



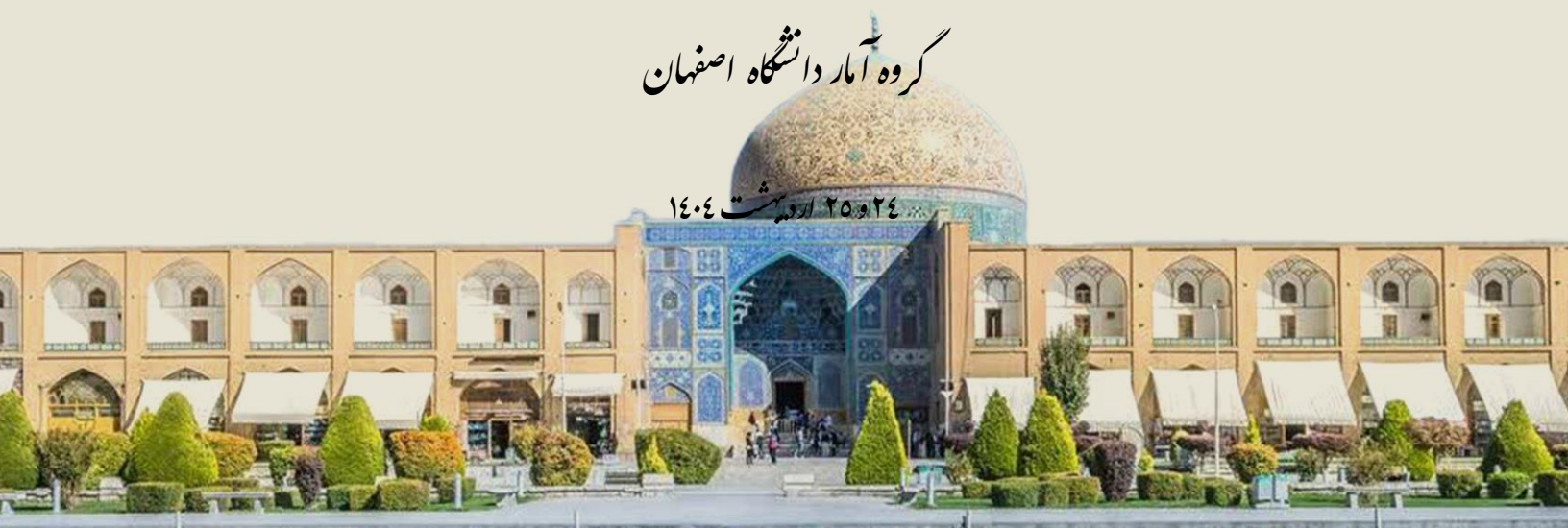
مجموعه مقالات

یازدهمین سمینار تخصصی

نظریه قابلیت اعتماد و کار بردهای آن

گروه آمار دانشگاه اصفهان

۲۵ و ۲۶ شهریور ۱۴۰۴



باسمه تعالی



مجموعه مقالات

یازدهمین سمینار تخصصی

نظریه قابلیت اعتماد و کاربردهای آن

گروه آمار، دانشکده ریاضی و آمار، دانشگاه اصفهان

اصفهان، ایران

۲۴ و ۲۵ اردیبهشت ۱۴۰۴

عنوان: مجموعه مقالات یازدهمین سمینار تخصصی نظریه قابلیت اعتماد و کاربردهای آن

تدوین: سمیه اشرفی

طراح جلد: علی معارفوند

تاریخ انتشار: شهریور ماه ۱۴۰۴

مقدمه

پیرو برگزاری ده دوره موفق سمینار تخصصی نظریه قابلیت اعتماد و کاربردهای آن، خداوند سبحان را شاکریم که توفیق داد تا یازدهمین دوره سمینار را توسط هسته پژوهشی قابلیت اعتماد و کاربردهای آن در دانشگاه اصفهان و با همکاری قطب علمی داده‌های ترتیبی، قابلیت اعتماد و وابستگی، انجمن آمار ایران و پایگاه استنادی علوم جهان اسلام در گروه آمار دانشگاه اصفهان به صورت حضوری برگزار نماییم. این سمینار دو روزه با هدف ایجاد فرصت مناسب برای ارائه جدیدترین دستاوردهای علمی و پژوهشی در شاخه‌های مختلف نظریه قابلیت اعتماد و تبادل نظر شرکت کنندگان روزهای ۲۴ و ۲۵ اردیبهشت در دانشگاه اصفهان برگزار گردید. امیدواریم این سمینار در دستیابی به اهداف خود موفق و موجبات ارتقای این شاخه از علم آمار در ایران را فراهم آورده باشد.

در این سمینار پس از فراخوان، مقالات دریافت شده توسط اعضای کمیته علمی و کمیته داوران مورد ارزیابی قرار گرفت و در نهایت ۳۰ مقاله به صورت سخنرانی و ۴ مقاله به صورت پوستر پذیرفته شد. لازم به ذکر است که در این سمینار ۶۰ نفر از اعضای هیات علمی، پژوهشگران و دانشجویان دانشگاه‌های مختلف کشور حضور دارند. در پایان لازم می‌دانیم از اعضای کمیته‌های علمی، اجرایی و کمیته داوران سمینار تشکر نماییم. همچنین مراتب قدردانی خود را از معاونت محترم پژوهش و فناوری دانشگاه اصفهان، رئیس محترم دانشکده ریاضی و آمار، مدیر محترم و سایر همکاران گروه آمار دانشگاه اصفهان اعلام نماییم.

از خداوند منان، آرزوی توفیق برای تمامی شرکت‌کنندگان محترم در این سمینار را داریم.

دبیر علمی سمینار: مهدی توانگر

دبیر اجرایی سمینار: سمیه اشرفی

محورهای سمینار

قابلیت اعتماد سیستم‌های منسجم و شبکه‌ها	قابلیت اعتماد نرم‌افزار
روش‌های بیزی در قابلیت اعتماد	کنترل کیفیت آماری
الگوهای تعمیر و نگهداری سیستم‌ها	مدل‌های ضمانت و داده‌های میدانی
ترتیب‌های تصادفی	قابلیت اعتماد در محیط فازی
تحلیل قابلیت اعتماد بر اساس داده‌های فرسایشی	مفاهیم سالخورده‌گی
داده‌کاوی در قابلیت اعتماد	مطالعات موردی در صنعت
آزمون‌های طول عمر تسریع‌یافته	مفاهیم وابستگی و توابع مفصل
مدل‌های شوک	ابر داده در قابلیت اعتماد
استنباط آماری و تحلیل داده‌های طول عمر	تحلیل ریسک در قابلیت اعتماد
تحلیل بقا	هوش مصنوعی در مهندسی قابلیت اعتماد

اعضای کمیته علمی (به ترتیب حروف الفبا)

دکتر رضا احمدی، دانشگاه علم و صنعت
دکتر مجید اسدی، دانشگاه اصفهان
دکتر سمیه اشرفی، دانشگاه اصفهان
دکتر اکبر اصغرزاده، دانشگاه مازندران
دکتر ابراهیم امینی سرشت، دانشگاه بوعلی سینا
دکتر حمزه ترابی، دانشگاه یزد
دکتر مهدی توانگر، دانشگاه اصفهان (دبیر علمی سمینار)
دکتر مجید چهکندی، دانشگاه بیرجند
دکتر نوشین حکمی‌پور، مرکز آموزش عالی بوئین زهرا
دکتر فیروزه حقیقی، دانشگاه تهران
دکتر مهدی دوست‌پرست، دانشگاه فردوسی مشهد
دکتر مصطفی رزمخواه، دانشگاه فردوسی مشهد
دکتر سمیه زارع‌زاده، دانشگاه شیراز

دکتر مریم شرفی، دانشگاه رازی کرمانشاه
دکتر فریبا عزیزی، دانشگاه الزهرا
دکتر مریم کلکین‌نما، دانشگاه صنعتی اصفهان
دکتر زهرا منصوروار، دانشگاه اصفهان

اعضای کمیته مشاورین علمی افتخاری (به ترتیب حروف الفبا)

دکتر جعفر احمدی، دانشگاه فردوسی مشهد
دکتر بهالدین خالدی، دانشگاه رازی کرمانشاه
دکتر احمد خدادادی، دانشگاه شهید بهشتی
دکتر اسماعیل خرم، دانشگاه صنعتی امیرکبیر
دکتر محمد خنجری صادق، دانشگاه فردوسی مشهد
دکتر عبدالحمید رضایی رکن‌آبادی، دانشگاه فردوسی مشهد
دکتر علی زینل همدانی، دانشگاه صنعتی اصفهان
دکتر هوشنگ طالبی، دانشگاه اصفهان
دکتر غلامرضا محتشمی‌برزادران، دانشگاه فردوسی مشهد

اعضای کمیته اجرایی (به ترتیب حروف الفبا)

دکتر جعفر احمدی، دانشگاه فردوسی مشهد
دکتر مجید اسدی، دانشگاه اصفهان
دکتر سمیه اشرفی، دانشگاه اصفهان (دبیر اجرایی سمینار)
دکتر مهدی توانگر، دانشگاه اصفهان
دکتر هادی جباری، دانشگاه فردوسی مشهد
دکتر احسان زمان‌زاده، دانشگاه اصفهان
دکتر زهرا منصوروار، دانشگاه اصفهان

کادر اجرایی سمینار

سمیرا نصرافصفهانی (دکتری آمار)، هبه عمار (دانشجوی دکتری)، سمیه اصلانی (دانشجوی کارشناسی ارشد)، صدیقه حاجی علی
اکبر (دانشجوی کارشناسی ارشد)

فهرست

- سیستم چه مدتی باید تحت آب‌بندی باشد؟
آزموده نژاد، م.، احمدی، ج. ۱
- تعداد مولفه‌های شکست خورده در یک سیستم منسجم فعال
اعلم یزدیان، ح. ۱۳
- مدلسازی فرسایش سیستم تعمیر پذیر تحت فرآیند گاما و برآورد پارامترهای آن
ایرانمنش، ف.، چهکندی، م.، اطمینان، ج. ۲۶
- رویکرد مبتنی بر یادگیری تقویتی عمیق برای برنامه‌ریزی پویا در تعمیر و نگهداری سیستم‌های تخریب‌پذیر
حسن‌تبار درزی، ف.، اسلامی مهدی‌آبادی، م. ۴۰
- نگهداری و تعمیر هوشمند با استفاده از یادگیری تقویتی عمیق
حقیقی، ف.، صفری، ع.، عیدی، ش. ۵۲
- محاسبه طول عمر باقی‌مانده در نرم‌افزار R
طباطبائی شیرازی، س.ا.، عمادی، م.، آرشی، م. ۶۰
- آزمون طول عمر شتابیده استرس پله‌ای ساده تحت سانسور فزاینده نوع دوم با مکانیسم خروج تصادفی وابسته
طیبی گیلانی، ن.، حقیقی، ف.، صفری، ع. ۷۸
- اثبات حفظ ویژگی IFR در آماره‌های ترتیبی گسسته به روشی جدید
علی‌محمدی، م.، قره‌باغی، ر. ۱۰۳

تحلیل خطرهای رقابتی در مسائل زیستی

هاشم آبادی، ع.، رزمخواه، م. ۱۱۱



یازدهمین سمینار تخصصی
نظریه قابلیت اعتماد و کاربردهای آن
۲۴ و ۲۵ اردیبهشت ۱۴۰۴



سیستم چه مدتی باید تحت آب‌بندی باشد؟

آزموده نژاد، م^۱ احمدی، ج^۲

گروه آمار، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد

چکیده

در این مقاله، مسئله تعیین زمان بهینه آب‌بندی در دوره‌های آب‌بندی و گارانتی مورد بررسی قرار گرفته و یک مدل ریاضی برای مینیمم‌سازی هزینه کلی سیستم ارائه شده است. در این مدل برای ساختن تابع هدف، هزینه‌های شامل آب‌بندی و گارانتی و میزان فرسایش قطعات به عنوان معیارهای اصلی به‌کار رفته‌اند. در نهایت نتایج بدست آمده، از طریق تحلیل حساسیت و بررسی تغییرات پارامترهای کلیدی بررسی شده است.

کلمات کلیدی: آب‌بندی، گارانتی، فرسایش، نرخ خطر

^۱ma.azmoudeh@gmail.com

^۲ahmadi-j@um.ac.ir

یکی از روش‌های مؤثر برای اطمینان از عملکرد مناسب یک محصول و حذف قطعات معیوب، فرآیند آب‌بندی (*Burn-in*) است. در این روش، محصول یا سیستم برای مدت مشخصی تحت شرایط کنترل‌شده قرار می‌گیرد تا نقص‌های احتمالی که ممکن است در مراحل اولیه عملکرد ظاهر شوند، شناسایی و برطرف شوند. تولیدکنندگان همواره تلاش می‌کنند تا خرابی‌های زود هنگام محصولات خود را کاهش دهند و عمر مفید آن‌ها را افزایش دهند. این مسئله نه تنها از دیدگاه اقتصادی و کاهش هزینه‌های گارانتی برای کارخانه‌ها اهمیت دارد، بلکه از منظر مشتریان نیز نقش کلیدی ایفا می‌کند. محصولاتی که در مدت کوتاهی پس از استفاده دچار نقص یا خرابی می‌شوند، موجب نارضایتی مشتریان و کاهش اعتبار برند تولیدکننده خواهند شد. از سوی دیگر، افزایش قابلیت اعتماد محصولات منجر به کاهش هزینه‌های خدمات پس از فروش و افزایش وفاداری مشتریان می‌شود.

فرآیند آب‌بندی عمدتاً برای سیستم‌هایی به کار می‌رود که نرخ خطر نزولی دارند، به این معنا که پس از حذف خرابی‌های اولیه، احتمال خرابی آن‌ها کاهش می‌یابد و وارد فاز عمر مفید خود می‌شوند. این ویژگی معمولاً در محصولاتی دیده می‌شود که در مرحله اولیه استفاده، به دلیل وجود نقص‌های ساخت، مشکلات طراحی یا خطاهای تولید، احتمال خرابی بالایی دارند. اما پس از گذر از این مرحله و حذف قطعات معیوب، سیستم پایداری بیشتری پیدا کرده و احتمال خرابی آن در آینده کاهش می‌یابد. بنابراین، آب‌بندی یک راهکار مؤثر برای افزایش قابلیت اعتماد محصولاتی است که رفتار خرابی اولیه قابل توجهی دارند.

اجرای موفقیت‌آمیز آب‌بندی نیازمند تعیین شرایط بهینه‌ای است که در آن، مدت زمان آزمایش به گونه‌ای انتخاب شود که بیشترین خرابی‌های زود هنگام را حذف کند، اما در عین حال، زمان اضافی برای محصولاتی که احتمال خرابی آن‌ها کم است، صرف نشود. اگر مدت زمان آب‌بندی بیش از حد کوتاه باشد، ممکن است برخی از قطعات معیوب شناسایی نشوند و وارد بازار شوند، که این امر منجر به افزایش نرخ خرابی محصولات در دست مشتریان خواهد شد. از سوی دیگر، اگر مدت زمان آب‌بندی بیش از حد طولانی باشد، هزینه‌های تولید افزایش می‌یابد و از طرفی ممکن است به برخی از قطعات آسیب وارد شود. بنابراین، یافتن مدت زمان بهینه آب‌بندی که بتواند تعادلی بین افزایش قابلیت اعتماد و کاهش هزینه‌های مرتبط برقرار کند، یکی از مسائل کلیدی در طراحی فرآیندهای تولیدی است. هدف از این فرآیند این است که با استفاده از یک بازه زمانی مناسب، نرخ خرابی محصولات در بازار را کاهش داده و از ورود قطعات معیوب به مرحله استفاده جلوگیری شود.

هنگام استفاده از روش آب‌بندی، علاوه بر طول دوره آب‌بندی، نوع تعمیرات در این دوره عوامل مهمی برای مدیریت هزینه‌ها از منظر تولیدکننده است که باید به آن‌ها پرداخته شود. کوون و همکاران [۱]، با در نظر گرفتن سیاست تعویض

بلوکی و تعمیر مینیمال، مدلی بهینه برای تعیین زمان آب‌بندی و فواصل بهینه تعمیر و تعویض قطعات ارائه کرده‌اند. در این مدل، ابتدا قطعات تحت فرآیند آب‌بندی قرار می‌گیرند تا خرابی‌های زود هنگام آن‌ها حذف شود. در صورت خرابی در این مرحله، قطعه به روش مینیمال تعمیر شده و فرآیند آب‌بندی ادامه می‌یابد. پس از پایان این مرحله، قطعه وارد فاز بهره‌برداری می‌شود و طبق سیاست تعویض بلوکی، در فواصل از پیش تعیین شده تعویض می‌گردد، در حالی که بین این تعویض‌های برنامه‌ریزی شده، تعمیرات مینیمال در صورت خرابی انجام می‌شود. مقیمی حاجی [۴]، به بهینه‌سازی طول دوره‌های آب‌بندی و گارانتی از دیدگاه تولیدکننده پرداخته است. او با استفاده از یک مدل ریاضی، هزینه‌های کل را در طول دوره‌های آب‌بندی و گارانتی محاسبه کرده است. در دوره آب‌بندی، خرابی‌های جزئی با تعمیرات مینیمال و خرابی‌های عمده با تعویض کامل محصول، اصلاح می‌شوند. در دوره گارانتی، تعمیرات مینیمال در محل مشتری یا تعمیرات کلی با هزینه‌های مختلف انجام می‌گیرد. هدف اصلی تحقیق، تعیین ترکیب بهینه‌ای از طول این دوره‌ها برای کاهش هزینه‌ها و افزایش رضایت مشتری است.

از آنجایی که در بسیاری از کاربردهای عملی، سیستم‌ها و تجهیزات یکسان نیستند، می‌توان آن‌ها را بر اساس ویژگی‌های قابلیت اطمینان به زیرجمعیت‌های مختلف طبقه‌بندی کرد. در چنین مواردی، در طراحی استراتژی آب‌بندی و نگهداری بهینه، در نظر گرفتن ناهمگونی جمعیت مساله مهمی است. لینگ و همکاران [۳] یک سیاست آب‌بندی را برای اجزایی از یک جمعیت ناهمگن که از n زیرجمعیت مرتب‌شده تشکیل شده‌اند، پیشنهاد می‌دهند. آن‌ها فرض می‌کنند که در صورت وقوع خرابی، اجزا به صورت مینیمال تعمیر می‌شوند. در این رویکرد، تعداد کل تعمیرات مینیمال مشاهده‌شده در طول فرآیند آب‌بندی به عنوان یک قاعده غربالگری برای شناسایی و حذف اجزای معیوب با عملکرد ضعیف مورد استفاده قرار می‌گیرد. ابتدا، این پژوهش با انجام تحلیل احتمالاتی، اثربخشی سیاست غربالگری را نشان می‌دهد. سپس، تنظیمات بهینه آب‌بندی به منظور به حداقل رساندن هزینه کل میانگین یا حداکثر کردن احتمال موفقیت در دوره مأموریت تعیین می‌شود. صفائی و تقی پور [۵] یک مدل یکپارچه مبتنی بر فرسایش برای بهینه‌سازی فرآیند آب‌بندی و نگهداری پیشگیرانه در اقلام ناهمگن و بسیار قابل اعتماد ارائه می‌دهد. با توجه به ناکارآمدی روش‌های سنتی مبتنی بر عمر در محصولات با قابلیت اطمینان بالا، این مدل از فرآیند فرسایش گاما برای ارزیابی سطح فرسایش و شناسایی واحدهای معیوب استفاده می‌کند. چهار پارامتر تصمیم‌گیری شامل مدت زمان آب‌بندی، نرخ استفاده، آستانه فرسایش و زمان‌بندی نگهداری در این مدل بهینه‌سازی می‌شوند. هدف، کاهش هزینه‌ها و افزایش دسترسی سیستم با رعایت محدودیت‌های ایمنی است.

در این مقاله تابع هزینه پیشنهادی در دوره‌های آب‌بندی و گارانتی ارائه شده و هدف اصلی این مقاله، تعیین زمان بهینه آب‌بندی با در نظر گرفتن محدودیت فرسایش است. ادامه مقاله به صورت زیر تنظیم شده است. در بخش ۲، نمادهای استفاده شده به طور مختصر معرفی و مدل مسئله شرح داده می‌شود. بخش ۳، توابع هدف شامل تابع هزینه و

تابع فرسایش مورد استفاده، معرفی می‌شود. بخش ۴، بررسی مسئله بهینه سازی بر اساس توابع هدف ارائه می‌شود. در بخش ۵، مسئله بهینه سازی مدل پیشنهادی و تحلیل حساسیت به صورت عددی انجام می‌شود. در انتها خلاصه‌ای از جمع‌بندی ارائه شده است.

۲ توصیف مدل

در اینجا نمادها و علائم استفاده شده در ادامه مقاله معرفی می‌شوند.

b : مدت زمان آب‌بندی

w : مدت زمان گارانتی

$C_{burn-in}(b)$: هزینه‌های مربوط به فرآیند آب‌بندی

$C_{warranty}(w)$: هزینه‌های مربوط به گارانتی

C_f : هزینه راه‌اندازی آب‌بندی

$C_b(b)$: هزینه عملیاتی آب‌بندی

C : هزینه عملیاتی آب‌بندی به‌ازای واحد زمان

C_r : هزینه تعمیر هر خرابی در آب‌بندی

$E[N_b(b)]$: متوسط تعداد خرابی‌ها در آب‌بندی

C_w : هزینه تعمیر هر خرابی در گارانتی

$E[N_w(w)]$: متوسط تعداد خرابی‌ها در گارانتی

$h_b(t)$: نرخ خرابی در آب‌بندی

$h_w(t)$: نرخ خرابی در گارانتی

γ : پارامتر مقیاس نرخ خرابی توزیع وایبل در آب‌بندی

β : پارامتر شکل نرخ خرابی توزیع وایبل در آب‌بندی

u : نرخ خرابی اولیه در گارانتی

s : نرخ خرابی افزایشی بر حسب زمان در گارانتی

α : ضریب کاهش نرخ خرابی بر اثر آب‌بندی

ϵ : نرخ فرسایش

θ : آستانه فرسایش

$D(b)$: تابع فرسایش در آب‌بندی

b_{deg} : حداکثر زمان مجاز آب‌بندی

b_{cost} : مدت زمان بهینه آب‌بندی از نظر هزینه

b^* : مدت زمان بهینه آب‌بندی با محدودیت فرسایش

در این مقاله، سیستمی را در نظر می‌گیریم که دارای یک دوره آب‌بندی است و پس از آن، تحت شرایط گارانتی قرار می‌گیرد. با استفاده از توابع هدف، مقدار بهینه b^* به گونه‌ای تعیین می‌شود که هزینه کلی سیستم مینیمم شود و همزمان محدودیت‌های فرسایش رعایت شود. برای حل این مسئله، ابتدا مدل‌سازی ریاضی هزینه‌ها و نرخ خرابی انجام شده و سپس با استفاده از تحلیل حساسیت، تاثیر پارامترهای کلیدی بر مقدار بهینه b^* بررسی خواهد شد.

۳ توابع هدف

در مسائل بهینه‌سازی، توابع هدف‌های مختلفی توسط پژوهشگران برحسب ماهیت مسئله تعریف شده است. در اینجا در ساختن تابع هدف، از توابع هزینه و فرسایش استفاده می‌کنیم.

۱.۳ تابع هزینه

در هر مسئله بهینه‌سازی باید تابع هدف کلی داشته باشیم، در اینجا $C_{total}(b)$ به عنوان تابع هدف در نظر گرفته شده است. هدف، یافتن مقدار بهینه است که هزینه کل را حداقل کند. این بهینه‌سازی با در نظر گرفتن اثرات متقابل هزینه‌های آب‌بندی، هزینه‌های گارانتی و نرخ خرابی می‌تواند انجام شود، به طوری که هزینه نهایی سیستم بهینه شود و از لحاظ اقتصادی مقرون به صرفه باشد.

تابع هزینه کلی به صورت

$$C_{total}(b) = C_{burn-in}(b) + C_{warranty}(w), \quad (1)$$

در نظر می‌گیریم. بر اساس مطالعات یه و همکاران [۶]، هزینه‌های آب‌بندی شامل هزینه‌های ثابت و متغیر است. این هزینه‌ها به صورت

$$C_{burn-in}(b) = C_f + C_b(b) + C_r E[N_b(b)], \quad (2)$$

مدل‌سازی شده‌اند. که در آن $C_b(b)$ به صورت خطی بیان می‌شود

$$C_b(b) = C \cdot b.$$

هزینه‌های گارانتی نیز شامل هزینه‌های تعمیر خرابی‌هایی است که پس از دوره آب‌بندی و در طول دوره گارانتی رخ می‌دهند. این هزینه‌ها به صورت زیر بیان می‌شوند

$$C_{warranty}(w) = C_w E[N_w(w)]. \quad (3)$$

برای بدست آوردن توابع $E[N_b(b)]$ و $E[N_w(w)]$ ، نیاز به توابع نرخ خرابی در دوره‌های آب‌بندی و گارانتی داریم. در اینجا فرض می‌کنیم نرخ خرابی در دوره آب‌بندی به صورت زیر باشد

$$h_b(t) = \gamma t^{\beta-1}, \quad 0 < \beta < 1. \quad (4)$$

در تابع (۴)، نرخ خرابی تابع نزولی بر حسب t است. همچنین فرض کنید نرخ خرابی در دوره گارانتی تابع خطی از t و به صورت زیر باشد

$$h_w(t) = u + st.$$

برای بررسی تاثیر آب‌بندی بر نرخ خرابی در دوره گارانتی از مدل خطی-نمایی به صورت

$$h_w(t|b) = ue^{-\alpha b} + st, \quad (5)$$

استفاده می‌کنیم.

با استفاده از (۴) و (۵)، می‌توانیم متوسط تعداد خرابی‌ها در دوره‌های آب‌بندی و گارانتی را بدست آوریم.

$$\begin{aligned} E[N_b(b)] &= \int_0^b h_b(t) dt \\ &= \frac{\gamma}{\beta} b^\beta, \end{aligned} \quad (6)$$

و

$$\begin{aligned} E[N_w(w)] &= \int_b^{b+w} h_w(t|b) dt \\ &= uwe^{-\alpha b} + swb + \frac{1}{\gamma} sw^2. \end{aligned} \quad (7)$$

با جایگذاری روابط (۶) و (۷) در (۱)، تابع هزینه کل برابر است با

$$C_{total}(b) = C_f + (C_r + sw)b + C_r \frac{\gamma}{\beta} b^\beta + C_w uwe^{-\alpha b} + \frac{1}{\gamma} sw^2. \quad (8)$$

۲.۳ تابع فرسایش

در طول فرآیند آببندی، ممکن است قطعات تحت فشار زیاد قرار گیرند و دچار فرسایش یا تخریب شوند. پژوهشگران توابع فرسایش‌های مختلفی بر اساس فرآیندهای تصادفی معرفی کرده‌اند. در اینجا برای جلوگیری از این مشکل، تابع فرسایش زیر را در نظر می‌گیریم

$$D(b) = 1 - e^{-\epsilon b}, \quad b > 0, \epsilon > 0.$$

حال یک آستانه فرسایش، θ ، در نظر می‌گیریم و شرط می‌گذاریم که اگر

$$D(b) > \theta,$$

باشد، آنگاه قطعه از فرآیند آببندی خارج شود.

بنابراین حداکثر مقدار زمان مجاز آببندی به صورت زیر بدست می‌آید

$$b_{cost} \leq b_{deg} = -\frac{1}{\epsilon} \ln(1 - \theta).$$

۴ بهینه‌سازی

برای بهینه‌سازی تابع هزینه با در نظر گرفتن محدودیت فرسایش، با توجه به اینکه تابع هزینه در (۸) نسبت به b مشتق پذیر است، بنابراین می‌توانیم از روش‌های مشتق‌گیری استفاده کنیم. اگر مقدار بهینه b_{cost} وجود داشته باشد، آنگاه باید در معادله $\frac{\partial}{\partial b} C_{total}(b) = 0$ صدق کند. از (۸) داریم

$$\frac{\partial}{\partial b} C_{total}(b) = C_0 + sw + \gamma C_r b^{\beta-1} - \alpha C_w u w e^{-\alpha b}.$$

اگر قرار دهیم

$$g_1(b) = C_0 + sw + \gamma C_r b^{\beta-1}, \quad (1)$$

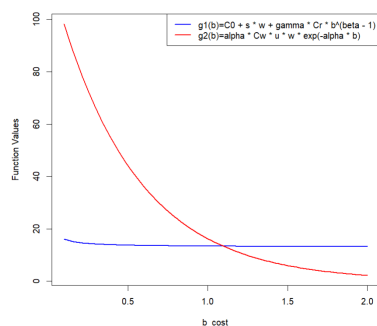
و

$$g_2(b) = \alpha C_w u w e^{-\alpha b}. \quad (2)$$

در این صورت مقدار b_{cost} ریشه معادله $g_1(b) - g_2(b) = 0$ است. پس کافی است محل تقاطع نمودار تابع $g_1(b)$ را با تابع $g_2(b)$ بررسی کنیم. برای این منظور باید مقادیر پارامترهای مدل داده شده باشند. ابتدا پارامترها را به صورت

$$\{C_0 = 4, s = 3, \gamma = 0.2, C_r = 2, \beta = 0.1, \alpha = 2, C_w = 4, u = 5, w = 3\}$$

در نظر می‌گیریم. سپس نمودار دو تابع $g_1(b)$ و $g_2(b)$ را در یک دستگاه مختصات رسم می‌کنیم.



شکل ۱: نمودار دو تابع $g_1(b)$ و $g_2(b)$

از شکل ۱، مشاهده می‌شود که دو نمودار $g_1(b)$ و $g_2(b)$ نقطه تقاطع دارند. با توجه به ۱ و ۲، مشاهده می‌شود که اگر نابرابری

$$C_0 + sw < \alpha uw C_w$$

برقرار باشد، آنگاه جواب مینیم یکتای متناهی برای مدل پیشنهادی داریم. بنابراین زمان بهینه نهایی برای آب‌بندی برابر است با

$$b^* = \min\{b_{cost}, b_{deg}\}.$$

با در نظر گرفتن الگوریتم زیر، مقدار زمان بهینه آب‌بندی b^* محاسبه می‌شود:

- ۱- قطعه وارد مرحله آب‌بندی می‌شود.
- ۲- زمان بهینه آب‌بندی از نظر هزینه محاسبه می‌شود.
- ۳- یک سطح فرسایش در مرحله آب‌بندی، در نظر می‌گیریم و زمان آن را محاسبه می‌کنیم.
- ۴- اگر زمان بهینه آب‌بندی کمتر از زمان فرسایش باشد، زمان بهینه نهایی همان زمان بهینه از نظر هزینه است.
- ۵- در غیر اینصورت زمان بهینه نهایی، زمان فرسایش خواهد بود.

۵ نتایج عددی

در این بخش، محاسبات و تحلیل عددی برای الگوی پیشنهاد شده، ارائه می‌شود تا رفتار زمان بهینه آب‌بندی نسبت به پارامترهای مدل را نشان دهد.

فرض می‌کنیم

$$\{C_o = 4, C_r = 2, C_w = 4, w = 3\}, \quad (1)$$

برای نشان دادن رفتار مقدار بهینه b^* نسبت به پارامترهای مدل، حالت‌های زیر را در نظر می‌گیریم.

• تحلیل حساسیت نسبت به پارامتر u

با فرض (۱)، مقدار u را برای مقادیر مختلف تغییر می‌دهیم تا تاثیر آن بر مقدار b^* بررسی شود.

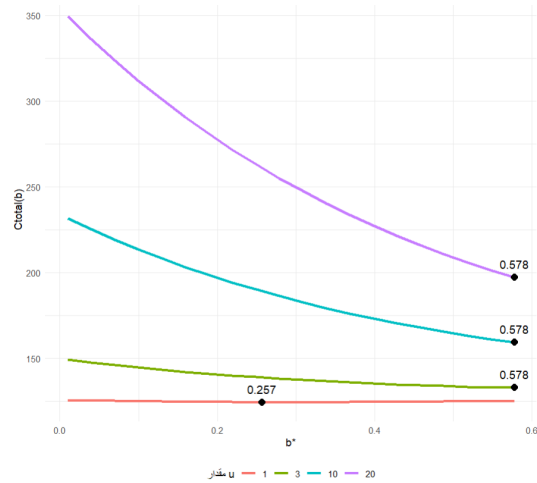
در جدول ۱ و شکل ۲، برای $u = 1$ مقدار بهینه برابر 0.257 است، که از حل معادله $\frac{\partial}{\partial b} C_{total}(b) = 0$ بدست آمده و کمتر از مقدار مجاز فرسایش b_{deg} است.

برای $u = 3, 10, 20$ با افزایش u ، مقدار b_{cost} افزایش می‌یابد. منطقی است زیرا با افزایش نرخ خرابی اولیه دوره گارانتی، هزینه کل سیستم افزایش می‌یابد. از طرفی چون مقدار b_{cost} بیشتر از حد مجاز فرسایش است بنابراین مقدار نهایی b^* برابر 0.578 می‌باشد.

لازم به ذکر است که مقدار b_{cost} از $u = 1$ تا $u = 3$ به شدت افزایش یافته است که این نشان می‌دهد سیستم به شدت نسبت به تغییرات نرخ خرابی اولیه در دوره گارانتی حساس است.

جدول ۱: مقدار بهینه b^* برای پارامتر u

γ	β	α	s	u	b^*
0.2	0.1	2	3	1	0.257
				3	0.578
				10	0.578
				20	0.578



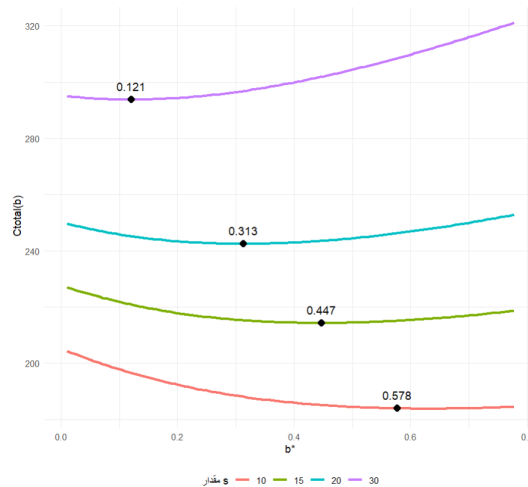
شکل ۲: مقدار بهینه b^* برای پارامتر u

• تحلیل حساسیت نسبت به پارامتر s

با فرض (۱)، مقدار s را برای مقادیر مختلف تغییر می‌دهیم تا تاثیر آن بر مقدار b^* بررسی شود. در جدول ۲ و شکل ۳، برای $s = 10$ مقدار b_{cost} بیشتر از مقدار حد مجاز فرسایش است بنابراین مقدار نهایی b^* برابر $b_{deg} = 0.578$ می‌باشد. برای $s = 15, 20, 30$ با افزایش s مقدار b^* کاهش می‌یابد. این منطقی است زیرا افزایش s به معنای رشد سریع خرابی‌ها و در نتیجه هزینه‌های بالاتر تعمیرات در دوره گارانتی است. برای کاهش این هزینه‌ها، زمان بهینه آب‌بندی b^* باید کوتاه‌تر شود تا قطعات معیوب زودتر حذف شوند.

جدول ۲: مقدار بهینه b^* برای پارامتر s

γ	β	α	u	s	b^*
0.2	0.1	2	5	10	0.578
				15	0.447
				20	0.313
				30	0.121



شکل ۳: مقدار بهینه b^* برای پارامتر s

علاوه بر دو پارامتر انتخاب شده برای تحلیل حساسیت، سایر پارامترهای موثر نیز می‌توانند مورد بررسی قرار گیرند تا دید جامع‌تری از تاثیر آن‌ها بر زمان بهینه آب‌بندی بدست آید.

۶ جمع‌بندی

در این مقاله، زمان بهینه آب‌بندی b^* در طول دوره‌های آب‌بندی و گارانتی مورد بررسی قرار گرفته است. هدف اصلی مقاله، تعیین مقدار بهینه b^* به گونه‌ای است که هزینه‌های کلی سیستم مینیمم شده و در عین حال، محدودیت‌های فرسایش نیز رعایت شوند. برای این منظور، یک مدل ریاضی برای مینیمم سازی هزینه کلی سیستم ارائه شده است. تابع هدف شامل، هزینه و فرسایش بوده و تاثیر پارامترهای مختلف در این توابع با استفاده از تحلیل حساسیت نشان داده شده است. همچنین تعادل بین هزینه‌های آب‌بندی و گارانتی و محدودیت‌های فرسایش، نقش کلیدی در تعیین مقدار بهینه b^* دارند.

در مطالعات آینده، می‌توان از مدل‌های غیرخطی و داده‌محور مانند مدل‌های یادگیری ماشین و شبکه‌های عصبی برای پیش‌بینی نرخ خرابی و بهینه‌سازی فرآیند آب‌بندی استفاده کرد.

مراجع

- [1] Huang, Y. S., Fang, C. C., and Kuo, C. Y. (2024). A study on two-dimensional warranty with a preventive maintenance policy under burn-in or run-in. *Applied Stochastic Models in Business*

and *Industry* , **40**(4), 979–995.

- [2] Kwon, Y. M., Wilson, R., and Na, M. H. (2010). Optimal burn-in with random minimal repair cost. *Journal of the Korean Statistical Society*, **39**, 245–249.
- [3] Ling, X., Yin, R., and Wang, L. (2025). Optimal burn-in for minimally repaired components from n sub-populations. *Proceedings of the Institution of Mechanical Engineers, Part O: Journal of Risk and Reliability*, **239**(1), 31–40.
- [4] MoghimiHadji, E. (2020). Optimal length of warranty and burn-in periods considering different types of repair. *Journal of Industrial and Systems Engineering*, **13**(2), 1–8.
- [5] Safaei, F., and Taghipour, S. (2024). Integrated degradation-based burn-in and maintenance model for heterogeneous and highly reliable items. *Reliability Engineering and System Safety*, **244**, 109942.
- [6] Ye, Z. S., Murthy, D. P., Xie, M., and Tang, L. C. (2013). Optimal burn-in for repairable products sold with a two-dimensional warranty. *IIE Transactions*, **45**(2), 164–176.



یازدهمین سمینار تخصصی
نظریه قابلیت اعتماد و کاربردهای آن
۲۴ و ۲۵ اردیبهشت ۱۴۰۴



تعداد مولفه‌های شکست خورده در یک سیستم منسجم فعال

اعلم یزدیان، ح^۱

گروه آمار، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد

چکیده

در این مقاله، تعداد مولفه‌های از کار افتاده در یک سیستم منسجم فعال مورد بررسی قرار می‌گیرد. فرض بر این است که مؤلفه‌های سیستم دارای توزیع‌های مختلفی از طول عمر هستند، به این معنا که هر مؤلفه زمان خرابی متفاوتی دارد. در ادامه، مسئله بهینه‌سازی برای تعیین مقادیر بهینه تعداد مؤلفه‌های هر نوع و نحوه آرایش آن‌ها در سیستم، فرمول بندی می‌شود.

کلمات کلیدی: سیستم منسجم، قابلیت اطمینان، بهینه سازی

۱ مقدمه

تاکنون تلاش‌های زیادی برای مطالعه سیستم‌های منسجم انجام شده است، زیرا این سیستم‌ها کاربردهای گسترده‌ای در جنبه‌های مختلف زندگی بشر دارند. از این سیستم‌ها در صنایع هوایی، سیستم‌های تأمین انرژی، سیستم‌های مخابراتی

^۱haniealam79@gmail.com

و خطوط تولید صنعتی استفاده می‌شود. همچنین به‌عنوان یک موضوع مهم در نظریه قابلیت اطمینان، تعداد مؤلفه‌های ازکارافتاده در یک سیستم (چه در حالت خرابی و چه در حال کار) توجه زیادی را به خود جلب کرده است. این مقدار در بسیاری از موقعیت‌های عملی مفید است. این اطلاعات به اپراتورهای سیستم کمک می‌کند تا هزینه‌ها را کاهش داده و با استفاده بهینه‌تر از منابع، سودآوری را بهبود بخشند. برای جلوگیری از هزینه‌های اضافی ناشی از تعویض غیرضروری، آن‌ها می‌توانند از تعمیر و نگهداری اصلاحی یا پیشگیرانه استفاده کنند. این استراتژی‌های نگهداری بسیار رایج بوده و به‌طور گسترده در مقالات علمی مورد مطالعه قرار گرفته‌اند. این روش‌ها به اپراتورها کمک می‌کنند تا بررسی‌های دوره‌ای را برنامه‌ریزی کنند، به‌طوری‌که در این بررسی‌ها، تمام مؤلفه‌های معیوب با نمونه‌های جدید جایگزین شده و مؤلفه‌های سالم نیز بازسازی شوند. در روش تعمیر اصلاحی، تعویض پس از خرابی سیستم انجام می‌شود، در حالی که در روش تعمیر پیشگیرانه، تعویض قبل از وقوع خرابی انجام می‌شود. دانستن تعداد مؤلفه‌های ازکارافتاده در یک سیستم (چه در حال کار و چه در حالت خرابی) یک شاخص بسیار مهم است که به اپراتورها کمک می‌کند تا اطلاعات لازم در مورد تعداد قطعات یدکی مورد نیاز بین دوره‌های تعمیر و نگهداری را به‌دست آورند. این موضوع در زمینه تخصیص بهینه افزونگی، برنامه‌های نگهداری بهینه و موارد مشابه بسیار مفید است.

توزیع تعداد مؤلفه‌های خراب‌شده زمانی که سیستمی شامل مؤلفه‌های تعویض‌پذیر از کار می‌افتد و خراب می‌شود، توسط راس^۱ و همکاران [۵] و اريلماز^۲ [۱] مورد بررسی قرار گرفته است. همان مشکل برای سیستم‌های k از n متوالی که به گفته نویسنده می‌تواند برای مدل‌سازی سیستم‌های مخبراتی و خطوط لوله نفتی استفاده شود، توسط پاپاستاوردیس^۳ [۴] حل شد. او فرض کرده بود که طول عمر مؤلفه‌ها متغیرهای تصادفی مستقل و هم‌توزیع هستند.

طراحان سیستم، در میان موضوعات دیگر، به جلوگیری از خرابی زودهنگام سیستم علاقه‌مند هستند، زیرا این امر ممکن است هزینه‌های غیرمنتظره بالایی برای کاربران ایجاد کند. در شرایط واقعی، ممکن است این طراحان تنها اطلاعات جزئی و ناقصی درباره طول عمر سیستم داشته باشند. بر اساس این اطلاعات جزئی، آنها به حفظ سیستم در شرایط کاری بهینه علاقه‌مند هستند. به همین دلیل، در سال‌های اخیر، تحقیقات متعددی تعداد مؤلفه‌های معیوب در یک سیستم در حال کار را مورد مطالعه قرار داده‌اند. معمولاً، موردی که بیشتر در مطالعات مورد توجه قرار گرفته است، مدل عمر مؤلفه‌های مستقل و هم‌توزیع است که طول عمر این مؤلفه‌ها به صورت $T_1, \dots, T_n$ نمایش داده می‌شود و معمولاً دارای توزیع پیوسته است زیرا تحلیل با فرض پیوستگی به مراتب ساده‌تر از حالت گسسته می‌باشد، که در آن احتمال هم‌زمانی خرابی مؤلفه‌ها با احتمال غیرصفر وجود دارد. با این حال تمام محاسبات انجام گرفته در این

Ross^۱

Eryilmaz^۲

Papastavridis^۳

مقاله برای حالت کلی ارائه می‌شود و در حالت پیوسته و گسسته قابل استفاده می‌باشد. اریلماز [۲] تعداد مولفه‌های خراب شده را زمانی که سیستم در حال کار است، برای ساختارهای معروف و کاربردی k از n در نظر گرفت. او فرض کرد که طول عمر مولفه‌ها متغیرهای تصادفی مستقل از K نوع ($K \geq 2$) هستند. او همچنین میانگین تعداد مولفه‌های شکست خورده از هر نوع را در حالی که سیستم k از n در یک نقطه زمانی دلخواه فعال است، مطالعه کرد. جاسینسکی^۴ [۳] نیز به بررسی تعداد مولفه‌های خراب در سیستم‌های منسجم برای سیستم پل پرداخت. هدف این مقاله گسترش این نتایج به سایر سیستم‌های منسجم غیر از ساختارهای k از n است. نتایج نظری در بخش ۳ ارائه می‌شود. بخش ۳ به مسئله بهینه‌سازی اختصاص دارد که به تعیین مقادیر بهینه تعداد مولفه‌های هر نوع همراه با آرایش آنها در سیستم مربوط می‌شود. در بخش ۴، نمونه‌ای از سیستم یک طرفه (یا سیستم جهت‌دار) برای نشان دادن نتایج نظری ارائه شده در این مقاله ارائه می‌شود.

