

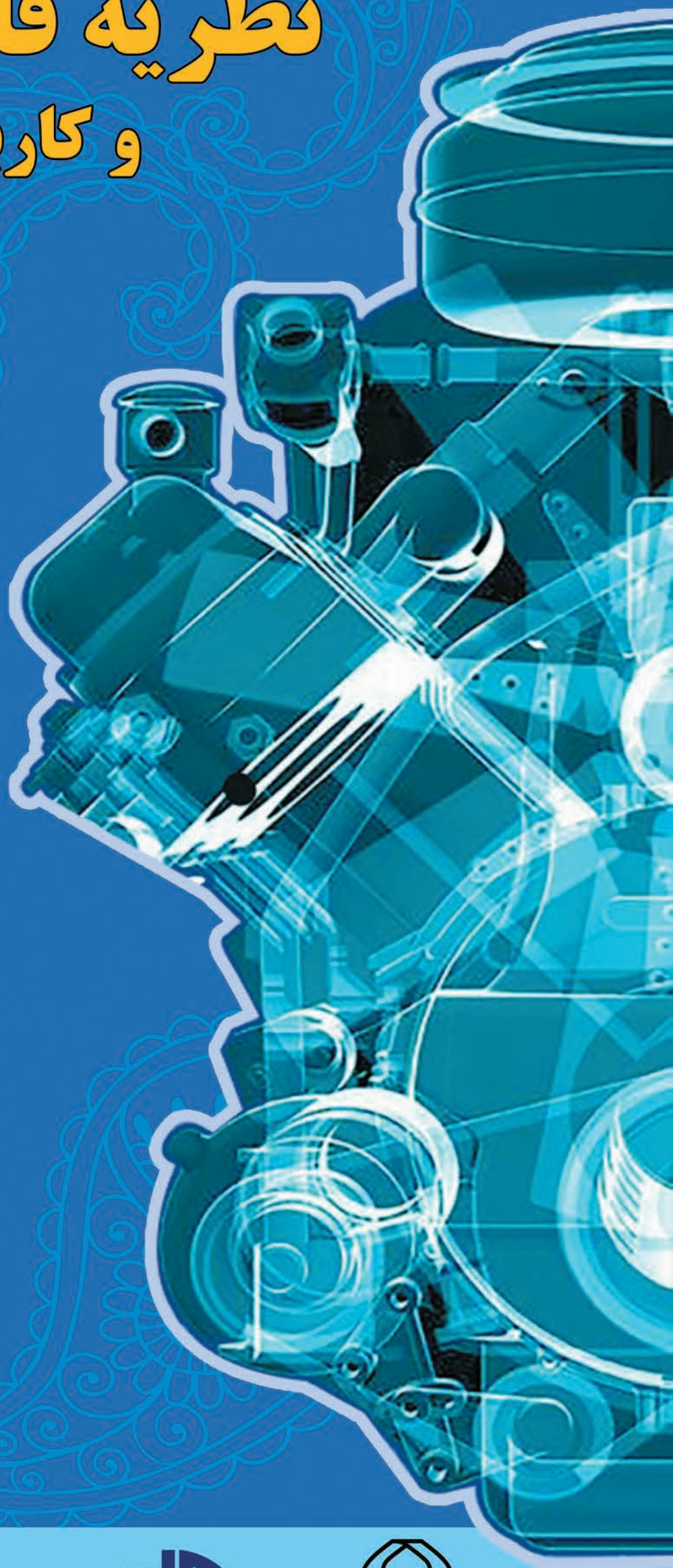
کتابچه راهنما و خلاصه مقالات

پنجمین سمینار تخصصی

نظریه قابلیت اعتماد و کاربردهای آن



۲۸ و ۲۹ فروردین ۱۳۹۸
دانشکده علوم ریاضی
دانشگاه یزد





مجموعه مقالات

پنجمین سمینار تخصصی نظریه قابلیت اعتماد و کاربردهای آن

بخش آمار دانشگاه سوادکوه

یزد، ایران

۲۸ و ۲۹ فروردین ۱۳۹۸

پیشگفتار

شکر و سپاس فراوان خداوند بزرگ را که نعمت خدمت‌گزاری به بخشی از جامعه علمی کشور را نصیب ما نمود تا پنجمین سمینار تخصصی نظریه قابلیت اعتماد و کاربردهای آن را در فروردین ماه ۱۳۹۸ با همکاری قطب علمی داده‌های ترتیبی و فضایی و حمایت مرکز پژوهشی دانشکده علوم ریاضی دانشگاه یزد، انجمن آمار ایران و پایگاه استنادی علوم اسلام در دانشگاه یزد برگزار نماییم.

این سمینار با هدف معرفی پیشرفت‌های جهانی در زمینه قابلیت اعتماد و کاربردهای آن و فراهم نمودن فرصتی در جهت همکاری و تبادل نظر درباره یافته‌های نوین پژوهشگران در این عرصه برگزار می‌گردد.

لازم به ذکر است که از بین مقالات دریافت شده پس از داوری، تعداد ۵۳ مقاله به‌صورت سخنرانی و تعداد ۱۳ مقاله به‌صورت پوستر پذیرفته گردید.

بایسته می‌دانیم از سوی کمیته‌های علمی و اجرایی سمینار از همه مسئولانی که به‌صورت مستقیم یا غیرمستقیم در برگزاری سمینار نقش داشته‌اند، از جمله ریاست محترم دانشگاه یزد، معاونت محترم پژوهشی دانشگاه یزد، همکاران هیات علمی، دانشجویان بخش آمار و کارکنان قسمت‌های مختلف دانشگاه یزد سپاس‌گزاری نماییم.

کلیه امور کامپیوتری دبیرخانه سمینار و همچنین تهیه و تنظیم خلاصه مقالات و راهنمای سمینار با همکاری و پشتکار بی‌دریغ آقایان دکتر رسول روزگار و دکتر علی دولتی، خانم‌ها رحمت‌السادات مشکوتی، خانم مرجان انتظاری و خانم رویا نظری صورت پذیرفت که بدینوسیله از این عزیزان نهایت تشکر و قدردانی را داریم. در پایان می‌بایست از اعضای کمیته علمی و سایر همکارانی که ما را در داوری مقالات یاری نموده‌اند و همچنین قطب علمی داده‌های ترتیبی و فضایی، انجمن آمار ایران و سایر سازمان‌های حمایت‌کننده (مادی و معنوی) و کلیه عزیزانی که ما را در برگزاری این سمینار یاری نموده‌اند، قدردانی و سپاس‌گزاری نماییم.

عیسی محمودی (دبیر)

فروردین ۱۳۹۸



اسامی اعضای کمیته‌های برگزاری، علمی و اجرایی پنجمین سمینار تخصصی نظریه قابلیت اعتماد و کاربردهای آن

اعضای کمیته برگزاری

۱. دکتر محمد صالح اولیاء رئیس دانشگاه یزد
۲. دکتر جعفر احمدی مدیر قطب داده‌های ترتیبی و فضایی
۳. دکتر محمد مهدی لطفی معاون پژوهش و فناوری دانشگاه یزد
۴. دکتر علی مروتی شریف‌آبادی معاون اداری مالی دانشگاه یزد
۵. دکتر قاسم برید لقمانی معاون آموزشی و تحصیلات تکمیلی دانشگاه یزد
۶. دکتر محمدرضا صادقیان شاهی معاون دانشجویی دانشگاه یزد
۷. دکتر حجت اله ذاکرزاده رئیس دانشکده علوم ریاضی دانشگاه یزد
۸. دکتر عیسی محمودی دبیر سمینار

اعضای کمیته اجرایی

۱. دکتر جعفر احمدی دانشگاه فردوسی مشهد
۲. دکتر حمزه ترابی دانشگاه یزد
۳. دکتر هادی جبباری نوقابی دانشگاه فردوسی مشهد
۴. دکتر علی اکبر جعفری دانشگاه یزد
۵. دکتر علی دست‌برآورده دانشگاه یزد
۶. دکتر علی دولتی دانشگاه یزد
۷. دکتر حجت‌اله ذاکرزاده دانشگاه یزد
۸. دکتر رسول روزگار دانشگاه یزد
۹. دکتر محمد صادق زمانی دانشگاه یزد
۱۰. دکتر عیسی محمودی دانشگاه یزد
۱۱. دکتر سید محسن میرحسینی دانشگاه یزد



اعضای کمیته علمی

۱. دکتر جعفر احمدی
 ۲. دکتر مجید اسدی
 ۳. دکتر اکبر اصغرزاده
 ۴. دکتر ابراهیم امینی سرشت
 ۵. دکتر محی الدین ایزدی
 ۶. دکتر حمزه ترابی
 ۷. دکتر مهدی توانگر
 ۸. دکتر علی اکبر جعفری
 ۹. دکتر فیروزه حقیقی
 ۱۰. دکتر بهاء الدین خالدي
 ۱۱. دکتر محمد خنجری صادق
 ۱۲. دکتر مهدی دوست پرست
 ۱۳. دکتر حجت اله ذاکرزاده
 ۱۴. دکتر سمیه زارع زاده
 ۱۵. دکتر محمد مهدی لطفی
 ۱۶. دکتر عیسی محمودی
 ۱۷. دکتر جواد بهبویان
 ۱۸. دکتر ماه بانو تاتا
 ۱۹. دکتر احمد خدادادی
 ۲۰. دکتر اسماعیل خرم
 ۲۱. دکتر علی زینل همدانی
- دانشگاه فردوسی مشهد
- دانشگاه اصفهان
- دانشگاه مازندران
- دانشگاه همدان
- دانشگاه رازی کرمانشاه
- دانشگاه یزد
- دانشگاه اصفهان
- دانشگاه یزد
- دانشگاه تهران
- دانشگاه رازی کرمانشاه
- دانشگاه بیرجند
- دانشگاه فردوسی مشهد
- دانشگاه یزد
- دانشگاه شیراز
- دانشگاه یزد
- دانشگاه یزد
- دانشگاه شیراز
- دانشگاه شهید باهنر کرمان
- دانشگاه شهید بهشتی
- دانشگاه صنعتی امیرکبیر
- دانشگاه صنعتی اصفهان

فهرست مندرجات

- ۱..... ویژگی‌هایی از تبدیل کل زمان آزمون
مجتبی اصفهانی، محمد امینی و غلامرضا محتشمی برزادران
- ۱۰..... برخی ویژگی‌های وابستگی و قابلیت اعتماد یک توزیع لیندلی دومتغیره
شادی جانقربان و سیدمحسن میرحسینی
- ۲۰..... بررسی قابلیت اطمینان پیش‌بینی زمان سفر با مدل GARCH
مریم جمالی و رسول روزگار
- ۲۹..... برآورد پارامترهای توزیع وایبل نمایی شده به روش درست‌نمایی ماکسیمم و بوت استرپ بر اساس نمونه‌گیری مجموعه‌ای رتبه‌دار
سعید حیدری، محمود افشاری، مراد علیزاده
- ۳۷..... بررسی ویژگی‌های قابلیت اعتماد در یک مدل نرخ خطر متناسب گسسته اصلاح شده
محدثه خیاط، رسول روزگار و قباد برمالزن
- ۴۷..... برآورد نیم‌پارامتری معیار قابلیت اعتماد مدل‌های تنش-مقاومت وابسته
علی دست برآورده و فاطمه محمودی‌فرد
- ۵۴..... برآورد نقطه‌ای، بازه‌ای پارامتر تنش-مقاومت در توزیع نمایی دو پارامتری براساس داده‌های رکوردی
صغری رحیمی
- ۶۴..... بررسی عوامل موثر بر امید به زندگی بیماران مسلول با داده‌های ناقص شهر زاهدان، سال‌های ۱۳۹۵-۱۳۹۳
مژگان سالاری و محمدحسین دهقان
- ۷۳..... تحلیل ریسک در فضای متغیر تصادفی هندسی مرکب
زهرا شریعتی و رسول روزگار
- ۸۱..... کاربرد توزیع پارتو نمایی شده در قابلیت اعتماد
سمیه شهرکی ده سوخته
- ۸۸..... مدل‌بندی و نگهداری سیستم‌های تصادفی چندجزئی با به‌کارگیری فرایند تولد-مرگ دوبعدی
فریبا عزیزی، فیروزه حقیقی و الهام جمالی
- ۹۳..... بررسی متوسط زمان تا خرابی در سیستم‌های مختلف با مولفه‌ی آماده به کار
عقیل علایی و نجمه شکیبایی زارع
- ۱۰۴..... مدل تطبیق بیزی موجب انقباضی در برآورد ناپارامتری تابع رگرسیون
الف

حمید کرمی کبیر، روح‌الله حقیقت‌طلب و محمود افشاری

۱۱۳..... معرفی توزیع جدید برگرفته از توزیع گرام چارلی و کاربرد آن در محاسبه معیارهای ریسک
علی‌اصغر لطیفی‌زاده، حجت‌الله ذاکرزاده و حمزه ترابی

۱۲۲..... بررسی زیان ناشی از بکارگیری توابع بقاء مشترک در بیمه‌های عمر و پس‌انداز
رحیم محمودوند

۱۳۰..... کاربردی از اطلاع کولبک- لیبلر تجمعی در پردازش تصاویر
فاطمه منصوری، سعید طهماسبی، فرهاد فرخ‌سرشت و صفیه دانشی

۱۳۹..... توزیع پارتو نمایی و کاربرد آن در قابلیت اعتماد
رویا نصیرزاده و زهرا نیکنام

ویژگی هایی از تبدیل کل زمان آزمون

مجتبی اصفهانی^۱، محمد امینی، غلامرضا محتشمی برزادران

گروه آمار، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد

چکیده: در این مقاله پس از معرفی مفهوم کل زمان آزمون، تبدیل کل زمان آزمون معرفی و ویژگی های آن مورد مطالعه قرار می گیرد. در ادامه نحوه رسم نمودار کل زمان آزمون بیان شده و در پایان ترتیب تصادفی کل زمان آزمون و ارتباط آن با برخی از ترتیب های تصادفی مورد بررسی قرار گرفته، شرایط حفظ این ترتیب تصادفی به وسیله متغیرهای وزنی بیان می شود و کاربردهایی از آن در برخی مفاهیم قابلیت اعتماد ارائه می گردد.

واژه های کلیدی: تبدیل کل زمان آزمون، ترتیب تصادفی TTT ، مشاهده TTT ، ترتیب محدب، $NBUT$.

۱ پیش گفتار

هنگامی که چند واحد به طور همزمان تحت آزمون قرار می گیرند (یا در یک فرآیند واقعی مشاهده می شوند)، شکست واحدهای مختلف در زمان های مختلف رخ می دهد. بنابراین، در هر لحظه از زمان در طی آزمایش، تعدادی از واحد ها هنوز در حال کار کردن هستند و تعدادی از واحد ها شکست خورده اند. آماره ی کل زمان آزمون نشانگر مجموع طول عمر واحدها یعنی واحدهایی که تکمیل شده (واحدهای شکست خورده) یا واحد های ناقص (واحدهای فعال) است. مفهوم کل زمان آزمون برای اولین بار به وسیله بارلو و داکسوم [۳] و بارلو و همکاران [۲] معرفی شد و در ادامه توسط بارلو و کامپو [۵] برای آنالیز داده های با نرخ شکست صعودی و نزولی مورد استفاده قرار گرفت. سپس بارلو و پروشان [۴] کاربردهایی از آن را در آزمون های طول عمر بیان کرد. در ادامه تحقیقات در این زمینه، تبدیل کل زمان آزمون و ترتیب تصادفی کل زمان آزمون توسط محققین معرفی و ویژگی ها و کاربرد های آنها مورد مطالعه قرار گرفت. از کاربردهای تبدیل کل زمان آزمون می توان کاربرد در زمینه نابرابری درآمد، جایگزینی عمر، مفاهیم گداختگی و تجزیه و تحلیل بقاء را نام برد.

نمودار تبدیل کل زمان آزمون یک معیار مناسب برای آنالیز داده های نامنفی است. این نمودار در انتخاب مدل مناسب برای داده ها و بهبود اطلاعات درباره نرخ شکست مفید است. برای مطالعه بیشتر به بارلو و کامپو [۵] مراجعه شود. با توجه به کاربرد فراوان مفهوم TTT در زمینه های مختلف آمار محققان مطالعات زیادی از دیدگاه های مختلف از دهه ۷۰ میلادی تا کنون در این رابطه انجام داده اند که به طور خلاصه به برخی از تحقیقات سال های اخیر اشاره می کنیم.

احمد و همکاران [۱] با معرفی کلاس $NBUT$ برخی از خواص متغیرهای تصادفی را در این کلاس مورد مطالعه قرار دادند. لی و شیکد [۱۱] با معرفی یک خانواده از ترتیب های تصادفی، تعمیم ترتیب کل زمان آزمون ($GTTT$) را بیان و ویژگی ها و ارتباط آن با سایر ترتیب های تصادفی و کاربرد آن در قابلیت اعتماد و بیمه را به دست آوردند. ویژگی ها و شرایط حفظ ترتیب ($GTTT$)

^۱مجتبی اصفهانی: m.esfahani@um.ac.ir

توسط آماره های مرتب و مشخصه سازی آن بر اساس مفاهیم $NBUT$ و $NWUT$ توسط فنگ ژاو [۹] مورد مطالعه قرار گرفته است. کارونی [۸] مفهوم TTT را از دو دیدگاه زمان آزمایش تصادفی و تعداد آزمایش تصادفی مورد بررسی قرار داده است. راوی و پراتیا [۱۵] به بررسی ارتباط بین کلاس $NBUT$ و سایر ترتیب های تصادفی پرداخته است. بیان ویژگیهای جدید از تبدیل TTT و ارتباط این تبدیل و برخی از مفاهیم قابلیت اعتماد توسط نایر و سانکاران [۱۳] مطالعه شده است. شرایط حفظ ترتیب های TTT و EW به وسیله آماره های ترتیبی تعمیم یافته (GOS) توسط بلزونس و ریکلمه [۶] مورد مطالعه قرار گرفته است. وینشکومار و همکاران [۱۸] نیز با در نظر گرفتن تبدیل TTT برخی از ترتیب های تصادفی را با استفاده از چندک ها به دست آورده اند. در ادامه سانکاران و همکاران [۱۴] با استفاده از تبدیل TTT یک آزمون ناپارامتری برای کران های تصادفی ارائه کرده اند. ریزوان و همکاران [۱۶] نیز با استفاده از ویژگی های تبدیل EW برخی از کلاس های سالخوردگی و طول عمر را مشخصه سازی کرده و ویژگی های آنها را به دست آورده اند.

در ادامه در بخش بعد پس از معرفی تبدیل TTT به بیان برخی از ویژگی های آن می پردازیم. در بخش ۳ نحوه رسم نمودار تبدیل TTT را بیان نموده و این نمودار را برای توزیع وایبول و همچنین برای یه سری داده مربوط به شکست مولفه های مکانیکی ارائه می کنیم. در ادامه ترتیب تصادفی TTT در بخش ۴ معرفی و ارتباط آن با سایر ترتیب های تصادفی و همچنین کاربردهایی از آن بیان می شود.

۲ تبدیل TTT

در این قسمت ابتدا تبدیل TTT را به دو روش مختلف معرفی و پس از بیان برخی ویژگی های آن، تبدیل TTT مقیاسی را معرفی می کنیم و ارتباط آنها را با یکدیگر مورد بررسی قرار می دهیم. الف) فرض کنید در یک آزمایش تصادفی n مولفه تحت آزمون هستند و در زمان های مختلف از کار می افتند. اگر زمان های شکست پی در پی مولفه ها را با $X_{1:n} \leq X_{2:n} \leq \dots \leq X_{n:n}$ نشان دهیم و $X_{r:n} < t \leq X_{r+1:n}$ ، که در آن $X_{r:n}$ ها آماره های مرتب از توزیع متغیر تصادفی طول عمر X با تابع توزیع مطلقاً پیوسته F است. آنگاه آماره کل زمان آزمون در بازه (t, ∞) به صورت زیر بیان می شود:

$$\tau(t) = nX_{1:n} + (n-1)(X_{2:n} - X_{1:n}) + \dots + (n-r+1)(X_{r:n} - X_{r-1:n}) + (n-r)(t - X_{r:n})$$

ب) فرض کنید F یک توزیع عمر با $F(\infty) = 1$ و میانگین متناهی μ باشد. در این صورت تبدیل TTT مرتبط با توزیع F را با H_F^{-1} نشان داده و به صورت زیر تعریف می شود.

$$H_F^{-1}(t) = \int_0^{F^{-1}(t)} \bar{F}(s) ds \quad \text{for } 0 \leq t \leq 1 \quad (1-2)$$

$$F^{-1}(t) = \inf\{x : F(x) \geq t\}$$

ج) تبدیل کل زمان آزمون مقیاسی را با $\varphi(u)$ نشان می دهیم و به صورت زیر تعریف می شود:

$$\varphi(u) = \frac{\int_0^{F^{-1}(u)} \bar{F}(s) ds}{\mu} \quad \text{for } 0 \leq u \leq 1 \quad (2-2)$$

$$\text{به عبارت دیگر } \varphi(u) = \frac{H_F^{-1}(u)}{\mu}$$

با استفاده از قسمت الف می توان تبدیل TTT مقیاسی را به صورت زیر نیز بازنویسی کرد:

$$\varphi_{r:n} = \frac{\tau(X_{r:n})}{\tau(X_{n:n})} = \frac{\sum_{j=1}^r (n-j+1)(X_{j:n} - X_{j-1:n})}{\sum_{j=1}^n (n-j+1)(X_{j:n} - X_{j-1:n})}, \quad X_{0:n} = 0$$

با انجام یک تغییر متغیر در رابطه (۱-۲) تبدیل کل زمان آزمون متغیر طول عمر X را می توان به صورت زیر نوشت

$$T(u) = \int_0^u (1-p)q(p)dp \quad (۳-۲)$$

که در آن $q(p)$ تابع چنک متغیر طول عمر X است.

برخی از ویژگی های تبدیل $T(u)$ به اختصار در زیر آمده است.

۱) $T(u)$ یک تابع صعودی است اگر و تنها اگر F پیوسته باشد.

$$۲) \varphi(u) = \frac{T(u)}{\mu}$$

$$۳) Q(u) = \int_0^u \frac{T'(p)}{1-p} dp$$

با توجه به تعریف بالا کل زمان صرف شده برای آزمون از لحظه ۰ تا زمان شکست $X_{1:n}$ برابر $nX_{1:n}$ و بین $X_{1:n}$ و $X_{2:n}$ برابر با $(n-1)(X_{2:n} - X_{1:n})$ است و به همین ترتیب در نهایت زمان آزمون بین $X_{r:n}$ و t برابر $(n-r)(t - X_{r:n})$ است.

بنابراین کل زمان آزمون برابر تا مشاهده r امین شکست برابر

$$\tau(X_{r:n}) = nX_{1:n} + (n-1)(X_{2:n} - X_{1:n}) + \dots + (n-r+1)(X_{r:n} - X_{r-1:n})$$

است.

$$\text{اگر } \bar{X}_n = \frac{\sum_{j=1}^n X_{j:n}}{n} \text{ آنگاه } \varphi_{r:n} = \frac{\tau(X_{r:n})}{n\bar{X}_n}$$

$H_F^{-1}(t)$ نشان دهنده میانگین زمانی است که یک مولفه برای آزمایش شدن تا لحظه از کار افتادن صرف می کند، در صورتی

که آزمایش تا زمان از کار افتادن t درصد از تمام مولفه های تحت آزمایش ادامه داشته باشد.

به راحتی می توان نشان داد که معکوس تبدیل TTT یعنی H_F یک تابع توزیع است. اگر F یک توزیع $IFR(DFR)$ باشد

آنگاه $\frac{d}{dt} H_F^{-1}(t)|_{t=F(x)} = \frac{1}{r(x)}$ صعودی (نزولی) برای هر x است. و در نتیجه H_F^{-1} مقعر (محدب) است، یعنی H_F محدب (مقعر) است.

۳ نمودار TTT

نمودار کل زمان آزمون یک ابزار مفید برای آنالیز داده های نامنفی است. این نمودار در انتخاب مدل مناسب برای داده ها و همچنین بهبود اطلاعات درباره نرخ شکست داده ها بسیار مفید است.

در ادامه نحوه رسم نمودار TTT را بیان می کنیم و این نمودار را برای توزیع وایبول در حالت های مختلف رسم و کاربرد آن را

مورد بررسی قرار می دهیم.

فرض کنید $0 = t(0) \leq t(1) \leq \dots \leq t(n)$ یک نمونه مرتب شده از زمان های شکست n مولفه مستقل غیر قابل تعمیر

با توزیع عمر F و تابع بقای $R(t) = 1 - F(t)$ باشند. نمودار TTT این مشاهدات به روش زیر رسم می شود.

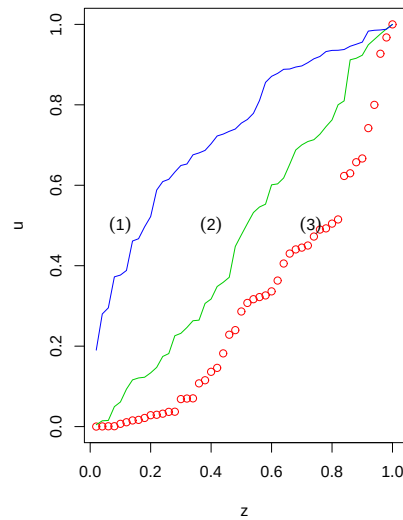
۱- فرض کنید $S_0 = 0$ و $S_j = S_{j-1} + (n-j+1)(t(j) - t(j-1))$ برای $j = 1, 2, \dots, n$ S_j که S_j ها همان

مقادیر TTT هستند.

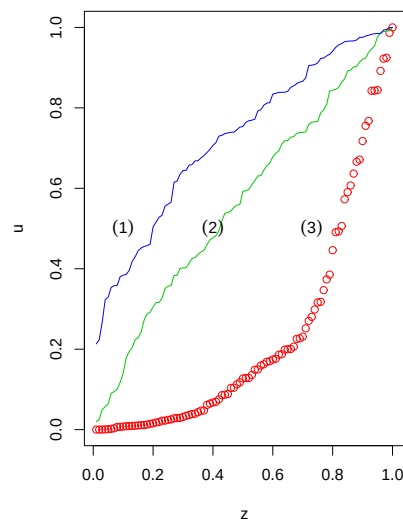
۲- نرمال سازی مقادیر TTT به وسیله رابطه $u_j = \frac{S_j}{S_n}$ برای $j = 0, 1, \dots, n$

۳- رسم نقاط $(\frac{j}{n}, u_j)$ و وصل کردن آنها به یکدیگر.

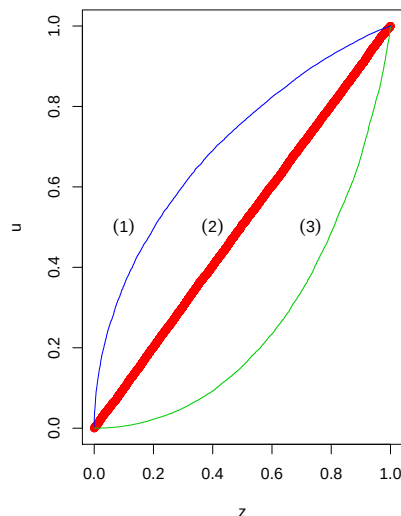
بر این اساس، نمودار TTT شامل منحنی است که از گوشه پایین سمت چپ شروع می شوند و انتهای آن در گوشه بالا سمت راست مربع واحد قرار دارد. هنگامی که حجم نمونه، n ، به بینهایت میل کند نمودار TTT به نمودار تبدیل TTT مقیاسی توزیع عمر F همگرا می شود. شکل های ۱ و ۲ و ۳ نمودار TTT را برای توزیع وایبول نشان می دهد.



شکل ۱: نمودار TTT داده های وایبول با $n = 50$ و $\beta = 5$ و $\alpha = 1$ ، (۱) $\alpha = 2$ ، (۲) $\alpha = 1$ و (۳) $\alpha = 0.5$.



شکل ۲: نمودار TTT داده های وایبول با $n = 100$ و $\beta = 5$ و $\alpha = 2$ ، (۱) $\alpha = 1$ و (۲) $\alpha = 0.5$ و (۳) $\alpha = 0.5$.



شکل ۳: نمودار TTT داده های وایبول با $n = 10000$ و $\beta = 5$ و $\alpha = 2$ ، (۱) $\alpha = 1$ و (۲) $\alpha = 0.5$ و (۳) $\alpha = 0.5$.

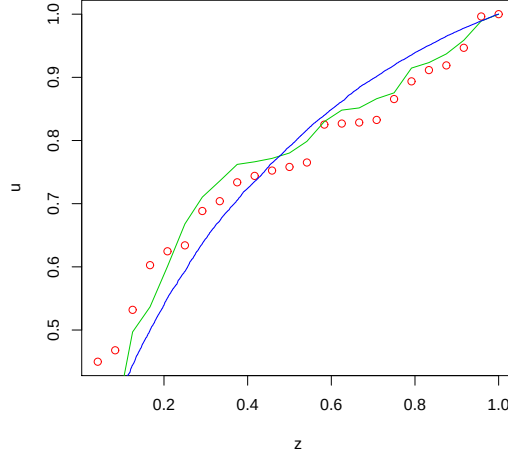
می دانیم که توزیع وایبول دارای نرخ شکست صعودی برای پارامتر شکل کمتر از ۱ و نرخ شکست نزولی برای پارامتر شکل بیشتر از ۱ و نرخ شکست ثابت برای پارامتر شکل ۱ است. از نمودار TTT برای داده های وایبول با پارامتر مقیاس ثابت ۵ مشاهده می شود که هنگامی که داده ها دارای نرخ شکست صعودی باشند نمودار TTT مقعر و هنگامی که نرخ شکست داده ها نزولی باشد نمودار TTT محدب می باشد. همچنین برای داده های با نرخ شکست ثابت نمودار TTT خط نیمساز مربع واحد است. (برای مطالعه بیشتر به [۷] مراجعه شود).

کاربرد برای داده های واقعی

همانطور که مشاهده شد یکی از کاربرد های نمودار TTT تشخیص مدل مناسب برای داده ها و همچنین تعیین نوع نرخ شکست آنها می باشد. در مطالعه داده های مربوط به زمان شکست ۲۴ مولفه مکانیکی از یک سیستم (داده های ۴۰۱ از مورتی و همکاران [۱۲]) پس از تجزیه و تحلیل آنها مدل وایبول با پارامتر شکل ۲۰۲۹ و پارامتر مقیاس ۲۵۰۸ به داده ها برازش داده شده است. با توجه به مقدار پارامتر شکل مدل برازش شده که بزرگتر از یک است انتظار می رود که نرخ شکست داده ها دارای نزولی باشد. نمودار TTT برای این داده ها و همچنین داده هایی از توزیع وایبول با پارامترهای بیان شده در حجم نمونه های ۲۴ و ۱۰۰۰۰ در شکل ۴ نشان داده شده است. همان طور که مشاهده می شود نمودار کل زمان آزمون برای ۲۴ قطعه مکانیکی (نقاط) و ۲۴ داده از توزیع وایبول (نمودار شکسته) و ۱۰۰۰۰ داده وایبول (منحنی) مقعر است و این یعنی داده ها دارای نرخ شکست صعودی IFR بوده که با توجه به ویژگی توزیع وایبول این نتیجه مورد انتظار می باشد.

۴ ترتیب تصادفی TTT

در این بخش پس از معرفی ترتیب تصادفی TTT برخی از ترتیب های تصادفی دیگر را معرفی کرده و ارتباط آنها را با ترتیب TTT به اختصار بیان می کنیم. فرض کنید X و Y دو متغیر تصادفی نا منفی باشند



شکل ۴: نمودار TTT داده های نرخ شکست قطعات مکانیکی

الف) اگر $p \in (0, 1)$, $H_F^{-1}(p) \leq H_G^{-1}(p)$ آنگاه X را کوچکتر از Y در ترتیب TTT گوئیم و با نمادهای $X \leq_{ttt} Y$ یا $G \leq_{ttt} F$ نشان می دهیم.

ب) اگر برای هر تابع h از مجموعه $\{h(u) > 0 \quad \forall u \in (0, 1), h(u) = 0 \quad \forall u \notin (0, 1)\}$

$$\int_{-\infty}^{F^{-1}(p)} h(F(x))dx \leq \int_{-\infty}^{G^{-1}(p)} h(G(x))dx \quad p \in (0, 1) \quad (1-4)$$

آنگاه متغیر تصادفی X را کوچکتر از Y در ترتیب $TTTT$ نامند و با نمادهای $X \leq_{ttt}^{(h)} Y$ یا $F \leq_{ttt}^{(h)} G$ نشان می دهند.

ج) اگر $p \in (0, 1)$, $W_X(p) = \int_{F^{-1}(p)}^{\infty} \bar{F}(x)dx \leq \int_{G^{-1}(p)}^{\infty} \bar{G}(x)dx = W_Y(p)$ آنگاه X را کوچکتر از Y در ترتیب EW گوئیم و با نمادهای $X \leq_{ew} Y$ یا $F \leq_{ew} G$ نشان می دهیم (برای اطلاع بیشتر به کوچار، لی و شیکد [۱۰] مراجعه شود). فرض کنید X و Y دو متغیر تصادفی با توابع توزیع F و G باشند و F^{-1} و G^{-1} به ترتیب توابع معکوس F و G باشند.

الف) اگر F^{-1} و G^{-1} از راست پیوسته باشند و $G^{-1}(p) - F^{-1}(p)$ تابعی صعودی باشد آنگاه X را کوچکتر از Y در ترتیب پراکندگی نامند و با نماد $X \leq_{disp} Y$ نشان می دهیم.

ب) اگر برای هر تابع مقعر صعودی ϕ که $\phi: R \rightarrow R$ داشته باشیم $E[\phi(X)] \leq E[\phi(Y)]$ آنگاه گوئیم X را کوچکتر از Y در ترتیب مقعر صعودی نامند و با نماد های $X \leq_{icv} Y$ یا $F \leq_{icv} G$ نشان می دهیم.

ج) اگر X و Y دو متغیر تصادفی نامنفی با توابع توزیع مطلقا پیوسته F و G باشند که دارای تکیه گاه های $[0, a]$ و $[0, b]$ هستند و a و b ثابتهای منتهای یا نامتهای اند و $G^{-1}(F)$ تابعی محدب باشد آنگاه X را کوچکتر از Y در ترتیب تبدیل محدب نامند و با نماد های $X \leq_c Y$ یا $F \leq_c G$ نشان می دهیم.

د) اگر $\frac{G^{-1}(F(x))}{x}$ نسبت به $x > 0$ صعودی باشد آنگاه X را کوچکتر از Y در ترتیب ستاره^۳ نامیم و با نماد $X \leq_* Y$ نشان می دهیم.

وقتی X یک متغیر تصادفی نامنفی با $\mu < \infty$ باشد مشاهده کل زمان آزمون وقتی X رخ داده را با X_{ttt} نشان می دهیم که به صورت $X_{ttt} = T_X(F(X))$ تعریف می شود و به راحتی ثابت می شود که تابع توزیع X_{ttt} معکوس تبدیل TTT است. متغیر تصادفی X_{ttt} که معرف مشاهده کل زمان آزمون است در مفاهیم IFR و $IFRA$ نیز صدق می کند اگر Y دارای توزیع پارتو باشد یعنی $\bar{F}(y) = \frac{1}{1+y}, y \geq 0$ ، به سادگی می توان نشان داد که X_{ttt} دارای خاصیت IFR است اگر و تنها اگر $X \leq_c Y$ و علاوه بر این اگر $X \leq_* Y$ آنگاه X_{ttt} دارای خاصیت $IFRA$ است. در ادامه ارتباط ترتیب تصادفی TTT با برخی از ترتیب های تصادفی و شرایط حفظ این ترتیب به وسیله متغیر های وزنی بیان می شود.

اگر X و Y دو متغیر تصادفی نامنفی باشند آنگاه:

$$X_{ttt} \leq_{st} Y_{ttt} \iff X \leq_{ttt} Y \quad (\text{الف})$$

$$X \leq_{st} Y \implies X_{ttt} \leq_{st} Y_{ttt} \text{ و } X \leq_{ttt} Y \implies X_{ttt} \leq_{ttt} Y_{ttt} \quad (\text{ب})$$

با توجه به این که X دارای خاصیت $NBUE$ است اگر و تنها اگر $X \leq_{ew} Y$ که در آن Y یک متغیر تصادفی نمایی با میانگین $E(X)$ است. از طرفی اگر $E(X) = E(Y)$ داریم

$$X \leq_{ew} Y \iff X \geq_{ttt} Y$$

کوچار و همکاران [۱۰].

فرض کنید X و Y دو متغیر تصادفی نامنفی باشند. آنگاه

$$X \leq_{icv} Y \implies X_{ttt} \leq_{icv} Y_{ttt} \quad (\text{الف})$$

$$X \leq_{disp} Y \iff X_{ttt} \leq_{disp} Y_{ttt} \iff X_{ew} \leq_{disp} Y_{ew} \quad (\text{ب})$$

شیکد شانتیکومار [۱۷] نشان دادند که متغیر تصادفی نامنفی X دارای خاصیت $HNBUE$ است اگر $X \geq_{icv} Y$ که در آن Y یک متغیر تصادفی نمایی با میانگین $E(X)$ است و $Y_{ttt} = U(0, E(X))$. بنابراین قضیه (۶.۴) نتیجه می دهد که اگر متغیر تصادفی نامنفی X دارای خاصیت $HNBUE$ باشد آنگاه $X_{ttt} \geq_{icv} U(0, E(X))$.

فرض کنید X یک متغیر تصادفی نامنفی با تابع چگالی $f(\cdot)$ و تابع توزیع مطلقا پیوسته $F(\cdot)$ باشد و همچنین $w(\cdot)$ یک تابع نامنفی تعریف شده روی محور اعداد حقیقی باشد و متغیر تصادفی X_w با تابع چگالی $f_w(x) = \frac{w(x)f(x)}{E(w(x))}$ متغیر تصادفی حالت وزنی متغیر X باشد.

فرض کنید X و Y دو متغیر تصادفی نامنفی با توابع توزیع مطلقا پیوسته F و G باشند و $\frac{\int_x^\infty w_1(t)f(t)dt}{\int_x^\infty w_2(t)g(t)dt} \leq \frac{E(w_1(X))}{E(w_2(X))}$ آنگاه

$$X_{w_1} \leq_{ttt} Y_{w_2}.$$

اگر شرایط گزاره قبل برقرار باشد آنگاه:

$$\frac{\int_x^\infty w_1(t)f(t)dt}{\int_x^\infty w_2(t)f(t)dt} \leq \frac{E(w_1(X))}{E(w_2(X))} \implies X_{w_1} \leq_{ttt} X_{w_2} \quad (۱)$$

$$\frac{\int_x^\infty w(t)f(t)dt}{\int_x^\infty w(t)g(t)dt} \leq \frac{E(w(X))}{E(w(Y))} \implies X_w \leq_{ttt} Y_w \quad (۲)$$

$$\frac{\int_x^\infty w(t)f(t)dt}{F(t)} \leq E(w(X)) \implies X \leq_{ttt} X_w \quad (۳)$$

فرض کنید برای هر متغیر تصادفی X تعریف کنیم:

$$X_t = [X - t | X > t], \quad t \in x : F(x) < 1$$

متغیر تصادفی نامنفی X را $NBUT(NWUT)$ گوئیم هرگاه:

$$X_t \leq_{ttt} (\geq_{ttt}) X, \quad t \geq 0$$

فنگ ژيانو [9] نشان داد که اگر X یک متغیر $NBUT$ و یک تابع نامنفی و نزولی روی بازه $[0, 1]$ باشد آنگاه X_{ttt} و $X_{ttt}^{(h)}$ نیز $NBUT$ هستند. همچنین اگر $X_1, X_2, \dots, X_n$ یک نمونه تصادفی مستقل و هم توزیع از توزیع متغیر X باشد و $X_{n:n}$ یک متغیر $NWUT$ باشد آنگاه X نیز $NWUT$ است.

۵ نتیجه گیری

در این مقاله ویژگی های تبدیل کل زمان آزمون و رفتار نمودار کل زمان آزمون برای برخی داده ها مورد مطالعه قرار گرفت. بررسی رفتار این نمودار برای سایر توزیع ها در آینده تحقیق مد نظر می باشد.

مراجع

- [1] Ahmad, I.A., Kayid, M. and Li, X. (2005), The NBUT class of life distributions, *IEEE Transactions on Reliability*, **54**(3), 396-401.
- [2] Barlow, R.E., Bartholomew, D.J., Bremner, J.M, and Brunk, H.D. (1972), *Statistical Inference under Order Restrictions*, John Wiley & Sons, New York.
- [3] Barlow, R.E. and Doksum, K. A. (1972), *Isotonic tests for convex orderings*. In *Proceedings of the Sixth Berkeley Symposium on Mathematical Statistics and Probability*, Volume 1: Theory of Statistics. The Regents of the University of California.
- [4] Barlow, R.E. and Proschan, F. (1975), , *Asymptotic Theory of Total Time on Test Processes, with Applications to Life Testing*(No. ORC-75-18). California Univ Berkeley Operations Research Center.
- [5] Barlow, R.E. and Campo, R. (1975), *Total time on test processes and applications to failure data analysis*, In: *Reliability and Fault Tree Analysis*. SIAM, Philadelphia, PA.
- [6] Belzunce, F. and Martínez-Riquelme, C. (2015), Some results for the comparison of generalized order statistics in the total time on test and excess wealth orders, *Statistical Papers*, **56**(4), 1175-1190.
- [7] Bergman, B. and Klefsjö, B. (1984) The total time on test concept and its use in reliability theory, *Operations Research*, **32**(3), 596-606.
- [8] Caroni, C. (2010), "Failure limited" data and TTT-based trend tests in multiple repairable systems, *Reliability Engineering & System Safety*, **95**(6), 704-706.

- [9] Fengxiao, W. (2009), Some results on the generalized TTT transform order for order statistics, $\square\square\square\square\square$, **25**(1).
- [10] Kochar, S.C., Li, X. and Shaked, M. (2002), The total time on test transform and the excess wealth stochastic orders of distributions, *Advances in Applied Probability*, **34**(4), 826-845.
- [11] Li, X. and Shaked, M. (2007), A general family of univariate stochastic orders, *Journal of statistical planning and inference*, **137**(11), 3601-3610.
- [12] Murthy, D.P., Xie, M. and Jiang, R. (2004), *Weibull models*, Vol. 505, John Wiley & Sons.
- [13] Nair, N.U. and Sankaran, P.G. (2013), Some new applications of the total time on test transforms, *Statistical Methodology*, **10**(1), 93-102.
- [14] Paduthol, S.G., Dewan, I. and Sreedevi, E.P. (2015), A nonparametric test for stochastic dominance using total time on test transform, *American Journal of Mathematical and Management Sciences*, **34**(2), 162-183.
- [15] Ravi, S. (2012), Relationship between certain classes of life distributions and some stochastic orderings, *ProbStat Forum*, **5**, Januar, 1-7.
- [16] Rizwan, U., Ahmed, B.T. and Hussainy, S.T. (2018), *Transformation Properties of Aging Classes under Excess Wealth Transform*.
- [17] Shaked, M. and Shanthikumar, J.G. (2007), Stochastic orders, *Springer Science & Business Media*.
- [18] Vineshkumar, B., Nair, N.U. and Sankaran, P.G. (2015), Stochastic orders using quantile-based reliability functions, *Journal of the Korean Statistical Society*, **44**(2), 221-231.

برخی ویژگی‌های وابستگی و قابلیت اعتماد یک توزیع لیندلی دومتغیره

شادی جانقریان^۱، سیدمحسن میرحسینی
گروه آمار، دانشکده علوم ریاضی، دانشگاه یزد

چکیده: یکی از توزیع‌های مهم در مبحث طول عمر، توزیع لیندلی است. اما این توزیع انعطاف‌پذیری کافی را ندارد، بنابراین ساخت توزیع‌های دومتغیره و چندمتغیره احساس می‌شود. در این مقاله یک توزیع طول عمر لیندلی دومتغیره با وابستگی منفی را ارائه می‌دهیم و به بررسی برخی از ویژگی‌های آن از جمله مفاهیم وابستگی، ضرایب همبستگی و معیارهای قابلیت اعتماد سیستم‌های سری و موازی با توجه به مدل پیشنهادی می‌پردازیم. واژه‌های کلیدی: طول عمر سیستم‌های سری و موازی، ضرایب همبستگی، مفاهیم وابستگی، وابستگی نسبت درست‌نمایی. کد موضوع‌بندی: 39B82، 34K20، 39B52، 47A55.

۱ پیش‌گفتار

از توزیع‌های دومتغیره و چندمتغیره در زمینه‌های مختلف پدیده‌های طبیعی و فیزیکی، همچون هیدرولوژی، علوم محیط زیست، اقتصاد، علوم پزشکی و فیزیولوژی استفاده شده است. می‌دانیم که اکثر توزیع‌های طول عمر باهم ارتباط مثبت دارند، در حالی‌که در عمل متغیرهای تصادفی زوجی نامنفی وجود دارند که به‌طور منفی وابسته اند. همچنین در بین متغیرهای اقلیمی، بارش و دما از اهمیت به‌سزایی برخوردار هستند و تخمین دقیق آن‌ها در مطالعات هواشناسی و کشاورزی از اهمیت زیادی برخوردار است. به عنوان مثال، در سال ۱۹۹۷ کورته و همکاران نشان دادند که شدت بارندگی و طول مدت بارش وابستگی منفی ایجاد می‌کنند که برای مطالعه توزیع فراوانی سیلاب از آن استفاده می‌گردد. در سال ۲۰۰۹ هانگ و همکاران در مطالعات خود به وجود همبستگی منفی بین بارش و دما در حوضه رودخانه زرد چین اشاره کردند.

یکی از توزیع‌های مهم و پرکاربرد در زمینه‌ی طول عمر، توزیع لیندلی است که در سال ۱۹۵۸ توسط لیندلی [۵] ارائه شد. قیطانی و همکاران [۲] در سال ۲۰۰۸ خواص آماری این توزیع را مورد بررسی قرار دادند و نشان دادند که این توزیع عملکرد مناسبی در تحلیل داده‌های طول عمر از خود نشان می‌دهد. از سال ۲۰۰۸ به بعد این توزیع مورد توجه محققان قرار گرفت و ما را نیز بر آن داشت که توزیع دومتغیره لیندلی را پیشنهاد دهیم. توزیع پیشنهادی وابستگی منفی دارد و برای تحلیل داده‌های طول عمر مناسب است. در این مقاله، ضمن معرفی توزیع لیندلی دومتغیره با وابستگی منفی، توزیع‌های حاشیه‌ای و شرطی، گشتاورها و ضرایب همبستگی، مفاهیم وابستگی، جداولی از ضرایب وابستگی و تابع وابستگی مکانی را بیان کرده و در ادامه طول عمر سیستم‌های سری و موازی دو مؤلفه‌ای با مؤلفه‌های وابسته برای توزیع لیندلی دومتغیره را مورد بررسی قرار داده و گشتاورهای ML پارامترها را نیز به‌دست می‌آوریم.

۲ توزیع لیندلی دومتغیره

روش‌های مختلفی برای ساخت توزیع‌های طول عمر وجود دارد. فرض کنید Z_1, Z_2, Z_3 متغیرهای تصادفی نامنفی مستقل باشند. می‌توان متغیرهای وابسته را با استفاده از $X = \text{Min}(Z_1, Z_2)$ و $Y = \text{Min}(Z_2, Z_3)$ یا $X = \text{Max}(Z_1, Z_3)$ و $Y = \text{Max}(Z_2, Z_3)$ تولید کرد. به نظر می‌رسد که امکان ایجاد وابستگی منفی با استفاده از این روش وجود ندارد بنابراین با استفاده از روش زیر سعی بر حل مشکلات فوق داریم. ابتدا به معرفی توزیع لیندلی می‌پردازیم. توزیع لیندلی در سال ۱۹۵۸ [۵] توسط لیندلی با تابع چگالی زیر ارائه شد:

$$f(x) = \frac{\theta^2}{\theta + 1} (1 + x) e^{-\theta x}, \quad x > 0, \theta > 0.$$

در ادامه، فرض کنید Z_1 و Z_2 متغیرهای تصادفی مستقل از توزیع لیندلی به ترتیب با میانگین‌های $\theta + 2/\theta(\theta + 1)$ و $2/\lambda(\lambda + 1)$ هستند. متغیر تصادفی دوتایی (X, Y) را به صورت زیر تعریف می‌کنیم:

$$X = Z_1, \quad Y = e^{Z_2 - Z_1}.$$

لازم به ذکر است که ایده، برای ساخت این توزیع لیندلی دومتغیره، برگرفته از بوهین و همکاران [۱] است. با استفاده از روش توزیع توابعی از متغیرهای تصادفی، می‌توان تابع چگالی احتمال پیوسته (X, Y) را به صورت زیر نوشت:

$$f(x, y) = \begin{cases} \frac{(\theta\lambda)^2}{(\theta+1)(\lambda+1)} (1+x)(1+x+\ln y) e^{-x(\theta+\lambda)} y^{-(1+\lambda)} & 0 < x < \infty, \quad e^{-x} < y < \infty \\ 0 & \text{سایر جاها.} \end{cases} \quad (1)$$

۱.۲ توزیع‌های حاشیه‌ای و شرطی

همان‌طور که ملاحظه می‌شود تابع چگالی حاشیه‌ای X ، با توجه به تعریف، توزیع لیندلی با میانگین $\theta + 2/\theta(\theta + 1)$ است. پس تابع چگالی شرطی $f_{Y|X}(y|x)$ از رابطه‌ی زیر پیروی می‌کند:

$$f_{Y|X}(y|x) = \begin{cases} \frac{\lambda^2}{(\lambda+1)} (1+x+\ln y) e^{-x\lambda} y^{-(1+\lambda)}, & y > e^{-x}, \\ 0 & \text{سایر جاها.} \end{cases}$$

حال تابع چگالی حاشیه‌ای Y با استفاده از (۱۴) و انجام محاسبات ساده به فرم زیر می‌شود:

$$f_Y(y) = \begin{cases} \int_{-\ln y}^{\infty} f(x, y) dx = \frac{(\theta\lambda)^2 y^{\theta-1}}{(\theta+1)(\lambda+1)(\theta+\lambda)} (1 - \ln y + \frac{2 \ln y + 2}{\theta + \lambda} \frac{2}{(\theta + \lambda)^2}), & 0 < y < 1, \\ \int_0^{\infty} f(x, y) dx = \frac{(\theta\lambda)^2}{(\theta+1)(\lambda+1)(\theta+\lambda)} (1 + \ln y + \frac{\ln y + 2}{\theta + \lambda} \frac{2}{(\theta + \lambda)^2}), & y > 1. \end{cases}$$

تابع چگالی شرطی $f_{X|Y}(x|y)$ عبارت است از:

$$f_{X|Y}(x|y) = \begin{cases} \frac{[1 + \ln y + (2 + \ln y)x + x^2]}{[1 - \ln y + \frac{2 \ln y + 2}{\theta + \lambda} + \frac{2}{(\theta + \lambda)^2}]} (\theta + \lambda) e^{-x(\theta + \lambda)} y^{-(\theta + \lambda)}, & 0 < y < 1, x > -\ln y, \\ \frac{[1 + \ln y + (2 + \ln y)x + x^2]}{[1 + \ln y + \frac{\ln y + 2}{\theta + \lambda} + \frac{2}{(\theta + \lambda)^2}]} (\theta + \lambda) e^{-x(\theta + \lambda)} y^{-(\theta + \lambda)}, & y > 1, x > 0. \end{cases}$$

۲.۲ گشتاورها و ضریب همبستگی

ابتدا به بررسی گشتاورهای توزیع حاشیه‌ای می‌پردازیم. چون توزیع چگالی شرطی X از توزیع لندلی پیروی می‌کند بنابراین داریم:

$$E(X^r) = \frac{r!(\theta + r + 1)}{\theta^r(\theta + 1)},$$

و

$$E(X) = \theta + 2/\theta(\theta + 1), \quad Var(X) = \frac{\theta^2 + 4\theta + 2}{\theta^2(\theta + 1)^2}.$$

برای مشاهده جزئیات بیشتر توزیع لندلی به قیطانی و همکاران [۲] مراجعه شود. حال برای به دست آوردن r -امین گشتاور توزیع $f_Y(y)$ می‌توان به صورت زیر عمل کرد:

$$\begin{aligned} E(Y) &= M_{Z_1}(1)M_{Z_2}(-1) \\ &= \frac{\lambda^2\theta^2(\theta + 2)}{(\lambda + 1)(\lambda - 1)^2(\theta + 1)^2}, \end{aligned} \quad (2)$$

که در آن $M_{Z_i}(t)$ تابع مولد گشتاور توزیع لندلی است. پس با توجه به (۱۵) داریم:

$$\begin{aligned} E(Y^r) &= E(e^{r(Z_2 - Z_1)}) \\ &= \frac{(\lambda\theta)^2(\lambda - r + 1)(\theta + r + 1)}{(\theta + 1)(\lambda + 1)(\theta + r)^2(\lambda - r)^2}. \end{aligned}$$

حال $Var(Y)$ را می‌توان به صورت زیر نوشت:

$$Var(Y) = \frac{(\lambda\theta)^2}{(\theta + 1)(\lambda + 1)} \left[\frac{(\lambda - 1)(\theta + 3)}{(\lambda - 2)^2(\theta + 2)^2} - \frac{\lambda^2\theta^2(\theta + 2)^2}{(\lambda + 1)(\lambda - 1)^4(\theta + 1)^5} \right].$$

همان‌طور که قبلاً بحث شد، انگیزه اصلی برای توزیع دوبعدی پیشنهاد شده این است که ضریب همبستگی بین X و Y می‌تواند هر مقداری در فاصله $(-1, 0)$ را برخلاف توزیع‌های دومتغیره وابسته منفی موجود اختیار کند. اکنون ضریب همبستگی X و Y را به دست می‌آوریم. با استفاده از استقلال Z_1 و Z_2 داریم:

$$\begin{aligned} Cov(X, Y) &= E(Z_1 e^{-Z_1})E(e^{Z_2}) - E(Z_1)E(e^{Z_2})E(e^{-Z_1}) \\ &= \frac{-\lambda^2\theta(\theta + 4)}{(\lambda + 1)(\lambda - 1)^2(\theta + 1)^4}. \end{aligned} \quad (3)$$

رابطه (۱۷) برای $\lambda > -1$ برقرار است. ضریب همبستگی بین X و Y به صورت زیر محاسبه می‌شود:

$$\rho(X, Y) = \frac{\frac{-\lambda^2\theta(\theta + 4)}{(\lambda + 1)(\lambda - 1)^2(\theta + 1)^4}}{\sqrt{\frac{\theta^2 + 4\theta + 2}{\theta^2(\theta + 1)^2} \frac{(\lambda\theta)^2}{(\theta + 1)(\lambda + 1)} \left[\frac{(\lambda - 1)(\theta + 3)}{(\lambda - 2)^2(\theta + 2)^2} - \frac{\lambda^2\theta^2(\theta + 2)^2}{(\lambda + 1)(\lambda - 1)^4(\theta + 1)^5} \right]}}. \quad (4)$$

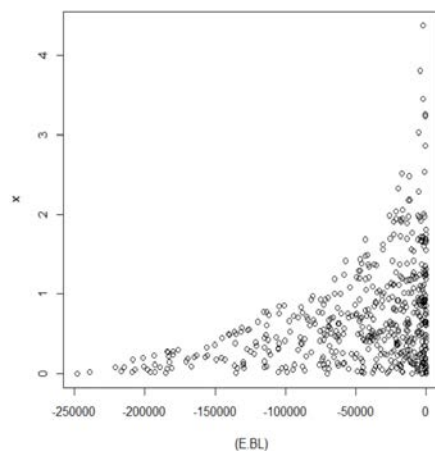
ضریب همبستگی بین X و Y می‌تواند مقادیر بین -۱ و ۰ را اختیار کند. اکنون، امید ریاضی و واریانس شرطی توزیع را محاسبه می‌کنیم. امید ریاضی شرطی و واریانس شرطی $Y|X = x$ به ترتیب عبارت‌اند از:

$$\begin{aligned} E(Y|X = x) &= \frac{\lambda^2}{\lambda + 1} e^{-\lambda x} \int_{e^{-x}}^{\infty} y^{-\lambda} (1 + x + \ln y) dy \\ &= \frac{\lambda^2}{(1 - \lambda^2)(\lambda - 1)} e^{-x}, \quad \lambda > 1, \end{aligned}$$

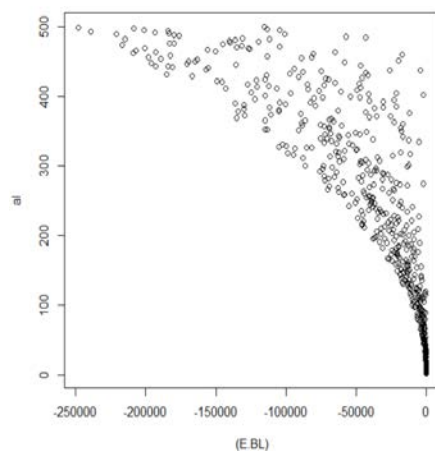
و

$$Var(Y|X = x) = \lambda^2 e^{-2x} \left(\frac{(\lambda - 1)}{(\lambda + 1)(\lambda - 2)(2 - \lambda)} - \frac{\lambda^4}{(1 - \lambda^2)^2 (\lambda - 1)^2} \right), \quad \lambda > 2.$$

با رسم نمودار $E(Y|X = x)$ برحسب تغییرات x و برحسب تغییرات λ متوجه می‌شویم که نسبت به x رفتاری صعودی و نسبت به λ رفتاری نزولی از خود نشان می‌دهد. معادله رگرسیونی Y روی X یک معادله خط لگاریتمی است و این به این معنا است



[نمودار $E(Y|X = x)$ برحسب تغییرات x]



[نمودار $E(Y|X = x)$ برحسب تغییرات λ]

شکل ۱: نمودار $E(Y|X = x)$ برحسب تغییرات x و برحسب تغییرات λ

که $\ln E(Y|X = x)$ تابع خطی از x است. همچنین جالب توجه است که $Var(Y|X = x)$ تابعی نزولی از x است که با

$$\lambda^2 \left(\frac{(\lambda-1)}{(\lambda+1)(\lambda-2)(2-\lambda)} - \frac{\lambda^4}{(1-\lambda^2)^2(\lambda-1)^2} \right) \text{ محدود شده است.}$$

۳ مفاهیم وابستگی منفی

در این بخش، چندین مفهوم وابستگی را معرفی می‌کنیم و این خواص را برای مدل پیشنهادی مورد بررسی قرار می‌دهیم. وابستگی ربعی منفی (مثبت): فرض کنید $F(x, y)$ تابع توزیع از متغیرهای تصادفی زوجی (X, Y) باشد که دارای توابع توزیع حاشیه‌ای $F_1(x)$ و $F_2(y)$ هستند. تابع توزیع F وابستگی ربعی منفی دارد اگر به ازای همه‌ی x و y ها

$$F(x, y) \leq F_1(x)F_2(y). \quad (5)$$

وابستگی به‌طور اکید است اگر نابرابری برای حداقل یک جفت (x, y) برقرار باشد. به‌طور مشابه وابستگی ربعی مثبت داریم اگر در نابرابری (۱۶) جهت نامساوی عوض شود. برای جزئیات بیشتر به لهن [۴] در سال ۱۹۶۶ رجوع شود. وابستگی رگرسیونی منفی (مثبت): $F(x, y)$ وابستگی رگرسیونی منفی دارند اگر و تنها اگر برای همه‌ی y ها و برای همه‌ی $x < x'$

$$F(y|x) \leq F(y|x'), \quad (6)$$

که $F(y|x) = P(Y \leq y | X = x)$. اگر جهت نامساوی (۱۹) عکس شود وابستگی رگرسیونی مثبت داریم. وابستگی نسبت درستنمایی منفی: گویم متغیرهای تصادفی X و Y وابستگی نسبت درستنمایی منفی دارند اگر به ازای هر $x_1 \leq x_2$ و $y_1 \leq y_2$ تابع چگالی احتمال توام $f(x, y)$ از رابطه‌ی زیر پیروی کند (کارلین ۱۹۶۸).

$$f(x_1, y_1)f(x_2, y_2) \leq f(x_1, y_2)f(x_2, y_1).$$

تابع چگالی (۱۴) دارای وابستگی نسبت درستنمایی منفی است. اثبات.

$$\begin{aligned} x_1 \leq x_2 &\Rightarrow x_1(\ln y_2 - \ln y_1) \leq x_2(\ln y_2 - \ln y_1) \\ &\Rightarrow x_1 \ln y_2 + x_2 \ln y_1 \leq x_1 \ln y_1 + x_2 \ln y_2 \\ &\Rightarrow 1 + x_2 + \ln y_2 + x_1 + x_1 x_2 + x_1 \ln y_2 + \ln y_1 + x_2 \ln y_1 + \ln y_1 \ln y_2 \\ &\leq 1 + x_2 + \ln y_1 + x_1 + x_1 x_2 + x_1 \ln y_1 + \ln y_2 + x_2 \ln y_2 + \ln y_1 \ln y_2 \\ &\Rightarrow (1 + x_1 + \ln y_1)(1 + x_2 + \ln y_2) \leq (1 + x_1 + \ln y_2)(1 + x_2 + \ln y_1) \\ &\Rightarrow \frac{(\theta\lambda)^2}{(\theta+1)(\lambda+1)}(1+x_1)(1+x_1+\ln y_1)e^{-x_1(\theta+\lambda)}y_1^{-(1+\lambda)} \\ &\quad \frac{(\theta\lambda)^2}{(\theta+1)(\lambda+1)}(1+x_2)(1+x_2+\ln y_2)e^{-x_2(\theta+\lambda)}y_2^{-(1+\lambda)} \\ &\leq \frac{(\theta\lambda)^2}{(\theta+1)(\lambda+1)}(1+x_1)(1+x_1+\ln y_2)e^{-x_1(\theta+\lambda)}y_2^{-(1+\lambda)} \\ &\quad \frac{(\theta\lambda)^2}{(\theta+1)(\lambda+1)}(1+x_2)(1+x_2+\ln y_1)e^{-x_2(\theta+\lambda)}y_1^{-(1+\lambda)} \end{aligned}$$

□

با توجه به این‌که وابستگی نسبت درستنمایی منفی برقرار است بقیه وابستگی‌ها یعنی وابستگی رگرسیونی منفی و وابستگی ربعی

منفی را نتیجه می‌دهد. با توجه به این‌که وابستگی نسبت درستنمایی منفی برقرار است و در نتیجه وابستگی ربعی منفی نیز نتیجه می‌شود، بنابراین کواریانس نیز منفی می‌شود که در (۱۷) بیان شد. در ادامه برخی از ضرایب همبستگی نظیر ضریب همبستگی کندال (τ_k) ، اسپیرمن (ρ_s) و پیرسون (ρ) را به دست می‌آوریم. خلاصه‌ای از نتایج آن‌ها در جداول ۱ و ۱ آمده است.

