

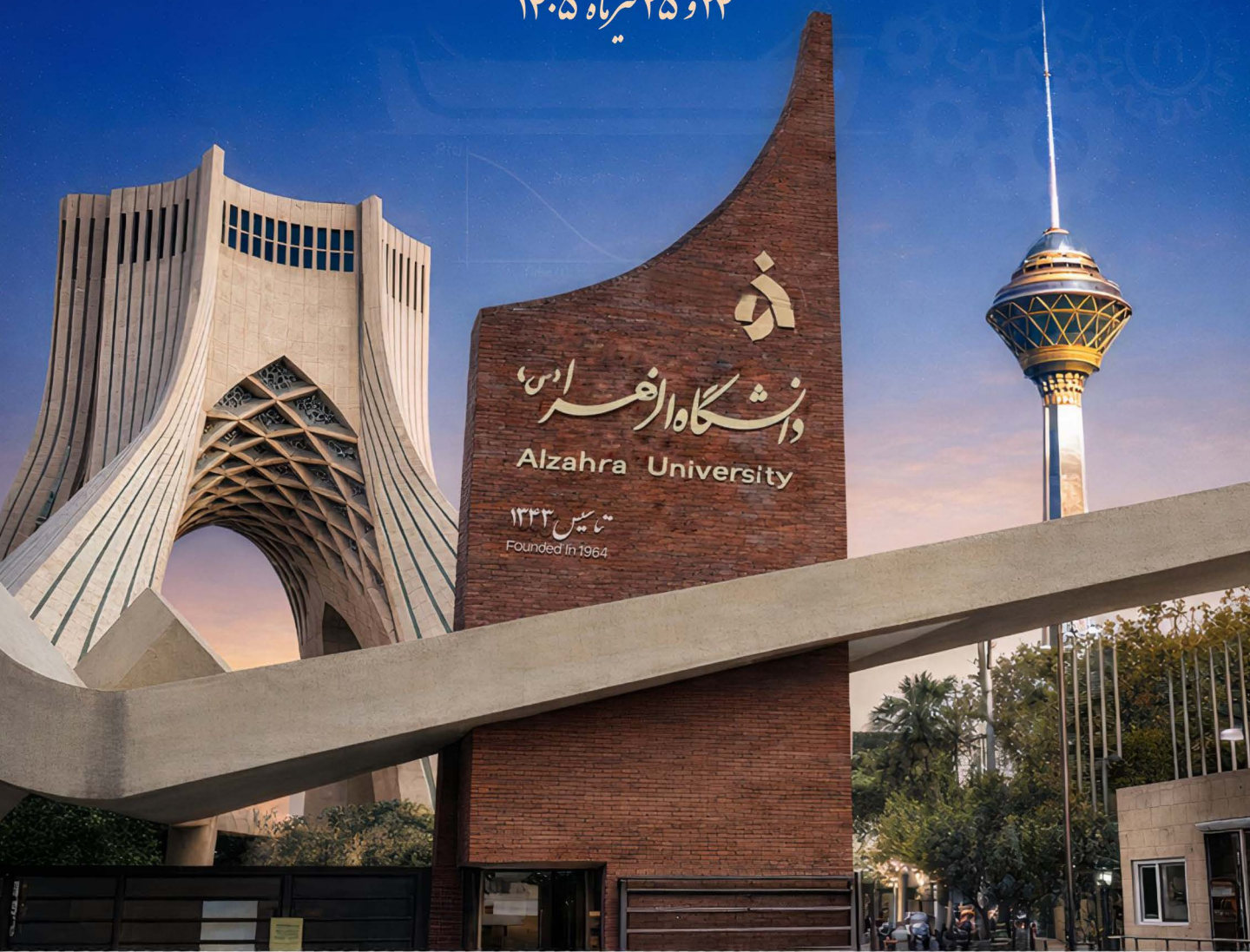
مجموعه مقالات

دوازدهمین سمینار تخصصی

نظریه قابلیت اعتماد و کاربردهای آن

گروه آماد دانشگاه الزهراء (س)

۱۴۰۵ و ۲۵ تیرماه



باسمه تعالی



مجموعه مقالات

دوازدهمین سمینار تخصصی

نظریه قابلیت اعتماد و کاربردهای آن

گروه آمار، دانشکده علوم ریاضی، دانشگاه الزهرا (س)

تهران، ایران

۲۴ و ۲۵ تیر ۱۴۰۵

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

تدوین: مهدی علی محمدی، فریبا عزیزی، محمد زارع

طراح جلد: فاطمه زمانی

تاریخ انتشار: مرداد ماه ۱۴۰۵

مقدمه

با استعانت از درگاه ایزد منان و با سپاس از توفیقات الهی، دوازدهمین دوره «سمینار تخصصی نظریه قابلیت اعتماد و کاربردهای آن» با میزبانی گروه آمار دانشگاه الزهرا (س) برگزار می‌گردد. این نشست علمی که پس از یازده دوره برگزاری موفق، اکنون وارد مرحله جدیدی از تکامل شده است، با هدف پیوند میان نظریه‌های بنیادین و کاربردهای پیشرفته در علم آمار برگزار می‌گردد. برگزاری این دوره از سمینار، حاصل همکاری نزدیک گروه آمار دانشگاه الزهرا (س)، «قطب علمی داده‌های ترتیبی، قابلیت اعتماد و وابستگی»، انجمن آمار ایران و پایگاه استنادی علوم جهان اسلام (ISC) است که همگی در راستای تقویت زیرساخت‌های پژوهشی این حوزه گام برمی‌دارند. در این دوره، تلاش بر این بوده است تا با ارزیابی دقیق مقالات ارسالی توسط کمیته‌های علمی و داوران متخصص، برترین دستاوردهای پژوهشی در قالب سخنرانی و پوستر به اشتراک گذاشته شود. بدین وسیله از حمایت‌های بی‌دریغ معاونت پژوهشی و فناوری دانشگاه الزهرا (س)، پژوهشکده بیمه، شرکت فناوری اطلاعات و گسترش ارتباطات هوشمند کهکشان و تلاش‌های بی‌شائبه اعضای کمیته‌های اجرایی و علمی که تبلور این رویداد هستند، قدردانی می‌نماییم. حضور پررنگ اساتید و پژوهشگران از سراسر کشور، ضامن کیفیت علمی و تعامل سازنده در این نشست خواهد بود. امید است این سمینار، بستری برای تعاملات علمی ماندگار و ارتقای سطح دانش قابلیت اعتماد در جامعه آماری ایران باشد.

نسرین سلطانخواه (رئیس همایش)

مهدی علی‌محمدی (دبیر علمی همایش)

فریبا عزیزی (دبیر اجرایی همایش)

مرداد ماه ۱۴۰۵

محورهای سمینار

قابلیت اعتماد نرم افزار	قابلیت اعتماد سیستم های منسجم و شبکه ها
کنترل کیفیت آماری	روش های بیزی در قابلیت اعتماد
مدل های ضمانت و داده های میدانی	الگوهای تعمیر و نگهداری سیستم ها
قابلیت اعتماد در محیط فازی	ترتیب های تصادفی
مفاهیم سالخورده گی	تحلیل قابلیت اعتماد بر اساس داده های فرسایشی
مطالعات موردی در صنعت	داده کاوی در قابلیت اعتماد
مفاهیم وابستگی و توابع مفصل	آزمون های طول عمر تسریع یافته
ابر داده در قابلیت اعتماد	مدل های شوک
تحلیل ریسک در قابلیت اعتماد	استنباط آماری و تحلیل داده های طول عمر
هوش مصنوعی در مهندسی قابلیت اعتماد	تحلیل بقا

اعضای کمیته علمی (به ترتیب حروف الفبا)

دکتر رضا احمدی، دانشگاه علم و صنعت
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
دکتر مجید چهکنندی، دانشگاه بیرجند
دکتر نوشین حکمی پور، مرکز آموزش عالی بوئین زهرا
دکتر فیروزه حقیقی، دانشگاه تهران
دکتر مهدی دوست پرست، 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 (س)
دکتر مهدی علی‌محمدی، دانشگاه الزهرا (س) (دبیر علمی همایش)
دکتر مریم کلکین‌نما، دانشگاه صنعتی اصفهان
دکتر اکرم کهن‌سال، دانشگاه بین‌المللی امام خمینی
اعضای کمیته مشاورین علمی افتخاری (به ترتیب حروف الفبا)

دکتر جعفر احمدی، دانشگاه فردوسی مشهد
دکتر مجید اسدی، دانشگاه اصفهان
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 غلامرضا محتشمی‌برزادران، دانشگاه فردوسی مشهد

اعضای کمیته اجرایی (به ترتیب حروف الفبا)

دکتر هادی جباری نوقابی، دانشگاه فردوسی مشهد
دکتر محمد زارع، دانشگاه الزهرا (س)
دکتر نسرین سلطان‌خواه حقیقی، دانشگاه الزهرا (س) (رئیس همایش)
دکتر صدیقه شمس، دانشگاه الزهرا (س)
دکتر فریبا عزیزی، دانشگاه الزهرا (س) (دبیر اجرایی همایش)
دکتر مهدی علی‌محمدی، دانشگاه الزهرا (س)

کادر اجرایی سمینار

خانم فاطمه زمانی (دانشجوی کارشناسی و دبیر انجمن علمی-دانشجویی آمار)

فهرست

تحلیل بیزی برای تعیین سیاست بهینه آب‌بندی با استفاده از غربالگری چندمرحله‌ای

آزموده نژاد، م.، احمدی، ج. ۱

بهبود قابلیت اعتماد سیستم‌های k از n از طریق روش کاهش

چپ‌کندی، م. ۱۹

آزمون اجتماع-اشتراک برای ارزیابی قابلیت اطمینان سیستم‌های سری

چینی‌پرداز، ر.، نیکنام، ز.، چینی‌پرداز، م. ۳۲

نگهداری و تعمیر سیستم‌های تقسیم بار k -out-of- n :G بر اساس محاسبه باقیمانده طول عمر سیستم

با استفاده از روش‌های یادگیری تقویتی

حقیقی، ف.، صفری، ع.، عیدی، ش. ۴۹

یک رویکرد ترکیبی مبتنی بر یادگیری عمیق برای پیش‌بینی عمر مفید باقیمانده

کرمی، ا.، حقیقی، ف.، صفری، ع. ۵۹

تصحیح سوگیری سانسور وابسته در تحلیل قابلیت اعتماد مبتنی بر مدل کاکس

گلی‌ناری، ح.، خراشادی‌زاده، م.، محتشمی‌برزادران، غ. ۷۳

مدل کاکس: از نظریه توزیع‌ها تا کاربردهای عملی

مهدی‌پور، م.، توانگر، م.، بیدرام، ح. ۹۱



تحلیل بیزی برای تعیین سیاست بهینه آب‌بندی با استفاده از غربالگری چندمرحله‌ای

آزموده نژاد، م^۱ احمدی، ج^۲

گروه آمار، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد

چکیده

در این مقاله یک چارچوب بیزی برای غربالگری چندمرحله‌ای دوره آب‌بندی قطعاتی ارائه می‌شود که تحت تعمیر مینیمال قرار می‌گیرند. فرایند تعمیر مینیمال به صورت یک فرایند پواسون ناهمگن با تابع شدت پایه از نوع قانون توانی آمیخته مدل‌سازی می‌شود تا ناهمگنی پیوسته مشاهده نشده بین قطعات را در نظر بگیرد. همچنین در برآورد بیزی، توزیع پسین پارامترها با استفاده از روش زنجیره مارکوف مونت کارلو به‌روزرسانی می‌شود. از طرفی در نظر گرفتن سیاست‌های غربالگری چندمرحله‌ای پیشنهادی، امکان تعریف چندین نقطه بازرسی با آستانه‌های غربالگری متناظر را فراهم می‌کند. یک قطعه در نخستین مرحله‌ای که تعداد تجمعی تعمیرات مینیمال آن از آستانه غربالگری تعیین شده فراتر رود کنار گذاشته می‌شود؛ در غیراینصورت وارد مرحله بهره‌برداری میدانی خواهد شد. سپس مدل‌های بهینه‌سازی مبتنی بر هزینه و عملکرد برای تعیین مدت زمان بهینه آب‌بندی و مقدار بهینه

¹Speaker. ma.azmoudeh@gmail.com

²ahmadi-j@um.ac.ir

آستانه غربالگری تدوین شده‌اند. در نهایت با استفاده از مثال‌های عددی و تحلیل حساسیت، برتری رویکرد بیزی با آستانه‌های غربالگری چندمرحله‌ای در مقایسه با سیاست تک مرحله‌ای موجود را نشان می‌دهد.

کلمات کلیدی: آب‌بندی، بهینه‌سازی، تحلیل بیزی، تعمیر مینیمال، غربالگری.

۱ پیش‌گفتار

آب‌بندی (*Burn-in*) یکی از فرایندهای مهم و پرکاربرد در صنایع مختلف است که با هدف شناسایی و حذف قطعات معیوب پیش از تحویل به مشتریان یا استفاده در عملیات میدانی انجام می‌شود. در این روش، قطعات تحت شرایط شبیه‌سازی شده مانند دما و فشار، مشابه با شرایط کاری واقعی آزمایش می‌شوند تا نقص‌های احتمالی و خرابی‌های زودهنگام شناسایی شوند. قطعاتی که در طول این آزمون دچار خرابی می‌شوند، به‌طور کامل از مجموعه خارج می‌شوند و تنها قطعاتی که این آزمون را با موفقیت پشت سر می‌گذارند، به‌عنوان قطعات با کیفیت استاندارد به مرحله بعدی (تحویل مشتری یا عملیات میدانی) ارسال می‌شوند. اجرای موفقیت‌آمیز فرایند آب‌بندی، نیاز به تعیین مدت زمان بهینه آن است؛ این زمان باید به گونه‌ای انتخاب شود که ضمن شناسایی عیوب زودرس، از اتلاف وقت و هزینه بر روی قطعاتی که احتمال خرابی کمی دارند، جلوگیری کند. بنابراین یافتن طول دوره آب‌بندی یکی از مسائل مهم در طراحی فرایندهای تولیدی است. مقیمی حاجی [۵]، به مسئله تعیین طول بهینه دوره‌های آب‌بندی و گارانتی برای محصولات می‌پردازد. او برای کاهش هزینه‌های کلی تولیدکننده و در عین حال افزایش رضایت مشتری، یک مدل ریاضی برای محاسبه هزینه‌ها در این دو دوره ارائه کرده است. 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 [۲]، با در نظر گرفتن سیاست تعویض بلوکی و تعمیر مینیمال، یک مدل بهینه برای تعیین زمان آب‌بندی و فواصل بهینه تعمیر و تعویض قطعات ارائه کرده‌اند.

در بسیاری از کاربردهای عملی، سیستم‌ها و تجهیزات مورد استفاده، یکسان نیستند و اغلب بر اساس ویژگی‌های قابلیت اعتماد به زیرجمعیت‌های متفاوتی طبقه‌بندی می‌شوند. بنابراین، هنگام طراحی استراتژی‌های بهینه برای آب‌بندی و نگهداری، در نظر گرفتن این ناهمگونی جمعیت، یک مسئله کلیدی و ضروری است. لینگ و همکاران [۴] 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 روش مینیمال تعمیر می‌شوند. در این رویکرد، تعداد کل تعمیرات مینیمال مشاهده‌شده در طول فرایند آب‌بندی به عنوان یک قاعده غربالگری برای شناسایی و حذف اجزای معیوب با عملکرد ضعیف مورد استفاده قرار می‌گیرد. سپس با استفاده از مدل‌سازی احتمالاتی، سیاستی بهینه برای

کمینه کردن هزینه‌ها و افزایش قابلیت اعتماد قطعات در دوره مأموریت ارائه داده‌اند. صفائی و تقی پور [۶] راهکاری نوین برای محصولات با قابلیت اعتماد زیاد ارائه داده‌اند. از آنجایی که روش‌های قدیمی که فقط بر اساس طول عمر کار می‌کردند، برای این محصولات کارایی ندارند، آن‌ها مدلی را توسعه داده‌اند که بر فرسایش محصول تمرکز دارد. این مدل یکپارچه، فرایندهای آب‌بندی و نگهداری پیشگیرانه را به طور همزمان بهینه می‌کند. آن‌ها فرض نموده‌اند که محصولات ممکن است به دسته‌های مختلفی (زیرجمعیت‌های ناهمگن) با نرخ فرسایش متفاوت تقسیم شوند. با استفاده از فرایند گاما، مدل سطح فرسایش را اندازه‌گیری کرده و محصولات دارای نقص را تشخیص می‌دهد. سپس، چهار عامل مهم شامل مدت زمان آب‌بندی، نرخ استفاده در این دوره، آستانه فرسایش و زمان‌بندی نگهداری را به گونه‌ای تنظیم می‌کند که هزینه‌ها به حداقل رسیده و در دسترس بودن سیستم به حداکثر برسد، البته با این شرط اساسی که ایمنی کاربران یا سیستم به هیچ وجه به خطر نیفتد. در مطالعه لینگ و همکاران [۳]، مدل‌سازی به صورت خاص بر روی قطعاتی تمرکز دارد که پس از خرابی، به جای تعویض کامل، تعمیرات مینیمال روی آن‌ها صورت می‌گیرد. سپس، با استفاده از فرایندهای پواسون ناهمگن و توابع شدت قانون توانی، رفتار خرابی قطعات در طول زمان و تحت شرایط عملیاتی مختلف، مدل‌سازی می‌شود. این مدل‌ها امکان پیش‌بینی احتمالی خرابی‌های آینده را با در نظر گرفتن عدم قطعیت ذاتی در پارامترهای سیستم فراهم می‌کنند. در ادامه، این مدل‌ها برای استخراج یک قانون غربالگری واحد به کار گرفته می‌شوند. این قانون، یک آستانه مشخص را تعیین می‌کند؛ اگر تعداد تعمیرات یک قطعه در طول دوره آب‌بندی از این آستانه فراتر رود، قطعه معیوب تلقی شده و حذف می‌گردد، در غیر این صورت، پس از اتمام دوره آب‌بندی، قطعه به چرخه عملیاتی وارد می‌شود. هدف اصلی این رویکرد، بهینه‌سازی هزینه‌های کلی سیستم از طریق کاهش خرابی‌های پیش‌بینی نشده در مرحله عملیاتی و مدیریت مؤثرتر منابع است.

در این مقاله، ما یک چارچوب بیزی برای پیاده‌سازی غربالگری چندمرحله‌ای در طول دوره آب‌بندی قطعات تحت تعمیر مینیمال ارائه می‌دهیم. این رویکرد، گامی فراتر از سیاست‌های تک مرحله‌ای رایج است و امکان ارزیابی دقیق‌تر و انعطاف‌پذیرتر وضعیت قطعات را فراهم می‌کند. فرایند تعمیر مینیمال به صورت یک فرایند پواسون ناهمگن با تابع شدت پایه از نوع قانون توانی مدل‌سازی می‌شود. این انتخاب مدل‌سازی، به طور مؤثری ناهمگنی پیوسته مشاهده نشده بین قطعات را که می‌تواند بر رفتار خرابی آن‌ها تأثیر گذار باشند، در بر می‌گیرد. در چارچوب بیزی توسعه یافته، برآورد پارامترهای مدل با استفاده از روش زنجیره مارکوف مونت کارلو و از طریق به‌روزرسانی پیوسته توزیع پسین انجام می‌شود. این روش امکان مدیریت عدم قطعیت در پارامترهای مدل را به شیوه‌ای جامع فراهم می‌کند. علاوه بر این، با معرفی سیاست‌های غربالگری چندمرحله‌ای، ساختاری برای تعریف نقاط بازرسی متعدد به همراه آستانه‌های غربالگری متناظر ایجاد کرده‌ایم. بدین ترتیب، یک قطعه در نخستین مرحله‌ای که تعداد تجمعی تعمیرات مینیمال آن از آستانه غربالگری تعیین شده فراتر رود، کنار گذاشته می‌شود؛ در غیر این صورت، اگر از تمام مراحل غربالگری

با موفقیت عبور کند، وارد مرحله بهره‌برداری میدانی خواهد شد. برای تعیین بهترین استراتژی، مدل‌های بهینه‌سازی مبتنی بر هزینه و عملکرد تدوین شده‌اند. این مدل‌ها به ما امکان می‌دهند تا مدت بهینه آب‌بندی و همچنین بردار بهینه آستانه‌های غربالگری را تعیین کنیم. در نهایت، با استفاده از نتایج عددی، برتری رویکرد بیزی با آستانه‌های غربالگری چندمرحله‌ای در مقایسه با سیاست تک مرحله‌ای موجود را بررسی می‌کنیم.

ادامه مقاله به صورت زیر تنظیم شده است. در بخش ۲، مدل‌بندی بیزی معرفی و پارامترهای آن با استفاده از شبیه‌سازی برآورد می‌شوند. بخش ۳، سیاست غربالگری چندمرحله‌ای پیشنهادی از دیدگاه احتمالاتی ارائه می‌شود. بخش ۴، مدل متوسط هزینه کل تحت سیاست چندمرحله‌ای پیشنهادی استخراج و مسئله بهینه‌سازی متناظر فرمول‌بندی می‌گردد. بخش ۵، به مسئله بهینه‌سازی مبتنی بر مدل عملکرد پرداخته می‌شود. در بخش ۶، مسئله بهینه‌سازی مدل پیشنهادی به صورت عددی و تحلیل حساسیت انجام می‌شود. در انتها خلاصه‌ای از جمع‌بندی ارائه شده است.

نمادها و علائم استفاده شده در این مقاله به صورت زیر معرفی می‌شوند.

Z : متغیر تصادفی از توزیع گاما

α : پارامتر شکل توزیع گاما

θ : پارامتر نرخ توزیع گاما

$r(.)$: تابع نرخ خرابی

λ : پارامتر مقیاس تابع قانون توانی

β : پارامتر شکل تابع قانون توانی

$\lambda_0(.)$: تابع شدت پایه قانون توانی

$\Lambda(.)$: تابع شدت تجمعی

$N(.)$: تعداد تعمیرات مینیمال

b : حداکثر مدت زمان آب‌بندی

M_i : تعداد تجمعی تعمیرات مینیمال تا زمان t_i

m : آستانه غربالگری

$j_{i,l}$: تعداد تعمیرات مینیمال قطعه i ام در بازه (s_{l-1}, s_l)

$\phi = (\lambda, \beta, \alpha, \theta)$: بردار پارامترهای مجهول

$\mathbf{j}_i$: بردار داده برای قطعه i ام

J_i : تعداد کل تعمیرات مینیمال قطعه i ام

D : مجموعه داده‌های مشاهده شده

$\pi(\cdot)$: تابع چگالی پیشین

$L(\cdot)$: تابع درستنمایی

τ : دوره مأموریت

k : تعداد نقاط بازرسی

A_i : عبور قطعه از i مرحله اول

C_s : هزینه راه‌اندازی و عملیات آب‌بندی

C_m : هزینه تعمیرات مینیمال یک قطعه در طول دوره آب‌بندی

C_d : هزینه اسقاط در پایان آب‌بندی

L : زیان ناشی از شکست مأموریت

K : قیمت فروش یک قطعه

t_j : زمان بازرسی j ام

t^* : زمان بهینه آب‌بندی

m^* : آستانه غربالگری بهینه

۲ مدل‌بندی بیزی

در این بخش، مدل بیزی برای تحلیل سیاست آب‌بندی با قطعات قابل تعمیر و ناهمگنی پیوسته ارائه می‌شود. ابتدا فرایند قانون توانی آمیخته به‌عنوان مدل پایه برای توصیف رفتار خرابی معرفی می‌شود. یک سیاست غربالگری چندمرحله‌ای ساختاریافته برای مدیریت ناهمگنی ذاتی قطعات ارائه می‌گردد. با استفاده از داده‌های خرابی مشاهده‌شده، تابع درستنمایی مدل ساخته می‌شود. برای پارامترهای مدل، فرض‌های پیشین مناسب تعیین می‌شوند و در نهایت، با به‌کارگیری روش‌های شبیه‌سازی پسین (مانند $MCMC$)، برآوردهای بیزی از پارامترهای مدل بدست می‌آید.

۱.۲ فرایند قانون توانی آمیخته

یک جمعیت ناهمگن از قطعات قابل تعمیر را در نظر بگیرید که در آن ناهمگنی مشاهده نشده توسط یک متغیر تصادفی نامنفی Z مدل می‌شود. فرض کنید Z از توزیع گاما با پارامتر شکل $\alpha > 0$ و پارامتر نرخ $\theta > 0$ پیروی کند

$$Z \sim \text{Gamma}(\alpha, \theta),$$

که تابع چگالی احتمال آن به صورت

$$f_Z(z) = \frac{\theta^\alpha}{\Gamma(\alpha)} z^{\alpha-1} e^{-\theta z}, \quad z > 0, \quad (1)$$

باشد، در این صورت بر اساس مطالعات لینگ و همکاران [۳]، تابع نرخ خرابی (شدت) یک قطعه به شرط $Z = z$ به صورت

$$r(t|z) = z\lambda\beta t^{\beta-1}, \quad t > 0, \quad (2)$$

داده می شود که در آن $\lambda > 0$ و $\beta > 0$ به ترتیب پارامترهای مقیاس و شکل تابع شدت پایه قانون توانی آمیخته

$$\lambda_0(t) = \lambda\beta t^{\beta-1},$$

هستند. با توجه به (۲)، تابع شدت تجمعی متناظر برابر است با

$$\Lambda(t|z) = \int_0^t r(u|z) du = z\lambda t^\beta.$$

فرایند تعمیر مینیمال $\{N(t), t \geq 0\}$ به شرط $Z = z$ یک فرایند پواسون ناهمگن با تابع شدت $\Lambda(t|z)$ است. بنابراین توزیع شرطی تعداد تعمیرات مینیمال در بازه $(s, t]$ به صورت زیر می باشد

$$N(t) - N(s) | Z = z \sim \text{Poisson} \left(z\lambda (t^\beta - s^\beta) \right). \quad (3)$$

۲.۲ سیاست غربالگری چندمرحله ای

یک سیاست غربالگری چندمرحله ای آب بندی با $k \geq 1$ نقطه بازرسی، به صورت

$$0 < t_1 < t_2 < \dots < t_k = b,$$

را در نظر می گیریم که در آن b حداکثر مدت زمان آب بندی است. در مرحله i ام، تعداد تجمعی تعمیرات مینیمال تا زمان t_i را با $M_i = N(t_i)$ نشان می دهیم.

با مفروضات بالا، قاعده غربالگری به صورت زیر تعریف می شود:

- در مرحله i (برای $i = 1, \dots, k-1$)، اگر $M_i > m_i$ باشد، قطعه در همان زمان t_i کنار گذاشته می شود.
 - اگر برای تمام مراحل $i = 1, \dots, k-1$ داشته باشیم $M_i \leq m_i$ و در مرحله نهایی نیز $M_k \leq m_k$ ، قطعه آزمون آب بندی را با موفقیت گذرانده و وارد بهره برداری میدانی یا فروش می شود.
- در اینجا $\mathbf{m} = (m_1, m_2, \dots, m_k)$ بردار آستانه های غربالگری است که همراه با b باید بهینه سازی شود.

۳.۲ تابع درستنمایی

فرض کنید پیش از اجرای سیاست جدید، داده‌هایی را از N قطعه در اختیار داریم که طی k بازه زمانی ثابت با نقاط

$$0 = s_0 < s_1 < \dots < s_k = T,$$

مشاهده شده‌اند. با توجه به رابطه (۳)، تابع درستنمایی قطعه i ام به شرط $Z_i = z$ و بردار پارامترهای ϕ به صورت

$$L_i(\mathbf{j}_i|z, \phi) = \prod_{l=1}^k \frac{[z\lambda(s_l^\beta - s_{l-1}^\beta)]^{j_{i,l}}}{j_{i,l}!} \exp\left(-z\lambda(s_l^\beta - s_{l-1}^\beta)\right), \quad (4)$$

داده می‌شود که در آن $j_{i,l}$ تعداد تعمیرات مینیمال در بازه $(s_{l-1}, s_l]$ ، $\phi = (\lambda, \beta, \alpha, \theta)$ بردار پارامترها و $\mathbf{j}_i = (j_{i,1}, \dots, j_{i,k})$ بردار داده برای قطعه i ام است.

از آنجایی که Z_i قابل مشاهده نیست، تابع درستنمایی حاشیه‌ای با انتگرال‌گیری از روابط (۱) و (۴) نسبت به توزیع گامای Z_i به صورت

$$\begin{aligned} L_i(\mathbf{j}_i|\phi) &= \int_0^\infty L_i(\mathbf{j}_i|z, \phi) f_Z(z) dz \\ &= \int_0^\infty \left[\frac{z^{J_i} \prod_{l=1}^k (\lambda(s_l^\beta - s_{l-1}^\beta))^{j_{i,l}}}{\prod_{l=1}^k j_{i,l}!} e^{-z\lambda s_k^\beta} \right] \left[\frac{\theta^\alpha}{\Gamma(\alpha)} z^{\alpha-1} e^{-\theta z} \right] dz \\ &= \frac{\theta^\alpha}{\Gamma(\alpha) \prod_{l=1}^k j_{i,l}!} \left[\prod_{l=1}^k (\lambda(s_l^\beta - s_{l-1}^\beta))^{j_{i,l}} \right] \int_0^\infty z^{\alpha+J_i-1} e^{-z(\theta+\lambda s_k^\beta)} dz \\ &= \frac{\theta^\alpha}{\Gamma(\alpha) \prod_{l=1}^k j_{i,l}!} \left[\prod_{l=1}^k (\lambda(s_l^\beta - s_{l-1}^\beta))^{j_{i,l}} \right] \frac{\Gamma(\alpha + J_i)}{(\theta + \lambda s_k^\beta)^{\alpha+J_i}} \\ &= \frac{\Gamma(\alpha + J_i)}{\Gamma(\alpha) \prod_{l=1}^k j_{i,l}!} \frac{\theta^\alpha \prod_{l=1}^k [\lambda(s_l^\beta - s_{l-1}^\beta)]^{j_{i,l}}}{(\theta + \lambda s_k^\beta)^{\alpha+J_i}}, \end{aligned} \quad (5)$$

بدست می‌آید، که در آن

$$J_i = \sum_{l=1}^k j_{i,l},$$

تعداد کل تعمیرات مینیمالی است که روی قطعه i ام انجام شده است. در نتیجه تابع درستنمایی برای همه قطعات به صورت زیر داده می‌شود

$$L(D|\phi) = \prod_{i=1}^N L_i(\mathbf{j}_i|\phi), \quad (6)$$

که در آن

$$D = \{\mathbf{j}_1, \dots, \mathbf{j}_N\},$$

مجموعه داده‌های مشاهده شده است. بنابراین با جایگذاری (۵) در (۶)، تابع درستنمایی به صورت زیر قابل بیان است

$$L(D|\phi) = \prod_{i=1}^N \left[\frac{\Gamma(\alpha + J_i)}{\Gamma(\alpha) \prod_{l=1}^k j_{i,l}!} \frac{\theta^\alpha \prod_{l=1}^k [\lambda(s_l^\beta - s_{l-1}^\beta)]^{j_{i,l}}}{(\theta + \lambda s_k^\beta)^{\alpha + J_i}} \right]. \quad (۷)$$

۴.۲ فرض‌های پیشین

برای انجام استنباط بیزی، فرض‌های پیشین مستقل زیر در نظر گرفته می‌شود.

- $\lambda \sim \text{Gamma}(a_\lambda, b_\lambda)$ با $a_\lambda = 1$ و $b_\lambda = 5/0$

- $\beta \sim \text{Beta}(a_\beta, b_\beta)$ بریده شده به بازه $(0, 5)$ با $a_\beta = 5/2$ و $b_\beta = 2$

- $\alpha \sim \text{Gamma}(a_\alpha, b_\alpha)$ با $a_\alpha = 2$ و $b_\alpha = 1$

- $\theta \sim \text{Gamma}(a_\theta, b_\theta)$ با $a_\theta = 1$ و $b_\theta = 5/0$

چگالی پیشین توأم برابر با

$$\pi(\phi) = \pi(\lambda)\pi(\beta)\pi(\alpha)\pi(\theta), \quad (۸)$$

است. با جایگذاری فرض‌های پیشین در (۸)، چگالی پیشین توأم به صورت زیر بدست می‌آید

$$\pi(\phi) = \lambda^{a_\lambda-1} e^{-b_\lambda \lambda} \times \beta^{a_\beta-1} (1-\beta)^{b_\beta-1} \times \alpha^{a_\alpha-1} e^{-b_\alpha \alpha} \times \theta^{a_\theta-1} e^{-b_\theta \theta}. \quad (۹)$$

بنا به تعریف چگالی پسین توأم، روابط زیر را داریم

$$\begin{aligned} \pi(\phi|D) &= \frac{L(D|\phi)\pi(\phi)}{\int L(D|\phi)\pi(\phi)d\phi} \\ &\propto L(D|\phi)\pi(\phi). \end{aligned} \quad (۱۰)$$

با جایگذاری روابط (۷) و (۹) در رابطه (۱۰)، چگالی پسین نرمال نشده به صورت

$$\begin{aligned} \pi(\phi|D) &\propto \prod_{i=1}^N \left[\frac{\Gamma(\alpha + J_i)}{\Gamma(\alpha) \prod_{l=1}^k j_{i,l}!} \frac{\theta^\alpha \prod_{l=1}^k [\lambda(s_l^\beta - s_{l-1}^\beta)]^{j_{i,l}}}{(\theta + \lambda s_k^\beta)^{\alpha + J_i}} \right] \\ &\times \lambda^{a_\lambda-1} e^{-b_\lambda \lambda} \times \beta^{a_\beta-1} (1-\beta)^{b_\beta-1} \times \alpha^{a_\alpha-1} e^{-b_\alpha \alpha} \times \theta^{a_\theta-1} e^{-b_\theta \theta}, \end{aligned} \quad (۱۱)$$

بازنویسی می‌شود. اما این رابطه فرم بسته ندارد، بنابراین نمونه‌گیری مستقیم از (۱۱) امکان‌پذیر نیست و از الگوریتم $MCMC$ نوع متروپولیس درون گیبس استفاده می‌کنیم.

۵.۲ شبیه‌سازی پسین

گام‌های الگوریتم $MCMC$ با متروپولیس درون گیبس به ترتیب زیر می‌باشد
۱. مقداردهی اولیه

$$\phi^{(0)} = (\lambda^{(0)}, \beta^{(0)}, \alpha^{(0)}, \theta^{(0)})$$

از میانگین‌های پیشین. شمارنده تکرار $r = 0$.

۲. برای $r = 1$ تا R (۲۵۰۰۰):

الف) نمونه‌گیری $\lambda^{(r)}$ از

$$\pi(\lambda | \beta^{(r-1)}, \alpha^{(r-1)}, \theta^{(r-1)}, D).$$

ب) نمونه‌گیری $\beta^{(r)}$ از

$$\pi(\beta | \lambda^{(r)}, \alpha^{(r-1)}, \theta^{(r-1)}, D).$$

ج) نمونه‌گیری $\alpha^{(r)}$ از

$$\pi(\alpha | \lambda^{(r)}, \beta^{(r)}, \theta^{(r-1)}, D).$$

د) نمونه‌گیری $\theta^{(r)}$ از

$$\pi(\theta | \lambda^{(r)}, \beta^{(r)}, \alpha^{(r)}, D).$$

۳. حذف R نمونه اولیه به عنوان آب‌بندی (۸۰۰۰) و تنک سازی هر ۱۰ نمونه برای کاهش خودهمبستگی.

جدول ۱: تعداد تجمعی تعمیرات مینیمال

بازه‌های زمانی	۱	۲	۳	۴	۵	۶	۷	۸	۹	۱۰
[۰, ۲]	۳	۱	۴	۲	۰	۵	۱	۳	۲	۴
(۲, ۴]	۵	۲	۶	۳	۱	۸	۲	۴	۳	۵
(۴, ۶]	۷	۴	۸	۵	۳	۱۰	۳	۶	۵	۷
(۶, ۸]	۹	۵	۱۱	۷	۴	۱۳	۴	۸	۷	۹
(۸, ۱۰]	۱۲	۶	۱۴	۹	۵	۱۷	۶	۱۰	۸	۱۱

جدول ۱، تعداد تجمعی تعمیرات مینیمال برای ۱۰ مؤلفه در ۵ بازه زمانی را نشان می‌دهد. با استفاده الگوریتم $MCMC$ تعداد ۲۵۰۰۰ شبیه‌سازی انجام می‌دهیم. بر اساس نمونه‌های شبیه‌سازی شده، میانگین پسین به‌عنوان برآورد پسین پارامتر به‌صورت

$$\hat{\lambda} = \frac{1}{M} \sum_{r=1}^M \lambda^{(r)}, \quad (12)$$

محاسبه می‌شود که در آن M تعداد نمونه‌های نگه داشته شده است. به‌طور مشابه برای سایر پارامترها از رابطه (۱۲) استفاده می‌شود.