۲ توصیف مدل و نتایج نظری

یک سیستم منسجم دلخواه متشکل از n مولفه را در نظر بگیرید که طول عمر آن‌ها $T_1, \dots, T_n$ متغیرهای تصادفی مستقل از چندین نوع هستند. به طور دقیق‌تر، این متغیرهای تصادفی از $K \geq 2$ نوع مختلف هستند طوری که دقیقاً n_i متغیر تصادفی از نوع i وجود دارند که تابع توزیع تجمعی آن‌ها $F_1, \dots, F_K$ دو به دو متفاوت هستند و $n_1 + n_2 + \dots + n_K = n$ می‌باشد. در ادامه فرض می‌شود

$$T_1^{(1)}, \dots, T_{n_1}^{(1)} \sim F_1,$$

$$T_1^{(2)}, \dots, T_{n_2}^{(2)} \sim F_2,$$

$$\vdots$$

$$T_1^{(K)}, \dots, T_{n_K}^{(K)} \sim F_K,$$

که در آن $T_j^{(i)}$ طول عمر مولفه j ام از نوع i را نشان می‌دهد. به طور خاص، طول عمر مولفه‌ها را می‌توان به طور گسسته توزیع کرد و مقادیر را در زیرمجموعه‌های متناهی یا نامتناهی از مجموعه اعداد صحیح غیرمنفی توزیع کرد. در ادامه، از نمایش $T_{1:n_i}^{(i)} \leq \dots \leq T_{n_i:n_i}^{(i)}$ برای آماره‌های مرتب طول عمرهای $T_1^{(i)}, \dots, T_{n_i}^{(i)}$ استفاده می‌شود. همچنین طول عمر سیستم توسط T نمایش داده می‌شود.

فرض کنید به ازای $i = 1, \dots, K$ ، تعداد مولفه‌های از کار افتاده از نوع i در زمان t را نشان می‌دهد در حالی

که سیستم در زمان t همچنان کار می‌کند (یعنی $T > t$). بنابراین، محاسبه احتمال شرطی زیر مورد توجه است

$$P(N_t^{(i)} = w | T > t) = \frac{P(N_t^{(i)} = w, T > t)}{P(T > t)}, w = 0, 1, \dots, n_i. \quad (1)$$

در ادامه از مفهوم مجموعه مسیرهای مینیمال سیستم استفاده می‌شود. اگر s مجموعه مسیر مینیمال به صورت $P_1, \dots, P_s$ برای یک سیستم وجود داشته باشد، آنگاه طول عمر سیستم را می‌توان به صورت زیر نشان داد

$$T = \max_{1 \leq j \leq s} \min_{p \in P_j} T_p, \quad (2)$$

به این معنی که یک سیستم در صورتی کار می‌کند که تمام مولفه‌های حداقل یکی از مسیرهای مینیمال آن کار کنند. اکنون با استفاده از (۲) می‌توان تابع قابلیت اعتماد سیستم را به صورت زیر نوشت:

$$P(T > t) = P(\max_{1 \leq j \leq s} \min_{p \in P_j} T_p > t).$$

حال با در نظر گرفتن $Y_j = \min_{p \in P_j} T_p$ داریم

$$\begin{aligned} P(T > t) &= P(\max Y_j \geq t) \\ &= P(Y_1 > t \cup \dots \cup Y_s > t). \end{aligned}$$

اکنون با استفاده از اصل شمول و عدم شمول تابع قابلیت اعتماد سیستم را به صورت زیر باز نویسی می‌کنیم

$$P(T > t) = \sum_{j=1}^s (-1)^{j+1} \sum_{\{k_1, \dots, k_j\} \subset \{1, \dots, s\}} P(\{Y_{k_1} > t\} \cap \dots \cap \{Y_{k_j} > t\}).$$

با توجه به این که در رابطه فوق همه Y_j ها از t بزرگتر هستند، پس کوچکترین آن‌ها نیز از t بزرگتر است بنابراین می‌توان نوشت

$$\begin{aligned} P(T > t) &= \sum_{j=1}^s (-1)^{j+1} \sum_{\{k_1, \dots, k_j\} \subset \{1, \dots, s\}} P(\min_{1 \leq l \leq j} Y_{k_l} > t) \\ &= \sum_{j=1}^s (-1)^{j+1} \sum_{\{k_1, \dots, k_j\} \subset \{1, \dots, s\}} P(\min_{1 \leq l \leq j} \min_{p \in P_{k_l}} T_p > t) \\ &= \sum_{j=1}^s (-1)^{j+1} \sum_{\{k_1, \dots, k_j\} \subset \{1, \dots, s\}} P(T_{\bigcup_{l=1}^j P_{k_l}} > t), \end{aligned}$$

که برای $1 \leq m \leq n$ و $C = \{i_1, \dots, i_m\} \subset \{1, \dots, n\}$ ، $1 \leq r \leq m$ نشان دهنده آماره مرتب r ام از $T_{i_1}, \dots, T_{i_m}$ است. از آن جایی که برای فعال بودن سیستم باید حداقل یک مسیر مینیمال فعال باشد، می‌توان فرمول

فوق را به صورت زیر بازنویسی نمود

$$P(T > t) = \sum_{m_1=0}^{n_1} \cdots \sum_{m_K=0}^{n_K} \alpha_{(m_1, \dots, m_K)} P(T_{p_{1,1}}^{(1)} \{k_1, \dots, k_j\} > t, \dots, T_{p_{1,m_1}}^{(1)} \{k_1, \dots, k_j\} > t, \dots, T_{p_{K,1}}^{(K)} \{k_1, \dots, k_j\} > t, \dots, T_{p_{K,m_K}}^{(K)} \{k_1, \dots, k_j\} > t), \quad (3)$$

که در آن $1 \leq m_1 + \dots + m_K$ است و ضریب $\alpha_{(m_1, \dots, m_K)}$ به صورت زیر تعریف می شود

$$\alpha_{(m_1, \dots, m_K)} = \sum_{j=1}^s (-1)^{j+1} \sum_{\{k_1, \dots, k_j\}} I(A), \quad (4)$$

که در آن A برابر است با $\cup_{l=1}^j P_{k_l}$ که دقیقاً شامل m_1 اندیس $p_{1,1}^{\{k_1, \dots, k_j\}}, \dots, p_{1,m_1}^{\{k_1, \dots, k_j\}}$ از طول عمر مولفه های نوع ۱، ... و دقیقاً شامل m_K اندیس $p_{K,1}^{\{k_1, \dots, k_j\}}, \dots, p_{K,m_K}^{\{k_1, \dots, k_j\}}$ از طول عمر مولفه های نوع K است. مشاهده می شود که ضریب $\alpha_{(m_1, \dots, m_K)}$ از طریق مجموعه مسیرهای مینیمال به ساختار سیستم بستگی دارد. تمام اجتماع های ممکن از مجموعه مسیرهای مینیمال به صورت $\cup_{l=1}^j P_{k_l}$ به ازای $j = 1, \dots, s$ در نظر گرفته می شود. اولین عبارت با جمله مثبت در سمت چپ رابطه نمایانگر تعداد مجموعه مسیرهای مینیمال تک مولفه ای است که شامل دقیقاً m_1 مولفه از نوع ۱، ... و m_K مولفه از نوع K می باشد. جمله دوم نشان می دهد که چند اجتماع از دو مجموعه مسیر مینیمال شامل دقیقاً m_1 مولفه از نوع ۱، ... و m_K مولفه از نوع K می باشد که با علامت منفی اضافه می شود. این روند تا اجتماع s مجموعه مسیر مینیمال ادامه می یابد. این روش در حالت کلی برای هر سیستم منسجمی قابل استفاده است. اما برای توضیح دقیق تر در بخش ۴ از آن برای سیستم یک طرفه به عنوان مثال استفاده می شود.

اگر چه در فرمول (۳) جمع فقط بر روی $m_1, \dots, m_K$ است، اما متغیرهای طول عمر موجود در این فرمول به $k_1, \dots, k_j$ بستگی دارد به همین دلیل از نماد $T_{p_{1,1}, \dots, p_{1,m_1}}^{(1)} \{k_1, \dots, k_j\}$ استفاده شده است که مقادیر طول عمر تحت تاثیر انتخاب های j مجموعه مسیر مینیمال قرار می گیرد. فرمول (۳) در مورد متغیرهای تصادفی مستقل به صورت زیر ساده می شود

$$P(T > t) = \sum_{m_1=0}^{n_1} \cdots \sum_{m_K=0}^{n_K} \alpha_{(m_1, \dots, m_K)} \bar{F}_1^{m_1}(t) \cdots \bar{F}_K^{m_K}(t). \quad (5)$$

اکنون به ازای $w = 0, \dots, n_i$ و با استفاده از فرمول (۲) داریم

$$\begin{aligned} P(N_t^{(i)} = w, T > t) &= P\left(T_{w:n_i}^{(i)} \leq t < T_{w+1:n_i}^{(i)}, T > t\right) \\ &= P\left(T_{w:n_i}^{(i)} \leq t < T_{w+1:n_i}^{(i)}, \max_{1 \leq j \leq s} \min_{p \in P_j} T_p > t\right) \\ &= P\left(T_{w:n_i}^{(i)} \leq t < T_{w+1:n_i}^{(i)}, \cup_{j=1}^s \{\min_{p \in P_j} T_p > t\}\right) \\ &= P\left(\cup_{j=1}^s \{T_{w:n_i}^{(i)} \leq t < T_{w+1:n_i}^{(i)}, T_{1:P_j} > t\}\right). \end{aligned}$$

بنابراین، با استفاده از اصل شمول و عدم شمول می‌توان فرمول فوق را به صورت زیر بازنویسی کرد.

$$\begin{aligned}
P(N_t^{(i)} = w, T > t) &= \sum_{j=1}^s (-1)^{j+1} \sum_{\{k_1, \dots, k_j\} \subset \{1, \dots, s\}} P(\{T_{w+1:n_i}^{(i)} \leq t < T_{\cup_{l=1}^j P_{k_l}}^{(i)} > t\}) \\
&= \sum_{j=1}^s (-1)^{j+1} \sum_{\{k_1, \dots, k_j\} \subset \{1, \dots, s\}} \sum_{m=1}^{n-w} P(B),
\end{aligned}$$

که در آن $|\cup_{l=1}^j P_{k_l}|$ نشان دهنده تعداد اعضای B و $\cup_{l=1}^j P_{k_l}$ پیشامد این است که دقیقاً $m = |\cup_{l=1}^j P_{k_l}|$ تا از $(T_1^{(1)}, \dots, T_{n_1}^{(1)}, \dots, T_1^{(K)}, \dots, T_{n_K}^{(K)})$ که اندیس‌های آن به $\cup_{l=1}^j P_{k_l}$ تعلق دارند، بزرگتر از t هستند و دقیقاً w تا از $(T_1^{(i)}, \dots, T_{n_i}^{(i)})$ که اندیس‌های آن به $\cup_{l=1}^j P_{k_l}$ تعلق ندارد، کوچکتر مساوی t هستند. در نتیجه،

$$\begin{aligned}
P(N_t^{(i)} = w, T > t) &= \sum_{j=1}^s (-1)^{j+1} \\
&\sum_{\{k_1, \dots, k_j\} \subset \{1, \dots, s\}} \sum_{m=1}^{n-w} P\left(\left\{\bigcup_{m_1=\cdot}^{n_1} \dots \bigcup_{m_i=\cdot}^{n_i-w} \dots \bigcup_{m_K=\cdot}^{n_K} A_{t,w,k_1, \dots, k_j}^{m, m_1, \dots, m_K}\right\}\right).
\end{aligned}$$

عبارت $A_{t,w,k_1, \dots, k_j}^{m, m_1, \dots, m_K}$ بیان می‌کند که دقیقاً m_1 تا از $(T_1^{(1)}, \dots, T_{n_1}^{(1)})$ که $p_{1,1}^{\{k_1, \dots, k_j\}}, \dots, p_{1,m_1}^{\{k_1, \dots, k_j\}}$ به $\cup_{l=1}^j P_{k_l}$ تعلق دارند، بزرگتر از t هستند، ...، دقیقاً m_i تا از $(T_1^{(i)}, \dots, T_{n_i}^{(i)})$ که $p_{i,1}^{\{k_1, \dots, k_j\}}, \dots, p_{i,m_i}^{\{k_1, \dots, k_j\}}$ به $\cup_{l=1}^j P_{k_l}$ تعلق دارند، بزرگتر از t هستند، دقیقاً w تا از $n_i - m_i$ تا از $T_p^{(i)}$ ها کوچکتر مساوی t هستند، بقیه $n_i - m_i - w$ ها از $T_p^{(i)}$ ها بزرگتر از t هستند، ...، دقیقاً $m_K = m - m_1 - \dots - m_{K-1}$ تا از $(T_1^{(K)}, \dots, T_{n_K}^{(K)})$ که $p_{K,1}^{\{k_1, \dots, k_j\}}, \dots, p_{K,m_K}^{\{k_1, \dots, k_j\}}$ به $\cup_{l=1}^j P_{k_l}$ تعلق دارند، بزرگتر از t هستند. باتوجه به این که به ازای $m_1 = \cdot, \dots, n_1, \dots, m_i = \cdot, \dots, n_i - w, \dots, m_K = \cdot$ در آن $m_1 + \dots + m_K = m$ به صورت دو به دو ناسازگار هستند، داریم

$$\begin{aligned}
P(N_t^{(i)} = w, T > t) &= \sum_{j=1}^s (-1)^{j+1} \sum_{\{k_1, \dots, k_j\} \subset \{1, \dots, s\}} \sum_{m=1}^{n-w} \sum_{m_1=\cdot}^{n_1} \dots \sum_{m_i=\cdot}^{n_i-w} \dots \sum_{m_K=\cdot}^{n_K} P(A_{t,w,k_1, \dots, k_j}^{m, m_1, \dots, m_K}).
\end{aligned}$$

توجه داشته باشید که مقادیر $m, m_1, \dots, m_K$ از انتخاب $k_1, \dots, k_j$ در مجموعه مسیرهای مینیمال $\cup_{l=1}^j P_{k_l}$ تعیین می‌شوند. بنابراین، برای اندیس‌های ثابت $k_1, \dots, k_j$ ، فقط یک عبارت غیر صفر در مجموع زیر وجود دارد

$$\sum_{m=1}^{n-w} \sum_{m_1=\cdot}^{n_1} \dots \sum_{m_i=\cdot}^{n_i-w} \dots \sum_{m_K=\cdot}^{n_K} P(A_{t,w,k_1, \dots, k_j}^{m, m_1, \dots, m_K})$$

بر اساس فرمول (۲) در حالت متغیرهای تصادفی مستقل، فرمول فوق دارای شکل بسته زیر است

$$\begin{aligned}
 & P\left(N_t^{(i)} = w, T > t\right) \\
 &= F_i^w(t) \bar{F}_i^{n_i-w}(t) \sum_{m=1}^{n-w} \sum_{m_1=1}^{n_1} \cdots \sum_{m_i=1}^{n_i-w} \cdots \sum_{m_K=1}^{n_K} \alpha_{(m_1, \dots, m_K)} \binom{n_i - m_i}{w} \\
 & \cdot \bar{F}_1^{m_1}(t) \cdots \bar{F}_{i-1}^{m_{i-1}}(t) \bar{F}_{i+1}^{m_{i+1}}(t) \cdots \bar{F}_K^{m_K}(t), \quad w = 1, \dots, n_i.
 \end{aligned} \quad (۶)$$

با استفاده از (۵) و (۶) می‌توان احتمال شرطی (۱) را به صورت زیر محاسبه نمود:

$$\begin{aligned}
 P(N_t^{(i)} = w \mid T > t) &= F_i^w(t) \bar{F}_i^{n_i-w}(t) \\
 & \cdot \frac{\sum_{m=1}^{n-w} \sum_{m_1=1}^{n_1} \cdots \sum_{m_i=1}^{n_i-w} \cdots \sum_{m_K=1}^{n_K} \alpha_{(m_1, \dots, m_K)} \binom{n_i - m_i}{w} \cdot \bar{F}_1^{m_1}(t) \cdots \bar{F}_{i-1}^{m_{i-1}}(t) \bar{F}_{i+1}^{m_{i+1}}(t) \cdots \bar{F}_K^{m_K}(t)}{\sum_{m_1=1}^{n_1} \cdots \sum_{m_K=1}^{n_K} \alpha_{(m_1, \dots, m_K)} \cdot \bar{F}_1^{m_1}(t) \cdots \bar{F}_{i-1}^{m_{i-1}}(t) \bar{F}_{i+1}^{m_{i+1}}(t) \cdots \bar{F}_K^{m_K}(t)}
 \end{aligned} \quad (۷)$$

برای $w = 1, \dots, n_i$ که در صورت $m_1 + \dots + m_i + \dots + m_K = m$ و در مخرج $m_1 + \dots + m_i + \dots + m_K \geq 1$ می‌باشد.

۳ مسئله بهینه سازی

بدیهی است که یک سیستم منسجم در حال کار به بازبینی‌های دوره‌ای در فواصل زمانی معین (مثلاً یک بار در روز، یک بار در هفته و غیره) نیاز دارد تا از شکست آن جلوگیری شود. پس از بررسی می‌توان تصمیم به جایگزینی آن با یک سیستم جدید گرفت. فرض کنید که سیستم منسجم در زمان t قبل از شکست آن جایگزین شده است، یعنی $T > t$. تمام مولفه‌های خراب با قطعات جدید جایگزین می‌شوند و قطعات شکست نخورده دوباره جوان می‌شوند. قابلیت اعتماد یک مولفه پس از جوانسازی مانند یک مولفه نو می‌شود. بنابراین سیستم پس از جایگزینی جدید می‌شود. c_i هزینه اکتساب یک مولفه از نوع i و c_i^* هزینه متحمل شده برای جوانسازی یک مولفه از نوع i را نشان می‌دهد، طوری که به ازای $c_i > c_i^*$ ، $i = 1, 2, \dots, K$ ، بنابراین، متوسط کل هزینه جایگزینی سیستم را می‌توان با فرمول زیر محاسبه کرد

$$\begin{aligned}
 C(t) &= \sum_{i=1}^K c_i E\left(N_t^{(i)} \mid T > t\right) + \sum_{i=1}^K c_i^* E\left(n_i - N_t^{(i)} \mid T > t\right) \\
 &= \sum_{i=1}^K (c_i - c_i^*) E\left(N_t^{(i)} \mid T > t\right) + \sum_{i=1}^K c_i^* n_i,
 \end{aligned} \quad (۱)$$

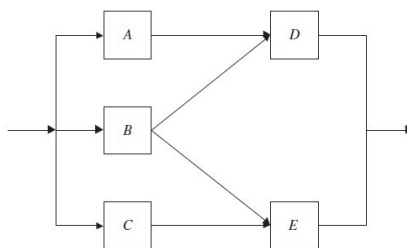
که در آن

$$E\left(N_t^{(i)} \mid T > t\right) = \sum_{w=1}^{n_i} w P\left(N_t^{(i)} = w \mid T > t\right)$$

میانگین تعداد مولفه‌های شکست خورده از نوع i ام در زمان t است در حالی که سیستم در زمان t کار می‌کند. توجه داشته باشید که این هزینه نه تنها به مقادیر $n_1, \dots, n_K$ بستگی دارد، بلکه به ساختار سیستم از طریق مجموعه مسیرهای مینیمال که برای محاسبه احتمالات $P(N_t^{(i)} = w | T > t)$ استفاده می‌شود، نیز بستگی دارد. زمانی که انواع مولفه‌هایی که اندیس‌های آن‌ها به مجموعه مسیرهای مینیمال تعلق دارند تغییر می‌کند بدون آن‌که مقادیر $n_1, \dots, n_K$ تغییر کند، هزینه کل $C(t)$ نیز تغییر خواهد کرد. برای n ، c_i و c_i^* داده شده، علاقه‌مند به تصمیم‌گیری بهینه هستیم که کجا و چه تعداد جزء از هر نوع باید در سیستم در مرحله طراحی قرار گیرد تا $C(t)$ در طول بررسی‌های دوره‌ای به حداقل برسد.

۴ مثال

در این بخش، سیستم یک طرفه ارائه شده در شکل ۱ مورد بررسی قرار می‌گیرد. در ادامه برای تشریح کامل و ساده‌تر تحلیل فرض می‌شود طول عمر مولفه‌ها از دو نوع متفاوت مشخص هستند.



شکل ۱: سیستم منسجم یک طرفه (یا جهت‌دار)

فرض کنید طول عمر مولفه‌ها $T_A, \dots, T_E$ متغیرهای تصادفی مستقل از $K = ۲$ نوع مختلف هستند طوری که

$$T_A, T_D \sim F_1 \quad \text{و} \quad T_B, T_C, T_E \sim F_2.$$

این سیستم منسجم یک طرفه دارای چهار مجموعه مسیر مینیمال است که عبارتند از:

$$P_1 = \{A, D\}, P_2 = \{B, D\}, P_3 = \{B, E\}, P_4 = \{C, E\}. \quad (۱)$$

با توجه به مسیرهای مینیمال معرفی شده در (۱) داریم

$$P_1 \cup P_2 = \{A, B, D\}$$

$$P_1 \cup P_3 = \{A, B, D, E\}$$

$$P_1 \cup P_4 = \{A, C, D, E\}$$

$$P_2 \cup P_3 = \{B, D, E\}$$

$$P_2 \cup P_4 = \{B, C, D, E\}$$

$$P_3 \cup P_4 = \{B, C, E\}$$

و

$$\bigcup_{l=1}^j P_{k_l} = \{A, B, C, D, E\}; \quad 1 \leq k_1 < \dots < k_j \leq 5, j \geq 3.$$

همچنین قابل توجه است که مجموعه مسیر مینیمال P_1 فقط از دو مولفه از نوع یک تشکیل شده است، در حالی که P_2 از یک مولفه از هر نوع می‌باشد. مجموعه مسیرهای مینیمال P_3 و P_4 از دو مولفه از نوع دو ساخته شده‌اند. به علاوه، اجتماع $P_1 \cup P_2$ از دو مولفه از نوع یک و یک مولفه از نوع دو است. اجتماع‌های $P_1 \cup P_3$ ، $P_1 \cup P_4$ ، $P_2 \cup P_3$ و $P_2 \cup P_4$ از دو مولفه از هر نوع تشکیل شده‌اند. اجتماع $P_3 \cup P_4$ از سه مولفه نوع دو هستند، در صورتی که اجتماع‌های $P_2 \cup P_3$ و $P_2 \cup P_4$ از یک مولفه نوع یک و سه مولفه نوع دو ساخته شده‌اند. در نهایت، $P_1 \cup P_2 \cup P_3$ ، $P_1 \cup P_2 \cup P_4$ و $\bigcup_{i=1}^4 P_i$ شامل دو مولفه نوع ۱ و سه مولفه از نوع ۲ هستند. از این رو داریم

$$\alpha_{(m_1, m_2)} = \begin{cases} 2, & (m_1, m_2) = (0, 2), \\ 1, & (m_1, m_2) = (1, 1), (m_1, m_2) = (2, 0), (m_1, m_2) = (2, 3), \\ -1, & (m_1, m_2) = (0, 3), (m_1, m_2) = (1, 2), (m_1, m_2) = (2, 1), \\ & (m_1, m_2) = (2, 2), \\ 0, & \text{در غیر اینصورت} \end{cases}$$

اکنون با استفاده از مقادیر فوق به محاسبه $P(T > t)$ می‌پردازیم. با استفاده از (۳) می‌توان نوشت

$$\begin{aligned} P(T > t) &= \sum_{m_1=0}^3 \sum_{m_2=0}^2 \alpha(m_1, m_2) \bar{F}_1^{m_1}(t) \bar{F}_2^{m_2}(t) \\ &= 2\bar{F}_2(t) - \bar{F}_2^2(t) + \bar{F}_1(t)\bar{F}_2(t) - \bar{F}_1(t)\bar{F}_2^2(t) \\ &\quad + \bar{F}_1^2(t) - \bar{F}_1^2(t)\bar{F}_2(t) \\ &\quad - \bar{F}_1^2(t)\bar{F}_2^2(t) + \bar{F}_1^2(t)\bar{F}_2^3(t) \end{aligned}$$

همچنین، با استفاده از (۷)، عبارت‌های زیر برای تابع جرم احتمال شرطی تعداد خرابی‌های نوع اول بدست می‌آید:

$$\begin{aligned} P(N_t^{(1)} = 0 | T > t) &= \frac{\bar{F}_1^2(t)}{P(T > t)}, \\ P(N_t^{(1)} = 1 | T > t) &= \frac{F_1(t)\bar{F}_1(t)\bar{F}_2(t)(2\bar{F}_2(t) - \bar{F}_2^2(t) + 1)}{P(T > t)}, \\ P(N_t^{(1)} = 2 | T > t) &= \frac{F_1^2(t)\bar{F}_2^2(t)(2 - \bar{F}_2(t))}{P(T > t)} \end{aligned} \quad (2)$$

به طور مشابه، برای تابع جرم احتمال شرط تعداد خرابی‌های نوع دوم احتمال‌های زیر بدست می‌آیند:

$$\begin{aligned} P(N_t^{(2)} = 0 | T > t) &= \frac{\bar{F}_2^3(t)}{P(T > t)}, \\ P(N_t^{(2)} = 1 | T > t) &= \frac{2F_2(t)\bar{F}_2^2(t)\bar{F}_1(t)}{P(T > t)}, \\ P(N_t^{(2)} = 2 | T > t) &= \frac{F_2^2(t)\bar{F}_2(t)\bar{F}_1(t)(1 + 2\bar{F}_1(t))}{P(T > t)}, \\ P(N_t^{(2)} = 3 | T > t) &= \frac{F_2^3(t)\bar{F}_1^2(t)}{P(T > t)}. \end{aligned} \quad (3)$$

اکنون فرض کنید که طول عمر مولفه‌ها به صورت متغیرهای تصادفی با توزیع هندسی است، طوری که برای $p \neq \theta$ ، $t = 0, 1, 2, \dots$ و $p, \theta \in (0, 1)$ داریم:

$$\begin{aligned} F_1(t) &= 1 - (1 - p)^t, \quad \bar{F}_1(t) = (1 - p)^t \\ F_2(t) &= 1 - (1 - \theta)^t, \quad \bar{F}_2(t) = (1 - \theta)^t. \end{aligned} \quad (4)$$

با جایگذاری (۴) در (۲) احتمال‌های شرطی $P(N_t^{(1)} = w | T > t)$ به ازای $w = 0, 1, 2$ قابل محاسبه هستند. نتایج برای $p = 0.7, \theta = 0.4$ و برخی مقادیر انتخاب شده t در جدول ۱ گزارش شده‌اند. به طور مشابه با جایگذاری (۴) در (۳) احتمال‌های شرطی $P(N_t^{(2)} = w | T > t)$ به ازای $w = 0, 1, 2, 3$ گزارش شده‌اند. توجه داشته باشید در سیستم یک طرفه (شکل ۱)، این که کدام مولفه‌ها از نوع ۱ و کدام مولفه‌ها از نوع ۲ هستند، رفتار احتمالات شرطی را تعیین می‌کنند.

جدول ۱: احتمال‌های شرطی $P(N_t^{(1)} = w | T > t)$ به ازای $w = 0, 1, 2$ برای سیستم یک طرفه با $p = 0.7$ ، $\theta = 0.4$ و t های مختلف

w	t						
	۰	۱	۲	۵	۸	۱۰	۳۰
۰	۱	۰/۱۵	۰/۰۳۴	5×10^{-4}	$7/68 \times 10^{-6}$	$4/78 \times 10^{-7}$	5×10^{-19}
۱	۰	۰/۴۴	۰/۲۲۶	۰/۰۲	$2/0.6 \times 10^{-3}$	$4/98 \times 10^{-4}$	5×10^{-10}
۲	۰	۰/۴۱	۰/۷۴	۰/۹۸	۰/۹۹۸	۰/۹۹۹۵	۱

جدول ۲: احتمال‌های شرطی $P(N_t^{(2)} = w | T > t)$ به ازای $w = 0, 1, 2$ برای سیستم یک طرفه با $p = 0.7$ ، $\theta = 0.4$ و t های مختلف

w	t						
	۰	۱	۲	۵	۸	۱۰	۳۰
۰	۱	۰/۳۶	۰/۱۹۶	۰/۰۳۹۸	8×10^{-3}	3×10^{-3}	1×10^{-7}
۱	۰	۰/۱۴	۰/۰۶۳	2×10^{-3}	$6/49 \times 10^{-5}$	$5/88 \times 10^{-6}$	2×10^{-16}
۲	۰	۰/۷۷	۰/۰۶۶	۰/۰۱۴	$1/9 \times 10^{-3}$	$4/8 \times 10^{-4}$	5×10^{-10}
۳	۰	$9/6 \times 10^{-3}$	$8/9 \times 10^{-3}$	$3/9 \times 10^{-4}$	$7/3 \times 10^{-6}$	$4/7 \times 10^{-7}$	5×10^{-19}

در جداول ۱ و ۲ می‌توان دید که احتمالات $i = 1, 2$ ، $P(N_t^{(i)} = 0 | T > 0)$ برابر یک هستند در صورتی که رخ داده‌های $\{N_t^{(i)} = 1 | T > 0\}$ و $\{N_t^{(i)} = 2 | T > 0\}$ غیرممکن هستند. به عبارت دیگر، هنگامی که سیستم یک طرفه در زمان $t = 0$ شروع به کار می‌کند، هیچ مولفه خراب از نوع ۱ و ۲ وجود ندارد که با شهود ما مطابقت دارد. علاوه بر این، رفتار احتمالات $P(N_t^{(i)} = w | T > t)$ به ازای $w = 0, 1, 2$ ، به عنوان توابعی از t در نظر گرفته شده است که به نوع مولفه‌ها نیز بستگی دارد. احتمالات $P(N_t^{(1)} = 0 | T > t)$ ، $P(N_t^{(1)} = 1 | T > t)$ ، $P(N_t^{(2)} = 0 | T > t)$ ، $P(N_t^{(2)} = 1 | T > t)$ ، $P(N_t^{(2)} = 2 | T > t)$ و $P(N_t^{(1)} = 2 | T > t)$ با افزایش t کاهش می‌یابد، در حالی که احتمال $P(N_t^{(1)} = 2 | T > t)$ با افزایش t افزایش می‌یابد.

اکنون به محاسبه متوسط هزینه تعویض سیستم زمانی که $(c_1, c_1^*) = (4, 1)$ ، $(c_2, c_2^*) = (3, 1/5)$ ، $p = 0.7$ و $\theta = 0.4$ با استفاده از (۱) می‌پردازیم.

$$\begin{aligned} C(t) &= \sum_{i=1}^2 (c_i - c_i^*) E(N_t^{(i)} | T > t) + \sum_{i=1}^2 c_i^* n_i \\ &= (4 - 1)E(N_t^{(1)} | T > t) + (3 - 1/5)E(N_t^{(2)} | T > t) + 2 + 4/5 \end{aligned}$$

جدول ۳: هزینه کل سیستم ۱ با $p = 0.7$ ، $\theta = 0.4$ و برای t های مختلف زمانی که $c_1 = 4$ ، $c_1^* = 1$ و $c_2 = 3$ ، $c_2^* = 1/5$.

t	۰	۱	۲	۵	۸	۱۰	۳۰
$C(t)$	۶/۵	۱۰/۷۷۶	۱۱/۹۵	۱۲/۴۸۵	۱۲/۴۹۹۶	۱۲/۴۹۹۹۶	۱۲/۵

در زمان $t = 0$ ، هزینه کل تنها ناشی از نوسازی تمامی مولفه‌ها است، زیرا در این زمان هیچ‌کدام از مولفه‌ها هنوز دچار شکست نشده‌اند. بنابراین، هزینه اولیه فقط به مقادیر پارامترهای p و θ بستگی ندارد. با توجه به نتایج بدست آمده مشاهده می‌شود که متوسط کل هزینه تعویض سیستم با افزایش زمان ($t = 2, 3, \dots$)، بیشتر می‌شود زیرا به مرور زمان مولفه‌ها دچار فرسودگی و نقص می‌شوند و نیاز به تعویض افزایش می‌یابد. در ادامه می‌توان به طور مشابه هزینه کل سیستم را برای حالت‌های مختلف که کدام مولفه‌ها از نوع ۱ و کدام مولفه‌ها از نوع ۲ باشند محاسبه نمود و با توجه به شرایط سیستم بهینه را انتخاب نمود.

۵ جمع‌بندی

در این مقاله به بررسی تعداد مولفه‌های خراب در یک سیستم منسجم در حال کار پرداخته شد. همان‌طور که مشاهده گردید، شناسایی و تحلیل این مولفه‌ها در طراحی بهینه سیستم و کاهش هزینه‌ها نقش مهمی دارد. در بسیاری از تحقیقات علمی و مقالات مرتبط، به تحلیل تعداد مولفه‌های خراب در سیستم‌های شکست‌خورده پرداخته شده است. این مقدار از آن جهت اهمیت دارد که تعیین می‌کند چه تعداد قطعه یدکی باید برای جایگزینی تمامی مولفه‌های خراب در زمان خرابی سیستم در دسترس باشد و این مسأله می‌تواند به بهینه‌سازی عملکرد سیستم‌های پیچیده کمک کند.

- [1] Eryilmaz, S. (2012), *The number of failed components in a coherent system with exchangeable components*, IEEE Transactions on Reliability 61. 1 , 203–207.
- [2] Eryilmaz, S. (2018), *The number of failed components in a k-out-of-n system consisting of multiple types of components*, Reliability Engineering & System Safety. 175, 246–250.
- [3] Jasinski, K. (2022), *On the number of failed components in a coherent system consisting of multiple types of components*, Journal of Computational and Applied Mathematics. 410, 114189.
- [4] Papastavridis, S. (1989), *The number of failed components in a consecutive-k-of-n:f system*, IEEE Transactions on Reliability 38. 3 , 338–340.
- [5] Ross, S. M., Shahshahani, M., and Weiss, G. (1980), *On the number of component failures in systems whose component lives are exchangeable.*, Mathematics of Operations Research 5. 3, 358–365.



یازدهمین سمینار تخصصی
نظریه قابلیت اعتماد و کاربردهای آن
۲۴ و ۲۵ اردیبهشت ۱۴۰۴



مدلسازی فرسایش سیستم تعمیر پذیر تحت فرآیند گاما و برآورد پارامترهای آن

ایرانمنش، ف^۱ چهکندی، م^۲ اطمینان، ج^۳

گروه آمار، دانشکده علوم ریاضی و آمار، دانشگاه بیرجند

چکیده

یک سیستم تعمیرپذیر، با تعمیر ناقص، که سطح فرسایش آن توسط فرآیند گاما مدل شده است را در نظر بگیرید. در هر دوره زمانی T ، عمل تعمیر و نگهداری ناقص انجام می‌شود و دو مدل مدل تعمیر ناقص که به مدل‌های کاهش حسابی فرسایش از مرتبه یک و بی‌نهایت معروفند، در نظر گرفته می‌شوند. با استفاده از روش‌های گشتاوری کلاسیک و اصلاح شده پیشنهادی پارامترهای مدل برآورد می‌شود و با شبیه‌سازی عملکرد برآوردگرهای دو روش مورد مقایسه قرار می‌گیرد.

کلمات کلیدی: مدل‌های کاهش حسابی فرسایش، تعمیر و نگهداری ناقص، برآورد گشتاوری، فرآیند گاما

^۱fatimazarandy@birjand.ac.ir

^۲mchahkandi@birjand.ac.ir

^۳jetminan@birjand.ac.ir

۱ پیش‌گفتار

بررسی طول عمر از گذشته یک روش مرسوم برای پیش بینی قابلیت اعتماد تولیدات مهندسی متفاوت بوده است ([۲]). مطالعات طول عمر در تجزیه و تحلیل بقا [۶]، زیست شناسی و سلامت [۵]، جمعیت شناسی، علوم اکچوئری، اقتصاد و غیره کاربرد دارد. با گسترش تولیدات با قابلیت اعتماد بالا (طول عمر بالا)، این بکارگیری با چالش‌های زیادی مانند، کمبود مشاهدات خرابی و زمان‌بر بودن و هزینه‌بر بودن آنها، مواجهه است. در مقایسه با رویکرد مطالعه طول عمر به کمک داده‌های طول عمر، استفاده از داده‌های فرسایش است که به طور مداوم بر مشخصه‌های اجرا که مستقیماً با کیفیت تولیدات در ارتباط است، نظارت می‌کنند. به عنوان نمونه‌هایی از داده‌های فرسایش، از فرآیندهای سایش، خوردگی، رشد ترک، لرزش و غیره می‌توان نام برد. نظارت بر مشخصه‌های اجرا، از نظر دقت و کارایی، روشی موثرتر برای ارزیابی قابلیت اعتماد سیستم می‌باشد ([۶]).

در طول سال‌های گذشته، تحقیقات بسیاری در مورد مدل‌سازی فرسایش انجام شده است. به طور کلی دو دیدگاه برای مدل‌سازی فرسایش به کار می‌رود. دیدگاه اول که دیدگاه مسیر کلی نام دارد، فرض بر آن است که فرسایش یک پدیده یکنوا و قطعی است و تغییرات تصادفی در بین مؤلفه‌ها وجود دارد. اما در دیدگاه دوم که برای مدل‌سازی فرسایش، با دامنه پیوسته زمان است، از فرآیندهای لوی و برخی از تعمیم‌های آن استفاده می‌شود.

این دیدگاه از اواسط دهه ۱۹۷۰ شروع شده و فرآیندهای متداول لوی که برای مدل‌سازی از آنها استفاده می‌شوند، فرآیند گاما و فرآیند وینر هستند. ویژگی‌های مستقل بودن و هم توزیعی نموها در این فرآیندها موجب تسهیل مدل‌سازی فرسایش می‌شود. فرآیند گاما برای مدل‌سازی انباشت مقدار زیادی از افزایش‌های کوچک در فرسایش غیر نزولی مناسب است، اما فرآیند وینر زمانی که یک فرآیند فرسایش پیوسته با روند کلی تصادفی وجود دارد، مناسب است. فرآیند گاوسی معکوس یکی دیگر از گزینه‌ها برای داده‌های فرسایش است که یک مسیر فرسایش یکنوا را ارائه می‌دهد.

وقتی سیستمی را از دیدگاه فرسایش مورد مطالعه قرار می‌دهیم، شکست سیستم زمانی رخ می‌دهد که سطح فرسایش از یک آستانه بالا رود. لذا عملیات تعمیر و نگهداری به اشکال مختلفی به منظور کاهش سطح فرسایش سیستم در بازه‌های زمانی متعدد انجام می‌پذیرد تا زمان شکست سیستم به تعویق افتد. بسته به اینکه پس از انجام عملیات تعمیر سطح فرسایش سیستم به چه شکلی کاهش می‌یابد انواع مختلفی از تعمیر را داریم. برای آشنایی با مدل‌های تعمیر و نگهداری مطالعه [۴] پیشنهاد می‌شود. مدل‌های کاهش حسابی فرسایش سیستم از مرتبه یک و بی‌نهایت به ترتیب ARD_1 و ARD_∞ از جمله معروف‌ترین انواع مدل‌های تعمیر ناقص هستند. در مدل ARD_1 ، بعد از هر عمل تعمیر و نگهداری، نسبتی از فرسایش انباشته شده پس از آخرین عمل تعمیر و نگهداری حذف می‌شود، و در ARD_∞ ، هر عمل تعمیر و نگهداری، نسبتی از کل فرسایش موجود را کاهش می‌دهد، یعنی فرسایش انباشته شده

از صفر تا زمان t . در این مقاله، به بررسی سیستم‌هایی که فرسایش آنها از مدل‌های ARD_1 و ARD_∞ تحت فرآیند گاما هستند، می‌پردازیم. در بخش ۲ به معرفی فرآیند گاما و مدل‌های ARD_1 و ARD_∞ می‌پردازیم. در بخش ۳ برآورد گشتاوری پارامترهای مدل با دو روش، مورد مطالعه قرار می‌گیرد و در نهایت در بخش ۴ به کمک شبیه‌سازی به بررسی عملکرد برآوردگرها می‌پردازیم.

۲ فرآیند گاما برای مدل‌های ARD_1 و ARD_∞

در این بخش به تشریح سطح فرسایش جمعیتی سیستم نگهداری شده براساس مدل‌های ARD_1 و ARD_∞ ، زمانی که فرسایش سیستم تحت فرآیند گاما است، می‌پردازیم. تابع چگالی متغیر تصادفی $X > 0$ دارای توزیع گاما با پارامتر شکل $a > 0$ و پارامتر مقیاس $\eta > 0$ است هرگاه

$$f(x; a, \eta) = \frac{\eta^a}{\Gamma(a)} x^{a-1} e^{-\eta x}, x > 0. \quad (1)$$

امید ریاضی و واریانس X به ترتیب برابر با $E(X) = \frac{a}{\eta}$ و $Var(X) = \frac{a}{\eta^2}$ است. در ادامه یک سیستم قابل تعمیر را در نظر می‌گیریم که فرسایش آن طبق یک فرآیند گاما مدل می‌شود. بنابراین، ابتدا تعریف فرآیند گاما را ارائه می‌کنیم.

تعریف ۱.۲. فرض کنید $a(t) : R^+ \rightarrow R^+$ یک تابع غیر نزولی است به طوری که $a(0) = 0$ و $\eta > 0$. $\{X_t; t \geq 0\}$ یک فرآیند گاما با تابع شکل $a(t)$ و پارامتر مقیاس η است اگر دارای ویژگی‌های زیر باشد:

۱. هر نمونه $X_t - X_s$ دارای توزیع گاما به شکل $\Gamma(a(t) - a(s), \eta)$ برای $s < t$ است؛

۲. نمونه‌های X_t مستقلند، یعنی، $X_{s_1} - X_{s_2}$ و $X_{t_1} - X_{t_2}$ برای هر $s_1 < s_2 < t_1 < t_2$ از هم مستقل هستند؛

۳. $X_0 = 0$ با احتمال یک.

سیستمی را در نظر بگیرید که فرآیند فرسایش آن $\{X_t; t \geq 0\}$ از فرآیند گاما (Γ) با تابع شکل $a(\cdot)$ و پارامتر مقیاس η پیروی می‌کند و در زمان فعالیت سیستم، $Y = (Y_t)_{t \geq 0}$ ، فرایند تصادفی است که سطح فرسایش جمعیتی سیستم را مطابق با مدل تعمیر نشان می‌دهد. سیستم تحت تعمیر و نگهداری ناقص در دوره‌های زمانی T قرار می‌گیرد. در مدل ARD_1 ، هر تعمیر از میزان فرسایش انباشته شده سیستم در فاصله زمانی دو تعمیر اخیر، ρ درصد می‌کاهد. به عبارت دیگر، بعد از تعمیر و نگهداری ناقص در زمان jT ، $j \in \{1, 2, \dots, n\}$ ، به نسبت ρ (پارامتر تعمیر) از فرسایش انباشته شده سیستم در فاصله زمانی $[(j-1)T, jT]$ کاسته می‌شود. این مدل را مدل کاهش حسابی فرسایش مرتبه

یک (ARD_1) می‌نامند. در زمان شروع به کار سیستم سطح فرسایش صفر فرض می‌شود، یعنی $Y_0 = X_0$. در فاصله زمانی $[0, T)$ ، قبل از اولین تعمیر، $Y_t = X_t$ ، و بعد از تعمیر اول $Y_T = (1 - \rho)(X_T^{(1)} - X_0^{(1)}) = (1 - \rho)X_T^{(1)}$. بعد از دومین تعمیر سطح فرسایش برابر است با

$$\begin{aligned} Y_{2T} &= Y_T + (X_{2T}^{(2)} - X_T^{(2)}) - \rho(X_{2T}^{(2)} - X_T^{(2)}) \\ &= (1 - \rho)((X_{2T}^{(2)} - X_T^{(2)}) + (X_T^{(1)} - X_0^{(1)})). \end{aligned}$$

بنابراین، برای هر $(n \in N^*)$ ، $t \in [nT, (n+1)T)$ ، میزان فرسایش تجمعی سیستم برابر است با

$$Y_t = Y_{nT} + (X_t^{(n+1)} - X_{nT}^{(n+1)}), \quad (2)$$

و

$$Y_{nT} = (1 - \rho) \sum_{j=1}^n (X_{jT}^{(j)} - X_{(j-1)T}^{(j)}). \quad (3)$$

میانگین و واریانس Y_t در لم زیر محاسبه شده‌اند، که در ادامه پژوهش مورد استفاده قرار می‌گیرند.