جدول ۱: ضرایب همبستگی کندال، اسپیرمن و پیرسون به ازای $\lambda = ۱۰/۵$ و $\theta = ۲$ و n های متفاوت

اندازه نمونه	τ_k	ρ_s	ρ
۲۰	-۰/۸۵۲۶	-۰/۹۴۸۸	-۰/۹۴۴۵
۵۰	-۰/۸۶۴۴	-۰/۹۶۷۶	-۰/۹۳۸۲
۱۰۰	-۰/۸۷۲۷	-۰/۹۷۴۱	-۰/۹۰۶۷
۲۰۰	-۰/۸۹۸۸	-۰/۹۸۲۹	-۰/۱۹۱۶

جدول ۲: ضرایب همبستگی کندال، اسپیرمن و پیرسون به ازای $n = ۱۰۰۰$ و مقادیر مختلف λ و θ

پارامترها	τ_k	ρ_s	ρ
$\lambda = ۱۰/۵, \theta = ۳/۵$	-۰/۷۷۱۳	-۰/۹۱۰۲	-۰/۹۷۴۸
$\lambda = ۸, \theta = ۲$	-۰/۸۱۸۸	-۰/۹۴۹۹	-۰/۸۵۶۲
$\lambda = ۶/۵, \theta = ۱/۵$	-۰/۸۳۳۰	-۰/۹۵۲۶	-۰/۸۳۴۱
$\lambda = ۵, \theta = ۱$	-۰/۸۶۶۳	-۰/۹۷۰۶	-۰/۷۰۳۱
$\lambda = ۴, \theta = ۰/۵$	-۰/۹۱۶۶	-۰/۹۸۷۶	-۰/۶۲۴۱
$\lambda = ۳, \theta = ۰/۲۵$	-۰/۹۴۵۳	-۰/۹۹۴۹	-۰/۳۵۴۴

۴ تابع وابستگی مکانی و ارتباط آن با وابستگی نسبت درستنمایی منفی

فرض کنید (X, Y) یک جفت تصادفی پیوسته دوبعدی با تابع چگالی احتمال $f(x, y)$ باشد. در تعریف تابع وابستگی مکانی، فرض کنید که تکیه گاه K از چگالی توأم $f(x, y)$ یک مجموعه از حاصلضرب دکارتی است. هنگامی که مجموعه K یک مجموعه حاصلضرب دکارتی نیست، شکل آن می‌تواند وابستگی دیگری بین X و Y ایجاد کند که منطقاً با وابستگی مکانی متفاوت است. تابع وابستگی مکانی (X, Y) به صورت زیر تعریف می‌شود:

$$\gamma_f(x, y) = \frac{\partial^2}{\partial x \partial y} \ln f(x, y),$$

در صورتی که مشتقات جزئی وجود داشته باشد. اکنون یک ویژگی بسیار مهم و جالب از تابع وابستگی مکانی ارائه می‌دهیم. فرض کنید $f(x, y)$ تابع چگالی جفت (X, Y) باشد که روی مجموعه K تعریف شده که K یک حاصلضرب دکارتی است. آنگاه $f(x, y)$ دارای وابستگی نسبت درستنمایی منفی (مثبت) است اگر و تنها اگر $\gamma_f(x, y) \leq (\geq) ۰$. برای اثبات می‌توان به هولاند و وانگ ۱۹۸۷ [۳] رجوع کرد. قضیه بالا ما را بر آن می‌دارد تا ویژگی به دست آمده در قضیه را برای توزیع لیندلی دومتغیره بررسی کنیم. با

استفاده از تعریف ۴ برای توزیع لیندلی دومتغیره پیشنهادی داریم:

$$\gamma_f(x, y) = -\frac{1}{y(1+x+\ln y)^2},$$

قابل ذکر است که چون $y > 0$ و عبارت داخل پرانتز هم به توان ۲ رسیده و مقداری مثبت می‌شود، در نتیجه در کل مقداری منفی خواهیم داشت. با توجه به این که $\gamma_f(x, y) \leq 0$ شد، بنابراین توزیع پیشنهادی دارای تایید گزاره ۹ است.

۵ توزیع‌های سیستم‌های سری و موازی

در این قسمت سیستم‌های دمولفه‌ای با احتمال‌های وابسته‌ای را داریم که طول عمر آن با استفاده از بردار تصادفی (X, Y) با قابلیت اعتماد پیوسته $R(x, y) = P(X > x, Y > y)$ است. فرض کنید طول عمر سیستم سری با $M_S = \min(X, Y)$ و سیستم موازی با $M_P = \max(X, Y)$ تعریف شود. در ادامه، توزیع متغیرهای M_S و M_P را برای توزیع لیندلی دومتغیره به دست می‌آوریم که به صورت زیر است:

$$\begin{aligned} F_{M_S}(s) &= 1 - \bar{F}_{Z_1}(s)F_{Z_2}(s + \ln s), \\ F_{M_P}(p) &= F_{Z_1}(p)\bar{F}_{Z_2}(p + \ln p), \end{aligned} \quad (7)$$

که Z_1 و Z_2 متغیرهای تصادفی مستقل از توزیع لیندلی با پارامترهای θ و λ هستند. توابع چگالی احتمال M_S و M_P به ترتیب به قرار زیر است:

$$\begin{aligned} f_{M_S}(s) &= f_{Z_1}(s)F_{Z_2}(s + \ln s) - \left(\frac{s+1}{s}\right)f_{Z_2}(s + \ln s)\bar{F}_{Z_1}(s), \\ f_{M_P}(p) &= f_{Z_1}(p)\bar{F}_{Z_2}(p + \ln p) - \left(\frac{p+1}{p}\right)f_{Z_2}(p + \ln p)F_{Z_1}(p). \end{aligned} \quad (8)$$

۱.۵ معیارهای قابلیت اعتماد

تابع نرخ شکست برای متغیر تصادفی پیوسته تک پارامتری به صورت $h(t) = \frac{f(t)}{R(t)}$ تعریف می‌شود که $R(t) = 1 - F(t)$ تابع قابلیت اعتماد است. تابع میانگین طول عمر باقی‌مانده نیز به صورت $m(t) = \frac{1}{R(t)} \int_t^\infty R(x)dx$ تعریف می‌شود. اگر بردار تصادفی (X, Y) از توزیع لیندلی دومتغیره پیشنهادی پیروی کند، آنگاه تابع نرخ شکست و تابع میانگین باقی‌مانده طول عمر برای سیستم‌های سری و موازی به ترتیب به صورت زیر به دست می‌آید. با استفاده از تعاریف تابع نرخ شکست، تابع قابلیت

اعتماد و روابط (۷) و (۸)، تابع نرخ خطر برای سیستم‌های سری و موازی به صورت زیر تعریف می‌شود:

$$\begin{aligned} h_{M_S}(s) &= \frac{f_{Z_1}(s)F_{Z_2}(s + \ln s) - (\frac{s+1}{s})f_{Z_2}(s + \ln s)\bar{F}_{Z_1}(s)}{\bar{F}_{Z_1}(s)F_{Z_2}(s + \ln s)}, \\ h_{M_P}(p) &= \frac{f_{Z_1}(p)\bar{F}_{Z_2}(p + \ln p) - (\frac{p+1}{p})f_{Z_2}(p + \ln p)F_{Z_1}(p)}{1 - F_{Z_1}(p)\bar{F}_{Z_2}(p + \ln p)}, \\ m_{M_S}(s) &= \frac{1}{\bar{F}_{Z_1}(s)F_{Z_2}(s + \ln s)} \int_s^\infty \bar{F}_{Z_1}(x)F_{Z_2}(x + \ln x)dx, \\ m_{M_P}(p) &= \frac{1}{1 - F_{Z_1}(p)\bar{F}_{Z_2}(p + \ln p)} \int_p^\infty (1 - F_{Z_1}(x)\bar{F}_{Z_2}(x + \ln x))dx. \end{aligned}$$

توجه شود که Z_1 و Z_2 متغیرهای تصادفی مستقل از توزیع لیندلی با پارامترهای θ و λ هستند.

۶ برآورد پارامترها

در این قسمت، به برآورد پارامترهای مدل پیشنهادی می‌پردازیم. برای برآورد پارامترهای نامعلوم θ و λ به روش ماکسیمم درستنمایی (MLE)، نمونه‌ی تصادفی دومتغیره $(X_1, Y_1), \dots, (X_n, Y_n)$ را در نظر می‌گیریم. تابع درستنمایی (θ, λ) به صورت زیر است:

$$\begin{aligned} L(\theta, \lambda) &= \prod_{i=1}^n \frac{(\theta\lambda)^2}{(\theta+1)(\lambda+1)} (1+x_i)(1+x_i+\ln y_i)e^{-x_i(\theta+\lambda)}y_i^{-(1+\lambda)} \\ &= \frac{(\theta\lambda)^{2n}}{(\theta+1)^n(\lambda+1)^n} \prod_{i=1}^n (1+x_i) \prod_{i=1}^n (1+x_i+\ln y_i)e^{-\sum_{i=1}^n x_i(\theta+\lambda)}e^{-(1+\lambda)\sum_{i=1}^n \ln y_i} \end{aligned}$$

لگاریتم تابع درستنمایی را می‌توان به صورت زیر نوشت:

$$\begin{aligned} l(\theta, \lambda) &= 2n \ln \theta + 2n \ln \lambda - n \ln(\theta+1) - n \ln(\lambda+1) + \sum_{i=1}^n \ln(1+x_i) \\ &\quad + \sum_{i=1}^n \ln(1+x_i+\ln y_i) - (\lambda+\theta) \sum_{i=1}^n x_i - (1+\lambda) \sum_{i=1}^n \ln y_i \end{aligned}$$

پس از مشتق‌گیری، برآوردهای θ و λ به ترتیب عبارت‌اند از:

$$\hat{\theta} = \frac{-(\sum_{i=1}^n x_i - n) + \sqrt{(\sum_{i=1}^n x_i - n)^2 + 4n \sum_{i=1}^n x_i}}{2 \sum_{i=1}^n x_i},$$

و

$$\hat{\lambda} = \frac{-(\sum_{i=1}^n x_i + \sum_{i=1}^n \ln y_i - n) + \sqrt{(\sum_{i=1}^n x_i + \sum_{i=1}^n \ln y_i - n)^2 + 4n \sum_{i=1}^n x_i + \sum_{i=1}^n \ln y_i}}{2(\sum_{i=1}^n x_i + \sum_{i=1}^n \ln y_i)}.$$

قابل ذکر است که برآورد ML پارامتر θ فقط تابعی از x_i ها است و همانند برآورد ML برای توزیع لیندلی با پارامتر θ است. با

مشتق‌گیری دوم نسبت به پارامترهای θ و α داریم:

$$\frac{\partial^2 l(\theta, \lambda)}{\partial \theta^2} = \frac{-2n}{\theta^2} + \frac{n}{(\theta+1)^2}, \quad \frac{\partial^2 l(\theta, \lambda)}{\partial \lambda^2} = \frac{-2n}{\lambda^2} + \frac{n}{(\lambda+1)^2}, \quad \frac{\partial^2 l(\alpha, \lambda)}{\partial \theta \partial \lambda} = 0. \quad (9)$$

فرض کنید $\hat{\theta}$ و $\hat{\lambda}$ برآوردهای ماکسیمم درستنمایی از θ و λ هستند و ماتریس اطلاع فیشر به صورت زیر است:

$$I(\hat{\theta}, \hat{\lambda}) = \begin{bmatrix} -\frac{\partial^2 l(\theta, \lambda)}{\partial \theta^2} & -\frac{\partial^2 l(\theta, \lambda)}{\partial \theta \partial \lambda} \\ -\frac{\partial^2 l(\theta, \lambda)}{\partial \theta \partial \lambda} & -\frac{\partial^2 l(\theta, \lambda)}{\partial \lambda^2} \end{bmatrix}_{\theta=\hat{\theta}, \lambda=\hat{\lambda}} \quad (10)$$

بنابراین با توجه به (۹) و (۱۰) مشاهده می شود که ماتریس به دست آمده معین منفی است، در این صورت برآوردهای به دست آمده MLE هستند.

۷ شبیه سازی توزیع لیندلی دومتغیره با وابستگی منفی

برای تولید بردار تصادفی (X, Y) که از توزیع لیندلی با وابستگی منفی پیروی کند، الگوریتم زیر دنبال می شود:

۱. تولید متغیر Z_1 از توزیع $Lindley(\theta)$ ؛

۲. تولید متغیر Z_2 از توزیع $Lindley(\lambda)$ ؛

۳. در نظر گرفتن $X = Z_1$ و $Y = Exp(Z_2 - Z_1)$ ؛

۴. پس زوج (X, Y) دارای توزیع لیندلی دومتغیره با پارامترهای θ و λ است.

برای بررسی برآوردکننده های ماکسیمم درستنمایی λ و θ نمونه ای از توزیع لیندلی دومتغیره با وابستگی منفی در نظر گرفته می شود. خلاصه ی نتایج در جدول ۵ آمده است. این جداول نشان می دهند که اریبی و میانگین مربعات خطا با افزایش n کاهش می یابد.

جدول ۳: برآورد ML پارامترهای $\lambda = 10/5$ و $\theta = 2$ توزیع لیندلی دومتغیره با وابستگی منفی

	$n = 20$	$n = 50$	$n = 100$	$n = 200$
$\lambda = 10/5$				
$\hat{\lambda}$	۱۰/۹۹۶۰	۱۰/۶۸۷۹	۱۰/۵۷۵۷	۱۰/۵۲۶۰
<i>Bias</i>	۰/۴۹۶۰	۰/۱۸۹۷	۰/۰۷۵۷	۰/۰۲۶۰
<i>MSE</i>	۵/۴۴۴۱	۲/۲۴۴۴	۰/۹۷۴۷	۰/۴۷۱۵
$\theta = 2$				
$\hat{\theta}$	۲/۰۸۷۷	۲/۰۳۶۹	۲/۰۱۷۶	۲/۰۱۶۲
<i>Bias</i>	۰/۰۸۷۷	۰/۰۳۶۹	۰/۰۱۷۶	۰/۰۱۶۲
<i>MSE</i>	۰/۱۴۷۲	۰/۰۵۴۳	۰/۰۲۴۷	۰/۰۱۳۱

مراجع

- [1] Bhuyan, P. S., Ghosh, Sh. and Majumder, P. (2018), *A Bivariate Life Distribution with Negative Dependence*, Murch.

- [2] Ghitany, ME., Atieh, B. and Nadarajah, S. (2008), Lindley distribution and its Application, *Mathematics Computing and Simulation*, **78**, 493-506.
- [3] Holland, P.W. and Wang, Y.J. (1987), Dependence function for continuous bivariate densities, *Communications in Statistics-Theory and Methods*, **16**, 863–887.
- [4] Lehmann, E.L. (1966), Some concepts of dependence, *The Annals of Mathematical Statistics*, **37**(5), 1137-1153.
- [5] Lindley, D.V. (1958), Fiducial distribution and Bayes theorem, *Journal of the Royal Statistical Society*, Series B, **20**, 102-107.

بررسی قابلیت اطمینان پیش‌بینی زمان سفر با مدل GARCH

مریم جمالی^۱، رسول روزگار

گروه آمار، دانشکده علوم ریاضی، دانشگاه یزد

چکیده: زمان سفر یک مقیاس اساسی در حمل و نقل است و پیش‌بینی دقیق زمان سفر در سیستم‌های حمل و نقل هوشمند بسیار مهم است. در حال حاضر تکنیک‌های بسیاری برای پیش‌بینی زمان سفر استفاده می‌شود؛ با این حال قابلیت اطمینان پیش‌بینی در این روش‌ها مورد بررسی قرار نگرفته است. قابلیت اطمینان پیش‌بینی زمان سفر با توجه به اینکه تغییراتش وابسته به زمان است، راه را برای اندازه‌گیری عملکرد سیستم‌ها فراهم می‌آورد. یکی از معیارها برای بررسی قابلیت اطمینان سفر، شناسایی فواصل پیش‌بینی (PI) است. اندازه‌گیری PI دارای اهمیت زیادی در توسعه سیستم‌هایی است که هدف آن‌ها انتشار اطلاعات ترافیکی واقعی به مسافران است. روش استفاده از مدل GARCH برای بررسی نوسانات زمان سفر و ارائه اطلاعات مربوط به قابلیت اطمینان برای پیش‌بینی زمان سفر بسیار اهمیت دارد. دو نمونه از داده‌های ترافیک شهری واقعی که روند کل مدل سازی را نشان می‌دهند برای بررسی قابلیت اطمینان پیش‌بینی در نظر گرفتیم، در آزمایشات از انحراف استاندارد پیش‌بینی شده شرطی (PSD) و فواصل پیش‌بینی (PI) برای بیان قابلیت اطمینان پیش‌بینی زمان سفر استفاده می‌کنیم. نتایج نشان می‌دهد که خطای میانگین ریشه (RMSE)، میانگین خطای مطلق (MAE) و میانگین خطای مطلق درصد (MAPE) با افزایش قابلیت اطمینان کاهش می‌یابد همچنین در مثال دیگر PI‌های معقولی بدست می‌آید. با توجه به این دو معیار ثابت می‌شود که مدل GARCH به خوبی قابلیت اطمینان پیش‌بینی زمان سفر را مشخص می‌کند و استفاده از روش‌های پیشنهادی امکان پذیر است.

واژه‌های کلیدی: پیش‌بینی زمان سفر، زمان سفر، سری زمانی، مدل GARCH، مدیریت ترافیک.

کد موضوع بندی: 39B82، 34K20، 39B52، 47A55.

۱ پیش‌گفتار

پیش‌بینی زمان سفر یکی از مهمترین و اساسی ترین مشکلات در سیستم‌های حمل و نقل هوشمند است. پیش‌بینی دقیق زمان سفر یکی از راه‌های به دست آوردن شرایط ترافیک در هنگام انجام برنامه ریزی برای یک مسیر است. به عنوان یک جنبه مهم از شرایط ترافیک، زمان سفر، دیدگاهی برای وضعیت ترافیک را نشان می‌دهد و پیش‌بینی زمان سفر، علاقه بسیاری از محققان را به خود جلب کرده است. یکی دیگر از مهمترین روش‌های پیش‌بینی مربوط به تحلیل سری‌های زمانی است که پیش‌بینی مقدار آینده یک متغیر بر اساس مقادیر گذشته همان متغیر است. در این روش‌ها یک سری زمانی برای زمان سفر به طور کلی به عنوان مجموعه‌ای از مشاهدات آماری است که به ترتیب زمان مرتب شده‌اند، در نظر گرفته می‌شود. یک سری مشاهده شده به لحاظ تئوری از دو بخش تشکیل می‌شود، سری تولید شده توسط فرایند واقعی و قسمت نویز یا خطا که نتیجه اختلالات بیرونی است و معمولاً به عنوان یک روند تصادفی در نظر گرفته می‌شود. بنابراین، حذف مولفه‌ی نویز قسمت اصلی برای مدل سازی و پیش‌بینی است. تکنیک‌های

^۱مریم جمالی: maryam.jamali@stu.ac.ir

متعددی مانند مدل ARIMA و مدل فصلی ARIMA توسط چندین محقق مورد استفاده قرار گرفته و تحت شرایط خاصی قابل اجرا هستند [۱].

قابلیت اطمینان یک معیار بسیار مهم است، به عنوان مثال در بهینه سازی مسیر ترافیک یا سیستم های کنترل ترافیک، اگر پیش بینی زمان سفر برخی از بخش ها قابل اعتماد نباشد، برنامه ریزی برای بهینه سازی مسیر یا کنترل ترافیک قابل اعتماد نیست. اگر اطمینان از پیش بینی زمان سفر در نظر گرفته شود، برنامه بهینه سازی مسیر یا کنترل ترافیک بسیار قابل اعتماد است. اگر می خواهیم اطمینان از پیش بینی ترافیک را بدانیم، رفتار و اندازه واریانس را نمی توان نادیده گرفت. هر دو واریانس شرطی و غیر شرطی مولفه ی نویز در حالت کلی یا در نظر گرفته نمی شود و یا به عنوان ثابت در نظر گرفته می شود، به عنوان مثال مقدار مورد انتظار از مربع هر یک از خطای داده شده معادل واریانس تمام خطاها با هم است. این فرضیه به نام ناهمسان نامیده می شود. به عبارت دیگر، اطمینان همه پیش بینی ها یکسان در نظر گرفته می شود. با این حال در حقیقت فرضیه این است که واریانس شرطی مولفه نویز ثابت است، به دلیل اثر رویدادهای غیر منتظره مانند حوادث ترافیکی و آب و هوایی، که در عمل یک امر عادی است است مناسب نیست. مقدار مورد انتظار خطا ممکن است برابر نباشد و انتظار می رود که شرایط خطا برای برخی از داده ها نسبت به دیگر داده ها بیشتر باشد. در اینجا، گفته می شود که داده ها از ناهمبستگی برخوردارند. اگر نویزها ناهمسان باشد، تخمین انحرافات استاندارد در مدل های فوق الذکر غیر قابل قبول است، اگر چه برآورد پارامترها بی فایده است. در این مثال، ما نمی توانیم اطلاعاتی درباره قابلیت اطمینان پیش بینی را بدست آوریم.

مدل خودبازگشتی با ناهمسانی در واریانس شرطی (ARCH) به طور خاص برای مدل سازی و پیش بینی واریانس های شرطی یا انحرافات استاندارد طراحی شده است. اگر چه مدل های ARCH مدلهایی ساده و در عمل آسان هستند، اما قادر به تشخیص خطاهای خوشه ای، غیر خطی و تغییرات در پیش بینی هستند. مدل ARCH توسط Engle [۲] معرفی شد و به وسیله Bollerslev و Taylor [۳] تعمیم یافته آن با عنوان مدل GARCH معرفی شد. Kamarianakis و همکاران [۴] روشی را که ترکیبی از فرآیند ARIMA برای مدلسازی میانگین شرطی و فرآیند GARCH برای مدلسازی واریانس شرطی است، به کار برده اند. بعدها Stathopoulos و Tsekeris روش Kamarianakis و همکاران را با استفاده از ARIMA-GARCH بهبود بخشیدند و مدل های ARFIMA-FIAPARCH را معرفی کردند و بعدها در برخی از مطالعات مدل GARCH در حوزه ترافیک، از سرعت نسبی پیش بینی، جریان ترافیکی و زمان سفر معرفی شده است. همه این مطالعات نشان می دهد که تغییرات ظاهری در نوسانات داده های ترافیکی قابل پیش بینی است و ممکن است از نوع خاصی از وابستگی غیر خطی حاصل شود. هدف از استفاده از مدل GARCH برای مطالعه و مدلسازی واریانس شرطی نویز است. بر اساس این مدل، ما می توانیم نوسان داده های زمان سفر را پیش بینی و برآورد کنیم. و از طریق این اطلاعات می توانیم قابلیت اطمینان را برای پیش بینی زمان سفر به دست آوریم. بر خلاف بعضی روش ها معادله میانگین روش پیشنهادی به مدل (Autoregressive Average-Moving) (ARMA) محدود نمی شود [۴]. به طور مشابه، این روش می تواند به دیگر مدل های معادله پیچیده تر هم اعمال شود. علاوه بر این، ما نقاط نمونه ای را که پیش بینی قابل اعتمادی می دهد را نمایش می دهیم. بر اساس این روش می توان دید که دقت نقاط نمایش داده شده بالاتر است [۴].

از آنجائیکه زمان سفر در دوره های مختلف و شرایط ترافیکی به طور قابل توجهی تغییر می کند، روش پیش بینی نقطه ای زمانی که با شرایط ترافیکی نامطلوب برخورد می شود کمتر قابل اعتماد و دقیق است. برای پیش بینی، یکی از شاخص های کارآیی، فواصل پیش بینی (PI) است که بر اساس آن می توان پایایی نتایج پیش بینی زمان سفر را ارزیابی کرد. PI یک فاصله زمانی محاسبه شده است که مقادیر زمان انتظار سفر را با احتمالات تعیین شده قبلی پوشش می دهد. در دسترس بودن PI ها، مسافران و مدیران ترافیک را قادر می سازد تا میزان عدم قطعیت پیش بینی زمان سفر را تعیین کنند و بنابراین راه های متعددی را برای انتخاب مسیر و زمان

خروج برای مقابله با بدترین و بهترین شرایط ایجاد می کنند. PI های کوچک بدان معنی است که وضعیت ترافیک نسبتاً پایدار است یکی دیگر از محدودیت های بسیاری از روش های پیش بینی فعلی نقطه ای این است که آنها واریانس داده ها را به عنوان مقادیر غیر وابسته و ثابت محسوب می کنند. با اوج این که در عمل واریانس تابعی از زمان است. اگرچه فرض واریانس ثابت ساده ساختار مدل نصب شده است، اما دقت و قابلیت اطمینان پیش بینی را به خطر می اندازد. [۵]

بقیه این مقاله به شرح زیر است: پس از معرفی مختصر مدل GARCH در بخش دوم، فرآیند ساخت مدل و پیش بینی واریانس شرطی و نتایج تجربی را در بخش های بعد توصیف می کنیم. در نهایت، بعضی از نتایج در بخش آخر ذکر شده می شود.

۲ مدل خودبازگشتی تعمیم یافته با ناهمسانی در واریانس شرطی

ایده اصلی GARCH این است که در یک زمان معین خطاها از توزیع گوسی با میانگین صفر پیروی می کنند، و تغییرات واریانس با تغییرات زمان (به عنوان مثال، heteroscedasticity مشروط) مطابقت دارد. واریانس یک ترکیب خطی از مقادیر متناهی قبلی، از جمله مربع خطاها و واریانس است. ارائه مختصری از فرم مدل GARCH ارائه شده است.

یک سری زمانی ε_t را از فرایند GARCH با $p \geq 1$ و $q \geq 0$ گویند اگر ε_t در شرایط زیر صدق کند:

$$\varepsilon_t = h_t \eta_t \quad (1)$$

$$h_t^2 = \alpha_0 + \sum_{i=1}^p \alpha_i \varepsilon_{t-i}^2 + \sum_{j=1}^q \beta_j h_{t-j}^2 \quad (2)$$

که ضرایب α_0 و α_i و β_j نامنفی هستند و $\eta_t \sim iidN(0, 1)$ و $\forall t, \eta_t$ مستقل از $\varepsilon_{t-k}, k \geq 0$ است. که برابر با مدل ARCH(p) اگر $q=0$ است. از نظر عملی، مدل های GARCH می توانند برآورد واریانس شرطی دقیق ارائه دهند.

از روابط فوق مشاهده می شود که واریانس خطای یک مدل رگرسیونی از مربع خطاها و واریانس قبلی است. بنابراین نوسانات خطاها دارای ویژگی حافظه تاریخی است. به عبارت دیگر، اگر واریانس چندین بار قبلی بزرگ باشد، واریانس زمان فعلی هم اغلب بزرگ است. در مقابل، اگر واریانس چندین بار گذشته کوچک باشد، واریانس زمان فعلی اغلب کوچک است. این مشخصه مدل ARCH تعمیم یافته است که خوشه بندی نوسانی نامیده می شود.

۳ ساختار مدل

نمودار جریان روش پیش بینی می تواند به شرح زیر نتیجه گیری شود: (۱) برای سری زمانی سفر و حذف بخش خطی تعیین شده؛ (۲) بررسی خطاهای باقی مانده برآورد مدل و بررسی اینکه آیا اثر ARCH وجود دارد؛ و (۳) اگر مجموعه خطاهای باقی مانده شرطی است که به صورت ناهمسان باشد، یک مدل ARCH (یا GARCH) را بسازیم.

۱.۳ ساخت یک مدل معیارشناسی

بسیاری از مدل‌ها می‌توانند به عنوان مدل اندازه‌گیری زمان سفر مانند ARMA، ARIMA و فصلی انتخاب شوند. پس از برآورد مدل، سری خطای باقی مانده ε_t را بدست می‌آوریم. گام بعدی بررسی اینکه آیا اثر ARCH برای سری خطای باقی مانده ε_t وجود دارد.

۲.۳ بررسی تأثیر مدل ARCH

روش‌های مختلفی برای بررسی تأثیر ARCH وجود دارد، مانند تست چندگانه LM [۲] و آزمایش Ljung-Box Q. آزمون LM به طور خلاصه به شرح زیر است که در این مقاله استفاده شده است:

$$1- \text{ساختن مدل } AR(r) \text{ برای مربع سری } \varepsilon_t^2 \text{ از باقیمانده خطاهای } \varepsilon_t \text{ به عنوان مثال:}$$

$$\varepsilon_t^2 = \varepsilon_0 + \alpha_1 \varepsilon_{t-1}^2 + \dots + \alpha_r \varepsilon_{t-r}^2 + e_t \quad (3)$$

که $e_t \sim iid(0, \delta^2)$ است.

۲- آزمون فرضیه:

فرضیه صفر: $\alpha_i = 0, i = 1, 2, \dots, r$ که در اینجا تأثیری از مدل ARCH وجود ندارد.

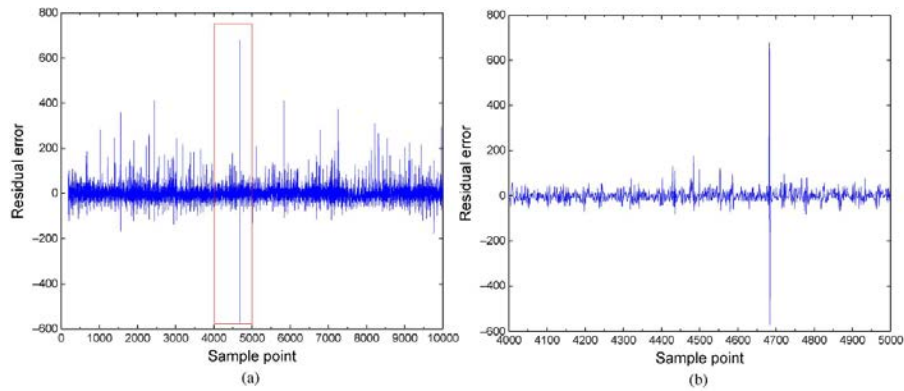
فرضیه جایگزین: $\alpha_i \neq 0$ برای هر i .

و آماره آزمون: $LM = TR^2$ که T سایز نمونه است و R^2 نیکویی برازش است. اگر فرض صفر درست باشد آنگاه $LM \sim \chi^2(r)$ است. به عبارت دیگر اگر فرض صفر را نتوان رد کرد به این معناست که سری یک فرایند ARCH نیست. [۱]

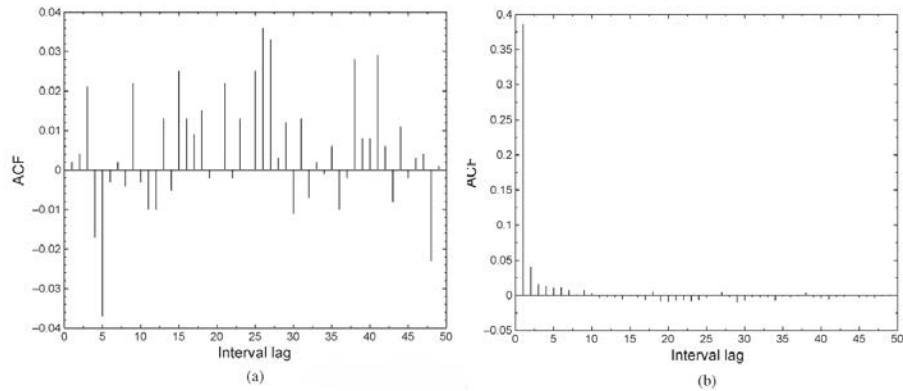
۳- ساختار یک مدل ARCH (یا GARCH)

اگر آزمون قبلی نشان داد که سری ARCH موجود از مرتبه پایین است، مدل ARCH با مرتبه به دست آمده می‌تواند تصویب شود. به طور خلاصه، اگر آزمون نشان می‌دهد که سری ARCH از مرتبه بالا است، مدل GARCH از مرتبه پایین‌تر مانند GARCH (2,1) و GARCH (1,2) و (1,1) و غیره نیز توصیه می‌شود. مرتبه مدل با تحلیل تابع خود همبستگی (ACF) و تابع خود همبستگی جزئی (PACF) از سری ε_t^2 نیز به دست می‌آید. علاوه بر این، مانند مدل ARMA، مرتبه نهایی مدل ARCH و GARCH نیز می‌تواند براساس معیار اطلاع آکاییک (AIC) یا معیار اطلاعات بیزی شاروٹ (BIC یا SBC) [۵] انتخاب شود و در نهایت پارامترهای با روش ماکسیمم درستنمایی برآورد می‌شوند.

پس از برآورد پارامترهای مدل GARCH، آزمون سازگاری مدل ضروری است. با استفاده از مدل برآورد شده GARCH، می‌توانیم واریانس شرطی سری $(\hat{h}_t^2)$ را برآورد کنیم. بر این اساس، مجموعه خطاهای باقی مانده $\hat{\eta}_t = \varepsilon_t / \hat{h}_t$ است. تست LM (یا آزمون Ljung-Box Q) می‌تواند دوباره برای بررسی اینکه آیا اثر خودهمبستگی در سری $\hat{\eta}_t^2$ از خطاهای باقی مانده وجود دارد یا خیر. مدل را می‌توان مناسب دانست در صورتی که خودهمبستگی $\hat{\eta}_t^2$ وجود نداشته باشد.



شکل ۱: (a) خطای باقیمانده از مدل SARIMA و (b) قسمت داخل مستطیل شکل a را به صورت واضح نمایش میدهد



شکل ۲: توابع خودهمبستگی باقیمانده خطاها و مربع باقیمانده خطاهای مدل SARIMA است

۴ پیشبینی از مدل خودبازگشتی تعمیم یافته با نا همسانی در واریانس

سری زمانی $\varepsilon_t, t = 1, 2, \dots, T$ را در نظر بگیرید، واریانس شرطی مدل GARCH در قسمت قبل نمایش داده شد. برای پیشبینی در زمان T با استفاده از مدل GARCH داریم

$$\hat{h}_{T+1}^2 = E(h_{T+1}^2 | \psi_T) \quad (4)$$

که ψ_T در بر دارنده ی اطلاعات قبل از زمان T است. در مدل GARCH(p,q) از آنجا که واریانس شرطی تنها به p فاصله ی زمانی گذشته و واریانس شرطی q فاصله زمانی گذشته وابسته است داریم:

$$\begin{aligned}\hat{h}_{T+1}^y &= E(\alpha_0 + \sum_{i=1}^p \alpha_i \varepsilon_{T-i+1}^y + \sum_{j=1}^q \beta_j h_{T-j+1}^y) \\ &= \alpha_0 + \sum_{i=1}^p \alpha_i \varepsilon_{T-i+1}^y + \sum_{j=1}^q \beta_j h_{T-j+1}^y\end{aligned}\quad (5)$$

۵ تجزیه تحلیل و نتایج آزمون

داده‌های مثال اول برای تجزیه و تحلیل از سیستم اندازه‌گیری زمان سفر در پکن چین گرفته شده است، داده از ۱ دسامبر ۲۰۰۶ تا ۲۰ ژوئن ۲۰۰۷ به مدت ۲۰۲ روز و ۱۹۳۹۲ عدد نمونه است. ۱۰۰۰۰ نمونه اولیه به عنوان داده‌های آزمایشی پردازش می‌شوند و بقیه برای داده‌های آزمون مورد استفاده قرار می‌گیرند. قابل توجه است که همه پیش‌بینی‌ها در آزمون آفلاین هستند. و داده‌های مثال بعد از کل سال ۲۰۰۸ است. داده‌های گم شده گاه به صورت میانگین سالانه از فواصل زمانی گم شده جایگزین شدند. در نتیجه ۱۷۹ روز کاری از اطلاعات زمان سفر شامل ۵۱۵۵۲ مشاهدات پنج دقیقه‌ای برای این مطالعه انتخاب شدند. فرکانس جمع‌آوری شده‌ی پنج دقیقه‌ای به اندازه کافی کوتاه است تا اطلاعات مهم درباره شرایط ترافیکی واقعی را نشان دهد. از آنجا که نوسانات در ساعت‌های ۳ تا ۷ قابل مشاهده‌تر از بقیه ساعات شبانه روز است بنابراین این مطالعه بر روی نوسانات داده‌های زمان سفر از ۳ تا ۷ صبح تمرکز دارد.

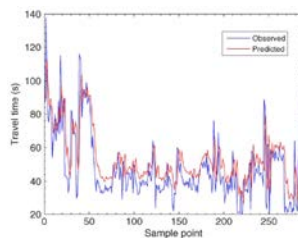
۱.۵ داده‌های جاده ریتان (۰.۷۵ کیلومتر)

در این مثال مدل ANN را برای جمع‌آوری داده‌ها در نظر می‌گیریم. مدل ANN از یک شبکه عصبی سه بعدی استفاده می‌کند که در لایه دوم ده سلول عصبی وجود دارد. داده‌های ورودی میانگین ۵ زمان سفر در گذشته است و خروجی، پیش‌بینی شده زمان سفر است. پس از آزمایش، خطاهای باقی مانده را به دست می‌آوریم، سپس مراحل ۲ (تست اثر ARCH) و ۳ (ساخت مدل GARCH) که در بخش ۳ توضیح داده شد را انجام می‌دهیم. در نهایت، مدل GARCH (۱،۱) برای سری خطاهای باقی مانده ساخته شده است و انحراف معیار استاندارد شرطی به صورت زیر به دست آمده است:

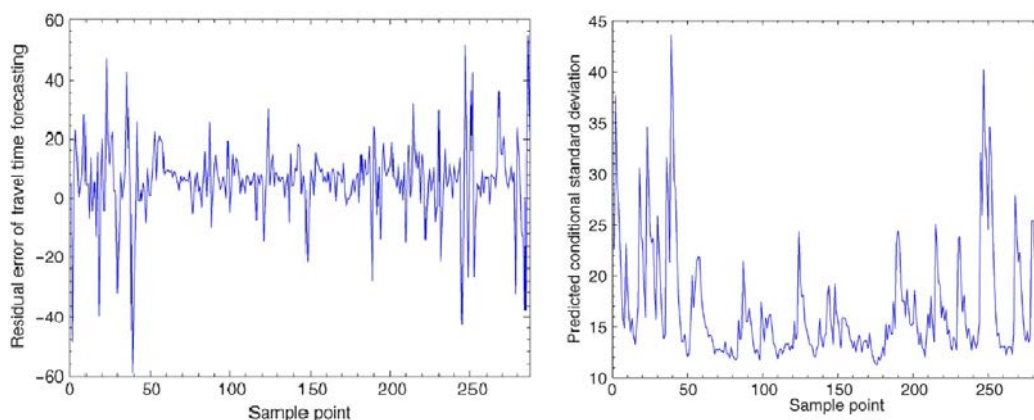
$$\hat{h}_t^y = 57.263 + 0.461 \varepsilon_{t-1}^y + 0.529 h_{t-1}^y \quad (6)$$

شکل ۳ نتایج پیش‌بینی زمان سفر سه روز را نشان می‌دهد. و شکل ۴ به ترتیب خطاهای احتمالی پیش‌بینی شده و انحراف معیار استاندارد هستند. قابلیت اطمینان پیش‌بینی زمان سفر توسط PSD بیان می‌شود، سپس صحت PSD ارزیابی می‌شود. نتایج در جدول زیر نشان داده شده است.

می‌توان دید که تمام مقادیر RMSE کمتر از مقادیر اولیه است و بعد از آزمایش، مقادیر RMSE، MAE و MAPE را می‌توان بسیار محدود کرد. پیش‌بینی انحراف استاندارد در این مثال بسیار دقیق است.



شکل ۳: نتایج پیش بینی برای سه روز آینده



شکل ۴: خطاهای احتمالی پیش بینی شده و انحراف معیار استاندارد

جدول نتایج مدل $GARCH(1,2) - ARIMA(1,1,1)$ را نشان می‌دهد. می‌توان نتیجه گرفت که عملکرد مدل $ARIMA-GARCH$ بسیار ضعیف است، به ویژه وقتی که تعداد اعضای نمونه کمتر از ۸۰۰۰ باشد. این مثال نشان می‌دهد که در روش پیشنهادی معادله می‌تواند آزادانه انتخاب شود. اگر مدل انتخاب شده مناسب باشد، می‌توان نتایج بسیار خوبی کسب کرد.

۲.۵ ایستگاه‌های تشخیص خودکار وسایل نقلیه واقع در امتداد بزرگراه هستون در U.S

در این مثال مدل $ARIMA$ به عنوان معادله میانگین به دلیل سهولت اجرای آن و کاربرد شناخته شده آن برای مدل سازی و پیش بینی پارامترهای ترافیک انتخاب شده است. دستورالعمل مناسب $ARIMA$ براساس معیار اطلاع AIC تعیین می‌شود. بهترین مدل آن است که دارای کمترین مقدار AIC است:

$$AIC = -2\log(L) + 2m \quad (7)$$

که در آن L احتمال داده برای مدل خاص است و m تعداد پارامترهای انتخاب شده برای مدل است. پارامترهای مدل با استفاده از روش حداکثر درستنمای توسط نرم افزار آماری R برآورد می‌شوند. دستورات بهینه و پارامترهای مدل $ARIMA$ برای هر مجموعه داده در جدول زیر ارائه شده است. ضرایب مربوط به قسمت های AR و MA از مدل های $AR(p)$ و $MA(q)$ نیز در جدول آمده است.

برای آزمون تاثیر ناهمسانی شرطی، که همچنین به عنوان اثر $ARCH$ شناخته می‌شود از روش $Ljung-Box$ استفاده می‌کنیم

جدول ۱: PSD

deviation standard predicted of threshold	point sample reliable of Number	RMSE	MAE	MAPE
Infinite	۹۳۹۲	۴۹.۶۸۹۷	۲۸.۹۸۲	۱۴.۰۹
۴۵	۷۵۸۷	۴۰.۲۲۶۸	۲۳.۹۶۹۶	۱۱.۸۵
۴۰	۶۹۷۶	۳۶.۳۴۰۷	۲۲.۷۴۶۵	۱۱.۴۱
۳۵	۶۰۸۱	۳۴.۰۶۳۴	۲۱.۴۹۹۴	۱۰.۹۲
۳۰	۴۴۱۳	۲۹.۶۳۷۸	۱۹.۸۳۹۵	۱۰.۱۵
۲۸	۳۴۷۶	۲۸.۵۰۷۵	۱۹.۲۵۴۲	۹.۸۱
۲۵	۱۴۶۰	۲۶.۱۱۸۱	۱۷.۶۱۱۲	۹.۰۳
۲۴	۶۷۳	۲۱.۴۲۲۱	۱۵.۹۰۲۱	۸.۲۳
۲۳	۵۸	۱۹.۵۶۹۷	۱۴.۰۹۳۶	۷.۱۶

جدول ۲: جدول برآورد توسط مدل ARIMA

Factor	Model	AR(۱)	AR(۲)	MA(۱)	MA(۲)	MA(۳)	MA(۴)	MA(۵)	MA(۶)
۱	ARIMA(۲,۲,۳)								
	Coefficients	۰.۳۹۳۱	۰.۳۲۴۳	۱.۰۷۰۷-	۰.۲۲۲۴-	۰.۳۰۲۲	na	na	na
	SE	۰.۲۷۸۷	۰.۱۹۳۲	۰.۲۷۴۹	۰.۳۶۷۳	۰.۱۱۳	na	na	na
۲	ARIMA(۰,۱,۶)								
	Coefficients	na	na	۰.۳۶۹۷	۰.۱۸۴۳	۰.۱۱۶۱	۰.۰۵۶۶	۰.۰۰۰۷	۰.۰۷۶۳
	SE	na	na	۰.۰۳۳۶	۰.۰۳۶۲	۰.۰۳۵۹	۰.۰۳۲۳	۰.۰۳۵۱	۰.۰۳۴۵
۳	ARIMA(۱,۱,۵)								
	Coefficients	۰.۷۹۷۸	na	۰.۳۲۰۳-	۰.۰۷۱۳-	۰.۱۴۰۳-	۰.۰۱۳۹	۰.۰۹۵۲	na
	SE	۰.۰۶	na	۰.۰۶۷۳	۰.۰۴۳۷	۰.۰۳۸۹	۰.۰۳۶	۰.۰۳۲۸	na
۴	ARIMA(۲,۱,۴)								
	Coefficients	۰.۴۷	۰.۲۳۷۲	۰.۳۰۳۸-	۰.۱۷۷۶	۰.۰۸۴۷۲	۰.۸۷۵	na	na
	SE	۰.۳۱۱۹	۰.۲۶۶۲	۰.۳۰۹۹	۰.۲۱۹۶	۰.۰۴۰۵	۰.۰۴۶۲	na	na
۵	ARIMA(۲,۱,۴)								
	Coefficients	na	na	na	na	na	۰.۲۷۰۲-	na	na
	SE	na	na	na	na	na	۰.۰۳۱۹	na	na

که فرضیه‌های آن به صورت زیر است اگر فرضیه صفر رد شود، اثر ARCH وجود دارد. سپس باید به مرحله بعدی برویم و یک مدل ARCH مناسب انتخاب کنیم.

$$H_0: \rho_1 = \rho_2 = \dots = \rho_m = 0 \quad (۸)$$

$$Q_m = N(N+2) \sum_{h=1}^m \rho_h^2 / (N-h)$$

که ρ_m^2 تابع خودهمبستگی نمونه ای و N سائز نمونه و Q_m تابعی از m و m عدد تاخیر مدل آزمون شده است. جدول زیر مقادیر p آماره Ljung-Box برای هر تاخیر تا ۱۰ را نشان می دهد. سطح معنا داری ۰.۰۵ است. تقریباً تمام مقادیر p از مقدار مجاز مربع کمتر از ۰.۰۵ است که نشان دهنده همبستگی قوی است و در نتیجه اثر ARCH وجود دارد. [۵]

جدول ۳: جدول برآورد توسط مدل ARIMA

Factor		lag ۱	lag ۲	lag ۳	lag ۴	lag ۵	lag ۶	lag ۷	lag ۸	lag ۹	lag ۱۰
۱	Residual	۰.۶۵۶	۰.۵۰۴	۰.۶۹۶	۰.۶۹۴	۰.۵۷۱	۰.۵۶۷	۰.۳۵۵	۰.۳۵۴	۰.۴۴۴	۰.۰۶۸
	residual Squared	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>
۲	Residual	۰.۵۳۹	۰.۸۱۸	۰.۹۳۶	۰.۹۲۷	۰.۹۳۱	۰.۹۶۹	۰.۶۶۷	۰.۵۸۷	۰.۱۴۱	۰.۱۶۸
	residual Squared	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>
۳	Residual	۰.۷۸۹	۰.۶۰۱	۰.۷۳۴	۰.۸۶۱	۰.۹۳۳	۰.۵۸۶	۰.۱۹	۰.۱۵۵	۰.۱۵۸	۰.۲۰۵
	residual Squared	۰.۰۱۸	۰.۰۱۸	۰.۰۱۳	۰.۰۲۸	۰.۰۱۷	۰.۰۲۴	۰.۰۳۵	۰.۰۵۸	۰.۰۸۵	۰.۱۱
۴	Residual	۰.۰۰۵	۰.۰۱۹	۰.۰۴۲	۰.۰۵۸	۰.۰۰۷	۰.۰۰۱	۰.۰۰۲	۰.۰۰۳	۰.۰۰۶	۰.۰۰۴
	residual Squared	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>
۵	Residual	۰.۰۰۲	۰.۰۰۴	۰.۰۱۱	۰.۰۲۲	۰.۰۰۴	۰.۰۰۶	۰.۰۹۴	۰.۱۳	۰.۱۸۱	۰.۱۹۱
	residual Squared	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>	۰.۰۰۱>

۶ نتیجه

در این مقاله، براساس مفهوم و روش شناسی مدل GARCH، روند مدل سازی نوسانات زمان سفر با استفاده از این مدل به طور دقیق معرفی شده است. سپس، ما از دو آزمایش مختلف بر روی سیستم حمل و نقل خودروهای شهری مختلف استفاده کردیم که در هر دو تاثیر مدل ناهمسانی واریانس شرطی پذیرفته شد. دیده می شود که اگر معادله میانگین به درستی انتخاب شود، روش پیشنهادی می تواند نتایج بسیار خوبی داشته باشد.

مراجع

- [1] Yang, M., Liu, Y. and You, Z. (2010), The reliability of travel time forecasting. *IEEE Transactions on Intelligent Transportation Systems*, **11**(1), 162-171.
- [2] Engle, R.F. (1982), Autoregressive conditional heteroscedasticity with estimates of the variance of United Kingdom inflation. *Econometrica: Journal of the Econometric Society*, 987-100
- [3] Taylor, S.J. (2008), *Modelling financial time series*. world scientific
- [4] Kamarianakis, Y., Kanas, A. and Prastacos, P. (2005), Modeling traffic volatility dynamics in an urban network. *Transportation Research Record, Journal of the Transportation Research Board*, **1923**, 18-27.
- [5] Zhang, Y., Sun, R., Haghani, A. and Zeng, X. (2013), Univariate volatility-based models for improving quality of travel time reliability forecasting. *Transportation Research Record: Journal of the Transportation Research Board*, **2365**, 73-81

برآورد پارامترهای توزیع وایبل نمایی شده به روش درستنمایی ماکسیمم و بوت استرپ بر اساس نمونه‌گیری مجموعه‌ای رتبه دار

سعید حیدری^۱، محمود افشاری، مراد علیزاده

گروه آمار، دانشگاه خلیج فارس بوشهر

چکیده: هدف اصلی تحلیل آماری استخراج حداکثر اطلاعات از داده‌ها و استنتاج ویژگی‌های جامعه است. در این مقاله ابتدا روش بوت استرپ و نمونه‌گیری مجموعه‌ای رتبه دار را معرفی و سپس برآورد درستنمایی ماکسیمم پارامترهای توزیع وایبل نمایی شده را بر اساس نمونه‌گیری مجموعه‌ای رتبه دار به‌دست می‌آوریم. در ادامه با استفاده از بوت استرپ فاصله اطمینان صدکی را برای پارامترها به‌دست می‌آوریم.

واژه‌های کلیدی: بوت استرپ، مدل تنش-مقاومت، تحلیل بقا.

کد موضوع بندی: 39B82، 34K20، 39B52، 47A55.

۱ پیش‌گفتار

افرون [۱] روش بوت استرپ را برای برآورد اندازه دقت و توزیع نمونه‌ای آماره‌ها ارائه کرد. این روش بر اساس ایده بازنمونه‌گیری از داده‌ها است که در چند سال اخیر با استفاده از کامپیوتر گسترش فراوانی یافته است. مباحث تحت بررسی بوت استرپ همان مسایل سنتی آمار هستند، تنها ابزار مورد استفاده تغییر یافته است. کامپیوترها این مباحث را قابل انعطاف، سریع، آسان و با حداقل فرضیات ریاضی ممکن می‌سازند [۲]. در آمار بوت استرپ یک متد کامپیوتری است برای نسبت دادن معیار دقت به تخمین‌های داده نمونه در این تکنیک تنها با یک روش خیلی ساده میتوان تقریباً هر آماره‌ای از توزیع داده‌های نمونه را تخمین زد. به طور عمومی این روش از روشهای بازنمونه‌گیری به حساب می‌آید. بوت استرپ در واقع ویژگیهای یک برآوردگر مانند واریانس را با استفاده از یک توزیع تقریبی از کل داده‌های نمونه تخمین میزند. یک انتخاب استاندارد برای توزیع تقریبی، توزیع تجربی داد‌های مشاهده شده است. در حالتی که بتوان فرض کرد مجموعه‌ای از مشاهده‌ها از جمعیتی مستقل و هم توزیع و به طور مساوی توزیع شده هستند، بوت استرپ میتواند با ساخت تعدادی باز نمونه پیاده سازی شود که هر کدام از این بازنمونه‌ها در واقع نمونه‌های تصادفی با جایگذاری از مجموعه داده‌های اصلی هستند. همچنین از بوت استرپ میتوان در ساخت آزمون فرض آماری استفاده کرد. از این روش معمولاً به عنوان جایگزینی برای متد‌های استنباطی بر پایه فرض‌های پارامتری هنگامی که در مورد این فرض‌ها شک داشته باشیم یا در مواردی که استنباط پارامتری غیرممکن باشد یا برای محاسبه خطای استاندارد، فرمول محاسباتی پیچیده شود از بوت استرپ استفاده میشود. یک فایده بزرگ بوت استرپ سادگی آن است. این روش برای تخمین خطای استاندارد و بازه اطمینان برای برآوردهای پیچیده پارامترهای توزیع مثل نقطه‌های صدکی، نسبت‌ها و ضرایب همبستگی بسیار ساده میباشد به علاوه روش مناسبی برای کنترل و بررسی پایداری نتایج است.

^۱ سعید حیدری: s_mheydari@yahoo.com

۲ نمونه گیری تصادفی ساده و نمونه گیری مجموعه ای رتبه دار

۱.۲ نمونه گیری تصادفی ساده

یکی از روش های متداول در نمونه گیری که برای جمع آوری اطلاعات بکار می رود، نمونه گیری تصادفی ساده است که به دو روش "بایدون جایگذاری" و "بدون جایگذاری" انجام می شود. نمونه گیری تصادفی ساده، روش انتخاب n واحد از جامعه ای به حجم N واحد است، به قسمی که همه ی $\binom{N}{n}$ نمونه ای که می توان انتخاب کرد، شانس یکسانی برای انتخاب شدن داشته باشند. در عمل، یک نمونه ی تصادفی ساده، واحد به واحد انتخاب می شود. واحدهای جامعه را از ۱ تا شماره گذاری می کنند، آنگاه با به وسیله ی جدول اعداد تصادفی و یا بوسیله ی برنامه های کامپیوتری که چنین جدولی را تولید می کنند، متوالیاً n عدد را که معرف n واحد نمونه هستند، انتخاب می نمایند. در هر انتخابی، فرآیندی که بکار می رود، باید شانس مساوی برای انتخاب هر واحد جامعه که قبلاً استخراج نشده است، فراهم کند. به راحتی می توان تحقیق کرد که تمام $\binom{N}{n}$ نمونه ی متمایز ممکن با این روش، شانس یکسان برای انتخاب شدن دارند. اگر در هر انتخاب، واحد انتخاب شده از جامعه کنار رود، روش نمونه گیری، نمونه گیری تصادفی ساده ی بدون جایگذاری نامیده می شود. اما اگر در هر انتخاب، واحد انتخاب شده، دوباره به جامعه برگردانده شود، روش نمونه گیری، نمونه گیری تصادفی ساده ی با جایگذاری نامیده می شود.

۲.۲ نمونه گیری مجموعه ای رتبه دار

این روش نمونه گیری اولین بار توسط مک اینتایر [۶] برای برآورد میانگین بازدهی های مرتع معرفی شد. در شرایطی که اندازه گیری واحدهای جامعه سخت و پرهزینه باشد اما بتوان واحدهای جامعه را به سادگی و با کمترین هزینه رتبه بندی کرد، نمونه گیری مجموعه ای رتبه دار روش های کارتری را نسبت به روش نمونه گیری تصادفی ساده برای استنباط درباره پارامترهای جامعه انجام می دهد. روش RSS به این صورت است که برای انتخاب یک نمونه ی k تایی، k نمونه k تایی از جامعه انتخاب می شود. هر نمونه ی k تایی، نسبت به متغیر مورد نظر، رتبه بندی شده و سپس از نمونه ی k تایی اول، واحد دارای کوچکترین رتبه، از نمونه ی k تایی دوم، واحد دارای دومین رتبه، ... و از نمونه ی k تایی m ، واحد دارای بزرگترین رتبه، اندازه گیری می شود. به منظور کاهش خطای رتبه بندی، اندازه ی k را معمولاً ۳ یا ۴ در نظر می گیرند. بر این اساس، جهت افزایش کارایی برآوردهای حاصل از نمونه گیری مجموعه ای رتبه دار، می توان فرآیند نمونه گیری مجموعه ای رتبه دار را چندین مرتبه تکرار کرد.

۳.۲ برآورد ماکزیمم درستنمایی پارامترهای توزیع وایبل نمایی شده با بکارگیری نمونه گیری تصادفی ساده

فرض کنید $Z_1, Z_2, \dots, Z_n$ یک نمونه ی تصادفی n تایی از توزیع EW با پارامترهای شکل (α, Y) و پارامتر مقیاسی λ باشند. پس لگاریتم تابع درستنمایی (LL) بصورت زیر خواهد بود

$$L(\alpha, \lambda, \gamma) = nLn\alpha + nLn\gamma + n\gamma \ln\lambda + (\gamma - 1) \sum Ln\zeta_i + (\alpha - 1) \quad (1)$$

جهت یافتن برآوردگر ماکزیمم درستنمایی (MLE) پارامترها، از لگاریتم تابع درستنمایی نسبت به تک تک پارامترها مشتق گرفته و روابط را مساوی صفر قرار می دهیم.

$$\frac{\partial L(\alpha, \lambda, \gamma)}{\partial \alpha} = \frac{n}{\alpha} + \sum L_n[1 - \exp(-\lambda \zeta_i)^\gamma] = 0 \quad (2)$$

$$\frac{\partial L(\alpha, \lambda, \gamma)}{\partial \lambda} = \frac{n\gamma}{\lambda} + (\alpha - 1)\gamma\lambda^{\gamma-1} \sum \frac{\exp(-\lambda \zeta_i)^\gamma}{1 - \exp(-\lambda \zeta_i)^\gamma} \zeta_i^\gamma - \gamma\lambda^{\gamma-1} \sum \zeta_i^\gamma = 0 \quad (3)$$

$$\begin{aligned} \frac{\partial L(\alpha, \lambda, \gamma)}{\partial \lambda} &= \frac{n}{\gamma} + nLn\lambda + \sum Ln\zeta_i \\ &+ (\alpha - 1)\lambda^\gamma \sum \frac{\exp(-\lambda \zeta_i)^\gamma}{1 - \exp(-\lambda \zeta_i)^\gamma} Ln(\lambda \zeta_i) - \lambda^\gamma \sum \zeta_i^\gamma Ln(\lambda \zeta_i) = 0 \quad (4) \end{aligned}$$

حل این معادلات به روشهای غیر عددی ناممکن است، بنابراین از روشهای عددی استفاده می کنیم.

۴.۲ برآورد ماکزیمم درستنمایی پارامترهای توزیع وایبل نمایی شده با بکارگیری نمونه گیری مجموعه ای رتبه دار

فرض کنید $\{Y_{ij} : i = 1, \dots, n; j = 1, \dots, k\}$ نمونه ای مجموعه ای رتبه دار به اندازه $q = kn$ از جامعه ی وایبل نمایی شده با پارامتر مقیاسی λ و پارامترهای شکل α و γ باشند و n ، سایز مجموعه و k ، تعداد چرخه ها (تکرارها) باشد. چون این روش در حقیقت از آماره های ترتیبی استفاده می کند، در این صورت تابع چگالی Y_{ij} بصورت زیر خواهد بود:

$$\begin{aligned} f_i(y_{ij}) &= \frac{n!}{(i-1)!(n-i)!} \alpha \gamma \lambda^\gamma y_{ij}^{\gamma-1} \exp(-\lambda y_{ij})^\gamma \\ &[1 - \exp(-\lambda y_{ij})^\gamma]^{i\alpha-1} * (1 - [1 - \exp(-\lambda y_{ij})^\gamma]^\alpha)^{n-i} \end{aligned}$$

بنابراین تابع درستنمایی بصورت زیر خواهد بود:

$$\begin{aligned} L(\alpha, \lambda, \gamma, Y) &= \prod_{j=1}^k \prod_{i=1}^n f_i(y_{ij}) \\ &= C \alpha^{nk} \gamma^{nk} \lambda^{nk\gamma} \prod_{j=1}^k \prod_{i=1}^n y_{ij}^{\gamma-1} \exp(-\lambda^\gamma \sum_{j=1}^k \sum_{i=1}^n y_{ij}^\gamma) \\ &* \prod_{j=1}^k \prod_{i=1}^n [1 - \exp(-\lambda y_{ij})^\gamma]^{i\alpha-1} \prod_{j=1}^k \prod_{i=1}^n (1 - [1 - \exp(-\lambda y_{ij}^{n-i})^\gamma]^\alpha) \end{aligned}$$

که C عدد ثابت است. اگر $LnL(\alpha, \lambda, \gamma, Y)$ را با L^* نشان دهیم، داریم:

$$\begin{aligned} L^* &= C^* + nk(Ln\alpha + Ln\gamma) + nk\gamma Ln\lambda \\ &+ (\gamma - 1) \sum_{j=1}^k \sum_{i=1}^n Lny_{ij} - \lambda^\gamma \sum_{j=1}^k \sum_{i=1}^n y_{ij}^\gamma + \sum_{j=1}^k \sum_{i=1}^n (i\alpha - 1) Ln[1 - \exp(-\lambda y_{ij})^\gamma] \\ &+ \sum_{j=1}^k \sum_{i=1}^n (n - i) Ln[1 - \exp(-\lambda y_{ij}^{n-i})^\gamma] \end{aligned}$$

که C^* عدد ثابت است. جهت یافتن برآوردگرهای MLE برای α ، λ و γ از L^* نسبت به این پارامترها مشتق گرفته و روابط را مساوی صفر قرار می دهیم.

$$\begin{aligned} \circ &= \frac{\partial L^*}{\partial \alpha} \\ &= \frac{nk}{\alpha} + \sum_{j=1}^k \sum_{i=1}^n i Ln[1 - \exp(-\lambda y_{ij})^\gamma] \\ &- \sum_{j=1}^k \sum_{i=1}^n (n - i) \frac{[1 - \exp(-\lambda y_{ij})^\gamma] Ln[1 - \exp(-\lambda y_{ij})^\gamma]}{(1 - [1 - \exp(-\lambda y_{ij}^{n-i})^\gamma])} \end{aligned}$$

$$\begin{aligned} \circ &= \frac{\partial L^*}{\partial \lambda} = \frac{nk\gamma}{\lambda} - \gamma \lambda^{\gamma-1} \sum_{j=1}^k \sum_{i=1}^n y_{ij}^\gamma \\ &+ \gamma \lambda^{\gamma-1} \sum_{j=1}^k \sum_{i=1}^n (i\alpha - 1) \frac{y_{ij}^\gamma \exp(-\lambda y_{ij})^\gamma}{1 - \exp(-\lambda y_{ij}^{n-i})^\gamma} \\ &- \alpha \gamma \lambda^{\gamma-1} \sum_{j=1}^k \sum_{i=1}^n (n - i) \frac{y_{ij}^\gamma \exp(-\lambda y_{ij})^\gamma [1 - \exp(-\lambda y_{ij}^{n-i})^\gamma]^{\alpha-1}}{(1 - [1 - \exp(-\lambda y_{ij}^{n-i})^\gamma])^\alpha} \end{aligned}$$

$$\begin{aligned} \circ &= \frac{\partial L^*}{\partial \gamma} \\ &= \frac{nk}{\gamma} + nk Ln\lambda + \sum_{j=1}^k \sum_{i=1}^n Lny_{ij} - \lambda^\gamma Ln\lambda \sum_{j=1}^k \sum_{i=1}^n y_{ij}^\gamma \\ &- \lambda^\gamma \sum_{j=1}^k \sum_{i=1}^n y_{ij}^\gamma Lny_{ij} \\ &+ \lambda^\gamma \sum_{j=1}^k \sum_{i=1}^n (i\alpha - 1) \frac{y_{ij}^\gamma Ln\lambda y_{ij} \exp(-\lambda y_{ij})^\gamma}{[1 - \exp(-\lambda y_{ij}^{n-i})^\gamma]} \\ &- \alpha \lambda^\gamma \sum_{j=1}^k \sum_{i=1}^n (n - i) \frac{y_{ij}^\gamma \exp(-\lambda y_{ij})^\gamma Ln\lambda y_{ij} [1 - \exp(-\lambda y_{ij}^{n-i})^\gamma]^{\alpha-1}}{(1 - [1 - \exp(-\lambda y_{ij}^{n-i})^\gamma])^\alpha} \end{aligned}$$

حل این معادلات به روشهای غیر عددی، ناممکن است. بنابراین از روشهای عددی استفاده می کنیم.

۳ فاصله اطمینان بوت استرپ برای R

در این بخش، دو فاصله اطمینان مبتنی بر روش‌های پارامتری بوت استرپ معرفی می‌شوند: روش صدکی بوت استرپ که از این به بعد boot-p خوانده می‌شود و بر پایه‌ی نظریه افرون [۲] می‌باشد. روش استیودنتیده بوت استرپ که boot-t خوانده می‌شود و بر پایه نظریه‌ی هال (۱۹۸۸) می‌باشد به طور خلاصه برآورد فاصله اطمینان برای R با هر دو روش را نشان می‌دهیم:

۱.۳ فاصله اطمینان صدکی بوت استرپ

الگوریتم برآورد فاصله اطمینان صدکی بوت استرپ به صورت زیر است:

گام اول: با استفاده از نمونه $Y_1, \dots, Y_n$ و $X_1, \dots, X_n$ مقادیر $\hat{\alpha}$ و $\hat{\beta}$ را محاسبه می‌کنیم.

گام دوم: با استفاده از $\hat{\alpha}$ و $\hat{\beta}$ نمونه‌ی بوت استرپ $X_1, \dots, X_n$ و به طور مشابه با استفاده از مقادیر $\hat{\beta}$ و $\hat{\alpha}$ نمونه‌ی بوت استرپ $Y_1, \dots, Y_n$ را تولید می‌کنیم. حال با استفاده از $X_1^*, \dots, X_n^*$ و $Y_1^*, \dots, Y_n^*$ برآورد بوت استرپ را برای R بدست آورده و با $\hat{R}^*$ نشان می‌دهیم.

گام سوم: گام دوم را N بار تکرار می‌کنیم.

گام چهارم: فرض کنید تابع توزیع تجمعی $\hat{R}^*$ به صورت $G_X(x) = P(\hat{R}^* \leq x)$ تعریف شود، حال اگر قرار دهیم: $G^{-1}(x) = R^{boot-p}(x)$. در این صورت برای هر x فاصله اطمینان $(1 - \gamma) \times 100$ برای R به صورت زیر تعریف می‌شود.

$$(\hat{R}_{boot-p}(\frac{\gamma}{2})), (\hat{R}_{boot-p}(1 - \frac{\gamma}{2}))$$

۲.۳ فاصله اطمینان استیودنتیده بوت استرپ برای R

الگوریتم برآورد فاصله اطمینان استیودنتیده بوت استرپ به صورت زیر می‌باشد:

گام اول: با استفاده از نمونه $Y_1, \dots, Y_n$ و $X_1, \dots, X_n$ مقادیر $\hat{\alpha}$ و $\hat{\beta}$ را محاسبه می‌کنیم.

