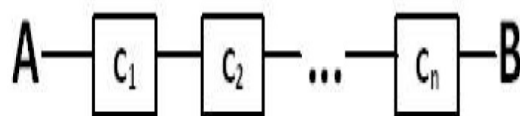
$$\begin{aligned} \lim_{\Delta t \rightarrow 0} \frac{P(t < X < t + \Delta t | X > t)}{\Delta t} &= \lim_{\Delta t \rightarrow 0} \frac{P(t < X < t + \Delta t)}{\Delta t P(X > t)} \\ &= \frac{1}{P(X > t)} \lim_{\Delta t \rightarrow 0} \frac{P(t < X < t + \Delta t)}{\Delta t} \\ &= \frac{1}{R(t)} \lim_{\Delta t \rightarrow 0} \frac{P(X > t) - P(X > t + \Delta t)}{\Delta t} \\ &= \frac{1}{R(t)} \lim_{\Delta t \rightarrow 0} \frac{R(t) - R(t + \Delta t)}{\Delta t} \\ &= \frac{1}{R(t)} \lim_{\Delta t \rightarrow 0} \frac{R(t + \Delta t) - R(t)}{\Delta t} \\ &= -\frac{R'(t)}{R(t)} = \frac{f(t)}{R(t)} \end{aligned}$$

تعریف فوق، تعریف کلی و از دید مهندسی، تابع نرخ خطر است. در این پایان‌نامه تابع نرخ خطر به صورت دیگر، که در ادامه خواهیم دید، بیان می‌شود. به عبارتی می‌توان گفت تعریف آماری نرخ خطر

مورد نظر است.

۸.۱ سیستم‌های سری و موازی

سیستم‌ها براساس نحوه اتصال مولفه‌ها (مقاومت‌ها) به چند نوع تقسیم می‌شوند، ولی هدف ما دو نوع سری^{۲۴} و موازی^{۲۵} است، و ما با این دو نوع سروکار داریم. زمانی که موفقیت یک سیستم وابسته به موفقیت تمام مولفه‌ها باشد سیستم را سری می‌نامند، که نحوه اتصال مولفه‌ها به صورت شکل زیر می‌باشد و با ازکار افتادن اولین مولفه سیستم متوقف می‌شود. در ادامه نحوی اتصال مولفه‌ها (مقاومت‌ها) را در یک سیستم سری مشاهده می‌کنیم. در سیستم‌های سری قابلیت اطمینان سیستم برابر است با حاصل ضرب



شکل ۲.۱: نحوی اتصال مولفه‌ها در سیستم‌های سری

قابلیت اطمینان‌های همه مولفه‌ها. با توجه به تعریف قابلیت اطمینان داریم:

$$\varphi(X_1, X_2, \dots, X_n) = 1 \iff X_1 = 1, \dots, X_n = 1 \iff \prod_{i=1}^n X_i = 1$$

$$\varphi(X_1, X_2, \dots, X_n) = 0 \iff \exists X_k = 0 \iff \prod_{i=1}^n X_i = 0$$

$$\implies R_{series}(X_1, \dots, X_n) = E(\varphi(X_1, X_2, \dots, X_n)) = E\left(\prod_{i=1}^n X_i\right) = \prod_{i=1}^n E(X_i) = \prod_{i=1}^n R_i$$

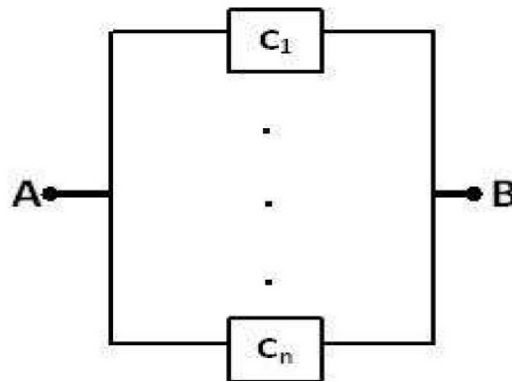
اگر موفقیت یک سیستم با موفقیت حداقل یک مولفه همراه باشد، گوییم مولفه‌ها به طور موازی به هم متصل شده‌اند، و سیستم را موازی نامیم. در سیستم‌ها موازی سیستم با یک مولفه درست، کارایی دارد و نیاز به درست بودن همه مولفه‌ها نیست. نحوه اتصال مولفه‌ها در سیستم‌های موازی به صورت شکل زیر است.

در این صورت:

$$\varphi(X_1, X_2, \dots, X_n) = 0 \iff X_1 = 0, \dots, X_n = 0 \iff 1 - X_1 = 1, \dots, 1 - X_n = 1$$

^{۲۴}Series

^{۲۵}Parallel



شکل ۳.۱: نجوی اتصال مولفه‌ها در سیستم‌های موازی

$$\Leftrightarrow \prod_{i=1}^n (1 - X_i) = 1 \Leftrightarrow 1 - \prod_{i=1}^n (1 - X_i) = 0$$

$$\varphi(X_1, X_2, \dots, X_n) = 1 \Leftrightarrow \prod_{i=1}^n (1 - X_i) = 1$$

$$\Rightarrow R_{Parallel}(X_1, \dots, X_n) = E(\varphi(X_1, X_2, \dots, X_n)) = 1 - \prod_{i=1}^n (1 - E(X_i)) = 1 - \prod_{i=1}^n (1 - R_i)$$

در حالت کلی می‌توان اثبات کرد

$$R_{series}(X_1, \dots, X_n) = \min(X_1, \dots, X_n)$$

$$R_{Parallel}(X_1, \dots, X_n) = \max(X_1, X_2, \dots, X_n)$$

۹.۱ برآورد قابلیت اطمینان

فرض کنید نمونه‌ای تصادفی به حجم n از جامعه $X_1, X_2, \dots, X_n$ از متغیر تصادفی طول عمر X با تابع اطمینان $R(t)$ در اختیار باشد. هدف برآورد $R(t)$ می‌باشد که در آن t زمان مشخص و معلوم است. معمولاً برآورد به دو شکل کلی تقسیم می‌شود:

الف) پارامتری: فرض می‌شود که $R(t)$ دارای شکل خاصی است مانند نمایی، وایبل، گاما و.... در این حالت کفایت پارامترهای جامعه برآورد شوند.

ب) ناپارامتری: هیچ گونه ساختار خاصی برای $R(t)$ در نظر گرفته نمی‌شود.

۱.۹.۱ برآورد پارامتری

روش‌های متفاوتی برای برآورد پارامتری تابع قابلیت اطمینان وجود دارد مانند روش گشتاوری، روش ماکزیمم درست‌نمایی، روش بیزی و... در این پایان نامه دو روش ماکزیمم درست‌نمایی و روش بیزی را

اسفاده می‌کنیم. بدلیل کاررایی و استفاده زیاد از روش ماکزیمم درستنمایی در ادامه مختصر توضیح از این روش ارائه می‌کنیم.

۱۰.۱ مدل فشار-مقاومت

در مدل‌های فشار-مقاومت^{۲۶} سیستم‌هایی هستند که در اثر مقاومتی که دارند (هر عضو سیستم دارای یک مقاومت ذاتی است) عمر بیشتری می‌کنند، این سیستم‌ها فقط یک سطح معینی از فشار را تحمل می‌کنند، یعنی اگر فشار روی سیستم از این حد معین و مجاز فراتر رود سیستم قادر به تحمل این فشار نخواهد بود و از کار می‌افتد.

فشار: متغیر وادار کردن به شکست است. چیزی که باعث شکست یک عضو، یک وسیله و یا یک کالا می‌شود، که به شکل‌های بار مکانیکی، محیط زیست، درجه حرارت و بار الکتریکی و غیره می‌تواند ظاهر شود.

مقاومت: توانایی یک عضو، یک وسیله و یا یک کالا را برای به انجام رساندن وظیفه خود به صورت رضایت بخش و بدون شکست در مقابل بار خارجی و محیط زیست مقاومت نامیده می‌شود. بنابراین مقاومت متغیر پایداری در مقابل شکست است.

تغییرات در «فشار» و «مقاومت» سیستم‌ها و میزان مقاومت مولفه‌های یک سیستم در برابر فشار وارد شده به آن‌ها باعث ایجاد توزیع آماری و پراکندگی طبیعی می‌شود، بنابراین یک سیستم دارای دو توزیع است، اگر مولفه‌ها دارای توابع چگالی احتمال معلوم باشند، قابلیت اطمینان یک عضو و قابلیت اطمینان کل سیستم به راحتی قابل دسترسی و محاسبه است. فرض کنید X و Y دو متغیر تصادفی باشند بطوریکه X «مقاومت» و Y متغیر «فشار» را نشان دهد، در این صورت اگر X و Y دارای چگالی احتمال توام $f(x, y)$ باشند، آنگاه قابلیت اطمینان سیستم به صورت زیر است

$$R = P[Y < X] = \int_{-\infty}^{+\infty} \int_{-\infty}^x f(x, y) dy dx$$

که در آن $P(Y < X)$ ، احتمال اینکه مقاومت بیشتر از فشار باشد را نشان می‌دهد و $f(x, y)$ تابع چگالی احتمال توام X و Y است.

۱.۱۰.۱ قابلیت اطمینان در مدل‌های فشار-مقاومت

همان‌گونه که از اسم مدل‌ها پیداست، اساس اصلی این مدل‌ها فشار وارده بر مولفه‌های سیستم می‌باشد. هدف اصلی ما در این پایان نامه نیز برآورد قابلیت اطمینان سیستم‌ها (سری و موازی) در این مدل‌ها می‌باشد، زمانی که مولفه‌ها همگی تحت فشار یکسان قرار گیرند. در این مدل‌ها مولفه‌ها مستقل از هم هستند و فشار متغیر است. فرض کنید در یک سیستم موازی دارای k مولفه مستقل، $Y_1, Y_2, Y_3, \dots, Y_k$ همگی تحت فشار یکسان X باشند. همان‌طور که قبلاً بیان کردیم در سیستم‌های موازی تا زمانی که

^{۲۶}Stress-Strength

تمامی مولفه‌ها از کار نیفتد سیستم کار می‌کند، بنابراین قابلیت اطمینان سیستم را به صورت زیر محاسبه می‌کنیم.

$$R = P[X < \max(Y_1, Y_2, Y_3, \dots, Y_k)]$$

یعنی احتمال اینکه فشار وارده کمتر از بیشترین مقاومت باشد. در سیستم‌های سری چون با از کار افتادن اولین مقاومت سیستم از کار می‌افتد بنابراین در رابطه فوق بجای $\max$ از $\min$ استفاده می‌کنیم.

$$R = P[X < \min(Y_1, Y_2, Y_3, \dots, Y_k)]$$

۲.۱۰.۱ روش ماکزیمم درستنمایی (روش MLE)

روش ماکزیمم درستنمایی یکی از قدیمیترین و بااهمیت‌ترین روش‌ها در نظریه برآوردهاست، این روش اولین بار توسط گوس در سال ۱۸۲۱ به کار گرفته شد و پس از آن به صورت گسترده‌تری در سال ۱۹۲۶ توسط فیشر در انتقاد از روش برآورد گشتاوری مورد استفاده قرار گرفت.

تابع درستنمایی

فرض کنید $X = (X_1, X_2, X_3, \dots, X_n)$ برداری از متغیرهای تصادفی با تابع چگالی جرم احتمال $f_\theta(X)$ باشند به طوری که $\theta \in \Theta$. برای هر مقدار داده شده $X = x$ ، تابع درستنمایی X را تابع چگالی احتمال توأم X ، یعنی $f_\theta(X)$ ، تعریف می‌کنیم که به صورت تابعی از θ در نظر گرفته می‌شود و آن را با نماد $L(\theta)$ نمایش می‌دهیم:

$$L(\theta) = f_\theta(X)$$

برای هر نمونه از مشاهده $x_1, x_2, x_3, \dots, x_n$ ، تابع درستنمایی $L(\theta)$ تابعی حقیقی مقدار که بر روی فضای پارامتر Θ ، تعریف می‌شود. توجه کنید که در اینجا فرض هم توزیع بودن و مستقل بودن متغیرهای تصادفی $X_1, X_2, X_3, \dots, X_n$ الزامی نمی‌باشد.

برآورد ماکزیمم درستنمایی

فرض کنید که برای یک نمونه $x_1, x_2, x_3, \dots, x_n$ ، تابع $L(\theta; X)$ بیشترین مقدار خود را بر اساس تابعی از θ روی فضای پارامتر Θ در $\theta = S(x)$ می‌گیرد، یعنی:

$$\sup_{\theta \in \Theta} L(\theta) = L[S(x)]; S(x) \in \Theta$$

سپس آماره $\hat{\theta} = S(x)$ ، برآورد ماکزیمم درستنمایی برای پارامتر θ نامیده می‌شود. توجه: اغلب در ماکزیمم کردن تابع $L(\theta)$ راحت‌تر است که Ln آن ماکزیمم شود زیرا مقداری که تابع

$L(\theta)$ را ماکزیمم کند آنگاه تابع $Ln(L)$ ، را نیز ماکزیمم می‌کند. از آنجایی که تابع $L(\theta)$ اغلب به صورت حاصلضرب می‌باشد و در نتیجه تابع $Ln(L)$ حاصلضرب را به مجموع تبدیل می‌کند و انتگرال‌گیری مجموع، نسبت به انتگرال‌گیری حاصلضرب آسان‌تر است. و تابع $Ln(L)$ را لگاریتم تابع درستنمایی می‌نامند که به صورت زیر است.

$$Ln(L) = l(\mathbf{X}|\theta) = \sum_{i=1}^n Ln[f(x_i|\theta)]$$

برای بدست آوردن برآورد ماکزیمم درستنمایی دو روش مجزا وجود دارد.

۱. ماکزیمم کردن مستقیم: امتحان کردن $L(\theta)$ ، بطور مستقیم، بطوریکه مقداری از θ را بیابیم که تابع $L(\theta)$ را برای نمونه مفروض $x_1, x_2, x_3, \dots, x_n$ ماکزیمم کند. این روش برای مواقعی که دامنه داده‌ها به پارامتر وابسته است مفید می‌باشد.

۲. معادلات درستنمایی: اگر دامنه داده‌ها به پارامتر وابسته نباشد، آنگاه فضای پارامتر Θ ، یک فضای باز نامیده می‌شود، و آنگاه تابع درستنمایی نسبت به پارامترها، روی فضای پارامتر Θ مشتق پذیر می‌باشد و برآوردگر ماکزیمم درستنمایی $\hat{\theta}$ در معادلات زیر صدق می‌کند.

$$\frac{\partial}{\partial \theta_k} Ln(L(\theta)) = 0; k = 1, 2, \dots, p$$

این معادلات، معادلات درستنمایی نامیده می‌شوند.

برآوردگرهای ماکزیمم درستنمایی دارای خواص مجانبی خاص می‌باشند. از این خواص می‌توانیم برای مشخص نمودن توزیع برآوردگرها استفاده کرد. فرض کنید $\hat{\beta}$ برآوردگر ماکزیمم درستنمایی پارامتر β باشد، می‌توان نشان داد که تحت شرایط خاصی

$$\hat{\beta} \sim N\left(\beta, \frac{1}{I(\beta)}\right)$$

نکته: در حالت کلی می‌توان گفت هر تابعی از $\hat{\beta}$ در توزیع به $N\left(\beta, \frac{1}{I(\beta)}\right)$ میل می‌کند. که در آن

$$I(\beta) = E\left(\frac{d}{d\beta} \log L(\beta)\right)^2 = E\left(\frac{d}{d\beta} l(\beta)\right)^2$$

تابع $I(\beta)$ به اطلاع فیشر معروف است، که در قسمت پیوست همین پایان به طور کامل بیان شده است.

اصل درستنمایی

فرض کنید $L(\theta)$ ، تابع درستنمایی برای θ بر اساس بردار $\mathbf{X}$ و $L^*(\theta)$ تابع درستنمایی برای θ بر اساس بردار $\mathbf{X}^*$ باشند. در صورتی که $L(\theta) = kL^*(\theta)$ (برای k هایی که به θ وابسته نباشند)، آنگاه باید استنباط یکسانی برای پارامتر θ بر اساس دو بردار مشاهده شده، داشته باشیم.

در برآورد نقطه‌ای اصل درستنمایی ایجاب می‌کند که اگر $\mathbf{X}$ و $\mathbf{X}^*$ دو بردار از متغیرها با توابع درستنمایی $L(\theta)$ و $L^*(\theta)$ ، و $L(\theta) = kL^*(\theta)$ (برای همه k ها) باشند. آنگاه دو برآورد نقطه‌ای درستنمایی ماکزیمم $\hat{\theta}$ و $\hat{\theta}^*$ باید با هم مساوی باشند.

قضیه ۱.۱۰.۱. قضیه سازگاری برآوردهای ماکزیمم درستنمایی: فرض کنید $x_1, x_2, x_3, \dots, x_n$ ، یک نمونه تصادفی مستقل و هم توزیع با تابع چگالی احتمال $f_\theta(x)$ باشند، و $L(\theta; \mathbf{X}) = \prod_{i=1}^n f_\theta(x_i)$ تابع درستنمایی، $\hat{\theta}$ برآوردهای ماکزیمم درستنمایی برای θ باشند، اگر $\tau(\theta)$ تابعی پیوسته از θ باشد، در این صورت، برای هر $\epsilon > 0$ و هر $\theta \in \Theta$ داریم:

$$\lim_{n \rightarrow \infty} P_\theta(|\tau(\hat{\theta}) - \tau(\theta)| \geq \epsilon) = 0$$

یعنی، $\tau(\hat{\theta})$ یک برآوردهای سازگار برای $\tau(\theta)$ است.

برهان. به پارسیان (۱۳۸۵) رجوع شود. $\square$

اگر $\hat{\theta}$ برآورد کننده به روش ماکزیمم درستنمایی برای θ باشد و $n \mapsto \infty$ ، آنگاه واریانس $\hat{\theta}$ برابر حد پایین نامساوی کرامر - رائو خواهد شد. نامساوی کرامر - رائو در فصل دوم به طور کامل بیان خواهد شد.

خاصیت پایایی برآوردهای ماکزیمم درستنمایی

بدون شک یک از مفیدترین خواص برآوردهای ماکزیمم درستنمایی، خاصیت پایایی می باشد. از این رو خاصیت پایایی یکی از پرکاربردترین مسائل در نگارش این پایان نامه است. خاصیت پایایی برآوردهای درسنمایی در حالت کلی در سال ۱۹۶۶ توسط زهنا اثبات شد.

فرض کنید توزیع مورد نظر ما دارای پارامتر θ باشد، و ما به برآورد ماکزیمم درستنمایی تابعی از θ مانند $\tau(\theta)$ ، نیاز داشته باشیم. خاصیت پایایی برآوردهای ماکزیمم درستنمایی بیان می کند که اگر $\hat{\theta}$ برآوردهای ماکزیمم درستنمایی برای پارامتر θ باشد، آنگاه $\tau(\hat{\theta})$ برآوردهای ماکزیمم درستنمایی برای $\tau(\theta)$ است.

اگر نگاشت $\tau(\theta) \mapsto \theta$ یک به یک باشد (یعنی، برای هر مقدار θ فقط یک مقدار یکتای $\tau(\theta)$ وجود داشته باشد و برعکس) هیچ مشکلی در استفاده از این قضیه وجود ندارد. در این نمونه، به آسانی می توان نشان داد که هیچ اختلافی بین اینکه ما تابع درستنمایی را بر حسب تابعی از θ ماکزیمم کنیم یا بر حسب تابعی از $\tau(\theta)$ ماکزیمم کنیم وجود ندارد و در هر دو روش ما پاسخ یکسانی را بدست می آوریم. فرض می کنیم که $\eta = \tau(\theta)$ ، سپس معکوس تابع $\tau^{-1}(\eta) = \theta$ بنابراین تابع درستنمایی بر حسب تابعی از η به شکل زیر است.

$$L^*(\eta; \mathbf{X}) = \prod_{i=1}^n f_{x_i}(\tau^{-1}(\eta)) = L(\tau^{-1}(\eta); \mathbf{X})$$

و

$$\sup_{\eta} L^*(\eta; \mathbf{X}) = \sup_{\eta} L(\tau^{-1}(\eta); \mathbf{X}) = \sup_{\theta} L(\theta; \mathbf{X})$$

بنابراین، تابع $L^*(\eta; \mathbf{X})$ به ازای $\eta = \tau(\theta) = \tau(\hat{\theta})$ ماکزیمم می شود. پس برآورد ماکزیمم درستنمایی برای $\tau(\theta)$ برابر $\tau(\hat{\theta})$ خواهد بود.

استفاده از خاصیت پایایی همیشه امکان پذیر نیست، گاهی اوقات ممکن است بسیاری از توابعی که مورد نیاز ما هستند یک به یک نباشند. مثلاً برای برآورد θ^2 (مربع میانگین توزیع نرمال) نداشت $\theta^2 \mapsto \theta$ یک به یک نیست،

اگر تابع $\tau(\theta)$ یک به یک نباشد، آنگاه به ازای یک مقدار معلوم η ممکن است که بیش از یک مقدار برای θ وجود داشته باشد بطوریکه $\tau(\theta) = \eta$ است. تابع درستنمایی L^* را به صورت زیر تعریف می کنیم.

$$L^*(\eta; \mathbf{X}) = \sup_{[\theta: \tau(\theta)=\eta]} L(\theta; \mathbf{X}) \quad (6.1)$$

مقدار $\hat{\eta}$ که $L^*(\eta; x)$ را ماکزیم می کند، برآوردگر ماکزیم درستنمایی $\eta = \tau(\theta)$ می باشد، و از رابطه (۱) مشخص است که ماکزیم L و L^* بر هم منطبق هستند.

۱۱.۱ برآورد بیزی

در حالت کلی استنباط بیزی دادای دو نوع برآورد است: ۱- برآورد نقطه ای ۲- برآورد فاصله ای.

۱.۱۱.۱ برآورد نقطه ای

در برآورد بیزی می خواهیم بر مبنای مشاهداتمان از یک متغیر تصادفی، متغیر تصادفی دیگری را تخمین بزنیم. هدف از برآورد نقطه ای، یافتن مقداری برای مقدار نامعلوم تابعی، مانند g ، تعریف شده بر روی فضای پارامتر در نقطه θ است. در تمام شاخه های مرتبط با برآورد نقطه ای، یافتن یک برآوردگر مناسب منجر به معیاری برای برآورد نقطه ای شده است. اما می توان به عنوان دسته بندی روش ها و معیارهایی که منجر به یک برآوردگر مناسب می شوند یک یا چند ملاک را به طور کلی در نظر گرفت. در این بخش مختصری از نحوه بدست آوردن یک برآوردگر نقطه ای مناسب با شرح داده می شود. برای ورود به مسئله یافتن یک برآوردگر مناسب ابتدا لازم است تعریف دقیقی از برآوردگر ارائه شود.

تعریف ۱.۱۱.۱. فرض کنید $X_1, X_2, \dots, X_n$ نمونه ای تصادفی از متغیرها با خانواده توابع چگالی احتمال $\{f(\theta; x), \theta \in \Theta\}$ باشد، در برآورد نقطه ای، هدف پیدا کردن آماره $\sigma(X_1, X_2, \dots, X_n)$ به گونه ای است که اگر $x_1, x_2, \dots, x_n$ مقادیر مشاهده شده، نمونه $X_1, X_2, \dots, X_n$ باشد. آنگاه مقدار آماره $\sigma(x_1, x_2, \dots, x_n)$ یک برآورد نقطه ای خوب باشد. بنابراین برآوردگر پارامتر θ ، برپایه نمونه تصادفی $X_1, X_2, \dots, X_n$ عبارت از آماره $\sigma(X_1, X_2, \dots, X_n)$ است که مقدار برآورد شده θ را براساس مقادیر ممکن $X_1, X_2, \dots, X_n$ مشخص می کند. چون پارامتر θ متعلق به مجموعه Θ است، انتظار داریم که مقادیر برآوردگر یعنی $\sigma(x)$ نیز متعلق به Θ باشد. یعنی

$$\forall \theta \in \Theta \quad P_\theta(\sigma(X_1, X_2, \dots, X_n) \in \Theta) = 1$$

در حالت کلی تر، تابع حقیقی مقدار $\delta(X)$ ، را برآوردگر تابعی از پارامترها مانند $g(\theta)$ می نامند. هرچند همواره انتظار می رود مقدار $\delta(X)$ به مقدار نامعلوم $g(\theta)$ نزدیک باشد. اما نیاز به چنین شرطی

در تعریف برآوردگر نیست. بدین ترتیب می توان انتظار داشت مقدار $\delta(X)$ در محدوده مقادیر ممکن $g(\theta)$ قرار داشته باشد.

اصولا می توان میزان نزدیکی $\delta(X)$ به مقدار $g(\theta)$ ، در مقدار d را، با تابع زیان $L(\theta, d)$ در نظر گرفت. به طوری که برای هر θ و d داریم: $L(\theta, g(\theta)) = 0$ و $L(\theta, d) \geq 0$. بنابراین نزدیکی یا دوری برآوردگر $\delta(X)$ با تابع زیان اندازه گیری می شود. برای غلبه بر تصادفی بودن Δ از متوسط زیان می توان استفاده نمود که تابع ریسک R نامیده می شود.

$$R(\theta, \delta) = E_{\theta}(L[\theta, \delta(X)])$$

حال می توان برآوردگری را جستجو نمود که تابع ریسک را به ازای هر θ مینیمم کند. واضح است که بهترین برآوردگر به طور یکنواخت با کمترین ریسک وجود ندارد. برای ساده تر کردن مسئله می توان یک کلاس از برآوردگرها با شرط نارایی یا پایایی در نظر گرفت.

فارغ از در نظر گرفتن چنین محدودیت هایی همواره بدنال برآوردگری هستیم که تابع ریسک $R(\theta, \delta)$ را مینیمم کند. بنابراین مسئله به صورت یافتن برآوردگر δ که منجر به کمترین مقدار تابع زیان وزن دار شده می شود.

$$r(\theta, \delta) = \int R(\theta, \delta) d\Lambda(\theta); \quad \int d\Lambda(\theta) = 1$$

برآوردگری که دارای چنین شرطی باشد برآوردگر بیز نامیده می شود. یافتن برآوردگر بیز را می توان به عنوان یک ابزار ریاضی که منجر به کمترین متوسط ریسک می شود یا یک روش برای بکارگیری تجربیات گذشته و اعمال یک گرایش خاص ذهنی خاص در حالتی که هیچ گونه اطلاعاتی وجود ندارد در نظر گرفت. یافتن برآوردگر بیز به سادگی امکان پذیر است. بدین منظور ابتدا فرض کنید قبل از جمع آوری هرگونه مشاهده ای پارامتر نامعلوم θ دارای توزیع احتمالی به شکل Λ است و برآوردگر بیز $g(\theta)$ مقدار d است به طوری که منجر به مینیمم شدن مقدار $E(L(\theta, d))$ می شود. پس از جمع آوری مشاهدات توزیع احتمال Λ با توزیع پسین جایگزین می شود. برآوردگر بیز هر تابع به صورت $\delta(X)$ است به طوریکه منجر به مینیمم شدن $E(L(\theta, \delta(x))|x)$ می شود. نهایتاً تعریف برآوردگر بیز را در قالب قضیه زیر بیان می کنیم.

قضیه ۲.۱۱.۱. فرض کنید پارامتر θ به عنوان یک متغیر تصادفی دارای تابع توزیع Λ است و $\Theta =$ و X دارای توزیع احتمال $f_{\theta}(X)$ باشد. همچنین فرض کنید حداقل یک برآوردگر با ریسک متناهی وجود دارد اگر به ازای تمام مقادیر X برآوردگر $\delta_{\Lambda}(X)$ وجود داشته باشد که دارای کمترین ریسک $E(L(\theta, \delta(x))|x)$ است در این صورت $\delta_{\Lambda}(X)$ یک برآوردگر بیز است.

□

برهان. رجوع شود به کتاب آمار ریاضی پاریسیان.

لم ۳.۱۱.۱. فرض کنید فرضیات قضیه ۲.۱۱.۱ برقرار است آنگاه

$$\delta_{\Lambda}(X) = E(g(\theta)|x) \text{ آنگاه } L(\theta, d) = (d - g(\theta))^2$$

$$\delta_{\Lambda}(X) = \frac{\int w(\theta)g(\theta)d\Lambda(\theta|x)}{\int w(\theta)d\Lambda(\theta|x)} \text{ آنگاه } L(\theta, d) = w(\theta)(d - g(\theta))^2$$

$$\delta_{\Lambda}(X) = |d - g(\theta)| \text{ آنگاه } L(\theta, d) = |d - g(\theta)| \text{ اگر } \delta_{\Lambda}(X) \text{ برابر است با میانه توزیع پسین}$$

اگر $L(\theta, d) = \begin{cases} 0; & |d - g(\theta)| < c \\ 1; & |d - g(\theta)| \geq c \end{cases}$ آنگاه $\delta_\lambda(X)$ نقطه وسط بازه I به طول $2c$ است به طوری که منجر به مینیم شدن مقدار $P(\theta \in I|x)$ می شود.