لم ۲.۲. برای هر $n \in N^*$ و $nT \leq t < (n+1)T$ ، میانگین و واریانس سطح فرسایش Y_t به ترتیب برابر است با

$$\begin{aligned} E(Y_t) &= \frac{1}{\eta}(a(t) - \rho a(nT)), \\ Var(Y_t) &= \frac{1}{\eta^2}(a(t) - \rho(2 - \rho)a(nT)). \end{aligned}$$

برهان. با توجه به اینکه $(X_t^{(n+1)} - X_{nT}^{(n+1)})$ و Y_{nT} از هم مستقلند و به ترتیب دارای توزیع گاما، $\Gamma(a(t) - a(nT), \eta)$ و $\Gamma(a(nT), \frac{\eta}{1-\rho})$ هستند و همچنین با توجه به رابطه ۲ لم اثبات می‌شود. $\square$

مدل ARD_∞ ، از نظر ساختاری با مدل قبل شباهت‌هایی دارد. تفاوت این مدل با مدل قبل در این است که اثر هر تعمیر بر روی فرسایش تجمعی سیستم از زمان ابتدایی $t = 0$ است. بنابراین در زمان jT ، $j \in N^*$ ، تعمیر فرسایش تجمعی را تا آن لحظه به اندازه ρ درصد کاهش می‌دهد.

فرآیند $(Y_t)_{t \geq 0}$ ، فرآیند تصادفی است که سطح فرسایش سیستم را مطابق با مدل ARD_∞ نشان می‌دهد. در مدل ARD_∞ ، همانند مدل ARD_1 سطح فرسایش اولیه و فرسایش تجمعی در فاصله $t \in [0, 2T)$ یکسان تعریف می‌شود. بعد از دومین تعمیر به نسبت ρ از فرسایش تجمعی از لحظه صفر کاسته می‌شود، یعنی $Y_{2T} = (1 - \rho)(Y_T + (X_{2T}^{(2)} - X_T^{(2)}))$ سطح فرسایش سیستم پس از n امین عمل تعمیر ناقص عبارت است از

$$Y_{nT} = \sum_{j=1}^n (1 - \rho)^{n-j+1} (X_{jT}^{(j)} - X_{(j-1)T}^{(j)}).$$

لذا برای هر $n \in N^*$ و $nT \leq t < (n+1)T$ داریم

$$Y_t = Y_{nT} + (X_t^{(n+1)} - X_{nT}^{(n+1)}) \quad (4)$$

با توجه به اینکه در فرآیند گاما نموها مستقل هستند و از توزیع گاما پیروی میکنند و از طرفی

$$\begin{cases} (X_t^{(n+1)} - X_{nT}^{(n+1)}) \sim \Gamma((a(t) - a(nT)), \eta), \\ (1 - \rho)^{n-j+1} (\Delta X^{(j)}) \sim \Gamma(\Delta a(jT), \frac{\eta}{(1-\rho)^{n-j+1}}) \end{cases}$$

که در آن $\Delta X^{(j)} = X_{jT}^{(j)} - X_{(j-1)T}^{(j)}$ و $\Delta a(jT) = a(jT) - a((j-1)T)$ را نشان می‌دهند. برای هر $n \in N^*$ و

$nT \leq t < (n+1)T$ امید ریاضی و واریانس سطح فرسایش Y_t به ترتیب برابر با

$$E(Y_t) = \frac{1}{\eta} \sum_{j=1}^n (1 - \rho)^{n-j+1} (a(jT) - a((j-1)T)) + (a(t) - a(nT)), \quad (5)$$

$$Var(Y_t) = \frac{1}{\eta^2} \left(\sum_{j=1}^n (1 - \rho)^{2(n-j+1)} (a(jT) - a((j-1)T))^2 + (a(t) - a(nT))^2 \right). \quad (6)$$

می‌باشند.

برای بهینه سازی سیاست‌های تعمیر و نگهداری، بدنبال برآورد پارمترهای مدل فرسایش و اثر تعمیر و نگهداری هستیم. با برآورد پارمترها رفتار آینده سیستم قابل پیش‌بینی می‌شود و امکان بهینه سازی هزینه و دوره اقدامات تعمیر و نگهداری وجود دارد. در ادامه، روش‌های برآورد گشتاوری (ME) برای مدل‌های ARD_1 و ARD_∞ ، زمانی که فرایند فرسایش سیستم توسط یک فرایند Γ مدل شود، مورد مطالعه قرار می‌گیرد. علاوه بر روش گشتاوری معمول، روش برآورد گشتاوری پیشنهادی برای برآورد پارمترهای مدل نیز ارائه خواهد شد.

۳ برآورد گشتاوری پارمترها

در این بخش، پارمترهای مدل‌های هدف براساس روش برآورد گشتاوری (ME) برآوردیابی می‌شوند. دوره زمانی T معلوم است و هدف برآورد پارمترهای تابع شکل $a(\cdot)$ ، پارمتر مقیاس η و اثر تعمیر ρ است. بدلیل سادگی محاسبات تابع $a(\cdot)$ به صورت $a : t \mapsto \alpha t^\beta$ تعریف می‌شود. بنابراین بردار پارامتری برابر با $\theta = (\alpha, \beta, \rho, \eta) \in \Theta$ با $\Theta = [0, \infty)^2 \times [0, 1) \times [0, \infty)$ است.

فرض کنید سطح فرسایش قبل از انجام عمل تعمیر n ام ($n \in N^*$) یعنی در زمان‌های $T^-, 2T^-, \dots, nT^-$ قابل اندازه‌گیری است، یعنی $Y_{T^-}, Y_{2T^-}, \dots, Y_{nT^-}$. سیستم هم‌توزیع و مستقل در نظر گرفته می‌شود. $\mathbf{y} = (y_j^{(i)})_{1 \leq j \leq n, 1 \leq i \leq s}$. جاییکه $y_j^{(i)}$ سطح فرسایش i امین سیستم را در زمان jT^- را برای $1 \leq j \leq n$ ، $1 \leq i \leq s$ ، نشان می‌دهد.

برای برآورد بردار پارامتر θ با روش گشتاوری دو حالت در نظر گرفته می‌شود. حالت اول، براساس روش تعمیم یافته گشتاورها برگرفته از [۸] و [۱۰] بررسی شده است. حالت دوم، نسخه‌ای اصلاح شده از برآورد گشتاوری پیشنهاد شده است که کارایی بیشتری نسبت به حالت اول دارد. برای برآورد پارامتر θ به روش گشتاوری باید فاصله بین گشتاورهای تجربی و گشتاورهای نظری براساس فاصله D مینیم شود.

$$D(\theta, \theta_0) = \sum_{k=1}^d \sum_{j=1}^n (m_k(\theta, jT^-) - m_k(\theta_0, jT^-))^2 \quad (۱)$$

که در آن θ_0 مقدار واقعی پارامتر، d حداکثر مرتبه گشتاورها، n تعداد تعصیرات و $m_k(\theta, jT^-)$ گشتاور k ام هستند. بنابراین

$$m_1(\theta, jT^-) = E(Y_j), m_k(\theta, jT^-) = E(Y_j - E(Y_j))^k; k \geq 2.$$

آنگاه برآورد θ_0 عبارت است از

$$\hat{\theta} = \arg \min_{\theta \in \Theta} \hat{D}(\theta)$$

که در آن $\hat{D}$ نسخه تجربی $D(\theta, \theta_0)$ است. $\hat{D}$ می‌تواند با جایگزینی $\hat{m}_k(jT^-)$ به جای $m_k(\theta_0, jT^-)$ در رابطه (۱) بدست آید، که در آن

$$\hat{m}_1(jT^-) = \frac{1}{s} \sum_{i=1}^s y_j^{(i)}, \hat{m}_k(jT^-) = \frac{1}{s} \sum_{i=1}^s (y_j^{(i)} - \hat{m}_1(jT^-))^k, k \geq 2.$$

برای برآورد بردار پارامتری $\theta = (\alpha, \beta, \rho, \eta)$ با روش گشتاوری برای مدل ARD_1 و مدل ARD_∞ ، ابتدا لازم است که بدانیم چه تعداد گشتاور d و چه تعداد مشاهده n نیاز است.

گزاره ۱.۳. پارامترهای مدل ARD_1 (یا ARD_∞) با روش گشتاوری ME شناسایی می‌شوند، اگر:

- حداقل در زمانهای $1T^-$, $2T^-$, $3T^-$ سطح فرسایش قابل مشاهده باشد، یعنی $n \geq 3$.
- حداقل ۲ گشتاور برای برآوردیابی استفاده شود، یعنی $d \geq 2$.

□

برهان. به [۱۰] مراجعه شود.

با توجه به گزاره ۱.۳ برای پیدا کردن برآوردگرها، در رابطه ۱، $d = 2$ در نظر گرفته می‌شود. با قرار دادن گشتاورهای مرکزی مرتبه ۱ و ۲، رابطه زیر بدست می‌آید.

$$\hat{D}(\theta) = \sum_{j=1}^n ((E(Y_{jT^-}) - \hat{m}_1(j))^2 + (Var(Y_{jT^-}) - \hat{m}_2(j))^2).$$

برای سادگی بردار پارامتر $\theta' = (\mu, \gamma, \rho, \eta)$ را با قرار دادن $\mu = \frac{\alpha}{\eta}$ و $\gamma = \frac{\alpha}{\eta^2}$ ، به جای θ جایگزین می‌شود.

۱.۳ برآورد گشتاوری برای مدل ARD_1

امید ریاضی و واریانس Y_{jT^-} برای هر $j = 1, \dots, n$ به ترتیب برابرند با

$$E(Y_{jT^-}) = \mu T^\beta (j^\beta - \rho(j-1)^\beta),$$

$$Var(Y_j) = \gamma T^\beta (j^\beta - \rho(2-\rho)(j-1)^\beta).$$

بنابراین، $\hat{D}(\theta)$ برابر است با

$$\hat{D}(\theta) = \mu^\gamma k_1(\beta, \rho) - 2\mu k_2(\beta, \rho) + \gamma^\gamma z_1(\beta, \rho) - 2\gamma z_2(\beta, \rho) + B \quad (2)$$

که در آن k_1, k_2, z_1, z_2 از μ و γ مستقلند و همچنین مقدار ثابت B از بردار پارامتر θ' مستقل است. مقدار B و توابع k_1, k_2, z_1, z_2 برابرند با

$$B = \sum_{j=1}^n (\hat{m}_1(j))^\gamma (jT^-) - \hat{m}_2(j)^\gamma,$$

$$k_1(\beta, \rho) = T^{\gamma\beta} \sum_{j=1}^n \left(j^{\gamma\beta} - 2\rho(j(j-1))^\beta + \rho^\gamma(j-1)^{\gamma\beta} \right)$$

$$k_2(\beta, \rho) = T^\beta \sum_{j=1}^n \left(j^\beta \hat{m}_1(jT^-) - \rho(j-1)^\beta \hat{m}_1(jT^-) \right),$$

$$z_1(\beta, \rho) = T^{\gamma\beta} \sum_{j=1}^n \left(j^{\gamma\beta} - 2\rho(2-\rho)(j(j-1))^\beta + \rho^\gamma(2-\rho)^\gamma(j-1)^{\gamma\beta} \right),$$

$$z_2(\beta, \rho) = T^\beta \sum_{j=1}^n \left(j^\beta \hat{m}_2(jT^-) - \rho(2-\rho)(j-1)^\beta \hat{m}_2(jT^-) \right).$$

روش I

برای پیدا کردن برآورگرها با روش I ، باید مشتقات جزئی $\hat{D}(\theta')$ را نسبت به μ و η برابر با ۰ قرار دهیم. بنابراین

$$\begin{cases} \frac{\partial}{\partial \mu} \hat{D}(\theta') = 0 \\ \frac{\partial}{\partial \gamma} \hat{D}(\theta') = 0 \end{cases}$$

که بدست می‌آید

$$\begin{cases} 2\mu k_1(\beta, \rho) - 2k_2(\beta, \rho) = 0 \\ 2\gamma z_1(\beta, \rho) - 2z_2(\beta, \rho) = 0 \end{cases}$$

بنابراین پارامترهای μ و η نسبت به β و ρ به صورت زیر بیان می‌شوند.

$$\mu = h_\mu(\beta, \rho) = \frac{k_2(\beta, \rho)}{k_1(\beta, \rho)}, \gamma = h_\gamma(\beta, \rho) = \frac{z_2(\beta, \rho)}{z_1(\beta, \rho)}$$

بنابراین رابطه (۲) به صورت رابطه زیر نوشته می‌شود

$$\tilde{D}(\beta, \rho) = \tilde{D}(h_\mu(\beta, \rho), \beta, \rho, h_\gamma(\beta, \rho)) = C - \left(\frac{k_2(\beta, \rho)}{k_1(\beta, \rho)} + \frac{z_2(\beta, \rho)}{z_1(\beta, \rho)} \right)$$

تابع $\tilde{D}(\beta, \rho)$ برای برآورد پارامترهای ρ و β به صورت عددی بهینه می‌شود. بنابراین، پارامترهای μ و γ با $\hat{\mu} = h_\mu(\hat{\beta}, \hat{\rho})$ و $\hat{\gamma} = h_\gamma(\hat{\beta}, \hat{\rho})$ برآورد می‌شوند. در پایان برآورد پارامتر θ به صورت زیر ارائه می‌شود

$$\hat{\theta} = \left(\frac{k_2(\hat{\beta}, \hat{\rho})z_1(\hat{\beta}, \hat{\rho})}{k_1(\hat{\beta}, \hat{\rho})z_2(\hat{\beta}, \hat{\rho})}, \hat{\beta}, \hat{\rho}, \frac{k_2(\hat{\beta}, \hat{\rho})z_1(\hat{\beta}, \hat{\rho})}{k_1(\hat{\beta}, \hat{\rho})z_2(\hat{\beta}, \hat{\rho})} \right), \quad (3)$$

که در آن

$$(\hat{\beta}, \hat{\rho}) = \arg \min_{(\beta, \rho) \in [0, \infty) \times [0, 1)} \tilde{D}(\beta, \rho).$$

روش II

در این روش، ابتدا برآوردگرهای اولیه α ، β و ρ با استفاده از اولین گشتاور تجربی و سطح فرسایش مشاهده شده برای s سیستم یافت می‌شوند. با توجه به اینکه

$$m_1(j) = E(Y_{jT-}) = \frac{\alpha T^\beta}{\eta} (j^\beta - (j-1)^\beta) \quad (4)$$

$$m_2(j) = Var(Y_{jT-}) = \frac{\alpha T^\beta}{\eta^2} (j^\beta - \rho(2-\rho)(j-1)^\beta) \quad (5)$$

می‌باشند. ابتدا با قرار دادن $j = 1$ با استفاده از داده های ستون اول ماتریس شبیه سازی شده، مقادیر $\frac{\alpha T^\beta}{\eta}$ و $\frac{\alpha T^\beta}{\eta^2}$ با استفاده از گشتاورهای تجربی به ترتیب $m_1(\theta, T^-)$ و $m_2(\theta, T^-)$ را برآورد می‌کنیم. با بهینه سازی تابع

$$\sum_{j=2}^n \left(\frac{\widehat{m}_1(jT^-)}{\widehat{m}_1(T^-)} - (j^\beta - \rho(j-1)^\beta) \right)^2$$

مقدار اولیه برای پارامترهای β و ρ برآورد می‌شود. سپس مقادیر $(j^{\hat{\beta}} - \hat{\rho}(j-1)^{\hat{\beta}})$ را برای $j = 1, \dots, n$ بدست می‌آوریم. با جایگزینی این مقادیر در رابطه $\frac{\widehat{m}_1(jT^-)}{(j^{\hat{\beta}} - \hat{\rho}(j-1)^{\hat{\beta}})}$ به تعداد n برآورد از $\frac{\alpha T^\beta}{\eta}$ بدست می‌آوریم. از آنجایی که مقدار T ثابت است لذا مقدار $T = 1$ را برای آن در همه جا در نظر میگیریم، بنابراین با میانگین گرفتن از این

روابط مقدار اولیه‌ای برای پارامتر $\hat{\alpha}_1 = \frac{1}{n} \sum_{j=1}^n \frac{\widehat{m}_1(jT^-)}{(j^{\hat{\beta}} - \hat{\rho}(j-1)^{\hat{\beta}})}$ برآورد می‌شود. با بهینه سازی

$$\sum_{j=2}^n \left(\frac{\widehat{m}_1(j)}{\hat{\alpha}_1} - (j^{\beta} - \rho(j-1)^{\beta})^2 \right) \quad (6)$$

برآوردهای جدید و دقیق‌تری برای ρ و β حاصل می‌شود. با قرار دادن $\hat{\rho}$ و $\hat{\beta}$ در رابطه

$$\frac{\widehat{m}_1(jT^-)}{(j^{\hat{\beta}} - \hat{\rho}(j-1)^{\hat{\beta}})}$$

میانگین گرفتن، مقدار برآورد دقیق‌تری برای پارامتر α نتیجه می‌شود. برای برآورد پارامتر η با توجه به اینکه $\frac{\alpha T^{\beta}}{\eta^2} = \left(\frac{\widehat{m}_2(jT^-)}{j^{\hat{\beta}} - \hat{\rho}(2-\hat{\rho})(j-1)^{\hat{\beta}}} \right)$ است، با جایگزینی برآوردهای α ، β و ρ در این رابطه و قرار دادن $j = 1, \dots, n$ به تعداد n برآورد داریم، که با میانگین گرفتن از آنها مقدار η برآورد می‌شود.

۲.۳ برآورد گشتاوری برای مدل ARD_{∞}

برای هر $j = 1, \dots, n$ امید و واریانس Y_{jT^-} به ترتیب برابند با $E(Y_{jT^-}) = \mu k_1^{(j)}(\beta, \rho)$ و $Var(Y_j) = \gamma k_2^{(j)}(\beta, \rho)$ که در آن

$$k_1^{(j)}(\beta, \rho) = T^{\beta} \left(j^{\beta} - (j-1)^{\beta} + \sum_{i=1}^{j-1} (1-\rho)^{j-i} (i^{\beta} - (i-1)^{\beta}) \right)$$

$$k_2^{(j)}(\beta, \rho) = T^{\beta} \left(j^{\beta} - (j-1)^{\beta} + \sum_{i=1}^{j-1} (1-\rho)^{2(j-i)} (i^{\beta} - (i-1)^{\beta}) \right).$$

بنابراین تابع $\hat{D}(\theta')$ برابر است با:

$$\hat{D}(\theta') = \sum_{j=1}^n \left(\mu k_1^{(j)}(\beta, \rho) - \hat{m}_1(j) \right)^2 + \sum_{j=1}^n \left(\gamma k_2^{(j)}(\beta, \rho) - \hat{m}_2(j) \right)^2 \quad (7)$$

روش I

برای پیدا کردن برآورگرها باید مشتقات جزئی $\hat{D}(\theta')$ را نسبت به μ و γ برابر با ۰ قرار می‌دهیم. بنابراین

$$\begin{cases} \frac{\partial}{\partial \mu} \hat{D}(\theta') = 0 \\ \frac{\partial}{\partial \gamma} \hat{D}(\theta') = 0 \end{cases}$$

که برابر است با

$$\begin{cases} \mu \sum_{j=1}^n (k_1^{(j)}(\beta, \rho))^2 - \sum_{j=1}^n k_1^{(j)}(\beta, \rho) \hat{m}_1(j) = 0 \\ \gamma \sum_{j=1}^n (k_2^{(j)}(\beta, \rho))^2 - \sum_{j=1}^n k_2^{(j)}(\beta, \rho) \hat{m}_2(j) = 0 \end{cases}$$

پارامترهای μ و γ نسبت به β و ρ به صورت زیر بیان می‌شوند.

$$\mu = h_\mu(\beta, \rho) = \frac{\sum_{j=1}^n k_1^{(j)}(\beta, \rho) \hat{m}_1(j)}{\sum_{j=1}^n (k_1^{(j)}(\beta, \rho))^2}$$

و

$$\gamma = h_\gamma(\beta, \rho) = \frac{\sum_{j=1}^n k_2^{(j)}(\beta, \rho) \hat{m}_2(j)}{\sum_{j=1}^n (k_2^{(j)}(\beta, \rho))^2}.$$

رابطه (۶) می‌تواند به صورت رابطه زیر نوشته شود

$$\begin{aligned} \tilde{D}(\beta, \rho) &= \tilde{D}(h_\mu(\beta, \rho), \beta, \rho, h_\gamma(\beta, \rho)) \\ &= C - \left(\frac{\left(\sum_{j=1}^n k_1^{(j)}(\beta, \rho) \hat{m}_1(j) \right)^2}{\sum_{j=1}^n (k_1^{(j)}(\beta, \rho))^2} + \frac{\left(\sum_{j=1}^n k_2^{(j)}(\beta, \rho) \hat{m}_2(j) \right)^2}{\sum_{j=1}^n (k_2^{(j)}(\beta, \rho))^2} \right) \end{aligned}$$

که در آن C مقدار ثابت است و تابع $\tilde{D}(\beta, \rho)$ برای برآورد پارامترهای ρ و β نسبت به ρ و β بصورت عددی بهینه می‌شود. بنابراین، پارامترهای μ و γ با $\hat{\mu} = h_\mu(\hat{\beta}, \hat{\rho})$ و $\hat{\gamma} = h_\gamma(\hat{\beta}, \hat{\rho})$ برآورد می‌شوند. در پایان برآورد پارامتر θ به صورت $\hat{\theta} = (\hat{\alpha}, \hat{\beta}, \hat{\rho}, \hat{\eta})$ بدست می‌آید، که در آن $\hat{\alpha}$ و $\hat{\eta}$ برابرند با

$$\hat{\alpha} = \frac{\left(\sum_{j=1}^n k_1^{(j)}(\beta, \rho) \hat{m}_1(j) \right)^2 \sum_{j=1}^n (k_2^{(j)}(\beta, \rho))^2}{\left(\sum_{j=1}^n (k_1^{(j)}(\beta, \rho))^2 \right)^2 \sum_{j=1}^n k_2^{(j)}(\beta, \rho) \hat{m}_2(j)}, \quad (8)$$

$$\hat{\eta} = \frac{\sum_{j=1}^n k_1^{(j)}(\beta, \rho) \hat{m}_1(j) \sum_{j=1}^n (k_2^{(j)}(\beta, \rho))^2}{\sum_{j=1}^n (k_1^{(j)}(\beta, \rho))^2 \sum_{j=1}^n k_2^{(j)}(\beta, \rho) \hat{m}_2(j)}. \quad (9)$$

و $\hat{\beta}$ و $\hat{\rho}$ از طریق

$$(\hat{\beta}, \hat{\rho}) = \arg \min_{(\beta, \rho) \in [0, \infty) \times [0, 1)} \tilde{D}(\beta, \rho).$$

بدست می‌آیند.

با به کارگرفتن روش II در مدل ARD_1 ، برآورد اولیه برای β و ρ با بهینه سازی تابع زیر حاصل می شود.

$$\sum_{j=2}^n \left(\frac{\widehat{m}_1(jT^-)}{\widehat{m}_1(T^-)} - (j^\beta - \rho(j-1)^\beta) \sum_{i=1}^{j-1} (1-\rho)^{j-i} (i^\beta - (i-1)^\beta) \right)^2 \quad (10)$$

سپس، برآوردی جدید برای $\frac{\alpha T^\beta}{\eta}$ با رابطه زیر کسب می شود

$$\frac{\widehat{\alpha T^\beta}}{\eta} = \frac{1}{n} \sum_{j=2}^n \frac{\widehat{m}_1(jT^-)}{(j^{\hat{\beta}} - \hat{\rho}(j-1)^{\hat{\beta}})} + \sum_{i=1}^{j-1} (1-\hat{\rho})^{j-i} (i^{\hat{\beta}} - (i-1)^{\hat{\beta}})$$

با جایگذاری $\frac{\widehat{\alpha T^\beta}}{\eta}$ در رابطه ۱۰ برآوردهای جدیدی برای β و ρ نتیجه می شود. سپس، با در نظر گرفتن $T=1$ داریم

$$\hat{\alpha} = \frac{1}{n} \sum_{j=1}^n \frac{\widehat{m}_1(jT^-)}{(j^{\hat{\beta}} - \hat{\rho}(j-1)^{\hat{\beta}}) \sum_{i=1}^{j-1} (1-\hat{\rho})^{j-i} (i^{\hat{\beta}} - (i-1)^{\hat{\beta}})}$$

با استفاده از دومین گشتاور تجربی $\hat{\eta}$ برابر است با

$$\hat{\eta} = \frac{\hat{\alpha}}{n} \sum_{j=1}^n \frac{(j^{\hat{\beta}} - \hat{\rho}(j-1)^{\hat{\beta}}) \sum_{i=1}^{j-1} (1-\hat{\rho})^{j-i} (i^{\hat{\beta}} - (i-1)^{\hat{\beta}})}{\widehat{m}_1(jT^-)}$$

در بخش بعد، به مقایسه عملکرد برآوردهای دو مدل ARD_1 و ARD_∞ با روش I و روش II با استفاده از شبیه سازی پرداخته می شود.

۴ دست آوردهای پژوهش

به منظور مطالعه مدل و مقایسه برآورد پارامترهای مدل با استفاده از دو روش ذکر شده، نیاز است که نمونه ای به حجم n که سطوح فرسایش را درست قبل از عمل تعمیر نشان دهد، برای s سیستم هم توزیع شبیه سازی شود. ابتدا، نموها $\Delta X^{(j)} = X_j^{(j)} - X_{j-1}^{(j)}$ را از توزیع گاما با پارامتر شکل $\mu = \alpha T^\beta (j^\beta - (j-1)^\beta)$ و پارامتر مقیاس η ، که در آن $j = 1, \dots, n$ ، برای s سیستم شبیه سازی می کنیم. آنگاه، سطوح فرسایش درست در مدل ARD_1 قبل از تعمیر از رابطه $Y_j = (1-\rho) \sum_{i=1}^{j-1} \Delta X^{(i)} + \Delta X^{(j)}$ برای هر $j = 1, 2, \dots, n$ بدست می آید. داده های شبیه سازی شده در

را در یک ماتریس $s \times n$ می‌ریزیم. چارچوب زیر در مورد پارامترهای مدل و ویژگی‌های مشاهدات ارائه می‌گردد. پارامترهای تابع شکل برابر با $\alpha = 1, \beta = 1/2$ و $\rho = \{0/2, 0/5, 0/8\}$ ، پارامتر مقیاس $\eta = 1$ ، دوره تعمیر $T = 1$ در نظر گرفته شدند. تعداد سطوح مشاهده شده فرسایش $n = \{5, 10\}$ و تعداد سیستم‌های هم توزیع و مستقل مشاهده شده برابر با $s = \{10, 50, 100\}$ می‌باشد.

برای هر مجموعه پارامتری، تعداد ۱۰۰۰۰ داده تولید و پارامتر θ با دو روش گشتاوری I و II برآورد شده است. برآوردیابی با روش I ، بر اساس بهینه سازی تابع دو متغیره است که در نتیجه آن دو پارامتر β و ρ برآورد می‌شوند. بهینه سازی با روش گردایان در محدوده پارامتر β ، $[0, 20]$ و پارامتر ρ ، $[0, 0.99]$ انجام می‌شود. سپس با جایگزینی مقادیر برآورد شده $\hat{\beta}$ و $\hat{\rho}$ در روابط ۳ پارامترهای α و η برآورد می‌شوند. برآوردیابی با روش II گشتاوری مطابق با آنچه که در بخش ۳ توضیح داده شد به دست می‌آیند. شبیه سازی برای مدل ARD_1 انجام شده و برای مقایسه دو روش MSE برآوردگرها محاسبه و در جدول ۱ آورده شده‌اند.

در جدول ۱ تاثیر تغییر s و n بر برآورد پارامترها مشهود است. در هر دو روش گشتاوری I و II ، هر چه تعداد سیستم‌ها افزایش یابد اریبی برآوردگرها کم می‌شود و عملکرد برآوردگرها بهتر می‌شود. اما افزایش n در همه موارد باعث بهبود عملکرد پارامترها نمی‌شود. در روش I ، اریبی برآورد پارامتر β با افزایش n افزایش می‌یابد. همانطور که در جدول مشاهده می‌کنید، برآورد پارامترها با روش II ، اریبی کوچکتری نسبت به برآوردگرها با روش I دارند. لذا، برآورد پارامترها با روش II عملکرد برآوردگرها را به روش I بهبود می‌دهد. در روش I ، برآورد پارامترهای β و b وقتی که $n = 10$ اریب هستند، برآورد پارامترهای ρ و α نسبتاً نا اریب هستند.

مراجع

- [۱] رزمخواه، م. و خطیب آستانه، ب. (۱۴۰۰) قابلیت اعتماد F سیستم فرسایشی k از n تسهیم بار، علوی، س. م. و تاج، م. (۱۳۹۴)، روش پاسخ تصادفیه و مقایسه آن با روش سیمونس، مجله مهندسی و مدیریت کیفیت، ۱۱، ۱۹۰-۱۷۷.
- [2] Barlow, R. E., and Proschan, F. (1965). *Mathematical theory of reliability*, Joun Wiley, New York.
- [3] Castro I. T. and Mercier S. Performance measures for a deteriorating system subject to imperfect maintenance and delayed repairs, *Proceedings of the Institution of Mechanical Engineers*, Part O: *Journal of Risk and Reliability*, **230**(4): 364-377, 2016.

جدول ۱: MSE برآورد پارامترها تحت مدل ARD_1 برای $\alpha = 1$, $\beta = 1/2$, $\eta = 1$

پیشنهادی				اول						
$\hat{\rho}$	$\hat{b}$	$\hat{\beta}$	$\hat{\alpha}$	$\hat{\rho}$	$\hat{b}$	$\hat{\beta}$	$\hat{\alpha}$	ρ	n	s
۰/۲۵۶۱	۰/۳۶۸۲	۰/۲۱۹۳	۰/۲۸۱۶	۰/۳۲۶۱	۲/۱۹۲۷	۱/۴۹۸۵	۰/۳۲۹۲	۰/۲	۵	۱۰
۰/۰۹۰۰	۰/۱۸۵۴	۰/۳۱۳۶	۰/۴۳۸۲	۰/۱۳۴۰	۰/۶۰۸۳	۲/۳۱۱۳	۰/۵۵۴۶	۰/۵		
۰/۰۸۸۹	۰/۱۳۶۳	۰/۲۲۰۵	۰/۲۲۴۱	۰/۱۵۰۱	۰/۹۶۲۴	۱/۲۲۹۸	۰/۲۵۷۵	۰/۸		
۰/۲۵۱۴	۰/۳۶۳۸	۰/۲۰۷۸	۰/۲۴۲۴	۰/۳۱۳۶	۱/۳۹۹۸	۳/۹۴۳۴	۰/۳۵۴۹	۰/۲	۱۰	
۰/۰۷۰۴	۰/۱۶۳۱	۰/۲۳۸۳	۰/۳۷۸۰	۰/۱۱۴۱	۱/۷۷۱۴	۴/۷۵۰۰	۰/۵۲۱۶	۰/۵		
۰/۰۷۹۶	۰/۱۰۱۷	۰/۱۲۳۳	۰/۱۷۴۳	۰/۱۲۰۲	۱/۵۶۱۷	۲/۷۰۹۲	۰/۳۲۹۸	۰/۸		
۰/۱۶۵۶	۰/۲۱۰۷	۰/۰۶۱۱	۰/۰۲۹۰	۰/۲۳۱۱	۰/۳۵۳۳	۰/۲۶۵۰	۰/۰۴۰۴	۰/۲	۵	۵۰
۰/۰۵۸۲	۰/۱۰۲۵	۰/۰۵۳۰	۰/۰۴۲۹	۰/۱۲۵۴	۰/۱۶۴۸	۰/۳۵۱۷	۰/۰۷۹۸	۰/۵		
۰/۱۲۴۲	۰/۰۸۴۸	۰/۰۵۵۲	۰/۰۴۲۶	۰/۱۳۱۰	۰/۱۲۷۸	۰/۲۸۱۱	۰/۰۵۸۰	۰/۸		
۰/۱۶۱۹	۰/۱۹۷۰	۰/۰۷۸۱	۰/۰۳۴۴	۰/۲۴۴۸	۰/۴۲۰۰	۰/۴۰۳۲	۰/۰۴۰۴	۰/۲	۱۰	
۰/۰۴۹۰	۰/۰۸۸۶	۰/۰۴۹۸	۰/۰۵۸۵	۰/۰۷۳۵	۰/۱۵۶۵	۰/۵۲۴۸	۰/۱۱۱۵	۰/۵		
۰/۰۵۸۱	۰/۰۵۰۸	۰/۰۳۸۲	۰/۰۲۹۳	۰/۰۹۵۳	۰/۰۷۵۳	۰/۶۴۱۰	۰/۰۹۰۷	۰/۸		
۰/۱۲۸۲	۰/۱۷۲۱	۰/۰۶۳۱	۰/۰۱۷۹	۰/۱۸۵۲	۰/۲۸۱۴	۰/۱۵۴۱	۰/۰۲۲۴	۰/۲	۵	۱۰۰
۰/۰۵۴۶	۰/۰۸۸۷	۰/۰۴۱۴	۰/۰۳۰۲	۰/۱۱۱۷	۰/۱۰۵۷	۰/۲۱۶۳	۰/۰۶۹۰	۰/۵		
۰/۱۲۲۳	۰/۰۶۶۸	۰/۰۴۲۵	۰/۰۳۰۳	۰/۱۲۹۴	۰/۰۹۹۵	۰/۲۰۴۴	۰/۰۴۳۷	۰/۸		
۰/۱۵۸۱	۰/۱۹۵۴	۰/۰۸۹۲	۰/۰۲۶۵	۰/۲۴۵۶	۰/۴۳۸۱	۰/۳۳۳۳	۰/۰۳۷۵	۰/۲	۱۰	
۰/۰۴۴۱	۰/۰۷۹۵	۰/۰۴۴۵	۰/۰۲۴۷	۰/۰۸۷۴	۰/۱۵۱۷	۰/۲۶۳۷	۰/۰۷۹۵	۰/۵		
۰/۰۷۸۷	۰/۰۴۹۲	۰/۰۳۳۹	۰/۰۲۷۱	۰/۱۰۵۶	۰/۱۰۲۳	۰/۵۱۱۱	۰/۰۹۷۴	۰/۸		

- [4] Doyen L. and Gaudoin O. Classes of imperfect repair models based on reduction of failure intensity or virtual age, *Reliability Engineering and System Safety*, **84**(1): 45-56, 2004.
- [5] Klein J. and Moeschberger M. Survival Analysis, Techniques for Censored and Truncated Data, *Springer-Verlag*, New York, 2003.
- [6] Meeker W. Q. and Escobar L. A. *Statistical methods for reliability data*, Wiley series in probability and statistics, *Journal Wiley & Sons*, New York, Chichester, 1998.
- [7] Mercier S. and Castro I. T. Stochastic comparisons of imperfect maintenance models for a gamma deteriorating system, *European Journal of Operational Research*, **273**(1): 237-248, 2019.
- [8] Liu R. and Lux T. Generalized method of moment estimation of multivariate multifractal models, *Economic Modelling*, **67**:136-148, 2017.

- [9] Pham H. and Wang H. Imperfect maintenance, *European Journal of Operational Research*, **94**: 438-52, 1996.
- [10] Salles G. *On the modelling and statistical analysis of a gamma deteriorating system with imperfect maintenance*, A thesis submitted for the degree of Doctor of Philosophy in Mathematics, 2020.
- [11] VAN NOORTWIJK J. A Survey of the application of Gamma processes in maintenance, *Reliability Engineering and System Safety*, **94**, 1, pp. 2–21, 2009.
- [12] Wang X. and Xu D. An inverse Gaussian process model for degradation data, *Technometrics* **52**(2):188–197, 2010.



یازدهمین سمینار تخصصی
نظریه قابلیت اعتماد و کاربردهای آن

۲۴ و ۲۵ اردیبهشت ۱۴۰۴



رویکرد مبتنی بر یادگیری تقویتی عمیق برای برنامه‌ریزی پویا در تعمیر و نگهداری سیستم‌های تخریب‌پذیر

حسن تبار درزی، ف^۱ اسلامی مهدی‌آبادی، م^۲

^۱ گروه آمار، دانشکده ریاضی، دانشگاه سیستان و بلوچستان

^۲ گروه برق و کامپیوتر، دانشکده مهندسی، دانشگاه ولایت

چکیده

استفاده از هوش مصنوعی در برنامه‌ریزی تعمیر و نگهداری به بسیاری از صنایع کمک می‌کند تا یک ابزار تصمیم‌گیری هوشمند داشته باشند که بهترین سیاست تعمیر و نگهداری را برای حداقل‌سازی هزینه‌های موردانتظار تعمیرات ارائه دهد. در این مطالعه، برای ارائه مدل تعمیر و نگهداری پویا در سیستم‌های تعمیرپذیر که تحت فرسایش قرار دارد از روش‌های یادگیری تقویتی استفاده شده است. برای مدل‌سازی فرآیند تخریب سیستم، از فرآیند گاما استفاده می‌گردد. مسئله تعمیر و نگهداری به صورت یک فرآیند تصمیم‌گیری مارکوفی مدل‌سازی شده و با استفاده از الگوریتم یادگیری عمیق حل گردیده است. در اغلب مدل‌های موجود در ادبیات، وضعیت تخریب سیستم باید گسسته‌سازی شود. با این حال، گسسته‌سازی منجر به کاهش دقت و کارایی مدل می‌شود. در این مطالعه، به جای گسسته‌سازی، سطح دقیق تخریب به عنوان وضعیت سیستم در نظر گرفته شده است. روش یادگیری عمیق Q به جای

^۱hasantabar_f@math.usb.ac.ir

^۲m.eslami@velayat.ac.ir

نگاشت هر وضعیت به بهترین اقدام ممکن، الگوهای موجود در داده‌ها را شناسایی می‌کند. در طول فرآیند یادگیری، یک شبکه عصبی آموزش داده می‌شود که می‌تواند به عنوان ابزار تصمیم‌گیری برای تیم تعمیر و نگهداری مورد استفاده قرار گیرد تا بهترین اقدام تعمیر و نگهداری را بر اساس سطح تخریب فعلی سیستم تعیین کند. در نهایت، با شبیه‌سازی نشان داده می‌شود که چگونه الگوریتم یادگیری تقویتی عمیق می‌تواند برای یافتن بهینه‌ترین اقدام تعمیر و نگهداری در هر سطح تخریب به کار گرفته شود.

کلمات کلیدی: الگوهای تعمیر و نگهداری، یادگیری تقویتی، یادگیری کمو، یادگیری عمیق Q.

۱ پیش‌گفتار

سیستم‌های صنعتی به طور کلی به دلیل استفاده و قرار گرفتن در معرض عوامل محیطی، دچار تخریب می‌گردند. این تخریب در نهایت منجر به خرابی سیستم شده و در نتیجه مشکلات ایمنی، آسیب به تجهیزات، مشکلات کیفیت و در دسترس نبودن غیرمنتظره دستگاه را به وجود می‌آورد. در دهه‌های گذشته تعمیر و نگهداری^۱ برای سیستم‌های تعمیرپذیر، بعد از خرابی سیستم در نظر گرفته می‌شد که در خیلی از موارد کار تشخیص و تعمیر سخت و غیرممکن می‌گردید. مدل تعمیر و نگهداری مبتنی بر شرایط^۲ (CBM) یکی از پرکاربردترین روش‌ها برای سیستم‌های در حال تخریب است که هدف آن کاهش هزینه‌های سیستم و بهینه‌سازی فراوانی فعالیت‌های تعمیر و نگهداری با توجه به وضعیت واقعی مؤلفه‌ها است. در بیشتر مطالعات تعمیر و نگهداری، تصمیمات نگهداری به صورت ایستا و بر اساس الگوهای تخریب از پیش تعیین شده اتخاذ می‌شوند. در بسیاری از مدل‌های CBM، آستانه‌های بهینه تعمیر و نگهداری تعیین شده و فعالیت‌های نگهداری بر اساس مقایسه سطح تخریب با این آستانه‌ها انجام می‌شود. با این حال، داشتن یک مدل تعمیر و نگهداری پویا به تیم نگهداری کمک می‌کند تا اقدامات مناسب را بر اساس وضعیت لحظه‌ای سیستم اجرا کند و از انجام تعمیرات غیرضروری جلوگیری کند، که این امر می‌تواند از نظر هزینه مقرون به صرفه باشد.

اخیراً استفاده از روش‌های یادگیری تقویتی^۳ (RL) در حل مسائل تعمیر و نگهداری، مورد توجه قرار گرفته است. برای آشنایی با RL و الگوریتم‌های مختلف آن، خوانندگان به کتاب ارزشمند [۱۱] ارجاع داده می‌شوند. از جمله مطالعاتی که از روش‌های RL در زمینه تعمیر و نگهداری سیستم‌ها استفاده کرده است می‌توان به موارد زیر اشاره کرد. یوسفی و همکاران [۱۳] مسئله CBM را به صورت یک فرایند تصمیم‌گیری مارکف مدل کرده سپس از الگوریتم یادگیری

^۱Maintenance

^۲Condition based maintenance

^۳Reinforcement Learning

Q در RL برای تصمیم‌گیری در مورد بهترین تعمیر و نگهداری سیستم استفاده کردند. همچنین در مطالعه دیگری یوسفی و همکاران [۱۴] از یادگیری تقویتی عمیق برای تعمیر و نگهداری پویا در سیستم‌های چند جزئی قابل تعمیر استفاده کردند. پنگ [۷] مسئله CBM را به عنوان یک فرآیند تصمیم‌گیری مارکوف پیوسته با فرآیند گاوسی مدل‌سازی کرد. در این مطالعه او بدون گسسته‌سازی حالت تخریب سیستم از الگوریتم‌های RL برای تصمیم‌گیری تعمیر و نگهداری استفاده کرد. الگوهای تعمیر و نگهداری سیستم‌ها با توجه به پیشرفت بشر در صنایع پیچیده مدت‌ها است مورد توجه می‌باشد که می‌توان به مطالعات [۸]، [۲] و [۴] و ارجاعات آن اشاره کرد. اما استفاده از روش‌های RL در الگوهای تعمیر و نگهداری اخیراً مورد توجه قرار گرفته است و توانسته تأثیرات شگرفی در این زمینه ایجاد کند. استفاده از RL در الگوهای تعمیر و نگهداری سیستم‌ها را می‌توان در کارهای [۶]، [۱۵] و [۱] و [۳] بررسی کرد.

یکی از الگوریتم‌های مؤثر در یادگیری تقویتی، الگوریتم یادگیری Q است که نیازی به دانش قبلی درباره محیط و انتقال حالات ندارد. در این روش، عامل یادگیرنده مسیرهای جدیدی را کشف کرده یا مسیرهای دارای بیشترین پاداش را دنبال می‌کند. با این حال، هنگامی که تعداد حالات و اقدامات افزایش یابد، کارایی الگوریتم یادگیری Q کاهش می‌یابد؛ زیرا هم نیاز به حافظه برای ذخیره جدول Q افزایش می‌یابد و هم زمان مورد نیاز برای یادگیری تمامی حالات به صورت عملیاتی غیرممکن می‌شود. به همین دلیل، در این مطالعه از یادگیری عمیق Q استفاده شده است. این روش ترکیبی از یادگیری Q و یادگیری عمیق است که در حل مسائل با تعداد حالات و اقدامات بالا کاربرد دارد. در این مطالعه، یادگیری عمیق Q برای بهینه‌سازی سیاست تعمیر و نگهداری سیستم‌های در حال فرسایش به کار گرفته شده است. فرآیند تخریب سیستم با استفاده از فرآیند گاما مدل‌سازی شده و در هر زمان بازرسی، اقدام بهینه با هدف حداقل‌سازی هزینه‌های تعمیر و نگهداری پیشنهاد می‌شود.

در ادامه در بخش ۲ مدل یادگیری عمیق Q در RL توضیح داده می‌شود. الگوریتم یادگیری عمیق Q برای سیستم‌های با وضعیت پیچیده در بخش ۳ بیان شده است. سپس مدل‌های تعمیر و نگهداری، به عنوان یک فرایند تصمیم‌گیری مارکف در بخش ۴ فرمول‌بندی می‌گردد تا با استفاده از روش‌های RL، بهترین تصمیم در هر بازرسی اتخاذ گردد. سپس در بخش ۵ نتایج مطالعات و شبیه‌سازی برای سیستم دو مؤلفه‌ای با پیکربندی موازی و سری ارائه می‌شود. نتیجه‌گیری و پیشنهادات در بخش ۶ ارائه می‌شود.