گام دوم: با استفاده از $\hat{\alpha}$ و $\hat{\beta}$ نمونه‌ی بوت استرپ $X_1, \dots, X_n$ و به طور مشابه با استفاده از مقادیر $\hat{\beta}$ و $\hat{\alpha}$ نمونه‌ی بوت استرپ $Y_1, \dots, Y_n$ را تولید می‌کنیم. حال با استفاده از $X_1^*, \dots, X_n^*$ و $Y_1^*, \dots, Y_n^*$ برآورد بوت استرپ را برای R بدست آورده و با $\hat{R}^*$ نشان می‌دهیم و آماره زیر را تعریف می‌کنیم:

$$T^* = \frac{\sqrt{m}(\hat{R}^* - \hat{R})}{\sqrt{var(\hat{R}^*)}}$$

گام سوم: گام دوم را N بار تکرار می‌کنیم.

گام چهارم: با استفاده از T^* که به تعداد N بار بدست آمده حدود بالا و پایین را برای فاصله اطمینان $(1 - \gamma) \times 100$ برای R به صورت زیر تعیین می‌کنیم:

تابع توزیع تجمعی T^* را به صورت $H(x) = P(T^* \leq x)$ در نظر می‌گیریم و برای هر x تعریف می‌کنیم:

$$\hat{R}_{boot-t} = \hat{R} + m^{-\frac{1}{2}} \sqrt{var(\hat{R})} H^{-1}(x)$$

فاصله اطمینان تقریبی $(1 - \gamma)100$ برای R به صورت زیر تعریف می‌شود:

$$(\hat{R}_{boot-t}(\frac{\gamma}{2})), (\hat{R}_{boot-t}(1 - \frac{\gamma}{2}))$$

۴ شبیه سازی

در این قسمت، با استفاده از نرم افزار *Maple* نمونه‌ی تصادفی ساده‌ی Y تا 20 از توزیع وایبل نمایی شده با پارامترهای $\alpha = 1/5$ ، $\gamma = 1/5$ و $\lambda = 1/5$ تولید می‌کنیم. سپس معادلات ۳، ۴ و ۵ را به صورت یک دستگاه سه معادله‌ای سه مجهولی در نرم افزار *Maple* وارد می‌کنیم، آنگاه از طریق تابع *fsolve*، ریشه‌های دستگاه را که $\hat{\alpha}$ ، $\hat{\gamma}$ و $\hat{\lambda}$ می‌باشند، می‌یابیم. لازم به ذکر است که تابع *fsolve*، برای حل دستگاه معادلات از انواع روشهای عددی استفاده می‌کند. تعداد تکرارها برای شبیه سازی، 1000 در نظر گرفته شده است.

جدول ۱: میزان اریبی و میانگین مربعات خطای حداکثر درست‌نمایی R به ازای α

$\alpha = 1/5$			
<i>SRS</i>			
کارایی	λ معلوم	λ مجهول	λ معلوم
$1/0.059$	$1/0.793$	$0/0.730(0/0.167)$	$0/0.597(0/0.126)$
$1/6282$	$1/1382$	$0/0.050(0/0.078)$	$0/0.394(0/0.094)$
$1/6944$	$1/7187$	$0/0.050(0/0.078)$	$0/0.202(0/0.032)$
$1/0.869$	$1/0.222$	$0/0.011(0/0.046)$	$0/0.011(0/0.045)$
<i>RSS</i>			
λ مجهول	λ معلوم	(n, m)	
$0/0.137(0/0.168)$	$0/0.058(0/0.136)$	$(5, 15)$	
$0/0.048(0/0.127)$	$0/0.039(0/0.107)$	$(10, 10)$	
$0/0.068(0/0.061)$	$0/0.024(0/0.055)$	$(15, 30)$	
$0/0.045(0/0.050)$	$0/0.006(0/0.046)$	$(20, 30)$	

۱.۴ نتایج شبیه‌سازی

در این بخش شبیه سازی برآورد حداکثر درست‌نمایی R با استفاده از بوت استرپ تحت شرایط مختلف پارامتری بررسی می‌گردد. برآورد حداکثر درست‌نمایی R زمانی که λ مجهول است

اریبی $\frac{1}{el} \sum_{i=1}^{el} (R^* - \hat{R}_i^*)^2$ و میانگین مربعات خطا $\frac{1}{el} \sum_{i=1}^{el} (R^* - \hat{R}_i^*)^2$ در برآورد حداکثر درست‌نمایی به ازای $n, m = 15, 20, 25, 30$ و $\alpha = 1/5, 2/5, 3/5$ و همچنین برای $\beta = 1/5, 2/5, 3/5$ و $\alpha = 1/5$ و $\beta = 1/5$ در جداول ۱ و ۲ ارائه شده است. بدون اینکه خدشه‌ای در شبیه سازی وارد شود لامبدا و دلتا را برابر یک در نظر گرفته شده اند. مقادیر

جدول ۲: میزان اریبی و میانگین مربعات خطای حداکثر درستمایی R به ازای β

$\alpha = 1/5$			
SRS		کارایی	
λ معلوم	λ مجهول	λ معلوم	λ مجهول
۰/۰۷۰۹(۰/۰۱۴۲)	۰/۰۶۹۵(۰/۰۱۵۸)	۱/۱۱۲۶	۱/۲۱۵۱
۰/۰۱۲۳(۰/۰۱۱۹)	۰/۰۰۲۱(۰/۰۰۹۹)	۱	۱/۴۲۴۲
۰/۰۲۶۱(۰/۰۰۴۷)	۰/۰۳۳۵(۰/۰۰۵۳)	۱/۳۴۰۴	۱/۳۰۱۸
۰/۰۰۳۳(۰/۰۰۴۳)	۰/۰۰۵۸(۰/۰۰۴۱)	۱/۲۰۹۳	۱/۳۶۵۸
(n, m)		RSS	
		λ معلوم	λ مجهول
(۵, ۱۵)		۰/۰۱۴۰(۰/۰۱۵۸)	۰/۰۱۶۸(۰/۰۱۹۲)
(۱۰, ۱۰)		۰/۰۰۱۷(۰/۰۱۱۹)	۰/۰۰۱۹(۰/۰۱۴۱)
(۱۵, ۳۰)		۰/۰۰۵۷(۰/۰۰۶۳)	۰/۰۰۵۹(۰/۰۰۶۹)
(۲۰, ۳۰)		۰/۰۰۳۳(۰/۰۰۵۲)	۰/۰۰۳۲(۰/۰۰۵۶)

میانگین مربعات خطا در داخل پارانتز نمایش داده شده اند.

برآورد حداکثر درستمایی λ به روش تکرار نقطه ثابت با مقدار اولیه ۱ محاسبه می شود و روند تکرار زمانی متوقف میشود که اختلاف دو تکرار متوالی کوچکتر از ۱۰-۶ باشد. پس از برآورد λ ، برآورد α ، β و γ و در نهایت برآورد حداکثر درستمایی R از رابطه $\hat{R} = \frac{\hat{\beta}}{\hat{\beta} + \hat{\alpha}}$ محاسبه می شود.

با توجه به جداول ۱ و ۲ مشاهده میشود که حتی برای اندازه نمونه کوچک، کارایی MLE بر اساس اریبی و میانگین مربعات خطا در هر دو روش رضایت بخش است. به ازای $n = m$ زمانی که n و m افزایش میابد، میانگین مربعات خطا کاهش میابد. این ویژگی سازگاری MLE پارامتر R را تایید می کند. همچنین به ازای $n(m)$ ثابت، با افزایش $n(m)$ میانگین مربعات خطا کاهش می یابد. همچنین در جدول ۱ مشاهده می شود که با افزایش β میانگین مربعات خطای برآورد حداکثر درستمایی R کاهش می یابد. در نتیجه می توان گفت برآورد R به ازای مقادیر بزرگتر β دقیق تر است. سایر پارامترهای مجهول نیز به روش مشابه محاسبه می شود و نیز با توجه به جداول مشخص میشود که با افزایش $m(n)$ میانگین مربعات خطا کاهش می یابد.

مراجع

- [1] Efron, B. (1979), Bootstrap Methods: Another Look At the Jackknife, *Ann. of Stat.*, **26**, 1-7.
- [2] Efron, B. (1982), *The jackknife, the bootstrap, and other resampling plans*, Philadelphia, Pa. :Society for Industrial and Applied Mathematics.

- [3] Efron, B. and Tibshirani R.J. (1993), *An introduction to the Bootstrap*, Chapman & Hall/CRC.
- [4] Dixit U.J. and Nasiri, P.F. (2001), Estimation of parameters of the exponential distribution in the presence of outliers generated from uniform distribution, *Metron*, **49**, 187-198.
- [5] Hall, P. (1988), Theoretical Comparison of Bootstrap Confidence Intervals, *The annals of statistics*, **16**, 927-953.
- [6] McIntyre, G.A. (1952), A method for unbiased selective sampling using ranked sets, *Australian Journal of Agricultural Research*, **3**, 385-390.

بررسی ویژگی‌های قابلیت اعتماد در یک مدل نرخ خطر متناسب گسسته اصلاح شده

محدثه خیاط^۱، رسول روزگار^۱، قباد برمال‌زن^۲

^۱ گروه آمار، دانشکده علوم ریاضی، دانشگاه یزد

^۲ گروه آمار، دانشکده علوم ریاضی و آمار، دانشگاه زابل

چکیده: مدل نرخ خطر به عنوان یکی از خانواده‌های انعطاف‌پذیر در قابلیت اعتماد توسط بالاکریشن و همکاران معرفی شده است. در این مقاله حالت گسسته این مدل در نظر گرفته می‌شود و به خواص سالخورده‌گی و حفظ شدن ترتیب تصادفی معمولی، نرخ خطر و نسبت درست‌نمایی در این خانواده از توزیع‌ها پرداخته شده است.

واژه‌های کلیدی: ترتیب تصادفی معمولی، ترتیب نرخ خطر، مدل نرخ خطر گسسته.

کد موضوع بندی: 39B82, 34K20, 39B52, 47A55.

۱ پیش‌گفتار

مارشال و ال‌کین (۱۹۹۷) با اضافه کردن یک پارامتر شکل به توزیع، یک خانواده کلی و انعطاف‌پذیر از توزیع‌ها را معرفی کردند. فرض کنید $\bar{F}$ بیانگر یک تابع بقای دلخواه روی $\mathbb{R}^+$ باشد. در این صورت خانواده توزیع‌های زیر، به خانواده توزیع‌های مارشال-ال‌کین معروف است:

$$\bar{G}(x; \alpha) = \frac{\alpha \bar{F}(x)}{1 - \alpha \bar{F}(x)}, \quad x \in \mathbb{R}^+, \bar{\alpha} = 1 - \alpha, \alpha > 0. \quad (1)$$

برای مشاهده جزئیات بیشتر در مورد خانواده توزیع‌های مارشال-ال‌کین به کرمانی و گوپتا (۲۰۰۱) و ناند و داس (۲۰۱۳) مراجعه نمایید.

اگر در تابع بقای $\bar{G}$ در رابطه (۱)، به جای $\bar{F}(x)$ از خانواده کلی $\bar{F}^\lambda(x)$ استفاده شود یک خانواده کلی و انعطاف‌پذیر از توزیع‌ها تحت عنوان مدل نرخ خطر متناسب اصلاح شده حاصل می‌شود (بالاکریشن و همکاران (۲۰۱۸) ببینید). نتایجی که تاکنون در رابطه با این خانواده از توزیع‌ها به دست آمده است با فرض پیوسته بودن تابع بقای $\bar{F}^\lambda(x)$ بوده است. در این مقاله، هدف بررسی و مطالعه خواص سالخورده‌گی و مقایسه تصادفی در این خانواده از توزیع‌ها با فرض گسسته بودن $\bar{F}^\lambda(x)$ است. در عالم واقعیت، مواردی وجود دارند که در آن متغیرهای طول عمر، به صورت گسسته در نظر گرفته می‌شوند. چندین مثال از این موارد را می‌توان به صورت زیر بیان کرد:

الف) گزارش‌هایی در مورد خرابی قطعات یا سیستم‌ها که به صورت هفته‌ای یا ماهیانه جمع‌آوری می‌شوند که در آن به جای زمان شکست‌ها، تعداد شکست‌ها ثبت شده است.

ب) قطعات یا سیستم‌هایی که به صورت گردشی کار می‌کنند و پژوهشگر تعداد گردش‌هایی که به طور کامل تا قبل از خرابی انجام می‌شود را ثبت می‌کند. مانند تعداد کل کپی‌های گرفته شده از یک دستگاه فتوکپی تا قبل از خراب شدن.

(ج) به دلیل سهولت در امر اندازه‌گیری، گاهی اوقات یک متغیر تصادفی پیوسته با یک متغیر تصادفی گسسته تقریب زده می‌شود. مانند تعداد کیلومترهایی که یک اتومبیل تا قبل از سائیدگی لاستیک‌های آن، پیموده است.

مطالعه متغیرهای طول عمر گسسته دیرتر از متغیرهای پیوسته آغاز گردید. مطالعه متغیرهای تصادفی گسسته، در ابتدا توسط بارلو و پروشان (۱۹۸۱) شروع شده و سپس توسط سایر نویسندگان مورد بحث و بررسی قرار گرفت.

فرض کنید متغیر تصادفی گسسته X با تکیه‌گاه $\mathbb{N} = \{1, 2, \dots\}$ دارای تابع جرم احتمال $P(X = k)$ و تابع بقای $\bar{F}(k) = P(X > k)$ باشد. در این صورت، متغیر تصادفی Y دارای توزیع نرخ خطر متناسب گسسته اصلاح شده با پارامترهای α و λ است $(Y \sim MDPHR(\alpha, \lambda))$ ، هرگاه دارای تابع بقا

$$\bar{H}(k) = \frac{\alpha \bar{F}^\lambda(k)}{1 - \alpha \bar{F}^\lambda(k)}, \quad k \in \mathbb{N}, \quad \lambda > 0, \quad \alpha > 0, \quad \bar{\alpha} = 1 - \alpha.$$

و تابع جرم احتمال

$$\begin{aligned} P(Y = k) &= \bar{H}(k-1) - \bar{H}(k) \\ &= \frac{\alpha(\bar{F}^\lambda(k-1) - \bar{F}^\lambda(k))}{(1 - \alpha \bar{F}^\lambda(k-1))(1 - \alpha \bar{F}^\lambda(k))}, \quad k \in \mathbb{N}, \end{aligned}$$

باشد که در آن $\alpha > 0, \lambda > 0$ است. در این توزیع، α به عنوان پارامتر تیل و λ به عنوان پارامتر شکل است. تابع نرخ خطر این توزیع عبارت است از:

$$\begin{aligned} r_Y(k) &= \frac{P(Y = k)}{\bar{H}(k-1)} \\ &= \frac{\bar{F}^\lambda(k-1) - \bar{F}^\lambda(k)}{\bar{F}^\lambda(k-1)(1 - \alpha \bar{F}^\lambda(k))}, \quad k \in \mathbb{N}. \end{aligned}$$

نتایج بیان شده در این مقاله در دو راستا است: ابتدا به بررسی حفظ شدن و عدم حفظ شدن خواص سالخوردگی متغیر اولیه X با تابع بقای $\bar{F}(k)$ ، نسبت به متغیر ثانویه Y با تابع بقای $\bar{H}(k)$ پرداخته شده است. سپس مقایسه تصادفی بین متغیرهای اولیه و ثانویه، بررسی شده است. به بیان دقیق‌تر، فرض کنید X_1 و X_2 دو متغیر تصادفی مستقل با توابع بقای به ترتیب $\bar{F}_1(k)$ و $\bar{F}_2(k)$ باشند. همچنین فرض کنید Y_1 و Y_2 دو متغیر تصادفی با توابع بقاهای به ترتیب

$$\bar{H}_1(k; \lambda_1) = \frac{\alpha \bar{F}_1^{\lambda_1}(k)}{1 - \alpha \bar{F}_1^{\lambda_1}(k)}, \quad k \in \mathbb{N}$$

و

$$\bar{H}_2(k; \lambda_2) = \frac{\beta \bar{F}_2^{\lambda_2}(k)}{1 - \beta \bar{F}_2^{\lambda_2}(k)}, \quad k \in \mathbb{N},$$

باشند. اگر $\alpha \geq \beta$ و $\lambda_1 \leq \lambda_2$ باشد آنگاه نشان داده می‌شود:

$$X_1 \geq_{st} X_2 \implies Y_1 \geq_{st} Y_2.$$

همچنین، اگر $\lambda_1 = \lambda_2$ و $\alpha \geq \beta$ آنگاه نتیجه زیر حاصل شده است:

$$X_1 \geq_{hr} X_2 \implies Y_1 \geq_{hr} Y_2,$$

و با دو مثال نقض نشان داده شده است که جهت ترتیب نسبت درستنمایی بین X_1 و X_2 ، توسط متغیرهای تصادفی Y_1 و Y_2 حفظ نمی‌شود. فرض کنید X و Y دو متغیر تصادفی گسسته با تکیه‌گاه‌های مشترک $\mathbb{N} = \{1, 2, \dots\}$ باشند که دارای توابع احتمال $P(X = k)$ ، $P(Y = k)$ و توابع بقای $\bar{F}(k) = P(X > k)$ ، $\bar{G}(k) = P(Y > k)$ و توابع نرخ خطر $r_X(k) = P(X = k)/P(X \geq k)$ و $r_Y(k) = P(Y = k)/P(Y \geq k)$ باشند. (شیکد و شانتیکومار، ۲۰۰۷). متغیر تصادفی X در ترتیب تصادفی معمولی^۴، بزرگتر از متغیر تصادفی Y است ($X \geq_{st} Y$)، هرگاه به ازای هر $k \in \mathbb{N}$ ، رابطه $\bar{F}(k) \geq \bar{G}(k)$ برقرار باشد. یا به عبارت دیگر، اگر برای هر تابع صعودی $\phi: \mathbb{R}^n \rightarrow \mathbb{R}$ و با شرط وجود امیدهای ریاضی، نابرابری زیر برقرار باشد:

$$\mathbb{E}(\phi(X)) \geq \mathbb{E}(\phi(Y)).$$

متغیر تصادفی X در ترتیب نرخ خطر^۵، بزرگتر از Y است ($X \geq_{hr} Y$)، هرگاه به ازای هر $k \in \mathbb{N}$ ، رابطه $r_Y(k) \geq r_X(k)$ برقرار باشد. به عبارت دیگر، $X \geq_{hr} Y$ است اگر و فقط اگر نسبت $\bar{F}(k)/\bar{G}(k)$ تابعی صعودی برحسب k باشد. رابطه زیر بین ترتیب‌های تصادفی فوق، برقرار است:

$$X \geq_{hr} Y \implies X \geq_{st} Y.$$

برای آگاهی بیشتر از جزئیات ترتیب‌های تصادفی، می‌توان به شیکد و شانتیکومار (۲۰۰۷) مراجعه نمود.

(الف) متغیر تصادفی X دارای خاصیت نسبت درستنمایی صعودی (ILR) است هرگاه نامساوی زیر برقرار باشد:

$$P(X = k + 2)P(X = k) \leq P^2(X = k + 1), \quad k \in \mathbb{N}. \quad (2)$$

به همین ترتیب، متغیر تصادفی X دارای خاصیت نسبت درستنمایی نزولی (DLR) است هرگاه جهت نامساوی (۲) عوض شود.

(ب) متغیر تصادفی X دارای خاصیت نرخ خطر صعودی (IFR) است هرگاه $r(k)$ تابعی صعودی برحسب $k \in \mathbb{N}$ باشد. یا بطور معادل، نسبت $\bar{F}(k+1)/\bar{F}(k)$ تابعی نزولی از k باشد.

به همین ترتیب، متغیر تصادفی X دارای خاصیت نرخ خطر نزولی (DFR) است هرگاه $r(k)$ تابعی نزولی برحسب $k \in \mathbb{N}$ باشد. یا بطور معادل، نسبت $\bar{F}(k+1)/\bar{F}(k)$ تابعی صعودی از k باشد.

(پ) متغیر تصادفی X دارای خاصیت میانگین نرخ خطر صعودی ($IFRA$) است هرگاه $[\bar{F}(k)]^{1/k}$ تابعی نزولی در k باشد. یا به عبارت دیگر، $[\bar{F}(k)]^{1/k} \geq [\bar{F}(k+1)]^{1/(k+1)}$.

به همین ترتیب، متغیر تصادفی X دارای خاصیت میانگین نرخ خطر نزولی ($DFRA$) است هرگاه $[\bar{F}(k)]^{1/k}$ تابعی صعودی در k باشد. یا به عبارت دیگر، $[\bar{F}(k)]^{1/k} \leq [\bar{F}(k+1)]^{1/(k+1)}$ باشد.

Usual stochastic order^۴

Hazard rate order^۵

(ت) متغیر تصادفی X دارای خاصیت نو بهتر از کارکرده (NBU) است هرگاه

$$\bar{F}(j+k) \leq \bar{F}(j)\bar{F}(k) \quad , \quad j, k \in \mathbb{N}. \quad (۳)$$

به همین ترتیب، متغیر تصادفی X دارای خاصیت نو بدتر از کارکرده (NWU) است هرگاه جهت نامساوی (۳) عوض شود.

(الف) اگر متغیر تصادفی X دارای خاصیت IFR باشد آنگاه برای $\alpha \geq ۱$ توزیع Y نیز دارای خاصیت IFR است.

(ب) اگر متغیر تصادفی X دارای خاصیت DFR باشد آنگاه برای $\alpha \leq ۱$ توزیع Y نیز دارای خاصیت DFR است.

اثبات.

(الف) تابع نرخ خطر متغیر تصادفی Y به صورت زیر است:

$$\begin{aligned} r_Y(k) &= \frac{P(Y=k)}{\bar{H}(k-۱)} \\ &= \frac{\bar{F}^\lambda(k-۱) - \bar{F}^\lambda(k)}{\bar{F}^\lambda(k-۱)(۱ - \bar{\alpha}\bar{F}^\lambda(k))} \quad , \quad \alpha > 0, \lambda > 0. \end{aligned}$$

برای اثبات صعودی بودن نرخ خطر Y ، کافی است نشان داده شود:

$$r_Y(k) \leq r_Y(k+۱) \quad , \quad k = ۱, ۲, \dots$$

یا به عبارت دیگر،

$$\frac{\bar{F}^\lambda(k-۱) - \bar{F}^\lambda(k)}{\bar{F}^\lambda(k-۱)(۱ - \bar{\alpha}\bar{F}^\lambda(k))} \leq \frac{\bar{F}^\lambda(k) - \bar{F}^\lambda(k+۱)}{\bar{F}^\lambda(k)(۱ - \bar{\alpha}\bar{F}^\lambda(k+۱))}.$$

با استفاده از این واقعیت که $\bar{F}(k) > \bar{F}(k+۱)$ و اینکه $۱ - \alpha \leq \bar{\alpha}$ است می‌توان نتیجه گرفت $۱/(۱ - \bar{\alpha}\bar{F}^\lambda(k))$ تابعی صعودی بر حسب k است. بنابراین برای $\alpha \geq ۱$ داریم:

$$\frac{۱}{۱ - \bar{\alpha}\bar{F}^\lambda(k)} \leq \frac{۱}{۱ - \bar{\alpha}\bar{F}^\lambda(k+۱)}. \quad (۴)$$

از طرف دیگر، چون متغیر تصادفی X دارای خاصیت نرخ خطر صعودی است مطابق بند (ب) تعریف ۱، به ازای $k = ۱, ۲, \dots$ داریم:

$$\frac{\bar{F}(k)}{\bar{F}(k-۱)} \geq \frac{\bar{F}(k+۱)}{\bar{F}(k)} \quad \text{یا} \quad \bar{F}^\lambda(k) \geq \bar{F}(k-۱)\bar{F}(k+۱)$$

با به توان λ رساندن کمیت $\bar{F}^\lambda(k) \geq \bar{F}(k-۱)\bar{F}(k+۱)$ و کم کردن $\bar{F}^\lambda(k)\bar{F}^\lambda(k-۱)$ از طرفین آن، داریم:

$$\bar{F}^\lambda(k)\bar{F}^\lambda(k-۱) - \bar{F}^{2\lambda}(k) \leq \bar{F}^\lambda(k)\bar{F}^\lambda(k-۱) - \bar{F}^\lambda(k-۱)\bar{F}^\lambda(k+۱),$$

که نتیجه می‌دهد

$$\frac{\bar{F}^\lambda(k-۱) - \bar{F}^\lambda(k)}{\bar{F}^\lambda(k-۱)} \leq \frac{\bar{F}^\lambda(k) - \bar{F}^\lambda(k+۱)}{\bar{F}^\lambda(k)} \quad (۵)$$

اکنون با ضرب طرفین روابط نامنفی ۱۴ و ۱۵، نتیجه مطلوب حاصل می‌شود.

(ب) اثبات این بند، همانند اثبات بند (الف) است. لذا از آوردن اثبات آن خودداری می‌شود.

□

مثال زیر نشان می‌دهد که به ازای $\alpha < 1$ خاصیت IFR متغیر X ، ممکن است به متغیر Y انتقال پیدا نکند. (سلویا و بولینگر، ۱۹۸۲). متغیر تصادفی گسسته X را با خاصیت نرخ خطر صعودی به صورت زیر در نظر بگیرید:

$$P(X = k) = \frac{(k - c) c^{k-1}}{k!}, \quad 0 < c \leq 1, \quad k \in \mathbb{N}.$$

با توجه به اینکه تابع بقای X به صورت $\bar{F}(k) = c^k/k!$ است در نتیجه تابع نرخ خطر متغیر تصادفی Y عبارت است از:

$$r_Y(k) = \frac{\left(\frac{c^{k-1}}{(k-1)!}\right)^\lambda - \left(\frac{c^k}{k!}\right)^\lambda}{\left(\frac{c^{k-1}}{(k-1)!}\right)^\lambda (1 - \bar{\alpha} \left(\frac{c^k}{k!}\right)^\lambda)}.$$

به ازای $\lambda = 2$ ، $\alpha = 0/1$ و $c = 9$ داریم:

$$r_Y(2) = 0/9356, \quad r_Y(3) = 0/9222, \quad r_Y(4) = 0/9500,$$

که نشان می‌دهد تابع نرخ خطر متغیر تصادفی Y نه صعودی است و نه نزولی. مثال زیر نشان می‌دهد که به ازای $\alpha > 1$ ، خاصیت DFR متغیر X ، ممکن است توسط متغیر Y حفظ نشود.

(ناکاگایا و اوساکی، ۱۹۷۵). فرض کنید متغیر تصادفی X دارای توزیع وایبول گسسته به صورت زیر باشد:

$$P(X = k) = q^{(k-1)\beta} - q^{k\beta}, \quad 0 < q \leq 1, \quad \beta > 0, \quad k \in \mathbb{N}.$$

تابع بقا و تابع نرخ خطر متناظر با متغیر تصادفی X به ترتیب عبارتند از:

$$\bar{F}(k) = q^{k\beta}, \quad r_X(k) = 1 - q^{k\beta - (k-1)\beta}.$$

می‌توان نشان داد برای $0 < \beta < 1$ ، این خانواده از توزیع‌ها دارای خاصیت DFR است. همچنین تابع نرخ خطر متغیر تصادفی Y عبارت است از:

$$r_Y(k) = \frac{q^{\lambda(k-1)\beta} - q^{\lambda k\beta}}{q^{\lambda(k-1)\beta} (1 - \bar{\alpha} q^{\lambda k\beta})}. \quad (6)$$

به ازای $\lambda = 2$ ، $\alpha = 5$ ، $q = 0/5$ و $\beta = 0/8$ داریم:

$$r_Y(2) = 0/4728, \quad r_Y(4) = 0/5458, \quad r_Y(13) = 0/4847,$$

که نشان می‌دهد تابع نرخ خطر متغیر تصادفی Y نه صعودی است و نه نزولی.

فرض کنید X دارای خاصیت NBU باشد آنگاه برای $\alpha \geq 1$ ، متغیر تصادفی Y دارای خاصیت NBU است. فرض کنید

X دارای خاصیت NWU باشد آنگاه برای $\alpha \leq 1$ ، متغیر تصادفی Y دارای خاصیت NWU است.

اثبات.

(الف) فرض کنید X دارای خاصیت NBU است. اگر قرار باشد Y نیز دارای خاصیت NBU باشد باید نابرابری زیر برقرار باشد:

$$\frac{\alpha \bar{F}^\lambda(j+k)}{1 - \alpha \bar{F}^\lambda(j+k)} \leq \frac{\alpha \bar{F}^\lambda(j)}{1 - \alpha \bar{F}^\lambda(j)} \frac{\alpha \bar{F}^\lambda(k)}{1 - \alpha \bar{F}^\lambda(k)}$$

یا به عبارت دیگر،

$$\frac{1 - \alpha \bar{F}^\lambda(j+k)}{\bar{F}^\lambda(j+k)} \geq \frac{(1 - \alpha \bar{F}^\lambda(j))(1 - \alpha \bar{F}^\lambda(k))}{\alpha \bar{F}^\lambda(j) \bar{F}^\lambda(k)}. \quad (7)$$

با اضافه کردن α به طرفین رابطه اخیر و ساده کردن محاسبات، می‌توان مشاهده کرد که رابطه (7)، معادل با رابطه زیر است:

$$\frac{1}{\bar{F}^\lambda(j+k)} \geq \frac{1 - \alpha(\bar{F}^\lambda(j) + \bar{F}^\lambda(j)) + \alpha \bar{F}^\lambda(j) \bar{F}^\lambda(k)}{\alpha \bar{F}^\lambda(j) \bar{F}^\lambda(k)}. \quad (8)$$

بنابراین برای اثبات NBU بودن متغیر تصادفی Y ، کافی است رابطه (8) اثبات شود.

چون X دارای خاصیت NBU است $\bar{F}(j) \bar{F}(k) \geq \bar{F}(j+k)$. یا به عبارت دیگر،

$$\frac{1}{\bar{F}^\lambda(j+k)} \geq \frac{1}{\bar{F}^\lambda(j) \bar{F}^\lambda(k)}. \quad (9)$$

از طرف دیگر، ادعا می‌کنیم که به ازای $\alpha \geq 1$ ، نامساوی زیر برقرار است:

$$1 - \alpha(\bar{F}^\lambda(j) + \bar{F}^\lambda(j)) + \alpha \bar{F}^\lambda(j) \bar{F}^\lambda(k) \leq \alpha \quad (10)$$

این ادعا با حد گرفتن از طرفین رابطه (10) نسبت به j و k و میل دادن آنها به سمت بینهایت، قابل حصول است. اکنون با استفاده روابط (9) و (10) می‌توان نتیجه گرفت:

$$\begin{aligned} \frac{1}{\bar{F}^\lambda(j+k)} &\geq \frac{1}{\bar{F}^\lambda(j) \bar{F}^\lambda(k)} \\ &= \frac{1 - \alpha(\bar{F}^\lambda(j) + \bar{F}^\lambda(j)) + \alpha \bar{F}^\lambda(j) \bar{F}^\lambda(k)}{\{1 - \alpha(\bar{F}^\lambda(j) + \bar{F}^\lambda(j)) + \alpha \bar{F}^\lambda(j) \bar{F}^\lambda(k)\} \bar{F}^\lambda(j) \bar{F}^\lambda(k)} \\ &\geq \frac{1 - \alpha(\bar{F}^\lambda(j) + \bar{F}^\lambda(j)) + \alpha \bar{F}^\lambda(j) \bar{F}^\lambda(k)}{\alpha \bar{F}^\lambda(j) \bar{F}^\lambda(k)}, \end{aligned}$$

و لذا نتیجه مطلوب حاصل می‌شود.

(ب) اثبات این بند، همانند اثبات بند (الف) است و لذا از آوردن اثبات آن خودداری می‌شود.

□

مثال زیر نشان می‌دهد که برای $\alpha < 1$ خاصیت NBU بودن متغیر تصادفی X ، ممکن است به متغیر تصادفی Y انتقال پیدا نکند. (براکومند و گادوین، ۲۰۰۳). فرض کنید متغیر تصادفی X دارای توزیع گسسته S ، با تابع جرم احتمال زیر باشد:

$$P(X = k) = p(1-a)^k \prod_{i=1}^{k-1} (1-p+pa^i) \quad , \quad 0 < p \leq 1, \quad 0 < a < 1, \quad k \in \mathbb{N}.$$

تابع بقا و تابع نرخ خطر متناظر با این توزیع به ترتیب عبارتند از:

$$\bar{F}(k) = \prod_{i=1}^k (1 - p + pa^i) \quad , \quad r_X(k) = p(1 - a)^k.$$

تابع بقای توزیع Y به صورت زیر است:

$$\bar{H}(k) = \frac{\alpha \prod_{i=1}^k (1 - p + pa^i)^\lambda}{1 - \bar{\alpha} \prod_{i=1}^k (1 - p + pa^i)^\lambda}.$$

برای $\lambda = 2$ و $\alpha = 0.2$, $a = 0.6$, $p = 0.3$, $k = 3$, $j = 2$ داریم:

$$\bar{H}(j+k) = 0.001433 \leq 0.00121 = \bar{H}(j)\bar{H}(k),$$

و در نتیجه خاصیت NBU بودن متغیر تصادفی Y نقض می‌شود.

در این دو قضیه، حفظ شدن مقایسه تصادفی بین X_i ها و Y_i ها از لحاظ ترتیب تصادفی معمولی، ترتیب نرخ خطر و ترتیب نسبت درستنمایی، مورد مطالعه قرار می‌گیرد.

فرض کنید X_1 و X_2 دو متغیر تصادفی مستقل با توابع بقای به ترتیب $\bar{F}_1(k)$ و $\bar{F}_2(k)$ باشند. همچنین فرض کنید Y_1 و Y_2 دو متغیر تصادفی مستقل دیگر با توابع بقای $\bar{H}_1(k; \lambda_1)$ و $\bar{H}_2(k; \lambda_2)$ باشند. اگر $\alpha \geq \beta$ و $\lambda_1 \leq \lambda_2$ باشد آنگاه

$$X_1 \geq_{st} X_2 \implies Y_1 \geq_{st} Y_2.$$

اثبات. فرض کنید $X_1 \geq_{st} X_2$ باشد. در این صورت $\bar{F}_1(k) \geq \bar{F}_2(k)$ است. برای رسیدن به نتیجه لازم کافی است نشان داده شود $\bar{H}_1(k; \lambda_1) \geq \bar{H}_2(k; \lambda_2)$ است. یا بطور معادل:

$$A = \frac{\alpha \bar{F}_1^{\lambda_1}(k)}{1 - \bar{\alpha} \bar{F}_1^{\lambda_1}(k)} / \frac{\beta \bar{F}_2^{\lambda_2}(k)}{1 - \bar{\beta} \bar{F}_2^{\lambda_2}(k)} \geq 1.$$

با ساده کردن عبارت A داریم:

$$\begin{aligned} A &= \frac{\alpha \bar{F}_1^{\lambda_1}(k) (1 - (1 - \beta) \bar{F}_2^{\lambda_2}(k))}{\beta \bar{F}_2^{\lambda_2}(k) (1 - (1 - \alpha) \bar{F}_1^{\lambda_1}(k))} \\ &= \frac{\alpha \bar{F}_1^{\lambda_1}(k) - \alpha \bar{F}_1^{\lambda_1}(k) \bar{F}_2^{\lambda_2}(k) + \alpha \beta \bar{F}_1^{\lambda_1}(k) \bar{F}_2^{\lambda_1}(k)}{\beta \bar{F}_2^{\lambda_2}(k) - \beta \bar{F}_1^{\lambda_1}(k) \bar{F}_2^{\lambda_2}(k) + \alpha \beta \bar{F}_1^{\lambda_1}(k) \bar{F}_2^{\lambda_2}(k)} \\ &= \frac{\alpha \bar{F}_1^{\lambda_1}(k) (1 - \bar{F}_2^{\lambda_2}(k)) + \alpha \beta \bar{F}_1^{\lambda_1}(k) \bar{F}_2^{\lambda_2}(k)}{\beta \bar{F}_2^{\lambda_2}(k) (1 - \bar{F}_1^{\lambda_1}(k)) + \alpha \beta \bar{F}_1^{\lambda_1}(k) \bar{F}_2^{\lambda_2}(k)} \end{aligned}$$

از بزرگتر یا مساوی یک بودن کمیت A باید داشته باشیم:

$$\alpha \bar{F}_1^{\lambda_1}(k) (1 - \bar{F}_2^{\lambda_2}(k)) \geq \beta \bar{F}_2^{\lambda_2}(k) (1 - \bar{F}_1^{\lambda_1}(k)), \quad (11)$$

□

که نتیجه فوق با استفاده از $\lambda_1 \leq \lambda_2$, $\alpha \geq \beta$ و $\bar{F}_1(k) \geq \bar{F}_2(k)$ برقرار است.

فرض کنید X_1 و X_2 دو متغیر تصادفی مستقل با توابع بقای به ترتیب $\bar{F}_1(k)$ و $\bar{F}_2(k)$ باشند. همچنین فرض کنید Y_1 و Y_2 دو متغیر تصادفی مستقل دیگر با توابع بقای $\bar{H}_1(k; \lambda)$ و $\bar{H}_2(k; \lambda)$ باشند. اگر $\alpha \geq \beta \geq 1$ باشد آنگاه

$$X_1 \geq_{hr} X_2 \implies Y_1 \geq_{hr} Y_2.$$

اثبات. توابع نرخ خطر متغیرهای تصادفی Y_1 و Y_2 به ترتیب عبارتند از:

$$r_{Y_1}(k) = \frac{\bar{F}_1^\lambda(k-1) - \bar{F}_1^\lambda(k)}{\bar{F}_1^\lambda(k-1)(1 - \alpha \bar{F}_1^\lambda(k))},$$

و

$$r_{Y_2}(k) = \frac{\bar{F}_2^\lambda(k-1) - \bar{F}_2^\lambda(k)}{\bar{F}_2^\lambda(k-1)(1 - \beta \bar{F}_2^\lambda(k))}.$$

برای رسیدن به نتیجه مطلوب، کافی است به ازای $k \in \mathbb{N}$ نشان داده شود $r_{Y_1}(k) \leq r_{Y_2}(k)$ است. چون ترتیب نرخ خطر، ترتیب تصادفی معمولی را نتیجه می‌دهد بنابراین از $X_1 \geq_{hr} X_2$ نتیجه می‌شود که به ازای هر $k \in \mathbb{N}$ ، رابطه $\bar{F}_1(k) \geq \bar{F}_2(k)$ برقرار است که به همراه $\alpha \geq \beta \geq 1$ می‌توان نتیجه گرفت:

$$\frac{1}{1 - \alpha \bar{F}_1^\lambda(k)} \leq \frac{1}{1 - \beta \bar{F}_2^\lambda(k)}. \quad (12)$$

از طرف دیگر، چون $X_1 \geq_{hr} X_2$ است در نتیجه نسبت $\bar{F}_1(k)/\bar{F}_2(k)$ تابعی صعودی از k است و این یعنی اینکه

$$\frac{\bar{F}_1(k)}{\bar{F}_2(k)} \geq \frac{\bar{F}_1(k-1)}{\bar{F}_2(k-1)} \quad \text{یا} \quad \frac{\bar{F}_1^\lambda(k)}{\bar{F}_1^\lambda(k-1)} \geq \frac{\bar{F}_2^\lambda(k)}{\bar{F}_2^\lambda(k-1)}.$$

با ضرب کردن عبارت $\frac{\bar{F}_1^\lambda(k)}{\bar{F}_1^\lambda(k-1)} \geq \frac{\bar{F}_2^\lambda(k)}{\bar{F}_2^\lambda(k-1)}$ در یک علامت منفی و جمع کردن آن با عدد یک، نتیجه می‌شود:

$$\frac{\bar{F}_1^\lambda(k-1) - \bar{F}_1^\lambda(k)}{\bar{F}_1^\lambda(k-1)} \leq \frac{\bar{F}_2^\lambda(k-1) - \bar{F}_2^\lambda(k)}{\bar{F}_2^\lambda(k-1)}. \quad (13)$$

اکنون از ضرب طرفین روابط نامنفی (12) و (13) نتیجه می‌شود که $r_{Y_1}(k) \leq r_{Y_2}(k)$ است.

□

فرض کنید متغیرهای تصادفی X_1 و X_2 دارای توابع جرم احتمال زیر باشند:

$$P(X_1 = k) = \begin{cases} 0 & k = 1 \\ 0/2 & k = 2 \\ 0/3 & k = 3 \\ 0/2 & k = 4 \\ 0/3 & k = 5. \end{cases}$$

$$P(X_2 = k) = \begin{cases} 0 & k = 1 \\ 0/3 & k = 2 \\ 0/4 & k = 3 \\ 0/2 & k = 4 \\ 0/1 & k = 5. \end{cases}$$

بسادگی می‌توان مشاهده نمود که $X_1 \geq_{lr} X_2$ است. اما باید توجه داشت که برای $\alpha = 5$ و $\lambda = 2$ داریم:

$$\frac{h_1(2)}{h_2(2)} = 0/5869, \quad \frac{h_1(3)}{h_2(3)} = 0/5512, \quad \frac{h_1(4)}{h_2(4)} = 1/0400, \quad \frac{h_1(5)}{h_2(5)} = 6/8870.$$

بنابراین چون نسبت $h_1(k)/h_2(k)$ نه صعودی است و نه نزولی، بنابراین $Y_1 \not\geq_{lr} Y_2$.

۲ دست‌آوردهای پژوهش

مدل نرخ خطر متناسب اصلاح شده به عنوان یکی از خانواده‌های انعطاف‌پذیر در قابلیت اعتماد و تحلیل بقا، و مقایسه‌های تصادفی سیستم‌های $(n - k + 1)$ از n در این خانواده از توزیع‌ها، توسط بالاکریشنان و همکاران (۲۰۱۸) معرفی شده است. در این مقاله، حالت گسسته این مدل در نظر گرفته می‌شود و به خواص سالخوردگی و حفظ شدن ترتیب تصادفی معمولی، نرخ خطر و نسبت درستنمایی در این خانواده از توزیع‌ها پرداخته شده است. نتایج اثبات شده، نتایج به دست آمده توسط کوندو و ناندا (۲۰۱۸) را تعمیم یا تکمیل می‌کند.

مراجع

- [1] Balakrishnan, N., Barmalzan, G. and Haidari, A. (2018), Modified proportional hazard rates and proportional reversed hazard rates models via Marshall-Olkin distribution and some stochastic comparisons, *Journal of the Korean Statistical Society*, **47**, 127-138.
- [2] Barlow R.E. and Proschan, F. (1981), *Statistical theory of reliability and life testing*, Holt, Rinehart and Winston, New York.
- [3] Caroni, C. (2010). Testing for the Marshall-Olkin extended form of the Weibull distribution, *Statistical Papers*, **51**, 325-336.
- [4] Kundu, P. and Nanda, A.K. (2018), Reliability study of proportional odds family of discrete distributions, *Communications in Statistics-Theory and Methods*, **47**, 1091-1103.

- [5] Kirmani, S.N.U.A. and Gupta, R.C. (2001), On the proportional odds model in survival analysis, *Annals of the Institute of Statistical Mathematics*, **53**, 203-216.
- [6] Marshall, A.W, and Olkin, I. (2007), *Life distributions*, New York, NY, Springer.
- [7] Nanda, A.K. and Das, S. (2013), Some ageing properties of Marshall-Olkin extended distribution. *International Journal of Mathematics and Statistics*, **13**, 93-107.
- [8] Shaked, M. and Shanthikumar, J.G. (2007), *Stochastic Orders*, Springer, New York.

برآورد نیم پارامتری معیار قابلیت اعتماد مدل‌های تنش-مقاومت وابسته

علی دست برآورده^۱، فاطمه محمودی فرد^۲

^۱ استادیار گروه آمار، دانشکده علوم ریاضی، دانشگاه یزد،

^۲ فارغ التحصیل کارشناسی ارشد آمار ریاضی، دانشگاه یزد

چکیده: در این مقاله، روشی جهت برآورد معیار قابلیت اعتماد مدل تنش-مقاومت وابسته با رویکرد مفصل مبنا پیشنهاد شده است. در روش پیشنهادی ابتدا توزیع‌های حاشیه‌ای به شیوه پارامتری و ساختار وابستگی به شیوه ناپارامتری برآورده شده، سپس براساس توزیع توأم به دست آمده معیار قابلیت اعتماد برآورد می‌شود. در نهایت عملکرد روش برآورد پیشنهادی به شیوه شبیه‌سازی مونت کارلو ارزیابی شده است.

واژه‌های کلیدی: آزمون نیکویی برازش، برآورد ماکسیمم درست‌نمایی، تابع مفصل، تنش-مقاومت وابسته، نیم پارامتری.
کد موضوع بندی: 39B82، 34K20، 39B52، 47A55.

۱ پیش‌گفتار

در برخی مسائل کاربردی به دنبال ارزیابی میزان تحمل یک مؤلفه تصادفی در برابر فشار وارده از سوی مؤلفه تصادفی دیگر هستیم. در مباحث قابلیت اعتماد این‌گونه مسائل را مدل‌های تنش-مقاومت گویند. در مدل تنش-مقاومت سیستم تا زمانی که مقاومت مؤلفه سیستم بیش از فشار وارد بر آن باشد سیستم به فعالیت خود ادامه می‌دهد. در این‌گونه مدل‌ها، معیاری برای سنجش میزان قابلیت اعتماد سیستم به صورت احتمال ادامه فعالیت سیستم یعنی $R = P(X > Y)$ در نظر گرفته می‌شود که در آن X و Y به ترتیب مقاومت و فشار تصادفی سیستم هستند.

قابلیت اعتماد در سیستم‌های مهندسی از اهمیت بالایی برخوردار هستند؛ اما کاربرد آن محدود به این علوم نیست. در حقیقت در زمینه‌های مختلفی از جمله کنترل کیفیت، ژنتیک، روانشناسی، تعلیم و تربیت، و بهداشت و درمان کاربرد دارد.

برخی از پژوهش‌های انجام شده در این زمینه، R را وقتی متغیرهای X و Y مستقل با توزیع‌های حاشیه‌ای مشترک از جمله، پاراتو تعمیم‌یافته [۱۵]، لجستیک تعمیم‌یافته [۳]، لندلی وزن‌دار [۱]، داگام [۲] هستند، محاسبه و برآورد کرده‌اند. با وجود اینکه در برخی مسائل کاربردی ممکن است تنش و مقاومت مستقل باشند، اما در بسیاری از مسائل کاربردی تنش و مقاومت وابسته‌اند و برای محاسبه R نیازمند توزیع توأم هستیم. برخی از پژوهش‌ها برای تعدادی توزیع‌های توأم خاص مانند نرمال دومتغیره [۸]، پاراتو دومتغیره [۱۰]، بتا دومتغیره [۱۲]، گاما دومتغیره [۱۳]، لگ نرمال دومتغیره [۹]، R را محاسبه و برآورد کرده‌اند؛ این پژوهش‌ها متغیرهای تنش و مقاومت را تنها در فرمت خاصی از وابستگی در نظر گرفته و در آنها هر دو توزیع حاشیه‌ای از یک خانواده هستند. توابع مفصل کاربرد وسیعی در مدل‌سازی توزیع‌های چند متغیره دارند. همچنین تابع مفصل ساختار وابستگی مختلف (از جمله خطی، غیر خطی، وابستگی دمی، ...) را در نظر می‌گیرد. استفاده از تابع مفصل برای محاسبه و برآورد R در مدل تنش-مقاومت وابسته رویکردی نوین و بسیار مفید در مسائل کاربردی است. این موضوع پژوهشی به تازگی مورد توجه محققان قرار گرفته است.

^۱ علی دست برآورده: dastbaravarde@yazd.ac.ir

به پژوهش‌های انجام شده در این زمینه می‌توان به دو پژوهش دوما و گیوردانو [۶] در سال ۲۰۱۲، رویکرد مفصل مبنا برای محاسبه R در مدل تنش-مقاومت وابسته و دوما و گیوردانو [۴] در سال ۲۰۱۳، کاربرد این رویکرد در برآورد معیار ورشکستگی خانواده اشاره کرد.

اسکلار در سال ۱۹۵۹ نشان داد که رابطه‌ای خاص بین تابع توزیع توأم و توابع توزیع حاشیه‌ای متغیرهای تصادفی وجود دارد. فرض کنید X و Y دو متغیر تصادفی با تابع توزیع توأم H و توابع توزیع حاشیه‌ای به‌ترتیب F و G باشند آنگاه، تابع حقیقی مقدار $[0, 1] \times [0, 1] \rightarrow [0, 1]$ وجود دارد به‌قسمی که برای هر دو مقدار حقیقی x و y

$$H(x, y) = C(F(x), G(y)).$$

تابع C را مفصل می‌نامند که دارای ویژگی‌های زیر است:

۱. به‌ازای کلیه u و v های متعلق به I ,

$$\begin{aligned} C(u, 0) &= C(0, v) = 0, \\ C(u, 1) &= u, \quad C(1, v) = v, \end{aligned}$$

۲. اگر $u_1 < u_2$ و $v_1 < v_2$ آنگاه

$$C(u_1, v_1) - C(u_1, v_2) - C(u_2, v_1) + C(u_2, v_2) \geq 0.$$

در صورتی که X و Y مطلقاً پیوسته بوده و C دارای مشتقات جزئی باشد، تابع چگالی توأم X و Y ، h ، را می‌توان به‌صورت زیر نوشت:

$$h(x, y) = c(F(x), G(y))f(x)g(y),$$

که در آن f و g به‌ترتیب تابع چگالی حاشیه‌ای X و Y ، و $c(u, v) = \frac{\partial^2 C(u, v)}{\partial u \partial v}$ تابع چگالی مفصل هستند. فرض کنید C یک تابع مفصل پیوسته مطلق باشد. در این صورت

$$\begin{aligned} \frac{\partial C(u, v)}{\partial u} &= P(V \leq v \mid U = u), \\ \frac{\partial C(u, v)}{\partial v} &= P(U \leq u \mid V = v). \end{aligned}$$

برای اطلاعات بیشتر درمورد تابع مفصل می‌توان به منبع [۱۴] مراجعه کرد.

در ادامه این مقاله، در بخش ۲ محاسبه پارامتر مدل تنش-مقاومت با رویکرد تابع مفصل ارائه می‌شود. در بخش ۵ روشی جدید برای برآورد قابلیت اعتماد پیشنهاد می‌شود، در بخش ۴ روش برآورد پیشنهادی را با استفاده شبیه‌سازی مونت‌کارلو مورد بررسی قرار می‌گیرد. در نهایت، جمع‌بندی نتایج در بخش ۵ ارائه می‌شود.

۲ محاسبه قابلیت اعتماد با استفاده از رویکرد مفصل مبنا

اگر X و Y متغیرهای تصادفی پیوسته با ساختار وابستگی C و توابع توزیع حاشیه‌ای به ترتیب F ، G و توابع چگالی حاشیه‌ای به ترتیب f و g باشند، آنگاه R به صورت زیر محاسبه می‌شود:

$$\begin{aligned} R = P(X > Y) &= \int_{-\infty}^{\infty} \int_{-\infty}^x h(x, y) dy dx \\ &= \int_{-\infty}^{\infty} \int_{-\infty}^x c(F(x), G(y)) g(y) f(x) dy dx. \end{aligned}$$

یا به طور معادل،

$$R = P(X > Y) = \int_0^1 \int_0^{G(F^{-1}(u))} c(u, v) dv du. \quad (1)$$

از طرف دیگر، با توجه به

$$G_{Y|X=x}(y) = \frac{\partial C(u, v)}{\partial u} \Big|_{u=F(x), v=G(y)}, \quad R = \int_{\mathbb{R}} G_{Y|X=x}(x) f(x) dx,$$

داریم:

$$R = E \left(\frac{\partial C(u, v)}{\partial u} \Big|_{u=F(X), v=G(X)} \right). \quad (2)$$

۳ روش برآورد پیشنهادی

روش ناپارامتری مرسوم برای برآورد R به صورت $r_o = \frac{1}{n} \sum_i I(x_i > y_i)$ است. اما در ادامه روشی برای برآورد نیم پارامتری R مبتنی بر مفصل پیشنهاد شده است. در این روش، فرض می‌کنیم توزیع‌های حاشیه‌ای دارای پارامتر(های) نامعلوم هستند. روش برآورد پیشنهادی، طی چهار مرحله به صورت زیر تشریح شده است.

۱. پارامتر توزیع‌های حاشیه‌ای را به روش ماکسیمم درست‌نمایی برآورد می‌کنیم $(F(x; \hat{\alpha}), G(y; \hat{\beta}))$.

۲. مفصل مناسب جهت تبیین ساختار وابستگی X و Y براساس آزمون نیکویی برازش، برازش داده می‌شود $(\tilde{C})$.

۳. محاسبه توزیع توأم X و Y با استفاده از توزیع‌های حاشیه‌ای با پارامتر برآورد شده و مفصل برازش داده شده،

$$\tilde{H}(x, y) = \tilde{C} \left(F(x; \hat{\alpha}), G(y; \hat{\beta}) \right). \quad (3)$$

۴. برآورد R با استفاده از توزیع توأم به دست می‌آید. یعنی،

$$\hat{R} = \int_{-\infty}^{+\infty} \int_{-\infty}^x \tilde{h}(x, y) dy dx,$$

که در آن $\tilde{h}(x, y)$ تابع چگالی توأم X و Y متناظر با $\tilde{H}$ است.

همچنین، مقدار $\hat{R}$ را می‌توان از رابطه زیر نیز به دست آورد.

$$\hat{R} = E \left(\frac{\partial \tilde{C}(u, v)}{\partial u} \Big|_{u=F(x; \hat{\alpha}), v=G(y; \hat{\beta})} \right).$$

۴ شبیه‌سازی

در این بخش عملکرد برآورد پیشنهادی برای قابلیت اعتماد مدل تنش-مقاومت وابسته را به روش شبیه‌سازی مونت کارلو مطالعه می‌کنیم.

شبیه‌سازی را برای نمونه‌هایی به حجم $n = 50, 100, 200$ با سه توزیع توأم مختلف انجام داده‌ایم. این سه توزیع توأم عبارتند از:

الف) توزیع توأم حاصل از مفصل فرانک با پارامتر وابستگی $\theta = 2$ و حاشیه‌ای نمایی با میانگین ۲ برای X و حاشیه‌ای نرمال با میانگین و واریانس به ترتیب ۱ و ۱۰ برای Y ،

ب) توزیع توأم حاصل از مفصل کلایتون با پارامتر وابستگی $\theta = 2$ و حاشیه‌ای نمایی با میانگین ۲ برای X و حاشیه‌ای نرمال با میانگین و واریانس به ترتیب ۱ و ۱۰ برای Y ،

پ) توزیع توأم حاصل از مفصل جو با پارامتر وابستگی $\theta = 2$ و حاشیه‌ای نمایی با میانگین ۲ برای X و حاشیه‌ای نرمال با میانگین و واریانس به ترتیب ۱ و ۱۰ برای Y .

در ادامه مراحل این شبیه‌سازی برای توزیع توأم در قسمت الف)، تشریح می‌شود. شبیه‌سازی برای توزیع‌های توأم دیگر به صورت مشابه انجام شده است.

۱. از مفصل دو بعدی فرانک با پارامتر وابستگی $\theta = 2$ ، نمونه‌ای تصادفی به حجم n تولید می‌کنیم؛ یعنی

$$(u_i, v_i) \stackrel{iid}{\sim} frank(2), \quad i = 1, \dots, n.$$

۲. در این مرحله نمونه تصادفی $(x_i, y_i), i = 1, \dots, n$ را به صورت زیر تولید می‌کنیم؛

$$x_i = F^{-1}(u_i), \quad y_i = G^{-1}(v_i),$$

که در آن $F^{-1}(s)$ چندک s -ام از توزیع نمایی با میانگین ۲ و $G^{-1}(t)$ چندک t -ام از توزیع نرمال با میانگین و واریانس ۱ و ۱۰ هستند.

۳. مقدار تجربی قابلیت اعتماد $(r_o = \frac{\sum_i I(x_i > y_i)}{n})$ ، را محاسبه می‌کنیم.

۴. در این مرحله U_i و V_i را به صورت زیر محاسبه می‌کنیم؛

$$U_i = F(x_i; \hat{\lambda}), \quad V_i = G(y_i; \hat{\mu}, \hat{\sigma}^2),$$

که در آن F و G به ترتیب توزیع‌های نمایی و نرمال با پارامترهای برآورد شده هستند که پارامترهای آنها به روش ماکسیم درست‌نمایی برآورد شده است. برای انتخاب مفصل مناسب نیز از معیار AIC استفاده می‌شود. برای این منظور در نرم افزار R از دستور $BiCopSelect$ از بسته $VineCopula$ استفاده کردیم. دستور $BiCopSelect$ از بین ۴۰ مفصل که شرح آنها در بسته مورد نظر آمده مفصلی که کمترین AIC را داشته باشد، پیشنهاد می‌دهد.

۵. در این مرحله با استفاده از حاشیه‌های نمایی و نرمال با پارامترهای برآورد شده و مفصل برازش داده شده (حاصل از مرحله ۴) توزیع توأم را به دست می‌آوریم، سپس با استفاده از رابطه (۲) قابلیت اعتماد برآورد می‌شود.

$$r_1 = E \left(\frac{\partial \tilde{C}(u, v)}{\partial u} \Big|_{u=F(X; \hat{\lambda}), v=G(Y; \hat{\mu}, \hat{\sigma}^2)} \right). \quad (4)$$

که $\tilde{C}$ مفصل برازش شده C در مرحله ۴، است.

۶. مراحل ۱ تا ۵ را ۱۰۰۰۰ بار تکرار می‌کنیم.

۷. در این مرحله اریبی و میانگین مربعات خطای روش‌های برآورد محاسبه می‌کنیم.

$$Bias(r_1) \approx \frac{1}{m} \sum_{i=1}^m r_{1i} - R,$$

$$MSE(r_1) \approx \frac{1}{m} \sum_{i=1}^m (r_{1i} - R)^2,$$

که در آن r_{1i} برآورد R تحت روش پیشنهادی در تکرار i -ام است. اریبی و میانگین مربعات خطای r_0 به صورت مشابه نیز محاسبه می‌گردد.

جدول ۱: نتایج شبیه‌سازی

n	R	اریبی r_0	اریبی r_1	$MSE(r_0)$	$MSE(r_1)$	کارایی r_0 نسبت r_1
مفصل فرانک						
۵۰		۰/۰۰۰۰۳	-۰/۰۰۰۸۸	۰/۰۰۰۴۶	۰/۰۰۰۳۳	۱/۳۳۵۹
۱۰۰	۰/۶۲۰۵	-۰/۰۰۰۰۱	-۰/۰۰۰۳۵	۰/۰۰۰۲۳	۰/۰۰۰۱۷	۱/۳۶۲۹
۲۰۰		-۰/۰۰۰۰۲	-۰/۰۰۰۱۵	۰/۰۰۰۱۲	۰/۰۰۰۰۸	۱/۴۰۸۰
مفصل کلایتون						
۵۰		۰/۰۰۰۰۱	-۰/۰۰۰۳۲	۰/۰۰۰۴۴	۰/۰۰۰۳۴	۱/۳۷۹۰
۱۰۰	۰/۶۶۵۹	۰/۰۰۰۰۴	-۰/۰۰۰۱۶	۰/۰۰۰۲۲	۰/۰۰۰۱۶	۱/۳۸۲۳
۲۰۰		-۰/۰۰۰۰۴	-۰/۰۰۰۰۹	۰/۰۰۰۱۱	۰/۰۰۰۰۸	۱/۳۸۸۷
مفصل جو						
۵۰		۰/۰۰۰۰۹	۰/۰۰۰۱۹	۰/۰۰۰۴۷	۰/۰۰۰۳۲	۱/۴۶۳۵
۱۰۰	۰/۶۲۸۲	۰/۰۰۰۰۵	۰/۰۰۰۱۹	۰/۰۰۰۲۴	۰/۰۰۰۱۶	۱/۴۶۸۷
۲۰۰		-۰/۰۰۰۰۵	۰/۰۰۰۱۲	۰/۰۰۰۱۲	۰/۰۰۰۰۸	۱/۴۷۵۸

همان‌طور که از نتایج مشاهده می‌شود، میانگین مربعات خطای r_1 نسبت به r_0 کمتر است و همچنین میانگین مربعات خطای r_0 و r_1 با افزایش حجم نمونه کاهش می‌یابد. بنابراین، می‌توان گفت که این روش برای برآورد قابلیت اعتماد مناسب‌تر است.

۵ دست‌آوردهای پژوهش

در این مقاله پارامتر قابلیت اعتماد مدل تنش-مقاومت را با رویکرد تابع مفصل مورد بررسی قرار دادیم. یک روش پیشنهادی برای برآورد آن ارائه دادیم و در آخر این روش پیشنهادی را برای داده‌های شبیه‌سازی شده مورد مطالعه قرار دادیم که با توجه به نتایج به دست آمده می‌توان گفت این رویکرد دقت بیشتری نسبت به برآورد کلاسیک دارد.

مراجع

- [1] Al-Mutairi, D.K., Ghitany, M.E. and Kundu, D. (2015), Inferences on stress-strength reliability from weighted Lindley distributions, *Communications in Statistics Theory and Methods*, **44**(19), 4096-4113.
- [2] Al-Zahrani, B. and Basloom, S. (2016), Estimation of the stress-strength reliability for the Dagum distribution, *Journal of Advanced Statistics*. **1**(3), 157.
- [3] Asgharzadeh, A., Valiollahi, R. and Raqab, M.Z. (2013), Estimation of the stress-strength reliability for the generalized logistic distribution, *Statistical Methodology*, **15**, 73-94.
- [4] Domma, F. and Giordano, S. (2012), A stress-strength model with dependent variables to measure household financial fragility, *Statistical Methods and Applications*, **21**(3), 375-389.
- [5] Dey, S., Al-Zahrani, B. and Basloom, S. (2017), Dagum distribution: Properties and different methods of estimation, *International Journal of Statistics and Probability*, **6**(2), 74.
- [6] Domma, F. and Giordano, S. (2013), A copula-based approach to account for dependence in stress-strength models, *Statistical Papers*, **54**(3), 807-826.
- [7] Guo, H. and Krishnamoorthy, K. (2004), New approximate inferential methods for the reliability parameter in a stress-strength model: The normal case. *Communications in Statistics-Theory and Methods*, **33**(7), 1715-1731.
- [8] Gupta, R.C. and Subramanian, S. (1998), Estimation of reliability in a bivariate normal distribution with equal coefficients of variation, *Communications in Statistics - Simulation and Computation*, **27**(3), 675-698.
- [9] Gupta, R.C., Ghitany, M.E. and Al-Mutairi, D.K. (2013), Estimation of reliability from a bivariate log-normal data, *Journal of Statistical Computation and Simulation*, **83**(6), 1068-1081.

- [10] Hanagal, D.D. (1997), Note on estimation of reliability under bivariate pareto stress-strength model, *Statistical Papers*, **38**, 453-459.
- [11] Huang, K., Mi, J. and Wang, Z. (2012), Inference about reliability parameter with gamma strength and stress, *Journal of tatistical Planning and Inference*, **142**(4), 848-854.
- [12] Nadarajah, S. (2005), Reliability for some Bivariate gamma distributions, *Mathematical Problems in Engineering*, **2005**(1), 151-163.
- [13] Nadarajah, S. (2005), Reliability for some bivariate beta distributions, *Mathematical Problems in Engineering*, **2005**(2), 101-11
- [14] Nelsen, R.B. (2007), *An Introduction to Copulas*, Springer Science and Business Media.
- [15] Rezaei, S., Tahmasbi, R. and Mahmoodi, M. (2010), Estimation of $P(Y < X)$ for generalized pareto distribution, *Journal of Statistical Planning and Inference*, **140**, 480-494.

برآورد نقطه‌ای، بازه‌ای پارامتر تنش-مقاومت در توزیع نمایی دو پارامتری براساس داده‌های رکوردی

صغری رحیمی^۱

فارغ التحصیل کارشناسی ارشد، دانشگاه علامه طباطبائی

چکیده: در این مقاله به برآورد نقطه‌ای و بازه‌ای پارامتر تنش-مقاومت $R = P(X < Y)$ در توزیع نمایی دو پارامتری براساس داده‌های کامل، می‌پردازیم. برای این منظور ابتدا برآورد ماکسیم درست‌نمایی پارامتر تنش-مقاومت را براساس داده‌های رکوردی با در نظر گرفتن این‌که پارامترهای مقیاسی و مکانی هر رو مجهول باشند محاسبه کرده و بازه‌ی اطمینان تعمیم‌یافته این پارامتر را به‌دست آورده و سپس، با استفاده از یک مطالعه شبیه‌سازی، به ارزیابی و مقایسه برآوردهای بازه‌ای معرفی شده نسبت به هم و با توجه به حجم نمونه‌های مختلف پرداخته می‌شود.