□

برهان. رجوع شود به کتاب آمار ریاضی پارسیان.

۲.۱۱.۱ برآورد فاصله‌ای

برآورد فاصله‌ای یا فاصله اطمینان برای یک پارامتر، محدودی است که با استفاده از آماره مورد نظر و با سطح اطمینان بالایی پارامتر جامعه را در بر می‌گیرد. ایجاد یک فاصله اطمینان یا برآورد فاصله‌ای برای یک پارامتر به عوامل زیر وابسته است:

۱- سطح اطمینانی که محقق را مجاب کند که فاصله ایجاد شده، پارامتر را در بر می‌گیرد.

۲- توزیع نمونه‌ای آماره که به کمک آن ضریب اعتماد تعیین می‌شود.

۳- واریانس آماره یا همان خطای معیار نمونه‌گیری (SE)

۴- حجم نمونه‌ای که برای برآورد انتخاب می‌شود. (n)

۱۲.۱ قابلیت اطمینان و دیدگاه بیزی

مدت زمانی که یک قطعه از یک سیستم کار می‌کند و یا پاسخ به این سوال که آیا قطعه کمتر از طول عمر استاندارد یا بیشتر از طول عمر استاندارد کار می‌کند، در قابلیت اطمینان مورد بحث قرار می‌گیرد. به عبارت دیگر در این مبحث، قابل اطمینان بودن قطعه مورد بررسی قرار می‌گیرد. با این تعریف از قابلیت اطمینان، کاملاً واضح است که پیشامدهای مورد بحث در برآورد قابلیت اطمینان طول عمر یا زمان شکست و پیشامدهای شکست و پیروزی است. تحلیل طول عمر یا زمان شکست شامل بررسی مدل‌های مربوط به متغیرهای تصادفی با فضای نمونه‌ای مثبت می‌شود. تحلیل طول عمر به دلیل ارتباط با پدیده‌های تصادفی نیاز به انتخاب یک دیدگاه استنباطی (درست‌نمایی، بیزی) برای تعیین دقیق زمان شکست دارد. تاکنون دیدگاه بیزی به شکل بسیار وسیعی در قابلیت اطمینان مورد استفاده قرار گرفته است. علت اصلی که محققان برای استفاده از دیدگاه بیز ارائه نموده اند، آن است که اغلب در مسائل قابلیت اطمینان اطلاعات اضافی و مرتبطی وجود دارند که می‌توان از آن‌ها در استنباط آماری استفاده نمود. دیدگاه درست‌نمایی بر پایه این اصل شکل گرفته است که تمام اطلاعات مجموعه داده‌ها در تابع چگالی داده‌های مشاهده شده قرار دارد، اما دیدگاه بیز امکان اضافه شدن اطلاعاتی خارج از مجموعه داده‌ها را می‌دهد و در مسائل قابلیت اطمینان اغلب تجربیات قبلی (بنا به وسعت داده‌های مربوط به کیفیت) و یا علوم فیزیک و شیمی اطلاعات اضافی را در اختیار پژوهشگر قرار می‌دهد. به طور کلی، این قابلیت دیدگاه بیز و کلیت مسائل قابل اعتماد منجر به استفاده روز افزون از استنباط بیزی در مسائل قابلیت اطمینان شده است.

فصل ۲

برآورد قابلیت اطمینان سیستم‌های سری در مدل‌های فشار-مقاومت

۱.۲ مقدمه

در این فصل، قابلیت اطمینان سیستم‌های سری را در صورتی که مولفه‌های فشارمقاومت دارای طول عمری، تحت توزیع‌های مختلف آماری (نمایی دو پارامتری، گاما، پاراتو، وایبل) باشند، به دو روش ماکزیمم درستنمایی و بیزی برآورد می‌کنیم. نهایتاً پس از حصول برآوردگر قابلیت اطمینان (به روش درستنمایی ماکزیمم)، توزیع برآوردگرها را به دست می‌آوریم. خواهیم دید که تحت بعضی از توزیع‌ها، برآوردگر قابلیت اطمینان دارای توزیع مشخصی است، و در مواردی به دلیل پیچیدگی برآوردگر و نداشتن توزیع معین نشان می‌دهیم برآوردگر دارای توزیع مجانبی است. توزیع مجانبی به طور کامل در بخش پیوست پایان نامه به طور کامل بیان شده است.

۲.۲ مدل فشار-مقاومت نمایی دو پارامتری

در این بخش، برآورد قابلیت اطمینان یک سیستم سری k مولفه‌ای، $X_1, X_2, \dots, X_k$ که با فشار مشترک تصادفی X_{k+1} روبرو است ارائه می‌شود. فرض کنید X_i ($i = 1, 2, \dots, k+1$)‌ها دارای توزیع مستقل نمایی دو پارامتری با پارامترهای μ و λ هستند. بدین ترتیب تابع چگالی احتمال متغیرهای تصادفی مدل به صورت زیر بیان می‌شود

$$f_i(x_i) = \frac{1}{\lambda_i} e^{-\frac{x_i - \mu_i}{\lambda_i}}, \quad x_i > \mu_i > 0, \lambda_i > 0$$

۱.۲.۲ قابلیت اطمینان سیستم‌های سری

اولین گام در محاسبه قابلیت اطمینان، حصول تابع توزیع بقا یا تابع طول عمر مولفه‌ها است. با توجه به فرضیات قسمت قبل، و بنا به تعریف قابلیت اطمینان سیستم‌های سری در فصل اول، فرض می‌کنیم

۲. $Z = \min(X_1, \dots, X_k)$ ، آنگاه تابع بقای سیستم به صورت زیر است

$$\begin{aligned}\bar{H}(z) &= \prod_{j=1}^r P[X_{i_j} > z] \quad \mu_{i_r} < z < \mu_{i_{r+1}}, \quad r = 1, 2, 3, \dots, k \\ &= e^{-\sum_{j=1}^r \frac{(z - \mu_{i_j})}{\lambda_{i_j}}}\end{aligned}$$

به طوری که $\mu_{i_{k+1}} = \infty$ و $i_1 \neq \dots \neq i_k = 1, 2, \dots, k$ برای درک بهتر موضوع اکنون تابع بقا را در صورتی که سیستم تنها شامل دو مولفه باشد، محاسبه می‌کنیم. بنابراین فرضیات فوق $Z > \min(\mu_1, \mu_2)$ ، لذا داریم:

$$\bar{H}(z) = \begin{cases} e^{-\frac{(z - \mu_1)}{\lambda_1}} & \mu_1 < z < \mu_2 \\ e^{-\frac{(z - \mu_1)}{\lambda_1}} - \frac{z - \mu_2}{\lambda_2} & \mu_2 < z < \infty \end{cases}$$

تابع چگالی احتمال تابع بقا برابر است با

$$\bar{h}(z) = \begin{cases} \frac{1}{\lambda_1} e^{-\frac{(z - \mu_1)}{\lambda_1}} & \mu_1 < z < \mu_2 \\ \left(\frac{1}{\lambda_1} + \frac{1}{\lambda_2}\right) e^{-\frac{(z - \mu_1)}{\lambda_1}} - \frac{z - \mu_2}{\lambda_2} & \mu_2 < z < \infty \end{cases}$$

تابع چگالی فوق تحت شرط $\mu_1 < \mu_2$ به دست آمده، اگر فرض کنیم که $\mu_2 < \mu_1$ باشد، آنگاه تابع چگالی تابع بقا با تابع چگالی فوق متفاوت خواهد بود. بنابراین نتیجه می‌گیریم به ازای k مولفه سیستم، تعداد $k!$ ، تابع چگالی احتمال برای تابع بقا وجود دارد. در نتیجه قابلیت اطمینان سیستم بنابر ارتباط بین پارامترها متفاوت خواهد شد. با توجه به تابع بقای به دست آمده از رابطه ۱.۲ قابلیت اطمینان سیستم، به صورت زیر حاصل می‌شود

$$\begin{aligned}R &= P[X_{k+1} < Z] = \int_{\mu_{k+1}}^{\infty} \bar{H}(x) dF_{k+1}(x) \quad \min\{\mu_1, \dots, \mu_k\} \\ &= \int_{\mu_{k+1}}^{\infty} e^{-\sum_{j=1}^r \frac{(x - \mu_{i_j})}{\lambda_{i_j}}} \frac{1}{\lambda_{k+1}} e^{-\left(\frac{x - \mu_{k+1}}{\lambda_{k+1}}\right)} dx \\ &= \int_{\mu_{k+1}}^{\infty} e^{-\sum_{j=1}^r \frac{(x - \mu_{k+1} + \mu_{k+1} - \mu_{i_j})}{\lambda_{i_j}}} \frac{1}{\lambda_{k+1}} e^{-\left(\frac{x - \mu_{k+1}}{\lambda_{k+1}}\right)} dx \\ &= e^{-\sum_{j=1}^r \frac{(\mu_{k+1} - \mu_{i_j})}{\lambda_{i_j}}} \frac{1}{\lambda_{k+1}} \int_{\mu_{k+1}}^{\infty} e^{-\sum_{j=1}^r \frac{(x - \mu_{k+1})}{\lambda_{i_j}}} e^{-\left(\frac{x - \mu_{k+1}}{\lambda_{k+1}}\right)} dx \\ &= e^{-\sum_{j=1}^r \frac{(\mu_{k+1} - \mu_{i_j})}{\lambda_{i_j}}} \frac{1}{\lambda_{k+1}} \int_{\mu_{k+1}}^{\infty} e^{-(x - \mu_{k+1}) \sum_{j=1}^r \frac{1}{\lambda_{i_j}}} e^{-\frac{(x - \mu_{k+1})}{\lambda_{k+1}}} dx\end{aligned}$$

$$= \frac{1}{\lambda_{k+1}} e^{-\sum_{j=1}^r \frac{(\mu_{k+1} - \mu_{i_j})}{\lambda_{i_j}}} \left(\frac{1}{\sum_{j=1}^r \frac{1}{\lambda_{i_j}} + \frac{1}{\lambda_{k+1}}} \right), \quad \mu_{i_r} \leq \mu_{k+1} \leq \mu_{i_{r+1}}. \quad (1.2)$$

۲.۲.۲ برآورد ماکزیمم درستنمایی

فرض کنید سیستم دارای k مولفه است. به طوریکه $X_1, \dots, X_{k_j}$ مولفه‌های سیستم و $X_{(k+1)_j}$ فشار وارده در j امین نمونه است. همچنین فرض کنید $(\lambda_1, \dots, \lambda_{k+1}, \mu_1, \dots, \mu_{k+1})$ بردار پارامترها $X = (x_1, \dots, x_n, \dots, x_{(k+1)_1}, \dots, x_{(k+1)_n})$ به طوریکه دارای $n(k+1)$ بعد می‌باشد. حالت اول: فرض می‌کنیم $\mu_{k+1} = \mu, \neq, \dots, \mu$ ، بنابراین تابع درستنمایی را می‌توان به صورت زیر نوشت

$$\begin{aligned} L(\lambda_1, \dots, \lambda_{k+1}, \mu_1, \dots, \mu_{k+1}; X) &= \prod_{i=1}^{k+1} \prod_{j=1}^n \frac{1}{\lambda_i} e^{-\frac{(x_{i_j} - \mu_i)}{\lambda_i}} I(x_{i_j} > \mu_i) \\ &= e^{-\sum_{i=1}^{k+1} \sum_{j=1}^n \frac{x_{i_j} - \mu_i}{\lambda_i}} \left(\prod_{i=1}^{k+1} \frac{1}{(\lambda_i)^n} \right) \prod_{i=1}^{k+1} I(x_{i_j} > \mu_i) \end{aligned}$$

با لگاریتم گرفتن از طرفین تساوی داریم

$$l(\lambda_1, \dots, \lambda_{k+1}, \mu_1, \dots, \mu_{k+1}) = \sum_{i=1}^{k+1} \sum_{j=1}^n \frac{X_{i_j} - \mu_i}{\lambda_i} - \sum_{i=1}^{k+1} n \ln(\lambda_i) I(x_{i_j} > \mu_i)$$

در روابط فوق منظور از x_{i_j} مینیمم مولفه‌ها است.

در نهایت با مشتق‌گیری نسبت به پارامترها، معادلات درستنمایی به صورت زیر است

$$\begin{aligned} \frac{\partial}{\partial \lambda_i} l(\lambda_1, \dots, \lambda_{k+1}, \mu_1, \dots, \mu_{k+1}) &= -\sum_{j=1}^n \frac{x_{i_j} - \mu_i}{\lambda_i^2} + n \frac{1}{\lambda_i} = 0, \\ \frac{\partial}{\partial \mu_i} l(\lambda_1, \dots, \lambda_{k+1}, \mu_1, \dots, \mu_{k+1}) &= -n I(x_{i_j} > \mu_i), \end{aligned}$$

بنابراین با حل معادلات درستنمایی، برآوردهای ماکزیمم درستنمایی پارامترها عبارتند از:

$$\begin{aligned} \hat{\lambda}_i &= \frac{\sum_{j=1}^n (X_{i_j} - X_{i(1)})}{n}, \\ \hat{\mu}_i &= X_{i(1)}; i = 1, 2, 3, \dots, k+1. \end{aligned}$$

حال با استفاده از خاصیت پایایی برآورد ماکزیمم درستنمایی، که در فصل اول بیان شد. می‌توان با قرار دادن برآوردهای فوق در ۱.۲، برآورد ماکزیمم درستنمایی قابلیت اطمینان را به دست می‌آوریم.

۳.۲.۲ برآورد بیزی

بدلیل نیاز به برآوردگر بیزی برای شبیه‌سازی در فصل چهارم، از بیان آن در این قسمت خود داری کرده، و برآوردگر بیزی را در فصل چهارم کامل بیان می‌کنیم.

۴.۲.۲ توزیع برآوردگر

قبلا از اینکه وارد مبحث توزیع برآوردگرها شویم، ابتدا قضیه زیر که مهمترین و کاربردی‌ترین قضیه در این فصل می‌باشد را بیان می‌کنیم.

قضیه ۱.۲.۲. فرض می‌کنیم $\{T_n\}$ دنباله‌ای از برآوردگرها با پارامتر θ باشد، که با توجه به تعریف توزیع مجانبی دارای توزیع زیر است

$$\sqrt{n}(T_n - \theta) \longrightarrow X \sim N[0, \sigma^2(\theta)]^2$$

آنگاه برای هر تابع مشتق پذیر در θ مانند g داریم:

$$\sqrt{n}(g(T_n) - g(\theta)) \longrightarrow X \sim N(0, (g'(\theta)\sigma(\theta))^2) \quad n = 1, 2, 3, \dots; \quad g'(\theta) \neq 0$$

برهان. کتاب راثو (۱۹۷۳) صفحه ۴۲۶ را ببینید. $\square$

پس از به دست آوردن برآوردگر درست‌نمایی قابلیت اطمینان، اکنون می‌خواهیم توزیع برآوردگر را مشخص کنیم.

جالت اول: فرض می‌کنیم $\mu_1 \neq \mu_2 \neq \dots \neq \mu_{k+1}$ از طرفی بنابر خاصیت پایایی داریم:

$$\hat{R} = R(\hat{\theta})$$

به‌طوریکه

$$\hat{\theta} = (\hat{\lambda}_1, \dots, \hat{\lambda}_{k+1}, \hat{\mu}_1, \dots, \hat{\mu}_{k+1})$$

قبلا نشان دادیم:

$$\hat{\lambda}_i = \frac{\sum_{j=1}^n (X_{ij} - X_{i(1)})}{n} \quad \hat{\mu}_i = X_{i(1)}; i = 1, 2, 3, \dots, k+1.$$

برای به‌دست آوردن توزیع برآوردگر قابلیت اطمینان گام اول مشخص نمودن توزیع برآورگرهای پارامترها است. لذا ابتدا باید تابع توزیع $\hat{\mu}_i, \hat{\lambda}_i$ را مشخص کنیم. بنابر آماره‌های ترتیبی $X_{i(1)}$ اولین آماره است، اگر $U_i = X_{i(1)}$ در نظر بگیریم، با توجه به تابع چگالی احتمال آماره‌های ترتیبی داریم

$$f_{U_i}(u_i) = \frac{n}{\lambda_i} e^{-\frac{n(u_i - \mu_i)}{\lambda_i}}, \quad i = 1, 2, \dots, k+1.$$

با توجه به تابع چگالی احتمال حاصله نتیجه می‌گیریم که $X_{i(1)}$ دارای توزیع نمایی دو پارامتری با پارامترهای μ و $\frac{\lambda}{n}$ است، از طرفی بنابر روابط بین توزیع‌ها، می‌دانیم $\frac{\lambda}{n} \chi_{2n-2}^2 \sim \hat{\lambda}_i$ ، بنابراین

واریانس آن‌ها برابر است با

$$V(\hat{\mu}_i) = \frac{\lambda_i^2}{n^2}, \quad V(\hat{\lambda}_i) = \frac{(n-1)\lambda_i^2}{n^2}, \quad i = 1, \dots, k+1$$

با توجه به تعریف سازگاری و توزیع مجانبی برآوردگرها، که در قسمت پیوست همین پایان‌نامه بیان شده، $\hat{\mu}_i$ و $\hat{\lambda}_i$ ها به ترتیب دارای توزیع مجانبی نرمال با میانگین‌های μ_i, λ_i و واریانس‌های به دست آمده در بالا می‌باشند. لذا بنابر قضیه ۱.۲.۲ قابلیت اطمینان سیستم‌ها دارای توابع توزیع مجانبی نرمال سازگار^۱ به صورت زیر می‌باشد

$$\begin{aligned} \hat{R} &\sim AN(R, B' \Lambda B) \\ B' &= \left(\frac{\partial R}{\partial \lambda_1}, \dots, \frac{\partial R}{\partial \lambda_{k+1}}, \dots, \frac{\partial R}{\partial \mu_1}, \dots, \frac{\partial R}{\partial \mu_{k+1}} \right) \\ \Lambda &= \frac{1}{n^2} \text{diag}((n-1)\lambda_1^2, \dots, (n-1)\lambda_{k+1}^2, \lambda_1^2, \dots, \lambda_{k+1}^2) \end{aligned}$$

حالت دوم. فرض کنید $\mu_1 = \dots = \mu_{k+1} = \mu$ ، در این صورت:

$$\hat{R} = R(\hat{\lambda}_1, \dots, \hat{\lambda}_{k+1})$$

بنابر قسمت قبل

$$\hat{\lambda}_i = \frac{\sum_{j=1}^n (X_{ij} - \hat{\mu})}{n}, \quad i = 1, 2, \dots, k+1$$

از آنجایی که $\hat{\mu} = \min X_{ij}$ داریم:

$$V(\hat{\lambda}_i) = \frac{\lambda_i^2}{n} \quad i = 1, 2, \dots, k+1$$

اکنون با توجه به این که $\hat{\mu}$ برآوردگری سازگار است و $\hat{\lambda}_i$ ها دارای توزیع مجانبی با میانگین λ_i و واریانس مشخص هستند، بنابر قضیه ۱.۲.۲ توزیع برآوردگر قابلیت اطمینان به صورت زیر است.

$$\begin{aligned} \hat{R} &\sim AN(R, B' \Lambda^* B) \\ B' &= \left(\frac{\partial R}{\partial \lambda_1}, \dots, \frac{\partial R}{\partial \lambda_{k+1}} \right) \\ \Lambda^* &= \frac{1}{n} \text{diag}(\lambda_1^2, \dots, \lambda_{k+1}^2) \end{aligned}$$

۳.۲ مدل فشار- مقاومت توزیع گاما

همانند بخش قبل فرض کنید در یک سیستم سری، k مولفه‌ای، $X_1, X_2, \dots, X_k$ ، مولفه‌ها همگی تحت فشار یکسان تصادفی X_{k+1} هستند. با توجه به تابع توزیع طول عمر مولفه‌ها، قابلیت اطمینان را برآورد می‌کنیم.

^۱Consistent Asymptotically Normal

۱.۳.۲ قابلیت اطمینان سیستم‌های سری

لم ۱.۳.۲. فرض کنید $X_i; (i = 1, 2, \dots, k+1)$ ها دارای توزیع مستقل گاما با پارامترهای α و μ هستند. بدین ترتیب تابع چگالی احتمال متغیرهای تصادفی مدل به صورت زیر بیان می‌شود:

$$f_i(x_i) = \mu_i^{\alpha_i} x_i^{\alpha_i-1} e^{(-\mu_i x_i)} / \Gamma(\alpha_i) \quad , x_i > 0, \quad \mu_i > 0, i = 1, 2, \dots, k+1.$$

که در آن α_i ثابت‌های صحیح مثبت و معلومی هستند، و اگر $\alpha_i = 1$ باشد، آنگاه توزیع طول عمر مولفه‌های فشار و مقاومت، به توزیع نمایی تبدیل خواهد شد. بنابراین، قابلیت اطمینان سیستم به صورت زیر به دست می‌آید:

$$\begin{aligned} R &= P[X_{k+1} < \min(X_1, X_2, \dots, X_k)] \\ &= \sum_{i_1=0}^{\alpha_1-1} \dots \sum_{i_k=0}^{\alpha_k-1} \int_0^\infty \frac{(\mu_1^{i_1} \dots \mu_k^{i_k}) \exp(\sum_{j=1}^{k+1} -\mu_j x_{k+1}) \mu_i^{\alpha_i}}{\Gamma(i_1+1) \dots \Gamma(i_k+1) \Gamma(\alpha_i)} dx_{k+1}. \end{aligned}$$

برهان. از آنجایی که فرض شد α_i ها مقادیر ثابت صحیح مثبتی هستند، بنابراین توزیع گاما به توزیع ارلانگ تبدیل می‌شود و می‌توان، تابع بقای آن را به فرمی که در ادامه می‌آید نوشت

$$P[X_j > z] = \sum_{i_j=0}^{\alpha_j-1} \mu_j^{i_j} e^{(-\mu_j z)} z^{i_j} / \Gamma(i_j+1), \quad j = 1, 2, \dots, k+1 \quad (2.2)$$

بنابراین اگر فرض کنیم $Z = \min(X_1, \dots, X_k)$ ، تابع بقای Z به صورت

$$\begin{aligned} \bar{H}(z) &= P[Z > z] \\ &= P[X_1 > z, X_2 > z, \dots, X_k > z] \\ &= \left(\sum_{i_1=0}^{\alpha_1-1} \frac{\mu_1^{i_1} \exp(-\mu_1 z) z^{i_1}}{\Gamma(i_1+1)} \right) \dots \left(\sum_{i_k=0}^{\alpha_k-1} \frac{\mu_k^{i_k} \exp(-\mu_k z) z^{i_k}}{\Gamma(i_k+1)} \right), \end{aligned}$$

خواهد بود. حال قابلیت اطمینان سیستم سری k مولفه‌ای عبارتست از

$$\begin{aligned} R &= P[X_{k+1} < \min(X_1, X_2, \dots, X_k)] \\ &= \int_0^\infty \bar{H}(x_{k+1}) dF_{k+1}(x_{k+1}) \\ &= \int_0^\infty \left(\sum_{i_1=0}^{\alpha_1-1} \frac{\mu_1^{i_1} \exp(-\mu_1 z) z^{i_1}}{\Gamma(i_1+1)} \right) \\ &\quad \dots \left(\sum_{i_k=0}^{\alpha_k-1} \frac{\mu_k^{i_k} \exp(-\mu_k z) z^{i_k}}{\Gamma(i_k+1)} \right) \mu_i^{\alpha_i} x_{k+1}^{\alpha_i-1} \exp(-\mu_i x_{k+1}) / \Gamma(\alpha_i) dx_{k+1}. \end{aligned}$$

با توجه به اینکه حدود انتگرال و حدود مجموع به هم وابسته نیستند، می‌توانیم جای مجموع و انتگرال

را عوض کنیم

$$R = \sum_{i_1=0}^{\alpha_1-1} \cdots \sum_{i_k=0}^{\alpha_k-1} \int_0^\infty \frac{(\mu_1^{i_1} \cdots \mu_k^{i_k}) \exp(\sum_{j=1}^{k+1} -\mu_j x_{k+1}) \mu_i^{\alpha_i}}{\Gamma(i_1+1) \cdots \Gamma(i_k+1) \Gamma(\alpha_i)} dx_{k+1} \quad (۳.۲)$$

$$= \sum_{i_1=0}^{\alpha_1-1} \cdots \sum_{i_k=0}^{\alpha_k-1} \left[\frac{\mu_1^{i_1} \cdots \mu_k^{i_k} \mu_{k+1}^{\alpha_{k+1}}}{(\mu_1 + \mu_2 + \cdots + \mu_k + \mu_{k+1})^{i_1+i_2+\cdots+i_k+\alpha_{k+1}}} \right] \quad (۴.۲)$$

$$\times \left[\frac{\Gamma(i_1+i_2+\cdots+i_k+\alpha_{k+1})}{\Gamma(i_1+1) \cdots \Gamma(i_k+1) \Gamma(\alpha_{k+1})} \right]. \quad (۵.۲)$$

□

۲.۳.۲ برآورد ماکزیمم درستنمایی

همانند قسمت قبل با توجه به فرضیات بخش ۱.۳.۲ ابتدا برآورد درستنمایی پارامترهای توزیع را، به دست می آوریم

$$L = \left[\prod_{i=1}^{k+1} \mu_i^{n\alpha_i} \left(\prod_{j=1}^n x_{ij}^{\alpha_i-1} \right) \right] \exp \left\{ - \sum_{i=1}^{k+1} \mu_i \sum_{j=1}^n x_{ij} \right\} / \left[\prod_{i=1}^{k+1} (\Gamma(\alpha_i))^n \right]$$

با لگاریتم گرفتن از تابع درستنمایی فوق، و مشتق گرفتن از تابع لگاریتمی داریم:

$$n\alpha_i/\mu_i - \sum_{j=1}^n x_{ij} = 0, \quad i = 1, 2, \dots, k+1$$

سپس با حل معادله درستنمایی فوق برآورگر درستنمایی μ_i ها به صورت زیر به دست می آید.

$$\hat{\mu}_i = \frac{n\alpha_i}{\sum_{j=1}^n x_{ij}}, \quad i = 1, 2, \dots, k+1.$$

اکنون بنابر خاصیت پایایی، با قرار دادن برآوردهای ماکزیمم درستنمایی $\hat{\mu}_i$ در رابطه (۳.۲)، برآورد قابلیت اطمینان سیستم به دست می آید.

۳.۳.۲ برآورد بیزی

با در نظر گرفتن فرضیات بخش ۱.۳.۲، پارامتر α_i معلوم و پارامترهای μ_i دارای توزیع پیشین زیر

$$\pi(\mu_i) = \theta \exp(-\mu_i \theta); \quad i = 1, 2, \dots, k+1,$$

می باشند. به طوری که θ معلوم باشد. ابتدا توزیع پسین را به دست می آوریم

$$\begin{aligned} \pi(\mu_i | \mathbf{X}) &\propto \prod_{j=1}^n \frac{\mu_i^{\alpha_i}}{\Gamma(\alpha_i)} x_j^{\alpha_i-1} \exp(-\mu_i x_j) \theta \exp(-\mu_i \theta) \\ &\propto \mu_i^{n\alpha_i} \exp(-\mu_i (\sum_{j=1}^n x_j + \theta)). \end{aligned}$$

بنابراین

$$\mu_i | \mathbf{X} \sim \Gamma(n\alpha_i + 1, \sum_{j=1}^n x_j + \theta), \quad i = 1, 2, \dots, k+1.$$

پس از مشخص کردن توزیع پسین پارامتر اکنون برآورد بیز قابلیت اطمینان را، به صورت زیر به دست می‌آوریم.

$$E(R|\mathbf{X}) = \int_0^\infty \cdots \int_0^\infty \sum_{i_1=0}^{\alpha_1-1} \cdots \sum_{i_k=0}^{\alpha_k-1} \left[\frac{(n\alpha_1 / \sum_{j=1}^n x_{1j})^{i_1} \cdots (n\alpha_k / \sum_{j=1}^n x_{kj})^{i_k} (n\alpha_{k+1} / \sum_{j=1}^n x_{k+1j})^{\alpha_{k+1}}}{((n\alpha_1 / \sum_{j=1}^n x_{1j}) + (n\alpha_2 / \sum_{j=1}^n x_{2j}) + \cdots + (n\alpha_{k+1} / \sum_{j=1}^n x_{k+1j}))^{i_1+i_2+\cdots+i_k+\alpha_{k+1}}} \right] \\ \times \left[\frac{\Gamma(i_1 + i_2 + \cdots + i_k + \alpha_{k+1})}{\Gamma(i_1 + 1) \cdots \Gamma(i_k + 1) \Gamma(\alpha_{k+1})} \right] \pi(\boldsymbol{\mu}|\mathbf{X}) d\mu_1 \cdots d\mu_{k+1}$$

به دلیل وجود انتگرال‌های چندگانه و پیچیده برای برآورد قابلیت اطمینان از روش‌های نمونه‌گیری MCMC استفاده می‌کنیم. از آنجایی که تابع داخل انتگرال مشخص است و می‌توان به راحتی از آن تولید نمونه کرد، از روش نمونه‌گیری گیبز برای محاسبه انتگرال فوق استفاده می‌کنیم.