۲ الگوریتم یادگیری Q

الگوریتم یادگیری Q یکی از روش‌های مهم یادگیری تقویتی است که برای حل مسائل تصمیم‌گیری بهینه در محیط‌های نامعین استفاده می‌شود. این الگوریتم بدون نیاز به مدل محیط، از طریق تجربه و تعامل با محیط، با استفاده از پاداش دریافتی از محیط به سیاست بهینه برای انتخاب اقدام‌های مناسب در هر وضعیت دست می‌یابد. در این الگوریتم، جدول Q به‌عنوان یک ماتریس ذخیره‌سازی مقدار پاداش مورد انتظار هر اقدام در هر وضعیت به‌روزرسانی می‌شود. در واقع این الگوریتم به دنبال یادگیری تابع Q بهینه (Q^*) با بیشینه کردن مقدار مورد انتظار از پاداش تجمعی آینده است. الگوریتم یادگیری Q برای مسائلی استفاده می‌شود که فضای حالات آنها محدود است و عامل می‌تواند مقدار Q ی هر عمل در هر حالت را در یک جدول نگهداری کند که در آن هر سطر متناظر با یک حالت s و هر ستون متناظر با یک عمل a است. هر خانه‌ی این جدول، مقدار Q برای آن زوج حالت-عمل را ذخیره می‌کند. عامل در هر گام زمانی t ، حالت فعلی s_t را مشاهده می‌کند، بر اساس سیاست اکتشاف، عمل a_t را انتخاب می‌کند، آن را در محیط اجرا می‌کند، پاداش r_{t+1} را دریافت می‌کند و به حالت بعدی s_{t+1} منتقل می‌شود. سپس، مقدار Q برای زوج حالت-عمل (s_t, a_t) را با استفاده از معادله‌ی به‌روزرسانی (۱) به‌روزرسانی می‌شود.

$$Q(s_t, a_t) \leftarrow Q(s_t, a_t) + \alpha [r_{t+1} + \gamma \max_{a'} Q(s_{t+1}, a') - Q(s_t, a_t)], \quad (1)$$

که در آن $Q(s_t, a_t)$ مقدار Q برای حالت s_t و عمل a_t ، α نرخ یادگیری، r_{t+1} پاداش دریافتی پس از انجام عمل a_t در حالت s_t ، γ ضریب تخفیف، s_{t+1} حالت بعدی و $\max_{a'} Q(s_{t+1}, a')$ بیشینه‌ی مقدار Q برای تمام عمل‌های ممکن در حالت s_{t+1} می‌باشند.

مهمترین محدودیت الگوریتم یادگیری Q این است که برای مسائلی که فضای حالت و عمل آنها زیاد و نامحدود است، نمایش جدولی امکان‌پذیر نبوده و این الگوریتم کارایی ندارد. بنابراین نیاز است از یک تخمین‌گر برای تخمین مقادیر جدول استفاده گردد که شبکه‌ی عمیق Q این محدودیت را برطرف می‌سازد.

۳ الگوریتم یادگیری عمیق Q

همان‌طور که بیان شد، الگوریتم یادگیری Q سنتی برای مسائلی با فضای حالت و عمل کوچک و گسسته مناسب است. با افزایش ابعاد فضای حالت و عمل، اندازه‌ی جدول Q به‌طور نمایی افزایش می‌یابد و ذخیره‌سازی و به‌روزرسانی آن غیرعملی می‌گردد. علاوه بر این، یادگیری Q سنتی قادر به تعمیم دانش بین حالت‌های مشابه نیست. توسعه الگوریتم

یادگیری Q و استفاده از شبکه عصبی برای تخمین تابع Q منجر به ارائه الگوریتم یادگیری عمیق Q ^۵ (DQN) شده است که اولین بار در یادگیری بازی‌های آتاری استفاده گردید و نتایج درخشانی به دنبال داشت.

در الگوریتم DQN، هر حالت محیط به عنوان ورودی شبکه و خروجی شبکه، مقدار Q برای هر عمل در آن حالت در نظر گرفته می‌شود. معماری شبکه‌ی عصبی مورد استفاده در DQN می‌تواند بسته به پیچیدگی مسئله متفاوت باشد. با این حال، معمولاً شامل چندین لایه‌ی پنهان است که می‌توانند از نوع تماماً متصل یا کانولوشنی باشد. لایه‌ی خروجی یک لایه‌ی تماماً متصل با تعداد نوروتهایی برابر با تعداد عمل‌های ممکن است. تابع فعال‌سازی لایه‌ی خروجی معمولاً خطی است، زیرا خروجی شبکه باید بتواند مقادیر Q را برای هر عمل تخمین بزند. DQN علاوه بر استفاده از شبکه‌ی عصبی عمیق، از دو تکنیک کلیدی دیگر برای بهبود پایداری و عملکرد آموزش بهره می‌برد، که عبارتند از:

- حافظه‌ی تکرار ^۶: تجربیات عامل (گذارها) به جای اینکه بلافاصله پس از دریافت، برای آموزش استفاده شوند، در یک حافظه‌ی تکرار ذخیره می‌شوند. در هر مرحله‌ی آموزش، یک دسته‌ی کوچک از گذارها به طور تصادفی از این حافظه نمونه‌برداری شده و برای به‌روزرسانی شبکه‌ی Q مورد استفاده قرار می‌گیرند. این امر به کاهش همبستگی بین نمونه‌ها و بهبود پایداری آموزش کمک می‌کند.

- شبکه‌ی هدف ^۷: برای کاهش ناپایداری ناشی از به‌روزرسانی‌های مکرر Q -value ها، از یک شبکه‌ی هدف جداگانه برای محاسبه‌ی مقادیر Q هدف در معادله‌ی به‌روزرسانی استفاده می‌شود. وزن‌های شبکه‌ی هدف به صورت دوره‌ای (مثلاً هر N گام) با کپی کردن وزن‌های شبکه‌ی Q اصلی به‌روزرسانی می‌شوند.

در شکل ۱ معماری شبکه‌ی عمیق Q با حافظه‌ی تکرار و شبکه‌ی هدف آمده است. برای آموزش شبکه‌ی Q ابتدا نمونه‌ای از گذرها از حافظه‌ی تکرار به تصادف انتخاب شده و سپس بوسیله‌ی رابطه‌ی (۱) میزان تابع زیان شبکه محاسبه می‌شود.

$$Loss = (Q_{predicted} - Q_{target})^2 \quad (1)$$

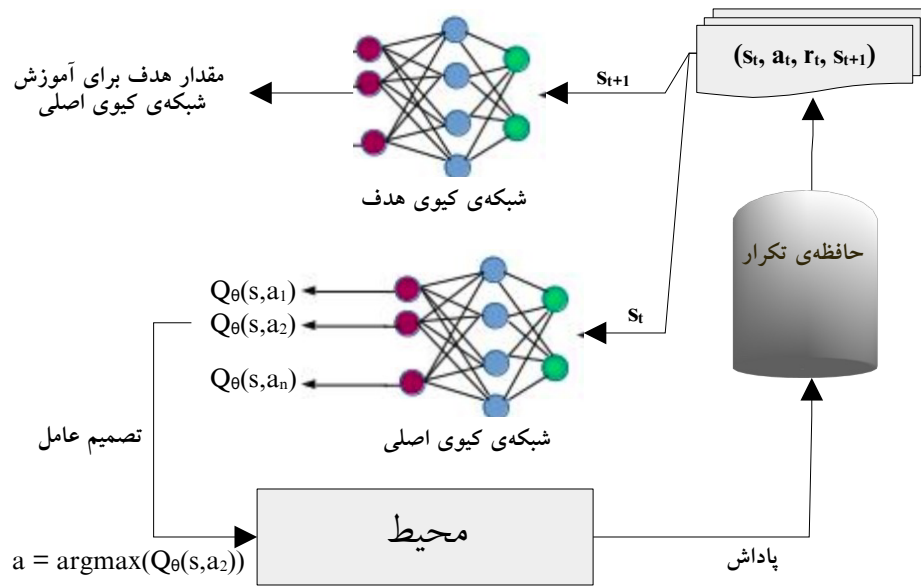
که، $Q_{predicted}$ ماکسیم خروجی شبکه عصبی اصلی است و شبکه کیوی هدف (Q_{target}) به صورت زیر بدست می‌آید،

$$Q_{target} = r_{t+1} + \gamma \max_{a'} Q(s_{t+1}, a'). \quad (2)$$

^۵Deep Q-learning

^۶Replay Buffer

^۷Target Network



شکل ۱: معماری شبکه‌ی عمیق Q

اگر پارامترهای شبکه عصبی θ در نظر گرفته شود با کمینه کردن رابطه‌ی (۳) و با استفاده از روش گرادیان نزولی پارامترهای مدل DQN به‌روز می‌شوند.

$$L(\theta) = \mathbb{E}_{(s,a,r,s') \sim D} [(r_{t+1} + \gamma \max_{a'} Q_{target}(s_{t+1}, a'; \theta^-) - Q(s, a; \theta))^2] \quad (3)$$

که، D حافظه‌ی تکرار است. دو چالش اصلی در آموزش DQN، همگرایی و پایداری مدل می‌باشد. در الگوریتم یادگیری Q سنتی، در هر تکرار، تنها مقدار یک خانه‌ی جدول Q به‌روزرسانی می‌شد و این تغییر، تأثیری بر سایر خانه‌های جدول نداشت. اما در DQN، با هر بار به‌روزرسانی پارامترهای شبکه، تخمین تمام مقادیر Q به طور ضمنی تغییر می‌کند، که این امر می‌تواند منجر به ناپایداری در فرآیند آموزش شود. از طرف دیگر، به دلیل استفاده از شبکه‌ی عصبی عمیق، DQN برای آموزش مؤثر به تعداد زیادی داده نیاز دارد و همگرایی آن می‌تواند بسیار چالش‌برانگیز به‌خصوص در محیط‌های پیچیده و پرنویز باشد. راهکارهای مختلفی برای بهبود همگرایی و پایداری DQN پیشنهاد شده است که یکی از موثرترین آن‌ها، شبکه‌ی عمیق Q دوگانه (DDQN) است که خوانندگان برای مطالعه آن می‌توانند به [۵] رجوع کنند.

۴ بیان مسئله CBM با استفاده از فرایند تصمیم‌گیری مارکف

مؤلفه‌های سیستم معمولاً با گذشت زمان دچار فرسودگی شده و کم‌کم بسوی زوال و تخریب پیش می‌روند. فرسودگی مؤلفه‌های سیستم را می‌توان با استفاده از فرایند گاما مدل کرد. به عبارتی اگر فرسایش مؤلفه i ام در طول زمان با $X_i[t]$ نمایش داده شود، فرض می‌گردد که دارای توزیع گاما بوده و بصورت رابطه (۱) بدست می‌آید،

$$g(x_i; \alpha_i(t-s), \beta_i) = \frac{\beta_i^{\alpha_i(t-s)} x_i^{\alpha_i(t-s)-1} e^{-\beta_i x_i}}{\Gamma(\alpha_i(t-s))}, \quad t > 0, s > 0 \quad (1)$$

که α_i و β_i به ترتیب پارامترهای شکل و مقیاس مؤلفه i ام هستند، و $\Gamma(\cdot)$ تابع گاما می‌باشد. به ازای هر $t > 0$ و $s > 0$ $(X(t) - X(s)) \sim \text{gamma}(\alpha(t) - \alpha(s), \beta)$ می‌باشد.

در الگوهای تعمیر و نگهداری، از سطح فرسایش مؤلفه‌های سیستم برای تصمیم‌گیری برای اقدامات تعمیر و نگهداری استفاده می‌شود. برای حل مسئله الگوهای تعمیر و نگهداری با استفاده از الگوریتم DQN باید آن را به صورت یک فرایند مارکف بیان کرد. لذا فرض می‌شود که مؤلفه‌های سیستم در طول زمان فرسوده شده و می‌تواند به صورت جداگانه نگهداری شود. وضعیت فرسایش هر مؤلفه در زمان بازرسی از قبل تعیین شده (روز، هفته، ماه و ...) مقدار $s_t \in S$ را خواهد داشت که مقداری پیوسته از فرسایش سیستم در زمان t است.

با توجه به اینکه فرایند تصمیم‌گیری مارکف به صورت چندگانه (S, A, P, R, γ) تعریف می‌گردد برای مدل کردن مسئله CMB بصورت فرایند تصمیم‌گیری مارکف داریم:

۱. فضای حالت سیستم S می‌باشد که در هر زمان بازرسی t ، $s_t \in S$.

۲. فضای عمل به صورت $A = \{a_1, a_2, a_3\}$ در نظر گرفته شده است و به صورت زیر تعریف می‌گردد،

$$A = \begin{cases} 0 & \text{هیچ اقدامی} \\ 1 & \text{تعمیر} \\ 2 & \text{جایگزینی} \end{cases}$$

فرض بر این است که تعمیر، فرسایش مؤلفه‌ام i ام، را کاهش می‌دهد. در واقع فرسایش سیستمی که تعمیر یافته شده در زمان $t+1$ بصورت $X_{(t+1)} = vX_t + \frac{\beta^{\alpha_i \tau}}{\Gamma(\alpha_i \tau)} X_t^{\alpha_i \tau} e^{(-\beta X_t)}$ محاسبه می‌گردد که v ، ضریبی برای کاهش فرسایش در اثر تعمیر، τ فاصله زمانی بین دو بازرسی می‌باشد. همچنین اگر جایگزینی را برای مؤلفه i ام انجام دهیم، سیستم به وضعیت اولیه (وضعیت صفر) با احتمال یک انتقال می‌یابد و همانند یک سیستم تازه نصب شده رفتار می‌کند. فرض می‌شود که سطح تخریب مؤلفه‌های سیستم نتواند در وضعیت صفر (اولیه) بماند.

۳. تابع پاداش بر اساس منفی تابع هزینه تعریف شده است. برای نمایش وضعیت خرابی مؤلفه، متغیر تصادفی باینری w در نظر گرفته می‌شود که در صورت تخریب و از کار افتادگی مؤلفه، مقدار آن یک و در صورت کارکرد سیستم مقدار آن صفر می‌باشد. برای مجموعه عمل‌های تعمیر و نگهداری مؤلفه i ام سیستم $(\sigma_{i0}, \sigma_{i1}, \sigma_{i2})$ متغیرهای تصادفی باینری می‌باشند که برای هیچ اقدامی، σ_{i0} ، تعمیر، σ_{i1} و جایگزینی، σ_{i2} که $\sum_{j=1}^2 \sigma_{ij} = 1$ می‌باشد. هزینه کل تعمیر و نگهداری (C_t) و تابع پاداش (R_t) به ترتیب در رابطه‌ی (۲) و (۱) آورده شده است؛

$$C_t = \sum_{i=1}^n (C_m \sigma_{i2} + C_r \sigma_{i1}) + C_p w, \quad (2)$$

$$R_t = -C_t, \quad (3)$$

که C_m هزینه تعمیر، C_r هزینه جایگزینی و C_p هزینه جریمه برای خرابی در نظر گرفته شده است.

ارائه بهترین سیاست تعمیر و نگهداری که حداقل هزینه را طی زمان در بر داشته باشد با استفاده از الگوریتم‌های DQN برای مدل CMB در ادامه انجام شده است.

۵ دست‌آوردهای پژوهش

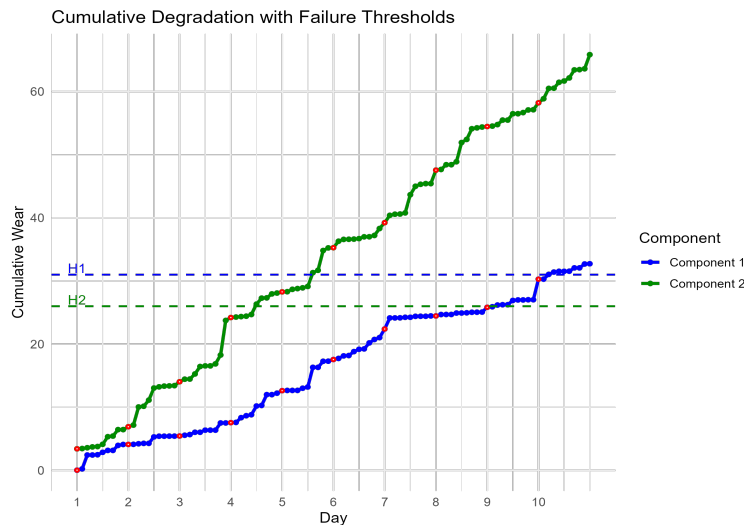
برای یافتن سیاست بهینه با الگوریتم DQN سیستمی دو مؤلفه‌ای با پیکربندی سری در نظر گرفته شده است. از طرفی در مسئله CBM فرض شده است که هر مؤلفه در طول زمان فرسوده می‌شود و می‌تواند به صورت جداگانه در سیستم جایگزین یا تعمیر شوند یا هیچ اقدامی برای آن انجام نگردد. همچنین بازه‌های بازرسی ثابت (به عنوان مثال ماه، هفته و ...) در نظر گرفته شده است. فرسایش سیستم در طول زمان با فرایند گاما مدل می‌گردد و پارامترهای فرایند گاما و آستانه تخریب برای دو مؤلفه سیستم در جدول ۱ داده شده است. همچنین پارامترهای تابع هزینه $C_p = 300, C_m = 100$ و $C_r = 1000$ قرار داده شده است. برای شبکه عصبی (اصلی و هدف)، دو لایه پنهان و هر لایه با ۱۲۸ و ۶۴ گره در نظر گرفته شده است. و پارامترهای شبکه عصبی هدف در هر ۵۰۰۰ تکرار به روز رسانی می‌شوند. هدف یافتن اقدامات تعمیر و نگهداری بهینه با الگوریتم DQN برای یک افق زمانی مشخص است.

شکل ۲، فرسایش تدریجی سیستم دو مؤلفه‌ای که از توزیع گاما با پارامترهای فرض شده در جدول ۱ ایجاد شده را نشان می‌دهد. سطح تخریب هر یک از مؤلفه‌ها نیز در نمودار مشخص شده است که در صورت رسیدن مؤلفه به آن مقدار مؤلفه از کار می‌افتد.

جدول ۲ ارزیابی مدل DQN برای سیستمی دو مؤلفه‌ای با پیکربندی سری می‌باشد. ورودی مدل میزان فرسایش مؤلفه‌ها با در نظر گرفتن دو نرون می‌باشد. برای لایه پنهان دو لایه در نظر گرفته شده است که لایه اول ۱۲۸ و لایه دوم

جدول ۱: پارامترهای فرض شده برای مؤلفه‌های دو سیستم

مؤلفه ۲	مؤلفه ۱	پارامتر مدل
۲۶	۳۱	$H[i]$
$(0.6, 1/1)$	$(0.3, 1/2)$	$X_i(t) \sim \Gamma(\alpha_i, \beta)$



شکل ۲: فرسایش سیستم دو مؤلفه‌ای در طول زمان با بازه‌های بازرسی و آستانه تخریب مشخص

۶۴ نرون دارد و لایه خروجی شامل ۹ نرون می‌باشد که هر نرون مقدار تابع ارزش حالت-عمل را نشان می‌دهد. عمل انتخابی بیشینه مقدار لایه آخر می‌باشد.

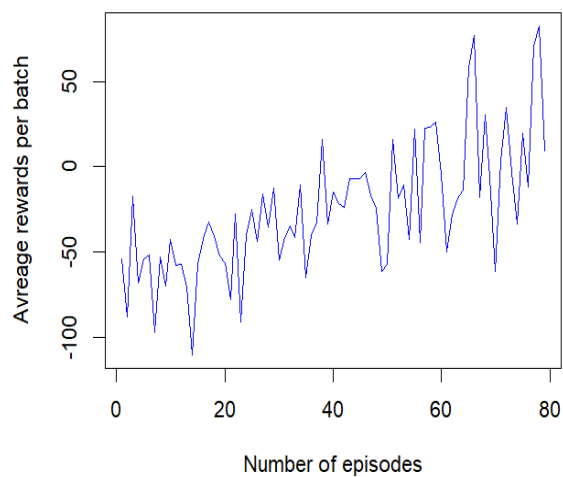
برای بررسی همگرایی مدل DQN شکل ۳ رسم شده است. همان‌طور که این نمودار نشان می‌دهد، پاداش تجمعی با گذشت اپیزودها صعودی می‌باشد که بیانگر یادگیری خوب الگوریتم می‌باشد.

۶ نتیجه‌گیری و پیشنهادات

در این مطالعه، یک سیاست تعمیر و نگهداری پویا برای یک سیستم دو مؤلفه‌ای در حال تخریب پیشنهاد شده است. در این سیستم، هر مؤلفه می‌تواند به صورت مستقل تعمیر شود و تحت یک فرآیند فرسایش قرار دارد. برای مدل‌سازی مسیر فرسایش تمامی مؤلفه‌ها از فرآیند گاما استفاده شده است. اگر میزان تخریب یک مؤلفه از یک آستانه از پیش تعیین‌شده فراتر رود، آن مؤلفه به عنوان خراب‌شده شناسایی می‌شود. سیستم در فواصل زمانی مشخص به‌طور دوره‌ای

جدول ۲: نتایج ارزیابی مدل

عمل انجام شده		سطح فرسایش		شماره
مؤلفه ۱	مؤلفه ۲	مؤلفه ۱	مؤلفه ۲	
۰	۰	$۲/۸۳E + ۰.۰$	$۶/۲۸E - ۰.۲$	۱
۰	۰	$۱/۶۷E - ۰.۱$	$۳/۶۵E + ۰.۰$	۲
۰	۰	$۱/۸۱E - ۲.۱$	$۷/۲۵E - ۳.۹$	۳
۱	۰	$۴/۵۷E + ۰.۱$	$۲/۴۳E - ۰.۵$	۴
۱	۰	$۲/۸۱E - ۰.۴$	$۷/۸۹E - ۲.۰$	۵
۰	۱	$۶/۳۱E + ۰.۰$	$۲/۳۲E - ۰.۶$	۶
۱	۱	$۱/۰۷E + ۰.۰$	$۲/۱۹E - ۰.۴$	۷
۲	۰	$۱/۳۱E + ۰.۱$	$۴/۲۱E - ۰.۱$	۸
۱	۰	$۱/۴۶E + ۰.۱$	$۱/۵۷E - ۰.۱$	۹
۲	۲	$۱/۸۰E + ۰.۱$	$۱/۱۹E + ۰.۱$	۱۰



شکل ۳: نمودار میانگین پاداش تجمعی هر اپیزود

بازرسی می‌شود و در هر بازرسی، بسته به سطح تخریب مؤلفه‌ها، اقدامات تعمیر و نگهداری و یا عدم انجام اقدامی می‌تواند اعمال شود.

از مزیت‌های مدل استفاده شده استفاده از سطح دقیق فرسایش سیستم می‌باشد. در واقع مدل‌های DQN زمانیکه تعداد حالت‌ها و اقدامات زیاد باشند عملکرد قابل قبولی دارد. الگوریتم DQN یکی از الگوریتم‌های پرکاربرد در یادگیری تقویتی است که با ترکیب یادگیری تقویتی و یادگیری عمیق، امکان حل مسائل با تعداد نامحدودی از حالات یا اقدامات را فراهم می‌کند. البته باید به این نکته اشاره کرد که پیاده‌سازی و اجرای مدل زمان زیادی را لازم دارد بطوریکه با کامپیوترهای خانگی زمان زیادی صرف یادگیری می‌گردد.

برای ادامه مطالعه الگوریتم در حالت های پیچیده‌تری با تعداد اقدامات بیشتر و در نظر گرفتن لایه‌های پنهان بالاتری مورد ارزیابی و بررسی قرار خواهد گرفت.

مراجع

- [1] Adsule, A., Kulkarni, M. and Tewari, A., (2020), Reinforcement learning for optimal policy learning in condition-based maintenance, *IET Collaborative Intelligent Manufacturing*, **2**(4), 182-188.
- [2] Chen N, Ye ZS, Xiang Y, Zhang L., (2015) Condition-based maintenance using the inverse Gaussian degradation model, *European Journal of Operational Research*, **243**(1), 190-9.
- [3] Dehghani Ghobadi, Z., Haghighi, F., and Safari, A., (2021), An overview of reinforcement learning and deep reinforcement learning for condition-based maintenance, *International Journal of Reliability, Risk and Safety: Theory and Application*, **4**(2), 81-89.
- [4] Gruber, A., Yanovski, S., & Ben-Gal, I., (2013), Condition-based maintenance via simulation and a targeted Bayesian network meta model, *Quality Engineering*, **25**(4), 370-384.
- [5] Van Hasselt, H., Guez, A., & Silver, D. (2016). Deep reinforcement learning with double q-learning. In Proceedings of the AAAI conference on artificial intelligence (Vol. 30, No. 1).
- [6] Huang, J., Chang, Q. and Arinez, J., (2020), Deep reinforcement learning based preventive maintenance policy for serial production lines, *Expert Systems with Applications*, **160**, 113701.

- [7] Peng, S., (2021), Reinforcement learning with Gaussian processes for condition-based maintenance, *Computers & Industrial Engineering*, **158**, 107321.
- [8] Rafiee K, Feng Q, Coit DW., (2015), Condition-based maintenance for repairable deteriorating systems subject to a generalized mixed shock model, *IEEE Transactions on Reliability*, **64**(4), 1164-1174.
- [9] Rausand, M., and Hoyland, A., (2003), *System reliability theory: models, statistical methods, and applications*, John Wiley & Sons.
- [10] Rasmekomen, N. and Parlikad, A. K. ,(2016), Condition-based maintenance of multi-component systems with degradation state-rate interactions, *Reliability Engineering & System Safety*, **148**, 1-10.
- [11] Sutton, R. S., and Barto, A. G., (2014), *Reinforcement learning: An introduction*, Cambridge: MIT press.
- [12] Vanli, O. A., (2014), A failure time prediction method for condition-based maintenance, *Quality Engineering*, **26**(3), 335-349.
- [13] Yousefi, N., Tsianikas, S. and Coit, D.W., (2020), Reinforcement learning for dynamic condition-based maintenance of a system with individually repairable components, *Quality Engineering*, **32**(3), 388-408.
- [14] Yousefi, N., Tsianikas, S. and Coit, D.W., (2022), Dynamic maintenance model for a repairable multi-component system using deep reinforcement learning. *Quality Engineering*, **34**(1), 16-35.
- [15] Zhang, N. and Si, W., (2020), Deep reinforcement learning for condition-based maintenance planning of multi-component systems under dependent competing risks, *Reliability Engineering & System Safety*, **203**, 107094.



یازدهمین سمینار تخصصی
نظریه قابلیت اعتماد و کاربردهای آن
۲۴ و ۲۵ اردیبهشت ۱۴۰۴



نگهداری و تعمیر هوشمند با استفاده از یادگیری تقویتی عمیق

حقیقی، ف^۱ صفری، ع^۲ عیدی، ش^۳

گروه آمار، دانشکده ریاضی، آمار و علوم کامپیوتر، دانشکده علوم، دانشگاه تهران

چکیده

در این پژوهش، یک چارچوب نوین برای بهینه‌سازی نگهداری و تعمیر سیستم‌های تک‌جزئی تحت تأثیر تخریب تدریجی و شوک‌های تصادفی غیرکشنده ارائه می‌شود. برخلاف روش‌های مرسوم، این روش ترکیبی از دو مؤلفه اصلی است: ۱. مدل تخریب، مبتنی بر سن مجازی است که پیری تسریع شده ناشی از فرسودگی مداوم و آسیب‌های ناشی از شوک را مدلسازی می‌کند. ۲. مکانیسم نگهداری و تعمیر ناقص پویا مبتنی بر عوامل بهبود وابسته به وضعیت سیستم هستند که با شرایط لحظه‌ای آن سازگار می‌شوند. فرآیند تصمیم‌گیری نگهداری و تعمیر به‌عنوان یک فرآیند تصمیم‌گیری مارکوف (MDP) با حالت پیوسته و سه اقدام ممکن مدلسازی شده است: عدم اقدام، تعمیر ناقص، یا تعویض سیستم. برای مدیریت کارآمد فضای حالت چندبُعدی و تعاملات پیچیده بین تخریب و شوک‌ها، از یادگیری تقویتی عمیق (DRL) با استفاده از یک شبکه عمیق Q (DQN) بهره گرفته‌ایم. مطالعات عددی نشان می‌دهند که این چارچوب در مقایسه با سیاست‌های سنتی، کاهش قابل توجهی در هزینه‌های بلندمدت نگهداری و

^۱fhaghighii@ut.ac.ir

^۲a.safari@ut.ac.ir

^۳shaqayq.eidi@gmail.com

تعمیر ایجاد می‌کند و همزمان قابلیت اطمینان سیستم را بهبود می‌بخشد. این پژوهش با معرفی یک رویکرد انطباقی و داده‌محور برای سیستم‌های تحت تأثیر تخریب چندحالتی، به حوزه نگهداری پیش‌بینانه کمک می‌کند.

کلمات کلیدی: فرسایش، نگهداری و تعمیر، تعمیر ناقص، فاکتور بهبود، فرایند تصمیم‌گیری مارکوف، یادگیری تقویتی عمیق

۱ مقدمه

در مسائل بهینه‌سازی نگهداری و تعمیرات با استفاده از الگوریتم‌های یادگیری تقویتی عمیق (DRL) یک عامل تصمیم‌گیر در هر گام زمانی یکی از سه عمل "عدم انجام عمل"، "تعمیر ناقص" یا "تعویض" را انتخاب می‌کند. عمل "عدم انجام عمل" به معنای عدم انجام هیچ‌گونه تعمیر، عمل "تعویض" به معنای جایگزینی کامل سیستم با یک سیستم جدید و عمل "تعمیر ناقص" به معنای بهبود نسبی وضعیت سیستم بین حالت نو و کهنه است.

برای مدل‌سازی اثر تعمیر ناقص بر وضعیت سیستم، مدل‌های مختلفی توسعه داده شده‌اند. مدل‌های کیجیما از جمله پرکاربردترین این مدل‌ها هستند که در آنها تعمیر به صورت یک متغیر تصادفی بین صفر و یک نشان داده می‌شود. در مدل اول کیجیما، تعمیر، سن اضافه شده به سیستم از آخرین تعمیر را کاهش می‌دهد، در حالی که در مدل دوم، تعمیر بر کل سن انباشته شده سیستم تأثیر می‌گذارد. مدل دیگر، مدل فاکتور بهبود مالیک است که در آن هنگامی که قابلیت اطمینان سیستم از حد مشخصی فراتر رود، یک عمل نگهداری و تعمیر انجام می‌شود که سن تقویمی سیستم را به سن مؤثر آن کاهش می‌دهد.

در بیشتر پژوهش‌های مبتنی بر مدل فاکتور بهبود، فرض بر این است که تأثیر عمل تعمیر ناقص بر سیستم توسط یک فاکتور بهبود توصیف می‌شود که از توزیعی با بازه [۰، ۱] مانند توزیع بتا، یکنواخت یا نرمال قطع شده پیروی می‌کند. بر اساس این فرض، مقادیر بهبود حاصل از اعمال تعمیر ناقص، تصادفی اما همگن هستند. با این حال، در کاربردهای عملی، واقع‌بینانه‌تر است که در نظر بگیریم میزان بهبود تعمیر به نیاز واقعی سیستم که توسط سطح تخریب آن در زمان تعمیر تعیین می‌شود، بستگی دارد.

در این مطالعه، سه مدل متمایز پیشنهاد می‌شود که هر یک برای نشان دادن تأثیر تعمیر و نگهداری بر سلامت و تخریب سیستم در شرایط خاص طراحی شده‌اند. مدل اول فرض می‌کند که یک تعمیر ثابت، بدون توجه به وضعیت واقعی سیستم انجام می‌شود و معمولاً زمانی استفاده می‌شود که ارزیابی آسیب سیستم امکان‌پذیر نباشد. مدل دوم و سوم اجازه می‌دهند تا اعمال تعمیر و نگهداری بر اساس وضعیت سیستم تنظیم شوند. به عنوان مثال، در مورد یک کولر آبی که صدای نامتعارف تولید می‌کند، ممکن است در برخی موارد فقط روغن‌کاری انجام شود که با مدل اول مطابقت

دارد، اما اگر بازرسی کامل امکان‌پذیر باشد و آسیب تسمه مشخص شود، ممکن است هم روغن‌کاری و هم تعویض تسمه انجام شود که با مدل دوم یا سوم مطابقت دارد. تفاوت مدل دوم و سوم در ماهیت اثرات تعمیر است: مدل دوم اثر قطعی دارد، در حالی که مدل سوم اثرات تصادفی را در نظر می‌گیرد. هدف اصلی این مطالعه تأکید بر اهمیت انتخاب مدل مناسب برای تصمیم‌گیری نگهداری در سناریویی است که اطلاعات وضعیت سیستم در دسترس است. ما نشان می‌دهیم که انتخاب نادرست مدل می‌تواند به پیامدهای قابل توجهی، به ویژه در زمینه افزایش هزینه‌های نگهداری، منجر شود. با مقایسه این سه مدل، نشان می‌دهیم که مدل‌های مبتنی بر وضعیت با اثرات قطعی یا تصادفی (مدل‌های دوم و سوم) عموماً مؤثرتر و مقرون به صرفه‌تر از مدل اول هستند، مشروط بر اینکه اطلاعات وضعیت سیستم در دسترس باشد.

۲ مدل فرسایش

در این پژوهش یک سیستم تک واحدی بررسی می‌شود که در معرض تخریب تدریجی و شوک‌های تصادفی قرار دارد. مدل‌سازی فرآیند تخریب سیستم بر اساس مدل مسیر تخریب تعمیم یافته ارائه شده توسط میکرو اسکوبار انجام می‌شود که در آن از مفهوم سن مجازی به جای سن تقویمی برای مدل‌سازی استفاده می‌شود.

شوک‌های تصادفی تأثیرگذار بر سیستم به عنوان غیرکشنده در نظر گرفته می‌شوند که هر کدام اثر تصادفی به اندازه W ایجاد می‌کنند، که در آن W از توزیع مشخصی به نام F پیروی می‌کند. این شوک‌ها بر اساس یک فرآیند پواسون همگن (HPP) با شدت ثابت λ رخ می‌دهند.

سن مجازی سیستم در هر زمان t که با T_v نشان داده می‌شود، به صورت ریاضی به این شکل بیان می‌شود:

$$T_v = t + \sum_{i=1}^{N(t)} W_i ,$$

که در آن $N(t)$ نشان دهنده تعداد تجمعی شوک‌ها تا زمان t است. شوک i ام به طور آنی سیستم را به اندازه W_i پیرتر می‌کند. در نتیجه، وضعیت تخریب سیستم در زمان t را می‌توان به صورت زیر نشان داد:

$$g(t, \Theta) = g\left(t + \sum_{i=1}^{N(t)} W_i, \Theta\right) = g(T_v, \Theta) , \quad (1)$$

در اینجا، g نشان‌دهنده یک تابع یکنوا و غیرنزولی و معکوس‌پذیر نسبت به زمان t است، در حالی که Θ نشان‌دهنده بردار پارامترهای مدل است.

۳ مدل نگهداری و تعمیر پویا

این مقاله یک سیاست بازرسی دوره‌ای برای سیستم‌های در معرض تخریب پیشنهاد می‌دهد که در آن بازرسی‌ها در فواصل ثابت $k\tau (k = 1, 2, \dots)$ انجام می‌شود. در هر نقطه بازرسی، وضعیت تخریب سیستم به طور کامل اندازه‌گیری می‌شود که منجر به یکی از دو اقدام تعمیر و نگهداری زیر می‌شود:

۱. نگهداری و تعمیر اصلاحی: تعویض کامل سیستم هنگامی که تخریب به آستانه بحرانی z_* می‌رسد
 ۲. نگهداری و تعمیر پیشگیرانه: تعمیر ناقص هنگامی که تخریب زیر z_* باقی می‌ماند
- در این پژوهش رویکردی پویا برای تعمیر ناقص ارائه می‌شود که در آن شدت تعمیر با سطح تخریب فعلی سیستم سازگار است. سه مدل تعمیر ناقص به شرح زیر معرفی می‌گردد:

• مدل I: تعمیر با اثر تصادفی

اثر تعمیر از طریق η تعریف می‌شود که η از توزیع بتا با پارامترهای (α, β) پیروی می‌کند و به صورت تصادفی سن سیستم را به $(1 - \eta)h(D^k, \Theta)$ کاهش می‌دهد. وضعیت تخریب بعدی به صورت زیر به‌روزرسانی می‌شود:

$$D^{k+1} = g \left((1 - \eta)h(D^k, \Theta) + \tau + \sum_{i=N(k\tau)}^{N((k+1)\tau)} W_i, \Theta \right), \quad (1)$$

که در آن h معکوس تابع g است و D^k و D^{k+1} مقدار فرسایش سیستم به ترتیب قبل و بعد از اعمال تعمیر است

• مدل II: تعمیر پویا با اثر مشخص

اثر تعمیر به طور دقیق توسط نسبت تخریب فعلی به آستانه $\frac{D^k}{z_*}$ ، تعیین می‌شود و سن سیستم را به $(1 - \frac{D^k}{z_*})h(D^k, \Theta)$ کاهش می‌دهد. قابل توجه است که این مدل به طور طبیعی هنگامی که $D^k = z_*$ باشد به تعویض کامل تبدیل می‌شود. به‌روزرسانی تخریب به این صورت است:

$$D^{k+1} = g \left((1 - \frac{D^k}{z_*})h(D^k, \Theta) + \tau + \sum_{i=N(k\tau)}^{N((k+1)\tau)} W_i, \Theta \right), \quad (2)$$

• مدل III: تعمیر پویا با اثر تصادفی

این مدل با ترکیب وابستگی به تخریب و عدم قطعیت ذاتی، از یک متغیر تصادفی α با مقدار مورد انتظار $E[\alpha] = D^k$ استفاده می‌کند که از توزیع گامای قطع‌شده پیروی می‌کند. کاهش سن به $(1 - \frac{\alpha}{z_*})h(D^k, \Theta)$ منجر

به این به روزرسانی می شود:

$$D^{k+1} = g \left(\left(1 - \frac{\alpha}{z}\right) h(D^k, \Theta) + \tau + \sum_{i=N(k\tau)}^{N((k+1)\tau)} W_i, \Theta \right). \quad (3)$$

مقایسه مدل ها:

- پیش بینی پذیری: مدل II نتایج از پیش تعیین شده ارائه می دهد، در حالی که مدل III عدم قطعیت عملیاتی را در نظر می گیرد

- انطباق پذیری: هر دو مدل پویا به طور خودکار شدت تعمیر را با پیشرفت تخریب تنظیم می کنند

- صرفه جویی در هزینه: تحلیل ها نشان می دهند که راهبردهای تعمیر متناسب منجر به کاهش هزینه های نگهداری می شوند

این چارچوب، رویکردی عملی برای پیاده سازی نگهداری مبتنی بر شرایط ارائه می دهد که در آن تصمیمات تعمیر به صورت پویا به تخریب اندازه گیری شده سیستم پاسخ می دهند و امکان صرفه جویی در هزینه را نسبت به راهبردهای مرسوم نگهداری فراهم می کنند.

۴ تعریف MDP برای مساله نگهداری و تعمیر

فرآیند تصمیم گیری نگهداری و تعمیر را به عنوان یک فرآیند تصمیم گیری مارکوف (MDP) مدل می کنیم که در آن یک عامل یاد می گیرد چگونه اقدامات بهینه را برای حداقل کردن هزینه های کلی نگهداری و تعمیر انتخاب کند. اجزای این مدل عبارتند از:

۱. فضای حالت (S_k) :

یک متغیر پیوسته یک بعدی است که نشان دهنده میزان تخریب سیستم در هر بازرسی است

۲. فضای اقدام (A_k) :

۰: "عدم انجام عمل"، ۱: "تعمیر ناقص" و ۲: "تعویض"

۳. تابع پاداش (R_k) :

$$\mathbf{R}_k = - \prod_{i=0}^2 \left((C_{ins} + C_i)^\gamma (C_{ins} + C^p)^{1-\gamma} \right)^{\delta(i)}, \quad k = 1, 2, \dots \quad (1)$$

که در آن:

• $\delta(i) = 1$ اگر اقدام i انجام شود و در غیر این صورت ۰ است.

• $\gamma = 1$ اگر سیستم سالم بماند، ۰ اگر خراب شود.

• C_{ins} : هزینه بازرسی

• C_1 : هزینه تعمیر ناقص

• C_2 : هزینه تعویض

• C^p : جریمه برای خرابی

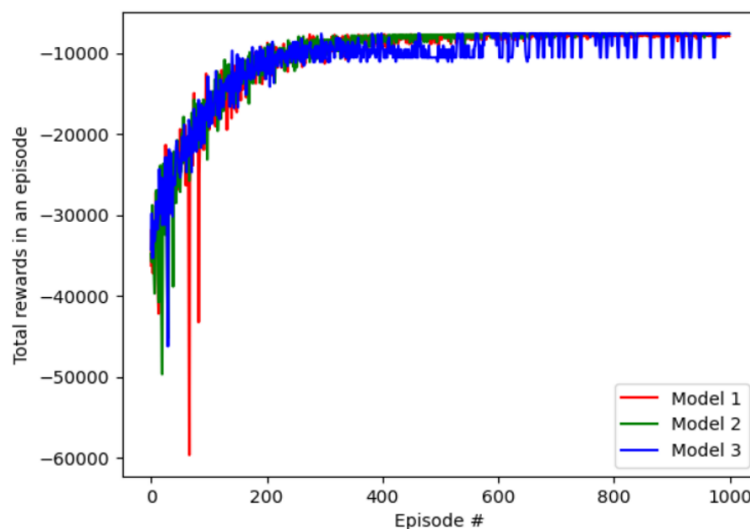
این فرمول‌بندی به عامل امکان می‌دهد تا از طریق یادگیری تقویتی، سیاست‌های نگهداری بهینه از نظر هزینه را یاد بگیرد و بین هزینه‌های نگهداری پیشگیرانه و ریسک خرابی تعادل برقرار کند. ساختار هزینه متناسب برای تعمیرات ناقص ($C_1 \propto D^k$) رویکرد پویای پیشنهادی برای نگهداری و تعمیر را منعکس می‌کند. در این مدل سیستم در هر مرحله بازرسی در یکی از حالت‌های تخریب ممکن قرار دارد و عامل بر اساس حالت فعلی سیستم، یکی از سه اقدام ممکن را انتخاب می‌کند و نهایتاً انتخاب هر اقدام منجر به دریافت یک پاداش (هزینه منفی) می‌شود هدف نهایی، یادگیری سیاستی است که در بلندمدت کمترین هزینه کل را ایجاد کند ویژگی کلیدی این روش، در نظر گرفتن هزینه تعمیر ناقص به صورت متناسب با میزان تخریب سیستم است که آن را از مدل‌های سنتی متمایز می‌کند. این رویکرد انعطاف‌پذیرتر بوده و می‌تواند به راهکارهای مقرون‌به‌صرفه‌تری منجر شود.

۵ مطالعات عددی

پیاده‌سازی مدل پیشنهادی نگهداری و تعمیر با پارامترهای زیر برای یک مدل فرسایش خطی انجام شده که در آن: $\tau = 3$ ضریب تخفیف و $C_{ins} = 50$ ، $C = 500$ ، $C_p = 1500$ ، $\gamma = 0.9$ و نرخ یادگیری 0.001 در نظر گرفته شده است. الگوریتم DQN در پایتون ۳.۶ با استفاده از کتابخانه PyTorch پیاده‌سازی شده که در آن یک شبکه عصبی با دو لایه پنهان (هر لایه ۱۲۸ گره) به عنوان تقریب‌زننده تابع مورد استفاده قرار گرفته شده است. از این طریق، سطح بهینه تخریب برای تعمیر (D^*)، زمان بهینه تعمیر (T^*) و هزینه متوسط (C^*) برای مدل‌های پیشنهادی نگهداری و تعمیر استخراج شد و در جدول (۱) گزارش شده اند هم چنین شکل (۱) نشان دهنده همگرایی الگوریتم به کار گرفته شده است. جزییات بیشتر در [۲] قابل رویت است.