واژه‌های کلیدی: بازه‌ی اطمینان تعمیم‌یافته، برآورد ماکسیم درست‌نمایی، پارامتر تنش-مقاومت، تابع قابلیت اعتماد، توزیع نمایی دو پارامتری، داده‌های رکوردی.

۱ پیش‌گفتار

مدل تنش-مقاومت ابتدا در داده‌های ناپارامتری توسط ویلکاگون (۱۹۴۵) و من و ویتنی (۱۹۴۷) مورد بررسی قرار گرفت. هدف اصلی از این بررسی، مقایسه‌ی دو متغیر تصادفی X و Y حاصل از نتایج دو نوع دارو در درمان بیماران بود. تلاش‌های اولیه برای مطالعه‌ی $P(X < Y)$ تحت فرض‌های پارامتری روی X و Y توسط اوون و همکاران (۱۹۶۴) انجام شد. آن‌ها حالتی که X و Y متغیرهای تصادفی مستقل یا وابسته از توزیع نرمال با انحراف استانداردهای برابر هستند را مورد بررسی قرار دادند و حدود اطمینان را برای $P(X < Y)$ برای اغلب توزیع‌ها به‌دست آوردند. چرچ و هریس (۱۹۷۰) (مقاله‌ای که در آن برای اولین بار $P(X < Y)$ با عنوان رسمی تنش-مقاومت مطرح شد) برآورد ماکسیم درست‌نمایی $P(X < Y)$ را که X و Y متغیرهای نرمال و مستقل هستند و توزیع X به‌طور کامل مشخص است، به‌دست آوردند.

در بیش‌تر مواقع داده‌ها زمانی ثبت می‌شوند که یک رکورد باشند، مانند داده‌هایی که در زمینه‌های ورزشی جمع‌آوری می‌شوند. اگر هر مشاهده بزرگ‌تر از تمام مشاهدات قبلی خود باشد یک رکورد بالا (یک رکورد) نامیده می‌شود. از طرف دیگر، اگر هر مشاهده کوچک‌تر از مشاهدات قبلی باشد یک رکود پایین نامیده می‌شود. اگر $X_1, X_2, \dots$ دنباله‌ای از متغیرهای تصادفی باشد آن‌گاه X_j یک رکورد بالا نامیده می‌شود اگر از همه‌ی مشاهدات قبلی خود بزرگ‌تر باشد.

در این مقاله متغیر Y مقاومت و متغیر X تنش در نظر گرفته می‌شود. برای مثال، در علم پزشکی X و Y به‌ترتیب طول عمر فرد بیمار تحت درمان و درمان سایر گروه‌ها است و این احتمال میزان درمان مؤثر را مشخص می‌کند. فرض کنید X و Y دو متغیر

^۱ صغری رحیمی: rahimi.soghra70@yahoo.com

تصادفی مستقل باشند به طوری که X و Y به ترتیب دارای توزیع نمایی دو پارامتری $Exp(\mu_1, \sigma_1)$ و $Exp(\mu_2, \sigma_2)$ با تابع چگالی احتمال زیر باشند

$$f(x; \mu_1, \sigma_1) = \frac{1}{\sigma_1} e^{-\frac{(x-\mu_1)}{\sigma_1}}, \quad x > \mu_1, \quad \mu_1 \in \mathbb{R}, \quad \sigma_1 \geq 0 \quad (1)$$

$$f(y; \mu_2, \sigma_2) = \frac{1}{\sigma_2} e^{-\frac{(y-\mu_2)}{\sigma_2}}, \quad y > \mu_2, \quad \mu_2 \in \mathbb{R}, \quad \sigma_2 \geq 0. \quad (2)$$

۲ پارامتر تنش-مقاومت

مسئله برآورد پارامتر تنش-مقاومت یکی از مسائل مهم و مورد توجه پژوهشگران در زمینه‌های علوم، صنعت، پزشکی و غیره است. اگر Y میزان مقاومت یک مؤلفه در یک سامانه و X مقدار تنش وارد شده به آن مؤلفه باشد، آنگاه احتمال آن که در یک فاصله زمانی مقاومت بیش‌تر از تنش باشد برابر $R = P(X < Y)$ است که به آن پارامتر تنش-مقاومت گویند. در واقع پارامتر R همان قابلیت اعتماد سامانه است. در حالت کلی، اگر X و Y دو متغیر تصادفی و به ترتیب دارای تابع توزیع تجمعی F و G باشند و فرض کنیم که Y در یک سامانه‌ی مورد نظر مقاومتی باشد در مقابل تنش X ، آنگاه R در آن سامانه به صورت زیر تعریف می‌شود

$$R = P(X < Y) = \int_0^\infty \bar{G}(x) dF(x),$$

به طوری که $\bar{G}(x) = 1 - G(x)$.

پارامتر تنش-مقاومت R در توزیع نمایی دو پارامتری به صورت زیر محاسبه می‌شود

$$R = P(X < Y) = \int_A \int \frac{1}{\sigma_1} e^{-\frac{(x-\mu_1)}{\sigma_1}} \frac{1}{\sigma_2} e^{-\frac{(y-\mu_2)}{\sigma_2}} dx dy,$$

اگر $\mu_1 > \mu_2$ باشد آنگاه ناحیه‌ی $A = \{(x, y) | y > x > \mu_1\}$ به صورت A خواهد بود و در نتیجه

$$\begin{aligned} R &= \int_{\mu_1}^\infty \frac{1}{\sigma_1} e^{-\frac{(x-\mu_1)}{\sigma_1}} \left[\int_x^\infty \frac{1}{\sigma_2} e^{-\frac{(y-\mu_2)}{\sigma_2}} dy \right] dx \\ &= \frac{\sigma_2}{\sigma_1 + \sigma_2} e^{-\frac{(\mu_1 - \mu_2)}{\sigma_2}}, \end{aligned}$$

حال اگر $\mu_1 \leq \mu_2$ ، آنگاه $A = \{(x, y) | \mu_2 < y < \infty, \mu_1 < x < y\}$ و در نتیجه

$$\begin{aligned} R &= \int_{\mu_2}^\infty \frac{1}{\sigma_2} e^{-\frac{(y-\mu_2)}{\sigma_2}} \left[\int_{\mu_1}^y \frac{1}{\sigma_1} e^{-\frac{(x-\mu_1)}{\sigma_1}} dx \right] dy \\ &= 1 - \frac{\sigma_1}{\sigma_1 + \sigma_2} e^{-\frac{(\mu_1 - \mu_2)}{\sigma_1}}, \end{aligned}$$

در آخر پارامتر تنش-مقاومت به صورت زیر حاصل می‌شود

$$R = \left(\frac{\sigma_2}{\sigma_1 + \sigma_2} e^{-\frac{(\mu_1 - \mu_2)}{\sigma_2}} \right) I(\mu_1 > \mu_2) + \left(1 - \frac{\sigma_1}{\sigma_1 + \sigma_2} e^{-\frac{(\mu_1 - \mu_2)}{\sigma_1}} \right) I(\mu_1 \leq \mu_2). \quad (3)$$

۳ برآورد نقطه‌ای پارامتر R براساس داده‌های رکوردی

فرض کنید $R_0, R_1, \dots, R_n$ مقدار رکورد بالای توزیع $Exp(\mu_1, \sigma_1)$ با مقادیر مشاهده شده‌ی $r_0, r_1, \dots, r_n$ باشند و همچنین $S_0, S_1, \dots, S_m$ مقدار رکورد بالا از توزیع $Exp(\mu_2, \sigma_2)$ با مقادیر مشاهده شده‌ی $s_0, s_1, \dots, s_m$ باشند که مستقل از نمونه‌های اول هستند. برای به دست آوردن برآورد ماکسیمم درست‌نمایی (ML) پارامتر R کافی است برآوردگر ML

هر یک از پارامترها به صورت جداگانه محاسبه شود. برآوردهای ML پارامترهای (μ_1, σ_1) و (μ_2, σ_2) براساس داده‌های رکوردی به صورت زیر محاسبه می‌شود.

تابع درست‌نمایی براساس مقادیر رکوردی بالا از توزیع $Exp(\mu_1, \sigma_1)$ با استفاده از (۱) به صورت زیر به دست می‌آید

$$L(\mu_1, \sigma_1; r) = \frac{\prod_{i=0}^n f(r_i)}{\prod_{i=0}^{n-1} [1 - F(r_i)]} = \frac{1}{\sigma_1^{n+1}} e^{-\frac{1}{\sigma_1}(r_n - \mu_1)}$$

با توجه به حدود $\mu_1 \leq r_0 < r_1 < \dots < r_n < \infty$ برآوردهای ML پارامتر μ_1 عبارت است از

$$\hat{\mu}_1 = R_0.$$

برآوردهای ML پارامتر σ_1 به صورت زیر به دست می‌آید

$$\hat{\sigma}_1 = \frac{R_n - R_0}{n + 1}.$$

با قرار دادن برآوردهای ML پارامترهای $\mu_1, \sigma_1, \mu_2, \sigma_2$ در (۳) برآوردهای ML پارامتر R به دست می‌آید، یعنی

$$\hat{R} = \begin{cases} \frac{\hat{\mu}_2 - \hat{\mu}_1}{\hat{\sigma}_1 + \hat{\sigma}_2} e^{-\frac{\hat{\mu}_2 - \hat{\mu}_1}{\hat{\sigma}_2}}, & \hat{\mu}_2 < \hat{\mu}_1 \\ 1 - \frac{\hat{\sigma}_1}{\hat{\sigma}_1 + \hat{\sigma}_2} e^{-\frac{\hat{\mu}_2 - \hat{\mu}_1}{\hat{\sigma}_1}}, & \hat{\mu}_2 \geq \hat{\mu}_1 \end{cases}$$

۴ به دست آوردن بازه‌ی اطمینان به کمک تابع محوری تعمیم یافته

کمیت محوری (متغیر) تعمیم یافته: فرض کنید $X = (X_1, X_2, \dots, X_n)$ دارای توزیعی باشد که به پارامترهای (θ, ζ) وابسته است و $x = (x_1, x_2, \dots, x_n)$ مقدار مشاهده شده‌ی آن باشد و همچنین فرض کنید $(\hat{\theta}, \hat{\zeta})$ برآوردی برای (θ, ζ) با مقدار مشاهده شده‌ی $(\hat{\theta}_0, \hat{\zeta}_0)$ باشد. کمیت $G_\theta = G_\theta(X, x, \hat{\theta}, \hat{\zeta})$ را یک کمیت محوری (متغیر) تعمیم یافته برای θ ، زمانی که ζ پارامتر مزاحم است، گویند هرگاه در شرایط زیر صدق کند

(۱) مقدار $G_\theta(X; x, \hat{\theta}, \hat{\zeta})$ در نقطه‌ی $(\hat{\theta}_0, \hat{\zeta}_0) = (\hat{\theta}, \hat{\zeta})$ باید برابر با خود پارامتر مورد نظر θ باشد.

(۲) برای مقدار داده شده‌ی $(\hat{\theta}_0, \hat{\zeta}_0)$ ، توزیع $G_\theta(X; x, \hat{\theta}, \hat{\zeta})$ به هیچ یک از پارامترهای نامعلوم $(\hat{\theta}, \hat{\zeta})$ وابسته نباشد.

در این بخش به محاسبه‌ی کمیت محوری تعمیم یافته برای پارامترهای μ_i و σ_i ، $i = 1, 2$ براساس داده‌های رکوردی می‌پردازیم. گرچه این پارامترها مورد نظر ما نمی‌باشند، اما با نتایجی که از این محاسبات به دست می‌آوریم برای محاسبه‌ی کمیت محوری تعمیم یافته‌ی پارامتر تنش-مقاومت استفاده می‌کنیم.

۱.۴ کمیت محوری تعمیم یافته برای μ_i ها

اگر $\hat{\mu}_{10}, \hat{\sigma}_{10}$ ، مقدارهای مشاهده شده برآوردهای $\hat{\mu}_1, \hat{\sigma}_1$ باشند، در این صورت کمیت محوری تعمیم یافته ی G_{μ} ارائه شده به وسیله ی باکلزی (۲۰۰۸) به صورت زیر است

$$\begin{aligned} G_{\mu_1} &= \hat{\mu}_{10} - \frac{\Upsilon(n+1)(\hat{\mu}_1 - \mu_1) \sigma_1}{\Upsilon(n+1)\sigma_1} \frac{\sigma_1}{\hat{\sigma}_1} \hat{\sigma}_{10} \\ &= \hat{\mu}_{10} - \frac{\Upsilon(\hat{\mu}_1 - \mu_1)}{\sigma_1} \frac{1}{\Upsilon(n+1)\hat{\sigma}_1/\sigma_1} (n+1)\hat{\sigma}_{10} \\ &\stackrel{D}{=} \hat{\mu}_{10} - \frac{\chi_{\Upsilon}^2}{\chi_{\Upsilon n}^2} (r_n - r_0), \end{aligned} \quad (۴)$$

که به معنای هم توزیع بودن است. کمیت محوری G_{μ_1} در شرایط کمیت محوری تعمیم یافته صدق می کند. یک بازه ی اطمینان تعمیم یافته $(1-\alpha) \times 100\%$ برای μ_1 عبارت است از $(h_{\alpha/2}, h_{1-\alpha/2}) \in G_{\mu_1}(X, x, \hat{\mu}_{10}, \hat{\sigma}_{10})$ که h_{β} صدک 100β توزیع G_{μ_1} است. از طرفی با استفاده از (۴)، داریم

$$\begin{aligned} 1-\alpha &= P(h_{\alpha/2} < G_{\mu_1} < h_{1-\alpha/2}) \\ &= P(h_{\alpha/2} < \hat{\mu}_{10} - \frac{\Upsilon(\hat{\mu}_1 - \mu_1)}{\sigma_1} \frac{1}{\Upsilon(n+1)\hat{\sigma}_1/\sigma_1} (n+1)\hat{\sigma}_{10} < h_{1-\alpha/2}) \\ &= P(h_{\alpha/2} < \hat{\mu}_{10} - \frac{n+1}{n} F_{\Upsilon, \Upsilon n} \hat{\sigma}_{10} < h_{1-\alpha/2}) \\ &= P\left(\frac{n}{(n+1)\hat{\sigma}_{10}} (\hat{\mu}_{10} - h_{1-\alpha/2}) < F_{\Upsilon, \Upsilon n} < \frac{n}{(n+1)\hat{\sigma}_{10}} (\hat{\mu}_{10} - h_{\alpha/2})\right). \end{aligned}$$

بنابراین $\frac{n}{(n+1)\hat{\sigma}_{10}} (\hat{\mu}_{10} - h_{\alpha/2}) = F_{\Upsilon, \Upsilon n, \alpha/2}$ و $\frac{n}{(n+1)\hat{\sigma}_{10}} (\hat{\mu}_{10} - h_{1-\alpha/2}) = F_{\Upsilon, \Upsilon n, 1-\alpha/2}$ در نتیجه

$$h_{\alpha/2} = \hat{\mu}_{10} - \frac{n+1}{n} \hat{\sigma}_{10} F_{\Upsilon, \Upsilon n, 1-\alpha/2}, \quad h_{1-\alpha/2} = \hat{\mu}_{10} - \frac{n+1}{n} \hat{\sigma}_{10} F_{\Upsilon, \Upsilon n, \alpha/2}.$$

در نتیجه یک بازه ی اطمینان تعمیم یافته $(1-\alpha) \times 100\%$ برای μ_1 عبارتست از

$$\mu_1 \in \left(\hat{\mu}_{10} - \frac{n+1}{n} \hat{\sigma}_{10} F_{\Upsilon, \Upsilon n, 1-\alpha/2}, \hat{\mu}_{10} - \frac{n+1}{n} \hat{\sigma}_{10} F_{\Upsilon, \Upsilon n, \alpha/2} \right).$$

به همین ترتیب کمیت محوری تعمیم یافته برای μ_2 عبارت است از

$$G_{\mu_2} \stackrel{D}{=} \hat{\mu}_{20} - \frac{\chi_{\Upsilon}^2}{\chi_{\Upsilon m}^2} (s_m - s_0).$$

و یک بازه ی اطمینان تعمیم یافته $(1-\alpha) \times 100\%$ برای μ_2 عبارت است از

$$\mu_2 \in \left(\hat{\mu}_{20} - \frac{m+1}{m} \hat{\sigma}_{20} F_{\Upsilon, \Upsilon m, 1-\alpha/2}, \hat{\mu}_{20} - \frac{m+1}{m} \hat{\sigma}_{20} F_{\Upsilon, \Upsilon m, \alpha/2} \right).$$

۲.۴ کمیت محوری تعمیم یافته برای σ_i ها

در این بخش کمیت محوری تعمیم یافته G_{σ_i} برای σ_i موقعی که $i = 1, 2, \mu_i$ پارامتر مزاحم است را معرفی می کنیم. با توجه به این که $\chi_{\Upsilon n}^2 \sim \frac{\Upsilon(n+1)\hat{\sigma}_1}{\sigma_1}$ ، باکلزی (۲۰۰۸) کمیت محوری تعمیم یافته برای σ_1 را به صورت زیر معرفی کرد

$$G_{\sigma_1} = G_{\sigma_1}(X, x, \hat{\mu}_1, \hat{\sigma}_1) = \frac{\sigma_1}{\Upsilon(n+1)\hat{\sigma}_1} \Upsilon(n+1)\hat{\sigma}_{10} \stackrel{D}{=} \frac{1}{\chi_{\Upsilon n}^2} \Upsilon(r_n - r_0).$$

شرایط کمیت محوری تعمیم یافته را دارد. از طرفی

$$\begin{aligned} 1 - \alpha &= P(h_{\alpha/2} < \frac{\sigma_1}{\sqrt{(n+1)\hat{\sigma}_1}} \sqrt{(n+1)\hat{\sigma}_1} < h_{1-\alpha/2}) \\ &= P(h_{\alpha/2} < \frac{1}{\chi_{2n}^2} (n+1)\hat{\sigma}_1 < h_{1-\alpha/2}) \\ &= P(\frac{(n+1)\hat{\sigma}_1}{h_{1-\alpha/2}} < \chi_{2n}^2 < \frac{(n+1)\hat{\sigma}_1}{h_{\alpha/2}}), \end{aligned}$$

بنابراین $\frac{(n+1)\hat{\sigma}_1}{h_{1-\alpha/2}} = \chi_{2n, \alpha/2}^2$ و $\frac{(n+1)\hat{\sigma}_1}{h_{\alpha/2}} = \chi_{2n, 1-\alpha/2}^2$ عبارتست از

$$\sigma_1 \in (\frac{(n+1)\hat{\sigma}_1}{\chi_{2n, 1-\alpha/2}^2}, \frac{(n+1)\hat{\sigma}_1}{\chi_{2n, \alpha/2}^2}).$$

به همین ترتیب یک کمیت محوری تعمیم یافته برای σ_2 عبارت است از

$$G_{\sigma_2} = G_{\sigma_2}(X, x, \hat{\mu}_2, \hat{\sigma}_2) = \frac{\sigma_2}{\sqrt{(m+1)\hat{\sigma}_2}} \sqrt{(m+1)\hat{\sigma}_2} \stackrel{D}{=} \frac{1}{\chi_{2m}^2} \chi^2(s_m - s_o).$$

و یک بازه اطمینان تعمیم یافته $(1-\alpha)\%$ برای σ_2 عبارت است از

$$\sigma_2 \in (\frac{(m+1)\hat{\sigma}_2}{\chi_{2m, 1-\alpha/2}^2}, \frac{(m+1)\hat{\sigma}_2}{\chi_{2m, \alpha/2}^2}).$$

۳.۴ کمیت محوری تعمیم یافته برای R

در رابطه (۳) مشاهده می شود که پارامتر تنش-مقاومت R تابعی از μ_1, μ_2, σ_1 و σ_2 است. بنابراین کمیت محوری تعمیم یافته برای R با جایگزینی کمیت های محوری $G_{\mu_1}, G_{\mu_2}, G_{\sigma_1}$ و G_{σ_2} بجای پارامترهای مربوطه در پارامتر R حاصل می شود. بنابراین به صورت زیر حاصل می شود

$$G_R = \begin{cases} \frac{G_{\sigma_2}}{G_{\sigma_1} + G_{\sigma_2}} e^{\frac{(G_{\mu_2} - G_{\mu_1})}{G_{\sigma_2}}} & G_{\mu_1} > G_{\mu_2} \\ 1 - \frac{G_{\sigma_1}}{G_{\sigma_1} + G_{\sigma_2}} e^{\frac{(G_{\mu_1} - G_{\mu_2})}{G_{\sigma_1}}} & G_{\mu_1} \leq G_{\mu_2}. \end{cases}$$

به طوری که

$$\begin{aligned} G_{\mu_1} &= \hat{\mu}_{1,o} - \frac{\chi_{2n}^2}{\chi_{2n}^2} (r_n - r_o), & G_{\mu_2} &= \hat{\mu}_{2,o} - \frac{\chi_{2m}^2}{\chi_{2m}^2} (s_m - s_o) \\ G_{\sigma_1} &= \frac{1}{\chi_{2n}^2} \chi^2(r_n - r_o), & G_{\sigma_2} &= \frac{1}{\chi_{2m}^2} \chi^2(s_m - s_o) \end{aligned}$$

با توجه به این که توزیع G_R یک توزیع پیچیده است، لذا نمی توان یک بازه اطمینان به فرم بسته برای R براساس کمیت محوری تعمیم یافته به دست آورد. از طرفی برای هر $(\hat{\mu}_{i,o}, \hat{\sigma}_{i,o})$ ، $i = 1, 2$ ، توزیع G_R به پارامترهای مجهول بستگی ندارد. بنابراین با استفاده از شبیه سازی مونت کارلو که توسط الگوریتم ۱ در زیر داده شده است، می توان حدود اطمینان R را به دست آورد.

الگوریتم ۱

۱. برای مجموعه داده های رکوردی داده شده، برآورد ML پارامترهای $\hat{\mu}_1, \hat{\mu}_2, \hat{\sigma}_1$ و $\hat{\sigma}_2$ را محاسبه می کنیم.

۲. مقادیر $\chi_1^2 \sim \chi_{2n}^2, Q_1 \sim \chi_1^2, Q_2 \sim \chi_1^2$ و $W_1 \sim \chi_{2m}^2$ و $W_2 \sim \chi_{2m}^2$ را تولید می کنیم.

۳. $G_{\sigma_1}, G_{\mu_1}, G_{\sigma_2}, G_{\mu_2}$ و G_R را محاسبه می‌کنیم.

۴. مرحله‌ی ۲ و ۳ را k بار تکرار می‌کنیم و مقادیر حاصل را به صورت $G_R^{(1)} \leq G_R^{(2)} \leq \dots \leq G_R^{(k)}$ مرتب می‌کنیم.

۵. حد پایین و بالا در این بازه به ترتیب $G_R^{(\lfloor (\alpha/2)k \rfloor)}$ و $G_R^{(\lfloor (1-\alpha/2)k \rfloor)}$ است، که در این‌ها $[x]$ بزرگ‌ترین عدد صحیح کوچک‌تر یا مساوی x است.

۵ مطالعه شبیه‌سازی

در این بخش با استفاده از یک مطالعه شبیه‌سازی براساس داده‌های رکوردی تولید شده از دو توزیع نمایی $Exp(\mu_1, \sigma_1)$ و $Exp(\mu_2, \sigma_2)$ ، به محاسبه‌ی بازه‌ی بیزی با دم‌های برابر $(EQBS)$ ، بازه‌ی تقریبی چگالی پسین رفیع $(AHPD)$ ارائه شده توسط چن و شائو (۱۹۹۹) و بازه‌ی اطمینان براساس کمیت محوری تعمیم‌یافته (GV) ، برای مقادیر مختلف پارامترها می‌پردازیم. در این شبیه‌سازی احتمال پوشش (CP) و طول مورد انتظار (EL) این بازه‌ها را با یکدیگر مقایسه می‌کنیم.

در این شبیه‌سازی از کل ترکیب‌های اندازه‌ی نمونه‌ی $n=4, 8, 12, 16$ و $m=4, 8, 12, 16$ با $n \leq m$ ، استفاده می‌کنیم. برای پارامتر R مقادیر $0/5, 0/65, 0/85$ و $0/95$ را در نظر می‌گیریم و پارامترهای $(\mu_1, \sigma_1, \mu_2, \sigma_2)$ را بگونه‌ای در نظر می‌گیریم که این مقادیر R توسط آن‌ها به دست آید. برای مثال با قرار دادن $\sigma_1 = 1, \sigma_2 = 2, \mu_1 = 1$ و $R = 0/65$ مقدار μ_2 از رابطه‌ی (۳) برابر و با قرار دادن $R = 0/95$ مقدار μ_2 برابر $2/897$ حاصل می‌شود. برای هر یک از این ترکیب پارامترها، 2000 داده‌ی رکوردی از توزیع‌های $Exp(\mu_1, \sigma_1)$ و $Exp(\mu_2, \sigma_2)$ تولید می‌کنیم. با قرار دادن ضریب اطمینان $0/9$ ، برای هر سری داده‌های رکوردی، بازه‌های اطمینان زیر را محاسبه می‌کنیم

(۱) $AHPD$: بازه‌ی تقریبی چگالی پسین رفیع ارائه شده توسط چن و شائو

(۲) $EQBS$: بازه‌ی اطمینان بیزی با دم‌های برابر

(۴) GV : بازه‌ی اطمینان براساس کمیت محوری تعمیم‌یافته که توسط الگوریتم

الگوریتم پیدا کردن بازه‌ی تقریبی چگالی پسین رفیع ارائه شده توسط چن و شائو $(AHPD)$ و بازه بیزی با دم‌های برابر $(EQBS)$:

عمومی‌ترین بازه‌ی بیزی با استفاده از بازه‌ی تقریبی چگالی پسین رفیع (HPD) است. این بازه، در میان بازه‌هایی با احتمال پوشش $1 - \alpha$ ، کوتاه‌ترین می‌باشد. به عبارت دیگر ناحیه‌ی پذیرش HPD آن ناحیه‌ای است که در دو شرط زیر صدق می‌کند (۱) احتمال پسین آن، $1 - \alpha$ باشد.

(۲) مقدار چگالی پسین هر نقطه در داخل آن ناحیه برابر یا بزرگ‌تر از مقدار چگالی پسین هر نقطه خارج از آن ناحیه باشد.

ساختن بازه‌ی HPD معمولاً دشوار است زیرا اغلب نقاط پایانی به شکل‌های بسته قابل محاسبه نیست. بنابراین روش‌های عددی که شامل انتگرال‌گیری عددی و یافتن ریشه‌های عددی است، استفاده می‌شود. چن و شائو (۱۹۹۹)، یک الگوریتم برای برآورد نقطه‌ی پایانی بازه‌های بیزی که در اجرا راحت است و به طور مجانبی دارای احتمال پوشش مورد نظر است، معرفی کردند. روش دیگر این است که بازه‌های اطمینان با دم‌های برابر را بکار ببریم که در آن در هر دو دنباله احتمال پوشش برابر $\frac{\alpha}{4}$ است. در این صورت تنها زمانی این بازه با بازه‌ی HPD برابر می‌شود که توزیع پسین پارامترها مورد نظر متقارن و تک‌مدی باشد.

بازه‌ی HPD را می‌توان بعد از مرتب کردن مقادیر تنش-مقاومت به صورت $R_{(1)} \leq R_{(2)} \leq \dots \leq R_{(M)}$ پیدا کرد.

براساس روش چن و شائو (۱۹۹۹)، بازه‌ی اطمینان بیزی با دم‌های برابر به شکل زیر است

$$(R_{[(\alpha/2)M]}, R_{[(1-\alpha/2)M]}).$$

همچنین بازه‌ی HPD به شکل کوتاه‌ترین بازه به شکل $(R_{(j)}, R_{(j+[(1-\alpha)M])})$ است که در شرط زیر صدق می‌کند

$$R_{(j+[(1-\alpha)M])} - R_{(j)} = \min_{1 \leq j^* \leq M - [(1-\alpha)M]} (R_{(j^*+[(1-\alpha)M])} - R_{(j^*)}).$$

در محاسبه‌ی بازه‌ی GV برای هر سری داده‌ی رکوردی الگوریتم ۱ را به تعداد $k=1000$ بار تکرار می‌کنیم و میانگین تعداد بازه‌هایی که R را در بردارند به عنوان احتمال پوشش (CP) بازه و میانگین طول بازه‌ها را به عنوان طول مورد انتظار (EL) بازه در نظر می‌گیریم. برای بازه‌های اطمینان AHPD و EQBS نیز تعداد تکرار تولید نمونه از توزیع‌های پسین μ_1, μ_2, σ_1 و σ_2 را برابر $M = 1000$ در نظر گرفته و همانند حالت قبل احتمال پوشش و طول مورد انتظار را براساس این ۱۰۰۰ تکرار به دست می‌آوریم. جدول‌های ۱ تا ۲ احتمال پوشش و طول مورد انتظار این سه بازه را برای مقدار سطح اطمینان ۰/۹۰ به دست می‌دهند.

با توجه به جدول‌های ۱ تا ۲ نتیجه می‌شود که احتمال پوشش تمام بازه‌ها با افزایش اندازه‌ی نمونه افزایش می‌یابد. طول مورد انتظار نیز در تمام بازه‌ها، با افزایش مقدار واقعی پارامتر قابلیت اطمینان تنش-مقاومت یا با افزایش اندازه‌ی نمونه، کاهش می‌یابد. با در نظر گرفتن احتمال پوشش بازه‌ها، بازه‌ی EQBS برای مقادیر R نزدیک به ۰/۵ محافظه‌کارانه است و مقدار احتمال پوشش را زیاد به دست می‌دهد. اما با حرکت به سمت مقادیر بزرگ‌تر R ، این احتمال پوشش به مقدار واقعی نزدیک‌تر می‌شود. بازه‌های HPD برای تمام مقادیر R تقریباً غیر محافظه‌کارانه هستند و احتمال پوشش کم‌تری را نسبت به بازه‌های EQBS و GV به دست می‌دهند.

همچنین با در نظر گرفتن طول مورد انتظار، بازه‌های HPD طول کوتاه‌تری نسبت به بازه‌های EQBS دارند و بازه‌های EQBS نیز طول کوتاه‌تری نسبت به بازه‌های GV دارند. در حالی که دو بازه‌ی EQBS و GV دارای احتمال پوشش تقریباً یکسانی هستند. در نتیجه، در حالت کلی، به نظر می‌رسد که بازه‌ی بیزی با دم‌های برابر بهترین عملکرد را از نظر احتمال پوشش و طول مورد انتظار دارند.

۶ خلاصه و نتیجه‌گیری

در این مقاله، با توجه به مفهوم داده‌های رکوردی، برآورد نقطه‌ای و بازه‌ای را برای پارامتر تنش-مقاومت در حالت‌هایی که پارامتر مکانی و مقیاسی بدون محدودیت باشند در توزیع نمایی دو پارامتری محاسبه شد. برای مقایسه‌ی کارایی بازه‌های مختلف شبیه‌سازی انجام داده و نتیجه گرفته شد که بازه اطمینان بیزی HPD دارای کوتاه‌ترین بازه است در حالی که بازه‌ی GV دارای طولانی‌ترین بازه است. همچنین نتیجه گرفته شد که احتمال پوشش تمام بازه‌ها با افزایش اندازه‌ی نمونه، افزایش می‌یابد و طول مورد انتظار کاهش می‌یابد.

مراجع

- [1] Baklizi, A. (2008a), Estimation of $P(X < Y)$ Using Record Values in the One and Two-Parameter Exponential Distributions, *Communications in Statistics-Theory and Methods*, **31**, 692-698.

جدول ۱: احتمال پوشش شبیه‌سازی شده، زمانی که $\alpha = 0.90$.

n	m	R	EQBS	HPD	GV
۴	۴	۰/۵۰	۰/۹۴۳۸	۰/۸۷۶۶	۰/۹۵۸۲
۴	۴	۰/۶۵	۰/۹۴۰۱	۰/۸۸۶۳	۰/۹۴۹۸
۴	۴	۰/۸۰	۰/۹۴۰۲	۰/۸۵۳۰	۰/۹۴۳۸
۴	۴	۰/۹۵	۰/۹۰۸۱	۰/۸۸۳۹	۰/۹۳۲۰
۴	۸	۰/۵۰	۰/۹۲۶۹	۰/۸۴۵۶	۰/۹۴۳۸
۴	۸	۰/۶۵	۰/۹۲۷۷	۰/۸۵۲۵	۰/۹۴۱۵
۴	۸	۰/۸۰	۰/۹۰۵۸	۰/۸۳۵۳	۰/۹۳۸۷
۴	۸	۰/۹۵	۰/۹۱۵۷	۰/۸۷۳۱	۰/۹۱۳۴
۴	۱۲	۰/۵۰	۰/۹۱۵۰	۰/۸۴۳۶	۰/۹۳۴۶
۴	۱۲	۰/۶۵	۰/۹۱۶۸	۰/۸۳۴۴	۰/۹۲۶۴
۴	۱۲	۰/۸۰	۰/۹۱۵۹	۰/۸۷۹۲	۰/۹۲۰۱
۴	۱۲	۰/۹۵	۰/۹۲۳۸	۰/۸۲۷۶	۰/۹۲۲۴
۴	۱۶	۰/۵۰	۰/۹۲۹۸	۰/۸۴۱۶	۰/۹۳۰۸
۴	۱۶	۰/۶۵	۰/۹۱۵۵	۰/۸۳۲۹	۰/۹۴۰۴
۴	۱۶	۰/۸۰	۰/۹۱۸۷	۰/۸۷۸۰	۰/۹۲۴۳
۴	۱۶	۰/۹۵	۰/۹۱۵۸	۰/۸۲۸۷	۰/۹۸۸۰
۸	۸	۰/۵۰	۰/۹۱۱۳	۰/۸۱۵۶	۰/۹۲۷۳
۸	۸	۰/۶۵	۰/۹۰۶۵	۰/۸۱۸۳	۰/۹۲۳۴
۸	۸	۰/۸۰	۰/۹۱۲۳	۰/۸۸۱۶	۰/۹۱۸۹
۸	۸	۰/۹۵	۰/۹۴۳۸	۰/۸۱۹۹	۰/۹۱۹۹

جدول ۲: طول مورد انتظار شبیه سازی شده، زمانی که $\alpha = 0.9$ - ۱.

n	m	R	EQBS	HPD	GV
۴	۴	۰/۵۰	۰/۷۷۸۶	۰/۷۰۴۱	۰/۹۰۴۴
۴	۴	۰/۶۵	۰/۷۷۳۷	۰/۶۹۴۷	۰/۸۹۶۶
۴	۴	۰/۸۰	۰/۷۶۴۳	۰/۶۶۲۰	۰/۸۹۱۶
۴	۴	۰/۹۵	۰/۷۴۷۹	۰/۶۱۶۸	۰/۸۸۲۶
۴	۸	۰/۵۰	۰/۷۶۷۷	۰/۶۹۷۹	۰/۸۸۰۴
۴	۸	۰/۶۵	۰/۷۶۶۶	۰/۹۸۱۶	۰/۸۸۰۲
۴	۸	۰/۸۰	۰/۷۶۰۱	۰/۹۴۹۲	۰/۸۸۴۶
۴	۸	۰/۹۵	۰/۷۴۸۶	۰/۶۱۶۸	۰/۸۸۳۰
۴	۱۲	۰/۵۰	۰/۷۷۴۳	۰/۶۹۴۷	۰/۸۶۶۷
۴	۱۲	۰/۶۵	۰/۷۵۷۰	۰/۶۶۴۳	۰/۸۶۷۹
۴	۱۲	۰/۸۰	۰/۷۴۸۹	۰/۶۴۲۰	۰/۸۷۲۵
۴	۱۲	۰/۹۵	۰/۷۵۰۱	۰/۶۱۸۹	۰/۸۸۲۶
۴	۱۶	۰/۵۰	۰/۷۷۰۹	۰/۶۸۷۶	۰/۸۵۹۱
۴	۱۶	۰/۶۵	۰/۷۶۸۷	۰/۶۷۶۸	۰/۸۷۷۸
۴	۱۶	۰/۸۰	۰/۷۵۶۶	۰/۶۵۰۳	۰/۸۷۶۹
۴	۱۶	۰/۹۵	۰/۷۴۹۲	۰/۶۱۵۷	۰/۸۸۳۵
۸	۸	۰/۵۰	۰/۷۷۴۶	۰/۶۹۲۳	۰/۸۵۳۱
۸	۸	۰/۶۵	۰/۷۷۰۸	۰/۶۸۲۴	۰/۸۴۸۳
۸	۸	۰/۸۰	۰/۷۷۲۸	۰/۶۷۰۱	۰/۸۴۶۶
۸	۸	۰/۹۵	۰/۷۸۲۷	۰/۶۵۶۸	۰/۸۵۸۷

- [2] Baklizi, A. (2008b), Likelihood and Bayesian Estimation of $\Pr(X < Y)$ Using Lower Record Values from the Generalized Exponential Distribution, *Computational Statistics and Data Analysis*, **52**, 3468-3473.
- [3] Baklizi, A. (2013), Interval Estimation of Stress-Strength Reliability in the Two-Parameter Exponential Distributions Based on Records, *Journal of Statistical Computation and Simulation*, 1-10.
- [4] Baklizi, A. (2014), Bayesian Inference for $P(X < Y)$ in the Exponential Distribution Based on Records, *Applied Mathematical Modelling*, **38**, 1698-1709.
- [5] Chen, M. and Shao, Q. (1999), Monte Carlo Estimation of Bayesian Credible and AHPD Intervals, *Journal of Computational and Graphical Statistics*, **8**(1), 69-92.

بررسی عوامل موثر بر امید به زندگی بیماران مسلول با داده‌های ناقص شهر زاهدان، سال‌های ۱۳۹۵-۱۳۹۳

مژگان سالاری^۱، محمدحسین دهقان

گروه آمار، دانشکده علوم ریاضی، دانشگاه سیستان و بلوچستان

چکیده: این مطالعه بنا به ضرورت منطقه ای و شیوع بیماری مسری سل مخصوصا اسمیر مثبت، بر روی ۲۵۰ بیمار مسلول انجام شده است. با توجه به همجواری این استان با کشورهای شرقی کم برخوردار از امکانات و اطلاعات بهداشتی، بیمارستان و کلینیک های پزشکی، شرایط رفاهی، ورود افراد خارجی غیر قانونی و بالا بودن ورودی-خروجی استان را می توان به عنوان ضرورت های انجام این تحقیق نام برد. این مطالعه با هدف تعیین میزان تابع بقا (احتمال زنده ماندن) بیماران مبتلا به سل ریوی اسمیر مثبت به شرط عوامل موثر بر آن، روی بیماران مبتلا به سل که به مرکز هماهنگ کننده سل شهرستان زاهدان مراجعه کرده اند، در فاصله سال های ۹۳ تا خرداد ۹۵ صورت گرفته است. به منظور نشان دادن معنی داری تفاوت بین توابع بقا افراد مسلول با توجه به بیماری با شرایط خاص دیگر، علاوه بر برآورد، رسم نمودارها، آزمون های آماری انجام شده است. در این بررسی وزن بیماران را به عنوان متغیر کمکی موثر در برآورد تابع بقا شرطی در نظر گرفته ایم و تاثیر عواملی هم چون جنسیت، ابتلا همزمان به دیابت، ایدز، سابقه قبلی بیماری، تماس با فرد مسلول و سابقه زندان را بر تابع بقا افراد مورد بررسی قرار گرفت که دارای اثر معنی دار می باشند. **واژه‌های کلیدی:** تابع بقاء شرطی، بیماری سل اسمیر مثبت، داده های ناقص، عوامل خطر. **کد موضوع بندی:** 39B82، 34K20، 39B52، 47A55.

۱ مقدمه

بیماری سل یک بیماری عفونی مسری می باشد که توسط میکروب مایکوباکتریوم توبرکلوزیس در انسان ایجاد می شود. باسیل سل می تواند در بافت های بدن به حالت خفته در آمده و برای سالها در همین وضعیت بماند. بیماری سل بزرگترین علت مرگ ناشی از بیماری های عفونی تک عاملی است (حتی بیشتر از ایدز، مالاریا و سرخک). مخزن و عامل انتشار بیماری، فرد مسلول ریوی اسمیر مثبت می باشد که با سرفه و عطسه قطرات ریز حاوی میکروب را در هوای پیرامون خویش پخش می نماید. هر سرفه قادر است تا ۳۰۰۰ ذره عفونی را تولید کند. این ذرات می توانند از طریق صحبت کردن، عطسه، تف کردن و آواز خواندن در هوا منتشر شده و مدتها به صورت معلق در هوا باقی بمانند و با استنشاق توسط افراد دیگر، داخل ریه، آنها شده و در صورت وجود زمینه مساعد در فرد، ایجاد آلودگی یا عفونت نماید [۱] [۱۰].

از دیگر راههای انتقال این بیماری، می توان از انتقال پوستی در حین کار با مواد آلوده (کار در مراکز بهداشتی)، انتقال از مادر به کودک و نیز انتقال از طریق شیر دام نام برد، ولی این راههای انتقال بیماری بسیار نادر بوده و در مقایسه با انتقال از طریق سرفه و عطسه فرد مسلول، اهمیت کم تری دارند. جالب اینکه بیماری سل ریوی از طریق غذا، آب، تماس جنسی، تزریق خون یا نیش حشرات منتقل نمی شود. مکان های پر ازدحام، سر بسته، کم نور، بدون تهویه مناسب و مرطوب، بهترین شرایط را برای تسهیل انتقال عفونت

ایجاد می کنند. این شرایط را می توان در زندان ها، آسایشگاه ها و پناهگاه ها به وفور یافت [۹].

میکروب سل، در گرد و غبار معلق ۸ تا ۱۰ روز، در خاک سرد و سایه دار ۶ ماه و در خلط در حال پوسیدن، هفته ها و ماهها مقاومت کرده و زنده می ماند. لذا هوای اتاق آلوده شده به میکروب سل توسط بیمار عفونت زا، می تواند حتی در غیاب بیمار نیز موجب انتقال میکروب گردد. تابش مستقیم نور آفتاب در عرض ۵ دقیقه، میکروب سل را از بین می برد، لذا در مناطق گرمسیری، تماس مستقیم اشعه آفتاب، روش مناسبی برای از بین بردن میکروب سل است.

سل می تواند کلیه ارگانهای بدن را درگیر نماید، ولی شایع ترین شکل بیماری، سل ریوی می باشد. راه انتقال تقریباً همیشه از راه تنفس است، ولی میکروب سل پس از ورود به ریه و ایجاد ضایعه اولیه، می تواند از طریق جریان خون، عروق لنفاوی و یا به علت مجاورت اعضا، مستقیماً به دیگر قسمت های بدن منتشر شود. بدین ترتیب بیماری به دو شکل در انسان ظاهر می شود [۳].

سل ریوی: ۸۰ درصد موارد ابتلا به سل را شامل می شود و دارای دو گونه است:

سل ریوی اسمیر مثبت: با توجه به خطر سرایت تنفسی بالا به اطرافیان در تماس، بیشترین اهمیت را در بحث مبارزه و کنترل بیماری سل دارد، چرا که یک بیمار اسمیر مثبت در صورتی که درمان نشود، طی سال اول قادر است ۱۵ تا ۲۰ نفر را آلوده سازد. سل ریوی منفی، خطر سرایت کمتری دارد.

سل خارج ریه: ابتلا سایر ارگانهای بدن بر اثر انتشار خونی، عروق لنفاوی و یا سرایت مجاورتی از یک عضو به عضو دیگر را شامل می شود. این فرم بیماری نیز به ندرت خطر سرایت از یک فرد به فرد دیگر را دارد.

بیماری سل در بار جهانی بیماری ها، دارای رتبه دهم می باشد و انتظار می رود تا سال ۲۰۲۰ تا رتبه هفتم بالا برود. بر طبق آمار سازمان جهانی بهداشت سالیانه تقریباً دو میلیون نفر در اثر بیماری جان خود را از دست می دهند [۱۵]. بیش از ۸۰ درصد مبتلایان، متعلق ۲۲ کشور در حال توسعه می باشد. میزان بروز و شیوع بیماری سل در ایران بالاست [۱۱]، [۸].

استان سیستان و بلوچستان به علت وجود عوامل و شرایط مساعد برای ماندگاری بیماری سل، از دیر باز میزبان بیماری سل بوده و علیرغم تمام تلاش های وسیع و ارزنده ای که توسط سیستم بهداشتی و درمانی صورت گرفته، هنوز بالاترین آمار بروز بیماری سل را در کشور دارا می باشد [۱۱]. عوامل موثری چون میزان بالای مهاجرت از کشورهای آلوده منطقه (افغانستان و پاکستان)، فقر و عوارض آن (سوء تغذیه)، اعتیاد، سطح بهداشت نامطلوب، افزایش آمار ایدز در سالهای اخیر و وضعیت جمعیتی و فرهنگی منطقه، نشان از مسیر طولانی در راستای کنترل بیماری سل دارد. به دلیل شیوع بالای سل ریوی اسمیر مثبت در بین انواع سل، ما در این پژوهش به مطالعه این نوع خاص از بیماری پرداخته ایم.

اهدافی که این مطالعه دنبال می کند عبارتند از برآورد و تحلیل تابع بقاء زمان بهبودی بیماری سل به شرط متغیر کمکی وزن که در بیماری سل موثر است و مقایسه آنها بین گروه مرد و زن، افرادی با سابقه بیماری و بدون سابقه بیماری سل، بیمار با افراد مسلول در تماس بوده یا نبوده، افراد سابقه زندان داشته یا نداشته اند و افراد با سابقه بیماری دیابت و بدون سابقه دیابت می باشد. به منظور دستیابی به این نتایج از نرم افزار R و بسته های نرم افزاری تابع بقا غیر شرطی کاپلان-مایر ($SurvivalFunction$) و بسته نرم افزاری جامع محاسبه تابع بقا ناپارامتری شرطی (GTE) و تابع کرنل نرمال استفاده شده است [۶].

۲ روش کار

این مطالعه با هدف تعیین میزان بقا (احتمال زنده ماندن) بیماران مبتلا به سل ریوی اسمیر مثبت و عوامل موثر بر آن صورت گرفته است و بر روی بیماران مبتلا به سل مراجعه کننده به مرکز سل شهرستان زاهدان، در فاصله سال های ۹۳ تا خرداد ۹۵ صورت گرفته

است. ابتدا اطلاعات مربوط به پرونده ۲۵۰ بیمار مبتلا به سل ریوی اسمیر مثبت جمع آوری گردیده و معیار تشخیص بیماری در افراد مورد مطالعه، مشاهده میکروسکوپی خلط بوده است و تمامی بیماران رادیوگرافی قفسه سینه داشته اند و توسط پزشک متخصص عفونی معاینه گردیده و بیماری سل ریوی اسمیر مثبت برای این بیماران تشخیص و تایید شده است. در این مطالعه برای تحلیل داده ها به متغیرهای اصلی و کمکی نیاز داریم، لذا متغیرهای اصلی و کمکی مورد مطالعه را در ذیل بررسی می کنیم:

زمان بهبودی فرد مسلول یعنی متغیر اصلی، ممکن است به متغیرهای کمکی مانند وزن بیمار در زمان تشخیص، گروه خونی، نژاد، جنسیت، ملیت، ابتلا به ویروس *HIV*، ابتلا به دیابت، سابقه تماس با فرد مسلول، سابقه ابتلا قبلی به بیماری سل، سابقه حضور بیمار در زندان یا آسایشگاه بستگی داشته باشند [۱۴]. وجود بعضی از متغیرهای کمکی در درمان و بهبود بیماری یا مقاومت در برابر بیماری موثر می باشد که به عنوان متغیر کمکی در نظر گرفته می شود، ما در این مطالعه وزن فرد بیمار را به عنوان متغیر کمکی در نظر گرفته ایم [۷].

۳ زمان بقا

در این مطالعه، فاصله زمانی بین شروع درمان تا قطع درمان به عنوان دوره مطالعه هر بیمار در نظر گرفته شده است، که با دو نوع سانسور مواجه هستیم:

- ۱- افرادی که در این فاصله زمانی به علتی غیر از بهبودی مثل مرگ، تغییر شهر زندگی و غیره، از دوره درمان خارج شده اند و درباره سرنوشت بیماری آنها اطلاعی نداریم، سانسور از راست در آخرین ویزیت مثبت (وجود بیماری) می نامیم.
- ۲- برای افرادی که بهبود یافته اند، آخرین ویزیت مثبت را L_i (وجود بیماری) و اولین نتیجه منفی R_i (عدم وجود بیماری) را بصورت سانسور فاصله ای در نظر گرفته و $(L_i, R_i; i: 1, \dots, 250)$ نشان می دهیم.

در مطالعات تابع بقاء، اگر داده های دقیق و سانسور از راست داشته باشیم عموماً برآوردگر کاپلان مایر برای برآورد تابع بقا استفاده می شود [۵]. در صورتی که داده ها شامل داده های دقیق، سانسور فاصله ای، سانسور از راست باشند این برآوردگر قادر به محاسبه تابع بقا نمی باشد، در نتیجه از برآوردگر جدیدی به نام برآوردگر ناپارامتری تابع بقا شرطی استفاده که در سال ۲۰۱۱، توسط م. ح. دهقان-ت دوشان (*GTE*) ارائه شد استفاده می کنیم [۵]. این برآوردگر از متغیر کمکی و تابع هسته برای افزایش دقت بیشتر برآورد تابع بقا استفاده می کند که نسبت به برآوردگر کاپلان مایر و برآوردگر ترن بول غیر شرطی، بسیار دقیق تر بوده و مناسب برای انواع داده های دقیق سانسور شده و یا مخلوطی از آن ها می باشد. علاوه بر این می توان اثر متغیرهای مختلف را بر روی برآوردگر تابع بقا مشاهده و توجه خودمان را بر نقاطی خاص از متغیر کمکی متمرکز کرده ایم. برآوردگر مذکور تابع بقا ناپارامتری شرطی و بصورت زیر تعریف می شوند:

$$\hat{S}_{Gkm}^h(t|(z_0)) = \prod_{\{i: L_i \leq t\}} \left\{ 1 - \frac{\mathbb{1}_{\{R_i < \infty\}} w_i^h(z_0)}{\sum_{k: L_k > L_K} w_k^h(z_0)} \right\} \quad (1)$$

که در آن $w_i^h(z_0)$ وزن متناسب با هر داده می باشد:

$$\hat{w}_i^h(z_0) = \frac{1}{h} * \left(\frac{k\left(\frac{z_i - z_0}{h}\right)}{\sum_{L=1}^k \frac{1}{h} k\left(\frac{z_i - z_0}{h}\right)} \right) \quad (2)$$

در روابط فوق T زمان (بر حسب روز، هفته، ماه، ...) z_0 مقداری خاص از متغیر کمکی، h پهنای باند و $k(\omega)$ تابع هسته (کرنل) می باشد [۵][۴]. در این مطالعه زمان بر حسب روز، متغیر کمکی وزن بر حسب کیلوگرم و تابع کرنل گوسین (هسته نرمال) استفاده شده است [۲].

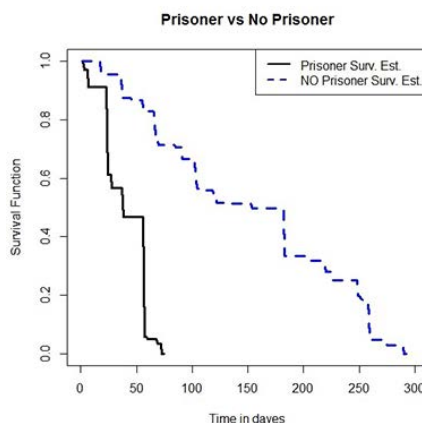
۴ شرایط یافته ها، تجزیه و تحلیل داده ها

تعداد بیماران مورد بررسی ۲۵۰ نفر می باشند که با استفاده از بانک اطلاعاتی معاونت بهداشتی علوم پزشکی زاهدان، مرکز هماهنگ کننده سل شهرستان زاهدان دارای پرونده درمانی استفاده شده است. تعداد ۴ بیمار سانسور از راست و باقی بیماران ما سانسور فاصله ای می باشند.

در مطالعه حاضر ۱۶ درصد از بیماران دارای سابقه زندان بوده اند. همانطور که مشاهده می شود بیماران غیر زندانی بدنشان از مقاومت بالاتری نسبت به افراد زندانی برخوردار است و در نتیجه از شانس بالاتری برای بهبود نسبت به افراد زندانی برخوردارند. در مطالعات انجام شده در سطح جهانی مشخص گردیده است که میزان بروز سل در زندان ها صد برابر میزان آن در جامعه عادی است، چرا که زندان ها به دلیل نداشتن فضای مناسب بهداشتی، تهویه و عدم رعایت بهداشت فردی انتقال و انتشار بیماری را به شدت تسهیل کرده است به طوری که زندان ها به صورت مخازنی از بیماری سل درآمده اند [۱۳].

همانطور که در شکل (۱) و جدول (۱) ملاحظه می شود، این تفاوت احتمال بقا این دو گروه کاملاً واضح است.

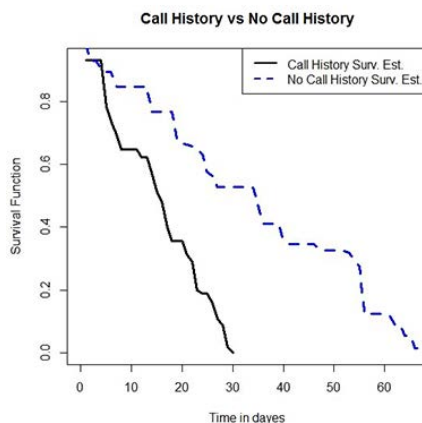
از سوی دیگر میزان بالای سل در زندان ها نقش بسزایی در انتشار این بیماری در سطح جامعه داشته است به دلیل تراکم بسیار زیاد جمعیت زندانی ها و نیز بازگشت زندانیان آزاد شده دوباره به زندان ها به همراه عوامل زمینه ساز چون عدم پایداری به رعایت شرایط بهداشتی و پیگیری سلامتی دست به دست هم داده اند تا سل حد و مرزهای ساخته شده در دو سوی زندان ها در نورد و زندان را به معضلی برای کنترل سل در همه کشورها تبدیل کند.



شکل ۱: برآورد تابع بقا شرطی بیماران دارای سابقه زندان و بیماران بدون سابقه

در این مطالعه ۱/۲ درصدی از افراد علاوه بر بیماری سل، به ایدز مبتلا هستند. به دلیل تضعیف سیستم دفاعی در بدن فرد مبتلا به ایدز، باکتری سل به دلیل مقاوم بودن به آنتی بیوتیک ها نسبت به بقیه باکتری ها، حالت فرصت طلبی پیدا کرده و در بدن افراد مبتلا به ایدز فعال شده و بیماری سل را تسریع می کند. لازم به توضیح است در گزارش مرکز بهداشت جهانی آمده است، بیش از ۵۰ درصد افرادی که قبلا به سل مبتلا و درمان شده اند بعدا به ویروس ایدز مبتلا شدند، مجددا بیماری سل آنها عود کرده است [۱۵]. این عارضه به ویژه در بین معتادان به مواد مخدر تزریقی شایع تر است.

۲۸/۲ درصد از افراد تحت بررسی این مطالعه با فرد آلوده به سل در ارتباط بوده اند، همانطور که قبلا گفته شد اصلی ترین عامل انتقال بیماری سل، فرد مسلول می باشد. در شکل (۲) برآورد تابع بقای شرطی بیمارانی که با فرد مسلول در ارتباط بوده اند و بیمارانی که در ارتباط نبوده اند رسم شده است.

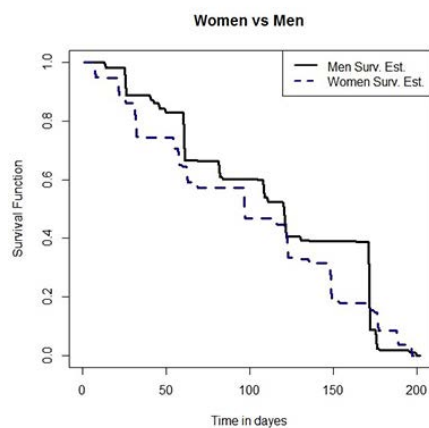


شکل ۲: برآورد تابع بقای شرطی بیمارانی مرتبط با فرد مسلول و بیمارانی بدون ارتباط

باید به این نکته توجه داشت که میزان ابتلاء زنان در تمام گروه های سنی مورد مطالعه، بیشتر از مردان بوده است. در این مطالعه اثر عواملی چون جنسیت، سابقه زندانی داشتن، سابقه ابتلاء قبلی، مبتلا به دیابت بودن و سابقه تماس با فرد مسلول مورد بررسی قرار داده و معنی داری آنها سنجیده شده است. ۵۱ درصد بیمارانی را زنان تشکیل می دهند و در شکل (۳) و جدول (۱) تفاوت تابع بقا برای مردان و زنان را ملاحظه خواهید کرد.

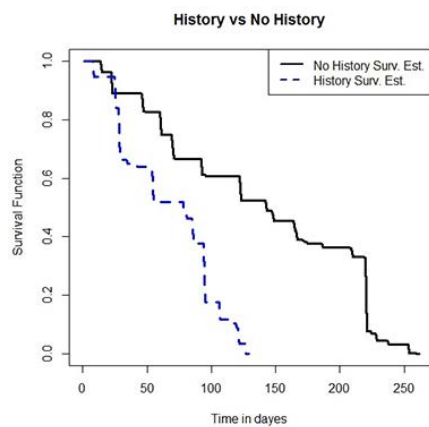
با توجه به میزان بالای ابتلای زنان به بیماری سل پیشنهاد می شود که در مناطق با شیوع بالای سل، اقدام به ساخت خانه ها و مدارس افتاب گیر به منظور استفاده از نور خورشید، برای از بین بردن میکروب مربوطه در نظر گرفته شود.

در این مطالعه ۲۹/۲ درصد از افراد تحت بررسی، سابقه قبلی ابتلا به بیماری سل را داشته اند. در شکل (۴) مشاهده می کنیم که



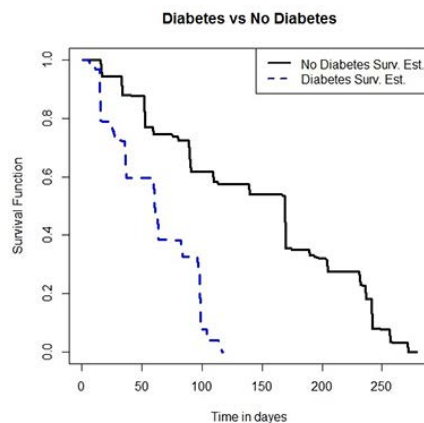
شکل ۳: برآورد تابع بقاشرطی زنان در برابر مردان

سابقه قبلی بیماری سل در روند درمان نسبت به افراد مسلولی که سابقه ابتلا قبلی ندارند تاثیر منفی دارد، چون افراد مسلول با سابقه قبلی بیماری بدنشان مقاومت کم تری در برابر بیماری از خود بروز می دهند.



شکل ۴: برآورد تابع بقا شرطی افراد مسلول با سابقه قبلی بیماری و بیماران بدون سابقه قبلی

از سال ۱۳۹۳، مراکز سل استان ها موظف شدند که علاوه بر آزمایش *HIV*، آزمایش دیابت نیز از بیماران سل گرفته شود که در این بررسی ۲۶ درصد از بیماران، مبتلا به بیماری دیابت بوده اند. در شکل (۵) برآورد تابع بقا شرطی بیماران مبتلا به دیابت و بیماران سل بدون دیابت رسم شده است.



شکل ۵: برآورد تابع بقا بیماران مبتلا به دیابت و بیماران عادی

در این مطالعه در بررسی نمودارهای توابع بقا اثر جنسیت، سابقه ابتلاء هم زمان به دیابت، سابقه قبلی بیماری، تماس با فرد مسلول و سابقه زندان داشتن، بر زمان بهبودی مشهود است، بعلاوه بررسی اثرات این عوامل بر زمان بهبودی، ابتدا یک دنباله مشترک منظم از دامنه تغییرات طول عمر در نظر گرفته با استفاده از آزمون استودنت تفاوت بین تابع بقا آن ها سنجیده شده است و نتایج را می توان در جدول شماره (۱) ملاحظه کرد.

جدول ۱: نتایج آزمون T

موضوع	T-Value	$CI_{0.85}(\mu_1 - \mu_2)$	P-Value	معنی داری
۱ مرد-زن	۱۳/۴۱	(۰/۰۸۹، ۰/۱۲۰)	$2/2e-16 < 0.001$	معنی داری
۲ با بیماری دیابت	۱۳/۹	(۰/۱۳۶، ۰/۱۸۰)	$2/2e-16 < 0.001$	معنی داری
۳ با سابقه بیماری	۴/۲۹	(۰/۰۴۶، ۰/۱۲۴)	$2/2e-5 < 0.001$	معنی داری
۴ تماس با مسلول	۱۹/۲۹	(۰/۳۳۳، ۰/۴۱۰)	$2/2e-16 < 0.001$	معنی داری
۵ با سابقه زندان	۵/۸۷	(۰/۰۶۳، ۰/۱۲۷)	$2/2e-8 < 0.001$	معنی داری

جدول ۲: توزیع وزن و درصد بیماران برحسب عوامل مذکور

۱	۲	۳	۴	۵
عامل	با دیابت-بدون دیابت	با سابقه بیماری-بدون سابقه	تماس با فرد مسلول-بدون تماس	سابقه زندان-بدون سابقه
میانگین وزن	۵۳/۵۵-۵۱/۱۹	۵۲/۱۸-۵۲/۷۸	۵۲/۳۴-۵۲/۳۲	۵۲/۶۵-۵۲/۰۹
انحراف معیار	۱۱/۳۴-۱۱/۹۴	۱۱/۸۳-۱۱/۳۸	۱۱/۶۱-۱۱/۹۷	۹/۶۲-۱۲/۰۵
توزیع درصدی	۴۹-۵۱	۷۴-۲۶	۷۱-۲۹	۸۴-۱۶

۵ دست‌آوردهای پژوهش

در این بررسی وزن بیماران را به عنوان متغیر کمکی موثر در برآورد تابع بقا شرطی (احتمال بقا شرطی) در نظر گرفته ایم و تاثیر عواملی هم چون جنسیت، ابتلاء هم زمان به دیابت، ایدز، سابقه بیماری قبلی، تماس با فرد مسلول و سابقه زندان بر تابع بقا افراد مورد بررسی قرار داده ایم. همانطور که ملاحظه می شود (با توجه به شکل های (۵-۱) و جدول شماره (۱)) جنسیت فرد مسلول، سابقه زندان فرد بیمار، سابقه ابتلاء هم زمان به دیابت، سابقه تماس با بیمار مسلول، سابقه ابتلاء قبلی بیماری، بر دوره درمان و زمان بهبودی بیماری بطور قابل ملاحظه ای موثر می باشند. با توجه به بررسی انجام شده مهم ترین عامل پیشگیری بیماری سل را آموزش جامعه در برابر این بیماری می توان دانست و شایسته است کلیه مدارس به ویژه در مدارس دخترانه کلاس های آموزشی-بهداشتی در رابطه با پیشگیری بیماری سل برای محصلین برگزار کنند. با توجه به نقش مراکز مذهبی در زندگی مردم این شهرستان از همکاری این مراکز با مراکز بهداشتی در ترویج راههای پیشگیری این بیماری می توان استفاده کرد. بیماران مبتلا به دیابت مستعد را شناسایی و آزمایش خلط را انجام داده و در صورت وجود بیماری اقدامات لازم را پیگیری شود. برای افرادی که در ارتباط نزدیک با فرد مسلول هستند، جلساتی جهت مشاوره برای پیشگیری در نظر گرفته شود. در زندان ها در بدو ورود از افراد آزمایش خلط گرفته شود و کلاس های آموزشی برای آنها برگزار گردد.

مراجع

- [1] Azizi, F., Hatami, H. and Janghorbani, M. (2004), *Epidemiology of Common disorders in Iran*, Tehran, second edition, 602-609.
- [2] Altman, N.S. (1990), An introduction to kernel and nearest neighbor nonparametric regression, *The American Statistics*, **46**, 175-185.
- [3] Arsng, Sh., Kazemnejad, A. and Amany, A. (1390), Epidemiology of TB in Iran, *Journal of Gorgan university of Medical Sciences*, **3**, 86-87.
- [4] Beran, R. (1980), *Nonparametric regression with randomly censored survival data*, Ams 5.
- [5] Dehghan, M.H. and Duchesne, T. (2011a), A generalization of Turnbulls estimator for nonparametric estimation of the conditional survival function with interval-censored data, *Lifetime data analysis*, 234-255.
- [6] Dehghan, M.H. and Duchesne, T. (2015), *GTE, R Packages*, <http://127.0.0.1:26911/library/gte.html>.