۴.۳.۲ توزیع برآوردگر

برای به دست آوردن توزیع برآوردگر قابلیت اطمینان ابتدا براساس تعریف توزیع مجانبی ثابت می‌کنیم که $\hat{\mu}_i$ ها دارای توزیع مجانبی نرمال سازگار (CAN) هستند. ثابت کردیم:

$$\hat{\mu}_i = \frac{n\alpha_i}{\sum_{j=1}^n x_{ij}}, \quad i = 1, 2, \dots, k+1.$$

از آنجایی که برآوردگرهای پارامترها در این قسمت مانند حالت نمایی دو پارامتری دارای تابع توزیع مشخصی نمی‌باشند، ابتدا با استفاده از ماتریس اطلاع فیشر که در قسمت پیوست پایان نامه بیان کردیم، واریانس برآوردگرها را مشخص سپس توزیع مجانبی آن‌ها را به دست می‌آوریم. لذا با توجه به تعریف ماتریس اطلاع فیشر داریم:

$$nI = n((I_{ij})); \quad i = 1, 2, \dots, k+1 \\ I_{ij} = \frac{\alpha_i}{\mu_i^2}$$

لذا

$$Var(\hat{\mu}_i) = \frac{\mu_i^2}{n\alpha_i}, \quad i = 1, 2, \dots, k+1$$

$$Cov(\hat{\mu}_i, \hat{\mu}_j) = 0$$

اکنون با مشخص شدن واریانس $\hat{\mu}_i$ ها، از آنجایی که این برآوردگر، برآوردگری سازگار است و براساس تعریف توزیع مجانبی نرمال سازگار می‌توان گفت:

$$\sqrt{n}(\hat{\mu} - \mu) \sim AN(0, \Lambda_1) \quad \Lambda_1 = \text{diag} \left(\frac{\mu_1^2}{\alpha_1}, \dots, \frac{\mu_{k+1}^2}{\alpha_{k+1}} \right)$$

با توجه به قضیه ۱.۲.۲ و اینکه $\hat{R}$ تابعی از $\hat{\mu}$ ها می‌باشد، داریم

$$\hat{R} = g(\hat{\mu}) \sim AN(g(\mu), G'_1 \Lambda_1 G_1 / n) \quad G'_1 = \left(\frac{\partial R}{\partial \mu_1}, \dots, \frac{\partial R}{\partial \mu_{k+1}} \right)$$

۴.۲ مدل فشار-مقاومت توزیع وایبل

۱.۴.۲ قابلیت اطمینان سیستم‌های سری

فرض کنید مولفه‌های فشار و مقاومت در یک سیستم سری، از توزیع وایبل پیروی کنند. در این صورت، تابع چگالی آن‌ها به فرم

$$f(x_i) = c x_i^{c-1} \frac{\exp(-(\frac{x_i}{\theta_i})^c)}{\theta_i^c},$$

خواهد بود که در آن پارامتر شکل برای تمام توزیع‌ها یکسان و برابر با c است. بنابراین با تعریف $Z = \min(X_1, \dots, X_k)$ ، تابع بقای Z به صورت

$$\begin{aligned} \bar{H}(z) &= P[Z > z] \\ &= \exp(-z^c (\frac{1}{\theta_1^c} + \dots + \frac{1}{\theta_k^c})), \end{aligned}$$

بیان می‌شود. بنابراین قابلیت اطمینان سیستم برابر است با

$$\begin{aligned} R &= P[X_{k+1} < \min(X_1, X_2, X_3, \dots, X_k)] \\ &= \int_0^\infty \bar{H}(x_{k+1}) dF_{k+1}(x_{k+1}) \\ &= \int_0^\infty \exp(-x_{k+1}^c (\frac{1}{\theta_1^c} + \dots + \frac{1}{\theta_k^c})) c x_{k+1}^{c-1} \frac{\exp(-(\frac{x_{k+1}}{\theta_{k+1}})^c)}{\theta_{k+1}^c} dx_{k+1} \\ &= \frac{1}{1 + (\frac{\theta_{k+1}}{\theta_1})^c + \dots + (\frac{\theta_{k+1}}{\theta_k})^c}. \end{aligned} \quad (6.2)$$

۲.۴.۲ برآورد ماکزیمم درستنمایی

با در نظر گرفتن فرضیات موجود در بخش ۲.۳.۲، تابع درستنمایی به صورت زیر است

$$L(\theta, c) = c^{n(k+1)} \exp(-\sum_{i=1}^{k+1} \sum_{j=1}^n (\frac{x_{ij}}{\theta_i})^c) \prod_{i=1}^{k+1} \prod_{j=1}^n x_{ij}^c \theta_i^{-c}.$$

با لگاریتم گرفتن از طرفین تساوی داریم

$$l(\theta, c) = n(k+1) \ln c - \sum_{i=1}^{k+1} \sum_{j=1}^n \left(\frac{x_{ij}}{\theta_i}\right)^c + \sum_{i=1}^{k+1} \sum_{j=1}^n (c \ln x_{ij} - c \ln \theta_i)$$

ضمن مشتق‌گیری از تابع درستمایی و حل معادلات درستمایی، برآورد پارامترهای توزیع با استفاده از روش ماکزیم درستمایی عبارتند از

$$\hat{\theta}_i = \left(\frac{\sum_{j=1}^n x_{ij}^{\hat{c}}}{n} \right)^{\frac{1}{\hat{c}}},$$

$$\hat{c} = \left[\frac{1}{k+1} \left(\sum_{i=1}^{k+1} \sum_{j=1}^n x_{ij}^{\hat{c}} \ln x_{ij} \right) \left[\sum_{j=1}^n x_{ij}^{\hat{c}} \right]^{-1} - \frac{1}{n(k+1)} \sum_{i=1}^{k+1} \sum_{j=1}^n \ln x_{ij} \right]^{-1}.$$

با قرار دادن برآوردهای ماکزیم درستمایی پارامترها در رابطه (۶.۲)، برآورد قابلیت اطمینان بدست می‌آید.

$$\hat{R} = \frac{1}{[1 + (\hat{\theta}_{k+1}/\hat{\theta}_1)^{\hat{c}}] + \dots + (\hat{\theta}_{k+1}/\hat{\theta}_k)^{\hat{c}}} \quad (7.2)$$

۳.۴.۲ توزیع برآوردگر

همانند قسمت قبل برای مشخص نمودن توزیع $\hat{R}$ ، ابتدا توزیع پارامترها را پیدا می‌کنیم. از آنجایی که برآوردگرهای پارامترها دارای توابع توزیع مشخص نیستند از روش ماتریس اطلاع فیشر استفاده می‌کنیم. بنابراین داریم:

$$\begin{aligned} I_{ij} &= \frac{c^2}{\theta_i^2} \quad i = 1, 2, \dots, k+1 \\ I_{k+2, k+2} &= (k+1) [\psi'(1) + \{\psi(2)\}^2] / c^2 \\ I_{ik+1} &= -\psi(2)/\theta_i, \quad i = 1, 2, \dots, k+1 \\ I_{ij} &= 0, \quad i \neq j = 1, 2, \dots, k+1 \end{aligned}$$

طبق تعریف $\psi(\cdot)$ و $\psi'(\cdot)$ به ترتیب توابع دای گاما و تری گاما هستند. که در پیوست تعریف کردیم.

گزاره ۱.۴.۲. براساس قضیه حد مرکزی $\sqrt{n}(\hat{\lambda} - \underline{\lambda})$ دارای توزیع مجانبی نرمال چند متغیره با میانگین صفر و ماتریس واریانس-کواریانس $I^{-1} = ((I^{ij}))$ می‌باشد. I^{ij} ها عناصر ماتریس معکوس ماتریس اطلاع فیشر می‌باشند.

اکنون می‌توان گفت که براساس قضیه (۱.۲.۲) و گزاره ۱.۴.۲ برآوردگر ماکزیم درستمایی قابلیت اطمینان دارای توزیع زیر است.

$$\hat{R} \sim AN(R, G'_2 \Lambda_2 G_2 / n) \quad G'_2 = \left(\frac{\partial R}{\partial \theta_1}, \dots, \frac{\partial R}{\partial \theta_{k+1}}, \frac{\partial R}{\partial c} \right), \quad \Lambda_2 = \frac{1}{n} I^{-1}$$

۴.۴.۲ برآورد بیزی

به دلیل نیاز به برآورد بیزی قابلیت اطمینان سیستم، در صورتی که مولفه‌ها دارای طول عمری با توزیع وایبل باشند، در فصل سوم و شبیه‌سازی برآوردگر از بیان آن در این قسمت خودداری کرده، و به طور کامل در فصل سوم بیان می‌کنیم.

۵.۴.۲ مدل نمایی

بنابر خواص توزیع وایبل، اگر $c = 1$ فرض شود، توزیع وایبل تبدیل به توزیع نمایی می‌شود. لذا توزیع برآوردگر درست‌نمایی قابلیت اطمینان را برای زمانی که سیستم دارای تنها دو مولفه X_1, X_2 است تحت فرض فوق محاسبه می‌کنیم.

$$R = P[X_2 < X_1] = \frac{1}{[1 + \theta_2/\theta_1]}.$$

فرض می‌کنیم $\hat{\theta}_1 = Z_1, \hat{\theta}_2 = Z_2$ ، در این صورت

$$\hat{R} = \frac{1}{[1 + Z_2/Z_1]}. \quad (۸.۲)$$

با توجه به برآورد ماکزیمم درست‌نمایی پارامترها، داریم $Z_1 = \sum_{j=1}^n X_{1j}$ و $Z_2 = \sum_{j=1}^n X_{2j}$ ، از طرفی چون (X_{1j}, X_{2j}) متغیرهای تصادفی از توزیع نمایی هستند، بنابر خواص توزیع نمایی داریم:

$$Z_1 \sim \Gamma(n, \theta_1) \quad Z_2 \sim \Gamma(n, \theta_2)$$

اکنون می‌توان با استفاده از روش ماتریس ژاکوبین، تابع چگالی احتمال $\hat{R}$ را بدست آورد.

$$R = \frac{1}{[1 + Z_2/Z_1]} = \frac{Z_1}{Z_1 + Z_2}$$

$$\begin{cases} t = Z_1 \\ w = \frac{Z_1}{Z_1 + Z_2} \end{cases} \Rightarrow$$

$$\begin{cases} Z_1 = t \\ Z_2 = \frac{t - tw}{w} \end{cases}$$

$$\begin{aligned}
 f_w &= \int_0^\infty \frac{t}{w^2} \frac{\theta_1^n}{\Gamma(n)} t^{n-1} e^{-\theta_1 t} \frac{\theta_2^n}{\Gamma(n)} \left(\frac{t-tw}{w}\right)^{n-1} e^{-\frac{\theta_2 t}{w} + \theta_2 t} e^{-\theta_2 \left(\frac{t-tw}{w}\right)} dt \\
 &= \frac{1}{w^2} \frac{\theta_1^n \theta_2^n}{\Gamma(n)\Gamma(n)} \int_0^\infty t^n \left(\frac{t-tw}{w}\right)^{n-1} e^{\{-\theta_1 t + \theta_2 t - \frac{\theta_2 t}{w}\}} dt \\
 &= \frac{(1-w)^{n-1}}{w^2 w^{n-1}} \frac{\theta_1^n \theta_2^n}{\Gamma(n)\Gamma(n)} \int_0^\infty t^{2n-1} e^{\{\theta_1 - \theta_2 + \frac{\theta_2}{w}\}t} dz \\
 &= \frac{(1-w)^{n-1}}{w^{n+1}} \frac{\theta_1^n \theta_2^n}{\Gamma(n)\Gamma(n)} \frac{\Gamma(2n)}{(\theta_1 - \theta_2 + \frac{\theta_2}{w})^{2n}}
 \end{aligned}$$

نهایتاً:

$$f(\hat{R}) = \frac{(1-\hat{R})^{n-1} \theta_1^n \theta_2^n \Gamma(2n)}{\hat{R}^{n+1} (\theta_1 - \theta_2 + \frac{\theta_2}{\hat{R}})^{2n} (\Gamma(n))^2}$$

* در صورتی که $\theta_1 = \theta_2$ باشد، برآوردگر $\hat{R}$ دارای توزیع $Beta(n, n)$ می‌باشد.

۵.۲ مدل فشار-مقاومت توزیع پاراتو

در این بخش، قابلیت اطمینان سیستم‌های سری در دو حالت ارائه می‌شود.

۱.۵.۲ قابلیت اطمینان سیستم‌های سری

حالت اول: فرض کنید مولفه‌های فشار و مقاومت از توزیع پارتو با تابع توزیع زیر پیروی کنند

$$F_i(x_i) = 1 - \left(\frac{\mu}{x_i}\right)^{a_i}; \quad x_i \leq \mu,$$

که در آن $\mu > 0$ پارامتر مشترک و $a_i > 0$ پارامتر غیر مشترک می‌باشد. همانند قسمت‌های قبل فرض می‌کنیم $Z = \min(X_1, X_2, \dots, X_k)$ ، آنگاه تابع بقا Z به صورت

$$\bar{H}(z) = \frac{\mu^{\sum_{i=1}^k a_i}}{z}; \quad z \leq 0,$$

می‌باشد به طوری که $\sum_{i=1}^k a_i > \mu$. بنابراین با تعریف قابلیت اطمینان سیستم سری به صورت زیر می‌توان با استفاده از رابطه فوق رابطه صریحی برای آن یافت. بدین ترتیب خواهیم داشت

$$\begin{aligned}
 R &= P[X_{k+1} < \min(X_1, X_2, \dots, X_k)] \\
 &= \int_0^\infty \bar{H}(x_{k+1}) dF_{k+1}(x_{k+1}) \\
 &= \int_\mu^\infty \left(\frac{\mu}{x_{k+1}}\right)^{\sum_{i=1}^k a_i} a_{k+1} \frac{\mu^{a_{k+1}}}{x_i^{a_{k+1}+1}} dx_{k+1} \\
 &= \frac{a_{k+1}}{\sum_{i=1}^k a_i}.
 \end{aligned} \tag{۹.۲}$$

۲.۵.۲ برآورد ماکزیمم درستنمایی

با در نظر گرفتن مفروضات بخش ۱.۵.۲، تابع درستنمایی برای این توزیع به صورت زیر بیان می‌شود

$$L = \prod_{i=1}^{k+1} a_i^n \mu^{n(\sum_{i=1}^{k+1} a_i)} \prod_{i=1}^{k+1} \prod_{j=1}^n \frac{I(\mu < X_{ij}) X_{ij}}{X_{ij}^{a_i+1}}.$$

از آنجایی که تابع درستنمایی نسبت به پارامتر μ است، بر این اساس برآورد درستنمایی μ برابر است با $\min_{i>0, j>0} (X_{ij})$ و برآورد سایر پارامترها با استفاده از معادلات درستنمایی

$$\frac{n}{a_i} + n \log \hat{\mu} - \log x_{ij} = 0, \quad i = 1, 2, \dots, k+1$$

به دست می‌آیند به طوری که

$$\hat{a}_i = \frac{1}{\log\left(\frac{\prod_{j=1}^n x_{ij}}{\hat{\mu}^n}\right)}, \quad \hat{\mu} = \min_{ij} (x_{ij}), \quad j = 1, 2, \dots, n$$

بنابر خاصیت پایایی با قردادن برآورد پارامترها در رابطه (۹.۲) برآورد قابلیت اطمینان به دست می‌آید.

$$\hat{R} = \frac{\hat{a}_{k+1}}{\sum_{i=1}^k \hat{a}_i}.$$

۳.۵.۲ برآورد بیزی

همانند بخش‌های قبل، بنابر مفروضات بخش ۱.۵.۲، اگر μ ثابت در نظر گرفته شود، و توزیع پسین پارامترهای α_i ، به صورت زیر

$$\pi(\alpha_i) = \lambda \exp(-\alpha_i \lambda); \quad i = 1, 2, \dots, k+1,$$

فرض شود که λ مقداری معلوم است، توزیع پسین را به صورت زیر به دست می‌آوریم

$$\begin{aligned} \pi(\alpha_i | \mathbf{X}) &\propto \prod_{j=1}^n \alpha_i \frac{\mu^{\alpha_i}}{x_j^{\alpha_i+1}} \lambda \exp(-\alpha_i \lambda) \\ &\propto \alpha_i^n \left(\frac{\mu^n}{\prod_{j=1}^n x_j} \right)^{\alpha_i} \exp(-\alpha_i \lambda) \\ &\propto \alpha_i^n \exp(-\alpha_i (-\ln \frac{\mu^n}{\prod_{j=1}^n x_j})) \exp(-\alpha_i \lambda) \\ &\propto \alpha_i^n \exp(-\alpha_i (\lambda - \ln \frac{\mu^n}{\prod_{j=1}^n x_j})). \end{aligned} \quad (10.2)$$

بنابراین می‌توان گفت

$$\alpha_i | \mathbf{X} \sim \Gamma(n+1, \lambda - \ln \frac{\mu^n}{\prod_{j=1}^n x_j}); \quad i = 1, 2, \dots, k+1.$$

برای بدست آوردن برآورد بیز قابلیت اطمینان با استفاده از رابطه ۹.۲ داریم

$$\begin{aligned} E(R|\mathbf{X}) &= E\left(\frac{a_{k+1}}{\sum_{i=1}^k a_i}|\mathbf{X}\right) \\ &= \int_0^\infty \cdots \int_0^\infty \frac{a_{k+1}}{\sum_{i=1}^k a_i} \pi(\alpha|\mathbf{X}) d\alpha_1 \cdots d\alpha_{k+1}, \end{aligned} \quad (11.2)$$

برای حل انتگرال فوق و برآورد قابلیت اطمینان، از روش‌های *MCMC* استفاده می‌کنیم.

۴.۵.۲ توزیع برآوردگر

در اینجا نیز برای بدست آوردن توزیع برآوردگر قابلیت اطمینان، ابتدا با استفاده از ماتریس اطلاع فشر، واریانس پارامترها را بدست آورده و پس از محاسبه توزیع پارامترها، توزیع برآوردگر را مشخص می‌کنیم. بنابر کران پایین کرامررئو عناصر ماتریس اطلاع برابر است با

$$I_{ij} = \frac{n}{a_i^2} \quad i = 1, 2, \dots, k+1$$

در نتیجه می‌توان واریانس مجانبی پارامترها را محاسبه کرد.

$$Var(\hat{a}_i) = \frac{a_i^2}{n}, \quad Cov(\hat{a}_i, \hat{a}_j) = 0, \quad i \neq j = 1, 2, \dots, k+1$$

داریم $\hat{\mu} \xrightarrow{p} \mu$ ، از طرفی بنابر قضیه حد مرکزی داریم:

$$\sqrt{n}(\hat{\underline{a}} - \underline{a}) \xrightarrow{d} AN(0, \Sigma)$$

$$\sqrt{n}(\hat{\mu} - \mu) \xrightarrow{p} 0.$$

در نتیجه

$$\sqrt{n}(\hat{\underline{a}} - \underline{a}) + \sqrt{n}(\hat{\mu} - \mu) \xrightarrow{d} AMVN(0, \Sigma) \quad \Sigma = diag(a_1^2, \dots, a_{k+1}^2)$$

با توجه به توزیع مجانبی پارامترها و بنابر قضیه (۱۰.۲.۲)، $\hat{R}$ دارای توزیع مجانبی زیر است.

$$\hat{R} \sim AN(R, B'_\Psi \Lambda_\Psi B_\Psi)$$

که در آن

$$B'_\Psi = (\partial R / \partial a_1, \dots, \partial R / \partial a_{k+1}) \quad \Lambda_\Psi = \frac{1}{n} diag(a_1^2, \dots, a_{k+1}^2).$$

حالت دوم: فرض می‌کنیم توزیع پارامترها دارای مقادیر مختلف $\mu_1, \dots, \mu_{k+1}$ باشد، در این صورت تابع چگالی احتمال برابر است با

$$f_i(x_i) = \frac{a_i \mu_i^{a_i}}{x_i^{a_i+1}} \quad \mu_i, x_i > 0, \quad x_i > \mu_i \quad i = 1, 2, \dots, k+1$$

همانند قسمت اول تابع بقا را به ازای $Z = \text{Min}(X_1, X_2, \dots, X_k)$ بدست می‌آوریم

$$\begin{aligned}\bar{H}(z) &= \prod_{j=1}^r P[X_{i_j} > z], \quad \mu_{i_r} < z < \mu_{i_{r+1}}; r = 1, 2, \dots, k \\ &= \prod_{j=1}^r \mu_{i_j}^{a_{i_j}} / z^{\sum_{j=1}^r a_{i_j}}, \quad \mu_{i_r} < z < \mu_{i_{r+1}}; r = 1, 2, \dots, k\end{aligned}$$

بطوری که $K = 2$ می‌کنیم $i_1 \neq \dots \neq i_k = 1, 2, \dots, k; \mu_{i_{k+1}} = \infty$ باشد، یعنی این که سیستم تنها دارای ۲ مولفه فشار-مقاومت باشد. در این صورت برای بررسی قابلیت اطمینان سیستم تابع بقا را برای $Z > \text{Min}(\mu_1, \mu_2)$ محاسبه می‌کنیم. با فرض این که $\mu_1 < \mu_2$ باشد، داریم

$$\bar{H} = \begin{cases} (\mu_1/z)^{a_1}, & \mu_1 < Z < \mu_2 \\ (\mu_1/z)^{a_1} (\mu_2/z)^{a_2}, & \mu_2 < Z < \infty \end{cases}$$

همچنین

$$h(z) = \begin{cases} a_1 \mu_1^{a_1} / z^{a_1+1}, & \mu_1 < Z < \mu_2 \\ (a_1 + a_2) \mu_1^{a_1} \mu_2^{a_2} / z^{a_1+a_2+1}, & \mu_2 < Z < \infty \end{cases}$$

تحت مفروضات فوق اگر $\mu_2 < \mu_1$ باشد، تابع بقا دارای فرم دیگری است، در نتیجه می‌توان گفت یک سیستم k مولفه‌ای بسته به نوع ارتباط بین پارامترهای خود دارای $k!$ تابع بقا، همچنین قابلیت اطمینان است. برآوردگر و توزیع برآوردگر قابلیت اطمینان سیستم را تحت فرض $\mu_1 < \mu_2$ ، به دست می‌آوریم.

$$R = \int_{\mu_{k+1}}^{\infty} \bar{H}(x) dF_{k+1}(x), \quad \mu_{(1)} \leq \mu_{\mu_{k+1}} \quad (12.2)$$

$$= 1 - \int_{\mu_{(1)}}^{\infty} \bar{F}_{k+1}(x) dH(x), \quad \mu_{k+1} \leq \mu_{(1)} \quad (13.2)$$

در نتیجه قابلیت اطمینان به صورت زیر محاسبه می‌شود.

$$\begin{aligned}
 R &= \int_{\mu_{k+1}}^{\infty} \bar{H}(x) dF_{k+1}(x) \\
 &= \int_{\mu_{k+1}}^{\infty} \frac{(\prod_{j=1}^r \mu_{i_j}^{a_{i_j}} / x^{\sum_{j=1}^r a_{i_j}}) a_{k+1} \mu_{k+1}^{a_{k+1}}}{x^{a_{k+1}+1}} \\
 &= \prod_{j=1}^r \mu_{i_j}^{a_{i_j}} a_{k+1} \mu_{k+1}^{a_{k+1}} \int_{\mu_{k+1}}^{\infty} \frac{1}{x^{\sum_{j=1}^r a_{i_j} + a_{k+1} + 1}} \\
 &= \prod_{j=1}^r \mu_{i_j}^{a_{i_j}} a_{k+1} \mu_{k+1}^{a_{k+1}} \left(\frac{1}{(\sum_{j=1}^r a_{i_j} + a_{k+1}) \mu_{k+1}^{\sum_{j=1}^r a_{i_j} + a_{k+1}}} \right) \\
 &= \frac{a_{k+1} \prod_{j=1}^r \mu_{i_j}^{a_{i_j}}}{(\sum_{j=1}^r a_{i_j} + a_{k+1}) \mu_{k+1}^{\sum_{j=1}^r a_{i_j} + a_{k+1}}} \\
 &= 1 - \sum_{r=1}^k \frac{(\sum_{j=1}^r a_{i_j}) \mu_{k+1}^{a_{k+1}} (\prod_{j=1}^r \mu_{i_j}^{a_{i_j}})}{(\sum_{j=1}^r a_{i_j} + a_{k+1})} \\
 &\quad \times \left[\frac{1}{\mu_{i_r}^{(\sum_{j=1}^r a_{i_j} + a_{k+1})}} - \frac{1}{\mu_{i_{r+1}}^{(\sum_{j=1}^r a_{i_j} + a_{k+1})}} \right] \quad \mu_{k+1} \leq \mu_{i_1} \leq \dots \leq \mu_{i_k}.
 \end{aligned}$$

برآوردگر ماکزیمم درست‌نمایی

بنابر فرضیات فوق و با توجه به تابع چگالی احتمال توزیع پاراتو در حالت دوم تابع ماکزیمم درست‌نمایی به صورت زیر می‌باشد.

$$L = \prod_{i=1}^{k+1} a_i^n \mu_i^{na_i} \prod_{i=1}^{k+1} \prod_{j=1}^n \frac{I(\mu < X_{ij}) X_{ij}}{X_{ij}^{a_i+1}}. \quad (14.2)$$

پس از لگاریتم‌گیری از تابع درست‌نمایی و مشتق‌گیری از آن نسبت به پارامترها، معادلات درست‌نمایی به صورت زیر حاصل می‌شود.

$$\frac{n}{a_i} + n \log \mu_i - \sum_{j=1}^n \log x_{ij} = 0, \quad i = 1, 2, \dots, k+1$$

با حل معادلات درست‌نمایی حاصله، برآورد ماکزیمم درست‌نمایی پارامترها برابر است با.

$$\hat{a}_i = 1 / \log(G_i / \hat{\mu}_i), \quad i = 1, \dots, k+1 \quad (15.2)$$

$$\hat{\mu}_i = X_{i(1)} = \text{Min}(X_{i1}, \dots, X_{in}), \quad (16.2)$$

بطوری که

$$G_i = [\prod_{j=1}^n x_{ij}]^{1/n}$$

همانند قسمت‌های قبل برای به دست آوردن توزیع برآوردگر قابلیت اطمینان، ابتدا با استفاده ماتریس اطلاع فشر واریانس $\hat{a}_i$ را مشخص می‌کنیم.

$$I_{ij} = n/a_i^2,$$

$$V(\hat{a}_i) = a_i^2/n \quad i = 1, \dots, k+1$$

از طرفی بنابر روابط بین توزیع‌ها، می‌دانیم $\hat{\mu}_i$ دارای توزیع پاراتو با پارامترهای (na_i, μ_i) می‌باشد. در نتیجه

$$V(\hat{\mu}_i) = na_i \mu_i^2 / [(na_i - 2)(na_i - 1)^2].$$

اکنون بنابر خاصیت پایایی، برآوردگر ماکزیم درستنمایی قابلیت اطمینان سیستم را با قرار دادن مقادیر $\hat{a}_i$ و $\hat{\mu}_i$ در رابطه ۱۴.۲، به دست می‌آوریم. از آنجایی که $\hat{\mu} \xrightarrow{p} \mu_i$ ، همچنین دارای تابع توزیع پاراتو با پارامترهای مشخص است. لذا بنابر قضیه (۱۰.۲.۲)، $\hat{R}$ دارای توزیع مجانبی زیر است.

$$\hat{R} \sim AN(R, B_{\Psi}' \Lambda_{\Psi} B_{\Psi} / n)$$

بطوری که:

$$B_{\Psi}' = (\partial R / \partial a_1, \dots, \partial R / \partial a_{k+1}, \partial R / \partial \mu_1, \dots, \partial R / \partial \mu_{k+1})$$

$$\Lambda_{\Psi} = \frac{1}{n} \text{diag} \left(a_1^2, \dots, a_{k+1}^2, \frac{n^2 a_1 \mu_1^2}{(na_1 - 2)(na_1 - 1)^2}, \dots, \frac{n^2 a_{k+1} \mu_{k+1}^2}{(na_{k+1} - 2)(na_{k+1} - 1)^2} \right). \quad 9$$

فصل ۳

برآورد قابلیت اطمینان سیستم‌های موازی

در این فصل به طور کاملتر به برآورد قابلیت اطمینان سیستم‌های موازی می‌پردازیم. ابتدا با فرض نمایی بودن توزیع طول عمر مولفه‌های فشار و مقاومت، برآورد قابلیت اطمینان را به دو روش ماکزیمم درستمایی و بیز محاسبه می‌کنیم. سپس این دو روش برآورد را در حالتی به کار می‌بریم که توزیع طول عمر مولفه‌ها وایبل در نظر گرفته شوند. در ادامه با استفاده از روش‌های انتگرال گیری عددی زنجیر مارکوفی مونت کارلو قابلیت اطمینان را شبیه‌سازی می‌کنیم. لازم بذکر است مقایسه این دو برآوردگر به دلیل نداشتن فرم بسته برآوردگرها، امکان‌پذیر نیست. لذا از مقایسه روش‌ها صرفه نظر می‌کنیم.