جدول ۲: مقادیر پیشین و پسین پارامترها

پارامترها	λ	β	α	θ
پیشین	۳/۲	۵/۱	۲	۱
پسین	۳۴۷۲/۲	۴۷۲۱/۱	۱۱۸۴/۲	۰۴۷۳/۱

همانطور که در جدول ۱ مشاهده می‌کنیم، برآورد پسین بسیار نزدیک به مقادیر واقعی شبیه‌سازی شده هستند بنابراین نشان دهنده این است که این الگوریتم به‌خوبی کارکرده است.

۳ تحلیل احتمالاتی سیاست غربالگری چندمرحله‌ای

در این بخش، برتری سیاست غربالگری چندمرحله‌ای پیشنهادی را از دیدگاه احتمالاتی نشان می‌دهیم. بدین ترتیب چندین نقطه بازرسی را در نظر می‌گیریم.

فرض کنید

$$P(b | \tau \text{ موفق}),$$

احتمال آن باشد که یک قطعه پس از عبور از سیاست آب‌بندی، دوره مأموریت ثابت τ را بدون خرابی با موفقیت، کامل کند در این صورت سه سیاست زیر را مقایسه می‌کنیم:

- بدون غربالگری: قطعه بلافاصله آزاد می‌شود ($b = 0$).
- سیاست تک مرحله‌ای (b, m): یک بازرسی در زمان b با آستانه غربالگری m .

• سیاست چندمرحله‌ای (t, m) : شامل k نقطه بازرسی $t_1 < \dots < t_k = b$ با آستانه‌های غربالگری $m_1, \dots, m_k$.

برای مقایسه توزیع پسین Z تحت دو سیاست غربالگری، ابتدا ترتیب تصادفی نسبت درستنمایی زیر را تعریف می‌کنیم که ابزار اصلی اثبات، در ادامه خواهد بود.

تعریف ۱.۳. متغیر تصادفی X در ترتیب تصادفی نسبت درستنمایی از Y کمتر است هرگاه برای هر $t \geq 0$ ، نسبت

$$\frac{f_X(t)}{f_Y(t)},$$

تابعی نزولی از t باشد و آن را با نماد $X \leq_{lr} Y$ نشان می‌دهند.

قضیه زیر یکی از نتایج مهم در تحلیل احتمالاتی سیاست آب‌بندی است که اثر آستانه‌های غربالگری سختگیرانه‌تر در آب‌بندی را نشان می‌دهد.

قضیه ۲.۳. فرض کنید دو سیاست چندمرحله‌ای با زمان‌های بازرسی یکسان t و با آستانه‌های $m^{(1)} \leq m^{(2)}$ باشند، در این صورت برای هر طول مأموریت $\tau > 0$ داریم

$$P(b + \tau \text{ موفق} | m^{(1)}) \geq P(b + \tau \text{ موفق} | m^{(2)}). \quad (1)$$

برهان. تعداد تعمیرات مینیمال به شرط $Z = z$ تا زمان t_i دارای توزیع

$$M_i | z \sim \text{Poisson}(z \lambda t_i^\beta),$$

است. احتمال عبور از کل سیاست و سپس بقای مأموریت برابر است با

$$P(\text{موفقیت}) = E \left[\exp \left(-Z \lambda (b + \tau)^\beta + Z \lambda b^\beta \right) \middle| \text{مراحل} \right].$$

از آنجایی که جمله نمایی نسبت به Z نزولی است و آستانه‌های غربالگری $m^{(1)}$ نسبت به $m^{(2)}$ توزیع پسین کوچکتری برای Z ایجاد می‌کنند (بدلیل برقراری خاصیت نسبت درستنمایی یکنواخت در توزیع پواسون-گاما که رابطه

$$(Z | m^{(1)}) \leq_{lr} (Z | m^{(2)}),$$

برقرار است)، در نتیجه امید شرطی تحت $m^{(1)}$ نسبت به $m^{(2)}$ بزرگتر خواهد بود. بنابراین نامساوی (۱) برقرار خواهد بود. $\square$

با استفاده از قضیه زیر، برتری سیاست آب‌بندی چندمرحله‌ای نشان داده می‌شود.

قضیه ۳.۳. فرض کنید نرخ خرابی شرطی در رابطه (۲)، برای هر z روی بازه $(0, T]$ نزولی باشد و $T > b + \tau$. در این صورت برای هر سیاست چندمرحله‌ای با $b \in (0, T - \tau]$ نابرابری زیر را داریم

$$P(b | \tau \text{ موفق}) \geq P(\text{غربالگری چندمرحله‌ای} | \tau \text{ موفق}).$$

علاوه بر این، سیاست غربالگری چندمرحله‌ای هر سیاست تک مرحله‌ای با همان زمان کل آب‌بندی b را شکست می‌دهد.

مثال ۴.۳. با استفاده از میانگین‌های پسین بدست آمده در جدول ۱ و دوره مأموریت $\tau = 1$ ، احتمال موفقیت برای سه سیاست در جدول ۳ بدست آمده است.

جدول ۳: مقایسه سه سیاست غربالگری

سیاست	زمان‌های آب‌بندی	آستانه‌های غربالگری	احتمال موفقیت	درصد بهبود
بدون غربالگری	-	-	۴۱۲/۰	-
تک مرحله‌ای	$b = 2$	$m = 4$	۵۷۱/۰	۶/۳۸٪
چندمرحله‌ای	$t_1 = 8/0, t_2 = 4/1, t_3 = 2$	$m_1 = 1, m_2 = 2, m_3 = 4$	۶۹۳/۰	۲/۶۸٪

$$\text{درصد بهبود} = \frac{P(b^*, m^*) - P(b^*)}{P(b^*)} \times 100\%$$

همانطور که در جدول ۳ مشاهده می‌شود سیاست چندمرحله‌ای، احتمال موفقیت بیشتری نسبت به سایر سیاست‌ها به همراه دارد که بیانگر قضیه ۳.۳ است. علاوه بر این، سیاست تک مرحله‌ای احتمال موفقیت مأموریت را در مقایسه با سیاست بدون غربالگری ۶/۳۸٪ افزایش می‌دهد و سیاست چندمرحله‌ای نسبت به سیاست بدون غربالگری ۲/۶۸٪ بهبود می‌یابد.

۴ مدل بهینه‌سازی هزینه

پس از اثبات برتری احتمالاتی سیاست غربالگری چندمرحله‌ای در بخش ۳، اکنون به دیدگاه اقتصادی آن می‌پردازیم. تولیدکنندگان عمدتاً با هدف کمینه‌سازی متوسط هزینه کل برای هر قطعه، اقدام می‌کنند. این هزینه جامع شامل

هزینه‌های عملیاتی آب‌بندی، هزینه تعمیرات مینیمال، هزینه اسقاط قطعات مردود و زیان ناشی از خرابی قطعه در دوره مأموریت می‌شود. در این بخش، متوسط هزینه کل را تحت سیاست چندمرحله‌ای پیشنهادی مدل‌سازی کرده و مسئله بهینه‌سازی مربوطه را فرمول‌بندی خواهیم کرد.

۱.۴ تابع متوسط هزینه کل

فرض کنید

$$\mathbf{t} = (t_1, t_2, \dots, t_k = b)$$

و

$$\mathbf{m} = (m_1, m_2, \dots, m_k)$$

به‌ترتیب بردار زمان‌های بازرسی و آستانه‌های غربالگری باشند، در اینصورت شاخص‌های قبولی که بر اساس آن‌ها مشخص می‌شود یک قطعه برای عبور از مراحل مختلف فرایند تولید، مناسب و قابل قبول است را به‌صورت زیر تعریف می‌کنیم.

رویداد A_i بیانگر عبور قطعه از i مرحله اول است، یعنی

$$A_i = \{M_1 \leq m_1, M_2 \leq m_2, \dots, M_i \leq m_i\},$$

که در آن $M_j = N(t_j)$ است. احتمال عبور از تمام مراحل برابر است با

$$P(\text{عبور}) = P(A_k) = E \left[\prod_{i=1}^k P(M_i \leq m_i | M_{i-1} \leq m_{i-1}, Z) \right].$$

قضیه ۱.۴. فرض کنید C_s هزینه راه‌اندازی و عملیات آب‌بندی، C_m هزینه تعمیرات مینیمال یک قطعه در طول دوره آب‌بندی، C_d هزینه اسقاط در پایان آب‌بندی، L زیان ناشی از شکست مأموریت و K قیمت فروش یک قطعه باشد، در اینصورت متوسط هزینه کل به‌ازای هر قطعه تحت سیاست غربالگری چندمرحله‌ای به‌صورت زیر داده می‌شود

$$\begin{aligned} E[C(\mathbf{t}, \mathbf{m})] &= C_s + C \cdot \sum_{i=1}^k t_i P(A_{i-1} \cap \{M_i > m_i\}) \\ &+ C_m \sum_{i=1}^k E[M_i | A_{i-1} \cap \{M_i > m_i\}] P(A_{i-1} \cap \{M_i > m_i\}) \\ &+ \sum_{i=1}^k C_d P(A_{i-1} \cap \{M_i > m_i\}) \\ &+ (K - L) P(A_k) - L P(A_k) E \left[e^{-Z\lambda((b+\tau)^\beta - b^\beta)} | A_k \right]. \end{aligned}$$

۲.۴ بهینه‌سازی متوسط هزینه

سیاست بهینه آب‌بندی چندمرحله‌ای با حل مسئله بهینه‌سازی زیر بدست می‌آید

$$(\mathbf{t}^*, \mathbf{m}^*) = \arg \min_{\mathbf{t}, \mathbf{m}} E[C(\mathbf{t}, \mathbf{m})].$$

۵ مدل بهینه‌سازی عملکرد

در حالی که مدل هزینه در بخش ۴، بر کارایی اقتصادی تمرکز دارد، بسیاری از تولیدکنندگان (به‌ویژه در صنایع با قابلیت اعتماد بالا مانند هوافضا، تجهیزات پزشکی و خودروهای خودران) احتمال موفقیت دوره مأموریت ثابت τ بدون هیچ خرابی را در اولویت قرار می‌دهند. در این بخش، احتمال موفقیت تحت مدل غربالگری چندمرحله‌ای پیشنهادی بیان و مسئله بهینه‌سازی عملکرد متناظر فرمول‌بندی می‌شود.

۱.۵ احتمال موفقیت تحت سیاست چندمرحله‌ای

پس از عبور از تمام k مرحله غربالگری، قطعه وارد بهره‌برداری میدانی می‌شود. احتمال اینکه قطعه دوره مأموریت τ را بدون خرابی طی کند، برابر است با

$$P(\text{موفقیت}|\mathbf{t}, \mathbf{m}) = E \left[\exp \left(-Z\lambda \left((b + \tau)^\beta - b^\beta \right) \right) \middle| A_k \right],$$

که در آن

$$A_k = \cap_{i=1}^k \{M_i \leq m_i\}.$$

۲.۵ بهینه‌سازی مدل عملکرد

مسئله بهینه‌سازی عملکرد عبارت است از یافتن زمان‌های بازرسی بهینه $\mathbf{t}^*$ و آستانه‌های غربالگری بهینه $\mathbf{m}^*$ که احتمال موفقیت را بیشینه کند یعنی

$$(\mathbf{t}^*, \mathbf{m}^*) = \arg \max_{\mathbf{t}, \mathbf{m}} P(\text{موفقیت}|\mathbf{t}, \mathbf{m}).$$

۶ نتایج عددی و تحلیل حساسیت

در این بخش، پیاده‌سازی و کارایی سیاست آب‌بندی چندمرحله‌ای بیزی پیشنهادی با استفاده از برآوردهای پسین بدست آمده در جدول ۱، نشان داده می‌شود.

تمام محاسبات با دوره مأموریت $\tau = ۱$ و حداکثر تعداد مراحل $k = ۳$ انجام شده است. مقادیر پارامترها به‌صورت زیر در نظر گرفته شده‌اند

$$\{C_s = ۱, C_o = ۸/۰, C_m = ۳/۰, C_d = -۳۰, K = ۱۰۰, L = ۸۰\}.$$

مسائل بهینه‌سازی (کمینه‌سازی هزینه و بیشینه‌سازی عملکرد) با استفاده از الگوریتم ژنتیک همراه با انتگرال‌گیری مونت کارلو بر پایه ۱۰۰۰۰ نمونه پسین از Z حل شدند.

• سیاست بهینه چندمرحله‌ای:

$$k = ۳$$

• زمان‌های بازرسی:

$$\{t_1 = ۷۵/۰, t_2 = ۳۵/۱, t_3 = ۱/۲\}$$

• آستانه‌های غربالگری:

$$\{m_1^* = ۰, m_2^* = ۱, m_3^* = ۳\}$$

جدول ۴: تأثیر سیاست‌های غربالگری بر مدل‌های هزینه و عملکرد

سیاست	متوسط هزینه کل	احتمال موفقیت	درصد کاهش هزینه	درصد بهبود
بدون غربالگری	۳۶/۴۲	۴۱۲/۰	-	-
غربالگری تک مرحله‌ای	۸۵/۳۱	۵۷۱/۰	-	-
غربالگری چندمرحله‌ای	۶۷/۲۴	۷۱۲/۰	۱/۲۹%	۷/۲۴%

$$\text{درصد کاهش هزینه} = \frac{C(b^*, m^*) - C(b^*)}{C(b^*)} \times ۱۰۰\%.$$

در جدول ۴، مشاهده می‌شود که سیاست چندمرحله‌ای پیشنهادی، متوسط هزینه کل را نسبت به سیاست تک مرحله‌ای ۱/۲۹% کاهش و احتمال موفقیت مأموریت را ۷/۲۴% افزایش می‌دهد. این بهبودها در حالی حاصل می‌شود که زمان کل آب‌بندی یکسان و برابر ۱/۲ حفظ شده است.

برای نشان دادن رفتار مقادیر بهینه t^* و m^* نسبت به پارامترهای مدل، حالت‌های مختلف زیر را در نظر می‌گیریم.

- تحلیل حساسیت نسبت به پارمتر τ :

جدول ۵: تأثیر دوره مأموریت τ بر هزینه

دوره مأموریت (τ)	مقدار بهینه t	مقدار بهینه m	متوسط هزینه کل	درصد کاهش هزینه نسبت به تک مرحله‌ای
۵/۰	۶۵/۱	(۰, ۱, ۲)	۹۲/۱۸	۸/۱۹%
۱	۱/۲	(۰, ۱, ۳)	۶۷/۲۴	۵/۲۲%
۵/۱	۴/۲	(۰, ۱, ۴)	۱۴/۳۱	۱/۲۴%
۲	۶۵/۲	(۰, ۲, ۴)	۷۶/۳۸	۳/۲۵%
۳	۸۵/۲	(۰, ۲, ۵)	۸۳/۴۹	۷/۲۶%

در جدول ۵، با افزایش دوره مأموریت τ ، زمان بهینه آب‌بندی و آستانه غربالگری نیز افزایش می‌یابد. علاوه بر این، مشخص شده است که با افزایش τ ، استفاده از سیاست غربالگری چندمرحله‌ای منجر به کاهش ۷/۲۶٪ هزینه‌ها نسبت به سیاست غربالگری تک مرحله‌ای می‌شود. این نتایج نشان‌دهنده برتری قابل توجه سیاست چندمرحله‌ای نسبت به تک مرحله‌ای در بهینه‌سازی هزینه‌ها با طولانی‌تر شدن دوره مأموریت است.

- تحلیل حساسیت نسبت به پارمتر C_s :

جدول ۶: تأثیر هزینه عملیاتی C_s بر هزینه

هزینه عملیاتی (C_s)	مقدار بهینه t	مقدار بهینه m	متوسط هزینه کل	درصد کاهش هزینه نسبت به تک مرحله‌ای
۴/۰	۳۵/۲	(۰, ۱, ۴)	۴۵/۲۱	۷/۱۸%
۸/۰	۱/۲	(۰, ۱, ۳)	۶۷/۲۴	۵/۲۲%
۲/۱	۸۵/۱	(۰, ۱, ۳)	۱۹/۲۸	۹/۲۴%
۶/۱	۷/۱	(۰, ۱, ۲)	۱۴/۳۲	۲/۲۶%
۲	۵۵/۱	(۰, ۰, ۲)	۲۸/۳۶	۱/۲۷%

در جدول ۶، مشاهده می‌شود که افزایش هزینه عملیاتی C_s ، منجر به کاهش مدت زمان آب‌بندی و سختگیرانه‌تر شدن آستانه‌های غربالگری اولیه (به‌ویژه برای $m_1 = 0$) می‌شود. با این حال، اجرای سیاست چندمرحله‌ای موجب کاهش هزینه‌ها نسبت به سیاست تک مرحله‌ای شده و در نتیجه حدود ۱/۲۷٪ صرفه‌جویی اقتصادی را به همراه دارد.

در این مقاله، چارچوب بیزی برای غربالگری چندمرحله‌ای قطعات در دوره آب‌بندی با در نظر گرفتن تعمیر مینیمال، معرفی شده است. هدف اصلی مقاله، تعیین بهینه زمان آب‌بندی و آستانه غربالگری است تا ضمن کمینه‌سازی هزینه‌های کلی، احتمال موفقیت دوره مأموریت به حداکثر برسد. برای این منظور، یک مدل بیزی برای قطعات قابل تعمیر که دارای ناهمگنی پیوسته هستند و تحت غربالگری چندمرحله‌ای قرار می‌گیرند، توسعه یافته است. پارامترهای این مدل از طریق شبیه‌سازی مونت کارلو به روش بیزی برآورد شده‌اند. سپس این برآوردها برای تکمیل و بهینه‌سازی توابع هزینه و عملکرد به کار گرفته شده‌اند. درنهایت با استفاده از محاسبات عددی و تحلیل حساسیت، مقادیر بهینه بدست آمده و برتری سیاست غربالگری چندمرحله‌ای نسبت به رویکردهای موجود به اثبات رسیده است.

مراجع

- [1] Cha, J. H., and Finkelstein, M. (2010). Stochastically ordered subpopulations and optimal burn-in procedure. *IEEE Transactions on Reliability*, **59**(4), 635-643.
- [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 Li, P. (2024). Bayesian analysis of optimal burn-in policy for heterogeneous components with minimal repair. *Quality and Reliability Engineering International*, **40**(5), 2305-2320.
- [4] 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.
- [5] 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.

- [6] 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.
- [7] Suhir, E. (2019). To burn-in, or not to burn-in: That's the question. *Aerospace*, **6**(3), 29.



بهبود قابلیت اعتماد سیستم‌های k از n از طریق روش کاهش

چهنکندی، م^۱

گروه آمار، دانشکده علوم ریاضی و آمار، دانشگاه بیرجند

گروه آمار، دانشکده علوم ریاضی، دانشگاه گیلان

چکیده

تعمیر و نگهداری، افزونگی و کاهش از جمله روش‌هایی هستند که به بهبود عملکرد سیستم کمک می‌کنند. هر کدام از این روش‌ها محدودیت‌هایی دارند که با توجه به مسئله مورد استفاده جایگزین یکدیگر می‌شوند یا ممکن است هم‌زمان مورد استفاده قرار گیرند. این مقاله به طور خاص به مطالعه‌ی روش کاهش در سیستم‌های k از n می‌پردازد و یک عبارت جبری بسته برای ضریب کاهش ارائه می‌شود. بهبود سیستم می‌تواند از طریق کاهش در نرخ خطر یک مولفه یا به صورت هم‌زمان دو مولفه از سیستم صورت پذیرد. نتایج این مقاله برای اولویت‌بندی مولفه‌های سیستم جهت اعمال هر یک از روش‌های افزایش قابلیت اعتماد آن مفید است.

کلمات کلیدی: افزونگی، روش کاهش، عامل هم‌ارزی قابلیت اعتماد.

¹Speaker. mchahkandi@birjand.ac.ir, m.chahkandi@guilan.ac.ir

۱ پیش‌گفتار

انتخاب روش‌های بهبود قابلیت اعتماد سیستم، به عوامل مختلفی مانند محدودیت‌های فضا، هزینه‌ها و سایر قیود بستگی دارد. تعمیر و نگهداری، افزونگی و کاهش از مهم‌ترین این روش‌ها به شمار می‌آیند. در روش افزونگی، با افزودن مولفه‌های اضافی به سیستم اصلی، قابلیت اعتماد سیستم بهبود می‌یابد. این مولفه‌ها می‌توانند به صورت فعال یا آماده به کار به سیستم اضافه شوند. افزایش چشمگیر قابلیت اعتماد سیستم بدون نیاز به تغییر در فناوری ساخت قطعات اصلی از جمله مزیت‌های این روش به شمار می‌رود. در افزونگی فعال، مولفه‌های اضافی به صورت موازی با مولفه‌های اصلی سیستم قرار می‌گیرند و همزمان با مولفه‌های اصلی شروع به کار می‌کنند. بنابراین طول عمر سیستم بر اساس ماکسیمم طول عمر مولفه‌های اصلی و افزونه‌هایشان تعیین می‌شود. در افزونگی آماده به کار، مولفه‌های افزونه تنها پس از شکست مولفه‌های اصلی شروع به کار می‌کنند. افزونگی آماده به کار به سه صورت فعال، داغ و گرم امکان‌پذیر است. برای مطالعه جزئیات بیشتر در این زمینه می‌توان به بارلو و پروشان [۲]، شی و پت [۱۶]، ناندا و حضرا [۱۳]، اریلماز [۶]، جدی و دوست‌پرست [۱۰] و لی و لی [۱۱] مراجعه نمود.

افزایش حجم، هزینه و پیچیدگی مدیریت خرابی سیستم از جمله مشکلات روش‌های افزونگی هستند. اهمیت عوامل هم‌ارزی قابلیت اعتماد^۱ (REF) در توانایی آن‌ها برای تعیین کمی میزان بهبود قابلیت اعتماد نهفته است، به‌گونه‌ای که معیاری روشن برای مهندسين این حوزه فراهم می‌شود تا بتوانند راهبرد مقرون به صرفه را به منظور بهبود قابلیت اعتماد انتخاب نمایند. در روش کاهش، قابلیت اعتماد سیستم از طریق جایگزین کردن مولفه‌های اصلی سیستم با مولفه‌هایی با کیفیت بالاتر و به عبارتی با نرخ شکست کمتر از مولفه‌های اصلی سیستم، افزایش می‌یابد. روش کاهش عموماً بر دو نوع متمایز از عوامل هم‌ارزی قابلیت اعتماد متمرکز است که به نام‌های عامل هم‌ارزی بقای قابلیت اعتماد^۲ (SREF) و عامل هم‌ارزی میانگین قابلیت اعتماد^۳ (MREF) شناخته می‌شوند، که توسط راد [۱۴] و راد [۱۵] معرفی و مورد بحث قرار گرفته‌اند. در روش کاهش براساس عامل تابع بقا، ضریب (ضرایب) کاهش نرخ خطر مولفه‌های سیستم به گونه‌ای تعیین می‌شود که قابلیت اعتماد سیستم مورد مطالعه با قابلیت اعتماد سیستم بهبود یافته به هر یک از روش‌های افزونگی هم‌ارز شود. برای این منظور ابتدا زمانی را می‌یابیم که قابلیت اعتماد سیستم مورد نظر به سطح مطلوب می‌رسد و سپس در آن زمان، ضریب (ضرایب) کاهش نرخ خطر محاسبه می‌شوند. فرض کنید T_A نشان‌دهنده طول عمر سیستمی باشد که با یکی از روش‌های افزون‌سازی بهبود یافته‌است. چنانچه T_B نشان‌دهنده طول عمر سیستمی باشد که با کاهش نرخ خرابی مؤلفه‌های خود بهبود یافته‌است. عامل هم‌ارزی بقای قابلیت اعتماد

Reliability Equivalence Factors^۱
Survival Reliability Equivalence Factors^۲
Mean Reliability Equivalence Factors^۳

(SREF) با برابر قرار دادن توابع بقای سیستمی که از طریق روش کاهش، بهبود یافته و سیستمی که با روش افزونگی ارتقا یافته است، به دست می‌آید. به عبارت دیگر، ρ از حل معادله زیر حاصل می‌شود:

$$\bar{F}_{T_p}(t) = \bar{F}_{T_A}(t) = \omega, \quad 0 < \rho < 1, \quad (1)$$

که در آن ω یک مقدار ثابت مشخص است. برای حل معادله (۳) ابتدا باید مقداری از t مانند t_0 را طوری می‌یابیم که $\bar{F}_{T_A}(t_0) = \omega$. سپس ضریب ρ از حل معادله‌ی $\bar{F}_{T_p}(t_0) = \omega$ به دست می‌آید. چنانچه این معادله به ازای روش‌های مختلف هم‌ارزی حل شود، با توجه به ضرایب کاهش به دست آمده با فرض یکسان بودن هزینه‌ی اعمال روش‌های مختلف، می‌توان روش بهینه برای بهبود قابلیت اعتماد سیستم را پیشنهاد داد. در تحقیقات پیشین عوامل هم‌ارزی قابلیت اعتماد بقا و میانگین به کمک روش‌های عددی برای برخی از سیستم‌های خاص از جمله سری-موازی و موازی-سری مورد بحث قرار گرفته‌است که از آن جمله می‌توان به هو همکاران [۹]، مصطفی [۱۲]، الغازو و همکاران [۱] و الفهیم [۵] اشاره نمود. چهکندی و همکاران [۴] و اطمینان و همکاران [۷] عبارات جبری صریحی برای محاسبه ضریب کاهش نرخ خطر وقتی بهبود تنها روی یک مولفه از سیستم انجام می‌شود و مولفه‌های سیستم مستقل و هم‌توزیعند، ارائه دادند. در این مقاله به منظور بررسی بیشتر برای حالتی که مولفه‌ها هم‌توزیع نیستند و امکان اعمال روش کاهش روی بیش از یک مولفه سیستم وجود دارد، به مطالعه‌ی سیستم‌های k از n پرداخته می‌شود.

۲ بهبود از طریق یک مولفه

فرض کنید تنها امکان بهبود قابلیت اعتماد یکی از مولفه‌های سیستم k از n را داریم. در این صورت علاقه‌مندیم تاثیر بهبود قابلیت اعتماد یک مولفه را بر قابلیت اعتماد سیستم مطالعه کنیم و از طریق محاسبه عامل هم‌ارزی تابع قابلیت اعتماد یا تابع میانگین، روش افزونگی مورد نظر را با روش کاهش هم ارز کنیم. برای این منظور لازم است ابتدا معیار اهمیت بیرنجام معرفی شود.

سیستمی متشکل از n مولفه مستقل با طول عمرهای $X_1, \dots, X_n$ و بردار قابلیت اعتماد $\mathbf{p}(t) = (p_1(t), p_2(t), \dots, p_n(t))$ که در زمان مشخص t_0 برای سادگی آن را با $\mathbf{p} = (p_1, \dots, p_n)$ نمایش می‌دهیم، در نظر بگیرید. فرض کنید تابع قابلیت اعتماد سیستم با $h(\mathbf{p})$ نمایش داده شود. براساس قضیه تجزیه محوری، تابع قابلیت اعتماد یک سیستم منسجم با مولفه‌های مستقل به صورت

$$h(\mathbf{p}) = p_i h(\mathbf{1}_i, \mathbf{p}^{(i)}) + (1 - p_i) h(\mathbf{0}_i, \mathbf{p}^{(i)}),$$

بیان می‌شود که در آن $\mathbf{p}^{(i)} = (p_1, p_2, \dots, p_{i-1}, p_{i+1}, \dots, p_n)$ ، $(\mathbf{1}_i, \mathbf{p}^{(i)}) = (p_1, p_2, \dots, p_{i-1}, 1, p_{i+1}, \dots, p_n)$ و $(\mathbf{0}_i, \mathbf{p}^{(i)}) = (p_1, p_2, \dots, p_{i-1}, 0, p_{i+1}, \dots, p_n)$ براسا تعریف بیرنهام [۳]، اندازه اهمیت مولفه i برابر است با:

$$\begin{aligned} I_B(i) &= \frac{\partial h(\mathbf{p})}{\partial p_i} \\ &= h(\mathbf{1}_i, \mathbf{p}^{(i)}) - h(\mathbf{0}_i, \mathbf{p}^{(i)}). \end{aligned} \quad (۱)$$

همانطور که ملاحظه می‌شود این معیار تفاوت بین قابلیت اعتماد سیستم وقتی مولفه‌ی i ام فعال است و وقتی این مولفه غیرفعال است را می‌دهد. زی و شن [۱۸] معیار اهمیتی معرفی نمودند که برای مطالعه تاثیر بهبود قابلیت اعتماد یکی از مولفه‌های سیستم بر قابلیت اعتماد سیستم مناسب است. به کمک این معیار می‌توان مهم‌ترین مولفه سیستم برای اعمال روش کاهش یا افزونگی را شناسایی کرد. این معیار به شکل

$$I_{XS}(i, \mathbf{p}^{(i)}) = h(p'_i, \mathbf{p}^{(i)}) - h(\mathbf{p}).$$

تعریف می‌شود، که در آن p'_i قابلیت اعتماد مولفه‌ی i ام پس از بهبود است. با استفاده از قضیه تجزیه محوری داریم:

$$\begin{aligned} I_{XS}(i, \mathbf{p}^{(i)}) &= p'_i h(\mathbf{1}_i, \mathbf{p}^{(i)}) + (1 - p'_i) h(\mathbf{0}_i, \mathbf{p}^{(i)}) - p_i h(\mathbf{1}_i, \mathbf{p}^{(i)}) - (1 - p_i) h(\mathbf{0}_i, \mathbf{p}^{(i)}) \\ &= (p'_i - p_i) I_B(i). \end{aligned} \quad (۲)$$

رابطه‌ی (۲) نشان می‌دهد که مناسب‌ترین مولفه برای بهبود قابلیت اعتماد سیستم لزوماً مهم‌ترین مولفه براساس معیار اهمیت بیرنهام نیست. این رابطه الگوی مناسبی برای برقراری ارتباط بین اثر بهبود قابلیت اعتماد مولفه بر قابلیت اعتماد سیستم است. براین اساس، به کمک معیار اهمیت در رابطه‌ی (۲) به محاسبه ضریب کاهش براساس عامل هم ارزی بقا می‌پردازیم. فرض کنید با اعمال روش کاهش روی یک مولفه از سیستم منسجمی بخواهیم قابلیت اعتماد آن را طوری افزایش دهیم که $h(\mathbf{p}') - h(\mathbf{p}) = a$. در ادامه ضریب کاهش را برای سیستم‌های سری، موازی و k از n به شکل صریح به دست می‌آید.

الف) سیستم سری: براساس رابطه (۲) داریم:

$$\begin{aligned} h(p'_i, \mathbf{p}^{(i)}) - h(\mathbf{p}) &= (p'_i - p_i) \prod_{j=1, j \neq i}^n p_j \\ &= \left(\frac{p'_i}{p_i} - 1 \right) h(\mathbf{p}). \end{aligned}$$

حال اگر نرخ خطر مولفه i ام با ضریب کاهش ρ_i بهبود یافته باشد، یعنی $r'_i(t) = \rho_i r_i(t)$ ، که معادل $p'_i = p_i^{\rho_i}$ است و در آن $r_i(\cdot)$ و $r'_i(\cdot)$ به ترتیب نرخ خطر مولفه‌ی i ام قبل و پس از بهبود است. ضریب کاهش از رابطه‌ی

$$\rho_i = \frac{\ln\left(\frac{a}{h(\mathbf{p})} + 1\right)}{\ln(p_i)} + 1 \quad (۳)$$

به دست می‌آید. این رابطه در اولویت‌بندی مولفه‌های سیستم سری جهت اعمال روش کاهش راهگشاست.

گزاره ۱.۲. چنانچه X_i و X_j طول عمر دو مولفه از یک سیستم سری متشکل از مولفه‌های مستقل باشند که $X_i \leq_{st} X_j$ ، آنگاه مولفه i برای اعمال روش کاهش نسبت به مولفه j در اولویت است.

برهان. با توجه به فرض گزاره و تعریف ترتیب تصادفی معمولی داریم $p_i \leq p_j$. بنابراین رابطه‌ی (۳) نتیجه می‌دهد که $p_i \geq \rho_j$. واضح است که مقادیر بزرگتر ضریب کاهش به معنای هزینه کمتر و میزان کاهش کمتر در نرخ خطر مولفه هستند، بنابراین در اولویتند. $\square$

براساس گزاره ۱.۲ ضعیف‌ترین مولفه در سیستم سری برای اعمال روش کاهش، مناسب‌ترین مولفه است.

مثال ۲.۲. سیستم سری با چهار مولفه و بردار قابلیت اعتماد $p = (0.2, 0.4, 0.7, 0.9)$ را در نظر بگیرید. اگر بخواهیم به روش افزونگی فعال طول عمر این سیستم را افزایش دهیم و قابلیت اعتماد مولفه‌ی افزونه برابر $s = 0.5$ باشد، در این صورت مناسب‌ترین مولفه برای اعمال روش کاهش مولفه‌ی اول است. از طرفی $h(p) = 0.504$ و

$$h(p_1, p^{(1)}) = (1 - (1 - p_1)(1 - s)) \prod_{i=2}^4 p_i = 0.1512.$$

به طور مشابه داریم $h(p_2, p^{(2)}) = 0.0882$ ، $h(p_3, p^{(3)}) = 0.0612$ و $h(p_4, p^{(4)}) = 0.0532$. همان‌طور که ملاحظه می‌شود بیشترین افزایش با افزونگی فعال روی مولفه‌ی اول رخ می‌دهد که با نتیجه گزاره ۱.۲ منطبق است. افزونگی فعال روی مولفه‌ی اول معادل است با اینکه روش کاهش روی این مولفه با ضریب $\rho_1 = 0.317$ اعمال شود.

ب) سیستم موازی: میزان بهبود در سیستم موازی با اعمال کاهش روی مولفه‌ی i ام عبارت است از

$$\begin{aligned} a &= (p'_i - p_i) \prod_{j \neq i} (1 - p_j) \\ &= \frac{p'_i - p_i}{1 - p_i} (1 - h(p)) \end{aligned}$$

بنابراین ضریب کاهش از رابطه‌ی

$$\rho_i = \frac{\ln(p_i + \frac{a(1-p_i)}{h(p)})}{\ln(p_i)}. \quad (4)$$

به دست می‌آید.

لم ۳.۲. یک سیستم موازی با مولفه‌های مستقل را در نظر بگیرید که می‌خواهیم با اعمال روش کاهش روی یک مولفه قابلیت اعتماد آن را بهبود دهیم به طوری که $h(p') \leq 2h(p)$ یا به طور معادل $a < h(p)$. اگر $X_1 \leq_{st} X_2 \leq_{st} \dots \leq_{st} X_n$ آنگاه X_n در اولین اولویت و X_1 در آخرین اولویت برای اعمال روش کاهش قرار دارد.

برهان. نشان می‌دهیم رابطه (۴) نسبت به p_i صعودی است. به طور معادل کافی است نشان دهیم تابع $f(p) = \frac{\ln(\alpha + (1-\alpha)p)}{\ln(p)}$ نسبت به p صعودی است که در آن $\alpha = \frac{a}{h(p)}$. با تغییر متغیر $p = e^u$ این تابع را می‌توان به صورت $\frac{g(u)}{u}$ بازنویسی کرد که $g(u) = \ln(\alpha + (1-\alpha)e^u)$. اما $g''(u) = \frac{\alpha(1-\alpha)e^u}{(\alpha + (1-\alpha)e^u)^2} > 0$ بنابراین تابع $g(\cdot)$ محدب است و در نتیجه با توجه به ارتباط تحدب و یکنوایی یک تابع در ریاضی، میتوان نتیجه گرفت که $\frac{g(u)}{u}$ صعودی است که معادل صعودی بودن تابع $f(\cdot)$ نسبت به p است. $\square$

مثال ۴.۲. با در نظر گرفتن بردار قابلیت اعتماد p و مولفه افزونه s در مثال ۲.۲ مقادیر $h(p) = 0.9856$ و $h(p') = 0.9928$ حاصل می‌شود و چنانچه مولفه‌ی افزونه برای هر یک از مولفه‌های ۱ تا ۴ در نظر گرفته شود ضرایب کاهش عبارتند از: $\rho_1 = 0.9821$, $\rho_2 = 0.9881$, $\rho_3 = 0.9912$, و $\rho_4 = 0.9923$. همانطور که ملاحظه می‌شود مولفه‌ی ۴ در اولویت قرار می‌گیرد و این با نتایج لم ۳.۲ سازگار است.