جدول ۱: مقایسه سطح فرسایش بهینه برای تعمیر، زمان بهینه برای تعمیر و هزینه برای سه مدل پیشنهادی تحت الگوریتم DQN

III Model	II Model	I Model	
$8.69 (sd = 0.13)$	$8.75 (sd = 0.21)$	$7.37 (sd = 0.30)$	D^*
$30.09 (sd = 1.97)$	$31.32 (sd = 1.46)$	$22.71 (sd = 1.32)$	T^*
$59.55 (sd = 2.44)$	$58.35 (sd = 1.85)$	$64.78 (sd = 2.16)$	C^*



شکل ۱: همگرایی الگوریتم DQN تحت سه مدل پیشنهادی

۶ بحث و نتیجه‌گیری

در این مقاله، یک استراتژی جدید نگهداری و تعمیر ارائه شده است که شامل سه اقدام «عدم انجام عملیات»، «تعمیر ناقص» و «تعویض» می‌باشد. نوآوری این استراتژی در وابستگی فاکتور بهبود تعمیر ناقص به سطح تخریب سیستم در زمان تعمیر است، برخلاف روش‌های مرسوم که این فاکتور ثابت یا تصادفی در نظر گرفته می‌شود. این رویکرد امکان تطبیق اقدامات تعمیر با شرایط خاص سیستم را فراهم کرده و از تعمیر بیش از حد یا ناکافی جلوگیری می‌کند. نتایج نشان می‌دهند که استراتژی پیشنهادی از لحاظ هزینه‌های نگهداری، عملکرد بهتری نسبت به روش‌های مرسوم دارد.

همچنین، استفاده از الگوریتم‌های یادگیری تقویتی عمیق مانند DQN نشان داده است که این روش‌ها در برابر تغییرات محیطی مقاوم‌تر و از نظر محاسباتی کارآمدتر هستند.

سه مدل مختلف نگهداری و تعمیر با فاکتورهای بهبود مستقل از سطح تخریب، ثابت و وابسته به تخریب، و تصادفی و وابسته به تخریب تعریف شده‌اند. مسئله به صورت فرآیند تصمیم‌گیری مارکوف فرموله شده و از الگوریتم‌های یادگیری تقویتی مدل-آزاد (DQN) برای یافتن سیاست‌های بهینه استفاده شده است. مطالعات عددی روی یک سیستم تک‌واحدی با مسیر تخریب خطی و شوک‌های تصادفی نشان داد که مدل‌های جدید با اقدامات تعمیر وابسته به حالت، می‌توانند طول عمر سیستم را بیش از ۳۷ درصد افزایش و هزینه نگهداری را حدود ۱۰ درصد کاهش دهند. از محدودیت‌های این مطالعه می‌توان به هزینه محاسباتی بالا و مسائل همگرایی اشاره کرد. همچنین، تغییرات جزئی در فاکتور بهبود ممکن است به مقادیر بهینه متفاوتی منجر شود. پیشنهاد می‌شود در تحقیقات آینده، عملکرد این روش‌ها روی مدل‌های غیرخطی پیچیده‌تر و فضاهاى عمل پیوسته بررسی شود. مقایسه جامع‌تر بین روش‌های مبتنی بر یادگیری تقویتی و روش‌های سنتی مبتنی بر تجدید نیز می‌تواند مفید باشد. به طور خلاصه، این مطالعه یک استراتژی انعطاف‌پذیر با قابلیت تطبیق اقدامات تعمیر با وضعیت سیستم ارائه کرده که پتانسیل کاهش هزینه‌ها و بهبود سلامت سیستم را دارد.

مراجع

- [1] Eidi, Sh., Haghighi, F, Safari, A Zio, E. (2025), *Optimal Dynamic State-Dependent Maintenance Policy by Deep Reinforcement Learning*. Quality and Reliability Engineering International, <https://doi.org/10.1002/qre.3806>.



یازدهمین سمینار تخصصی
نظریه قابلیت اعتماد و کاربردهای آن
۲۴ و ۲۵ اردیبهشت ۱۴۰۴

محاسبه طول عمر باقی مانده در نرم افزار R

طباطبائی شیرازی، س.ا. عمادی، م. آرشی، م.^۲

دانشگاه فردوسی مشهد، دانشکده علوم ریاضی، گروه آمار

چکیده

میانگین طول عمر باقی مانده معیار مهمی است که در زمینه های مختلف از جمله شرکت های داروسازی، شرکت های تولیدی و شرکت های بیمه برای تجزیه و تحلیل بقا استفاده می شود. با این حال، محاسبه میانگین طول عمر باقی مانده می تواند پر زحمت و چالش برانگیز باشد. برای پرداختن به این مسئله، ما *reslife* را معرفی می کنیم: *reslife* یک بسته نرم افزاری در نرم افزار R است که به طور خاص برای ساده سازی و تسریع محاسبه طول عمر باقی مانده متوسط طراحی شده است. جنبه منحصر به فرد این بسته این است که به جای تکیه بر انتگرال گیری عددی، طول عمر باقی مانده متوسط را بر اساس راه حل های تحلیلی برای رایج ترین توزیع های بقا، مانند توزیع نمایی، وایبل، گاما، گامای تعمیم یافته و توزیع فیشر تعمیم یافته و... محاسبه می کند. بسته *reslife* این قابلیت را ارائه می دهد که از نتایج رگرسیون بقا یا پارامترهای ارائه شده توسط کاربر برای محاسبه طول عمر باقی مانده درصدی نیز ارائه می دهد. علاوه بر این، بسته *reslife* قابلیت پیش بینی طول عمر باقی مانده متوسط برای داده های

^۱ tabatabaei_amirhossein@yahoo.com

^۲ emadi@um.ac.ir

^۲ arashi@um.ac.ir

مشاهده‌نشده را نیز فراهم می‌کند. در این مقاله، ما بسته reslife را ارائه می‌کنیم؛ اصول ریاضی اساسی آن را توضیح می‌دهیم، عملکرد آن را نشان می‌دهیم، و مثال‌هایی در مورد نحوه استفاده از بسته ارائه می‌دهیم. هدف این است که استفاده از میانگین طول عمر باقیمانده را تسهیل کند و آن را برای پزشکان در رشته‌های مختلف، به ویژه آنهایی که در تجزیه و تحلیل بقا در صنعت داروسازی دخیل هستند، در دسترس و کارآمدتر کند.

کلمات کلیدی: میانگین طول عمر باقی‌مانده، تحلیل بقا، رگرسیون بقا، تابع باقی‌مانده عمر

۱ پیش‌گفتار

تحلیل بقا مجموعه‌ای از روش‌های آماری برای تحلیل داده‌هاست که در آن متغیر اصلی، زمان وقوع یک پدیده است. این روش‌ها به طور گسترده در حوزه‌های بهداشت و درمان، اقتصاد و علوم اجتماعی استفاده می‌شوند و در مهندسی با عنوان نظریه قابلیت اطمینان شناخته می‌گردند. ویژگی بارز تحلیل بقا، توانایی آن در استفاده از داده‌های سانسور شده است؛ به این معنا که حتی زمانی که برخی واحدها در زمان ارزیابی همچنان زنده یا فعال هستند، اطلاعات آن‌ها در مدل لحاظ می‌شود.

یکی از مفاهیم بنیادین در این حوزه، طول عمر باقی‌مانده به ویژه میانگین طول عمر باقی‌مانده است. این معیار طول عمر مورد انتظار را در صورتی اندازه‌گیری می‌کند که فرد یا سیستم تا یک زمان مشخص زنده مانده باشد. همان‌طور که [۲] اشاره می‌کند، MRL نقش مهمی در نظریه قابلیت اطمینان دارد و ابزاری کلیدی برای مدل‌سازی و توصیف توزیع‌های آماری است. این معیار در کاربردهای عملی متنوعی همچون صنایع داروسازی، بیم‌سنجی، اقتصاد و مهندسی مورد استفاده قرار می‌گیرد.

در دو دهه اخیر، توجه به مفهوم MRL به شکل چشمگیری افزایش یافته و پژوهشگران نتایج متنوعی از آن استخراج کرده‌اند. برای مثال، [۱] از تابع میانگین عمر باقی‌مانده کاهشی به عنوان معیاری برای سنجش پیری استفاده کرده و چارچوبی از مفاهیم مرتبط با معیارهای مختلف پیری ارائه دادند. [۳] یک آماره آزمون برای تابع MRL کاهشی توسعه دادند. همچنین [۶] کلاس‌هایی از توزیع‌ها با میانگین عمر باقی‌مانده خطی را توصیف کردند. در ادامه، [۷] رفتار کلی MRL را در توزیع‌های طول عمر پیوسته و گسسته با توجه به نرخ‌های شکست بررسی کردند، و [۵] مفهوم حالت دو متغیره را بسط داد و ارتباط بین قابلیت اطمینان و تابع MRL را نشان داد. با وجود این دستاوردها، محاسبه میانگین عمر باقی‌مانده در عمل می‌تواند پیچیده و زمان‌بر باشد. روش‌های مرسوم غالباً بر انتگرال‌گیری عددی تکیه دارند که با بزرگ شدن حجم داده‌ها محاسبات را دشوار و گاه غیرعملی می‌سازد. در همین راستا، پژوهش‌های جدید به دنبال روش‌های کاراتر بوده‌اند؛ برای نمونه، [۹] روشی نوین مبتنی بر نمونه‌گیری رتبه‌ای ارائه کردند که با کاهش

هزینه‌های نمونه‌گیری، برآورد دقیق‌تری از تابع MRL به دست می‌دهد و می‌تواند چالش‌های محاسباتی در مجموعه داده‌های بزرگ را برطرف کند.

در این مقاله در بخش ۳ میانگین طول عمر باقی‌مانده در مدل‌های پارامتری ارائه می‌کنیم. بخش ۴ را به اجرا بسته reslife اختصاص داده‌ایم؛ در هر بخش نتایج را کامل تفسیر و تحلیل می‌نماییم و در انتها خلاصه‌ای از جمع بندی ارائه شده است.

۲ میانگین طول عمر باقیمانده در مدل‌های پارامتری

۱.۲ ویژگی‌های تابع طول عمر باقیمانده:

یکی از پارامترهای مهم تحلیل بقا میانگین طول عمر باقیمانده در زمان x است. برای افراد با سن x ، این پارامتر مدت زمان مورد انتظار زندگی باقیمانده را اندازه‌گیری می‌کند. این پارامتر به صورت $mrl(x) = E(X - x | X > x)$ تعریف می‌شود. میانگین طول عمر باقیمانده برابر مساحت زیر منحنی بقا در سمت راست x تقسیم بر $S(x)$ است. توجه داشته باشید که میانگین طول عمر کل، $mrl(0)$ برابر مساحت کل زیر منحنی بقا است مشروط بر اینکه فرد تا زمان معینی x زنده مانده باشد. به طور خاص، برای یک متغیر پیوسته، میانگین عمر باقی‌مانده می‌تواند به صورت زیر نمایش داده شود:

$$MRL(x) = \frac{\int_x^\infty (t-x)f(t)dt}{S(x)} = \frac{\int_x^\infty S(t)dt}{S(x)} \quad (1)$$

و میانگین و واریانس مربوط به تابع به صورت زیر است:

$$\mu = E(X) = \int_0^\infty tf(t)dt = \int_0^\infty S(t)dt. \quad (2)$$

$$Var(x) = 2 \int_0^\infty tS(t)dt - [\int_0^\infty S(t)dt]^2 \quad (3)$$

به آسانی می‌توان نشان داد که وقتی $x = 0$ باشد، تابع میانگین عمر باقی‌مانده معادل امید ریاضی T است. این رابطه ریاضی را می‌توان به صورت زیر نمایش داد:

$$MRL(\cdot) = \frac{\int_x^\infty (t - \cdot) f(t) dt}{S(\cdot)} = \frac{\int_x^\infty (t - \cdot) f(t) dt}{1} = E[T] \quad (4)$$

می‌توانیم از روش‌های دیگری برای محاسبه میانگین عمر باقی‌مانده استفاده کرد:

$$MRL(x) = \frac{\int_x^\infty (t - x) f(t) dt}{S(x)} = \frac{\int_x^\infty t f(t) dt}{S(x)} - x \quad (5)$$

$$MRL(x) = \frac{\int_x^\infty (t - x) f(t) dt}{S(x)} = \frac{E[T] - \int_x^\infty S(t) dt}{S(x)} \quad (6)$$

توزیع‌های بقای پارامتری داخلی در بسته reslife شامل چندین توزیع آماری استاندارد است که برای تحلیل بقا استفاده می‌شوند. این توزیع‌ها به صورت زیر دسته‌بندی می‌شوند:

۲.۲ میانگین طول عمر باقیمانده برای توزیع بقا پارامتری:

توزیع‌های پارامتری که در حال حاضر برای محاسبه طول عمر باقیمانده در این پکیج گنجانده شده است، در جدول ۱ لیست شده‌اند. تمام توزیع‌های پارامتری ذکر شده در جدول فوق دارای فرم بسته میانگین طول عمر باقیمانده هستند. اثبات‌های دقیق برای دو توزیع گاما تعمیم‌یافته و دو توزیع فیشر تعمیم‌یافته در پیوست [۸] موجود است.

۳ اجرا بسته reslife

در نرم‌افزار R، بسته‌های متعددی برای تحلیل بقا وجود دارند که امکان مدلسازی و تحلیل داده‌های طول عمر را فراهم می‌کنند. برای مثال، بسته‌ی نرم افزاری که توسط [۴] معرفی شد؛ یک ابزار جامع پارامتری است که اجازه می‌دهد توزیع‌های مختلف بقا را برازش داده، احتمال بقا، خطر و طول عمر باقی‌مانده را محاسبه و تحلیل‌های پیشرفته را برای متغیر توضیحی انجام داد. در همین راستا، بسته‌ی نرم افزاری reslife توسط [۸] به عنوان یک ابزار تخصصی برای تحلیل طول عمر باقی‌مانده معرفی شده است. این بسته امکان انجام تحلیل‌های پارامتری و غیرپارامتری MLR و دیگر توابع مرتبط را فراهم می‌کند. در ادامه، جزئیات نحوه استفاده از این بسته‌ها در R و کاربردهای عملی آن‌ها در تحلیل‌های پیشرفته طول عمر باقی‌مانده بررسی می‌شود. بسته reslife در نرم افزار R برای محاسبه و تحلیل تابع بقای باقی‌مانده استفاده می‌شود. این تابع در زمینه‌های مختلفی مانند تحلیل بقا، مهندسی قابلیت اطمینان و سایر حوزه‌هایی که با زمان تا وقوع یک رویداد سروکار دارند، کاربرد دارد. یکی از مزیت‌های قابل توجه این بسته از نظر

جدول ۱: توزیع‌های پارامتری داخلی در بسته reslife

اسم توزیع	پارامتر	آرگومان	میانگین طول عمر باقی‌مانده
Exponential	rate	"exponential"	λ
Weibull	shape, scale	"weibull"	$\frac{\lambda}{\alpha} \Gamma(\frac{1}{\alpha}) S(x^\alpha) / S(x)$
Gamma	shape, rate	"gamma"	$x^\alpha \exp(-\frac{x}{\lambda}) / [\lambda^{\alpha-1} \Gamma(\alpha) S(x)]$
Gompertz	shape, rate	"gompertz"	$\exp(\frac{\lambda}{\alpha} e^{\alpha x}) (\frac{1}{\alpha}) \Gamma(\cdot, \frac{\lambda}{\alpha e^{\alpha x}})$
Log-normal	meanlog, sdlog	"lnorm"	$\frac{\exp(\mu + \frac{\sigma^2}{2}) [1 - \Phi(\frac{\ln(x) - (\mu + \frac{\sigma^2}{2})}{\sigma})]}{S(x) - x}$
Log-logistic	shape, scale	"llogis"	$\frac{\lambda}{\alpha} \Gamma(1 - \frac{1}{\alpha}) \Gamma(\frac{1}{\alpha}) S(z(x); 1 - \frac{1}{\alpha}, \frac{1}{\alpha}) (1 + \frac{x}{\lambda})^\alpha$ (where $z(x) = (\frac{x}{\lambda})^\alpha / (1 + (\frac{x}{\lambda})^\alpha)$)
GGs (۱۹۶۲)	shape, scale, k	"gengamma.orig"	$\alpha \frac{\Gamma(k + \frac{1}{b}, (\frac{x}{\alpha})^b)}{\Gamma(k, (\frac{x}{\alpha})^b)} - x$
GGP (۱۹۷۵)	mu, sigma, Q	"gengamma"	$\frac{e^\mu (Q^{-1})^{\frac{\sigma}{Q}} \gamma(\frac{\sigma}{Q} + Q^{-1}, A) - x \gamma(Q^{-1}, A)}{\gamma(Q^{-1}, A)}$
GFP (۱۹۷۵)	s1, s2, sigma, mu,	"genf.orig"	$\exp(\beta) (\frac{m_1}{m})^\sigma \frac{B(m_1 + \sigma, m_1 - \sigma)}{B(m_1, m_1)}$ $\exp(\beta) (\frac{m_1}{m})^\sigma \frac{B(m_1 + \sigma, m_1 - \sigma)}{B(m_1, m_1)}$ $m_1 = 1(Q^2 + 2P + Q\delta)^{-1}$ $m_2 = 1(Q^2 + 2P - Q\delta)^{-1}$ $\delta = (Q^2 + 2P)^{1/2}$
GFP (۱۹۷۵)	Q, P, sigma, mu,	"genf"	$\delta = (Q^2 + 2P)^{1/2}$

سرعت و دقت به ویژه برای مجموعه داده‌های بزرگ مناسب می‌سازد. یک جدول مقایسه‌ای بین reslife با قابلیت‌های منحصر به فرد و سایر بسته‌های معروف تحلیل بقا در نرم‌افزار R ارائه می‌شود. این جدول بر اساس ویژگی‌های کلیدی مانند پشتیبانی از توزیع‌ها، سرعت، و کاربردهای خاص طراحی شده است:

با وجود قابلیت‌های منحصر به فردی چون پشتیبانی از توزیع‌های غیر معمول و بهینه‌سازی سرعت محاسبات، این بسته با محدودیت‌های قابل توجهی روبه‌رو می‌باشد که می‌تواند کاربردپذیری آن را تحت تأثیر قرار دهد. پشتیبانی محدود از توزیع‌های آماری یکی از این محدودیت‌ها می‌باشد به طوریکه برخی توزیع‌های پرکاربرد به صورت پیش‌فرض در این بسته نرم‌افزاری وجود ندارد مانند Tweedie. همچنین می‌توان عدم پشتیبانی از پردازش موازی و عدم بهینه‌سازی در داده‌های حجیم (> یک میلیون مشاهده) را به عنوان محدودیت‌های این بسته نرم‌افزاری بشمار آورد.

در این بخش، ما بسته reslife را با مثال‌های قابل باز تولید نشان می‌دهیم و نحوه استفاده از آن را با داده‌های

جدول ۲: جدول مقایسه‌ای بسته reslife با سایر بسته‌های تحلیل بقا در نرم‌افزار R

ویژگی	reslife	survival	flexsurv	parsnip	rstpm2
توزیع‌های پارامتری	دارد (گسترده)	دارد (محدود)	دارد (گسترده)	ندارد	دارد
توزیع‌های غیرمعمول	دارد (پشتیبانی کامل)	ندارد	دارد (محدود)	ندارد	ندارد
داده‌های سانسور شده	دارد (همه نوع)	دارد	دارد	دارد	دارد
محاسبات سریع	دارد	ندارد	ندارد	ندارد	ندارد
پشتیبانی از tidyverse	دارد	ندارد	ندارد	دارد	ندارد
کاربردهای صنعتی	دارد (بهینه‌شده)	ندارد	دارد	ندارد	ندارد
یادگیری ماشین	ندارد	ندارد	ندارد	دارد	ندارد

دنیای واقعی توضیح می‌دهیم. بسته reslife شامل دو تابع است: residlife و reslifefsr. تابع reslifefsr به خروجی شیء رگرسیون بقا وابسته است، درحالی‌که residlife پارامترهای ورودی کاربر را می‌پذیرد. یک مرور کلی و مثال از هر دو ارائه داده شده است.

۱.۳ مرور کلی reslifefsr

تابع reslifefsr برای استفاده پس از ایجاد شیء خروجی رگرسیون بقا است و به شکل زیر تعریف می‌شود:

```
reslifefsr (obj, life, p = 0.5, type = "mean", newdata = data.frame())
```

پارامترهای ورودی

- obj: مدل بقای قبلی که باید با استفاده از survreg() ایجاد شده باشد.
- life: متغیری که می‌خواهید بر اساس آن پیش‌بینی انجام شود.
- p: احتمال مورد نظر برای محاسبه محدوده اطمینان (پیش‌فرض: ۰/۵)
- type: نوع روش پیش‌بینی "percentile"، "median"، "mean" و "all" (پیش‌فرض "mean")
- newdata: فریم داده‌ای جدید برای پیش‌بینی (پیش‌فرض: یک فریم داده‌ای خالی)

۲.۳ نمایش reslifefsr:

برای نشان دادن نحوه استفاده از reslifefsr، مثال زیر را بررسی می‌کنیم. ابتدا نشان می‌دهیم که چگونه خروجی تابع با نتیجه انتگرال‌گیری عددی مقایسه می‌شود. ما با استفاده از تابع بقا بر روی مجموعه داده‌های "ovarian" شروع می‌کنیم که درون بسته survival تعبیه شده است. بیایید به مجموعه داده‌های ovarian نگاهی بیندازیم.

```
> library(survival)
```

```
> library(reslife)
```

```
> library(flexsurv)
```

```
> head(ovarian)
```

	futime	fustat	age	resid.ds	rx	ecog.ps
1	59	1	72.3315	2	1	1
2	115	1	74.4932	2	1	1
3	156	1	66.4658	2	1	2
4	421	0	53.3644	2	2	1
5	431	1	50.3397	2	1	1
6	448	0	56.4301	1	1	2

داده‌های ovarian یکی از مجموعه داده‌های معروف در تحلیل بقا در نرم‌افزار R است. این داده‌ها مربوط به بیماران مبتلا به سرطان تخمدان است و اطلاعاتی درباره زمان بقا و عوامل مؤثر بر آن را ارائه می‌دهد.

```
> fs1 <- flexsurvreg(Surv(futime, fustat) ~ 1, data = ovarian, dist = "gengamma.orig")
```

```
> reslifefsr(obj = fs1, life = 1, type = "mean")
```

```
[1] 1608.511
```

بر اساس نتایج تحلیل، میانگین نحوه برازش مدل داده‌های سرطان تخمدان یک واحد پس از شروع مطالعه، حدود ۱۶۰۸/۵۱ روز (یا تقریباً ۴/۴۱ سال) است. اکنون می‌توانیم میانگین عمر باقیمانده را با استفاده از انتگرال عددی

محاسبه کنیم تا نسبت به خروجی تابع `reslifefsr` اطمینان حاصل کنیم.

```
> b <- exp(as.numeric(fs1$coefficients[1]))
> a <- exp(as.numeric(fs1$coefficients[2]))
> k <- exp(as.numeric(fs1$coefficients[3]))
> f <- function(x) {return(pgengamma.orig(x, shape = b, scale = a, k = k, lower.tail=FALSE))}
> integrate(f, 1, Inf)$value/pgengamma.orig(1, shape = b, scale = a, k = k, lower.tail = FALSE)

[1] 1608.511
```

مقادیر یکسان هستند. حالا می‌توانیم یک نمایش با چندین متغیر کمکی انجام دهیم. ابتدا در مجموعه داده‌های "ovarian" از متغیر پیوسته به نام `age` هم استفاده می‌کنیم. سپس یک گرسیون بقا جدید ایجاد می‌شود. برای این رگرسیون بقا دوم از توزیع وایبل استفاده می‌کنیم.

```
> fsr2 <- flexsurvreg(Surv(futime,fustat)~ age, data =+ ovarian, dist = "weibull")
```

حالا می‌توانیم از تابع `reslifefsr` برای یافتن میانگین عمر باقی‌مانده استفاده کنیم. مقدار `life` را برابر با ۴ تنظیم می‌کنیم و برای صرفه‌جویی در فضا، پاسخ‌ها را به ۸ مورد محدود می‌کنیم. نتیجه کامل شامل ۲۶ مقدار است، یکی برای هر مشاهده در مجموعه داده.

```
> reslifefsr (obj = fsr2, life = 4)
```

```
[1] 202.1022 163.4903 358.0965 1272.5397 1703.6646 946.5172
901.2830 679.7562
```

می‌توانیم این کار را برای میانه عمر باقی‌مانده نیز تکرار کنیم:

```
> reslifefsr (obj = fsr2, life = 4, type = "median")
```

```
[1] 180.0515 145.5628 319.3829 1136.1163 1521.1714 844.9317
804.5310 606.6751
```

و در نهایت برای هر سه نوع عمر باقی‌مانده. مقدار p برابر با ”۸” تنظیم شده است، بنابراین به صدک ۸۰ام نگاه می‌کنیم. از آنجا که ”all” برای type مشخص شده است، خروجی یک چارچوب داده است که هر ستون آن نشان‌دهنده یک معیار مختلف از عمر باقی‌مانده است.

```
> reslifefsr (obj = fsr2, life = 4, p = .8, type = "all")
```

	mean	median	percentile
1	202.1022	180.0515	303.8587
2	163.4903	145.5628	246.1030
3	358.0965	319.3829	537.1315
4	1272.5397	1136.1163	1904.2490
5	1703.6646	1521.1714	2548.7607
6	946.5172	844.9317	1416.8544
7	901.2830	804.5310	1349.2299
8	679.7562	606.6751	1018.0454

یکی دیگر از قابلیت‌های reslifefsr توانایی آن برای کار با داده‌های جدید ارائه‌شده توسط کاربر است. می‌توانیم یک چارچوب داده نمونه ایجاد کرده و آن را به تابع ارائه دهیم تا پیش‌بینی‌ها را تولید کند.

```
> age <- c(30, 65, 70)
> newdata <- data.frame(age)
> reslifefsr (obj = fsr2, life = 4, p = 0.6, type = 'all', newdata= newdata)
```

	mean	median	percentile
1	1086.696	970.131	1151.374

```
2 9050.825 8083.199 9588.044
3 4613.486 4120.045 4887.431
```

تابع قادر است برخی تغییرات در چارچوب داده‌های ورودی را مدیریت کند. تغییر ترتیب ستون‌ها تأثیری بر نتایج ندارد، و داشتن ستون‌های بیشتر از نیاز نیز مشکلی ایجاد نمی‌کند. تنها نیازمندی‌ها این است که تمام متغیرهای کمکی موجود در تابع بقا باید در چارچوب داده وجود داشته باشند و اگر متغیرها فاکتور باشند، هیچ سطح فاکتور خارج از محدوده مجموعه داده اصلی مجاز نیست.

```
> group <- c("Medium", "Good", "Poor")
> newdata2 = data.frame(group)
> reslifefsr (obj = fsr2, life = 4, p = 0.6, type = "all", newdata = newdata2)
```

```
Error in weibull_mlr(obj, life, p, type, newdata) :
```

```
Wrong columns in inputted data
```

تابع یک خطا برمی‌گرداند زیرا ستون‌های کمتری نسبت به آنچه لازم است ارائه شده است.

۳.۳ بررسی کلی residlife:

تابع residlife برای پارامترهای وارد شده توسط کاربر است و دارای فرم زیر است:

```
residlife (values, distribution, parameters, p = 0.5, type = "mean")
```

values: محدوده مقادیری که در آن عمر باقیمانده محاسبه می‌شود. معمولاً به صورت بردار ارائه می‌شود.

distribution: نام توزیع مورد نظر که می‌تواند یکی از موارد زیر باشد:

• "Weibull"

• "Gompertz"

• "Gamma"

• "Gengamma.orig"

• "Exponential"

• "Lnorm"

• "Llogis"

• "Genf"

• "Genf.orig"

parameters: پارامترهای تابع بقا که باید به صورت بردار و به ترتیب صحیح وارد شوند.

p: درصد مورد نظر برای محاسبه عمر باقیمانده. مقدار پیش فرض ۰/۵ است.

type: نوع خروجی عمر باقیمانده که می‌تواند یکی از موارد زیر باشد:

• "mean" (میانگین)

• "median" (میانه)

• "percentile" (درصد)

• "item" (همه موارد بالا)

یک تفاوت عمده بین reslife و residlife این است که residlife یک متغیره است، زیرا مقادیر دقیق پارامترها باید توسط کاربر ارائه شوند.

۴.۳ نمایش residlife:

اکنون نحوه عملکرد تابع residlife را نشان می‌دهیم. برای مثال‌های این بخش، از خروجی یک تابع بقا با استفاده از مجموعه داده myeloma استفاده می‌کنیم.

```
> head(myeloma)
```

```
id year entry futime death
1  1   57     0   1431     1
```

2	2	61	0	686	1
3	3	53	0	6270	1
4	4	66	0	365	1
5	5	67	0	1340	1
6	6	88	0	1567	1

داده‌های myeloma یکی از مجموعه داده‌هایی است که در مطالعات پزشکی استفاده می‌شود. این داده‌ها مربوط به بیماران سرطان مغز استخوان است.

```
> fit <- flexsurvreg(Surv(entry, futime, death) ~ 1, data = myeloma, dist = "weibull")
> fit$res
```

	est	L95%	U95%	se
shape	0.8888754	0.8615703	0.917046	0.01414992
scale	1126.3205165	1077.3592490	1177.506860	25.53992073

با توجه به پارامترهای شکل و مکان در بالا، برای این مثال دنباله‌ای از اعداد ۱ تا ۱۰ را با فاصله ۰/۶ تولید می‌کنیم:

```
> residlife(values = seq(1,10,0.6), distribution = "weibull", parameters = c(shape = 0.8888754,
scale = 1126.3205165))
```

```
[1] 1194.617 1195.218 1195.769 1196.287 1196.778 1197.248 1197.701
[8] 1198.139 1198.564 1198.978 1199.381 1199.776 1200.162 1200.540
[15] 1200.911 1201.276
```

نام پارامترهای وارد شده در آرگومان

```
> residlife(values = seq(1,10,0.6), distribution = "weibull", parameters = c(shape = 0.8888754,
not_scale = 1126.3205165))
```

```
Error in residlife(values = seq(1,10,0.6), distribution = "weibull",:
  incorrect parameters entered. Parameters for weibull are
  shape and scale
```

همانطور که در بالا ذکر شد، تابع می‌تواند پارامترهای مختلفی را برای توزیع‌ها با گزینه‌های متعدد داشته باشد
 بعنوان مثال در زیر از توزیع گاما با پارامترهای شکل و مقیاس استفاده می‌شود.

```
> residlife (values = seq(1,10,0.6), distribution = "gamma", parameters = c(shape = 0.8888754,
  scale = 1126.3205165))
```

```
[1] 1002.187 1002.640 1003.050 1003.431 1003.788 1004.128 1004.453
[8] 1004.766 1005.067 1005.359 1005.641 1005.916 1006.184 1006.446
[15] 1006.701 1006.951
```

می‌توان از پارامترهای شکل به جای نرخ استفاده کنیم، نتیجه یکسان را مشاهده کنید:

```
> residlife (values = seq(1,10,0.6), distribution = "gamma", parameters = c(shape = 0.8888754,
  rate = 1/1126.3205165))
```

```
[1] 1002.187 1002.640 1003.050 1003.431 1003.788 1004.128 1004.453
[8] 1004.766 1005.067 1005.359 1005.641 1005.916 1006.184 1006.446
[15] 1006.701 1006.951
```

مانند تابع `residlife`، `reslifefsr` توانایی نشان دادن میانگین، میانه، صدک، هر سه نوع عمر باقیمانده در زیر نمونه‌ای
 از تابع با نوع تنظیم شده روی "all" آمده است.

```
> residlife (values = seq(1,10,0.5), distribution = "weibull", parameters = c(shape = 0.8888754,
  scale = 1126.3205165), p = 0.65, type = "all")
```

	mean	median	percentile
1	1194.617	747.0875	1191.117
2	1195.122	747.6060	1191.690
3	1195.590	748.0869	1192.222
4	1196.032	748.5407	1192.727
5	1196.453	748.9735	1193.209
6	1196.857	749.3893	1193.673
7	1197.248	749.7907	1194.122
8	1197.626	750.1798	1194.558
9	1197.994	750.5580	1194.983
10	1198.353	750.9267	1195.397
11	1198.703	751.2867	1195.803
12	1199.045	751.6388	1196.200
13	1199.381	751.9839	1196.589
14	1199.710	752.3223	1196.972
15	1200.034	752.6546	1197.348
16	1200.352	752.9813	1197.718
17	1200.664	753.3027	1198.082
18	1200.972	753.6191	1198.441
19	1201.276	753.9309	1198.796

مواردی وجود دارد که تابع یک NaN برمیگرداند، اما این باید بیشتر با مقادیر ورودی نسبت به خود تابع ترکیبی از ارزش‌های زندگی و توزیع پارامترها قابل محاسبه نیستند؛ در زیر یک نمونه را نشان خواهیم داد. پارامتر شکل دارد به ۲/۵، مقیاس به ۲/۲ تغییر یافته و مقادیر بردار از ۲۰ تا ۴۰ است.

```
> residlife (values = seq(20,40), distribution = "weibull", parameters = c(shape = 2.5, scale = 2.2), p = 0.65, type = "all")
```

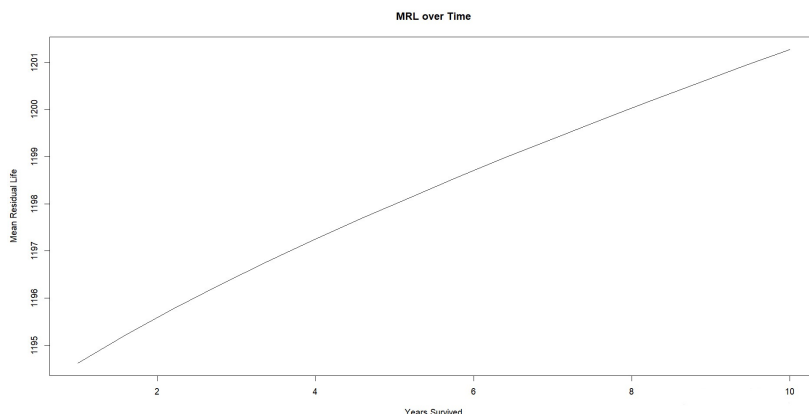
	mean	median	percentile
1	0.03202811	0.02223490	0.03366196
2	0.02977600	0.02066773	0.03129092
3	0.02777551	0.01927626	0.02918545
4	0.02598904	0.01803412	0.02730575
5	0.02438592	0.01691982	0.02561934
6	0.02294087	0.01591567	0.02409952
7	0.02163293	0.01500701	0.02272416
8	0.02044459	0.01418161	0.02147473
9	0.01936109	0.01342916	0.02033569
10	0.01836993	0.01274096	0.01929384
11	0.01746047	0.01210957	0.01833798
12	NaN	Inf	Inf
13	NaN	Inf	Inf
14	NaN	Inf	Inf
15	NaN	Inf	Inf
16	NaN	Inf	Inf
17	NaN	Inf	Inf
18	NaN	Inf	Inf
19	NaN	Inf	Inf
20	NaN	Inf	Inf
21	NaN	Inf	Inf

مقادیر عمر پس از ۱۱ امکان پذیر نیستند، با تابع میانگین یک مقدار NaN و مقدار را برمیگرداند عمر باقی مانده میانه و صدک که مقدار Inf را برمی گرداند. پس خروجی های این تابع را می توان به راحتی گرفته و در نمودار یا هر نوع دیگری از آن استفاده کرد ابزار تحلیل بیایید به نمونه ای از چنین نموداری در شکل زیر نگاه کنید که میانگین عمر باقی مانده را نشان می دهد.

فرض کنید می‌خواهیم میانگین طول عمر باقی‌مانده را در طول زمان رسم کنیم. ابتدا یک برداری برای زمان ایجاد می‌کنیم، آن را از طریق تابع `residlife()` با نوع "mean" اجرا کردیم و خروجی را بر اساس رسم کنیم.

زمان پس از استخراج پارامترها، می‌توانیم `residlife()` را اجرا کرده و یک بردار خروجی تولید کنیم. توجه داشته باشید از آنجایی که `residlife()` میتواند یک ورودی برداری بگیرد، ما مجبور نیستیم از یک حلقه استفاده کنیم.

```
> time = seq(1,10,0.6)
> life = residlife(values = time, distribution = "weibull", parameters = c(shape = 0.8888754,
  scale = 1126.3205165))
> plot(time, life, xlab = "Years_Survived", ylab = "Mean_Residual_Life", main =
  "MRL_over_Time", type = "l")
```



نمودار نشان می‌دهد که چگونه امید به زندگی باقی‌مانده در طول زمان تغییر می‌کند. این اطلاعات برای درک الگوی بقا در طول زمان مفید است.

۴ نتیجه‌گیری

بسته `reslife` به عنوان یک ابزار ارزشمند برای محاسبه میانگین طول عمر باقیمانده و طول عمر باقیمانده درصدی عمل می‌کند. هر دوی این معیارها در تحلیل بقا در زمینه‌های مختلف از اهمیت بالایی برخوردار هستند. تا جایی که می‌دانیم، هیچ بسته‌ای در `R` وجود ندارد که بتواند این محاسبات را انجام دهد. با ساده‌سازی محاسبه میانگین طول عمر باقیمانده، این بسته به محققان و متخصصان امکان می‌دهد تا بینش‌های ارزشمند را به صورت کارآمدتری به دست آورند و به پیشرفت‌ها در تحلیل بقا در رشته‌های مختلف، از جمله صنایع دارویی، کمک کنند.

مراجع

- [1] Bryson, M. C., Siddiqui, M. M. (1969). Some criteria for aging. *Journal of the American Statistical Association*, 64(328), 1472-1483.
- [2] Clark, T. G. (2003). Survival analysis Part I: Basic concepts and first analyses. PMC. <https://pmc.ncbi.nlm.nih.gov/articles/PMC2394262/>
- [3] Hall, W. J., Wellner, J. A. (1981). Mean residual life. In M. Csörgő, D. A. Dawson, J. N. K. Rao, A. K. Md. E. Saleh (Eds.), *Statistics and related topics* (pp. 169–184). North-Holland.
- [4] Jackson, C. H. (2016). `flexsurv`: A Platform for Parametric Survival Modeling in `R`. *Journal of Statistical Software*, 70(8), 1-33. <https://doi.org/10.18637/jss.v070.i08> *Journal of Statistical Software*+2PMC+2
- [5] Nair, N. U. (1989). Bivariate mean residual life. *IEEE Transactions on Reliability*, 38(3), 362–364.
- [6] Oakes, D., Dasu, T. (1990). Relative mean residual life models. *Statistics in Medicine*, 25(2), 1–12.
- [7] Tang, L. C., Lu, Y., Chew, E. P. (1999). Mean residual life of lifetime distributions. *IEEE Transactions on Reliability*, 48(1), 73–78.

- [8] Wang, Z., Crawford, A., Lee, K. L., Jaganathan, S. (2023). reslife: Residual Lifetime Analysis Tool in R [Preprint]. arXiv.
- [9] Zamanzade, E., Parvardeh, A., Asadi, M. (2019). Estimation of mean residual life based on ranked set sampling. Computational Statistics Data Analysis, 135, 35-55.



یازدهمین سمینار تخصصی
نظریه قابلیت اعتماد و کاربردهای آن
۲۴ و ۲۵ اردیبهشت ۱۴۰۴



آزمون طول عمر شتابیده استرس پله‌ای ساده تحت سانسور فزاینده نوع دوم با مکانیسم خروج تصادفی وابسته

طیبی گیلانی، ن^۱ حقیقی، ف^۲ صفری، ع^۳

گروه آمار، دانشکده ریاضی، آمار و علوم کامپیوتر، دانشگاه تهران

چکیده

در این مقاله یک مکانیسم خروج تصادفی برای سانسور تصادفی فزاینده نوع دوم در یک آزمون طول عمر شتابیده استرس پله‌ای ساده با تأکید بر توزیع نمایی مورد مطالعه قرار گرفته است. احتمال‌های خروج سانسور تصادفی فزاینده نوع دوم از طریق تابع ربط لوجیت به شرایط آزمایش (فاصله‌های شکست و سطوح استرس) مرتبط شده‌اند. مکانیسم خروج شامل پارامترهایی است که شدت و نحوه تأثیر شرایط آزمایش بر احتمال خروج تصادفی را مشخص می‌کنند و توسط آزمایشگر تعیین می‌گردند. نتایج، کاهش زمان آزمون و در نتیجه کاهش هزینه‌ها را نشان می‌دهند. در راستای بررسی این رویکرد، برآوردهای حداکثر درستنمایی پارامترهای مدل ارائه شده است. همچنین یک الگوریتم شبیه‌سازی برای تولید نمونه‌ها در یک آزمون طول عمر شتابیده استرس پله‌ای ساده تحت رویکرد جدید پیشنهاد داده شده است. سپس چندین طرح شبیه‌سازی شده برای بررسی کارایی و عملکرد

^۱n.tabibi@ut.ac.ir

^۲fahaghi@ut.ac.ir

^۳a.safari@ut.ac.ir

رویکرد پیشنهادی ارائه گردیده است. در پایان نیز حساسیت مکانیسم سانسور تصادفی نسبت به پارامترهای تنظیم مورد مطالعه قرار گرفته است.

کلمات کلیدی: آزمون طول عمر شتابیده، تابع ربط لوجیت، مکانیسم خروج تصادفی، سانسور فزاینده نوع دوم، توزیع نمایی

۱ مقدمه

امروزه اغلب محصولات و تولیدات دارای قابلیت اعتماد بالایی هستند، بنابراین میانگین زمان شکست آن‌ها بسیار طولانی است. در نتیجه استنباط آماری در مورد توزیع‌های طول عمر آن‌ها با استفاده از آزمون‌های طول عمر معمول بسیار زمان‌بر و دشوار است. در چنین شرایطی آزمون‌های طول عمر شتابیده (ALT)^۱ مورد استفاده قرار می‌گیرند. در آزمون‌های طول عمر شتابیده، زمان‌های شکست تحت استرس‌هایی بالاتر از شرایط عملیاتی نرمال بسیار سریع‌تر رخ می‌دهند. دما، ولتاژ، رطوبت و فشار نمونه‌هایی از این استرس‌ها می‌باشند. برای نمونه [۳]، [۴] و [۵] را مشاهده کنید. یک کلاس خاص از ALT ، آزمون استرس پله‌ای نامیده می‌شود که به آزمایشگر اجازه می‌دهد یک یا چند عامل استرس را در یک آزمون طول عمر انتخاب کند [۶]. آزمون طول عمر شتابیده استرس پله‌ای با دو سطح استرس، آزمون طول عمر شتابیده استرس پله‌ای ساده ($Simple SSALT$)^۲ نامیده می‌شود. برای تحلیل داده‌ها تحت یک $SSALT$ ، نیاز به یک مدل است که تابع توزیع تجمعی طول عمر تحت استرس‌های مختلف را به تابع توزیع تجمعی طول عمر محصولات تحت شرایط عملیاتی نرمال مرتبط کند. یک مدل متداول و شناخته شده مدل نلسون^۳ است که در این مقاله نیز به کار گرفته شده است.

علاوه بر این، یکی از موضوعات مهم در مباحث آزمون‌های طول عمر، استنباط براساس داده‌های سانسور شده است. یکی از سانسورهای پرکاربرد و محبوب سانسور فزاینده نوع دوم^۴ است که به آزمایشگر اجازه می‌دهد که واحدهای آزمون را در زمانی غیر از پایان آزمایش از آزمون خارج کند که این امر منجر به کاهش هزینه‌های آزمون می‌شود. علاوه بر این یکی از مسائل مهم در مبحث سانسور فزاینده، چگونگی تعیین تعداد واحدهای سالم خارج شده در مراحل مختلف شکست می‌باشد. محققان تعداد واحدهای خارج شده در سانسور فزاینده نوع دوم را مقداری ثابت یا متغیر تصادفی

^۱ Accelerated Life Test

^۲ simple step-stress accelerated life test

^۳ Nelson

^۴ type-II progressive censoring

از توزیعی با پارامترهای ثابت در نظر گرفته‌اند. اما در عمل به نظر می‌رسد تعداد واحدهای خارج شده در هر شکست تابعی از شرایط آزمایش است، به گونه‌ای که تغییر شرایط آزمایش باعث شدن تعداد متفاوتی واحد در هر مرحله از شکست می‌گردد. حسن تبار و همکاران [۲]، مکانیسم سانسور تصادفی که قابلیت به‌روز کردن احتمال خروج در هر مرحله از آزمایش (از جمله فواصل شکست و اطلاعات کمکی موجود) را دارد مطرح نمودند که در آن برای وابسته کردن تعداد واحدهای خارج شده در هر مرحله از شکست به شرایط آزمایش از توابع ربط مربوط به مدل‌های GLM ^۵ از جمله تابع ربط لجیت استفاده کرده‌اند. مکانیسم سانسور تصادفی وابسته پیشنهادی آن‌ها دارای این قابلیت است که تعداد واحدهای سانسور شده در هر مرحله از شکست را یک متغیر تصادفی و تابعی از شرایط آزمایش در نظر می‌گیرد. در واقع در این جا، تعداد واحدهای خارج شده در مرحله i ام از توزیع دوجمله‌ای شرطی با احتمال موفقیت (سانسور) p_i پیروی می‌کند که در آن p_i تابعی از فاصله زمان‌های شکست سپری شده و دیگر اطلاعات کمکی موجود تا مرحله i ام است. در این مقاله سانسور فزاینده نوع دوم تصادفی وابسته به شرایط آزمایش (به اختصار *Type-II PCRD*) با تأکید بر آزمون شتابیده استرس پله‌ای ساده مورد مطالعه قرار گرفته است. ادامه مقاله به شرح زیر است: ابتدا مکانیسم سانسور تصادفی وابسته به شرایط آزمون در آزمون شتابیده استرس پله‌ای ساده بیان شده است. سپس مدل با تأکید بر توزیع نمایی بیان گردیده است و برآوردهایی حداکثر درست‌نمایی پارامترهای مدل محاسبه شده است. همچنین دو مثال شبیه‌سازی شده برای بررسی عملکرد رویکرد جدید ارائه گردیده است. در انتها نیز تحلیل حساسیت مکانیسم سانسور تصادفی نسبت به پارامترهای تنظیم مورد مطالعه قرار گرفته است.