۱.۳ مدل فشار-مقاومت توزیع نمایی

در این بخش فرض می‌کنیم سیستم موازی بوده و مولفه‌های فشار و مقاومت دارای طول عمری با توزیع نمایی باشند. لذا ابتدا تحت لم زیر تابع بقای سیستم را به دست می‌آوریم.

لم ۱.۱.۳. فرض کنید در یک سیستم موازی با k مولفه، که همگی تحت فشار یکسان (X_{k+1}) قرار دارند طول عمر مولفه‌ها دارای توزیع نمایی با پارامتر λ باشد. طبق تعریفی که از سیستم‌های موازی در فصل اول داشتیم، عملکرد سیستم‌های موازی در گرو عملکرد حداقل مولفه‌ها است، یعنی تازمانی که کلیه مولفه از کار نیفتد، سیستم کار می‌کند، لذا برای محاسبه قابلیت اطمینان سیستم، آستانه تحمل مولفه‌ای با ماکزیمم قدرت در برابر فشار وارده را محاسبه می‌کنیم بنابراین متغیر Y را به صورت $Y = \max(X_1, X_2, X_3, \dots, X_n)$ تعریف کنیم، آنگاه قابلیت اطمینان سیستم موازی k مولفه‌ای به صورت زیر تعریف می‌شود

$$\begin{aligned} R &= P(Y < X_{k+1}) \\ &= \int_0^\infty \bar{G}(x_{k+1}) dF_{k+1}(x_{k+1}) \\ &= \sum_{r=1}^k (-1)^{r-1} \sum_{i_1 > \dots}^k \dots \sum_{> i_r} \frac{\lambda_{k+1}}{(\lambda_{k+1} + \sum_{j=1}^r \lambda_{i_j})}. \end{aligned}$$

برهان. ابتدا داریم

$$\begin{aligned} G(y) &= P(Y \leq y) \\ &= P(X_1 \leq y, X_2 \leq y, \dots, X_k \leq y) \\ &= (1 - e^{-\lambda_1 y})(1 - e^{-\lambda_2 y}) \dots (1 - e^{-\lambda_k y}) \\ &= 1 - \sum_{r=1}^k (-1)^{r-1} \sum_{i_1 > \dots > i_r} e^{-\sum_{j=1}^r \lambda_{i_j} y}. \end{aligned}$$

برای اثبات این رابطه، مشابه اثبات رابطه (۱.۴) عمل می‌کنیم. بنابراین، قابلیت اطمینان به صورت زیر است.

$$\begin{aligned} R &= \int_0^\infty \bar{G}(x_{k+1}) dF_{k+1}(x_{k+1}) \\ &= \int_0^\infty \sum_{r=1}^k (-1)^{r-1} \sum_{i_1 > \dots > i_r} e^{-\sum_{j=1}^r \lambda_{i_j} x_{k+1}} \lambda_{k+1} e^{-\lambda_{k+1} x_{k+1}} dx_{k+1} \\ &= \sum_{r=1}^k (-1)^{r-1} \sum_{i_1 > \dots > i_r} \int_0^\infty e^{-\sum_{j=1}^r \lambda_{i_j} x_{k+1}} \lambda_{k+1} e^{-\lambda_{k+1} x_{k+1}} dx_{k+1} \\ &= \sum_{r=1}^k (-1)^{r-1} \sum_{i_1 > \dots > i_r} \frac{\lambda_{k+1}}{(\lambda_{k+1} + \sum_{j=1}^r \lambda_{i_j})}. \end{aligned}$$

□

۱.۱.۳ برآورد ماکزیمم درستنمایی

فرض کنید سیستم دارای n نمونه داریم، با K مولفه موازی اصلی $n, \dots, 1, j, X_{kj}, \dots, X_{1j}$ ، که تحت فشار یکسان $X_{(k+1)j}$ (فشار روی j امین نمونه) است، به‌طوری‌که $\lambda = (\lambda_1, \dots, \lambda_{k+1})$ بردار پارامترها و $\mathbf{X} = (x_{11}, \dots, x_{1n}, \dots, x_{(k+1)1}, \dots, x_{(k+1)n})$ بردار نمونه باشد که دارای $(k+1)n$ بعد می‌باشد.

تابع درستنمایی برای توزیع نمایی به صورت

$$\begin{aligned} L(\lambda; \mathbf{X}) &= \prod_{j=1}^n \prod_{i=1}^{k+1} \lambda_i e^{-\lambda_i x_{ij}} \\ &= \prod_{j=1}^n (e^{-\sum_{i=1}^{k+1} \lambda_i x_{ij}} \cdot \prod_{i=1}^{k+1} \lambda_i) \\ &= e^{-\sum_{i=1}^{k+1} \sum_{j=1}^n \lambda_i x_{ij}} \cdot \left(\prod_{i=1}^{k+1} \lambda_i \right)^n \\ &= e^{-\sum_{i=1}^{k+1} \sum_{j=1}^n \lambda_i x_{ij}} \cdot \prod_{i=1}^{k+1} \lambda_i^n, \end{aligned}$$

است. با لگاریتم گرفتن از طرفین رابطه فوق، داریم

$$\begin{aligned} l(\lambda; \mathbf{X}) &= \ln(e^{-\sum_{i=1}^{k+1} \sum_{j=1}^n \lambda_i x_{ij}} \cdot \prod_{i=1}^{k+1} \lambda_i^n) \\ &= -\sum_{i=1}^{k+1} \sum_{j=1}^n \lambda_i x_{ij} + n \sum_{i=1}^{k+1} \ln(\lambda_i). \end{aligned}$$

با مشتق‌گیری از رابطه بالا نسبت به پارامتر λ ، برآورد ماکزیمم درستنمایی این پارامتر برابر است با

$$\hat{\lambda}_i = \frac{n}{\sum_{j=1}^n x_{ij}}.$$

بنابراین برآورد درستنمایی قابلیت اطمینان بنابر خاصیت پایایی برابر است با

$$\hat{R} = \sum_{r=1}^k (-1)^{r-1} \sum_{i_1 > \dots > i_r} \frac{\hat{\lambda}_{k+1}}{(\hat{\lambda}_{k+1} + \sum_{j=1}^r \hat{\lambda}_{i_j})}.$$

۲.۱.۳ برآورد بیزی

برای یافتن برآورد بیزی قابلیت اطمینان، ابتدا باید توزیع پسین پارامترها $\lambda = (\lambda_1, \dots, \lambda_{k+1})$ را مشخص نماییم. فرض کنید λ_i ها متغیرهای تصادفی مستقل و هم‌توزیع با توزیع پیشین $\Gamma(a_i, b_i)$ باشند. در این صورت فرم عمومی تابع توزیع پسین به صورت زیر است

$$\begin{aligned} \pi(\lambda|\mathbf{X}) \propto f(\mathbf{X}|\lambda) \cdot \pi(\lambda) &= L(\lambda; \mathbf{X}) \cdot \pi(\lambda) \\ &\propto e^{-\sum_{i=1}^{k+1} \lambda_i \sum_{j=1}^n x_{ij}} \left(\prod_{i=1}^{k+1} \lambda_i^n \right) \cdot \prod_{i=1}^{k+1} \lambda_i^{a_i-1} \cdot e^{-b_i \lambda_i} \\ &= e^{-\sum_{i=1}^{k+1} \lambda_i \sum_{j=1}^n x_{ij}} \cdot \prod_{i=1}^{k+1} e^{-b_i \lambda_i} \cdot \prod_{i=1}^{k+1} \lambda_i^{n+a_i-1} \\ &= e^{-\sum_{i=1}^{k+1} \lambda_i \sum_{j=1}^n x_{ij} - \sum_{i=1}^{k+1} b_i \lambda_i} \cdot \prod_{i=1}^{k+1} \lambda_i^{n+a_i-1} \\ &= e^{-\sum_{i=1}^{k+1} \lambda_i (\sum_{j=1}^n x_{ij} + b_i)} \cdot \prod_{i=1}^{k+1} \lambda_i^{n+a_i-1} \\ &= \prod_{i=1}^{k+1} e^{\lambda_i (\sum_{j=1}^n x_{ij} + b_i)} \cdot \lambda_i^{n+a_i-1}. \end{aligned}$$

با توجه به فرم عمومی تابع توزیع پسین، می‌توان نتیجه گرفت که این توزیع پسین از توزیع گاما پیروی می‌کند. بدین ترتیب برآورد بیزی λ_i به صورت

$$E(\lambda_i|\mathbf{X}) = \frac{n + a_i}{b_i + \sum_{j=1}^n x_{ij}},$$

می‌باشد. به همین ترتیب برآورد بیزی قابلیت اطمینان عبارتست از

$$\hat{R}_b = E(R|\mathbf{X})$$

$$\begin{aligned}
&= E \left(\sum_{r=1}^k (-1)^{r-1} \sum_{i_1 >}^k \cdots \sum_{i_r >}^k \frac{\lambda_{k+1}}{(\lambda_{k+1} + \sum_{j=1}^r \lambda_{i_j})} | \mathbf{X} \right) \\
&= \int_0^\infty \cdots \int_0^\infty \left(\sum_{r=1}^k (-1)^{r-1} \sum_{i_1 >}^k \cdots \sum_{i_r >}^k \frac{\lambda_{k+1}}{(\lambda_{k+1} + \sum_{j=1}^r \lambda_{i_j})} \right) \\
&\quad \cdot \pi(\lambda_1, \dots, \lambda_{k+1} | X_1, \dots, X_{k+1}) d\lambda_1 \cdots d\lambda_{k+1} \\
&= \int_0^\infty \cdots \int_0^\infty \left(\sum_{r=1}^k (-1)^{r-1} \sum_{i_1 >}^k \cdots \sum_{i_r >}^k \frac{\lambda_{k+1}}{(\lambda_{k+1} + \sum_{j=1}^r \lambda_{i_j})} \right) \\
&\quad \lambda_1^{n+a_1-1} e^{b_1 + \sum_{j=1}^n x_{1j}} \cdots \lambda_k^{n+a_{k+1}-1} e^{b_{k+1} + \sum_{j=1}^n x_{k+1j}} d\lambda_1, \dots, d\lambda_{k+1}, \quad (1.3)
\end{aligned}$$

برای محاسبه رابطه (۱.۳) به دلیل وجود انتگرال‌های چندگانه، و پیچیدگی رابطه از الگوریتم‌های عددی بر پایه روش‌های MCMC استفاده می‌کنیم. در این مورد، به خاطر این که تابع داخل انتگرال دارای فرم مشخص است و می‌توان از آن به راحتی تولید نمونه کرد روش نمونه‌گیری گیبز را به کار می‌بریم.

۲.۳ مدل فشار-مقاومت توزیع وایبل

در این بخش فرض کنید سیستم موازی، با k مولفه، مولفه‌ها دارای طول عمری با توزیع وایبل، با پارامتر شکل a و پارامترهای مقیاس مختلف $b_i; (i = 1, 2, \dots, k)$ هستند. تابع چگالی احتمال متغیرهای فشار و مولفه‌ها به صورت زیر نمایش داده می‌شود

$$f_i(x_i) = ab_i x_i^{a-1} e^{-b_i x_i^a}; i = 1, 2, \dots, k+1.$$

بنابر تعریف سیستم‌های موازی، برای برآورد قابلیت اطمینان این سیستم‌ها ابتدا تابع بقا را به صورت لم زیر بیان می‌کنیم.

لم ۱.۲.۳. اگر متغیر Y را به صورت $Y = \max(X_1, X_2, X_3, \dots, X_k)$ تعریف کنیم، آنگاه تابع بقا سیستم موازی k مولفه ای عبارتست از

$$G_k(y) = 1 - \sum_{r=1}^k (-1)^{r-1} \sum_{i_1 >}^k \cdots \sum_{i_r >}^k e^{-\sum_{j=1}^r b_{i_j} y^a}.$$

برهان. برای اثبات درستی رابطه از استقرا استفاده می‌شود. ابتدا نشان می‌دهیم رابطه به ازای $k = 1$ برقرار است.

$$\begin{aligned}
G_1(y) &= 1 - \sum_{r=1}^1 (-1)^{r-1} \sum_{i_1 >}^1 \cdots \sum_{i_r >}^1 e^{-\sum_{j=1}^r b_{i_j} y^a} \\
&= 1 - (-1)^0 e^{-b_1 y^a} \\
&= 1 - e^{-b_1 y^a}.
\end{aligned}$$

از طرفی خواهیم داشت

$$G_{n+1}(y) = 1 - \sum_{r=1}^{n+1} (-1)^{r-1} \sum_{i_1 >}^{n+1} \cdots \sum_{i_r >}^{n+1} e^{-\sum_{j=1}^r b_{i_j} y^a}.$$

برای ادامه اثبات، ابتدا مقادیر مختلف r را در $G_{n+1}(y) - 1$ قرار می‌دهیم. برای $r = 1$ داریم

$$\begin{aligned} A &= (-1)^{1-1} \sum_{i_1=1}^{n+1} e^{-\sum_{j=1}^1 b_{i_j} y^a} \\ &= (-1)^{1-1} \sum_{i_1=1}^{n+1} e^{-b_{i_1} y^a} \\ &= e^{-b_{n+1} y^a} + \sum_{i_1=1}^n e^{-b_{i_1} y^a}. \end{aligned} \quad (2.3)$$

برای $r = 2$ خواهیم داشت

$$\begin{aligned} B &= (-1)^{2-1} \sum_{i_1 > i_2}^{n+1} \sum_{i_2=1}^{n+1} e^{-\sum_{j=1}^2 b_{i_j} y^a} \\ &= (-1)^{2-1} \left(e^{-b_{n+1} y^a} \cdot \sum_{i_2=1}^n e^{-b_{i_2} y^a} + \sum_{i_1 > i_2}^n \sum_{i_2=1}^n e^{-(b_{i_1} + b_{i_2}) y^a} \right). \end{aligned} \quad (3.3)$$

و به همین ترتیب برای $r = 3$ و $r = n + 1$ می‌توان نوشت

$$\begin{aligned} C &= (-1)^{3-1} \sum_{i_1 > i_2 > i_3}^{n+1} \sum_{i_2=1}^{n+1} \sum_{i_3=1}^{n+1} e^{-(b_{i_1} + b_{i_2} + b_{i_3}) y^a} \\ &= (-1)^{3-1} e^{b_{n+1} y^a} \sum_{i_2 > i_3} \sum_{i_3=1}^n e^{-(b_{i_2} + b_{i_3}) y^a} + (-1)^{3-1} \sum_{i_1 > i_2 > i_3}^n \sum_{i_2=1}^n \sum_{i_3=1}^n e^{-(b_{i_1} + b_{i_2} + b_{i_3}) y^a}. \end{aligned} \quad (4.3)$$

$r = n + 1$

$$\begin{aligned} D &= (-1)^{n+1-1} \sum_{i_1 > i_2 > \dots > i_{n+1}}^{n+1} \sum_{i_2=1}^{n+1} \dots \sum_{i_{n+1}=1}^{n+1} e^{-\sum_{j=1}^{n+1} b_{i_j} y^a} \\ &= (-1)^{n+1-1} e^{b_{n+1} y^a} \sum_{i_2 > i_3 > \dots > i_{n+1}} \sum_{i_3=1}^n \dots \sum_{i_{n+1}=1}^n e^{-\sum_{j=2}^{n+1} b_{i_j} y^a} \\ &\quad + (-1)^{n+1-1} \sum_{i_1 > i_2 > \dots > i_{n+1}}^n \sum_{i_2=1}^n \dots \sum_{i_{n+1}=1}^n e^{-\sum_{j=1}^{n+1} b_{i_j} y^a}. \end{aligned} \quad (5.3)$$

فرض می‌کنیم به ازای $k = n$ برقرار است ثابت می‌کنیم که به ازای $k = n + 1$ نیز برقرار است. با استفاده از تساوی‌های (۲.۳) تا (۵.۳) داریم

$$\begin{aligned} G_{n+1}(y) &= 1 - \sum_{r=1}^{n+1} (-1)^{r-1} \sum_{i_1 > i_2 > \dots > i_r}^{n+1} \sum_{i_r=1}^{n+1} e^{-\sum_{j=1}^r b_{i_j} y^a} \\ &= A + B + C + \dots + D \end{aligned}$$

$$\begin{aligned}
&= 1 - \left(e^{-b_{n+1}y^a} (1 + (-1)^{2-1} \sum_{i_2} e^{-b_{i_2}y^a} \right. \\
&\quad \left. + \dots + (-1)^{n+1-1} \sum_{i_2 > i_3} \sum_{i_3 > i_{n+1}} \dots \sum_{i_{n+1}} e^{-\sum_{j=2}^{n+1} b_{i_j} y^a} \right) \\
&\quad - \left(\sum_{i_1=1}^n e^{-b_{i_1}y^a} + \sum_{i_1 > i_2}^n \sum_{i_2 > i_3} e^{-(b_{i_1}+b_{i_2})y^a} \right. \\
&\quad \left. + \dots + (-1)^{n+1-1} \sum_{i_1 > i_2}^n \sum_{i_2 > i_3} \dots \sum_{i_{n+1}} e^{-\sum_{j=1}^{n+1} b_{i_j} y^a} \right) \\
&= 1 - e^{-b_{n+1}y^a} G_n(y) - (1 - G_n(y)) \\
&= G_n(y)(1 - e^{-b_{n+1}y^a}).
\end{aligned}$$

□

بنابراین با استفاده از رابطه فوق می‌توان قابلیت اطمینان را بدست آورد

$$\begin{aligned}
R &= \int_0^\infty \bar{G}(x_{k+1}) dF_{k+1}(x_{k+1}) \\
&= \int_0^\infty \sum_{r=1}^k (-1)^{r-1} \sum_{i_1 > i_2}^k \dots \sum_{i_r}^k e^{-\sum_{j=1}^r b_{i_j} x_{k+1}^a} a b_{k+1} x_{k+1}^{a-1} e^{-b_{k+1} x_{k+1}^a} dx_{k+1} \\
&= \sum_{r=1}^k (-1)^{r-1} \sum_{i_1 > i_2}^k \dots \sum_{i_r}^k \int_0^\infty e^{-\sum_{j=1}^r b_{i_j} x_{k+1}^a} a b_{k+1} x_{k+1}^{a-1} e^{-b_{k+1} x_{k+1}^a} dx_{k+1} \\
&= \sum_{r=1}^k (-1)^{r-1} \sum_{i_1 > i_2}^k \dots \sum_{i_r}^k \frac{b_{k+1}}{b_{k+1} + \sum_{j=1}^r b_{i_j}}.
\end{aligned}$$

۱.۲.۳ برآورد ماکزیمم درست‌نمایی

همانند حالت نمایی، فرض کنید سیستم دارای n نمونه، با K مولفه موازی اصلی $X_{1j}, \dots, X_{kj}$ که تحت فشار یکسان $X_{(k+1)j}$ (فشار روی j امین نمونه) هستند، باشد، $j = 1, \dots, n$ و $\mathbf{b} = (b_1, \dots, b_{k+1})$ بردار پارامترهای مقیاس هستند. تابع درست‌نمایی برای توزیع وایبل عبارتست از

$$\begin{aligned}
L(a, \mathbf{b}, \mathbf{X}) &= \prod_{j=1}^n \prod_{i=1}^{k+1} f(x_i; a, \mathbf{b}) \\
&= \prod_{j=1}^n \prod_{i=1}^{k+1} a b_i x_{i_j}^{a-1} \cdot e^{-b_i x_{i_j}^a} \\
&= a^{n(k+1)} \cdot \left(\prod_{i=1}^{k+1} b_i^n \right) \cdot \left(\prod_{j=1}^n \prod_{i=1}^{k+1} x_{i_j}^{a-1} \right) e^{-\sum_{i=1}^{k+1} b_i \sum_{j=1}^n x_{i_j}^a}.
\end{aligned}$$

با لگاریتم گرفتن از رابطه بدست آمده، داریم

$$l = \ln L(a, \mathbf{b}, \mathbf{X}) = n(k+1) \ln a + \sum_{i=1}^{k+1} n \ln b_i + \sum_{j=1}^n \sum_{i=1}^{k+1} (a-1) \ln x_{i_j} - \sum_{i=1}^{k+1} b_i \sum_{j=1}^n x_{i_j}^a,$$

برای بدست آوردن برآورد درستنمایی پارامترها، از رابطه فوق نسبت به پارامترها مشتق می‌گیریم

$$\frac{\partial l}{\partial a} = \frac{n(k+1)}{a} + \sum_{i=1}^{k+1} \sum_{j=1}^n \ln x_{ij} - \sum_{i=1}^{k+1} \sum_{j=1}^n b_i x_{ij}^a \ln x_{ij} = 0,$$

$$\frac{\partial l}{\partial b_i} = \frac{n}{b_i} - \sum_{j=1}^n x_{ij}^a = 0, \quad i = 1, 2, \dots, k+1.$$

این معادلات را می‌توان با استفاده از روش نیوتون-رافسون حل کرد. با استفاده از برآوردهای به‌دست آمده، برآورد درستنمایی ماکزیمم قابلیت اطمینان بنابر خاصیت پایایی، به صورت زیر به‌دست می‌آید

$$\hat{R} = \sum_{r=1}^k (-1)^{r-1} \sum_{i_1 >}^k \dots \sum_{i_r >}^k \frac{\hat{b}_{k+1}}{\hat{b}_{k+1} + \sum_{j=1}^r \hat{b}_{i_j}}$$

۲.۲.۳ برآورد بیزی

برای محاسبه برآورد بیزی قابلیت اطمینان، ابتدا فرض کنید $\beta_i = \log b_i$ ، $i = 1, 2, \dots, k+1$ ، آنگاه تابع چگالی وایبل به فرم زیر خواهد شد

$$f_i(x_i|a, \beta_i) = ax^{a-1}e^{\beta_i - e^{\beta_i}x^a}, \quad x_i > 0,$$

به طوری که $a > 0$ و $-\infty < \beta_i < \infty$ ، $i = 1, \dots, k+1$. فرض کنید $\beta = (\beta_1, \dots, \beta_{k+1})$ متغیرهای مستقل و هم‌توزیع و همچنین مستقل از a با توزیع‌های پیشین زیر باشند

$$a \sim \Gamma(\lambda, \gamma), \quad \beta_i \sim N(\mu_i, \sigma_i^2).$$

تابع چگالی پسین این متغیرها عبارتند از

$$\pi(a) = \frac{\gamma^\lambda}{\Gamma(\lambda)} a^{\lambda-1} e^{-\gamma a},$$

$$\pi(\beta_i) = \frac{1}{\sqrt{2\pi}\sigma_i} e^{-\frac{1}{2\sigma_i^2}(\beta_i - \mu_i)^2}, \quad \sigma_i^2 > 0.$$

به دلیل استقلال متغیرها، چگالی توام پیشین به صورت

$$\begin{aligned} \pi(a, \beta) &= \pi(a) \prod_{i=1}^{k+1} \pi(\beta_i) \\ &= \frac{\gamma^\lambda}{\Gamma(\lambda)} a^{\lambda-1} e^{-\gamma a} \prod_{i=1}^{k+1} \frac{1}{\sqrt{2\pi}\sigma_i} e^{-\frac{1}{2\sigma_i^2}(\beta_i - \mu_i)^2}, \end{aligned}$$

و چگالی توام پسین $a, \beta | \mathbf{X}$ به فرم

$$\pi(a, \beta | \mathbf{X}) \propto L(a, \beta_1, \dots, \beta_{k+1}) \pi(a) \prod_{i=1}^{k+1} \pi(\beta_i)$$

$$\propto a^{n(k+1)+\lambda-1} e^{\sum_{i=1}^{k+1} (n\beta_i - \sum_{j=1}^n e^{\beta_i} X_{ij}^a) + a(\sum_{i=1}^{k+1} \sum_{j=1}^n \log X_{ij} - \gamma)} e^{-\sum_{i=1}^{k+1} \frac{1}{\sigma_i^2} (\beta_i - \mu_i)^2}, \quad (6.3)$$

است. از رابطه (۶.۳)، هسته اصلی چگالی‌های پسین حاشیه‌ای را بدست می‌آوریم

$$\begin{aligned} \pi(a|\beta_i, \mathbf{X}) &\propto a^{n(k+1)+\lambda-1} e^{a \sum_{i=1}^{k+1} \sum_{j=1}^n \log X_{ij} - \sum_{i=1}^{k+1} e^{\beta_i} \sum_{j=1}^n X_{ij}^a - \gamma a}, \\ \pi(\beta_i|a, \mathbf{X}) &\propto e^{n\beta_i - e^{\beta_i} \sum_{j=1}^n X_{ij}^a - \frac{1}{2\sigma_i^2} (\beta_i - \mu_i)^2}. \end{aligned}$$

همانگونه که مشاهده می‌شود، این چگالی‌های پسین به هیچ توزیع شناخته‌شده‌ای تعلق ندارند. برآوردگر بیز قابلیت اطمینان حاصل عبارت زیر است

$$\begin{aligned} \hat{R}_b &= E(R|\mathbf{X}) \\ &= E \left(\sum_{r=1}^k (-1)^{r-1} \sum_{i_1 > \dots > i_r} \dots \sum_{i_r} \frac{b_{k+1}}{(b_{k+1} + \sum_{j=1}^r b_{i_j})} | \mathbf{X} \right) \\ &= \int_0^\infty \dots \int_0^\infty \left(\sum_{r=1}^k (-1)^{r-1} \sum_{i_1 > \dots > i_r} \dots \sum_{i_r} \frac{b_{k+1}}{(b_{k+1} + \sum_{j=1}^r b_{i_j})} \right) \\ &\quad \pi(a, b_1, \dots, b_{k+1} | \mathbf{X}) da db_1 \dots db_{k+1} \quad (7.3) \end{aligned}$$

برای حل انتگرال رابطه (۷.۳)، به دلیل وجود چگالی‌های حاشیه‌ای پسین با توزیع نامعلوم، از روش نمونه‌گیری متروپولیس-هستینگز، که از روش‌های انتگرال‌گیری MCMC می‌باشد، استفاده می‌کنیم.

۳.۳ مطالعه شبیه‌سازی

در این بخش با استفاده از یک مطالعه شبیه‌سازی، برآوردهای درست‌نمایی ماکزیم و بیز را برای هر دو توزیع نمایی و وایبل مورد بررسی قرار می‌دهیم. برای این منظور، یک سیستم سه مولفه‌ای موازی را با مقادیر متفاوت پارامترها برای توزیع نمایی و مقادیر b_i متفاوت به عنوان پارامتر مقیاس ولی پارامتر شکل یکسان a برای توزیع وایبل در نظر می‌گیریم. همچنین متغیر فشار برای توزیع نمایی، نمایی و برای توزیع وایبل، وایبل در نظر گرفته می‌شود.

۱.۳.۳ توزیع نمایی

با در نظر گرفتن چهار حجم نمونه $n = \{25, 50, 75, 100\}$ ، سه حالت متفاوت را برای هر حجم نمونه در نظر می‌گیریم. این سه حالت در جدول ۱.۳ نشان داده شده است.

برآورد درست‌نمایی ماکزیم

همانطور که در بخش ۱.۱.۳ بیان شد، ابتدا برآوردهای درست‌نمایی ماکزیم پارامترها رو می‌یابیم آن‌گاه برآورد درست‌نمایی ماکزیم قابلیت اطمینان را به دست می‌آوریم. نتایج شبیه‌سازی در جدول ۲.۳ نشان

جدول ۱۰.۳: پارامترهای توزیع نمایی

حالت	λ_1	λ_2	λ_3	λ_4
۱	۰/۱	۰/۲	۰/۳	۰/۲
۲	۰/۱۵	۰/۱	۰/۳	۰/۲
۳	۰/۱	۰/۲	۰/۱۵	۰/۲

داده شده است. این جدول شامل مقادیر برآورد درستنمایی ماکزیمم، انحراف معیار و فواصل اطمینان است.