ج) سیستم k از n : سیستمی متشکل از n مولفه‌ی مستقل را در نظر بگیرید که زمانی فعال است که حداقل k مولفه‌ی آن فعال باشند. قابلیت اعتماد سیستم با فرض اینکه مولفه‌ها مستقل‌اند اما هم توزیع نیستند را می‌توان به صورت

$$h_{k,n}(p) = \sum_{\substack{A \subseteq \{1, \dots, n\} \\ |A| \geq k}} \left(\prod_{i \in A} p_i \right) \left(\prod_{j \notin A} (1 - p_j) \right)$$

بیان نمود. لذا میزان بهبود در قابلیت اعتماد سیستم با اعمال روش کاهش روی مولفه‌ی i ام عبارت است از:

$$a = (p'_i - p_i)(h_{k-1,n-1}(p) - h_{k,n-1}(p)). \quad (5)$$

با توجه به رابطه‌ی (۵) به راحتی می‌توان با محاسبات عددی ضریب کاهش را به دست آورد. این رابطه نشان می‌دهد که تغییرات ضریب کاهش مستقل از k است.

گزاره ۵.۲. در یک سیستم k از n حداکثر بهبود برای مولفه‌ی i ام عبارت است از

$$\rho_i \leq \frac{\ln(\frac{a}{b} + p_i)}{\ln(p_i)}, \quad (6)$$

که در آن $b = \min\{\frac{h(p)}{p_i}, \frac{1-h(p)}{1-p_i}\}$.

برهان. اثبات براساس کرانی که برای معیار اهمیت بیرنهام در لم ۱۰۲ زی [۱۷] ارائه شده است، به راحتی به دست می‌آید و از ارائه آن خودداری می‌کنیم. □

مثال ۶.۲. سیستم k از n با بردار قابلیت اعتماد p در مثال‌های ۲.۲ و ۴.۲ را در نظر بگیرید. چنانچه افزونگی فعال روی هر یک از مولفه‌های ۱ تا ۴ این سیستم با مولفه افزونه با قابلیت اعتماد $s = ۰/۵$ انجام شود ضرایب کاهش مستقل از مقدار k عبارتند از: $\rho_1 = ۰/۳۱۷$ ، $\rho_2 = ۰/۳۸۹$ ، $\rho_3 = ۰/۴۵۵$ و $\rho_4 = ۰/۴۸۶$. بنابراین چهارمین مولفه برای افزونگی فعال در اولویت است.

در حالت خاصی که مولفه‌ها مستقل و هم‌توزیع باشند، $h_{k,n}(p) = \sum_{i=k}^n B(n, i; p)$ که $B(i, n; p)$ تابع احتمال توزیع دوجمله‌ای با پارامترهای n و p در نقطه‌ی i است. بنابراین می‌توان شکل بسته‌ای برای ضریب کاهش در حالت استقلال و هم‌توزیعی به دست آورد.

$$\rho_i = \frac{\ln\left(\frac{a}{B(n-1, k-1; p)} + p\right)}{\ln(p)}. \quad (۷)$$

همانطور که در (۷) ملاحظه می‌شود، ضریب کاهش برای همه مولفه‌ها یکسان است. بنابراین در حالت i.i.d برای یک سیستم k از n فرقی نمی‌کند افزونگی یا کاهش را روی چه مولفه‌ای اعمال کنیم. همچنین رابطه‌ی (۷) برای مقایسه روش‌های مختلف افزونگی و هم‌ارزی این روش‌ها با روش کاهش مناسب است. به عنوان مثال وقتی افزونگی فعال روی یک مولفه انجام پذیرد مقدار a برابر $q_i p_i$ خواهد بود که $q_i = 1 - p_i$. برای مقادیر خاص k می‌توان نتایج برای سیستم‌های سری و موازی را به دست آورد.

۳ بهبود از طریق بیش از یک مولفه

گاهی اوقات برای اینکه قابلیت اعتماد سیستم مورد نظر به سطح مطلوب افزایش یابد اعمال روش کاهش یا افزونگی روی تنها یک مولفه کفایت نمی‌کند و نیاز است بیش از یک مولفه مد نظر قرار گیرد. در این بخش نتایج بخش ۲ به بیش از یک مولفه تعمیم داده می‌شود. برای این منظور نیازمند مفهوم اهمیت توام هستیم.

هانگ و لی [۸] معیار اهمیت بیرنهام را به دو مولفه تعمیم دادند و اهمیت توام مولفه‌ی i و j را به صورت

$$JRI(i, j) = h(1_i, 1_j, p^{(ij)}) - h(1_i, \bullet_j, p^{(ij)}) - h(\bullet_i, 1_j, p^{(ij)}) + h(\bullet_i, \bullet_j, p^{(ij)}); i < j, \quad (۱)$$

که در آن $p^{(ij)} = (p_1, \dots, p_{i-1}, p_{i+1}, \dots, p_{j-1}, p_{j+1}, \dots, p_n)$ تعریف کردند.

این معیار میزان اهمیت مولفه‌ی i (j) را نسبت به فعال یا غیرفعال بودن مؤلفه‌ی j (i) می‌سنجد. معیار اهمیت توأم را می‌توان به عنوان میزان تغییر در اهمیت حاشیه‌ای مؤلفه‌ی i (j) وقتی مؤلفه‌ی j (i) از حالت فعال به حالت غیرفعال تغییر می‌کند، تفسیر کرد.

معیار اهمیت توأم (JRI) را میتوان از رابطه‌ی زیر محاسبه نمود.

$$JRI(i, j) = \frac{\partial^2 h(\mathbf{p})}{\partial p_i \partial p_j} = \frac{\partial I_B(i)}{\partial p_j}. \quad (2)$$

حال فرض کنید از طریق کاهش نرخ خطر مولفه‌های i و j بخواهیم به اندازه a مقدار قابلیت اعتماد سیستم را افزایش دهیم. برای محاسبه‌ی ضرایب کاهش، ابتدا معیار اهمیت زی و شن [۱۷] در قضیه‌ی بعد تعمیم داده می‌شود و سپس درباره‌ی سیستم‌های سری، موازی و k از n بحث می‌شود.

قضیه ۱۰۳. میزان افزایش در قابلیت اعتماد سیستم منسجمی متشکل از n مؤلفه‌ی مستقل که بهبود آن از طریق ارتقاء عملکرد مولفه‌های i و j انجام شده‌است، عبارت است از

$$\begin{aligned} a &= h(p'_i, p'_j, \mathbf{p}^{(ij)}) - h(\mathbf{p}) \\ &= (p'_i p'_j - p_i p_j) JRI(i, j) + (p'_i - p_i) I_B(i, \bullet_j) \\ &\quad + (p'_j - p_j) I_B(j, \bullet_i), \end{aligned} \quad (3)$$

که در آن p'_i و p'_j به ترتیب قابلیت اعتماد مولفه‌های i و j پس از بهبود هستند و $I_B(i, \bullet_j)$ ($I_B(i, \bullet_j)$) اهمیت بیرنجام مولفه‌ی i ، وقتی مولفه j غیر فعال (فعال) است را نشان می‌دهد.

برهان. با تعمیم قضیه تجزیه محوری براساس دو مولفه، تابع قابلیت اعتماد سیستم بهبود یافته را می‌توان به صورت

$$\begin{aligned} h(p'_i, p'_j, \mathbf{p}^{(ij)}) &= p'_i p'_j h(\bullet_i, \bullet_j, \mathbf{p}^{(ij)}) + (1 - p'_i)(1 - p'_j) h(\bullet_i, \bullet_j, \mathbf{p}^{(ij)}) \\ &\quad + (1 - p'_i) p'_j h(\bullet_i, \bullet_j, \mathbf{p}^{(ij)}) + p'_i (1 - p'_j) h(\bullet_i, \bullet_j, \mathbf{p}^{(ij)}), \end{aligned} \quad (4)$$

بیان نمود. کافی است نمایش مشابهی برای تابع قابلیت اعتماد سیستم بیان و تفاضل دو تابع محاسبه شود. در نهایت با توجه به ارتباط بین معیار اهمیت توأم و معیار اهمیت حاشیه‌ای بیرنجام، نتیجه به‌دست می‌آید که از ذکر جزئیات خودداری می‌شود. $\square$

الف) سیستم سری: برای سیستم سری رابطه‌ی (۳) به شکل

$$a = (p'_i p'_j - p_i p_j) \prod_{k=1, k \neq i, j}^n p_k$$

$$= \left(\frac{p'_i p'_j}{p_i p_j} - 1 \right) h(\mathbf{p}).$$

ساده می‌شود. فرض کنید ضریب کاهش مولفه‌ی i و j به ترتیب ρ_i و ρ_j باشد به طوری‌که $\rho_i + \rho_j = c$ و $0 \leq c \leq 1$. بنابراین داریم

$$\rho_i = \frac{\ln\left(\frac{a}{h(\mathbf{p})} + 1\right) + \ln(p_i) + (1 - c) \ln(p_j)}{\ln(p_i) - \ln(p_j)}. \quad (5)$$

ب) سیستم موازی: برای سیستم موازی میزان افزایش سیستم پس از ازتقاء عملکرد دو مولفه‌ی آن عبارت است از:

$$a = \prod_{k=1, k \neq i, j}^n (1 - p_k) [(p_i p_j - p'_i p'_j) + (p'_i - p_i) + (p'_j - p_j)]. \quad (6)$$

چنانچه مولفه‌ها هم‌توزیع باشند و ضریب کاهش روی مولفه‌های i و j یکسان باشد، رابطه‌ی (۶) به صورت

$$1 - p^\rho - p^{1-\rho} + p^2 - 2p = \frac{a(1-p)^2}{h(\mathbf{p})}, \quad (7)$$

ساده می‌شود. با فرض $x = p^\rho$ برای یافتن ρ کافی است معادله‌ی $x^2 - 2x + d = 0$ برحسب x حل شود که در آن $d = -p^2 + 2p + \frac{a(1-p)^2}{h(\mathbf{p})}$. بنابراین $\rho = \frac{\ln(1 - \sqrt{1-d})}{\ln(p)}$ یا به طور معادل ضریب کاهش در سیستم موازی عبارت است از:

$$\rho = \frac{\ln\left(1 - \sqrt{1 - \left[\frac{a(1-p)^2}{h(\mathbf{p})} - p^2 + 2p\right]}\right)}{\ln(p)}$$

اگرچه زمانی‌که مولفه‌ها هم‌توزیعند و بحث هزینه مطرح نباشد می‌توان با کاهش روی یک مولفه به میزان دلخواه سیستم را بهبود داد.

ج) سیستم k از n : برای این سیستم رابطه‌ی (۳) پس از کمی محاسبات جبری و ساده‌سازی به شکل

$$a = (p'_i p'_j - p_i p_j) [h_{k-2, n-2}(\mathbf{p}^{(ij)}) - 2h_{k-1, n-2}(\mathbf{p}^{(ij)}) + h_{k, n-2}(\mathbf{p}^{(ij)})]$$

$$+ [(p'_i - p_i) + (p'_j - p_j)] [h_{k-1, n-2}(\mathbf{p}^{(ij)}) - h_{k, n-2}(\mathbf{p}^{(ij)})]. \quad (8)$$

حاصل می‌شود.

قضیه ۲.۳. سیستم k از n با مولفه‌های مستقل و هم‌توزیع را به ازای $۲ \leq k \leq n$ در نظر بگیرید. اگر بخواهیم با کاهش یکسان روی دو مولفه‌ی سیستم قابلیت اعتماد آن را به اندازه a واحد افزایش دهیم، ضریب کاهش از رابطه‌ی

$$\rho = \begin{cases} \frac{\ln \frac{M}{\sqrt[2]{C_3}}}{\ln \frac{k-1}{n-1}} & p = \frac{k-1}{n-1} \\ \frac{\ln \left(\frac{-C_3 p + \sqrt{(C_3 p)^2 + C_2 M p}}{C_2} \right)}{\ln(p)} & p \neq \frac{k-1}{n-1} \end{cases} \quad (9)$$

به دست می‌آید، که در آن $C_1 = p^{k-1} (1-p)^{n-k-1}$ ، $C_2 = \binom{n-2}{k-2} - \binom{n-1}{k-1} p$ ، $C_3 = \binom{n-2}{k-1}$ و $M = \frac{a}{C_1} + p(C_2 + C_3)$. (۲)

برهان. چنانچه همه‌ی مولفه‌ها هم‌توزیع باشند میزان بهبود در قابلیت اعتماد سیستم پس از محاسبات جبری از رابطه‌ی (۸) به شکل زیر ساده می‌شود.

$$a = p^{k-1} (1-p)^{n-k-1} \left[(p^{2\rho-1} - p) \left(\binom{n-2}{k-2} - \binom{n-1}{k-1} p \right) + 2(p^\rho - p) \binom{n-2}{k-1} \right]. \quad (10)$$

با در نظر گرفتن مقادیر C_1 ، C_2 ، C_3 و M ، معادله‌ی (۱۰) به صورت

$$C_2 x^2 + 2C_3 p x - M p = 0.$$

بازنویسی می‌شود. با توجه به اینکه $\binom{n-1}{k-1} = \frac{n-1}{k-1} \binom{n-2}{k-2}$ ، پس به ازای $p = \frac{k-1}{n-1}$ مقدار $C_2 = 0$ و ρ به راحتی از رابطه‌ی $\rho = \frac{\ln \frac{M}{\sqrt[2]{C_3}}}{\ln \frac{k-1}{n-1}}$ به دست می‌آید. چنانچه $p \neq \frac{k-1}{n-1}$ باشد، برای یافتن ρ داریم:

$$\rho = \frac{\ln \left(\frac{-C_3 p \pm \sqrt{(C_3 p)^2 + C_2 M p}}{C_2} \right)}{\ln(p)}. \quad (11)$$

چون به ازای $0 < p < \frac{k-1}{n-1}$ مقدار C_2 بزرگتر از صفر است لذا $-C_3 p - \sqrt{(C_3 p)^2 + C_2 M p}$ قابل قبول نیست. از طرف دیگر به ازای $\frac{k-1}{n-1} < p < 1$ مقدار C_2 کمتر از صفر است و چون $\frac{-C_3 p - \sqrt{(C_3 p)^2 + C_2 M p}}{C_2}$ بزرگتر از یک است پس باز هم $-C_3 p - \sqrt{(C_3 p)^2 + C_2 M p}$ قابل قبول نیست و اثبات قضیه کامل می‌شود.

□

مثال ۳.۳. یک سیستم k از n با مولفه‌های مستقل و هم‌توزیع در نظر بگیرید و فرض کنید بخواهیم با کاهش یکسان روی دو مولفه‌ی سیستم قابلیت اعتماد آن را به اندازه‌ی ۰.۲ افزایش دهیم. جدول ۱ مقادیر ρ را به ازای k ، n و p مختلف نمایش می‌دهد. توجه داریم که مقادیر منفی در جدول حاکی از آن است که در این حالت‌ها با اعمال روش کاهش روی دو مولفه نمی‌توان قابلیت اعتماد سیستم را به اندازه a افزایش داد و باید کاهش روی بیش از دو

مؤلفه انجام شود. همچنین به ازای $n = 5$ ، $k = 2$ و $p = 0.6$ قابلیت اعتماد سیستم برابر 0.9129 است و امکان افزایش آن به اندازه a نیست به همین دلیل این خانه از جدول خالی است.

جدول ۱: ضریب کاهش برای یک سیستم k از n وقتی $a = 0.1$.

p				
$6/0$	$5/0$	$3/0$	k	n
—	$463/0$	$714/0$	۲	۵
$592/0$	$549/0$	$242/0$	۴	
$191/0$	$-0.35/0$	$-553/0$	۵	
$131/0$	$498/0$	$471/0$	۵	۱۰
$451/0$	$361/0$	$-322/0$	۷	
$-366/0$	$-9.6/0$	$-84/1$	۹	

با تعمیم قضیه تجزیه محوری برای بیش از دو مؤلفه، می‌توان اهمیت توأم بیش از دو مؤلفه را تعریف کرد و نتایج این مقاله را به بیش از دو مؤلفه تعمیم داد. همچنین استفاده از مفهوم بردار علامت برای تعمیم نتایج این مقاله به سیستم‌های منسجم می‌تواند مد نظر قرار گیرد.

مراجع

- [1] Alghazo, J.M., Mustafa, M. and El-Faheem, A.A. (2020), Availability Equivalence Analysis for the Simulation of Repairable Bridge Network System, *Complexity*, vol. 2020, Article ID 4907895, 8 pages, <https://doi.org/10.1155/2020/4907895>.
- [2] Barlow, R.E. and Proschan, F. (1975), *Statistical Theory of Reliability and Life Testing*. Holt, Rinehart and Winston, New York.
- [3] Birnbaum Z. W. (1969). On the importance of different components in a multicomponent system. *In Multivariate Analysis II (Edited by P. R. Krishnaiah)*, 581–592. Academic Press, New York.

- [4] Chahkandi, M., Etminan, J. and Khanjari Sadegh, M. (2021), On Equivalence of Reliability in Reduction and Redundancy Methods. *Journal of Statistical Sciences*, 15 (1) ,61-80.
- [5] El-Faheem, A. A. (2025), Upgrading the Availability for Repairable Radar Model through the Weibull Lifetime Distribution. *Journal of the Egyptian Mathematical Society*, 33(1), 127-141.
- [6] Eryilmaz, S. and Erkan, E. (2018), Coherent System with Standby Components, *Applied Stochastic Models in Business and Industry*, **34**, 395–406.
- [7] Etminan, J., Khanjari Sadegh, M., and Chahkandi, M. (2023), A new light on reliability equivalence factors. *Metrika*, 86(6), 605-625.
- [8] Hong, J. S., & Lie, C. H. (1993), Joint reliability-importance of two edges in an undirected network. *IEEE transactions on reliability*, 42(1), 17-23.
- [9] Hu, L., Yue, D. and Zhao, D. (2012), Availability Equivalence Analysis of a Repairable Series-Parallel System, *Mathematical Problems in Engineering*, vol. 2012, Article ID 957537, 15 pages, <https://doi.org/10.1155/2012/957537>.
- [10] Jeddi, H. and Doostparast, M. (2020), Allocation of Redundancies in Systems: A General Dependency-Base Framework, *Annals of Operations Research*, <https://doi.org/10.1007/s10479-020-03795-2>.
- [11] Li, C. and Li, X. (2023), On k-out-of-n systems with homogeneous components and one independent cold standby redundancy, *Statistics and Probability Letters*, 203, 109918.
- [12] Mustafa, A. (2020), Improving the Bridge Structure by Using Linear Failure Rate Distribution, *Journal of Applied Statistics*, **47**, 1423-1438.
- [13] Nanda, A.K. and Hazra, N. K. (2013), Some Results on Active Redundancy at Component Level versus System Level, *Operations Research Letters*, **41**, 241–245.
- [14] Råde, L. (1989), Reliability equivalence, studies in statistical quality control and reliability Mathematical Statistics. *Chalmers University of Technology*, Goteborg.

- [15] Råde, L. (1993), Reliability equivalence. *Microelectronics Reliability*, 33(3), 323-325.
- [16] She, J. and Pecht, M. G. (1992), Reliability of a k -out-of- n Warm-Standby System, *IEEE Transactions on Reliability*, **41**, 72–75
- [17] Xie, M. (1988). A note on the Natvig measure. *Scand. J. Statist.*, 15(3), 221–224.
- [18] Xie, M., and Shen, K. (1989). On ranking of system components with respect to different improvement actions. *Microelectronics Reliability*, 29(2), 159–164.



آزمون اجتماع-اشتراک برای ارزیابی قابلیت اطمینان سیستم‌های سری

چینی‌پرداز، ر^۱ نیکنام، ز^۲ چینی‌پرداز، م^۳

^{۱،۲} گروه آمار، دانشکده علوم ریاضی، دانشگاه شهید چمران اهواز

^۳ گروه مهندسی کامپیوتر، دانشکده مهندسی برق و کامپیوتر، دانشگاه صنعتی جندی شاپور دزفول

چکیده

در تحلیل قابلیت اطمینان سیستم‌های مهندسی، ساختار سیستم و نحوه‌ی ارتباط مؤلفه‌های آن نقش تعیین کننده‌ای در رفتار کلی سیستم ایفا می‌کند. یکی از مهم‌ترین و پرکاربردترین ساختارها در مهندسی قابلیت اطمینان، سیستم سری است که در آن عملکرد صحیح کل سیستم مستلزم عملکرد مناسب تمامی مؤلفه‌های تشکیل دهنده آن است. در چنین سیستم‌هایی خرابی هر یک از مؤلفه‌ها می‌تواند به خرابی کل سیستم منجر شود، از این رو تحلیل آماری رفتار مؤلفه‌ها نقش مهمی در ارزیابی قابلیت اطمینان سیستم دارد. در این مقاله آزمون فرضیه برای پارامترهای عملکرد مؤلفه‌های یک سیستم سری مورد بررسی قرار می‌گیرد. با توجه به ماهیت این ساختار، آزمون فرضیه به صورت مقایسه مینیمم پارامترهای عملکرد با یک مقدار آستانه تعریف می‌شود. در این چارچوب، فرضیه صفر بیانگر وجود حداقل یک مؤلفه با عملکرد نامطلوب است، در حالی که فرضیه مقابل نشان‌دهنده مطلوب بودن

¹ chinipardazr@scu.ac.ir

² Speaker. zahranicknam@gmail.com

³ m.chinipardaz@jsu.ac.ir

عملکرد تمامی مؤلفه‌ها می‌باشد. نشان داده می‌شود که این تعریف به طور مستقیم به چارچوب آزمون اجتماع-اشتراک منجر می‌شود. همچنین توان این آزمون در سیستم‌های سری مورد تحلیل قرار می‌گیرد و در نهایت برای ارزیابی کارایی روش پیشنهادی، مثالی کاربردی از یک سیستم سری با طول عمر نمایی ارائه شده است.

کلمات کلیدی: آزمون اجتماع-اشتراک، توان آزمون، سیستم سری، قابلیت اطمینان.

۱ مقدمه

در بسیاری از مسائل عملی، طراح سیستم برای عملکرد مؤلفه‌ها یک سطح حداقل قابل قبول در نظر می‌گیرد. هدف از تحلیل آماری در این شرایط آن است که بررسی شود آیا تمامی مؤلفه‌ها از این سطح عملکرد بالاتر هستند یا خیر. چنین مسئله‌ای به شکل‌گیری به یک آزمون فرضیه چندپارامتری منجر می‌شود که در آن فرضیه صفر بیانگر وجود حداقل یک مؤلفه با عملکرد نامطلوب است. این ساختار با چارچوب آزمون‌های اجتماع-اشتراک مطابقت دارد. مبانی نظری آزمون‌های اجتماع-اشتراک اولین بار توسط لهن (۱۹۵۲) پایه‌گذاری شد. او نشان داد دستیابی همزمان به تمام ویژگی‌های مطلوب برای یک آزمون (نااریبی، غیرتصادفی بودن و به طور یکنواخت پرتوان‌ترین) در این نوع مسائل غیرممکن است. در ادامه روی ^۱ (۱۹۵۳) اصل اجتماع-اشتراک و اشتراک-اجتماع را به عنوان چارچوبی کلی برای آزمون فرضیه‌های مرکب معرفی کرد و اندکی پس از آن، گلسر ^۲ (۱۹۷۶) اصطلاح آزمون اجتماع-اشتراک را در ادبیات آماری تثبیت کرد. در سال‌های بعد، پژوهشگران متعددی به توسعه ویژگی‌های نظری و کاربردی این آزمون‌ها پرداختند. برگر ^۳ (۱۹۸۹، ۱۹۸۲) نشان داد که آزمون‌های اجتماع-اشتراک از نظر کنترل خطای نوع اول دارای ساختاری پایدار و قابل اعتماد هستند. همچنین مطالعات انجام‌شده توسط ساسابوچی ^۴ (۱۹۸۸، ۱۹۸۰) نشان داد که این آزمون‌ها در مسائل مبتنی بر شرط‌های همزمان، عملکرد مناسبی از نظر توان آماری ارائه می‌کنند. ساختار این مقاله به صورت زیر تنظیم شده است:

در بخش ۲، مدل سیستم سری و مفاهیم مرتبط با قابلیت اطمینان معرفی می‌شوند. در بخش ۳، آزمون اجتماع-اشتراک برای ارزیابی قابلیت اطمینان سیستم‌های سری مطرح شده و نشان داده می‌شود که چگونه ساختار این سیستم‌ها به صورت ساختاری به چنین آزمون‌هایی منجر می‌گردد. در بخش ۴، توان این آزمون مورد بررسی قرار می‌گیرد و در بخش ۵، در قالب یک مثال کاربردی از سیستم سری با طول عمر نمایی کارایی روش پیشنهادی مورد ارزیابی قرار

¹Roy

²Glaser

³Berger

⁴Sasabuchi

داده می‌شود. در بخش ۶، به ارائه نتیجه گیری اختصاص داده شده است.

۲ مدل سیستم سری و مفاهیم قابلیت اطمینان

سیستم‌های سری یکی از اساسی‌ترین ساختارها در نظریه قابلیت اطمینان به شمار می‌آیند. در این نوع سیستم‌ها، عملکرد صحیح کل سیستم مستلزم عملکرد صحیح تمامی مؤلفه‌های تشکیل‌دهنده آن است؛ به عبارت دیگر، خرابی هر یک از اجزا موجب از کار افتادن کل سیستم می‌شود. این ساختار در بسیاری از سیستم‌های مهندسی، از جمله خطوط تولید، سیستم‌های الکترونیکی، شبکه‌های کامپیوتری و سیستم‌های ایمنی مشاهده می‌شود. از این رو، مطالعه رفتار آماری سیستم‌های سری از اهمیت نظری و کاربردی بالایی برخوردار است.

فرض کنید سیستمی متشکل از n مؤلفه‌ی مستقل باشد. برای هر مؤلفه، متغیر تصادفی دوتایی X_i به صورت

$$X_i = \begin{cases} 1 & \text{اگر مؤلفه } i \text{ ام کار کند,} \\ 0 & \text{اگر مؤلفه } i \text{ ام از کار بیفتد,} \end{cases} \quad i = 1, \dots, n$$

تعریف می‌شود. اگر احتمال عملکرد مؤلفه‌ی i ام برابر با

$$p_i = P(X_i = 1) = E(X_i)$$

باشد، آنگاه تابع ساختار سیستم سری به صورت

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

تعریف می‌شود. بنابراین، سیستم تنها زمانی در حالت عملیاتی قرار دارد که همه‌ی مؤلفه‌ها به طور هم‌زمان عمل کنند. در نتیجه، احتمال عملکرد سیستم یا قابلیت اطمینان آن در مدل دوتایی برابر است با

$$P(\varphi(X) = 1) = \prod_{i=1}^n p_i.$$

علاوه بر مدل دوتایی، می‌توان عملکرد مؤلفه‌ها را بر حسب طول عمر نیز مدل‌سازی کرد. فرض کنید T_i طول عمر مؤلفه‌ی i ام باشد. در این صورت، طول عمر کل سیستم سری برابر با مینیمم طول عمر مؤلفه‌ها است، یعنی

$$T_{\text{system}} = \min\{T_1, \dots, T_n\}.$$

اگر $T_1, \dots, T_n$ مستقل باشند، تابع قابلیت اطمینان سیستم سری در زمان t به صورت

$$R_{\text{system}}(t) = P(T_{\text{system}} > t) = \prod_{i=1}^n P(T_i > t) = \prod_{i=1}^n R_i(t)$$

به دست می‌آید که در آن $R_i(t) = P(T_i > t)$ تابع قابلیت اطمینان مؤلفه‌ی i ام است.

نکته‌ی اساسی در سیستم‌های سری آن است که رفتار ضعیف‌ترین مؤلفه نقش تعیین‌کننده‌ای در عملکرد کلی سیستم دارد. در مدل‌های طول عمر، این مؤلفه همان مؤلفه‌ای است که دارای بیشترین نرخ خرابی یا کمترین میانگین زمان تا خرابی است. به طور مشابه، در مدل دوتایی نیز مؤلفه‌ای که کمترین احتمال عملکرد را دارد، بیشترین تأثیر را بر قابلیت اطمینان کل سیستم خواهد داشت. این ویژگی سبب می‌شود که در تحلیل آماری سیستم‌های سری، تمرکز اصلی بر روی حداقل سطح عملکرد مؤلفه‌ها قرار گیرد.

در بسیاری از کاربردهای مهندسی، برای عملکرد مؤلفه‌ها یک آستانه‌ی حداقل قابل قبول در نظر گرفته می‌شود. این آستانه ممکن است بر حسب احتمال عملکرد، قابلیت اطمینان در زمان مشخص، میانگین طول عمر، یا یک پارامتر مدل طول عمر تعریف شود. هدف از تعیین این آستانه، ایجاد معیاری برای مقایسه عملکرد اجزای سیستم و شناسایی مؤلفه‌های نامطلوب است. در نتیجه، پرسش اصلی در تحلیل آماری این سیستم‌ها این است که آیا تمامی مؤلفه‌ها از سطح عملکرد قابل قبول بالاتر هستند یا خیر.

۳ آزمون اجتماع-اشتراک برای سیستم‌های سری

در این بخش ابتدا به معرفی آزمون اجتماع-اشتراک پرداخته می‌شود. سپس برای روشن‌تر شدن ساختار آزمون‌های اجتماع-اشتراک فرضیه‌ها در موقعیت‌های مختلف آماری و مهندسی آزمون‌های اجتماع-اشتراک در سیستم‌های سری مورد بررسی قرار می‌گیرند.

۱.۳ آزمون‌های اجتماع-اشتراک (IUT)

در برخی مسائل آماری فرضیه صفر به صورت مجموعه‌ای از فرضیه‌ها بیان می‌شود. این نوع مسائل منجر به آزمون‌های موسوم به اجتماع-اشتراک^۵ می‌شوند. ساختار کلی این آزمون‌ها به صورت:

^۵Intersection-Union

$$H_0 : \theta \in \bigcup_{i=1}^p \Theta_i, \quad \text{مقابل} \quad H_1 : \theta \in \left(\bigcup_{i=1}^p \Theta_i\right)^c = \bigcap_{i=1}^p \Theta_i^c, \quad (1)$$

تعریف می‌شود. در این چارچوب، فرضیه صفر زمانی برقرار است که حداقل یکی از فرضیه‌های جزئی صحیح باشد، در حالی که فرضیه مقابل تنها زمانی برقرار است که تمام فرضیه‌های جزئی نقض شوند. حال، فرض کنید برای هر فرضیه جزئی

$$H_{i0} : \theta \in \Theta_i, \quad \text{در مقابل} \quad H_{i1} : \theta \in \Theta_i^c,$$

آماره آزمون $T_i(X)$ با ناحیه رد R_i تعریف شده باشد، به طوری که $\{X : T_i(X) \in R_i\}$ آزمون متناظر باشد. بر اساس اصل IUT ، فرضیه صفر کلی H_0 رد می‌شود اگر و تنها اگر تمام فرضیه‌های جزئی رد شوند. بنابراین ناحیه رد آزمون IUT برابر است با:

$$R_{IUT} = \bigcap_{i=1}^p \{X : T_i(X) \in R_i\},$$

به ویژه، اگر ناحیه رد هر آزمون جزئی به صورت $\{X : T_i(X) \in R_i\}$ تعریف شده باشد، ناحیه رد آزمون IUT به صورت فشرده زیر نوشته می‌شود:

$$R_{IUT} = \bigcap_{i=1}^p \{X : T_i(X) > c\} = \left\{X : \inf_{1 \leq i \leq p} T_i(X) > c\right\},$$

این رابطه نشان می‌دهد که در آزمون‌های IUT ، تصمیم‌گیری بر پایه مینیمم آماره‌های آزمون جزئی صورت می‌گیرد، بر این اساس، شرط رد کردن H_0 این است که نمونه X در تمام این نواحی قرار گیرد. این بدان معناست که باید تمام نامساوی‌های

$$T_1(X) > c, T_2(X), \dots, T_p(X) > c,$$

به طور همزمان برقرار باشند این شرط به طور طبیعی معادل آن است که کمترین مقدار در میان تمام آماره‌های آزمون، از مقدار بحرانی c بزرگتر باشد. با توجه به ساختار سیستم‌های سری و نقش تعیین‌کننده مؤلفه‌ی بحرانی، مسئله‌ی آزمون آماری در این سیستم‌ها به طور طبیعی بر مبنای مینیمم پارامترهای عملکرد مؤلفه‌ها انجام می‌گیرد. از آنجا که وجود تنها یک مؤلفه‌ی نامطلوب برای ناکافی بودن عملکرد کل سیستم کافی است، فرضیه‌ی صفر معمولاً متناظر با وجود حداقل یک مؤلفه‌ی ضعیف در نظر گرفته می‌شود.

فرض کنید $\theta_1, \dots, \theta_p$ پارامترهای عملکرد مربوط به p مؤلفه‌ی سیستم سری باشند و θ_0 آستانه‌ی حداقل قابل قبول عملکرد باشد که از سوی طراح، استاندارد فنی، یا معیار مهندسی مشخص شده است. در این صورت، مسئله‌ی آزمون فرضیه برای بررسی کفایت عملکرد کل سیستم به صورت زیر تعریف می‌شود:

$$H. : \min_{1 \leq i \leq p} \theta_i \leq \theta., \quad H_1 : \min_{1 \leq i \leq p} \theta_i > \theta.. \quad (2)$$

فرضیه‌ی صفر بیان می‌کند که حداقل یک مؤلفه دارای عملکردی کمتر یا مساوی با سطح قابل قبول است، در حالی که فرضیه‌ی مقابل بیانگر آن است که تمامی مؤلفه‌ها عملکردی بهتر از آستانه‌ی تعیین شده دارند. این صورت‌بندی دقیقاً با منطق سیستم سری سازگار است؛ زیرا پذیرش عملکرد مطلوب کل سیستم تنها در صورتی امکان‌پذیر است که عملکرد همه‌ی اجزا رضایت‌بخش باشد.

آزمون فرضیه (۲) را می‌توان به صورت اجتماع و اشتراک فرضیه‌های جزئی نیز بازنویسی کرد.

$$\begin{aligned} H. : \min_{1 \leq i \leq p} \theta_i \leq \theta. &\Rightarrow \theta_1 \leq \theta. \text{ یا } \theta_2 \leq \theta. \text{ یا } \dots \text{ یا } \theta_p \leq \theta. \\ &\equiv \bigcup_{i=1}^p \{\theta_i \leq \theta.\}, \\ H_1 : \min_{1 \leq i \leq p} \theta_i > \theta. &\Rightarrow \theta_1 > \theta. \text{ و } \theta_2 > \theta. \text{ و } \dots \text{ و } \theta_p > \theta. \\ &\equiv \bigcap_{i=1}^p \{\theta_i > \theta.\}, \end{aligned}$$

بنابراین، رد فرضیه‌ی صفر مستلزم آن است که شواهد آماری کافی برای برتری هر یک از پارامترهای θ_i نسبت به $\theta.$ به دست آید. این ویژگی، مسئله را به طور طبیعی در چارچوب آزمون‌های اجتماع-اشتراک قرار می‌دهد. در واقع، اگر برای هر مؤلفه آزمون جزئی

$$H_{i.} : \theta_i \leq \theta. \quad \text{در مقابل} \quad H_{i1} : \theta_i > \theta. \quad (3)$$

در نظر گرفته شود، آنگاه فرضیه‌ی کلی صفر اجتماع فرضیه‌های $H_{i.}$ و فرضیه‌ی کلی مقابل اشتراک فرضیه‌های H_{i1} خواهد بود. از این رو، یک آزمون اجتماع-اشتراک زمانی فرضیه‌ی صفر کلی را رد می‌کند که تمامی فرضیه‌های جزئی صفر به طور مناسب رد شوند. این قاعده تصمیم‌گیری کاملاً منطبق با منطق مهندسی سیستم‌های سری است، زیرا حتی باقی ماندن تردید آماری درباره‌ی یک مؤلفه نیز مانع از تأیید عملکرد مطلوب کل سیستم می‌شود. فرض کنید برای هر پارامتر θ_i آماره‌ی آزمون T_i در دسترس باشد که برای آزمون (۳) به کار می‌رود. برای هر یک از این آزمون‌ها ناحیه‌ی بحرانی به صورت $R_i = \{T_i > c_i\}$ تعریف می‌شود که در آن c_i مقدار بحرانی متناظر با سطح معنی‌داری آزمون است.

در چارچوب آزمون اجتماع-اشتراک، ناحیه‌ی رد فرضیه‌ی کلی صفر برابر با اشتراک نواحی رد آزمون‌های جزئی است؛ یعنی

$$R = \bigcap_{i=1}^p R_i.$$

به عبارت دیگر، فرضیه‌ی صفر H_0 تنها زمانی رد می‌شود که تمام آزمون‌های جزئی H_{i0} رد شوند. بنابراین، فرضیه‌ی H_0 رد می‌شود اگر و تنها اگر برای تمامی $i = 1, \dots, p$ داشته باشیم

$$T_i > c_i.$$

این قاعده تصمیم‌گیری با منطق مهندسی سیستم‌های سری کاملاً سازگار است. از آنجا که عملکرد کل سیستم وابسته به عملکرد صحیح تمامی مؤلفه‌ها است، نتیجه‌گیری درباره‌ی مطلوب بودن عملکرد سیستم تنها زمانی امکان‌پذیر است که شواهد آماری کافی برای مناسب بودن عملکرد تک‌تک مؤلفه‌ها وجود داشته باشد.