- [7] Shoraka, H.R., Hosseini, SH. and Alizadeh, H. (2011), Epidemiology of tuberculosis and other related factors in the province of North KHorasan, Iran, 2005-2010, *Nkh Uni Med Sci*, **3**(3), 43-50.
- [8] Sofia, M., Zarinfar, N., Mirzaee, M. and Moosavinejad, A. (2009), *Epidemiology of tuberculosis in Aeak*, Iran, Koomesh, **10**(4), 261-266.
- [9] Gholami, A., Gharehaghaji, R., Moosavi-Jahromi, L. and Sadaghiyanifar, A. (2009), Epidemiologic Survey of Pulmonary Tuberculosis in Urmia city during 2004-2007, *Knowledge & Health* , **4**(3), 19-23.
- [10] Hatami, H., Razavi, S.M. and Eftekhari-Ardabili, H. (2009), *Text book of public health*, 1121-1129, Volume 2, Second edition, Arjmand press, Tehran.
- [11] Heshmati, H., Ravanbakhsh, K., Khajavi, S. and Behnampour, N. *Epidemiological study of TB in the city Galkysh between 1385-90*, **1**, 61-65.
- [12] Kaplan, E. and Meier, P. (1958), Nonparametric estimation from incomplete observation, *J. Anner. Statist*, **53**, 475-481.
- [13] Mtant, M., Sharifimood, B., Alavi Nani, R. and Aminian Fard, M. (1390), Epidemiology of TB in recent years and an overview of the situation in south est Iran, *Research Journal of Medical Sciences Zahedan*, **13**.
- [14] Panahi Myshkar, A. 50 basic and important thing about TB, *Quarterly golden horizon Health*, Deputy Zabol university of Medical Sciences.
- [15] Rasolinejad, M., Haddadi, A., Hedayat Yaghoubi, M., Moradmand badi, B. and Alijani, N. Review the outcomes of TB in AIDS patients treated with the standard regimen, *Journal of Medical Sciences*, Tehran university of Medical Sciences.

تحلیل ریسک در فضای متغیر تصادفی هندسی مرکب

زهرا شریعتی^۱، رسول روزگار

گروه آمار، دانشکده علوم ریاضی، دانشگاه یزد

چکیده: در علم احتمال، توزیع هندسی مرکب کاربردهای فراوانی دارد. از جمله کاربردهای آن، ایجاد رابطه‌ای بین احتمال بقای متغیر تصادفی هندسی مرکب و احتمال ورشکستگی یک شرکت در زمان نامتناهی است. در این مقاله ابتدا به معرفی توزیع هندسی مرکب و کاربرد آن در زمینه قابلیت اعتماد و نظریه ریسک پرداخته شده است و سپس با به کارگیری برخی ویژگی‌های آن، کران بالا برای آن محاسبه گردیده است. پس از آن با توجه به کاربرد این توزیع در نظریه ریسک، کران بالا و همچنین مقدار مجانبی برای احتمال ورشکستگی (خرابی) ارائه شده است. سپس با محاسبه مقدار واقعی احتمال ورشکستگی، مقدار مجانبی و کران بالای به دست آمده مورد بررسی قرار گرفته است.

واژه‌های کلیدی: توزیع هندسی مرکب، احتمال ورشکستگی، نظریه ریسک، قابلیت اعتماد.

کد موضوع بندی: 39B82، 34K20، 39B52، 47A55.

۱ مقدمه

توزیع هندسی مرکب نقش مهمی در حل مسائل قابلیت اعتماد، نظریه صف و بیمه دارد. از جمله کاربردهای آن، محاسبه کران بالا برای احتمال بقای متغیر تصادفی هندسی مرکب است، که در ادامه به آن‌ها پرداخته شده است. (متغیر تصادفی هندسی مرکب) متغیر تصادفی W را دارای توزیع هندسی مرکب گویند، و به صورت زیر تعریف می‌شود:

$$W = \begin{cases} \sum_{k=1}^M X_k, & M > 0, \\ 0, & M = 0, \end{cases}$$

به طوری که M دارای توزیع هندسی با تابع جرم احتمال $p_M(x) = \gamma^x(1-\gamma)$ و تابع مولد احتمال $G_M(z) = (1-\gamma)/(1-\gamma z)$ است. همچنین $X_1, X_2, \dots$ از M مستقل هستند. متغیر تصادفی هندسی مرکب دارای تابع مولد احتمال به فرم زیر است:

$$G_W(z) = \sum_{n=0}^{\infty} f_W(n)z^n = \frac{1-\gamma}{1-\gamma G_X(z)}, \quad |z| < \tau, \quad (1)$$

که در آن $G_X(z) = \sum_{n=0}^{\infty} f_X(n)z^n$ تابع مولد احتمال متغیر تصادفی رخ داده‌ها، X_i ها می‌باشد. با فرض اینکه احتمال بقای متغیر تصادفی رخ داده‌ها با نماد $\bar{F}_X(n) = \sum_{j=n+1}^{\infty} f_X(j)$ برای هر $n = 0, 1, 2, \dots$ نشان داده شود. در زیر ابتدا تحت قضیه‌ای ارتباط متغیر تصادفی هندسی مرکب با سیستم‌های DS-NWU بیان شده است. فرض کنید متغیر تصادفی W دارای توزیع هندسی مرکب با تابع چگالی $\{f_W(n); n = 0, 1, \dots\}$ دارای توزیع هندسی مرکب باشد. آن‌گاه $\{f_W(n); n = 0, 1, \dots\}$

دارای ویژگی DS-NWU می‌باشد. اثبات. بنابر ادعای برون [۴] در سال ۱۹۹۰ و کای [۴] در سال ۲۰۰۰، تابع چگالی هندسی مرکب دارای ویژگی NWU است. بنابراین با توجه به این ویژگی نتیجه می‌شود:

$$\begin{aligned}\bar{f}_W(n+m+1) &= P(W > m+n+1) \\ &\geq P(W > m+1/2)P(W > n+1/2) \\ &= P(W > m)P(W > n) \\ &= \bar{f}_W(m)\bar{f}_W(n)..\end{aligned}$$

□

حال متغیر تصادفی W را در نظر بگیرید. در این صورت کران بالا برای احتمال بقای متغیر تصادفی هندسی مرکب W با توجه به مطالب گفته شده، از رابطه‌ی زیر به دست می‌آید:

$$\bar{F}_W(n) \leq \left(\frac{1}{\tau}\right)^{n+1}.$$

در ابتدا بیان شد که کاربرد دیگر توزیع هندسی مرکب در مسائل مربوط به شرکت‌های بیمه‌گذار است. از جمله‌ی این مسائل نظریه ورشکستگی می‌باشد. یک شرکت بیمه زمانی ورشکست می‌شود که سرمایه آن منفی یا صفر شود. مقدار سرمایه شرکت که در دوره زمانی k ام با U_k نشان داده می‌شود و همواره از آن با عنوان فرآیند مازاد دوره k ام یاد می‌شود، از رابطه‌ی زیر به دست می‌آید:

$$U_k = u + ck - \sum_{i=1}^k Y_i, \quad k \geq 0, \quad (2)$$

به طوری که u سرمایه اولیه، c حق بیمه دریافتی در هر دوره، k دوره‌ی زمانی گذشته و Y_i مقدار خسارت وارد شده در دوره i ام می‌باشد. لازم به ذکر است که در این جا، $c = 1$ در نظر گرفته می‌شود.

(متغیر تصادفی حوادث ادعایی) متغیرهای تصادفی $\{I_k, k = 0, 1, \dots\}$ را در نظر بگیرید به گونه‌ای که $I_k = 1$ است اگر در دوره k ام خسارتی وارد شود و در غیر این صورت ۰ خواهد بود. در یک مدل ریسک، I_k را متغیر تصادفی حوادث ادعایی می‌نامند. با توجه به تعریف متغیر تصادفی حوادث ادعایی، معادله‌ی (۲) به صورت زیر می‌شود:

$$U_k = u + k - \sum_{i=1}^{N_k} B_j, \quad k \geq 0, \quad (3)$$

به طوری که B_j متغیر تصادفی مقدار رخداد می‌باشد که دارای مقادیر مثبت و صحیح هستند (تنها زمان‌هایی را در نظر می‌گیرد که خسارت وجود داشته باشد).

دیگر پارامترهای مربوط به نظریه ریسک که در این جا نیاز به معرفی دارند، احتمال‌های ورشکستگی و بقای شرطی و متغیر تصادفی زمان ورشکستگی است، که در زیر به آن‌ها پرداخته می‌شود. احتمالات ورشکستگی و بقای شرطی زمان-نامتناهی به صورت زیر هستند:

$$\psi(u|i) = P_u(T < \infty | I_0 = i)$$

و

$$\phi(u|i) = 1 - \psi(u, |i) = P_u(U_k \geq 0, \forall K \in \{1, 2, 3, \dots | I_0 = i\}).$$

برای تحلیل کردن دربارهی معادله‌ی سرمایه (۳) حالت‌های گوناگونی وجود دارد. یکی از این حالت‌ها، وابستگی متغیرهای تصادفی حوادث ادعایی در قالب زنجیر مارکف است. بدین گونه که متغیرهای $\{I_k; k = 0, 1, \dots\}$ تشکیل یک زنجیر مارکف همگن با فضای حالت $\{0, 1\}$ و ماتریس احتمال انتقال زیر را می‌دهند:

$$P = \begin{bmatrix} (1-q) + \pi q & q - \pi q \\ (1-q) - \pi(1-q) & q + \pi(1-q) \end{bmatrix} = \begin{bmatrix} p_{00} & p_{01} \\ p_{10} & p_{11} \end{bmatrix}.$$

مدل ریسک تحت این شرط را مدل ریسک دوجمله‌ای مرکب مارکف می‌نامند. حال براساس مطالب بیان شده، متغیر تصادفی W_k را به صورت $W_k = S_k - k$ در نظر بگیرید. بنابر تعریف ۱، می‌توان احتمال ورشکستگی را به صورت زیر نیز بیان نمود:

$$\psi(u|i) = P(\sup_{k \in \mathbb{N}^+} W_k > u | I_0 = i),$$

طبق عبارت بیان شده، می‌توان احتمال ورشکستگی را به عنوان تابع بقای یک متغیر تصادفی بیان کرد. بر این اساس، ادعا می‌شود که احتمال ورشکستگی همان احتمال بقای متغیر تصادفی هندسی مرکب است، که طی قضیه‌ای که در ادامه آمده این مطلب ثابت شده است. در مدل دوجمله‌ای مرکب مارکف، تابع مولد احتمال ورشکستگی شرطی زمان-نامتناهی به صورت زیر بیان می‌شود:

$$\tilde{\psi}(z|i) = \frac{1 - R_i(z)}{1 - z},$$

که $R_i(z)$ برای $i = 0$ و $i = 1$ ، تابع مولد احتمال توزیع هندسی مرکب است.

اثبات. ابتدا تابع مولد احتمال ورشکستگی را به صورت زیر در نظر بگیرید:

$$\tilde{\psi}(z|i) = \sum_{u=0}^{\infty} z^u \psi(u|i) = \sum_{u=0}^{\infty} z^u (1 - \phi(u|i)) = \frac{1 - (1-z) \tilde{\phi}(z|i)}{1 - z} = \frac{1 - R_i(z)}{1 - z}, \quad (4)$$

برای اثبات اینکه احتمال ورشکستگی $\{\psi(u|i), u \in \mathbb{N}\}$ در مدل مارکف احتمال بقای متغیر تصادفی هندسی مرکب است بنابر نکته ۱ کافی است نشان داده شود، $R_i(z)$ تابع مولد احتمال توزیع هندسی مرکب است. برای این منظور پس از در نظر گرفتن پارامترهای مورد نیاز و ساده کردن عبارت‌ها نتیجه‌ی زیر حاصل می‌شود:

$$\begin{aligned} R_0(z) &= \frac{1 - \psi(0|0)}{(1/(1-z))(1 - ((z(p_{11}G_B(z) - z))/(\pi G_B(z) - zp_{00})))} \\ &= \frac{1 - \psi(0|0)}{1 - (z/(1-z))p_{01}((G_B(z) - z)/(\pi G_B(z) - zp_{00}))} \end{aligned}$$

و

$$\begin{aligned} R_1(z) &= \frac{1 - \psi(0|1)}{(1/(1-z))(1 - ((p_{11} \tilde{f}_B(z) - z) + \pi f_B(1) - \pi(G_B(z)/z))/(\pi f_B(1) - p_{00}))} \\ &= \frac{1 - \psi(0|1)}{1 - (1/(1-z))((1-z)\pi f_B(1) - zp_{01} + (p_{11} - (\pi/z))G_B(z))/(\pi f_B(1) - p_{00})} \\ &= \frac{1 - \psi(0|1)}{1 - (\pi f_B(1) - (z/(1-z))p_{01} + ((p_{11} - \pi/z)/(1-z))G_B(z))/\pi f_B(1) - p_{00}}. \end{aligned}$$

حال فرض کنید:

$$G_{\circ}(z) = \frac{z}{1-z} \frac{p_{\circ 1}}{\psi(\circ|\circ)} \frac{G_B(z) - z}{\pi G_B(z) - z p_{\circ\circ}} \quad (5)$$

$$G_1(z) = \frac{\pi f_B(1) - z/((1-z)p_{\circ 1} + ((p_{11} - \pi/z)/(1-z))G_B(z))}{\psi(\circ|1)(\pi f_B(1) - p_{\circ\circ})}, \quad (6)$$

از روابط (5) و (6) می‌توان $R_i(z)$ را به شکل زیر بازنویسی کرد؛

$$R_i(z) = \frac{1 - \psi(\circ|i)}{1 - \psi(\circ|i)P_i(z)}, \quad |z| < z_i \quad i \in \{\circ, 1\}.$$

حال برای اثبات اینکه $R_i(z)$ تابع مولد احتمال متغیر تصادفی هندسی مرکب است، کافی است نشان داده شود $G_i(z)$ برای $i \in \{\circ, 1\}$ یک تابع مولد احتمال است.

ابتدا $i = 1$ در نظر بگیرید. پس از ساده کردن $G_1(z)$ به صورت زیر به دست می‌آید:

$$G_1(z) = \sum_{k=1}^{\infty} z^k \frac{p_{\circ 1}(1 - F_B(k)) + \pi f_B(k+1)}{\psi(\circ|1)(p_{\circ\circ} - \pi f_B(1))}.$$

با توجه به رابطه بالا، $p_1(k)$ به صورت زیر در نظر گرفته می‌شود:

$$p_1(k) = \begin{cases} \frac{p_{\circ 1}(1 - F_B(k)) + \pi f_B(k+1)}{\psi(\circ|1)(p_{\circ\circ} - \pi f_B(1))} & k \in \mathbb{N}^+, \\ \circ & \text{دیگر جاها} \end{cases}$$

برای اینکه ثابت شود $G_1(z)$ تابع مولد احتمال است، باید نشان داده شود که $p_1(k)$ تابع جرم احتمال است. با توجه به تعریف

$p_1(k)$ ، به ازای هر $x \in \mathbb{N}^+$ ، $p_1(k) \geq \circ$ است. پس شرط اول تابع چگالی را دارد. بررسی شرط دوم به صورت زیر است:

$$\begin{aligned} \sum_{k=1}^{\infty} p_1(k) &= \sum_{k=1}^{\infty} \frac{p_{\circ 1}(1 - F_B(k)) + \pi f_B(k+1)}{\psi(\circ|1)(p_{\circ\circ} - \pi f_B(1))} \\ &= \frac{p_{\circ 1}(\mu_B - 1) + \pi(1 - f_B(1))}{1 - \phi(\circ|1)(p_{\circ\circ} - \pi f_B(1))} \\ &= \frac{p_{1\circ}\psi(\circ|\circ) + \pi(1 - f_B(1))}{(1 - \phi(\circ|1))(p_{\circ\circ} - \pi f_B(1))} \\ &= \frac{p_{1\circ}(1 - \phi(\circ|\circ)) + \pi(1 - f_B(1))}{(1 - \phi(\circ|1))(p_{\circ\circ} - \pi f_B(1))} \\ &= \frac{(p_{\circ 1} + \pi - \pi f_B(1)) - (p_{\circ\circ} - \pi f_B(1))\phi(\circ|1)}{(1 - \phi(\circ|1))(p_{\circ\circ} - \pi f_B(1))} \\ &= \frac{(p_{\circ\circ} - \pi f_B(1))(1 - \phi(\circ|1))}{(1 - \phi(\circ|1))(p_{\circ 1} + \pi - \pi f_B(1))} = 1. \end{aligned}$$

که در نتیجه اثبات قضیه برای $i = 1$ کامل می‌شود.

این بار $i = \circ$ در نظر گرفته و پس از انجام محاسبات نتیجه می‌شود:

$$G_{\circ}(z) = \sum_{k=1}^{\infty} z^k \frac{1}{\pi} \frac{p_{\circ 1}}{\psi(\circ|\circ)} (1 - G(k-1)).$$

تابع جرم احتمال $p_{\circ}(k)$ را با استفاده از عبارت (7)، به صورت زیر در نظر بگیرید:

$$p_{\circ}(k) = \begin{cases} \frac{1}{\pi} \frac{p_{\circ 1}}{\psi(\circ|\circ)} (1 - G(k-1)) & k \in \mathbb{N}^+, \\ \circ & \text{دیگر جاها} \end{cases} \quad (7)$$

با توجه به عبارت (۷)، به ازای هر $x \in \mathbb{N}^+$ همواره $p_{\circ}(x) \geq 0$ است. پس فقط باید نشان داده شود که، $\sum_{k=1}^{\infty} p_{\circ}(k) = 1$.
با توجه به تعریف امید ریاضی نتیجه می‌شود:

$$\begin{aligned} \sum_{k=1}^{\infty} p_{\circ}(k) &= \sum_{k=1}^{\infty} \frac{1}{\pi} \frac{p_{\circ 1}}{\psi(\circ|\circ)} (1 - G(k-1)) \\ &= \frac{1}{\pi} \frac{p_{\circ 1}}{\psi(\circ|\circ)} \sum_{k=1}^{\infty} (1 - G(k-1)) \\ &= \frac{p_{\circ 1}}{\pi} \frac{p_{1\circ}}{p_{\circ 1}(\mu_B)} \frac{\pi}{p_{1\circ}} \mu_B = 1. \end{aligned}$$

بنابراین ثابت شد که در مدل دوجمله‌ای مرکب مارکف، $\{\psi(u|i), u \in \mathbb{N}\}$ احتمال بقای متغیر تصادفی هندسی مرکب است. $\square$
لم زیر کار را برای محاسبه‌ی کران بالا و مقدار مجانبی احتمال ورشکستگی زمان-نامتناهی هموار می‌کند. اگر $z_{\circ}$ و z_1 به ترتیب کران همگرایی دو تابع مولد احتمال $R_{\circ}(z)$ و $R_1(z)$ باشند، آن‌گاه:

$$1. z_{\circ} = z_1 = z^*.$$

۲. اگر $q\mu_B < 1$ ، آن‌گاه z^* که تنها کران همگرایی توابع مولد احتمال فوق است، همیشه از ۱ بزرگ‌تر است.

مطلب دیگری که برای توزیع هندسی مرکب بیان می‌شود، این است که برای احتمال بقای متغیر تصادفی هندسی مرکب می‌توان یک مقدار مجانبی به دست آورد. قضیه زیر گویای همین مطلب است. در مدل دوجمله‌ای مرکب مارکف، مقدار مجانبی برای $\psi(u|\circ)$ و $\psi(u|1)$ به ترتیب به فرم زیر هستند:

$$\psi(u|\circ) \sim \frac{(1 - \psi(\circ|\circ))(P_{\circ\circ} - \pi(G_B(z^*)/z^*))}{(p_{11}G'_B(z^*) - 1) - (\pi/z^*)(G'_B(z^*) - G_B(z^*)/z^*)} \left(\frac{1}{z^*}\right)^{u+1} \quad (8)$$

و

$$\psi(u|1) \sim \frac{(1 - \psi(\circ|1))(P_{\circ\circ} - \pi f_B(1))}{(p_{11}G'_B(z^*) - 1) - (\pi/z^*)(G'_B(z^*) - G_B(z^*)/z^*)} \left(\frac{1}{z^*}\right)^{u+1}. \quad (9)$$

اثبات. با توجه به قضیه ۲ از مقاله ویلموت [۳] می‌دانیم که:

$$\psi(u|i) \sim \frac{1 - \psi(\circ|i)}{\psi(\circ|i)(z^* - 1)G'_i(z^*)} \left(\frac{1}{z^*}\right)^{u+1}. \quad (10)$$

بنابراین با محاسبه‌ی $G'_i(z^*)$ به ازای $i = \circ$ و $i = 1$ عبارت‌های (۸) و (۹) حاصل می‌شوند. $\square$

۲ محاسبه‌ی مقدار دقیق احتمال ورشکستگی

بعد از محاسبه‌ی مقدار مجانبی برای احتمال ورشکستگی، این سؤال مطرح می‌شود که تا چه اندازه این مقدار مجانبی به مقدار واقعی احتمال ورشکستگی نزدیک است. برای این منظور، پس از در نظر گرفتن معادله‌ی احتمال ورشکستگی و بقا، چون محاسبه‌ی احتمال بقا با به کار بردن تابع توزیع متغیر S_k صورت می‌گیرد، لذا ابتدا احتمال بقا محاسبه شده است. تابع چگالی متغیر S_k با در نظر گرفتن این ویژگی که متغیر تصادفی حوادث ادعایی تشکیل زنجیر مارکف می‌دهند، به صورت زیر محاسبه می‌شود:

$$p_{s_k}(j|i) = p_{i\circ}p_{s_{k-1}}(j|\circ) + p_{i1} \sum_{l=1}^j p_{s_{k-1}}(j-l|1)f_B(l).$$

لم زیر برای محاسبه‌ی احتمال بقای زمان-نامتناهی شرطی با استفاده از توزیع S_k بیان شده است. احتمال بقای شرطی زمان-نامتناهی برای حالت‌های خاص به صورت زیر است:

۱.

$$\phi(\circ|\circ) = \frac{\mathbb{1} - qE[B]}{\mathbb{1} - q},$$

۲.

$$\phi(\circ|\mathbb{1}) = \frac{p_{\mathbb{1}\circ}}{p_{\circ\circ} - (p_{\circ\circ}p_{\mathbb{1}\mathbb{1}} - p_{\circ\mathbb{1}}p_{\mathbb{1}\circ})} \phi(\circ|\circ),$$

و برای $u \in \mathbb{N}^+$ ؛

۳.

$$\phi(u|\circ) = \phi(\circ|\circ) + \frac{p_{\circ\mathbb{1}}}{p_{\mathbb{1}\circ}} \sum_{j=\circ}^{u-\mathbb{1}} \phi(j|\mathbb{1})(\mathbb{1} - F_B(u-j)),$$

۴.

$$\phi(u|\mathbb{1}) = \frac{p_{\mathbb{1}\circ}\phi(u|\circ) + \{p_{\circ\circ}p_{\mathbb{1}\mathbb{1}} - p_{\circ\mathbb{1}}p_{\mathbb{1}\circ}\} \sum_{j=\mathbb{1}}^{u+\mathbb{1}} \phi(u+\mathbb{1}-j|\mathbb{1})f_B(j)}{p_{\circ\circ} - (p_{\circ\circ}p_{\mathbb{1}\mathbb{1}} - p_{\circ\mathbb{1}}p_{\mathbb{1}\circ})f_B(\mathbb{1})}.$$

از آن‌جا که احتمال بقا و ورشکستگی متمم یک‌دیگر هستند، نتیجه‌ی زیر برای محاسبه‌ی احتمال ورشکستگی زمان-نامتناهی بیان می‌شود. برای $u = \circ$ نتیجه می‌شود:

۱.

$$\psi(\circ|\circ) = \frac{q}{\mathbb{1} - q} (E[B] - \mathbb{1}) \quad (11)$$

۲.

$$\psi(\circ|\mathbb{1}) = \frac{\pi(\mathbb{1} - f_B(\mathbb{1})) + p_{\mathbb{1}\circ}\psi(\circ|\circ)}{p_{\circ\circ} - \pi f_B(\mathbb{1})}.$$

و برای $u \in \mathbb{N}^+$ ؛

۳.

$$\psi(u|\circ) = \frac{q}{\mathbb{1} - q} \sum_{k=\mathbb{1}}^u \phi(u-k|\mathbb{1})(\mathbb{1} - F_B(k)) + \frac{q}{\mathbb{1} - q} \sum_{k=u+\mathbb{1}}^{\infty} (\mathbb{1} - F_B(k)),$$

۴.

$$\begin{aligned} \psi(u|\mathbb{1}) &= \sum_{k=\mathbb{1}}^u \frac{p_{\circ\mathbb{1}}(\mathbb{1} - F_B(k)) + \pi f_B(k+\mathbb{1})}{p_{\circ\circ} - \pi f_B(\mathbb{1})} \psi(u-k|\mathbb{1}) \\ &+ \sum_{k=u+\mathbb{1}}^{\infty} \frac{p_{\circ\mathbb{1}}(\mathbb{1} - F_B(k)) + \pi f_B(k+\mathbb{1})}{p_{\circ\circ} - \pi f_B(\mathbb{1})}. \end{aligned} \quad (12)$$

فرض کنید که متغیر تصادفی B دارای توزیع هندسی بریده شده در صفر، با تابع چگالی احتمال $f_B(i) = (1 - \alpha)\alpha^{i-1}$ و تابع توزیع $F_B(i) = 1 - \alpha^i$ برای هر $i \in N^+$ باشد. مقدار دقیق احتمال ورشکستگی برای این حالت، از رابطه‌ی زیر به دست می‌آید:

$$\psi(u|i) = \psi(0|i) \left(\frac{\alpha}{p_{00} - \pi(1 - \alpha)} \right)^u, \quad i \in \{0, 1\}, \quad (13)$$

و

$$\psi(u) = \psi(0) \left(\frac{\alpha}{p_{00} - \pi(1 - \alpha)} \right)^u.$$

اثبات. با جای‌گذاری کردن تابع چگالی و تابع توزیع متغیر تصادفی هندسی بریده شده در صفر، در عبارت‌های (۱۷) و (۱۶) عبارت‌های (۱۳) و (۲) حاصل می‌شوند. $\square$

با توجه به گزاره قبل و همچنین مقدار مجانبی به دست آمده برای احتمال ورشکستگی، در قضیه زیر ثابت خواهد شد زمانی که متغیرهای تصادفی $\{B_j; j = 1, 2, \dots\}$ دارای توزیع هندسی بریده شده در صفر باشند، مقدار مجانبی همان مقدار دقیق برای آن‌هاست. در مدل دوجمله‌ای مرکب مارکف، وقتی متغیر تصادفی دارای توزیع هندسی بریده شده در صفر باشد، مقادیر مجانبی (۸) و (۹)، مقادیر دقیق (۱۳) را می‌دهند.

۳ شبیه‌سازی

جداول زیر درستی گزاره فوق را ابتدا برای $i = 0$ و سپس برای $i = 1$ نشان می‌دهند.

جدول ۱: مقادیر دقیق، مجانبی و کران بالا برای احتمال ورشکستگی به شرط $i = 0$

u	$\psi(u 0)$	مقدار مجانبی	کران بالا
۰	$1/0.71429 \times 10^{-1}$	$1/0.71429 \times 10^{-1}$	$4/444444 \times 10^{-1}$
۵	$1/858021 \times 10^{-3}$	$1/858021 \times 10^{-3}$	$7/707347 \times 10^{-3}$
۱۰	$3/222093 \times 10^{-5}$	$3/222093 \times 10^{-5}$	$1/336572 \times 10^{-4}$
۱۵	$5/587602 \times 10^{-7}$	$5/587602 \times 10^{-7}$	$2/317820 \times 10^{-6}$
۲۰	$9/689756 \times 10^{-9}$	$9/689756 \times 10^{-9}$	$4/019455 \times 10^{-8}$
۳۰	$2/913987 \times 10^{-12}$	$2/913987 \times 10^{-12}$	$1/208765 \times 10^{-11}$
۴۰	$8/763195 \times 10^{-16}$	$8/763195 \times 10^{-16}$	$3/635103 \times 10^{-15}$
۵۰	$2/635344 \times 10^{-19}$	$2/635344 \times 10^{-19}$	$1/093180 \times 10^{-18}$
۶۰	$7/925235 \times 10^{-23}$	$7/925235 \times 10^{-23}$	$3/287505 \times 10^{-22}$
۷۰	$2/383345 \times 10^{-26}$	$2/383345 \times 10^{-26}$	$9/886469 \times 10^{-26}$
۸۰	$7/167402 \times 10^{-30}$	$7/167402 \times 10^{-30}$	$2/973145 \times 10^{-29}$
۹۰	$2/155443 \times 10^{-33}$	$2/155443 \times 10^{-33}$	$8/941098 \times 10^{-33}$
۱۰۰	$6/482036 \times 10^{-37}$	$6/482036 \times 10^{-37}$	$2/688844 \times 10^{-36}$
۱۲۵	$1/016596 \times 10^{-45}$	$1/016596 \times 10^{-45}$	$4/216991 \times 10^{-45}$
۱۵۰	$1/594357 \times 10^{-54}$	$1/594357 \times 10^{-54}$	$6/613628 \times 10^{-54}$

و

جدول ۲: مقادیر دقیق، مجانبی و کران بالا برای احتمال ورشکستگی به شرط $i = 1$

u	$\psi(u 1)$	مقدار مجانبی	کران بالا
۰	$3/055556 \times 10^{-1}$	$3/055556 \times 10^{-1}$	$4/444444 \times 10^{-1}$
۵	$5/298801 \times 10^{-3}$	$5/298801 \times 10^{-3}$	$7/707347 \times 10^{-3}$
۱۰	$9/188931 \times 10^{-5}$	$9/188931 \times 10^{-5}$	$1/336572 \times 10^{-4}$
۱۵	$1/593501 \times 10^{-6}$	$1/593501 \times 10^{-6}$	$2/317820 \times 10^{-6}$
۲۰	$2/763375 \times 10^{-8}$	$2/763375 \times 10^{-8}$	$4/019455 \times 10^{-8}$
۳۰	$8/310261 \times 10^{-12}$	$8/310261 \times 10^{-12}$	$1/208765 \times 10^{-11}$
۴۰	$2/499134 \times 10^{-15}$	$2/499134 \times 10^{-15}$	$3/635103 \times 10^{-15}$
۵۰	$7/515611 \times 10^{-19}$	$7/515611 \times 10^{-19}$	$1/093180 \times 10^{-18}$
۶۰	$2/260160 \times 10^{-22}$	$2/260160 \times 10^{-22}$	$3/287505 \times 10^{-22}$
۷۰	$6/796948 \times 10^{-26}$	$6/796948 \times 10^{-26}$	$9/886469 \times 10^{-26}$
۸۰	$2/044037 \times 10^{-29}$	$2/044037 \times 10^{-29}$	$2/973145 \times 10^{-29}$
۹۰	$6/147005 \times 10^{-33}$	$6/147005 \times 10^{-33}$	$8/941098 \times 10^{-33}$
۱۰۰	$1/848581 \times 10^{-36}$	$1/848581 \times 10^{-36}$	$2/688844 \times 10^{-36}$
۱۲۵	$2/899182 \times 10^{-45}$	$2/899182 \times 10^{-45}$	$4/216991 \times 10^{-45}$
۱۵۰	$4/546869 \times 10^{-54}$	$4/546869 \times 10^{-54}$	$6/613628 \times 10^{-54}$

جداول فوق نشان می‌دهند زمانی که متغیر تصادفی مقادیر رخدادها از توزیع هندسی بریده شده در صفر تبعیت کنند، آنگاه مقدار واقعی احتمال ورشکستگی $(\psi(u|i), \forall i \in \{0, 1\})$ و مقدار مجانبی آن با هم برابر است. همچنین مقدار کران بالای به‌دست آمده برای احتمال ورشکستگی شرطی نشان می‌دهد که از دو مقدار دیگر بزرگ‌تر می‌باشد.

مراجع

- [1] Cossette, H., Landriault, D. and Marceau, E. (2003), Ruin probabilities in the compound Markov binomial model, *Scandinavian Actuarial Journal*, **4**, 301-323.
- [2] Cossette, H., Landriault, D. and Marceau, E. (2004), Exact expression and upper bound for ruin probabilities in the compound Markov binomial model, *Insurance: Mathematics and Economics*, **34**, 449-466.
- [3] Willmot, G.E. (1989), Limiting tail behavior of some discrete compound distributions, *Insurance: Mathematics and Economics*, **8**, 175-185.
- [4] Willmot, G.E. and Lin, S. (2001), *Lundberg Approximations for compound Distributions with Insurance Applications*, Lecture Notes in Statistics, New York.

کاربرد توزیع پارتو نمایی شده در قابلیت اعتماد

سمیه شهرکی ده سوخته^۱

گروه آمار، دانشکده علوم ریاضی و آمار، دانشگاه زابل، زابل، ایران

چکیده: توزیع پارتو نمایی شده یکی از توزیع‌های مهم طول عمر است. در این مقاله ابتدا به معرفی توزیع پارتو نمایی شده می‌پردازیم و ویژگی‌های قابلیت اعتمادی آن را مورد بررسی قرار می‌دهیم. در نهایت به مقایسه مدل معرفی شده در برازش داده‌های واقعی با سایر مدل‌ها خواهیم پرداخت. واژه‌های کلیدی: قابلیت اعتماد، برآورد حداکثر درست‌نمایی، توزیع پارتو نمایی شده.

۱ پیش‌گفتار

فرض کنید $F(x)$ یک تابع توزیع باشد. ساختار تابع توزیع خانواده توزیع‌های ارتجاعی به صورت

$$F_{\alpha}(x) = [F(x)]^{\alpha} \quad (1)$$

است که در آن $\alpha > 0$ پارامتر ارتجاعی نامیده می‌شود. در واقع تابع توزیع خانواده توزیع‌های ارتجاعی علاوه بر، بردار پارامترهای دیگر حاوی پارامتر جدید α نیز می‌باشد. $(F(x))^{\alpha}$ تعمیم نمایی شده نامیده می‌شود. تابع چگالی مربوط به تابع توزیع (۲) به صورت زیر است:

$$f_{\alpha}(x) = \alpha f(x)(F(x))^{\alpha-1} \quad (2)$$

که $F(x)$ ، تابع توزیع احتمال هر توزیعی که دارای فرم بسته است می‌تواند باشد. گوپتا و کوندا ([۱۰]، [۶]، [۷] و [۸]) و ناداراجا [۱۰] با در نظر گرفتن $F(x)$ تابع توزیع نمایی به معرفی توزیع نمایی تعمیم یافته و توزیع نمایی شده پرداختند. علی و همکاران [۲] برخی از توزیع‌های نمایی شده را معرفی کردند. عقیقی و همکاران [۲] با در نظر گرفتن $F(x)$ تابع توزیع وایبل-پارتو که توسط طاهیر و همکاران [۱۱] معرفی شده است به معرفی توزیع وایبل-پارتو نمایی شده پرداختند. در بخش ۲ به معرفی توزیع پارتو نمایی شده و ویژگی‌های آن می‌پردازیم و در بخش ۳ رفتار توابع نرخ خطر، نرخ خطر معکوس، میانگین باقی‌مانده عمر و گذشته عمر در توزیع پارتو نمایی شده را مطالعه می‌کنیم. در بخش ۴ با یک مثال کاربردی توانایی این مدل در تحلیل داده‌های واقعی در مقایسه با سایر مدل‌های رقیب مورد بررسی قرار می‌دهیم. در نهایت نتیجه‌گیری را ارائه می‌دهیم.

۲ توزیع پارتو نمایی شده و ویژگی‌های آن

متغیر تصادفی X دارای توزیع پارتو نمایی شده است هرگاه تابع چگالی احتمال آن به صورت زیر باشد:

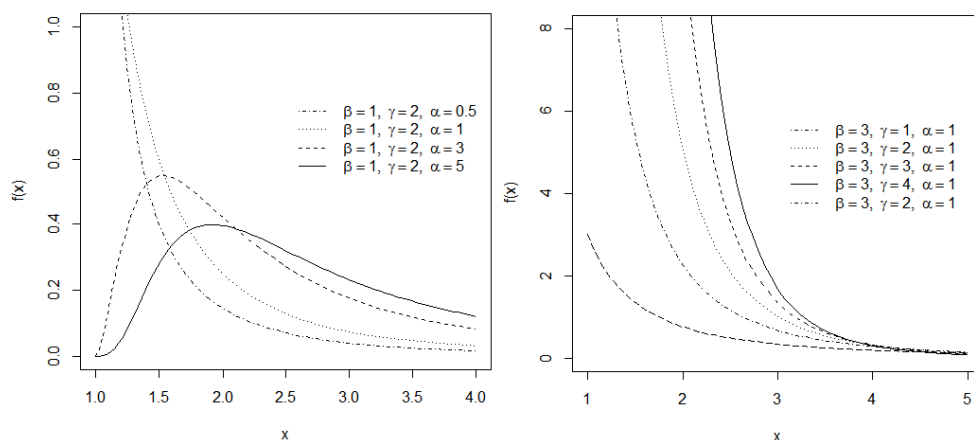
$$f(x) = \alpha \gamma \beta^{\gamma} x^{-\gamma-1} \left(1 - \left(\frac{\beta}{x}\right)^{\gamma}\right)^{\alpha-1}, \quad x \geq \beta, \alpha, \beta, \gamma > 0$$

^۱ سیمیه شهرکی ده سوخته: Sshahraki7@gmail.com

با نماد $X \sim EDPD(\alpha, \beta, \gamma)$ نشان می دهیم. (علی و همکاران [۲])
فرض کنید متغیر تصادفی X دارای توزیع پارتو نمایی شده باشد، تابع توزیع احتمال آن به صورت زیر است:

$$F(x) = [1 - (\frac{\beta}{x})]^\alpha \quad (۳)$$

شکل (۱) نمودار تابع چگالی احتمال توزیع پارتو نمایی به ازای مقادیر مختلف پارامترها را نشان می دهد.



شکل ۱: نمودار تابع چگالی احتمال توزیع پارتو نمایی به ازای مقادیر مختلف پارامترها

گشتاورهای n -ام متغیر تصادفی X ، به صورت زیر به دست می آید:

$$E(X^k) = \begin{cases} \alpha\beta^k\gamma \sum_{i=0}^{\alpha-1} \binom{\alpha-1}{i} \frac{(-1)^i}{\gamma(i+1)-k}, & \text{if } \gamma > k, \alpha \in N \\ \alpha\beta^k\gamma \sum_{i=0}^{\infty} \frac{(\alpha-1)^{Pi}}{i!} \binom{\alpha-1}{i} \frac{(-1)^i}{\gamma(i+1)-k}, & \text{if } \gamma > k, \alpha \notin N \end{cases}$$

تابع مولد گشتاور آن به صورت زیر است:

$$E(e^{tX}) = \begin{cases} \alpha\gamma \sum_{i=0}^{\alpha-1} \binom{\alpha-1}{i} (-1)^i (-\beta t)^{\gamma(i+1)} \Gamma(-\gamma(i+1), -\beta t), & \text{if } t < 0, \alpha \in N \\ \alpha\gamma \sum_{i=0}^{\infty} \frac{(\alpha-1)^{Pi}}{i!} (-1)^i (-\beta t)^{\gamma(i+1)} \Gamma(-\gamma(i+1), -\beta t), & \text{if } t < 0, \alpha \notin N \end{cases}$$

$$\Gamma(\alpha, x) = \int_x^\infty e^{-y} y^{\alpha-1} dy, \quad \alpha > 0$$

۳ رفتار توابع قابلیت اعتماد در توزیع پارتو نمایی شده

فرض کنید $F(x)$ یک تابع توزیع تجمعی و α یک عدد حقیقی مثبت باشد. با تعریف

$$F_\alpha(x) = [F(x)]^\alpha$$

داریم

الف) اگر $\alpha > 1$ و F یک تابع توزیع IFR باشد، در این صورت $F_\alpha(x)$ نیز یک تابع توزیع IFR است.
 ب) اگر $\alpha < 1$ و F یک تابع توزیع DFR باشد، در این صورت $F_\alpha(x)$ نیز یک تابع توزیع DFR است. (گوپتا و همکاران [۹])

توابع قابلیت اعتماد، نرخ خطر، نرخ خطر معکوس، میانگین باقی مانده عمر، میانگین گذشته عمر توزیع پارتو نمایی شده به ترتیب عبارتند از:

$$\bar{F}(x) = 1 - F(x) = 1 - [1 - (\frac{\beta}{x})^\gamma]^\alpha$$

$$h_\alpha(x) = \frac{f_\alpha(x)}{\bar{F}_\alpha(x)} = \frac{\alpha\gamma\beta^\gamma x^{-\gamma-1} (1 - (\frac{\beta}{x})^\gamma)^{\alpha-1}}{1 - [1 - (\frac{\beta}{x})^\gamma]^\alpha}$$

$$h^*(x) = \frac{f(x)}{F(x)} = \frac{\alpha\gamma\beta^\gamma x^{-\gamma-1} (1 - (\frac{\beta}{x})^\gamma)^{\alpha-1}}{[1 - (\frac{\beta}{x})^\gamma]^\alpha} = \frac{\alpha\gamma\beta^\gamma x^{-\gamma-1}}{[1 - (\frac{\beta}{x})^\gamma]}$$

$$m(x) = \frac{\int_x^\infty 1 - [1 - (\frac{\beta}{u})^\gamma]^\alpha du}{1 - [1 - (\frac{\beta}{x})^\gamma]^\alpha}$$

$$m^*(x) = \frac{\int_\beta^x [1 - (\frac{\beta}{u})^\gamma]^\alpha du}{[1 - (\frac{\beta}{x})^\gamma]^\alpha}$$

برای بررسی رفتار تابع نرخ خطر توزیع پارتو نمایی شده از قضیه (۳) استفاده می‌کنیم.

تابع نرخ خطر توزیع پارتو نمایی شده به ازای $\alpha < 1$ نزولی است.

اثبات. فرض کنید $F(x)$ ، $f(x)$ و $h(x)$ به ترتیب تابع توزیع تجمعی و تابع چگالی و تابع نرخ خطر پارتو و $f_\alpha(x)$ ، $F_\alpha(x)$ و $h_\alpha(x)$ به ترتیب تابع توزیع و تابع چگالی و تابع نرخ خطر پارتو نمایی شده باشند. داریم:

$$\begin{aligned} h_\alpha(x) &= \frac{f_\alpha(x)}{\bar{F}_\alpha(x)} = \alpha \frac{\gamma\beta^\gamma x^{-\gamma-1}}{1 - [1 - (\frac{\beta}{x})^\gamma]} \frac{[1 - (\frac{\beta}{x})^\gamma]^{\alpha-1} - [1 - (\frac{\beta}{x})^\gamma]^\alpha}{1 - [1 - (\frac{\beta}{x})^\gamma]^\alpha} \\ &= \alpha h(x) \psi(x; \alpha) \end{aligned}$$

که در آن

$$\psi(x; \alpha) = \frac{[1 - (\frac{\beta}{x})^\gamma]^{\alpha-1} - [1 - (\frac{\beta}{x})^\gamma]^\alpha}{1 - [1 - (\frac{\beta}{x})^\gamma]^\alpha}$$

اکنون رفتار $\psi(x; \alpha)$ را بررسی می‌کنیم. با تغییر متغیر $y = 1 - (\frac{\beta}{x})^\gamma$ داریم

$$g(y, \alpha) = \frac{y^{\alpha-1} - y^\alpha}{1 - y^\alpha}$$

اکنون با مشتق گیری داریم

$$\begin{aligned}\frac{\partial}{\partial y} g(y, \alpha) &= \frac{\partial}{\partial y} \frac{y^{\alpha-1} - y^\alpha}{1 - y^\alpha} \\ &= \frac{y^{\alpha-2}}{(1 - y^\alpha)^2} (y^\alpha - \alpha y + \alpha - 1) < 0 \quad \text{if} \quad \alpha < 1\end{aligned}$$

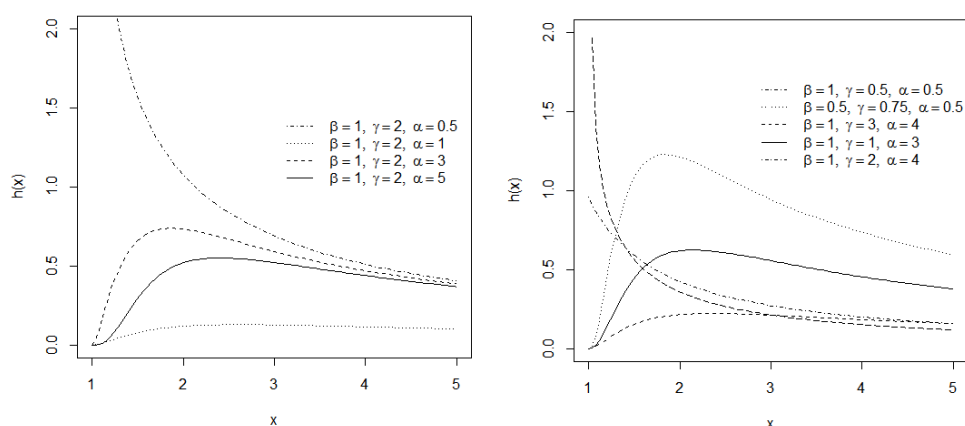
در نتیجه

$$\begin{aligned}\frac{\partial}{\partial x} h_\alpha(x) &= \alpha h'(x) \frac{\partial}{\partial x} \psi(x; \alpha) \\ &= \alpha h'(x) \frac{[1 - (\frac{\beta}{x})^\gamma]^{\alpha-2}}{(1 - [1 - (\frac{\beta}{x})^\gamma]^\alpha)^2} \left([1 - (\frac{\beta}{x})^\gamma]^\alpha - \alpha [1 - (\frac{\beta}{x})^\gamma] + \alpha - 1 \right)\end{aligned}$$

بنابراین با توجه به اینکه $h(x)$ تابع نرخ خطر توزیع پارتو است و نزولی است. اگر $\alpha < 1$ باشد در این صورت با توجه به قضیه (۳) داریم $\frac{\partial}{\partial x} h_\alpha(x) < 0$ و نتیجه مطلوب حاصل شد.

□

شکل (۲) نمودار تابع نرخ خطر توزیع پارتو نمایی به ازای مقادیر مختلف پارامترهای آن را نشان می‌دهد.



شکل ۲: نمودار تابع نرخ خطر توزیع پارتو نمایی به ازای مقادیر مختلف پارامترها

۴ مثال کاربردی

در این بخش یک مجموعه داده که نشان دهنده زمان انتظار (برحسب دقیقه) قبل از ارائه خدمات به مشتری بانک را تحت هر یک از مفروضات پارتو نمایی شده، وایبل-پارتو نمایی شده، وایبل-پارتو و وایبل مورد تجزیه و تحلیل قرار می‌دهیم. این مجموعه داده شامل ۱۰۰ مشاهده است و اولین بار توسط قیتانی و همکاران [۴] تحلیل گردید. این داده‌ها در جدول (۱) ارائه شده‌اند.

جدول ۱: داده‌های مربوط به زمان انتظار (برحسب دقیقه) قبل از ارائه خدمات به مشتری بانک

۳/۲	۳/۱	۲/۹	۲/۷	۲/۶	۲/۱	۱/۹	۱/۹	۱/۸	۱/۵	۱۰/۳	۰/۸	۰/۸
۴/۹	۴/۹	۴/۸	۴/۷	۴/۷	۴/۶	۴/۴	۴/۴	۴/۳	۴/۳	۴/۲	۴/۲	۴/۱
۷/۴	۷/۱	۷/۱	۷/۱	۷/۱	۶/۹	۶/۷	۶/۳	۶/۲	۶/۲	۶/۲	۶/۱	۵/۷
۱۰/۹	۱۰/۷	۹/۸	۹/۷	۹/۶	۹/۵	۸/۹	۸/۹	۸/۸	۸/۸	۸/۶	۸/۶	۸/۶
۱۴/۱	۱۳/۹	۱۳/۷	۱۳/۶	۱۳/۳	۱۳/۱	۱۳/۰	۱۲/۹	۱۲/۵	۱۲/۴	۱۱/۹	۱۱/۵	۱۱/۲
۱۱/۲	۱۱/۱	۲۳/۰	۲۱/۹	۲۱/۴	۲۱/۳	۲۰/۶	۱۹/۹	۱۹/۰	۱۸/۹	۱۸/۴	۱۸/۲	۱۸/۱
۱۱/۰	۸/۲	۸/۰	۷/۷	۷/۶	۵/۷	۵/۵	۵/۳	۵/۰	۴/۰	۳/۶	۳/۵	۳/۳
					۱۱/۰	۳۸/۵	۳۳/۱	۳۱/۶	۲۷/۰	۱۷/۳	۱۵/۴	۱۵/۴

جدول (۱) خلاصه‌ای از آماره‌های محاسبه شده برای داده‌ها را نشان می‌دهد.

جدول ۲: شاخص‌های آماری مجموعه داده‌های مربوط به زمان انتظار (برحسب دقیقه) قبل از ارائه خدمات به مشتری بانک

تعداد (n)	چارک اول	چارک سوم	دامنه	ضریب تغییرات	انحراف معیار	میان	میانگین
۱۰۰	۴/۶۷۵	۱۳/۰۲۵	۳۷/۷	۰/۷۳۲	۷/۲۳۶	۸/۱۰۰	۹/۸۷۷

جدول (۲) برآورد حداکثر درست‌نمایی پارامترهای مدل‌های مطرح شده و انحراف استاندارد پارامترها و همچنین معیارهای اطلاع آکائیک (AIC)، بیزی (BIC) و لگاریتم تابع درست‌نمایی را برای مدل پارتو نمایی شده و سه مدل دیگر وایبل و وایبل نمایی شده (عفی‌فی و همکاران [۲]) و وایبل-پارتو (طاهیر و همکاران [۱۱]) نشان می‌دهد. همانطور که مشاهده می‌کنید مقادیر AIC و BIC برای مدل پارتو نمایی شده، کمتر از مدل‌های وایبل و وایبل-پارتو نمایی شده و وایبل-پارتو می‌باشد همچنین مقدار $loglike$ برای مدل پارتو نمایی شده بیشتر از سه مدل دیگر است.

شکل (۳) نمودار بافت‌نگار مجموعه داده‌های مربوط به زمان انتظار (برحسب دقیقه) قبل از ارائه خدمات به مشتری بانک به همراه چگالی‌های برازش داده شده به آنها را نشان می‌دهد. این شکل به نوعی نتیجه به دست آمده توسط جدول (۲) را تأیید می‌کند و نشان می‌دهد که توزیع پارتو نمایی برازش بهتری به داده‌ها ارائه کرده است.

بحث و نتیجه گیری

در این مقاله به معرفی توزیع پارتو نمایی شده و ویژگی‌های آن پرداختیم. همچنین رفتار توابع قابلیت اعتماد این توزیع مورد مطالعه قرار دادیم. یک مجموعه داده را تحت هر یک از مفروضات پارتو نمایی شده، وایبل-پارتو نمایی شده، وایبل-پارتو و وایبل مورد تجزیه و تحلیل قرار دادیم و مشاهده کردیم که توزیع پارتو نمایی شده نسبت به سایر مدل‌ها عملکرد بهتری به داده‌ها ارائه کرده است.

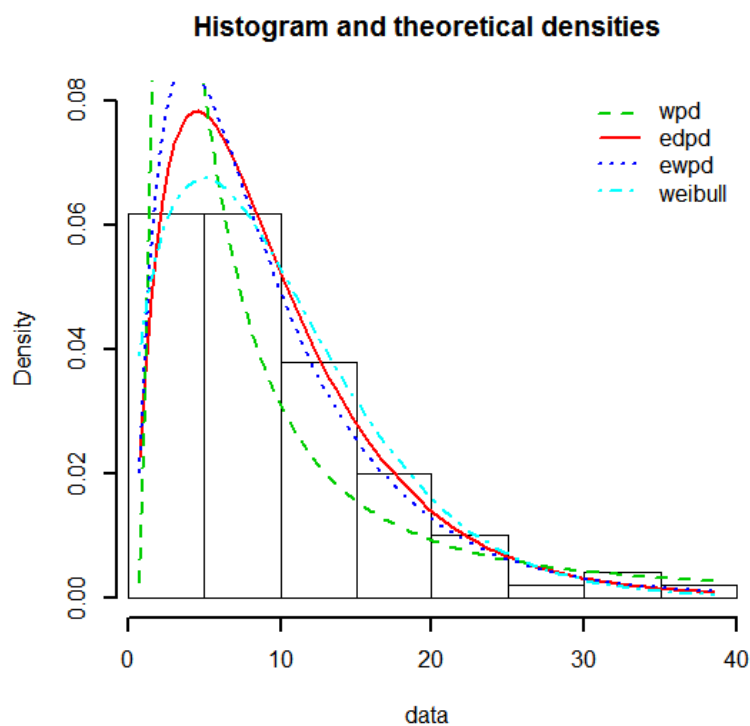
جدول ۳: مقادیر MLE و AIC و BIC برای سه مدل پارتو نمایی شده، وایبل پارتو نمایی شده، وایبل پارتو، وایبل (عبارت داخل پرانتز انحراف استاندارد پارامترها را نشان می دهد).

پارتو نمایی شده	وایبل پارتو نمایی شده	وایبل پارتو	وایبل
$\alpha = 13/66$	$a = 0/212$	$c = 7/16$	$\alpha = 1/4458$
(4/70)	(0/168)	(2/72)	(0/1097)
$\beta = 0/384$	$b = 2/86$	$\beta = 0/197$	$\gamma = 10/95$
(0/097)	(0/617)	(0/070)	(0/794)
$\gamma = 1/051$	$\theta = 0/358$	$\theta = 0/066$	
(0/078)	(1/390)	(0/11)	
	$\alpha = 1/175$		
	(2/341)		
AIC	639/86	643/694	681/106
BIC	647/67	653/194	2673/406
$loglike$	- 316/930	- 317/617	- 337/684
			- 318/730

مراجع

[۱] آ. عبدالحی نانواپیشه، توزیع پارتو یکنواخت و توزیع پارتو یکنواخت نمایی شده و کاربردهای آن‌ها در داده‌های درآمد، اندیشه آماری، ۱۲(۱۳۹۷)، جلد ۲۳، شماره ۲

- [2] Afify, A.Z., Yousof, H. M., Hamedani, G.G. and Aryal, G. (2016), The Exponentiated Weibull Pareto Distribution with Application, *Journal of Data Science*, **6**, 1-19.
- [3] Ali, M.M., Manisha, P. and Jungsoo, W. (2007), Some Exponentiated Distributions, *the Korean Communications in Statistics*, **14**, 93-109.
- [4] Ghitany, M.E., Atieh, B. and Nadarajah, S. (2008), Lindley distribution and its application, *Mathematics and Computers in Simulation*, **78**, 493-506 .
- [5] Gupta, R.D. and Kunda D. (1999), Generalized exponential distributions, *Australian and New Zeland Journal of Statistics*, **41**, 173-188.
- [6] Gupta, R.D. and Kundu, D. (2001a), Exponentiated exponential family: an alternative to gamma and Weibull distributions, *Biometrical Journal*, **43**, 117-130.



شکل ۳: نمودار بافتنگار مجموعه داده‌های مربوط به زمان انتظار (برحسب دقیقه) قبل از ارائه خدمات به مشتری بانک و برازش آن‌ها با توزیع‌های پارتو نمایی، پارتو نمایی شده و پارتو

- [7] Gupta, R.D. and Kundu, D. (2001b), Generalized exponential distribution: different method of estimations, *Journal of Statistical Computation and Simulation*, **69**, 315-337.
- [8] Gupta, R.D. and Kunda, D. (2007), Generalized Exponential Distribution: Existing results and some recent developments, *Journal of Statistical Planning and Inference*, **137**, 3537-3547.
- [9] Gupta, R.C., Gupta, P.L. and Gupta, R.D. (1998), Modeling failure time data by Lehman alternatives, *Communications in Statistics-Theory and methods*, **27**(4), 887-904.
- [10] Nadarajah, S. and Haghighi, F. (2011), An extension of the exponential distribution, *Statistics*, **45**(6), 543-558.
- [11] Tahir, M.H., Cordeiro, G.M., Alzaatreh, A., Mansoor, M. and Zubair, M. (2015), A new Weibull-Pareto distribution: properties and applications, *Communications in Statistics - Simulation and Computation*, **45**, 3548-3567.

مدل‌بندی و نگهداری سیستم‌های تصادفی چندجزئی با به‌کارگیری فرایند تولد-مرگ دوبعدی

فریبا عزیزی^۱، فیروزه حقیقی، الهام جمالی

گروه آمار، دانشکده علوم ریاضی، دانشگاه تهران

چکیده: در سال‌های اخیر، علاقه‌ی فزاینده‌ای به مطالعه‌ی سیستم‌های پیچیده که حالت آن‌ها در طی زمان تغییر می‌کند، به وجود آمده است. در این مقاله، سیستم‌های چندجزئی تخریب‌پذیر که در معرض پایش وضعیت مستمر قرار دارند، مورد بررسی قرار می‌گیرند که در آن وضعیت هر جزء سیستم به سه حالت: سالم، هشدار و شکست طبقه‌بندی شده است. از آنجایی که سیستم پیشنهادی دارای ساختار پیچیده‌ای است، به منظور مدل‌بندی تغییرات تصادفی آن از فرایند تولد-مرگ دوبعدی استفاده شده و توزیع زمان اصابت سیستم به مجموعه‌ای از حالت‌های خاص جهت نگهداری و تعمیرات پیشگیرانه محاسبه شده است. در این مطالعه، ابتدا با استفاده از شبیه‌سازی به برآورد پارامترهای نامعلوم پرداخته و سپس احتمال‌های انتقال را به روش یکنواخت‌سازی ارائه کرده‌ایم. در نهایت به بررسی توزیع زمان اصابت سیستم به مجموعه‌ای از حالت‌های خاص پرداخته‌ایم.

واژه‌های کلیدی: نگهداری سیستم‌های چندجزئی، فرایندهای تصادفی، احتمال‌های انتقال، زمان اصابت.

۱ پیش‌گفتار

نگهداری و تعمیر سیستم‌های پیشرفته همواره یکی از دغدغه‌های مدیران صنایع بوده است. در برخی از صنایع سیستم‌های بسیار حیاتی و پراهمیت وجود دارد که نبض اصلی آن صنعت را در دست دارند. این سیستم‌ها بی‌وقفه مشغول کار هستند و توقف و خرابی آنها هزینه‌ی سنگینی برای شرکت ایجاد می‌کند. علوم آماری از جمله قابلیت اعتماد [۴] و نگهداری و تعمیرات سیستم‌ها [۶] همیشه کمک شایانی در بهبود راندمان کاری بوده‌اند. در سیستم‌های مورد نظر زمان خرابی و توقف دستگاه بسیار مهم است. به همین دلیل ما به کمک روش‌های نگهداری و تعمیرات سیستم‌ها به دنبال بهترین زمان برای تعمیر و یا تعویض سیستم‌ها هستیم.

برای سیستم‌های چندجزئی، برنامه‌ریزی برای تعمیر و نگهداری سیستم با توجه به حالت کلی سیستم ارائه می‌شود. در چنین مواردی هنگامی که سیستم برای تعمیر تعدادی از اجزاء متوقف می‌شود، معمولاً به دلیل هزینه‌ی زیاد ناشی از دمنوتاژ سیستم، صرفه‌ی اقتصادی در این است که دیگر اجزاء را نیز تعمیر و سرویس نمود [۱]. مطالعات بسیاری در زمینه‌ی نگهداری و تعمیرات پیشگیرانه سیستم‌های چندجزئی سری با در نظر گرفتن اهداف و ویژگی‌های مختلف صورت گرفته است. در این راستا، می‌توان به [۲] اشاره کرد. در برخی از مطالعات نیز فرض شده است که اجزاء سیستم به طور موازی به یکدیگر متصل هستند. برای مثال، نوشین و مکیس [۵] سیستم موازی را در نظر گرفته‌اند که از تعداد معدودی از اجزاء تشکیل شده است به طوری که هر جزء دارای سه حالت: سالم، هشدار و شکست می‌باشد و حالت کلی سیستم از طریق پایش وضعیت در دوره‌های منظم نمونه‌گیری به دست می‌آید. نویسندگان با

^۱فریبا عزیزی: fariba.azizi@khayam.ut.ac.ir

هدف مینیم کردن متوسط هزینه در واحد زمانی، سعی در به دست آوردن سطوح تعویض فرصت طلبانه و پیشگرا نه در چارچوب فرایند تصمیم‌گیری نیمه‌مارکف را داشته‌اند.

ما نیز در این مقاله سیستم خاصی با اجزاء فراوان را در نظر گرفته‌ایم که به‌طور موازی و مستقل از یکدیگر در حال فعالیت هستند و در آن وضعیت هر جزء سیستم به سه حالت: سالم، هشدار و شکست طبقه‌بندی شده است. توربین‌های بادی می‌تواند مثال خوبی از این نوع سیستم‌ها باشد. از آنجایی که سیستم مورد نظر دارای ساختار پیچیده‌ای است، به منظور مدل‌بندی تغییرات تصادفی آن از فرایند تولد-مرگ دوبعدی [۳] استفاده کرده و سپس زمان اصابت سیستم به مجموعه‌ای از حالت‌ها که در آن حداقل r واحد به شکست رسیده باشد، را محاسبه خواهیم کرد.

۲ مدل و فرضیات

فرض کنید:

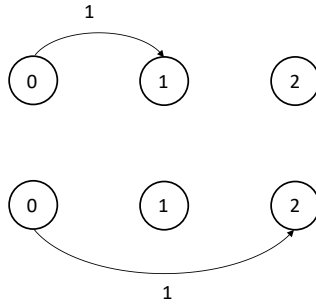
- سیستم دارای تعداد زیاد M جزء تخریب‌پذیر می‌باشد که به طور موازی و مستقل از یکدیگر مشغول به کار هستند.
- حالت هر جزء سیستم به سه رده طبقه‌بندی می‌شود:
حالت سالم (۰): حالتی که در آن یک جزء به خوبی کار می‌کند،
حالت هشدار (۱): حالتی که در آن یک جزء دچار عیب‌هایی شده است اما هنوز مشغول به کار است،
حالت شکست (۲): حالتی که در آن یک جزء کاملاً از کار افتاده و فعالیت آن متوقف شده است.
- هر جزء سیستم می‌تواند از حالت ۰ به ۱ با احتمال p_{01} ، از حالت ۱ به ۲ با احتمال p_{12} و از حالت ۰ به ۲ با احتمال p_{02} انتقال یابد که $p_{01} + p_{02} = 1$ و $p_{12} = 1$.
- زمان باقی‌ماندن در حالت i برای $i = \{0, 1\}$ دارای توزیع نمایی با پارامتر λ_i می‌باشد.

۱.۲ فرایند تولد/تولد-مرگ

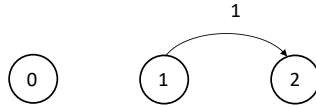
فرض کنید $N_0(t)$ ، $N_1(t)$ و $N_2(t)$ به ترتیب تعداد واحدها در حالت ۰، تعداد واحدها در حالت ۱ و تعداد واحدها در حالت ۲ باشد. با دلیل اینکه مجموع تعداد واحدها در سه حالت برابر با مقدار مشخص M می‌باشد، تنها به ردیابی تعداد واحدها در حالت ۱ و ۲ به کمک فرایند دوبعدی $\{N(t) = (N_1(t), N_2(t)), t \geq 0\}$ با فضای حالت $0 \leq i, j \leq M$ ، $S = \{(i, j) : i + j \leq M\}$ می‌پردازیم. فرایند فوق یک فرایند مارکف دوبعدی می‌باشد که برخلاف فرایند یک‌بعدی تولد-مرگ که تعداد یک جمعیت را ردیابی می‌کند، این فرایند تعداد اعضای دو جمعیت که در تعامل با یکدیگر هستند ردیابی می‌کند. واضح هست که مؤلفه‌ی اول این فرایند یعنی تعداد واحدها در حالت ۲ در حال افزایش می‌باشد از این رو این فرایند مارکف دوبعدی را فرایند تولد/تولد-مرگ می‌نامند.