همانطور که در جدول مشاهده می‌شود روش درستنمایی ماکزیمم برآوردهای بسیار دقیقی ارائه کرده است

جدول ۲۰.۳: برآورد درستنمایی ماکزیمم

حالت	n	R	$\hat{R}$	انحراف معیار	۹۵٪ فاصله اطمینان
۱	۲۵	۰/۷۹۷۶۱۹	۰/۷۹۴۰۸۹۱	۰/۰۰۲۵۶۴۶۶۶	۰/۷۸۹۰۶۲۵ – ۰/۷۹۹۱۱۵۸
۲	۲۵	۰/۸۱۹۲۹۱۸	۰/۸۱۲۲۳۰۷	۰/۰۰۱۸۹۱۷۷۷	۰/۸۰۸۵۲۲۹ – ۰/۸۱۵۹۳۸۵
۳	۲۵	۰/۸۳۷۷۰۶۷	۰/۸۴۰۷۶۷۷	۰/۰۰۱۸۳۵۴۹۴	۰/۸۳۷۱۷۰۲ – ۰/۸۴۴۳۶۵۲
۱	۵۰	۰/۷۹۷۶۱۹	۰/۷۹۲۰۲۲۷	۰/۰۰۱۱۶۸۸۴۲	۰/۷۸۹۷۳۱۸ – ۰/۷۹۴۳۱۳۵
۲	۵۰	۰/۸۱۹۲۹۱۸	۰/۸۲۳۱۵۴۲	۰/۰۰۱۱۱۸۵۰۰	۰/۸۲۰۹۶۲ – ۰/۸۲۵۳۴۶۴
۳	۵۰	۰/۸۳۷۷۰۶۷	۰/۸۴۵۳۹۵۶	۰/۰۰۱۲۲۰۳۴۶	۰/۸۴۳۰۰۳۸ – ۰/۸۴۷۷۸۷۴
۱	۷۵	۰/۷۹۷۶۱۹	۰/۸۰۳۵۸۹۷	۰/۰۰۰۸۴۱۱۹۱	۰/۸۰۱۹۴۱ – ۰/۸۰۵۲۳۸۴
۲	۷۵	۰/۸۱۹۲۹۱۸	۰/۸۲۴۳۱۲۸	۰/۰۰۰۵۹۵۰۰۵۵	۰/۸۲۳۱۴۶۶ – ۰/۸۲۵۴۷۹
۳	۷۵	۰/۸۳۷۷۰۶۷	۰/۸۴۰۶۱۲۵	۰/۰۰۰۸۱۳۸۱۲۹	۰/۸۳۹۰۱۷۵ – ۰/۸۴۲۲۰۷۶
۱	۱۰۰	۰/۷۹۷۶۱۹۰	۰/۷۹۴۱۴۵	۰/۰۰۰۷۲۶۴۸۱۳	۰/۷۹۲۷۲۱۱ – ۰/۷۹۵۵۶۸۸
۲	۱۰۰	۰/۸۱۹۲۹۱۸	۰/۸۱۸۱۲۱۸	۰/۰۰۰۷۱۴۶۷۲۰	۰/۸۱۶۷۲۱۱ – ۰/۸۱۶۷۲۱۱
۳	۱۰۰	۰/۸۳۷۷۰۶۷	۰/۸۳۲۷۸۳۴	۰/۰۰۰۸۴۱۷۰۷۹	۰/۸۳۱۱۳۳۷ – ۰/۸۳۴۴۳۳۱

و خطاهای بسیار ناچیزی دارد. با افزایش حجم نمونه از ۲۵ تا ۱۰۰ برآوردها دقیق‌تر شده و فواصل اطمینان کوتاه‌تر شده‌اند.

برآورد بیزی

همانطور که در بخش ۲.۱.۳ ذکر شد، برای به‌دست آوردن برآورد بیزی، به الگوریتم‌های MCMC نیاز داریم. با توجه به بخش مذکور، توزیع‌های پسین حاشیه‌ای و پارامترها دارای توزیع گاما هستند. بنابراین به راحتی و به طور مستقیم می‌توان از این توزیع‌ها شبیه‌سازی کرد و با استفاده از روش مونت کارلو، قابلیت اطمینان را برآورد کرد. برای به‌دست آوردن برآورد بیزی، دو حالت از توزیع‌های پیشین

آگاهی‌بخش^۱ و کم آگاهی‌بخش^۲ را در نظر می‌گیریم.

حالت اول، توزیع پیشین آگاهی‌بخش

در این حالت، توزیع‌های گامای پیشین با واریانس کم را در نظر می‌گیریم که باعث می‌شود توزیع‌های پیشین آگاهی‌بخش داشته باشیم. جدول ۳.۳ پارامترهای توزیع‌های پیشین آگاهی‌بخش گاما را در ۳ حالت نشان می‌دهد. برآوردهای بیز قابلیت اطمینان تحت تابع زیان درجه دوم خطا برای حجم نمونه‌های متفاوت به همراه مقادیر انحراف معیار و نواحی اطمینان در جدول ۴.۳ آمده است. با توجه به جدول

جدول ۳.۳: پارامترهای توزیع‌های پیشین

حالت	a_1	b_1	a_2	b_2	a_3	b_3	a_4	b_4
۱	۳۰	۳۰۰	۳۰	۱۵۰	۶۰	۲۰۰	۲۰	۱۰۰
۲	۳۰	۲۰۰	۳۵	۳۵۰	۶۰	۲۰۰	۲۰	۱۰۰
۳	۳۰	۳۰۰	۳۰	۱۵۰	۳۰	۲۰۰	۲۰	۱۰۰

جدول ۴.۳: برآورد بیز

حالت	n	R	$\hat{R}$	انحراف معیار	۹۵٪ فاصله اطمینان
۱	۲۵	۰/۷۹۷۶۱۹	۰/۷۹۵۸۰۱۹	۰/۰۰۰۰۸۲۶۴۱۳۹	۰/۷۹۴۱۸۲۱ – ۰/۷۹۷۴۲۱۶
۲	۲۵	۰/۸۱۹۲۹۱۸	۰/۸۱۳۰۲۱۲	۰/۰۰۰۰۶۶۷۲۸۱۱	۰/۸۱۱۷۱۳۴ – ۰/۸۱۴۳۲۹۱
۳	۲۵	۰/۸۳۷۷۰۶۷	۰/۸۳۸۳۶۹۴	۰/۰۰۰۰۶۴۶۶۶۸۹	۰/۸۳۷۱۰۲ – ۰/۸۳۹۶۳۶۹
۱	۵۰	۰/۷۹۷۶۱۹	۰/۷۹۳۸۹۹۳	۰/۰۰۰۰۶۷۸۳۰۳۲	۰/۷۹۲۵۶۹۹ – ۰/۷۹۵۲۲۸۸
۲	۵۰	۰/۸۱۹۲۹۱۸	۰/۸۲۰۷۸۳۶	۰/۰۰۰۰۶۴۲۲۴۸۹	۰/۸۱۹۵۲۴۹ – ۰/۸۲۲۰۴۲۴
۳	۵۰	۰/۸۳۷۷۰۶۷	۰/۸۴۲۷۸۶۴	۰/۰۰۰۰۶۳۱۸۵۰۴	۰/۸۴۱۵۴۸ – ۰/۸۴۴۰۲۴۸
۱	۷۵	۰/۷۹۷۶۱۹	۰/۸۰۰۷۶۷	۰/۰۰۰۰۴۷۶۲۷۶۷	۰/۷۹۹۸۳۳۵ – ۰/۸۰۱۷۰۰۵
۲	۷۵	۰/۸۱۹۲۹۱۸	۰/۸۲۲۲۱۸۴	۰/۰۰۰۰۴۰۰۰۹۱۵	۰/۸۲۱۴۳۴۳ – ۰/۸۲۳۰۰۲۶
۳	۷۵	۰/۸۳۷۷۰۶۷	۰/۸۳۹۷۷۶۴	۰/۰۰۰۰۵۸۸۴۱۷۱	۰/۸۳۸۶۲۳۱ – ۰/۸۴۰۹۲۹۶
۱	۱۰۰	۰/۷۹۷۶۱۹	۰/۷۹۳۱۱۳۷	۰/۰۰۰۰۵۸۱۵۵۹۷	۰/۷۹۱۹۷۳۹ – ۰/۷۹۴۲۵۳۶
۲	۱۰۰	۰/۸۱۹۲۹۱۸	۰/۸۱۷۵۶۱	۰/۰۰۰۰۴۹۶۸۶۵	۰/۸۱۶۵۸۷۲ – ۰/۸۱۸۵۳۴۹
۳	۱۰۰	۰/۸۳۷۷۰۶۷	۰/۸۳۴۴۳۲	۰/۰۰۰۰۴۹۵۹۹۸۷	۰/۸۳۳۴۵۹۸ – ۰/۸۳۵۴۰۴۱

۴.۳، مشاهده می‌شود که بر خلاف برآورد درست‌نمایی ماکزیمم، با افزایش حجم نمونه، برآوردهای بیز با سرعت کمتری نسبت به برآوردهای درست‌نمایی ماکزیمم بهبود می‌یابند. همچنین فواصل اطمینان نیز با افزایش حجم نمونه تغییر محسوسی نمی‌کنند. در این روش نیز برآوردها دقیق و مقدار خطاها کوچک می‌باشند.

^۱Informative

^۲Less informative

حالت دوم، توزیع پیشین کم آگاهی‌بخش

در این حالت توزیع‌های پیشینی را در نظر می‌گیریم که تنها اطلاعات خیلی کمی را ارائه می‌دهند. پارامترهای توزیع پیشین به گونه‌ای انتخاب شده‌اند که توزیع پیشین پراکندگی زیادی داشته باشد. مقادیر پارامترها در جدول ۵.۳ و برآوردهای بیز، انحراف معیار و نواحی اطمینان در جدول ۶.۳ گزارش شده‌اند. با توزیع‌های پیشین با اطلاعات کم هم، با توجه به جدول ۶.۳، برآوردها مانند برآوردهای حاصل از

جدول ۵.۳: پارامترهای توزیع‌های پیشین

حالت	a_1	b_1	a_2	b_2	a_3	b_3	a_4	b_4
۱	۲	۰/۰۱	۲/۵	۰/۰۲	۲	۰/۰۲۵	۱/۵	۰/۰۱۵
۱	۲	۰/۰۱	۲/۵	۰/۰۲	۲	۰/۰۲۵	۱/۵	۰/۰۱۵
۱	۲	۰/۰۱	۲/۵	۰/۰۲	۲	۰/۰۲۵	۱/۵	۰/۰۱۵

جدول ۶.۳: برآورد بیزی

حالت	n	R	$\hat{R}$	انحراف معیار	۹۵٪ فاصله اطمینان
۱	۲۵	۰/۷۹۷۶۱۹	۰/۷۷۷۰۳۹	۰/۰۰۳۵۸۲۹۱۸	۰/۷۷۰۰۱۶۶ – ۰/۷۸۴۰۶۱۴
۲	۲۵	۰/۸۱۹۲۹۱۸	۰/۸۱۵۴۸۲۹	۰/۰۰۲۷۱۶۳۰۳	۰/۸۱۰۱۵۹ – ۰/۸۲۰۸۰۶۷
۳	۲۵	۰/۸۳۷۷۰۶۷	۰/۸۲۴۶۶۸۲	۰/۰۰۲۲۹۱۱۶۴	۰/۸۲۰۱۷۷۶ – ۰/۸۲۹۱۵۸۸
۱	۵۰	۰/۷۹۷۶۱۹	۰/۷۹۵۵۶۸۶	۰/۰۰۱۲۲۲۷۵۳	۰/۷۹۳۱۷۲۱ – ۰/۷۹۷۹۶۵۲
۲	۵۰	۰/۸۱۹۲۹۱۸	۰/۸۱۰۱۵۸۹	۰/۰۰۱۵۳۰۸۱۴	۰/۸۰۷۱۵۸۵ – ۰/۸۱۳۱۵۹۲
۳	۵۰	۰/۸۳۷۷۰۶۷	۰/۸۳۵۳۸۵۳	۰/۰۰۰۹۶۷۳۲۶	۰/۸۳۳۴۸۹۴ – ۰/۸۳۷۲۸۱۳
۱	۷۵	۰/۷۹۷۶۱۹	۰/۷۹۷۶۹۹۳	۰/۰۰۱۱۶۸۷۰۹	۰/۷۹۵۴۰۸۷ – ۰/۷۹۹۹۹
۲	۷۵	۰/۸۱۹۲۹۱۸	۰/۸۲۳۶۵۷۱	۰/۰۰۰۹۰۸۰۶۲۷	۰/۸۲۱۸۷۷۴ – ۰/۸۲۵۴۳۶۹
۳	۷۵	۰/۸۳۷۷۰۶۷	۰/۸۲۹۲۱۵۹	۰/۰۰۰۹۴۳۴۶۰۶	۰/۸۲۷۳۶۶۷ – ۰/۸۳۱۰۶۵
۱	۱۰۰	۰/۷۹۷۶۱۹	۰/۷۹۳۷۴۸۲	۰/۰۰۰۹۰۴۷۵۱۱	۰/۷۹۱۹۷۴۹ – ۰/۷۹۵۵۲۱۵
۲	۱۰۰	۰/۸۱۹۲۹۱۸	۰/۸۱۴۷۷۸۲	۰/۰۰۰۶۹۹۶۳۹۳	۰/۸۱۳۴۰۷ – ۰/۸۱۶۱۴۹۵
۳	۱۰۰	۰/۸۳۷۷۰۶۷	۰/۸۳۸۴۸۷۵	۰/۰۰۰۹۵۲۳۴۳	۰/۸۳۶۶۲۰۹ – ۰/۸۴۰۳۵۴

توزیع‌های پیشین با اطلاعات بالا، با افزایش حجم نمونه با سرعت کمی بهبود می‌یابند. از طرفی اختلاف بین برآوردهای حاصل از توزیع‌های پیشین آگاهی‌بخش و کم آگاهی‌بخش، به طور تعجب‌برانگیزی کوچک است.

۲.۳.۳ توزیع وایبل

همانطور که در بخش ۲.۲.۳ اشاره شد، ابتدا با روش نیوتن-رافسون برآورد درست‌نمایی ماکزیمم پارامترها را می‌یابیم آن‌گاه برآورد قابلیت اطمینان را به دست می‌آوریم. همانند شبیه‌سازی برای توزیع نمایی، سه

حالت متفاوت با پارامترهای متفاوت را برای چهار حجم نمونه $n = \{25, 50, 75, 100\}$ در نظر می‌گیریم. این مقادیر در جدول ۷.۳ ذکر شده‌اند.

جدول ۷.۳: پارامترهای توزیع وایبل

حالت	a	b_1	b_2	b_3	b_4
۱	۰/۵	۲/۵	۵	۴	۳
۲	۰/۶	۲/۵	۴/۵	۲	۳
۳	۰/۴	۳/۵	۴	۴/۵	۳

جدول ۸.۳: برآورد درستنمایی ماکزیم

حالت	a	n	R	$\hat{R}$	انحراف معیار	۹۵٪ فاصله اطمینان
۱	۰/۵	۲۵	۰/۷۰۴۴۱۸۸	۰/۷۰۵۲۴۳۱	۰/۰۰۰۴۲۹۱۷۵۵	۰/۶۴۹۷۵۵۴ – ۰/۷۰۹۰۵۳۵
۲	۰/۶	۲۵	۰/۷۷۹۶۶۵۱	۰/۷۸۷۹۹۷۹	۰/۰۰۰۴۴۸۰۰۳۷	۰/۷۷۹۲۱۷۲ – ۰/۷۹۶۷۷۸۶
۳	۰/۴	۲۵	۰/۶۷۰۷۹۸۸	۰/۶۶۷۲۰۱۹	۰/۰۰۰۴۸۲۲۴۶۳	۰/۶۵۷۷۵ – ۰/۶۷۶۶۵۳۷
۱	۰/۵	۵۰	۰/۷۰۴۴۱۸۸	۰/۷۰۲۳۱۴۹	۰/۰۰۰۲۴۰۳۹۰۶	۰/۶۹۷۶۰۳۳ – ۰/۷۰۷۰۲۶۴
۲	۰/۶	۵۰	۰/۷۷۹۶۶۵۱	۰/۷۷۸۷۰۷۷	۰/۰۰۰۱۵۸۱۴۵۴	۰/۷۷۵۶۰۸۱ – ۰/۷۸۱۸۰۷۳
۳	۰/۴	۵۰	۰/۶۷۰۷۹۸۸	۰/۶۷۰۷۹۸۸	۰/۰۰۰۱۹۵۳۹۲۵	۰/۶۷۵۳۷۹ – ۰/۶۸۳۰۳۸۳
۱	۰/۵	۷۵	۰/۷۰۴۴۱۸۸	۰/۷۱۰۷۴۶۳	۰/۰۰۰۱۰۵۱۵۵۳	۰/۷۰۸۶۸۵۳ – ۰/۷۱۲۸۰۷۳
۲	۰/۶	۷۵	۰/۷۷۹۶۶۵۱	۰/۷۷۵۳	۰/۰۰۰۱۲۲۵۸۷۲	۰/۷۷۲۸۹۷۳ – ۰/۷۷۷۷۰۲۶
۳	۰/۴	۷۵	۰/۶۷۰۷۹۸۸	۰/۶۷۶۷۴۹۱	۰/۰۰۰۱۷۱۵۰۱۹	۰/۶۷۳۳۸۷۷ – ۰/۶۸۰۱۱۰۴
۱	۰/۵	۱۰۰	۰/۷۰۴۴۱۸۸	۰/۶۹۹۳۹۱۶	۰/۰۰۰۱۳۹۵۵۵۹	۰/۶۹۶۶۵۶۴ – ۰/۷۰۲۱۲۶۸
۲	۰/۶	۱۰۰	۰/۷۷۹۶۶۵۱	۰/۷۸۱۶۸۵۳	۰/۰۰۰۰۶۶۴۲۳۵	۰/۷۸۰۳۸۳۴ – ۰/۷۸۲۹۸۷۲
۳	۰/۴	۱۰۰	۰/۶۷۰۷۹۸۸	۰/۶۶۷۷۶۲۱	۰/۰۰۰۱۲۳۹۷۴۹	۰/۶۶۵۳۳۲۲ – ۰/۶۷۰۱۹۱۹

همانند نتایج شبیه‌سازی برآورد درستنمایی ماکزیم برای توزیع نمایی، برآوردها در توزیع وایبل نیز با افزایش حجم نمونه بهبود پیدا کرده‌اند و خطای برآورد و همچنین طول فاصله اطمینان کاهش پیدا کرده است. همانطور که از خطاها پیداست، روش درستنمایی ماکزیم حتی برای حجم نمونه کوچک ۲۵ نیز بسیار خوب عمل کرده است. این مزیت در موقعیت‌های عملی بسیار مهم است چرا که در واقعیت، ممکن است با حجم نمونه کوچکی برخورد کنیم.

برآورد بیزی

برای بدست آوردن برآورد بیزی برای توزیع وایبل همانطور که در بخش ۲.۲.۳ بیان شد، به دلیل نامعلوم بودن توزیع‌های پسین حاشیه‌ای، از روش متروپولیس-هستینگز استفاده می‌کنیم. همانند بدست آوردن برآورد بیزی برای توزیع نمایی، دو حالت توزیع پیشین آگاهی‌بخش و کم آگاهی‌بخش را در نظر می‌گیریم.

حالت اول، توزیع پیشین آگاهی‌بخش

در این حالت توزیع‌های پسینی در نظر می‌گیریم که در همسایگی میانگین توزیع متمرکز شده باشند. همانطور که قبلاً ذکر شد،

$$a \sim \Gamma(\lambda, \gamma), \quad \beta_i = \log b_i \sim N(\mu_i, \sigma_i); i = 1, \dots, 4.$$

جدول ۹.۳ شامل مقادیر پارامترها برای حالت مذکور می‌باشد. نتایج شبیه‌سازی، در جدول ۱۰.۳ ارائه شده‌اند.

جدول ۹.۳: پارامترهای توزیع‌های پیشین

حالت	λ	γ	μ_1	σ_1	μ_2	σ_2	μ_3	σ_3	μ_4	σ_4
۱	۵	۱۰	۲/۵	۱۰	۵	۱۰	۴	۱۰	۳	۱۰
۲	۸	۱۲	۲/۵	۱۰	۴/۵	۱۰	۲	۱۰	۳	۱۰
۳	۴	۱۰	۳/۵	۱۰	۴	۱۰	۴/۵	۱۰	۳	۱۰

جدول ۱۰.۳: پیرآورد بیز

حالت	n	R	$\hat{R}$	انحراف معیار	۹۵٪ فاصله اطمینان
۱	۲۵	۰/۷۰۴۴۱۸۸	۰/۷۱۰۱۵۵۲	۰/۰۰۵۴۴۸۲۴۱	۰/۶۹۹۴۷۶۸ – ۰/۷۲۰۸۳۳۵
۲	۲۵	۰/۷۷۹۶۶۵۱	۰/۸۰۰۹۵۶	۰/۰۰۳۸۴۴۹۶۶	۰/۷۹۳۴۲ – ۰/۸۰۸۴۹۲
۳	۲۵	۰/۶۷۰۷۹۸۸	۰/۶۸۵۶۹۳۷	۰/۰۰۲۴۴۳۵۶۸	۰/۶۸۰۹۰۴۴ – ۰/۶۹۰۴۸۳
۱	۵۰	۰/۷۰۴۴۱۸۸	۰/۷۲۶۶۱۸۶	۰/۰۰۱۳۸۷۸۹۷	۰/۷۲۳۸۹۸۴ – ۰/۷۲۹۳۳۸۸
۲	۵۰	۰/۷۷۹۶۶۵۱	۰/۷۸۲۰۶۶۵	۰/۰۰۱۷۲۶۵۰۶	۰/۷۷۸۶۸۲۶ – ۰/۷۸۵۴۵۰۴
۳	۵۰	۰/۶۷۰۷۹۸۸	۰/۶۶۰۴۲۷	۰/۰۰۳۰۸۴۵۴۲	۰/۶۵۴۳۸۱۴ – ۰/۶۶۶۴۷۲۶
۱	۷۵	۰/۷۰۴۴۱۸۸	۰/۷۲۳۶۵۹	۰/۰۰۱۳۲۲۰۸۹	۰/۷۲۱۰۶۷۷ – ۰/۷۲۶۲۵۰۲
۲	۷۵	۰/۷۷۹۶۶۵۱	۰/۷۸۷۰۳۲۲	۰/۰۰۱۷۸۸۲۴۹	۰/۷۸۳۵۲۷۳ – ۰/۷۹۰۵۳۷۱
۳	۷۵	۰/۶۷۰۷۹۸۸	۰/۶۹۳۷۱۳۸	۰/۰۰۱۷۲۸۱۷۸	۰/۶۹۰۳۲۶۷ – ۰/۶۹۷۱۰۱
۱	۱۰۰	۰/۷۰۴۴۱۸۸	۰/۷۲۴۷۹۶۹	۰/۰۰۰۹۹۸۳۱۵۷	۰/۷۲۲۸۴۰۲ – ۰/۷۲۶۷۵۳۵
۲	۱۰۰	۰/۷۷۹۶۶۵۱	۰/۷۸۱۳۵۴۲	۰/۰۰۰۴۸۴۵۰۶۳	۰/۷۸۰۴۰۴۶ – ۰/۷۸۲۳۰۳۸
۳	۱۰۰	۰/۶۷۰۷۹۸۸	۰/۶۷۷۴۳۵۶	۰/۰۰۰۶۴۳۶۵۳۵	۰/۶۷۶۱۷۴۱ – ۰/۶۷۸۶۹۷۲

با توجه به اینکه نتایج شبیه‌سازی قابل قبول می‌باشد، اما همانند نتایج شبیه‌سازی برای مدل فشار مقاومت نمایی، برآوردها و فواصل با سرعت کمی بهبود می‌یابند.

حالت دوم، توزیع پیشین کم آگاهی‌بخش

برای تولید توزیع‌های پیشینی که اطلاعات کمی دارند، پارامترهای توزیع گاما و نرمال به گونه‌ای انتخاب می‌شوند که این توزیع‌های پیشین پراکندگی زیادی داشته باشند. مقادیر پارامترهای این توزیع‌ها در جدول ۱۱.۳ آمده‌اند. همانگونه که مشاهده می‌شود، مقادیر انحراف معیار در حالت کم آگاهی‌بخش، ۱۰۰ برابر

بزرگتر از حالتی است که توزیع‌های پیشین آگاهی‌بخش استفاده شد. در این حالت تقریباً اطلاعاتی وجود ندارد. جدول ۱۲.۳، مقادیر شبیه‌سازی را همراه با انحراف معیار و نواحی اطمینان برای حالت‌های

جدول ۱۱.۳: پارامترهای توزیع‌های پیشین

حالت	λ	γ	μ_1	σ_1	μ_2	σ_2	μ_3	σ_3	μ_4	σ_4
۱	۰/۵	۰/۰۱	۲/۵	۱۰۰۰	۵	۱۰۰۰	۴	۱۰۰۰	۳	۱۰۰۰
۲	۰/۶	۰/۰۱	۲/۵	۱۰۰۰	۴/۵	۱۰۰۰	۲	۱۰۰۰	۳	۱۰۰۰
۳	۰/۴	۰/۰۱	۳/۵	۱۰۰۰	۴	۱۰۰۰	۴/۵	۱۰۰۰	۳	۱۰۰۰

مختلفی که در جدول ۱۱.۳ ذکر شد، نشان می‌دهد. در این حالت نیز با افزایش حجم نمونه، برآوردها

جدول ۱۲.۳: پیرآورد بیز

حالت	n	R	$\hat{R}$	انحراف معیار	۹۵٪ فاصله اطمینان
۱	۲۵	۰/۷۰۴۴۱۸۸	۰/۷۱۳۷۵۹	۰/۰۰۰۳۰۸۷۴۱۶	۰/۷۰۷۷۰۷۸ – ۰/۷۱۹۸۱۰۳
۲	۲۵	۰/۷۷۹۶۶۵۱	۰/۷۹۸۵۱۴۷	۰/۰۰۰۲۲۶۵۵۴۷	۰/۷۹۴۰۷۴۳ – ۰/۸۰۲۹۵۵۱
۳	۲۵	۰/۶۷۰۷۹۸۸	۰/۶۸۳۱۵۱۹	۰/۰۰۰۳۶۰۰۳۲۷	۰/۶۷۶۰۹۵۴ – ۰/۶۹۰۲۰۸۴
۱	۵۰	۰/۷۰۴۴۱۸۸	۰/۷۲۴۲۲۹۶	۰/۰۰۰۱۴۹۳۹۱۳	۰/۷۲۱۳۰۱۶ – ۰/۷۲۷۱۵۷۶
۲	۵۰	۰/۷۷۹۶۶۵۱	۰/۷۹۵۱۶۴۵	۰/۰۰۰۴۲۳۸۰۵۲	۰/۷۸۶۸۵۸۱ – ۰/۸۰۳۴۷۰۹
۳	۵۰	۰/۶۷۰۷۹۸۸	۰/۶۸۷۷۰۷۵	۰/۰۰۰۲۷۷۲۴۸۸	۰/۶۸۲۲۷۳۵ – ۰/۶۹۳۱۴۱۴
۱	۷۵	۰/۷۰۴۴۱۸۸	۰/۷۰۸۱۸۶۱	۰/۰۰۰۱۰۷۵۳۸۴	۰/۷۰۶۰۷۸۴ – ۰/۷۱۰۲۹۳۹
۲	۷۵	۰/۷۷۹۶۶۵۱	۰/۷۹۳۱۳۸۵	۰/۰۰۰۰۷۴۴۶۶۶۵	۰/۷۹۱۶۷۹ – ۰/۷۹۴۵۹۸
۳	۷۵	۰/۶۷۰۷۹۸۸	۰/۶۹۳۱۳۳۷	۰/۰۰۰۱۲۶۷۷۷۲	۰/۶۹۰۶۴۸۹ – ۰/۶۹۵۶۱۸۵
۱	۱۰۰	۰/۷۰۴۴۱۸۸	۰/۷۱۶۹۸۵۳	۰/۰۰۰۱۱۷۴۰۴۵	۰/۷۱۴۶۸۴۲ – ۰/۷۱۹۲۸۶۴
۲	۱۰۰	۰/۷۷۹۶۶۵۱	۰/۷۸۸۹۰۳۶	۰/۰۰۰۱۰۱۰۱۸۷	۰/۷۸۶۹۲۳۶ – ۰/۷۹۰۸۸۳۵
۳	۱۰۰	۰/۶۷۰۷۹۸۸	۰/۶۸۳۰۱۵۶	۰/۰۰۰۱۲۶۴۴۳۷	۰/۶۸۰۵۳۷۴ – ۰/۶۸۵۴۹۳۹

با سرعت کمی بهبود پیدا می‌کنند. همانند برآورد بیزی برای مدل فشار-مقاومت نمایی، اختلاف بین برآوردهای حاصل از توزیع‌های پیشین آگاهی‌بخش و کم آگاهی‌بخش اندک است.