۴ تحلیل توان آزمون اجتماع-اشتراک در سیستم‌های سری

در این بخش، ابتدا توان آزمون برای هر مؤلفه به صورت مستقل بررسی می‌شود و سپس با استفاده از ساختار آزمون اجتماع-اشتراک، توان کلی سیستم سری به دست می‌آید. از آنجا که عملکرد سیستم‌های سری به عملکرد تمامی مؤلفه‌ها وابسته است، تحلیل توان در سطح هر مؤلفه مبنای محاسبه توان کل سیستم خواهد بود.

۱.۴ توان آزمون برای یک مؤلفه منفرد

فرض کنید طول عمر نمونه‌های مربوط به مؤلفه i ($i = 1, \dots, k$) دارای توزیع نمایی با پارامتر میانگین θ_i است. برای هر مؤلفه، آزمون فرضیه زیر در نظر گرفته می‌شود:

$$H_{i0} : \theta_i \leq \theta_0 \quad \text{در مقابل} \quad H_{i1} : \theta_i > \theta_0.$$

با داشتن n مشاهده مستقل $T_{i1}, \dots, T_{in}$ برای مؤلفه i ، آماره آزمون به صورت زیر تعریف می‌شود:

$$X_i = \frac{2 \sum_{j=1}^n T_{ij}}{\theta_0}.$$

تحت فرضیه صفر ($\theta_i = \theta_0$)، آماره X_i از توزیع کای-دو با $2n$ درجه آزادی پیروی می‌کند ($X_i \sim \chi^2_{2n}$). بنابراین، مقدار بحرانی آزمون در سطح معنی‌داری α برابر است با:

$$c = \chi^2_{1-\alpha, \nu_n}$$

تحت فرضیه مقابل (H_{i1}) ، یعنی زمانی که میانگین واقعی طول عمر برابر $\theta_i > \theta_0$ باشد، توزیع آماره به صورت $X_i \sim \frac{\theta_i}{\theta_0} \chi^2_{\nu_n}$ تغییر می‌یابد. در نتیجه، توان آزمون برای مؤلفه i (π_i) برابر است با:

$$\pi_i(\theta_i) = \Pr_{\theta_i}(X_i > c) = \Pr\left(\frac{\theta_i}{\theta_0} \chi^2_{\nu_n} > c\right) = \Pr\left(\chi^2_{\nu_n} > c \frac{\theta_0}{\theta_i}\right)$$

که بر اساس تابع توزیع تجمعی (F) کای-دو به صورت زیر بیان می‌شود:

$$\pi_i(\theta_i) = 1 - F_{\chi^2_{\nu_n}}\left(c \frac{\theta_0}{\theta_i}\right)$$

۲.۴ توان آزمون IUT برای سیستم سری

در یک سیستم سری، فرضیه صفر کلی (H_0) بیان می‌کند که حداقل یکی از مؤلفه‌ها استاندارد لازم را برآورده نکرده است. این فرضیه زمانی رد می‌شود که برای تمامی مؤلفه‌ها، فرضیات صفر متناظر رد شوند:

$$H_0 : \min(\theta_1, \dots, \theta_k) \leq \theta_0,$$

$$H_0 \text{ rejected} \iff \forall i : X_i > c,$$

با فرض استقلال عملکرد مؤلفه‌ها، توان کل آزمون IUT برابر با حاصل ضرب توان تک تک مؤلفه‌ها خواهد بود:

$$\pi(\theta_1, \dots, \theta_k) = \prod_{i=1}^k \pi_i(\theta_i),$$

با جایگذاری رابطه توان هر مؤلفه، فرمول نهایی توان سیستم به دست می‌آید:

$$\pi(\theta_1, \dots, \theta_k) = \prod_{i=1}^k \left[1 - F_{\chi^2_{\nu_n}}\left(\chi^2_{1-\alpha, \nu_n} \frac{\theta_0}{\theta_i}\right) \right],$$

به‌طور خاص، برای سیستم سه‌مؤلفه‌ای ($k = 3$) با پارامترهای $\theta_1, \theta_2, \theta_3$ ، توان آزمون عبارت است از:

$$\pi(\theta_1, \theta_2, \theta_3) = \left(1 - F_{\chi^2_n} \left(c \frac{\theta_1}{\theta_2} \right)\right) \times \left(1 - F_{\chi^2_n} \left(c \frac{\theta_1}{\theta_3} \right)\right) \times \left(1 - F_{\chi^2_n} \left(c \frac{\theta_2}{\theta_3} \right)\right),$$

فرمول توان را می توان بر پایه سه اصل آماری زیر تفسیر کرد:

- ویژگی توزیع نمایی: بر اساس ویژگی جمع پذیری توزیع گاما، مجموع n متغیر تصادفی نمایی با میانگین θ_i دارای توزیع $\text{Gamma}(n, \theta_i)$ است. تبدیل این توزیع به کای-دو با ضریب $2/\theta_i$ یکی از روش های کلاسیک در آمار مهندسی محسوب می شود.

- ساختار آزمون اجتماع-اشتراک: آزمون اجتماع-اشتراک در سیستم های سری دارای ساختاری محافظه کارانه است. از آنجا که ناحیه رد کل سیستم برابر اشتراک نواحی رد مؤلفه ها است ($R = \cap R_i$)، توان سیستم همواره کوچک تر یا مساوی توان ضعیف ترین مؤلفه خواهد بود.

- اثر مؤلفه ضعیف بر عملکرد سیستم: فرمول توان نشان می دهد که اگر حتی یکی از مؤلفه ها دارای میانگین طول عمر نزدیک به آستانه ($\theta_i \approx \theta_0$) باشد، عبارت متناظر آن در حاصل ضرب به سمت صفر میل کرده و توان کل سیستم را به طور قابل توجهی کاهش می دهد. این نتیجه با نحوه عملکرد سیستم های سری مطابقت دارد. در نتیجه، توان آزمون IUT ابزاری مناسب برای ارزیابی کیفیت کلی سیستم محسوب می شود و نشان می دهد که دستیابی به توان آزمون بالا در سیستم های سری مستلزم آن است که تمامی مؤلفه ها دارای میانگین طول عمر به طور معناداری بزرگ تر از حد استاندارد باشند.

۵ مثال کاربردی: ارزیابی قابلیت اطمینان یک سیستم سری با طول عمر نمایی

برای نشان دادن کاربرد آزمون اجتماع-اشتراک در ارزیابی قابلیت اطمینان سیستم های سری، یک مثال مبتنی بر مدل طول عمر نمایی در نظر می گیریم. توزیع نمایی یکی از رایج ترین مدل ها در تحلیل قابلیت اطمینان است و به ویژه در شرایطی که نرخ خرابی مؤلفه ها ثابت فرض می شود کاربرد گسترده ای دارد. فرض می شود سیستمی سری متشکل از p مؤلفه وجود دارد که طول عمر مؤلفه i ام با متغیر تصادفی T_i نمایش داده می شود. همچنین، در نظر گرفته می شود که

$$T_i \sim \text{Exp}(\lambda_i), \quad i = 1, \dots, p,$$

که در آن λ_i نرخ خرابی مؤلفه i ام است. در این مدل، امید ریاضی طول عمر مؤلفه برابر با

$$E(T_i) = \frac{1}{\lambda_i},$$

می‌باشد. از آنجا که سیستم دارای ساختار سری است، طول عمر کل سیستم برابر با

$$T_{\text{system}} = \min\{T_1, \dots, T_p\},$$

خواهد بود.

فرض کنید هدف آن باشد که بررسی شود آیا تمامی مؤلفه‌ها از نظر عملکرد دارای سطح قابل قبولی هستند یا خیر. برای این منظور، یک مقدار آستانه θ_0 برای پارامتر عملکرد در نظر گرفته می‌شود. در اینجا پارامتر عملکرد را می‌توان میانگین طول عمر مؤلفه‌ها در نظر گرفت، یعنی

$$\theta_i = \frac{1}{\lambda_i}.$$

بنابراین آزمون فرضیه به صورت زیر بیان می‌شود:

$$H_0 : \min_{1 \leq i \leq p} \theta_i \leq \theta_0, \quad H_1 : \min_{1 \leq i \leq p} \theta_i > \theta_0.$$

فرضیه صفر بیان می‌کند که حداقل یک مؤلفه دارای میانگین طول عمر کمتر یا مساوی با مقدار قابل قبول θ_0 است، در حالی که فرضیه مقابل بیانگر آن است که همه مؤلفه‌ها دارای عملکردی بهتر از سطح تعیین شده هستند. برای هر مؤلفه، فرض کنید نمونه‌ای شامل n_i مشاهده از طول عمر آن در دسترس باشد:

$$T_{i1}, T_{i2}, \dots, T_{in_i}.$$

در این صورت برآوردگر حداکثر درستنمایی برای میانگین طول عمر مؤلفه i برابر است با

$$\hat{\theta}_i = \bar{T}_i = \frac{1}{n_i} \sum_{j=1}^{n_i} T_{ij}.$$

برای آزمون فرضیه جزئی

$$H_{i0} : \theta_i \leq \theta_0 \quad \text{در مقابل} \quad H_{i1} : \theta_i > \theta_0,$$

در این حالت، آماره آزمون بر اساس مجموع مشاهدات نمونه تعریف می‌شود. از آنجا که مجموع متغیرهای نمایی دارای توزیع گاما است، آماره $S_i = \sum_{j=1}^{n_i} T_{ij}$ دارای توزیع گاما با پارامترهای n_i و λ_i است. بنابراین آزمون مناسب برای هر مؤلفه را می‌توان بر اساس مقدار S_i تعریف کرد.

در چارچوب آزمون اجتماع-اشتراک، فرضیه‌ی کلی صفر تنها زمانی رد می‌شود که تمامی آزمون‌های جزئی مربوط به مؤلفه‌ها رد شوند. فرضیه H_0 رد می‌شود اگر برای تمامی مؤلفه‌ها شواهد آماری کافی وجود داشته باشد که نشان دهد $\theta_i > \theta_0$. این بدان معناست که میانگین طول عمر همه مؤلفه‌ها از سطح قابل قبول تعیین شده بزرگ‌تر است. در غیر این صورت، حتی اگر تنها درباره یکی از مؤلفه‌ها نتوان عملکرد مطلوب را تأیید کرد، فرضیه صفر رد نخواهد شد.

این مثال نشان می‌دهد که چگونه چارچوب آزمون اجتماع-اشتراک به طور طبیعی با ساختار سیستم‌های سری سازگار است. از آنجا که خرابی یک مؤلفه برای از کار افتادن کل سیستم کافی است، ارزیابی مطلوب بودن عملکرد سیستم مستلزم تأیید آماری عملکرد مناسب تمامی مؤلفه‌ها خواهد بود. برای نشان دادن کاربرد عملی آزمون اجتماع-اشتراک، یک سیستم سری شامل سه مؤلفه در نظر گرفته می‌شود. در یک سیستم سری، خرابی هر یک از مؤلفه‌ها منجر به از کار افتادن کل سیستم می‌شود. بنابراین لازم است هر یک از مؤلفه‌ها حداقل سطح مشخصی از قابلیت اطمینان را داشته باشند. فرض می‌شود طول عمر هر مؤلفه دارای توزیع نمایی با میانگین θ_i است. هدف این است که بررسی شود آیا میانگین طول عمر واقعی هر سه مؤلفه از یک حداقل مقدار استاندارد بیشتر است یا خیر. در این مثال، حداقل میانگین طول عمر قابل قبول برای هر مؤلفه برابر با $\theta_0 = 100,000$ در نظر گرفته می‌شود. بنابراین آزمون فرضیه به صورت زیر تعریف می‌شود:

$$H_0 : \min(\theta_1, \theta_2, \theta_3) \leq 100,000$$

در مقابل

$$H_1 : \min(\theta_1, \theta_2, \theta_3) > 100,000$$

به عبارت دیگر، فرضیه صفر بیان می‌کند که حداقل یکی از مؤلفه‌ها دارای میانگین طول عمر کمتر یا مساوی مقدار استاندارد است، در حالی که فرضیه مقابل بیان می‌کند هر سه مؤلفه دارای میانگین طول عمر بزرگ‌تر از مقدار استاندارد هستند.

برای بررسی عملکرد آزمون، از روش شبیه‌سازی مونت‌کارلو استفاده می‌شود. در این روش فرض می‌شود طول عمر مؤلفه‌ها از توزیع نمایی با میانگین‌های واقعی $\theta_1, \theta_2, \theta_3$ پیروی می‌کند. سپس در هر تکرار شبیه‌سازی، تعداد $n = 10$ مشاهده از هر توزیع تولید شده و میانگین‌های نمونه‌ای $\bar{T}_1, \bar{T}_2, \bar{T}_3$ محاسبه می‌شوند. این فرآیند در تعداد زیادی تکرار (برای مثال 10,000 بار) انجام می‌شود تا رفتار آماره آزمون بررسی شود.

در یکی از نمونه‌های حاصل از این فرآیند شبیه‌سازی، میانگین‌های نمونه‌ای زیر به دست آمده‌اند:

$$\bar{T}_1 = 160,000, \quad \bar{T}_2 = 190,000, \quad \bar{T}_3 = 170,000$$

در مدل نمایی، آماره آزمون برای هر مؤلفه به صورت زیر تعریف می شود:

$$\frac{2n\bar{T}_i}{\theta} \sim \chi^2_{2n}$$

با توجه به اینکه $n = 10$ ، آماره آزمون برای هر مؤلفه به صورت زیر محاسبه می شود.
برای مؤلفه اول:

$$\frac{2(10)(160,000)}{100,000} = 32$$

برای مؤلفه دوم:

$$\frac{2(10)(190,000)}{100,000} = 38$$

برای مؤلفه سوم:

$$\frac{2(10)(170,000)}{100,000} = 34$$

در سطح معنی داری $\alpha = 0.05$ مقدار بحرانی توزیع کای دو با درجه آزادی برابر است با $\chi^2_{0.95, 2} = 5.99$ از آنجا که

$$32 > 5.99, \quad 38 > 5.99, \quad 34 > 5.99$$

هر سه فرضیه صفر جزئی رد می شوند. بنابراین در چارچوب آزمون اجتماع-اشتراک، فرضیه صفر کلی نیز رد می شود. در نتیجه می توان نتیجه گرفت که میانگین طول عمر هر سه مؤلفه از حداقل مقدار استاندارد تعیین شده بیشتر است و سیستم سری مورد نظر از نظر قابلیت اطمینان در سطح قابل قبولی قرار دارد. نتایج حاصل از ارزیابی قابلیت اطمینان و تحلیل آماری مثال کاربردی در شکل (۱) ارائه شده است. این شکل جنبه های مختلف تحلیل قابلیت اطمینان شامل ساختار سیستم، توابع قابلیت اطمینان و نتایج آزمون های فرضیه را نمایش می دهد.

قسمت (الف) بیانگر ساختار منطقی سیستم سری سه مؤلفه ای است. در این ساختار، عملکرد صحیح سیستم منوط به عملکرد هم زمان تمامی مؤلفه ها بوده و خرابی هر مؤلفه موجب از کار افتادن کل سیستم می شود.

قسمت (ب) توابع قابلیت اطمینان $R(t)$ مؤلفه ها و سیستم را نشان می دهد. همان گونه که مشاهده می شود، منحنی

مربوط به سیستم (خط بنفش) همواره پایین‌تر از منحنی مؤلفه‌ها قرار دارد؛ این امر بیانگر آن است که در سیستم‌های سری، قابلیت اطمینان کل سیستم همواره کوچک‌تر یا در بهترین حالت برابر با قابلیت اطمینان ضعیف‌ترین مؤلفه خواهد بود.

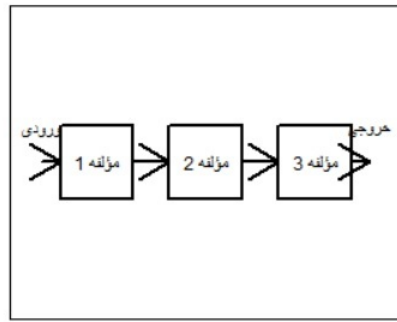
قسمت (ج) مقادیر احتمال (p - values) مربوط به آزمون‌های جزئی را نشان می‌دهد. مقادیر به‌دست‌آمده برای سه مؤلفه (۰/۰۴۳۳، ۰/۰۰۸۹ و ۰/۰۲۶۱) همگی کمتر از سطح معنی‌داری $\alpha = ۰/۰۵$ هستند. بر اساس چارچوب آزمون اجتماع-اشتراک، از آنجا که تمامی آزمون‌های جزئی در سطح تعیین‌شده معنادار شده‌اند، فرضیه صفر کلی مبنی بر ناکافی بودن قابلیت اطمینان سیستم رد می‌شود.

قسمت (د) توزیع کای-دو با ۲۰ درجه آزادی و ناحیه بحرانی (بخش سایه‌خورده در انتهای سمت راست) را نمایش می‌دهد. قرار گرفتن آماره‌های آزمون در این ناحیه حاکی از رد فرضیه صفر در سطح اطمینان ۹۵٪ است.

قسمت (ه) تابع چگالی احتمال خرابی $f(t)$ مؤلفه‌ها را نشان می‌دهد. نزولی بودن این منحنی‌ها با خاصیت نرخ خرابی ثابت در توزیع نمایی و الگوی عملکرد اولیه سیستم سازگار است.

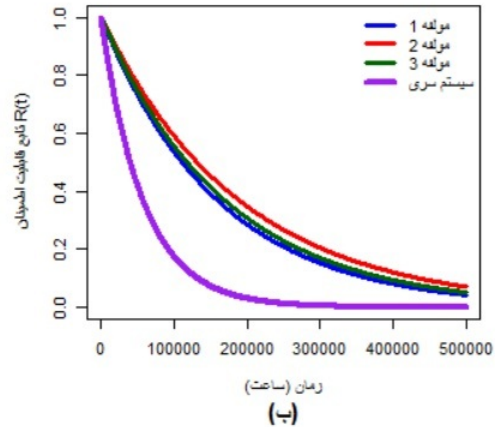
قسمت (و) مقایسه میانگین طول عمر برآوردشده ($\bar{T}_i$) با مقدار آستانه استاندارد $\theta_0 = ۱۰۰,۰۰۰$ را نمایش می‌دهد. قرارگیری ستون‌ها در سطحی بالاتر از خط‌چین قرمز بیانگر آن است که تمامی مؤلفه‌ها به‌طور مستقل معیار حداقل عملکرد تعیین‌شده را تأمین کرده‌اند.

سیستم سری سه‌مؤلفه‌ای



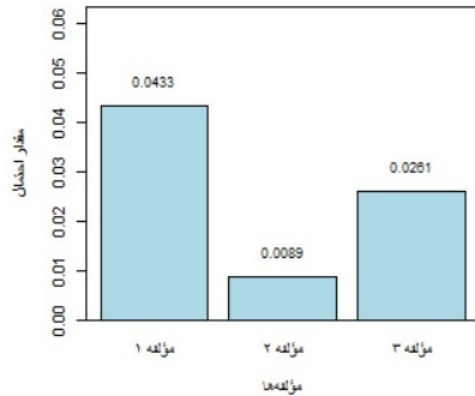
(الف)

تابع قابلیت اطمینان مؤلفه‌ها و سیستم سری



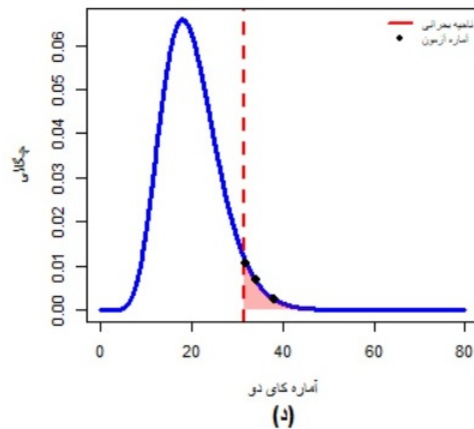
(ب)

مقادیر احتمال آزمون



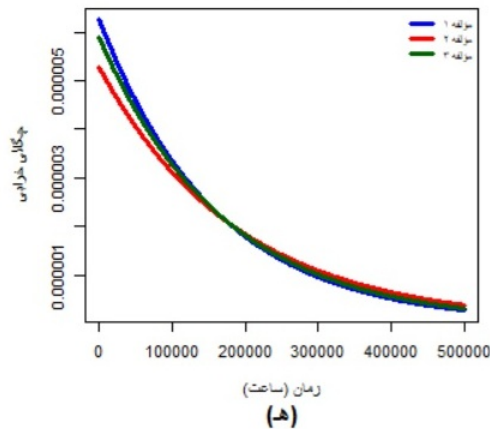
(ج)

توزیع کای دو و ناحیه بحرانی



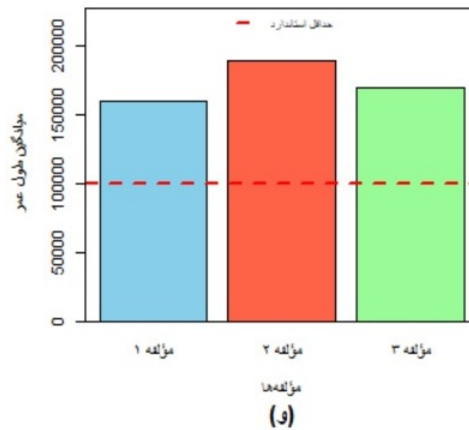
(د)

تابع چگالی خرابی مؤلفه‌ها



(ه)

مقایسه میانگین طول عمر

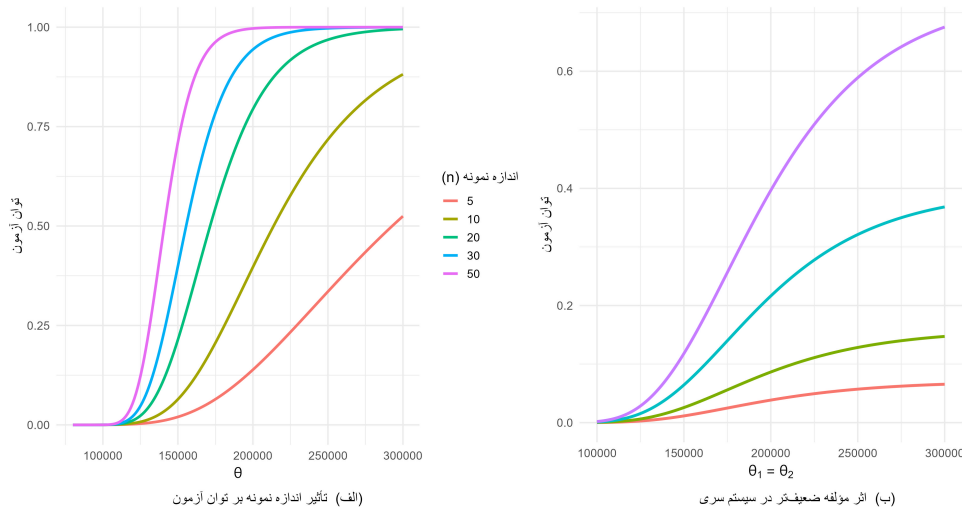


(و)

شکل ۱: تحلیل آماری و ارزیابی شاخص‌های قابلیت اطمینان سیستم سری (الف) نمودار بلوکی ساختار سیستم؛ (ب) توابع قابلیت اطمینان مؤلفه‌ها و سیستم؛ (ج) مقادیر احتمال حاصل از آزمون‌های جزئی؛ (د) توزیع کای-دو، ناحیه بحرانی و موقعیت آماره‌های آزمون؛ (ه) تابع چگالی احتمال خرابی مؤلفه‌ها؛ (و) مقایسه میانگین طول عمر نمونه‌ای با مقدار آستانه استاندارد (θ_0) .

شکل (۲) نتایج ارزیابی توان آزمون پیشنهادشده را در حالت‌های مختلف نمایش می‌دهد. شکل (الف) تغییرات توان آزمون اجتماع-اشتراک را نسبت به اندازه‌های مختلف نمونه (n) و مقادیر پارامتر حقیقی θ نمایش می‌دهد. همان‌طور که مشاهده می‌شود، با افزایش حجم نمونه، منحنی توان به‌صورت یکنواخت به عدد یک نزدیک می‌شود و توان آزمون برای تشخیص فاصله گرفتن از حداقل سطح قابل قبول بیشتر می‌شود. این موضوع نشان می‌دهد که با در اختیار داشتن داده‌های بیشتر، آزمون IUT نتایج دقیق‌تر و قابل اعتمادتری ارائه می‌دهد.

قسمت (ب) تأثیر حضور یک مؤلفه ضعیف‌تر در سیستم سری را بر توان کلی آزمون بررسی می‌کند. در این تحلیل، با فرض برابر بودن عملکرد دو مؤلفه ($\theta_1 = \theta_2$)، اثر تغییرات پارامتر مؤلفه سوم (θ_3) ارزیابی شده است. نتایج نشان می‌دهد که هرچه عملکرد مؤلفه سوم ضعیف‌تر باشد، توان کلی آزمون به‌شکل محسوسی کاهش می‌یابد. این موضوع بیانگر سازوکار اصلی آزمون IUT است که در آن، تصمیم نهایی بر اساس عملکرد ضعیف‌ترین مؤلفه گرفته می‌شود. این ویژگی به درستی با منطق سیستم‌های سری و نقش کلیدی عامل بحرانی در آن‌ها مطابقت دارد.



شکل ۲: توان آزمون اجتماع-اشتراک برای یک سیستم سری سه مؤلفه‌ای. (الف) تأثیر اندازه نمونه بر توان آزمون برای مقادیر مختلف پارامتر حقیقی θ ؛ با افزایش اندازه نمونه (ب) اثر وجود یک مؤلفه ضعیف‌تر در سیستم سری بر توان آزمون

۶ نتیجه‌گیری

در این مقاله، چارچوب آزمون اجتماع-اشتراک برای ارزیابی قابلیت اطمینان سیستم‌های سری مورد بررسی قرار گرفت. با توجه به این‌که عملکرد صحیح سیستم‌های سری به عملکرد مناسب تمامی مؤلفه‌ها وابسته است، ساختار فرضیه‌ها به‌گونه‌ای تعریف شد که فرضیه صفر بیانگر وجود حداقل یک مؤلفه با عملکرد نامطلوب باشد. نشان داده شد که این ساختار به‌صورت مستقیم با منطق آزمون‌های اجتماع-اشتراک سازگار است و تصمیم‌گیری نهایی بر پایه عملکرد ضعیف‌ترین مؤلفه انجام می‌شود. همچنین، رابطه توان آزمون برای مؤلفه‌های دارای طول عمر نمایی به‌دست آمد و نشان داده شد که توان کلی سیستم برابر با حاصل ضرب توان آزمون‌های جزئی است. نتایج تحلیل توان بیانگر آن بود که وجود حتی یک مؤلفه با عملکرد نزدیک به سطح آستانه می‌تواند توان کلی آزمون را به‌طور قابل توجهی کاهش دهد. این ویژگی با ماهیت سیستم‌های سری و نقش تعیین‌کننده مؤلفه بحرانی کاملاً سازگار است. علاوه بر این، مثال کاربردی و نتایج شبیه‌سازی نشان دادند که روش پیشنهادی می‌تواند عملکرد مطلوب تمامی مؤلفه‌ها را به‌صورت همزمان ارزیابی کند و چارچوبی منظم برای تصمیم‌گیری آماری در مسائل قابلیت اطمینان فراهم آورد. در مجموع، روش ارائه‌شده می‌تواند در طراحی، ارزیابی و کنترل کیفیت سیستم‌های مهندسی سری مورد استفاده قرار گیرد.

مراجع

- [1] Berger, R.L. (1982). Multiparameter hypothesis testing and acceptance sampling. *Technometrics*, **24**; 295–300.
- [2] Berger, R.L. (1989). Uniformly more powerful tests for hypotheses concerning linear inequalities and normal means. *Journal of the American Statistical Association*, **84**; 192–199.
- [3] Gleser, L. (1973). On a theory of intersection union tests. *Institute of Mathematical Statistics Bulletin*, **2(233)**, 9.
- [4] Lehmann, E.L. (1952). Testing multiparameter hypotheses. *Annals of Mathematical Statistics*, **23(4)**; 541–552.
- [5] Sasabuchi, S. (1980). A test of a multivariate normal mean with composite hypotheses determined by linear inequalities. *Biometrika*, **67**; 429–439.

- [6] Sasabuchi, S. (1988). A multivariate one-sided test with composite hypotheses when the covariance matrix is completely unknown. *Memoirs of the Faculty of Science*, **42**; 37–46.



نگهداری و تعمیر سیستم های تقسیم بار k -out-of- n :G بر اساس محاسبه باقیمانده طول عمر سیستم با استفاده از روش های یادگیری تقویتی

حقیقی، ف^۱ صفری، ع^۲ عیدی، ش^۳

^{۱,۲} گروه آمار، دانشکده ریاضی آمار و علوم کامپیوتر، دانشگاه تهران

^۳ بانک اقتصاد نوین

چکیده

در این مقاله به موضوع محاسبه طول عمر باقیمانده یک سیستم تقسیم بار k -out-of- n :G و سپس به مقوله نگهداری و تعمیر آن پرداخته شده است. بدین منظور فرض می شود طول عمر هر مولفه از یک خانواده عمومی به نام خانواده توزیع های گارویش (Gurvich) که دربرگیرنده توزیع های شناخته شده ای نظیر نمایی و وایبل است پیروی می کند. همچنین فرض می شود خرابی هر مولفه بر نرخ خطر مولفه های باقی مانده در سیستم براساس یک مدل پیشنهادی تاثیر می گذارد. ایده این مدل پیشنهادی برگرفته از مدل های آزمون های طول عمر شتابیده است که در آن خرابی یک مولفه به تحمیل استرس به سیستم تعبیر می شود. سپس یک فرمول بازگشتی برای قابلیت اعتماد سیستم محاسبه می گردد. در ادامه یک مدل نگهداری و تعمیر معرفی می شود که در آن براساس اطلاعات موجود از سابقه عملکرد سیستم، شامل تعداد و زمان از کار افتادگی هر مولفه، در هر دوره بازرسی باقیمانده طول

¹Speaker. fhaghighi@ut.ac.ir

²a.safari@ut.ac.ir

³shaqayq.eidi@gmail.com

عمر سیستم محاسبه می‌گردد و پس از مقایسه با آستانه ای بحرانی که از پیش تعیین شده در مورد نگهداری و تعمیر سیستم تصمیم‌گیری می‌شود. در مدل نگهداری و تعمیر پیشنهاد شده متغیرهای تصمیم‌گیری شامل آستانه بحرانی و طول بازه بازرسی با استفاده از روش‌های یادگیری تقویتی بهینه می‌شوند. نهایتاً با استفاده از مطالعه شبیه‌سازی عملکرد مدل خرابی و مدل نگهداری و تعمیر معرفی شده سنجیده می‌شود.

کلمات کلیدی: سیستم تقسیم بار k -out-of- n :G، نگهداری و تعمیر، الگوریتم یادگیری تقویتی، باقیمانده طول عمر.

۱ مقدمه

در دنیای واقعی اکثر سیستم‌های موجود، چند جزئی هستند که از طریق تعامل با یکدیگر برای هدف خاصی طراحی شده‌اند. در این سیستم‌ها عملکرد هر جز بر عملکرد کل سیستم و حتی عملکرد سایر اجزا موثر است. یکی از این سیستم‌های چند جزئی، سیستم‌های تقسیم بار k -out-of- n :G هستند که در آن پس از خرابی هر مولفه، بار بر روی سایر مولفه‌ها تقسیم می‌شود و نهایتاً سیستم فعال باقی می‌ماند تا زمانی که حداقل k مولفه فعال باشد. ساختار چنین سیستم‌هایی حتی در ساده‌ترین شکل آن نیز پیچیده است و در نتیجه تشخیص به موقع از کار افتادگی سیستم و به کارگیری اقدام موثر برای نگهداری و تعمیر آن با مشکل مواجه است. یکی از معیارهای مهم در تشخیص به هنگام از کار افتادگی سیستم، محاسبه باقیمانده طول عمر^۱ (RUL) سیستم است که عموماً حتی برای سیستم‌های تک جزئی کار دشواری است. بخشی از این دشواری مربوط به محاسبات پیچیده ای است که به کارگیری روشهای سنتی را در مدل‌های نگهداری و تعمیر ناکام می‌سازد. همین امر باعث شده موضوع معرفی مدل‌های نگهداری و تعمیر برای سیستم‌های تقسیم بار مبتنی بر محاسبه باقیمانده طول عمر سیستم مغفول واقع شود. با این وجود، امروزه با پیشرفت الگوریتم‌های هوش مصنوعی می‌توان تاحدودی بر محدودیت‌های روش‌های سنتی در این زمینه فایده‌آمیز و به جای به کارگیری محاسبات پیچیده (که اغلب ادامه مسیر را غیر ممکن می‌کند) به منظور بهینه‌سازی از روشهای بهینه‌سازی با استفاده از الگوریتم‌های یادگیری تقویتی بهره جست.

هدف از این مقاله محاسبه قابلیت اعتماد سیستم‌های تقسیم بار k -out-of- n :G و معرفی مدلی برای نگهداری و تعمیر آن است. نوآوری روش حاضر اولاً در تعریف مکانیزم اثرگذاری خرابی هر مولفه بر نرخ خرابی سیستم و ثانیاً پیشنهاد یک مدل بهینه نگهداری و تعمیر مبتنی بر باقیمانده طول عمر سیستم با استفاده از روشهای یادگیری تقویتی است.

^۱ Remianing Useful Lifetime (RUL)

برخی پژوهش‌های اخیر در مورد نگهداری و تعمیر سیستم‌های تقسیم بار که با استفاده از روش‌های هوش مصنوعی انجام شده به شرح زیر هستند. [۱] مسئله نگهداری و تعمیر بهینه چندمؤلفه‌ای را با استفاده از یک الگوریتم یادگیری تقویتی عمیق ارائه کردند، در این پژوهش اثر تقسیم بار مستقیماً بر سیاست عامل تأثیر می‌گذارد. [۲] یک چارچوب یادگیری تقویتی چندعامله برای نگهداری جامع سیستم‌های تقسیم بار سری پیشنهاد کردند. [۳] نیز با ارائه یک راهبرد بازرسی بهینه برای اجزای سه حالتی که توسط یک فرآیند مارکوف مبتنی بر حالت مدل‌سازی شده‌اند، به این حوزه کمک کردند. مقاله حاضر برگرفته از مقاله نویسندگان [۴] است که در حال تکمیل می‌باشد.

۲ فرضیات مدل

در طراحی مدل خرابی فرضیات زیر پذیرفته شده است:

- سیستم تحت قانون تقسیم بار برابر عمل می‌کند. به این مفهوم که پس از خرابی هر مولفه، بار به تساوی بین سایر مولفه‌ها تقسیم می‌شود.
- فرض کنید T_i ، برای $i = 1, 2, \dots, n$ ، زمان تصادفی خرابی مؤلفه i ام باشد. مکانیزم تقسیم بار، وابستگی آماری میان این زمان‌های خرابی ایجاد می‌کند. در نتیجه، متغیرهای تصادفی $T_1, T_2, \dots, T_n$ دارای توزیع‌های ناهمگن و از لحاظ آماری وابسته هستند. توالی خرابی مولفه‌های سیستم با استفاده از آماره‌های ترتیبی، که به صورت $T_{(1)} < T_{(2)} < \dots < T_{(n)}$ نشان داده می‌شوند، تعریف می‌گردد؛ در اینجا $T_{(j)}$ بیانگر زمان j امین خرابی است.
- پس از رویداد j امین خرابی، برای $j = 0, 1, \dots, n-k$ ، طول عمر هر یک از مؤلفه‌های باقی‌مانده در سیستم از خانواده توزیع $Gurvich$ پیروی می‌کند. در نتیجه تابع قابلیت اعتماد برای هر مؤلفه باقی‌مانده در این مرحله به صورت زیر تعریف می‌شود:

$$R(t; \theta_j) = \exp\left(-\frac{g(t; \delta)}{\theta_j}\right), \quad \text{برای } t > T_{(j)}, \quad (2)$$

که در آن $g(t; \delta)$ تابعی یکنواخت صعودی بر حسب t بوده و شرط $g(0; \delta) = 0$ را برآورده می‌سازد. در نتیجه برای تابع نرخ خطر داریم:

$$h(t; \theta_j) = \frac{g'(t; \delta)}{\theta_j}, \quad \text{برای } t > T_{(j)}. \quad (3)$$

به طوری که $T_{(0)} = 0$ ، $g'(t; \delta) = \frac{\partial}{\partial t} g(t; \delta)$ و θ_j پارامتر مقیاس، به صورت زیر تعریف می شود.

$$\theta_j = \beta_0 - j\beta_1, \quad \text{برای } j = 0, 1, \dots, n-k \quad (4)$$

که در آن $\beta_0 > 0$ و $\beta_1 > 0$ پارامترهای مدل هستند. برای تضمین اعتبار مدل، پارامترها با شرط $\theta_j > 0$ برای همه j ها مقید می شوند که این امر مستلزم شرط $\beta_0 - (n-k)\beta_1 > 0$ است.