در هر لحظه یکی از سه انتقال زیر امکان‌پذیر است:

- تغییر وضعیت یک جزء از حالت ۰ به ۱
- تغییر وضعیت یک جزء از حالت ۰ به ۲



• تغییر وضعیت یک جزء از حالت ۱ به ۲



هر یک از این انتقال‌های با احتمال‌های زیر رخ می‌دهد که آن‌ها نرخ انتقال لحظه‌ای می‌نامیم.

$$\begin{aligned} P(N(t+dt) = (i+1, j) | N(t) = (i, j)) &= \lambda_{ij}^{(1)} dt (1 - \lambda_{ij}^{(2)} dt) (1 - \gamma_{ij} dt) + o(dt) \\ &= \lambda_{ij}^{(1)} dt + o(dt), \end{aligned}$$

$$\begin{aligned} P(N(t+dt) = (i, j+1) | N(t) = (i, j)) &= \lambda_{ij}^{(2)} dt (1 - \lambda_{ij}^{(1)} dt) (1 - \gamma_{ij} dt) + o(dt) \\ &= \lambda_{ij}^{(2)} dt + o(dt), \end{aligned}$$

$$\begin{aligned} P(N(t+dt) = (i+1, j-1) | N(t) = (i, j)) &= \gamma_{ij} dt (1 - \lambda_{ij}^{(1)} dt) (1 - \lambda_{ij}^{(2)} dt) + o(dt) \\ &= \gamma_{ij} dt + o(dt), \end{aligned}$$

$$P(N(t+dt) = (i, j) | N(t) = (i, j)) = 1 - (\lambda_{ij}^{(1)} + \lambda_{ij}^{(2)} + \gamma_{ij}) dt + o(dt),$$

حال با توجه به نرخ‌های انتقال لحظه‌ای، احتمال‌های انتقال یک مرحله‌ای فرایند زمان پیوسته $\{N(t)\}$ به صورت زیر محاسبه می‌شود.

• برای $1 \leq j \leq M-i-1$, $0 \leq i \leq M-1$ داریم:

$$p_{(i,j),(i',j')} = \begin{cases} \frac{\lambda_{ij}^{(1)}}{\lambda_{ij}^{(1)} + \lambda_{ij}^{(2)} + \gamma_{ij}}, & i' = i + 1, j' = j \\ \frac{\lambda_{ij}^{(2)}}{\lambda_{ij}^{(1)} + \lambda_{ij}^{(2)} + \gamma_{ij}}, & i' = i, j' = j + 1 \\ \frac{\gamma_{ij}}{\lambda_{ij}^{(1)} + \lambda_{ij}^{(2)} + \gamma_{ij}}, & i' = i + 1, j' = j - 1 \end{cases}$$

• برای $0 \leq i \leq M - 1, j = 0$ داریم:

$$p_{(i,j),(i',j')} = \begin{cases} \frac{\lambda_{ij}^{(1)}}{\lambda_{ij}^{(1)} + \lambda_{ij}^{(2)}}, & i' = i + 1, j' = j \\ \frac{\lambda_{ij}^{(2)}}{\lambda_{ij}^{(1)} + \lambda_{ij}^{(2)}}, & i' = i, j' = j + 1 \end{cases}$$

• برای $0 \leq i \leq M - 1, j = M - i$ داریم:

$$p_{(i,j),(i',j')} = 1$$

که در آن $i' = i + 1, j' = j - 1$.

در ادامه به محاسبه‌ی احتمال‌های انتقال فرایند در مدت زمان t می‌پردازیم. فرض کنید:

$$P_{(i,j),(i',j')}(t) = Pr(N_1(t) = i', N_2(t) = j' | N(0) = (i, j))$$

احتمال آن است که فرایندی که در حال حاضر در حالت (i, j) است در مدت زمان t واحد بعدی در حالت (i', j') باشد. با بهره‌گیری از روش یکنواخت‌سازی که در نرم‌افزارهای ریاضی از سرعت محاسباتی بالایی برخوردار است، $P_{(i,j),(i',j')}(t)$ به صورت زیر به دست می‌آید.

$$P_{(i,j),(i',j')}(t) = \sum_{n=0}^{\infty} e^{-\nu t} \frac{(\nu t)^n}{n!} \bar{p}_{(i,j),(i',j')}^{(n)}, \quad (1)$$

که در آن $\bar{p}_{(i,j),(i',j')}^{(n)}$ به صورت بازگشتی از رابطه‌ی زیر به دست می‌آید.

$$\bar{p}_{(i,j),(i',j')}^{(n)} = \sum_{(k,h) \in S} \bar{p}_{(i,j),(k,h)}^{(n-1)} \bar{p}_{(k,h),(i',j')}, \quad n = 1, 2, \dots$$

۳ زمان اصابت

زمان اصابت که یکی از مفاهیم مطرح شده در فرایندهای تصادفی است، یک متغیر تصادفی است و بیانگر اولین زمانی است که فرایند وارد مجموعه‌ای از حالت‌ها می‌شود. این متغیر در فرایندهای تصادفی نقش اساسی دارد و می‌توان از آن برای نگهداری و تعمیرات پیشگیرانه یک سیستم تصادفی بهره گرفت.

فرض کنید τ_H اولین زمانی باشد که فرایند $\{N(t)\}$ به $H = \{(i, j) : i = r, 0 \leq j \leq M - r\}$ اصابت می‌کند. با در نظر گرفتن فرایند جدید $\{Z(t)\}$ که در آن $\lambda_{rj}^{(1)} = \lambda_{rj}^{(2)} = \gamma_{rj} = 0$ است، $Pr(\tau_H \leq t)$ را به صورت زیر محاسبه خواهیم کرد.

$$Pr(\tau_H \leq t) = \sum_{(r,j) \in H} P_{(0,0),(r,j)}(t), \quad (2)$$

که در آن $P_{(0,0),(r,j)}(t) = Pr(Z(t) = (r, j) | Z(0) = (0, 0))$.

۴ برآوردگر ماکسیمم درست‌نمایی

در این بخش به برآورد پارامترهای مجهول $\lambda_0, \lambda_1, p_{01}, p_{02}$ می‌پردازیم. با در نظر گرفتن $\theta = (\lambda_0, \lambda_1, p_{01}, p_{02})$ به عنوان برداری از پارامترهای مجهول، تابع درست‌نمایی برای برداری از مشاهدات به صورت زیر است:

$$l(\theta|D) = \sum_{i=0}^{n-1} \log P_{N(t_i), N(t_{i+1})}(t_{i+1} - t_i; \theta), \quad (3)$$

که در آن $D = \{N(t_0), N(t_1), \dots, N(t_n)\}$ برداری از مشاهدات به دست آمده از وضعیت سیستم در زمان‌های $t_0 = 0 < t_1 < \dots < t_n$ است.

مراجع

[۱] ف. کلاهان، م. دوست پرست و م. ماموریان، تعیین نوع و زمان‌بندی بهینه‌ی نگهداری و تعمیرات پیشگیرانه سیستم‌های چندجزئی بر اساس قابلیت اطمینان، نشریه دانشکده فنی، جلد ۴۱، شماره ۴، ۱۳۸۶.

- [2] Cho, D.I. and Parlar, M. (1991), A survey of maintenance models for multi-unit systems, *European journal of operational research*, **51**(1), 1-23.
- [3] Ho, L.S.T., Xu, J., Crawford, F.W., Minin, V.N. and Suchard, M.A. (2018), Birth/birth-death processes and their computable transition probabilities with biological applications, *Journal of mathematical biology*, **76**(4), 911-944.
- [4] Ramakumar, R. (1993), *Reliability Engineering: Fundamentals and Applications*.
- [5] Salari, N. and Makis, V. (2017), Comparison of two maintenance policies for a multi-unit system considering production and demand rates, *International Journal of Production Economics*, **193**, 381-391.
- [6] Wang, H. (2002), A survey of maintenance policies of deteriorating systems, *European journal of operational research*, **139**(3), 469-489.

بررسی متوسط زمان تا خرابی در سیستم‌های مختلف با مولفه‌ی آماده به کار

عقیل علایی^۱، نجمه شکیبایی زارع^۲

اگره آمار، دانشکده علوم ریاضی، دانشگاه صنعتی اصفهان

^۲ کمیته تحقیق دانشجویی، گروه آمار زیستی، دانشکده علوم پزشکی دانشگاه اصفهان، اصفهان

چکیده: همان‌طور که می‌دانیم بررسی سیستم‌های منسجم^۶ در حالتی که طول عمر مولفه‌های آن وابسته‌اند، حائز اهمیت است. در این تحقیق ابتدا تابع میانگین زمان رسیدن به خرابی در حالت وابسته را شرح می‌دهیم و کران بالا و پایینی برای تابع مذکور در حالت وابستگی بیان می‌کنیم و در آخر به بررسی رفتار^۳ نوع سیستم منسجم با مولفه‌ی آماده به کار می‌پردازیم. **واژه‌های کلیدی:** متوسط زمان تا خرابی، وابستگی، سیستم منسجم. کد موضوع بندی: 60K10، 90B25، 47A55.

۱ پیش‌گفتار

همان‌طور که می‌دانیم بررسی سیستم‌های منسجم^۷ در حالتی که طول عمر مولفه‌های آن وابسته‌اند حائز اهمیت است. در بسیاری از مقالاتی مانند [۲]، [۷] و [۴]، به بررسی ویژگی‌های سالخوردگی سیستم‌های منسجم با فرض استقلال بین طول عمر مولفه‌ها پرداخته شده که از نتایج آن‌ها می‌تواند در برخی از مسائل در حالت وابستگی، استفاده شود. به طور مثال در سیستم منسجمی که از چندین زیرسیستم و هر زیرسیستم از چندین مولفه تشکیل شده است، می‌توان طول عمر زیر سیستم‌ها را مستقل از هم در نظر گرفت. در بخش ۲ با استفاده از مفصل^۸ FGM کران بالا و پایینی برای میانگین مدت زمان تا خرابی^۹ سیستم بدست خواهیم آورد که از این کران‌ها، می‌توان مدت زمانی که برای هر سیستم ضمانت عدم خرابی شده است را نتیجه گرفت. [۱۱] و [۱۱] به بررسی کران‌های زمان خرابی سیستم با فرض استقلال و ناهم‌توزیع بودن طول عمر مولفه‌های سیستم اشاره کردند. در بخش آخر با استفاده از مولفه‌های آماده به کار^{۱۰} به عنوان یک روش برای افزایش قابلیت اطمینان سیستم، به بررسی معیارهای قابلیت اطمینان می‌پردازیم.

۲ متوسط زمان خرابی سیستم‌های منسجم دو مولفه‌ای با مولفه‌های وابسته

شاخص متوسط زمان تا خرابی یکی از پایه‌ای‌ترین معیارهای سنجش عملکرد سیستم‌های تعمیرناپذیر، که توسط بسیاری از دانشمندان به عنوان اولین معیار بررسی سیستم‌های منسجم مورد مطالعه قرار گرفته است، می‌باشد. برخی از این مطالعه‌ها با فرض وجود استقلال

^۱ a.alaei@math.iut.ac.ir

^۶ Coherent

^۷ Coherent

^۸ Copula

^۹ Mean Time To Failure

^{۱۰} Standby

یا ناهم‌توزیعی بین طول‌عمر مولفه‌های سیستم است که می‌توان به [۹] و [۱۰] اشاره کرد. فرض کنید X و Y طول‌عمر مولفه‌های یک سیستم منسجم با توابع توزیع حاشیه‌ای به ترتیب F_X و G_Y و تابع توزیع توام $H(x, y)$ است. تابع قابلیت اطمینان توام برای هر $x, y \in \mathbb{R}^+$ ، برابر با،

$$\bar{H}(x, y) = 1 - F_X(x) - G_Y(y) + H(x, y). \quad (۱)$$

در ادامه با توجه به ساختار سیستم، به بررسی ویژگی‌های تابع قابلیت اطمینان می‌پردازیم.

۱.۲ سیستم موازی با دومولفه‌ی وابسته

طول‌عمر یک سیستم موازی برابر با $T = \max(X, Y)$ است، پس متوسط زمان تا خرابی به صورت زیر تعریف می‌شود.

$$MTTF = E(T) = E(X) + E(Y) - \int_0^\infty \bar{H}(t, t) dt. \quad (۲)$$

با توجه به کران فرجیت-هوفیدینگ^{۱۱}، یک کران برای متوسط زمان تا خرابی سیستم دومولفه‌ای موازی به صورت زیر بدست می‌آید.

$$E(X) + E(Y) - \int_0^\infty \bar{M}(t, t) dt \leq E(T) \leq E(X) + E(Y) - \int_0^\infty \bar{W}(t, t) dt, \quad (۳)$$

که $W(t, t) = \max(F(t) + G(t) - 1, 0)$ و $M(t, t) = \min(F(t), G(t))$ است. اثبات. با توجه به کران فرجیت-هوفیدینگ برای هر تابع توزیع توام $H(x, y)$ و هر $x, y \in \mathbb{R}^+$ ،

$$W(x, y) \leq H(x, y) \leq M(x, y). \quad (۴)$$

بنابراین با توجه به رابطه‌ی (۱)،

$$\bar{W}(x, y) \leq \bar{H}(x, y) \leq \bar{M}(x, y).$$

که $\bar{W}(x, y)$ ، $\bar{H}(x, y)$ و $\bar{M}(x, y)$ توابع قابلیت اطمینان توام متعلق به توابع توزیع توام مربوطه است. با انتگرال‌گیری از رابطه‌ی بالا و با توجه به این‌که توابع توزیع حاشیه‌ی $M(x, y)$ و $W(x, y)$ ، برابر با $F(x)$ و $G(y)$ است داریم،

$$\int_0^\infty \bar{W}(t, t) dt \leq \int_0^\infty \bar{H}(t, t) dt \leq \int_0^\infty \bar{M}(t, t) dt.$$

پس،

$$E(X) + E(Y) - \int_0^\infty \bar{M}(t, t) dt \leq E(T) \leq E(X) + E(Y) - \int_0^\infty \bar{W}(t, t) dt.$$

□

بنابراین برای تمامی سیستم‌های دومولفه‌ای موازی می‌توان مقدار $E(X) + E(Y) - \int_0^\infty \bar{M}(t, t) dt$ را به عنوان متوسط ضمانت عدم خرابی سیستم در نظر گرفت. در حالتی که $F(x) = G(x)$ برای هر $x \in \mathbb{R}^+$ و با توجه به این‌که با این فرض $M(t, t) = F(t)$ ، آنگاه کران بالا به صورت زیر بازنویسی می‌شود،

$$E(X) \leq E(T) \leq 2E(X) - \int_0^\infty \bar{W}(t, t) dt. \quad (۵)$$

^{۱۱}Frechet-hoeffding

۲.۲ سیستم سری با دو مولفه‌ی وابسته

باتوجه به این‌که طول عمر یک سیستم سری برابر با $T^* = \min(X, Y)$ است، متوسط زمان تا خرابی به صورت زیر تعریف می‌شود.

$$MTTF = \int_0^\infty \bar{H}(t, t) dt. \quad (۶)$$

با استفاده از کران فرچیت-هوفیدینگ مشابه حالت موازی می‌توان یک کران برای متوسط زمان تا خرابی سیستم دومولفه‌ای سری به صورت زیر بدست آورد.

$$\int_0^\infty \bar{W}(t, t) dt \leq E(T^*) \leq \int_0^\infty \bar{M}(t, t) dt, \quad (۷)$$

با توجه به رابطه‌ی (۴) و با انتگرال گیری از دو طرف این رابطه، به راحتی می‌توان رابطه‌ی بالا را بدست آورد. در حالتی که برای هر $x \in \mathbb{R}^+$ ، $F(x) = G(x)$ ، آن‌گاه کران بالایی برای متوسط زمان تا خرابی سیستم دومولفه‌ای سری به صورت، $E(X)$ است. بدین معنی که برای هر سیستم دومولفه‌ای سری، حداکثر زمان برای کارکردن سیستم به طور متوسط برابر با $E(X)$ است. فرض کنید X و Y طول عمر مولفه‌های یک سیستم دومولفه‌ای سری با تابع توزیع مشترک $F(x) = 1 - e^{-\left(\frac{x}{\beta}\right)^\alpha}$ ، است. آن‌گاه،

$$E(X) = \int_0^\infty e^{-\left(\frac{t}{\beta}\right)^\alpha} dt = \frac{\beta}{\alpha} \Gamma\left(\frac{1}{\alpha}\right),$$

و،

$$\int_0^\infty \bar{W}(t, t) dt = \frac{\beta}{\alpha} \Gamma_{\ln 2}\left(\frac{1}{\alpha}\right) - \beta (\ln 2) \frac{1}{\alpha}.$$

بنابراین یک کران برای $MTTF$ سیستم با توجه به رابطه‌ی (۷) برابر است با

$$\frac{\beta}{\alpha} \Gamma_{\ln 2}\left(\frac{1}{\alpha}\right) - \beta (\ln 2) \frac{1}{\alpha} \leq E(T^*) \leq \frac{\beta}{\alpha} \Gamma\left(\frac{1}{\alpha}\right).$$

اگر در رابطه‌ی بالا $\alpha = 1$ یا به عبارت دیگر توزیع مشترک طول عمر مولفه‌های سیستم برابر با توزیع نمایی با پارامتر $\frac{1}{\beta}$ باشد، آن‌گاه یک کران برای متوسط زمان تا خرابی این سیستم به صورت زیر بدست می‌آید.

$$\beta(1 - \ln 2) \leq E(T^*) \leq \beta. \quad (۸)$$

اگر طرفین معادله‌های (۸) و نتیجه بدست آمده در حالت موازی را جمع کنیم، آن‌گاه همان‌طور که [۵] به آن اشاره کرد، این مجموع مقداری ثابتی است یا به عبارت دیگر مجموع میانگین طول عمر سیستم‌های موازی و سری همواره ثابت است.

مثال ۲.۲ را در نظر بگیرید به طوری که تابع توزیع توام از مفصل FGM بدست آید و همچنین توابع توزیع حاشیه‌ای توزیع وایبل با پارامترهای α و β ، باشد، آن‌گاه متوسط زمان تا خرابی یک سیستم دومولفه‌ای سری برابر است با

$$MTTF = \frac{\beta}{\alpha} \Gamma\left(\frac{1}{\alpha}\right) \left(2^{-\frac{1}{\alpha}} + \theta \left\{ 2^{-\frac{1}{\alpha}} + 4^{-\frac{1}{\alpha}} - 2(3^{-\frac{1}{\alpha}}) \right\} \right). \quad (۹)$$

با توجه به مفصل قابلیت اطمینان FGM برای توابع توزیع حاشیه‌ای وایبل داریم،

$$\begin{aligned}\bar{H}(t, t) &= e^{-\left(\frac{t}{\beta}\right)^{\alpha}} \left(1 + \theta \left(1 - e^{-\left(\frac{t}{\beta}\right)^{\alpha}} \right)^2 \right) \\ &= e^{-\left(\frac{t}{\beta}\right)^{\alpha}} + \theta \left\{ e^{-\left(\frac{t}{\beta}\right)^{\alpha}} + e^{-\left(\frac{t}{\beta}\right)^{\alpha}} - 2e^{-\left(\frac{t}{\beta}\right)^{\alpha}} \right\},\end{aligned}\quad (10)$$

حال متوسط زمان تا خرابی سیستم را به صورت زیر نوشته می‌شود،

$$\begin{aligned}\int_0^{\infty} \bar{H}(t, t) dt &= \int_0^{\infty} e^{-\left(\frac{t}{\beta}\right)^{\alpha}} dt + \theta \left\{ \int_0^{\infty} e^{-\left(\frac{t}{\beta}\right)^{\alpha}} dt + \int_0^{\infty} e^{-\left(\frac{t}{\beta}\right)^{\alpha}} dt - 2 \int_0^{\infty} e^{-\left(\frac{t}{\beta}\right)^{\alpha}} dt \right\} \\ &= \frac{\beta}{\alpha} \Gamma\left(\frac{1}{\alpha}\right) \left(2^{-\frac{1}{\alpha}} + \theta \left\{ 2^{-\frac{1}{\alpha}} + 2^{-\frac{1}{\alpha}} - 2 * 3^{-\frac{1}{\alpha}} \right\} \right).\end{aligned}$$

در حالتی که $\alpha = 1$ باشد، متوسط زمان تا خرابی برابر است با،

$$MTTF = \frac{\beta}{2} - \frac{\theta\beta}{12} = \frac{\beta(6 + \theta)}{12}.$$

همان‌طور که در رابطه‌ی بالا مشاهده می‌شود، رابطه‌ی مستقیمی بین متوسط زمان تا خرابی سیستم دومولفه‌ای سری و پارامتر وابستگی FGM وجود دارد، به این معنی که با افزایش مقدار پارامتر وابستگی، متوسط زمان تا رسیدن به خرابی برای سیستم دومولفه‌ای موازی نیز افزایش می‌یابد، که این امر نیز با ارتباط بین خانواده توزیع‌های دارای وابستگی مربعی مثبت و سیستم‌های سری مطابق است. داده‌های جدول ۱ مربوط به طول عمرهای دو نوع باتری در موتور تولید برق است، که به صورت سری به یکدیگر متصل شده‌اند. با استفاده از آزمون کلموگروف اسمیرنوف توزیع طول عمرهای این دو نوع باتری دارای توزیع مشترک وایبل با پارامترهای

باتری نوع اول	۳/۸	۳/۱	۴/۳	۱/۶	۴/۱	۴/۷	۳/۳	۲/۵	۳/۴	۲/۲
باتری نوع دوم	۳/۳	۴/۱	۳/۳	۳/۱	۳/۷	۳/۹	۳/۲	۲/۹	۳/۸	۳/۲
باتری نوع اول	۲/۶	۴/۴	۳/۶	۳/۳	۴/۵	۳/۲	۳/۷	۳/۴	۳/۱	۳/۸
باتری نوع دوم	۳/۵	۳/۴	۳/۴	۳/۷	۲/۶	۴/۲	۱/۹	۳/۹	۴/۷	۳/۳

جدول ۱: طول عمر دونوع باتری در موتورهای تولید برق بر حسب ساعت

اطلاعات مربوط به این آزمون در جدول ۲ آورده شده است. $\alpha = 5/24$ و $\beta = 3/73$ است.

با استفاده از آزمون آکائیک و مفصل‌های پیشنهادی FGM ، AMH ، 12 فرانک^{۱۳} و کلایتون^{۱۴}، تابع توزیع توأم برازش شده برای این داده‌ها از مفصل FGM تبعیت می‌کنند. اطلاعات مربوط به این آزمون نیز در جدول ۳ آورده شده است. بنابراین مقدار

^{۱۲}Ali-Mikhael-Haq

^{۱۳}Frank

^{۱۴}Clynton

مقدار خطای برآورد α	مقدار خطای برآورد β	α	β	p -مقدار	مقدار آماره آزمون	باطری نوع اول
۰/۹۳	۰/۱۶	۵/۲۴	۳/۷۳	۰/۹۵۱۸	۰/۱۱۵	
مقدار خطای برآورد α	مقدار خطای برآورد β	α	β	p -مقدار	مقدار آماره آزمون	باطری نوع دوم
۱/۰۷	۰/۱۳	۵/۲۴	۳/۷۳	۰/۸۶۱۴	۰/۱۳۴	

جدول ۲: مقادیر آزمون کلموگروف برای داده‌های جدول ۱

کلايتون	فرانک	FGM	AMH	مفصل
۸۵.۹۴	۸۵.۷۱	۸۵.۶۰۶۴	۸۵.۶۲۲۲	مقدار آکائیک
-۰.۰۵۳	-۰.۸۱	-۰.۴۱	-۰.۶۴۴	مقدار برآورد پارامتر وابستگی

جدول ۳: مقادیر آزمون آکائیک برای داده‌های جدول ۱

متوسط مدت زمان تا خرابی برای این سیستم برابر است با $MTTF = ۲/۹۹۷$ ، که در مقایسه با متوسط طول عمر هر کدام از باطری‌ها مقدار کمتری است ($E(X) = ۳/۴۳$)، زیرا مقدار پارامتر وابستگی عددی منفی است و تاثیر معکوسی در عملکرد سیستم سری دارد.

۳ سیستم سری-موازی دو مولفه‌ای وابسته همراه با مولفه‌ی آماده به کار

حالت آماده به کار فزونگی^{۱۵} یک روش موثر برای افزایش و بهبود بخشیدن به قابلیت اطمینان سیستم است. در مطالعه‌ی قابلیت اطمینان سیستم‌ها، اغلب فرض شده است که خرابی مولفه‌های موجود در سیستم بر عملکرد دیگر مولفه‌ها تاثیری ندارند. در دنیای واقعی با حالت‌هایی مواجه می‌شویم که این‌گونه فرضیه‌ها جایگاه کارآمدی ندارند و می‌بایست برای تحقیق و مطالعه‌ی بر روی سیستم‌ها، فرض وجود وابستگی بین مولفه‌ها یک فرض اساسی و اصلی در نظر گرفته شود. نکته قابل توجه دیگر برای واقعی سازی مطالعه بر روی سیستم‌ها، در نظر گرفتن یک مولفه‌ی آماده به کار در سیستم است. مطالعات بسیاری بر روی سیستم‌ها با فرض وابستگی بین مولفه‌ها بدون در نظر گرفتن مولفه‌ی آماده به کار انجام شده است که می‌توان به مقالاتی چون [۱۲] و [۱۳] اشاره کرد. در حالی که در بسیاری از سیستم‌های حساس در دنیای واقعی، مولفه‌ی آماده به کار یکی از مهمترین عضو سیستم است، چرا که هم از لحاظ امنیتی و هم از لحاظ بهبود بخشیدن به قابلیت اطمینان سیستم، مولفه‌ی کارآمدی است. لازم به ذکر است که، به دلیل این که اغلب سیستم‌های موجود در صنعت، به صورت سیستم‌های موازی، سری و یا ترکیبی از این دو نوع هستند، مفید است که به مطالعه‌ی سیستم‌های سری و موازی با فرض‌های ذکر شده بپردازیم. در این‌گونه سیستم‌ها فرض کردیم، مولفه‌ی آماده به کار تا زمانی که در حالت آماده به کار است هرگز غیرفعال نمی‌شود و با توجه به ساختار سیستم وقتی که سیستم غیرفعال یا خراب شد، به سرعت مولفه آماده به کار جایگزین مولفه‌ی غیرفعال می‌شود و به طور مشابه وظیفه مولفه‌ی غیرفعال را تا زمانی که مولفه‌ی غیرفعال تعمیر شود، در سیستم انجام می‌دهد. برای مطالعه بیشتر و دقیق‌تر در مورد عمل‌کرد مولفه‌های آماده به کار در سیستم می‌توان به مقالات [۱۵] و [۱۴] مراجعه کرد.

در ادامه این بخش، ابتدا متغیر تصادفی طول عمر سیستم وقتی که سیستم دومولفه‌ای سری یا موازی با یک مولفه آماده به کار است، تعریف و سپس به بررسی شاخص سالخوردگی متوسط زمان تا خرابی با این فرضیه‌ها اشاره شده است. همچنین یک کران

^{۱۵}Standby Redundancy

برای تابع متوسط زمان تا خرابی سیستم وقتی که مولفه‌ها نسبت به هم وابستگی مثبت یا منفی مربعی دارند، معرفی می‌شود. لازم به ذکر است که این کران در عمل می‌تواند حائز اهمیت باشد؛ چرا که برای محقق مفید است که بداند سیستم مورد نظر حداقل تا چه واحدی از زمان خراب یا غیرفعال نمی‌شود.

۱.۳ توصیف سیستم

یک سیستم دومولفه‌ای وابسته با یک مولفه آماده به کار به طوری که مولفه‌ها به صورت سری و یا موازی به یکدیگر متصل شده‌اند را در نظر بگیرید. در زمان $t = 0$ هر دو مولفه اصلی سیستم، فعال‌اند و زمانی که سیستم غیرفعال شود مولفه‌ی آماده به کار جایگزین مولفه‌ی خراب می‌شود. در حالت سری، مولفه‌ی آماده به کار می‌تواند به دو صورت سری یا موازی با سایر مولفه‌ها سیستم متصل شود، بنابراین با توجه به ساختار کل سیستم به بررسی ویژگی‌های سیستم می‌پردازیم.

۲.۳ سیستم سری دومولفه‌ای با مولفه‌های آماده به کار سری

اگر طول عمر متناظر با مولفه‌های اصلی را با X_1 و X_2 و طول عمر مولفه آماده به کار را با Z نشان دهیم، آنگاه [۶] طول عمر سیستم را به صورت زیر توصیف کرد.

$$T = \min(X_1, X_2) + \min(X_1^*, Z), \quad (11)$$

که X_1^* بیانگر طول عمر باقی‌مانده از مولفه‌ای که بعد از اولین خرابی سیستم، غیرفعال می‌شود و به صورت زیر تعریف می‌شود.

$$X_1^* = (X_1 - \min(X_1, X_2)) | X_1 > \min(X_1, X_2). \quad (12)$$

اگر فرض کنیم طول عمر مولفه‌های سیستم هم‌توزیع است، آنگاه تابع توزیع متغیر تصادفی X_1^* به صورت زیر بدست می‌آید،

$$\begin{aligned} F^*(t) &= P(X_1^* \leq t) = P(X_1 - \min(X_1, X_2) \leq t | X_1 > \min(X_1, X_2)) \\ &= \frac{P(\min(X_1, X_2) < X_1 \leq t + \min(X_1, X_2))}{P(X_1 > \min(X_1, X_2))} \\ &= \frac{P(X_1 \leq t + \min(X_1, X_2)) - P(X_1 \leq \min(X_1, X_2))}{P(X_1 > \min(X_1, X_2))} \\ &= \frac{P(X_1 \leq t + X_2) - P(X_1 \leq X_2)}{P(X_1 > X_2)}. \end{aligned}$$

فرض کنید X_1 و X_2 متغیرهای تصادفی تعویض‌پذیری^{۱۶} است پس،

$$P(X_1 \leq X_2) = P(X_1 > X_2) = \frac{1}{2},$$

و تابع توزیع متغیر تصادفی X_1^* به صورت زیر است.

$$\begin{aligned} F^*(t) &= 2P(X_1 \leq t + X_2) - 1 \\ &= 2 \int_0^\infty \int_0^{y+t} c(F(x), F(y)) dF(x) dF(y) - 1, \end{aligned} \quad (13)$$

^{۱۶}Exchangable

$$c(u, v) = \frac{\partial^2}{\partial u \partial v} C(u, v) \text{ به طوری که}$$

در دنیای واقعی معمولاً مولفه‌ی آماده به کار در یک سیستم عملکرد مشابهی با مولفه‌های اصلی سیستم دارد به همین دلیل منطقی است اگر توزیع طول عمر مولفه آماده به کار هم توزیع با مولفه‌های اصلی سیستم در نظر گرفت. بنابراین با فرض هم توزیع بودن طول عمر مولفه آماده به کار با مولفه‌های اصلی سیستم برای تمامی $t_1, t_2 \geq 0$ ، توابع قابلیت اطمینان (X_1, X_2) و (X_1^*, Z) به ترتیب برابر است با

$$P\{X_1 > t_1, X_2 > t_2\} = 1 - F(t_1) - F(t_2) + C(F(t_1), F(t_2)), \quad (14)$$

و،

$$P\{X_1^* > t_1, Z > t_2\} = 1 - F^*(t_1) - F(t_2) + C(F^*(t_1), F(t_2)). \quad (15)$$

با استفاده از رابطه‌های (14) و (15) می‌توان متوسط زمان تا خرابی سیستم تعریف شده را به صورت زیر بدست آورد.

$$MTTF(T) = E(\min(X_1, X_2) + \min(X_1^*, Z)) \quad (16)$$

$$= \int_0^\infty 1 - 2F(t) + C(F(t), F(t))dt + \int_0^\infty 1 - F(t) - F^*(t) + C(F^*(t), F(t))dt.$$

فرض کنید (X_1, X_2) و (X_1^*, Z) بردارهای تصادفی که دارای وابستگی مثبت مربعی است، آن‌گاه می‌توان برای متوسط زمان تا خرابی سیستم، یک کران پایین به صورت زیر بدست آورد.

$$E(T) \geq \int_0^\infty (1 - F(t))^2 dt + \int_0^\infty (1 - F(t))(1 - F^*(t))dt, \quad (17)$$

واضح است که فرمول سمت راست رابطه‌ی (17)، متوسط زمان تا خرابی یک سیستم دو مولفه‌ای سری با مولفه آماده به کار است وقتی که طول عمر مولفه‌ها از هم دیگر مستقل باشند. به طور مشابه اگر (X_1, X_2) و (X_1^*, Z) بردارهای تصادفی که دارای وابستگی منفی هستند، باشند آن‌گاه،

$$E(T) \leq \int_0^\infty (1 - F(t))^2 dt + \int_0^\infty (1 - F(t))(1 - F^*(t))dt. \quad (18)$$

قابل توجه است که چنین کرانی توسط [11] برای سیستم‌های موازی دو مولفه‌ای بدون مولفه‌ی آماده به کار نیز بدست آمده است. اگر تابع توزیع توأم طول عمر مولفه‌های سیستم از مفصل FGM مدل‌بندی شده باشند، آن‌گاه تابع توزیع X_1^* ، برای $t \geq 0$ ، با توجه به رابطه‌ی (13) به صورت زیر است،

$$\begin{aligned} F^*(t) &= 2 \int_0^\infty \int_0^{y+t} 1 + \alpha(1 - 2F(x))(1 - 2F(y))dF(x)dF(y) - 1 \\ &= 2 \int_0^\infty F(y+t)(1 + \alpha(1 - F(y+t))(1 - 2F(y))dF(y) - 1. \end{aligned} \quad (19)$$

فرض کنید توزیع طول عمر مولفه‌های سیستم داری توزیع نمایی با پارامتر λ است، پس برای $t \geq 0$ ،

$$F^*(t) = 1 - e^{-\lambda t} + \frac{\alpha}{3} e^{-\lambda t} (1 - e^{-\lambda t}). \quad (20)$$

تابع چگالی X_1^* با مشتق گرفتن از رابطه (20) نسبت به t به صورت زیر بدست می‌آید.

$$\begin{aligned} f^*(t) &= \frac{d}{dt} F^*(t) \\ &= \lambda e^{-\lambda t} (1 - \frac{\alpha}{3}) + \frac{2\lambda\alpha}{3} e^{-2\lambda t}. \end{aligned} \quad (21)$$

از رابطه (۲۰) می‌توان نتیجه گرفت که اگر $\alpha = 0$ ، یا به عبارت دیگر طول عمر مولفه‌ها مستقل از هم باشد، $F^*(t) = 1 - e^{-\lambda t}$ است. این موضوع به دلیل خاصیت عدم حافظه در توزیع نمایی و بیانگر این است که خرابی مولفه اول تاثیری در مدت زمان باقی‌مانده طول عمر مولفه دوم ندارد. برای $t \geq 0$ ، تابع قابلیت اطمینان متغیر تصادفی $\min(X_1, X_2)$ وقتی که متغیرهای X_1 و X_2 دارای توزیع نمایی با پارامتر λ باشد، برابر است با،

$$\begin{aligned} P\{\min(X_1, X_2) > t\} &= 1 - 2F(t) + C(F(t), F(t)) \\ &= (1 - \alpha)e^{-2\lambda t} - 2\alpha e^{-3\lambda t} + \alpha e^{-4\lambda t}. \end{aligned} \quad (22)$$

اگر برای بدست آوردن تابع قابلیت اطمینان $\min(X_1^*, Z)$ برای تمامی $t \geq 0$ از مفصل قابلیت اطمینان FGM استفاده شود آن‌گاه،

$$\begin{aligned} P\{\min(X_1^*, Z) > t\} &= \bar{F}^*(t)\bar{F}(t)\{1 + \alpha F^*(t)F(t)\} \\ &= (1 - \frac{\alpha}{3}(1 - e^{-t\lambda}))e^{-2\lambda t}\{1 + \alpha(1 + \frac{\alpha}{3}e^{-t\lambda})(1 - e^{-t\lambda})^2\}. \end{aligned} \quad (23)$$

در نهایت با استفاده از روابط (۱۶)، (۲۲) و (۲۳)، متوسط زمان تا خرابی سیستم برای $\alpha \in [-1, 1]$ ، به صورت،

$$E(T) = \frac{1}{\lambda} \left\{ 1 + \frac{\alpha}{9} - \frac{\alpha^2}{180} - \frac{\alpha^3}{540} \right\},$$

است. که نمودار این تابع در ۱ نمایش داده شده است.

۳.۳ سیستم سری با مولفه‌های وابسته و مولفه‌ی آماده به کار موازی

طول عمر یک سیستم موازی با یک مولفه آماده به کار موازی برابر است با

$$T_1 = \min(X_1, X_2) + \max(X_1^*, Z). \quad (24)$$

شاخص، متوسط زمان تا خرابی سیستم با استفاده از رابطه‌های (۱۴) و (۱۵) به صورت زیر است.

$$\begin{aligned} E(T_1) &= E(\min(X_1, X_2) + \max(X_1^*, Z)) \\ &= \int_0^\infty 1 - 2F(t) + C(F(t), F(t))dt + \int_0^\infty 1 - C(F^*(t), F(t))dt. \end{aligned} \quad (25)$$

فرض کنید (X_1, X_2) و (X_1^*, Z) بردارهای تصادفی دارای وابستگی مثبت مربعی، باشند آن‌گاه برای متوسط زمان تا خرابی سیستم، یک کران پایین به صورت زیر می‌توان بدست آورد.

$$E(T_1) \geq \int_0^\infty (1 - F(t))^2 dt + \int_0^\infty 1 - F(t)F^*(t)dt. \quad (26)$$

سمت راست رابطه‌ی (۲۶)، متوسط زمان تا خرابی یک سیستم دومولفه‌ای سری با مولفه آماده به کار موازی است وقتی که طول عمر مولفه‌ها از هم دیگر مستقل باشد. به طور مشابه اگر (X_1, X_2) و (X_1^*, Z) بردارهای تصادفی که با وابستگی منفی باشند آن‌گاه،

$$E(T_1) \leq \int_0^\infty (1 - F(t))^2 dt + \int_0^\infty 1 - F(t)F^*(t)dt. \quad (27)$$

همچنین تابع قابلیت اطمینان $\max(X_1^*, Z)$ برای تمامی $t \geq 0$ با استفاده از مفصل قابلیت اطمینان FGM ، برابر است با

$$\begin{aligned} P\{\max(X_1^*, Z) > t\} &= 1 - [F^*(t)F(t) \{1 + \alpha \bar{F}^*(t)\bar{F}(t)\}] \\ &= 1 - \left[\left(1 + \frac{\alpha}{3} e^{-\lambda t}\right) (1 - e^{-\lambda t})^2 \left\{ 1 + \alpha e^{-2\lambda t} \left(1 - \frac{\alpha}{3} (1 - e^{-\lambda t})\right) \right\} \right] \end{aligned} \quad (28)$$

در نهایت با استفاده از (۱۶)، (۲۲) و (۲۸)، متوسط زمان تا خرابی سیستم برای $\alpha \in [-1, 1]$ ، به صورت زیر بدست می‌آید.

$$E(T_1) = \frac{1}{\lambda} \left\{ \frac{3}{2} - \frac{\alpha}{36} + \frac{\alpha^2}{20} - \frac{\alpha^3}{540} \right\}.$$

نمودار تابع بالا در ۱ نمایش داده شده است.

۴.۳ سیستم موازی با مولفه‌های وابسته و مولفه‌ی آماده به کار

در این حالت، زمانی که هر دو مولفه‌ی اصلی سیستم غیرفعال شوند، مولفه‌ی آماده به کار، شروع به انجام وظیفه‌ی خود می‌کند. بنابراین طول عمر سیستم به صورت زیر تعریف می‌شود.

$$T_2 = \max(X_1, X_2) + Z. \quad (29)$$

همان‌طور که واضح است در رابطه‌ی بالا دیگر نیازی به محاسبه‌ی توزیع X^* نیست. بنابراین اگر همانند حالت‌های قبلی توزیع طول عمر مولفه‌ها دارای توزیع نمایی با پارامتر λ و تابع توزیع توأم از مفصل FGM پیروی کنند، متوسط زمان تا خرابی به صورت زیر بدست می‌آید.

$$E(T_2) = E(\max(X_1, X_2)) + E(Z),$$

که برای توزیع نمایی با پارامتر λ داریم،

$$\begin{aligned} E(T_2) &= \frac{18 - \alpha}{12\lambda} + \frac{1}{\lambda} \\ &= \frac{30 - \alpha}{12\lambda}. \end{aligned}$$

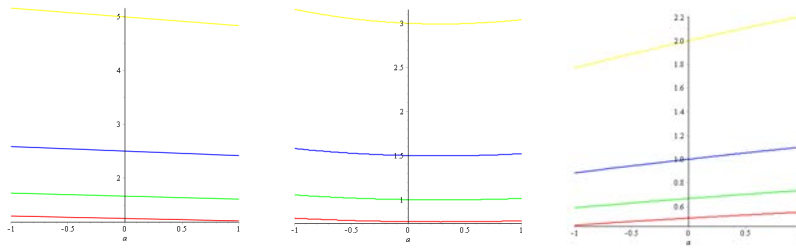
نمودار این تابع برای برخی از پارامترها در (۱) رسم شده است. همچنین در جدول ۴ برخی از مقادیر مختلف برای متوسط کارکرد سیستم در حالت‌های مختلف را بیان کردیم که با توجه به این جدول همواره برای هر α و λ ،

$$MTTF(T) \leq MTTF(T_1) \leq MTTF(T_2). \quad (30)$$

بنابراین با توجه به نتایج بدست آمده پیشنهاد می‌شود ساختار سیستم به صورت حالت سوم متصل شود، اما در برخی از موارد متصل کردن مولفه‌ها به صورت موازی امکان‌پذیر نیست برای رفع این مشکل بهترین نوع اتصال مولفه‌ها، به صورت حالت دوم است.

مراجع

- [1] Aitken, A.C. and Gonin, H.T. (1935), On fourfold sampling with and without replacement, in *Proc. Roy. Soc. Edinburgh*, **55**, 114-125.



شکل ۱: نمودار $MTTF(T)$ (راست)، $MTTF(T_1)$ (وسط) و $MTTF(T_2)$ (چپ) برای $\lambda = 2$ (قرمز)، $\lambda = 1$ (آبی)، $\lambda = 1/5$ (سبز) و $\lambda = 2$ (قرمز)

$(\lambda, \alpha) \mapsto$	$(2, -0.5)$	$(1, -0.5)$	$(0.5, -0.5)$	$(0.5, 0)$	$(0.5, 1)$	$(0.5, 2)$	$(0.5, 2)$	$(0.5, 1)$	$(0.5, 0.5)$
$MTTF(T)$	0.47	0.94	1.88	0.50	1.00	2.00	0.52	1.05	2.10
$MTTF(T_1)$	0.76	1.52	3.05	0.75	1.50	3.00	0.74	1.49	2.99
$MTTF(T_2)$	1.27	2.54	5.08	1.25	2.50	5.00	1.22	2.45	4.91

جدول ۴: مقادیر متوسط زمان کارکرد برای حالت‌های متفاوت

- [2] Asadi, M., Bairamov. I.G. (2005), A note on the mean residual life function of a parallel system, *Communication in Statistics-theory and Methods*, **34**(2), 475-484,
- [3] Bayramoglu, (2014), Reliability and Mean Residual Life of Complex Systems With Two Dependent Components Per Element, *IEEE Transactions on Reliability*, **62**(1).
- [4] Charles, E. (1997), An Introduction to Reliability and Maintainability Engineering, *Boston: McGraw-Hill*, 23-32.
- [5] Chin Diew Lai and Min Xie. *Stochastic Ageing and Dependence for Reliability*, **Springer**.
- [6] Eryilmaz, S. (2014), Reliability properties of consecutive k-out-of-n systems of arbitrarily dependent components, *Reliability Engineering and System Safety*, **94**, 350-356.
- [7] Guess, F. and Proschan, F. (1988), Mean residual life. Theory and Applications, *Handbook of Statistics*, **7**, 215-224
- [8] Kotz, S., Lai, C.D. and Xie, M. (2003), On the effect of redundancy for systems with dependent components, *IIE Transactions*, **35**, 1103-1110.
- [9] Mitchell, J. (2002), Approximation of Mean Time Between Failure When a System has Periodic Maintenance, *IEEE Transactions on Reliability*, **51**.
- [10] Modarres, M. and Kaminskiy, M. (2010), Reliability Engineering and Risk Analysis, *A Practical Guide (2nd ed.)*. CRC Press.

- [11] Navarro, J. and Rychlik, T. (2007), Reliability and expectation bounds for coherent systems with exchangeable components, *Journal of Multivariate Analysis*, **98**, 102-113.
- [12] Navarro, J., Ruiz, J.M. and Sandoval, C.J. (2007), Properties of coherent systems with dependent components, *Communications in Statistics Theory and Methods*, **36**, 175-191.
- [13] Navarro, J. and Spizzichino, F. (2010), Comparisons of series and parallel systems with components sharing the same copula, *Applied Stochastic Models in Business and Industry*, **26**, 775-791.
- [14] Wang, G.J. and Zhang, Y.L. (2011), A bivariate optimal replacement policy for a cold standby repairable system with preventive repair, *Applied Mathematics and Computation*, **218**, 3158-3165.
- [15] Wu, Q. and Wu, S. (2011), Reliability analysis of two-unit cold standby repairable systems under Poisson shocks, *Applied Mathematics and Computation*, **218**, 171-182.

مدل تطبیق بیزی موجک انقباضی در برآورد ناپارامتری تابع رگرسیون

حمید کرمی کبیر^۱، روح الله حقیقت طلب، محمود افشاری

گروه آمار، دانشگاه خلیج فارس بوشهر

چکیده: آنالیز موجکی یکی از تکنیک‌های مفید در ریاضی است که در سال‌های اخیر نیز در علم آمار بسیار مورد استفاده قرار گرفته است. یکی از روش‌هایی که در نظریه برآوردها مورد استفاده قرار می‌گیرد، روش موجک است. در این مقاله پس از بیان مختصری از موجک انقباضی بیزی، مدل تطبیق بیزی موجکی را در برآورد ناپارامتری تابع $f(x)$ معرفی نموده و سپس توزیع ضرایب آن محاسبه می‌گردد.

واژه‌های کلیدی: آستانه، برآورد بیزی، بیز تطبیقی، مدل‌های آمیخته، موجک انقباضی، نوفه‌زدایی.

کد موضوع بندی: 60T65, 08G62, 05G62.

۱ مقدمه

در برازش مدل‌های مبتنی بر موجک، انقباض ضرایب تجربی موجک، یک ابزار موثر برای نوفه‌زدایی^{۱۷} داده‌ها محسوب می‌شود. یک رویکرد انقباضی بیزی، با در نظر گرفتن توزیع‌های پیشین مختلف روی ضرایب موجک به دست می‌آید. این ضرایب از دو توزیع نرمال با انحراف استانداردهای متفاوتی تشکیل شده‌اند. روش‌های استاندارد موجک، موج‌ها را با نوفه اضافی در نظر می‌گیرند و از طریق انقباض ضرایب تجربی موجک به دنبال نوفه‌زدایی هستند. در این زمینه می‌توان به تحقیقات [۳]، [۵] و [۱] مراجعه کرد. یکی از روش‌های طبیعی برای به دست آوردن برآوردهای انقباضی ضرایب صحیح، استفاده از روش‌های بیزی می‌باشد. در رویکرد بیزی، یک توزیع پیشین روی هر ضریب فرض می‌شود. در این مقاله پس از بیان مختصری از موجک انقباضی بیزی، فرض می‌کنیم برآورد درستی از سطح نوفه وجود دارد و به همین دلیل آن را یک پارامتر معلوم در نظر می‌گیریم. این کار امکان آن را می‌دهد که فرم‌های بسته‌ای برای میانگین‌های پسین و واریانس‌های ضرایب صحیح موجک به دست بیاوریم.

۲ موجک انقباضی بیزی

موجک‌ها یک ابزار قوی برای تقریب توابع هستند. برآوردهای موجک را می‌توان به عنوان یک نوع خاص، از برآوردهای سری‌های متعامد استفاده کرد زیرا موجک‌ها می‌توانند از پایه‌های یکامتعامد برای توابع فضایی گوناگون استفاده کنند. فرض کنید $\{X_n, n \geq 1\}$ یک دنباله از متغیرهای تصادفی روی فضای احتمال $(\Omega, \mathcal{N}, P)$ باشند. همچنین فرض کنید که X_i ها کران‌دار هستند و تکیه‌گاه فشرده چگالی حاشیه‌ای $f(x)$ ، نسبت به اندازه لیگ، به i وابسته نباشند. این چگالی را با n مشاهده

^۱ حمید کرمی کبیر: h_karamikabir@yahoo.com

$i = 1, \dots, n, X_i$ برای هر تابع $f \in L^2(\mathbb{R})$ می‌توان یک بسط قراردادی نوشت:

$$\hat{f} = \sum_{k \in \mathbb{Z}} \hat{\alpha}_{m,k} \phi_{m,k} + \sum_{j=1}^{\infty} \sum_{k \in \mathbb{Z}} \hat{\delta}_{j,k} \psi_{j,k} \quad (1)$$

که $\psi_{j,k}$ و $\phi_{j,k}(x)$ به ترتیب تابع مقیاس و موجک متعامد هستند. طوری که

$$\phi_{j,k}(x) = 2^j I\left[\frac{k}{2^j}, \frac{k+1}{2^j}\right](x), \quad \psi_{j,k}(x) = 2^{j-1} I\left[\frac{k}{2^j}, \frac{k+1}{2^j}\right](x) - 2^{j-1} I\left[\frac{k-1}{2^j}, \frac{k}{2^j}\right](x),$$

برای همین هر تابع $f \in L^2(\mathbb{R})$ را می‌توان در شرایطی از موجک‌های یکامتعامد $\{\psi_{j,k}(x)\}$ توسط معادله زیر نمایش داد:

$$f(x) = \sum_{j \in \mathbb{Z}} \sum_{k \in \mathbb{Z}} d_{jk} \psi_{j,k}(x) \quad (2)$$

که $d_{jk} = \int_{\mathbb{R}} f(x) \psi_{j,k}(x) dx$ ضرایب موجک از تابع f هستند. برای محاسبه ضرایب موجک در یک نمونه‌ی داده شده به اندازه‌ی $2^j = n$ می‌توان از تبدیل گسسته موجک استفاده کرد. فرض کنید معادله زیر را داشته باشیم:

$$y_i = f(t_i) + \epsilon_i, \quad i = 1, 2, \dots \quad (3)$$

که ϵ_i تعدادی نوفه با واریانس σ^2 است و $t_i = \frac{i}{n}$ یک تابع می‌باشد. تبدیل موجک گسسته می‌تواند به وسیله ماتریس متعامد W معرفی شود. سپس $d_i = W y_i$ تبدیل موجک روی داده‌های نوفه‌دار را اجرا می‌کند. که ضرایب موجک با روش‌هایی به فرم یک آرایه از ضرایب، $\hat{d}$ را تعدیل و اصلاح می‌کنند و سپس تبدیل معکوس W^T برای محاسبه،

$$\hat{f} = W^T \hat{d} \quad (4)$$

به کار می‌رود. که $\hat{f}$ یک برآورد از f در نقاط t_i است. جزئیات بیشتر از این تبدیل را می‌توان در [5] یافت.

روش موجک انقباضی مبتنی بر دو پارامتر اساسی مقدار آستانه و تابع حذف است. با مشخص شدن مقدار آستانه و تابع حذف، مراحل الگوریتم موجک انقباضی بصورت زیر خواهد بود:

(۱) ابتدا از مشاهدات اغتشاشی تبدیل موجک گسسته گرفته می‌شود. به عبارت دیگر، با فرض اینکه ماتریس تبدیل موجکی گسسته W ، و مدل مشاهدات به صورت $y = f + \epsilon$ باشد ابتدا ضرایب موجکی را از تبدیل موجکی زیر بدست می‌آوریم.

$$W y = W f + W \epsilon \Rightarrow d^* = d + \epsilon,$$

(۲) با استفاده از مقدار آستانه، ضرایب موجکی به دو دسته ضرایب پر اهمیت و ضرایب کم اهمیت دسته‌بندی می‌شوند به اینصورت که اگر ضریب موجکی بیشتر از مقدار آستانه باشد، متعلق به دسته ضرایب پر اهمیت است و در غیر این صورت جزو دسته ضرایب کم اهمیت خواهد بود. سپس ضرایب کم اهمیت حذف شده و ضرایب پر اهمیت با توجه به تابع حذف، به صورت زیر اصلاح می‌شوند: در تابع حذف سخت، ضرایب کمتر از مقدار آستانه برابر صفر و باقی ضرایب بدون تغییر باقی می‌مانند. در تابع حذف نرم، ضرایب کمتر از مقدار آستانه برابر صفر و باقی ضرایب به اندازه مقدار آستانه به صفر نزدیک می‌شوند.

(۳) سیگنال حاصل را با استفاده از تبدیل موجک معکوس بازسازی می‌کنیم.

روش‌های مختلفی برای تعیین آستانه وجود دارد که از آن جمله می‌توان به آستانه جهانی، اعتبار متقابل و قطعی اشاره نمود.

سوال کلیدی در موجک انقباضی این است که چگونه ضرایب موجک d ، باید به فرم $\hat{d}$ اصلاح و تعدیل شوند؟ روش موجک انقباضی را برای مشکلات کلی برآورد منحنی در [۴] ارائه شده است. دلایل منطقی وجود دارد که چرا موجک انقباضی برای برآورد توابع مفید است. دلایل اصلی که موجک انقباضی برآوردکنندهای خوبی هستند این است که، برای یک برد گسترده از توابع زیان و برای کلاس توابع کلی خواص تقریباً بهینه را دارند. در [۳] پیشنهاد داده شد که آستانه‌ی جهانی می‌تواند در فرآیند انقباض ویسو^{۱۸} ثبت شود. دیگر آستانه‌ی انتخاب‌کننده بر اساس برآوردگر نااریب ریسک استاین^{۱۹} ($SURE$) که توسط [۴] پیشنهاد شد و آن را انقباض قطعی^{۲۰} نامیدند. آستانه انقباض قطعی انتخاب‌کننده مقدار آستانه‌ی t_j است که هر وضوح^{۲۱} سطح j در تبدیل موجک را تعیین می‌کند. موجک بیزی، همواره برای موجک انقباضی رایج است.

ابتدا برای ضرایب صحیح موجک $d_{j,k}$ یک توزیع پیشین تعیین می‌کنیم، این توزیع پیشین مخصوص در معرفی موجک به طور ذاتی تنک و پراکنده در نظر گرفته شده است. سپس با استفاده از قضیه بیز، توزیع پسین ضرایب موجک را با استفاده از توزیع ضرایب موجک نوفه‌ای که معمولاً معلوم فرض شده‌اند محاسبه می‌کنیم. به طور کلی، می‌توان یک توزیع پسین برای تابع ناشناخته را با به کار بردن تبدیل معکوس موجک گسسته، روی توزیع پسین ضرایب موجک برآورد کرد.

۳ مدل تطبیق بیزی موجک انقباضی

در این بخش، به بررسی مدل تطبیق بیزی موجک انقباضی^{۲۲} پرداخته شده است. فرض کنید داده‌ها به فرم زیر باشد:

$$y_i = f\left(\frac{i}{n}\right) + \sigma z_i \quad i = 0, 1, \dots, n-1 \quad (5)$$

که z_i مشاهدات مستقل و هم‌توزیع از نرمال استاندارد و σ معلوم فرض شده و n مضربی از دو و صحیح می‌باشد. موجک‌ها می‌توانند از پایه‌های متعامد برای فضای توابع گوناگون استفاده کنند. بنابراین فرض می‌کنیم که $\{\psi_{j,k}\}_{j,k \in \mathbb{Z}}$ یک موجک یک‌متعامد بر اساس فضای $L^2(\mathbb{R})$ باشد که:

$$\psi_{j,k}(x) = 2^{j/2} \psi(2^j x - k) \quad (6)$$

یک اتساع در مقیاس j و یک انتقال به وسیله $k/2^j$ از تابع موجک مادر ψ دارد. ضریب متناظر موجک را می‌توان، برای $f \in L^2(\mathbb{R})$ مانند $\theta_{j,k} = \langle f, \psi_{j,k} \rangle$ نوشت. در مفهوم داده‌های گسسته یک مورد مشابه وجود دارد. به عنوان مثال، با استفاده از موجک‌های دایچیز که تکیه‌گاه فشرده دارند ([۲]) با یک تبدیل گسسته موجک بردار مشاهدات y به طول n را می‌گیرند و به بردار w با طول مشابه می‌برند. بردار w حاوی ضرایب تجربی موجک است که این فرآیند را می‌توان به صورت یک ضرب در ماتریس متعامد W نشان داد:

$$w \equiv Wy = Wf + \sigma Wz \equiv \theta + \sigma z^* \quad (7)$$

که $\theta = Wf$ برداری به طول n ، با ضرایب گسسته f است. z^* بردار مشاهدات به طول n که مستقل و هم‌توزیع از نرمال استاندارد می‌باشد.

^{۱۸} Visu shrink

^{۱۹} Stein unbiased risk estimator

^{۲۰} Sure shrink

^{۲۱} Resolution

^{۲۲} Adaptive Bayesian Wavelet Shrinkage (ABWS)

دنباله‌ی $d_i = Wy_i$ از مجموع تبدیل موج $\theta_i = Wf_i$ و تبدیل نوفه $\eta_i = W\epsilon_i$ نتیجه می‌شود. این یک نتیجه منطقی از خطی بودن تبدیل W ، که به صورت زیر می‌باشد.

$$d_i = \theta_i + \eta_i, \quad i = 1, \dots, N = 2^n \quad (8)$$

اگر دنباله‌ی نویزی ϵ_i توسط توزیع‌های نرمال با میانگین صفر و واریانس σ^2 مدل‌بندی شده باشند، از تبدیل متعامد W نتیجه می‌شود که دنباله‌ی η_i هم، خواص توزیعی یکسان همانند دنباله‌ی ϵ_i را دارد. یعنی اگر:

$$\epsilon_i \sim N(0, \sigma^2) \Rightarrow \eta_i \sim N(0, \sigma^2) \quad (9)$$

بنابراین به جای برآورد مستقیم موج، تبدیل موجک آن را برآورد می‌کنیم. به طور معادل، θ_i را در $d_i \sim N(\theta_i, \sigma^2)$ برآورد می‌کنیم. هدف نمایش یک مدل پارامتری بیزی است:

$$d|\theta, \sigma^2 \sim N(\theta, \sigma^2) \quad (10)$$

که σ^2 نامعلوم است و اولین مشکل انتخاب پیشین روی σ^2 می‌باشد. به دلایل کاربردی و خواص آگاهی بخش توزیع پیشین روی σ^2 را توزیع نمایی انتخاب می‌کنند. طوری که توسط [۶] نشان داده شد، توزیع نمایی بیشینه‌ی بی‌نظمی را در میان تمام توزیع‌ها بر روی تکیه‌گاه $(0, \infty)$ با یک گشتاور اولیه ثابت دارد. بنابراین

$$\sigma^2 \sim E(\mu), \quad (f(\sigma^2|\mu) = \mu e^{-\mu\sigma^2}) \quad (11)$$

و توزیع حاشیه‌ای، نمایی دوگانه^{۲۳} نیز می‌باشد که در واقع توزیع نمایی دوگانه مقیاس آمیخته‌ای از توزیع‌های نرمال هستند:

$$d|\theta \sim DE\left(\theta, \frac{1}{\sqrt{2\mu}}\right), \quad f(d|\theta) = \frac{1}{2} \sqrt{2\mu} e^{-\sqrt{2\mu}|d-\theta|} \quad (12)$$

در حالت کلی، قواعد بیزی منقبض کننده هستند و قالب آن‌ها در موارد بسیاری یک خاصیت مطلوب برای موجک انقباضی دارند.

فرض کنید $d|\theta \sim DE\left(\theta, \frac{1}{\sqrt{2\mu}}\right)$ و فرض کنید $\pi(\theta)$ یک توزیع پیشین متقارن روی θ باشد:

$$\pi(\theta) = \pi(-\theta), \quad \theta \in \mathbb{R}$$

آن‌گاه قاعده بیزی نسبت به زیان مربع خطا به صورت زیر است:

$$\delta(d) = d - \frac{\Pi_1'(c) - \Pi_2'(c)}{\Pi_1'(c) + \Pi_2'(c)} \quad (13)$$

که Π_1 و Π_2 تبدیلات لاپلاس از توابع $\pi(\theta + d)$ و $\pi(\theta - d)$ و $\theta \in (0, \infty)$ و $c = \sqrt{2\mu}$. اثبات. قاعده بیزی نسبت به زیان مربع خطا به صورت زیر است:

$$\delta(d) = \frac{\int_{\mathbb{R}} \theta f(d|\theta) \pi(\theta) d\theta}{\int_{\mathbb{R}} f(d|\theta) \pi(\theta) d\theta}$$

^{۲۳}Double exponential

بنابراین در مخرج کسر داریم :

$$\begin{aligned}
 \int_{\mathbb{R}} f(d|\theta)\pi(\theta)d\theta &= \int_{-\infty}^{\infty} e^{-c|\theta-d|}\pi(\theta)d\theta \\
 &= \int_{-\infty}^{\circ} e^{cu}\pi(u+d)du + \int_{\circ}^{\infty} e^{-cu}\pi(u+d)du \\
 &= \int_{\circ}^{\infty} e^{-cu} [\pi(u+d) + \pi(u-d)] du \\
 &= \Pi_{\downarrow}(c) + \Pi_{\uparrow}(c) = A
 \end{aligned}$$

و نیز در صورت کسر داریم:

$$\begin{aligned}
 \int_{\mathbb{R}} \theta f(d|\theta)\pi(\theta)d\theta &= \int_{\circ}^{\infty} (d-t)e^{-ct}\pi(t-d)dt \\
 &\quad + \int_{\circ}^{\infty} (d+t)e^{-ct}\pi(t+d)dt \\
 &= d(\Pi_{\downarrow}(c) + \Pi_{\uparrow}(c)) - (\Pi'_{\downarrow}(c) - \Pi'_{\uparrow}(c)) = B
 \end{aligned}$$

□

طبق روابط A و B اثبات کامل می‌شود.

توابع انقباضی با اعمال یک ساختار پیشین خاص درون مدل در فضای ضرایب موجک به دست آمده‌اند. در این ساختار پیشین، ضرایب نسبت به هم استقلال دو جانبه دارند و هر ضریب آمیخته یا ترکیبی از دو توزیع نرمال است، بدین صورت:

$$\theta_{j,k}|\gamma_{j,k} \sim \gamma_{j,k}N(\circ, c_j^{\downarrow}\tau_j^{\downarrow}) + (1 - \gamma_{j,k})N(\circ, \tau_j^{\downarrow}) \quad (14)$$

پارامتر آمیخته $\gamma_{j,k}$ توزیع پیشین خود را دارد که با این فرمول به دست می‌آید:

$$P(\gamma_{j,k} = 1) = 1 - P(\gamma_{j,k} = \circ) \equiv p_j \quad (15)$$

زمانی که $\gamma_{j,k} = \circ$ است $\theta_{j,k} \sim N(\circ, \tau_j^{\downarrow})$ و زمانی که $\gamma_{j,k} = 1$ است $\theta_{j,k} \sim N(\circ, c_j^{\downarrow}\tau_j^{\downarrow})$ می‌باشد. برای به دست آوردن معادله (۱۴)، به عنوان مثال برای $\theta_{j,k}|\gamma_{j,k}$ از نرمال چندمتغیره پیشین استفاده می‌کنیم

$$\theta|\gamma \sim N_p(\circ, D_{\gamma}RD_{\gamma})$$

که $\gamma = (\gamma_1, \dots, \gamma_p)$ و R یک ماتریس همبسته پیشین است و

$$D_{\gamma} \equiv \text{diag}[a_1\tau_1, a_2\tau_2, \dots, a_p\tau_p] \quad (16)$$

یک ماتریس قطری است با $a_j = c_j$ اگر $\gamma_j = 1$ و $a_j = 1$ اگر $\gamma_j = \circ$ باشد.

p_j, c_j, τ_j پارامترهای پیشینی هستند که باید انتخاب شوند. انتخاب τ_j در (و) باید طوری باشد که اگر $\theta_{j,k} \sim N(\circ, \tau_j^{\downarrow})$ باشد، سپس $\theta_{j,k}$ بتواند با صفر به طور مطمئن جایگزین شود، زیرا با احتمال بالا $|\theta_{j,k}| \leq 3\tau_j$ است. همچنین انتخاب $c_j (> 1)$ در (و) باید طوری باشد که اگر $\theta_{j,k} \sim N(\circ, c_j^{\downarrow}\tau_j^{\downarrow})$ ، سپس یک برآورد غیرصفر از $\theta_{j,k}$ باید در مدل نهایی قرار بگیرد. از یک دیدگاه معقولانه بیزی، اگر بخواهند یک c_j نسبتاً بزرگ انتخاب کنند که تکیه‌گاه آن روی مقادیر $\theta_{j,k}$ باشد، باید پارامتری غیر از صفر انتخاب کنند.