فصل ۴

برآورد قابلیت اطمینان سیستم تحت توزیع نمایی دو پارامتری

۱.۴ مدل فشار - مقاومت نمایی دو پارامتری

در این فصل، قابلیت اطمینان را برای سیستم‌های موازی هنگامی که مولفه‌های فشار - مقاومت، دارای طول عمری با توزیع نمایی دو پارامتری هستند، به دو روش درست‌نمایی ماکزیم و بیزی به دست می‌آوریم. این مبحث، جدید بوده و بوسیله نویسنده پایان‌نامه انجام شده است.

۱.۱.۴ قابلیت اطمینان سیستم‌های موازی

در این بخش، برآورد قابلیت اطمینان یک سیستم موازی، با k مولفه $X_1, X_2, \dots, X_k$ ، که با فشار مشترک تصادفی X_{k+1} روبرو است ارائه می‌شود. فرض کنید $(i = 1, 2, \dots, k+1)$ X_i ها دارای توزیع مستقل نمایی دو پارامتری با پارامترهای μ و λ هستند. بدین ترتیب تابع چگالی احتمال متغیرهای تصادفی مدل به صورت زیر بیان می‌شود

$$f_i(x_i) = \frac{1}{\lambda_i} e^{-\frac{x_i - \mu_i}{\lambda_i}}, \quad x_i > \mu_i > 0, \lambda_i > 0.$$

در ادامه روند برآورد قابلیت اطمینان، ابتدا تابع بقای سیستم را تحت لم زیر به دست می‌آوریم.

لم ۱.۱.۴. فرض کنید $X_i; (i = 1, 2, \dots, k)$ ها (متغیرهای طول عمر) دارای توزیع نمایی دو پارامتری هستند. آنگاه به ازای $y > \mu^*$ که $\mu^* = \max(\mu_1, \mu_2, \mu_3, \dots, \mu_k)$ ، توزیع $Y = \max(X_1, X_2, X_3, \dots, X_k)$ برابر عبارت زیر است

$$\begin{aligned} G(y) &= P(Y \leq y) \\ &= P(X_1 \leq y, X_2 \leq y, X_3 \leq y, \dots, X_k \leq y) \\ &= (1 - e^{-\frac{y - \mu_1}{\lambda_1}})(1 - e^{-\frac{y - \mu_2}{\lambda_2}})(1 - e^{-\frac{y - \mu_3}{\lambda_3}}) \dots (1 - e^{-\frac{y - \mu_k}{\lambda_k}}) \end{aligned}$$

$$= 1 - \sum_{r=1}^k (-1)^{r-1} \sum_{i_1=1}^k \dots \sum_{i_r=1}^k e^{-\sum_{j=1}^r (\frac{y-\mu_{i_j}}{\lambda_{i_j}})}.$$

برهان. ابتدا به ازای $k=3$ رابطه فوق را اثبات می‌کنیم. بنابراین

$$\begin{aligned} G(y) &= 1 - \sum_{r=1}^3 (-1)^{r-1} \sum_{i_1=1}^3 \sum_{i_2=1}^3 \sum_{i_3=1}^3 e^{-\sum_{j=1}^r (\frac{y-\mu_{i_j}}{\lambda_{i_j}})} \\ &= 1 - (-1)^0 \sum_{i_1=1}^3 e^{-\frac{y-\mu_{i_1}}{\lambda_{i_1}}} - (-1)^1 \sum_{i_1=1}^3 \sum_{i_2=1}^3 e^{-\sum_{j=1}^2 (\frac{y-\mu_{i_j}}{\lambda_{i_j}})} \\ &\quad - (-1)^2 \sum_{i_1=1}^3 \sum_{i_2=1}^3 \sum_{i_3=1}^3 e^{-\sum_{j=1}^3 (\frac{y-\mu_{i_j}}{\lambda_{i_j}})} \\ &= 1 - A - B - C. \end{aligned} \quad (1.4)$$

مقادیر A ، B و C به صورت عبارات زیر هستند

$$A = e^{-\frac{y-\mu_1}{\lambda_1}} + e^{-\frac{y-\mu_2}{\lambda_2}} + e^{-\frac{y-\mu_3}{\lambda_3}}. \quad (2.4)$$

$$B = -e^{-\frac{y-\mu_1}{\lambda_1} - \frac{y-\mu_2}{\lambda_2}} - e^{-\frac{y-\mu_2}{\lambda_2} - \frac{y-\mu_3}{\lambda_3}} - e^{-\frac{y-\mu_3}{\lambda_3} - \frac{y-\mu_1}{\lambda_1}}. \quad (3.4)$$

$$C = e^{-\frac{y-\mu_1}{\lambda_1} - \frac{y-\mu_2}{\lambda_2} - \frac{y-\mu_3}{\lambda_3}}. \quad (4.4)$$

با قرار دادن روابط (۲.۴-۴.۴) در رابطه (۱.۴) خواهیم داشت

$$G(y) = (1 - e^{-\frac{y-\mu_1}{\lambda_1}})(1 - e^{-\frac{y-\mu_2}{\lambda_2}})(1 - e^{-\frac{y-\mu_3}{\lambda_3}}),$$

بنابراین نشان می‌دهیم رابطه به ازای $k=3$ اثبات شده، و به ازای مقدار کلی k اثبات مشابهی انجام می‌شود. $\square$

اکنون قابلیت اطمینان سیستم مورد نظر را براساس تابع بقا فوق به صورت زیر محاسبه می‌کنیم.

$$\begin{aligned} R &= P[X_{k+1} < Y] \\ &= \int_{\mu_{k+1}}^{\infty} \bar{G}(x_{k+1}) dF_{k+1}(x_{k+1}) \\ &= \int_{\mu_{k+1}}^{\infty} \sum_{r=1}^k (-1)^{r-1} \sum_{i_1=1}^k \dots \sum_{i_r=1}^k e^{-\sum_{j=1}^r (\frac{x-\mu_{i_j}}{\lambda_{i_j}})} \frac{1}{\lambda_{k+1}} e^{-\frac{(x-\mu_{k+1})}{\lambda_{k+1}}} dx \\ &= \sum_{r=1}^k (-1)^{r-1} \sum_{i_1=1}^k \dots \sum_{i_r=1}^k \int_{\mu_{k+1}}^{\infty} e^{-\sum_{j=1}^r (\frac{x-\mu_{i_j}}{\lambda_{i_j}})} \frac{1}{\lambda_{k+1}} e^{-\frac{(x-\mu_{k+1})}{\lambda_{k+1}}} dx \\ &= \sum_{r=1}^k (-1)^{r-1} \sum_{i_1=1}^k \dots \sum_{i_r=1}^k \int_{\mu_{k+1}}^{\infty} e^{-\sum_{j=1}^r (\frac{x-\mu_{k+1} + \mu_{k+1} - \mu_{i_j}}{\lambda_{i_j}})} \frac{1}{\lambda_{k+1}} e^{-\frac{(x-\mu_{k+1})}{\lambda_{k+1}}} dx \\ &= \sum_{r=1}^k (-1)^{r-1} \sum_{i_1=1}^k \dots \sum_{i_r=1}^k \frac{1}{\lambda_{k+1}} e^{-\sum_{j=1}^r (\frac{\mu_{k+1} - \mu_{i_j}}{\lambda_{i_j}})} \int_{\mu_{k+1}}^{\infty} e^{-\sum_{j=1}^r (\frac{x-\mu_{k+1}}{\lambda_{i_j}})} e^{-\frac{(x-\mu_{k+1})}{\lambda_{k+1}}} dx \\ &= \sum_{r=1}^k (-1)^{r-1} \sum_{i_1=1}^k \dots \sum_{i_r=1}^k \frac{1}{\lambda_{k+1}} e^{-\sum_{j=1}^r (\frac{\mu_{k+1} - \mu_{i_j}}{\lambda_{i_j}})} \int_{\mu_{k+1}}^{\infty} e^{-(x-\mu_{k+1}) \sum_{j=1}^r \frac{1}{\lambda_{i_j}}} e^{-\frac{(x-\mu_{k+1})}{\lambda_{k+1}}} dx \end{aligned}$$

$$\begin{aligned}
&= \sum_{r=1}^k (-1)^{r-1} \sum_{i_1=1}^k \cdots \sum_{i_r > i_{r-1}} \frac{1}{\lambda_{k+1}} e^{-\sum_{j=1}^r (\frac{\mu_{k+1}-\mu_{i_j}}{\lambda_{i_j}})} \int_{\mu_{k+1}}^{\infty} e^{-(x-\mu_{k+1})(\sum_{j=1}^r \frac{1}{\lambda_{i_j}} + \frac{1}{\lambda_{k+1}})} dx \\
&= \sum_{r=1}^k (-1)^{r-1} \sum_{i_1=1}^k \cdots \sum_{i_r > i_{r-1}} \frac{1}{\lambda_{k+1}} e^{-\sum_{j=1}^r (\frac{\mu_{k+1}-\mu_{i_j}}{\lambda_{i_j}})} \frac{1}{(\sum_{j=1}^r \frac{1}{\lambda_{i_j}} + \frac{1}{\lambda_{k+1}})}, \quad (5.4)
\end{aligned}$$

به طوری که $\mu_{k+1} \leq \mu^*$ ، $\mu_{i_j} \leq \dots \mu_{i_k}$ و $\bar{G}(t) = P(Y > t)$ اگر در رابطه فوق، تساوی $\mu_1 = \dots = \mu_{k+1} = \mu$ برقرار باشد، آنگاه قابلیت اطمینان به صورت زیر باز نویسی می شود

$$R = \sum_{r=1}^k (-1)^{r-1} \sum_{i_1=1}^k \cdots \sum_{i_r > i_{r-1}} \frac{1}{\lambda_{k+1}} \frac{1}{(\sum_{j=1}^r \frac{1}{\lambda_{i_j}} + \frac{1}{\lambda_{k+1}})}. \quad (6.4)$$

۲.۱.۴ برآورد ماکزیمم درستنمایی

تحت مفروضات بخش ۱.۱.۴ تابع ماکزیمم درستنمایی را می توان به صورت زیر نوشت

$$\begin{aligned}
L(\lambda_1, \dots, \lambda_{k+1}, \mu_1, \dots, \mu_{k+1}; \mathbf{X}) &= \prod_{i=1}^{k+1} \prod_{j=1}^n \lambda_i e^{-\lambda_i (x_{i_j} - \mu_i)} I(x_{i_j} | (\mu_i, \infty)) \\
&= e^{-\sum_{i=1}^{k+1} \sum_{j=1}^n \lambda_i (X_{i_j} - \mu_i)} \left(\prod_{i=1}^{k+1} (\lambda_i)^n \right) \prod_{i=1}^{k+1} I(x_{i(1)} | (\mu_i, \infty)),
\end{aligned}$$

با لگاریتم گرفتن از طرفین تساوی داریم

$$l(\lambda_1, \dots, \lambda_{k+1}, \mu_1, \dots, \mu_{k+1}) = - \sum_{i=1}^{k+1} \sum_{j=1}^n \lambda_i (X_{i_j} - \mu_i) + \sum_{i=1}^{k+1} n \ln(\lambda_i) I(x_{i(1)} > \mu_i),$$

در نهایت با مشتق گیری نسبت به پارامترهای مقیاس، معادله درستنمایی به صورت زیر نوشته می شود

$$\frac{\partial}{\partial \lambda_i} l(\lambda_1, \dots, \lambda_{k+1}, \mu_1, \dots, \mu_{k+1}) = -\lambda_i \sum_{j=1}^n (x_{i_j} - \mu_i) + n \frac{1}{\lambda_i} = 0,$$

بنابراین با حل معادله درستنمایی، برآورد ماکزیمم درستنمایی برای پارامتر مقیاس عبارتند از

$$\hat{\lambda}_i = \frac{n}{\sum_{j=1}^n (X_{i_j} - X_{i(1)})}.$$

از سوی دیگر چون تابع درستنمایی بر حسب μ صعودی است، بنابراین این تابع بیشترین مقدار خود را در بیشترین مقدار μ که $X_{(1)}$ باشد به دست می آورد. بنابراین برآورد ماکزیمم درستنمایی برای μ عبارتند از

$$\hat{\mu}_i = X_{i(1)}; i = 1, 2, 3, \dots, k+1.$$

حال با استفاده از خاصیت پایایی برآورد ماکزیمم درستنمایی، می‌توان با قرار دادن برآوردهای فوق در روابط ۵.۴ و ۶.۴، برآورد ماکزیمم درستنمایی قابلیت اطمینان را در دو حالت ذکر شده به‌دست آورد. در این فصل با انجام مطالعه شبیه‌سازی، برآوردهای درستنمایی ماکزیمم و بیزی را برای قابلیت اطمینان سیستم فشار مقاومت موازی دارای توزیع نمایی دو پارامتری، به‌دست می‌آوریم.

۳.۱.۴ برآورد بیزی

برای به‌دست آوردن برآورد بیزی قابلیت اطمینان، دو حالت را در نظر می‌گیریم. در حالت اول فرض می‌کنیم پارامتر مقیاس معلوم و برابر مقدار معلوم $\lambda_i; i = 1, 2, \dots, k+1$ می‌باشد. برای به‌دست آوردن برآورد بیزی پارامتر شکل، ابتدا تابع پیشین زیر را در نظر می‌گیریم

$$\pi(\mu_i) = a_i e^{-a_i \mu_i}, \quad i = 1, 2, \dots, k+1.$$

بنابراین فرم عمومی توزیع پسین با فرض معلوم بودن a_i ها، به‌صورت زیر است

$$\begin{aligned} \pi(\boldsymbol{\mu}|\mathbf{X}) &\propto f(\mathbf{X}|\boldsymbol{\mu}).\pi(\boldsymbol{\mu}) = L(\boldsymbol{\mu}; \mathbf{X}).\pi(\boldsymbol{\mu}) \\ &\propto e^{-\sum_{i=1}^{k+1} \sum_{j=1}^n \lambda_i (x_{ij} - \mu_i)} \prod_{i=1}^{k+1} I(x_{i(1)} > \mu_i) \cdot \prod_{i=1}^{k+1} a_i e^{-a_i \mu_i} \\ &\propto e^{n \sum_{i=1}^{k+1} \lambda_i \mu_i} e^{-\sum_{i=1}^{k+1} a_i \mu_i} \prod_{i=1}^{k+1} I(x_{i(1)} > \mu_i) \\ &\propto e^{-\sum_{i=1}^{k+1} (n\lambda_i - a_i) \mu_i} \prod_{i=1}^{k+1} I(x_{i(1)} > \mu_i) \\ &= \prod_{i=1}^{k+1} e^{-(n\lambda_i - a_i) \mu_i} I(x_{i(1)} > \mu_i) \\ &\propto \prod_{i=1}^{k+1} e^{-(n\lambda_i - a_i) \mu_i} e^{(n\lambda_i - a_i) x_{i(1)}} \\ &\propto \prod_{i=1}^{k+1} (n\lambda_i - a_i) e^{-(n\lambda_i - a_i) \mu_i} e^{(n\lambda_i - a_i) x_{i(1)}} \\ &= \prod_{i=1}^{k+1} (n\lambda_i - a_i) e^{-(n\lambda_i - a_i)(\mu_i - x_{i(1)})}, \end{aligned}$$

با توجه به فرم عمومی توزیع پسین می‌توان گفت که این توزیع پسین دارای توزیع نمایی دو پارامتری با پارامتر مقیاس $n\lambda_i + a_i$ و پارامتر شکل $x_{i(1)}$ است. بنابراین برآورد بیزی μ_i به صورت

$$E(\mu_i|\mathbf{X}) = x_{i(1)} + \frac{1}{n\lambda_i - a_i},$$

می‌باشد. به این ترتیب برآورد بیزی قابلیت اطمینان به صورت زیر محاسبه می‌شود

$$\begin{aligned} \hat{R}_b &= E(R|\mathbf{X}) \\ &= E \left(\sum_{r=1}^k (-1)^{r-1} \sum_{i_1=1}^k \dots \sum_{i_r=1}^k \lambda_{k+1} e^{-\sum_{j=1}^r \lambda_{i_j} (\mu_{k+1} - \mu_{i_j})} \frac{1}{(\sum_{j=1}^r \lambda_{i_j} + \lambda_{k+1})} | \mathbf{X} \right) \end{aligned}$$

$$\begin{aligned}
&= \int_0^\infty \cdots \int_0^\infty \left(\sum_{r=1}^k (-1)^{r-1} \sum_{i_1=1}^k \cdots \sum_{i_r > i_{r-1}} \lambda_{k+1} e^{-\sum_{j=1}^r \lambda_{i_j} (\mu_{k+1} - \mu_{i_j})} \frac{1}{(\sum_{j=1}^r \lambda_{i_j} + \lambda_{k+1})} \right) \\
&\quad \cdot \pi(\mu_1, \dots, \mu_{k+1} | X_1, \dots, X_{k+1}) d\mu_1 \cdots d\mu_{k+1} \\
&= \int_0^\infty \cdots \int_0^\infty \left(\sum_{r=1}^k (-1)^{r-1} \sum_{i_1=1}^k \cdots \sum_{i_r > i_{r-1}} \lambda_{k+1} e^{-\sum_{j=1}^r \lambda_{i_j} (\mu_{k+1} - \mu_{i_j})} \frac{1}{(\sum_{j=1}^r \lambda_{i_j} + \lambda_{k+1})} \right) \\
&\quad (n\lambda_1 - a_1) e^{-(n\lambda_1 - a_1)(\mu_1 - x_{1(1)})} \cdots (n\lambda_{k+1} - a_{k+1}) e^{-(n\lambda_{k+1} - a_{k+1})(\mu_{k+1} - x_{k+1(1)})} \\
&\quad d\mu_1 \cdots d\mu_{k+1},
\end{aligned} \tag{۷.۴}$$

به دلیل پیچیده بودن انتگرال (۷.۴)، همان گونه که در فصل ۳ همین پایان نامه بیان کردیم در این گونه موارد از الگوریتم‌های عددی بر پایه روش‌های MCMC استفاده می‌کنیم. در این مورد، به دلیل مشخص بودن توزیع پسین پارامترها و سهولت در نمونه‌گیری روش نمونه‌گیری گیبز را به کار می‌بریم.

در حالت دوم: فرض می‌کنیم پارامتر شکل ثابت و برابر با مقدار معلوم $\mu_i; i = 1, 2, \dots, k+1$ می‌باشد. در این حالت نیز توزیع پیشین پارامتر مقیاس را به صورت زیر در نظر می‌گیریم

$$\pi(\lambda_i) = b_i e^{-b_i(\lambda_i - \mu_i)}, \quad i = 1, 2, \dots, k+1.$$

آنگاه برای توزیع پسین با فرض معلوم بودن b_i ها، داریم

$$\begin{aligned}
\pi(\lambda_i | \mathbf{X}_i) &\propto \prod_{j=1}^n f(x_{ij}; \lambda_i) \pi(\lambda_i) \\
&\propto \prod_{j=1}^n \lambda_i e^{-\lambda_i(x_{ij} - 1)} \cdot b_i e^{-b_i(\lambda_i - \mu_i)} \\
&\propto \lambda_i^n e^{-\lambda_i \sum_{j=1}^n (x_{ij} - 1)} \cdot e^{-b_i \lambda_i} \\
&\propto \lambda_i^n e^{-\lambda_i(n\bar{x}_i - (n - b_i))}.
\end{aligned}$$

در نتیجه توزیع پسین عبارتست از

$$\lambda_i | \mathbf{X}_i \sim \Gamma(n+1, n\bar{x}_i - (n - b_i)).$$

بنابراین تحت تابع زیان درجه دوم خطا، برآوردگر بیزی بصورت

$$E(\lambda_i | \mathbf{X}_i) = \frac{n+1}{n\bar{x}_i - (n - b_i)},$$

محاسبه می‌شود. برای محاسبه برآورد بیزی قابلیت اطمینان، همانند رابطه (۷.۴) عمل کرده و از روش‌های عددی MCMC استفاده می‌کنیم.

۴.۱.۴ مطالعه شبیه‌سازی

در این فصل نیز با انجام یک مطالعه شبیه‌سازی، برآوردگرهای درست‌نمایی ماکزیم و بیزی را برای توزیع نمایی دو پارامتری ارزیابی می‌کنیم. همانند فصل گذشته و مطالعه شبیه‌سازی برای توزیع وایبل و نمایی، در این فصل نیز یک سیستم سه مولفه‌ای موازی را با مقادیر متفاوت پارامترهای شکل و مقیاس برای

توزیع نمایی دو پارامتری در نظر می‌گیریم. متغیر فشار نیز برای این توزیع، نمایی دو پارامتری در نظر گرفته می‌شود.

همانند فصل قبل، چهار حجم نمونه $n = \{25, 50, 75, 100\}$ و سه حالت متفاوت برای هر حجم نمونه در نظر می‌گیریم.

برآورد درست‌نمایی ماکزیم

برای به‌دست آوردن برآورد درست‌نمایی ماکزیم، مقادیر جداول ۱.۴ و ۲.۴ را به ترتیب بعنوان مقادیر پارامتر شکل و مقیاس توزیع در نظر می‌گیریم.

جدول ۱.۴: پارامترهای شکل توزیع نمایی دو پارامتری

حالت	μ_1	μ_2	μ_3	μ_4
۱	۱	۲	۳	۲
۲	۱/۵	۱	۳	۲
۳	۱	۲	۱/۵	۲

جدول ۲.۴: پارامترهای مقیاس توزیع نمایی دو پارامتری

حالت	λ_1	λ_2	λ_3	λ_4
۱	۰/۱	۰/۲	۰/۳	۰/۲
۲	۰/۱۵	۰/۱	۰/۳	۰/۲
۳	۰/۱	۰/۲	۰/۱۵	۰/۲

با توجه به بخش ۱.۱.۴، ابتدا برآوردهای درست‌نمایی ماکزیم پارامترها رو به‌دست می‌آوریم، آنگاه با قرار دادن مقادیر این برآوردها در رابطه (۵.۴)، برآورد درست‌نمایی قابلیت اطمینان را محاسبه می‌کنیم. نتایج شبیه‌سازی در جدول ذکر شده‌اند. با توجه به نتایج جدول ۳.۴، مشاهده می‌شود که برآوردها به روش درست‌نمایی ماکزیم بسیار دقیق می‌باشند و خطاهای بسیار ناچیزی دارند. همچنین افزایش حجم نمونه تاثیر مثبتی بر بهبود برآوردها دارند و با افزایش حجم نمونه، خطاها کمتر و فواصل اطمینان کوچکتر شده‌اند. با توجه به اینکه مقادیر پارامترهای مقیاس در توزیع نمایی دو پارامتری در این فصل و توزیع نمایی در فصل قبل با هم برابر هستند، اما نتایج و خطاهای تقریباً مشابهی در برآورد قابلیت اطمینان به‌دست آمده است.

برآورد بیزی

مطابق آنچه در بخش ۳.۱.۴ بیان شد، با در نظر گرفتن پارامتر مقیاس برابر مقدار معلوم λ_i ، در این فصل تنها حالت توزیع پیشین کم‌آگاهی‌بخش را در نظر می‌گیریم. اگر نتایج قابل قبول باشند، نتایج

جدول ۳.۴: برآورد درست‌نمایی ماکزیمم قابلیت اطمینان

حالت	n	R	$\hat{R}$	انحراف معیار	۹۵٪ فاصله اطمینان
۱	۲۵	۰/۷۹۳۷۷۶	۰/۷۹۴۳۳۷	۰/۰۰۲۳۸۳۷۹۸	۰/۷۸۹۶۶۱۶ – ۰/۷۹۹۰۰۵۹
۲	۲۵	۰/۸۰۹۹۲۴۲	۰/۸۰۸۹۷۹۶	۰/۰۰۳۱۱۶۰۱۵	۰/۸۰۲۸۷۲۳ – ۰/۸۱۵۰۸۶۹
۳	۲۵	۰/۸۱۹۲۷۰۳	۰/۸۱۸۹۱۱۸	۰/۰۰۳۳۸۴۱۴۱	۰/۸۱۲۲۷۹ – ۰/۸۲۵۵۴۴۶
۱	۵۰	۰/۷۹۳۷۷۶	۰/۷۹۵۰۲۸۵	۰/۰۰۱۵۱۹۴۲۲	۰/۷۹۲۰۵۰۴ – ۰/۷۹۸۰۰۶۵
۲	۵۰	۰/۸۰۹۹۲۴۲	۰/۸۰۹۱۱۶۸	۰/۰۰۱۳۴۱۱۵۷	۰/۸۰۶۴۸۸۲ – ۰/۸۱۱۷۴۵۴
۳	۵۰	۰/۸۱۹۲۷۰۳	۰/۸۱۹۹۸۱	۰/۰۰۱۳۷۰۴۰۱	۰/۸۱۷۲۹۵۱ – ۰/۸۲۲۶۶۷
۱	۷۵	۰/۷۹۳۷۷۶	۰/۷۹۲۱۳۱۹	۰/۰۰۱۱۲۰۵۵۱	۰/۷۸۹۹۳۵۶ – ۰/۷۹۴۳۲۸۱
۲	۷۵	۰/۸۰۹۹۲۴۲	۰/۸۱۰۶۲۴۹	۰/۰۰۰۸۹۷۴۱۶۷	۰/۸۰۸۸۶۶ – ۰/۸۱۲۳۸۳۸
۳	۷۵	۰/۸۱۹۲۷۰۳	۰/۸۲۰۰۱۷۴	۰/۰۰۰۹۰۹۸۴۹۵	۰/۸۱۸۲۳۴۲ – ۰/۸۲۱۸۰۰۷
۱	۱۰۰	۰/۷۹۳۷۷۶	۰/۷۹۳۵۹۰۲	۰/۰۰۰۷۷۴۹۹۸۱	۰/۷۹۲۰۷۱۳ – ۰/۷۹۵۱۰۹۲
۲	۱۰۰	۰/۸۰۹۹۲۴۲	۰/۸۱۰۷۹۱۹	۰/۰۰۰۶۸۰۸۴۱	۰/۸۰۹۴۵۷۴ – ۰/۸۱۲۱۲۶۳
۳	۱۰۰	۰/۸۱۹۲۷۰۳	۰/۸۱۹۶۹۳۷	۰/۰۰۰۶۲۷۶۶۹	۰/۸۱۸۴۶۳۴ – ۰/۸۲۰۹۲۳۹

مشابهی را می‌توان در حالتی که توزیع پیشین آگاهی‌بخش انتخاب می‌شود متصور بود. با در نظر گرفتن توزیع‌های پیشینی که در همسایگی میانگین توزیع متمرکز نیستند و در واقع حاوی اطلاعات زیادی نمی‌باشند، جدول ۴.۴ که شامل مقادیر پارامترهای a_i برای توزیع پیشین $\pi(\mu_i) = a_i e^{-a_i \mu_i}$ است را در نظر می‌گیریم. برآوردهای قابل اطمینان در حالت مذکور در جدول ۵.۴ ذکر شده‌اند.

جدول ۴.۴: پارامترهای توزیع پیشین پارامتر شکل

حالت	a_1	a_2	a_3	a_4
۱	۰/۱	۰/۳	۰/۲	۰/۳
۲	۰/۱۵	۰/۱	۰/۲	۰/۳
۳	۰/۱	۰/۳	۰/۱۵	۰/۳

با وجود توزیع پیشینی با اطلاعات کم اما مقادیر خطاها و برآوردها نشان از برآوردی بسیار دقیق دارند. از طرفی با افزایش حجم نمونه، برآوردها دقیق‌تر شده‌اند. در حالتی که پارامتر شکل مقداری معلوم است، جدول ۶.۴ را مقادیر پارامترهای b_i توزیع پیشین $\pi(\lambda_i) = b_i e^{-b_i(\lambda_i - \mu_i)}$ در نظر می‌گیریم. برآورد بیزی قابلیت اطمینان در جدول ۷.۴ گزارش شده‌اند. همانطور که در جدول ۷.۴ مشاهده می‌شود، برآوردهای بیزی پارامتر مقیاس از دقت خوبی برخوردار هستند اما در مقایسه به برآورد بیزی پارامتر شکل در جدول ۵.۴، این میزان دقت کمتر است. اما به‌طور مشابه با نتایج جدول ۵.۴، با افزایش حجم نمونه، دقت برآوردها افزایش یافته‌اند.