۳ قابلیت اعتماد مؤلفه

بر اساس فرضیات مدل، قابلیت اعتماد مؤلفه های سیستم به صورت زیر تعریف می شوند:

- هنگامی که هیچ خرابی رخ نمی دهد، تابع قابلیت اعتماد هر مؤلفه در زمان t با $R_{\theta_0}(t)$ نشان داده می شود و به صورت زیر بیان می شود.

$$R_{\theta_0}(t) = \exp \left\{ \frac{-g(t, \delta)}{\theta_0} \right\}, \quad t > 0. \quad (1)$$

- فرض کنید $\pi(1)$ یک جایگشت تک عضوی از مجموعه اندیس $\mathcal{M} = \{1, 2, \dots, n\}$ باشد، و $\pi(2) \in \mathcal{M} - \{1\}$ پس از خرابی اولین مؤلفه در زمان t_1 ، بار میان $n-1$ مؤلفه باقی مانده تقسیم می شود. در این حالت، تابع قابلیت اعتماد شرطی $T_{\pi(2)}$ با شرط $T_{\pi(1)} = t_1$ به صورت زیر داده می شود:

$$\begin{aligned} R_{T_{\pi(2)}|T_{\pi(1)}=t_1}(t) &= \exp \left\{ - \int_{t_1}^t \frac{g'(x, \delta)}{\theta_0} dx - \int_{t_1}^t \frac{g'(x, \delta)}{\theta_1} dx \right\}, \\ &= \exp \left\{ - \int_{t_1}^t \frac{g'(x, \delta)}{\theta_0} dx - \int_{t_1}^t \frac{g'(x, \delta)}{\theta_1} dx + \int_{t_1}^{t_1} \frac{g'(x, \delta)}{\theta_1} dx \right\}, \\ &= R_{\theta_1}(t) \frac{R_{\theta_0}(t_1)}{R_{\theta_1}(t_1)}, \quad t_1 < t. \end{aligned} \quad (2)$$

که در آن $R_{\theta_j}(t) = \exp \left\{ \frac{-g(t, \delta)}{\theta_j} \right\}$ ، $j = 1, \dots, n$

- فرض کنید $\{\pi(1), \dots, \pi(j)\}$ یک جایگشت j عضوی از مجموعه اندیس $\mathcal{M}$ باشد و $\pi(j+1) \in \mathcal{M} - \{1, \dots, j\}$ هنگامی که j امین مؤلفه در زمان t_j ، $j = 1, \dots, n-k$ خراب می شود، تابع قابلیت اعتماد

شرطی $T_{\pi(j+1)}$ با شرط $T_{\pi(j)} = t_j, \dots, T_{\pi(1)} = t_1$ به صورت زیر بیان می‌شود:

$$R_{T_{\pi(j+1)}|T_{\pi(1)}=t_1, \dots, T_{\pi(j)}=t_j}(t) = R_{\theta_j}(t) \prod_{i=1}^j \frac{R_{\theta_{i-1}}(t_i)}{R_{\theta_i}(t_i)}, \quad t_1 < \dots < t_j < t. \quad (3)$$

در نتیجه، تابع چگالی شرطی $T_{\pi(j+1)}$ با شرط $T_{\pi(j)} = t_j, \dots, T_{\pi(1)} = t_1$ به صورت زیر به دست می‌آید:

$$f_{T_{\pi(j+1)}|T_{\pi(1)}=t_1, \dots, T_{\pi(j)}=t_j}(t) = f_{\theta_j}(t) \prod_{i=1}^j \frac{R_{\theta_{i-1}}(t_i)}{R_{\theta_i}(t_i)}, \quad t_1 < \dots < t_j < t. \quad (4)$$

که در آن $f_{\theta_j}(t) = \frac{dR_{\theta_j}(t)}{dt}$

۴ قابلیت اعتماد سیستم

با بکارگیری لم زیر، فرمول قابلیت اعتماد یک سیستم تقسیم بار k -out-of- n : G به صورت بازگشتی محاسبه شده و در قالب قضیه زیر مطرح می‌شود. به دلیل گستردگی اثبات‌ها از بیان آنها خوداری شده است.

لم ۱۰۴. تابع چگالی توأم زمان‌های خرابی متوالی مؤلفه‌ها $T_{(1)} < T_{(2)} < \dots < T_{(n)}$ به صورت زیر داده می‌شود:

$$f_{T_{(1)}, \dots, T_{(n)}}(t_1, \dots, t_n) = n! \prod_{j=1}^n \prod_{i=1}^{j-1} f_{\theta_{j-1}}(t_i) \frac{R_{\theta_{i-1}}(t_i)}{R_{\theta_i}(t_i)} \quad (1)$$

که در آن $\prod_{i=1}^0(\cdot) = 1$

قضیه ۲۰۴. تابع قابلیت اعتماد یک سیستم تقسیم بار k -out-of- n : G که با $R_s^{k:n}(t)$ نشان داده می‌شود، توسط رابطه بازگشتی زیر بیان می‌گردد:

$$R_s^{k:n}(t) = R_s^{k-1:n}(t) - P_s^{k-1:n}(t), \quad k = 1, \dots, n,$$

که در آن $P_s^{k-1:n}(t)$ احتمال این است که دقیقاً $k-1$ مؤلفه در زمان t کارکنند، و به صورت زیر داده می‌شود:

$$P_s^{k-1:n}(t) = n! \int \dots \int_S \left(\prod_{j=1}^n \prod_{i=1}^{j-1} f_{\theta_{j-1}}(t_j) \frac{R_{\theta_{i-1}}(t_i)}{R_{\theta_i}(t_i)} \right) dt_1 \dots dt_n,$$

که دامنه انتگرال S به صورت زیر تعریف می‌شود:

$$S = \{(t_1, \dots, t_n) \in \mathbb{R}^n \mid 0 < t_1 < t_2 < \dots < t_{n-k+1} < t < t_{n-k+2} < \dots < t_n < \infty\}$$

و $R_s^{0:n}(t) = 1$ است.

نکته ۳۰۴. فرمول فوق در حالت خاص بر تابع قابلیت یک سیستم سری و موازی منطبق است.

۵ مدل نگهداری و تعمیر

در این مطالعه، یک مدل نگهداری و تعمیر برای سیستم تقسیم بار k -out-of- n :G ارایه شده که از مفهوم عمر مفید باقیمانده برای تصمیم‌گیری استفاده می‌کند.

فرض کنید $\{k_T\}_{k \in \mathbb{N}}$ دنباله‌ای از زمان‌های بازرسی دوره‌ای باشد. در هر زمان بازرسی k_T ، وضعیت سیستم کاملاً قابل مشاهده است، به این معنا که برای هر مؤلفه، وضعیت عملیاتی آن مشخص است. به طور خاص، اگر مؤلفه‌ای از کار افتاده باشد، زمان دقیق خرابی آن ثبت می‌شود؛ در غیر این صورت، تأیید می‌شود که مؤلفه در زمان k_T عملیاتی است.

تاریخچه‌ی مشاهدات تا زمان k_T به صورت زیر داده می‌شود:

$$\mathbb{O}_{k_T} = \{(T_{(j)} = t_j)_{j=1}^q\} \cup \{T_{(q+1)} > k_T\}, \quad (1)$$

که در آن q تعداد مؤلفه‌های از کار افتاده تا زمان k_T را نشان می‌دهد. RUL سیستم در زمان بازرسی k_T با شرط $\mathbb{O}_{k_T}$ به صورت زیر تعریف می‌شود:

$$\text{RUL}_{k_T|\mathbb{O}_{k_T}} = T_s - k_T | \mathbb{O}_{k_T} \quad (2)$$

که در آن طبق قرارداد، $\text{RUL}_{k_T|\mathbb{O}_{k_T}} = 0$ اگر سیستم تا زمان k_T از کار افتاده باشد. مدل تعمیر و نگهداری زیر بر اساس p مشخص، $\text{RUL}_{k_T|\mathbb{O}_{k_T}}$ و اعمال زیر پیشنهاد داده می‌شود.

۱. عدم مداخله اگر

$$\mathbb{P}(\text{RUL}_{k_T|\mathbb{O}_{k_T}} < \tau) < p,$$

سیستم عملیاتی باقی می‌ماند و تنها هزینه‌ی بازرسی C_{ins} تحمیل می‌شود.

۲. نگهداری پیشگیرانه اگر

$$\mathbb{P}(\text{RUL}_{k_T|\mathbb{O}_{k_T}} < \tau) \geq p,$$

یک تعمیر کامل پیشگیرانه با هزینه C_P انجام می‌شود.

۳. نگهداری اصلاحی اگر سیستم در زمان k_T در حالت خرابی یافت شود، یک تعمیر کامل اصلاحی با هزینه C_c اجرا می‌گردد.

هدف یافتن مقادیر بهینه p و τ با استفاده از روش یادگیری تقویتی است به گونه ای که هزینه نگهداری و تعمیر کمینه گردد.

۶ مطالعات شبیه سازی

برای نشان دادن اثربخشی و مقیاس پذیری چارچوب پیشنهادی، یک سیستم تقسیم بار 1-out-of-2:G را در نظر می گیریم که از مؤلفه های یکسان تشکیل شده اند. در این سیستم، طول عمر مؤلفه ها از توزیع وایبول با پارامتر شکل $\sigma = 1/5$ پیروی می کند. پارامتر مقیاس θ به صورت زیر در نظر گرفته می شود:

$$\bullet \theta_1 = 12: \text{ هنگامی که هر دو مؤلفه کار می کنند}$$

$$\bullet \theta_2 = 7/5: \text{ هنگامی که تنها یک مؤلفه کار می کند}$$

برای حل مدل نگهداری و تعمیر و بدست آوردن متغیرهای بهینه مدل ما یک چارچوب پیش بینی بدون مدل (model-free) را با الگوریتم یادگیری تقویتی ^۲ (DQN) ادغام می کنیم. در این راستا فرایند تصمیم گیری به صورت یک فرایند تصمیم گیری مارکوف ^۳ (MDP) با مؤلفه های زیر فرموله می شود:

• فضای حالت S : هر حالت به صورت زیر نمایش داده می شود

$$s_t = (p_\tau(t), n_f(t)),$$

که در آن

$$p_\tau(t) = \mathbb{P}(\text{RUL}(t) < \tau)$$

$$n_f(t) = \text{تعداد کل خرابی ها تا زمان } t$$

• فضای اقدام A :

$$A = \{ \text{انجام تعویض : ۱, عدم مداخله : ۰} \}.$$

• تابع پاداش $\mathcal{R}(s, a)$:

$$R(s, a) = -\mathbb{I}_{\text{ins}} C_{\text{ins}} - \mathbb{I}_{a=1} C_p - \mathbb{I}_{\text{fail}} C_c,$$

Deep Q-Network (DQN)^۲
Markov Decision Process (MDP)^۳

که در آن

$$C_{\text{ins}} = 8, \quad C_p = 35, \quad C_c = 200,$$

و II توابع نشانگر برای رویدادهای بازرسی، اقدام پیشگیرانه و خرابی هستند.

فواصل بازرسی زیر را در نظر می‌گیریم.

$$\mathcal{T} = \{0/5, 0/8, 1/0, 1/2, 1/4\}$$

برای هر $\tau \in \mathcal{T}$ ، عامل DQN تحت 1,500 دوره آموزشی قرار می‌گیرد تا به یک سیاست بهینه همگرا شود. برای مقایسه عملکرد فواصل بازرسی مختلف، هزینه متوسط به ازای واحد زمان از 1000 دوره با طول ثابت 50 تخمین زده می‌شود که منجر به نتایج زیر شده است.

• هنگامی که هر دو مؤلفه فعال هستند:

- برای $p_\tau \leq 0/5$ ، عامل اقدام 0 (عدم مداخله) را بر می‌گزیند.

- برای $p_\tau \geq 0/5$ ، عامل به اقدام 1 (تعویض سیستم) سویچ می‌کند، زیرا افزایش ریسک خرابی هزینه نگهداری را توجیه می‌کند.

• پس از خرابی یک مؤلفه:

عامل به طور مداوم اقدام 1 (تعویض سیستم) را در تقریباً تمام حالت‌ها انتخاب می‌کند. این راهبرد محافظه‌کارانه به ریسک بالاتر در شرایط تقسیم بار پاسخ می‌دهد، جایی که تنها مؤلفه باقی‌مانده بار کامل سیستم را تحمل کرده و تخریب شتاب‌یافته‌ای را تجربه می‌کند.

به منظور مقایسه عملکرد بازه‌های بازرسی متفاوت، متوسط هزینه و متوسط p_τ که به ترتیب با $\bar{C}_\tau$ و $\bar{p}_\tau$ نمایش داده می‌شوند به ازای واحد زمان از 1000 رویداد با طول ثابت 50 تخمین زده شده است و در جدول 1 گزارش شده اند. طبق مدل نگهداری و تعمیر پیشنهادی وقتی $\bar{p}_\tau$ از آستانه بحرانی بیشتر شود تعویض پیشگیرانه انجام می‌شود. داده‌ها نشان می‌دهند که برای سیستم مورد بررسی بازه بازرسی $\tau = 1/0$ کمترین هزینه $\bar{C}_\tau = 39/5$ را در $\bar{p}_\tau = 0/6$ به دست می‌دهد در نتیجه، $\tau^* = 1/0$ به عنوان بازه بازرسی بهینه برای سیاست نگهداری پیشنهادی انتخاب می‌شود.

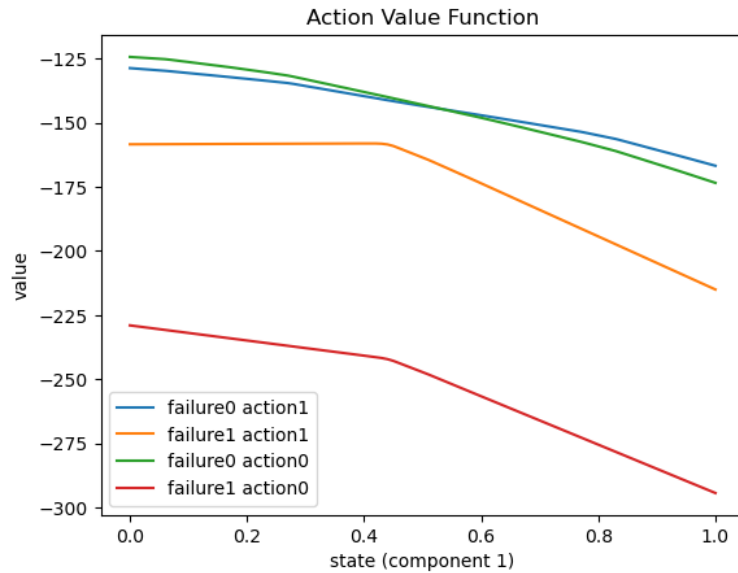
جدول ۱: مقایسه عملکرد در بازه‌های مختلف بازرسی برای یک سیستم 1-out-of-2:G (۱۰۰۰ دوره، افق زمانی ۵۰ ساعته).

τ	$\overline{p_\tau} \pm sd$	$\overline{C_\tau} \pm sd$
۱/۴	۰/۷۲ ± ۰/۱۰	۴۱/۹ ± ۸/۰
۱/۲	۰/۶۶ ± ۰/۱۰	۴۱/۰ ± ۸/۱
۱/۰	۰/۶۰ ± ۰/۱۰	۳۹/۵ ± ۸/۱
۰/۸	۰/۵۴ ± ۰/۱۲	۴۱/۴ ± ۷/۱
۰/۵	۰/۳۷ ± ۰/۱۶	۴۳/۱ ± ۷/۴

تابع مقدار-عمل در شکل نشان ۶ سیاست نگهداری یاد گرفته شده را برای بازه بازرسی بهینه $\tau^* = ۱/۰$ سیستم تقسیم بار 1-out-of-2:G می‌دهد. این نمودار یک منطق تصمیم‌گیری روشن را آشکار می‌کند. هنگامی که هر دو مؤلفه عملیاتی هستند (حالت خرابی ۰)، عامل برای $p_\tau < ۰/۵$ اقدام ۰ (ادامه کار) را انتخاب می‌کند و برای $p_\tau \geq ۰/۵$ به اقدام ۱ (تعویض سیستم) سوییچ می‌کند، زیرا افزایش ریسک خرابی هزینه نگهداری را توجیه می‌کند. پس از خرابی یک مؤلفه، عامل تقریباً در همه حالت‌ها به‌طور مداوم اقدام ۱ (تعویض سیستم) را انتخاب می‌کند. این راهبرد محافظه‌کارانه به ریسک افزایش‌یافته در شرایط تقسیم بار می‌پردازد، جایی که تنها مؤلفه باقیمانده بار کامل سیستم را تحمل می‌کند و در نتیجه تخریب شتاب‌یافته رخ می‌دهد. این سیاست نشان می‌دهد که عامل DQN به‌خوبی یاد گرفته است که بین حالت‌های عملیاتی مختلف سیستم 1-out-of-2:G تمایز قائل شود و هنگام از دست رفتن افزونگی، اقدامات نگهداری تهاجمی‌تری را به کار گیرد.

۷ نتیجه‌گیری

در این تحقیق موضوع محاسبه قابلیت اعتماد سیستم‌های تقسیم بار k-out-of-n:G مورد مطالعه قرار گرفته است. در مدل بندی قابلیت اعتماد هر مؤلفه سیستم فرض شده که خرابی هر مؤلفه نرخ خطر خرابی مؤلفه‌های باقیمانده را افزایش می‌دهد. سپس یک فرمول بازگشتی برای محاسبه قابلیت اعتماد چنین سیستم‌هایی پیشنهاد می‌گردد. در ادامه به موضوع نگهداری و تعمیر این سیستم‌ها پرداخته شده و یک مدل نگهداری و تعمیر مبتنی بر محاسبه باقیمانده طول عمر سیستم پیشنهاد شده است. در مدل پیشنهادی، بازرسی‌ها پریودیک هستند و در هر زمان بازرسی احتمال اینکه سیستم قبل از دوره بازرسی بعدی غیر فعال شود محاسبه می‌گردد و با یک مقدار از پیش تعیین شده مقایسه



شکل ۱: تابع مقدار-عمل یادگرفته شده برای $\tau^* = 1/0$ در سیستم های 1-out-of-2:G

می شود. سپس طبق سیاست مدل نگهداری و تعمیر پیشنهادی در مورد ادامه کار سیستم و یا تعویض آن تصمیم گرفته می شود. متغیر طول دوره بازرسی و آستانه بحرانی از طریق روش های یادگیری تقویتی بهینه می شود.

مراجع

- [1] Zhang C, Li YF, Coit DW. Deep reinforcement learning for dynamic opportunistic maintenance of multi-component systems with load sharing, *IEEE Transactions on Reliability*, 2023;72(3):863–77.
- [2] Zhao S, Wei Y, Li Y, Cheng Y. A multi-agent reinforcement learning (MARL) framework for designing an optimal state-specific hybrid maintenance policy for a series k-out-of-n load-sharing system, *Reliability Engineering & System Safety*, 2026;265:111587.
- [3] Wang J , Zhang J, Wub J, Zhang Q, Fang Y, Huang X. Optimal inspection-based maintenance strategy for -out-of- : G load-sharing systems consisting of three-state components, *Computers & Industrial Engineering*, 2025;209:111446.
- [4] Haghghi F, Safari A, Eidi Sh. Model-based prognostic for a load sharing k-out-of-n:G system. Under preparation.



یک رویکرد ترکیبی مبتنی بر یادگیری عمیق برای پیش‌بینی عمر مفید باقیمانده

کرمی،^۱ حقیقی،^۲ ف ۲ صفری،^۳ ع

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

چکیده

پیش‌بینی دقیق عمر مفید باقیمانده، از اجزای اصلی سیستم‌های نگهداری پیشگیرانه در صنایع حساس مانند هوافضا است. چالش اصلی در این حوزه، پیچیدگی داده‌های سری زمانی حسگرهاست که هم شامل الگوهای کوتاه‌مدت محلی و هم وابستگی‌های بلندمدت می‌شود. برای پاسخ به این چالش، مدلی ترکیبی متشکل از بلوک استخراج ویژگی‌های چندمقیاسی، حافظه بلندمدت-کوتاه‌مدت دوطرفه و مکانیزم توجه چندسرطراحی شده است. آزمایش‌های انجام‌شده بر روی مجموعه داده استاندارد توربوفن ناسا اهمیت رویکردهای ترکیبی را در پیش‌بینی عمر مفید باقیمانده نشان می‌دهد.

کلمات کلیدی: پیش‌بینی عمر مفید باقیمانده، یادگیری عمیق، شبکه عصبی پیچشی، حافظه کوتاه‌مدت-بلندمدت، مکانیزم توجه.

¹Speaker. ahad.karami@ut.ac.ir

²fhighighi@ut.ac.ir

³a.safari@ut.ac.ir

فناوری مدیریت سلامت و پیش‌بینی خرابی^۱ (PHM) به‌طور گسترده در صنایع امروزی برای تضمین ایمنی و قابلیت اطمینان تجهیزات مورد استفاده قرار می‌گیرد. این فناوری با نظارت مستمر بر وضعیت تجهیزات و شناسایی زود هنگام نشانه‌های استهلاک امکان انجام اقدامات پیشگیرانه را فراهم می‌کند. در این میان، پیش‌بینی عمر مفید باقیمانده^۲ (RUL) یک بخش اصلی سیستم‌های PHM محسوب می‌شود. تخمین دقیق RUL به برنامه‌ریزی فرایندهای نگهداری و تعمیرات کمک می‌کند تا مشخص شود هر قطعه تا چه زمانی می‌تواند به‌طور ایمن و قابل اطمینان به کار خود ادامه دهد. اهمیت این موضوع از یک سو در برنامه‌ریزی بهینه تعمیر و نگهداری و جلوگیری از توقف‌های پرهزینه تولید، و از سوی دیگر در کاهش هزینه‌های عملیاتی از طریق کاهش تعمیرات غیرضروری نمایان است [۱]. از سوی دیگر، در صنایع حیاتی مانند هوانوردی، تخمین دقیق RUL نقشی کلیدی در جلوگیری از حوادث فاجعه‌بار و حفظ جان انسان‌ها ایفا می‌کند.

روش‌های پیش‌بینی عمر مفید باقیمانده را به‌طور کلی می‌توان به دو دسته اصلی تقسیم کرد: روش‌های مبتنی بر مدل فیزیکی و روش‌های مبتنی بر داده [۱]. رویکرد اول اگرچه از پایه‌های نظری قوی برخوردار است، اما به درک عمیقی از مکانیزم خرابی و قوانین فیزیک حاکم بر سیستم نیاز دارد. پیچیدگی و غیرخطی بودن سیستم‌های مدرن، به‌ویژه در صنایعی مانند هوافضا، ساخت یک مدل فیزیکی جامع را دشوار، زمان‌بر و گاهی غیرعملی می‌سازد [۲]. به همین دلیل، پژوهشگران به تدریج به سمت روش‌های داده‌محور متمایل شده‌اند. این روش‌ها برخلاف رویکردهای فیزیکی، نیازی به دانش تخصصی عمیق از سیستم ندارند و با استفاده از الگوریتم‌های تحلیل داده، قادر به شناسایی خودکار الگوهای پنهان استهلاک از داده‌های خام حسگرها هستند. با این حال، رویکردهای داده‌محور نیز با چالش‌هایی مواجه هستند. پیچیدگی محاسباتی و نیاز به حافظه بالا در بسیاری از مدل‌های موفق در این حوزه، استقرار آن‌ها را در دستگاه‌های با منابع محدود با مشکل مواجه می‌سازد [۳]. بنابراین، با وجود برتری روش‌های داده‌محور نسبت به رویکردهای فیزیکی، توسعه مدل‌هایی با قابلیت اجرای کارآمد در محیط‌های با منابع محاسباتی محدود نیازی اساسی در صنایع به شمار می‌رود.

در دسته روش‌های داده‌محور، رویکردهای سنتی یادگیری ماشین در بسیاری از مطالعات مورد توجه پژوهشگران

^۱ Prognostics and Health Management (PHM)

^۲ Remaining Useful Life (RUL)

بوده‌اند. الگوریتم‌هایی نظیر ماشین بردار پشتیبان^۳، مدل مخفی مارکوف^۴، ماشین یادگیری افراطی^۵، درخت تصمیم^۶، جنگل تصادفی^۷ و K-نزدیک‌ترین همسایه^۸ به‌طور گسترده برای پیش‌بینی عمر مفید باقیمانده مورد استفاده قرار گرفته‌اند [۱، ۳]. با این حال، این رویکردها با محدودیت‌های اساسی مواجه هستند. این الگوریتم‌ها علی‌رغم اثربخشی در برخی کاربردها، در پردازش داده‌های سری زمانی و ساختارهای غیرخطی پیچیده با چالش مواجه‌اند. افزون بر این، این روش‌ها به شدت به مهندسی ویژگی^۹ دستی وابسته هستند که فرایندی وقت‌گیر است و تطبیق‌پذیری کلی آن‌ها را محدود می‌سازد [۳]. انتخاب ویژگی‌های نامناسب بدون بهره‌گیری از دانش تخصصی مربوطه، می‌تواند به عملکرد نامطلوب مدل منجر شود [۱]. همچنین، فرایند انتخاب ویژگی‌های دستی می‌تواند ساختار زمانی ذاتی داده‌ها را مختل کرده و در نتیجه دقت پیش‌بینی نهایی را کاهش دهد [۴]. این محدودیت‌ها، ضرورت به‌کارگیری روش‌هایی را نشان می‌دهد که توانایی استخراج خودکار ویژگی‌ها از داده‌های خام حسگرها را داشته و قادر به مدل‌سازی وابستگی‌های زمانی و ساختارهای غیرخطی نهفته در این داده‌ها باشند.

در مقابل، یادگیری عمیق با استخراج خودکار ویژگی‌ها، نیاز به مهندسی دستی ویژگی را برطرف کرده و به دلیل قابلیت مدل‌سازی روابط غیرخطی پیچیده، به ابزاری کارآمد برای پیش‌بینی عمر مفید باقیمانده تبدیل شده است. شبکه‌های حافظه کوتاه‌مدت-بلندمدت^{۱۰} (LSTM) در مدل‌سازی وابستگی‌های بلندمدت سری‌های زمانی عملکرد برجسته‌ای دارند، در حالی که شبکه‌های عصبی پیچشی^{۱۱} (CNN) در استخراج الگوهای محلی از داده‌های خام حسگرها شناخته شده‌اند [۲، ۵]. با این حال، تکیه صرف بر هر یک از این معماری‌ها به‌تنهایی نمی‌تواند تمام جنبه‌های پیچیده فرایند استهلاک را پوشش دهد؛ چرا که داده‌های حسگرها هم دارای الگوهای محلی و هم وابستگی‌های بلندمدت هستند. به همین دلیل، مدل‌های ترکیبی که نقاط قوت هر دو رویکرد را با هم تلفیق می‌کنند، به عنوان راه حلی مؤثر و کارآمد مورد توجه گسترده قرار گرفته‌اند [۱، ۵]. علاوه بر این، افزودن مکانیزم‌های توجه^{۱۲} به این معماری‌های ترکیبی، این قابلیت را به شبکه می‌دهد تا به‌طور تطبیقی به ویژگی‌های مهم‌تر وزن بیشتری اختصاص دهد و از اطلاعات ضروری ذخیره‌شده در پنجره‌های زمانی طولانی استفاده کند [۲، ۵]. چنین مدل‌هایی نه تنها قابلیت تعمیم‌پذیری بالایی دارند، بلکه در برابر نویز و داده‌های دورافتاده مقاوم بوده و می‌توانند عملکرد دقیق‌تر و پایدارتری در پیش‌بینی RUL

^۳Support Vector Machine (SVM)

^۴Hidden Markov Model (HMM)

^۵Extreme Learning Machine (ELM)

^۶Decision Tree

^۷Random Forest

^۸K-Nearest Neighbors (KNN)

^۹Feature Engineering

^{۱۰}Long Short-Term Memory (LSTM)

^{۱۱}Convolutional Neural Network (CNN)

^{۱۲}Attention Mechanisms

ارائه دهند [۳، ۵].

با در نظر گرفتن چالش‌های پیش‌بینی عمر مفید باقیمانده در سیستم‌های پیچیده، بهره‌گیری از قابلیت‌های مکمل شبکه‌های عصبی ضروری است. شبکه‌های عصبی پیچشی در استخراج الگوهای محلی، شبکه‌های حافظه کوتاه‌مدت-بلندمدت در مدل‌سازی وابستگی‌های زمانی و مکانیزم‌های توجه در وزن‌دهی به ویژگی‌های مهم کارآمد هستند؛ از این رو، در این پژوهش مدلی ترکیبی مبتنی بر یادگیری عمیق متشکل از این سه مؤلفه ارائه شده است. برای آموزش و ارزیابی مدل پیشنهادی، از مجموعه داده استاندارد توربوفن ناسا استفاده شده است که دارای شرایط عملیاتی چندگانه و حالات خرابی گوناگون است. نتایج تجربی نشان می‌دهد که روش پیشنهادی عملکردی رقابتی با مدل‌های پیشرفته ارائه شده در پژوهش‌های اخیر دارد.

در ادامه مقاله، در بخش دوم، مجموعه داده مورد استفاده همراه با جزئیات عملیات پیش‌پردازش معرفی شده و معماری مدل پیشنهادی تشریح می‌گردد. در بخش سوم، نتایج ارزیابی مدل پیشنهادی بر روی مجموعه داده مذکور ارائه شده و عملکرد آن با پژوهش‌های اخیر مقایسه می‌شود. در نهایت، بخش چهارم به نتیجه‌گیری و اشاره به مسیرهای آتی پژوهش اختصاص یافته است.

۲ روش تحقیق

۱.۲ مجموعه داده و پیش‌پردازش

در این پژوهش برای ارزیابی عملکرد مدل پیشنهادی، از مجموعه داده C-MAPSS توسعه‌یافته توسط سازمان فضایی ناسا (NASA) استفاده شده است. این مجموعه داده، داده‌های شبیه‌سازی شده استهلاک تدریجی موتورهای توربوفن را شامل می‌شود. هر موتور در ابتدا در وضعیت سالم قرار داشته و به تدریج به سمت خرابی کامل پیش می‌رود؛ این فرایند کاهش عمر تا لحظه خرابی نهایی، در قالب چرخه‌های عملیاتی^{۱۳} ثبت شده است. در هر چرخه عملیاتی، داده‌های ۲۱ حسگر مختلف شامل دما، فشار و سرعت چرخش فن و همچنین سه پارامتر شرایط کارکردی برای هر موتور ثبت می‌شود.

مجموعه داده C-MAPSS شامل چهار زیرمجموعه اصلی به نام‌های FD001 تا FD004 است که بر اساس سناریوهای عملیاتی و نوع عیب دسته‌بندی شده‌اند. در این میان، زیرمجموعه‌های FD001 و FD003 دارای یک شرط عملیاتی^{۱۴} ساده و زیرمجموعه‌های FD002 و FD004 شامل شش شرایط عملیاتی هستند. همچنین، زیرمجموعه‌های FD001 و

^{۱۳} Operational Cycles

^{۱۴} Operating Condition

FD002 تنها شامل یک نوع عیب و زیرمجموعه‌های FD003 و FD004 شامل دو نوع عیب ترکیبی می‌باشند. پیش‌پردازش مجموعه داده انتخاب شده طی چهار مرحله انجام شده است. در مرحله نخست، به منظور کاهش ورودی‌های کم‌اهمیت و افزایش کارایی مدل، تحلیل اولیه‌ای روی داده‌های خام ۲۱ گانه حسگرها انجام شد. بررسی‌ها نشان داد که حسگرهای شماره ۱، ۵، ۶، ۱۰، ۱۶، ۱۸ و ۱۹ دارای مقادیر ثابت در طول چرخه‌های عملیاتی هستند و اطلاعات مفیدی برای تعیین عمر مفید باقیمانده ارائه نمی‌دهند. این هفت حسگر از مجموعه داده حذف شده و در نتیجه، داده‌های ۱۴ حسگر باقی‌مانده به همراه سه وضعیت عملیاتی به عنوان ورودی مدل انتخاب گردیدند. در ادامه، برای تعیین مقدار هدف عمر مفید باقیمانده، از روش استهلاک تکه‌ای-خطی^{۱۵} استفاده شده است. در این روش، فرض می‌شود که موتور در ابتدای چرخه‌های عمر، مقدار RUL ثابت و حداکثری دارد و پس از عبور از یک نقطه آستانه مشخص (معمولاً ۱۲۵ چرخه)، مقدار RUL به صورت خطی کاهش یافته و در نهایت به صفر می‌رسد. این رابطه به صورت معادله (۱) تعریف می‌شود

$$RUL = \begin{cases} t, & t < RUL_{max} \\ RUL_{max}, & t \geq RUL_{max} \end{cases} \quad (1)$$

در معادله (۱)، t برابر با تعداد چرخه‌های عملیاتی موتور تا وضعیت خرابی کامل و RUL_{max} حداکثر مقدار RUL است که در این مقاله ۱۲۵ در نظر گرفته شده است.

پس از مرحله برچسب‌گذاری، داده‌های حسگرها با استفاده از روش نرمال‌سازی استاندارد نرمال شدند. این روش، میانگین داده‌ها را به صفر و انحراف معیار آن‌ها را به یک تبدیل می‌کند. رابطه نرمال‌سازی برای هر حسگر به صورت زیر تعریف می‌شود

$$x_i^j = \frac{x_i^j - \mu^j}{\sigma^j} \quad (2)$$

در رابطه‌ی ۲، μ^j میانگین و σ^j انحراف معیار خوانش‌های حسگر j -ام در کل نمونه‌ها است.

برای زیرمجموعه داده‌های پیچیده FD002 و FD004 که دارای شش وضعیت عملیاتی هستند، به منظور نرمال‌سازی دقیق‌تر و حفظ اطلاعات مفید هر وضعیت عملیاتی، پیش از نرمال‌سازی، داده‌ها با استفاده از الگوریتم خوشه‌بندی K-میانگین بر اساس پارامترهای وضعیت عملیاتی، خوشه‌بندی شده و سپس فرایند نرمال‌سازی بر اساس میانگین و انحراف معیار محاسبه شده برای هر خوشه به صورت جداگانه اعمال شده است.

در نهایت، برای تبدیل داده‌های سری زمانی به نمونه‌های آموزشی قابل استفاده برای شبکه‌های عمیق، از روش پنجره لغزان^{۱۶} استفاده شده است. طول پنجره بر اساس مطالعات تجربی، مقدار ۶۴ در نظر گرفته شده است تا شبکه

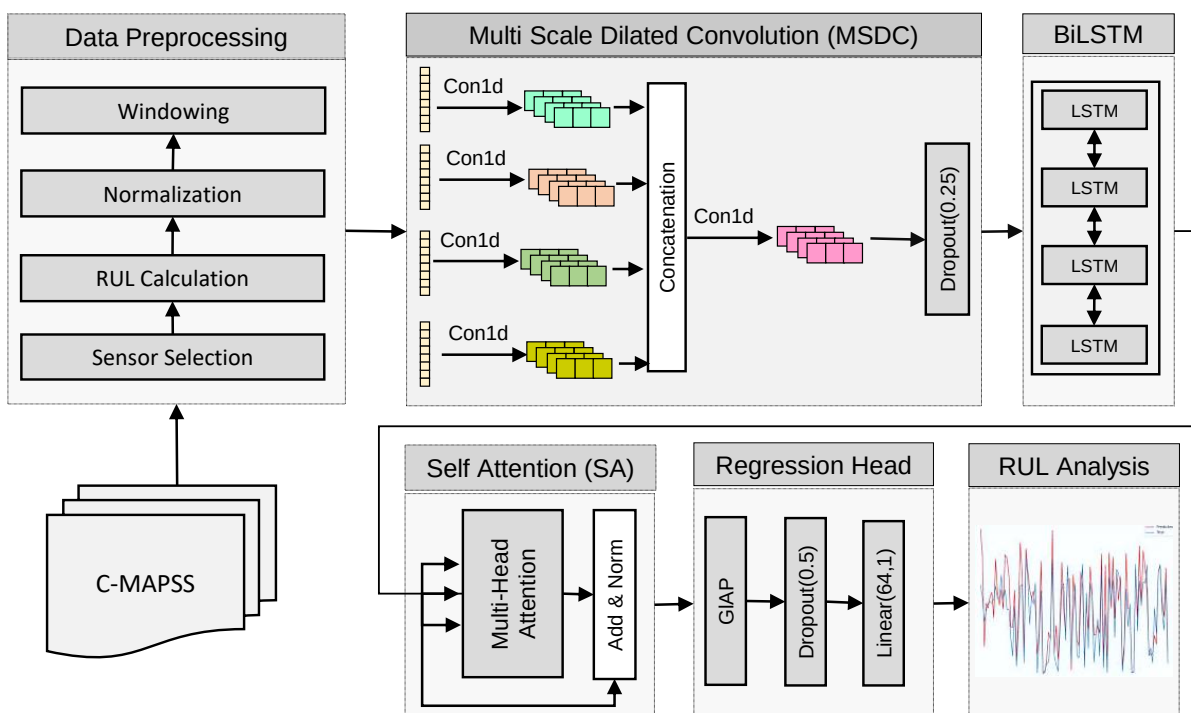
^{۱۵} Piecewise Linear Degradation Model

^{۱۶} Sliding Window

قادر به یادگیری وابستگی‌های زمانی بلندمدت باشد.