از همین پارامترهای پیشین برای تمامی ضرایب در سطحی مفروض از j استفاده می‌کنیم. مولفه‌ی $N(\circ, \tau_j^2)$ به ما این امکان را می‌دهد که مقداری از جرم را در نزدیکی صفر متمرکز کنیم، این در حالی است که مولفه‌ی $N(\circ, c_j^2 \tau_j^2)$ بقیه جرم را در اطراف مقادیر بزرگ‌تر، پراکنده می‌کند.

با شرط داشتن مقادیر $\theta_{j,k}$ و σ^2 ضرایب تجربی موجک دارای توزیع نرمال هستند:

$$\omega_{j,k} | \theta_{j,k}, \sigma^2 \sim \sigma N(\theta_{j,k}, \sigma^2) \quad (17)$$

زمانی که داده‌ها مشاهده شدند، ضرایب تجربی موجک محاسبه می‌شوند. به دنبال توزیع پسین در $\theta_{j,k}$ مشاهده نشده هستیم و بر مقدار پیش‌بینی شده و واریانس این توزیع تمرکز می‌کنیم. با داشتن مقدار پیش‌بینی شده، f را می‌توان از طریق $\hat{f} = W^T E(\theta | \omega)$ برآورد کرد.

فرض کنید $\theta_{j,k}$ دارای توزیع پیشین به فرم زیر باشد. رابطه زیر وابسته به ثابت‌های τ_j ، c_j و p_j است. برای این که بتوان از توزیع پیشین آمیخته زیر استفاده کنیم باید مقادیر این ثابت‌ها انتخاب شوند.

$$\theta_{j,k} | \gamma_{j,k} \sim \gamma_{j,k} N(\circ, c_j^2 \tau_j^2) + (1 - \gamma_{j,k}) N(\circ, \tau_j^2) \quad (18)$$

در این بخش، به تشریح نحوه‌ی انتخاب این مقادیر و ارتباط آن‌ها با شیوه‌ی انقباض ضرایب تجربی موجک، می‌پردازیم. با داشتن این انتخاب‌های قراردادی، نرمال آمیخته پیشین، رویکردی ساده و اتوماتیک برای موجک انقباضی فراهم می‌کند. ابتدا نقش عمومی τ_j ، c_j و p_j را در تعیین ضرایب موجک انقباضی بیان می‌کنیم. سپس انتخاب‌های قراردادی را ارائه می‌دهیم.

روند برآورد موجک انقباضی زمانی خوب عمل می‌کند که مجموعه ضرایب صحیح $\theta_{j,k}$ پراکنده و تنک (نسبت نوفه به موج پایین باشد) باشند که فقط چند مورد ضریب بزرگ وجود دارد، و بقیه همه کوچک‌اند. مولفه نرمال $N(\circ, \tau_j^2)$ آمیخته، برای توصیف یک ضریب کوچک، در نظر گرفته شده است. τ_j را انتخاب می‌کنیم تا اگر ضریب در فاصله‌ی $(-\tau_j, \tau_j)$ قرار داشته باشد، آن قدر کوچک است که برای اهداف عملی به سمت صفر متمایل شود. مولفه‌ی $N(\circ, c_j^2 \tau_j^2)$ نیز برای توصیف ضریب بزرگ در نظر گرفته شده است. این مولفه نرمال باید به اندازه کافی پهن و گسترده باشد تا دامنه کامل ضرایب قابل استفاده را دربرگیرد. پارامتر p_j را می‌توان به عنوان احتمال غیر قابل اغماض یک ضریب تعبیر کرد. انتظار می‌رود که درصد ضرایب در سطح j تفاوت معناداری با صفر داشته باشند و مقادیر کوچک p_j پراکندگی ضرایب را نشان می‌دهند.

توجه داشته باشید که مقادیر مختلفی از پارامترها را در سطوح متفاوتی انتخاب می‌کنیم. برای τ_j و c_j این روند طبیعی و نرمال است، چون کوچک یا بزرگ بودن یک ضریب در یک بخشی وابسته به ارتفاع و عرض موجک مربوطه است، که این‌ها همه به سطح وضوح بستگی دارد. یعنی هرچه p_j کوچک‌تر شود سطح وضوح یا قدرت تشخیص بالاتر می‌رود. برای درک این که چگونه انتخاب ابرپارامترها، به انقباض ضرایب تجربی موجک ارتباط پیدا می‌کند این رابطه را می‌توانیم بنویسیم:

$$\begin{aligned} E(\theta_{j,k} | \omega_{j,k}) &= E_{\gamma_{j,k} | \omega_{j,k}} E(\theta_{j,k} | \omega_{j,k}, \gamma_{j,k}) \\ &= P(\gamma_{j,k} = 1 | \omega_{j,k}) E(\theta_{j,k} | \omega_{j,k}, \gamma_{j,k} = 1) \\ &\quad + P(\gamma_{j,k} = 0 | \omega_{j,k}) E(\theta_{j,k} | \omega_{j,k}, \gamma_{j,k} = 0) \\ &= P(\gamma_{j,k} = 1 | \omega_{j,k}) \frac{(c_j \tau_j)^2}{\sigma^2 + (c_j \tau_j)^2} \omega_{j,k} \\ &\quad + P(\gamma_{j,k} = 0 | \omega_{j,k}) \frac{\tau_j^2}{\sigma^2 + \tau_j^2} \omega_{j,k} \end{aligned} \quad (19)$$

تابع انقباض را می‌توان به عنوان یک تناسب هموار بین دو خط شیب‌دار، $\frac{\tau_j^2}{\sigma^2 + \tau_j^2}$ و $\frac{(c_j \tau_j)^2}{\sigma^2 + (c_j \tau_j)^2}$ تعبیر کرد. زمانی که $\omega_{j,k}$ کوچک است، یعنی $\theta_{j,k}$ کوچک است پس $P(\gamma_{j,k} = \circ | \omega_{j,k})$ بزرگ می‌شود. در این مورد، تابع انقباض تقریباً از یک خط راست با تقاطع صفر و شیب $\frac{\tau_j^2}{\sigma^2 + \tau_j^2}$ تبعیت می‌کند. زمانی که τ_j کوچک باشد شیب همین است. بنابراین مقادیر نسبتاً کوچک τ_j ، بخش مسطح و همواری از تابع انقباض را پیرامون صفر به ما می‌دهد. زمانی که $\omega_{j,k}$ بزرگ است، تابع انقباض تقریباً از یک خط راست با تقاطع صفر و شیب $\frac{(c_j \tau_j)^2}{\sigma^2 + (c_j \tau_j)^2}$ تبعیت می‌کند. برای c_j بزرگ، این شیب به یک نزدیک خواهد شد. بنابراین تابع انقباض به صورت یک میانگین وزن‌دار از یک تابع خطی با شیب کوچک و یک تابع خطی با یک شیب نزدیک به یک به دست می‌آید. با افزایش $\omega_{j,k}$ ، وزن روی خطی که شیب بزرگتری دارد، افزایش یافته و به یک نزدیک خواهد شد. پارامترهای τ_j و c_j شیب‌های دو خط را محاسبه و تعیین می‌کنند. مقادیر کوچک p_j عرض فاصله تا صفر را افزایش می‌دهد. در اینجا تابع انقباض به خطی متصل می‌شود که شیب کمتری دارد. با افزایش c_j فاصله ای که در طی آن تابع انقباض از خط کم شیب‌تر به خط پرشیب‌تر صعود می‌کند، کوچکتر می‌شود زیرا با افزایش c_j دو مولفه فرعی آمیخته، تمایز و اختلاف بیشتری با هم پیدا می‌کنند.

۴ توزیع ضرایب موجکی

در این بخش، جزئیات محاسبه‌ی میانگین‌ها و واریانس‌های پسین را ارائه می‌دهیم. با فرض این‌که σ معلوم باشد، فرمول‌هایی برای $E(\theta_{j,k} | \omega_{j,k})$ و $\text{Var}(\theta_{j,k} | \omega_{j,k})$ به فرم بسته‌ی ساده وجود دارند. ماتریس کوواریانس پسین بردار $\theta_{j,k}$ ، یک ماتریس قطری است که آن را با D نمایش می‌دهیم. میانگین پسین f از تبدیل خطی موجک به صورت زیر به دست می‌آید:

$$E(f|y) = W^T E(\theta|\omega)$$

ماتریس کوواریانس f ، به صورت $W^T D W$ می‌باشد. برای محاسبه این ماتریس، در ادامه یک راهبرد کارآمد شرح می‌دهیم. برای این‌که ارائه و تشریح فرمول‌ها را ساده‌تر کنیم، زیرنویس‌های k و j را حذف می‌کنیم. ابتدا با پسین حاشیه‌ای γ ، پسین مربوط به θ را به دست می‌آوریم و سپس به محاسبه‌ی $\theta | \gamma$ می‌پردازیم. پسین حاشیه‌ای مربوط به γ را می‌توان به صورت زیر نشان داد:

$$P(\gamma = 1 | \omega) = \frac{A}{A+1}, \quad A = \frac{p\pi(\omega | \gamma = 1)}{(1-p)\pi(\omega | \gamma = \circ)} \quad (20)$$

که در رابطه (۲۰)، $\pi(\omega | \gamma = 1) \sim N(\circ, \sigma^2 + c^2 \tau^2)$ و $\pi(\omega | \gamma = \circ) \sim N(\circ, \sigma^2 + \tau^2)$ هستند. توزیع پسین برای θ را می‌توان به صورت زیر بیان کرد:

$$F(\theta | \omega) = F(\theta | \omega, \gamma = 1) \cdot \frac{A}{A+1} + F(\theta | \omega, \gamma = \circ) \cdot \frac{1}{1+A}, \quad (21)$$

که در آن $F(\theta | \omega, \gamma = 1)$ یک تابع توزیع برای

$$\theta | \omega, \gamma = 1 \sim N\left(\frac{(c\tau)^2}{\sigma^2 + (c\tau)^2} \omega, \frac{\sigma^2 (c\tau)^2}{\sigma^2 + (c\tau)^2}\right) \quad (22)$$

و $F(\theta | \omega, \gamma = \circ)$ یک تابع توزیع برای

$$\theta | \omega, \gamma = \circ \sim N\left(\frac{\tau^2}{\sigma^2 + \tau^2} \omega, \frac{\sigma^2 \tau^2}{\sigma^2 + \tau^2}\right) \quad (23)$$

می‌باشند، حال میانگین پسین θ به صورت زیر به دست می‌آید:

$$E[\theta | \omega] = \left[\frac{(c\tau)^2}{\sigma^2 + (c\tau)^2} \cdot \frac{A}{A+1} + \frac{\tau^2}{\sigma^2 + \tau^2} \cdot \frac{1}{A+1} \right] \cdot \omega \quad (24)$$

که حاصل آن مضربی از ω ، با عامل انقباضی مشترک زیر است:

$$s(\omega) = \left[\frac{(c\tau)^2}{\sigma^2 + (c\tau)^2} \cdot \frac{A}{A+1} + \frac{\tau^2}{\sigma^2 + \tau^2} \cdot \frac{1}{A+1} \right] \quad (25)$$

در این فرمول، $s(\omega) \leq 1$ است که خود تابعی از ω است. برآورد θ از طریق میانگین پسین $E(\theta|\omega)$ معادل با، استفاده از یک تابع انقباض غیرخطی است. میانگین پسین مشابه ضرایب انقباضی است. واریانس پسین را با استفاده از فرمول زیر محاسبه می‌کنیم:

$$\begin{aligned} \text{Var}(\theta|\omega) &= E_\gamma [\text{Var}(\theta|\omega, \gamma)] + \text{Var}_\gamma [E(\theta|\omega, \gamma)] \\ &= E_\gamma [\text{Var}(\theta|\omega, \gamma)] + E_\gamma [E(\theta|\omega, \gamma)^2] \\ &\quad - (E_\gamma [E(\theta|\omega, \gamma)])^2 \end{aligned} \quad (26)$$

نتایج به دست آمده از این فرمول، عبارات زیر را حاصل می‌دهند:

$$E_\gamma [\text{Var}(\theta|\omega, \gamma)] = \left[\frac{\sigma^2(c\tau)^2}{\sigma^2 + \sigma^2(c\tau)^2} \cdot \frac{A}{A+1} + \frac{\sigma^2\tau^2}{\sigma^2 + \tau^2} \cdot \frac{1}{A+1} \right] \quad (27)$$

$$(E_\gamma [E(\theta|\omega, \gamma)])^2 = \left[\left(\frac{(c\tau)^2}{\sigma^2 + (c\tau)^2} \omega \right)^2 \cdot \frac{A}{A+1} + \left(\frac{\tau^2}{\sigma^2 + \tau^2} \omega \right)^2 \cdot \frac{1}{A+1} \right]$$

و البته $(E_\gamma [E(\theta|\omega, \gamma)])^2 = (E[\theta|\omega])^2$ می‌باشد. برای مقادیر کوچک ω واریانس پسین به طور تقریبی برابر با $\text{Var}(\theta|\omega) \approx \text{Var}(\theta)$ است و برای مقادیر بزرگ ω برابر با σ^2 می‌باشد. در بین این دو مقدار، اکسترم جایی است که در آن واریانس بزرگترین مقدار را می‌گیرد. این فاصله دقیقاً همان دامنه‌ای است که کمترین اطمینان را داریم که آیا ضریب اصلی شامل موج می‌شود یا این که موج را در بر نمی‌گیرد. ساختار جدید f یعنی $\hat{f}$ ، حاصل معکوس بردار میانگین پسین $E(\theta|\omega)$ است، که به صورت $\hat{f} = W^T E(\theta|\omega)$ می‌باشد.

ماتریس کوواریانس برای این بردار بازسازی شده، دقیقاً به صورت $\Delta \equiv W^T D W$ است، که در آن $D \equiv \text{diag} \{\text{Var}(\theta|\omega)\}$ می‌باشد. با گرفتن تبدیل معکوس موجک برای ستون‌های ماتریس D و همچنین تکرار عمل ترانزاده ماتریس قطری Δ بصورت زیر به دست می‌آید.

$$\Delta \equiv W^T D W = W^T (W^T D)^T. \quad (28)$$

مراجع

- [1] Afshari, M., Lak F. and Gholizadeh, B. (2017), A new Bayesian wavelet thresholding estimator of nonparametric regression, *Journal of Applied Statistics*, **44**(4), 649-666.
- [2] Daubechies, I. (1992), *Ten Lectures on Wavelets*, Philadelphia, Pennsylvania: SIAM.
- [3] Donoha, D.L. and Johnstone, I.M. (1994), *Minimax estimation by wavelet shrinkage*, Technical Report, Department of Statistics, Stanford University.

- [4] Donoho, D.L., Johnstone, I.M., Kerkyacharian, G. and Picard, D. (1995), Wavelet shrinkage: asymptopia?, *Journal of the Royal Statistical Society, Series B (Methodological)*, 301-369.
- [5] Nason, G.P. and Silverman, B.W. (1994), The discrete wavelet transform in S, *Journal of Computational and Graphical Statistics*, **3**(2), 163-191.
- [6] Zellner, A. (1996), *Bayesian method of moments (BMOM) analysis of mean and regression models*, In *Modelling and Prediction Honoring Seymour Geisser* (pp. 61-72), Springer New York.

معرفی توزیع جدید برگرفته از توزیع گرام چارلی و کاربرد آن در محاسبه معیارهای ریسک

علی اصغر لطیفی زاده^۱، حجت الله ذاکرزاده، حمزه ترابی

گروه آمار، دانشکده علوم ریاضی، دانشگاه یزد

چکیده: در این مقاله توزیع گرام چارلی سکانت هایپربولیک را معرفی می‌کنیم و از آن برای اندازه‌گیری معیارهای ریسک استفاده می‌کنیم. فرم واضحی برای محاسبه کسری مورد انتظار این توزیع به کمک تابع فوق هندسی تعمیم یافته ارائه می‌دهیم. همچنین به منظور بررسی کاربرد این توزیع یک مجموعه داده واقعی که مربوط به شرکت ادوبی است را بررسی می‌کنیم. **واژه‌های کلیدی:** تابع فوق هندسی تعمیم یافته، توزیع گرام چارلی هایپربولیک سکانت، توزیع های دم پهن، کسری مورد انتظار، مقدار در معرض خطر.

کد موضوع بندی: 60E05، 91B30، 91B82، 62P20.

۱ مقدمه

شواهد زیادی موجود است که نشان می‌دهد، داده‌های مالی و بازده دارایی، کشیده^{۲۴} و یا چوله^{۲۵} هستند که این امر در مباحث اقتصادی و علوم مالی به رسمیت شناخته شده است [۱]. کشیده بودن، چوله و دم پهن بودن سه ویژگی ذاتی بسیاری از داده‌های مالی است. این سه ویژگی در مدل سازی نوسانات و زیان سبدهای سرمایه‌گذاری بسیار تاثیرگذار هستند، زیرا اندازه‌گیری زیان به شدت تحت تاثیر ضخامت دم‌های توزیع و چولگی آن است [۲].

برای سال‌های متعددی مدل‌های گوسی، در علم اقتصاد و امور مالی و بخصوص مدل سازی داده‌های بازده دارایی، مورد استفاده قرار گرفته شده است [۳]. اما بحران‌های اقتصادی اخیر باعث برهم زدن فرض نرمال بودن بازده دارایی شده است. فرض نرمال بودن داده‌های بازده دارایی، مدل سازی این داده‌ها را ساده تر می‌کند اما زمانی که توزیع داده‌های بازده دارایی رفتاری دم‌پهن نشان‌دهند، این مدل کارایی خود از دست می‌دهد [۴]. در سال‌های اخیر برای مدل‌سازی هرچه بهتر داده‌های مالی محققان روش‌های متفاوتی ارائه داده‌اند. به عنوان مثال پژوهشگران می‌توانند به جای مدل‌سازی داده‌ها با استفاده از توزیع نرمال، از توزیع‌های دیگری مانند توزیع لوی، توزیع تی-استیودنت یا توزیع‌های دم‌پهن دیگر، که در مدل سازی داده‌های مالی و اندازه‌گیری زیان، بسیار پرکاربرد هستند، استفاده کنند. روش‌های دیگر نیز برای تبدیل توزیع نرمال برای انطباق با ویژگی‌های مورد نظر (چوله بودن، کشیده و دم‌پهن بودن) ارائه شده است. به عنوان مثال زویا با تغییر شکل توزیع نرمال توسط چند جمله‌ای‌های متعامد هرمیت و استفاده از توزیع گوسی به عنوان تابع وزن به بررسی داده‌های دم‌پهن و نامتقارن پرداخت [۵]. اخیراً زویا و همکاران رویکرد جدیدی برای مدل‌سازی توزیع سبدهای مالی متشکل از بازده‌های دارایی و زیان‌های بیمه با به کارگیری بسط شبه گرام چارلی^{۲۶} ارائه داده است [۶].

^۱ علی اصغر لطیفی زاده: Ali_a.zadeh@hotmail.com

^{۲۴} Leptokurtic

^{۲۵} Skew

^{۲۶} Gram-Charlier-like Expansions

در این مقاله با استفاده از چند جمله‌ای‌های متعامد و بسط شبه گرام چارلی برای توزیع سکانت هایپربولیک، به بررسی معیار های ریسک مانند ارزش در معرض خطر ^{27}VaR و کسری مورد انتظار ^{28}ES برای یک سبد سرمایه می‌پردازیم. در بخش دوم نگاهی کوتاه به معیارهای ریسک خواهد شد و در بخش سوم توزیع گرام چارلی سکانت هایپربولیک ($GCHS$) معرفی می‌شود و یک فرم صریح برای محاسبه کسری مورد انتظار زمانی که توزیع زیان گرام چارلی سکانت هایپربولیک است، ارائه می‌شود. در بخش چهارم به بررسی داده‌های مالی واقعی می‌پردازیم و در بخش پنجم به نتیجه‌گیری از این مقاله خواهیم پرداخت.

۲ نگاهی به معیار های ریسک

عبارت "ارزش در معرض خطر" که به اختصار VaR می‌گویند، تا اوایل دهه ۱۹۹۰ وارد ادبیات مالی نشده بود. نقطه آغازین توجه به این روش به سال ۱۹۹۲ بازمی‌گردد که در آن سال بورس اوراق بهادار نیویورک برای اولین بار به طور غیر رسمی سرمایه شرکت‌های عضو را موضوع آزمون قرار داد [۷]. ارزش در معرض خطر بیانگر حداکثر زیان مورد انتظار که یک پرتفوی یا یک دارایی در طول دوره نگهداری معین (دوره‌ای که در طول آن ترکیب پرتفوی بدون تغییر باقی می‌ماند) و با احتمال مشخص، متحمل می‌شود. به بیانی ساده‌تر این معیار مبلغی از ارزش پرتفوی که انتظار می‌رود در یک دوره معین (افق زمانی مشخص) با احتمال ثابت از بین برود، مشخص می‌کند. فرض کنید $F_X(x)$ تابع توزیع متغیر تصادفی X باشد که این متغیر مقدار زیان است، بر این اساس ارزش در معرض خطر به صورت

$$VaR_X(q) = \inf \{x : F_X(x) \geq q\} = \nu_q = F_X^{-1}(q), \quad (1)$$

تعریف می‌شود که در آن $q \in (0, 1)$.

هرچند این معیار بسیار پر کاربرد است اما محدودیت هایی نیز دارد به عنوان مثال، ارزش در معرض خطر تنها به مبلغ ریالی ریسک پرتفوی توجه دارد و تبادل ریسک و بازده مورد انتظار را نادیده می‌گیرد. این امر باعث حذف بخش قابل توجهی از اطلاعات پرتفوی می‌شود. از طرفی دیگر، روش منحصر به فردی برای محاسبه آن وجود ندارد که این امر موجب می‌شود تا تحلیل نتایج حاصله دستخوش قضاوت های ذهنی تحلیل گران قرار گیرد [۸]. از طرفی دیگر ارزش در معرض خطر یک نوع تحلیل آماری است که عمدتاً در تعیین کمی ریسک بازار کاربرد دارد و از ارزیابی ریسک مربوط به رویدادهای کیفی ناتوان است [۹]. یکی دیگر از معیار های جالب در بررسی ریسک امید شرطی دمی (TCE)^{۲۹} است که به صورت

$$TCE_X(q) = E[X|X \geq \nu_q], \quad (2)$$

تعریف می‌شود که E نشان دهنده امید ریاضی است.

به دلیل اینکه امید شرطی دمی، خاصیت زیر جمع پذیری ندارد، به طور کلی نمی‌تواند به عنوان یک معیار منسجم اندازه گیری ریسک به شمار آید. این موضوع در توزیع‌های ناپیوسته مشهود است زیرا این معیار به شدت تحت تاثیر تغییرات کوچک در سطح اطمینان است. برای متغیر تصادفی حقیقی مقدار و میانگین-متناهی با توزیع پیوسته مطلق و تابع توزیع اکیدا صعودی، امید شرطی

^{۲۷} Value at Risk

^{۲۸} Expected Shortfall

^{۲۹} Tail Conditional Expectation

دمی و کسری مورد انتظار با هم برابرند [۱۰]. فرض کنید $f(\cdot)$ تابع چگالی احتمال متغیر تصادفی X باشد، کسری مورد انتظار را می توان به صورت

$$ES_X(q) = \frac{\int_{\nu_q}^{\infty} x f(x) dx}{\int_{\nu_q}^{\infty} f(x) dx}, \quad (3)$$

نمایش داد.

۳ توزیع گرام چارلی سکانت هایربولیک

در این فصل متغیر تصادفی گرام چارلی سکانت هایربولیک معرفی و با استفاده از تابع فوق هندسی تعمیم یافته^{۲۰}، برای معیار کسری مورد انتظار این متغیر فرم واضحی ارایه شده است. متغیر تصادفی X را دارای توزیع گرام چارلی سکانت هایربولیک با پارامتر β گوئیم و آن را با نماد $X \sim GCHS(\beta)$ نشان می دهیم، هرگاه دارای تابع چگالی به صورت

$$f_X(x; \beta) = \left(1 + \frac{\beta}{4}(x^2 - 1)\right)^{-\frac{1}{4}} \operatorname{sech}\left(\frac{\pi x}{4}\right) \quad x \in (-\infty, \infty), \quad 0 \leq \beta \leq 4. \quad (4)$$

باشد. این توزیع با استفاده از بسط گرام چارلی و استفاده از توزیع سکانت هایربولیک به عنوان تابع وزن به دست آمده است. شکل ۱ نمودار تابع چگالی (۴) را به ازای ۰، ۱، ۳، β نشان می دهد. نتایج زیر یک فرم صریح برای محاسبه کسری مورد انتظار توزیع گرام چارلی سکانت هایربولیک ارایه می دهد. فرض کنید متغیر تصادفی $X \sim GCHS(\beta)$ آنگاه

$$ES_X(q) = -\frac{\left(1 - \frac{\beta}{4}\right)G_{\frac{\pi}{4}}^1(\nu_q) + \frac{\beta}{4}G_{\frac{\pi}{4}}^3(\nu_q)}{2 - \left(\left(1 - \frac{\beta}{4}\right)G_{\frac{\pi}{4}}^0(\nu_q) + \frac{\beta}{4}G_{\frac{\pi}{4}}^2(\nu_q)\right)}, \quad (5)$$

که $G_a^n(z)$ برابر است با $\int z^n \operatorname{sech}(az) dz$. اثبات. یک تقریب برای $G_a^n(z)$ براساس تابع فوق هندسی تعمیم یافته به صورت

$$\int z^n \operatorname{sech}(az) dz = n! e^{az} \sum_{j=0}^n \frac{(-1)^j z^{n-j}}{(n-j)! a^{j+1}} {}_jF_{j+1}\left(\frac{1}{4}, \dots, \frac{1}{4}; 1; \frac{3}{4}, \dots, \frac{3}{4}; -e^{2az}\right), \quad n \in N^+ \quad (6)$$

است که در رابطه (۶)، $F(\cdot)$ تابع فوق هندسی تعمیم یافته و N^+ مجموعه اعداد حسابی است [۱۱]. سمت راست رابطه (۳) را کسری مورد انتظار در نظر می گیریم و صورت را با نماد A و مخرج را با نماد B نشان می دهیم. با جایگذاری تابع چگالی $f_X(x)$ ، که در (۴) مشخص شده است در A داریم

$$A = \int_{\nu_q}^{\infty} x f_X(x) dx = \frac{1}{4} \left(\left(1 - \frac{\beta}{4}\right) G_{\frac{\pi}{4}}^1(x) + \frac{\beta}{4} G_{\frac{\pi}{4}}^3(x) \right) \Big|_{x=\nu_q}^{\infty} \quad (7)$$

به دلیل اینکه

$$\lim_{x \rightarrow \infty} G_{\frac{\pi}{4}}^1(x) = 0, \quad \lim_{x \rightarrow \infty} G_{\frac{\pi}{4}}^3(x) = 0.$$

پس A را می توان به صورت

$$A = -\frac{1}{\gamma} \left(\left(1 - \frac{\beta}{\gamma}\right) G_{\frac{\gamma}{\gamma}}^1(\nu_q) + \frac{\beta}{\gamma} G_{\frac{\gamma}{\gamma}}^2(\nu_q) \right) \quad (8)$$

نوشت. به طور مشابه برای B داریم

$$B = \int_{\nu_q}^{\infty} f_X(x) dx = \frac{1}{\gamma} \left(\left(1 - \frac{\beta}{\gamma}\right) G_{\frac{\gamma}{\gamma}}^0(x) + \frac{\beta}{\gamma} G_{\frac{\gamma}{\gamma}}^2(x) \right) \Big|_{x=\nu_q}^{\infty} \quad (9)$$

حال با توجه به اینکه

$$\lim_{x \rightarrow \infty} G_{\frac{\gamma}{\gamma}}^0(x) = 2, \quad \lim_{x \rightarrow \infty} G_{\frac{\gamma}{\gamma}}^2(x) = 2.$$

پس B برابر است با

$$B = 1 - \frac{1}{\gamma} \left(\left(1 - \frac{\beta}{\gamma}\right) G_{\frac{\gamma}{\gamma}}^0(\nu_q) + \frac{\beta}{\gamma} G_{\frac{\gamma}{\gamma}}^2(\nu_q) \right) \quad (10)$$

در نتیجه

$$ES_X(q) = -\frac{\left(1 - \frac{\beta}{\gamma}\right) G_{\frac{\gamma}{\gamma}}^1(\nu_q) + \frac{\beta}{\gamma} G_{\frac{\gamma}{\gamma}}^3(\nu_q)}{2 - \left(\left(1 - \frac{\beta}{\gamma}\right) G_{\frac{\gamma}{\gamma}}^0(\nu_q) + \frac{\beta}{\gamma} G_{\frac{\gamma}{\gamma}}^2(\nu_q) \right)}. \quad (11)$$

□

۴ تحلیل داده های واقعی

در این فصل به بررسی لگاریتم بازده داده های شاخص سهام شرکت ادوبی^{۳۱} در بازه زمانی یکم، ژانویه سال ۲۰۰۰ تا یکم، ژانویه سال ۲۰۱۲ می پردازیم. این بازه شامل بحران های مالی مهمی است. مانند انفجار حباب دات-کام در سال های ۲۰۰۰ و ۲۰۰۱ (حباب اینترنت) و همچنین بحران های مالی سال های ۲۰۰۷ و ۲۰۰۸. این مجموعه داده از وب سایت finance.yahoo.com قابل دریافت است.

برای برازش توزیع گرام چارلی سکانت هایبربولیک بر این داده ها ابتدا پارامتر β را با استفاده از روش حداکثر درستنمایی برآورد می کنیم، به این منظور ما از تابع mle در نرم افزار متلب استفاده کرده ایم. مقدار برآورد شده این پارامتر ۲/۹۰۹۴ است. به منظور بررسی توصیفی این مجموعه داده، نمودارهای p-p و نمودار بافت نگار این داده ها را به همراه برازش توزیع های نرمال، گرام چارلی سکانت هایبربولیک و توزیع لجستیک در شکل ۲ و شکل ۳ به ترتیب، رسم کرده ایم. همان طور که در شکل ۳ مشخص است، توزیع گرام چارلی سکانت هایبربولیک مدل بهتری برای توصیف این داده ها است. نتایج آزمون کلموگروف و اسمیرنوف به منظور بررسی نیکویی برازش توزیع گرام چارلی سکانت هایبربولیک بر این مجموعه داده در جدول ۱ گزارش شده است، براساس این جدول، فرض اینکه داده ها از توزیع گرام چارلی سکانت هایبربولیک پیروی می کنند را نمی توانیم رد کنیم.

ADOBE^{۳۱}

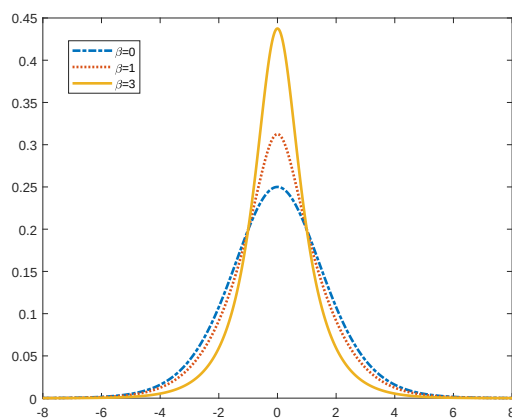
با استفاده از برآورد حداکثر درستنمایی پارامتر β مقادیر ارزش در معرض خطر و کسری مورد انتظار برای توزیع گرام چارلی سکانت هایپربولیک محاسبه شده است و نتایج آن در جدول ۲ ارائه شده است. با توجه به جدول ۲ در می‌یابیم که، نه تنها معیار ارزش در معرض خطر میزان ریسک را نسبت به معیار کسری مورد انتظار کم برآورد می‌کند بلکه کسری مورد انتظار و ارزش در معرض خطر توزیع گرام چارلی سکانت هایپربولیک به کسری مورد انتظار تجربی و ارزش در معرض خطر تجربی نزدیک تر است. به علاوه به ازای چندک های بالا ($q = 0.99$) توزیع گرام چارلی سکانت هایپربولیک دارای کسری مورد انتظار بیشتری نسبت به توزیع نرمال و لجستیک است و این یعنی توزیع گرام چارلی سکانت هایپربولیک در مقایسه با توزیع نرمال و لجستیک برآوردگر محافظه‌کارانه‌ای برای ریسک ارائه می‌دهد (شکل ۴ را ملاحظه فرمایید).

۵ نتیجه‌گیری

در این مقاله توزیع گرام چارلی سکانت هایپربولیک را معرفی کرده‌ایم و فرم بسته‌ای برای محاسبه کسری مورد انتظار آن به کمک تابع فوق هندسی تعمیم یافته ارائه داده‌ایم. این توزیع را به داده‌های سهام شرکت ادوبی برازش داده‌ایم و مزیت‌های آن را نسبت به توزیع پر کاربرد نرمال و لجستیک بررسی کرده‌ایم. همچنین با مقایسه کسری مورد انتظار و ارزش در معرض خطر توزیع تجربی با کسری مورد انتظار و ارزش در معرض خطر توزیع های برازش داده‌شده نشان دادیم که برای این داده‌ها، کسری مورد انتظار و ارزش در معرض خطر توزیع گرام چارلی سکانت هایپربولیک به کسری مورد انتظار و ارزش در معرض خطر توزیع تجربی نزدیک تر است.

مراجع

- [1] Szegö, G.P. (2004), *Risk measures for the 21st century*, Chichester, Hoboken, NJ, Wiley.
- [2] Liu, S. and Heyde, C.C. (2008), On estimation in conditional heteroskedastic time series models under non-normal distributions, *Statistical Papers*, **49**, 455-469.
- [3] Ivanovski, Z., Stojanovski, T. and Narasanov, Z. (2015), Volatility and Kurtosis of Daily Stock Returns at MSE, *UTMS Journal of Economics*, **6**, 209-221.
- [4] Hellman, A. (2015), *Estimating value at risk - an extreme value approach*, Mathematical Statistics Stockholm University Bachelor Thesis 28.
- [5] Zoia, M.G. (2009), Tailoring the Gaussian law for excess kurtosis and skewness by Hermite polynomials, *Communications in Statistics-Theory and Methods*, **39**, 52-64.
- [6] Zoia, M. G., Biffi, P. and Nicolussi, F. (2018), Value at Risk and Expected Shortfall based on Gram-Charlier-like expansions, *Journal of Banking & Finance*, **93**, 92-104.
- [7] Jorion, P. (2001), *Value at risk: the new benchmark for managing financial risk*, volume 2. McGraw-Hill.



شکل ۱: نمودار تابع چگالی احتمال توزیع گرام چارلی هایپرولیک سکانت برای $\beta = 0, 1, 3$

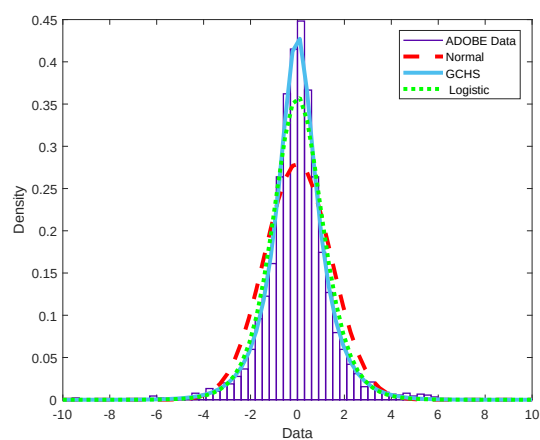
- [8] Uryasev, S. (2000), Conditional value at risk: optimization algorithms and applications, *Financial Engineering*, **14**, 1-5.
- [9] Artzner, P. , Delbaen, F. , Eber, J.-M. and Heath, D. (1999), Coherent measures of risk, *Mathematical Finance*, **9**, 203-228 .
- [10] Acerbi, C. and Tasche, D. (2002), On the coherence of expected shortfall, *Journal of Banking & Finance*, **26**(7), 1487-1503 .
- [11] Wolfram Research, Inc. (2018), Mathematica, Version 11.3, Champaign, IL .
<http://functions.wolfram.com/01.24.21.0027.01>

جدول ۱: آماره آزمون کلموگروف و اسمیرنوف به همراه p -مقدار و نقطه بحرانی برای داده های شرکت ادوبی.

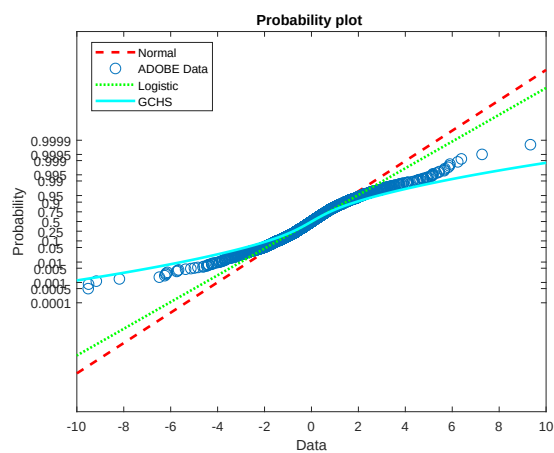
آماره	مقدار بحرانی ۵%	p -مقدار
۰/۰۱۵۲	۰/۰۲۴۷	۰/۴۸۳۰

جدول ۲: مقایسه ارزش در معرض خطر و کسری مورد انتظار برای توزیع تجربی داده‌ها، توزیع نرمال، توزیع گرام چارلی سکانت هایپربولیک و توزیع لجستیک .

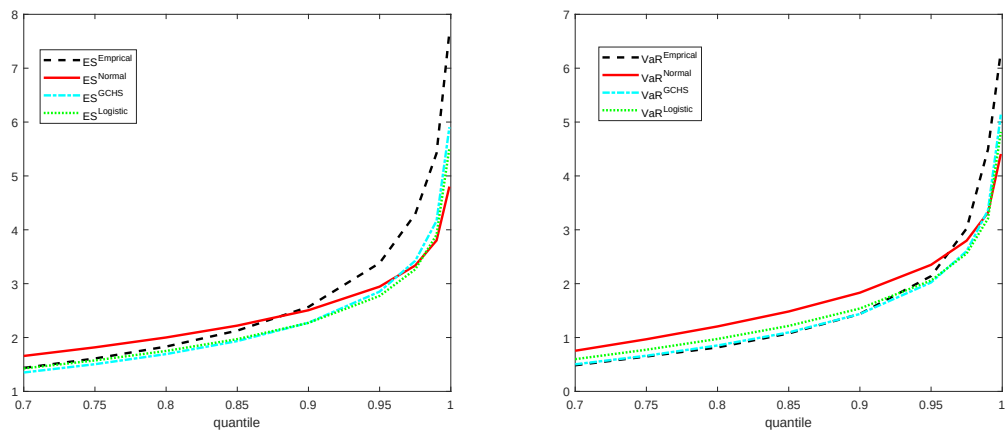
چندک	توزیع تجربی	نرمال	GCHS	لجستیک
$Var_X(q)$				
۰/۷۰۰	۰/۴۸۵	۰/۷۵۴	۰/۵۰۳	۰/۵۹۸
۰/۷۵۰	۰/۶۴۸	۰/۹۶۸	۰/۶۶۲	۰/۷۷۳
۰/۸۰۰	۰/۸۱۷	۱/۲۰۶	۰/۸۵۲	۰/۹۷۴
۰/۸۵۰	۱/۰۷۶	۱/۴۸۴	۱/۰۹۵	۱/۲۱۶
۰/۹۰۰	۱/۴۳۷	۱/۸۳۳	۱/۴۳۷	۱/۵۳۸
۰/۹۵۰	۲/۱۴۰	۲/۳۵۰	۲/۰۲۲	۲/۰۵۹
۰/۹۷۵	۳/۰۲۴	۲/۷۹۹	۲/۶۰۵	۲/۵۵۹
۰/۹۹۰	۴/۴۹۵	۳/۳۲۰	۳/۳۶۵	۳/۲۰۸
۰/۹۹۹	۶/۳۰۸	۴/۴۰۸	۵/۱۳۸	۴/۸۱۸
$ES_X(q)$				
۰/۷۰۰	۱/۴۳۶	۱/۶۵۸	۱/۳۵۱	۱/۴۲۶
۰/۷۵۰	۱/۶۱۱	۱/۸۱۸	۱/۵۰۵	۱/۵۷۵
۰/۸۰۰	۱/۸۳۲	۲/۰۰۱	۱/۶۹۳	۱/۷۵۱
۰/۸۵۰	۲/۱۲۸	۲/۲۲۱	۱/۹۳۴	۱/۹۷۱
۰/۹۰۰	۲/۵۶۸	۲/۵۰۷	۲/۲۷۴	۲/۲۷۲
۰/۹۵۰	۳/۳۸۳	۲/۹۴۵	۲/۸۵۲	۲/۷۷۳
۰/۹۷۵	۴/۲۹۱	۳/۳۳۷	۳/۴۲۴	۳/۲۶۵
۰/۹۹۰	۵/۴۲۲	۳/۸۰۳	۴/۱۶۶	۳/۹۰۸
۰/۹۹۹	۷/۶۶۴	۴/۸۰۲	۵/۹۰۵	۵/۵۱۴



شکل ۲: نمودار بافت نگار لگاریتم بازده داده‌های شرکت ادوبی به همراه برازش توزیع های نرمال، گرام چارلیر سکانت هایپربولیک و توزیع لجستیک.



شکل ۳: نمودار pp برای داده های لگاریتم بازده روزانه‌ی شرکت ادوبی.



شکل ۴: مقایسه مقدار در معرض خطر (راست) و کسری مورد انتظار (چپ) برای توزیع نرمال، توزیع گرام چارلی سکانت هایپربولیک و توزیع تجربی داده‌ها.

بررسی زیان ناشی از بکارگیری توابع بقاء مشترک در بیمه‌های عمر و پس‌انداز

رحیم محمودوند^۱

گروه آمار، دانشکده علوم پایه، دانشگاه بوعلی سینا

چکیده: بیمه‌های عمر و پس‌انداز، با هدف پوشش برخی از ریسک‌های زندگی، در شکل‌های مختلفی معرفی و عرضه می‌شوند. مزایای این بیمه‌ها با توجه به فوت یا حیات بیمه‌گذاران تعریف می‌شوند. در نتیجه توزیع طول عمر نقش کلیدی در محاسبات فنی آنها دارد. معمولاً شرکت‌های بیمه از جدول‌های عمر و توابع بقاء مشترک برای همه اعضای هم‌سن در یک پرتفوی استفاده می‌کنند. هرچند برای افزایش دقت محاسبات همواره به شرکت‌های بیمه توصیه می‌شود که جداول عمر را به روزرسانی کنند. در این مقاله ابتدا، با استفاده از معیار ضریب تغییرات، دو کرانگین برای احتمال بقاء معرفی می‌کنیم. سپس با استفاده از ابزارهای احتمالی نشان می‌دهیم که استفاده از توابع بقاء مشترک به لحاظ نظری باعث زیان شرکت‌های بیمه خواهد شد. **واژه‌های کلیدی:** بیمه‌های عمر و پس‌انداز، بیشانیدن، جدول عمر، ضریب تغییرات. **کد موضوع بندی:** 91B30، 62M86، 60G51.

۱ مقدمه

بشر از بدو تولد تا لحظه مرگ با ریسک‌های مختلفی مواجه می‌شود. اما بدون شک مهم‌ترین ریسک پیش‌روی بشر همان مرگ است. مرگ هر انسانی علاوه بر ناراحتی‌های روحی، پیامدهای مالی برای بازماندگان دارد. بیمه عمر (زندگی) در شکل اولیه برای جبران پیامدهای مالی ناشی از مرگ افراد طراحی شد. اما در حال حاضر بیمه‌های عمر موارد دیگری نظیر نیاز به مراقبت‌های پزشکی، تامین هزینه درمانی بیماری‌های صعب‌العلاج، نیاز به پس‌انداز و تامین آتیه را نیز پوشش می‌دهند. با این کارکردهای اخیر شاید اصطلاح «بیمه زندگی» گویاتر از «بیمه عمر» باشد. از لحاظ نوع قرارداد، بیمه‌های زندگی معمولاً بلندمدت هستند. در نتیجه بیمه‌گر برای استفاده از مبالغ حق بیمه در چرخه اقتصاد فرصت بیشتری در اختیار دارد.

مطابق با آمارهای رسمی در سال ۲۰۱۵ بیش از ۵۳ درصد از کل حق بیمه‌های تولیدی صنعت بیمه در جهان مربوط به بیمه‌های زندگی است (سویس ری، ۲۰۱۸) در حالی که این رقم برای ایران طبق سالنامه بیمه مرکزی ۱۳۹۴ حدود ۱۲ درصد است. بر این اساس می‌توان نتیجه گرفت که بازار بیمه‌های زندگی ظرفیت زیادی برای شکوفا شدن دارد. بدون شک استقبال عمومی و افزایش ضریب نفوذ بیمه‌های زندگی نیازمند بررسی‌های علمی است. در سال‌های گذشته مقاله‌های زیادی درباره بیمه‌های زندگی توسط پژوهشگران و صاحب‌نظران بیمه به چاپ رسیده است. چنانکه مهدوی و ماجد (۱۳۹۳) نیز اشاره کرده‌اند، پژوهش‌های انجام شده درباره بیمه‌های زندگی در ایران در یکی از دسته‌های زیر قرار می‌گیرد:

- نقش بیمه‌های زندگی در اقتصاد و جامعه

- نقش متغیرهای اقتصادی، اجتماعی و شخصیتی در تقاضای بیمه‌های زندگی

^۱ رحیم محمودوند: r.mahmoudvand@basu.ac.ir

• بررسی علل عدم رشد بیمه های زندگی و روش های توسعه آنها.

این دسته بندی، پژوهش های سال های اخیر را نیز دربر می گیرد. برای مثال اسدی و همکاران (۱۳۹۶) عوامل اقتصادی- اجتماعی موثر بر توسعه بیمه های عمر را مطالعه کرده است. در این پژوهش تاثیر متغیرهای اقتصادی و اجتماعی مانند نرخ تورم، سطح تحصیلات، امید به زندگی، نرخ شهرنشینی، نرخ بهره و ... بر ضریب نفوذ بیمه با استفاده از مدل های رگرسیونی بررسی شده است. جاهد و همکاران (۱۳۹۶) وجود گزگزینی در بازار بیمه عمر ایران را با استفاده از اطلاعات اقتصادی- اجتماعی به دست آمده از داده های هزینه-درآمد خانوار بررسی کرده است. بسیاری از پژوهش های انجام شده در ایران که در دسته بندی بالا نیز قرار دارند، مبتنی بر داده های ناکامل و در برخی موارد به صورت پرسشنامه ای و کیفی هستند. بنابراین قابلیت تعمیم و کاربرست آنها برای صنعت بیمه جای تردید بسیار دارد. این در حالی است که در بسیاری از روش های آماری پیشرفته و مدل های روز دنیا صحبت به میان آمده است. به نظر می رسد که مهم ترین زیرساخت برای توسعه بیمه های زندگی در ایران، تهیه بانک های اطلاعاتی درست، دقیق، کارآمد و جامع است که بسیار مغفول مانده است. در این مقاله نشان می دهیم که کمبود اطلاعات و داده های نامناسب چگونه می تواند باعث زیان بیمه گر شود.

۲ معرفی دو پرتفوی کرانگین

ابتدا یادآور می شویم که جدول عمر شامل احتمال های مرگ و میر یک ساله است که بر اساس برآورد تابع q_x ساخته می شود. در صورتی که متغیر تصادفی T طول عمر یک انسان فرضی باشد، تابع q_x عبارت است از:

$$q_x = P(T \leq x + 1 | T > x).$$

با وجود آنکه به نظر می رسد برآورد این تابع کار دشواری نباشد، اما نداشتن داده های مناسب باعث شده است تا در شرکت های بیمه ایرانی محاسبات بیمه ای بر مبنای جداول عمر قدیمی سایر کشورها باشد. به عنوان نمونه در حال حاضر از جدول عمر سال ۱۹۹۳ کشور فرانسه در برخی از شرکت ها استفاده می شود. این در حالی است که جداول عمر با توجه به تحولات جمعیتی هر چند سال یکبار نیاز به بازنگری و اصلاح دارد. جدای از این موضوع یکی از معایب جداول عمر آن است که در یک زمان شانس بقای برابری برای افراد هم سن در نظر می گیرد. در حالی که این موضوع ممکن است با توجه به عوامل مختلف برای شرکت بیمه هزینه زیادی به همراه داشته باشد.

فرض کنید n بیمه گزار در پرتفوی محصولی از یک نوع بیمه زندگی قرار دارند. به علاوه فرض کنید متغیر تصادفی T_j طول عمر بیمه گزار j باشد. در این صورت احتمال های بقاء یک ساله را به صورت زیر تعریف می کنیم:

$$p_{x,j} = P(T_j > x + 1 | T_j > x), \quad j = 1, \dots, n. \quad (1)$$

بردار احتمال های بقاء را به صورت $\mathbf{p}_x = (p_{x,1}, \dots, p_{x,n})$ تعریف می کنیم و برای سادگی فرض می کنیم $p_{x,1} \geq p_{x,2} \geq \dots \geq p_{x,n}$ می توان دو حالت کرانگین در نظر گرفت:

الف- مدل با کمترین تنوع احتمال های بقاء،

ب- مدل با بیشترین تنوع احتمال های بقاء.

در مدل با کمترین تنوع، از احتمال بقاء مشترک استفاده می‌شود؛ یعنی بردار $\mathbf{p}_x = (\bar{p}_x, \dots, \bar{p}_x)$ که در آن

$$\bar{p}_x = \frac{1}{n} \sum_{j=1}^n p_{x,j}$$

مبنای محاسبات است. برای مدل (ب) فرض کنید Π_k به صورت زیر تعریف شده است:

$$\Pi_k = \left\{ (p_{x,1}, \dots, p_{x,n}) : k \leq \sum_{j=1}^n p_{x,j} \leq k+1 \right\}, \quad k = 0, 1, \dots, n-1. \quad (2)$$

در این صورت بردار $\mathbf{p}_x$ که درایه $k+1$ ام آن برابر k و $\sum_{j=1}^n p_{x,j}$ درایه‌های قبل از آن، در صورت وجود، برابر یک و درایه‌های بعد از آن، در صورت وجود، برابر صفر هستند، بیشترین تنوع را فراهم می‌کند. نماد $(\bar{p}_{x,1}, \dots, \bar{p}_{x,n})$ را به عنوان مقادیر مشخص شده تحت مدل (ب) در نظر می‌گیریم. در قضیه زیر میزان تنوع مدل‌های (الف) و (ب) را بررسی می‌کنیم.

بردارهای احتمال بقاء در مدل‌های (الف) و (ب) به ترتیب کمترین و بیشترین ضریب تغییرات را به ازای مقدار ثابتی از $\sum_{j=1}^n p_{x,j}$ فراهم می‌کنند. اثبات. ابتدا توجه کنید که $0 \leq p_{x,j} \leq 1$. پس ضریب تغییرات برای بردار احتمالات بقاء همواره نامنفی است. در این صورت واضح است که تحت مدل (الف) ضریب تغییرات برابر صفر، یعنی کمترین مقدار ممکن، است. برای قسمت (ب) با توجه به نامنفی بودن مقادیر درایه‌های بردار احتمال و ثابت بودن مجموع آنها، کافیسیت روی مجموع توان دوم درایه‌های بردار احتمال متمرکز شویم. تحت شرایط مجموعه Π_k به سادگی داریم:

$$\sum_{j=1}^n p_{x,j}^2 \leq k + \sum_{j=k+1}^n p_{x,j}^2 \quad (3)$$

برابری در رابطه (۳) برای حالتی که k درایه نخست بردار احتمال‌های بقاء برابر ۱ باشند، حاصل می‌شود. اکنون توجه داریم که:

$$\sum_{j=k+1}^n p_{x,j}^2 \leq \left(\sum_{j=k+1}^n p_{x,j} \right)^2 \quad (4)$$

به سادگی می‌توان دید که برابری در رابطه (۴) به ازای حالتی که $\sum_{j=k+1}^n p_{x,j} = p_{x,k+1}$ و مولفه‌های بعدی صفر باشند، حاصل می‌شود. با تلفیق دو رابطه (۳) و (۴) رابطه زیر نتیجه می‌شود:

$$\sum_{j=1}^n p_{x,j}^2 \leq k + \left(\sum_{j=k+1}^n p_{x,j} \right)^2 \quad (5)$$

که در آن برابری به ازای بردار $\mathbf{p}_x$ حاصل می‌شود که درایه $k+1$ ام آن برابر k و $\sum_{j=1}^n p_{x,j}$ درایه‌های قبل از آن، در صورت وجود، برابر یک و درایه‌های بعد از آن، در صورت وجود، برابر صفر هستند. با فرض آنکه $B = \sum_{j=1}^n p_{x,j}$ بیشترین مقدار ضریب تغییرات مربوط به مدل (ب) برابر است با:

$$\max cv^2 = n + \frac{nk(k+1-2B)}{B^2} \quad (6)$$

پیش‌تر در مطالعه محمودوند و حسنی (۲۰۰۹) نشان داده شده است که بیشترین مقدار توان دوم ضریب تغییرات، با فرض نامنفی بودن مشاهده‌ها برابر n است که در رابطه بالا به ازای $k=0$ حاصل می‌شود. $\square$

ضریب تغییرات پیش‌تر برای اندازه‌گیری ریسک تمرکز اعتبار در پرتفوی وام‌های بانکی به کار گرفته شده است (محمودوند و اولیویرا، ۲۰۱۸). در اینجا نیز می‌توان از این معیار به عنوان شاخصی که تنوع سلامتی اعضای یک پرتفوی را نشان می‌دهد استفاده کرد. علاوه بر این با استفاده از مفهوم بیشانیدن کرانگین‌های معرفی شده تحت مدل‌های (الف) و (ب) را می‌توان به صورت ترتیب زیر نوشت. برای هر بردار دلخواه از احتمال‌های بقاء $\mathbf{p}_x = (p_{x,1}, \dots, p_{x,n})$ همواره داریم:

$$(\bar{p}_x, \dots, \bar{p}_x) \prec (p_{x,1}, \dots, p_{x,n}) \prec (\tilde{p}_{x,1}, \dots, \tilde{p}_{x,n}) \quad (7)$$

۳ کاربرد نتایج در بیمه عمر و پس انداز

به منظور استفاده از نتایج بالا قرارداد بیمه عمر و پس انداز یک ساله‌ای را در نظر بگیرید که در آن پرداخت حق بیمه ρ به صورت ماهانه و در ابتدای هر ماه انجام می‌شود. در صورت فوت، سرمایه δ در پایان ماه فوت و در صورت زنده ماندن، سرمایه σ در پایان سال به بیمه‌گذار پرداخت می‌شود. با این تعاریف، ارزش فعلی حق بیمه‌های پرداختی، سرمایه فوت و سرمایه بقاء برای یک پرتفوی شامل n بیمه‌گذار به صورت زیر است (داهان و همکاران، ۲۰۰۲):

$$P = \rho \sum_{j=1}^n \sum_{t=0}^{11} p_{x,j} \nu^t \quad (8)$$

$$D = \delta \sum_{j=1}^n \sum_{t=0}^{11} \left(p_{x,j} \nu^t - p_{x,j} \nu^{t+1} \right) \nu^{t+1} \quad (9)$$

$$S = \sigma \sum_{j=1}^n p_{x,j} \nu^{12} \quad (10)$$

که در آن $\nu = (1+i)^{-1/12}$ مقدار تنزیل ماهانه با در نظر گرفتن نرخ بهره سالانه i است. بر این اساس بایستی مقادیر ρ ، δ و σ به گونه‌ای تعیین شوند که برابری زیر در زمان صدور قرارداد برقرار باشد:

$$P = D + S \quad (11)$$

رابطه D را می‌توان به صورت زیر بازنویسی کرد:

$$\begin{aligned} D &= \delta \sum_{j=1}^n \sum_{t=0}^{11} p_{x,j} \nu^{t+1} - \delta \sum_{j=1}^n \sum_{t=0}^{11} p_{x,j} \nu^{t+1} \\ &= \delta \sum_{j=1}^n \sum_{t=0}^{11} p_{x,j} \nu^{t+1} - \delta \sum_{j=1}^n \sum_{t=1}^{12} p_{x,j} \nu^t \\ &= \delta \sum_{j=1}^n \sum_{t=0}^{11} p_{x,j} \nu^{t+1} - \delta \sum_{j=1}^n \sum_{t=0}^{11} p_{x,j} \nu^t + \delta n - \delta \sum_{j=1}^n p_{x,j} \nu^{12} \\ &= \delta n - \delta (1 - \nu) \sum_{j=1}^n \sum_{t=0}^{11} p_{x,j} \nu^t - \delta \sum_{j=1}^n p_{x,j} \nu^{12} \end{aligned} \quad (12)$$

در نتیجه رابطه (۱۱) را می‌توان به صورت زیر نوشت:

$$\rho \sum_{j=1}^n \sum_{t=0}^{11} \frac{t}{11} p_{x,j} \nu^t = \delta n - \delta(1-\nu) \sum_{j=1}^n \sum_{t=0}^{11} \frac{t}{11} p_{x,j} \nu^t - (\delta - \sigma) \sum_{j=1}^n p_{x,j} \nu^{12} \quad (13)$$

با توجه به رابطه (۱۳) داریم:

$$\rho = -\delta(1-\nu) + \frac{\delta n + (\sigma - \delta) \sum_{j=1}^n p_{x,j} \nu^{12}}{\sum_{j=1}^n \sum_{t=0}^{11} \frac{t}{11} p_{x,j} \nu^t}. \quad (14)$$

به طور مشابه داریم:

$$\delta = \frac{\rho \sum_{j=1}^n \sum_{t=0}^{11} \frac{t}{11} p_{x,j} \nu^t - \sigma \sum_{j=1}^n p_{x,j} \nu^{12}}{n - (1-\nu) \sum_{j=1}^n \sum_{t=0}^{11} \frac{t}{11} p_{x,j} \nu^t - \sum_{j=1}^n p_{x,j} \nu^{12}}, \quad (15)$$

و

$$\sigma = \delta + \frac{(\rho + \delta(1-\nu)) \sum_{j=1}^n \sum_{t=0}^{11} \frac{t}{11} p_{x,j} \nu^t - \delta n}{\sum_{j=1}^n p_{x,j} \nu^{12}}. \quad (16)$$

در روابط (۱۴)، (۱۵) و (۱۶) می‌توانیم از تقریب نمایی، بلدوجی یا خطی برای مولفه‌های $\frac{t}{11} p_{x,j}$ استفاده کنیم (گربر، ۲۰۱۳). در این صورت برای نشان دادن نحوه وابستگی مقادیر حق بیمه و مزایای فوت و بقاء به مقادیر احتمال بقاء از نمادهای $\rho(\mathbf{p}_x)$ ، $\delta(\mathbf{p}_x)$ و $\sigma(\mathbf{p}_x)$ استفاده می‌کنیم. تحت چنین شرایطی داهان و همکاران (۲۰۰۲) با استفاده از مفهوم بیشانیدن و توابع شور-محدب و شور-مقعر نشان داده‌اند که

$$\rho(\bar{p}_{x,1}, \dots, \bar{p}_{x,n}) \leq \rho(p_{x,1}, \dots, p_{x,n}) \leq \rho(\tilde{p}_{x,1}, \dots, \tilde{p}_{x,n}) \quad (17)$$

$$\delta(\bar{p}_{x,1}, \dots, \bar{p}_{x,n}) \geq \delta(p_{x,1}, \dots, p_{x,n}) \geq \delta(\tilde{p}_{x,1}, \dots, \tilde{p}_{x,n}) \quad (18)$$

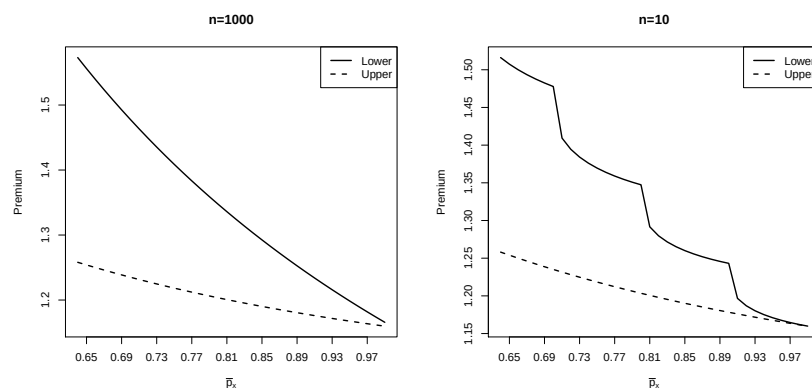
$$\sigma(\bar{p}_{x,1}, \dots, \bar{p}_{x,n}) \geq \sigma(p_{x,1}, \dots, p_{x,n}) \geq \sigma(\tilde{p}_{x,1}, \dots, \tilde{p}_{x,n}) \quad (19)$$

این بدان معناست که مدل (الف) با وجود دریافت حق بیمه‌ی کمتر، سرمایه‌ی فوت و پس انداز بیشتری می‌پردازد. بیمه عمر و پس انداز یک ساله‌ای را در نظر بگیرید که سرمایه فوت آن برابر ۱۰ میلیون تومان و سرمایه حیات برابر ۱۵ میلیون تومان باشد. به علاوه فرض کنید $i = 0.15$. تحت این شرایط کران‌های بالایی و پایینی برای حق بیمه در شکل ۱ برای یک پرتفوی با تعداد بیمه‌گذاران مختلف رسم شده است.

همانگونه که ملاحظه می‌کنیم:

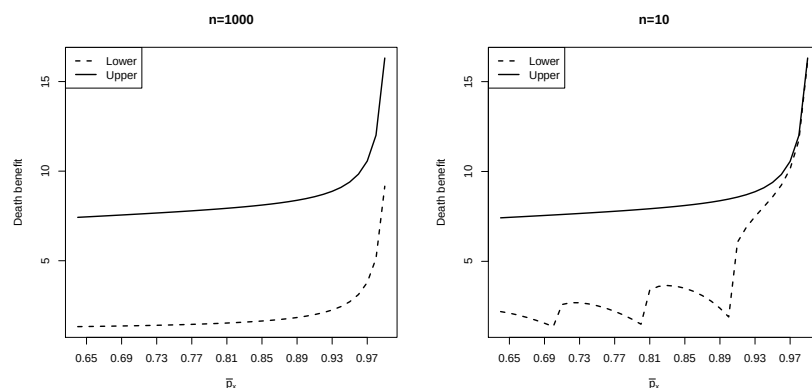
- با افزایش تعداد بیمه‌گذاران کران بالایی افزایش پیدا می‌کند اما کران پایینی ثابت است.
- با افزایش متوسط احتمال‌های بقاء، کران پایینی و بالایی کاهش پیدا می‌کنند.
- با افزایش متوسط احتمال‌های بقاء اختلاف بین کران‌های بالایی و پایینی کم می‌شود.
- با افزایش همزمان متوسط احتمال و تعداد بیمه‌گذاران اختلاف کران‌های بالایی و پایینی افزایش پیدا می‌کند.

شکل‌های ۲ و ۳ نیز به ترتیب مقادیر سرمایه فوت و سرمایه پس‌انداز را به ازای شرایط مختلف نشان می‌دهند. همانگونه که ملاحظه می‌کنیم، هر دو شکل نشان می‌دهند که با افزایش متوسط احتمال بقاء، کران‌ها افزایشی هستند. با این وجود اختلاف بین دو کران با افزایش متوسط احتمال بقاء کاهش پیدا می‌کند. هر چند با افزایش همزمان تعداد بیمه‌گذاران و متوسط احتمال بقاء، اختلاف



شکل ۱: مقایسه کران پایینی و بالایی حق بیمه در دو گروه کوچک و بزرگ

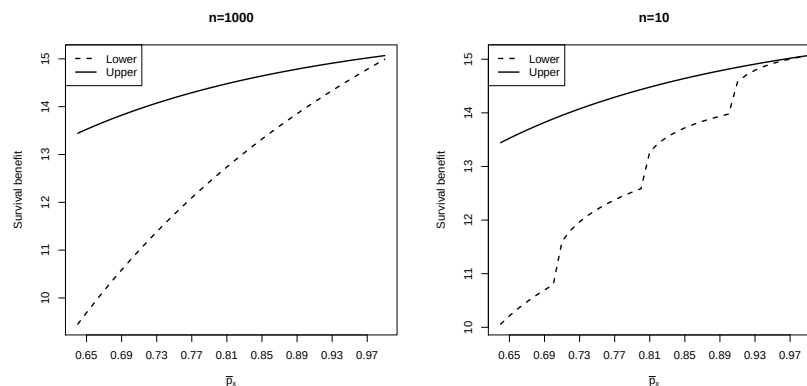
کران‌های پایینی و بالایی اندکی افزایش پیدا می‌کند. در ضمن برخلاف حق بیمه‌ها، در هر دو شکل ۱ و ۲ کران‌های بالایی و پایینی تغییر می‌کنند.