جدول ۵.۴: برآورد بیزی قابلیت اطمینان

حالت	n	R	$\hat{R}$	انحراف معیار	۹۵٪ فاصله اطمینان
۱	۲۵	۰/۷۹۳۷۷۶	۰/۷۹۹۰۴۴۴	۰/۰۰۰۰۲۲۹۵۷۸۱	۰/۷۹۵۵۹۴۴ – ۰/۷۹۹۴۹۴۴
۲	۲۵	۰/۸۰۹۹۲۴۲	۰/۸۱۶۹۹۹۱	۰/۰۰۰۰۱۶۱۴۳۸۹	۰/۸۰۶۶۸۲۷ – ۰/۸۱۷۳۱۵۵
۳	۲۵	۰/۸۱۹۲۷۰۳	۰/۸۲۵۷۹۱۷	۰/۰۰۰۰۱۴۳۱۰۱۹	۰/۸۱۵۵۱۱۲ – ۰/۸۲۶۰۷۲۲
۱	۵۰	۰/۷۹۳۷۷۶	۰/۷۹۶۱۶۴۱	۰/۰۰۰۰۰۳۳۰۶۸۵۵	۰/۷۹۱۰۹۹۳ – ۰/۷۹۶۲۲۸۹
۲	۵۰	۰/۸۰۹۹۲۴۲	۰/۸۱۲۷۴۶۶	۰/۰۰۰۰۰۴۴۶۸۱۸۳	۰/۸۰۲۶۵۹ – ۰/۸۱۲۸۳۴۲
۳	۵۰	۰/۸۱۹۲۷۰۳	۰/۸۲۳۰۶۰۸	۰/۰۰۰۰۰۳۴۵۴۰۸۷	۰/۸۱۲۹۹۳۱ – ۰/۸۲۳۱۲۸۵
۱	۷۵	۰/۷۹۳۷۷۶	۰/۷۹۵۱۱۲۴	۰/۰۰۰۰۰۱۹۰۹۵۳۵	۰/۷۹۰۰۷۵ – ۰/۷۹۵۱۴۹۸
۲	۷۵	۰/۸۰۹۹۲۴۲	۰/۸۱۱۷۰	۰/۰۰۰۰۰۱۲۸۰۶۱۳	۰/۸۰۱۶۷۷۹ – ۰/۸۱۱۷۲۸۱
۳	۷۵	۰/۸۱۹۲۷۰۳	۰/۸۲۱۳۷۹۶	۰/۰۰۰۰۰۲۰۴۰۰۷۲	۰/۸۱۱۳۳۹۶ – ۰/۸۲۱۴۱۹۶
۱	۱۰۰	۰/۷۹۳۷۷۶	۰/۷۹۵۲۵۶	۰/۰۰۰۰۰۰۹۶۴۸۸۰۴	۰/۷۹۰۲۳۷۱ – ۰/۷۹۵۲۷۴۹
۲	۱۰۰	۰/۸۰۹۹۲۴۲	۰/۸۱۰۸۵۸۱	۰/۰۰۰۰۰۱۱۲۳۴۳۹	۰/۸۰۰۸۳۶۱ – ۰/۸۱۰۸۸۰۱
۳	۱۰۰	۰/۸۱۹۲۷۰۳	۰/۸۲۱۰۶۶۷	۰/۰۰۰۰۰۰۹۹۷۵۴۷۹	۰/۸۱۱۰۴۷۱ – ۰/۸۲۱۰۸۶۲

جدول ۶.۴: پارامترهای توزیع پیشین پارامتر مقیاس

حالت	۱	۲	۳	۴
۱	۳۰۰	۴۰۰	۲۵۰	۳۲۰
۲	۳۵۰	۳۰۰	۲۵۰	۳۲۰
۳	۳۰۰	۴۰۰	۳۵۰	۳۲۰

جدول ۷.۴: برآورد بیزی قابلیت اطمینان

حالت	n	R	$\hat{R}$	انحراف معیار	۹۵٪ فاصله اطمینان
۱	۲۵	۰/۷۵۳۷۷۶	۰/۷۵۰۲۵۸۵	۰/۰۰۲۱۶۴۴۳۹	۰/۷۴۶۰۱۶۳ – ۰/۷۵۴۵۰۰۷
۲	۲۵	۰/۷۴۹۹۲۴۲	۰/۷۳۷۸۲۱۷	۰/۰۰۵۶۴۴۷۷۴	۰/۷۲۶۷۵۸۲ – ۰/۷۴۹۹۸۵۳
۳	۲۵	۰/۷۸۹۲۷۰۳	۰/۷۷۵۸۰۹	۰/۰۰۲۱۹۰۲۵۶	۰/۷۷۱۵۱۶۲ – ۰/۷۹۰۱۰۱۹
۱	۵۰	۰/۷۵۳۷۷۶	۰/۷۵۲۰۸۲۱	۰/۰۰۱۸۹۲۶۰۹	۰/۷۴۸۳۷۲۷ – ۰/۷۵۵۷۹۱۶
۲	۵۰	۰/۷۴۹۹۲۴۲	۰/۷۳۶۲۰۹۷	۰/۰۰۵۶۶۴۴۸۹	۰/۷۳۵۱۰۷۵ – ۰/۷۵۷۳۱۱۹
۳	۵۰	۰/۷۸۹۲۷۰۳	۰/۷۷۴۴۶۹۹	۰/۰۰۲۳۶۲۵۶۲	۰/۷۶۹۸۳۹۴ – ۰/۷۸۹۴۰۰۵
۱	۷۵	۰/۷۵۳۷۷۶	۰/۷۵۳۱۷۸۴	۰/۰۰۱۹۸۳۴۱۳	۰/۷۴۹۲۹۱ – ۰/۷۵۷۰۶۵۸
۲	۷۵	۰/۷۴۹۹۲۴۲	۰/۷۴۰۷۸۴۲	۰/۰۰۵۰۱۸۸۰۳	۰/۷۳۰۹۴۷۵ – ۰/۷۵۰۶۲۰۹
۳	۷۵	۰/۷۸۹۲۷۰۳	۰/۷۷۹۹۳۲۵	۰/۰۰۱۷۸۸۹۵۵	۰/۷۷۶۴۲۶۳ – ۰/۷۸۹۹۳۸۸
۱	۱۰۰	۰/۷۵۳۷۷۶	۰/۷۵۸۶۴۶۶	۰/۰۰۱۴۴۳۳۵۴	۰/۷۵۱۸۱۷۷ – ۰/۷۶۱۴۷۵۵
۲	۱۰۰	۰/۷۴۹۹۲۴۲	۰/۷۴۱۶۹۲۲	۰/۰۰۴۸۸۵۱۴۲	۰/۷۳۲۱۱۷۵ – ۰/۷۵۱۲۶۶۹
۳	۱۰۰	۰/۷۸۹۲۷۰۳	۰/۷۸۸۲۳۸۳	۰/۰۰۱۱۴۱۴۵۲	۰/۷۸۶۰۰۱۱ – ۰/۷۹۰۴۷۵۵

نتیجه‌گیری و آینده تحقیق

قابلیت اطمینان سیستم‌ها در انتخاب سیستم و میزان رضایت‌مندی از سیستم‌ها در صورتی که مولفه‌های سیستم دارای طول عمری تحت توزیع‌های مختلف آماری باشند، بسیار با اهمیت است. لذا برآورد قابلیت اطمینان سیستم‌ها امری مهم است که محققان زیادی تاکنون به آن پرداخته‌اند و روش‌های برآورد متعددی را مورد بحث و بررسی قرار داده‌اند. دو روش برآوردیابی ماکزیمم درستنمایی و بیزی که از کاربردیترین روش‌ها در برآورد قابلیت اطمینان سیستم‌ها هستند را در صورتی که مولفه‌ها دارای طول عمری تحت توزیع‌های مختلف آماری باشند، برای سیستم‌های سری و موازی به‌دست آورده، همچنین تابع توزیع برآوردگرهای قابلیت اطمینان را نیز به‌دست می‌آوریم و پس از شبیه‌سازی به نتایج زیر می‌رسیم:

- با فرض این‌که مولفه‌های سیستم‌ها (سری و موازی) دارای طول عمری تحت توزیع‌های مختلف آماری از جمله نمایی، نمایی دوپارامتری، گاما، وایبل و... باشند، می‌توان با استفاده از ماتریس اطلاع فیشرفر و قضیه ۱۰۲۰۲، تابع توزیع را برای برآوردگرهای ماکزیمم درستنمایی قابلیت اطمینان به‌دست آورد، و اثبات کرد که این برآوردگرها دارای توزیع مجانبی نرمال هستند.
- پس از انجام شبیه‌سازی نتایج به‌دست آمده در جداول بیان شده در فصول مختلف به این نتیجه می‌رسیم که برای نمونه‌هایی با حجم زیاد استفاده از برآوردگر ماکزیمم درستنمایی در بحث برآورد به دلیل داشتن خطای کمتر و فواصل اطمینان کوچکتر برای برآوردگر، به برآورد بیزی ترجیح داده می‌شود.
- با مقایسه جداول مربوط به برآوردگرهای بیزی به ازای توزیع‌های پیشین مختلف، (آگاهی بخش و کم آگاهی بخش) به این نتیجه می‌رسیم که نتایج تفاوت چندان ندارند. یعنی انتخاب توابع توزیع پیشین مختلف (از نظر میزان آگاهی بخشی) در برآورد بیزی تفاوت چشمی‌گیری ندارد.
- در ادامه روند برآورد قابلیت اطمینان سیستم‌ها، می‌توان قابلیت اطمینان را تحت توزیع‌های مختلف آماری و به روش‌های گوناگون برآوردیابی، برآورد کرده و بهترین مدل را انتخاب کرد.
- می‌توانیم قابلیت اطمینان سیستم‌های سری و موازی دارای مولفه‌های یکسان، با طول عمری تحت توزیع‌های یکسان برآورد کرده و باهم مقایسه کنیم.

پیوست آ

نمادها و تعریفها

برخی از کمیت‌ها در متن پایان‌نامه مورد استفاده قرار گرفته‌اند که تعریف‌های آن‌ها به صورت زیر می‌باشند

آ.۱ توزیع‌های پارامتری مهم در قابلیت اطمینان.

در این بخش چند توزیع مهم را که در مطالعات طول عمر و تحلیل بقا نقش اساسی ایفاء می‌کنند بیان می‌کنیم.

تعریف آ.۱.۱. توزیع نمایی

توزیع نمایی یکی از مهمترین توزیع‌های آماری است که به دلیل خواص جالب، نقش مهمی در مدل سازی داده‌های قابلیت اطمینان دارد. توزیع نمایی دارای دو نوع تک پارامتری و دو پارامتری است. اگر متغیر تصادفی X دارای تابع چگالی احتمال زیر باشد

$$f_{\lambda}(x) = \lambda e^{-\lambda x} \quad x > 0, \lambda > 0$$

گوییم X دارای توزیع نمایی تک پارامتری با پارامتر λ است و آن را با نماد $X \sim E(\lambda)$ نشان می‌دهیم. اکنون فرض می‌کنیم X دارای توزیع نمایی با پارامتر λ باشد در این صورت میانگین و واریانس توزیع بصورت زیر است:

$$E_{\lambda}(x) = \frac{1}{\lambda} \quad Var_{\lambda}(x) = \frac{1}{\lambda^2}$$

توجه کنید که تابع قابلیت اطمینان این توزیع مساوی است با:

$$\bar{F}(x) = \int_x^{\infty} \lambda e^{-\lambda x} dx = e^{-\lambda x} \quad x > 0, \lambda > 0$$

یکی از خواص مهم توزیع نمایی خاصیتی است به نام فقدان حافظه^۱. متغیر تصادفی پیوسته و نامنفی

^۱Memoryless Property

T را دارای خاصیت فقدان حافظه گوئیم هرگاه به ازای هر t_1 و t_2 نامنفی داشته باشیم:

$$P(T > t_1 + t_2 | T > t_1) = P(T > t_2)$$

تعریف ۲.۱.۲. توزیع گاما: دستگاهی را در نظر بگیرید که از K مولفه مشابه تشکیل شده است، که به صورت دنباله‌ای به هم متصل شده‌اند. هرگاه مولفه شماره ۱ از کار می‌افتد، مولفه شماره ۲ به صورت خودکار شروع به کار می‌کند و ... تا این که مولفه k ام از کار بیفتد. فرض می‌کنیم مولفه‌ها دارای توزیع نمایی با پارامتر λ هستند و این که مولفه‌ها مستقل از یکدیگر کار می‌کنند. آنگاه طول عمر دستگاه T_D ، برابر خواهد شد با مجموع طول عمر همه مولفه‌ها که با $T_1, T_2, \dots, T_n$ نشان می‌دهیم. یعنی

$$T_D = \sum_{i=1}^n T_i$$

با توجه به تابع توزیع طول عمر مولفه‌ها و بنابر روابط بین توزیع‌ها، به راحتی می‌توان نشان داد که توزیع T_D دارای چگالی به فرم زیر می‌باشد:

$$f_{T_D}(t) = \frac{t^{k-1}}{(k-1)!} \lambda^k e^{-kt} \quad t > 0, \lambda > 0, k = 1, 2, \dots \quad (1.7)$$

این توزیع به توزیع ارلانگ معروف است که خود حالت خاصی از توزیع گاما است که در ادامه به معرفی آن می‌پردازیم. تابع توزیع ارلانگ برابر است با:

$$\begin{aligned} F(t) &= \frac{1}{(k-1)!} \lambda^k \int_0^t x^{k-1} e^{-\lambda x} dx \\ &= 1 - e^{-\lambda t} \sum_{j=1}^{k-1} \frac{(\lambda t)^j}{j!} \quad t > 0 \end{aligned}$$

چنانچه در توزیع ارلانگ که پارامتر k مقادیر $1, 2, \dots$ را اختیار می‌کند، بتواند هر مقدار مثبتی بین صفر و بی‌نهایت را بگیرد توزیع را گاما می‌نامند.

اگر متغیر تصادفی X دارای تابع چگالی احتمال زیر باشد

$$f_{\alpha, \lambda}(x) = \frac{\lambda^\alpha}{\Gamma(\alpha)} x^{\alpha-1} e^{-\lambda x} \quad , \quad x > 0 \quad , \quad \alpha > 0 \quad , \quad \lambda > 0$$

گوئیم X دارای توزیع گاما با پارامترهای α و λ است و آن را با نماد $X \sim \Gamma(\alpha, \lambda)$ یا $X \sim G(\alpha, \lambda)$ نشان می‌دهیم. حالت خاص. اگر $\alpha = 1$ باشد توزیع گاما تبدیل به توزیع نمایی تک پارامتری می‌شود. اگر متغیر تصادفی X دارای توزیع گاما با پارامترهای فوق باشد آنگاه میانگین و واریانس آن به صورت زیر است.

$$E_{\alpha, \lambda}(X) = \frac{\alpha}{\lambda} \quad Var_{\alpha, \lambda}(X) = \frac{\alpha}{\lambda^2}$$

اکنون می‌خواهیم دو توزیع جدید از مشتقات توزیع گاما، که در فصل دوم استفاده می‌شود بیان کنیم. توزیع دی گاما و توزیع تری گاما.

توزیع دی گاما: فرض می‌کنیم $\Gamma(x)$ توزیع گاما با پارامترهای مورد نظر باشد. آنگاه توزیع دی گاما به صورت

زیر بیان می‌شود

$$\psi(x) = \frac{d}{dx} \ln \Gamma(x) = \frac{\Gamma'(x)}{\Gamma(x)}$$

توزیع تری گاما:

$$\psi_1(z) = \frac{d^2}{dz^2} \ln \Gamma(z)$$

تعریف ۳.۱.۰. توزیع نرمال: توزیع نرمال با توجه به این‌که در فاصله $(-\infty, \infty)$ مقدار می‌گیرد به ندرت در مباحث قابلیت اطمینان به عنوان یک مدل طول طول عمر به کار می‌رود. (تنها در حالتی می‌توان آن را به عنوان مدل طول عمر به کار برد که میانگین آن خیلی بزرگتر از صفر باشد) اگر متغیر تصادفی X دارای تابع چگالی احتمال زیر باشد

$$f_{\mu, \sigma^2}(x) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left\{-\frac{1}{2\sigma^2}(x - \mu)^2\right\} \quad x \in R, \quad \mu \in R, \quad \sigma > 0$$

گوئیم X دارای توزیع نرمال با پارامترهای μ و σ^2 است و آن را با نماد $X \sim N(\mu, \sigma^2)$ نشان می‌دهیم. میانگین و واریانس توزیع نرمال را به صورت زیر قابل دستیابی است.

$$E_{\mu, \sigma^2}(x) = \mu \quad Var_{\mu, \sigma^2}(x) = \sigma^2$$

تعریف ۴.۱.۰. توزیع وایبل: وایبل یکی دیگر از مهمترین توزیع‌های طول عمر در مطالعات قابلیت اطمینان و تحلیل بقا است. اگر متغیر تصادفی X دارای تابع چگالی احتمال زیر باشد

$$f_{\alpha, \lambda}(x) = \alpha \lambda x^{\alpha-1} e^{-\lambda x^\alpha} \quad x > 0, \quad \alpha > 0, \quad \lambda > 0$$

در این صورت X دارای توزیع وایبل با پارامترهای α و λ است و دارای نماد $X \sim W(\alpha, \lambda)$ می‌باشد. در توزیع وایبل اگر $\alpha = 2$ فرض شود، یعنی $W(2, \lambda)$ ، به توزیع رای لی تبدیل می‌شود. بدلیل عدم نیاز به توزیع رای لی از بیان کلی آن صرف نظر می‌کنیم. همچنین اگر $\alpha = 1$ باشد، توزیع $W(1, \lambda)$ همان توزیع نمایی تک پارامتری با پارامتر λ است. اگر X دارای توزیع وایبل با پارامترهای فوق باشد، آن‌گاه داریم:

$$E_{\alpha, \lambda}(X) = \frac{\Gamma(1 + \frac{1}{\alpha})}{\lambda^{\frac{1}{\alpha}}} \quad Var_{\alpha, \lambda}(X) = \frac{\Gamma(1 + \frac{2}{\alpha}) - \Gamma^2(1 + \frac{1}{\alpha})}{\lambda^{\frac{2}{\alpha}}}$$

تابع نرخ در توزیع وایبل برابر است با

$$h(t) = \frac{f(t)}{\bar{F}(t)} = \lambda^\alpha \alpha t^{\alpha-1}$$

در نتیجه برای $\alpha > 1$ ، $h(t)$ تابعی صعودی از t است و برای $\alpha < 1$ ، $h(t)$ تابعی نزولی از t است و برای $\alpha = 1$ مقداری ثابت خواهد داشت.

از جمله دلایل اهمیت توزیع وایبل در تحلیل قابلیت اطمینان می‌توان به موارد زیر اشاره کرد.
 الف) نرخ خطر توزیع وایبل دارای فرم بسته است $(h(t) = \lambda^\alpha \alpha t^{\alpha-1})$ و همین امر باعث می‌شود تحلیل در مورد داده‌ها با این توزیع راحت‌تر صورت گیرد.
 ب) توزیع وایبل مدلی است که انواع داده‌ها را می‌توان با آن تفسیر کرد. داده‌های با نرخ خطر صعودی $(\alpha > 1)$ ، نرخ خطر نزولی $(\alpha < 1)$ و نرخ خطر ثابت $(\alpha = 1)$.
 ج) تحت شرایطی، ثابت می‌شود که مینم یک مجموعه از متغیرهای تصادفی وایبل، خود دارای توزیع وایبل است.

تعریف ۵.۱.۵. توزیع پاراتو: اگر متغیر تصادفی X دارای تابع چگالی احتمال زیر باشد

$$f_{\alpha,\beta}(x) = \frac{\alpha\beta^\alpha}{x^{\alpha+1}}, \quad x \geq \beta, \quad \beta > 0, \quad \alpha > 0$$

گوئیم X دارای توزیع پاراتو با پارامترهای α و β است و آن را با نماد $X \sim Pa(\alpha, \beta)$ نشان می‌دهیم. اگر متغیر تصادفی X دارای توزیع پاراتو با پارامترهای بیان شده باشد، میانگین و واریانس آن به صورت زیر است.

$$E_{\alpha,\beta}(x) = \frac{\alpha\beta}{\alpha-1}, \quad \alpha > 1 \quad Var_{\alpha,\beta}(X) = \frac{\alpha\beta^2}{(\alpha-1)^2(\alpha-2)}; \quad \alpha > 2$$

۲.۲. توزیع مجانبی.

در بحث مشخص کردن توزیع برآوردگرهای قابلیت اطمینان، یکی از مهمترین و پرکاربردترین توزیع‌ها، توزیع مجانبی است. توزیع مجانبی، توزیع احتمالاتی است که هرگاه اندازه نمونه به بی‌نهایت نزدیک شود داده آماری بدان می‌گراید. در ریاضیات و آمار، توزیع مجانبی توزیع محدود از دنباله‌ای از توزیع‌ها است. یکی از استفاده‌های مهم این توزیع در ارائه تقریب به توابع توزیع تجمعی آماری است. فرض می‌کنیم Z_i دنباله‌ای از متغیرهای تصادفی است. در ساده‌ترین حالت دنباله مورد نظر دارای توزیع مجانبی است، اگر Z_i ها دارای همگرایی در احتمال باشند. یعنی زمانی که $i \rightarrow \infty$ ، دنباله فوق به ۰ میل کند. مفهوم کلی‌تر توزیع مجانبی مطرح شده برای متغیرهای تصادفی اصلاح شده است. یعنی اگر متغیرهای تصادفی مورد نظر را اصلاح کنیم می‌توان گفت تابع توزیع آن‌ها توزیع مجانبی است. متغیرهای تصادفی $\{Z_i\}$ را به صورت زیر اصلاح می‌کنیم.

$$Y_i = \frac{Z_i - a_i}{b_i}$$

در صورتی که دنباله‌های a_i و b_i دارای شرط همگرایی در توزیع باشند، متغیرهای Y_i ، اصلاح شده و غیر تصادفی هستند. بنابراین تابع توزیع آن‌ها مجانبی است، اگر

ط

$$P\left(\frac{Z_n - a_n}{b_n} \leq x\right) \approx F(x) \quad P(Z_n \leq z) \approx F\left(\frac{Z - a_n}{b_n}\right)$$

که یکی از پرکاربردترین توزیع‌ها در بحث مجانبی است بیان می‌کنیم.

تعریف ۱.۲.آ. توزیع مجانبی نرمال سازگار توزیع مجانبی نرمال سازگار اصولاً بر اساس کارایی و سازگاری برآوردها بیان می‌شود، یعنی برآوردگری دارای این توزیع است که هم کارا باشد و هم سازگار. برآوردگر سازگار: گوییم دنباله برآوردهای $\{T_n\}$ برای $\gamma(\theta)$ سازگار است اگر n به سمت بی‌نهایت میل کند، آن‌گاه

$$T_n \longrightarrow \gamma(\theta)$$

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

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

یا برای هر $\epsilon > 0$ ،

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

تعریف ۲.۲.آ. کارایی منظور از کارایی یک برآوردهگر سنجش و درجه دقت برآوردهگر است. کارایی یک برآوردهگر دارای رابطه معکوس با واریانس آن است. یعنی هرچه واریانس برآوردهگر کمتر باشد، برآوردهگر کاراتر است. پس هدف ما پیدا کردن برآوردهگرهای با کمترین واریانس است. یکی از مهمترین مباحث آماری در زمینه کمترین واریانس کران پایین کران-رئو یا نامساوی کران-رئو است. براین اساس برآوردهگر کارا است که واریانسش برابر با کران پایین کران-رئو باشد. ابتدا تعریف داریم از این نامساوی.

تعریف ۳.۲.آ. نامساوی کران-رئو که بعضاً نامساوی اطلاع نیز نامیده می‌شود به دلیل تحقیق مستقل کران^۲ (۱۹۲۶) و رئو^۳ (۱۹۲۵) در مورد واریانس برآوردهگرها بدین نام معروف شده است. قبل از بیان نامساوی ابتدا شرایطی که برای برقراری نامساوی لازم هستند را بیان می‌کنیم:

اگر X یک متغیر تصادفی با توزیعی از خانواده چگالیهای $\{P_\theta(x) : \theta \in \Theta\}$ باشد. فرض می‌کنیم:

۱. Θ یک زیر فاصله باز از عداد حقیقی است، یعنی $\Theta \in R$

۲. مشتق تابع چگالی نسبت به پارامتر، یعنی $\frac{d}{d\theta} P_\theta(X)$ وجود دارد.

۳. جابه‌جایی عملگرهای مشتق و انتگرال مجاز است، یعنی

$$\frac{d}{d\theta} \int P_\theta(x) d\mu(x) = \int \frac{d}{d\theta} P_\theta(x) d\mu(x)$$

^۲ H. Cramer

^۳ C.R. Rao

۴. برای هر $\theta \in \Theta$

$$E_{\theta}\left[\frac{d}{d\theta} \ln P_{\theta}(x)\right]^2 > 0$$

۵. برای هر برآوردگر T پارامتر $\gamma(\theta)$

$$\frac{d}{d\theta} E_{\theta}[T(x)] = \frac{d}{d\theta} \int t(x) P_{\theta}(x) d\mu(x) = \int t(x) \frac{d}{d\theta} P_{\theta}(x) d\mu(x)$$

بنابر شرایط مذکور، نتیجه می‌گیریم که برای هر $\theta \in \Theta$

$$E_{\theta}\left[\frac{d}{d\theta} \ln P_{\theta}(x)\right] = 0$$

$$\frac{d}{d\theta} E_{\theta}(T(x)) = E_{\theta}(T(x) \frac{d}{d\theta} \ln P_{\theta}(x)) = \text{cov}_{\theta}(T(x), \frac{d}{d\theta} \ln P_{\theta}(x))$$

قضیه ۴.۲. فرض کنید $X_1, X_2, \dots, X_n$ دارای توزیع توأم با تابع چگالی احتمال توأم $P_{\theta}(x_1, x_2, \dots, x_n)$ ، $\theta \in \Theta$ باشند. اگر T برآوردگر پارامتر $\gamma(\theta)$ باشد، آنگاه تحت شرایط فوق کران پایین کرامر-رائو.

$$V_{\theta}(T) \geq \frac{[\gamma'(\theta) + b'(\theta)]^2}{E_{\theta}\left[\frac{d}{d\theta} \ln P_{\theta}(X_1, X_2, \dots, X_n)\right]^2} \quad (2.1)$$

در رابطه فوق $b(\theta) = E_{\theta}(T) - \gamma(\theta)$ نماینگر اریبی برآوردگر T در برآورد پارامتر $\gamma(\theta)$ است.

□

برهان. به پارسیان (۱۳۸۵) رجوع شود.

تعریف ۵.۲. ماتریس اطلاع فشر فرض کنید $\{X_{ij}\}$ دنباله‌ای از متغیرهای تصادفی باشد با تابع چگالی احتمال $f_{X_{ij}}(\theta)$ به طوری که $k = 1, 2, \dots, n$ و $j = 1, 2, \dots, n$ در این صورت ماتریس اطلاع فشر یک ماتریس nk بعدی است.

$$\begin{bmatrix} I_{11} & I_{12} & \cdots & I_{1n} \\ I_{21} & I_{22} & \cdots & I_{2n} \\ \vdots & & & \end{bmatrix}$$

در ماتریس فوق

$$I_{ii} = \text{Var}(X_{ii}) \quad I_{ij} = \text{cor}(X_i, X_j)$$

تعریف ۶.۲. روش عددی نیوتون-رافسون^۴

قسمت‌هایی از مسایل استنباط آماری به پیدا کردن ریشه یک معادله، می‌نیم‌سازی یک تابع یا یافتن برآوردهای ماکسیمم درست‌نمایی منتهی می‌شوند. این‌گونه مسایل، حل یک معادله یا یک دستگاه

^۴Newton-Raphson

معادلات را شامل می‌شوند که حل آن‌ها به یک فرم تحلیلی منتهی نمی‌شوند. در چنین مواردی، باید جواب‌های تقریبی برای آن‌ها به دست آورد. یک رهیافت برای حصول جواب‌های تقریبی، حل عددی معادله‌ها با روش‌های عددی است. یکی از این روش‌های عددی، روش نیوتون-رافسون است. در این روش با داشتن مقدار اولیه x_0 ، خطی بر تابع f در این نقطه مماس می‌کنیم. خط مماس f در x_0 عبارتست از:

$$f(x_0) + f'(x_0)(x - x_0).$$

نقطه صفر این خط مماس را برابر x_1 قرار می‌دهیم. نقطه صفر (ریشه) این خط، x_1 ، در معادله زیر صدق می‌کند:

$$f(x_0) + f'(x_0)(x_1 - x_0) = 0.$$

اکنون می‌توان جای x_0 را با x_1 عوض کرد و عمل را تکرار کرد. تکرار را تا آن‌جا ادامه می‌دهیم که ریشه به دست آمده به مقدار واقعی نزدیک باشد. این عمل منتهی به رابطه زیر می‌شود:

$$x_{i+1} = x_i - \frac{f(x_i)}{f'(x_i)}.$$

به طور مشابه برای ابعاد بالاتر از یک، خواهیم داشت

$$x_{i+1} = x_i - f'(x_i)^{-1} f(x_i),$$

که در آن $f'(x)$ ماتریس مشتق‌های جزئی است.