۲ مکانیسم سانسور تصادفی وابسته به شرایط آزمایش در *simple SSALT*

یک آزمون طول عمر شتابیده استرس پله‌ای ساده تحت سانسور فزاینده نوع دوم با n واحد و m شکست، با بردار زمان‌های شکست تصادفی $T = (T_{1:m:n}, \dots, T_{m:m:n})$ در نظر بگیرید. علاوه بر این، سطح استرس از پیش تعیین شده S_1 و S_2 را در نظر بگیرید که $S_1 < S_2$. در ابتدا n واحد مستقل در سطح استرس اولیه S_1 در آزمایش قرار می‌گیرند، آنگاه در زمان از پیش تعیین شده τ سطح استرس از S_1 به S_2 تغییر می‌کند. فرض کنید تعداد شکست‌ها قبل از τ را با N_1 نمایش دهیم؛ آنگاه

$$t_{1:m:n} < \dots < t_{N_1:m:n} < \tau < t_{N_1+1:m:n} < \dots < t_{m:m:n}$$

^۵Generalized linear model

^۶Type-II Progressive Censoring with Random Dependent removal

اگر $S(t)$ نشان دهنده سطح استرس در زمان t باشد، در این صورت

$$S(t) = \begin{cases} S_1 & \text{if } 0 \leq t < \tau \\ S_2 & \text{if } \tau < t < T_{m:m:n} \end{cases}$$

همچنین فرض کنید S_H حد بالایی سطح استرس و S_U حد نرمال سطح استرس باشد در این صورت مقدار استرس استاندارد شده، $x(t)$ ، به صورت زیر تعریف می شود:

$$x(t) = \frac{S(t) - S_U}{S_H - S_U} = \begin{cases} x_1 & \text{if } 0 \leq t < \tau \\ x_2 & \text{if } \tau < t < T_{m:m:n} \end{cases}$$

که $x(t) \in [0, 1]$.

فرض کنید R_i (تعداد واحدهای خارج شده در مرحله i ام)، $i = 1, \dots, m-1$ دارای توزیع دوجمله ای شرطی با احتمال موفقیت وابسته به شرایط آزمایش باشد؛ یعنی

$$R_1 | T_{1:m:n}, x(T_{1:m:n}) \sim \text{bin}(n-m, p_1)$$

$$R_i | R_1, \dots, R_{i-1}, T_{1:m:n}, \dots, T_{i:m:n}, x(T_{i:m:n}) \sim \text{bin}(n-m - \sum_{j=1}^{i-1} R_j, p_i)$$

$$R_m = n - \sum_{i=1}^{m-1} R_i - m$$

که در آن p_i ، احتمال خروج تصادفی وابسته در مرحله i ام است که با استفاده از تابع ربط لوجیت به شرایط آزمایش، فاصله شکست $FD_i = t_{i:m:n} - t_{i-1:m:n}$ و مقدار استرس در مرحله i ام $x(t_{i:m:n})$ وابسته می گردد. بنابراین اگر

$$\text{احتمال خروج در مرحله } i \text{ام تنها به فاصله شکست بستگی داشته باشد؛ داریم}$$

$$p_i = \frac{e^{\alpha_0 + \alpha_1 FD_i}}{1 + e^{\alpha_0 + \alpha_1 FD_i}}, i = 2, \dots, m-1$$

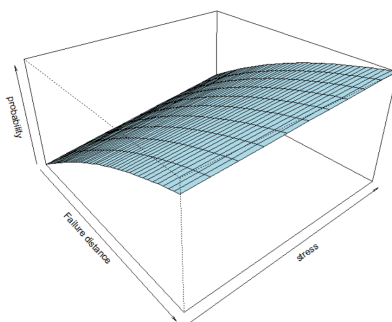
همچنین در صورتی که احتمال خروج در مرحله i ام علاوه بر فاصله شکست به مقدار استرس در مرحله i ام نیز بستگی داشته باشد، داریم:

$$p_i = \frac{e^{\alpha_0 + \alpha_1 FD_i + \alpha_2 x(t_{i:m:n})}}{1 + e^{\alpha_0 + \alpha_1 FD_i + \alpha_2 x(t_{i:m:n})}}, i = 2, \dots, m-1 \quad (1)$$

در این جا α_0 و α_1 و α_2 به عنوان پارامترهای تنظیم مکانیسم سانسور تصادفی می‌باشند که با تنظیم مقادیر آنها می‌توان میزان و نحوه تأثیر شرایط آزمایش بر روی احتمال‌های خروج را کنترل کرد. حسن‌تبار و همکاران [۲] فاصله شکست را به صورت فاصله شکست کوچک ($FD_i \leq 0.1$)، فاصله شکست متوسط ($0.1 < FD_i \leq 2$) و فاصله شکست بزرگ ($FD_i > 2$) در نظر گرفتند و پارامترهای تنظیم در تابع ربط لوجیت را مورد مطالعه قرار دادند. (برای اطلاعات بیشتر [۱] را مشاهده کنید). تعیین و پیشنهاد پارامترهای تنظیم در تابع ربط لوجیت زمانی که احتمال خروج هم به فاصله شکست‌ها و هم به مقدار استرس در هر مرحله بستگی دارد، با توجه به این که در این جا سه پارامتر تنظیم α_0 و α_1 و α_2 داریم، در حالت کلی امکان‌پذیر نیست و محقق در هر آزمایش براساس شرایط آزمون طول عمر خود و معیارهای مورد نظر می‌تواند α_0 و α_1 و α_2 را تنظیم کند. اما در حالت کلی می‌توان نشان داد که:

- زمانی که $\alpha_1 > 0$ باشد، با افزایش فاصله شکست‌ها، احتمال خروج تصادفی افزایش می‌یابد.
- زمانی که $\alpha_1 < 0$ باشد، با افزایش فاصله شکست‌ها، احتمال خروج تصادفی کاهش می‌یابد.
- زمانی که $\alpha_2 > 0$ باشد، با افزایش استرس، احتمال خروج تصادفی افزایش می‌یابد.
- زمانی که $\alpha_2 < 0$ باشد، با افزایش استرس، احتمال خروج تصادفی کاهش می‌یابد.
- هر چه α_0 ، α_1 و α_2 بزرگ‌تر باشد، احتمال خروج تصادفی بزرگ‌تر است.

شکل (۱) مقدار احتمال خروج را برای $\alpha_0 = 0$ ، $\alpha_1 = \alpha_2 = 0.5$ در مقابل استرس (۱، ۰) و فواصل شکست (۴، ۰) نشان می‌دهد.



شکل ۱: احتمال خروج تصادفی وابسته به شرایط آزمایش Type – II PCRD

۳ آزمون شتابیده استرس پله‌ای ساده تحت سانسور $Type - II$ PCRD با تأکید بر توزیع نمایی

فرض کنید در یک simple SSALT طول عمر واحدهای آزمایش از توزیع نمایی با میانگین θ_i در هر سطح استرس S_i ، $i = 1, 2$ پیروی می‌کند. بنابراین تابع توزیع تجمعی طول عمر واحدهای آزمایش در سطح استرس S_i به صورت:

$$F_i(t) = \begin{cases} 0 & t < 0 \\ 1 - e^{-\frac{t}{\theta_i}} & t \geq 0 \end{cases} \quad (1)$$

و تابع چگالی احتمال آن به صورت:

$$f_i(t) = \begin{cases} 0 & t < 0 \\ \frac{1}{\theta_i} e^{-\frac{t}{\theta_i}} & t \geq 0 \end{cases} \quad (2)$$

می‌باشد. برای تحلیل داده‌ها تحت یک simple SSALT نیاز به مدلی است که تابع توزیع تجمعی طول عمر تحت سطوح استرس مختلف را به تابع توزیع تجمعی طول عمر واحد در شرایط نرمال مرتبط کند. در این جا، از مدل پرکاربرد به نام مدل نلسون استفاده می‌کنیم. بنابراین تابع توزیع تجمعی و تابع چگالی احتمال طول عمر یک simple SSALT تحت مدل نلسون به ترتیب به صورت زیر است:

$$G(t) = \begin{cases} G_1(t) = 1 - \exp\left(\frac{-t}{\theta_1}\right) & 0 < t < \tau \\ G_2(t) = 1 - \exp\left(\frac{-1}{\theta_2}(t - \tau) - \frac{1}{\theta_1}\tau\right) & \tau \leq t < \infty \end{cases} \quad (3)$$

$$g(t) = \begin{cases} g_1(t) = \frac{1}{\theta_1} \exp\left(\frac{-t}{\theta_1}\right) & 0 < t < \tau \\ g_2(t) = \frac{1}{\theta_2} \exp\left(\frac{-1}{\theta_2}(t - \tau) - \frac{1}{\theta_1}\tau\right) & \tau \leq t < \infty \end{cases} \quad (4)$$

بنابراین در حالتی که N_1 تعداد شکست‌ها تحت استرس S_1 و N_2 تعداد شکست‌ها تحت استرس S_2 باشد و $N_1 \neq 0$ و $N_2 \neq 0$ باشند، تابع درستنمایی simple SSALT تحت سانسور $Type - II$ PCRD با تابع ربط لوجیت برای توزیع نمایی به صورت زیر خواهد بود:

$$L(\theta_1, \theta_2 | r_m, t_m) = C \prod_{i=1}^{N_1} g_1(t_{i:m:n}) [1 - G_1(t_{i:m:n})]^{r_i} \prod_{i=N_1+1}^m g_2(t_{i:m:n}) [1 - G_2(t_{i:m:n})]^{r_i} \prod_{i=1}^{m-1} p_i^{r_i} (1 - p_i)^{n-m-\sum_{j=1}^{i-1} r_j} \quad (5)$$

که در آن $t_m = (t_{1:m:n}, \dots, t_{m:m:n})$ داده‌های طول عمر مشاهده شده در آزمون طول عمر طراحی شده و $r_m = (r_1, \dots, r_m)$ بردار خروج مشاهده شده تحت مکانیسم تصادفی بیان شده است و همچنین

$$C = n \prod_{i=2}^m (n - (i - 1) - \sum_{j=1}^{i-1} r_j) = \prod_{j=1}^m \sum_{k=j}^m (r_k + 1)$$

با جایگذاری روابط (۶)، (۳) و (۱) در رابطه (۵) تابع درستنمایی به فرم زیر خواهد بود:

$$L(\theta_1, \theta_2 | r_m, t_m) = C \left(\frac{1}{\theta_1} \right)^{N_1} \left(\frac{1}{\theta_2} \right)^{N_2} \exp\left(-\frac{A}{\theta_1}\right) \exp\left(-\frac{B}{\theta_2}\right) \prod_{i=1}^{m-1} \binom{n-m-\sum_{j=1}^{i-1} r_j}{r_i} \exp(r_i(\alpha_0 + \alpha_1 F D_i + \alpha_2 x(t_{i:m:n}))) (1 + \exp(\alpha_0 + \alpha_1 F D_i + \alpha_2 x(t_{i:m:n})))^{-(n-m-\sum_{j=1}^{i-1} r_j)}$$

$$, 0 < t_{1:m:n} < \dots < t_{N_1+m:m:n} < \tau < t_{N_1+1:m:n} < t_{m:m:n}$$

که در آن

$$A = \sum_{i=1}^{N_1} (r_i + 1) t_{i:m:n} + \tau \sum_{i=N_1+1}^m (r_i + 1)$$

$$B = \sum_{i=N_1+1}^m (r_i + 1)(t_{i:m:n} - \tau)$$

در حالت خاص اگر $N_1 = m$ و $N_2 = 0$ باشد، تابع درستنمایی به صورت زیر خواهد بود:

$$L(\theta_1, \theta_2 \mid r_m, t_m) = C\left(\frac{1}{\theta_1}\right)^m \exp\left(-\frac{1}{\theta_1} \sum_{i=1}^m (r_i + 1)t_{i:m:n}\right) \\ \prod_{i=1}^{m-1} \binom{n-m-\sum_{j=1}^{i-1} r_j}{r_i} \exp(r_i(\alpha_0 + \alpha_1 F D_i + \alpha_2 x(t_{i:m:n}))) \\ (1 + \exp(\alpha_0 + \alpha_1 F D_i + \alpha_2 x(t_{i:m:n})))^{-(n-m-\sum_{j=1}^{i-1} r_j)} \\ , 0 < t_{1:m:n} < \dots < t_{m:m:n} < \tau$$

همچنین در حالتی که $N_1 = 0$ و $N_2 = m$ ، تابع درستنمایی به صورت زیر خواهد بود:

$$L(\theta_1, \theta_2 \mid r_m, t_m) = C\left(\frac{1}{\theta_2}\right)^m \exp\left(-\frac{1}{\theta_2} \sum_{i=1}^m (r_i + 1)(t_{i:m:n} - \tau)\right) \exp\left(-\frac{\tau}{\theta_1} \sum_{i=1}^m (R_i + 1)\right) \\ \prod_{i=1}^{m-1} \binom{n-m-\sum_{j=1}^{i-1} r_j}{r_i} \exp(r_i(\alpha_0 + \alpha_1 F D_i + \alpha_2 x(t_{i:m:n}))) \\ (1 + \exp(\alpha_0 + \alpha_1 F D_i + \alpha_2 x(t_{i:m:n})))^{-(n-m-\sum_{j=1}^{i-1} r_j)} \\ , \tau < t_{1:m:n} < \dots < t_{m:m:n} < \infty$$

۴ استنباط آماری

زمانی که حداقل یک شکست قبل از τ و حداقل یک شکست بعد از τ رخ داده باشد (حالتی که $N_1 \neq 0$ و $N_2 \neq 0$)، برآوردهای ماکزیم درستنمایی پارامترهای θ_1 و θ_2 ، یعنی $\hat{\theta}_1$ و $\hat{\theta}_2$ به صورت زیر می باشد:

$$\hat{\theta}_1 = \frac{A}{N_1}$$

و

$$\hat{\theta}_2 = \frac{B}{N_2}$$

زمانی که $N_1 = m$ و $N_2 = 0$ باشد؛ در این صورت:

$$\hat{\theta}_1 = \frac{1}{m} \sum_{i=1}^m (r_i + 1) t_{i:m:n}$$

در این حالت $\hat{\theta}_2$ وجود ندارد. همچنین در صورتی که $N_1 = 0$ و $N_2 = m$ باشد؛ در این صورت:

$$\hat{\theta}_2 = \frac{1}{m} \sum_{i=1}^m (r_i + 1) (t_{i:m:n} - \tau)$$

در این حالت $\hat{\theta}_1$ وجود ندارد.

همچنین درایه‌های ماتریس اطلاع فیشر مشاهده شده که شامل مشتقات جزئی مرتبه دوم لگاریتم درستنمایی هستند برای حالتی که $N_1 \neq 0$ و $N_2 \neq 0$ است به صورت زیر هستند:

$$\frac{\partial^2}{\partial \theta_1^2} \ell(\theta_1, \theta_2 | r_m, t_m) = \frac{N_1}{\theta_1^3} - \frac{2}{\theta_1^3} \left(\sum_{i=1}^m (r_i + 1) t_{i:m:n} + \tau \sum_{i=N_1+1}^m (r_i + 1) \right)$$

$$\frac{\partial^2}{\partial \theta_1 \partial \theta_2} \ell(\theta_1, \theta_2 | r_m, t_m) = \frac{\partial^2}{\partial \theta_2 \partial \theta_1} \ell(\theta_1, \theta_2 | r_m, t_m) = 0$$

$$\frac{\partial^2}{\partial \theta_2^2} \ell(\theta_1, \theta_2 | r_m, t_m) = \frac{N_2}{\theta_2^3} - \frac{2}{\theta_2^3} \sum_{i=N_1+1}^m (r_i + 1) (t_{i:m:n} - \tau)$$

بنابراین ماتریس اطلاع فیشر مشاهده شده در نقطه‌ی براورد ماکزیم درستنمایی پارامترهای θ_1 و θ_2 به صورت زیر است:

$$I(\hat{\theta}_1, \hat{\theta}_2) = - \begin{pmatrix} \frac{\partial^2}{\partial \theta_1^2} \ell(\theta_1, \theta_2 | r_m, t_m) & \frac{\partial^2}{\partial \theta_1 \partial \theta_2} \ell(\theta_1, \theta_2 | r_m, t_m) \\ \frac{\partial^2}{\partial \theta_2 \partial \theta_1} \ell(\theta_1, \theta_2 | r_m, t_m) & \frac{\partial^2}{\partial \theta_2^2} \ell(\theta_1, \theta_2 | r_m, t_m) \end{pmatrix} \bigg|_{\theta_1 = \hat{\theta}_1, \theta_2 = \hat{\theta}_2}$$

ماتریس واریانس-کوواریانس براوردگرهای ماکزیم درستنمایی پارامترهای $(\hat{\theta}_1, \hat{\theta}_2)$ زمانی که اندازه‌ی نمونه به اندازه کافی بزرگ باشد، $v = \lim I^{-1}(\hat{\theta}_1, \hat{\theta}_2)$ تقریب قابل قبولی ارائه می‌کند. بنابراین توزیع توأم براوردگرهای ماکزیم درستنمایی $(\hat{\theta}_1, \hat{\theta}_2)$ تقریباً توزیع نرمال دو متغیره به صورت زیر خواهد بود:

$$\begin{pmatrix} \hat{\theta}_1 - \theta_1 \\ \hat{\theta}_2 - \theta_2 \end{pmatrix} \sim N \left(\begin{pmatrix} 0 \\ 0 \end{pmatrix}, V \right) = N \left(\begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} \hat{v}ar(\hat{\theta}_1) & \hat{c}ov(\hat{\theta}_1, \hat{\theta}_2) \\ \hat{c}ov(\hat{\theta}_2, \hat{\theta}_1) & \hat{v}ar(\hat{\theta}_2) \end{pmatrix} \right)$$

برقراری رابطه فوق، در *Simple SSALTs* تحت *Type-II PCRD* استنباط‌های لازم برای پارامترهای توزیع را ساده خواهد کرد [۱].

۵ مطالعات عددی

۱.۵ الگوریتم تولید داده‌ای حاصل از آزمون طول عمر شتابیده استرس پله‌ای ساده تحت $Type -$

II PCRD

الگوریتم تولید داده‌های سانسور فزاینده نوع دوم با بردار خروج ثابت توسط بالاگرایش و سند [۷] ارائه شده است. این الگوریتم برای تولید داده‌های طول عمر از یک آزمون طول عمر شتابیده استرس پله‌ای ساده تحت سانسور فزاینده نوع دوم با بردار خروج تصادفی وابسته مبتنی بر تابع ربط لوجیت قابل تعمیم است و مراحل آن به شرح زیر می‌باشد:

۱. نمونه تصادفی مستقل $W_1, \dots, W_m$ را از توزیع یکنواخت استاندارد تولید کنید.

۲. قرار دهید $change = 1 - F_1(\tau)$ (با کمک رابطه (۱))

۳. قرار دهید $V_m = W_m^{\frac{1}{n}}$

۴. قرار دهید $V_{-m} = 1$

۵. برای $i = 1, \dots, m$ مراحل زیر را تکرار کنید

۶. $V_{-m} = V_{-m} * V_{m-i+1}$

۷. $U_i = 1 - V_{-m}$

۸. اگر $U_i \leq change$

۹. $T_{i:m:n} = G_1^{-1}(U_i)$ (با کمک رابطه (۳))

۱۰. p_1 را برحسب رابطه (۶) محاسبه کنید.

۱۱. در غیر این صورت

۱۲. $T_{i:m:n} = G_1^{-1}(U_i)$ (با کمک رابطه (۳))

۱۳. p_i را برحسب رابطه (۶) محاسبه کنید.

$$14. \text{ اگر } i < m$$

$$15. R_i \sim \text{binomial}(n - m - \sum_{j=1}^{i-1} R_j, p_i)$$

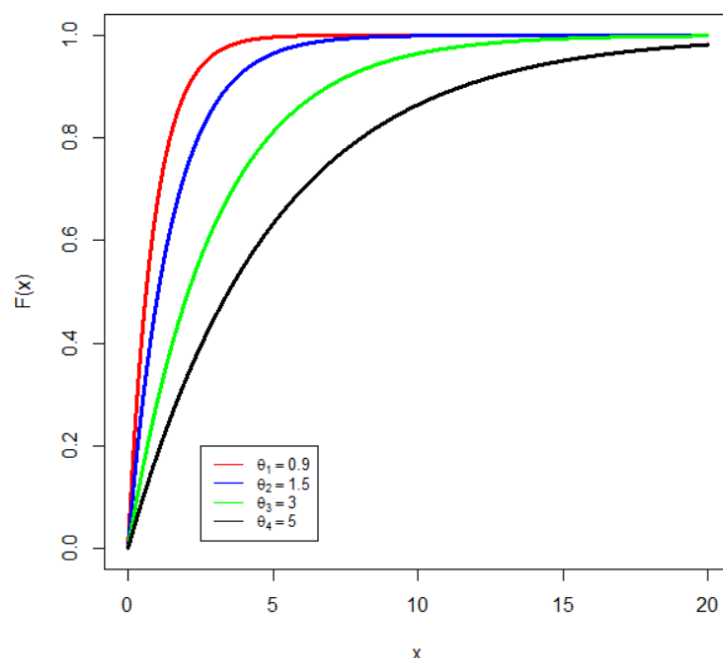
$$16. \text{ قرار دهید } V_{m-i} = W_{m-i}^{\left(\frac{1}{(n-i-\sum_{j=1}^i R_j)}\right)}$$

$$17. \text{ در غیر این صورت}$$

$$18. R_m = n - m - \sum_{j=1}^{m-1} R_j$$

برای تولید نمونه‌های Type-II PCRD در یک SSALT از توزیع نمایی با توجه به این که در هر مرحله باید تعداد واحدهای سانسور شده براساس مکانیسم سانسور مشخص گردد، باید پارامترهای تنظیم مکانیسم سانسور تصادفی وابسته $(\alpha_0, \alpha_1, \alpha_2)$ توسط آزمایشگر تعیین گردد. تعیین این پارامترها برای داشتن خروج‌های تصادفی مطابق با فواصل شکست و همچنین مقدار استرس‌ها بسیار حائز اهمیت می‌باشد. منطقی است که انتظار داشته باشیم آزمایشگر α_2 را به گونه‌ای انتخاب کند که با افزایش استرس، جهت صرفه‌جویی اقتصادی، احتمال خروج‌های تصادفی افزایش یابد، بنابراین $\alpha_2 > 0$ منطقی به نظر می‌رسد. از طرفی نوع توزیع داده‌های طول عمر می‌تواند تا حدودی رفتار فواصل شکست احتمالی را در آزمایشات طول عمر پیش‌بینی‌پذیر کند. بنابراین برای انتخاب پارامترهای تنظیم در مکانیسم سانسور تصادفی نیاز داریم در ابتدا ایده‌ای در مورد فواصل شکست احتمالی داشته باشیم تا با استفاده از آن بتوانیم پارامترهای تنظیم را تعیین و در مکانیسم سانسور تصادفی وابسته قرار دهیم. در توزیع نمایی، توزیع مورد بحث در این مقاله، نرخ شکست ثابت است و تابعی از زمان نمی‌باشد. همان‌طور که در شکل (۲) مشاهده می‌کنید با افزایش مقدار مقیاس θ (میانگین) احتمال مشاهده شکست کوچک کاهش می‌یابد.

در شکل (۳)، احتمال خروج در مقابل FD_i ها در یک Simple SSALT تحت توزیع نمایی با پارامترهای $\theta_1 = 0.9$ و $\theta_2 = 0.5$ برای $\tau = 0.2$ ، $x_1 = 0.1$ ، $x_2 = 0.5$ ، $n = 30$ و $m = 15$ می‌باشد. در این حالت بیشتر فاصله شکست‌ها متعلق به فاصله شکست کوچک می‌باشد. بنابراین برای این که احتمال خروج تصادفی متأثر از فاصله‌های شکست باشد نیاز به یک α_1 بزرگ است. در شکل (۳)، با انتخاب $\alpha_0 = -2$ با α_1 های متفاوت، تأثیر فاصله‌های شکست بر روی احتمال‌های خروج تصادفی مشخص است. شکل (۳) ب برای $\alpha_1 = 10$ و $\alpha_2 = 0$ به ازای α_0 های متفاوت است. اگر محقق علاقه‌مند به خروج واحدها در مراحل ابتدایی آزمایش جهت جلوگیری از ضرر بیشتر باشد

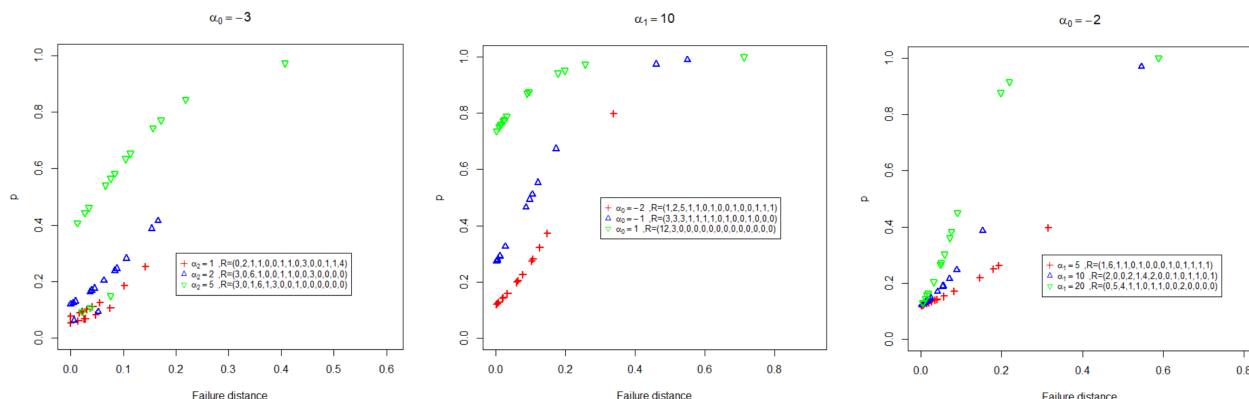


شکل ۲: تابع توزیع تجمعی توزیع نمایی برای میانگین‌های مختلف

نیاز به ایجاد احتمال خروج‌های بزرگ در مکانیسم سانسور تصادفی وابسته دارد. برای این منظور می‌تواند از $\alpha_0 = 1$ جهت ایجاد احتمال‌های خروج بزرگ‌تر استفاده کند. شکل (۳) ج برای $\alpha_0 = -3$ و $\alpha_1 = 10$ به ازای α_2 ‌های متفاوت است، همان‌طور که از شکل پیداست زمانی که احتمال‌های خروج علاوه بر فاصله شکست، به استرس‌ها نیز بستگی دارند، احتمال‌های خروج تحت تأثیر قرار می‌گیرند و فواصل شکست یکسان تحت استرس‌های متفاوت دارای احتمال‌های خروج متفاوت هستند.

۲.۵ تنظیم پارامترهای ورودی در شبیه‌سازی و ارائه نتایج

در این بخش مطالعه‌ای از یک Simple SSALT تحت Type- II PCRD برای توزیع نمایی در دو مثال ارائه شده است. تعداد نمونه‌های تولید شده $(n = 20 \text{ و } m = 10)$ و $(n = 30 \text{ و } m = 10)$ و استرس‌ها $x_1 = 0/1$ و $x_2 = 0/5$ در نظر گرفته شده است. معیارهای گزارش داده شده برای برآوردها، مقدار اریبی ($Bias$)، میانگین مجموع توان دوم خطا (MSE)، میانگین زمان آزمایش ($E(T_m)$) و نرخ همگرایی (CR) می‌باشد. نتایج براساس ۱۰۰۰۰ تکرار مطالعات شبیه‌سازی ارائه شده است.



(آ) احتمال خروج وابسته به فواصل شکست (ب) احتمال خروج وابسته به فواصل شکست (ج) احتمال خروج وابسته به فواصل شکست
 $\alpha_2 = 0$ $\alpha_2 = 0$ و استرس و $\alpha_1 = 10$

شکل ۳: احتمال‌های خروج وابسته به مکانیسم تصادفی

مثال ۱۰۵. در این مثال نمونه‌های تولید شده (زمان‌های شکست) از یک *Simple SSALT* تحت *Type-II PCRD* و *Type-II PCR^y* برای توزیع نمایی با پارامترهای $\theta_1 = 0.9$ و $\theta_2 = 0.5$ می‌باشد که زمان تغییر استرس S_1 به S_2 ، $\tau = 0.2$ ، در نظر گرفته شده است. در این حالت حدود ۸۰ درصد داده‌ها متعلق به زمان‌های شکست کوچک می‌باشند. بنابراین با انتخاب α_1 بزرگ (در این جا $\alpha_1 = 10$) احتمال‌های خروج به فواصل شکست وابسته می‌شوند. نتایج در جدول (۱) ارائه شده است. طول مدت زمان آزمایش در مدل *Type-II PCRD* با توجه به پارامترهای تنظیم در نظر گرفته شده، از مدل *Type-II PCR*، با توزیع دو جمله‌ای با احتمال‌های خروج ثابت با پارامترهای ($p = 0.25, 0.5, 0.75$) و همچنین از توزیع یکنواخت کم‌تر شده است. مقدار آریبی و میانگین مجموع توان دوم خطاها برای برآوردهای ماکزیم درست‌نمایی در هر دو مدل با هم تفاوت زیادی ندارند. در واقع مکانیسم سانسور و توزیع تعداد خروج‌ها در سانسور فزاینده روی برآوردهای ماکزیم درست‌نمایی تأثیری ندارد که این موضوع توسط کرامر و لیپولس [۸] اثبات شده است. بنابراین انتظار می‌رود که برآوردهای ماکزیم درست‌نمایی ویژگی یکسانی در سانسور فزاینده نوع دوم با خروج‌های متفاوت (خروج ثابت، تصادفی و تصادفی وابسته) داشته باشند. نرخ همگرایی ارائه شده برای θ_1 و θ_2 به سطح اطمینان اسمی (۰/۹۰) بسیار نزدیک است، نرخ همگرایی، نسبت دفعاتی است که فاصله اطمینان برآوردها، مقدار واقعی پارامتر را در ۱۰۰۰۰ تکرار مطالعه شبیه‌سازی در بر گرفته است.

^yType-II Progressive Censoring with Random Removal

مثال ۲.۵. در این مثال نمونه‌های تولید شده (زمان‌های شکست) از یک *Simple SSALT* تحت *II PCRD* *Type- II PCR* و *Type-* برای توزیع نمایی با پارامتر $\theta_1 = 3$ و $\theta_2 = 1/5$ می‌باشد که زمان تغییر استرس S_1 به S_2 ، $\tau = 0/5$ در نظر گرفته شده است. در این حالت تقریباً احتمال دیدن زمان‌های شکست با فاصله زمان‌های شکست کوچک و متوسط یکسان است. بنابراین انتخاب α_1 متوسط (مانند ۵ یا ۲) می‌تواند مناسب باشد. نتایج در جدول (۲) ارائه شده است. در این مثال همانند مثال قبل طول مدت زمان آزمایش تحت *Type II PCR* با توزیع دو جمله‌ای با احتمال‌های خروج ثابت با پارامترهای $(p = 0/25, 0/5, 0/75)$ و همچنین از توزیع یکنواخت کمتر شده است. مقدار آریبی و میانگین مجموع توان دوم خطا برای برآوردهای ماکزیم درست‌نمایی در هر دو مدل با هم تفاوت زیادی ندارد. نرخ همگرایی ارائه شده برای θ_1 و θ_2 به سطح اطمینان اسمی $(0/90)$ بسیار نزدیک است.

جدول ۱: مقایسه برآوردهای ماکسیمم درستمایی و میانگین زمان آزمایش در سانسورهای فزاینده نوع دوم با مکانیسم تصادفی برای $\theta_1 = 0.9$, $\theta_2 = 0.5$ و $x_1 = 0.5$ و $x_2 = 0.5$

n	m	رویکرد	پارامتر تنظیم	$\hat{\theta}_1$	$Bias(\hat{\theta}_1)$	$MSE(\hat{\theta}_1)$	$\hat{\theta}_2$	$Bias(\hat{\theta}_2)$	$MSE(\hat{\theta}_2)$	CR_{θ_1}	CR_{θ_2}	$E(T_m)$
۲۰	۱۰	Type - II PCRD	$\alpha_0 = -2, \alpha_1 = 1^0, \alpha_2 = 0$	۰/۹۶۸۲	۰/۰۶۸۲	۰/۱۹۳۱	۰/۵۰۱۰	۰/۰۰۱۰	۰/۰۴۲۳	۰/۸۹۶۲	۰/۸۸۲۰	۰/۹۹۹۹
		Type - II PCRD	$\alpha_0 = -1, \alpha_1 = -1^0, \alpha_2 = 0$	۰/۹۶۹۱	۰/۰۶۹۱	۰/۲۰۰۹	۰/۵۰۲۰	۰/۰۰۲۰	۰/۰۴۱۲	۰/۸۹۲۲	۰/۸۸۰۷	۰/۸۷۵۲
		Type - II PCRD	$\alpha_0 = -2/5, \alpha_1 = 1^0, \alpha_2 = 1$	۰/۹۸۳۶	۰/۰۸۳۶	۰/۲۰۱۹	۰/۵۰۴۷	۰/۰۰۴۷	۰/۰۴۲۴	۰/۸۹۷۳	۰/۸۸۵۴	۰/۹۰۳۸
		Type - II PCRD	$\alpha_0 = -2/5, \alpha_1 = -1^0, \alpha_2 = 3$	۰/۹۸۶۵	۰/۰۸۶۵	۰/۲۰۸۴	۰/۴۹۸۲	-۰/۰۰۱۸	۰/۰۴۳۲	۰/۸۷۰۸	۰/۸۹۹۳	۰/۶۷۵۲
	۱۵	Type - II PCR	$p = 0.25$ با جمله‌ای	۰/۹۵۶۳	۰/۰۵۶۳	۰/۱۸۸۱	۰/۵۰۰۵	۰/۰۰۰۵	۰/۰۴۰۹	۰/۸۸۹۰	۰/۸۸۲۰	۱/۰۸۵۶
		Type - II PCR	$p = 0.5$ با جمله‌ای	۰/۸۹۴۵	-۰/۰۰۵۵	۰/۱۵۵۶	۰/۴۹۸۷	-۰/۰۰۱۳	۰/۰۳۸۲	۰/۸۷۴۷	۰/۸۸۱۶	۱/۴۴۶۷
		Type - II PCR	$p = 0.75$ با جمله‌ای	۰/۸۵۳۲	-۰/۰۴۶۸	۰/۱۳۱۶	۰/۴۹۹۲	-۰/۰۰۰۸	۰/۰۳۶۲	۰/۸۶۲۰	۰/۸۸۷۶	۱/۴۸۹۸
		Type - II PCR	یکنواخت	۰/۸۹۸۱	-۰/۰۰۱۹	۰/۱۵۸۱	۰/۴۹۶۴	-۰/۰۰۳۶	۰/۰۳۷۱	۰/۸۷۳۷	۰/۸۸۰۵	۱/۴۴۲۱
	۳۰	Type - II PCRD	$\alpha_0 = -2, \alpha_1 = 1^0, \alpha_2 = 0$	۰/۹۴۶۲	۰/۰۴۶۲	۰/۱۳۴۵	۰/۵۰۰۳	۰/۰۰۰۳	۰/۰۲۶۳	۰/۹۰۳۸	۰/۹۰۲۲	۱/۲۷۴۵
		Type - II PCRD	$\alpha_0 = -1, \alpha_1 = -1^0, \alpha_2 = 0$	۰/۹۳۰۰	۰/۰۳۰۰	۰/۱۳۲۸	۰/۵۰۲۹	۰/۰۰۲۹	۰/۰۲۵۳	۰/۹۰۱۳	۰/۹۰۶۶	۱/۲۷۸۴
		Type - II PCRD	$\alpha_0 = -2/5, \alpha_1 = 1^0, \alpha_2 = 1$	۰/۹۷۰۱	۰/۰۷۰۱	۰/۱۵۴۰	۰/۵۰۲۸	۰/۰۰۲۸	۰/۰۲۸۴	۰/۹۱۱۹	۰/۹۰۱۸	۰/۶۵۵۸
		Type - II PCRD	$\alpha_0 = -2/5, \alpha_1 = -1^0, \alpha_2 = 3$	۰/۹۶۹۳	۰/۰۶۹۳	۰/۱۴۹۳	۰/۴۹۹۸	-۰/۰۰۰۲	۰/۰۲۷۷	۰/۹۱۳۱	۰/۸۹۶۹	۰/۹۶۹۳
۳۰	۱۵	Type - II PCR	$p = 0.25$ با جمله‌ای	۰/۹۲۲۲	۰/۰۲۲۲	۰/۱۲۷۱	۰/۵۰۱۷	۰/۰۰۱۷	۰/۰۲۵۸	۰/۸۹۷۴	۰/۹۰۱۷	۱/۴۷۰۲
		Type - II PCR	$p = 0.5$ با جمله‌ای	۰/۸۶۹۶	-۰/۰۳۰۴	۰/۰۹۹۸	۰/۴۹۹۳	-۰/۰۰۰۷	۰/۰۲۴۰	۰/۸۹۲۴	۰/۹۰۶۱	۱/۶۷۸۸
		Type - II PCR	$p = 0.75$ با جمله‌ای	۰/۸۱۴۰	-۰/۰۸۶۰	۰/۰۸۷۲	۰/۵۰۱۵	۰/۰۰۱۵	۰/۰۲۳۹	۰/۸۸۲۲	۰/۹۰۸۷	۱/۷۰۵۱
		Type - II PCR	یکنواخت	۰/۸۶۱۱	-۰/۰۳۸۹	۰/۱۰۲۵	۰/۴۹۸۶	-۰/۰۰۱۴	۰/۰۲۳۹	۰/۸۸۵۷	۰/۹۰۷۱	۱/۶۸۱۶

جدول ۲: مقایسه برآوردهای ماکسیمم درستمایی و میانگین زمان آزمایش در سانسورهای فزاینده نوع دوم با مکانیسم تصادفی برای $\theta_1 = 3$, $\theta_2 = 1.5$ و $\tau = 0.5$, $x_1 = 0.1$ و $x_2 = 0.5$

n	m	دیکرد	پارامتر تنظیم	$\hat{\theta}_1$	$Bias(\hat{\theta}_1)$	$MSE(\hat{\theta}_1)$	$\hat{\theta}_2$	$Bias(\hat{\theta}_2)$	$MSE(\hat{\theta}_2)$	CR_{θ_1}	CR_{θ_2}	$E(T_m)$
۲۰	۱۰	Type - II PCRD	$\alpha. = -4, \alpha_1 = 2, \alpha_2 = 0$	۲/۹۹۰۰	-۰/۰۰۱۰	۱/۴۰۲۸	۱/۵۰۵۴	۰/۰۰۵۴	۰/۳۶۹۸	۰/۹۰۱۴	۰/۸۸۰۹	۱/۳۲۶۴
		Type - II PCRD	$\alpha. = -3, \alpha_1 = 5, \alpha_2 = 0$	۲/۹۳۲۸	-۰/۰۰۶۷۲	۱/۳۱۵۸	۱/۴۹۷۷	-۰/۰۰۲۳	۰/۳۴۳۳	۰/۸۸۲۵	۰/۸۸۶۵	۲/۰۴۹۲
		Type - II PCRD	$\alpha. = -4, \alpha_1 = 2, \alpha_2 = 1$	۲/۹۸۷۴	-۰/۰۰۱۲۶	۱/۳۷۹۶	۱/۴۹۷۶	-۰/۰۰۲۴	۰/۳۵۲۱	۰/۹۰۵۴	۰/۸۸۴۰	۱/۳۶۷۳
		Type - II PCRD	$\alpha. = -3, \alpha_1 = -5, \alpha_2 = 1$	۲/۹۸۳۲	-۰/۰۰۱۶۸	۱/۳۹۸۴	۱/۴۹۴۰	-۰/۰۰۰۶۰	۰/۳۵۶۶	۰/۸۹۲۶	۰/۸۷۹۵	۱/۳۹۸۷
	۱۵	Type - II PCR	$p = 0.25$ با جمله‌ای	۲/۷۷۸۴	-۰/۰۰۱۶	۱/۲۶۰۰	۱/۴۹۸۵	-۰/۰۰۱۵	۰/۳۳۴۳	۰/۸۶۴۵	۰/۸۸۳۶	۳/۲۳۸۲
		Type - II PCR	$p = 0.5$ با جمله‌ای	۲/۶۰۰۶	-۰/۰۰۳۹۹۳	۱/۱۵۲۲	۱/۵۰۶۹	۰/۰۰۰۶۹	۰/۳۲۴۲	۰/۸۵۳۰	۰/۸۸۹۹	۴/۳۰۴۱
		Type - II PCR	$p = 0.75$ با جمله‌ای	۲/۴۱۳۰	-۰/۰۰۵۸۷۰	۱/۱۵۰۴	۱/۴۹۶۳	-۰/۰۰۳۷	۰/۳۰۹۸	۰/۸۴۷۵	۰/۸۸۷۳	۴/۴۱۹۴
		Type - II PCR	یکتاخت	۲/۶۱۳۷	-۰/۰۰۳۸۶۳	۱/۲۱۷۲	۱/۵۰۳۵	۰/۰۰۳۵	۰/۳۲۰۰	۰/۸۵۶۴	۰/۸۹۰۲	۴/۳۱۰۰
۳۰	۱۵	Type - II PCRD	$\alpha. = -4, \alpha_1 = 2, \alpha_2 = 0$	۳/۴۲۵۷	۰/۴۲۵۷	۲/۷۶۶۶	۱/۴۸۷۵	-۰/۰۰۱۲۵	۰/۲۲۴۹	۰/۹۱۷۴	۰/۸۹۴۷	۱/۴۰۶۸
		Type - II PCRD	$\alpha. = -3, \alpha_1 = 5, \alpha_2 = 0$	۳/۳۹۳۵	۰/۳۹۳۵	۲/۶۷۱۹	۱/۵۰۸۷	۰/۰۰۸۷	۰/۲۲۰۵	۰/۹۰۶۹	۰/۹۰۷۱	۲/۳۳۹۷
		Type - II PCRD	$\alpha. = -4, \alpha_1 = 2, \alpha_2 = 1$	۳/۳۹۰۱	۰/۳۹۰۱	۲/۶۵۰۶	۱/۴۹۸۲	-۰/۰۰۱۸	۰/۲۲۲۰	۰/۹۲۰۰	۰/۹۰۵۶	۱/۴۸۳۲
		Type - II PCRD	$\alpha. = -3, \alpha_1 = -5, \alpha_2 = 1$	۳/۴۱۳۸	۰/۴۱۳۸	۲/۷۲۳۵	۱/۴۹۵۵	-۰/۰۰۴۵	۰/۲۱۷۳	۰/۹۱۴۷	۰/۹۰۱۷	۱/۶۱۳۳
	۱۰	Type - II PCR	$p = 0.25$ با جمله‌ای	۳/۳۲۷۰	۰/۳۲۷۰	۲/۵۶۷۱	۱/۴۹۹۵	-۰/۰۰۰۵	۰/۲۱۱۲	۰/۸۹۶۳	۰/۹۰۲۷	۴/۳۷۲۶
		Type - II PCR	$p = 0.5$ با جمله‌ای	۳/۲۲۷۹	۰/۲۲۷۹	۲/۱۷۰۸	۱/۴۹۴۲	-۰/۰۰۵۸	۰/۱۹۰۸	۰/۸۹۶۶	۰/۹۰۹۳	۵/۰۱۲۹
		Type - II PCR	$p = 0.75$ با جمله‌ای	۳/۰۵۵۱	۰/۰۵۵۱	۱/۷۸۴۳	۱/۴۹۸۶	-۰/۰۰۱۴	۰/۱۹۶۲	۰/۸۸۴۲	۰/۹۰۸۷	۵/۰۱۲۷
		Type - II PCR	یکتاخت	۳/۱۶۰۹	۰/۱۶۰۹	۲/۱۵۵۶	۱/۵۰۲۹	۰/۰۰۲۹	۰/۱۹۹۶	۰/۸۸۵۴	۰/۹۰۶۲	۵/۰۵۹۳

جدول (۳) و (۴) میانگین زمان آزمایش را در مدل Type-II PCRD تحت توزیع نمایی با پارامترهای $\theta_1 = 0.9$ و $\theta_2 = 0.5$ زمانی که $\tau = 0.2$ است نشان می‌دهد:

• جدول (۳):

$$\alpha_0 = \{-4, 1\}, \alpha_1 = \{-20, -10, 0, 10, 20\}, \alpha_2 = \{0\}$$

• جدول (۴):

$$\alpha_0 = \{-4\}, \alpha_1 = \{-10, 10\}, \alpha_2 = \{0.5, 2\}$$

اندازه نمونه‌ها $n = \{6, 20, 30, 50\}$ در نظر گرفته شده است که برای هر یک از آن‌ها m به گونه‌ای انتخاب شده است که نمونه‌های مشاهده شده شامل ۵۰ درصد، ...، ۹۰ درصد و ۱۰۰ درصد از نمونه‌های موجود است. همان‌طور که در جدول (۳) نشان داده شده است میانگین زمان آزمایش نسبت به نمونه‌های شکست کامل ($n = m$) کاهش داشته است اما این کاهش در جدول (۴) یعنی در حالتی که احتمال‌های خروج به استرس‌ها وابسته است بیشتر است.