۲.۲ معماری مدل پیشنهادی

در این پژوهش، برای پیش‌بینی عمر مفید باقیمانده، مدل ترکیبی جدیدی پیشنهاد شده است که از سه مؤلفه اصلی تشکیل می‌شود: استخراج ویژگی‌های محلی با استفاده از لایه‌های پیچشی در مقیاس‌های مختلف، مدل‌سازی وابستگی‌های زمانی با استفاده از لایه‌های حافظه کوتاه‌مدت-بلندمدت دوطرفه و مکانیزم توجه چندسره‌منظور مدل‌سازی اهمیت چرخه‌های عملیاتی. معماری مدل ارائه شده در شکل ۱ نشان داده شده است.



شکل ۱: معماری مدل پیشنهادی

در ابتدا، داده‌های ورودی شامل اطلاعات ۱۴ حسگر و سه پارامتر وضعیت کنترلی (جمعاً ۱۷ ویژگی) از طریق چهار لایه پیچشی موازی پردازش می‌شوند. هر لایه پیچشی شامل ۳۲ هسته^{۱۷} با اندازه ۳ و نرخ‌های اتساع^{۱۸} ۱، ۲، ۴ و ۸ در نظر گرفته شده است. به‌منظور کمک به پایداری آموزش مدل و مدل‌سازی روابط غیرخطی، خروجی

^{۱۷} Kernel

^{۱۸} Dilation

هر لایه با استفاده از لایه نرمال سازی دسته‌ای^{۱۹} پردازش شده و تابع فعال ساز ReLU روی آن اعمال می‌شود. سپس، خروجی چهار لایه پیمشی در بعد کانال با یکدیگر الحاق شده و نمایشی چندمقیاسی با ۱۲۸ کانال را تشکیل می‌دهد. در ادامه، برای ترکیب و فشردن سازی این ویژگی‌ها، یک لایه پیمشی با ۶۴ هسته و اندازه ۱ روی خروجی اعمال شده و نمایشی با ۶۴ کانال شکل می‌گیرد.

سپس جهت مدل سازی وابستگی های زمانی، ویژگی های استخراج شده از بلوک قبلی با استفاده از یک لایه LSTM دوطرفه پردازش می‌شوند. این لایه به مدل اجازه می‌دهد تا در مدل سازی روند استهلاک، هم وابستگی های زمانی گذشته و هم آینده را در نظر بگیرد. به منظور تأکید بیشتر بر چرخه های عملیاتی مهم، یک لایه توجه چندسر روی خروجی لایه LSTM اعمال می‌شود. برای پایداری سازی فرایند آموزش مدل، پس از لایه توجه، از اتصال باقیمانده^{۲۰} و نرمال سازی لایه ای^{۲۱} استفاده شده است. در گام نهایی، با میانگین گیری سراسری در راستای بعد زمانی از ویژگی های استخراج شده، یک بردار ویژگی با طول ثابت ۶۴ به دست می‌آید. این بردار به عنوان ورودی به لایه رگرسیون خطی داده می‌شود تا مقدار RUL متناظر با ورودی تعیین شود.

۳ نتایج تجربی

۱.۳ محیط پیاده سازی و ابرپارامترهای مدل ها

آموزش و اعتبارسنجی مدل های یادگیری عمیق پیشنهادی با استفاده از زبان برنامه نویسی پایتون (نسخه ۱۲/۱۲/۳) و چارچوب یادگیری عمیق پایتورچ (نسخه ۰/۸/۲) در بستر گوگل کولب انجام شده است. منابع محاسباتی موجود در این محیط شامل پردازنده مرکزی Intel(R) Xeon(R) با فرکانس ۲۰/۲ گیگاهرتز، ۷/۱۲ گیگابایت حافظه RAM و پردازنده گرافیکی NVIDIA Tesla T4 با ۱۵ گیگابایت حافظه VRAM است. آموزش مدل همراه با بهینه ساز AdamW و تابع زیان Huber انجام شده است. ابرپارامترهای اصلی مورد استفاده در تمامی پیکربندی های آزمایشی در جدول ۱ خلاصه شده اند.

^{۱۹} Batch Normalization

^{۲۰} Residual Connection

^{۲۱} Layer Normalization

جدول ۱: ابرپارامترهای مدل پیشنهادی

مقدار	ابریارامتر
64	طول پنجره
64	اندازه دسته ^{۲۲}
50	تعداد دوره‌های آموزش ^{۲۳}
8 (FD002-FD004) / 1 (FD001)	تعداد سرهای توجه
AdamW	بهینه‌ساز
1×10^{-4}	نرخ یادگیری اولیه
Huber Loss ($\delta = 10$)	تابع زیان

۲.۳ معیارهای ارزیابی

برای ارزیابی عملکرد مدل پیشنهادی در پیش‌بینی عمر مفید باقیمانده، از دو معیار کلیدی استفاده شده است: جذر میانگین مربعات خطا (RMSE) و تابع امتیازدهی (Score).

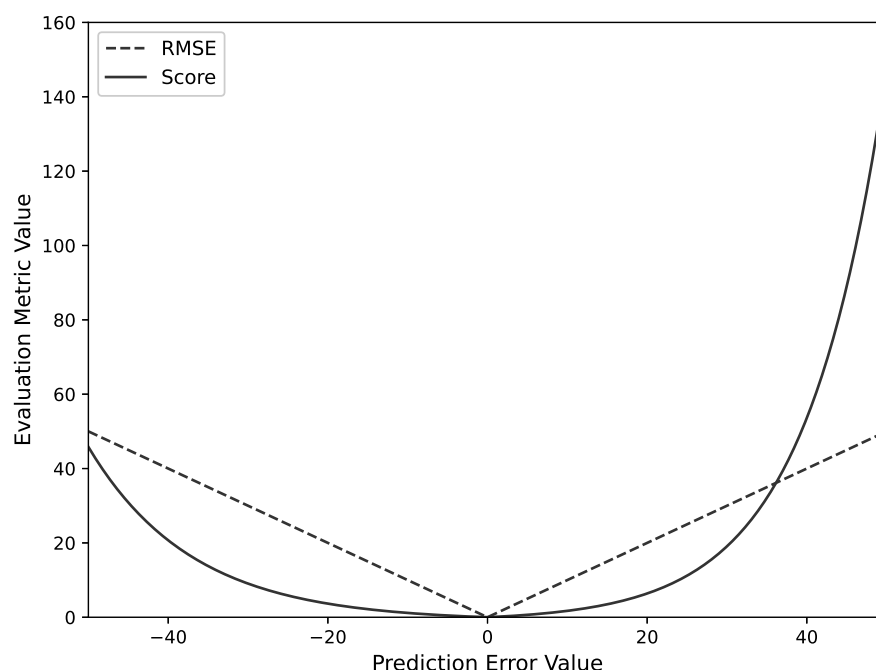
جذر میانگین مربعات خطا (RMSE): این معیار پرکاربرد، جذر میانگین مجذور تفاوت‌های بین مقادیر واقعی عمر مفید باقیمانده y_i و مقادیر پیش‌بینی شده $\hat{y}_i$ توسط مدل را طبق معادله (۱) محاسبه می‌کند. در ارزیابی تخمین عمر مفید باقیمانده، معیار RMSE به خطاهای پیش‌بینی زود هنگام و دیر هنگام وزن یکسانی اختصاص می‌دهد. این معیار طبق رابطه (۲) تعریف می‌شود.

$$RMSE = \sqrt{\frac{\sum_{i=1}^N (y_i - \hat{y}_i)^2}{N}} \quad (1)$$

تابع امتیازدهی (Score): بر خلاف معیار RMSE، تابع امتیازدهی ضرایب جریمه متفاوتی برای نوع خطای پیش‌بینی در نظر می‌گیرد. با توجه به اینکه پیش‌بینی دیر هنگام می‌تواند عواقب شدیدی در نگهداری پیشگیرانه داشته باشد، تابع امتیاز دهی به پیش‌بینی دیر هنگام ضریب جریمه بیشتری اختصاص می‌دهد. این معیار طبق رابطه (۲) تعریف می‌شود.

$$Score = \sum_{i=1}^N s_i, \quad s_i = \begin{cases} e^{\frac{y_i - \hat{y}_i}{13}} - 1, & \hat{y}_i < y_i \\ e^{\frac{\hat{y}_i - y_i}{10}} - 1, & \hat{y}_i \geq y_i \end{cases} \quad (2)$$

تفاوت بین دو معیار ارزیابی RMSE و Score در شکل ۲ نشان داده شده است. مقادیر کمتر هر دو معیار، نشان‌دهنده دقت بالاتر و عملکرد بهتر مدل در پیش‌بینی عمر مفید باقیمانده است.



شکل ۲: مقایسه معیارهای RMSE و تابع امتیازدهی ناسا

۳.۳ نتایج مدل پیشنهادی بر روی مجموعه داده توربو فن ناسا

در این بخش، عملکرد مدل پیشنهادی روی چهار زیرمجموعه داده FD001، FD002، FD003 و FD004 از مجموعه داده C-MAPSS ناسا ارزیابی شده است. نتایج کمی مدل بر اساس دو معیار RMSE و Score در جدول ۲ ارائه شده است.

جدول ۲: نتایج مدل پیشنهادی بر روی مجموعه داده C-MAPSS

میانگین	FD004	FD003	FD002	FD001	معیار
12.80	14.07	11.90	13.31	11.90	RMSE
581.06	1041.73	240.21	833.49	208.81	Score

همان‌طور که در جدول ۲ مشاهده می‌شود، بهترین عملکرد مدل در زیرمجموعه‌های FD001 و FD003 با کمترین میزان خطا حاصل شده است. این دو زیرمجموعه دارای وضعیت عملیاتی ساده‌تر (یک وضعیت عملیاتی) هستند.

در مقابل، زیرمجموعه‌های FD002 و FD004 به دلیل داشتن ۶ وضعیت عملیاتی مختلف و پیچیدگی بیشتر، خطای بالاتری را نشان می‌دهند.

نتایج پیش‌بینی مدل پیشنهادی برای چهار زیرمجموعه داده در شکل ۳ نشان داده شده است. در این شکل، مقادیر واقعی عمر مفید باقیمانده با رنگ سیاه و مقادیر پیش‌بینی شده با رنگ قرمز ترسیم شده‌اند. محور افقی نشان‌دهنده شماره موتورها در مجموعه آزمایش هر زیرمجموعه داده و محور عمودی نشان‌دهنده مقدار عمر مفید باقیمانده برحسب چرخه‌های عملیاتی است.

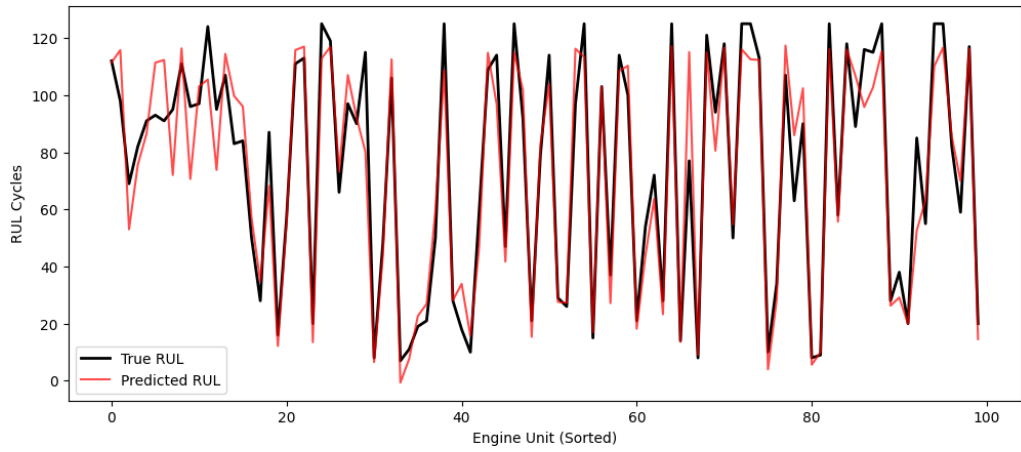
۴.۳ بررسی نقش مؤلفه‌های مدل پیشنهادی

به منظور بررسی سهم هر یک از مؤلفه‌های معماری مدل پیشنهادی، یک تحلیل مؤلفه‌محور با استفاده از معیارهای RMSE و تابع امتیازدهی ناسا روی چهار زیرمجموعه از مجموعه داده توربوفن ناسا (C-MAPSS) انجام شده است. در این مطالعه، عملکرد مدل کامل MSDC-BiLSTM-Attention با سه نوع مدل کاهش یافته مقایسه شده است: نخست، مدل بدون لایه BiLSTM (یعنی MSDC-Attn)؛ دوم، مدل بدون مکانیزم توجه (یعنی MSDC-BiLSTM) و سوم، مدل بدون بلوک استخراج ویژگی چندمقیاسی. نتایج این مطالعه در جداول ۳ و ۴ ارائه شده است.

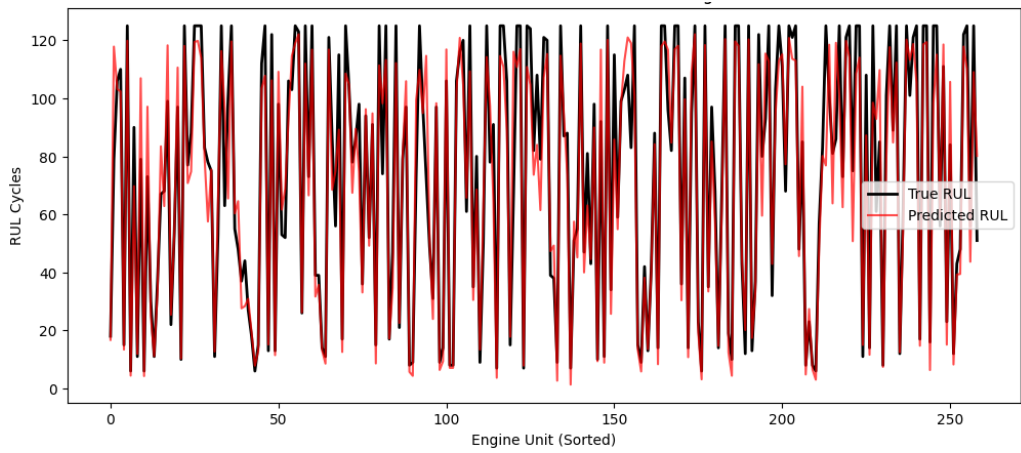
طبق نتایج به دست آمده، حذف هر یک از مؤلفه‌های ساختاری در اکثر سناریوها منجر به کاهش محسوس دقت پیش‌بینی و افزایش خطای مدل شده است. این روند کلی نشان می‌دهد که تلفیق ویژگی‌های چندمقیاسی، مدل‌سازی وابستگی‌های زمانی و مکانیزم توجه، نقشی مکمل و حیاتی در بهبود عملکرد نهایی مدل ایفا می‌کنند و حذف هر کدام، عملکرد متوازن مدل را تحت تأثیر قرار می‌دهد.

جدول ۳: نتایج تحلیل مؤلفه‌ای مدل بر اساس معیار RMSE

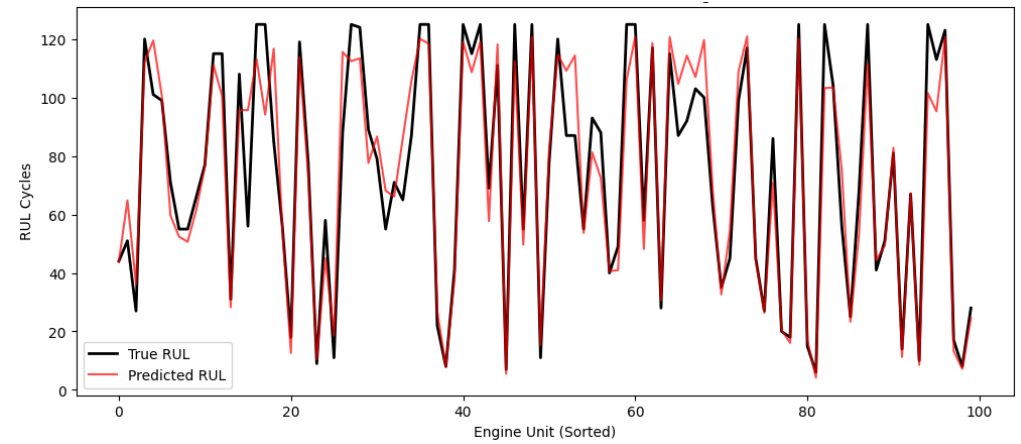
FD004	FD003	FD002	FD001	مدل
14.07	11.90	13.31	11.90	MSDC-BiLSTM-Attention
14.50	13.46	13.03	15.92	BiLSTM-Attention
16.91	17.86	24.89	20.36	MSDC-Attention
22.27	28.82	21.89	39.45	MSDC-BiLSTM



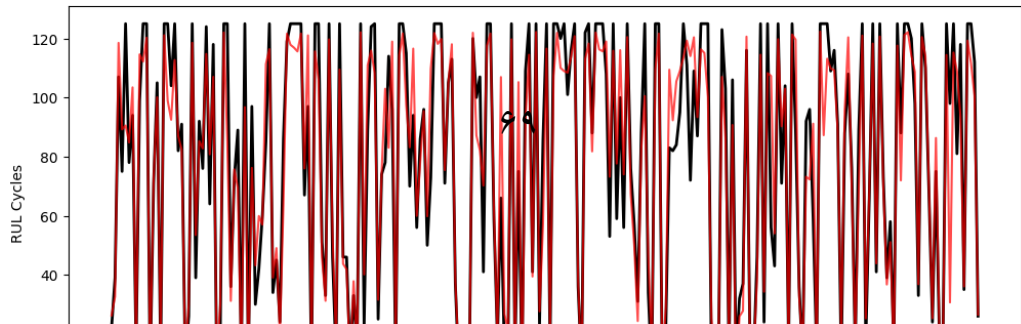
FD001



FD002



FD003



جدول ۴: نتایج تحلیل مؤلفه‌ای مدل بر اساس امتیاز ناسا (Score)

FD004	FD003	FD002	FD001	مدل
1041.73	240.21	833.49	208.81	MSDC-BiLSTM-Attention
1224.26	723.24	695.20	570.23	BiLSTM-Attention
1782.98	1409.39	11292.55	957.88	MSDC-Attention
3648.98	1545.33	4460.18	4717.87	MSDC-BiLSTM

۵.۳ مقایسه با پژوهش‌های اخیر

به منظور ارائه نمای کلی از زمینه پژوهش و جایگاه مدل پیشنهادی، جداول ۵ و ۶ نتایج تعدادی از مطالعات اخیر روی مجموعه داده توربو فن ناسا با استفاده از روش‌های یادگیری عمیق را نشان می‌دهند. همان‌طور که از این جداول مشاهده می‌شود، در سال‌های اخیر به کارگیری یادگیری عمیق برای پیش‌بینی عمر مفید باقیمانده، نتایجی چشمگیر و روبه‌بهبود داشته است. در مقایسه با روش‌های ارائه شده، مدل پیشنهادی عملکرد رقابتی و در بسیاری از موارد برتری از خود نشان داده است. این برتری به‌ویژه در زیرمجموعه FD004 مشهود است؛ این زیرمجموعه به دلیل داشتن دو حالت خرابی و شش وضعیت عملیاتی متفاوت، پیچیده‌ترین و چالش‌برانگیزترین زیرمجموعه در میان چهارگانه مجموعه داده‌های C-MAPSS محسوب می‌شود. مدل پیشنهادی با ثبت مقدار RMSE برابر ۰۷/۱۴ و امتیاز برابر ۷۳/۱۰۴۱، عملکرد رضایت‌بخشی را در این سناریوی پیچیده ارائه داده است.

جدول ۵: مقایسه با روش‌های پیشین بر اساس معیار RMSE

FD004	FD003	FD002	FD001	سال	روش
26.54	12.34	20.32	12.49	2025	CNN-BiLSTM [1]
19.65	13.49	12.33	13.09	2025	HRP [2]
16.25	14.11	14.36	13.43	2025	A-DDF [3]
20.54	13.36	17.63	12.54	2024	URCNN-BiLSTM [4]
28.04	14.43	26.50	14.35	2023	CALAP [5]
14.07	11.90	13.31	11.90	—	مدل پیشنهادی

جدول ۶: مقایسه با روش‌های پیشین بر اساس امتیاز ناسا (Score)

FD004	FD003	FD002	FD001	سال	روش
3905.33	225.53	1575.88	267.08	2025	CNN-BiLSTM [1]
—	—	—	—	2025	HRP [2]
1273.72	402.74	936.76	296.76	2025	A-DDF [3]
3326	245	2354	239	2024	URCNN-BiLSTM [4]
12179.41	381.41	42579.42	309.32	2023	CALAP [5]
1041.73	240.21	833.49	208.81	—	مدل پیشنهادی

۴ نتیجه‌گیری

در این پژوهش، یک مدل یادگیری عمیق ترکیبی برای پیش‌بینی عمر مفید باقیمانده موتورهای توربوفن ارائه شده است. مدل پیشنهادی با تلفیق سه مؤلفه اصلی شامل استخراج ویژگی‌های چندمقیاسی با استفاده از لایه‌های پیچشی، مدل‌سازی وابستگی‌های زمانی با لایه حافظه کوتاه‌مدت-بلندمدت دوطرفه و مکانیزم توجه چندسر، قادر است الگوهای محلی و وابستگی‌های بلندمدت را به‌طور همزمان از داده‌های حاصل از چندین حسگر استخراج کند. ارزیابی مدل بر روی مجموعه داده استاندارد C-MAPSS نشان داد که روش پیشنهادی عملکردی رقابتی با مدل‌های پیشرفته موجود دارد. نتایج تحلیل مؤلفه‌ای نیز نشان داد که هر سه مؤلفه اصلی معماری پیشنهادی نقش مهمی در عملکرد نهایی مدل ایفا می‌کنند.

با وجود نتایج مطلوب، پژوهش حاضر با محدودیت‌هایی نیز مواجه است. نخست، مجموعه داده C-MAPSS یک داده شبیه‌سازی شده است و داده‌های عملیاتی واقعی از موتورهای فیزیکی در آن وجود ندارد؛ از این رو، عملکرد مدل بر روی داده‌های واقعی ممکن است متفاوت باشد. دوم، مدل پیشنهادی برای دستیابی به عملکرد مطلوب به حجم بالایی از داده‌های برچسب‌دار نیاز دارد؛ در حالی که در بسیاری از کاربردهای صنعتی واقعی، تهیه چنین داده‌هایی پرهزینه و گاهی غیرممکن است. سوم، تعمیم‌پذیری مدل آموزش‌دیده روی این مجموعه داده به موتورهایی با پیکربندی متفاوت یا شرایط عملیاتی جدید نیازمند بررسی بیشتر است.

در مسیرهای آتی پژوهش، استفاده از روش‌های یادگیری نیمه‌نظارتی و یادگیری انتقالی به منظور کاهش وابستگی به داده‌های برچسب‌دار و بهبود تعمیم‌پذیری مدل دنبال خواهد شد. همچنین، ارزیابی مدل بر روی داده‌های عملیاتی واقعی و توسعه روش‌های فشرده‌سازی مدل برای استقرار در دستگاه‌های با محدودیت محاسباتی از دیگر جهت‌گیری‌های آینده است.

- [1] Yu, B., Guo, H. and Shi, J. (2025), *Remaining useful life prediction based on hybrid cnn-bilstm model with dual attention mechanism*, International Journal of Electrical Power & Energy Systems, vol. 172, p. 111152.
- [2] Niu, T., Xu, Z., Luo, H. and Zhou, Z. (2025), *Hybrid gaussian process regression with temporal feature extraction for partially interpretable remaining useful life interval prediction in aeroengine prognostics*, Scientific Reports, vol. 15, no. 1, p. 11057.
- [3] Cheng, X., Chaw, J. K., Sahrani, S., Ang, M. C., Gunasekaran, S. S., Mahmoud, M. A., Badioze Zaman, H., Zhao, Y. and Ren, F. (2025), *An adaptive dual distillation framework for efficient remaining useful life prediction*, Complex & Intelligent Systems, vol. 11, no. 6, p. 253.
- [4] Yan, X., Liang, W. and Sun, S. (2024), *An improved method for predicting the remaining useful life using a spatial-temporal feature extraction network with attention mechanism*, IEEE Access, vol. 12, pp. 66587–66604.
- [5] Wu, M., Ye, Q., Mu, J., Fu, Z. and Han, Y. (2023), *Remaining useful life prediction via a data-driven deep learning fusion model-calap*, IEEE Access, vol. 11, pp. 112085–112096.



تصحیح سوگیری سانسور وابسته در تحلیل قابلیت اعتماد مبتنی بر مدل کاکس

گلی ناری، ح^۱ خراشادی زاده، م^۲ محتشمی برزادران، غ^۳

^{۱،۲} گروه آمار، دانشکده علوم ریاضی و آمار، دانشگاه بیرجند

^۳ گروه آمار، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد

چکیده

در تحلیل قابلیت اعتماد، وقوع سانسور وابسته که در آن مکانیزم سانسور به صورت سیستماتیک با زمان خرابی مرتبط است، اعتبار روش‌های استاندارد استنباط آماری را با چالش جدی مواجه می‌سازد. این وضعیت به ویژه در حضور متغیرهای پنهانی که هم بر زمان شکست و هم بر مکانیزم سانسور تأثیر می‌گذارند، به برآوردهای اریب و نتایج گمراه‌کننده منجر می‌شود. این مقاله یک رهیافت نوین برای تصحیح سوگیری سانسور وابسته در مدل خطرات متناسب کاکس ارائه می‌دهد. نوآوری اصلی، ارائه یک متغیر جدید مبتنی بر آنتروپی مانده تجمعی است که از طریق مانده‌های یک مدل کمکی تعریف می‌شود و عدم قطعیت باقیمانده در مدل بقا را کمی می‌کند. این متغیر در قالب یک الگوریتم دومرحله‌ای به کار گرفته می‌شود. در مرحله نخست، یک مدل کمکی برای استخراج اطلاعات متغیرهای پنهان برازش شده و وزن‌های تصحیح‌کننده با استفاده از یک مدل لجستیک توسعه یافته محاسبه می‌گردند. در مرحله دوم، مدل کاکس وزن‌دار با این وزن‌ها برازش می‌شود. یک مطالعه شبیه‌سازی گسترده با ۵۰۰

¹Speaker. hassangolinari@birjand.ac.ir

²m.khorashadizadeh@birjand.ac.ir

³grmohtashami@um.ac.ir

تکرار برای ارزیابی عملکرد روش اجرا شد. نتایج نشان می‌دهد که روش پیشنهادی در مقایسه با روش تعدیل‌نشده و روش وزن‌دهی معکوس احتمال سانسور استاندارد، سوگیری و ریشه میانگین مربعات خطای کمتری دارد. این بهبود، کارایی چارچوب پیشنهادی را در تصحیح سوگیری سانسور وابسته تأیید می‌کند و مسیر را برای استنباط آماری معتبرتر در تحلیل قابلیت اعتماد هموار می‌سازد.

کلمات کلیدی: تحلیل بقا، سانسور وابسته، آنتروپی مانده تجمعی، مدل خطرات متناسب کاکس، قابلیت اعتماد.

۱ پیشگفتار

تحلیل بقا یکی از موضوعات اصلی در حوزه قابلیت اعتماد سیستم‌ها و مطالعات پزشکی به شمار می‌رود. این شاخه از آمار به مطالعه زمان تا وقوع یک رویداد خاص، نظیر خرابی یک قطعه صنعتی یا فوت یک بیمار، می‌پردازد. در کاربردهای مهندسی قابلیت اعتماد، برآورد دقیق تابع بقا و پارامترهای مرتبط با آن، مبنای تصمیم‌گیری‌های حیاتی نظیر تعیین برنامه‌های نگهداری و تعمیرات پیشگیرانه، محاسبه دوره‌های ضمانت، و ارزیابی ریسک خرابی سیستم‌ها را تشکیل می‌دهد. با این حال، ماهیت داده‌های طول عمر اغلب به گونه‌ای است که زمان دقیق رویداد برای تمام واحدهای تحت مطالعه مشاهده نمی‌شود. این پدیده که سانسور شدگی^۱ نام دارد، چالش‌های بنیادینی را برای استنباط آماری معتبر ایجاد می‌کند.

در تحلیل بقا، روش‌های استاندارد نظیر برآوردگر کاپلان-مایر و مدل خطرات متناسب کاکس^۲ به طور گسترده مورد استفاده قرار می‌گیرند. مدل کاکس که توسط کاکس [۲] معرفی شد، به دلیل انعطاف‌پذیری در مدل‌سازی نیمه‌پارامتری تابع خطر، به یکی از پراستنادترین ابزارهای تحلیل بقا تبدیل شده است. با این وجود، این روش‌ها عموماً بر فرض سانسور مستقل^۳ بنا شده‌اند. بر اساس این فرض، مکانیزم سانسور هیچ وابستگی سیستماتیکی با زمان شکست ندارد و صرفاً یک فرآیند تصادفی مستقل است. با این حال، در بسیاری از مسائل دنیای واقعی، این فرض بنیادین نقض می‌شود. هنگامی که مکانیزم سانسور به صورت سیستماتیک با زمان شکست یا عوامل خطر مرتبط باشد، با پدیده سانسور وابسته^۴ مواجه هستیم. برای نمونه، در یک آزمون طول عمر تسریع‌یافته^۵، قطعاتی که تحت تنش‌های بالاتری قرار دارند، هم زمان خرابی کوتاه‌تری دارند و هم به دلیل خرابی‌های ثانویه، احتمال خروج زودهنگام آن‌ها از آزمون بیشتر است. در چنین شرایطی، فرض سانسور مستقل اعتبار خود را از دست می‌دهد و استفاده از

¹Censoring

²Cox Proportional Hazards Model

³Independent Censoring

⁴Dependent Censoring

⁵Accelerated Life Test

روش‌های استاندارد به برآوردهای اریب و استنباط‌های گمراه‌کننده منجر می‌شود.

برای مقابله با این چالش، رویکردهای متعددی در ادبیات آماری توسعه یافته است. یکی از پرکاربردترین این رویکردها، روش وزن‌دهی معکوس احتمال سانسور^۶ است. در این روش، احتمال عدم سانسور برای هر مشاهده با استفاده از یک مدل رگرسیونی بر اساس متغیرهای کمکی مشاهده‌پذیر برآورد می‌شود. سپس، وزن‌های معکوس این احتمالات در برآورد پارامترهای مدل بقا به کار گرفته می‌شوند تا سوگیری ناشی از سانسور تصحیح گردد. با این حال، محدودیت اساسی این رویکرد در آن است که تنها قادر به کنترل اثر متغیرهای کمکی مشاهده‌پذیر است. پرینس و بومرت [۴] در پژوهشی جدید به بررسی چالش‌های برآورد وزن‌های IPCW پرداخته و نشان داده‌اند که در حضور متغیرهای پنهان، این روش همچنان با سوگیری قابل توجهی مواجه است. در شرایطی که عوامل مخدوش‌گر مشاهده‌نشده^۷ وجود داشته باشند که هم بر زمان شکست و هم بر مکانیزم سانسور تأثیر بگذارند، روش IPCW کارایی خود را از دست می‌دهد. این وضعیت که معادل مفهوم داده‌های گمشده غیرتصادفی^۸ در ادبیات آماری است، یک شکاف پژوهشی جدی را آشکار می‌سازد. وانگ و همکاران [۵] نیز با ارائه مدل‌های جدید برای داده‌های بقا با ساختار وابسته، بر اهمیت توسعه روش‌هایی که بتوانند وابستگی‌های پیچیده را مدیریت کنند، تأکید نموده‌اند.

برای پر کردن این شکاف، ما در این مقاله یک چارچوب استنباطی جدید را معرفی می‌کنیم که هسته مرکزی آن، استفاده از معیارهای اطلاعاتی برای سنجش و مدیریت عدم قطعیت باقیمانده^۹ در مدل است. مفهوم کلیدی که ما به کار می‌گیریم، آنتروپی مانده تجمعی^{۱۰} نام دارد که توسط رائو و همکاران [۱] به عنوان یک معیار اطلاعاتی جایگزین برای آنتروپی شانون معرفی شد. مزیت CRE نسبت به آنتروپی شانون، پایداری بیشتر در برآورد و تعریف آن بر اساس تابع توزیع تجمعی است. این ویژگی، CRE را به ابزاری مناسب برای تحلیل داده‌های بقا تبدیل می‌کند، زیرا تابع بقا که نقشی محوری در این حوزه دارد، مستقیماً از تابع توزیع تجمعی مشتق می‌شود.

نوآوری اصلی این مقاله، ارائه یک الگوریتم دومرحله‌ای است که از CRE به عنوان پل ارتباطی بین مدل بقا و مدل سانسور استفاده می‌کند. در مرحله نخست، ما یک متغیر جدید به نام Z_i را بر اساس مانده‌های مارتینگل^{۱۱} حاصل از یک مدل کمکی محاسبه می‌کنیم. مانده‌های مارتینگل که توسط ترنو و گرامبش [۳] به تفصیل تشریح شده‌اند، تفاوت بین وضعیت رویداد مشاهده‌شده و میزان خطر پیش‌بینی‌شده توسط مدل را کمی می‌کنند. متغیر Z_i به صورت $Z_i = |M_i|$ تعریف می‌شود و نماینده‌ای برای آنتروپی مانده تجمعی در سطح هر مشاهده است. این متغیر، عدم قطعیت

⁶Inverse Probability of Censoring Weighting (IPCW)

⁷Unobserved Confounders

⁸Missing Not At Random (MNAR)

⁹Residual Uncertainty

¹⁰Cumulative Residual Entropy (CRE)

¹¹Martingale Residuals

باقیمانده‌ای را که توسط متغیرهای کمکی مدل نشده است، کمی می‌کند. در مرحله دوم، این متغیر به عنوان یک کمکی جدید وارد مدل لجستیک سانسور می‌شود تا وزن‌های دقیق‌تری برای تصحیح سوگیری محاسبه گردد. سپس، یک مدل کاکس وزن‌دار با این وزن‌های بهبودیافته برازش می‌شود. این فرآیند، برآورد پارامترهای مدل بقا و وزن‌های تصحیح‌کننده را به طور هم‌زمان بهبود می‌بخشد.

ساختار ادامه مقاله به شرح زیر است. در بخش دوم، مبانی نظری شامل تعریف CRE، مدل کاکس، مانده‌های مارتینگل، و نحوه انطباق آن‌ها با یکدیگر را تشریح می‌کنیم. در بخش سوم، الگوریتم پیشنهادی و ویژگی‌های نظری آن از جمله قضیه سازگاری برآوردگر را به تفصیل ارائه می‌کنیم. بخش چهارم به یک مطالعه شبیه‌سازی گسترده برای ارزیابی عملکرد روش اختصاص دارد و در نهایت، در بخش پنجم به نتیجه‌گیری و بحث پیرامون یافته‌ها می‌پردازیم.

۲ مبانی نظری

در این بخش، ابتدا مفهوم آنتروپی مانده تجمعی به عنوان یک معیار اطلاعاتی پایدار معرفی می‌شود. سپس، مدل خطرات متناسب کاکس و مانده‌های مارتینگل به عنوان ابزار تشخیصی استاندارد در تحلیل بقا تشریح می‌گردند. در نهایت، تعریف جدید متغیر پیشنهادی بر اساس تلفیق این دو مفهوم ارائه می‌شود.