تعریف ۷.۲.۰. قضیه حد مرکزی چندمتغیره

فرض کنید $\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_n$ بردارهای تصادفی مستقل و هم‌توزیع در R^k با بردار میانگین $\mu = E(\mathbf{X}_i)$ و ماتریس کوواریانس Σ باشند. این قضیه بیان می‌کند که اگر مجموع این بردارها را استاندارد کنیم، آن‌گاه آماره به دست آمده به توزیع نرمال چندمتغیره همگراست. اگر تعریف کنیم

$$\begin{pmatrix} X_{1(1)} \\ \vdots \\ X_{1(k)} \end{pmatrix} + \begin{pmatrix} X_{2(1)} \\ \vdots \\ X_{2(k)} \end{pmatrix} + \dots + \begin{pmatrix} X_{n(1)} \\ \vdots \\ X_{n(k)} \end{pmatrix} = \begin{pmatrix} \sum_{i=1}^n X_{i(1)} \\ \vdots \\ \sum_{i=1}^n X_{i(k)} \end{pmatrix} = \sum_{i=1}^n \mathbf{X}_i,$$

آن‌گاه می‌توان بردار میانگین را به صورت زیر تعریف کرد:

$$\frac{1}{n} \sum_{i=1}^n \mathbf{X}_i = \frac{1}{n} \begin{pmatrix} \sum_{i=1}^n X_{i(1)} \\ \vdots \\ \sum_{i=1}^n X_{i(k)} \end{pmatrix} = \begin{pmatrix} \bar{X}_{i(1)} \\ \vdots \\ \bar{X}_{i(k)} \end{pmatrix} = \bar{\mathbf{X}}_n.$$

بنابراین داریم

$$\frac{1}{\sqrt{n}} \sum_{i=1}^n [\mathbf{X}_i - E(\mathbf{X}_i)] = \frac{1}{\sqrt{n}} \sum_{i=1}^n [\mathbf{X}_i - \mu] = \sqrt{n}(\bar{\mathbf{X}}_n - \mu).$$

قضیه حد مرکزی چندمتغیره بیان می‌کند

$$\sqrt{n}(\bar{\mathbf{X}}_n - \mu) \xrightarrow{D} N(0, \Sigma).$$

به طوری که

$$\Sigma = \begin{pmatrix} Var(X_{1(1)}) & Cov(X_{1(1)}, X_{1(2)}) & Cov(X_{1(1)}, X_{1(3)}) & \cdots & Cov(X_{1(1)}, X_{1(k)}) \\ Cov(X_{1(2)}, X_{1(1)}) & Var(X_{1(2)}) & Cov(X_{1(2)}, X_{1(3)}) & \cdots & Cov(X_{1(2)}, X_{1(k)}) \\ Cov(X_{1(3)}, X_{1(1)}) & Cov(X_{1(3)}, X_{1(2)}) & Var(X_{1(3)}) & \cdots & Cov(X_{1(3)}, X_{1(k)}) \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ Cov(X_{1(k)}, X_{1(1)}) & Cov(X_{1(k)}, X_{1(2)}) & Cov(X_{1(k)}, X_{1(3)}) & \cdots & Var(X_{1(k)}) \end{pmatrix},$$

و منظور از $\xrightarrow{D}$ ، همگرایی در توزیع است در ماتریس فوق منظور از $Var(X_{1(1)})$ واریانس نمونه اول است.

پیوست ب

کد برنامه R برای اجرای فصل سوم و چهارم

دستورات و کدهای مربوط به شبیه‌سازی فصول ۳ و ۴.
برآورد درست‌نمایی ماکزیمم توزیع نمایی تک پارامتری فصل ۳
• کد زیر مربوط به جدول ۲.۳ می‌باشد.

```
X=function(n,lambada,a,b){
K=length(lambada)#in hamun k+1 tu maghale hast
k=K-1
x=matrix(0,n,K)
  for (ii in 1:K){
    x[,ii]=rexp(n,lambada[ii])
  }
aa=apply(x,2,sum)
lambda_hat=n/aa
lambdahat_last=lambda_hat[K]
lambda_last=lambda[K]
r=1
t=c();w=c();
while (r<=k){
comp=combn(1:k,r)
col=ncol(comp)
  v=c();q=c();
  for (j in 1:col){
    u=comp[,j]
    lam=lambda[u]
```

```

    lam_hat=lambda_hat[u]
    v[j]=lambda_last/(lambda_last+sum(lam))
    q[j]=lambdahat_last/(lambdahat_last+sum(lam_hat))
  }
  t[r]=((-1)^(r-1))*sum(v)
  w[r]=((-1)^(r-1))*sum(q)
  r=r+1
}
Rp=sum(t)
Rp_hat=sum(w)
b=list(Rp,Rp_hat,x)
return(b)
}

```

- برآورد بیزی توزیع نمایی تک پارامتری فصل ۳
دستورات مربوط به جدول ۴.۳

```

Bayes=function(n,Nsim,a,b,x){
  K=ncol(x)
  lambda_bayes=matrix(0,Nsim,K)
  for (i in 1:K){
    lambda_bayes[,i]=rgamma(Nsim,shape=n+a[i],scale=1/(b[i]+sum(x[,i])))
  }
  Lam=expand.grid(lambda_bayes[,1],lambda_bayes[,2],
    lambda_bayes[,3],lambda_bayes[,4])
  Lam=as.matrix(Lam)
  row=nrow(Lam)
  estim=numeric(row)
  for (m in 1:row){
    lambda=Lam[m,]
    lambda_last=lambda[K]
    k=K-1;r=1;
    t=c()
    while (r<=k){

```

```
comp=combn(1:k,r)
col=ncol(comp)
v=c()
for (j in 1:col){
  u=comp[,j]
  lam=lambda[u]
  v[j]=lambda_last/(lambda_last+sum(lam))
}
t[r]=((-1)^(r-1))*sum(v)
r=r+1
}
Rp=sum(t)
estim[m]=Rp
}
Estim=mean(estim)
return(Estim)
}
RR=X(n,lambda,a,b)
x=RR[[3]]
Nsim=10
bb=Bayes(n,Nsim,a,b,x)
```

- برای نمونه، با استفاده از دستورات زیر می‌توان انحراف از معیار برآوردهای درست‌نمایی قابلیت اطمینان را، که در جداول ۲.۳ از آن استفاده شده، به دست آورد. به عنوان مثال برای نمونه به حجم 10^6 کد زیر را می‌نویسیم.

```
iter=60
n=100
tt=numeric(iter)
for (i in 1:iter){
  result=X(n,lambda,a,b)
  tt[i]=result[[2]]
  x=result[[3]]
}
```



```
mm_hat=mean(tt)
sd_hat=var(tt)
U=mm_hat+qnorm(0.975)*sd_hat
L=mm_hat-qnorm(0.975)*sd_hat
```

- در زیر دستورات مربوط به برآوردهای درسنمایی و بیزی قابلیت اطمینان سیستم‌های موازی را در شرایطی که مولفه‌ها تحت توزیع وایبل باشند، بیان می‌کنیم.
برآورد ماکزیمم درسنمایی توزیع وایبل فصل ۳
دستور زیر مربوط به جدول ۸.۳ می‌باشد.

```
library(miscTools)
library(maxLik)
y=function(n,a,b){
  K=length(b)
  k=K-1
  x=matrix(0,n,K)
  for (i in 1:K){
    x[,i]=rweibull(n,a,(b[i])^(-1/a))
  }
  pa=c(a,b)
  Like=function(pa,x){
    a=pa[1]
    a=abs(a)
    b=pa[2:length(pa)]
    b=abs(b)
    g=n*(k+1)*log(a)+n*sum(log(b))+(a-1)*(sum(apply(log(x),2,sum)))
      -(sum(b*apply(x^a,2,sum)))
    g
  }
  m=maxLik(Like,start=c(a,b),x=x,method="BFGS")
  b_last=b[K]
  b_hat=abs(coef(m))[-1]
  bhat_last=b_hat[K]
  r=1
  t=c();w=c()
  while (r<=k){
```

```

comp=combn(1:k,r)
col=ncol(comp)
v=c();q=c()
for (j in 1:col){
  u=comp[,j]
  B=b[u]
  B_HAT=b_hat[u]
  v[j]=b_last/(b_last+sum(B))
  q[j]=bhat_last/(bhat_last+sum(B_HAT))
}
t[r]=((-1)^(r-1))*sum(v)
w[r]=((-1)^(r-1))*sum(q)
r=r+1
}
Rp=sum(t)
Rp_hat=sum(w)
l=list(Rp,Rp_hat,x)
return(l)
}

```

برآورد بیزی توزیع وایبل فصل ۳

- دستورات زیر برای بدست آوردن برآورد بیزی قابلیت اطمینان در حالتی که مولفه‌ها دارای توزیع وایبل هستند به کار می‌روند.

```

library(miscTools)
library(maxLik)
library(MASS)
library(MHadaptive)
func=function(n,Nsim,a,b,lambda,gamma,mu,sigma2){
x=matrix(0,n,K)
  for (i in 1:K){
    x[,i]=rweibull(n,a,(b[i])^(-1/a))
  }
pa=c(a,b)
Like=function(pa,x){

```

```

a=pa[1]
a=abs(a)
b=pa[2:length(pa)]
b=abs(b)
g=n*(k+1)*log(a)+n*sum(log(b))+(a-1)*(sum(apply(log(x),2,sum)))
  -(sum(b*apply(x^a,2,sum)))
g
}
m=maxLik(Like,start=c(a,b),x=x,method="BFGS")
a1=abs(coef(m))[1]
b1=log(abs(coef(m))[2])
b2=log(abs(coef(m))[3])
b3=log(abs(coef(m))[4])
b4=log(abs(coef(m))[5])
prior_reg<-function(pars){
  a<-abs(pars[1])
  beta<-pars[2:length(pars)]
  prior_a<-dgamma(a,shape=lambda,scale=1/gamma,log=TRUE)
  prior_beta1<-dnorm(beta[1],mu[1],sigma2[1],log=TRUE)
  prior_beta2<-dnorm(beta[2],mu[2],sigma2[2],log=TRUE)
  prior_beta3<-dnorm(beta[3],mu[3],sigma2[3],log=TRUE)
  prior_beta4<-dnorm(beta[4],mu[4],sigma2[4],log=TRUE)
  return(prior_a+prior_beta1+prior_beta2+prior_beta3+prior_beta4)
}
loglike=function(pars,data){
x=data
a=abs(pars[1])
beta=pars[2:length(pars)]
g=(n*(k+1))*log(a)+sum(n*beta)+(a-1)*sum(apply(log(x),2,sum))
  -sum(exp(beta)*apply(x^a,2,sum))
prior=prior_reg(pars)
return(g+prior)
}
mcmc_r<-Metro_Hastings(li_func=loglike,pars=c(a1,b1,b2,b3,b4))

```

^)

```
,par_names=c('a','beta1','beta2','beta3','beta4'),iterations=2000,data=x)
mcmc_r<-mcmc_thin(mcmc_r)

Bayes=function(n,Nsim,x){
MH=mcmc_r$trace[1:Nsim,]
Lam=expand.grid(MH[,2],MH[,3],MH[,4],MH[,5])
Lam=as.matrix(Lam)
row=nrow(Lam)
estim=numeric(row)
for (m in 1:row){
  lambda=Lam[m,]
  lambda_last=lambda[K]
  k=K-1;r=1;
  t=c()
  while (r<=k){
    comp=combn(1:k,r)
    col=ncol(comp)
    v=c()
    for (j in 1:col){
      u=comp[,j]
      lam=lambda[u]
      v[j]=lambda_last/(lambda_last+sum(lam))
    }
    t[r]=((-1)^(r-1))*sum(v)
    r=r+1
  }
  Rp=sum(t)#*(max(MH[,1])-min(MH[,1]))
  estim[m]=Rp
}
Estim=mean(estim)
return(Estim)
}
estim=Bayes(n,Nsim,x)
return(estim)
```

}

برآورد درستنمایی توزیع نمایی دو پارامتری فصل ۴

- از دستورات زیر برای به دست آوردن برآورد درستنمایی ماکزیمم قابلیت اطمینان در حالتی که مولفه‌ها دارای توزیع نمایی دو پارامتری هستند استفاده می‌کنیم. نتایج آن در جدول ۳.۴ نشان داده شده است.

```

rexp2 <- function(n,rate,loc) {
  rexp(n,rate)+loc
}
X=function(n,lambd,mu){
K=length(lambd)
k=K-1
x=matrix(0,n,K)
  for (ii in 1:K){
    x[,ii]=rexp2(n,lambd[ii],mu[ii])
  }
mu_hat=apply(x,2,min)
y=matrix(mu_hat,n,K,byrow=T)
z=x-y
lambda_hat=n/(apply(z,2,sum))
lambdahat_last=lambda_hat[K]
lambda_last=lambda[K]
muhat_last=mu_hat[K]
mu_last=mu[K]
r=1
t=c();w=c();
while (r<=k){
comp=combn(1:k,r)
col=ncol(comp)
  v=c();q=c();
  for (j in 1:col){
    u=comp[,j]
    lam=lambda[u]
    muu=mu[u]
    lam_hat=lambda_hat[u]
    muu_hat=mu_hat[u]

```

```

e=numeric(length(u))
e_hat=numeric(length(u))
  for (l in 1:length(u)){
    e[l]=(mu_last-muu[l])*lam[l]
    e_hat[l]=(muhat_last-muu_hat[l])*lam_hat[l]
  }
v[j]=(lambda_last)
      *exp(-sum(e))/(sum(lam)+lambda_last)
q[j]=(lambdahat_last)
      *exp(-sum(e_hat))/(sum(lam_hat)+lambdahat_last)
}
t[r]=((-1)^(r-1))*sum(v)
w[r]=((-1)^(r-1))*sum(q)
r=r+1
}
Rp=sum(t)
Rp_hat_parallel=sum(w)
b=list(Rp,Rp_hat_parallel,x)
return(b)
}

```

برآورد بیزی توزیع نمایی دو پارامتری فصل ۴

- برای محاسبه برآورد بیزی قابلیت اطمینان هنگامی که مولفه‌ها دارای توزیع نمایی دو پارامتری هستند، از دستورات زیر استفاده می‌کنیم. دو حالت زیر را در محاسبه برآورد بیزی داریم. حالت اول، زمانی که پارامتر شکل معلوم است.

```

rexp2 <- function(n,rate,loc) {
  rexp(n,rate)+loc
}
X=function(n,Nsim,mu,lambda,a){
K=length(lambda)
k=K-1
x=matrix(0,n,K)
  for (ii in 1:K){
    x[,ii]=rexp2(n,lambda[ii],mu[ii])
  }
}

```

```

mu_bay=matrix(0,Nsim,K)
  for (i in 1:K){
    mu_bay[,i]=rexp2(Nsim,n*lambda[i]-a[i],min(x[,i]))
  }
Mu=expand.grid(mu_bay[,1],mu_bay[,2],
               ,mu_bay[,3],mu_bay[,4])
Mu=as.matrix(Mu)
row=nrow(Mu)
estim=numeric(row)
for (m in 1:row){
  mu_bayes=Mu[m,]
  mu_last_bayes=mu_bayes[K]
  mu_last=mu[K]
  lambda_last=lambda[K]
  k=K-1;r=1;
  t=c();w=c();
while (r<=k){
  comp=combn(1:k,r)
  col=ncol(comp)
  v=c();q=c();
  for (j in 1:col){
    u=comp[,j]
    lam=lambda[u]
    muu=mu[u]
    muu_bayes=mu_bayes[u]
    e=numeric(length(u))
    e_bayes=numeric(length(u))
    for (l in 1:length(u)){
      e[l]=(mu_last-muu[l])*lam[l]
      e_bayes[l]=(mu_last_bayes-muu_bayes[l])*lam[l]
    }
    v[j]=(lambda_last)
      *exp(-sum(e))/(sum(lam)+lambda_last)
    q[j]=(lambda_last)
  }
  r=r+1;
}
estim[m]=sum(v)/sum(q)
}

```

```

        *exp(-sum(e_bayes))/(sum(lam)+lambda_last)
    }
    t[r]=((-1)^(r-1))*sum(v)
    w[r]=((-1)^(r-1))*sum(q)
    r=r+1
}
Rp=sum(t)
Rp_Bayes=sum(w)
estim[m]=Rp_Bayes
}
Estim=mean(estim)
b=list(Rp,Estim)
return(b)
}

```

- دستور زیر برای محاسبه برآورد بیزی قابلیت اطمینان زمانی که مولفه‌ها دارای توزیع نمایی دو پارامتری و پارامتر مقیاس معلوم است.

```

rexp2 <- function(n,rate,loc) {
    rexp(n,rate)+loc
}
X=function(n,Nsim,mu,lambda,b){
K=length(lambda)
k=K-1
x=matrix(0,n,K)
for (ii in 1:K){
    x[,ii]=rexp2(n,lambda[ii],mu[ii])
}
lambda_bay=matrix(0,Nsim,K)
for (i in 1:K){
    lambda_bay[,i]=rgamma(Nsim,shape=n+1,scale=1/(n*mean(x[,i])-(n-b[i])))
}
LAM=expand.grid(lambda_bay[,1]
    ,lambda_bay[,2],lambda_bay[,3],lambda_bay[,4])
LAM=as.matrix(LAM)
row=nrow(LAM)

```



```

estim=numeric(row)
for (m in 1:row){
  lambda_bayes=LAM[m,]
  lambda_last_bayes=lambda_bayes[K]
  lambda_last=lambda[K]
  mu_last=mu[K]
  k=K-1;r=1;
  t=c();w=c();
while (r<=k){
  comp=combn(1:k,r)
  col=ncol(comp)
  v=c();q=c();
  for (j in 1:col){
    u=comp[,j]
    lam=lambda[u]
    muu=mu[u]
    lamm_bayes=lambda_bayes[u]
    e=numeric(length(u))
    e_bayes=numeric(length(u))
    for (l in 1:length(u)){
      e[l]=(mu_last-muu[l])*lam[l]
      e_bayes[l]=(mu_last-muu[l])*lamm_bayes[l]
    }
    v[j]=(lambda_last)
      *exp(-sum(e))/(sum(lam)+lambda_last)
    q[j]=(lambda_last_bayes)
      *exp(-sum(e_bayes))/(sum(lamm_bayes)+lambda_last_bayes)
  }
  t[r]=((-1)^(r-1))*sum(v)
  w[r]=((-1)^(r-1))*sum(q)
  r=r+1
}
Rp=sum(t)
Rp_Bayes=sum(w)

```

```

    estim[m]=Rp_Bayes
}
Estim=mean(estim)
b=list(Rp,Estim)
return(b)
}

```


مراجع

- [۱] اسدی،م. (۱۳۹۲). ”آشنایی با نظریه قابلیت اعتماد”. چاپ اول. مرکز نشر دانشگاهی.
- [۲] پارسیان،اح. (۱۳۸۵). ”مبانی آمار ریاضی”. دیبایی نیاب. چاپ دوم. مرکز نشر دانشگاه صنعتی اصفهان.
- [3] Asmussen,Søren. ,Glynn, Peter, W . (2007). *Stochastic Simulation: Algorithms and Analysis*. Stochastic Modelling and Applied Probability. 57. 242-258.
- [4] Bandyopadhyay,S. P. and Forster, R. M. (2011). *Philosophy of Statistics*. New Hampshire. 7. 354-361.
- [5] Bayes, T. (1763). *A letter from the late Reverend Mr. Thomas Bayes, FRS to John Canton, MA and FRS*. Philosophical Transactions. 1683-1775 . 269-271.
- [6] Berger,J. O. and Wolpert, R. L. (1988). The Likelihood Principle (2nd ed.). *Haywood. CA: The Institute of Mathematical Statistics, New York*.
- [7] Beg,M. A. (1980). *On the estimation of $P(Y<X)$ for the two parameter exponential distribution*. Metrika. 27. 29-34.
- [8] Bhattacharya,G. K. and Johnson,R. A. (1974) . *Estimation of Reliability in a Multicomponent Stress-Strength Model* . Journal of American Statistical Association. 69. 966–970.
- [9] Church,J. D. and Harris,B. (1970). *The Estimation of Reliability from Stress-Strength Relationships*. Technometrics. 12. 49–54.
- [10] Cramer,H. (1946). *A Contribution to the Theory of Statistical Estimation*. Skandinavisk Akluarietetidskrift. 29. 85-94
- [11] Draper,N. R. and Guttman, I. (1978). *Bayesian Analysis of Reliability in Multicomponent Stress-Strength Models*. *Communications in Statistics*. Theory and Methods. 7. 441–451.
- [12] Duran,Monica. , Penaand,Alexis and Victor, Salin,H. (2007). *A Bayesian approach to parallel stress-strength models*. Economic Quality Control. 22 . 1. 71–85.
- [13] Ebrahimi,N. (1982). *Estimation of reliability for a series stress-strength system*. Institute of Electrical and Electronics Engineers Transactions on Reliability. 31. 202-205.

- [14] Enis,P. and Geisser,S. (1971). *Estimation of probability that $Y < X$* . Journal of the American Statistical Association 66. 333 . 162-168
- [15] Fisher,R. A. (1926). The arrangement of field experiments. *In Breakthroughs in Statistics. 82-91. New York.*
- [16] Gauss, C. F. *Theory of the Combination of Observations which leads to the smallest Errors*. Gauss Werke 4 (1821). 1-93.
- [17] Gelfand,A. E. and Smith,A. F. M. (1990). *Sampling Based Approaches to Calculating Marginal Densities*. Journal of the American Statistical Association. 85. 398–409.
- [18] Gelman,A. and Rubin,D. B. (1992). *Inference from Iterative Simulation using Multiple Sequences* . Statistic Science. 7. 457–511.
- [19] Geman. S. and Geman. D. (1984). *Stochastic Relaxation, Gibbs Distributions, and the Bayesian Restoration of Images*. Institute of Electrical and Electronics Engineers Transactions on Pattern Analysis and Machine Intelligence. 6. 721–741.
- [20] Green,P. J. (1995). *Reversible-jump Markov chain Monte Carlo computation and Bayesian model determination*. Biometrika. 82(4). 711-732.
- [21] Hanagal,D. D. (1996). *Estimation of Reliability from Stress-Strength Relationship. Communication in Statistics. Theory and Methods*. 25. 1783–1797
- [22] Hanagal. D. D. (1996a). *A multivariate Pareto distribution*. Communications in Statistics. Theory and Methods. 25(7). 1471-88.
- [23] Hanagal,D. D. (1996b). *Estimation of system reliability from stress-strength relationship*. Communications in Statistics, Theory and Methods. 25(8). 1783-17 97.
- [24] Hanagal,D. D. (1998). *On the Estimation of Reliability in Stress-Strength Models*. Economic Quality Control. 13. 17-22.
- [25] Hanagal,D. D. (2003). *Estimation of system reliability in multicomponent series stressstrength models*. Journal of the Indian Statistical Association. 41. 1–7.
- [26] Ibraim, J. G. , Chen, M. and Sinha, D. (2001). *Bayesian Survival Analysis*. . Springer,New York.
- [27] Johnson,R. A. (1988). *Stress-Strength Models for Reliability. Handbook of Statistics. Eds Krishnaiah. P. R. and Rao. C. R.* 7. 27-54.
- [28] Johnson,N. L. and Kotz,S. (1970). *Continuous Univariate Distributions. John Wiley and Sons,Inc. New York*
- [29] Kelley,G. D. , Kelley,J. A. and Schucany,W. R. (1976). *Efficient estimation of $PY < X$ in the exponential case*. Technometrics. 18(3). 359-360.

- [30] Kotz,S. , Lumelskii. Y. and Pensky,M. (2003). The Stress-Strength Model and its Generalizations:Theory and Applications. *Imperial College Press. London.*
- [31] Kotz,S. , Lumelskii,Y and Pensky,M. (2003) The Stress-Strength Model and its Generalizations . *John Wily and Sons, Inc. London*
- [32] Kunchur,S. H. and Munoli,S. B. (1993). Estimation of Reliability for a multicomponent survival Stress-Strength model based on exponential distributions. *Communications on Statistics - Theory and Methods.* 22. 767–779.
- [33] Murray, I. , Ghahramani,Z. and MacKay, D. (2006). *MCMC for doubly-intractable distributions*. In Proceedings of the 22nd Annual Conference on Uncertainty in Artificial Intelligence (UAI-06). Arlington. Virginia. AUAI Press (2006). No 359-366
- [34] Neyman,J. and Pearson,E. S. (1933). *On the Problem of the Most Efficient Tests of Statistical Hypotheses. Philosophical Transactions of the Royal Society A: Mathematical.* Physical and Engineering Sciences. 231. 694-706.
- [35] Raftery,A. E. and Lewis,S. (1992). *How Many Iterations in the Gibbs Sampler*. In Bayesian Statistics. 4. 763-773. (eds. J. M. Bernardo, J. Berger. A. P. Dawid and A. F. M. Smith). Oxford University Press, Oxford
- [36] Rao,C. R. (1945). *Information and the accarocy attainable in the Estimation of Statistical Parameters*. Bull,Calcutta. Math,Soc. 37. 81-91
- [37] Rao,C. R. (1973). Linear Statistical Inference and Its Applications. *Wiley Eastern Ltd.*
- [38] Raqab,M. Z. (2009) *Statistics and Probability Letters*. Statistics and Probability Letters . 79. 1839-1846.
- [39] Richey,M. (2010). *The Evolution of Markov Chain Monte Carlo Methods*. The American Mathematical Monthly. Mathematical Association of America. 117(5). 383-413.
- [40] Robert,C. and Casella,G. (2004). Monte Carlo Statistical Methods. *2nd edn. Springer. New York.*
- [41] Shewhar,Walter, A. (1994). Economy in producing Quality. *New York.*
- [42] Smith,R. L. (1984). *Efficient Monte Carlo procedures for generating points uniformly distributed over bounded regions*. Operations Research. 32. 1296–1308.
- [43] Sverdrup, E. (1965). *Estimates and test procedures in connection with stochastic models for deaths, recoveries and transfers between different states of health*. Scandinavian Actuarial Journal 1965. 3-4 . 184-211.

-
- [44] Tierney, L. and Kadane, J. (1986). *Accurate approximations for posterior moments and marginal densities*. *Journals of the American Statistical Association*. 81(393). 82–86.
- [45] Tong, H. (1974). *A Note on the Estimation of $P(Y < X)$ in the Exponential Case*. *Technometrics*. 16. 625,
- [46] Weier, D. R. (1981). *Bayes estimation for a bivariate survival model based on exponential distributions*. *Communications on Statistics-Theory and Methods* 10, 1415–1427.
- [47] Zacks, S. (1971). *The Theory of Statistical Inference*. *Wiley, New York*

Aabstract

Nowadays, reliability and performance of the systems are the most important principles in using systems. Our goal is the estimation of the system reliability in two cases, parallel and series systems, in stress-strength models. Therefore, we estimated the system reliability given the independence of the components (strength) under different distributions (Exponential, Weibull, Gamma, Pareto,...) by maximum likelihood and Bayesian approach in order to achieve better estimates, useful results and to raise working efficiency. The results are compared for some cases. In addition to the estimation of reliability, we derive the estimator distributions. Also, we simulated some of estimators to get optimal estimator. Under some of distributions, the estimator of reliability has complex form. So in such cases for simulating, using the Monte Carlo Markov Chain (MCMC) sampling method is recommended.

Keywords Reliability, Series and parallel systems, Stress-Strength models, Monte Carlo Markov Chain sampling methods.



University of Shahrood

Faculty Of Mathematical Sciences

**A Bayesian approach for the reliability of a
parallel system in stress-strength models**

Mohsen Mehdizadeh

Supervisor

dr.Ahmad Nezakati

Advisor

dr.Hossein Baghishani

September 2015