اگر میانگین زمان آزمایش برای داده‌های کامل را با $E(T^*)$ نشان دهیم، در این صورت میانگین زمان آزمایش نسبی ($REET$) داده‌ها تحت سانسور فزاینده نوع دوم به صورت زیر حاصل می‌گردد

$$REET = \frac{E(T_{m:m:n})}{E(T^*)}$$

این معیار برای اندازه‌گیری میزان کاهش زمان آزمایش در سانسور فزاینده نوع دوم نسبت به داده‌های شکست کامل به کار می‌رود. با توجه به این‌که این معیار به واحد اندازه‌گیری بستگی ندارد بسیار مفید و کاربردی است. شکل‌های (۴) و (۵) نمودارهای میانگین زمان آزمایش برای Simple SSALT را تحت سانسور Type-II PCRD نسبت به طرح شکست کامل را تحت عنوان نمودارهای میانگین زمان آزمایش نسبی نشان می‌دهند. شکل (۴) میانگین زمان آزمایش نسبی در مقابل m را برای n و α_1 ‌های مختلف به ازای $\alpha_0 = \{-4, 1\}$ برای توزیع نمایی با پارامتر $\theta_1 = 0.9$ و $\theta_2 = 0.5$ نشان می‌دهد. هرچه میانگین نسبی زمان آزمایش به یک نزدیک‌تر باشد بیان‌گر یکسان بودن میانگین زمان آزمایش در طرح سانسور و مشاهدات داده‌های شکست کامل می‌باشد و هرچه این فاصله از یک بیشتر و به صفر نزدیک‌تر باشد بیانگر تفاوت زمانی بیشتر در دو نوع طرح و در واقع کارا بودن سانسور است برای $\alpha_0 = -4$ در $n = 6$ اگر $\alpha_1 \neq 0$ هرچه α_1 و m کوچک‌تر هستند طرح سانسور به کار گرفته شده کارا تر است. برای $\alpha_0 = 1$ در $n = 6$ هرچه α_1 و m کوچک‌تر هستند طرح سانسور به کار گرفته شده کارا تر است. در سایر حالات میانگین زمان آزمایش نسبی تنها رابطه مستقیم با m دارد.

جدول ۳: میانگین زمان آزمایش در آزمون طول عمر شتابیده استرس پله‌ای ساده تحت سانسورهای فزاینده نوع دوم با مکانیسم تصادفی برای $\alpha_1 = 0.9, \theta_1 = 0.5, \theta_2 = 0.2, \tau = 0, \alpha_2 = 0, x_1 = 0.1$ و $x_2 = 0.5$

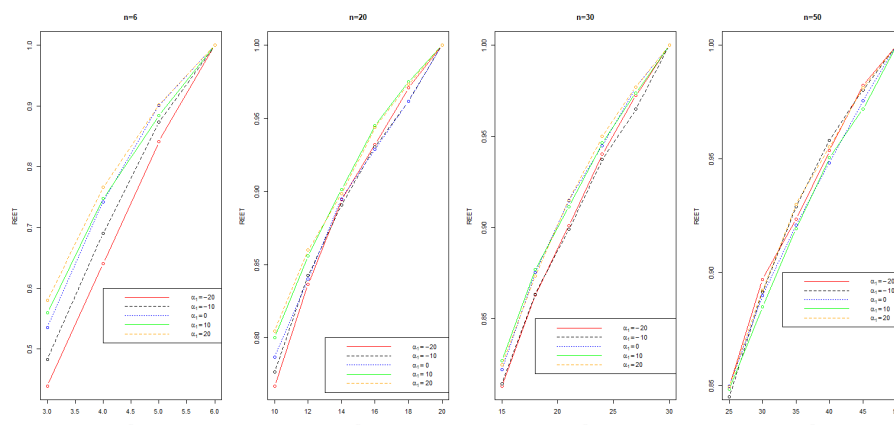
$\alpha_0 = -4$							$\alpha_0 = 1$					
n	m	α_1	-20	-10	0	10	20	-20	-10	0	10	20
6	6		1/1965	1/2061	1/1933	1/2032	1/1992	1/1905	1/1873	1/1893	1/1969	1/1836
6	5		1/0071	1/0535	1/0752	1/0642	1/0815	1/0093	1/0499	1/0609	1/0648	1/0768
6	4		0/7662	0/8323	0/8858	0/8999	0/9184	0/7644	0/8404	0/8865	0/8933	0/9173
6	3		0/5249	0/5816	0/6388	0/6731	0/6961	0/5219	0/5778	0/6507	0/6735	0/6911
20	20		1/8837	1/8938	1/8933	1/8747	1/8755	1/8833	1/8771	1/8778	1/8813	1/8799
20	18		1/8290	1/8212	1/8211	1/8278	1/8253	1/8166	1/8111	1/8344	1/8217	1/8289
20	16		1/7560	1/7622	1/7588	1/7720	1/7702	1/7555	1/7623	1/7508	1/7642	1/7672
20	14		1/6848	1/6867	1/6944	1/6898	1/6854	1/6743	1/6807	1/6868	1/6813	1/7047
20	12		1/5753	1/5952	1/5906	1/6043	1/6124	1/5814	1/5939	1/6034	1/5958	1/6026
20	10		1/4443	1/4707	1/4894	1/4997	1/5087	1/4420	1/4674	1/4871	1/5006	1/5057
30	30		2/0897	2/1003	2/0788	2/0804	2/0805	2/0892	2/0754	2/0862	2/0785	2/0881
30	27		2/0325	2/0267	2/0309	2/0262	2/0331	2/0313	2/0319	2/0228	2/0330	2/0189
30	24		1/9651	1/9689	1/9645	1/9693	1/9767	1/9784	1/9601	1/9671	1/9700	1/9711
30	21		1/8834	1/8883	1/9017	1/8965	1/9040	1/8935	1/8923	1/9028	1/8994	1/8923
30	18		1/8035	1/8136	1/8200	1/8245	1/8173	1/8060	1/8062	1/8214	1/8147	1/8227
30	15		1/6993	1/7106	1/7095	1/7206	1/7159	1/7016	1/7116	1/7100	1/7178	1/7225
50	50		2/3261	2/3256	2/3422	2/3422	2/3318	2/3339	2/3345	2/3394	2/3300	2/3435
50	45		2/2850	2/2797	2/2848	2/2763	2/2884	2/2850	2/2859	2/2914	2/2935	2/2877
50	40		2/2184	2/2284	2/2209	2/2265	2/2273	2/2317	2/2176	2/2234	2/2374	2/2223
50	35		2/1479	2/1605	2/1567	2/1531	2/1686	2/1444	2/1412	2/1641	2/1535	2/1488
50	30		2/0862	2/0731	2/0839	2/0725	2/0801	2/0727	2/0645	2/0816	2/0751	2/0675
50	25		1/9768	1/9656	1/9905	1/9870	1/9812	1/9809	1/9695	1/9825	1/9886	1/9842

جدول ۴: میانگین زمان آزمایش در آزمون طول عمر شتابیده استرس پله‌ای ساده تحت سانسورهای فزاینده نوع دوم با

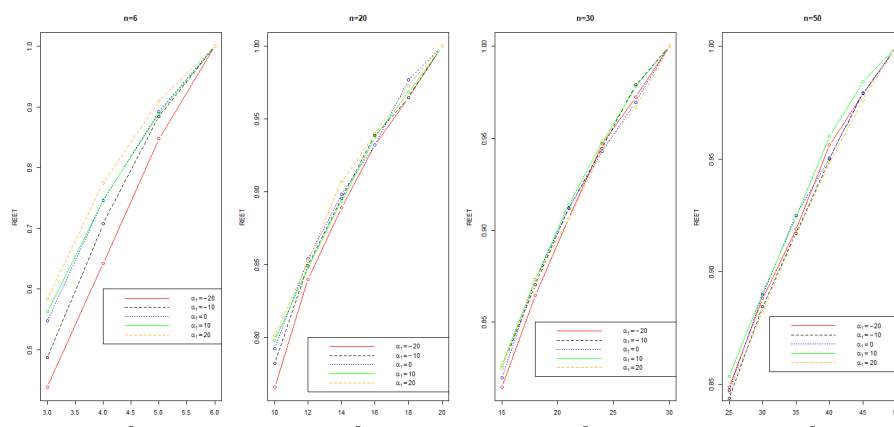
مکانیسم تصادفی برای $\theta_1 = 0.9, \theta_2 = 0.5, \tau = 0.2, \alpha_0 = -4$ ، $x_1 = 0.1$ و $x_2 = 0.5$

$\alpha_T = 0.5$				$\alpha_T = 2$		$\alpha_T = 5$		
n	m	α_1	-1.0	1.0	-1.0	1.0	-1.0	1.0
6	6		1/2.084	1/2.041	1/19.08	1/2.050	1/19.17	1/18.889
6	5		0.9349	0.9902	0.9915	0.8812	0.8853	0.9356
6	4		0.6419	0.6930	0.6915	0.5880	0.5900	0.6435
6	3		0.3557	0.3659	0.3659	0.3503	0.3516	0.3550
20	20		1/18747	1/18843	1/188.08	1/18823	1/18715	1/18857
20	18		1/7350	1/8068	1/80.94	1/5698	1/5655	1/7212
20	16		1/5565	1/7118	1/70.17	1/2643	1/2569	1/5484
20	14		1/3236	1/5907	1/58.52	0.9631	0.9679	1/3270
20	12		1/0.411	1/40.57	1/41.92	0.7226	0.7228	1/0.352
20	10		0.7560	1/1603	1/16.17	0.5377	0.5334	0.7579
30	30		2/0.877	2/0.941	2/0.797	2/0.925	2/0.826	2/0.840
30	27		1/9323	2/0071	2/0.216	1/6967	1/7007	1/9368
30	24		1/7483	1/9143	1/91.90	1/3017	1/3131	1/7568
30	21		1/4927	1/8208	1/80.61	0.9739	0.9739	1/4885
30	18		1/1522	1/6640	1/65.91	0.7216	0.7168	1/1521
30	15		0.8106	1/4224	1/42.46	0.5314	0.5328	0.8089
50	50		2/3316	2/3409	2/3334	2/3458	2/3380	2/3155
50	45		2/2152	2/2801	2/2676	1/8631	1/8675	2/2026
50	40		2/0271	2/1797	2/1757	1/3843	1/3903	2/0133
50	35		1/7527	2/0841	2/0868	1/0041	0.9973	1/7586
50	30		1/3740	1/9597	1/9513	0.7334	0.7303	1/3847
50	25		0.9356	1/7696	1/7628	0.5433	0.5493	0.9372

شکل (۵) میانگین زمان آزمایش نسبی در مقابل m را برای n ، $\alpha_0 = -2$ ، $\alpha_1 = \{-10, 10\}$ و $\alpha_2 = 0.5$ و $\{0.5, 5\}$ برای توزیع نمایی با پارامتر $\theta_1 = 0.9$ و $\theta_2 = 0.5$ زمانی که $\tau = 0.2$ است نشان می‌دهد زمانی که $\alpha_1 = -10$ است در تمام حجم n ، طرح سانسور به کار گرفته شده متأثر از α_2 و m می‌باشد و هرچه α_2 و m کوچک‌تر هستند طرح سانسور به کار گرفته شده کارا تر است. همان‌طور که از شکل پیداست کم‌ترین زمان آزمایش مربوط به $\alpha_2 = 0.5$ برای حجم نمونه‌های مختلف است. نتایج برای $\alpha_1 = 10$ قابل تکرار است.

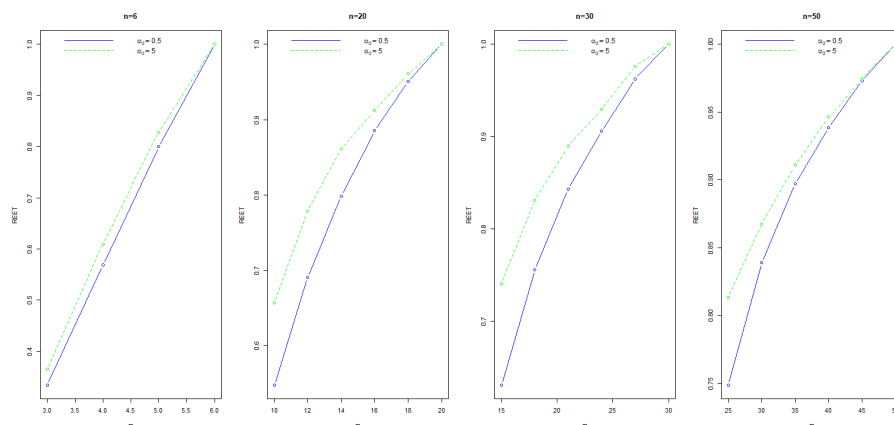


(آ) میانگین زمان آزمایش نسبی برای $\alpha_0 = -4$

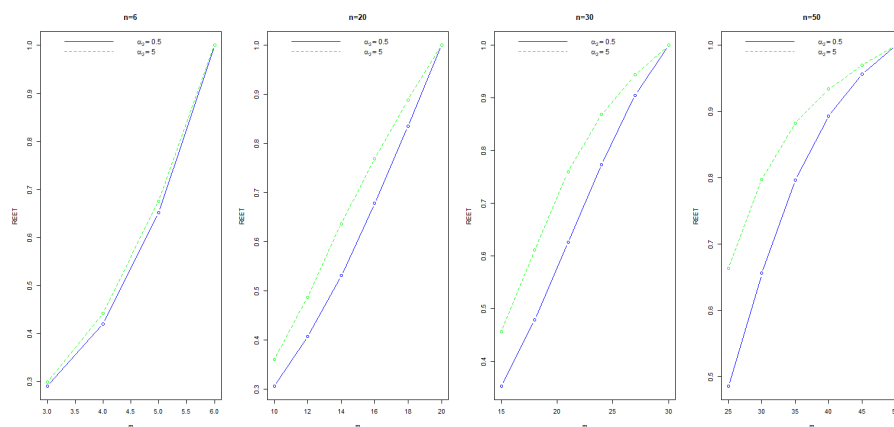


(ب) میانگین زمان آزمایش نسبی برای $\alpha_0 = 1$

شکل ۴: میانگین زمان آزمایش نسبی در توزیع نمایی با پارامتر $\theta_1 = 0.9$ و $\theta_2 = 0.5$ با فرض مقادیر مختلف برای n ، m و $\alpha_2 = 0$.



(آ) میانگین زمان آزمایش نسبی برای $\alpha_1 = 10$



(ب) میانگین زمان آزمایش نسبی برای $\alpha_1 = -10$

شکل ۵: میانگین زمان آزمایش نسبی در توزیع نمایی با پارامتر $\theta_1 = 0.9$ و $\theta_2 = 0.5$ با فرض مقادیر مختلف برای m, n و $\alpha_0 = -2$.

۶ تحلیل حساسیت مکانیسم سانسور تصادفی نسبت به پارامترهای تنظیم

تحلیل حساسیت به مطالعه تأثیرپذیری مدل از متغیرهای ورودی می‌پردازد. به عبارت دیگر در این مطالعه، با تغییر در مقادیر پارامترهای تنظیم به صورت سیستماتیک، اثر این تغییرها بر مکانیسم سانسور تصادفی وابسته بررسی می‌گردد. معمولاً برای تحلیل حساسیت مدل، از اندازه کارایی نسبی استفاده می‌شود. اندازه کارایی نسبی به مقایسه مجموعه‌ای از مقادیر پارامتر، $\theta^{(1)}$ ، با تغییر در یک متغیر ورودی از همان مجموعه مقادیر پارامترها $\theta^{(0)}$ می‌پردازد. این مقایسه با توجه به معیار و شاخص مد نظر محقق، در این جا میانگین زمان آزمایش، می‌تواند تعریف گردد. بنابراین اندازه کارایی

نسبی (RE) برای تحلیل حساسیت به صورت زیر تعریف می‌گردد،

$$RE(\theta) = \frac{E(T_{m:m:n})_{under\theta^{(0)}}}{E(T_{m:m:n})_{under\theta^{(1)}}}$$

مقدار RE ، حساسیت مکانیسم سانسور تصادفی را نسبت به پارامترهای ورودی نشان می‌دهد. در واقع هرچه مقادیر به دست آمده از یک فاصله بیشتری داشته باشند به این معناست که مکانیسم سانسور تصادفی وابسته به پارامترهای ورودی مدل حساسیت بیشتری دارد.

تحلیل حساسیت را برای سانسور $Type-II PCRD$ به کار برده شده در آزمون طول عمر شتابیده استرس پله‌ای ساده با 10000 بار شبیه‌سازی از توزیع نمایی با پارامترهای $\theta_1 = 3$ و $\theta_2 = 1/5$ با نمونه‌هایی به حجم $n = 20, m = 10$ و $\tau = 0/5$ انجام می‌دهیم. در این حالت تقریباً احتمال مشاهده زمان‌های شکست با فاصله کوچک و متوسط یکسان است، بنابراین انتخاب مقادیر متوسط برای α_1 مناسب می‌باشد. نتایج در جدول‌های (۵) و (۶) ارائه شده است که انتخاب مقادیر متوسط برای α_1 و اثربخشی پارامتر تنظیم α_2 را با توجه به این که اندازه کارایی نسبی (RE) از مقدار یک فاصله می‌گیرد، تأیید می‌کند.

۷ بحث و نتیجه‌گیری

در این مقاله، مکانیسم سانسور تصادفی براساس تابع ربط لوجیت برای سانسور فزاینده نوع دوم در یک آزمون طول عمر شتابیده استرس پله‌ای ساده با تأکید بر توزیع نمایی مورد مطالعه قرار گرفته است. رویکرد جدید شامل پارامترهای تنظیم است که خروج تصادفی را با کمک تابع ربط لوجیت به شرایط آزمایش تطبیق می‌دهد. مطالعات شبیه‌سازی نشان می‌دهد که میانگین زمان آزمایش محاسبه شده در دو روش $Type-II PCR$ و $Type-II PCRD$ متفاوت است. میانگین زمان آزمایش در مدل $Type-II PCRD$ با توجه به پارامترهای تنظیم در نظر گرفته شده از مدل $Type-II PCR$ با توزیع دوجمله‌ای با احتمال ثابت $0/25, 0/5, 0/75$ و همچنین توزیع یکنواخت کمتر است. این امر منجر به کاهش هزینه تمام شده برای آزمایش می‌شود. تعیین و پیشنهاد پارامترهای تنظیم در تابع ربط لوجیت زمانی که احتمال خروج هم به فاصله شکست‌ها و هم مقدار استرس بستگی دارد یک امر چالش برانگیز برای آزمایشگر محسوب می‌شود که موضوع پژوهش آینده نویسندگان این مقاله می‌باشد.

جدول ۵: تحلیل حساسیت مکانیسم سانسور تصادفی در آزمون طول عمر شتابیده استرس پله‌ای ساده تحت $Type - II$ نسبت به پارامترهای تنظیم $PCRD$

$\alpha_T = *$	$\alpha_* = -2$	$\theta_T = 1/5$	$\theta_1 = 3$		
	$R = (R_1, \dots, R_m)$	$RE(\alpha_1)$	$E(T_m)$	$\alpha_1^{(1)}$	$\alpha_1^{(*)}$
	$(*, *, *, *, *, 1, *, 2, *, 8)$	1/0527	1/3726	-20	
	$(*, 1, 1, *, *, *, 2, *, *, 6)$	1/0233	1/4120	-15	
	$(*, *, 1, *, 1, *, *, 1, *, 7)$	1/0000	1/4449	12	-12
	$(1, 1, *, 1, *, *, *, 1, *, 6)$	0/9791	1/4758	-10	
	$(*, *, 3, 2, *, *, *, *, 5)$	0/9525	1/5169	-8	
	$(*, *, 1, *, 1, 2, *, *, 2, 4)$	1/1722	1/5676	-6	
	$(1, 1, 2, 1, 1, *, *, 1, *, 3)$	1/0501	1/7500	-2	
	$(*, 2, 2, *, 1, *, *, 1, *, 4)$	1/0000	1/8376	-1	-1
	$(3, 2, 1, 2, 1, *, *, *, 1)$	0/8525	2/1555	1	
	$(3, *, 1, 2, *, 1, *, *, *, 3)$	0/7550	2/4338	2	
	$(1, 9, *, *, *, *, *, *, *)$	0/4905	3/7465	6	
	$(2, 1, 3, 2, *, 1, 1, *, *, *)$	1/0563	4/0449	8	
	$(4, 4, 1, *, 1, *, *, *, *, *)$	1/0214	4/1831	10	
	$(3, 6, *, *, 1, *, *, *, *, *)$	1/0000	4/2725	12	12
	$(2, 8, *, *, *, *, *, *, *, *)$	0/9786	4/3658	15	
	$(6, 4, *, *, *, *, *, *, *, *)$	0/9741	4/3863	20	

جدول ۶: تحلیل حساسیت مکانیسم سانسور تصادفی در آزمون طول عمر شتابیده استرس پله‌ای ساده تحت $Type - II$ PCRD نسبت به پارامترهای تنظیم

$\alpha_1 = 5$	$\alpha_0 = -3$	$\theta_2 = 1/5$	$\theta_1 = 3$	
$R = (R_1, \dots, R_m)$	$RE(\alpha_1)$	$E(T_m)$	$\alpha_2^{(1)}$	$\alpha_2^{(0)}$
(۰, ۲, ۱, ۰, ۰, ۰, ۱, ۰, ۰, ۶)	۱/۳۱۰۴	۱/۶۳۹۰	-۵	
(۱, ۲, ۰, ۱, ۲, ۰, ۰, ۴, ۰, ۰)	۱/۲۶۲۲	۱/۷۰۱۶	-۴	
(۰, ۱, ۱, ۱, ۰, ۲, ۰, ۰, ۰, ۵)	۱/۲۰۹۰	۱/۷۷۶۵	-۳	
(۱, ۱, ۳, ۰, ۱, ۱, ۰, ۰, ۰, ۳)	۱/۱۴۹۰	۱/۸۶۹۲	-۲	
(۲, ۱, ۱, ۱, ۰, ۱, ۰, ۰, ۰, ۴)	۱/۱۱۰۰	۱/۹۳۵۰	-۱	
(۳, ۲, ۰, ۱, ۱, ۰, ۰, ۱, ۱, ۱)	۱/۰۰۰۰	۲/۱۴۷۸	۱	۱
(۱, ۱, ۰, ۱, ۰, ۰, ۳, ۱, ۱, ۲)	۰/۹۳۳۷	۲/۳۰۰۴	۲	
(۲, ۰, ۳, ۰, ۰, ۰, ۰, ۱, ۱, ۳)	۰/۸۸۱۸	۲/۴۳۵۶	۳	
(۱, ۰, ۱, ۱, ۳, ۱, ۳, ۰, ۰, ۰)	۰/۸۲۹۸	۲/۵۸۸۲۵	۴	
(۲, ۰, ۱, ۲, ۰, ۰, ۰, ۰, ۳, ۲۲)	۰/۷۹۱۱	۲/۷۱۴۹	۵	

مراجع

- [۱] ف. حسن تبار، سانسور فزاینده نوع دوم با خروج تصادفی وابسته به فواصل شکست، رساله دکتری (۱۴۰۱)، دانشگاه تهران، دانشکده ریاضی، آمار و علوم کامپیوتر.
- [2] Darzi, F. H., Mahabadi, S. E., Haghighi, F. (2022), *Type-II progressive censoring with GLM-based random removal mechanism dependent on the experimental conditions*. Journal of Applied Statistics, 50(16), 3199–3228. <https://doi.org/10.1080/02664763.2022.2104230>
- [3] Nelson, W. B. (1980), *Accelerated life testing: step-stress models and data analysis*. IEEE Transactions on Reliability, 29, pp. 103–108.
- [4] Balakrishnan, N. Xie, Q. (2007), *Exact inference for a simple step-stress model with Type-I hybrid censored data from the exponential distribution*, vol. 137, pp. 3268–3290.
- [5] Balakrishnan, N. Xie, Q. (2007), *Exact inference for a simple step-stress model with Type-II hybrid censored data from the exponential distribution*, vol. 137, pp. 2543–2563.

- [6] Xie, Q., Balakrishnan, N., Han, Dh. (2008), *Exact Inference and Optimal Censoring Scheme for a Simple Step-Stress Model Under Progressive Type-II Censoring*, Advances in Mathematical and Statistical Modeling. Statistics for Industry and Technology. Birkhäuser Boston.
https://doi.org/10.1007/978-0-8176-4626-4_9

- [7] Balakrishnan, N., Sandhu, R. A. (1995), *A Simple Simulation Algorithm for Generating Progressive Type-II Censored Samples*, The American Statistician, 49(2), 229–230.
<https://doi.org/10.1080/00031305.1995.10476150>

- [8] E. Cramer and G. Iliopoulos, (2010), *Adaptive progressive Type-II censoring*, TEST, 19, no. 2, pp. 342-358.



یازدهمین سمینار تخصصی
نظریه قابلیت اعتماد و کاربردهای آن
۲۴ و ۲۵ اردیبهشت ۱۴۰۴



اثبات حفظ ویژگی IFR در آماره‌های ترتیبی گسسته به روشی جدید

علی محمدی، م^۱ قره‌باغی، ر^۲

گروه آمار، دانشکده علوم ریاضی، دانشگاه الزهراء، تهران، ایران

چکیده

حدود ۶۰ سال قبل اثبات شد که اگر متغیر تصادفی پیوسته X دارای نرخ خطر صعودی (IFR) باشد آنگاه آماره‌های ترتیبی آن نیز IFR خواهند بود و این مسئله برای حالت گسسته بدون اثبات باقی مانده بود تا اینکه اخیراً یک روش اثبات با استفاده از یک نامساوی انتگرالی ارائه شد. در این مقاله، یک روش کاملاً متفاوت برای حل این مسئله ارائه می‌دهیم.

کلمات کلیدی: آماره‌های ترتیبی، تابع نرخ شکست، تحذب

۱ پیش‌گفتار

آماره‌های ترتیبی نقش بسیار مهمی در شاخه‌های مختلف آمار و احتمال بازی می‌کنند. همچنین بررسی رفتار شاخص‌هایی نظیر تابع نرخ شکست (خطر) در مباحثی مانند تحلیل بقا در آمار زیستی و نظریه قابلیت اعتماد در آمار مهندسی

^۱m.alimohammadi@alzahra.ac.ir

^۲rezvangharebaghi6773@gmail.com

از اهمیت ویژه‌ای برخوردار است. بررسی این شاخص و حفظ آن در آماره‌های مختلف چندین دهه است که مورد توجه محققان قرار دارد.

در مبحث قابلیت اعتماد به سیستمی که کارکردن آن مستلزم کار کردن حداقل k جز از n جز آن باشد، سیستم k از n می‌گویند. آماره ترتیبی $(n - k + 1)$ -ام طول عمر چنین سیستمی را نشان می‌دهد. بررسی حفظ ویژگی نرخ خطر صعودی (IFR) در این سیستم‌ها از اهمیت ویژه‌ای برخوردار است. در حالت پیوسته، [۷] اثبات کردند که اگر متغیر تصادفی X دارای ویژگی IFR باشد آنگاه آماره‌های ترتیبی آن نیز IFR خواهند بود. روش اثبات در این حالت سراسر است می‌باشد زیرا در حالت پیوسته تابع چگالی آماره‌های ترتیبی فرم بسته‌ای دارد و مسئله مهم در اینجا همین است که تابع جرم احتمال آماره‌های ترتیبی در حالت گسسته فرم بسته‌ای ندارد. از طرفی ویژگی IFR در سیستم‌های منسجم که حالت کلی تر از سیستم‌های k از n هستند (برای تعریف و ویژگی‌های سیستم‌های منسجم می‌توان به [۴] مراجعه کرد) حفظ نمی‌شود یعنی اگر اجزا IFR باشند آنگاه سیستم منسجم IFR نخواهد بود. بنابراین ویژگی دیگری به نام نرخ خطر صعودی در میانگین (IFRA) معرفی شد و [۴] نشان دادند که ویژگی IFRA در سیستم‌های منسجم زمانی که طول عمر اجزا متغیرهای تصادفی پیوسته باشند حفظ می‌شود. بعد از آن، [۹] اثبات کردند که زمانی که طول عمر اجزا گسسته باشد آنگاه ویژگی IFRA در سیستم‌های منسجم حفظ می‌شود. در نتیجه مشاهده می‌شود که حفظ ویژگی IFR در سیستم‌های k از n برای حالت گسسته همچنان بدون اثبات باقی مانده بود.

فرض کنید $X_1, X_2, \dots, X_m$ یک نمونه تصادفی به حجم m از متغیر تصادفی نامنفی X باشد. تابع بقای آماره ترتیبی k -ام $X_{k:m}$ به صورت زیر است D

$$\begin{aligned}\bar{F}^{k:m}(x) &= \sum_{j=0}^{k-1} \binom{m}{j} F^j(x) \bar{F}^{n-j}(x) \\ &= k \binom{n}{k} \int_0^{\bar{F}(x)} y^{m-k} (1-y)^{k-1} dy.\end{aligned}\quad (1)$$

لازم به ذکر است که رابطه (۱) ارتباطی با گسسته یا پیوسته بودن توزیع X نداشته و برای هر دو حالت استفاده می‌شود. برای ادامه کار به چند تعریف نیاز داریم (می‌توان به [۸] مراجعه کرد). گوییم تابع پیوسته f محدب است اگر

$$f(\alpha x + (1 - \alpha)y) \leq \alpha f(x) + (1 - \alpha)f(y), \quad \forall 0 < \alpha < 1, \quad (2)$$

و لگ- مقعر است اگر

$$f(\alpha x + (1 - \alpha)y) \geq f(x)^\alpha f(y)^{1-\alpha}, \quad \forall 0 < \alpha < 1. \quad (3)$$

همچنین گوییم دنباله a_n محدب است اگر

$$2a_n \leq a_{n-1} + a_{n+1}, \quad (4)$$

و لگ- مقعر است اگر

$$a_n^2 \geq a_{n-1} a_{n+1}. \quad (5)$$

از این به بعد برای تفاوت گذاشتن بین تابع توزیع پیوسته و گسسته به ترتیب از نماد F و F_n استفاده می کنیم. یک متغیر تصادفی پیوسته IFR است اگر $\bar{F}$ تابعی لگ-مقعر باشد و یک متغیر تصادفی گسسته IFR است اگر $\bar{F}_n$ دنباله ای لگ-مقعر باشد [۴].

[۳] نشان دادند که اگر متغیر تصادفی گسسته X ، IFR باشد آنگاه به توجه به روابط (۱) و (۵)، رابطه زیر برقرار است:

$$\begin{aligned} (\bar{F}_n^{k:m})^2 &= \left(k \binom{n}{k} \int_0^{\bar{F}_n} y^{m-k} (1-y)^{k-1} dy \right)^2 \\ &\geq \left(k \binom{n}{k} \int_0^{\bar{F}_{n-1}} y^{m-k} (1-y)^{k-1} dy \right) \left(k \binom{n}{k} \int_0^{\bar{F}_{n+1}} y^{m-k} (1-y)^{k-1} dy \right) \\ &= \bar{F}_{n-1}^{k:m} \bar{F}_{n+1}^{k:m}. \end{aligned}$$

در ادامه بر اساس روند این اثبات، [۲] حدس زدند که رابطه

$$\left(\int_a^b f(y) dy \right)^2 \geq \left(\int_a^b f(\epsilon y) dy \right) \left(\int_a^b f(\epsilon^{-1} y) dy \right) \quad (6)$$

باید برای هر تابع f که در آن $f(e^x)$ لگ- مقعر است برقرار باشد و ثابت کردند که $f(e^x)$ لگ- مقعر اگر و تنها اگر (۶) به ازای هر $a < b$ و $0 < \epsilon$ برقرار باشد. همچنین آن‌ها ثابت کردند که $f(x)$ لگ- مقعر اگر و تنها اگر

$$\left(\int_a^b f(y) dy \right)^2 \geq \left(\int_a^b f(y - \epsilon) dy \right) \left(\int_a^b f(y + \epsilon) dy \right)$$

به ازای هر $a < b$ و $0 < \epsilon$ برقرار باشد.

در این مقاله یک روش اثبات کاملاً متفاوت بر اساس رابطه (۴) و گسسته سازی ابتکاری روش به کار رفته در [۴] ارائه می دهیم که روشی ساده تر نسبت به آنچه که [۳] و [۲] انجام داده اند، می باشد.

۲ دست‌آوردهای پژوهش

در این بخش فرض می‌کنیم F و G نشان دهنده توابع توزیع متغیرهای تصادفی پیوسته و E نشان دهنده توزیع نمایی استاندارد (پارامتر توزیع نمایی تاثیری در روند کار ندارد) باشد. ترکیب دو تابع را نیز برای سادگی با نماد $FG (= F \circ G)$ نشان می‌دهیم. ابتدا تعریف زیر را از [۴] بیان می‌کنیم.

تعریف ۱.۰۲. گوییم توزیع F در ترتیب‌بندی محدب از G کوچکتر است و با نماد $F \leq_c G$ نشان می‌دهیم هرگاه $G^{-1}F$ تابعی محدب باشد.

اکنون برای شروع کار بر اساس این تعریف و رابطه (۴)، تعریف زیر را ارائه می‌دهیم.

تعریف ۲.۰۲. گوییم توزیع F_n در ترتیب‌بندی محدب از G کوچکتر است و با نماد $F_n \leq_c G$ نشان می‌دهیم هرگاه $G^{-1}F_n$ دنباله‌ای محدب باشد.

اثبات ۳ لم زیر مشابه [۴] است و برای تکمیل کار در این جا آورده می‌شود.

لم ۳.۰۲. اگر f تابعی محدب و غیر نزولی و a_n دنباله‌ای محدب باشد آنگاه $f a_n$ دنباله‌ای محدب است.

برهان. از محدب بودن a_n داریم $a_n \leq \frac{1}{2}(a_{n-1} + a_{n+1})$ و چون f غیر نزولی است داریم $f(a_n) \leq f(\frac{a_{n-1}}{2} + \frac{a_{n+1}}{2})$. حال چون f محدب است طبق رابطه (۲) به ازای $\alpha = \frac{1}{2}$ داریم $f(a_n) \leq \frac{1}{2}f(a_{n-1}) + \frac{1}{2}f(a_{n+1})$. در نتیجه

$$2fa_n \leq fa_{n-1} + fa_{n+1}.$$

□

لم ۴.۰۲. اگر $F_n \leq_c F$ و $F \leq_c G$ ، آنگاه $F_n \leq_c G$.

برهان. اگر در لم ۳.۰۲ به جای f و a_n به ترتیب قرار دهیم $G^{-1}F$ و $F_n^{-1}F_n$ داریم

$$G^{-1}FF_n^{-1}F_n = G^{-1}F_n,$$

□

که دنباله‌ای محدب است

لم ۵.۰۲. اگر $F_n \leq_c G$ ، آنگاه $F_n^{k:m} \leq_c G^{k:m}$.

برهان. بر اساس رابطه (۱) می‌توان نوشت

$$F_n^{k:m} = B^{k:m} F_n,$$

که در آن

$$B^{k:m}(x) = k \binom{n}{k} \int_0^x y^{k-1} (1-y)^{m-k} dy.$$

در نتیجه

$$G^{k:m-1} F_n^{k:m} = (B^{k:m} G)^{-1} B^{k:m} F_n = G^{-1} B^{k:m-1} B^{k:m} F_n = G^{-1} F_n.$$

□

در ادامه به قضیه کلیدی زیر نیاز داریم.

قضیه ۶.۲. *IFR بودن متغیر تصادفی گسسته X با $F_n \leq_c E$ معادل است.*

برهان. چون $E^{-1}(x) = -\ln(1-x)$ ، طبق تعریف ۲.۲ و رابطه (۴) داریم

$$-2 \ln(1 - F_n) \leq -\ln(1 - F_{n-1}) - \ln(1 - F_{n+1})$$

$$\Leftrightarrow \ln(\bar{F}_n)^2 \geq \ln(\bar{F}_{n-1}) + \ln(\bar{F}_{n+1})$$

$$\Leftrightarrow \ln(\bar{F}_n)^2 \geq \ln(\bar{F}_{n-1} \bar{F}_{n+1})$$

$$\Leftrightarrow (\bar{F}_n)^2 \geq \bar{F}_{n-1} \bar{F}_{n+1},$$

□

و اثبات تمام است.

اکنون آماده‌ایم تا قضیه اصلی را اثبات کنیم.

قضیه ۷.۲. *اگر متغیر تصادفی گسسته X ، IFR باشد آنگاه آماره ترتیبی k -ام $X_{k:m}$ به ازای همه $1 \leq k \leq m$ ، IFR است.*

برهان. طبق قضیه ۶.۲ کافی است نشان دهیم که $F_n^{k:m} \leq_c E$. طبق قضیه ۶.۲، IFR بودن X معادل است با $F_n \leq_c E$ و در نتیجه بر اساس لم ۵.۲ داریم $F_n^{k:m} \leq_c E^{k:m}$. از طرفی تابع چگالی آماره‌های ترتیبی از توزیع نمایی، لگ-مقعر هستند و چون لگ-مقعر بودن، IFR بودن را نتیجه می‌دهد (می‌توان به [۱] مراجعه کرد) داریم $E^{k:m} \leq_c E$. در نهایت طبق لم ۴.۲ نتیجه مورد نظر به دست می‌آید و اثبات تمام است.

□

۳ یک کاربرد از قضیه

طول عمر یک سیستم k از m توسط آماره ترتیبی $-(m - k + 1)$ ام تعیین می شود [۴]. تابع نرخ خطر در حالت گسسته عبارت است از

$$h_n = \frac{\bar{F}_{n-1} - \bar{F}_n}{\bar{F}_{n-1}} = 1 - \frac{\bar{F}_n}{\bar{F}_{n-1}}, \quad n \in \mathbb{N}.$$

بنابراین نرخ خطر سیستم k از m عبارت است از

$$h_n^{m-k+1:m} = 1 - \frac{\bar{F}_n^{m-k+1:m}}{\bar{F}_{n-1}^{m-k+1:m}}, \quad n \in \mathbb{N}. \quad (1)$$

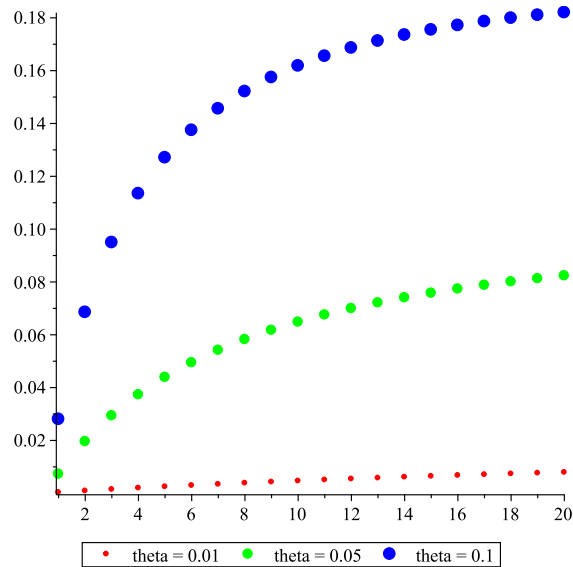
ابتدا سیستمی را مثال می زنیم که در آن طول عمر اجزا، تعداد دفعات روشن و خاموش شدن تا زمان خرابی است [۶]. حال یک سیستم ۲ از ۳ که طول عمر اجزای آن دارای توزیع هندسی با پارامتر θ است را در نظر بگیرید. در اینجا متغیر تصادفی هندسی تعداد دفعات انجام صحیح وظیفه محوله تا زمان خرابی است (به عبارت دیگر موفقیت در اینجا خراب شدن جز است). طبق رابطه (۱) و (۱) داریم:

$$\begin{aligned} h_n^{2:3} &= 1 - \frac{\bar{F}_n^{2:3}}{\bar{F}_{n-1}^{2:3}} \\ &= 1 - \frac{(1-\theta)^{3n} + 3(1-(1-\theta)^n)(1-\theta)^{2n}}{(1-\theta)^{3(n-1)} + 3(1-(1-\theta)^{n-1})(1-\theta)^{2(n-1)}}. \end{aligned}$$

در شکل ۱ تابع نرخ خطر سیستم به ازای θ های مختلف رسم شده است که همگی نشان دهنده IFR بودن آن است.

لازم به ذکر است که توزیع هندسی هم IFR و هم DFR است (چون نرخ خطر آن ثابت است). اما همانطور که در شکل دیده می شود، نرخ خطر سیستم IFR است.

به عنوان مثالی دیگر، در مبحث قابلیت اعتماد سلاح ها، تعداد گلوله های شلیک شده تا خرابی، مهم تر از عمر سلاح است [۱۰]. در نتیجه در اینجا نیز با یک متغیر تصادفی طول عمر گسسته سر و کار داریم. فرض کنید یک اسلحه سبک (در حالت ساده) از دو جز ماشه و گلنگدن تشکیل شده باشد که طول عمر هر کدام تعداد دفعاتی است که وظیفه خود را به درستی انجام می دهند. از آنجایی که عمل شلیک گلوله زمانی با موفقیت انجام می شود که هر دو جز به درستی کار کنند، با یک سیستم سری و یا به عبارتی با یک سیستم ۲ از ۲ سر و کار داریم. طبق رابطه (۱) و (۱)



شکل ۱: تابع نرخ خطر سیستم ۲ از ۳ با اجزای هندسی به ازای $\theta = 0.01, 0.05, 0.1$

داریم:

$$\begin{aligned}
 h_n^{1:2} &= 1 - \frac{\bar{F}_n^{1:2}}{\bar{F}_{n-1}^{1:2}} \\
 &= 1 - \frac{(1-\theta)^{2n}}{(1-\theta)^{2(n-1)}} = 1 - (1-\theta)^2,
 \end{aligned}$$

که مقداری ثابت (و در نتیجه IFR) است. البته این مسئله اتفاقی نبوده و به این خاطر است که طبق یک مشخصه سازی، متغیر تصادفی گسسته X دارای توزیع هندسی است اگر و تنها اگر $X_{1:m}$ دارای توزیع هندسی باشد [۵].

بحث و نتیجه گیری

یکی از مسائل مهم در تحلیل بقا و قابلیت اعتماد بسته بودن ویژگی‌های سالخوردگی در آماره‌های گوناگون از جمله آماره‌های ترتیبی است. بسته بودن ویژگی IFR در آماره‌های ترتیبی از مسائل جالبی است که فاصله اثبات آن برای حالت پیوسته و گسسته به بیش از شش دهه می‌رسد. در این مقاله یک اثبات متفاوت از آنچه که اولین بار [۳] مطرح کرده بودند با استفاده از تعریف ترتیب‌بندی محدب، اثبات چندین لم در این زمینه و تعریف IFR بودن متغیر تصادفی گسسته بر اساس ترتیب‌بندی محدب با متغیر تصادفی نمایی ارائه دادیم.