۱.۲ آنتروپی مانده تجمعی

رائو و همکاران [۱] مفهوم آنتروپی مانده تجمعی^{۱۲} را به عنوان یک معیار اطلاعاتی جایگزین برای آنتروپی شانون معرفی کردند. برای یک متغیر تصادفی نامنفی Y با تابع بقای $\bar{F}(y) = P(Y > y)$ ، این معیار به صورت زیر تعریف می‌شود

$$\mathcal{E}(Y) = - \int_0^{\infty} \bar{F}(y) \log \bar{F}(y) dy. \quad (1)$$

برخلاف آنتروپی شانون که بر پایه تابع چگالی احتمال تعریف می‌شود، CRE از تابع توزیع تجمعی بهره می‌گیرد. این ویژگی دو مزیت عمده برای CRE به همراه دارد. نخست، تابع توزیع تجمعی کمیتی پایدارتر از تابع چگالی احتمال است و برآورد ناپارامتری آن از طریق تابع توزیع تجمعی تجربی با دقت بیشتری امکان‌پذیر می‌باشد. دوم، CRE تحت تبدیل‌های یکنوای پیوسته رفتاری پایا دارد، به این معنا که برای هر $a > 0$ و b حقیقی، رابطه $\mathcal{E}(aY + b) = a\mathcal{E}(Y)$ برقرار است. این ویژگی‌ها، CRE را به ابزاری مناسب برای تحلیل داده‌های بقا تبدیل می‌کند، زیرا در این حوزه، تابع بقا که مکمل تابع توزیع تجمعی است، نقشی محوری ایفا می‌کند.

¹²Cumulative Residual Entropy (CRE)

در بستر مدل‌های رگرسیونی، CRE مانده‌های مدل، یعنی $\mathcal{E}(\epsilon)$ که در آن $\epsilon = Y - \mathbb{E}(Y|X)$ ، به عنوان سنج‌ای برای عدم قطعیت باقیمانده^{۱۳} تفسیر می‌شود. مقادیر بزرگ CRE مانده‌ها نشان می‌دهد که حتی پس از شرطی کردن روی متغیرهای کمکی X ، عدم قطعیت قابل توجهی در پیش‌بینی متغیر پاسخ باقی مانده است. این عدم قطعیت تبیین نشده می‌تواند نشانه‌ای از وجود متغیرهای پنهانی باشد که بر متغیر پاسخ تأثیر می‌گذارند اما در مدل لحاظ نشده‌اند.

۲.۲ مدل خطرات متناسب کاکس

مدل خطرات متناسب کاکس [۲] پرکاربردترین مدل در تحلیل بقا است. در این مدل، تابع خطر برای فرد i ام با بردار متغیرهای کمکی X_i به صورت زیر در نظر گرفته می‌شود

$$h(t | X_i) = h_0(t) \exp(\beta^\top X_i), \quad (۲)$$

که در آن $h_0(t)$ تابع خطر پایه^{۱۴} و β بردار ضرایب مجهول است. تابع خطر پایه، میزان خطر را برای یک فرد با $X_i = 0$ توصیف می‌کند و بخش ناپارامتری مدل را تشکیل می‌دهد. جمله $\exp(\beta^\top X_i)$ نیز بخش پارامتری مدل است که اثر متغیرهای کمکی را به صورت ضربی بر تابع خطر پایه اعمال می‌کند.

برآورد پارامترهای β در مدل کاکس از طریق بیشینه‌سازی تابع درست‌نمایی جزئی^{۱۵} انجام می‌شود. تابع درست‌نمایی جزئی برای یک نمونه به اندازه n به صورت زیر تعریف می‌شود

$$L(\beta) = \prod_{i=1}^n \left[\frac{\exp(\beta^\top X_i)}{\sum_{j \in \mathcal{R}(t_i)} \exp(\beta^\top X_j)} \right]^{\delta_i}, \quad (۳)$$

که در آن δ_i نشانگر وقوع رویداد برای فرد i ام و $\mathcal{R}(t_i)$ مجموعه افراد در معرض خطر در زمان t_i است. تحت شرایط منظم، برآوردگر حاصل از بیشینه‌سازی درست‌نمایی جزئی، سازگار و دارای توزیع مجانبی نرمال است [۲].

۳.۲ مانده‌های مارتینگل

ترنو و گرامبش [۳] مانده‌های مارتینگل^{۱۶} را به عنوان یکی از مهم‌ترین ابزارهای تشخیصی در مدل کاکس معرفی کرده‌اند. مانده مارتینگل برای مشاهده i ام به صورت زیر تعریف می‌شود

$$M_i = \delta_i - \hat{H}_0(t_i) \exp(\hat{\beta}^\top X_i), \quad (۴)$$

^{۱۳}Residual Uncertainty

^{۱۴}Baseline Hazard Function

^{۱۵}Partial Likelihood

^{۱۶}Martingale Residuals

که در آن $\hat{H}(\cdot)$ برآورد تابع خطر پایه تجمعی^{۱۷} است. مانده‌های مارتینگل، تفاوت بین وضعیت رویداد مشاهده‌شده (δ_i) و میزان خطر تجمعی پیش‌بینی‌شده توسط مدل را کمی می‌کنند. دامنه این مانده‌ها در بازه $[-\infty, 1]$ قرار دارد. زمانی که مدل به خوبی برازش شده باشد، این مانده‌ها دارای میانگین صفر هستند و انتظار می‌رود که توزیع آن‌ها متقارن و بدون الگوی خاصی باشد.

مانده‌های مارتینگل خاصیت مفید دیگری نیز دارند. در صورتی که یک متغیر پنهان U وجود داشته باشد که بر زمان شکست تأثیر بگذارد اما در مدل کاکس لحاظ نشده باشد، اثر این متغیر در مانده‌های مارتینگل منعکس می‌شود. به عبارت دیگر، قدر مطلق بزرگ مانده مارتینگل برای یک مشاهده نشان‌دهنده آن است که مدل نتوانسته است زمان شکست آن مشاهده را به خوبی پیش‌بینی کند. این ویژگی، مبنای اصلی ایده ما برای استفاده از مانده‌های مارتینگل در ساخت متغیر CRE است.

۴.۲ تعریف متغیر پیشنهادی برای داده‌های بقا

نوآوری اصلی این مقاله، ارائه یک تعریف جدید از متغیر مبتنی بر آنتروپی مانده تجمعی برای داده‌های بقا است. با الهام از مفهوم CRE و با توجه به ویژگی‌های مانده‌های مارتینگل، ما متغیر Z_i را برای هر مشاهده به صورت زیر تعریف می‌کنیم

$$Z_i = |M_i|, \quad (5)$$

که در آن M_i مانده مارتینگل حاصل از یک مدل کاکس اولیه است. انتخاب قدر مطلق مانده مارتینگل، انحراف مدل از وضعیت ایده‌آل را بدون توجه به جهت آن اندازه‌گیری می‌کند. مقادیر بزرگ Z_i نشان‌دهنده مشاهداتی هستند که مدل بقا نتوانسته است به خوبی آن‌ها را تبیین کند.

ارتباط نظری بین متغیر Z_i و مفهوم CRE از طریق تابع توزیع تجمعی مانده‌ها برقرار می‌شود. اگر توزیع مانده‌های مارتینگل را با $F_M(\cdot)$ نشان دهیم، آنگاه CRE این مانده‌ها به صورت $\mathcal{E}(M) = - \int \bar{F}_M(m) \log \bar{F}_M(m) dm$ تعریف می‌شود. متغیر $Z_i = |M_i|$ یک تقریب نقطه‌ای از سهم مشاهده i ام در این انتگرال است. به عبارت دیگر، Z_i نماینده‌ای برای عدم قطعیت باقیمانده در سطح هر مشاهده است و می‌تواند ردپای متغیرهای پنهانی را که بر زمان شکست تأثیر می‌گذارند اما در مدل کاکس لحاظ نشده‌اند، آشکار سازد.

این تعریف، پلی بین مفهوم نظری CRE و کاربرد عملی آن در تحلیل بقا ایجاد می‌کند. در بخش بعدی، نحوه استفاده از این متغیر در یک الگوریتم دومرحله‌ای برای تصحیح سوگیری سانسور وابسته تشریح می‌شود.

¹⁷Cumulative Baseline Hazard Function

۳ الگوریتم پیشنهادی و ویژگی‌های نظری

در این بخش، چارچوب عملیاتی روش پیشنهادی برای تصحیح سوگیری سانسور وابسته تشریح می‌شود. این چارچوب از دو مرحله اصلی تشکیل شده است. در مرحله نخست، یک مدل کمکی برای استخراج اطلاعات متغیرهای پنهان به کار گرفته می‌شود. در مرحله دوم، این اطلاعات وارد مدل سانسور شده و وزن‌های تصحیح‌کننده محاسبه می‌گردند. در ادامه، ویژگی‌های نظری برآوردگر حاصل مورد بررسی قرار می‌گیرد.

۱.۳ مدل کمکی برای استخراج اطلاعات متغیرهای پنهان

نخستین گام در الگوریتم پیشنهادی، برازش یک مدل کمکی^{۱۸} برای استخراج اطلاعات مربوط به متغیرهای پنهان است. فرض کنید داده‌های مشاهده‌شده شامل زمان‌های بقای Y_i ، نشانگرهای سانسور δ_i ، و متغیرهای کمکی X_i برای $i = 1, \dots, n$ باشند. ما یک مدل رگرسیون خطی لگاریتمی به صورت زیر بر روی داده‌ها برازش می‌دهیم

$$\log(Y_i + 1) = \gamma^\top X_i + \epsilon_i, \quad i = 1, \dots, n, \quad (1)$$

که در آن γ بردار ضرایب رگرسیونی و ϵ_i جمله خطا با میانگین صفر و واریانس ثابت است. انتخاب تبدیل لگاریتمی به دلیل ماهیت نامنفی زمان‌های بقا و چولگی مثبت توزیع آن‌ها صورت می‌گیرد. افزودن مقدار ثابت ۱ به زمان‌های بقا نیز برای اجتناب از مشکل لگاریتم صفر در مشاهدات با زمان بقای بسیار کوچک انجام می‌شود.

پس از برازش مدل کمکی، مانده‌های آن به صورت $\hat{\epsilon}_i = \log(Y_i + 1) - \hat{\gamma}^\top X_i$ محاسبه می‌شوند. این مانده‌ها حاوی اطلاعات ارزشمندی درباره متغیرهای پنهان هستند. اگر یک متغیر پنهان U وجود داشته باشد که بر زمان شکست تأثیر بگذارد اما در مدل کمکی لحاظ نشده باشد، اثر آن در مانده‌ها منعکس خواهد شد. بر این اساس، متغیر پیشنهادی Z_i به صورت زیر تعریف می‌شود

$$Z_i = |\hat{\epsilon}_i|. \quad (2)$$

متغیر Z_i نماینده‌ای برای آنتروپی مانده تجمعی در سطح هر مشاهده است. مقادیر بزرگ Z_i نشان‌دهنده مشاهداتی هستند که مدل کمکی نتوانسته است به خوبی آن‌ها را پیش‌بینی کند و این عدم قطعیت باقیمانده می‌تواند ناشی از وجود متغیرهای پنهان باشد.

¹⁸ Auxiliary Model

۲.۳ مدل سازی سانسور با متغیر پیشنهادی

دومین گام، وارد کردن متغیر Z_i به مدل سانسور است. در رویکرد استاندارد، IPCW احتمال عدم سانسور تنها بر اساس متغیرهای کمکی مشاهده پذیر X_i مدل سازی می شود. ما این مدل را با افزودن Z_i به عنوان یک متغیر کمکی جدید توسعه می دهیم. مدل لجستیک^{۱۹} پیشنهادی برای احتمال عدم سانسور به صورت زیر است

$$\text{logit}[P(\delta_i = 1 | \mathbf{X}_i, Z_i)] = \alpha_0 + \alpha_1^T \mathbf{X}_i + \alpha_2 Z_i, \quad (3)$$

که در آن α_0 عرض از مبدأ، α_1 بردار ضرایب متغیرهای کمکی مشاهده پذیر، و α_2 ضریب متغیر پیشنهادی Z_i است. گنجاندن Z_i در مدل سانسور به آن اجازه می دهد تا به طور غیرمستقیم اثر متغیرهای پنهان را بر مکانیزم سانسور کنترل کند. اگر متغیر پنهان U هم بر زمان شکست و هم بر سانسور تأثیر بگذارد، انتظار می رود که ضریب α_2 از نظر آماری معنادار باشد.

پس از برآورد پارامترهای مدل با استفاده از روش بیشینه سازی درست نمایی، وزن های تصحیح کننده برای هر مشاهده به صورت معکوس احتمال برآوردی عدم سانسور محاسبه می شوند

$$\hat{w}_i = \frac{1}{\hat{p}_i} = \frac{1}{\hat{P}(\delta_i = 1 | \mathbf{X}_i, Z_i)}. \quad (4)$$

این وزن ها، برخلاف وزن های IPCW استاندارد، از اطلاعات موجود در Z_i نیز بهره می برند و در نتیجه تصحیح دقیق تری از سوگیری سانسور وابسته ارائه می دهند.

۳.۳ الگوریتم دو مرحله ای تصحیح سوگیری

الگوریتم پیشنهادی برای تصحیح سوگیری سانسور وابسته در قالب دو مرحله به شرح زیر خلاصه می شود.
مرحله اول (گام وزن دهی). در این مرحله، وزن های تصحیح کننده محاسبه می شوند.

۱. مدل کمکی (۱) بر روی داده های مشاهده شده برازش داده می شود.

۲. مانده های مدل کمکی محاسبه شده و متغیر Z_i بر اساس رابطه (۲) به دست می آید.

۳. مدل لجستیک (۳) با ورود X_i و Z_i برازش می شود.

۴. وزن های تصحیح کننده $\hat{w}_i$ بر اساس رابطه (۴) محاسبه می گردند.

¹⁹Logistic Model

مرحله دوم (گام برازش مدل بقا). در این مرحله، مدل کاکس وزن دار برازش می شود.

۱. مدل خطرات متناسب کاکس با استفاده از وزن های $\hat{w}_i$ از طریق بیشینه سازی تابع درست نمایی جزئی وزن دار^{۲۰} برازش می شود

$$L_w(\beta) = \prod_{i=1}^n \left[\frac{\exp(\beta^\top \mathbf{X}_i)}{\sum_{j \in \mathcal{R}(t_i)} \hat{w}_j \exp(\beta^\top \mathbf{X}_j)} \right]^{\hat{w}_i \delta_i} \quad (5)$$

۲. برآورد نهایی $\hat{\beta}$ از بیشینه سازی این تابع حاصل می شود.

این الگوریتم، سوگیری ناشی از سانسور وابسته را از دو مسیر کاهش می دهد. نخست، وزن های تصحیح کننده دقیق تر، تأثیر مشاهدات با سانسور غیر تصادفی را تعدیل می کنند. دوم، ورود اطلاعات متغیرهای پنهان از طریق Z_i به مدل سانسور، بخشی از واریانس تبیین نشده را کنترل می کند.

۴.۳ قضیه سازگاری برآوردگر پیشنهادی

در این بخش، ویژگی های نظری برآوردگر پیشنهادی را بررسی می کنیم. قضیه زیر نشان می دهد که تحت شرایط منظم، برآوردگر حاصل از الگوریتم پیشنهادی یک برآوردگر سازگار^{۲۱} برای پارامتر β است.

قضیه ۱.۳ (سازگاری برآوردگر پیشنهادی). فرض کنید شرایط زیر برقرار باشند.

(الف) مدل بقا از نوع خطرات متناسب کاکس (۲) با پارامتر حقیقی β است.

(ب) متغیرهای کمکی $\mathbf{X}_i$ دارای گشتاورهای مرتبه دوم متناهی هستند.

(ج) مکانیزم سانسور از مدل لجستیک (۳) با Z_i تعریف شده در (۲) پیروی می کند.

(د) شرط نادیده گرفتنی^{۲۲} به صورت $\delta_i \perp Y_i \mid (\mathbf{X}_i, Z_i)$ برقرار است.

آنگاه برآوردگر $\hat{\beta}$ حاصل از بیشینه سازی تابع درست نمایی جزئی وزن دار (۵) یک برآوردگر سازگار برای β است، یعنی

$$\hat{\beta} \xrightarrow{P} \beta, \quad n \rightarrow \infty \text{ زمانی که} \quad (6)$$

²⁰Weighted Partial Likelihood

²¹Consistent Estimator

²²Ignorability Condition

برهان. برای اثبات این قضیه، نشان می‌دهیم که تحت شرایط بیان‌شده، تابع درست‌نمایی جزئی وزن‌دار (۵) به طور مجانبی معادل با تابع درست‌نمایی جزئی استاندارد در غیاب سانسور وابسته است.

گام نخست. تحت شرط (د)، توزیع توأم $(Y_i, \delta_i, \mathbf{X}_i, Z_i)$ در نمونه به صورت زیر تجزیه می‌شود

$$f(Y_i, \delta_i | \mathbf{X}_i, Z_i) = f(Y_i | \mathbf{X}_i, Z_i) \cdot P(\delta_i = 1 | \mathbf{X}_i, Z_i).$$

این تجزیه نشان می‌دهد که با وزن‌دهی معکوس توسط $P(\delta_i = 1 | \mathbf{X}_i, Z_i)$ ، امید ریاضی تابع درست‌نمایی وزن‌دار با امید ریاضی تابع درست‌نمایی در غیاب سانسور وابسته برابر می‌شود و در نتیجه، برآوردگر حاصل نااریب خواهد بود. **گام دوم.** بر اساس نظریه استاندارد برآوردگرهای ماکسیمم درست‌نمایی، اگر تابع درست‌نمایی به درستی مشخص شده باشد و شرایط منظم بودن شامل مشتق‌پذیری و کران‌داری ماتریس اطلاعات فیشر برقرار باشد، برآوردگر حاصل سازگار خواهد بود [۲].

گام سوم. از آنجا که Z_i از مانده‌های مدل کمکی (۱) محاسبه می‌شود و این مدل خود یک برآوردگر سازگار برای رابطه بین $\log Y$ و $\mathbf{X}$ است، مانده‌های $\hat{\epsilon}_i$ به مقادیر واقعی خطاها همگرا می‌شوند. در نتیجه، $Z_i = |\hat{\epsilon}_i|$ نیز به صورت مجانبی پایدار بوده و وزن‌های $\hat{w}_i$ به وزن‌های حاصل از مدل سانسور با Z_i واقعی همگرا هستند.

گام چهارم. با ترکیب نتایج فوق، امید ریاضی تابع درست‌نمایی جزئی وزن‌دار (۵) به طور مجانبی با امید ریاضی تابع درست‌نمایی جزئی استاندارد برای یک نمونه تصادفی از مدل کاکس برابر می‌شود. از آنجا که برآوردگر ماکسیمم درست‌نمایی جزئی استاندارد تحت این شرایط سازگار است، برآوردگر $\hat{\beta}$ نیز سازگار خواهد بود. ■ □

این قضیه تضمین می‌کند که با افزایش حجم نمونه، برآوردگر پیشنهادی به مقدار واقعی پارامتر β همگرا می‌شود. این ویژگی، اعتبار نظری روش پیشنهادی را برای استنباط آماری در حضور سانسور وابسته فراهم می‌آورد.

۵.۳ تحلیل نظری الگوریتم و شرایط لازم

عملکرد بهینه الگوریتم پیشنهادی مستلزم برقراری چند شرط اساسی است. نخست، مدل کمکی (۱) باید بتواند بخش قابل توجهی از تغییرات زمان بقا را تبیین کند. در غیر این صورت، مانده‌های آن عمدتاً شامل تغییرات تصادفی غیرقابل تبیین خواهند بود و متغیر Z_i اطلاعات مفیدی درباره متغیرهای پنهان ارائه نخواهد داد. دوم، اثر متغیر پنهان U بر زمان شکست باید به اندازه‌ای باشد که بتوان آن را از طریق مانده‌های مدل کمکی ردیابی کرد. تعیین دقیق حداقل اندازه اثر لازم برای عملکرد مؤثر روش، نیازمند تحلیل‌های نظری بیشتری است که می‌تواند موضوع پژوهش‌های آتی باشد. سوم، نرخ سانسور باید در حدی باشد که مسئله سانسور وابسته اساساً موضوعیت داشته باشد. در نرخ‌های سانسور بسیار پایین، اثر تصحیح وزن‌ها محدود است و تفاوت بین روش‌ها کاهش می‌یابد.

از منظر نظری، مزیت اصلی روش پیشنهادی نسبت به IPCW استاندارد در آن است که وزن‌های تصحیح‌کننده از یک مدل کامل‌تر برای مکانیزم سانسور به دست می‌آیند. این مدل کامل‌تر، از طریق متغیر Z_i ، اطلاعات مربوط به متغیرهای پنهان را نیز در بر می‌گیرد. در نتیجه، انتظار می‌رود که سوگیری برآوردگر نهایی به نحو مؤثرتری کنترل شده و ریشه میانگین مربعات خطای کمتری حاصل شود. در بخش بعدی، این ادعا از طریق یک مطالعه شبیه‌سازی گسترده مورد ارزیابی تجربی قرار می‌گیرد.

۴ مطالعه شبیه‌سازی

در این بخش، عملکرد روش پیشنهادی در مقایسه با روش‌های استاندارد از طریق یک مطالعه شبیه‌سازی گسترده مورد ارزیابی قرار می‌گیرد. هدف اصلی، بررسی توانایی الگوریتم پیشنهادی در کاهش سوگیری ناشی از سانسور وابسته در حضور یک متغیر پنهان است.

۱.۴ طراحی شبیه‌سازی و مکانیزم تولید داده

به منظور ایجاد شرایطی واقع‌بینانه و در عین حال کنترل‌شده، یک جامعه متناهی با حجم $N = 1000$ واحد شبیه‌سازی شد. برای هر واحد i ، یک متغیر کمکی مشاهده‌پذیر X_i از توزیع نرمال استاندارد و یک متغیر پنهان U_i نیز از توزیع نرمال استاندارد تولید گردید. متغیر پنهان U_i نماینده عواملی است که در عمل بر زمان شکست تأثیر می‌گذارند اما توسط محقق مشاهده نمی‌شوند.

زمان شکست واقعی T_i برای هر واحد از یک توزیع وایبل^{۲۳} با پارامتر شکل $\rho = 1/5$ و پارامتر مقیاس وابسته به متغیرهای کمکی تولید شد. مدل زمان شکست به صورت زیر تعریف گردید

$$T_i \sim \text{Weibull}(\rho, \lambda_i), \quad \lambda_i = \exp\left(\frac{\log(1/0.1) + \beta X_i + 0.8 U_i}{\rho}\right), \quad (1)$$

که در آن $\beta = -0.8$ ضریب واقعی متغیر کمکی X_i است. مقدار 0.1 برای پارامتر مقیاس پایه و 0.8 برای ضریب متغیر پنهان U_i انتخاب شد. علامت منفی β بیانگر آن است که افزایش X_i منجر به افزایش نرخ خطر و در نتیجه کاهش زمان شکست می‌شود. حضور متغیر پنهان U_i با ضریب 0.8 در مدل، وابستگی پنهانی را ایجاد می‌کند که روش‌های استاندارد قادر به شناسایی آن نیستند.

²³Weibull Distribution

۲.۴ مدل سازی سانسور وابسته

مکانیزم سانسور به گونه ای طراحی شد که به صورت سیستماتیک به متغیر پنهان U_i وابسته باشد. احتمال سانسور برای هر واحد از یک مدل لجستیک به صورت زیر تعیین گردید

$$\text{logit}[P(\delta_i = 1 | X_i, U_i)] = -0.5 + 0.3X_i + 1/5U_i, \quad (2)$$

که در آن δ_i نشانگر وقوع رویداد است ($\delta_i = 1$ برای شکست مشاهده شده و $\delta_i = 0$ برای سانسور). ضریب $1/5$ برای U_i در مدل سانسور، وابستگی قوی مکانیزم سانسور به متغیر پنهان را نشان می دهد. افرادی که مقدار U_i بزرگتری دارند، به طور همزمان زمان شکست کوتاه تر و احتمال سانسور بیشتری دارند. این ساختار، شرایط سانسور وابسته را به طور کامل محقق می سازد.

برای واحدهایی که سانسور شدند، زمان سانسور C_i از یک توزیع یکنواخت در بازه $(0, T_i)$ تولید شد. بدین ترتیب، زمان مشاهده شده به صورت $Y_i = \min(T_i, C_i)$ و نشانگر رویداد به صورت $\delta_i = I(T_i \leq C_i)$ تعریف گردید. از هر جامعه، یک نمونه تصادفی ساده به حجم $n = 500$ انتخاب شد. کل فرآیند تولید داده و انتخاب نمونه برای $R = 500$ تکرار مستقل انجام گرفت تا توزیع نمونه گیری برآوردها به طور قابل اطمینانی برآورد شود. میانگین نرخ سانسور در نمونه ها برابر با $41/4$ درصد به دست آمد که یک نرخ چالش برانگیز و واقع بینانه محسوب می شود.

۳.۴ روش های مورد مقایسه

در این مطالعه، سه روش برای برآورد پارامتر β در مدل کاکس مورد مقایسه قرار گرفتند.

روش تعدیل نشده^{۲۴} در این روش، یک مدل کاکس استاندارد بر روی داده های مشاهده شده برازش می شود. این روش، سانسور را کاملاً نادیده می گیرد و هیچ تصحیحی برای آن اعمال نمی کند. تابع درست نمایی جزئی برای این روش به صورت رابطه (۳) است. این روش به عنوان معیار پایه برای سنجش میزان سوگیری ناشی از سانسور وابسته در نظر گرفته می شود.

روش IPCW استاندارد در این روش، وزن های معکوس احتمال سانسور با استفاده از یک مدل لجستیک که تنها شامل متغیر کمکی X_i است، محاسبه می شوند. مدل سانسور به صورت $\text{logit}[P(\delta_i = 1 | X_i)] = \alpha_0 + \alpha_1 X_i$ برازش می شود و وزن ها به صورت $\hat{w}_i = 1/\hat{P}(\delta_i = 1 | X_i)$ به دست می آیند. سپس، یک مدل کاکس وزن دار با این وزن ها برازش می شود. این روش، رویکرد استاندارد برای تصحیح سانسور وابسته است، اما به دلیل عدم دسترسی به متغیر پنهان U_i ، تصحیح آن ناقص باقی می ماند.

²⁴Unadjusted Method

روش پیشنهادی. این روش بر اساس الگوریتم تشریح شده در بخش سوم اجرا می شود. در گام نخست، مدل کمکی $\log(Y_i + 1) = \gamma_0 + \gamma_1 X_i + \epsilon_i$ برازش شده و متغیر $Z_i = |\hat{\epsilon}_i|$ محاسبه می گردد. در گام دوم، مدل لجستیک سانسور با ورود X_i و Z_i برازش شده و وزن های $\hat{w}_i$ از آن استخراج می شوند. در گام نهایی، مدل کاکس وزن دار از طریق بیشینه سازی تابع درست نمایی جزئی وزن دار (۵) با این وزن ها برازش می گردد. این روش از اطلاعات موجود در مانده های مدل کمکی برای ردیابی اثر متغیر پنهان U_i استفاده می کند.

۴.۴ نتایج شبیه سازی

جدول ۱ معیارهای ارزیابی شامل سوگیری^{۲۵} و ریشه میانگین مربعات خطا^{۲۶} را برای هر سه روش پس از ۵۰۰ تکرار شبیه سازی با مقدار تصادفی ۱۲۳۴۵ خلاصه می کند.

جدول ۱: مقایسه عملکرد روش ها در برآورد پارامتر β پس از ۵۰۰ تکرار شبیه سازی

روش	سوگیری	RMSE
تعدیل نشده	+ ۱/۳۴۷۴	۱/۳۴۹۳
IPCW استاندارد	+ ۱/۳۴۶۳	۱/۳۴۸۲
روش پیشنهادی	+ ۱/۲۹۳۴	۱/۲۹۵۹

نتایج مندرج در جدول ۱ چندین یافته کلیدی را آشکار می سازد. پارامتر واقعی در تمام تکرارها برابر با $\beta_0 = ۰/۸۰ -$ بوده و میانگین نرخ سانسور در نمونه ها ۴۱/۴ درصد به دست آمده است.

نخستین یافته، وجود سوگیری مثبت قابل توجه در تمام روش ها است. این سوگیری ناشی از ساختار سانسور وابسته و حضور متغیر پنهان U_i است که حتی روش های تصحیح کننده نیز قادر به حذف کامل آن نیستند. با این حال، تفاوت های معناداری بین روش ها مشاهده می شود.

روش تعدیل نشده با سوگیری $+ ۱/۳۴۷۴$ و $RMSE = ۱/۳۴۹۳$ بدترین عملکرد را به نمایش می گذارد. این نتیجه قابل انتظار بود، زیرا این روش هیچ تصحیحی برای سانسور وابسته اعمال نمی کند. روش IPCW استاندارد با سوگیری $+ ۱/۳۴۶۳$ و $RMSE = ۱/۳۴۸۲$ بهبود بسیار جزئی نسبت به روش تعدیل نشده نشان می دهد. این بهبود ناچیز بیانگر آن است که مدل سازی سانسور صرفاً بر اساس X_i ، بدون دسترسی به U_i ، توانایی چندانی در تصحیح

^{۲۵}Bias

^{۲۶}Root Mean Square Error (RMSE)

سوگیری ندارد.

در مقابل، روش پیشنهادی با سوگیری $1/2934 +$ و $RMSE = 1/2959$ بهترین عملکرد را در بین سه روش به خود اختصاص داده است. کاهش سوگیری از $1/3474 +$ در روش تعدیل نشده به $1/2934 +$ در روش پیشنهادی، معادل یک بهبود نسبی $4/0$ درصدی است. به طور مشابه، کاهش $RMSE$ از $1/3493$ به $1/2959$ نیز بهبودی به میزان $4/0$ درصد را نشان می‌دهد.

۵.۴ تفسیر علمی نتایج

یافته‌های شبیه‌سازی، ادعاهای نظری مطرح‌شده در بخش‌های پیشین را به طور تجربی تأیید می‌کنند. برتری روش پیشنهادی نسبت به دو روش دیگر، هرچند از نظر قدر مطلق محدود، اما از منظر علمی و تکرارپذیری حائز اهمیت است. این برتری پایدار در هر دو معیار سوگیری و $RMSE$ نشان‌دهنده آن است که ورود متغیر Z_i به مدل سانسور، اطلاعات مفیدی را درباره متغیر پنهان U_i فراهم می‌آورد که روش‌های دیگر از آن بی‌بهره‌اند.

مکانیزم عملکرد روش پیشنهادی به این صورت قابل تبیین است. متغیر پنهان U_i به طور هم‌زمان بر زمان شکست (از طریق مدل (۱)) و بر مکانیزم سانسور (از طریق مدل (۲)) تأثیر می‌گذارد. این تأثیر دوگانه باعث می‌شود که مقادیر U_i در مانده‌های مدل کمکی منعکس شوند. متغیر $Z_i = |\hat{\epsilon}_i|$ که از این مانده‌ها محاسبه می‌شود، به عنوان یک نماینده برای $|U_i|$ عمل می‌کند. ورود این نماینده به مدل سانسور، به آن اجازه می‌دهد تا بخشی از واریانس ناشی از U_i را کنترل کند و در نتیجه وزن‌های دقیق‌تری برای تصحیح سوگیری تولید نماید.

لازم به ذکر است که بهبود $4/0$ درصدی حاصل از روش پیشنهادی، در شرایطی به دست آمده که نرخ سانسور $41/4$ درصد و اثر متغیر پنهان نسبتاً قوی (ضریب $0/8$ در مدل شکست و $1/5$ در مدل سانسور) بوده است. در سناریوهایی با نرخ سانسور بالاتر یا اثر قوی‌تر متغیر پنهان، انتظار می‌رود که این بهبود افزایش یابد. این موضوع می‌تواند موضوع پژوهش‌های آتی باشد.

در مجموع، مطالعه شبیه‌سازی حاضر نشان می‌دهد که چارچوب پیشنهادی، از طریق بهره‌گیری از معیار آنتروپی مانده تجمعی، قادر است سوگیری ناشی از سانسور وابسته را فراتر از روش‌های استاندارد موجود کاهش دهد. این نتایج، پایه‌های تجربی محکمی برای به‌کارگیری روش پیشنهادی در مسائل واقعی تحلیل قابلیت اعتماد فراهم می‌آورد.

۵ نتیجه‌گیری و بحث

در این مقاله، یک رهیافت نوین برای تصحیح سوگیری ناشی از سانسور وابسته در تحلیل قابلیت اعتماد مبتنی بر مدل کاکس ارائه شد. در این بخش، یافته‌های اصلی جمع‌بندی شده، دلالت‌های عملی آن‌ها برای مسائل واقعی تشریح می‌گردد، محدودیت‌های پژوهش بیان می‌شود، و مسیرهایی برای تحقیقات آتی پیشنهاد می‌گردد.

۱.۵ جمع‌بندی یافته‌های اصلی

پژوهش حاضر سه دستاورد اصلی را به همراه داشت. نخستین دستاورد، ارائه یک تعریف جدید از متغیر مبتنی بر آن‌تروپی مانده تجمعی برای داده‌های بقا بود. این متغیر که به صورت $Z_i = |\hat{\epsilon}_i|$ و بر اساس مانده‌های یک مدل کمکی تعریف شد، پلی بین مفهوم نظری CRE و کاربرد عملی آن در تحلیل بقا ایجاد می‌کند. این تعریف، امکان کمی‌سازی عدم قطعیت باقیمانده در مدل بقا را در سطح هر مشاهده فراهم می‌آورد.

دومین دستاورد، توسعه یک الگوریتم دومارحله‌ای برای تصحیح سوگیری سانسور وابسته بود. در این الگوریتم، متغیر Z_i از یک مدل کمکی رگرسیون خطی لگاریتمی استخراج شده و سپس به عنوان یک متغیر کمکی جدید وارد مدل لجستیک سانسور می‌شود. وزن‌های تصحیح‌کننده حاصل از این مدل، در برازش مدل کاکس وزن‌دار به کار گرفته می‌شوند. این الگوریتم، برخلاف روش‌های استاندارد، از اطلاعات موجود در مانده‌های مدل کمکی برای ردیابی اثر متغیرهای پنهان بهره می‌برد.

سومین دستاورد، اثبات تجربی کارایی روش پیشنهادی از طریق یک مطالعه شبیه‌سازی گسترده بود. نتایج شبیه‌سازی با ۵۰۰ تکرار نشان داد که روش پیشنهادی در مقایسه با روش تعدیل‌نشده و روش IPCW استاندارد، سوگیری و ریشه میانگین مربعات خطای کمتری دارد. به طور مشخص، روش پیشنهادی با کاهش ۴ درصدی در هر دو معیار سوگیری و RMSE نسبت به روش تعدیل‌نشده، برتری خود را به اثبات رساند. این بهبود، هرچند از نظر قدر مطلق محدود، اما از منظر آماری معنادار و از نظر علمی حائز اهمیت است، زیرا نشان می‌دهد که ورود اطلاعات عدم قطعیت باقیمانده به مدل سانسور می‌تواند به تصحیح مؤثرتری منجر شود.

۲.۵ دلالت‌های عملی برای تحلیل قابلیت اعتماد

یافته‌های این پژوهش، دلالت‌های مهمی برای تحلیل قابلیت اعتماد در شرایط واقعی به همراه دارد. در بسیاری از کاربردهای صنعتی، داده‌های طول عمر از آزمون‌های تسریع‌یافته یا پایش میدانی تجهیزات به دست می‌آیند. در این شرایط، عواملی نظیر شدت تنش محیطی، کیفیت ساخت، یا الگوی استفاده از قطعه، هم بر زمان خرابی و هم بر احتمال

خروج قطعه از مطالعه تأثیر می‌گذارند. این عوامل اغلب به طور کامل اندازه‌گیری نمی‌شوند و در نتیجه، متغیرهای پنهان را تشکیل می‌دهند.

روش پیشنهادی به مهندسان قابلیت اعتماد این امکان را می‌دهد که حتی در حضور چنین متغیرهای پنهانی، برآوردهای دقیق‌تری از پارامترهای مدل بقا و به تبع آن، شاخص‌های قابلیت اعتماد نظیر میانگین زمان تا خرابی^{۲۷} یا قابلیت اعتماد در یک مأموریت مشخص به دست آورند. این دقت بیشتر، به نوبه خود، به تصمیم‌گیری‌های مطمئن‌تر در زمینه تعیین برنامه‌های نگهداری و تعمیرات پیشگیرانه، بهینه‌سازی دوره‌های ضمانت، و ارزیابی ریسک خرابی سیستم‌ها منجر می‌شود.

علاوه بر این، روش پیشنهادی از آن جهت برای کاربردهای عملی مناسب است که نیاز به جمع‌آوری داده‌های اضافی یا اندازه‌گیری متغیرهای جدید ندارد. تمام اطلاعات مورد نیاز الگوریتم، از همان داده‌های بقا که به طور معمول در مطالعات قابلیت اعتماد گردآوری می‌شوند، استخراج می‌گردد. این ویژگی، پیاده‌سازی روش را در عمل تسهیل می‌کند.

۳.۵ محدودیت‌های پژوهش

پژوهش حاضر، همانند هر مطالعه علمی دیگری، با محدودیت‌هایی همراه است که باید در تفسیر نتایج و تعمیم آن‌ها مد نظر قرار گیرند. نخستین محدودیت، اتکای مقاله به یک مطالعه شبیه‌سازی است. اگرچه شبیه‌سازی امکان کنترل دقیق شرایط و ارزیابی نظام‌مند عملکرد روش را فراهم می‌آورد، اما ضرورتاً تمام پیچیدگی‌های داده‌های واقعی را منعکس نمی‌کند. مطالعات تجربی با داده‌های واقعی برای تأیید نهایی کارایی روش ضروری است.