شکل ۲: مقایسه کران پایینی و بالایی سرمایه فوت در دو گروه کوچک و بزرگ

۴ نتیجه‌گیری و مطالعات آتی

در این مقاله درباره اهمیت تابع بقاء در بیمه‌های زندگی بحث شد. مطابق با نتایج به دست آمده، مقدار حق بیمه‌ای که از بیمه‌گذاران در پرتفوی بیمه زندگی دریافت می‌شود، در صورت استفاده از تابع بقاء مشترک کمتر از مقدار حق بیمه‌ای است که در صورت نابرابری توابع بقاء محاسبه و دریافت شود. در عین حال مزایایی که پرداخت می‌شود با فرض برابری تابع بقاء بیشتر از حالت نابرابر است. اما نکته‌ای که از حل مثال عددی به دست آمد نشان داد که میزان اختلاف برای مقادیر احتمال بقاء بالا در حق بیمه و سرمایه حیات بسیار ناچیز است. در حالی که اختلاف سرمایه فوت در گروه‌های بزرگ‌تر برای همه مقادیر احتمال بقاء نسبتاً بالاست. به منظور تکمیل



شکل ۳: مقایسه کران پایینی و بالایی سرمایه پس انداز در دو گروه کوچک و بزرگ

نتایج محاسبات برای بیمه‌های عمر با مدت بیشتر از یکسال و بررسی میزان زیان ناشی از بکارگیری توابع بقاء مشترک بایستی در مطالعه‌های بعدی انجام شود.

مراجع

- [۱] طاهره جاهد، قدرت‌الله امام‌وردی و علیرضا دقیقی اصلی. (۱۳۹۶). تحلیل وجود گزگزینی در بازار بیمه عمر ایران با استفاده از مدل‌های لوجیت. پژوهشنامه بیمه، شماره ۱۲۵، صص ۴۳ تا ۶۴.
- [۲] غدیر مهدوی و وحید ماجد. (۱۳۹۳). اثر عوامل کمی و کیفی موثر بر تقاضای بیمه عمر در کشور. پژوهشنامه بیمه، شماره ۱۱۴، صص ۳۷ تا ۶۶.
- [۳] سعید اسدی قراگوز، علی ارشدی و غلامعلی حاجی. (۱۳۹۶). عوامل اقتصادی-اجتماعی موثر بر توسعه بیمه عمر، مطالعه مقایسه‌ای بین ایران و کشورهای توسعه یافته در طول دوره ۱۹۸۵-۲۰۱۴: رویکرد گشتاورهای تعمیم‌یافته. پژوهشنامه بیمه، شماره ۱۲۷، صص ۲۱ تا ۴۰.
- [۴] سایت بیمه مرکزی ج.ا.ا. به آدرس <http://www.centinsur.ir>

- [5] Dahan, M., Frostig, E. and Longberg, N.A. (2002), Life Insurance Policies with Statistical Heterogeneous Populations, *Scandinavian Actuarial Journal*, **3**, 212-222.
- [6] Gerber, H.U. (2013), *Life insurance mathematics*, Springer Science & Business Media.
- [7] Mahmoudvand, R. and Hassani, H. (2009), Two New Confidence Interval for the Coefficient of Variation in a Normal Distribution, *Journal of Applied Statistics*, **36**(4), 429-442.

- [8] Mahmoudvand, R. and Oliveira, T.A. (2018), *On the application of sample coefficient of variation for managing loan portfolio risks*, In Recent Studies on Risk Analysis and Statistical Modeling (pp. 87-97), Springer, Cham.
- [9] Swiss Re, sigma database, sigma No. 3/2018.

کاربردی از اطلاع کولبک-لیبلر تجمعی در پردازش تصاویر

فاطمه منصوری^۱، سعید طهماسبی^۱، مراد علیزاده^۱، فرهاد فرخ سرشت^۱، صفیه دانشی^۲

^۱ گروه آمار، دانشکده علوم ریاضی، دانشگاه خلیج فارس

^۲ گروه آمار، دانشکده علوم ریاضی، دانشگاه صنعتی شاهرود

چکیده: امروزه معیار اطلاع تشخیص کولبک-لیبلر جایگاه ویژه‌ای در مبحث قابلیت اعتماد پیدا کرده است. اطلاع کولبک-لیبلر تجمعی به تازگی بعنوان یک تعمیم جدید از معیار کولبک-لیبلر براساس تابع توزیع تجمعی پیشنهاد شده‌است. در این مقاله ابتدا به معرفی اطلاع کولبک-لیبلر تجمعی می‌پردازیم، سپس ویژگی‌ها و کران‌های آن را بیان کرده و در ادامه به معرفی کولبک-لیبلر تجمعی پویا برای طول عمرهای گذشته می‌پردازیم و در نهایت نشان می‌دهیم که با استفاده از کولبک-لیبلر تجمعی تجربی می‌توان سطح خاکستری تصاویر دیجیتالی شده را مقایسه کرد و میزان تفاوت و وضوح تصاویر را نشان داد. واژه‌های کلیدی: اطلاع کولبک-لیبلر، تابع توزیع تجمعی، پردازش تصاویر، آنتروپی تجمعی. کد موضوع بندی: 60e15, 62n05, 94a17.

۱ مقدمه ای بر معیار اطلاع کولبک-لیبلر تجمعی

نظریه اطلاع، شاخه نسبتاً جدیدی از علم ریاضی است که کمیت مفهوم اطلاع را تعیین می‌کند. آنتروپی شانون^{۳۲} به عنوان زیربنای نظریه اطلاع، برای اولین بار توسط شانون (۱۹۴۸) مطرح شد. به دلیل چالش‌ها و نواقص آنتروپی شانون در مباحث قابلیت اطمینان، تلاش‌های زیادی به منظور تعریف معیارهای جدید صورت گرفته‌است. از جمله این معیارها، معیار آنتروپی باقی‌مانده تجمعی^{۳۳} است که توسط راثو و همکاران (۲۰۰۴) معرفی شد. همچنین آنتروپی تجمعی توسط دی کرشنزو و لونگوبردی (۲۰۰۹) مطرح و ناوارو و همکاران (۲۰۱۰) خواص آن در مباحث قابلیت اعتماد را بیان کردند. تابع اطلاع تشخیص کولبک-لیبلر^{۳۴} که اولین بار توسط کولبک و لیبلر (۱۹۵۱) معرفی شده یکی از اندازه‌های اطلاع برای مقایسه دو توزیع است. بررسی خواص و ویژگی‌های این معیار در نظریه قابلیت اعتماد، محاسبه کران‌های بالا و پایین این معیار اطلاع و محاسبه مقادیر تجربی آن در مطالعات طول عمر سیستم‌ها از جمله مسائلی است که به آن پرداخته شده است. در این بخش، ابتدا به معرفی معیار اطلاع کولبک-لیبلر تجمعی که توسط دی کرشنزو و لونگوبردی (۲۰۱۵) مطرح شده، می‌پردازیم و در ادامه خواص آن را بیان می‌کنیم. فرض کنید X و Y متغیرهای تصادفی با میانگین‌های متناهی باشند، همچنین فرض کنید که $r_X = \sup\{t \in \mathbb{R} : F_X(t) < 1\}$ و $l_X = \inf\{t \in \mathbb{R} : F_X(t) > 0\}$ و $l = l_X = l_Y$ آنگاه اطلاع کولبک-لیبلر تجمعی X از Y به صورت زیر تعریف

^۱ فاطمه منصوری: fatemehmansouri1993@gmail.com

^{۳۲} Shannon

^{۳۳} Residual cumulative entropy

^{۳۴} Kullback-Leibler

می‌شود:

$$C_{KL}(X, Y) = \int_l^{\max(r_X, r_Y)} F(t) \log\left(\frac{F(t)}{G(t)}\right) dt + E(X) - E(Y), \quad (۱)$$

به شرطی که انتگرال سمت راست متناهی باشد. اگر X, Y متغیرهای تصادفی کراندار با تابع توزیع F و G باشند، آنگاه برای $a, c \geq 0$ و $b \geq 0$ داریم

$$\begin{aligned} C_{KL}(aX + b, cY + b) &= aC_{KL}\left(X, \frac{cY}{a}\right) + aE(X) - cE(Y) \\ &= cC_{KL}\left(\frac{aX}{c}, Y\right) + aE(X) - cE(Y). \end{aligned}$$

اثبات. ابتدا براساس رابطه (۱) داریم

$$\begin{aligned} C_{KL}(aX + b, cY + b) &= \int_l^{\max(r_X, r_Y)} F_{aX+b}(x) \log\left(\frac{F_{aX+b}(x)}{G_{cY+b}(x)}\right) dx + E(aX + b) - E(cY + b) \\ &= \int_l^{\max(r_X, r_Y)} F_X\left(\frac{x-b}{a}\right) \log\left(\frac{F_X\left(\frac{x-b}{a}\right)}{G_Y\left(\frac{x-b}{c}\right)}\right) dx + aE(X) - cE(Y) \\ &= aC_{KL}\left(X, \frac{cY}{a}\right) + aE(X) - cE(Y) \end{aligned}$$

در نتیجه اثبات با تغییر متغیر $z = \frac{x-b}{a}$ و به طور مشابه با تغییر متغیر $t = \frac{x-b}{c}$ می‌توان نوشت

$$\begin{aligned} C_{KL}(aX + b, cY + b) &= \int_l^{\max(r_X, r_Y)} F_X\left(\frac{x-b}{a}\right) \log\left(\frac{F_X\left(\frac{x-b}{a}\right)}{G_Y\left(\frac{x-b}{c}\right)}\right) dx + aE(X) - cE(Y) \\ &= \int_l^{\max(r_X, r_Y)} F_X\left(\frac{at}{c}\right) \log\left(\frac{F_X\left(\frac{at}{c}\right)}{G_Y(t)}\right) cdt + aE(X) - cE(Y) \\ &= cC_{KL}\left(\frac{aX}{c}, Y\right) + aE(X) - cE(Y). \end{aligned}$$

□

در نتیجه اثبات کامل است.

برای هر جفت از متغیرهای تصادفی کراندار X, Y ، معیار عدم دقت تجمعی چنین تعریف می‌شود:

$$K(X, Y) = - \int_l^{\max(r_X, r_Y)} F(t) \log(G(t)) dt. \quad (۲)$$

برای هر جفت از متغیرهای تصادفی X, Y با توابع توزیع F و G و تعریف $C_{KL}(X, Y)$ یک عبارت تحلیلی به صورت زیر داریم

$$C_{KL}(X, Y) = K(X, Y) - C\mathcal{E}(X) + E(X) - E(Y), \quad (۳)$$

به طوری که $C\mathcal{E}(X)$ آنتروپی تجمعی X است. فرض کنید X و X_θ بیانگر متغیرهای تصادفی در مدل مخاطره معکوس با توابع توزیع $F(x)$ و $F_\theta(x) = (F(x))^\theta$ باشند، آنگاه داریم

$$C_{KL}(X, X_\theta) = (\theta - 1)C\mathcal{E}(X) + E(X) - E(X_\theta)$$

۲ نابرابری و محاسبه کران‌ها برای $C_{KL}(X, Y)$

دو متغیر تصادفی X و Y که دارای توابع توزیع $F(t)$ و $G(t)$ هستند را در نظر بگیرید، متغیر تصادفی X بزرگتر از Y باشد،
 الف) در ترتیب تصادفی معمولی ($X \geq_{st} Y$) اگر $F(t) \leq G(t)$ برای هر $t \in \mathcal{R}$ ،
 ب) در ترتیب تصادفی نرخ خطر معکوس ($X \geq_{rh} Y$) اگر $\tau_X(t) \geq \tau_Y(t)$ یا برای هر $s \leq t$ ، $F(s)G(t) \leq F(t)G(s)$ ، اگر $X \leq_{st} Y$ باشد، آنگاه

$$\mathcal{CE}(X) - E(X) \leq \mathcal{CE}(Y) - E(Y). \quad (۴)$$

الف) فرض کنید X و Y متغیرهای تصادفی مطلقاً پیوسته باشند به طوریکه مقادیرش را در بازه $[0, r]$ می‌گیرد. آنگاه داریم

$$C_{KL}(X, Y) \geq [r - E(X)] \log\left(\frac{r - E(X)}{r - E(Y)}\right) + E(X) - E(Y). \quad (۵)$$

با استفاده از نامساوی $x \log(x) \geq x - 1$ به آسانی می‌توان دید که سمت راست رابطه (۵) غیر منفی است و وقتی $EX = EY$ باشد، مقدار $C_{KL}(X, Y)$ برابر صفر می‌شود.

ب) اگر X و Y متغیرهای تصادفی با میانگین متناهی و دارای نقاط سمت چپ $l_X = l_Y = l$ باشد، آنگاه داریم

$$C_{KL}(X, Y) \geq \frac{1}{2} \int_l^{\max(r_X, r_Y)} \frac{[F(u) - G(u)]^2}{F(u) + G(u)} du. \quad (۶)$$

فرض کنید X و Y متغیرهای تصادفی با میانگین های متناهی و با فرض این که حد پایین $l = l_x = l_y$ باشد، آنگاه
 الف)

$$C_{KL}(X, Y) \leq \int_l^{\max(r_X, r_Y)} F(u) \left[\frac{F(u)}{G(u)} - 1 \right] du + EX - EY, \quad (۷)$$

البته فرض می‌شود انتگرال سمت راست متناهی باشد.

ب) اگر $X \geq_{st} Y$ باشد، آنگاه خواهیم داشت

$$C_{KL}(X, Y) \leq \frac{1}{2} \int_l^{\max(r_X, r_Y)} F(u) \left[\frac{F(u)}{G(u)} - 1 \right] \left[3 - \frac{F(u)}{G(u)} \right] du + EX - EY. \quad (۸)$$

۳ کولبک-لیبلر تجمعی پویا برای طول عمرهای گذشته

دو متغیر تصادفی نامنفی X ، Y را در نظر بگیرید، که دارای طول عمر گذشته در بازه $(0, r_X)$ و $(0, r_Y)$ باشند. اگر برای هر $t > 0$ ، طول عمر گذشته آن‌ها به فرم $X_{(t)} = [X | X \leq t]$ و $Y_{(t)} = [Y | Y \leq t]$ باشد، آنگاه برای هر $t > 0$ ، توابع توزیع $X_{(t)}$ و $Y_{(t)}$ بصورت زیر تعریف می‌شوند:

$$F_{(t)}(x) = \frac{F(x)}{F(t)}; \quad G_{(t)}(x) = \frac{G(x)}{G(t)}, \quad 0 \leq x \leq t, \quad (۹)$$

این مفهوم از تعریف متغیر $X(t)$ و $Y(t)$ در زمینه تحلیل بقا و قابلیت اطمینان سیستم‌ها کاربرد فراوانی دارد. در بخش های قبل معیار کولبک-لیبلر تجمعی را ارائه کردیم و کران بالا و پایین را برای آن بیان نمودیم. اکنون می‌خواهیم یک نسخه پویا از اطلاع کولبک-لیبلر تجمعی از طول عمر گذشته $X(t)$ و $Y(t)$ ارائه دهیم که شامل میانگین طول عمر گذشته X, Y است که بصورت زیر تعریف می‌شود:

$$\mu_X(t) = E[X(t)] = \int_0^t \left[1 - \frac{F(x)}{F(t)}\right] dx, \quad \mu_Y(t) = E[Y(t)] = \int_0^t \left[1 - \frac{G(y)}{G(t)}\right] dy, t \in D_{X,Y}$$

در صورتی که $D_{X,Y} := (\circ, \min(r_x, r_y))$.

اطلاع کولبک-لیبلر تجمعی بین دو متغیر طول عمر گذشته تصادفی $X(t)$ و $Y(t)$ بصورت زیر تعریف می‌شود:

$$C_{KL}(X(t), Y(t)) := \int_0^t \frac{F(x)}{F(t)} \log\left(\frac{F(x)}{F(t)} \frac{G(t)}{G(x)}\right) dx + \mu_X(t) - \mu_Y(t), \quad (10)$$

در صورتی که $C_{KL}(X(t), Y(t))$ نزدیکی دو توزیع طول عمر را اندازه‌گیری می‌کند. نکته ۱: از تعریف (۱۰) و (۱) به آسانی می‌توان دید که $C_{KL}(X(t), Y(t))$ برای هر $t \in D_{X,Y}$ نامنفی است. همچنین می‌توان نوشت

$$\lim_{x \rightarrow \circ^+} C_{KL}(X(t), Y(t)) = \circ, \quad \lim_{x \rightarrow \sup D_{X,Y}} C_{KL}(X(t), Y(t)) = C_{KL}(X, Y).$$

نکته ۲: اگر $t \in D_{X,Y}$ بطور مشابه از معادله (۳) و رابطه (۱۰) اطلاع کولبک-لیبلر تجمعی X و Y را می‌توان بصورت زیر بیان کرد

$$C_{KL}(X(t), Y(t)) := K[X(t), Y(t)] - C\mathcal{E}(X(t)) + \mu_X(t) - \mu_Y(t),$$

که طبق رابطه (۲) و (۹)، $K[X(t), Y(t)]$ عدم دقت تجمعی بین دو متغیر تصادفی طول عمر گذشته است که بصورت زیر تعریف می‌شود:

$$K[X(t), Y(t)] := - \int_0^t \frac{F(x)}{F(t)} \log\left(\frac{G(x)}{G(t)}\right) dx, \quad (11)$$

همچنین

$$C\mathcal{E}(X(t)) = - \int_0^t \frac{F(x)}{F(t)} \log\left(\frac{F(x)}{F(t)}\right) dx, \quad (12)$$

آنتروپی تجمعی پویای X است. نکته آن‌که معیار عدم دقت تجمعی داده‌شده در رابطه (۱۱) یک معیار پویا بین دو توزیع طول عمر گذشته است که توسط کومار و همکاران (۲۰۱۱) پیشنهاد شده است.

اکنون یک کران بالایی برای $C_{KL}(X(t), Y(t))$ بیان می‌کنیم. فرض کنید X و Y متغیرهای تصادفی نامنفی با توابع توزیع F و G باشند.

اگر $X \geq_{rh} Y$ آنگاه یک کران بالا برای $C_{KL}(X(t), Y(t))$ بصورت زیر است:

$$C_{KL}(X(t), Y(t)) \leq \frac{1}{\sqrt{t}} \int_0^t F(t)(u) \left[\frac{F(t)(u)}{G(t)(u)} - 1 \right] \left[\sqrt{t} - \frac{F(t)(u)}{G(t)(u)} \right] du + \mu_X(t) - \mu_Y(t).$$

۴ کاربرد اطلاع کولبک-لیبلر تجمعی در تجزیه و تحلیل تصاویر

در این بخش سعی داریم که بر روی کاربرد مناسب اطلاعات کولبک-لیبلر تجمعی در تجزیه و تحلیل تصویر تمرکز کنیم؛ برای این منظور ابتدا به مسئله تخمین اطلاع کولبک-لیبلر تجمعی با استفاده از یک اندازه‌گیری تجربی از تشخیص میپردازیم. تعریف: فرض کنید X, Y متغیرهای تصادفی با توابع توزیع G و F باشند، اطلاع کولبک-لیبلر تجمعی تجربی بین X, Y بصورت زیر محاسبه می‌شود

$$C_{KL}(\hat{F}_n, \hat{G}_m) = \int_0^{+\infty} \hat{F}_n(x) \log\left(\frac{\hat{F}_n(x)}{\hat{G}_m(x)}\right) dx + \bar{X}_n - \bar{Y}_m,$$

که در آن $\hat{F}_n(x) = \frac{1}{n} \sum_{i=1}^n I_{X_i \leq x}$ و $\hat{G}_m(x) = \frac{1}{m} \sum_{i=1}^m I_{Y_i \leq x}$ توزیع تجربی متغیرهاست و $\bar{X}_n, \bar{Y}_m$ میانگین‌های نمونه هستند. برای بدست آوردن یک فرمول موثر برای $C_{KL}(\hat{F}_n, \hat{G}_m)$ آماره‌های ترتیبی X را بصورت $X_{(1)} \leq X_{(2)} \leq \dots \leq X_{(n)}$ و بطور مشابه آماره‌های ترتیبی Y نیز به صورت $Y_{(1)} \leq Y_{(2)} \leq \dots \leq Y_{(m)}$ تعریف می‌شود. نکته: طبق رابطه (۳) می‌دانیم که اطلاع کولبک-لیبلر تجمعی را می‌توان به این صورت زیر بازنویسی کرد

$$C_{KL}(\hat{F}_n, \hat{G}_m) = K[\hat{F}_n, \hat{G}_m] - \mathcal{CE}(\hat{F}_n) + \bar{X}_n - \bar{Y}_m. \quad (۱۳)$$

مثال کاربردی ۱: حال می‌خواهیم اطلاع کولبک-لیبلر تجمعی تجربی را برای اندازه‌گیری اختلافات در سطح خاکستری تصاویر دیجیتالی شده در شکل ۱ بکار بگیریم.

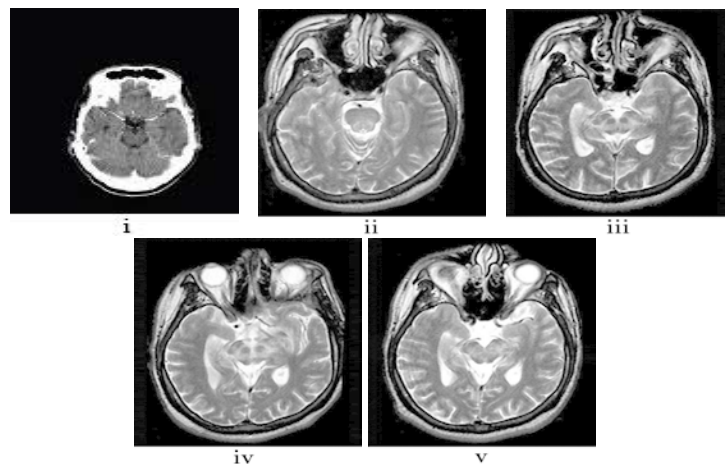
تصاویر توسط ۱۶۰۰ سلول تشکیل شده‌است. سطح خاکستری از هر سلول توسط یک عدد واقعی از صفر (سیاه) تا یک (سفید) اندازه‌گیری شده‌است. شکل ۲ سطح خاکستری مرتب شده در جهت افزایش را نشان می‌دهد. بنابراین میانگین سطح خاکستری به سمت مقدار یک، وضوح تصویر بیشتر و به سمت صفر عدم وضوح را نشان می‌دهد. علاوه بر این متوسط سطح خاکستری از پنج تصویر در جدول ۱ ارائه شده‌است.

اطلاعات کولبک-لیبلر تجمعی تجربی از سطوح خاکستری برای جفت تصاویر با استفاده از معادله (۱۳) برای $n = m = ۱۶۰۰$ جدول ۲ ارزیابی می‌شود. حداکثر اختلاف بین شکل‌های $(i), (v)$ مشاهده می‌شود، در حالیکه حداقل مقدار $C_{KL}(\hat{F}_n, \hat{G}_m)$ مربوط به شکل‌های $(v), (vi)$ هستند.

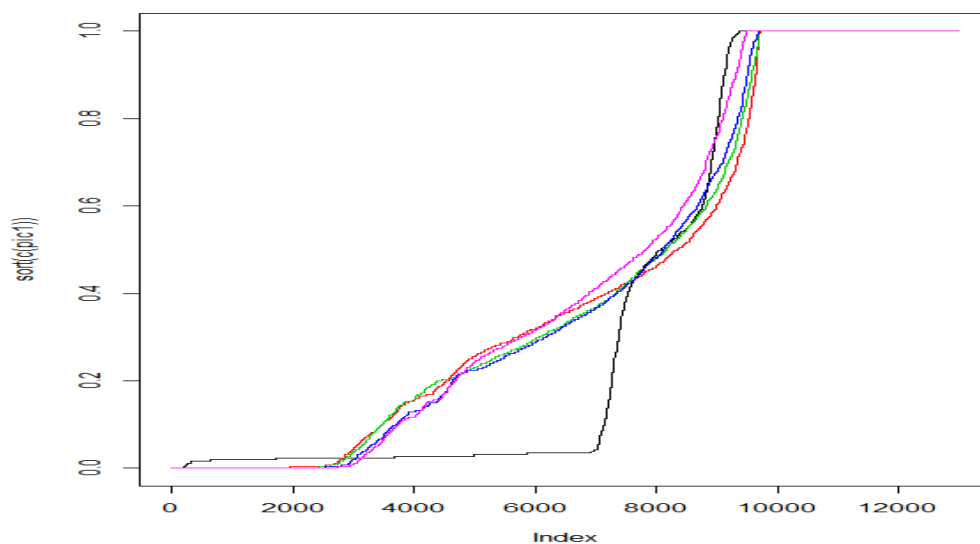
یک کاربرد مناسب از اطلاع کولبک-لیبلر تجمعی تجربی در تجزیه و تحلیل تصویر را نشان دادیم. در این زمینه، با شروع از سطوح خاکستری تصاویر در نمونه‌های تصادفی مقادیر بالاتر از $C_{KL}(\hat{F}_n, \hat{G}_m)$ تفاوت قابل توجهی در شکل‌ها داشتند. بنابراین یک روش خودکار که اطلاعات کولبک-لیبلر تجمعی تجربی را ارزیابی می‌کند قادر است تصاویر را براساس سطح خاکستری یا (رنگی) طبقه‌بندی کند.

برای درک بهتر کاربرد اطلاع کولبک-لیبلر تجمعی تجربی در اندازه‌گیری اختلافات سطح خاکستری تصاویر به مثال دیگری می‌پردازیم.

مثال کاربردی ۲: همانند مثال قبل سطح خاکستری از هر سلول توسط یک عدد واقعی از صفر (سیاه) تا یک (سفید) اندازه‌گیری شده‌است. هرچه میانگین سطح خاکستری بیشتر باشد اختلاف بین دو تصویر کمتر است به این معنا که هر دو تصویر از نظر وضوح تصویر و سطح خاکستری نزدیک به هم هستند. در جدول ۳ همانطور که مشاهده می‌کنیم شکل (i) دارای بیشترین سطح خاکستری می‌باشد. به این معنی است که وضوحی شکل (i) به سلول‌های سفید مربوط می‌شود، همچنین اطلاع کولبک-لیبلر تجمعی تجربی از



شکل ۱: تصاویر اسکن مغز



شکل ۲: میزان سطح خاکستری مرتب‌شده در شکل ۱

سطح خاکستری برای جفت تصاویر با استفاده از معادله ۱۳ ارزیابی شده است. از جدول ۴ می‌بینیم که حداکثر اختلاف بین شکل‌های (i) ، (ii) می‌باشد، در حالیکه حداقل مقدار $C_{KL}(\hat{F}_n, \hat{G}_m)$ مربوط به شکل‌های (ii) و (iii) هستند. به این منظور که اگر اختلاف کولبک-لیبلر تجمعی تجربی بین دو شکل زیاد باشد یعنی تفاوت زیادی در سطح خاکستری تصاویر و وضوح تصویر دارند و اگر تفاوت بین دو تصویر کم باشد بدین معنی است که سطح خاکستری در دو تصویر نزدیک به هم است و وضوح هر دو تصویر یکسان است.

شکل	میانگین سطح خاکستری
(i)	۰/۲۸۰۷۱
(ii)	۰/۵۶۲۹۷
(iii)	۰/۵۹۲۳۴
(iv)	۰/۶۳۲۱۹
(v)	۰/۷۸۵۰۲

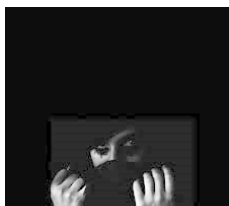
جدول ۱: میانگین سطح خاکستری برای شکل (۱)

X, Y	(i)	(ii)	(iii)	(iv)	(v)
(i)	۰	۰/۰۹۳۷۲	۰/۱۲۸۵۹	۰/۲۱۸۴۲	۰/۴۰۲۳۵
(ii)	۰/۰۷۵۵۰	۰	۰/۰۷۲۹۵	۰/۰۸۹۹۰	۰/۰۵۴۳۵
(iii)	۰/۱۸۲۹۱	۰/۳۴۴۱۵	۰	۰/۰۱۲۸۷	۰/۰۷۲۹۵
(iv)	۰/۱۲۷۵۱	۰/۰۲۱۶۸	۰/۱۴۹۴۲	۰	۰/۰۰۱۴۸
(v)	۰/۳۳۲۱۷	۰/۰۶۶۲۴	۰/۳۳۲۱۷	۰/۱۵۱۲	۰

جدول ۲: کولبک-لیبلر تجمعی تجربی برای شکل (۱)



i



ii

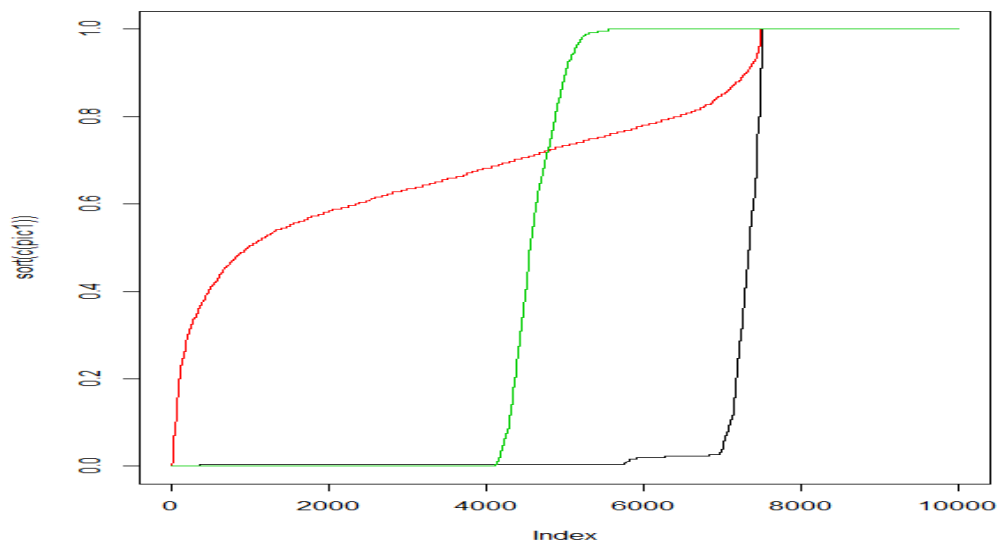


iii

شکل ۳: تصاویر منتخب

شکل	میانگین سطح خاکستری
(i)	۰/۷۴۰۶۱۸۸
(ii)	۰/۲۷۳۷۰۳۵
(iii)	۰/۵۳۹۰۸۲۴

جدول ۳: میانگین سطح خاکستری برای شکل (۳)



شکل ۴: سطح خاکستری مرتب شده در شکل ۳

X, Y	(i)	(ii)	(iii)
(i)	°	°/۳۵۷۷۷۶۵	°/۲۰۱۴۲۶۲
(ii)	۱/۱۱۴۲۵۶	°	°/۶۷۵۵۰۱۷
(iii)	°/۵۷۳۷۰۷۰	°/۰۵۹۳۲۵۸۷	°

جدول ۴: کولبک-لیبلر تجمعی تجربی برای شکل (۳)

مراجع

- [1] Di Crescenzo, A. and Longobardi, M. (2015), Some properties and applications of cumulative Kullback-Leibler information. *Applied Stochastic Models in Business and Industry*, **31**(6), 875-891.
- [2] Di Crescenzo, A. and Longobardi, M. (2009), On cumulative entropies. *Journal of Statistical Planning and Inference*, **139**(12), 4072-4087.
- [3] Kullback, S. and Leibler, R.A. (1951), On information and sufficiency. *The Annals of Mathematical Statistics*, **22**(1), 79-86.
- [4] Kumar, V., Taneja, H.C. and Srivastava, R. (2011), A dynamic measure of inaccuracy between two past lifetime distributions. *Metrika*, **74**, 1-10.

- [5] Navarro, J., del Aguila, Y. and Asadi, M. (2010), Some new results on the cumulative residual entropy. *Journal of Statistical Planning and Inference*, **140**(1), 310-322.
- [6] Rao, M., Chen, Y., Vemuri, B.C. and Wang, F. (2004), Cumulative residual entropy: a new measure of information. *IEEE transactions on Information Theory*, **50**(6), 1220-1228.
- [7] Shannon, C. (1948). A mathematical theory of communication. *Bell System Technical Journal* , **27**, 379-432.

توزیع پارتو نمایی و کاربرد آن در قابلیت اعتماد

رویا نصیرزاده^۱، زهرا نیکنام^۲

^۱ گروه آمار، دانشگاه فسا، فسا، ایران

^۲ گروه آمار، دانشکده علوم ریاضی و آمار، دانشگاه اهواز، اهواز، ایران

چکیده: توزیع پارتو نمایی یکی از مهمترین توزیع‌ها در تحلیل داده‌های طول عمر می‌باشد. در این مقاله ابتدا مروری بر توزیع پارتو نمایی خواهیم کرد. سپس به اثبات تعدادی از ویژگی‌های آن در قابلیت اعتماد خواهیم پرداخت. برآورد حداکثر درستنمایی پارامتر تنش-مقاومت در توزیع پارتو نمایی را به دست می‌آوریم. در نهایت به کمک یک مثال کاربردی به مقایسه قابلیت مدل معرفی شده در برازش داده‌های واقعی با سایر مدل‌های رقیب خواهیم پرداخت. واژه‌های کلیدی: قابلیت اعتماد، تابع نرخ خطر، برآورد حداکثر درستنمایی، توزیع پارتو نمایی. کد موضوع بندی: 39B82، 34K20، 39B52، 47A55.

۱ پیش‌گفتار

بسیاری از آماردانان علاقه‌مند هستند تا با اضافه کردن پارامتر به مدل، انعطاف پذیری توزیع‌های آماری را افزایش دهند. برای انجام این امر روش‌های متفاوتی مانند روش ترکیبی (توزیع‌های مرکب) (بارتو-سوزا و همکاران [۶]، روزگار و ناداراجا [۱۷]، طهماسبی و جعفری [۱۸])، خانواده مارشال و الکین (مارشال و الکین [۱۳])، خانواده توزیع‌های ارتجاعی (گوپتا و کوندا [۱۰])، خانواده توزیع‌های بتا (اوجن و همکاران [۹]) وجود دارد. با استفاده از روش‌های مذکور می‌توان تعمیم‌هایی برای توزیع‌های آماری ارائه داد. نتایج مربوط به این کار در زمینه‌های کاربردی مانند مدل‌سازی داده‌ها نتایج مطلوبی ارائه می‌دهند. تابع توزیع یک کلاس از این تعمیم‌ها، به صورت زیر می‌باشد

$$F(x) = \int_a^{\frac{1}{1-F^\#(x)}} f(y) dy \quad (1)$$

که در آن $f(x)$ تابع چگالی و $F^\#(x)$ تابع توزیع هر متغیر تصادفی دلخواه باشند. اگر $f(\cdot)$ تابع چگالی توزیع نمایی و $F^\#(\cdot)$ تابع توزیع پارتو باشد، آنگاه $F(x)$ دارای توزیع پارتو نمایی می‌باشد. (آل-کادیم و همکاران [۳]).

در بخش ۲ به معرفی توزیع پارتو نمایی و ویژگی‌های آن می‌پردازیم. در بخش ۳ رفتار توابع قابلیت اعتماد مانند توابع نرخ خطر، نرخ خطر معکوس، میانگین باقی‌مانده عمر و میانگین گذشته عمر را برای این توزیع به دست می‌آوریم. در بخش ۴ برآورد پارامتر تنش-مقاومت در توزیع پارتو نمایی را مطالعه می‌کنیم. در بخش ۵ به کمک یک مثال کاربردی، توانایی این مدل در تحلیل داده‌های واقعی در مقایسه با سایر مدل‌های رقیب را مورد بررسی قرار می‌دهیم.

۲ معرفی توزیع پارتو نمایی

گوییم متغیر تصادفی X دارای توزیع پارتو نمایی است، اگر تابع توزیع احتمال آن به صورت

$$F(x) = \int_0^{\sqrt[p]{1-F^\sharp(x)}} f(y) dy \quad (۲)$$

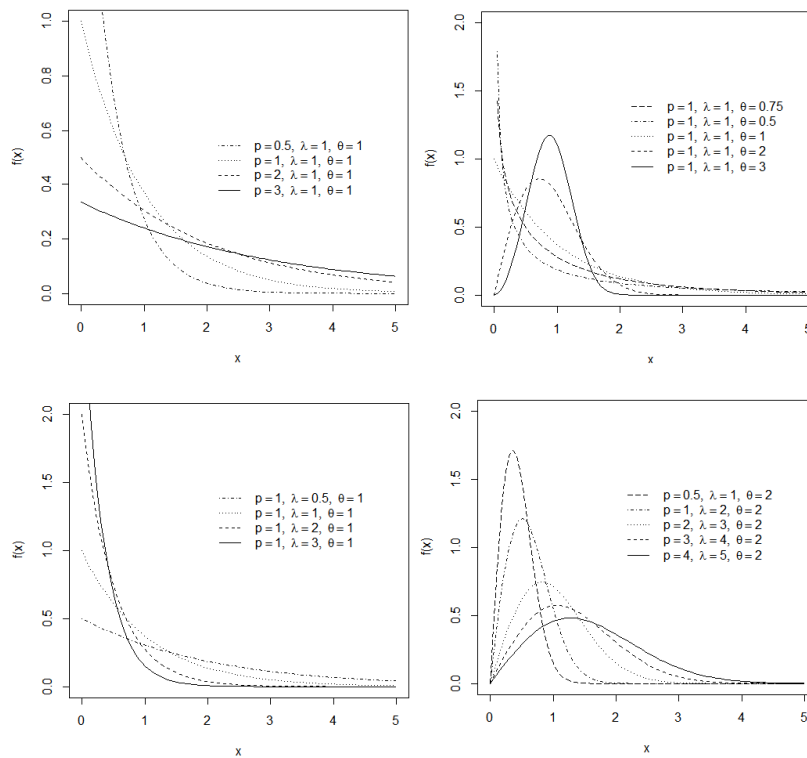
که در آن $F^\sharp(x)$ تابع توزیع پارتو و $f(y)$ تابع چگالی نمایی است، باشد. با جایگذاری تابع چگالی نمایی و تابع توزیع پارتو در رابطه (۲) داریم:

$$F(x) = \int_0^{\sqrt[p]{1-(\frac{x}{p})^\theta}} \lambda e^{-\lambda y} dy = \int_0^{(\frac{x}{p})^\theta} \lambda e^{-\lambda y} dy = 1 - e^{-\lambda(\frac{x}{p})^\theta} \quad (۳)$$

با مشتق‌گیری از رابطه (۳)، تابع چگالی پارتو نمایی به صورت زیر به دست می‌آید:

$$f(x; p, \lambda, \theta) = \frac{\lambda \theta}{p} \left(\frac{x}{p}\right)^{\theta-1} e^{-\lambda(\frac{x}{p})^\theta} \quad x > 0$$

به طور خلاصه، توزیع پارتو نمایی با پارامترهای p ، λ و θ را با نماد $EPD(p, \lambda, \theta)$ نشان می‌دهیم. شکل (۱) نمودار تابع چگالی احتمال توزیع پارتو نمایی به ازای مقادیر مختلف پارامترها را نشان می‌دهد.



شکل ۱: نمودار تابع چگالی احتمال توزیع پارتو نمایی به ازای مقادیر مختلف پارامترها

گشتاورهای n -ام متغیر تصادفی X ، به صورت زیر به دست می‌آید

$$E(X^n) = \left(\frac{p}{\lambda^{\frac{1}{\theta}}}\right)^n \Gamma\left(\frac{n}{\theta} + 1\right).$$

بنابراین واریانس آن به صورت زیر می‌باشد

$$\begin{aligned} Var(X) &= E(X^2) - (E(X))^2 = \left(\frac{p}{\lambda^{\frac{1}{\theta}}}\right)^2 \Gamma\left(\frac{2}{\theta} + 1\right) - \left[\left(\frac{p}{\lambda^{\frac{1}{\theta}}}\right) \Gamma\left(\frac{1}{\theta} + 1\right)\right]^2 \\ &= \left(\frac{p}{\lambda^{\frac{1}{\theta}}}\right)^2 \left\{ \Gamma\left(\frac{2}{\theta} + 1\right) - \left[\Gamma\left(\frac{1}{\theta} + 1\right)\right]^2 \right\}. \end{aligned}$$

۳ رفتار توابع قابلیت اعتماد در توزیع پارتو نمایی

در این بخش به بررسی رفتار توابع نرخ خطر، نرخ خطر معکوس، میانگین باقی‌مانده عمر، میانگین گذشته عمر در توزیع پارتو نمایی می‌پردازیم. قبل از بیان این مطالب ابتدا به تعریف لگ مقعر بودن تابع چگالی احتمال و لم (۳) و (۴) که در ادامه مطالب این بخش مورد استفاده قرار می‌گیرد، می‌پردازیم.

یکی از مفاهیم آماری که با مفاهیم سالخوردگی در قابلیت اعتماد مرتبط است، لگ مقعر بودن تابع چگالی احتمال است. رویدن [۱۶] به بیان لگ مقعر بودن یک تابع حقیقی پرداخته است. تابع غیر منفی $g: \mathbb{R} \rightarrow \mathbb{R}$ ، لگ مقعر است، اگر به ازای هر $x, y \in \mathbb{R}$ و $\alpha \in (0, 1)$ داشته باشیم

$$g(\alpha x + (1 - \alpha)y) \geq [g(x)]^\alpha [g(y)]^{1-\alpha}.$$

اگر به ازای هر $x \in \mathbb{R}$ ، تابع $g(x) > 0$ باشد، آنگاه تعریف لگ مقعر، معادل است با

$$Lng(\alpha x + (1 - \alpha)y) \geq \alpha Lng(x) + (1 - \alpha)Lng(y). \quad (۴)$$

به عبارت دیگر، در صورتی که $Lng(x)$ دو بار مشتق‌پذیر باشد و $-\frac{d^2}{dx^2} Lng(x) \geq 0$ برقرار باشد، شرط تعریف لگ مقعر بودن به صورت رابطه (۴) ساده می‌شود. فرض کنید X متغیر تصادفی طول عمر باشد، اگر X دارای نرخ خطر صعودی باشد، آنگاه X دارای نرخ خطر معکوس و میانگین عمر باقی‌مانده نزولی و میانگین گذشته عمر صعودی است. (می [۱۴])

فرض کنید X یک متغیر تصادفی با تابع چگالی $f(x)$ باشد. اگر $f(x)$ لگ مقعر باشد، آنگاه X دارای نرخ خطر صعودی است. (آن [۴]، بگنولی و برگستر [۵]) اگر توابع قابلیت اعتماد، نرخ خطر، نرخ خطر معکوس، میانگین باقی‌مانده عمر، میانگین گذشته عمر توزیع پارتو نمایی را به ترتیب با $\bar{F}(\cdot)$ ، $h(\cdot)$ ، $h^*(\cdot)$ ، $m(\cdot)$ و $m^*(\cdot)$ نشان دهیم، آنگاه

$$\bar{F}(x) = 1 - F(x) = 1 - \left[1 - e^{-\lambda\left(\frac{x}{p}\right)^\theta}\right] = e^{-\lambda\left(\frac{x}{p}\right)^\theta}$$

$$h(x) = \frac{f(x)}{\bar{F}(x)} = \frac{\frac{\lambda\theta}{p}\left(\frac{x}{p}\right)^{\theta-1} e^{-\lambda\left(\frac{x}{p}\right)^\theta}}{e^{-\lambda\left(\frac{x}{p}\right)^\theta}} = \frac{\lambda\theta}{p}\left(\frac{x}{p}\right)^{\theta-1}$$

$$h^*(x) = \frac{f(x)}{F(x)} = \frac{\frac{\lambda\theta}{p}\left(\frac{x}{p}\right)^{\theta-1} e^{-\lambda\left(\frac{x}{p}\right)^\theta}}{1 - e^{-\lambda\left(\frac{x}{p}\right)^\theta}}$$

$$m(x) = \frac{\int_x^\infty e^{-\lambda(\frac{u}{p})^\theta} du}{e^{-\lambda(\frac{x}{p})^\theta}} = \frac{\int_x^\infty \sum_{j=0}^\infty \frac{(-1)^j}{j!} (\frac{u}{p})^{j\theta}}{\sum_{j=0}^\infty \frac{(-1)^j}{j!} (\frac{x}{p})^{j\theta}}$$

$$m^*(x) = \frac{\int_x^\infty (1 - e^{-\lambda(\frac{u}{p})^\theta}) du}{1 - e^{-\lambda(\frac{x}{p})^\theta}}$$

تابع میانگین باقی‌مانده عمر و تابع میانگین گذشته عمر توزیع پارتو نمایی، فرم بسته ای ندارند. در نتیجه بررسی مستقیم این توابع به آسانی امکان پذیر نیست. در ادامه با توجه به روابطی که بین توابع نرخ خطر و نرخ خطر معکوس و میانگین باقی‌مانده عمر و میانگین گذشته عمر وجود دارد، رفتار این دو تابع را به طور غیرمستقیم بررسی می‌کنیم.

همانگونه که مشخص است رفتار تابع نرخ خطر تنها به پارامتر θ بستگی دارد و پارامترهای دیگر در آن نقشی ندارند. به منظور مطالعه تابع نرخ خطر رفتار مشتق آن را مورد بررسی قرار می‌دهیم. در سه حالت $\theta = 1$ و $0 < \theta < 1$ و $\theta > 1$ رفتار تابع نرخ خطر متفاوت است.

رفتار تابع نرخ خطر توزیع پارتو نمایی به صورت زیر است:

(الف) تابع نرخ خطر به ازای $\theta > 1$ ، صعودی است.

(ب) تابع نرخ خطر به ازای $0 < \theta < 1$ ، نزولی است.

(ج) تابع نرخ خطر به ازای $\theta = 1$ ، ثابت است.

اثبات. [برهان الف.] مشتق تابع $h(x)$ به صورت زیر است:

$$h'(x) = \frac{\lambda\theta(\theta - 1)}{p^\theta} x^{\theta-2}$$

می‌باشد. بدیهی است که به ازای $x > 0$ و $\theta > 1$ ، $h'(x) > 0$ ، می‌باشد. بنابراین در این حالت تابع نرخ خطر صعودی است. اثبات. [برهان ب.] $h'(x) > 0$ به ازای $x > 0$ و $0 < \theta < 1$ ، همواره منفی می‌باشد. لذا تابع نرخ خطر نزولی در حالت $0 < \theta < 1$ ، است. اثبات. [برهان ج.] تابع نرخ خطر به ازای $\theta = 1$ برابر $h(x) = \frac{\lambda}{p}$ است. بنابراین در حالت $\theta = 1$ ، تابع نرخ خطر، ثابت می‌باشد. $\square$

تابع چگالی توزیع پارتو نمایی، لگ مقعر است. اثبات. لگاریتم نسبت درست‌نمایی تابع چگالی پارتو نمایی به صورت

$$\psi(x; p, \lambda, \theta) = \ln f(x; p, \lambda, \theta) = \ln\left(\frac{\lambda\theta}{p}\right) + (\theta - 1)\ln\left(\frac{x}{p}\right) - \lambda\left(\frac{x}{p}\right)^\theta$$

می‌باشد. با مشتق‌گیری از این تابع داریم

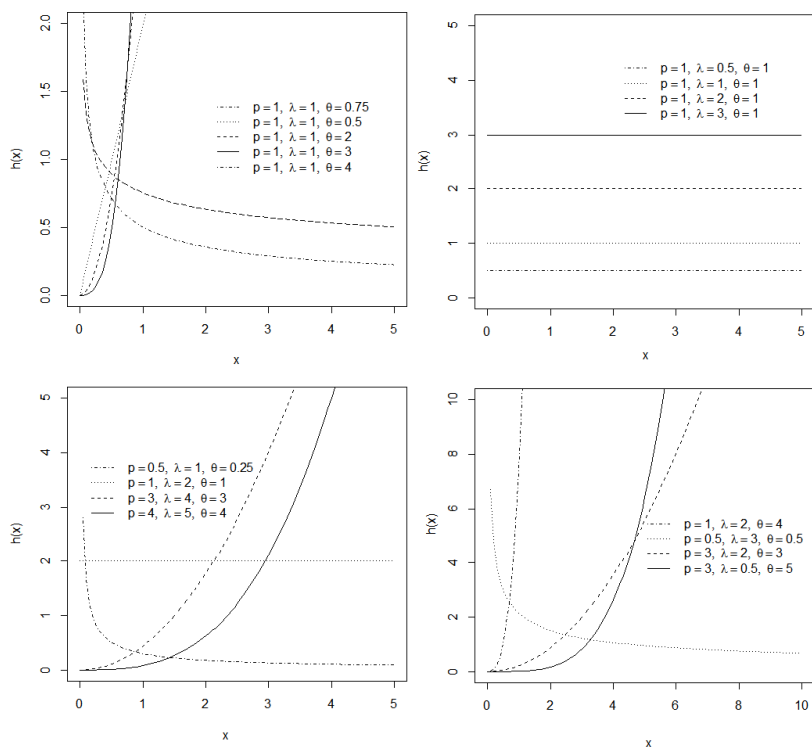
$$\frac{\partial \psi(x; p, \lambda, \theta)}{\partial x} = (\theta - 1)\frac{1}{x} - \frac{\lambda\theta}{p^\theta} x^{\theta-1}$$

$$\frac{\partial^2 \psi(x; p, \lambda, \theta)}{\partial x^2} = \frac{-(\theta - 1)}{x^2} - \frac{\lambda \theta}{p^\theta} (\theta - 1) x^{\theta-2} = \begin{cases} \leq 0, & \theta > 1 \\ > 0, & 0 < \theta < 1 \end{cases}$$

□

بنابراین تابع چگالی پارتو نمایی به ازای $\theta > 1$ لگ مقعر است.

شکل (۲) نمودار تابع نرخ خطر توزیع پارتو نمایی را به ازای مقادیر مختلف پارامترهای آن نشان می‌دهد.



شکل ۲: نمودار تابع نرخ خطر توزیع پارتو نمایی به ازای مقادیر مختلف پارامترها

۴ برآورد پارامتر تنش-مقاومت

فرض کنید X و Y دو متغیر تصادفی دلخواه باشند. همچنین متغیر تصادفی Y میزان فشار وارد بر یک عنصر در یک آزمایش طول عمر و متغیر تصادفی X میزان تحمل آن عنصر را بیان می‌کنند. بدیهی است که یک عنصر، زمانی دچار شکست می‌شود که

$X < Y$ باشد. پارامتر تنش-مقاومت $R = P(X > Y)$ یکی از مهمترین پارامترها در مبحث قابلیت اعتماد است، بنابراین برآورد آن از اهمیت خاصی برخوردار می باشد. برآورد پارامتر R زمانی که متغیرهای تصادفی X و Y مستقل و هم توزیع هستند، مورد مطالعه محققان زیادی قرار گرفته است. چارچ و هریس [۷] و دانتون [۸] فرض کرده اند که X و Y دارای توزیع نرمال هستند و مدل تنش-مقاومت برای آن را به دست آورده اند. کوندا و گوپتا [۱۲] توزیع نمایی تعمیم یافته را بررسی کردند. بررسی به مراتب گسترده تر مدل تنش-مقاومت توسط کاتز و همکاران [۱۱] ارائه شده است. در این بخش به بحث پیرامون پارامتر تنش-مقاومت در توزیع پارتو نمایی می پردازیم.

فرض کنید $X \sim EPD(p_1, \lambda_1, \theta_1)$ و $Y \sim EPD(p_2, \lambda_2, \theta_2)$ متغیرهای تصادفی مستقل باشند. آنگاه

$$\begin{aligned} R &= P(X > Y) = \int_0^{\infty} f_X(x) F_Y(x) dx \\ &= 1 - \sum_{j=0}^{\infty} \frac{(-1)^j}{j!} \lambda_2^j \left(\frac{p_1}{p_2}\right)^{\theta_2 j} \Gamma\left(\frac{\theta_2 j}{\theta_1} + 1\right) \end{aligned} \quad (5)$$

اثبات.

$$\begin{aligned} R &= P(X > Y) = \int_0^{\infty} \left(1 - e^{-\lambda_2 \left(\frac{x}{p_2}\right)^{\theta_2}}\right) \frac{\lambda_1 \theta_1}{p_1} \left(\frac{x}{p_1}\right)^{\theta_1 - 1} e^{-\lambda_1 \left(\frac{x}{p_1}\right)^{\theta_1}} dx \\ &= 1 - \frac{\lambda_1 \theta_1}{p_1^{\theta_1}} \int_0^{\infty} x^{\theta_1 - 1} e^{-\lambda_1 \left(\frac{x}{p_1}\right)^{\theta_1}} e^{-\lambda_2 \left(\frac{x}{p_2}\right)^{\theta_2}} dx \end{aligned}$$

با استفاده از بسط سری توانی $e^{-\theta u} = \sum_{j=0}^{\infty} \frac{(-\theta u)^j}{j!} = \frac{(-1)^j}{j!} \theta^j u^j$ داریم:

$$R = 1 - \frac{\lambda_1 \theta_1}{p_1^{\theta_1}} \sum_{j=0}^{\infty} \frac{(-1)^j}{j!} \frac{\lambda_2^j}{p_2^{\theta_2 j}} \int_0^{\infty} x^{\theta_1 + \theta_2 j - 1} e^{-\lambda_1 \left(\frac{x}{p_1}\right)^{\theta_1}} dx$$

با تغییر متغیر $u = \lambda_1 \left(\frac{x}{p_1}\right)^{\theta_1}$ داریم

$$\begin{aligned} R &= 1 - \sum_{j=0}^{\infty} \frac{(-1)^j}{j!} \frac{\lambda_2^j}{p_2^{\theta_2 j}} \int_0^{\infty} \left(u^{\frac{1}{\theta_1}} p_1\right)^{\theta_2 j} e^{-u} du \\ &= 1 - \sum_{j=0}^{\infty} \frac{(-1)^j}{j!} \frac{\lambda_2^j}{p_2^{\theta_2 j}} p_1^{\theta_2 j} \int_0^{\infty} u^{\frac{\theta_2 j}{\theta_1}} e^{-u} du. \\ &= 1 - \sum_{j=0}^{\infty} \frac{(-1)^j}{j!} \lambda_2^j \left(\frac{p_1}{p_2}\right)^{\theta_2 j} \Gamma\left(\frac{\theta_2 j}{\theta_1} + 1\right). \end{aligned}$$

□

حال فرض کنید $X_1, X_2, \dots, X_n$ و $Y_1, Y_2, \dots, Y_m$ دو نمونه تصادفی مستقل از توزیع های $X \sim EPD(p_1, \lambda_1, \theta_1)$ و $Y \sim EPD(p_2, \lambda_2, \theta_2)$ باشند. تابع درستنمایی به صورت زیر به دست می آید

$$L(\theta) = \lambda_1^n \theta_1^n \frac{\left(\prod_{i=1}^n x_i\right)^{\theta_1 - 1} e^{-\frac{\lambda_1}{\theta_1} \left(\sum_{i=1}^n x_i^{\theta_1}\right)}}{p_1^{n\theta_1}} \lambda_2^m \theta_2^m \frac{\left(\prod_{j=1}^m y_j\right)^{\theta_2 - 1} e^{-\frac{\lambda_2}{\theta_2} \left(\sum_{j=1}^m y_j^{\theta_2}\right)}}{p_2^{m\theta_2}}$$

لگاریتم تابع درستنمایی برای این مشاهدات عبارت است از

$$\begin{aligned} \ell = LnL(\theta) &= nLn\lambda_1 + mLn\lambda_2 + nLn\theta_1 + mLn\theta_2 \\ &+ (\theta_1 - 1) \sum_{i=1}^n Lnx_i + (\theta_2 - 1) \sum_{j=1}^m Lny_j - n\theta_1 Lnp_1 - m\theta_2 Lnp_2 \\ &- \frac{\lambda_1}{p_1} \left(\sum_{i=1}^n x_i^{\theta_1} \right) - \frac{\lambda_2}{p_2} \left(\sum_{j=1}^m y_j^{\theta_2} \right) \end{aligned}$$

برای دستیابی به برآورد حداکثر درستنمایی پارامترهای بردار $\theta = (p_1, p_2, \lambda_1, \lambda_2, \theta_1, \theta_2)$ لازم است معادله غیر خطی $\left(\frac{\partial \ell}{\partial p_1}, \frac{\partial \ell}{\partial p_2}, \frac{\partial \ell}{\partial \lambda_1}, \frac{\partial \ell}{\partial \lambda_2}, \frac{\partial \ell}{\partial \theta_1}, \frac{\partial \ell}{\partial \theta_2} \right)$ حل شود. با جایگزینی برآوردهای حداکثر درستنمایی این پارامترها در معادله (۵) پارامتر تنش-مقاومت نیز برآورد خواهد شد.

۵ مثال کاربردی

در این بخش توانایی برازش مدل پارتو نمایی را بر روی داده‌های واقعی در مقایسه با سایر مدل‌های رقیب نشان می‌دهیم. مجموعه داده‌ای که در این مثال مورد بررسی قرار می‌دهیم شامل زمان شکست متوالی یک سیستم تهویه هوا در ناوگان هواپیمایی بوئینگ ۷۲۰ می‌باشد. این مجموعه داده شامل ۲۱۳ مشاهده است و اولین بار توسط پروشان [۱۵] تحلیل گردید. آدامیدیس و لوکاس [۱] و بارتو-سوزا و همکاران [۶] بررسی‌های بیشتری بر روی این داده‌ها انجام دادند. مقادیر این داده‌ها در جدول ارائه شده‌اند. جدول (۱) خلاصه‌ای از آماره‌های محاسبه شده برای داده‌ها را نشان می‌دهد.

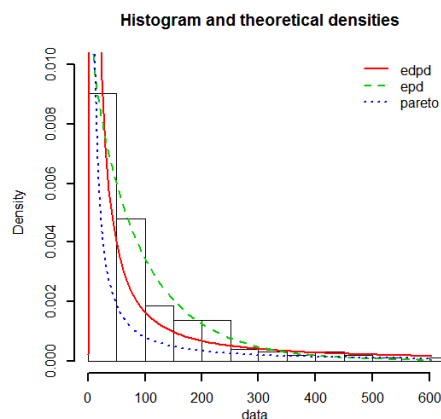
جدول ۱: شاخص‌های آماری داده‌های زمان شکست متوالی سیستم تهویه هوا در ناوگان هواپیمایی بوئینگ ۷۲۰

تعداد (n)	چارک اول	چارک سوم	دامنه	ضریب تغییرات	انحراف معیار	میانه	میانگین
۲۱۳	۲۲	۱۱۸	۶۰۲	۱/۱۴۶	۱۰۶/۷۶۳۶	۵۷	۹۳/۱۴

جدول (۲) برآورد حداکثر درستنمایی پارامترهای مدل‌های مطرح شده و انحراف استاندارد پارامترها و همچنین معیارهای اطلاع آکائیک (AIC)، بیزی (BIC) و لگاریتم تابع درستنمایی را برای مدل پارتو نمایی و دو مدل رقیب پارتو و پارتو نمایی شده (علی و همکاران [۲]) نشان می‌دهد. همانطور که مشاهده می‌کنید مقادیر AIC و BIC برای مدل پارتو نمایی، کمتر از مدل‌های پارتو نمایی شده و پارتو می‌باشد همچنین مقدار $loglike$ برای مدل پارتو نمایی بیشتر از دو مدل دیگر است. شکل (۳) نمودار بافت‌نگار مجموعه داده‌های مربوط به زمان شکست متوالی یک سیستم تهویه هوا در ناوگان هواپیمایی بوئینگ ۷۲۰ به همراه چگالی‌های برازش داده شده به آنها را نشان می‌دهد. در حقیقت شکل (۳) به نوعی نتیجه به دست آمده توسط جدول (۲) را تأیید می‌کند و نشان می‌دهد که توزیع پارتو نمایی برازش بهتری به داده‌ها ارائه کرده است.

جدول ۲: مقادیر MLE و AIC و BIC برای سه مدل پارتو نمایی، پارتو نمایی شده، پارتو (عبارت داخل پرانتز انحراف استاندارد پارامترهاست).

پارتو	پارتو نمایی شده	پارتو نمایی	
$\alpha = 0/256$	$\alpha = 2/125$	$\theta = 0/924$	
(0/017)	(0/044)	(0/408)	
$\beta = 1$	$\beta = 0/997$	$\lambda = 0/01$	برآورد حداکثر درستنمایی
(0/0001)	(0/001)	(0/009)	
	$\gamma = 0/268$	$p = 0/895$	
	(0/018)	(0/702)	
2670/044	2391/038	2361/17	AIC
2673/406	2401/121	2371/254	BIC
- 1334/02	- 1192/519	- 1177/585	$loglike$



شکل ۳: نمودار بافتنگار مجموعه داده‌های مربوط به زمان شکست متوالی یک سیستم تهویه هوا در ناوگان هواپیمایی بوئینگ ۷۲۰ و برازش آن‌ها با توزیع‌های پارتو نمایی، پارتو نمایی شده و پارتو

بحث و نتیجه گیری

در این مقاله به معرفی توزیع پارتو نمایی و بررسی ویژگی‌های آن پرداختیم. همچنین رفتار توابع قابلیت اعتماد این توزیع مورد مطالعه قرار دادیم. توانایی این توزیع را در تحلیل داده‌های واقعی در مقایسه با مدل‌های رقیب مورد بررسی قرار دادیم و دیدیم که توزیع پارتو نمایی نسبت به سایر مدل‌ها برازش بهتری به داده‌ها ارائه کرده است.

- [1] Adamidis, K. and Loukas, S. (1998), A lifetime distribution with decreasing failure rate, *Statistics & Probability Letters*, 39, 35-42.
- [2] Ali, M.M. and Manisha, P. and Jungsoo, W. (2007), Some Exponentiated Distributions, *the Korean Communications in Statistics*, 14, 93-109.
- [3] Al-Kadim, K.R. and Boshi, M.A. (2013), *Exponential Pareto Distribution*, Mathematical Theory and Modeling. Vol.3, No.5, 135-146.
- [4] An, Mark Yuying. (1998), Logconcavity versus logconvexity: a complete characterization, *Journal of economic theory*, 80(2), 350-369.
- [5] Bagnoli, M. and Bergstrom, T. (2005), Log-concave probability and its applications, *Economic theory*, 26(2), 445-469.
- [6] Barreto-Souza, W., de Morais, A.L. and Cordeiro, G.M. (2011), The Weibull-geometric distribution, *Journal of Statistical Computation and Simulation*, 81(5), 645-657.
- [7] Church, J.D. and Harris, B.(1970), The estimation of reliability from stress strength relationship, *Technometrics*, 12, 49-54.
- [8] Downton, F. (1973), The estimation of $\Pr(Y < X)$ in the normal case, *Technometrics*, 15, 551-558.
- [9] Eugene, N., Lee, C and Famoye, F. (2002), Beta-normal distribution and its applications, *Commun Stat Theory Methods*, 31, 497-512.
- [10] Gupta, R.D. and Kunda D. (1999), Generalized exponential distributions, *Australian and New Zealand Journal of Statistics*, 41, 173-188.
- [11] Kotz, S. and Lumelskii, Y. and Pensky, M. (2003), *The Stress-Strength Model and its Generalizations*, theory and applications, World Scientific.
- [12] Kundu, D. and Gupta, R.D. (2005), Estimation of $P[Y < X]$ for generalized exponential distribution, *Metrika*, 61(3), 291-308.
- [13] Marshall, A.N. and Olkin, I. (1997), A new method for adding a parameter to a family of distributions with applications to the exponential and Weibull families, *Biometrika*, 84, 641-652.
- [14] Mi, J. (1995), Bathtub failure rate and upside-down bathtub mean residual life, *IEEE Transactions on Reliability*, 44(3), 388-391.

- [15] Proschan, F. (1963), Theoretical explanation of observed decreasing failure rate, *Technometrics*, **5**, 375-383.
- [16] Royden, H.L. (1968), *Real analysis*. Prentice Hall of India, New Delhi.
- [17] Roozegar, R. and Nadarajah, S. (2017), The Quadratic Hazard Rate Power Series Distribution, *Journal of Testing and Evaluation*, **45**, 1072-1058.
- [18] Tahmasebi, S. and Jafari, A.A. (2016), Generalized Gompertz-Power Series Distributions, *Hacettepe Journal of Mathematics and Statistics*, **45**, 1604-1579.

Seminar Schedule & Abstract Booklet
of
5th Seminar on Reliability Theory
and its Applications

پنجمین سمینار تخصصی
نظریه قابلیت اعتماد
و کاربردهای آن
۲۸ و ۲۹ فروردین ماه ۱۳۹۸



5th Seminar on
Reliability Theory
and its Applications
17-18 April 2019

۰۳۵-۳۸۲۰۹۸۲۶
stat@confs.yazd.ac.ir
http://wosdce.um.ac.ir

نشانی: یزد، صفائییه، دانشگاه یزد، دانشکده علوم ریاضی



Yazd University
17-18 April 2019