- [1] Alimohammadi, M., Alamatsaz, M.H. and Cramer, E. (2016). Convolutions and generalization of logconcavity: Implications and applications. *Naval Research Logistics (NRL)*, **63**(2), 109-123.
- [2] Alimohammadi, M., Balakrishnan, N. and Simon, T. (2024). An inequality for log-concave functions and its use in the study of failure rates. *Probability in the Engineering and Informational Sciences*, **38**(4), 695–704.
- [3] Alimohammadi, M. and Navarro, J. (2024). Resolving an old problem on the preservation of the IFR property under the formation of systems with discrete distributions. *Journal of Applied Probability*, **61**(2), 644-653.
- [4] Barlow, R.E. and Proschan, F. (1975). *Statistical theory of reliability and life testing*. New York: Holt, Rinehart and Winston.
- [5] David, H.A. and Nagaraja, H.N. (2004). *Order statistics*. John Wiley Sons.
- [6] Dembińska, A. (2018). On reliability analysis of k-out-of-n systems consisting of heterogeneous components with discrete lifetimes. *IEEE Transactions on Reliability*, **67**(3), 1071-1083.
- [7] Esary, J.D. and Proschan, F. (1963). Relationship between system failure rate and component failure rates. *Technometrics*, **5**(2), 183-189.
- [8] Rockafellar, R. (1970). *Convex Analysis*. Princeton university press. New Jersey.
- [9] Roy, D. and Gupta, R.P. (1992). Classifications of discrete lives. *Microelectronics Reliability*, **32**(10), 1459-1473.
- [10] Shaked, M., Shanthikumar, J.G., and Valdez-Torres, J.B. (1995). Discrete hazard rate functions. *Computers & Operations Research*, **22**(4), 391-402.



یازدهمین سمینار تخصصی
نظریه قابلیت اعتماد و کاربردهای آن
۲۴ و ۲۵ اردیبهشت ۱۴۰۴



تحلیل خطرهای رقابتی در مسائل زیستی

هاشم آبادی، ع^۱ رزمخواه، م^۲

گروه آمار، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد

چکیده

در این مقاله، به بررسی جامع و کاربردی روش‌های تحلیل بقا در حضور خطرهای رقابتی پرداخته شده است. هدف اصلی، ارزیابی اثر عوامل مختلف (به‌ویژه نوع درمان) بر رویدادهای مرتبط با مرگ ناشی از سرطان و بر اثر دلایل خارجی، با استفاده از مدل‌های غیرپارامتری (مانند برآوردگر کاپلان-مایر) و مدل‌های رگرسیونی نسبتی خطر (کاکس) می‌باشد. به‌علاوه، با به‌کارگیری مدل‌های تحلیل خطرهای رقابتی، سعی شده است تا مشکلات ناشی از فرض استقلال رویدادها رفع شده و تصویر دقیق‌تری از فرآیند شکست بیماران ارائه شود. برای بررسی کاربردی این مدل‌ها، تحلیل‌های آماری بر روی داده‌های واقعی مربوطه به سرطان مثانه انجام شده است. نتایج به‌دست آمده نشان می‌دهد که شیوه‌های مختلف درمان به طور معناداری بر خطر مرگ ناشی از سرطان تأثیرگذار بوده و مدل‌های رقابتی قادر به ارائه بینشی دقیق‌تر از روند بقا در مواجهه با خطرهای متعدد می‌باشند. یافته‌های این پژوهش، علاوه بر ارتقای درک نظری از مسائل مرتبط با تحلیل بقا و اعتمادپذیری، می‌تواند مبنایی جهت بهبود تصمیم‌گیری‌های بالینی و طراحی مطالعات آتی فراهم آورد.

کلمات کلیدی: خطرهای رقابتی، تحلیل بقا، قابلیت اعتماد، پیش‌بینی

^۱emad.h.1381@gmail.com

^۲razmkhah_m@um.ac.ir

۱ مقدمات

در مطالعات اولیه در مورد این موضوع، پرنیتیس و همکاران [۲] با ارائه روش‌های نوین در تحلیل زمان‌های شکست در حضور خطرهای رقابتی، پایه‌ای برای توسعه مدل‌های رگرسیونی نسبتی خطر فراهم کردند. این رویکرد به ویژه در حوزه‌های پزشکی و بایواستاتستیک کاربرد فراوانی یافته است. از سوی دیگر، گولی و همکاران [۳] با معرفی نمایی جدید از برآوردهای احتمال شکست، سعی در رفع برخی از چالش‌های موجود در ارزیابی خطرهای رقابتی داشتند. در پژوهش حاضر، با استناد به این دستاوردهای علمی، تحلیل‌های انجام شده بر روی مجموعه داده bladder1 نشانگر تأثیر معنادار نوع درمان بر بقا بیماران بوده و مسیر پیشرفت روش‌های آماری در این حوزه را بهبود بخشیده است.

داده‌های بقای استاندارد، طول بازه زمانی را از زمان مبدا تا وقوع رویداد مورد نظر اندازه‌گیری می‌کنند. در فرایند بهبود یا بیماری، در بیشتر مواقع بیش از یک نوع رویداد نقش دارد که معمولاً یک نوع آن را می‌توان به عنوان رویداد اصلی مورد توجه در نظر گرفت. گاهی اوقات ممکن است رویدادهای دیگر از رخ دادن رویداد مورد علاقه جلوگیری کنند، یا زودتر از آن رخ دهند. در احتمال وقوع رویداد مورد توجه در حضور این خطرهای رقابتی احتیاط لازم است. فرض کنیم که افراد هنوز زنده هستند، اما زمان پیگیری به پایان رسیده است، برآوردگر کاپلان-مایر، که یک روش رایج برای تخمین منحنی بقا است، ممکن است دچار سوگیری شود. دلیل این سوگیری فرض کلیدی برآوردگر کاپلان-مایر می‌باشد: زمان وقوع رویداد و توزیع سانسور باید مستقل از یکدیگر باشند. به عبارت دیگر، سانسور شدن یک فرد نباید با احتمال وقوع رویداد مورد نظر مرتبط باشد. اگر علل رقیب وجود داشته باشند، این فرض ممکن است نقض شود. به عنوان مثال، اگر افراد مسن بیشتر در معرض خطر مرگ ناشی از بیماری قلبی باشند و همچنین بیشتر در معرض خطر مرگ ناشی از حادثه، فرض استقلال نقض می‌شود.

در این آزمایشات بقا نقش اصلی را زمان دارد که ویژگی‌های خاصی را به داده‌ها می‌دهد. یک رویداد ممکن است بین دو زمان مشاهده متوالی در بازه مورد انتظار رخ دهد، که منجر به داده‌های سانسور شده فاصله‌ای^۱ می‌شوند. زمانی که ما می‌دانیم یک فرد یا شیء تا یک زمان مشخص در مطالعه حضور داشته و رویداد مورد نظر برای او رخ نداده است،

^۱Interval Censored

اما نمی‌دانیم بعد از آن زمان چه اتفاقی افتاده است، گوییم سانسور از راست^۲ رخ داده است. به عنوان مثال فرض کنید مطالعه‌ای برای بررسی اثر یک دارو بر بقای بیماران مبتلا به سرطان انجام می‌شود. اگر بیماری بعد از سه سال هنوز زنده باشد و مطالعه در آن زمان به پایان برسد، داده‌های مربوط به این بیمار به صورت سانسور از راست در نظر گرفته می‌شود. ما می‌دانیم که بیمار حداقل تا سه سال زنده بوده، اما نمی‌دانیم بعد از آن چه مدت زنده می‌ماند.

سانسور از چپ^۳ زمانی اتفاق می‌افتد که ما می‌دانیم یک فرد یا شیء قبل از یک زمان مشخص رویداد مورد نظر را تجربه کرده است، اما زمان دقیق وقوع رویداد را نمی‌دانیم. به عبارت دیگر، ما می‌دانیم که رویداد قبل از آن زمان رخ داده است. به عنوان مثال فرض کنید مطالعه‌ای برای بررسی زمان شروع اعتیاد به مواد مخدر در افراد انجام می‌شود. اگر فردی به یاد نیاورد که دقیقاً چه زمانی شروع به مصرف مواد کرده است، اما می‌داند که قبل از ۱۸ سالگی این کار را انجام می‌داده، داده‌های مربوط به این فرد به صورت سانسور از چپ در نظر گرفته می‌شود. ما می‌دانیم که فرد قبل از ۱۸ سالگی معتاد شده، اما زمان دقیق آن را نمی‌دانیم.

۲ توصیف مدل

در این تحقیق فرض می‌شود که تمام توزیع‌های زمان شکست، پیوسته هستند. در مدل برای داده‌های از راست سانسور شده هر فرد i دارای یک زمان رویداد t_i و یک زمان سانسور c_i است. مشاهدات به صورت $x_i = \min(t_i, c_i)$ هستند. یعنی از بین زمان سانسور و رویداد هر کدام زودتر رخ دهد آن را به عنوان مشاهده در نظر می‌گیریم. داده‌های بقا یک بیمار با سه متغیر بیان می‌شود: زمانی که فرد وارد خطر می‌شود، زمانی که فرد سانسور می‌شود یا رویداد رخ می‌دهد و متغیر دیگر، متغیری است که نشان‌دهنده سانسور شده یا رخ دادن رویداد است. زمان رخداد رویداد و زمان سانسور برای هر فرد به صورت یک نمونه تصادفی $(t_1, c_1), \dots, (t_n, c_n)$ که t دارای تابع بقا S است، به گونه‌ای که $S(t) = \text{Prob}(T > t)$ است و $c_i \sim G$.

با توجه به مدل‌های استاندارد برای داده‌های از راست سانسور شده توزیع زمان رخداد و زمان سانسور باید از یکدیگر مستقل باشند. همچنین با توجه فرض استقلال، نرخ خطر افرادی که سانسور شده‌اند، برابر است با نرخ خطر

Right Censored^۲

Left Censored^۳

افرادی که در آزمایش باقی مانده‌اند، بنابراین تابع نرخ خطر برای توزیع‌های پیوسته به صورت زیر تعریف می‌شوند:

$$\lambda(t) = \lim_{\Delta t \rightarrow 0} \frac{Prob(t \leq T \leq t + \Delta t | T \geq t)}{\Delta t} \quad (1)$$

برآوردگر کاپلان-مایر زمانی استفاده می‌شود که داده‌ها سانسور از راست شده باشند. اگر حجم نمونه افزایش یابد، تعداد زمان‌های رخداد نیز افزایش می‌یابد و برآوردگر کاپلان-مایر به یک توزیع پیوسته نزدیک می‌شود. تاثیر متغیرهای کمکی بر پیشرفت بیماری اغلب با استفاده از مدل خطرات متناسب کاکس مدل سازی می‌شود. در ساده‌ترین شکل آن، تابع نرخ خطر برای یک موضوع با متغیرهای کمکی $Z = (Z_1, \dots, Z_p)$ بدین گونه در نظر گرفته می‌شود:

$$\lambda(t|Z) = \lambda_0(t) \exp(\beta^T Z) \quad (2)$$

که در آن β برداری از ضرایب رگرسیونی و $\lambda_0(t)$ نرخ خطر پایه است. اگر فرض کنیم تمام زمان‌های رخدادها متمایز هستند، برآورد بردار β با استفاده از ماکسیم سازی تابع درستنمایی زیر بدست می‌آید:

$$L(\beta) = \prod_{j=1}^N \frac{\exp(\beta^T Z_j)}{\exp(\sum_{l \in R_j} \beta^T Z_l)} \quad (3)$$

که در آن R_j مجموعه خطر لحظه‌ای می‌باشد. در واقع شامل تمام افرادی است که درست قبل از زمان رخداد j هنوز در حال پیگیری بوده و رویداد مورد توجه را تجربه نکرده‌اند. در ادامه، برای کنترل اثر متغیرهای طبقه‌ای بر نرخ خطر پایه، می‌توان از مدل کاکس طبقه‌بندی شده استفاده کرد. این مدل به نرخ خطر پایه امکان می‌دهد در سطوح مختلف متغیرهای طبقه‌ای، تغییر کند. در این مدل، نرخ خطر پایه می‌تواند در زیرگروه‌های مختلف $(h = 1, \dots, m)$ ، متفاوت باشد، در حالی که اثر سایر متغیرهای مستقل به صورت مشترک در تمام طبقات برآورد می‌شود.

$$\lambda_h(t|Z) = \lambda_{h,0}(t) \exp(\beta^T Z) \quad (4)$$

همچنین پارامترهای مدل با حداکثرسازی تابع درستنمایی جزئی برای هر طبقه برآورد می‌شوند:

$$\begin{aligned} L(\beta) &= \prod_{h=1}^m L_h(\beta) \\ &= \prod_{h=1}^m \prod_{j=1}^N \frac{\exp(\beta^T Z_j)}{\exp(\sum_{l \in R_{hj}} \beta^T Z_l)} \end{aligned} \quad (5)$$

که در آن R_{hj} شامل تمام افرادی از طبقه h است که درست قبل از زمان t_j هنوز "در معرض خطر" رخداد رویداد هستند. این افراد می‌توانند در زمان t_j یا زمان‌های بعد رویداد را تجربه کنند.

در مدل کاکس طبقه‌بندی شده، داده‌ها بر اساس ویژگی‌هایی به چند لایه تقسیم می‌شوند. تابع احتمال $L_h(\beta)$ برای هر لایه جداگانه محاسبه می‌شود و در صورتی که پارامترهای β در هر لایه متفاوت باشند، این مدل معادل با برازش یک مدل کاکس جداگانه برای هر لایه خواهد بود. نتایج حاصل از یک مدل کاکس، که اثر متغیرهای مستقل را بر نرخ خطر مدل می‌کند، می‌تواند برای توصیف اثرات تجمعی نیز استفاده شود. به عبارت دیگر، مدل کاکس نه تنها اثر لحظه‌ای متغیرها بر خطر را نشان می‌دهد، بلکه می‌تواند اثرات انباشته شده در طول زمان را نیز تخمین بزند. اگر یک بیمار دارای مقادیر متغیرهای مستقل Z باشد، به صورت منحنی بقای زیر تخمین زده می‌شود.

$$\hat{S}(t) = \exp\{-\hat{\Lambda}_\bullet(t)e^{\hat{\beta}^T Z}\} = \hat{S}_\bullet(t)^{\exp(\hat{\beta}^T Z)} \quad (6)$$

که با $\hat{S}_\bullet(t) = \exp(\hat{\beta}^T Z)$ منحنی بقای پایه برآورد می‌شود.

۳ فرض استقلال

اکثر اوقات استقلال بین رخ دادن رویداد و سانسور شدن آن بدون اثبات پذیرفته می‌شود، البته گاهی اوقات خیلی ساده این فرض رد می‌شود. به دلایل زیر یک رویداد از راست سانسور می‌شود:

- تمام شدن مطالعه: چون زمان تقویمی، مشاهده را به رویدادهایی که در گذشته رخ داده‌اند محدود می‌کند، ممکن است زمان وقوع یک رویداد به دلیل اینکه فرد به اندازه کافی تحت نظر نبوده است، «سانسور از راست» باشد. به این نوع سانسور، «سانسور اداری»^۴ نیز گفته می‌شود. اگر دلیل سانسور تمام شدن مطالعه باشد، به طور کلی می‌توانیم فرض کنیم که مکانیسم سانسور مستقل از پیشرفت بیماری است. در دیگر موارد باید محتاط‌تر بود.
- از دست دادن بررسی: دلایلی مانند مهاجرت یا خستگی می‌توانند باعث شوند فردی از مطالعه خارج شود. علاوه بر این، ممکن است رویداد مورد بررسی برای او اتفاق افتاده باشد، اما ما از آن بی‌اطلاع باشیم. زمان

^۴ Administrative Censoring

سانسور به دلیل از دست دادن بررسی با زمان رخداد هنگامی که فرد حس کند نیازی به پیگیری ندارد همبستگی منفی دارد و بنابراین از مطالعه خارج می‌شود. سانسور شدن این افراد هنگامی که از مطالعه خارج می‌شوند، باعث سوگیری رو به پایین منحنی بقا می‌شود.

زمان سانسور با زمان رویدادی که افراد مبتلا به پیشرفت بیماری پیشرفته احتمال بیشتری برای ترک مطالعه دارند، همبستگی مثبت دارد. یک دلیل ممکن است این باشد که آنها برای پیگیری بیشتر، بیش از حد بیمار شده‌اند یا اینکه به کشور محل تولد خود بازگردند تا آخرین دوره را با خانواده خود بگذرانند. در اینجا سانسور کردن این افراد باعث سوگیری رو به بالا منحنی بقا می‌شود.

وقتی فردی از مطالعه خارج می‌شود (سانسور می‌شود)، معمولاً اطلاعات مربوط به او تا زمان خروج ثبت می‌شود. اما گاهی اوقات، پس از خروج فرد، اطلاعات اضافی در مورد او به دست می‌آید. این اطلاعات می‌توانند مربوط به وضعیت سلامتی، رخداد رویداد مورد نظر یا سایر متغیرهای مرتبط با مطالعه باشند. استفاده از این اطلاعات اضافی می‌تواند به کاهش سوگیری ناشی از سانسور شدن کمک کند. به عنوان مثال، اگر فردی به دلیل مهاجرت از مطالعه خارج شود و بعداً مشخص شود که او فوت کرده است، در نظر نگرفتن این اطلاعات می‌تواند منجر به تخمین نادرست احتمال بقا شود.

گاهی اوقات نیز اطلاعات اضافه‌ای پس از ترک مطالعه باقی می‌ماند. به عنوان مثال از طریق ثبت نام، اگر انتخاب اطلاعات پس از ترک مطالعه به روشی مناسب انجام شود، استفاده از این اطلاعات ممکن است سوگیری را کاهش دهد یا به طور کلی از بین ببرد. هوور و همکاران [۴] در مطالعه‌ای، به بررسی دقیق‌تر راهبردهای سانسور داده‌ها با پیگیری افراد پس از ترک مطالعه و همچنین، بررسی اریبی (سوگیری) ناشی از این روش‌ها در تخمین پارامترها پرداختند. اگرچه این مطالعه در زمینه بررسی هم‌گروهی‌های مبتلا به HIV/AIDS انجام شده، اما نتایج آن به صورت کلی‌تر قابل تعمیم و استفاده است.

۴ خطرهای رقابتی

در موقعیت‌هایی که بیش از یک علت می‌تواند منجر به وقوع رویداد (شکست) شود، با پدیده‌ای به نام ”خطرهای رقابتی“^۵ مواجه هستیم. برای مثال، اگر شکست به معنای فوت باشد و دلایل متعددی برای آن وجود داشته باشد، تنها

^۵Competing Risks

اولین دلیل فوت که رخ می‌دهد قابل مشاهده است. در شرایط دیگر، ممکن است رویدادهای بعد از اولین شکست قابل رصد باشند، اما مورد توجه ما نباشند.

در تحلیل بقا، گاه رویدادهای رقیبی رخ می‌دهند که می‌توانند جایگزین یا مانع وقوع رویداد اصلی شوند؛ مانند مرگ بر اثر علل گوناگون در مطالعات سرطان. در چنین شرایطی، استفاده از روش‌های سنتی نظیر برآوردگر کاپلان-مایر، به دلیل فرض نادرست استقلال بین سانسور و رخداد اصلی، ممکن است منجر به برآوردهای جانبدارانه شود. به‌ویژه، افراد مسن یا دارای بیماری‌های زمینه‌ای، علاوه بر مرگ ناشی از بیماری اصلی، در معرض رویدادهای رقیبی متعددی قرار دارند که احتمال تجربه علت اصلی را کاهش می‌دهد.

موضوع خطرهای رقابتی به قرن هجدهم و زمانی باز می‌گردد که برنولی [۵] پیامدهای احتمالی ریشه‌کن کردن آبله بر نرخ مرگ و میر را بررسی کرد. در واقع، مسئله تخمین احتمالات شکست پس از حذف یکی از خطرهای رقابتی، از اهمیت زیادی برخوردار بوده و در دهه ۱۹۷۰ [۶، ۲] موضوع بحث‌های زیادی بوده است.

برخلاف سانسور معمولی که صرفاً به معنای عدم مشاهده رویداد در بازه مطالعه است، در خطرهای رقابتی وقوع یک علت رقیب به منزله حذف قطعی امکان رخداد علت اصلی است. در نتیجه، نمی‌توان این افراد را مانند سانسور معمولی در نظر گرفت.

وقتی خطرهای رقابتی وجود دارند، استفاده از برآوردگر کاپلان-مایر به صورت ساده می‌تواند منجر به تخمین بیش از حد احتمال شکست و در نتیجه، تخمین کمتر از حد احتمال بقا شود. دلیل این امر این است که این برآوردگر به طور ضمنی فرض می‌کند که افراد سانسور شده به دلیل پایان مطالعه یا از دست دادن پیگیری، هنوز در معرض خطر وقوع رویداد مورد توجه هستند، در حالی که افرادی که به دلیل خطرهای رقابتی سانسور شده‌اند، دیگر هرگز رویداد مورد توجه را تجربه نخواهند کرد.

برای برآورد دقیق‌تر احتمال وقوع هر علت، به جای استفاده از روش‌های ساده، باید از توابعی چون «تابع تجمعی بروز علت اختصاصی»^۶ یا مدل‌های خطر علت‌ویژه استفاده کرد. در این رویکردها، هر علت به‌طور جداگانه مدل‌سازی می‌شود و رخداد زودهنگام یک علت، مانع وقوع علت دیگر تلقی می‌گردد. این تحلیل جامع، تصویری کامل‌تر از روند شکست و نقش عوامل مؤثر در بروز رویدادهای مختلف ارائه می‌دهد و در حوزه‌های بالینی و اپیدمیولوژیک به تصمیم‌گیری دقیق‌تر کمک می‌کند.

این موضوع به دهه ۷۰ میلادی باز می‌گردد، زیرا اغلب اوقات یک دلیل بیولوژیکی مشترک وجود دارد که باعث وقوع هر دو رویداد اصلی و رقابتی می‌شود. بنابراین، اگر بخواهیم مکانیسم رویداد رقابتی را تغییر دهیم به عنوان مثال در نظر داشته باشیم جلوی مرگ ناشی از بیماری قلبی را بگیریم، در واقع احتمال رخ دادن رویداد مورد توجه مرگ

ناشی از سرطان را هم تغییر می‌دهیم و این یعنی زمان وقوع این دو رویداد مستقل از هم نیست. در نتیجه، این یک فرضیه کاملاً متفاوت خواهد بود که ما نمی‌توانیم در مورد آن چیزی بگوییم. برای درک بهتر خطای روش کاپلان-مایر ساده، می‌توانید به [۳] مراجعه کنید.

داده‌های مشاهده شده در مدل‌های خطرهای رقابتی با زمان شکست T ، دلیل شکست D و احتمالاً یک بردار کمکی Z (آن را نادیده خواهیم گرفت) نمایش داده می‌شوند. بنابراین استنتاج باید بر اساس توزیع مشترک T و D باشد که احتمالاً Z داده می‌شود. در مدل‌های خطر رقابتی، تمرکز بر علت‌های مختلف وقوع رویداد (شکست) است که تابع نرخ خطر علت K ام، مفهوم کلیدی برای تحلیل این مدل‌ها می‌باشد. این تابع نرخ لحظه‌ای وقوع رویداد به دلیل یک علت خاص را نشان می‌دهد.

$$\lambda_k(t) = \lim_{\Delta t \rightarrow 0} \frac{Prob(t \leq T \leq t + \Delta t, D = k | T \geq t)}{\Delta t} \quad (1)$$

برای هر بیمار، زمان بالقوه شکست برای هر علت شکست در نظر گرفته می‌شود. اما فقط اولین شکست (رویدادی که زودتر رخ می‌دهد) مشاهده می‌شود و بقیه زمان‌های شکست بالقوه پنهان می‌مانند. اگر بیماری به دلیل پایان مطالعه یا از دست دادن پیگیری سانسور شود، نشان داده می‌شود که به دلیل هیچ علت خاصی شکست نخورده است ($D = 0$). برای مدل‌سازی این نوع سانسور، یک توزیع سانسور اضافی معرفی می‌شود. رویکرد مبتنی بر زمان‌های شکست پنهان، بر توزیع احتمال همزمان همه زمان‌های شکست تمرکز دارد، که توسط تابع بقای توأم توصیف می‌شود. در اینجا D یک متغیر اشاره‌گر است که نمایانگر شماره اولین دلیلی است که رویداد موجب شکست شده است. رویکرد زمان شکست‌های پنهان به توزیع زمان بین k رویداد مختلف می‌پردازد.

اگر توابع بقای مشترک مختلفی با وابستگی‌های متفاوت بین زمان‌های شکست (زمان وقوع رویداد) داشته باشیم، همیشه می‌توانیم یک تابع بقای مشترک دیگر پیدا کنیم که در آن زمان‌های شکست مستقل از هم باشند، اما همچنان خطرهای مربوط به هر علت (یا رویداد) یکسان بمانند. اگر بتوان یک ویژگی را به طور دقیق بر اساس خطرهای مربوط به هر علت (رویداد) مشخص کرد، آن ویژگی قابل محاسبه و تخمین است. تابع نرخ خطر تجمعی و تابع بقا برای هر علت را به ترتیب می‌توانیم به صورت

$$\Lambda_k(t) = \int_0^t \lambda_k(s) ds \quad (2)$$

و

$$S_k(t) = \exp(-\Lambda_k(t)) \quad (3)$$

بنویسیم.

فرمول زیر نشان دهنده تابع بقای کلی است که برابر است با نمایی منفی مجموع توابع نرخ خطر تجمعی خاص علت برای همه علت‌ها. این تابع احتمال عدم شکست از هیچ علتی در زمان t را نشان می‌دهد. به عبارت دیگر، $S(t)$ احتمال زنده ماندن فرد تا زمان t بدون تجربه هیچکدام از رویدادهای رقابتی را نشان می‌دهد.

$$S(t) = \exp(-\sum_{k=1}^K \Lambda_k(t)) \quad (4)$$

تابع توزیع تجمعی زمان رخداد شکست به واسطه علت k ام، با $I_k(t)$ نشان داده می‌شود و به صورت زیر تعریف می‌گردد:

$$I_k(t) = \text{Prob}(T \leq t, D = k) \quad (5)$$

از آنجا که خطرهای رقابتی احتمال زنده ماندن را کاهش می‌دهند، تابع بقای کلی همیشه کمتر یا مساوی تابع بقای ویژه علت است $(S(t) \leq S_k(t))$ ، نتیجه می‌گیریم که $I_k(t) \leq 1 - S_k(t)$. این عدم تساوی نشان می‌دهد که تابع توزیع تجمعی همیشه کمتر یا مساوی احتمال شکست برآورد شده با استفاده از کاپلان-مایر است. برابری بین $I_k(t)$ و $1 - S_k(t)$ تنها در صورتی برقرار می‌شود که خطرهای رقابتی وجود نداشته باشند، یعنی اگر $\sum_{j=1, j \neq k}^K \Lambda_j(t) = 0$ باشد (مجموع توابع نرخ خطر تجمعی برای همه دلایل غیر از k برابر با صفر باشد). این نتیجه نشان می‌دهد که استفاده از برآوردگر کاپلان-مایر ساده در حضور خطرهای رقابتی، منجر به سوگیری می‌شود. تابع نرخ خطر خاص علت k ام نیز به صورت زیر است:

$$\lambda_k(t) = \lim_{\Delta t \rightarrow 0} \frac{P(t \leq T \leq t + \Delta t, D = k | T \geq t)}{\Delta t} \quad (6)$$

در ادامه با توجه به (5) و (6)، تابع توزیع تجمعی زمان رخداد شکست به واسطه علت k ام بر حسب تابع بقا و تابع نرخ خطر بدین صورت است:

$$I_k(t) = \int_0^t \lambda_k(s) S(s) ds \quad (7)$$

اکنون به تخمین توابع رویداد تجمعی می پردازیم. فرض کنید $(0 < t_1 < t_2 < \dots < t_N)$ زمان های گسسته ای هستند که در آنها شکست (از هر علتی) رخ می دهد. d_{ki} نمایانگر تعداد بیمارانی است که در زمان t_i به دلیل علت k شکست می خورند که مقدار صفر یا یک می گیرد. همچنین $d_j = \sum_{k=1}^K d_{kj}$ نشان دهنده تعداد کل شکست ها به هر علتی در زمان t_i است. n_i را نیز تعداد بیماران در معرض خطر در زمان t_i در نظر گرفتیم. در واقع آنهایی که هنوز تحت نظر هستند و به هر دلیلی شکست نخورده اند.

اگر توزیع داده ها گسسته باشد، تابع نرخ خطر علت k ام بدین صورت خواهد بود:

$$\lambda_k(t) = \text{Prob}(T = t_i, D = k | T > t_{i-1}) \quad (8)$$

که بر اساس نسبت افراد در معرض خطر که از علت k ام شکست خورده اند، به صورت زیر برآورد می شود:

$$\widehat{\lambda}_k(t_j) = \frac{d_{ki}}{n_i}. \quad (9)$$

برآورد تابع بقا کل را به صورت زیر نیز می توان نوشت:

$$\widehat{S}(t) = \prod_{j: t_j \leq t} (1 - \sum_{k=1}^K \widehat{\lambda}_k(t_j)) \quad (10)$$

احتمال شکست علت k ام از حاصلضرب تابع خطر و تابع بقا بدست می آید:

$$\widehat{p}_k = \widehat{\lambda}_k(t_j) \widehat{S}(t_{j-1}) \quad (11)$$

در نهایت تابع تجمعی رخداد براساس علت k ام در زمان t به صورت مجموع این مقادیر برای تمام نقاط زمانی قبل از زمان t برآورد می شود.

$$\hat{I}_k(t) = \Sigma_{j:t_j \leq t} (\hat{p}_k(t_j)) \quad (۱۲)$$

۵ تحلیل داده‌های واقعی

در ادامه به بررسی داده‌های سرطان مثانه (bladder1) که در بسته آماری survival در نسخه R ۲.۳.۴ وجود دارد، می‌پردازیم. این مجموعه داده شامل ۲۹۴ مشاهده در ۱۱ ستون است که سه ستون treatment، status و stop آن در این تحقیق استفاده شده است.

treatment یک متغیر مستقل که نشان‌دهنده نوع درمان بیماران سرطانی است، شامل سه نوع "placebo"، "pyridoxine" و "thiotepa" می‌باشد. status در داده‌های ما نشان‌دهنده نوع داده است که اگر $t_i \leq c_i$ ، یعنی رویدادی زودتر از زمان سانسور رخ داده است و داده ما مشاهده شده است.

$$\delta = ۱, ۲, ۳$$

که به ترتیب نشان‌دهنده بازگشت دوباره به بیماری، مرگ بر اثر سرطان، مرگ بر اثر دلایل خارجی است. در اینجا $status = ۲$ رویداد مورد توجه تحقیق و $status = ۳$ خطر رقابتی آن است. همچنین اگر $t_i \geq c_i$ یعنی زمان سانسور از زمان رویداد زودتر رخ داده است و داده سانسور شده است.

$$\delta = ۰$$

زمان رخداد رویداد در این مجموعه داده همان stop در نظر گرفته شده است که زمان رخ دادن رویداد را نمایش می‌دهد.

در تحلیل‌های زیر از مدل رگرسیون خطر نسبی کاکس استفاده شده است. در این مدل، نوع درمان بیماران (treatment) به عنوان تنها متغیر مستقل در نظر گرفته شد. به این منظور، دو تحلیل مجزا انجام گردید. در تحلیل اول، وضعیت مرگ بر اثر سرطان مثانه به عنوان رویداد مورد توجه ($status = ۲$) و در تحلیل دوم، مرگ بر اثر دلایل

خارجی به عنوان رویداد مورد توجه تعریف شده است ($status = 3$). نکته قابل توجه در این تحلیل‌ها این است که تاثیر نوع درمان بر وضعیت مرگ بیماران در این تحقیق محاسبه شده است.

Call:

```
coxph(formula = Surv(stop, status == 2) ~ treatment, data = new_data)
```

	coef	exp(coef)	se(coef)	z	p
treatmentpyridoxine	-1.910e+01	5.052e-09	8.589e+03	-0.002	0.9982
treatmentthiotepa	1.429e+00	4.177e+00	6.308e-01	2.266	0.0234

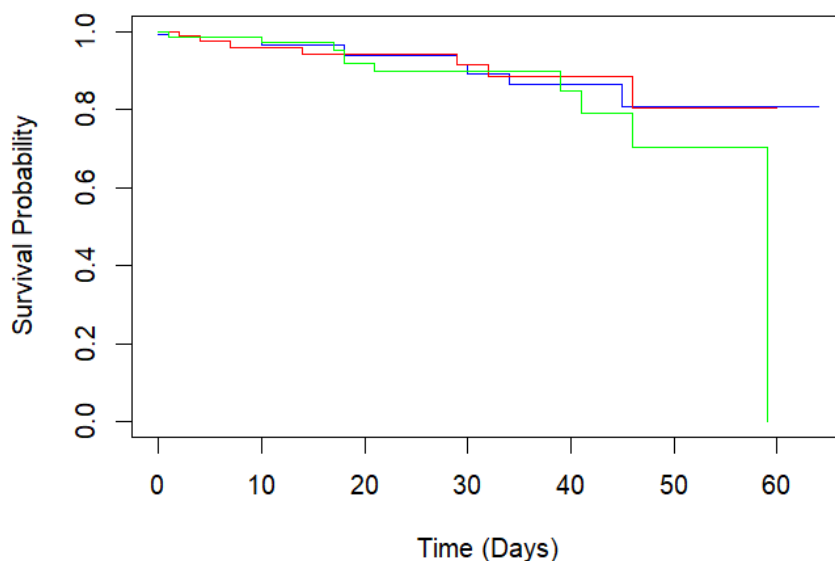
Likelihood ratio test=12.89 on 2 df, p=0.001585

n= 29, number of events= 17

خروجی بالا شامل دو سطر از متغیر مستقل treatment است که اثر تخمینی سطوح مختلف درمان به صورت نسبت لگاریتمی خطر را نسبت به گروه پایه دارونما (placebo) نشان می‌دهند. سطر اول متغیر مستقل روش pyridoxine را با نسبت خطر نزدیک به صفر و p -مقدار، ۰/۹۹۸۲ نمایش می‌دهد. به دلیل این که مقدار P بیشتر از سطح معناداری ۰/۰۵ است، می‌توان گفت که این گروه تاثیری بر روی نرخ خطر رخ دادن مرگ بر اثر سرطان مثانه ندارد. اما p -مقدار در گروه دوم درمان یعنی thiotepa کمتر از سطح معناداری است ($۰/۰۲۳ > ۰/۰۵$). همچنین نسبت خطر در این گروه برابر ۴/۱۷۷ می‌باشد، بدین معنا که افراد تحت درمان با thiotepa تقریباً چهار برابر گروه پایه در معرض خطر مرگ بر اثر سرطان مثانه می‌باشند.

نمودار کاپلان-مایر بالا نشان‌دهنده بقای بیماران در برابر «مرگ ناشی از بیماری مثانه» است. محور عمودی، احتمال زنده ماندن بیماران تا هر زمان مشخص بدون تجربه مرگ از بیماری مثانه را نشان می‌دهد و محور افقی، زمان را بر حسب روز بیان می‌کند. هر افت در نمودار بیانگر وقوع یک یا چند مورد مرگ ناشی از بیماری مثانه در آن زمان است. رنگ‌های مختلف نیز نمایانگر گروه‌های متفاوت (مانند گروه‌های درمانی) هستند. مقایسه این منحنی‌ها کمک می‌کند تا تأثیر عوامل یا درمان‌های گوناگون بر زمان تا مرگ ناشی از بیماری مثانه بررسی شود.

Kaplan-Meier Curve for Death from Bladder Disease



Call:

```
coxph(formula = Surv(stop, status == 3) ~ treatment, data = new_data)
```

	coef	exp(coef)	se(coef)	z	p
treatmentpyridoxine	1.7742	5.8953	0.7271	2.440	0.0147
treatmentthiotepa	1.6648	5.2848	0.9117	1.826	0.0678

Likelihood ratio test=7.36 on 2 df, p=0.02522

n= 29, number of events= 12

مدل بالا نیز اثر تخمینی سطوح مختلف درمان به صورت نسبت لگاریتمی خطر را نسبت به گروه پایه دارونما (placebo) نشان می‌دهند. در این مورد، رویداد زمانی تعریف می‌شود که وضعیت (status) برابر با ۳ باشد یعنی مرگ بر اثر دلایل خارجی رخ دهد. همانطور که در خروجی مشخص است سطر اول اثرات گروه pyridoxine را نمایش می‌دهد که نسبت خطر برای این گروه برابر ۵/۸۹ می‌باشد، یعنی افرادی که پیریدوکسین دریافت می‌کنند، تقریباً ۶ برابر بیشتر از گروه پایه (placebo) در معرض خطر مرگ بر اثر دیگر دلایل قرار دارند که به دلیل این که p -مقدار کمتر

از سطح معناداری می‌باشد ($0/014 > 0/05$)، از نظر آماری این گروه معنادار هستند و تاثیر بر روی رخ دادن مرگ بر اثر دلایل خارجی دلایل دارند. با وجود اینکه نسبت خطر تخمینی برای گروه thiotepa نیز نسبتاً بالاست (تقریباً $5/3$)، مقدار p به دست آمده ($0/0678$) نشان می‌دهد که شواهد فقط تا حدی معنادار هستند و نمی‌توان با قطعیت در مورد اثر این گروه بر مرگ بر اثر دلایل خارجی اظهار نظر کرد.

در ادامه بسته cmprsk که برای آنالیز خطرهای رقابتی طراحی شده است را فراخوانی می‌کنیم. یک مدل رگرسیون خطرهای رقابتی با استفاده از تابع `crr()` از این بسته برازش داده شده است. ضروری است که برای مدل رگرسیونی متغیر مستقل را به صورت عددی تبدیل کنیم. این مدل، رابطه بین متغیر مستقل `treatment` و خطر توزیع زیرین برای شکست‌های نوع اول را ارزیابی می‌کند.

```
convergence: TRUE
coefficients:
new_data$treatment1
0.19
standard errors:
[1] 0.2947
two-sided p-values:
new_data$treatment1
0.52
```

اگر در یک الگوریتم بهینه‌سازی مورد استفاده برای تخمین پارامترهای مدل به یک جواب پایدار برسد، آن الگوریتم همگرا می‌نامیم.^۷ زمانی که این مقدار برابر True می‌شود نشان می‌دهد الگوریتم با موفقیت یک مجموعه پارامتر ثابت پیدا کرده است. در خطرهای رقابتی این مقدار نمایانگر این است که برازش مدل با موفقیت انجام شده و نتایج قابل اعتماد هستند. با استفاده از محاسبه نمایی ضریب می‌توانیم نسبت خطر را بدست آوریم:

$$HR = \exp(0/19) \approx 1/21$$

Convergence^۷

با افزایش ضریب، نسبت خطر حدود ۱/۲۱ واحد بدست می‌آید که به معنای افزایش ۲۱ درصدی خطر است. با این حال این اثر برآورد شده نسبتاً کوچک است و از نظر آماری معنادار نیست. مدل خطرهای رقابتی نشان می‌دهد که یک ارتباط مثبت کوچک بین متغیر مستقل درمان (treatment) و نرخ خطر توزیع زیرین برای رخداد مرگ بر اثر سرطان مثانه (failcode2) وجود دارد. با این حال، از آنجایی که مقدار p برابر ۰/۵۲ است، این اثر از نظر آماری معنی‌دار نیست. در عمل، نتیجه می‌گیرید که هیچ مدرکی از این مدل وجود ندارد که نشان دهد درمان به طور قابل توجهی توزیع تجمعی رویداد مورد توجه که مرگ بر اثر سرطان است را در مقایسه با گروه مرجع تغییر می‌دهد.

```
convergence: TRUE

coefficients:

new_data$treatment1

0.08237

standard errors:

[1] 0.2668

two-sided p-values:

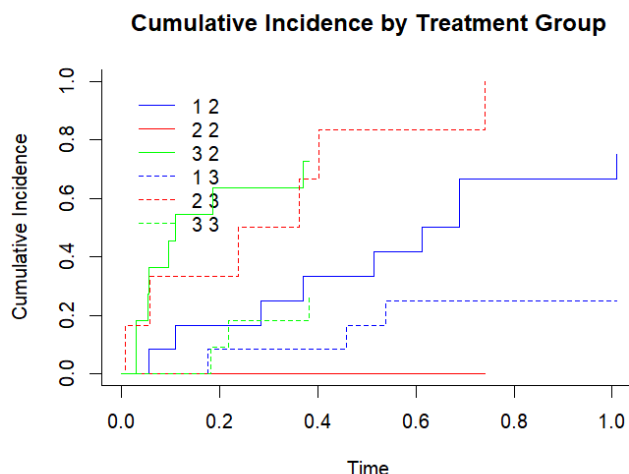
new_data$treatment1

0.76
```

نسبت خطر نیز برای رویداد مرگ بر اثر دلایل خارجی به صورت زیر است:

$$HR = \exp(0.082) \approx 1.08$$

در این مدل نیز با افزایش ضریب، نسبت خطر حدود ۱/۰۸ واحد بدست می‌آید که به معنای افزایش ۸ درصدی خطر است که در نتیجه این اثر برآورد شده باز هم کوچک است و از نظر آماری معنادار نیست. مدل بالا نشان می‌دهد که یک ارتباط مثبت کوچک بین متغیر مستقل درمان (treatment) و نرخ خطر توزیع زیرین برای رخداد مرگ بر اثر دلایل خارجی (failcode3) وجود دارد. با این حال، از آنجایی که مقدار p برابر ۰/۷۶ است، این اثر نیز از نظر آماری معنی‌دار نیست.



در بالا نمودار توابع تجمعی رخداد برای هر ترکیب از گروه‌های درمانی و نوع رخداد ترسیم شده است که محور افقی زمان را نشان می‌دهد و محور عمودی احتمال بروز تجمعی رخداد (از ۰ تا ۱) را در طی زمان مشخص می‌کند. هر منحنی پله‌ای نشان‌دهنده توزیع تجمعی رویداد مورد نظر است؛ به عنوان مثال، منحنی با برچسب "۱۲" نمایانگر احتمال تجمعی رویداد نوع ۲ در گروه درمانی ۱ است.

هرگاه منحنی یک گروه نسبت به گروه دیگر در ارتفاع بالاتری قرار بگیرد، بدین معناست که آن گروه درصد بیشتری از رویداد را تا آن زمان تجربه کرده است. همچنین، اگر پله‌های یک منحنی زودتر بالا بروند، نشان می‌دهد که رویداد زودتر یا سریع‌تر در آن گروه رخ داده است. در مقابل، هم‌پوشانی یا نزدیکی دو منحنی می‌تواند بیانگر شباهت در الگوی بروز رخداد میان آن دو گروه باشد.

در پایان، می‌توان گفت که تحلیل خطرهای رقابتی به وسیله مدل‌های کاپلان-مایر و کاکس نسبی خطر، به بهبود تبیین فرآیند شکست در داده‌های زیستی کمک شایانی کرده است. یافته‌های به دست آمده نشان می‌دهد که متغیرهای مورد بررسی تأثیر معناداری بر بقا دارند. با این حال، برخی فرضیات (مانند استقلال سانسور) ممکن است محدودیت‌هایی ایجاد کنند. در ادامه، پیشنهاد می‌شود که پژوهش‌های آتی با استفاده از داده‌های گسترده‌تر و با در نظر گرفتن مدل‌های جایگزین، به بررسی دقیق‌تر این اثرات پرداخته و کاربردهای عملی آن در حوزه‌های مختلف بالینی و صنعتی مورد ارزیابی قرار گیرند.

- [1] Putter, H. (2007), Tutorial in biostatistics: Competing risks and multi-state models, *Statistics in Medicine*, **26**, 2389–2430.
- [2] Prentice, R.L., Kalbfleisch, J.D., Peterson, A.V., Flournoy, N., Farewell, V.T. and Breslow, N.E. (1978), The analysis of failure times in the presence of competing risks, *Biometrics*, **34**, 541–554.
- [3] Gooley, T.A., Leisenring, W., Crowley, J. and Storer, B.E. (1999), Estimation of failure probabilities in the presence of competing risks: new representations of old estimators, *Statistics in Medicine*, **18**, 695–706.
- [4] Hoover, D.R., Munoz, A., Carey, V., Taylor, J.M.G., Vanraden, M., Chmiel, J.S. and Kingsley, L. (1993), Using events from dropouts $\sim$ in nonparametric survival function estimation with application to incubation of AIDS, *Journal of the American Statistical Association*, **8**, 37–43.
- [5] Bernoulli, D. (1760), Essai d’une nouvelle analyse de la mortalite causee par la petite Verole, et des avantages de l’inoculation pour la prevenir, *Memoires de l’Academie Royal des Sciences*, 1–45.
- [6] Gail, M. (1975), A review and critique of some models used in competing risk analysis, *Biometrics*, **31**, 209–222.