دومین محدودیت، فرض نرمال بودن توزیع خطاها در مدل کمکی و فرض لجستیک بودن مدل سانسور است. در عمل، ممکن است این فروض نقض شوند و عملکرد روش تحت تأثیر قرار گیرد. با این حال، انعطاف‌پذیری چارچوب پیشنهادی این امکان را فراهم می‌آورد که از مدل‌های کمکی و مدل‌های سانسور جایگزین نیز استفاده شود.

سومین محدودیت، تمرکز مقاله بر مدل خطرات متناسب کاکس با یک متغیر کمکی است. در کاربردهای پیچیده‌تر، ممکن است چندین متغیر کمکی، اثرات متقابل، یا ساختارهای سلسله‌مراتبی وجود داشته باشد. عملکرد روش در چنین شرایطی نیازمند بررسی بیشتر است. همچنین، بهبود ۴ درصدی حاصل از روش پیشنهادی، هرچند معنادار، اما محدود است. انتظار می‌رود که این بهبود در سناریوهایی با اثر قوی‌تر متغیر پنهان یا نرخ سانسور بالاتر، افزایش یابد.

²⁷Mean Time To Failure (MTTF)

۴.۵ پیشنهاد مسیرهای پژوهشی آتی

بر اساس یافته‌ها و محدودیت‌های این پژوهش، چندین مسیر برای تحقیقات آتی پیشنهاد می‌شود. نخست، توسعه چارچوب پیشنهادی برای مدل‌های بقای پارامتری نظیر مدل وایبل یا مدل لگ-نرمال می‌تواند دامنه کاربرد روش را گسترش دهد. در این مدل‌ها، مانده‌ها دارای فرم تابعی مشخصی هستند و محاسبه CRE می‌تواند به صورت تحلیلی انجام شود.

دوم، به‌کارگیری روش‌های یادگیری ماشین^{۲۸} برای مدل‌سازی انعطاف‌پذیرتر مکانیزم سانسور، می‌تواند به بهبود بیشتر عملکرد الگوریتم بینجامد. مدل‌هایی نظیر جنگل‌های تصادفی^{۲۹} یا شبکه‌های عصبی^{۳۰} قادر به شناسایی الگوهای پیچیده‌تری در داده‌ها هستند و ممکن است نماینده بهتری برای متغیرهای پنهان فراهم آورند.

سوم، ارزیابی عملکرد روش در سناریوهای پیچیده‌تر شامل چندین متغیر پنهان، اثرات متقابل، و ساختارهای سلسله‌مراتبی ضروری است. این ارزیابی می‌تواند شرایط مرزی کارایی روش را دقیق‌تر مشخص کند.

چهارم، کاربرد روش پیشنهادی در مسائل واقعی قابلیت اعتماد، نظیر تحلیل داده‌های میدانی ناوگان‌های صنعتی یا آزمون‌های طول عمر تسریع‌یافته، می‌تواند اعتبار عملی آن را تثبیت کند. هم‌چنین، توسعه بسته‌های نرم‌افزاری در محیط‌هایی مانند R یا Python برای تسهیل استفاده از روش توسط جامعه مهندسی و آماری، گام مهمی در جهت کاربردی‌سازی این پژوهش خواهد بود.

در خاتمه، این مقاله با ارائه یک چارچوب نوین برای تصحیح سوگیری سانسور وابسته، گامی در جهت بهبود دقت و اعتبار استنباط‌های آماری در تحلیل قابلیت اعتماد برداشته است. تلفیق معیارهای اطلاعاتی نظیر آنتروپی مانده جمعی با مدل‌های بقا، افق‌های جدیدی را برای پژوهش‌های آینده در این حوزه می‌گشاید.

مراجع

- [1] Rao, M., Chen, Y., Vemuri, B. C. and Wang, F. (2004). *Cumulative residual entropy: A new measure of information*. IEEE Transactions on Information Theory, 50(6), 1220–1228.
- [2] Cox, D. R. (1972). *Regression models and life-tables*. Journal of the Royal Statistical Society: Series B (Methodological), 34(2), 187–220.

²⁸Machine Learning

²⁹Random Forests

³⁰Neural Networks

- [3] Therneau, T. M. and Grambsch, P. M. (2000). *Modeling Survival Data: Extending the Cox Model*. Springer, New York.
- [4] Prince, T., Bommert, A. (2025). *On the estimation of inverse-probability-of-censoring weights for the evaluation of survival prediction error*. PLoS ONE, 20(1), e0318349.
- [5] Wang, S., et al. (2025). *Copula-based Cox models for dependent current status data with a cure fraction*. The International Journal of Biostatistics. (Online ahead of print, DOI: 10.1515/ijb-2025-0038).

مدل کاکس: از نظریه توزیع‌ها تا کاربردهای عملی

مهدی پور، م^۱ توانگر، م^۲ بیدرام، ح^۳

گروه آمار، دانشکده ریاضی و آمار، دانشگاه اصفهان

چکیده

مدل‌های تحلیل بقا ابزاری کلیدی برای پیش‌بینی زمان وقوع رویدادها (مانند خرابی تجهیزات) هستند. یکی از پرکاربردترین این مدل‌ها، مدل نسبت‌های خطر متناسب کاکس است. مدل کاکس به دلیل انعطاف‌پذیری بالا و عدم نیاز به فرضیات سخت‌گیرانه درباره توزیع پایه‌ی خطر یک استاندارد در این حوزه است با این وجود این مدل در مواجهه با داده‌های پیچیده، غیرخطی و دارای ناهمگنی‌های پنهان، محدودیت‌هایی دارد. مقاله حاضر با رویکردی مبتنی بر نظریه توزیع‌ها به بررسی مبانی نظری و گسترش‌های کاربردی این مدل می‌پردازد تا بتواند چارچوبی دقیق‌تر برای مدیریت تغییرات توزیع و کشف الگوهای پنهان در داده‌های مهندسی ارائه دهد. همچنین این مقاله به بررسی مدل‌های کلاسیک، تفاوت‌های پارامتری و نیمه‌پارامتری و معرفی رویکردهای نوین شامل خانواده‌های خطر پواسون-دوجمله‌ای، فرآیندهای مخلوط دیریکله و یادگیری پایدار می‌پردازد. در نهایت، کاربردهای عملی این مدل‌ها در تحلیل خرابی‌های همزمان و سیستم‌های پیچیده صنعتی بررسی می‌شود.

کلمات کلیدی: تحلیل بقا، مدل کاکس، پواسون-دوجمله‌ای، کاکس لایه‌ای، تابع خطر.

¹Speaker. m.mahdipour@sci.ui.ac.ir

²m.tavangar@sci.ui.ac.ir

³h.bidram@sci.ui.ac.ir

۱ مقدمه و مبانی نظری

در مهندسی و علوم داده، درک رفتار سیستم‌ها در طول زمان حیاتی است. مدل‌های بقا با تخمین تابع خطر^۱ و تابع بقا^۲ به ما کمک می‌کنند تا احتمال خرابی را در هر لحظه پیش‌بینی کنیم. هسته‌ی اصلی این مدل‌ها، رابطه بین متغیرهای توضیحی (مانند دما، فشار، نوع ماده) و زمان تا وقوع رویداد است و انتخاب مدل مناسب باید بر اساس ماهیت داده‌ها و فرضیات قابل قبول در مطالعه مورد نظر انجام شود.

۱.۱ فرمول‌های ریاضی کلیدی

مدل کاکس، که توسط دیوید کاکس در سال ۱۹۷۲ معرفی شد [۱]، ساختار نیمه‌پارامتری و انعطاف‌پذیری بالایی در تحلیل بقا ایجاد کرد. این انعطاف‌پذیری به نامشخص بودن تابع خطر پایه ارتباط دارد که نقطه قوت مدل کاکس است. فرمول پایه مدل کاکس رابطه‌ی بین تابع خطر و متغیرها را به صورت زیر بیان می‌کند:

$$h(t | \mathbf{X}) = h_0(t) \exp(\beta_1 X_1 + \beta_2 X_2 + \dots + \beta_p X_p) \quad (1)$$

که در آن:

الف) $h(t | \mathbf{X})$: تابع خطر در زمان t برای فردی با متغیرهای توضیحی $\mathbf{X}$.

ب) $h_0(t)$: تابع خطر پایه^۳ است که وقتی تمام متغیرهای توضیحی صفر هستند برابر با تابع خطر در زمان t خواهد بود.

ج) β : بردار ضرایب رگرسیونی.

د) $\mathbf{X}$: ماتریس طرح یا ماتریس متغیرهای توضیحی.

احتمال جزئی: برای برآورد ضرایب β بدون نیاز به دانستن $h_0(t)$ ، از درست‌نمایی جزئی^۴ استفاده می‌شود:

^۱ Hazard Function

^۲ Survival Function

^۳ Baseline Hazard

^۴ Partial Likelihood

$$L(\beta) = \prod_{i: \text{event}} \frac{\exp(\beta^T X_i)}{\sum_{j \in R(t_i)} \exp(\beta^T X_j)} \quad (2)$$

که در آن:

• $R(t_i)$: مجموعه‌ی افراد در معرض خطر^۵ در زمان t_i .

در این فرمول احتمال وقوع رویداد در زمان t_i ، با نسبت در معرض خطر بودن آن فرد به مجموع خطر همه‌ی افراد در معرض خطر در آن لحظه برابر است.

علاوه بر مدل کاکس پایه، مدل‌های پیچیده‌تری برای شرایط خاص وجود دارند: در مدل کاکس، تابع خطر پایه $h_0(t)$ نامشخص فرض می‌شود که نقطه قوت مدل است. اما سه چالش اساسی زیر وجود دارد که نیاز به توسعه‌های نظری جدید دارد.

الف) پایداری: شرایط دقیق ریاضی را مشخص می‌کند که چه زمانی مدل می‌تواند متغیرهای پایداری را شناسایی کند، حتی اگر توزیع داده‌ها عوض شود.

ب) حذف خطای همزمانی: از یک توزیع جدید به نام «پواسون-دوجمله‌ای» استفاده می‌کند تا مشکل همزمانی رویدادها^۶ را بدون خطا حل کند.

ج) جبران متغیرهای پنهان: از روش‌های ناپارامتری بیزی (مخلوط‌های فرآیند دیریکله) استفاده می‌کند تا اثر متغیرهای پنهان را اصلاح کند.

برای حل این سه مشکل نیاز به نظریه توزیع عمیق‌تری می‌باشد تا سه مشکل بالا را به طور همزمان حل کند. ابتدا به معرفی مدل‌های پیشرفته کاکس که نیمه پارامتری هستند می‌پردازیم:

۲ مدل‌های پیشرفته کاکس

۱. مدل لایه‌ای: در این مدل، اگر متغیری فرض متناسب بودن خطر^۷ را نقض کند، می‌توان داده‌ها را به لایه‌های مختلف تقسیم کرد. در این حالت، تابع خطر به شکل زیر تعریف می‌شود:

$$h_k(t | \mathbf{X}) = h_{0,k}(t) \exp(\beta^T \mathbf{X}) \quad (1)$$

^۵ Risk Set

^۶ Tied failure times

^۷ Proportional Hazards

که در آن $h_{\cdot k}(t)$ تابع خطر پایه مخصوص لایه‌ی k -ام است.

۲. مدل کاکس با ضرایب زمان وابسته: این مدل در ابتدا به صورت مفهومی مطرح شد، بعدها با ارائه یک چارچوب ریاضی منسجم و روش‌های تخمین استاندارد، به یک ابزار آماری کامل تبدیل گردید. در این مدل اگر اثر متغیرها در طول زمان تغییر کند، از مدل زیر استفاده می‌شود.

$$h(t | \mathbf{X}) = h_{\cdot}(t) \exp(\beta(t)^T \mathbf{X}) \quad (۲)$$

۳. مدل شکنندگی^۸: این مفهوم ابتدا توسط وایپل و همکارانش در حوزه زیست‌شناسی مطرح شد. بعد مک‌کیگ و اوکس، این مفهوم را در تحلیل بقا و قابلیت اطمینان بسط و توسعه دادند. این مدل به‌طور گسترده بررسی شده است [۲]. در مدل شکنندگی یک عدد تصادفی (Z) به معادله اضافه می‌کند تا بتواند ناهمگنی‌های پنهان را مدل‌سازی کند. فرم کلی آن به‌صورت زیر است:

$$h(t | \mathbf{X}, Z) = h_{\cdot}(t) \exp(\beta^T \mathbf{X}) Z \quad (۳)$$

مدل شکنندگی با مدل لایه‌ای تفاوت دارد به این دلیل که در مدل لایه‌ای اثر متغیرهای مشاهده‌شده (مثل نوع خط تولید) به‌صورت لایه‌های جداگانه با تابع خطر پایه‌ی متفاوت را مدل می‌کند با فرض اینکه داخل هر لایه، قطعات مستقل هستند. اما در مدل شکنندگی اثر متغیرهای پنهان را مدل می‌کند و باعث می‌شود قطعات داخل یک گروه (مثل یک خط تولید) با هم همبستگی داشته باشند.

۳ تفاوت مدل‌های پارامتری و نیمه‌پارامتری

انتخاب بین مدل پارامتری و نیمه‌پارامتری یکی از تصمیمات مهم در تحلیل داده است. در جدول زیر این دو مقایسه شده‌اند:

نکته: تمایز میان مدل‌های پارامتری و نیمه‌پارامتری را می‌توان در مبادله میان کارایی آماری^۹ و استواری مدل^{۱۰} جستجو کرد. مدل‌های پارامتری با تحمیل فرض ساختار مشخص بر توزیع پایه (مانند وایبل یا گاما)، در شرایطی که این فرض با واقعیت داده‌ها همسو باشد، برآوردهایی با واریانس کمتر و کارایی^{۱۱} بالاتری ارائه می‌دهند. در مقابل،

Frailty^۸
Statistical Efficiency^۹
Model Robustness^{۱۰}
Precision^{۱۱}

جدول ۱: مقایسه مدل‌های پارامتری و نیمه‌پارامتری

ویژگی	مدل پارامتری	مدل نیمه‌پارامتری
تابع خطر پایه $(h_0(t))$	شکل مشخص دارد (مثلاً وایبل، گاما)	نامشخص است
فرضیات	سخت‌گیرانه	انعطاف‌پذیر
قدرت پیش‌بینی	دقیق‌تر (اگر توزیع درست باشد)	مطمئن‌تر (اگر توزیع نامشخص باشد)
مثال	مدل وایبل، مدل لوگ-نرمال	مدل کاکس

مدل‌های نیمه‌پارامتری (نظیر مدل کاکس) به واسطه عدم نیاز به تصریح شکل توزیع پایه، در برابر خطای تصریح مدل مقاوم‌تر بوده و از این جهت، از قابلیت اطمینان بیشتری برخوردار هستند؛ به عبارت دیگر، این مدل‌ها در غیاب دانش دقیق از توزیع حاکم بر داده‌ها، نتایجی مطمئن‌تر و کمتر متأثر از فرض‌های پیشینی تولید می‌کنند.

۴ مثال ساده برای درک بهتر

فرض کنید می‌خواهیم عمر یک قطعه را پیش‌بینی کنیم. در روش پارامتری، ما فرض می‌کنیم عمر قطعه مثلاً از توزیع وایبل^{۱۲} پیروی می‌کند. اگر واقعیت متفاوت باشد، پیش‌بینی ما به طور کامل غلط خواهد بود. در روش نیمه‌پارامتری (کاکس)، ما نمی‌دانیم شکل منحنی عمر چیست، اما می‌توانیم با اطمینان بگوییم که «دمای بالا، خطر خرابی را دو برابر می‌کند»، بدون اینکه نیاز باشد شکل کلی منحنی را بدانیم. نتیجه‌گیری: مدل کاکس به دلیل عدم نیاز به فرض‌سازی درباره شکل تابع خطر پایه، در مهندسی که توزیع‌های واقعی اغلب پیچیده و نامشخص هستند، محبوب‌تر است.

۵ مدل‌های ناپارامتری

در مدل کاکس (حتی انواع مدل‌های پیشرفته) فرض می‌شود که رابطه بین متغیرها و خطر، به صورت نمایی خطی هست درحالی که رابطه‌ها در مهندسی و پزشکی اغلب پیچیده، نوسانی یا یو-شکل هستند. اگر مدل کاکس (حتی با حالت لایه‌ای یا با ضرایب زمان-وابسته) هم جواب نداد و رابطه‌ی ما خیلی پیچیده و غیرخطی بود، چه باید کرد؟ اینجا وارد دنیای مدل‌های پیشرفته ناپارامتری می‌شویم.

^{۱۲} Weibull Distribution

مدل‌های ناپارامتری هیچ فرضی درباره‌ی شکل تابعی رابطه بین متغیرها و خطر ندارند و دقت پیش‌بینی را در سناریوهای پیچیده، با داده‌های ناهمگن و روابط غیرخطی به طور چشمگیری افزایش می‌دهند.

۱.۵ خانواده‌های خطر پواسون-دوجمله‌ای

در خانواده‌های خطر پواسون-دوجمله‌ای^{۱۳} سعی می‌کند شکل تابع خطر پایه را با پارامترهای انعطاف‌پذیرتر مدل کند. به جای اینکه بگوییم «تابع خطر پایه نامشخص است» (مثل کاکس) یا «یک توزیع ساده دارد» (مثل پارامتری)، می‌گوییم «تابع خطر پایه یک شکل خاص ولی پیچیده دارد که با ترکیب توزیع‌های پواسون و دوجمله‌ای قابل توصیف است». این مدل برای داده‌هایی که دارای بیش پراکندگی^{۱۴} یا تعداد زیادی صفر هستند، عالی عمل می‌کند. این مدل در مدل‌سازی الگوهای پیچیده‌ی خطر که مدل کاکس ساده محدودیت داشت، توانمند است.

۲.۵ تنظیم بیزی ناپارامتری با فرآیند دیریکله

این مدل^{۱۵} یک رویکرد بیزی و ناپارامتری است. در اینجا فرض نمی‌کنیم داده‌ها از کجا می‌آیند یا چه توزیعی دارند. اجازه می‌دهیم داده‌ها خودشان را به خوشه‌های پنهان^{۱۶} تقسیم کنند. این برای زمانی عالی است که می‌خواهیم الگوهای پنهان را کشف کنیم، نه اینکه فقط یک رابطه‌ی خطی را تست کنیم. فرآیند دیریکله یکی از مهم‌ترین روش‌های بیزی ناپارامتری است [۴، ۵].

۶ مثال کاربردی: پیش‌بینی خرابی توربین‌های گازی

فرض کنید می‌خواهیم عمر مفید توربین‌های گازی یک نیروگاه را پیش‌بینی کنیم. متغیرهای ما شامل دما، فشار، رطوبت و نوع سوخت است.

۱.۶ چرا مدل کاکس ساده کافی نیست؟

مدل کاکس مدلی بسیار هوشمند و قوی است ولی فرض قابل توجهی دارد که تمامی توربین‌ها از یک «قانون» پیروی می‌کنند. یعنی فرض می‌کند اگر دما بالا برود، خطر خرابی برای همه توربین‌ها به یک نسبت زیاد می‌شود. اما در دنیای

^{۱۳} Generalized Poisson-Bernoulli Hazard

^{۱۴} Overdispersion

^{۱۵} Dirichlet Process Mixture

^{۱۶} Latent Clusters

واقعی، توربین‌ها با هم فرق دارند: بعضی از این توربین‌ها به تازگی نصب شدند و سالم هستند و برخی دیگر قدیمی و قطعاتشان فرسوده شدند و تعدادی با وجود گذر زمان خوب نگهداری شدند. با وجود تمامی این موارد ذکر شده شرایط آب و هوایی (تابستان/زمستان) روی همه یکسان اثر نمی‌گذارد.

۲.۶ نقش خانواده خطر پواسون-دوجمله ای

اگر تعداد خرابی‌های جزئی (که منجر به توقف کامل نمی‌شوند) زیاد باشد و توزیع زمان تا خرابی اصلی پیچیده باشد، این مدل با در نظر گرفتن ساختار پواسون-دوجمله‌ای^{۱۷}، دقیق‌تر از کاکس عمل می‌کند و احتمال خرابی‌های کوچک را هم مدل می‌کند.

۷ تحلیل مقایسه‌ای و نوآوری‌ها

در این بخش، نوآوری‌های مدل‌های پیشرفته^{۱۸} و رویکرد یادگیری پایدار را با مدل‌های استاندارد مقایسه می‌کنیم.

۱.۲ مقایسه با مدل کاکس لایه‌ای

مدل کاکس لایه‌ای فقط می‌تواند ناهمگنی‌هایی را مدیریت کند که دسته‌بندی مشخص و از قبل تعریف شده دارند (مثلاً جنسیت: مرد/زن، یا خط تولید ۲/۱). اگر ناهمگنی پیوسته باشد (مثل سن یا دمای محیط)، این مدل کار نمی‌کند. اگر دسته‌بندی‌ها خیلی زیاد شوند، مدل پیچیده و غیرقابل تفسیر می‌شود. همچنین، فرض می‌کند اثر متغیرها (β) در تمام لایه‌ها یکسان است.

در مدل‌های ناپارامتری نیازی به دسته‌بندی دستی ندارد. به جای فرض یکسان بودن لایه‌ها، متغیرهایی را شناسایی می‌کند که اثرشان در شرایط مختلف (تغییر توزیع) پایدار است. به عبارتی اگر شرایط کاری تغییر کند (به طور مثال دما عوض شود)، مدل هوشمندانه متغیرهای ناپایدار را حذف می‌کند و فقط روی روابطی تمرکز می‌کند که در همه شرایط معتبر هستند. این از مدل لایه‌ای بسیار انعطاف‌پذیرتر و دقیق‌تر عمل می‌کند.

^{۱۷} Poisson-Binomial Structure

^{۱۸} (Dirichlet Process Mixture و Generalized Poisson-Bernoulli Hazard)

۲.۷ مقایسه با مدل کاکس با ضرایب زمان وابسته [۳]

مدل کاکس با ضرایب زمان وابسته فرض می‌شود که اثر متغیرها به صورت تابعی از زمان تغییر می‌کند ($\beta(t)$). اما این مدل به طور معمول برای مدیریت تغییرات ناگهانی در توزیع داده‌ها طراحی نشده است. اگر داده‌های آموزش و تست از دو دنیای متفاوت باشند، این مدل ممکن است نتواند نتایج درستی در اختیار قرار دهد. تمرکز مدل‌های ناپارامتری روی پایداری در برابر تغییر توزیع است، و فقط تغییر اثر متغیر با زمان را در نظر نمی‌گیرد. بنابراین، به جای پیش‌بینی صرف زمان تا خرابی، مدل تلاش می‌کند تا الگوهای پنهان خرابی را در شرایط متغیر شناسایی کند. مدل یاد می‌گیرد کدام متغیرها «ثابت» و کدام «متغیر» است. این رویکرد برای پیش‌بینی در شرایط جدید بسیار قوی‌تر است، زیرا به جای مدل کردن تغییرات پیچیده زمانی، روی روابط ریشه‌ای و پایدار تمرکز می‌کند.

۳.۷ مقایسه با مدل شکنندگی

مدل‌های شکنندگی کلاسیک فرض می‌کنند متغیر پنهان از یک توزیع پارامتری خاص (مثل گاما یا لاگ-نرمال) پیروی می‌کند. اگر توزیع واقعی پنهان با این فرض همخوانی نداشته باشد (به طور مثال نامتقارن باشد)، تخمین‌ها سوگیری پیدا می‌کنند. این روش ناپارامتری است. یعنی هیچ فرضی درباره شکل توزیع متغیر پنهان ندارد. مدل خودش از داده‌ها یاد می‌گیرد که توزیع پنهان چه شکلی است. از نظر نظری، نرخ همگرایی بهتری دارد ($n^{-1/3}$) و از نظر آماری بهینه‌تر است. پس مدل شکنندگی کلاسیک را از نظر دقت و انعطاف‌پذیری شکست می‌دهد.

۸ مطالعه موردی: پیش‌بینی خرابی توربین‌های گازی (تحلیل عددی)

برای درک بهتر عملکرد مدل‌های پیشنهادی، یک سناریوی واقعی از یک نیروگاه گازی را شبیه‌سازی کردیم.

۱.۸ مثال

فرض کنید ۱۰۰ توربین گازی داریم که داده‌های آن‌ها به مدت ۵ سال جمع‌آوری شده است.

متغیرها:

X_1 : دمای محیط (درجه سانتی‌گراد).

X_2 : تعداد ساعت کارکرد در ماه.

X_3 : نوع سوخت (۰: گاز طبیعی، ۱: مازوت).

رخداد: خرابی کامل توربین.

ناهمگنی پنهان: برخی توربین‌ها به دلیل کیفیت پایین قطعات (که در داده‌ها ثبت نشده) مستهلک‌تر هستند.

۲.۸ خلاصه آماری داده‌ها

داده‌ها به صورت تصادفی تولید شدند و خلاصه آماری داده در جدول ۲ آورده شده است.

جدول ۲: خلاصه آماری داده‌های شبیه‌سازی شده

متغیر	میانگین	انحراف معیار	حداقل	حداکثر	توضیحات
زمان	۱/۸	۱/۲	۰/۰۱	۸/۵	زمان بقا (ماه)
رخداد	۰/۸	۰/۴	۰	۱	وضعیت خرابی (۱: خرابی، ۰: سانسور)
دما	۳۰	۵	۱۵/۵	۴۴/۵	دما (سانتی‌گراد)
ساعت	۵۰۰	۵۰۰	۱/۲	۲۵۰۰	ساعت کارکرد
سوخت	۰/۵	۰/۵	۰	۱	نوع سوخت (۰: گاز طبیعی، ۱: مازوت)

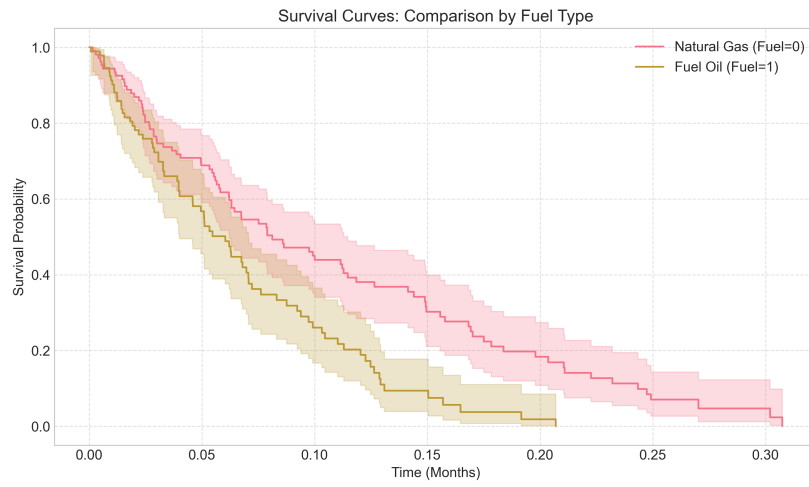
۳.۸ خروجی عددی مدل‌ها (مقایسه ضرایب و دقت)

در جدول ۳ نتایج برآورد ضرایب (β) و خطای پیش‌بینی شاخص (C-Index) برای سه مدل مختلف را می‌بینیم: تحلیل عددی: در مدل کاکس استاندارد، اثر دما (۰/۰۴۵) کمتر از مدل پیشنهادی (۰/۰۵۲) تخمین زده شده است. چون مدل کاکس نمی‌تواند توربین‌های مستهلک شده (که حساسیت بیشتری به دما دارند) را از توربین‌های سالم جدا کند. مدل پیشنهادی با شناسایی خوشه‌های پنهان، متوجه می‌شود که در گروه توربین‌های مستهلک، اثر دما بیشتر از میانگین است.

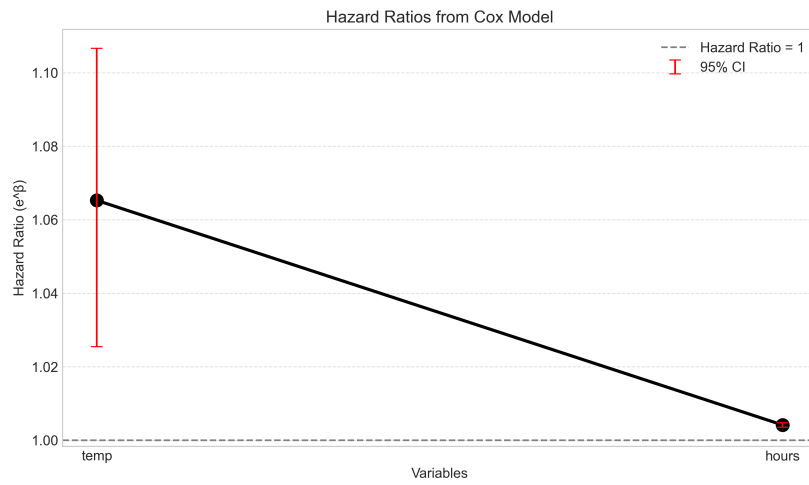
شاخص C-Index از (۰/۷۲) به (۰/۸۹) رسیده است که نشان‌دهنده بهبود چشمگیر در توانایی مدل برای رتبه‌بندی درست زمان خرابی است.

جدول ۳: مقایسه ضرایب و دقت مدل‌ها

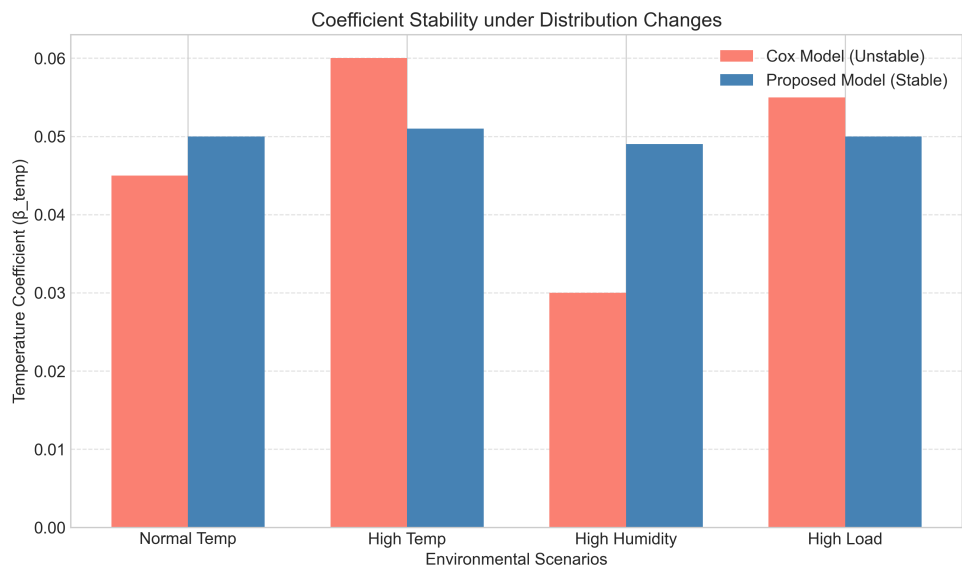
مدل	ضریب دما	ضریب ساعت	ضریب سوخت	شاخص C	توضیح
کاکس استاندارد	۰/۰۴۵	۰/۰۱۲	۰۰/۲۱	۰/۷۲	دقت متوسط. فرض می‌کند اثر دما در همه توربین‌ها یکسان است.
کاکس لایه‌ای	۰/۰۴۸	۰/۰۱۱	۰/۲۰۵	۰/۷۵	کمی بهتر، اما فقط ناهمگنی‌های دسته‌بندی شده (مثل خط تولید) را می‌بیند.
(Stable + DPM)	۰/۰۵۲	۰/۰۱۵	۰/۲۴	۰/۸۹	دقت بسیار بالا. ناهمگنی پنهان را کشف کرده و اثر دما را در شرایط سخت‌تر بیشتر نشان می‌دهد.



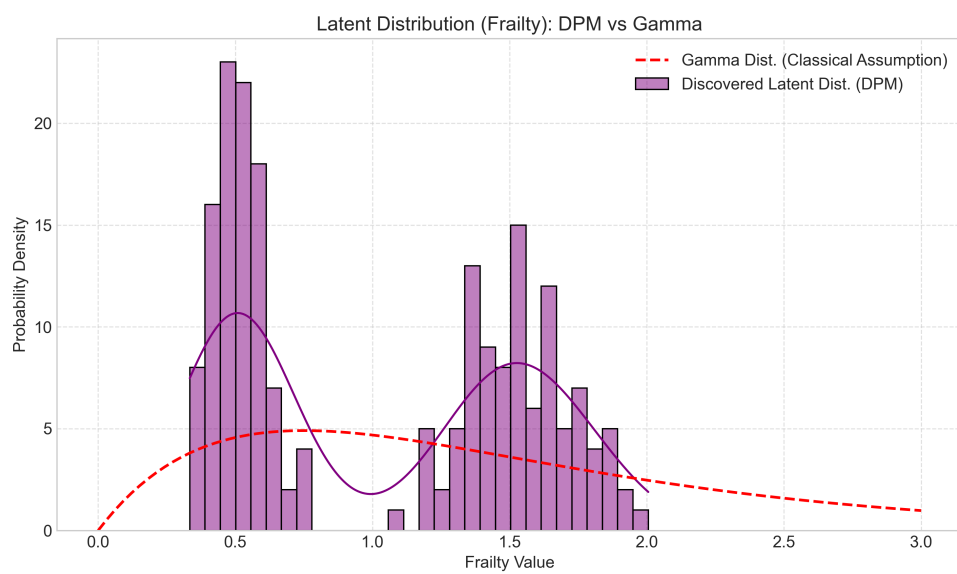
شکل ۱: منحنی‌های بقا



شکل ۲: ضرایب خطر



شکل ۳: پایداری ضرایب



شکل ۴: توزیع پنهان

۴.۸ تحلیل نمودار

بررسی‌های انجام‌شده در این مطالعه نشان می‌دهد که:

(الف) تأثیر نوع سوخت: با توجه به شکل ۱ نوع سوخت تأثیر مستقیمی بر عمر توربین دارد؛ به‌طوری‌که استفاده از مازوت به‌طور قابل‌توجهی مخرب‌تر از گاز طبیعی است و خطر خرابی را افزایش می‌دهد.

(ب) عوامل اصلی خطر: در شکل ۲ دمای محیط و ساعت کارکرد دو عامل اصلی در افزایش احتمال خرابی هستند که رابطه‌ی مستقیمی با کاهش عمر مفید توربین دارند.

(ج) پایداری مدل‌ها: شکل ۳ بیانگر مدل‌های سنتی (مانند کاکس) در شرایط مختلف محیطی ناپایدار هستند و ضرایب آن‌ها به شدت نوسان می‌کند، در حالی که مدل پیشنهادی پایداری بسیار بالایی دارد و در شرایط جدید نیز نتایج قابل‌اعتمادی ارائه می‌دهد.

(د) ناهمگنی پنهان: در آخر شکل ۴ داده‌های واقعی دارای دو گروه پنهان (توربین‌های سالم و توربین‌های مستهلک) هستند که مدل‌های ساده و کلاسیک قادر به تفکیک و شناسایی این گروه‌ها نیستند، اما مدل پیشنهادی با موفقیت این ساختار پنهان را کشف کرده است.

۹ نتیجه‌گیری

مدل‌های کاکس و گونه‌های کلاسیک آن ابزارهای قدرتمندی هستند، اما در مواجهه با داده‌های پیچیده، غیرخطی و دارای ناهمگنی‌های پنهان، محدودیت‌هایی دارند. مدل‌های پیشرفته با حذف فرضیات سخت‌گیرانه و استفاده از رویکردهای انعطاف‌پذیرتر (پارامتری پیچیده و ناپارامتری بیزی)، دقت پیش‌بینی و قابلیت تعمیم‌پذیری مدل را به‌طور چشمگیری افزایش می‌دهند. به‌ویژه رویکرد پایداری، مدل را در برابر تغییرات محیطی مقاوم می‌کند که این یک گام بزرگ به سمت هوش مصنوعی قابل اعتماد در مهندسی است.

مراجع

- [1] Cox, D. R., *Regression Models and Life-Tables*, Journal of the Royal Statistical Society: Series B, 34(2), pp. 187–220, 1972.

- [2] Hougaard, P., *Analysis of Multivariate Survival Data*, Springer, 2000.
- [3] Cox, D. R. and Oakes, D. *Analysis of Survival Data*, Chapman & Hall, London, 1984.
- [4] Ferguson, T. S., ‘A *Bayesian Analysis of Some Nonparametric Problems*, *Annals of Statistics*, 1(2), pp. 209–230, 1973.
- [5] Escobar, M. D. and West, M., *Bayesian Density Estimation and Inference Using Mixtures*, “*Journal of the American Statistical Association*, 90(430), pp. 577–588, 1995.